

CAPITOLO 5

CIRCUITI IN REGIME SINUSOIDALE

COMPLEMENTI

PREMESSA: I NUMERI COMPLESSI (forma rettangolare e forma polare)

CONVERSIONE DALLA FORMA RETTANGOLARE ALLA FORMA POLARE

CONVERSIONE DALLA FORMA POLARE ALLA FORMA RETTANGOLARE

RELAZIONI NOTEVOLI (formula di Eulero e applicazioni)

CONIUGATO DI UN NUMERO COMPLESSO

OPERAZIONI SUI NUMERI COMPLESSI

SOMMA E DIFFERENZA DI NUMERI COMPLESSI

PRODOTTO E RAPPORTO DI NUMERI COMPLESSI

PREMESSA: I NUMERI COMPLESSI

I numeri complessi rappresentano un fondamentale strumento matematico per la trattazione dei circuiti in regime sinusoidale. E' perciò opportuno puntualizzare alcune definizioni e alcune operazioni relative a questi numeri.

Un numero complesso può essere indicato come:

$$\overline{N} = a + jb$$

Quella che abbiamo scritto è l'espressione "binomia" o "rettangolare" del numero complesso, nella quale si ha:

$$a = \operatorname{Re}\left(\overline{N}\right) = \text{parte reale di } \overline{N}$$

$$j = \text{unità immaginaria} = \sqrt{-1}$$

$$b = \operatorname{Im}\left(\overline{N}\right) = \text{parte immaginaria o coefficiente dell'immaginario di } \overline{N}$$

Dato per esempio il numero complesso:

$$\bar{N} = 3 - j4$$

si ha:

$$3 = \operatorname{Re}\left(\bar{N}\right) = \text{parte reale di } \bar{N}$$

$$-4 = \operatorname{Im}\left(\bar{N}\right) = \text{parte immaginaria o coefficiente dell'immaginario di } \bar{N}$$

Vediamo, nella tabella che segue, alcuni esempi di individuazione della parte reale e della parte immaginaria di alcuni numeri complessi:

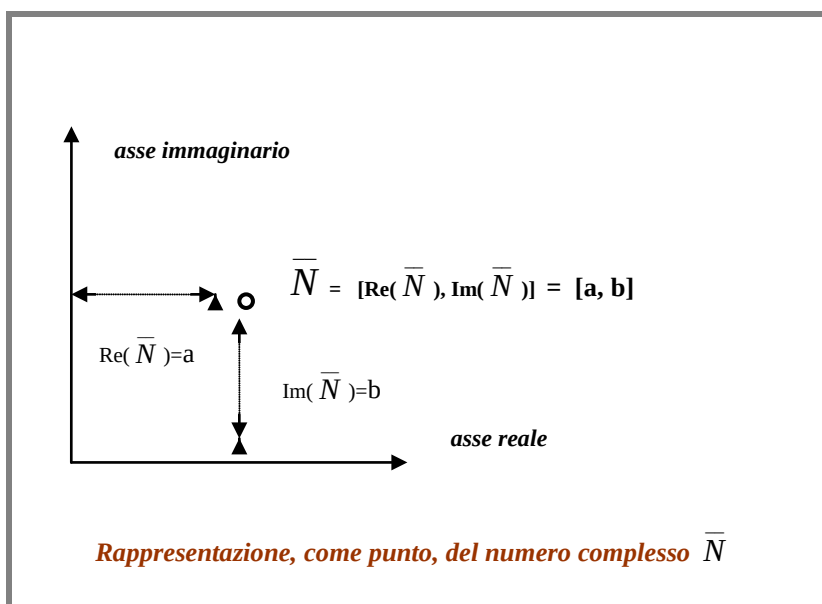
NUMERO COMPLESSO	PARTE REALE	PARTE IMMAGINARIA
$-4+j5$	-4	+5
$2,2-j3$	2,2	-3
77	77	0
$-j3$	0	-3
J	0	1

Da quanto si è detto si vede che i numeri reali sono un sottoinsieme dei numeri complessi e cioè il sottoinsieme dei numeri complessi con parte immaginaria nulla.

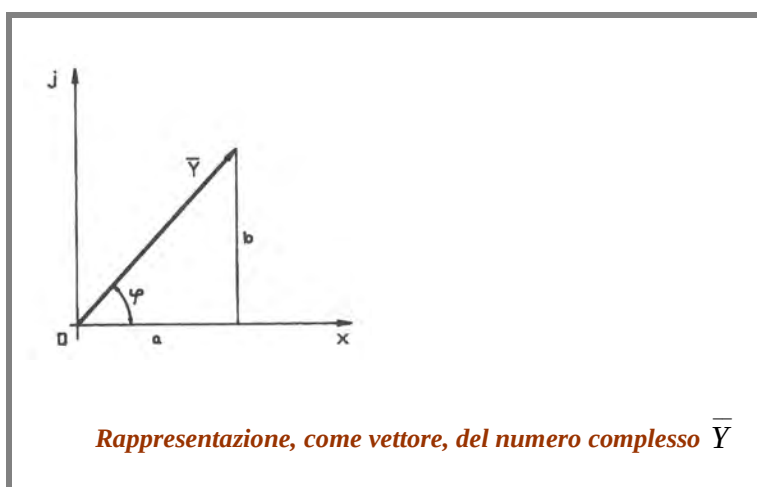
I numeri complessi possono essere rappresentati come vettori o come punti di un piano, detto “piano di Gauss”, nel quale:

- l’asse orizzontale è formato da numeri puramente reali (per esempio: +2, -5,7), cioè numeri con parte immaginaria nulla, perciò è chiamato “asse reale”
- l’asse verticale è formato da numeri puramente immaginari (per esempio: (+j2, -j5 ecc.), cioè numeri con parte reale nulla, perciò è chiamato “asse immaginario”

Se rappresentiamo il numero complesso, come punto, sul piano di Gauss, la parte reale e la parte immaginaria saranno le coordinate del punto.



Se rappresentiamo il numero complesso, come vettore, sul piano di Gauss, la parte reale e la parte immaginaria saranno le coordinate della “punta” della freccia rappresentativa del vettore. L’altro estremo del vettore è l’origine degli assi del piano.



I numeri complessi possono anche essere rappresentati nella forma “esponenziale” o “polare”

$$\overline{N} = |\overline{N}| \cdot e^{j\phi}$$

nella quale:

$$N = |\overline{N}| = \text{modulo di } \overline{N} = \sqrt{a^2 + b^2} = \text{lunghezza del vettore } \overline{N}$$

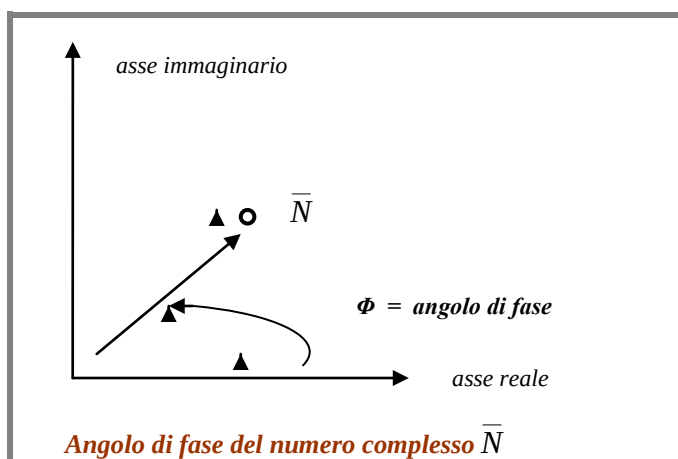
$$\phi = \text{fase di } \overline{N} = \arctg\left(\frac{b}{a}\right) = \text{angolo formato dal vettore } \overline{N} \text{ con il semiasse reale positivo}$$

Abbiamo detto che il modulo è la lunghezza del vettore rappresentativo del numero complesso,

(oppure la distanza del punto rappresentativo del numero dall’origine del piano di Gauss).

La fase Φ è l’angolo (descritto in senso antiorario) che il vettore rappresentativo del numero

complesso forma con il semiasse reale positivo.



Riguardo alla fase, elenchiamo nella tabella seguente i casi nei quali è immediato determinarla:

TIPI DI NUMERI	ESEMPI	FASE
Reali positivi (giacciono sul semiasse orizzontale positivo del piano di Gauss)	+3, +5, +0,22	0°
Reali negativi (giacciono sul semiasse orizzontale negativo del piano di Gauss)	-13, -5, -0,22	±180°
Puramente immaginari positivi (giacciono sul semiasse verticale positivo del piano di Gauss)	+j3, +j, +j0,8922	+90°
Puramente immaginari negativi (giacciono sul semiasse verticale negativo del piano di Gauss)	-j3, -j, -j0,89	-90°

CONVERSIONE DALLA FORMA RETTANGOLARE ALLA FORMA POLARE

Se conosciamo un numero complesso in forma rettangolare (e quindi se conosciamo la parte reale “a” e la parte immaginaria “b”), per convertire il numero in forma polare, dobbiamo determinarne il modulo e la fase. Il calcolo del modulo e della fase si effettua mediante le relazioni:

$$N = \left| \overline{N} \right| = \sqrt{a^2 + b^2}$$

$$\phi = \arctg \left(\frac{b}{a} \right)$$

CONVERSIONE DALLA FORMA POLARE ALLA FORMA RETTANGOLARE

Se conosciamo un numero complesso in forma polare (e quindi se ne conosciamo il modulo N e la fase ϕ), per convertire il numero in forma polare, dobbiamo determinarne la parte reale a e la parte immaginaria b . Il calcolo della parte reale e della parte immaginaria si effettua mediante le relazioni:

$$a = \left| \overline{N} \right| \cdot \cos(\phi) \quad \text{"a" è la proiezione del vettore sull'asse orizzontale}$$

$$b = \left| \overline{N} \right| \cdot \sin(\phi) \quad \text{"b" è la proiezione del vettore sull'asse verticale}$$

Oppure, in modo equivalente, mediante la relazione:

$$\overline{N} = \left| \overline{N} \right| \cdot [\cos(\phi) + j \cdot \sin(\phi)]$$

RELAZIONI NOTEVOLI

Formula di Eulero

$$e^{+j\Phi} = \cos\Phi + j\sin\Phi$$

oppure:

$$e^{-j\Phi} = \cos\Phi - j\sin\Phi$$

- Se valutiamo la formula di Eulero per $\Phi=90^\circ$, otteniamo:

$$e^{+j90^\circ} = \cos 90^\circ + j\sin 90^\circ = 0 + j \cdot 1 = j$$

quindi, in definitiva:

$$e^{+j90^\circ} = j$$

- Se valutiamo la formula di Eulero per $\Phi = -90^\circ$, otteniamo:

$$e^{-j90^\circ} = \cos 90^\circ - j\sin 90^\circ = 0 - j \cdot 1 = -j$$

pertanto, in definitiva:

$$e^{-j90^\circ} = -j$$

- Se valutiamo la formula di Eulero per $\Phi=0^\circ$, otteniamo:

$$e^{+j0^\circ} = \cos 0^\circ + j \sin 0^\circ = 1 + j \cdot 0 = +1$$

quindi, in definitiva:

$$e^{j0^\circ} = +1$$

- Se valutiamo la formula di Eulero per $\Phi=+180^\circ$ e per $\Phi=-180^\circ$, otteniamo:

$$e^{+j180^\circ} = \cos 180^\circ + j \sin 180^\circ = -1 + j \cdot 0 = -1$$

perciò, in definitiva:

$$e^{\pm j180^\circ} = -1$$

- **Quadrato dell'unità immaginaria**

$$j \cdot j = j^2 = (\sqrt{-1})^2 = -1$$

ossia:

$$j^2 = -1$$

- **Reciproco dell'unità immaginaria**

$$\frac{1}{j} = \frac{1 \cdot j}{j \cdot j} = \frac{j}{-1} = -j$$

pertanto, in sintesi:

$$\frac{1}{j} = -j$$

CONIUGATO DI UN NUMERO COMPLESSO

Dato un numero complesso

$$\overline{N} = a + jb$$

il suo **coniugato**

$$\overline{N}^* = a - jb$$

è il numero che ha la **stessa parte reale** e la **parte immaginaria uguale** ma **opposta in segno**.

PROPRIETA'

Considerato il numero complesso

$$\overline{N} = a + jb$$

il **prodotto** di un **numero** per il suo **coniugato** è pari al **quadrato del modulo** del numero:

$$\overline{N} \cdot \overline{N}^* = \left| \overline{N} \right|^2 = a^2 + b^2$$

infatti:

$$\overline{N} \cdot \overline{N}^* = (a + jb) \cdot (a - jb) = a^2 - ja \cdot b + jb \cdot a - j \cdot j \cdot b^2 = a^2 - (-1) \cdot b^2 = a^2 + b^2$$

OPERAZIONI SUI NUMERI COMPLESSI

SOMMA E DIFFERENZA

Devono essere eseguite su numeri espressi in forma binomia o rettangolare.

La somma di due numeri complessi è un numero complesso che ha come parte reale la somma delle parti reali e come parte immaginaria la somma delle parti immaginarie.

La differenza di due numeri complessi è un numero complesso che ha come parte reale la differenza delle parti reali e come parte immaginaria la differenza delle parti immaginarie.

ESEMPI

Dati:

$$\overline{N} = 3 - j4$$

$$\overline{M} = 7 + j14$$

La somma è:

$$\overline{N} + \overline{M} = (3 + 7) + j(-4 + 14) = 10 + j10$$

La differenza è:

$$\overline{N} - \overline{M} = (3 - 7) + j(-4 - 14) = -4 - j18$$

PRODOTTO E RAPPORTO

E' preferibile eseguirli su numeri espressi in forma esponenziale o polare.

Il **prodotto** di due numeri complessi è un numero complesso che ha come **modulo** il **prodotto dei moduli** e come **fase** la **somma** delle fasi.

Il **rapporto** di due numeri complessi è un numero complesso che ha come **modulo** il **rapporto dei moduli** e come **fase** la **differenza delle fasi**.

ESEMPI

Dati:

$$\overline{N} = 8 \cdot e^{j45^\circ}$$

$$\overline{M} = 2 \cdot e^{j15^\circ}$$

Il prodotto è:

$$\overline{M} \cdot \overline{N} = 8 \cdot 2 \cdot e^{j(45^\circ+15^\circ)} = 16 \cdot e^{j60^\circ}$$

Il rapporto è:

$$\frac{\overline{M}}{\overline{N}} = \frac{8}{2} \cdot e^{j(45^\circ-15^\circ)} = 4 \cdot e^{j30^\circ}$$

Prodotto e rapporto possono anche essere eseguiti su numeri in forma binomia secondo le regole del prodotto e del rapporto fra binomi e tenendo conto che:

- $j^2 = jj = -1$
- $1/j = -j$.
- Per **rendere reale il denominatore** del rapporto fra due numeri complessi, basta **moltiplicare, per il coniugato del numero a denominatore**, sia il numero a numeratore che il numero a denominatore

ESEMPIO

Dati:

$$\bar{N} = 3 - j4$$

$$\bar{M} = 7 + j14$$

Il prodotto si può eseguire nel modo seguente:

$$\bar{M} \cdot \bar{N} = (3 - j4) \cdot (7 + j14) = 3 \cdot 7 + 3 \cdot j14 - j4 \cdot 7 - j4 \cdot j14 \rightarrow$$

$$\bar{M} \cdot \bar{N} = 21 + j42 - j28 - j \cdot j56 = 21 + j42 - j28 + 56 \rightarrow$$

$$\bar{M} \cdot \bar{N} = 21 + 56 + j(42 - 28) = 77 + j14$$

CAPITOLO 6

MOTORI E GENERATORI ELETTRICI

COMPLEMENTI

FUNZIONAMENTO DEL MOTORE C.C. SU QUATTRO QUADRANTI

REOSTATO DI AVVIAMENTO

DIFFERENTI CONFIGURAZIONI DEL MOTORE C.C. (motore a eccitazione separata o indipendente; motore a eccitazione derivata; motore a eccitazione-serie; motore a eccitazione composta)

CONTROLLO DEI MOTORI C.C. CON TECNICA PWM

CONTROLLO A QUATTRO QUADRANTI

FUNZIONAMENTO DEL MOTORE C.C. SU QUATTRO QUADRANTI

Quando si parla di **funzionamento su quattro quadranti** di una macchina elettrica si intende dire che la macchina può funzionare:

- ❖ sia come **MOTORE**, ASSORBENDO energia ELETTRICA
- ❖ sia come **generatore di energia elettrica**, assorbendo energia meccanica.

Inoltre, **in entrambi i tipi di funzionamento**, è possibile la **rotazione dell'albero sia in senso orario che in senso antiorario**.

In questo modo, considerando le possibilità “motore/generatore” e le possibilità “verso orario/verso antiorario”, si ottengono quattro condizioni di funzionamento che possiamo rappresentare sui **quattro quadranti di un piano** avente:

- ❖ l'asse *orizzontale* rappresentativo del *verso n di rotazione*
- ❖ l'asse *verticale* rappresentativo della *coppia c* .

Precisiamo che:

- considereremo convenzionalmente **positivo** il **verso orario** di rotazione
- considereremo convenzionalmente **positiva la potenza P** , se questa è **assorbita dalla macchina**.

Vedremo che:

- la macchina si comporta da **MOTORE** quando la rotazione e la coppia sono **CONCORDI**
- la macchina si comporta da ***GENERATORE*** quando la rotazione e la coppia sono ***DISCORDI***

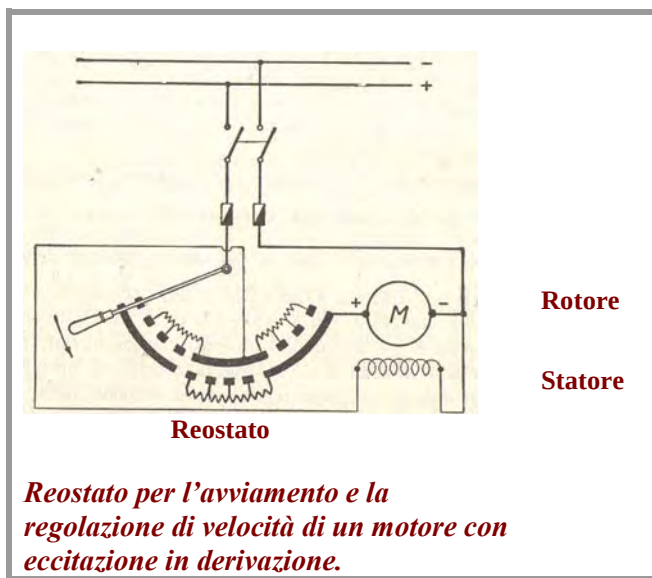
quadrante IV n = negativa = rotazione antioraria c = positiva = coppia oraria DISCORDE alla rotazione GENERATORE in senso antiorario FRENATURA del MOTORE in marcia indietro <i>Asse orizzontale negativo: n (rotazione)</i>	<i>asse verticale positivo:</i> C <i>(coppia)</i> quadrante I n = positiva = rotazione oraria c = positiva = coppia oraria CONCORDE alla rotazione funzionamento diretto o <u>MOTORE</u> in marcia avanti <i>Asse orizzontale positivo: n (rotazione)</i>
quadrante III n = negativa = rotazione antioraria c = negativa = coppia antioraria CONCORDE alla rotazione funzionamento inverso o <u>MOTORE</u> in marcia indietro	quadrante II n = positiva = rotazione oraria c = negativa = coppia antioraria DISCORDE alla rotazione GENERATORE in senso orario FRENATURA del MOTORE in marcia avanti <i>asse verticale negativo:</i> C <i>(coppia)</i>

REOSTATO DI AVVIAMENTO

Nel momento in cui il motore c.c. si avvia, il rotore è fermo e di conseguenza la forza controelettromotrice è nulla.

Di conseguenza il motore assorbe dalla linea di alimentazione una corrente molto intensa che potrebbe non essere sopportata dalla macchina.

Proprio per **limitare questa corrente** si collega **in serie** al motore un **reostato di avviamento**, cioè una resistenza a più sezioni, che viene gradualmente disinserita man mano che il motore aumenta la sua velocità.



DIFFERENTI CONFIGURAZIONI DEL MOTORE C.C.

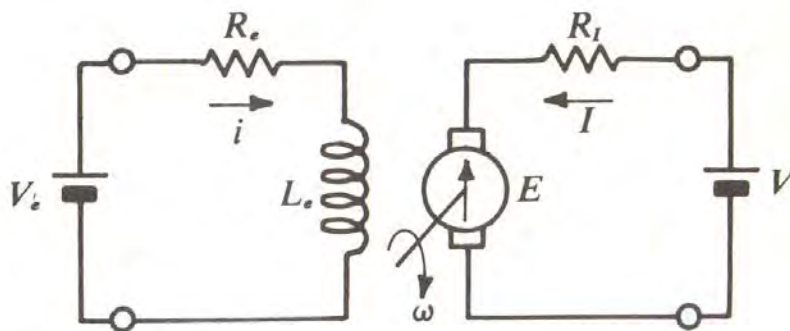
MOTORE C.C. A ECCITAZIONE SEPARATA O INDIPENDENTE

Richiede l'applicazione di **due distinte tensioni di alimentazione**: una al circuito di statore e una al circuito di rotore

La velocità del motore a eccitazione separata o indipendente può essere variata in tre modi:

1. agendo sulla tensione di alimentazione dell'armatura, cioè dell'avvolgimento di rotore; con questa tecnica (regolazione di armatura) la velocità può essere variata di un fattore 30;
2. agendo sul flusso e quindi sullo statore; con questa tecnica (regolazione di campo) la velocità può essere variata di un fattore 5;
3. combinando le due tecniche precedenti e ottenendo una fattore di variazione della velocità pari a 150.

Il motore a eccitazione separata o indipendente è indicato in quelle **applicazioni** che richiedono **regolazioni di velocità ampie e precise**.



*Circuito equivalente per l'eccitazione SEPARATA o INDIPENDENTE
L'induttanza L_s e la resistenza R_s (a sinistra) schematizzano lo statore.
 E e la resistenza R_l (a destra) schematizzano il rotore.*

MOTORE C.C. A ECCITAZIONE DERIVATA

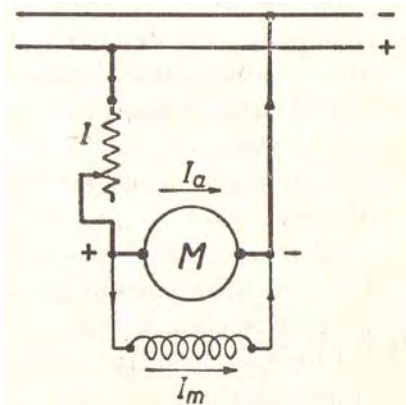
Richiede l'applicazione di una sola **tensione di alimentazione**.

I circuiti di **rotore** e di **statore** sono in **parallelo**.

Se il motore, come di norma avviene, è **alimentato a tensione costante**, allora:

- assorbe dalla linea una corrente proporzionale alla coppia resistente che deve vincere;
- presenta una **velocità praticamente costante al variare del carico**.

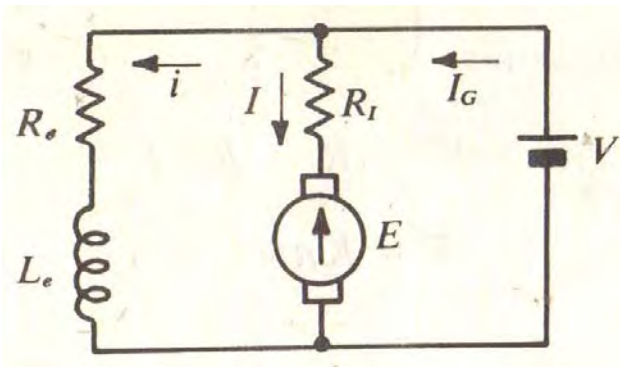
Può passare dal funzionamento da motore al funzionamento da generatore continuando a girare nello stesso verso.



Schema di inserzione di motore eccitato IN DERIVAZIONE

L'induttanza schematizza lo statore.

M schematizza il rotore.



Circuito equivalente per l'eccitazione IN DERIVAZIONE

L'induttanza L_e e la resistenza R_e (a sinistra) schematizzano lo statore.

E e la resistenza R_l (a destra) schematizzano il rotore.

MOTORE C.C. A ECCITAZIONE-SERIE

Richiede l'applicazione di una sola **tensione di alimentazione**.

I circuiti di **rotore** e di **statore** sono collegati in **serie**.

Il flusso aumenta all'aumentare della corrente di carico e quindi il numero di giri diminuisce fortemente al crescere del carico.

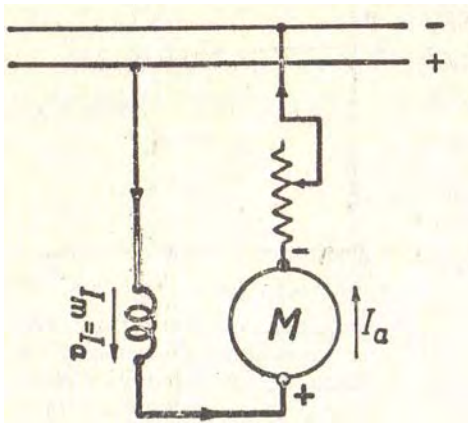
La velocità a vuoto assume valori molto elevati.

E' in genere **alimentato a tensione costante**.

E' indicato per:

- **trazione elettrica (locomotive ecc.)**
- **sollevamento (argani, gru, montacarichi)**

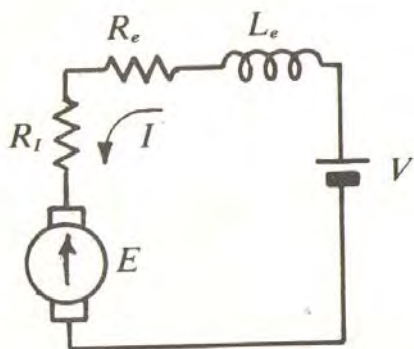
Può passare dal funzionamento da motore al funzionamento da generatore invertendo il senso di rotazione.



Inserzione di motore eccitato IN SERIE

L'induttanza schematizza lo statore.

M schematizza il rotore.



Circuito equivalente per l'eccitazione SERIE

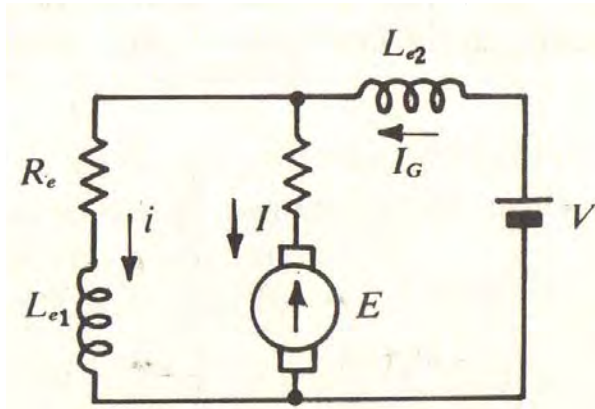
L'induttanza L_e e la resistenza R_e (in alto) schematizzano lo statore.

E e la resistenza R_i (a sinistra) schematizzano il rotore.

MOTORE C.C. A ECCITAZIONE COMPOSTA

Richiede l'applicazione di una sola **tensione di alimentazione**.

L'avvolgimento di eccitazione di **statore** è **in parte in serie** e **in parte in parallelo** a quello di **rotore**.



Circuito equivalente per l'eccitazione COMPOSTA

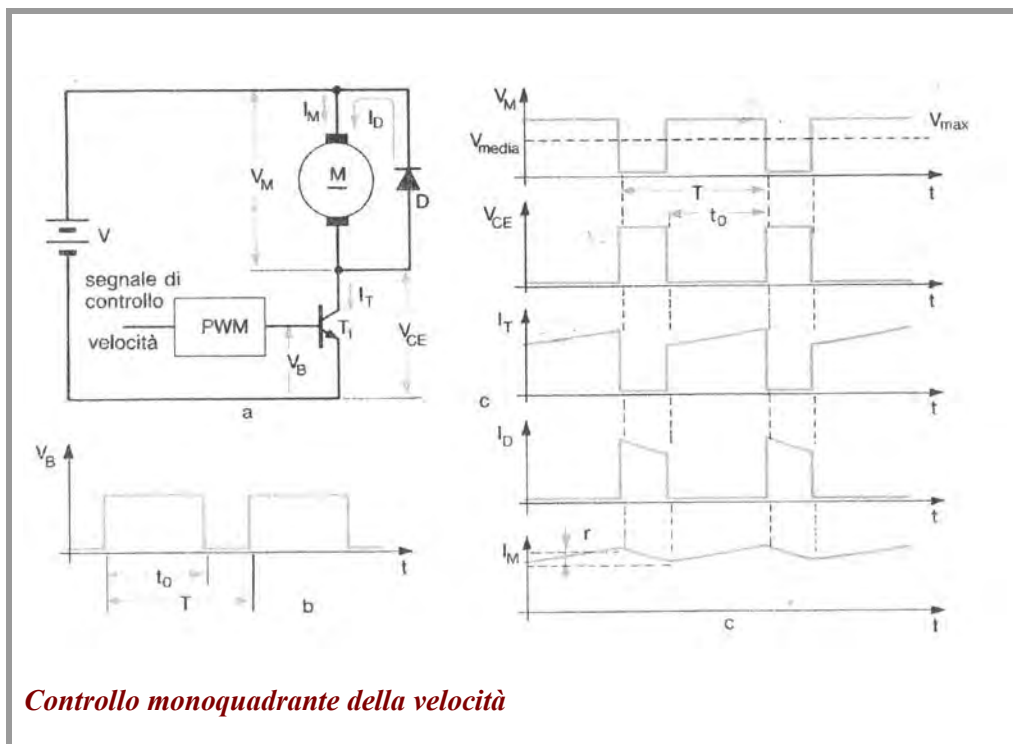
Le induttanze L_{e1} ed L_2 e la resistenza R_e schematizzano lo statore.

CONTROLLO DEI MOTORI C.C. CON TECNICA PWM

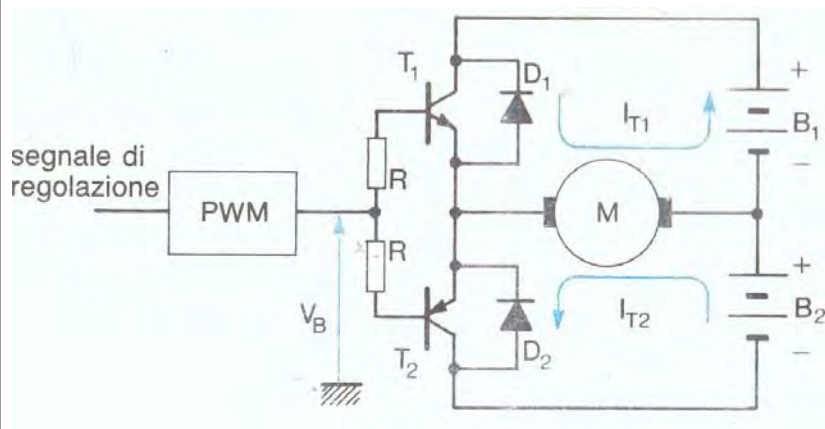
Siccome la velocità “n” del motore c.c. è proporzionale alla tensione di alimentazione, si potrebbe pensare di inserire un reostato per regolare la velocità. Si avrebbero però i seguenti inconvenienti:

- ❖ dissipazione di energia sulla resistenza e conseguente diminuzione del rendimento;
- ❖ riduzione della corrente e quindi della coppia.

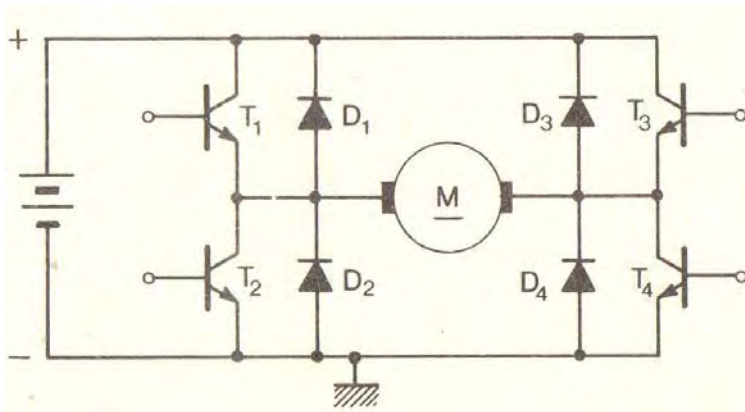
Perciò si ricorre a una tecnica diversa, basata sulla **modulazione PWM**.



CONTROLLO A QUATTRO QUADRANTI



*Circuito di regolazione a "T" a quattro quadranti
con due tensioni di alimentazione*



*Circuito di regolazione ad "H" a quattro quadranti
con una sola tensione di alimentazione*

CAPITOLO 7

SISTEMI TRIFASI

OBIETTIVI

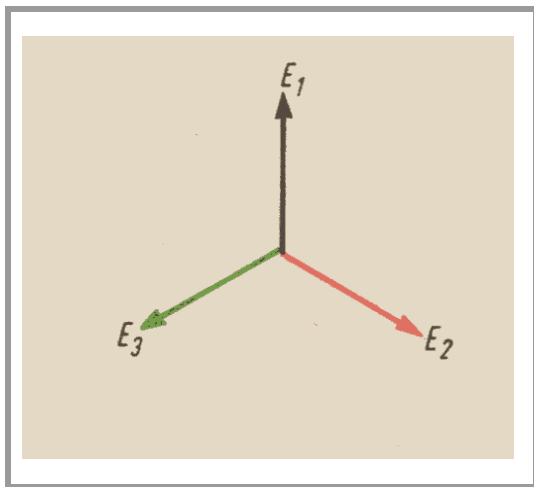
- **comprendere il funzionamento e l'utilità dei sistemi trifasi**
- **conoscere i parametri caratteristici dei sistemi trifasi**
- **conoscere le configurazioni, le tipologie e i parametri dei sistemi trifasi**

PREREQUISITI

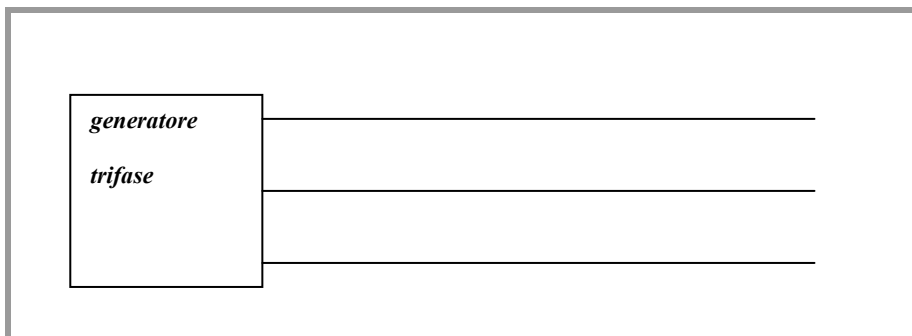
- **Capitoli 2, 3, 5, 6**

SISTEMI TRIFASI SIMMETRICI

Un sistema trifase è un insieme di tre tensioni sinusoidali, con la stessa frequenza, sfasate di 120° una rispetto all'altra.



I tre “**fili**” del sistema, essendo interessati ciascuno da una tensione caratterizzata da una propria fase, sono chiamati “**fasi**” del sistema.



Se, come nella rappresentazione vettoriale precedente, **le tre tensioni** hanno **UGUALE VALORE EFFICACE**, si parla di **sistema trifase SIMMETRICO**. Noi ci occuperemo solo di sistemi SIMMETRICI¹.

¹ E' invece DISSIMMETRICO un sistema trifase di tre tensioni sinusoidali, con la stessa

I sistemi trifase sono tipicamente utilizzati nella **produzione e distribuzione di energia elettrica**, in impianti che coinvolgono **potenze elettriche di migliaia di megawatt**.

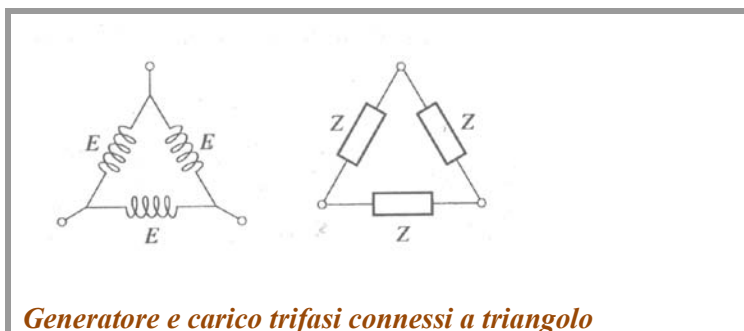
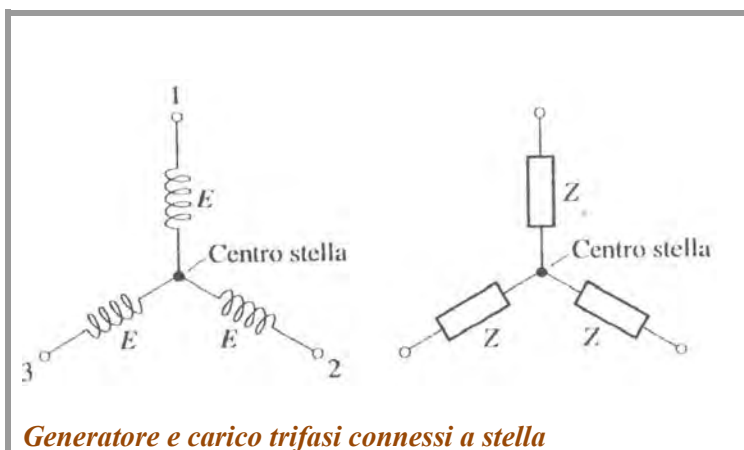
La generazione di un sistema di tensioni trifase è ottenuta mediante un'unica macchina: **l'alternatore trifase**.

Questo alternatore genera un campo magnetico induttore che agisce su **tre avvolgimenti indotti uguali**, situati sulla periferia della macchina e **distanziati tra loro di un angolo corrispondente allo sfasamento** che si vuole ottenere fra le forze elettromotrici che vengono indotte nei singoli avvolgimenti.

Sia i generatori che gli utilizzatori (o carichi) trifasi hanno una struttura interna che può essere basata su due differenti tipi di collegamento:

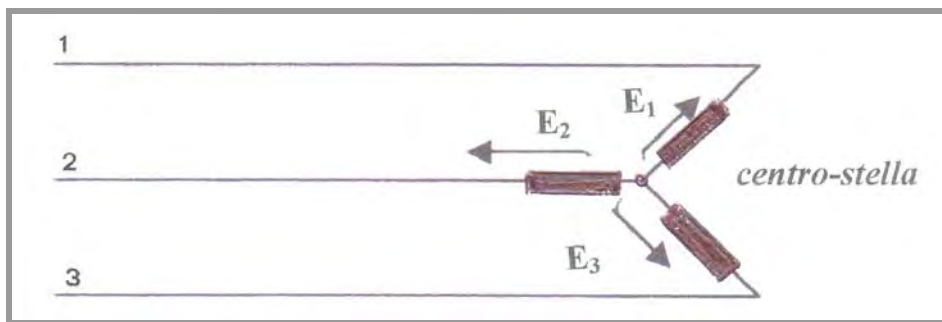
- collegamento **a stella**
- collegamento **a triangolo**.

frequenza ma con **valori efficaci differenti una dall'altra** o sfasamenti diversi **120°** una rispetto all'altra.



COLLEGAMENTO A STELLA, TENSIONI STELLATE E TENSIONI CONCATENATE

Premettiamo che possono essere collegati a stella (oppure, come vedremo, "a triangolo") sia gli utilizzatori che i generatori, occupiamoci ora degli utilizzatori.



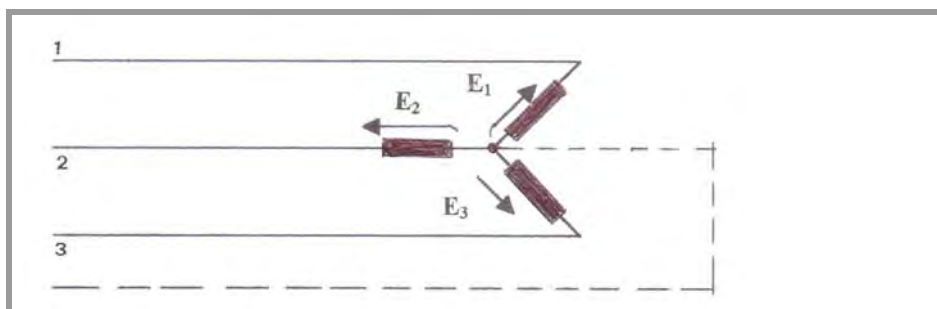
Le **tensioni** presenti fra ciascun filo e il **centro-stella** sono chiamate **tensioni stellate** o tensioni **di fase** e sono indicate con E_1 , E_2 , E_3 .

La loro somma vettoriale è nulla:

$$\overline{E_1} + \overline{E_2} + \overline{E_3} = 0$$

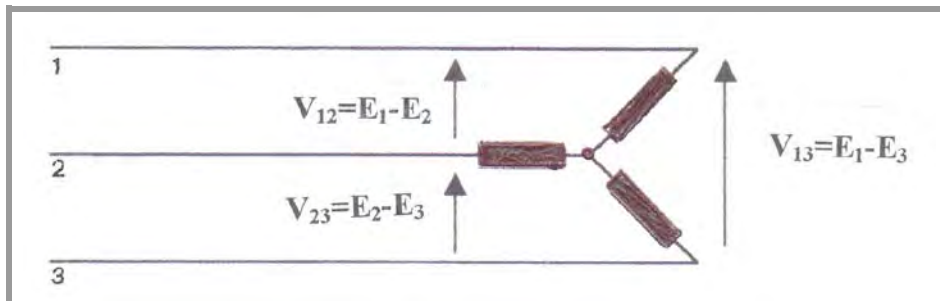
Il **centro-stella** è il **punto attraverso il quale avviene il passaggio di corrente da una fase (filo) a un'altra**, in modo tale che, per ciascuna fase, le altre due costituiscano il circuito di ritorno.

Un sistema trifase può anche essere a quattro fili, *derivando un filo dal centro stella*, filo che viene detto "**neutro**".



Le tensioni stellate possono essere misurate mediante tre voltmetri, collegati ciascuno fra una delle fasi e il centro-stella.

Per il collegamento a stella, oltre alle tensioni E (tensioni stellate o “di fase”), si definiscono anche le **tensioni CONCATENATE “ V ”** che sono le tensioni misurate **FRA due FASI del sistema** (e non fra ciascuna fase e il centro-stella, come le tensioni concatenate).



Ciascuna **tensione concatenata** è la **differenza vettoriale di due tensioni stellate o di fase**. (I vettori sono indicati i grassetto e senza il trattino sopra)

Per esempio:

$$V_{12} = E_1 - E_2$$

Osserviamo che, data una coppia di fasi (cioè di fili), fra esse si possono misurare due tensioni concatenate uguali e opposte. Per esempio:

$$V_{12} = E_1 - E_2$$

$$V_{21} = E_2 - E_1$$

CORRENTI NEL SISTEMA TRIFASE SIMMETRICO

Se il sistema è **simmetrico**, le espressioni delle correnti sono:

$$I_1 = \frac{E_1}{Z_1}$$

$$I_2 = \frac{E_2}{Z_2}$$

$$I_3 = \frac{E_3}{Z_3}$$

(dove I_i , E_i e Z_i sono vettori) e sono sfasate, rispetto alle relative tensioni stellate, degli angoli φ così definiti:

$$\varphi_1 = \frac{\text{Im}(Z_1)}{\text{Re}(Z_1)} = \frac{X_1}{R_1}$$

$$\varphi_2 = \frac{\text{Im}(Z_2)}{\text{Re}(Z_2)} = \frac{X_2}{R_2}$$

$$\varphi_3 = \frac{\text{Im}(Z_3)}{\text{Re}(Z_3)} = \frac{X_3}{R_3}$$

SISTEMA EQUILIBRATO

Se l'**utilizzatore trifase** è formato da **tre impedenze uguali** tra loro, collegate a stella, si parla di **carico equilibrato**.

RELAZIONE FRA TENSIONI DI FASE (TENSIONI “E” STELLATE) E TENSIONI “V” CONCATENATE

$$\begin{aligned} V &= \sqrt{3} \cdot E \\ V &= 1,73 \cdot E \end{aligned}$$

Le tensioni concatenate V sono 1,73 volte più grandi delle tensioni E stellate. Il fattore 1,73 deriva dal calcolo:

$$V = E \cdot \cos 30^\circ + E \cdot \cos 30^\circ = 2E \cdot \cos 30^\circ = 2E \cdot \frac{\sqrt{3}}{2} = E \cdot \sqrt{3}.$$

Se ipotizziamo un sistema trifase a 230 V (sistema con tensioni stellate, o “di fase” di 230 V), le tensioni concatenate risultano di 400 V in quanto:

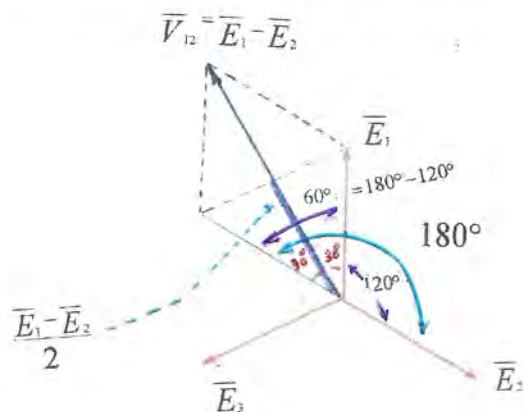
$$V = \sqrt{3} \cdot E = 1,73 \cdot 230 \cong 400V.$$

Se ipotizziamo un sistema trifase a 220 V (sistema con tensioni stellate o “di fase” di 220 V), le tensioni concatenate risultano di 380 V in quanto:

$$V = \sqrt{3} \cdot E = 1,73 \cdot 220 \cong 380V.$$

APPROFONDIMENTO

COME SI RICAVALA LA RELAZIONE FRA TENSIONI DI FASE (TENSIONI "E" STELLATE)
E TENSIONI "V" CONCATENATE



$$\frac{\overline{E}_1 - \overline{E}_2}{2} = \overline{E}_1 \cos(30^\circ) = \overline{E}_1 \cdot 0,866$$

applicando l'espressione appena ricavata si ha:

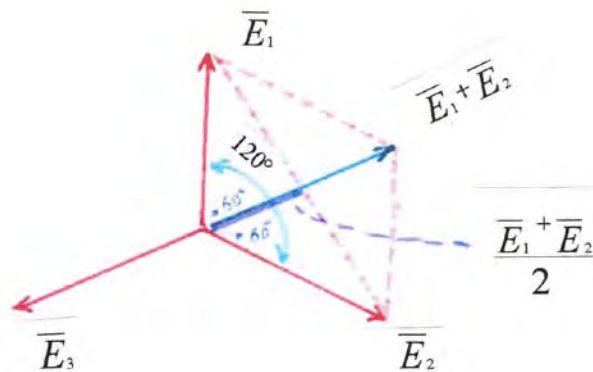
$$\overline{V}_{12} = \overline{E}_1 - \overline{E}_2 = 2 \cdot \frac{\overline{E}_1 - \overline{E}_2}{2} = 2 \cdot \overline{E}_1 \cdot 0,866 = 1,732 \cdot \overline{E}_1$$

ed essendo le tre tensioni stellate uguali e pari ad E:

$$\overline{V}_{12} = 1,732 \cdot \overline{E} = \sqrt{3} \cdot \overline{E}$$

APPROFONDIMENTO

**PERCHE' LA SOMMA DELLE TENSIONI STELLATE E' NULLA
(SE LE TENSIONI SONO UGUALI)**



$$\frac{\vec{E}_1 + \vec{E}_2}{2} = \vec{E}_1 \cdot \cos(60^\circ) = \vec{E}_1 \cdot 0,5$$

$$\vec{E}_1 + \vec{E}_2 = \vec{E}_1 \cdot 0,5 \cdot 2 = \vec{E}_1$$

ed essendo la somma vettoriale $\vec{E}_1 + \vec{E}_2$ (coincidente con $\vec{E}_1$), uguale e opposta al vettore $\vec{E}_3$, si ha $\vec{E}_3 = -\vec{E}_1$ e quindi:

$$\vec{E}_1 + \vec{E}_2 + \vec{E}_3 = \vec{E}_1 + \vec{E}_3 = \vec{E}_1 - \vec{E}_1 = 0$$

$$\vec{E}_1 + \vec{E}_2 + \vec{E}_3 = 0$$

CORRENTI IN UN SISTEMA TRIFASE A TRE FILI

Le tre correnti che percorrono i conduttori si chiamano “correnti di fase” e la loro somma vettoriale, per il primo principio di Kirchhoff applicato al centro-stella, deve risultare **nulla**:

$$I_1 + I_2 + I_3 = 0$$

Il sistema trifase a tre fili è adatto per l'alimentazione di un carico trifase a stella **equilibrato**, mentre non può alimentare un carico squilibrato.

CORRENTI IN UN SISTEMA TRIFASE A QUATTRO FILI

Se il sistema trifase è a quattro fili, cioè se ha un quarto conduttore, detto NEUTRO (derivato dal centro-stella), allora la somma vettoriale delle tre correnti che percorrono i conduttori risulta **DIVERSA da ZERO**:

$$I_1 + I_2 + I_3 = I_0$$

Il sistema trifase a quattro fili è adatto per l'alimentazione di un carico trifase a stella **non equilibrato**.

Il sistema trifase a quattro fili è usato nella **distribuzione cittadina di energia elettrica**.

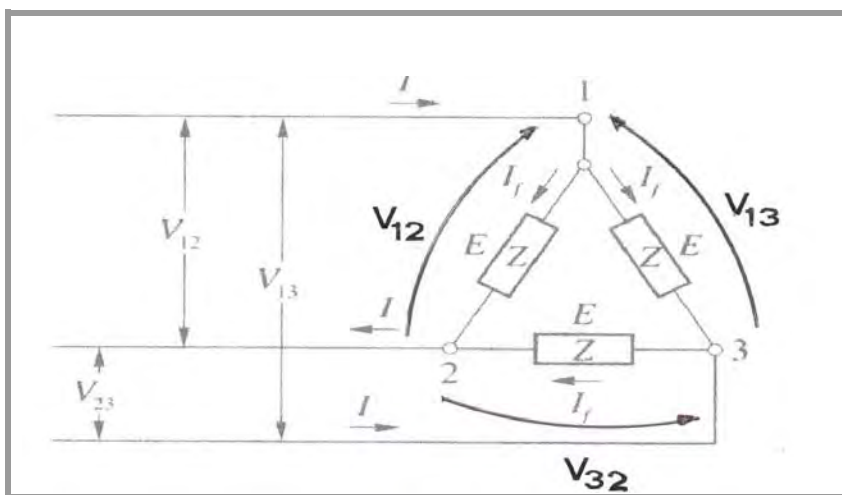
Dalla linea trifase vengono poi derivate le linee monofase delle utenze domestiche.

Se le tre impedenze del carico trifase simmetrico sono poco diverse tra loro (e quindi risultano tali anche le correnti), allora la corrente sul filo neutro risulta molto piccola e pertanto trascurabile.

Da ciò si comprende come sia praticamente possibile realizzare un sistema trifase a soli tre fili sopprimendo o interrompendo il filo neutro.

COLLEGAMENTO A TRIANGOLO

I tre elementi, cioè le tre impedenze del sistema sono collegati uno di seguito all'altro in modo da formare un triangolo dai vertici del quale si dipartono le tre fasi, cioè le tre linee o fili del sistema trifase.



In questo sistema le tensioni V “**CONCATENATE**” possono essere definite come le **ddp fra i capi di ciascuna delle impedenze del triangolo**.

Anche stavolta ogni tensione concatenata V è la differenza di sue tensioni di fase E .

Le **correnti che scorrono nelle fasi o linee** (che portano, per esempio, energia dal generatore al triangolo di carico) vengono chiamate **correnti di linea** (I_{L1} , I_{L2} , I_{L3}).

Invece le **correnti che scorrono nelle tre impedenze** che formano il triangolo vengono chiamate **correnti di fase**.

Siccome ciascuna corrente di fase nasce dalla suddivisione di una corrente di linea, le **correnti di fase** sono sempre **più piccole** delle correnti di linea, il che equivale a dire che **le correnti I_f che scorrono nelle impedenze del triangolo sono più piccole** delle correnti I_L che scorrono nelle linee o fasi del sistema. Quindi nel collegamento a triangolo bisogna fare sempre attenzione alla differenza fra le correnti di fase (più piccole), che scorrono nelle singole impedenze del triangolo, e le correnti di linea (più grandi) che scorrono sui conduttori che hanno una terminazione sui vertici del triangolo.

In particolare, se il sistema è **simmetrico** (**tensioni uguali**) e se sono **uguali** fra loro anche le **impedenze** del triangolo (sistema **equilibrato**) si ha:

$$I_L = \sqrt{3} \cdot I_f$$

$$I_L = 1,73 \cdot I_f$$

Cioè:

- **corrente di linea = 1,73 * corrente di fase**
- **corrente di fase = corrente di linea / 1,73**

Inoltre:

- **La somma vettoriale delle 3 correnti di fase è sempre *nulla***
 $I_{12} + I_{23} + I_{31} = 0$
In generale invece (cioè in sistemi non simmetrici) queste correnti di fase sono diverse fra loro, in quanto il valore di ciascuna di esse dipende dalla tensione concatenata presente su lato del triangolo e dall'impedenza del lato stesso.

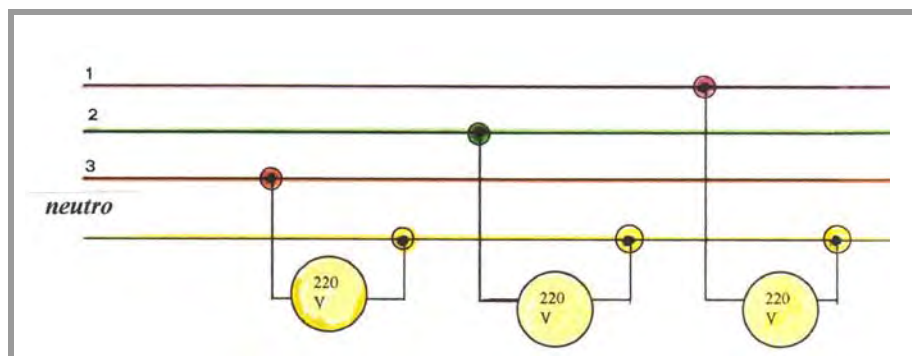
Si ha cioè, per esempio:

$$I_{12} = \frac{V_{12}}{Z_{12}}$$

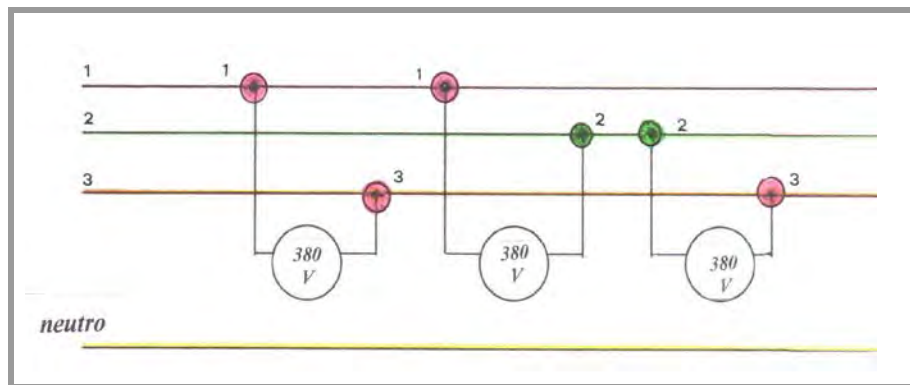
DERIVAZIONE DI COLLEGAMENTI MONOFASI DA UN SISTEMA TRIFASE

Dato un sistema trifase con tensioni “E” (tensioni “di fase” o stellate) di 220 V (e quindi con tensioni concatenate di 380 V), da esso si possono derivare due differenti tensioni monofase:

- 1) una tensione monofase di 220 V ottenuta derivando (cioè collegando) **una qualsiasi fase** con il **NEUTRO** (il che si fa per le **utenze DOMESTICHE**);



- 2) una tensione *monofase di 380 V* ottenuta derivando (cioè collegando) ***DUE FASI tra LORO*** e quindi usufruendo delle *tensioni concatenate “V”* del sistema (il che si fa per i *carichi INDUSTRIALI*).



TRE FILI O QUATTRO FILI?

VINCOLI DELLA SCELTA “TRE FILI/QUATTRO FILI” SUL TIPO DI COLLEGAMENTO STELLA-TRIANGOLO DEI GENERATORI E DEGLI UTILIZZATORI

Se si vuole realizzare una distribuzione di energia elettrica trifase a **TRE fili**, cioè con tre fasi **SENZA neutro**, **NON ci sono vincoli** sul tipo di configurazione **stella-triangolo**, né riguardo alla terna di generatori, né riguardo alla terna di impedenze utilizzatrici.

Se infatti la terna generatrice è a stella a tre fili, allora l'utilizzatore potrà essere a triangolo (nel qual caso non c'è bisogno del neutro) o a stella senza neutro.

Analogamente, se la terna generatrice è a triangolo, allora l'utilizzatore potrà essere a triangolo (nel qual caso non c'è bisogno del neutro) o a stella senza neutro.

Se si vuole realizzare una distribuzione di energia elettrica trifase a **quattro fili**, cioè con tre fasi più il filo **neutro** allora è necessario che le fasi generatrici (cioè la terna dei **generatori**) siano collegate **a stella** (solo in questo modo infatti è possibile realizzare il filo neutro, che ha origine nel centro stella) .

La configurazione trifase con neutro, cioè la configurazione **trifase a quattro fili**, è quella **utilizzata** nella **distribuzione di energia a scopo misto di illuminazione, forza motrice e riscaldamento**.

UTILIZZATORI TRIFASI

L'energia elettrica trifase può essere utilizzata:

- da **apparecchi trifasi propriamente detti**, cioè concepiti specificamente per essere alimentati **direttamente** dalla linea trifase;
- da apparecchi **monofase** che possono essere derivati:
 - singolarmente ***fra due fili*** dalla linea trifase in modo da costituire il lato di un ***triangolo***
 - fra ***un filo di linea e un punto comune*** in modo da dare luogo al ramo di una ***stella***.

UTILIZZATORI TRIFASI PROPRIAMENTE DETTI

Le macchine e gli **apparecchi trifasi** sono **tipicamente** dei **carichi equilibrati** perché sono costituiti da tre circuiti interni uguali. Questi circuiti, essendo alimentati ai loro capi da tensioni simmetriche, assorbono tre correnti uguali fra loro e ugualmente sfasate.

Per queste apparecchiature utilizzatrici sono possibili due tipi di collegamento: **“stella senza neutro”** oppure **“triangolo”**.

La scelta può essere fatta a seconda delle caratteristiche elettriche richieste all'apparecchiatura:

- Collegamento a triangolo: ogni fase dell'utilizzatore (ogni lato del triangolo) deve essere alimentata dalla una tensione concatenata e deve essere percorsa da una corrente 1,73 volte più piccola di quella di linea;
- Collegamento a stella senza filo neutro: ogni fase (ogni lato del triangolo) deve essere alimentata da una tensione stellata (1,73 volte più piccola della concatenata) e deve essere percorsa dall'intera corrente di linea (1,73 volte più grande di quella che attraverserebbe il lato del triangolo).

STELLA O TRIANGOLO?

Per **sistemi trifasi senza neutro**, cioè a **tre fili**, i collegamenti a **stella** e a **triangolo** sono **equivalenti**.

Nel caso che i sistemi siano anche simmetrici ed equilibrati, essi (sia che vengano collegati a stella sia che vengano connessi a triangolo) sono caratterizzati da:

- stessa potenza attiva e reattiva e quindi stesso fattore di potenza
- stesse tensioni di linea
- stesse correnti di linea.

Pertanto per un sistema trifase a tre fili (ossia senza neutro), simmetrico ed equilibrato, la scelta fra collegamento a triangolo o a stella può essere fatta in base ai criteri che seguono.

COLLEGAMENTO A STELLA: QUANDO SI USA

Nel collegamento **a stella** (rispetto a quello a triangolo) ogni ramo deve sopportare una **tensione più bassa** (tensione stellata che è $\sqrt{3}$ volte più piccola della tensione concatenata cui è sottoposto il ramo del triangolo), mentre deve essere percorso dall'intera corrente di linea, cioè da una **corrente più grande** (corrente che è $\sqrt{3}$ volte più grande della corrente -detta "corrente di fase"- che scorre nel ramo del triangolo).

Perciò si usa il collegamento a stella **quando le tensioni di linea del sistema da realizzare sono alte** (in modo che il collegamento a stella, essendo caratterizzato da tensioni più basse sui rami, **determini una riduzione delle tensioni** cui sono sottoposti i rami del sistema, con i relativi vantaggi tecnologici che questa riduzione comporta).

Nel caso dei motori e dei generatori elettrici, per esempio, il fatto di poterli far lavorare a tensioni più basse permette di realizzare gli avvolgimenti con un **minor numero di spire** e di avere meno difficoltà nell'effettuazione dell'isolamento.

COLLEGAMENTO A TRIANGOLO: QUANDO SI USA

Nel collegamento **a triangolo** (rispetto a quello a stella) ogni ramo deve sopportare una **tensione più alta** (tensione concatenata che è $\sqrt{3}$ volte più grande della tensione stellata cui è sottoposto il ramo della stella), mentre deve essere percorso da una **corrente più piccola** (corrente -detta "corrente di fase"- che è $\sqrt{3}$ volte più piccola della corrente che scorre nel ramo della stella).

Perciò si usa il collegamento a triangolo *quando le tensioni di linea del sistema da realizzare sono basse* (in quanto in questa situazione non c'è bisogno della "riduzione di tensione" sui rami che potrebbe essere offerta dal collegamento a stella).

COMMUTAZIONE TRIANGOLO-STELLA

Molti dispositivi sono dotati di una morsettiera, cui fanno capo i terminali degli avvolgimenti di fase, che permette di commutare facilmente il collegamento di questi avvolgimenti da triangolo a stella e viceversa.

OSSERVAZIONE

Uno stesso dispositivo, con le fasi collegate a **triangolo** e funzionante alla tensione **V** può lavorare anche alla tensione di **$1,73V$** se le fasi vengono **collegate a stella**.

Viceversa un dispositivo collegato a stella e funzionante alla tensione **V** può lavorare anche alla tensione di **$V/1,73$** se le fasi vengono **collegate a triangolo**.

PREGI E UTILIZZAZIONE DEI SISTEMI TRIFASE

I pregi dei sistemi trifase sono i seguenti:

- Un sistema trifase (che impiega 3 conduttori) è in grado di **trasportare una potenza elettrica che è 1,73 volte più grande** della potenza trasportabile da una linea monofase (a 2 conduttori) impiegando una **quantità di metallo che è solo 1,5 volte (3/2 volte) maggiore²**; tale aspetto è particolarmente utile nella realizzazione delle **lunghe linee di trasporto dell'energia**;
- possibilità di creare un campo magnetico rotante senza bisogno di ruotare meccanicamente un elettromagnete; questa possibilità (scoperta da Galileo Ferraris) è sfruttata per la realizzazione dei motori asincroni;
- maggiore semplicità e razionalità costruttiva di alternatori (generatori di alternata) e trasformatori trifase rispetto alle analoghe apparecchiature monofase.

² Infatti, per realizzare un sistema monofase, occorrono due fili.

Per realizzare un sistema trifase (senza neutro) servono tre fili.

Se eseguiamo sia il collegamento monofase che il trifase con conduttori identici, il sistema trifase richiederà una quantità di conduttore 3/2 più elevata (1,5 volte più elevata) di quella richiesta dal monofase.

D'altra parte la potenza attiva trasportata dal sistema monofase è $P_1 = V I \cos \Phi$ mentre la potenza attiva trasportata dal sistema trifase è $P_3 = \sqrt{3} V I \cos \Phi = 1,73 V I \cos \Phi$, il che significa che il sistema trifase trasporta (a parità di V e di I) una potenza attiva che è 1,73 volte più grande della potenza del monofase.

Pertanto il sistema trifase risulta più conveniente del monofase perché trasporta una potenza 1,73 volte più grande con una quantità di conduttore che è solo 1,5 volte più grande.

In conclusione il collegamento trifase è più vantaggioso perché consente un aumento di potenza (**rispetto al monofase**) maggiore dell'aumento della quantità di conduttore richiesta.

Per i motivi ora esposti i sistemi trifasi sono **UTILIZZATI** per:

- **generatori di alternata (alternatori)**
- **linee di trasporto dell'energia**
- **trasformatori elevatori e riduttori fra centrali e linee e fra sottostazioni e linee**
- **carichi industriali (motori elettrici).**

CAPITOLO 9

INTRODUZIONE AGLI IMPIANTI ELETTRICI

OBIETTIVI

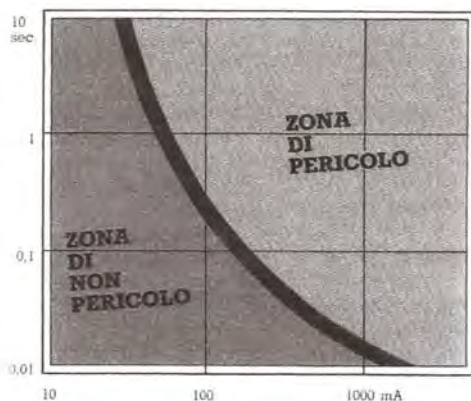
- conoscere le problematiche relative alla pericolosità della corrente elettrica
- conoscere i sistemi di protezione delle persone e degli impianti
- conoscere le linee-guida della progettazione di un impianto elettrico

PREREQUISITI

- Capitoli 2, 3, 4, 5, 6, 7, 8

PERICOLOSITA' DELLA CORRENTE E SISTEMI DI PROTEZIONE

Il pericolo
non è sempre uguale...



... aumenta
con il valore
della corrente
e il tempo
di contatto

La corrente elettrica
provoca nel corpo
umano **danni** che sono
tanto **più gravi**:

- quanto **maggiore**
è l'**intensità**
della corrente
- quanto **maggiore**
è il **tempo** di
contatto

Valori di corrente I in mA	Effetti della corrente sul corpo umano
meno di 10 mA 10 ÷ 50 mA	nessun pericolo di morte possibile morte per asfissia
50 ÷ 500 mA	possibile morte per danni al cuore
oltre 500 mA	possibile morte per paralisi dei centri nervosi

RESISTENZA DEL CORPO UMANO

Il corpo umano ha una sua resistenza elettrica che varia da persona a persona. Possiamo comunque assumere come resistenza del corpo umano, nel caso più sfavorevole, il valore $R_u = 2K\Omega$

SOGLIA DI PERICOLOSITA' DELLA CORRENTE

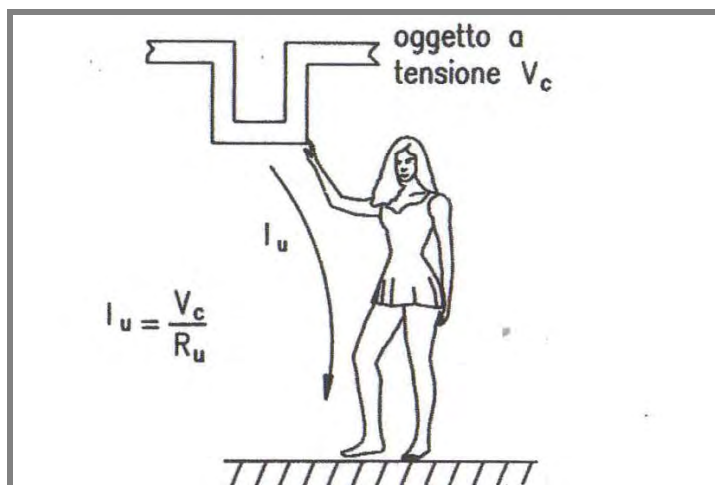
Possiamo considerare come **soglia di pericolosità della corrente** il valore di **30mA** (che è sopportabile per 0,5 sec).

Perciò, un interruttore che apra automaticamente il circuito appena il corpo umano viene sottoposto a una corrente maggiore di 30mA, costituisce una protezione fondamentale.

Questo dispositivo è l'**interruttore differenziale o salvavita**.

TENSIONE DI CONTATTO

La tensione di contatto V_c è la tensione dell'oggetto (cavo o dispositivo elettrico) con il quale il corpo entra in contatto.



La corrente I_U che attraversa il corpo umano quando è sottoposto alla tensione di contatto è data da:

$$I_U = \frac{V_c}{R_U}$$

$$I_U = \frac{V_{RETE}}{R_U} = \frac{230}{2000} = 115 \text{ mA} \quad \text{valore pericoloso!}$$

Quando la tensione di contatto è la tensione di rete, il corpo viene attraversato da una corrente di intensità mortale anche se agisce per pochi decimi di secondo. Nasce quindi la necessità di ricorrere a **sistemi di protezione**

SISTEMI DI PROTEZIONE DELLE PERSONE

- **isolamento** delle **parti attive** dell'impianto
- protezione mediante **involucri** e barriere
- impianti di **messa a terra**
- **interruttori differenziali**¹

SOGLIA DI PERICOLOSITA' DELLA TENSIONE

A una corrente di 30 mA, ritenuta non pericolosa, corrisponde una tensione di contatto² di 60V.

Tuttavia, per motivi precauzionali, viene assunto come **massimo valore di tensione alternata non pericolosa** quello di **25V**

Per le tensioni continue invece il massimo valore di tensione non pericolosa è fissato a 60V.

La corrente **alternata** alle **frequenze industriali** adottate di 50 e 60 costituisce un **pericolo maggiore** in quanto è in grado di indurre **spasmi muscolari** e cardiaca, mentre gli effetti neuroeccitanti (causa di contrazioni e percezione della *scossa*) della corrente continua sono inferiori³. Perciò l'organismo umano può tollerare valori più elevati di corrente continua.

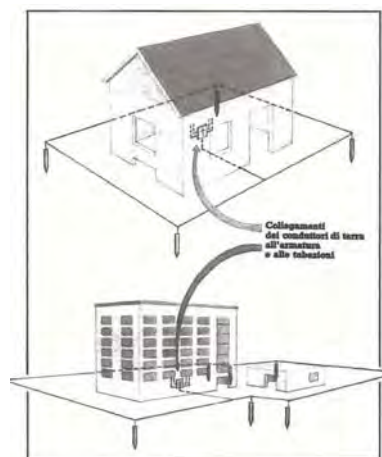
..

¹ L'interruttore magnetotermico è un organo di protezione non delle persone, ma dell'impianto

² $V_c = R_U \cdot I_{Unp} = 2000 \cdot 0,03 = 60V$

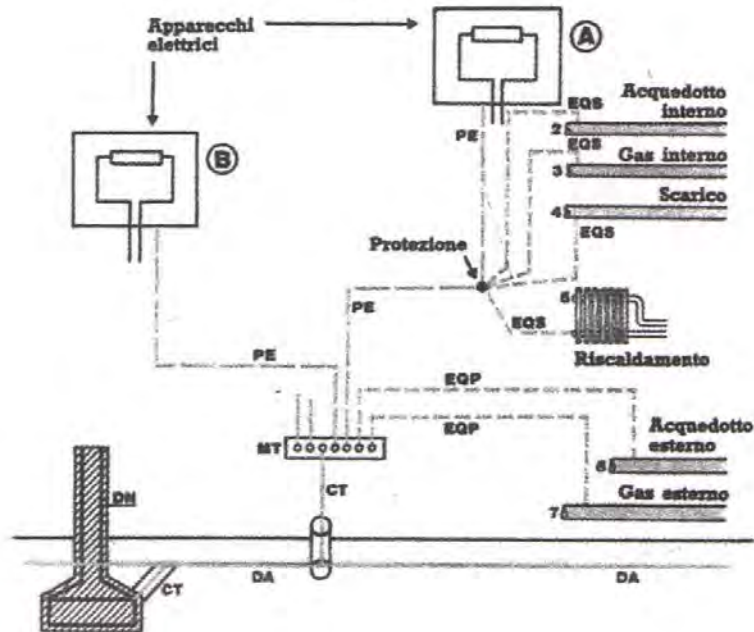
³ poiché le cellule sottoposte ad uno stimolo continuo manifestano un effetto di accomodamento, per cui la soglia di eccitabilità aumenta con il tempo

IMPIANTO DI MESSA A TERRA



*L'impianto
di messa a terra
ha il compito
di disperdere verso terra
le correnti
che si manifestano
in caso di guasto.*

Esempio dei collegamenti di un impianto di terra



LEGENDA

- DA: Dispersore (artificiale)
 DN: Dispersore (naturale)
 CT: Conduttore di terra
 Nota: tratto di conduttore nudo —
 tratto di conduttore non in
 intimo contatto col terreno —
 MT: Collettore di terra
 PE: Conduttore di protezione
 EQP: Conduttori equipotenziali: principali
 EQS: Conduttori equipotenziali: supplementari
 (in locale da bagno)
 A-B: Masse
 2,3,4,5,6,7: Masse estranee

ORGANI DELL'IMPIANTO ELETTRICO

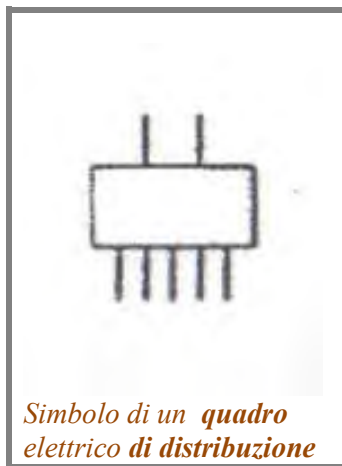
QUADRO ELETTRICO

E' un **nodo di distribuzione** dell'energia elettrica.

Può essere racchiuso in armadi, in banchi, in piccoli contenitori ed è dotato di pulsanti, spie, indicatori e strumenti di misura.

Si configura come un **complesso di apparecchiature** che ha le funzioni di:

- **smistamento** delle linee di alimentazione verso gli utilizzatori
- **centralizzazione** dei dispositivi di comando e misura.



ORGANI DI INTERRUZIONE

SEZIONATORE

E' un organo che ha la funzione **di effettuare l'apertura di un circuito**, attraversato da **corrente trascurabile**, **separando due punti** elettricamente connessi, in modo che, nella posizione di "aperto":

- **non vi sia più continuità metallica** fra i due punti
- il sezionamento sia **visibilmente evidente**
- siano **escluse le richiusure involontarie**

Il sezionatore quindi deve essere in grado di:

- **condurre** in modo *continuativo* un determinato valore della **corrente** di normale *funzionamento*
- **condurre** per un tempo specificato un determinato valore della corrente di funzionamento anormale

	
<i>Simbolo circuitale del sezionatore</i>	<i>Interruttore di manovra-sezionatore</i>

I sezionatori sono impiegati per isolare parti dell'impianto elettrico al fine di effettuare **in completa sicurezza riparazioni e manutenzione.**

Il sezionatore differisce dall'interruttore perché **non è in grado di mantenere la corrente di carico.**

Il sezionatore quindi effettua l'**interruzione** di un circuito elettrico quando il circuito è **senza carico** cioè quando **non** è percorso da corrente (o è percorso da correnti trascurabili).

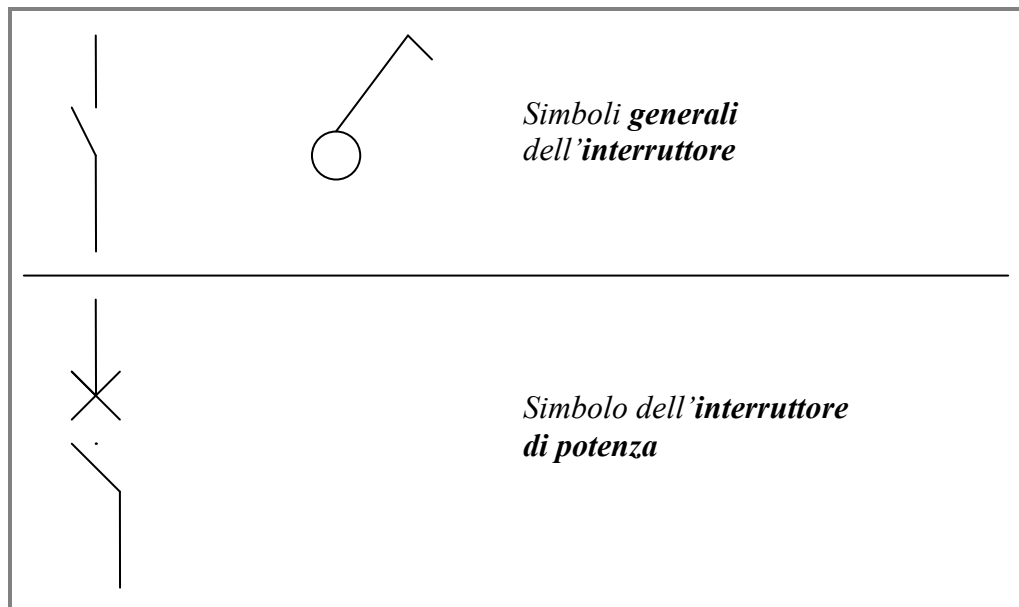
INTERRUTTORE

E' un organo che effettua **l'interruzione o la chiusura di un circuito** elettrico **SOTTO CARICO**, cioè quando il circuito è percorso da corrente sia in **condizioni normali** che in **condizioni di guasto** (fino a uno specificato valore della corrente di guasto).

L'interruttore inoltre deve essere in grado di **condurre continuativamente corrente** fino a un determinato valore.

L'interruttore deve essere quindi in grado di **sopportare i fenomeni energetici** causati dall'**apertura di un circuito percorso da corrente**.

In particolare deve essere in grado di spegnere l'arco elettrico⁴:



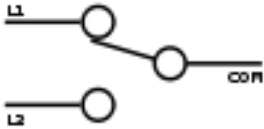
4 L'arco elettrico è l'indesiderata conduzione elettrica che si manifesta in un materiale normalmente isolante (come, per esempio l'aria) posto fra due elettrodi, per effetto dei seguenti fattori:

- diminuzione della superficie di contatto (dell'interruttore che si apre) con conseguente aumento di resistenza,
- riscaldamento locale per effetto Joule, fusione e vaporizzazione del metallo degli elettrodi
- emissione di elettroni da parte degli elettrodi riscaldati
- ionizzazione (acquisizione o perdita di un elettrone periferico da parte dell'atomo) dell'aria in vicinanza degli elettrodi e conseguente presenza di cariche elettriche libere e ionizzazione degli atomi del metallo vaporizzato

DEVIATORE

E' un dispositivo che, invece di interrompere il **flusso di corrente** in un cavo elettrico, come fa l'interruttore, lo **devia su un altro cavo**

Il deviatore è quindi un organo che permette di **indirizzare la corrente** che lo attraversa **su due uscite alternative**.



Simbolo circuitale del deviatore

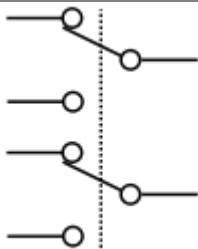


Simbolo elettrico del deviatore

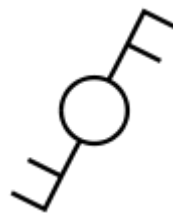
Il deviatore può servire, per esempio, a **spegnere e accendere un utilizzatore** (es. lampadina) **da due punti diversi** di un impianto.

Un'installazione che preveda appunto di comandare gli utilizzatori da due punti viene detta "impianto deviato".

Il deviatore più semplice è quello a una via. Esistono anche deviatori a due o più vie

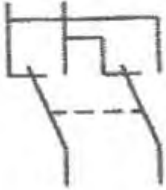


Simbolo circuitale del deviatore a due vie (doppio deviatore)



Simbolo elettrico del deviatore doppio

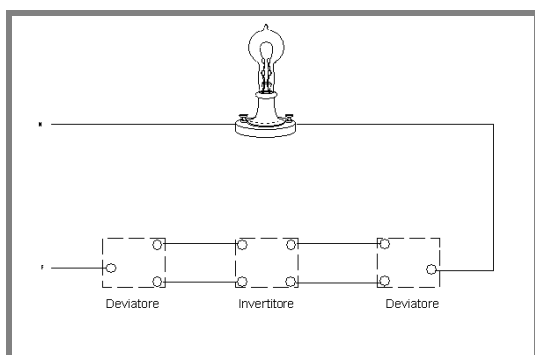
INVERTITORE

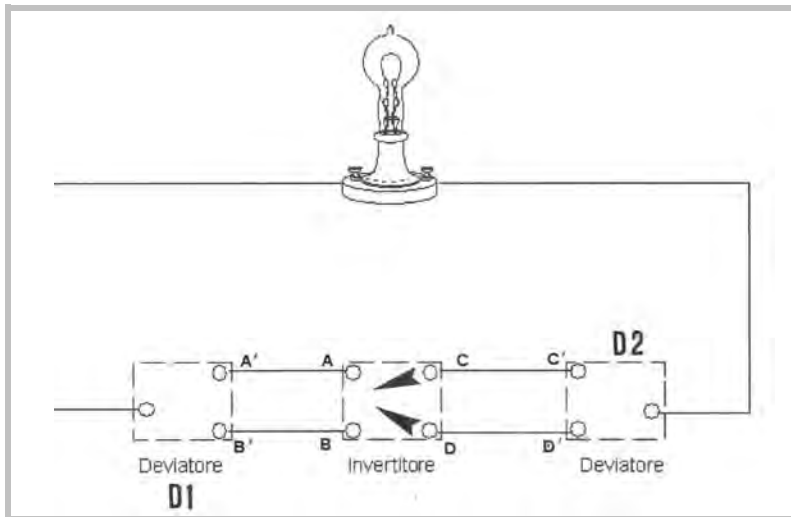
 The image shows two physical components of a switch: a white plastic rocker switch and a metal terminal block with four screw terminals. Below these is a graphic symbol consisting of a circle with four lines extending from it at 45-degree angles, labeled "segno grafico".	 The image shows an alternative electrical symbol for a switch, consisting of a horizontal line with two vertical lines crossing it, and two diagonal lines extending from the crossing points.
<i>Possibile aspetto esterno e simbolo circuitale dell'invertitore</i>	<i>Altro simbolo elettrico</i>

L'invertitore è un apparecchio che serve a **comandare** l'utilizzatore **da almeno tre punti** e che viene installato sempre **in combinazione con due deviatori** (e posizionato fra gli stessi due deviatori).

Un impianto si dice appunto "invertito" quando prevede il comando di un utilizzatore da almeno tre punti, ossia da più di due punti.

Per ogni ulteriore punto di comando (oltre i primi due) bisogna aggiungere un invertitore alla coppia di deviatori.





L'invertitore svolge la funzione d'invertire le due linee d'ingresso con le due linee d'uscita. Più esattamente l'invertitore inverte le linee A'A e B'B che entrano in esso, con le linee CC' e DD' che escono da esso⁵ stabilendo, come nella figura seguente:

i collegamenti CA e DB (orizzontali nel disegno)	i collegamenti DA e CB (incrociati nel disegno)

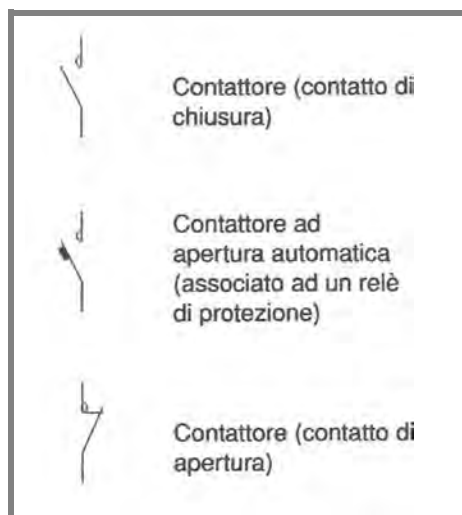
Quindi azionando l'invertitore, cioè “spostandolo” da una configurazione all'altra, si può invertire la situazione della lampadina (lampadina ON o lampadina OFF) creata dai due deviatori.

⁵ L'invertitore può collegare la linea CC' con la linea A'A oppure con la linea B'B e può collegare la linea DD' sempre con la linea A'A oppure con la linea B'B

CONTATTORE o TELERUTTORE

Il contattore è un dispositivo destinato ad **aprire e chiudere il circuito** sia in **condizioni normali** (fino a un certo valore di corrente), sia in condizioni di **sovraccarico**, ma non in condizioni di guasto. E' caratterizzato dal fatto che (contrariamente all'interruttore) può permanere nella **posizione di chiuso** solo in presenza dell'azione di comando, che generalmente è di tipo elettromagnetico. La posizione di **chiuso** cioè è instabile perché si conserva finché dura l'azione di comando, dopodiché i contatti si aprono. Ha *una sola posizione stabile* che è quella di **aperto** (mentre l'interruttore ha due posizioni stabili: aperto e chiuso).

E' in grado inoltre di lavorare con **elevata frequenza di manovra**.

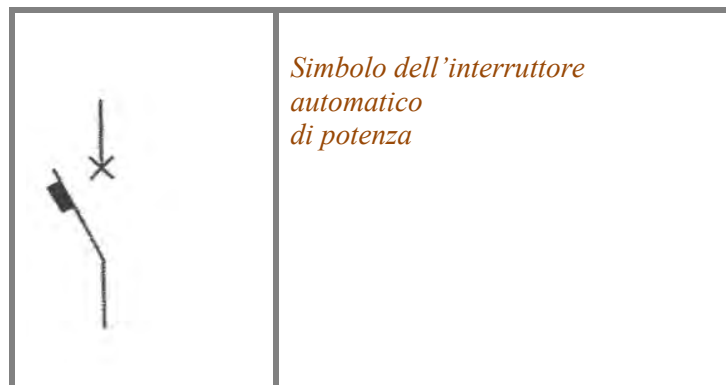


ORGANI DI PROTEZIONE (INTERR. AUTOMATICI E FUSIBILI)

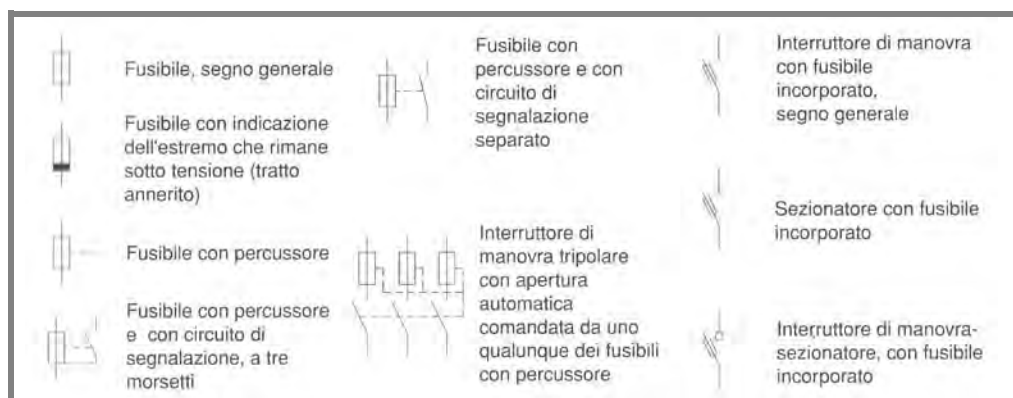
Gli organi di protezione sono dispositivi che intervengono **quando si manifestano anomalie** di funzionamento del circuito stesso.

Gli organi di protezione sono essenzialmente:

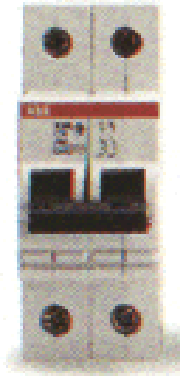
- ❖ gli **INTERRUTTORI AUTOMATICI**, che sono anche organi di MANOVRA e che, negli impianti domestici, sono:
 - l'interruttore magnetotermico (contro cortocircuiti e sovraccarichi)
 - l'interruttore differenziale o salvavita (contro le correnti di dispersione verso terra)



- ❖ i **FUSIBILI**, o valvole fusibili, (che agiscono, come l'interruttore magnetotermico, contro cortocircuiti e sovraccarichi)



INTERRUTTORE AUTOMATICO MAGNETOTERMICO, CORTOCIRCUITO E SOVRACCARICO

 Aspetto esterno	 Simbolo circuitale	<i>Nel simbolo circuitale:</i> <i>il rettangolo superiore rappresenta il funzionamento termico</i> <i>il rettangolo inferiore rappresenta il funzionamento magnetico</i>
---	--	---

Il nome di interruttore magnetotermico deriva dal fatto che quest'organo protegge (come il **fusibile**) nei confronti di fenomeni che determinano un surriscaldamento dei cavi (fenomeni contro i quali l'interruttore differenziale o salvavita non agisce). Questi fenomeni sono:

- il **CORTOCIRCUITO**, che è un **contatto accidentale** fra due **conduttori a potenziale diverso**, contatto che può verificarsi per effetto dell'**isolamento danneggiato**. Si ha cortocircuito quando i due conduttori (es. fase e neutro), fra i quali in condizioni normali c'è una ddp non nulla, entrano in contatto attraverso un'impedenza trascurabile. La corrente allora assume un valore molto maggiore della portata dei cavi. Questa corrente di corto circuito assume valori elevati perché è limitata solo dall'impedenza (molto piccola) del tratto di linea a monte del punto di guasto. La corrente di cc determina una **sovratemperatura** che **danneggia i cavi e le apparecchiature poste a monte del punto di guasto**: la corrente di cortocircuito è una corrente **notevolmente maggiore** della corrente nominale.
- il **SOVRACCARICO** si determina per il collegamento simultaneo di **troppi utilizzatori**, che richiede una **corrente superiore alla portata** dei cavi, che perciò si **surriscaldano**; la corrente di sovraccarico è una corrente **maggiore** della corrente nominale.

L'interruttore magnetotermico è un apparecchio di manovra, per la protezione dei circuiti⁶, che effettua le operazioni di apertura e chiusura di un circuito. L'apertura dell'interruttore è effettuata da uno di due **relé** (facente parte dell'apparecchiatura):

- un relé magnetico (a bobina su nucleo ferromagnetico) che agisce in caso di cortocircuito
- un relé termico (a lamina bimetallica) che agisce in caso di sovraccarico

L'interruttore è caratterizzato dalla **tensione nominale**, cioè dalla tensione del suo normale utilizzo (assegnata dal costruttore). Per i circuiti domestici è di 230 volt. La sua **corrente nominale (I_n)** è invece quella che può circolare senza problemi a una certa temperatura ambiente (indicata sulla targa se diversa da 30°C).

Le **correnti nominali** in uso hanno i seguenti valori espressi in *ampere*:

6	10	13	16	20	25	32	40	50	63	80	100	125
---	----	----	----	----	----	----	----	----	----	----	-----	-----

Le modalità di intervento magnetico sono tre in base ai limiti della corrente di intervento (riferiti alla corrente nominale I_n) in caso di cortocircuito:

TIPO	LIMITI DELLA CORRENTE DI INTERVENTO
B	$3I_n$ --- $5I_n$
C	$5I_n$ --- $10I_n$
D	$10I_n$ --- $20I_n$

Il tipo B interviene per più basse correnti.

⁶ Mentre l'interruttore differenziale o salvavita è per la protezione delle persone

INTERRUTTORE DIFFERENZIALE (SALVAVITA) E CORRENTI DI DISPERSIONE VERSO TERRA

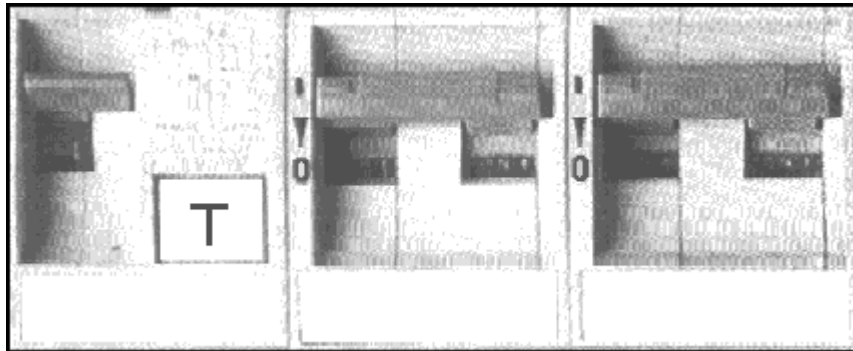
 <i>Simbolo circuitale</i>	<i>L'interruttore differenziale ha la funzione di interrompere il circuito quando si verifica una dispersione verso terra. L'interruzione deve avvenire prima che si determinino danni alle persone.</i> <i>Questo interruttore è sensibile allo squilibrio di corrente che si produce in occasione di una dispersione verso terra.</i> <i>Un interruttore differenziale è caratterizzato da una corrente differenziale di intervento.</i> <i>Se, per esempio la corrente differenziale di intervento è di 30mA, esso apre il circuito appena la corrente di dispersione verso terra raggiunge i 30 mA</i>
---	---

L'interruttore differenziale ha la funzione di interrompere il circuito quando si verifica una dispersione verso terra. L'interruzione deve avvenire prima che si determinino danni alle persone.

Questo interruttore è sensibile allo squilibrio di corrente che si produce in occasione di una dispersione verso terra.

Un interruttore differenziale è caratterizzato da una corrente differenziale di intervento. Se, per esempio la corrente differenziale di intervento è di 30mA, esso apre il circuito appena la corrente di dispersione verso terra raggiunge i 30 mA.

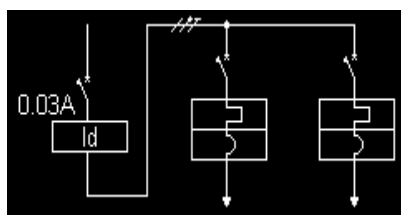
INTERRUTTORE DIFFERENZIALE CON DUE MAGNETOTERMICI



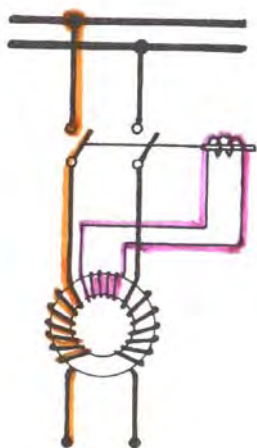
Il primo interruttore, generale (a sinistra nella foto), è il **T**, conosciuto comunemente come **salvavita**, e si individua facilmente per la **presenza di un pulsante**, utile per la **manutenzione**, contrassegnato con la lettera **T**

interruttore di tipo **1**, con cui si comandano e si proteggono i circuiti luce e i circuiti che alimentano le

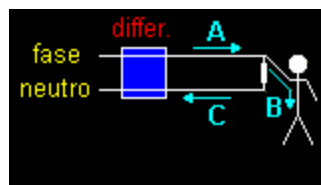
interruttore di tipo **2**, con cui si comandano e si proteggono i circuiti luce e i circuiti che alimentano le



schema elettrico corrispondente ai tre interruttori.



*schema di principio
dell'interruttore
differenziale*



INTERRUTTORI MAGNETOTERMICI-DIFFERENZIALI



Esistono anche interruttori **magnetotermici-differenziali** che racchiudono in un solo componente anche gli **magnetici e termici**.

TIPOLOGIE COSTRUTTIVE

TIPO	DESCRIZIONE
AC	solo per correnti di guasto sinusoidali
A	anche per correnti di guasto pulsanti
B	anche per correnti di guasto continue

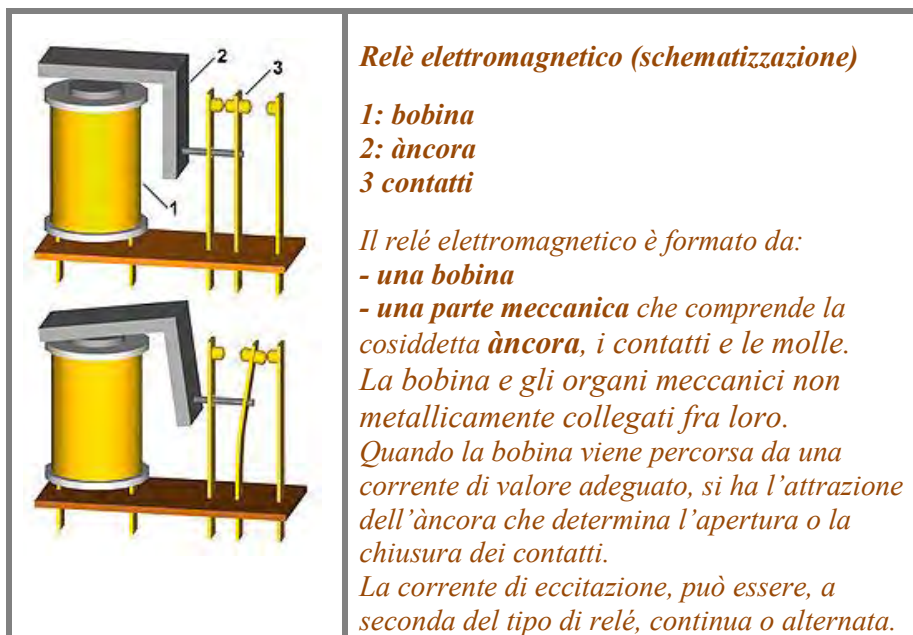
RELE'

Sono organi che permettono di **azionare in modo automatico** gli **interruttori** per effetto di una **grandezza di comando**.

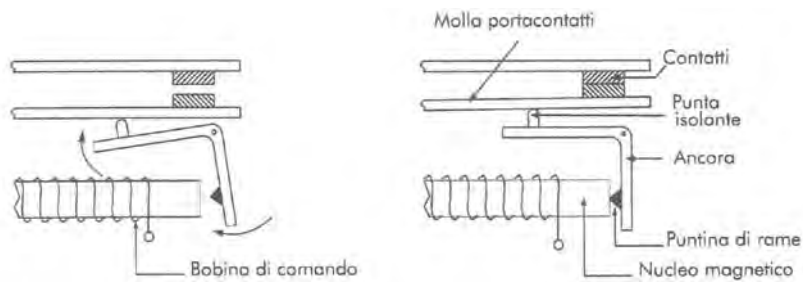
La grandezza di comando, normalmente, è la **corrente** che percorre una bobina di eccitazione oppure la **tensione** che viene applicata ai capi della bobina (si parla in questo caso di **relé elettromagnetici** o **elettromeccanici**).

Esistono però anche **relé termici**, per i quali la **grandezza di comando** è la **dilatazione termica**, cioè la dilatazione per effetto del calore, di un elemento percorso da corrente.

Il complesso di un relé elettromagnetico e di un relé termico prende il nome di relé magnetotermico.



Relè elettromagnetico



A sinistra: la bobina non è percorsa da corrente, l'ancora è in posizione di riposo e i contatti sono aperti

A destra: la bobina è percorsa da corrente, l'ancora spinge le molle che chiudono i contatti

CLASSIFICAZIONE DEI RELE'

(in relazione al verso della grandezza di eccitazione della bobina)

- relè **NEUTRO**: il **passaggio** dei contatti da una **posizione all'altra** è **indipendente** dal verso della grandezza di eccitazione della bobina
- relè **POLARIZZATO**: il **passaggio** dei contatti da una posizione all'altra **dipende** dal verso della grandezza di eccitazione della bobina.

CLASSIFICAZIONE DEI RELE'

(in relazione al diverso comportamento in presenza della corrente di eccitazione della bobina)

- relè **PASSO-PASSO o CICLICI**
- relè **MONOSTABILI**
- relè **BISTABILI**
- relè **A TEMPO o RITARDATO** (temporizzatore).

RELE' PASSO-PASSO o CICLICO

Questo relè commuta ad ogni impulso che arriva dalla bobina e i contatti rimangono nella posizione assunta all'arrivo dell'ultimo impulso fino a all'applicazione di un successivo impulso.

Se, per esempio, il contatto è chiuso, con l'invio di un impulso alla bobina, il contatto si apre e rimane aperto anche una volta che l'impulso sia terminato. All'arrivo di un ulteriore impulso, il contatto si richiude.

La struttura del relé passo-passo prevede la presenza di un albero a camme che, ad ogni eccitazione della bobina, apre i contatti se sono chiusi e li chiude se sono aperti.

Esistono più tipologie di relé passo-passo:

- relé passo-passo a interruttore unipolare
- relé passo-passo a interruttore bipolare
- relé passo-passo a deviatore
- relé passo-passo a commutatore con sequenza a quattro passi (un passo per ogni impulso)
- relé passo-passo a deviatore con sequenza a quattro passi (un passo per ogni impulso).

RELE' MONOSTABILE

E' caratterizzato dal fatto che (contrariamente a quanto avviene per il passo-passo) i contatti rimangono nella posizione commutata, cioè nella posizione determinata dall'impulso, solo durante l'applicazione dell'impulso, cioè solo durante l'eccitazione della bobina.

Appena l'impulso cessa, e termina lo stato di eccitazione della bobina, i contatti ritornano nella posizione di riposo, cioè nella posizione che avevano prima dell'applicazione dell'impulso.

Il funzionamento di questo relé è caratterizzato quindi:

- da una posizione instabile, o posizione "commutata" (contatti chiusi), che viene mantenuta soltanto nel corso della durata dell'impulso di eccitazione
- da una posizione di riposo (contatti aperti), che è la posizione stabile sulla quale i contatti si portano al termine degli impulsi di eccitazione e sulla quale permangono in assenza di impulsi.

RELE' BISTABILE

Il relé monostabile possiede, come indica il nome, un solo stato stabile, cioè un solo stato in grado di perdurare anche dopo la fine dell'impulso che lo ha determinato, stato stabile che è quello di riposo (contatto aperto), mentre lo stato commutato (contatti chiusi) è instabile in quanto cessa appena l'impulso che lo ha determinato si esaurisce. Nel relé bistabile invece entrambi gli stati (contatto aperto e contatto chiuso) sono stabili e quindi perdurano anche una volta esauritosi l'impulso che li ha determinati. Nel relé bistabile vi sono due bobine, ciascuna delle quali, eccitandosi, può determinare uno solo dei due stati e non l'altro:

- lo "stato1" nasce solo quando viene applicato un impulso di eccitazione alla bobina B1, si mantiene inalterato anche dopo che l'impulso di B1 si è esaurito e ha termine solo quando un impulso applicato all'altra bobina (B2) porta il relé nello stato 2; se, mentre il relé è nello stato 2 , viene nuovamente applicato un impulso a B1, ritorna lo stato 1.
- lo "stato 2" ha inizio solo quando viene applicato un impulso di eccitazione alla bobina B2, persiste anche dopo che l'impulso di B2 si è estinto e ha termine solo quando un impulso applicato alla prima bobina (B1) porta il relé nello stato 1; se, mentre il relé è in stato 1, viene nuovamente applicato un impulso a B2, ritorna lo stato 2.

RELE' A TEMPO o RITARDATO (temporizzatore)

E' caratterizzato dal fatto che i tempi di attrazione o di rilascio dell'âncora sono resi più lunghi di quanto necessario al normale funzionamento.

L'UTENZA DOMESTICA: FASE E NEUTRO

Il trasporto dell'energia elettrica con il sistema trifase a stella può essere effettuato in due modi:

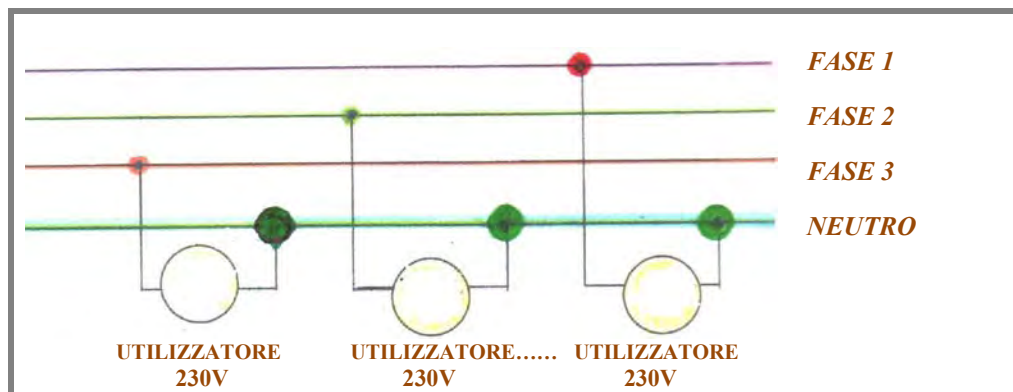
- o usando solo i tre terminali di fase (L1, L2, L3)
- o a quattro fili usando sia i tre terminali di fase (L1, L2, L3) sia il neutro, collegato al centro-stella

Nella distribuzione di energia elettrica a scopo misto di illuminazione, forza motrice e riscaldamento (e quindi fino alle cabine MT/BT) è utilizzata la configurazione trifase con neutro (cioè a quattro fili).

La distribuzione dell'energia elettrica per usi industriali utilizza il sistema trifase perché le grosse macchine elettriche usano questo tipo di sistema.

La distribuzione dell'energia elettrica per **usi civili** usa il sistema **MONOFASE** che viene ottenuto dal “sistema trifase con neutro” collegando il circuito utilizzatore **fra** una **fase** (una qualunque delle tre⁷) e il **neutro**. Tra la fase e il neutro c'è una tensione alternata, detta tensione “stellata”, o tensione “di fase” con valore efficace di **230V**.

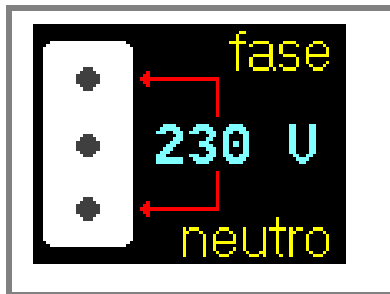
Nelle nostre case la linea di distribuzione arriva quindi con soli **due fili**.



⁷ Le forniture agli utenti monofase sono distribuite tra le tre fasi in modo da equilibrare statisticamente gli assorbimenti ed ottimizzare il trasporto. Le correnti di ritorno dai neutri delle abitazioni si compensano mediamente in modo da fare tendere a zero la corrente di neutro verso il trasformatore presente in cabina.

INDIVIDUAZIONE DEL FILO NEUTRO

- il neutro è un **conduttore** di colore **azzurro**.
- Per individuare il neutro si può usare un cercafase. Il **neutro** è **il solo** conduttore che, testato, **non** fa **illuminare** il cercafase (il filo di fase invece lo fa illuminare).



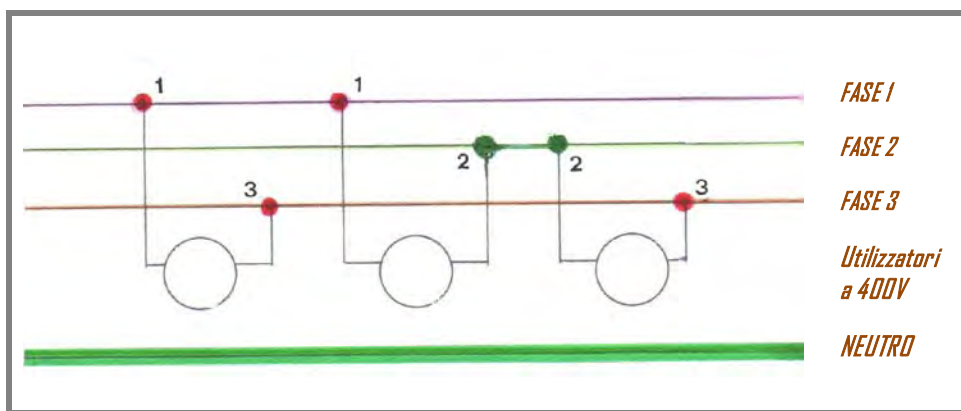
Se accadesse che il **cercafase s'illumina** su **entrambi i cavi**, significa che **non c'è neutro** e quindi in nessun caso verrà usato un conduttore di colore azzurro. Quando si verifica questa situazione? In alcune zone di Italia si utilizza un tipo di distribuzione trifase con tensioni concatenate, cioè con ddp tra fase e fase⁸, di 230 volt (anziché 380V), per cui la **tensione fase-neutro** risulta di **133 Volt** (anziché 230V). In questi casi, per fornire comunque 230 volt all'utenza, nelle abitazioni vengono portati due conduttori di fase (tensione concatenata) invece di una fase e neutro.

⁸ e non fra fase e neutro

DISTRIBUZIONE ELETTRICA PER USO INDUSTRIALE

La distribuzione dell'energia elettrica per usi industriali utilizza il sistema trifase perché le grosse macchine elettriche usano questo tipo di sistema.

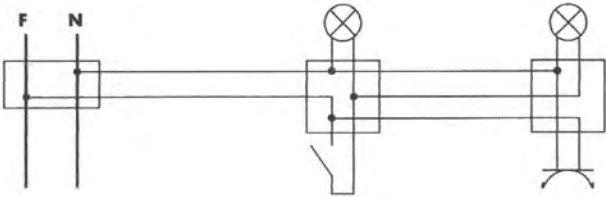
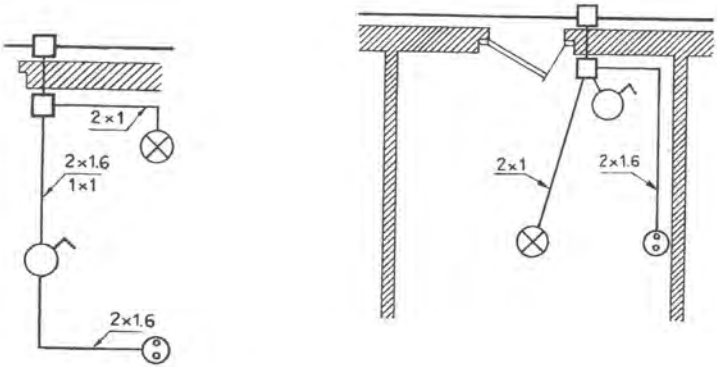
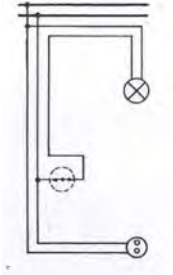
Gli utilizzatori sono collegati **fra due fasi** e ai capi di ciascuno di essi è presente una tensione **concatenata** di **400V**



COME SI DESCRIVE UN IMPIANTO ELETTRICO

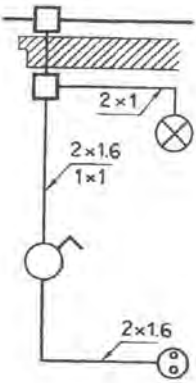
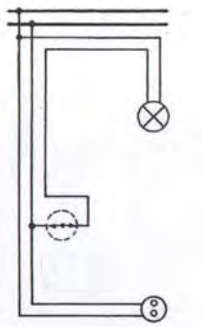
CLASSIFICAZIONE DEI POSSIBILI SCHEMI IN BASE AL MODO DI RAPPRESENTAZIONE

- **Schema ORDINARIO:** le apparecchiature **non sono disposte in base all'effettiva disposizione topografica**, ma **nel modo che rende più semplice la rappresentazione** delle connessioni.
- **Schema DI MONTAGGIO:** mostra le diverse **connessioni** fra gli elementi di un impianto rispettando le **posizioni reciproche**, indicando la **distribuzione e il tipo dei conduttori** e mettendo in rilievo i **terminali** e le **morsettiere**.
- **Schema TOPOGRAFICO:** i vari elementi dell'impianto o della parte di impianto sono rappresentati rispettando la loro **posizione reciproca**. I **conduttori** sono raggruppati in **cavetti**. Per ogni cavetto sono indicati il **numero** e la **sezione dei conduttori**; serve all'installatore per **conoscere quali apparecchiature** occorrono e la **sezione dei fili**.
- **Schema FUNZIONALE o di PRINCIPIO:** descrive **concettualmente** l'impianto, raffigurando – se possibile - i vari elementi nell'ordine in cui intervengono nella normale sequenza di funzionamento.

Tipo di schema	ESEMPI
Schema DI MONTAGGIO	
Schemi TOPOGRAFICI	
Schema FUNZIONALE o di PRINCIPIO	

CLASSIFICAZIONE DEGLI SCHEMI IN BASE ALLA RAPPRESENTAZIONE DEI CONDUTTORI

- Schema UNIfilare
- Schema MULTIfilare

Tipo di schema	Utilizzazione
	Schema UNIFILARE Tutti i conduttori di un sistema a due o più linee sono rappresentati con una sola linea. Serve all'installatore per conoscere quali apparecchiature occorrono e la sezione dei fili
	Schema BIFILARE o MULTIFILARE Ciascun conduttore è indicato con una sua specifica linea

IL PROGETTO DI UN IMPIANTO ELETTRICO

PREMESSE

1. FATTORE DI CONTEMPORANEITA' K_c

Non tutti gli utilizzatori (lavatrice, ferro da stiro ecc.) dell'impianto elettrico sono simultaneamente attivi, perciò è definito un "coefficiente di contemporaneità":

$$K_c = \text{fattore di contemporaneità} = \frac{\text{effettiva potenza totale assorbita (P}_t\text{)}}{\text{somma delle potenze dei singoli utilizzatori (potenza massima)}}$$

$$K_c = \frac{P_t}{\sum_{i=1}^n P_i}$$

E' sempre: $K_c \leq 1$

Nel caso limite in cui tutte le utenze funzionano contemporaneamente: $K_c = 1$

2. FATTORE DI UTILIZZAZIONE K_U

Alcuni utilizzatori possono non essere utilizzati alla loro piena potenza (cioè alla potenza nominale P_N), ma alla “potenza massima richiesta” che è minore, o, al limite, uguale alla potenza nominale.

Perciò viene definito un “fattore di utilizzazione”:

$$K_U = \text{fattore di utilizzazione} = \frac{\text{effettiva potenza massima richiesta } (P_M < P_N)}{\text{potenza nominale } (P_N) \text{ installata}}$$

$$K_U = \frac{P_M}{P_N}$$

E' sempre: $K_U \leq 1$

IL DIMENSIONAMENTO

- Calcolo della **potenza installata presunta**
- Calcolo della **potenza convenzionale o potenza impegnata**
- Determinazione della **potenza contrattuale**
- Scelta del **numero dei *circuiti* di distribuzione**
- Dimensionamento del **centralino di distribuzione**
- Dimensionamento della **portata dei cavi**
- Determinazione delle **dotazioni impiantistiche** previste
- **Disegno degli schemi elettrici** e di piani di installazione
- **Verifica** del corretto dimensionamento

CALCOLO DELLA POTENZA INSTALLATA PRESUNTA

La potenza installata presunta può essere calcolata:

- o come **somma** delle **potenze nominali** assorbite dalle **utenze** previste per l'appartamento
- o in base a **tabelle** (come quella seguente) che forniscono la potenza attiva non in relazione al funzionamento effettivo, ma in base:
 1. alla **superficie** dell'appartamento (ai fini della potenza destinata all'illuminazione)
 2. al **numero dei locali** (ai fini della potenza destinata, per esempio, agli scaldacqua)
 3. alla **dislocazione dell'appartamento in zona urbana o rurale** (ai fini della potenza richiesta da servizi vari)

Potenza INSTALLATA PRESUNTA

Per l'illuminazione	10 W per m ² di superficie dell'appartamento col minimo di 500 W
Scaldacqua	1000 W per appartamenti fino a quattro locali 2000 W per appartamenti con oltre quattro locali
Cucina	da considerare solo se ne è prevista esplicitamente l'installazione
Servizi vari ⁽¹⁾	40 W per m ² di superficie dell'appartamento in zone urbane 20 W per m ² di superficie dell'appartamento in zone rurali
⁽¹⁾ La diffusione di nuovi tipi di utilizzatori tende a elevare e unificare questi livelli verso 50+60 W/m ² per appartamenti in zone urbane e oltre 20 per appartamenti in zone rurali.	

CALCOLO DELLA POTENZA CONVENZIONALE O POTENZA IMPEGNATA

La potenza **convenzionale** o potenza **impegnata** di un circuito è la **potenza attiva** per la quale deve essere dimensionato il circuito *non in base al suo effettivo funzionamento*, ma in base :

- ai **dati di targa (potenza NOMINALE) degli utilizzatori** fissi alimentati
- alla **potenza massima dei gruppi di prese** a spina collegati al circuito
- a opportuni **coefficienti** mediante i quali si tiene conto:
 - della potenza mediamente richiesta dai singoli carichi (fattore di utilizzazione k_U)
 - della **contemporaneità** di funzionamento ipotizzabile nel caso dell'alimentazione di più carichi (fattore di contemporaneità k_C)

Sussiste la relazione:

$$P_{\text{tot-conv}} = K_c \cdot \sum_{i=1}^n (P_i)$$

Potenza totale convenzionale = (fattore di contemporaneità) · (somma delle potenze dei singoli utilizzatori)

In sostanza la potenza convenzionale o potenza impegnata si ottiene dalla **potenza installata** presunta **moltiplicandola** per il **fattore** di contemporaneità **K_c** (si vedano le due tabelle seguenti).

Essendo K_C minore o uguale ad 1, la potenza convenzionale, o potenza impegnata, è **minore** della potenza installata presunta.

FATTORE K_c DI CONTEMPORANEITA'

	Servizio	Apparecchi o circuiti	Abitazione	Uffici	Negozi	Alberghi
Campo di applicazione:	Illuminazione	Punti luce	0,65	0,90	0,90	0,75
edifici civili (Norme di riferimento CEI 11-11)	Servizi vari	Prese a spina	0,25	0,50	0,50	0,50
	Scaldacqua	Bollitore più potente	1	1	1	1
		2° bollitore	0,75	0,75	0,75	0,75
		3° bollitore	0,50	0,50	0,50	0,50
		Altri	0,50	0,25	0,25	0,25
	Cucina	Apparecchio più potente	1	–	–	1
		Altri	1	–	–	0,75
	Ascensori	Motore più potente	–	3	3	3
		2° ascensore	–	1	1	1
		Altri	–	0,7	0,7	0,7
	Colonne montanti	1 unità d'impianto	1	–	–	–
		2-4 unità d'impianto	0,8	–	–	–
		5-10 unità d'impianto	0,5	–	–	–
		Più di 10 unità d'impianto	–	–	–	–

FATTORE K_c DI CONTEMPORANEITA'

Campo di applicazione	Servizio	Apparecchi o circuiti	Destinazione edifici			
			abitazioni	uffici	laboratori	negozi
Edifici civili	Illuminazione Servizi vari, scaldacqua	punti luce	0,65	0,90	—	0,90
		prese a spina	0,25	0,50	—	0,50
		bollitore più potente	1	1	—	1
		secondo bollitore	0,75	0,75	—	0,75
		terzo bollitore	0,50	0,50	—	0,50
		altri	0,50	0,25	—	0,25
	Cucina	apparecchio più potente	1	—	—	—
		altri	1	—	—	—
	Ascensori	motore più potente	—	3	—	3
		secondo ascensore	—	1	—	1
		altri	—	0,70	—	0,70
	Colonne montanti	1 unità d'impianto	1	—	—	—
		da 2 a 4 unità d'impianto	0,8	—	—	—
		da 5 a 10 unità d'impianto	0,5	—	—	—
		più di 10 unità d'impianto	0,4	—	—	—
Edifici prefabbricati e costruzioni modulari	illuminazione	punti luce	0,75	0,80	0,80	0,90
		prese a spina	0,10	0,10	0,10	0,30
	Piccola forza motrice ed usi domestici	utilizzatori fissi	0,70	0,80	0,80	0,90
		prese a spina 10A	0,20	0,10	0,10	0,10
		prese a spina > 10A	0,15	0,05	0,05	0,05
	Condizionamento d'aria centralizzata	circuito di potenza	1	1	1	1
	Servizi termici centralizzati	centrale termica, centrale idrica, ecc.	0,80	0,60	0,60	0,60
	Servizi logistici	cucina, lavanderia, stireria, ecc.	—	0,70	0,70	0,70
	Elevatori (ascensori, montacarichi, ecc.)	motore più potente	3	3	3	3
		secondo motore	1	1	1	1
		altri	0,70	0,70	0,70	0,70

CALCOLO DELLA POTENZA CONTRATTUALE

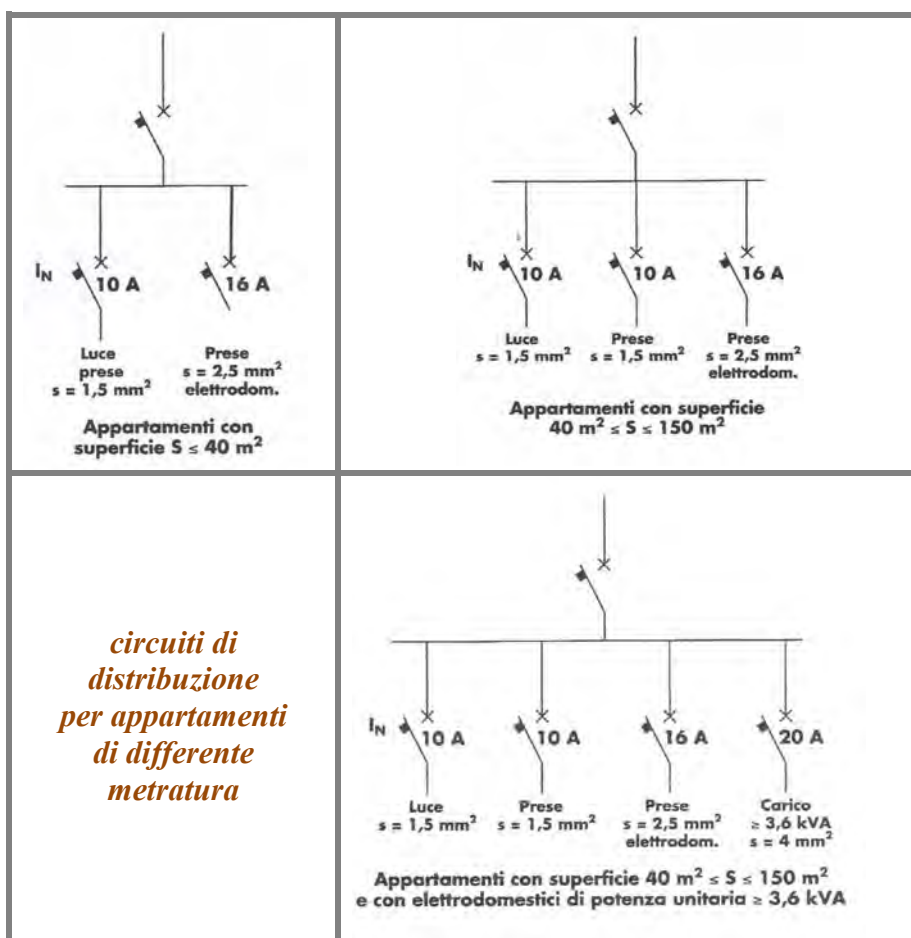
La potenza contrattuale relativa a un appartamento è la potenza massima per la quale è stipulato il contratto di fornitura dell'energia elettrica dell'appartamento e che non deve essere superato nel funzionamento dell'impianto.

Se, in un certo istante, il consumo di potenza supera la potenza contrattuale, interviene un limitatore.

Il valore di potenza contrattuale deve essere scelto, fra quelli proposti dalla società fornitrice di energia elettrica, maggiore del valore della potenza convenzionale.

SCELTA DEL NUMERO DEI CIRCUITI DI DISTRIBUZIONE (nei quali l'impianto potrà essere suddiviso)

Il **circuito di distribuzione** è quella **parte** dell'impianto elettrico utilizzatore che **ha un'unica alimentazione** ed è protetta da **uno specifico dispositivo di protezione** contro le sovracorrenti (per es: interruttore automatico **magnetotermico**).

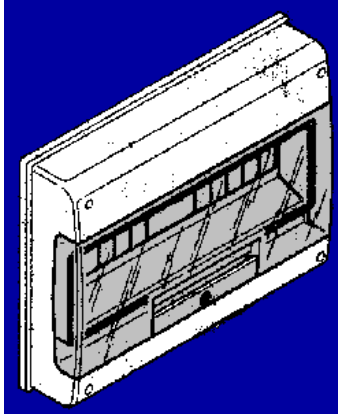


Riguardo al numero dei circuiti di distribuzione, la Guida CEI 64-50 consiglia di orientarsi in funzione della superficie dell'appartamento, come è mostrato nella tabella seguente

SCELTA DEL NUMERO DEI CIRCUITI DI DISTRIBUZIONE

SUPERFICIE dell'appartamento	NUMERO TOTALE MINIMO di CIRCUITI DI DISTRIBUZIONE	DESTINAZIONE dei circuiti	SEZIONE dei conduttori in mm ²
Fino a 40 m ²	2	<u>Un circuito</u> per punti-luce e prese di corrente da 10A (piccoli elettrodomestici e lampade da lavoro)	1,5
		<u>Un circuito</u> per prese di corrente da 16A (elettrodomestici fissi)	2,5
Fra 40 m ² e 150 m ²	3	<u>Due circuiti</u> per punti-luce e prese di corrente da 10A (piccoli elettrodomestici e lampade da lavoro)	1,5
		Un <u>circuito</u> per prese di corrente da 16A (elettrodomestici fissi)	2,5

DIMENSIONAMENTO DEL CENTRALINO DI DISTRIBUZIONE (o QUADRO ELETTRICO)



In ogni impianto elettrico, a valle del *contatore*, viene installato un quadro di distribuzione. I più piccoli sono in materiale plastico autoestinguente a isolamento e possono essere sia incassati che a muro. Gli altri sono di tipo metallico.

All'interno vengono alloggiati gli apparecchi, che hanno le seguenti funzioni:

- **protezione dei circuiti**
- **protezione delle persone**, in grado di interrompere, entro 5 secondi, le correnti di guasto verso terra
- **sezionamento**, ovvero interruzione dell'alimentazione dei circuiti, ad esempio per compiere lavori sull'impianto elettrico in tutta sicurezza.
- **funzioni relative a impianti “particolari”**, e cioè agli impianti **telecomandati**, o a **inserimento programmato o temporizzato** o a **bassissima tensione** (suonerie); in quest'ipotesi il centralino di distribuzione diventa un centro di protezione, comando e controllo automatico

Il **numero di interruttori** installati deriva principalmente da **considerazioni di tipo funzionale**. Ad esempio, in un normale appartamento, è possibile proteggere l'impianto elettrico con un solo interruttore - , ma in caso di guasto o di lavoro su una sola parte dell'impianto, verrà a mancare l'alimentazione a tutto l'appartamento.

DIMENSIONAMENTO DEI CAVI

PREMESSE

MASSIMA TEMPERATURA (t_z) SOPPORTABILE DALL'ISOLANTE DEL CAVO

Dipende ovviamente dal materiale di realizzazione dell'isolamento. L'isolamento in **gomma "G5"** (gomma etilen-propilenica) sopporta una temperatura massima di **90°**, **maggiore** di quella (**70°**) sopportata dall'isolamento in **PVC** (PoliVinCloruro).

PORTATA DEL CAVO (I_z)

E' il **valore di corrente** che **porta** il cavo alla **massima temperatura sopportabile dall'isolante**. Il valore di portata **non deve mai essere superato** dal valore della corrente trasportata dal cavo. La portata:

- **aumenta** all'aumentare della **sezione (mm^2)** del cavo
- **aumenta** all'aumentare della **massima temperatura sopportata dall'isolante**. Nelle condizioni più sfavorevoli (isolante in PVC, sei cavi unipolari in tubo protettivo, temperatura ambiente di 40°C) si ha:

$$I_z = 8 \cdot S^{0,625}$$

dove:

S= **sezione** in mm^2 del cavo, **I_z** = **portata in A**

8= **portata "caratteristica"**, cioè portata in A di un cavo di **sezione unitaria (1mm^2)**

- diminuisce all'aumentare del numero dei cavi posti nella stessa canaletta o nello stesso tubo protettivo
- **è maggiore** se il cavo è posato in **canaletta** invece che in tubo, in quanto, nella canaletta, i cavi risultano più distanziati o comunque meno pressati uno contro l'altro; in ogni caso si può dire che la portata dipende dal "tipo di posa" dei cavi
- diminuisce all'aumentare della temperatura-ambiente
- diminuisce all'aumentare della "polarità", cioè del numero di conduttori che formano il cavo

CORRENTE DI IMPIEGO (I_B) DEL CAVO

E' la corrente che **il cavo trasporta in condizioni normali**, cioè in assenza di sovraccarichi e di cortocircuiti.

La **I_B** deve sempre essere **minore della portata I_Z**

MARGINE DI SICUREZZA ($I_Z - I_B$) DELLA CORRENTE

Si fa in modo che sia rispettata la condizione **$I_Z - I_B \neq 0$** per evitare che la temperatura del cavo possa superare la temperatura massima t_Z .

TABELLA PER LA DETERMINAZIONE DELLA *PORTATA I_Z* DEI CAVI
In base a polarità, materiale isolante, tipo di posa, sezione (tab. UNEL 35024-70)

A Tipi di posa	B Tipi di cavi	C Isolante	D Numero dei conduttori									
cavi unipolari o multipolari — entro tubi — sotto modanatura	unipolari* senza guaina**	PVC R o RF; gomma G		4	3	2						
	multipolari ed unipolari con guaina	gomma G2				4	3	2				
		PVC R o RF; gomma G		4	3	2						
cavi multipolari distanziati - fissati alle pareti	multipolari	gomma G2				4	3	2				
		PVC R o RF; gomma G			4	3	2					
		gomme G2 o G5 Polietilene reticolato					4	3	2			
cavi unipolari non distanziati — fissati alle pareti - su passerelle	unipolari senza guaina	PVC R o RF; gomma G				4	3	2				
		gomme G2 o G5 Polietilene reticolato						4	3	2		
		PVC R o RF; gomma G								n		
cavi unipolari distanziati - su passerelle o su supporti analoghi	unipolari senza guaina	gomme G2 o G5 Polietilene reticolato									n	
		PVC R o RF; gomma G										
		gomme G2 o G5 Polietilene reticolato										n
cavi unipolari non distanziati - in canale su più strati secondo IEC 364-5-523 ***	unipolari senza guaina	PVC R o RF; gomma G	4	3	2							
		gomme G2 o G5 Polietilene reticolato				4	3	2				
		F Sezione nominale dei conduttori mm ²	E Portata in regime permanente in Amper									
		1	9,6	10,5	12	13,5	15	17	19	21	23	25
		1,5	12,4	14	15,5	17,5	19,5	22	24	27	29	32
		2,5	17	19	21	24	26	30	33	37	40	44
		4	23	25	28	32	35	40	45	50	55	60
		6	29,4	32	36	41	46	52	58	64	70	77
		10	40	44	50	57	63	71	80	88	97	107
		16	54	59	68	76	85	96	107	119	130	143
		25	72	75	89	101	112	127	142	157	172	190
		35	88	97	111	125	138	157	175	194	213	235
		50	110	115	134	151	168	190	212	235	257	284
		70	137	147	171	192	213	242	270	299	327	362
		95	165	178	207	232	258	293	327	362	396	440
		120	191	209	239	269	299	339	379	419	458	507
		150	219	242	275	309	344	390	435	481	527	584
		185	250	276	314	353	392	444	496	549	602	666
		240	295	322	369	415	461	522	584	645	707	781

* Tra i cavi unipolari si distinguono:
— i cavi unipolari senza guaina, che hanno praticamente soltanto il rivestimento isolante e che generalmente non possono essere posati che entro tubi o sotto modanature;
— i cavi unipolari con guaina, che sono muniti di un rivestimento esterno e che possono essere posati come i cavi multipolari, potendo p.e. essere fissati direttamente alle pareti o sospesi a fune portante.
Una guaina è un rivestimento continuo che assicura la protezione dei cavi contro l'azione dell'ambiente.
* Per cavi unipolari isolati con gomma di qualità G, G2, G5 le Norme CEI prescrivono sempre una guaina od almeno una treccia tessile.
**n significa numero qualsiasi di cavi unipolari.
***Per più di 4 cavi adottare i seguenti coefficienti da applicare alla portata relativa a 2 cavi:

n° cavi unipo. raggruppati in canale su più strati	6	8	10	12	14	15÷20
K	0,7	0,65	0,6	0,55	0,55	0,5

DIMENSIONAMENTO DEI CAVI

DIMENSIONAMENTO DI UNA LINEA (di cavo)

Il progettista ricava prima la corrente di impiego I_B e, soltanto dopo, la sezione del cavo. La I_B viene ricavata in funzione delle potenze dei carichi e del fattore K_C di contemporaneità.

I **passi** del dimensionamento della linea sono:

1. determinare la corrente assorbita da ciascun utilizzatore
2. stabilire il valore del coefficiente K_C di contemporaneità dei carichi
3. ricavare il valore della corrente di impiego I_B in base ai dati del punto 1 (correnti assorbite) e del punto 2 (coefficiente di contemporaneità)
4. Cercare sulla tabella delle portate dei cavi unipolari un valore di portata I_Z maggiore di I_B
5. individuare, sulla stessa tabella, la sezione corrispondente alla I_Z scelta.

TABELLA DELLE PORTATE DEI CAVI UNIPOLARI

S [mm ²]	I_Z [A] PVC	I_Z [A] G2, G5
1	10,5	13,5
1,5	14	17,5
2,5	19	24
4	25	32
6	32	41
10	44	57
16	59	76
25	75	101
35	97	125

CAPITOLO 10

MOTORI IN ALTERNATA ASINCRONI E SINCRONI

COMPLEMENTI

PREMESSA: I CONVERTITORI DI FREQUENZA

***TECNICHE DI CONTROLLO DELLA VELOCITA’
DEL MOTORE ASINCRONO***

***PREGI E INCONVENIENTI
DEL MOTORE ASINCRONO***

CONFRONTO CON IL MOTORE SINCRONO

***CONTROLLO DI VELOCITA’
DEL MOTORE SINCRONO***

***CONTROLLO DI VELOCITA’
DEL MOTORE SINCRONO
PER VARIAZIONE DI FREQUENZA***

PROPULSIONE ELETTRICA NAVALE

IL CONTROLLO DI VELOCITA’ DEI MOTORI IN ALTERNATA E I TIRISTORI

MOTORE SINCRONO O ASINCRONO?

PREMESSA

I CONVERTITORI DI FREQUENZA

Il controllo di velocità dei motori elettrici in alternata è effettuato mediante convertitori statici di frequenza (detti anche convertitori ac/ac).

Il convertitore statico di frequenza converte tensioni e correnti ad ampiezza e frequenza costante in tensioni e correnti ad ampiezza e frequenza variabili.

Ricordiamo infatti che la velocità di un motore elettrico in alternata dipende essenzialmente dalla frequenza della tensione e della corrente con le quali è alimentato.

Il numero di giri al minuto cresce al crescere della frequenza.

Se poi si vuole che la variazione di velocità avvenga a coppia massima costante, bisogna variare non solo la frequenza (della tensione di alimentazione), ma anche l'ampiezza, in misura proporzionale alla frequenza stessa.

I convertitori statici di frequenza sono detti:

- di tipo “indiretto” se, nella conversione alternata-alternata, effettuano una conversione intermedia in continua
- “cicloconvertitori” se, nella conversione alternata-alternata, non effettuano la conversione intermedia in continua

TECNICHE DI CONTROLLO DELLA VELOCITA' DEL MOTORE ASINCRONO

REGOLAZIONE DI VELOCITA' DI UN MOTORE ASINCRONO MEDIANTE VARIAZIONE DI FREQUENZA A COPPIA MASSIMA COSTANTE

Se si vuole variare la velocità del motore mantenendo **costante la coppia motrice**, **bisogna far variare la tensione di alimentazione proporzionalmente alla frequenza** di alimentazione¹, ossia bisogna mantenere **costante il flusso magnetico induttore**².

La regolazione a coppia costante comporta una mera traslazione della caratteristica elettromeccanica e viene effettuata per basse frequenze (generalmente non oltre i 50 Hz)³.

¹ infatti, riguardo alla coppia massima, generalmente chiamata C_M oppure T_M , si ha

$$C_m = T_m = \frac{k \cdot V^2}{\omega X}$$

con V = tensione di alimentazione

ω ed f = pulsazione e frequenza di sincronismo, cioè del campo rotante

X = reattanza associata al flusso disperso di rotore

Essendo sia ω che X proporzionali alla frequenza si ha:

$$C_m = T_m = \frac{k \cdot V^2}{k_1 \cdot f \cdot k_2 \cdot f} = \frac{k}{k_1 \cdot k_2} \cdot \frac{V^2}{f^2}$$

se poi facciamo variare la tensione di alimentazione in maniera proporzionale alla frequenza:

$$C_m = T_m = \frac{k}{k_1 \cdot k_2} \cdot \frac{(k_3 \cdot f)^2}{f^2} = \frac{k \cdot k_3}{k_1 \cdot k_2} = \text{cost}$$

² “mantenere costante il flusso magnetico equivale a far variare la tensione V di alimentazione nella stessa misura della frequenza” per il seguente motivo:

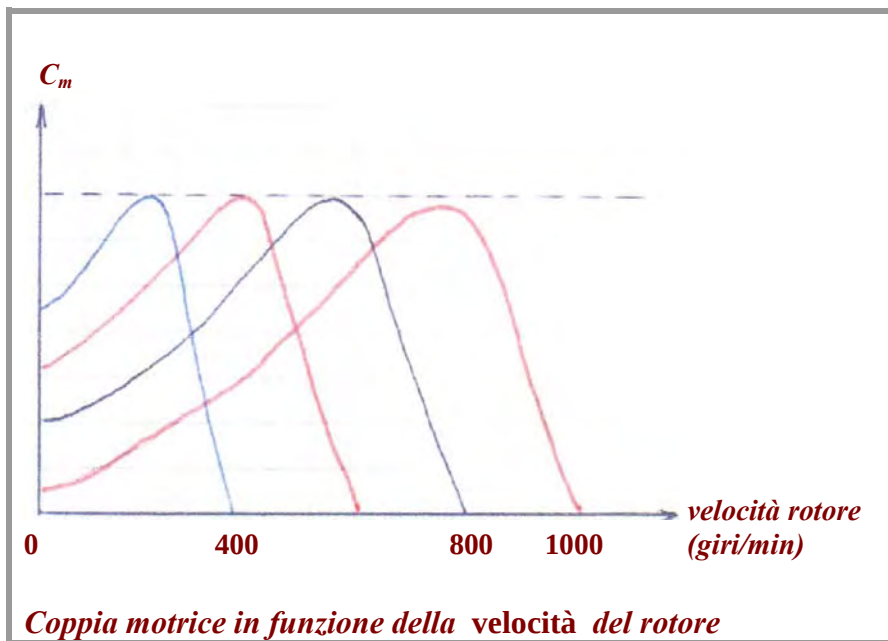
$$v = -\frac{d\phi}{dt} \quad (\text{legge dell'induzione, con } v = \text{f.e.m.})$$

se il flusso è sinusoidale: $V = -j\omega\phi$ e in modulo: $V = 2\pi f\phi$

questa relazione ci dice che, se facciamo variare la tensione proporzionalmente a f , cioè se facciamo in modo che si abbia $V = kf$,

si ha $K \cdot f = 2\pi \cdot f \cdot \phi$ da cui $K = 2\pi \cdot \phi$ Allora il flusso rimane costante: $\phi = \frac{k}{2\pi} = \text{costante}$

³ L'aumento della velocità a coppia costante comporta un aumento della potenza, infatti, man mano che spostiamo la caratteristica velocità-coppia verso destra, aumentando la velocità, agli stessi valori di coppia max vengono a corrispondere velocità maggiori, per cui:
potenza meccanica = coppia * velocità = (quantità costante) * (quantità crescente) =
= (quantità crescente)



Il motivo è che, a frequenze più elevate, il mantenimento del flusso a un valore costante farebbe aumentare notevolmente le perdite nel ferro.

Volendo mantenere costante la potenza (e non la coppia motrice) deve mantenersi costante anche la tensione.

REGOLAZIONE DI VELOCITA' MEDIANTE VARIAZIONE DI FREQUENZA A POTENZA COSTANTE

Se, a tensione costante, facciamo crescere la frequenza, la velocità cresce, ma diminuisce il flusso⁴ e di conseguenza diminuisce anche la coppia motrice (in quanto a parità di corrente rotorica I_2 , la coppia è proporzionale al flusso:

$$C_m = k \Phi I_2).$$

Siccome la frequenza aumenta e la coppia diminuisce nella stessa misura, allora (essendo $P_T = 2\pi f c$) questa regolazione avviene a potenza costante.⁵

La regolazione in esame determina, nella caratteristica elettromeccanica, un abbassamento delle curve al crescere della velocità⁶.

$$^4 \quad v = -\frac{d\phi}{dt} \quad (\text{legge dell'induzione con } v = \text{f.e.m}),$$

se il flusso è sinusoidale: $V = -j\omega\phi$ e in modulo: $V = 2\pi f\phi$

E' evidente che se, con $V=\text{cost}$, facciamo crescere f , il flusso deve diminuire

⁵ Osserviamo che nella $P_T = 2\pi f c$ è:

P_T = potenza trasmessa al rotore

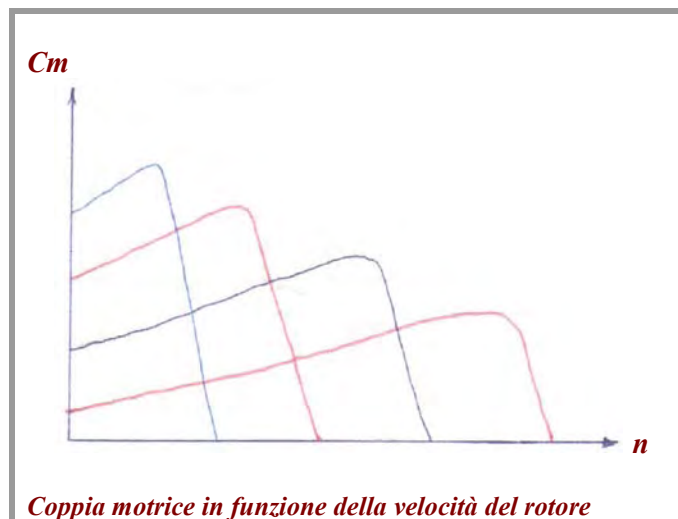
c = coppia trasmessa al rotore

⁶ infatti, man mano che spostiamo la caratteristica velocità-coppia verso destra, aumentiamo la velocità, il che comporta, se la potenza rimane costante, una diminuzione della coppia, infatti:

potenza meccanica = coppia*velocità;

coppia = potenza / velocità = (quantità costante) / (quantità crescente) =
= (quantità decrescente).

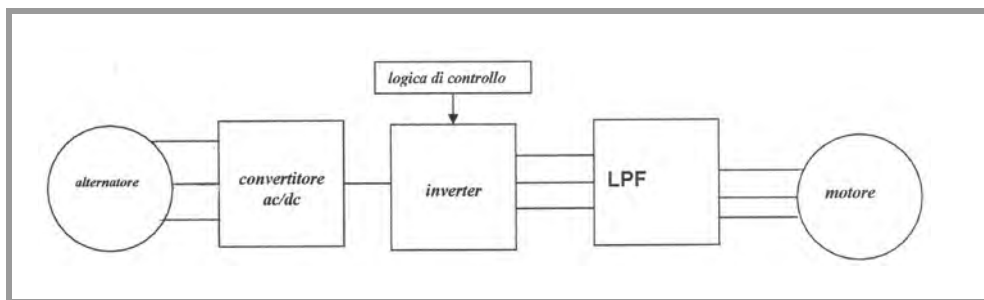
La diminuzione della coppia comporta, graficamente, l'abbassamento della curva.



CONTROLLO CONTINUO DI VELOCITA' DEL MOTORE ASINCRONO TRIFASE MEDIANTE CONVERTITORE AC/DC E INVERTER

Questa tecnica prevede che il motore non sia alimentato direttamente dalle tensioni trifase della rete di distribuzione, ma da un sistema di tre tensioni a frequenza e⁷ ampiezza variabili, prodotte da un inverter pilotato dal sistema di controllo.

Ciò comporta che la tensione trifase di rete prima sia resa continua da un convertitore ca/cc e poi nuovamente trasformata in alternata (ma con frequenza e⁸ ampiezza regolabili) da un inverter.



Siccome è:

$$n = \frac{f_s \cdot 60}{p} \cdot (1 - s)$$

si può far crescere la velocità facendo crescere la frequenza f delle tensioni di alimentazione.

E' opportuno però che la regolazione avvenga senza che il flusso diminuisca rispetto al suo valore nominale (per sfruttare meglio la macchina).

Lavorare a flusso costante però significa lavorare a coppia max costante, dal momento che la coppia max dipende dal flusso (infatti la coppia massima è proporzionale al flusso⁹).

Come si fa a mantenere costante la coppia max?

Dobbiamo far entrare in gioco anche l'ampiezza della tensione di alimentazione. Infatti la coppia max è funzione sia della frequenza che dell'ampiezza della tensione di alimentazione, secondo una legge del tipo:

$$C_{m_max} = k \cdot \frac{V^2}{f^2}$$

Ora, siccome noi (per aumentare la velocità) stiamo facendo crescere la frequenza, vediamo (dall'ultima relazione) che, per mantenere costante la coppia, dobbiamo far crescere anche l'ampiezza della tensione di alimentazione simultaneamente alla frequenza, in modo che risulti $C_{m_max} = \text{costante}$ ¹⁰.

Fino a quando possiamo far crescere la velocità facendo crescere simultaneamente frequenza e ampiezza della tensione di alimentazione?

⁷ eventualmente

⁸ eventualmente

⁹ $C_{m_max} = k \cdot \phi \cdot I_2 \cdot \cos(\varphi_2)$ con φ fattore di potenza del circuito di rotore

¹⁰ Osserviamo che far crescere la velocità a coppia max costante significa far traslare la caratteristica elettromeccanica parallelamente a se stessa

Fino a quando non viene raggiunta la tensione nominale dell'inverter: a quel punto l'ampiezza V non può più essere aumentata e deve rimanere costante.

Se imponiamo un ulteriore aumento di frequenza, questo determinerà una diminuzione del flusso (in quanto $V = k \cdot f \cdot \phi$, per cui, se V è costante, un aumento di f determina una riduzione di ϕ).

La diminuzione del flusso comporterà a sua volta una diminuzione di coppia massima¹¹ e un mantenimento della potenza a livello costante

(in quanto la potenza trasmessa al rotore è: $P_T = 2\pi f \cdot c_{m_max}$).

Dunque questa seconda “fase” dell'aumento di velocità avviene a tensione costante e a potenza costante.¹²

Abbiamo visto, in sintesi, che l'aumento di velocità mediante aumento della frequenza di alimentazione si sviluppa in due fasi:

- una prima fase, che avviene a **coppia massima costante** (e a **flusso** magnetico costante) e che richiede anche un **aumento** dell'**ampiezza** delle tensioni di alimentazione, simultaneo all'aumento della frequenza
- una seconda fase, che avviene a **potenza costante** (e ad **ampiezza** della tensione di alimentazione costante) e che richiede una **diminuzione del flusso magnetico**.

CONCLUSIONE

*In sintesi, ciò che noi vorremmo sarebbe una variazione di velocità **a coppia costante**, ma il controllo a coppia costante risulta **possibile solo per velocità non molto elevate**.*

*Infatti, affinché il motore funzioni a coppia costante, deve mantenersi **costante il rapporto V/f** , però, se la velocità desiderata è alta, deve essere alta anche frequenza f , e ciò richiederebbe (per la costanza della coppia) un proporzionale aumento anche di V , cioè della tensione di alimentazione dell'inverter.*

E' evidente quindi che, se la velocità richiesta è tanto elevata da richiedere, per la costanza della coppia, una tensione di alimentazione maggiore di quella massima fornita dall'inverter, una volta raggiunta la V_{MAX} dell'inverter, l'aumento di velocità non avviene più a coppia costante, ma a potenza costante.

¹¹ $c_{m_max} = k \cdot \phi \cdot I_2 \cdot \cos(\varphi_2)$ con fattore di potenza del circuito di rotore.

¹² comporta una diminuzione della coppia max e quindi una traslazione con abbassamento delle caratteristiche elettromeccaniche.

ALTRE TECNICHE DI CONTROLLO DELLA VELOCITA' DEL MOTORE ASINCRONO

CONTROLLO A SCATTI DELLA VELOCITA' DEL MOTORE ASINCRONO MEDIANTE CAMBIAMENTO DEL NUMERO DEI POLI

E' adottata per alcune macchine utensili.

Si può ottenere:

- con due avvolgimenti di rotore distinti, ma è una soluzione costosa perché richiede un aumento delle dimensioni del motore
- mediante speciali avvolgimenti commutabili che però permettono il passaggio soltanto fra numeri di poli che stanno nel rapporto 1:2

CONTROLLO A SCATTI DELLA VELOCITA' DEL MOTORE ASINCRONO MEDIANTE COLLEGAMENTO IN CASCATA DI DUE MOTORI CON UGUALI CARATTERISTICHE

E' stato impiegato in alcuni tipi di locomotive italiane.

I rotori dei due motori sono collegati meccanicamente fra loro.

Dal punto di vista elettrico, il rotore di un motore alimenta lo statore dell'altro.

In questo modo la velocità sincrona del gruppo diventa la metà di quella che avrebbe ciascun motore se venisse alimentato direttamente dalla linea.

In generale infatti, detti p_1 e p_2 i numeri di coppie polari dei due motori, la velocità sincrona del gruppo è data da:

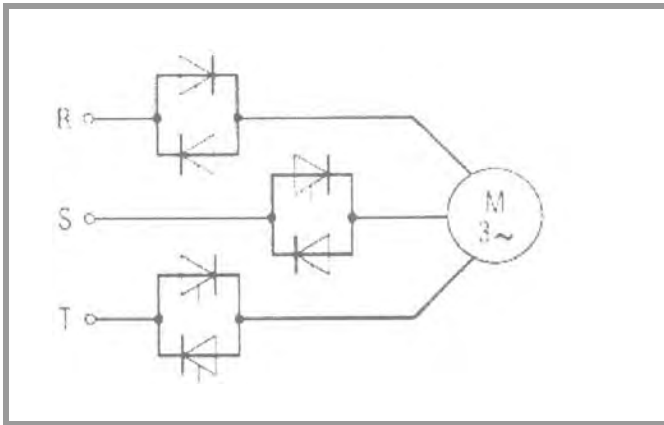
$$n_s = \frac{60 \cdot f}{p_1 + p_2}$$

CONTROLLO DELLA VELOCITA' DEL MOTORE ASINCRONO MEDIANTE VARIAZIONE CONTINUA DELLA SOLA AMPIEZZA DELLA TENSIONE DI ALIMENTAZIONE

Con questa tecnica si possono ottenere solo piccole variazioni continue di velocità.

In particolare sono possibili riduzioni di velocità (rispetto al valore a pieno carico e a tensione nominale) entro il 10% di questo valore.

Il sistema di controllo può essere realizzato variando il valore efficace delle tensioni applicate allo statore mediante tre coppie di SCR, con gli SCR di ciascuna coppia posti in antiparallelo e controllati in fase.



Un circuito di controllo viene collegato ai terminali di gate degli SCR e impone ai componenti un opportuno ritardo di fase.

Nel caso di motori di piccola potenza le coppie di SCR possono essere sostituite da TRIAC.

Questo sistema è detto “**variatore o regolatore di tensione a.c.**”

La riduzione della tensione di alimentazione non agisce direttamente sulla coppia motrice, ma determina un aumento dello scorrimento s e quindi una diminuzione della velocità n di rotore che, a regime, il motore deve raggiungere per sviluppare la coppia richiesta.

Dato cioè un certo valore di coppia, la velocità che il motore deve raggiungere per svilupparla è tanto più piccola quanto minore è la tensione di alimentazione.

CONTROLLO CONTINUO DELLA VELOCITA' DEL MOTORE ASINCRONO MEDIANTE INSERIZIONE DI RESISTORI VARIABILI SUL CIRCUITO DI ROTORE

E' utilizzato prevalentemente per apparecchiature di sollevamento e permette la regolazione di velocità, riducendola entro il 20% del valore nominale.

Presenta alcuni inconvenienti:

- ❖ dà luogo a un funzionamento poco stabile
- ❖ dissipa in calore, sulle resistenze di regolazione, tutta la potenza corrispondente alla riduzione di velocità
- ❖ è applicabile solo a motori con **rotore avvolto**.

PREGI E INCONVENIENTI DEL MOTORE ASINCRONO

APPLICAZIONI

I motori **asincroni** sono i motori **elettrici industriali più diffusi**.

Presentano però qualche problema nell'avviamento e nella regolazione della velocità, che in passato ne hanno limitato l'utilizzazione al settore ferroviario.

Possono essere usati per:

- ❖ pompe centrifughe e alternative
- ❖ presse
- ❖ ventilatori
- ❖ montacarichi
- ❖ elevatori (ascensori)
- ❖ gru
- ❖ compressori per frigoriferi
- ❖ alcune macchine utensili
- ❖ alcuni elettrodomestici.

Sono usati come motori e raramente come generatori.

Sono destinati a funzionare su reti polifasi e soprattutto trifasi.

Sono anche realizzati, mediante opportuni accorgimenti, motori asincroni monofasi per piccole potenze e nel settore degli elettrodomestici.

CARATTERISTICHE GENERALI

La velocità della macchina (se non la si varia deliberatamente con una delle possibili tecniche di regolazione) è praticamente costante:

- ❖ dipende dal numero di poli dell'avvolgimento di statore
- ❖ al variare del carico subisce solo piccole variazioni dovute alla variazione di scorrimento
- ❖ campo di potenze: (100W- qualche decina di MW)
- ❖ campo di velocità: (750-3000) giri/min

PREGI IN ASSOLUTO

- ❖ Semplicità di costruzione
- ❖ Economia di costo
- ❖ Robustezza e resistenza a sovraccarichi di potenza
- ❖ Semplicità di avviamento e possibilità di manovra da parte di personale non specializzato.

PREGI RISPETTO AL MOTORE c.c.

- ❖ assenza di spazzole e di conseguente scintillio
- ❖ la velocità massima è limitata solo da elementi meccanici
- ❖ dimensioni ridotte grazie al rotore monolitico
- ❖ maggiore semplicità e minor costo
- ❖ possibilità di alimentare il motore direttamente dalla rete se si guasta l'inverter
- ❖ rendimento più elevato
- ❖ assenza quasi totale di disturbi nel caso di adozione del sistema "raddrizzatore-inverter-PWM".

INCONVENIENTI RISPETTO AI MOTORI C.C. DOTATI DI TRASDUTTORI

- ❖ minor campo di variazione di velocità rispetto ai motori c.c.
- ❖ minor precisione della regolazione a bassa velocità
- ❖ minor campo di funzionamento a potenza costante.

CONFRONTO CON IL MOTORE SINCRONO

La **velocità meccanica**, cioè la velocità del rotore, del motore **sincrono** è **univocamente legata** alla **frequenza** delle tensioni di alimentazione.

Se la frequenza di alimentazione è costante, la velocità del motore sincrono è rigorosamente costante.

Siccome la velocità del motore sincrono è univocamente determinata dalla frequenza di alimentazione, per questo motore **si definisce una sola velocità**, mentre per il motore asincrono si fa una distinzione fra la velocità del campo rotante, o velocità di sincronismo, e la velocità meccanica, che è la velocità del rotore.

La **velocità meccanica** del motore **asincrono** dipende anch'essa **essenzialmente dalla frequenza** della tensione di alimentazione, **ma anche**, sia pure **leggermente**, dall'eventuale **variazione del carico**.

A frequenza di alimentazione costante, la velocità meccanica del motore asincrono è “quasi costante” nel senso che **subisce leggere variazioni al variare del carico o, più precisamente, al variare della coppia richiesta dal carico**.

La **velocità meccanica di regime** è quella che si ha **nel momento in cui c'è uguaglianza fra la coppia resistente del carico e la coppia motrice del motore**.

La **variabilità della velocità meccanica** (rispetto alla velocità del campo) del motore asincrono **per effetto di fattori come la variazione di carico** è espressa dal parametro **s** di **scorrimento**: la velocità meccanica “scorre” cioè diminuisce¹³ rispetto alla velocità del campo rotante per effetto delle eventuali variazioni di carico.

Da queste considerazioni emerge che il motore sincrono va preferito quando non sono tollerate variazioni di velocità per effetto delle variazioni di carico.

¹³ tipicamente in una misura che va dal 2% al 4%

TABELLE DI CONFRONTO ASINCRONO/SINCRONO

PARAMETRO	motore ASINCRONO	motore SINCRONO
POTENZA	media o piccola (centinaia di W - qualche decina di MW)	grossa (dalle centinaia di KW alle centinaia di MW)
AVVIAMENTO	autoavviante (come i motori c.c.)	NON autoavviante: richiede un motore o altri dispositivi di “lancio”
INVERSIONE DEL SENSO DI MARCIA	Si ottiene scambiando fra loro due conduttori di linea sui morsetti del motore e lasciando il terzo al suo posto	Si ottiene scambiando fra loro due qualsiasi fasi della linea trifase di alimentazione

PREGI E INCONVENIENTI	motore ASINCRONO	motore SINCRONO
PREGI	Semplicità economia robustezza resistenza a sovraccarichi di potenza attualmente non richiede manovre complicate per l'avviamento	Velocità univocamente determinata dalla frequenza della tensione di alimentazione e rigorosamente costante se la frequenza di alimentazione è costante elevato rapporto potenza / peso bassa inerzia del rotore e conseguente facilità di accelerazione
INCONVENIENTI	Utilizzabile solo per potenze non elevatissime Presenta alcuni problemi relativi all'avviamento e alla regolazione di velocità, che in passato ne hanno limitato l'impiego calore generato anche sul rotore per effetto delle correnti indotte elevata inerzia del rotore e conseguenti difficoltà di accelerazioni	➤ Costo più elevato rispetto agli asincroni ➤ esigenza di un motore di lancio o di altri dispositivi di avviamento ➤ necessità di corrente continua per l'eccitazione dell'induttore ➤ il funzionamento avviene a velocità rigorosamente costante per cui il motore sincrono non può essere usato in quelle applicazioni che richiedono frequenti cambiamenti di velocità ➤ facilità con la quale "perde il passo" cioè esce di sincronismo per effetto di variazioni della coppia resistente

APPLICAZIONI E CARATTERISTICHE	motore ASINCRONO	motore SINCRONO
APPLICAZIONI	Pompe compressori per frigoriferi montacarichi e ascensori ponti mobili gru elevatori ventilatori motori laterali per navi	Forni rotativi per cemento frantoi pompe a vuoto laminatoi compensatori o condensatori sincroni per rifasamento di impianti elettrici motori per navi a propulsione elettrica Sono talvolta usati anche in applicazioni di piccola potenza e per gruppi elettrogeni di emergenza
CARATTERISTICHE SALIENTI	velocità praticamente (ma non esattamente) costante al variare del carico	per una data frequenza, la velocità di rotazione è rigorosamente costante: l'esigenza di mantenere il sincronismo fra rotore e campo rotante porta questo motore a ad avere una velocità di rotazione rigorosamente costante inoltre se la coppia resistente supera il valore limite il motore rallenta perdendo il sincronismo e si ferma

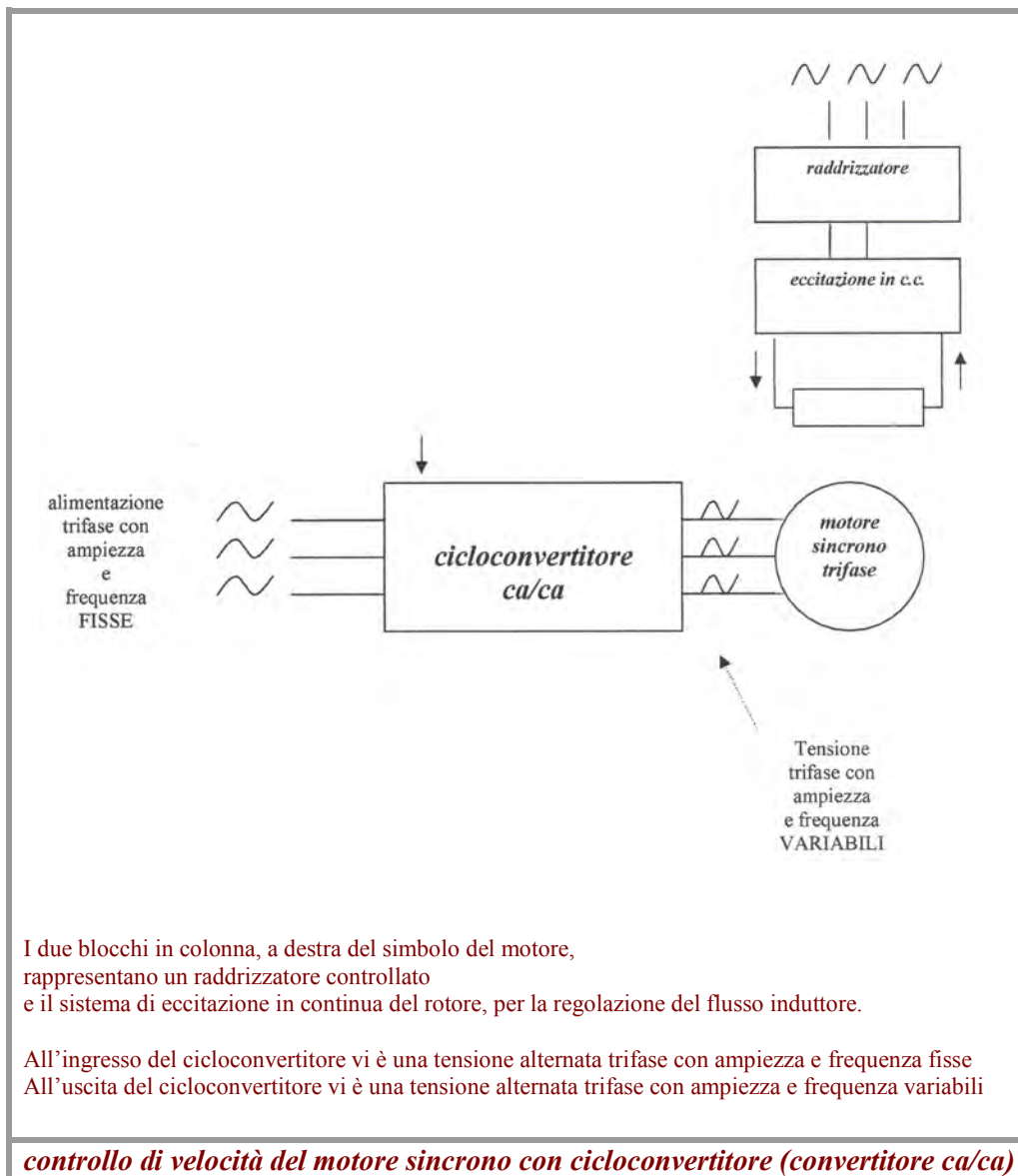
CONTROLLO DI VELOCITA' DEL MOTORE SINCRONO

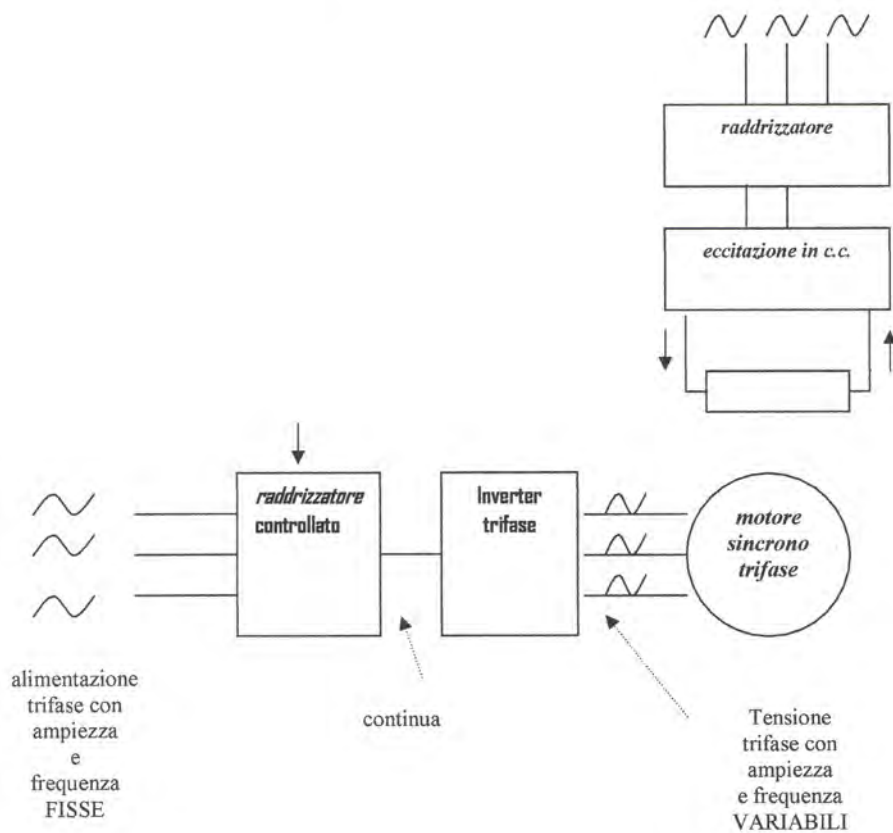
VARIAZIONE DI VELOCITA' A COPPIA COSTANTE

Per avere una variazione di velocità a coppia costante bisogna:

- ❖ **variare la frequenza** e, **proporzionalmente** alla frequenza, anche l'**ampiezza** (come per il motore asincrono)
- ❖ mantenere costante il flusso di eccitazione
- ❖ limitare la massima corrente assorbita dal motore, per evitare perdite eccessive negli avvolgimenti
- ❖ mantenere costante l'angolo di coppia.

CONTROLLO DI VELOCITA' DEL MOTORE SINCRONO PER VARIAZIONE DI FREQUENZA





I due blocchi in colonna, a destra in alto, rappresentano un raddrizzatore controllato e il sistema di eccitazione in continua del rotore, per la regolazione del flusso induttore.

Il primo blocco da sinistra è un **raddrizzatore** controllato, il secondo è un **inverter trifase**
 All'ingresso del raddrizzatore vi è una tensione alternata trifase a frequenza e ampiezza fisse
 All'uscita del raddrizzatore (e quindi all'ingresso dell'inverter) vi è una tensione continua
 All'uscita dell'inverter vi è una tensione alternata trifase a frequenza variabile

**controllo di velocità del motore sincro
 con raddrizzatore (convertitore ca/cc) e inverter (convertitore cc/ca)**

PROPULSIONE ELETTRICA NAVALE

La propulsione **elettrica** riveste grande importanza per le **grosse navi passeggeri** e per le **navi-traghetto** (per il trasporto di persone e veicoli)

La propulsione elettrica in **corrente continua** attualmente è di fatto **inutilizzata a causa dei suoi inconvenienti, che sono essenzialmente:**

- la presenza di spazzole, soggette a usura e bisognose di manutenzione
- lo scintillio
- i picchi di corrente in commutazione
- la minore potenza, a parità di dimensioni, rispetto al motore in alternata
- la “parte di potenza” più complessa

Sono invece sicuramente interessanti la propulsione in alternata mediante motore c.a. asincrono e mediante motore c.a. sincrono, tipicamente trifasi.

In una nave con propulsione elettrica l’energia elettrica prodotta mediante gruppi elettrogeni è utilizzata non solo per la propulsione, ma anche per l’alimentazione delle utenze di bordo.

Ciascun gruppo elettrogeno è formato da un motore diesel (generalmente a media velocità) e da un alternatore.

I vantaggi offerti dalla propulsione elettrica in alternata rispetto alla propulsione non elettrica sono:

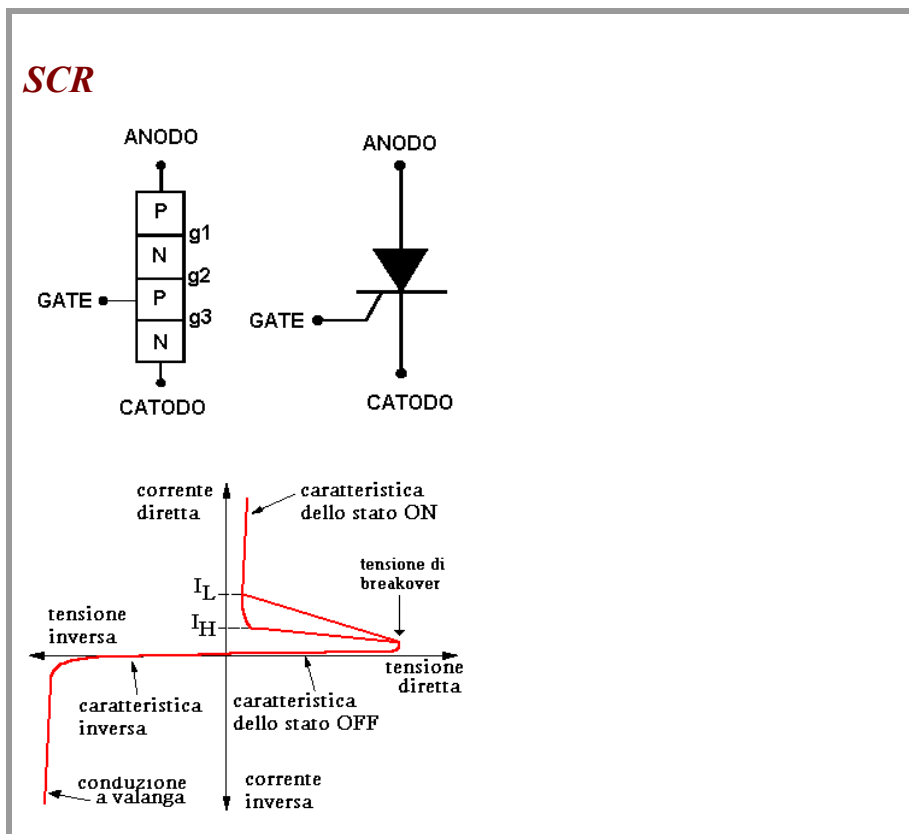
- la riduzione del rumore e delle vibrazioni
- la semplicità di manutenzione della macchine elettriche in alternata (motori e alternatori)
- la possibilità di ridurre il numero dei motori diesel a bordo
- la possibilità di usare i diesel alternatori semiveloci (poco ingombranti)
- un controllo di velocità che è diventato abbastanza agevole grazie allo sviluppo dei componenti elettronici al silicio per il controllo di potenza.

IL CONTROLLO DI VELOCITA' DEI MOTORI IN ALTERNATA E I TIRISTORI

Un punto di svolta nell'evoluzione dei motori elettrici nautici è stata la diffusione dei componenti elettronici allo stato solido per il controllo della potenza elettrica: i cosiddetti tiristori (si veda il relativo capitolo).

Nella terminologia tecnica la parola "tiristore" è usata in qualche caso per indicare specificamente il più noto e diffuso fra tutti i componenti al silicio per il controllo di potenza (e cioè il Raddrizzatore Controllato al Silicio o SCR), in altri casi per indicare genericamente tutti i componenti per il controllo della potenza (SCR, TRIAC, DIAC ecc.).

Solo la disponibilità di dispositivi basati su tiristori ha reso agevole la regolazione di velocità dei motori in alternata.



MOTORE SINCRONO O ASINCRONO?

Il motore sincrono è caratterizzato da una velocità che, una volta fissato il numero di coppie polari, dipende esclusivamente dalla frequenza della tensione di alimentazione.

La velocità meccanica (cioè il numero di giri n del rotore) è esattamente uguale alla velocità di sincronismo n_s con la quale ruota il campo .

La velocità di SINCRONISMO, cioè la velocità del CAMPO del motore è data da:

$$n_s = \frac{60 \cdot f}{p} \quad (\text{velocità in giri/min})$$

con:

p = numero di coppie polari per ogni fase

f = frequenza di alimentazione

La velocità del motore *asincrono* invece ***non dipende esclusivamente dalla frequenza, ma anche dalla coppia resistente applicata all'albero: nel senso che la velocità tende a diminuire se la coppia resistente aumenta.***

La velocità meccanica (cioè il numero di giri n del rotore) è sempre leggermente più bassa della velocità di sincronismo n_s del campo magnetico rotante.

L'effetto della coppia resistente sulla velocità del motore asincrono viene descritto matematicamente mediante un parametro chiamato scorrimento o slittamento:

$$\text{scorrimento percentuale del motore asincrono } s_{\%} = \frac{n_s - n}{n_s} \cdot 100$$

(tipicamente compreso fra il 2% e il 4%)

Nel momento in cui lo scorrimento si annulla, il motore si arresta.

La velocità meccanica del motore asincrono si esprime come: $n = n_s \cdot (1 - s)$

In base a quali criteri si decide di equipaggiare una nave a propulsione elettrica con un motore sincro o piuttosto che con un asincro?

In generale possiamo dire che:

- ***per potenze più basse***, e cioè fino a **8 MW** si preferisce il motore ***asincrono***.
- ***per potenze più alte*** (tipicamente fino a **30 MW**) si preferisce il motore **sincro**.

Le ragioni tecniche che portano a questa scelta sono sinteticamente queste:

- il motore sincro è caratterizzato da un rapporto potenza/peso più elevato (per cui risulta più conveniente nel caso di grosse navi e quindi di grosse potenze elettriche)
- il motore **sincro** va **preferito** quando **non sono tollerate variazioni di velocità** per effetto delle **variazioni di carico**.

D'altra parte il motore sincro presenta rispetto all'asincro, alcuni inconvenienti:

- non è autoavviante, nel senso che per "partire" ha bisogno di un sistema di eccitazione (in corrente continua) del rotore
- è più costoso
- non è adatto a quelle applicazioni che richiedono frequenti cambiamenti di velocità.

CAPITOLO 11

I SEMICONDUTTORI PER L'ELETTRONICA

OBIETTIVI

- conoscere le differenze fra semiconduttori, conduttori e isolanti
- comprendere i motivi per i quali, attualmente, i componenti elettronici sono basati sui semiconduttori
- conoscere il significato e le finalità del drogaggio
- comprendere i concetti di giunzione e polarizzazione

PREREQUISITI

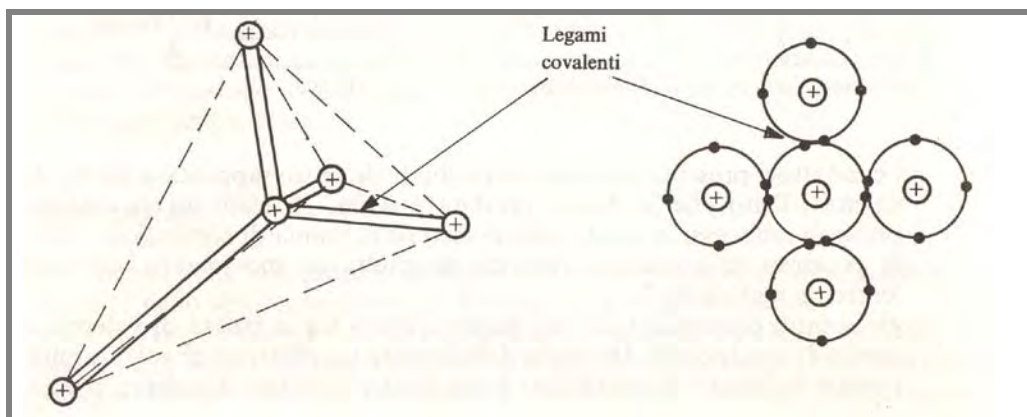
- Nozioni di base sui legami chimici
- Capitolo 3

SEMICONDUTTORI E DROGAGGIO

COMPONENTI “ALLO STATO SOLIDO” O “A SEMICONDUTTORE”

Gli odierni componenti a semiconduttore sono detti “allo stato solido”, in contrapposizione ai vecchi tubi a vuoto (valvole) costituiti da filamenti e griglie metalliche incapsulati in ampole di vetro sotto vuoto.

I semiconduttori attualmente più usati per i componenti elettronici sono in primo luogo il silicio e poi il germanio e l'arseniuro di gallio



semiconduttore intrinseco: rappresentazione semplificata

L'atomo di silicio (Si) possiede 14 elettroni.

L'atomo di germanio (Ge) possiede 32 elettroni.

Il **germanio e il silicio** sono **tetravalenti**.

L'atomo di questi materiali ha quindi quattro elettroni di valenza, cioè **quattro elettroni** (appartenenti all'orbita più esterna) **utilizzabili per legarsi ad altri atomi**.

Ogni atomo di Si (o di Ge) si lega ad altri quattro atomi di Si (o di Ge), mettendo in “comune” i suoi quattro elettroni di valenza con i quattro elettroni appartenenti, ciascuno, a uno dei quattro atomi di contorno. Si forma così nello spazio un tetraedro con un atomo al centro e quattro ai vertici.

La sequenza spaziale di tetraedri forma la struttura cristallina del semiconduttore.

Si realizza così attorno all'atomo “centrale” un “OTTETTO” di elettroni, sul quale si basa il legame covalente

CARATTERISTICHE DEI SEMICONDUTTORI

CONDUCIBILITA' DIPENDENTE DALLA TEMPERATURA

Contrariamente ai conduttori (es. metalli) che hanno tutti gli elettroni in condizioni di "libertà" dai vincoli del legame chimico già alla temperatura di 0°K (cioè di -273°C), i semiconduttori, a 0°K, presentano solo pochissimi elettroni liberi.

Gli elettroni del semiconduttore, per liberarsi dal legame chimico, hanno bisogno che venga loro fornita energia termica o luminosa.

Al crescere della temperatura cresce la concentrazione di elettroni liberi.

CARICHE LIBERE NEGATIVE E POSITIVE

Nel semiconduttore sono presenti anche **cariche elettriche di segno positivo**, dette "**lacune**".

Quando un elettrone riesce a liberarsi dal legame chimico e si allontana dall'atomo di appartenenza, si crea in tale atomo un "vuoto di carica".

Questa mancanza di carica negativa equivale a una carica positiva (la lacuna).

Anche le lacune si spostano: il "vuoto di carica" di un atomo A può essere riempito da un altro elettrone (che sia riuscito al liberarsi da un atomo B), ciò equivale a dire che la lacuna si è spostata dall'atomo A all'atomo B.

INADEGUATEZZA DEL SEMICONDUTTORE PURO

ALLA REALIZZAZIONE DI COMPONENTI: IL DROGAGGIO

Il "drogaggio" è una tecnica che consiste nell'introdurre, all'interno della struttura cristallina del semiconduttore, atomi di altri materiali, detti "materiali droganti" in modo da aumentare la concentrazione di cariche libere e ridurre la dipendenza delle correnti dalla temperatura.

Il semiconduttore puro o intrinseco, cioè non drogato, è inadeguato alla realizzazione dei componenti elettronici per i seguenti motivi:

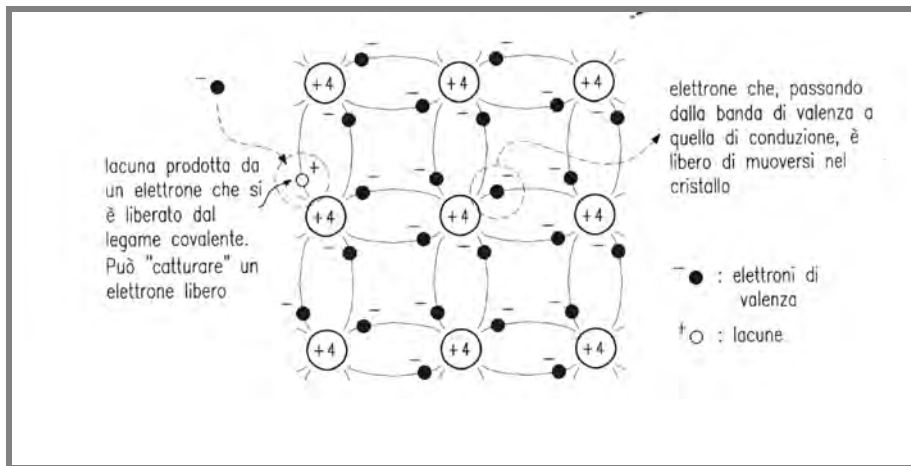
- la conducibilità è troppo bassa per le normali applicazioni elettroniche
- le correnti sono troppo dipendenti dalla temperatura.

Si ricorre allora alla tecnica del drogaggio.

PERCHE' IL DROGAGGIO AUMENTA LA CONDUCIBILITA' E RENDE LA CORRENTE MENO DIPENDENTE DALLA TEMPERATURA

Il drogaggio aumenta la conducibilità del semiconduttore e rende la corrente meno dipendente dalla temperatura perché gli elettroni dovuti alla presenza di atomi dei materiali droganti (inseriti nella struttura cristallina del semiconduttore) non risultano sottoposti alla forza del legame chimico covalente del semiconduttore e quindi non hanno bisogno di energia termica (calore) per poter partecipare alla formazione di correnti.

Questi elettroni provenienti dal materiale drogante esterno (non essendo vincolati al legame chimico covalente) sono quindi immediatamente disponibili per la formazione di correnti elettroniche senza bisogno di energia termica, il che spiega la minore dipendenza della conducibilità dalla temperatura.



DROGAGGIO DI TIPO “P” (POSITIVO)

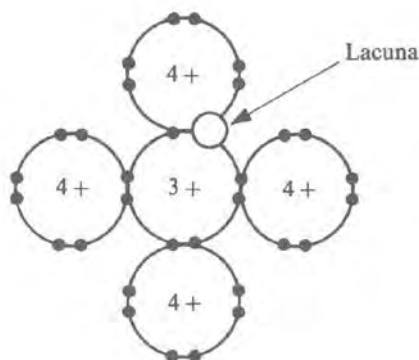
Il drogaggio è di **tipo P** quando la maggior parte delle **cariche elettriche libere** che esso produce nel semiconduttore è formata da **cariche positive**, cioè da **lacune**. Si usa dire perciò che in un semiconduttore con drogaggio P le lacune sono le cariche libere **maggioritarie**.

Nel semiconduttore P vi sono anche elettroni liberi, ma sono in minoranza.

Il drogaggio P è ottenuto introducendo nel semiconduttore **atomi di sostanze trivalenti**, cioè con **tre** elettroni di valenza (**alluminio, boro, gallio, indio**).

Questi atomi, introdotti nel semiconduttore, diventano ioni negativi e si comportano come cariche elettriche fisse, le quali, ovviamente, non contribuiscono alla formazione di correnti.

Il semiconduttore drogato di tipo P (come anche quello di tipo N che vedremo) è **eletttricamente neutro** perché la somma algebrica delle cariche positive e negative risulta nulla: per ogni ione negativo c'è una lacuna.



semiconduttore con drogaggio “p”: rappresentazione semplificata

Come si vede dalla figura, attorno all’atomo trivalente **non può costituirsi l’“ottetto” completo** di elettroni sul quale si basa il legame covalente. Per la realizzazione dell’ottetto **manca un elettrone**, e quindi una carica negativa. Questa **“mancanza” di carica negativa** è una carica positiva, che viene chiamata **“lacuna”**

Gli elementi chimici **trivalenti** solitamente usati per il drogaggio “n” sono: **alluminio, boro, gallio, indio**.

DROGAGGIO DI TIPO “N” (NEGATIVO)

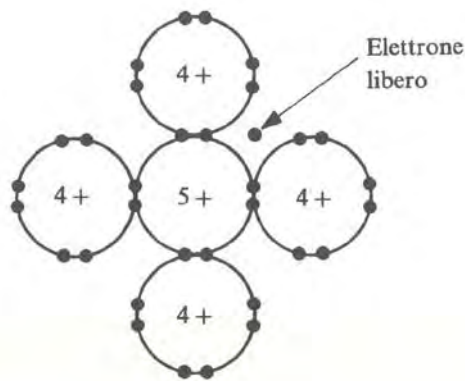
Il drogaggio è di **tipo N** quando la maggior parte delle **cariche elettriche libere** che esso produce nel semiconduttore è formata da **cariche negative**, cioè da **elettroni**. Si usa dire perciò che in un semiconduttore con drogaggio N gli elettroni sono le cariche libere **maggioritarie**.

Nel semiconduttore N vi sono anche lacune libere, ma sono in minoranza.

Il drogaggio N viene ottenuto introducendo nel semiconduttore **atomi di sostanze pentavalenti**, cioè con **cinque** elettroni di valenza (**fosforo, arsenico, antimonio, bismuto**).

Questi atomi, introdotti nel semiconduttore, diventano ioni positivi e si comportano come cariche elettriche fisse, le quali, ovviamente, non contribuiscono alla formazione di correnti.

Il semiconduttore drogato di tipo N (come anche quello di tipo P) è **eletttricamente neutro** perché la somma algebrica delle cariche positive e negative risulta nulla: per ogni ione positivo c'è un elettrone.

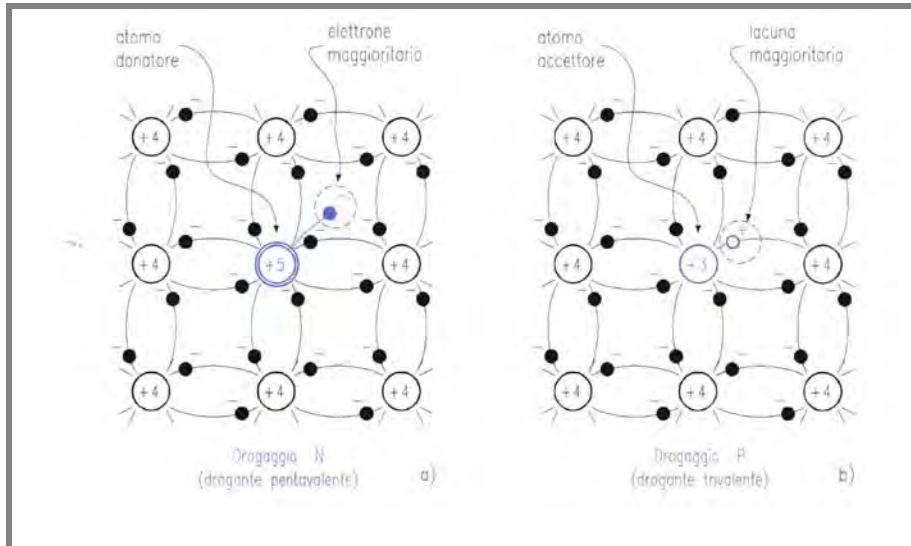


semiconduttore con drogaggio “n”: rappresentazione semplificata

Come si vede dalla figura, attorno all'atomo **pentavalente**, oltre all'“**ottetto**” **completo di elettroni sul quale si basa il legame covalente**, vi è un ulteriore elettrone che, proprio perché è in eccesso (rispetto all'ottetto), è debolmente vincolato al legame e rappresenta una “carica libera” negativa.

Gli elementi chimici **pentavalenti** solitamente usati per il drogaggio “n” sono: **fosforo, arsenico, antimonio, bismuto**.

DROGAGGI DI TIPO “P” ED “N” A CONFRONTO



GIUNZIONE (JUNCTION) PN

La giunzione PN è la superficie di “confine” (di “separazione”) fra due zone di semiconduttore diversamente drogate.

E' essenziale analizzare la possibilità di passaggio di corrente attraverso una giunzione PN.

La giunzione PN è una struttura elettronica fondamentale ed è alla base del funzionamento di quasi tutti i componenti elettronici allo stato solido.

Il diodo, per esempio, è, in linea di principio, una semplice giunzione PN con due contatti metallici. Il transistor BJT è basato su due giunzioni PN.

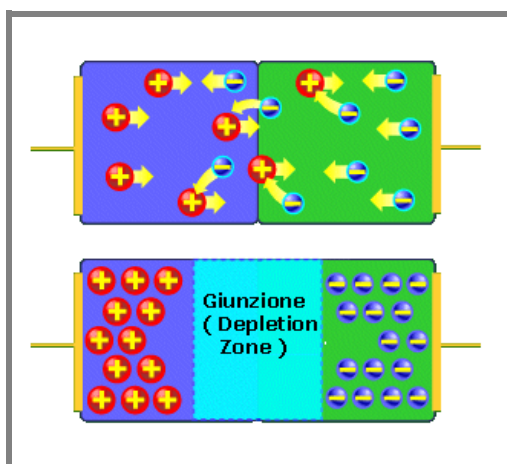
Esaminiamo ora cosa succede appena si crea una giunzione (senza applicazione di campi dall'esterno).

Gli elettroni, che inizialmente si trovano **affollati** nella zona N, nella quale sono maggioritari, tendono, per un fenomeno di **diffusione**, a oltrepassare la giunzione e a portarsi nella regione P.

Per lo stesso motivo le lacune, maggioritarie nella zona P, tendono a portarsi nella zona N.

In questo modo, immediatamente a sinistra e a destra della superficie di giunzione, elettroni e lacune si neutralizzano elettricamente fra di loro creando una zona di svuotamento o di carica spaziale.

In questa zona di svuotamento rimangono “elettricamente scoperti” gli ioni positivi e negativi, per cui si crea un campo elettrico e quindi una d.d.p., che costituisce una “barriera di potenziale”.



POLARIZZAZIONE DI UNA GIUNZIONE PN

Polarizzare la giunzione significa applicare una tensione (continua) ai suoi capi attraverso un generatore di tensione (costante), per esempio una batteria.

Si parla di polarizzazione DIRETTA quando viene portata a potenziale PIU' ALTO la ZONA P della giunzione, mentre si parla di polarizzazione inversa quando viene portata a potenziale più alto la zona n della giunzione.

POLARIZZAZIONE DIRETTA

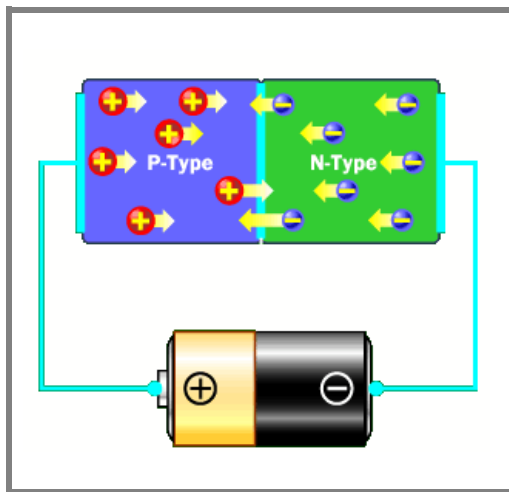
Il polo positivo della batteria è applicato alla zona p.

Cosa fanno le cariche libere di maggioranza?

Il polo positivo della batteria attira gli elettroni spingendoli ad attraversare la giunzione.

Il polo negativo della batteria attira le lacune spingendoli ad attraversare la giunzione.

Si determina quindi una corrente di attraversamento della giunzione convenzionalmente diretta dalla zona P alla zona N. Questa corrente si chiama "corrente DIRETTA".



POLARIZZAZIONE INVERSA

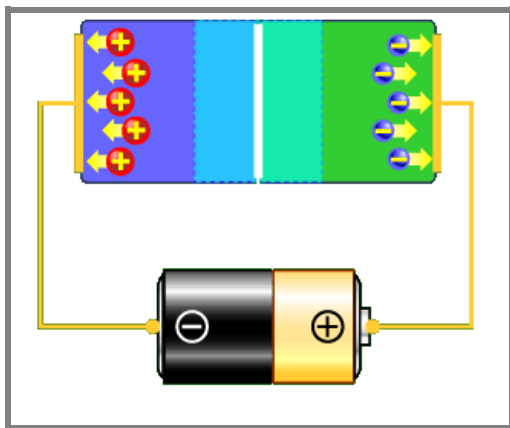
Il polo positivo della batteria è applicato alla zona N.

Cosa fanno le cariche libere di maggioranza?

Non attraversano la giunzione, anzi si allontanano da essa senza attraversarla.

La corrente diretta è quindi nulla e possiamo dire che la giunzione non conduce.

A rigore c'è però una piccola corrente inversa formata da cariche minoritarie (elettroni della zona P e lacune della zona N). Questa piccola corrente inversa è convenzionalmente diretta da N verso P ed è – per i componenti al silicio – spesso trascurabile.



CAPITOLO 12

IL DIODO

COMPLEMENTI

SOTTOCLASSIFICAZIONE DEI DIODI NORMALI

DIODI DI SEGNALE E DIODI DI POTENZA

PRINCIPI DI FUNZIONAMENTO DEL DIODO LASER

PRESTAZIONI DEI FOTOEMETTITORI

FOTORIVELATORI

PARAMETRI DI UN FOTORIVELATORE

FOTODIODO PIN (Positivo-Intrinseco-Negativo)

FOTODIODO APD

PRESTAZIONI DEI FOTORIVELATORI

SOTTOCLASSIFICAZIONE DEI DIODI “NORMALI”:

NOME	NOME INGLESE	DEFINIZIONE	ESEMPI COMMERCIALI
DIODI PER USI GENERALI	General purpose diodes	Sono quei diodi che NON presentano l’ottimizzazione di nessun parametro elettrico Tensioni di rottura tipiche: decine di V Correnti dirette massime: centinaia di mA	BA 128 BAV 17
DIODI PER COMMUTAZIONE	Switching diodes	Sono quei diodi che risultano adatti al funzionare con segnali dalle commutazioni veloci Presentano piccoli valori dei tempi di commutazione e in particolare del tempo di recupero inverso	1N 4148 1N 914B
DIODI SPECIFICAMENTE RADDRIZZATORI (e diodi raddrizzatori “veloci”)	Rectifier diodes (e “fast rectifiers – soft recovery”)	Sono chiamati in questo modo quei diodi concepiti per raddrizzare la TENSIONE di RETE (o, comunque, tensioni alternate di notevole ampiezza e correnti notevoli)	1N4001 1N4007
ponti integrati di diodi raddrizzatori			FSA 27 04 M

DIODI DI SEGNALE E DIODI DI POTENZA



Si usa anche suddividere i diodi in due grandi categorie:

1. diodi “di **SEGNALE**” (impiegati in circuiti interessati da correnti dell’ordine delle **DECINE di MILLIampère**)
2. diodi “di **POTENZA**” (impiegati in circuiti interessati da correnti dell’ordine delle **DECINE o CENTINAIA di AMPERE**)

Per renderci conto della differenza facciamo un confronto tra i parametri di due diodi appartenenti alle due diverse categorie.

	DIODO DI SEGNALE BAX 13	DIODO DI POTENZA BY 127
Massimo valore di corrente continua	75 mA	1 A
<i>Massimo valore di picco della corrente a 50 Hz</i>	<i>150 mA</i>	<i>10 A</i>
Massima tensione inversa	50 V	1250 V
<i>Massima caduta di tensione diretta</i>	<i>< 1,5V (per $I = 75\text{ mA}$)</i>	<i>< 1,5V (per $I = 1\text{ A}$)</i>
Massimo valore della temperatura di giunzione	200 °C	150 °C

PRINCIPI DI FUNZIONAMENTO DEL DIODO LASER

EMISSIONE STIMOLATA

Il diodo laser è basato sul fenomeno dell'emissione stimolata (di fotoni, cioè di luce) che viene di seguito descritta.

Quando un **elettrone**, che si trova a un livello di energia superiore, viene **colpito da un fotone** dotato di una quantità di energia pari al gap energetico (esistente fra il livello in cui l'elettrone si trova e il livello sottostante), allora avviene la transizione forzata dell'elettrone al livello sottostante. Durante questa transizione forzata in discesa, l'elettrone non assorbe il fotone incidente, anzi **emette un secondo fotone**. Il secondo fotone, avendo energia pari al gap, è caratterizzato dalla **stessa frequenza**, dalla **stessa fase** e dalla **stessa direzione** del fotone incidente. Si ha quindi una emissione di **luce coerente**;

E' proprio l'emissione stimolata a consentire al laser di emettere una potenza ottica molto maggiore di quella del led.

INVERSIONE DI POPOLAZIONE

Affinché possa verificarsi l'emissione stimolata di cui abbiamo parlato prima, è necessario che si verifichi nel materiale una inversione di popolazione.

L'inversione di popolazione consiste nel fatto che, contrariamente alla norma, vi siano **piu' elettroni** nella **banda di energia superiore** che in quella inferiore.

Il fenomeno si chiama "inversione" perché costituisce un capovolgimento della situazione abituale: normalmente infatti è maggiore il numero di elettroni che si trovano nella banda inferiore.

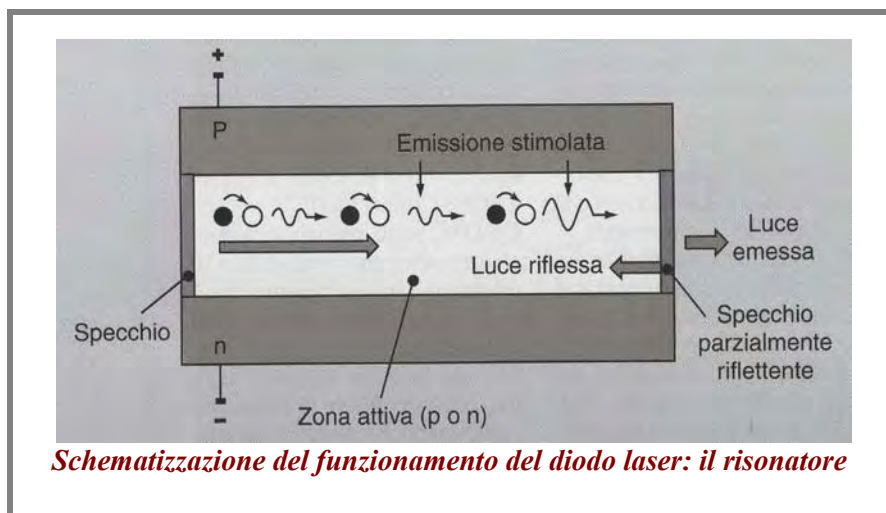
FORTE POLARIZZAZIONE DIRETTA

E' necessaria per imporre l'inversione di popolazione. La forte polarizzazione diretta determina una corrente di intensità sufficiente a iniettare nel diodo un gran numero di cariche libere.

CONCENTRAZIONE DELLA LUCE EMESSA IN UNO STRETTO FASCIO

La luce emessa dal laser rimane concentrata in un angolo di pochi gradi.

Questa caratteristica del diodo laser viene ottenuta realizzando all'interno del componente un **risonatore**.



Il risonatore è caratterizzato dal fatto di presentare, una opposta all'altra, una superficie riflettente al 100% (come uno specchio) e una superficie parzialmente riflettente.

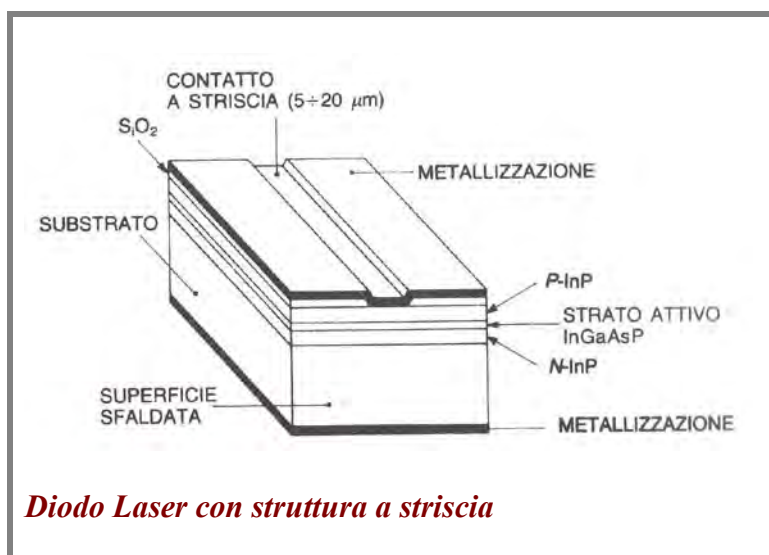
La luce viene confinata tra le due superfici in modo che:

- parte di essa percorra ripetutamente, in avanti e all'indietro, il materiale "eccitato" situato fra le due superfici, determinando, in questo modo, un'amplificazione della luce stessa, ossia l'emissione stimolata di fotoni in una sola direzione
- una parte dell'energia luminosa venga emessa attraverso la superficie parzialmente riflettente.

Il fatto che l'emissione stimolata avvenga in una sola direzione determina la formazione di un fascio molto ristretto nello spazio.

La luce emessa è inoltre molto "ristretta" anche nel dominio della frequenza, nel senso che il suo spettro ha una banda molto stretta (3nm).

(Questo fatto è dovuto proprio al meccanismo di emissione stimolata di luce coerente, per il quale il fotone emesso ha la stessa frequenza di quello incidente).



PRESTAZIONI DEI FOTOEMETTITORI

PARAMETRO	SIMBOLO e unità di misura	LED	LASER
POTENZA iniettata in fibra	P [mW]	0,02 – 0,1	1 – 3
<i>LARGHEZZA SPETTRALE</i>	$\Delta\lambda$ [nm]	30 - 50	<4
CORRENTE DI PILOTAGGIO	[mA]	10 - 200	10 - 200
<i>MASSIMA LUNGHEZZA D'ONDA (di funzionamento)</i>	λ_{max} [nm]	1550	1550
TEMPO DI VITA	[ore]	10^5 - 10^6	10^4 – 10^5
<i>COSTO</i>		<i>BASSO</i>	<i>ALTO</i>

FOTORIVELATORI: PRINCIPI DI FUNZIONAMENTO

La funzione essenziale del ricevitore ottico è la **fotorivelazione**, che consiste nel **convertire** il **segnale ottico** ricevuto in una **corrente elettrica** proporzionale, istante per istante, all'intensità della radiazione luminosa (proveniente, per esempio, da una fibra ottica).

Sono fotorivelatori i fotodiodi e i fototransistor.



Questo funzionamento può essere realizzato in base al fatto che, **quando un elettrone assorbe un fotone (cioè quando assorbe luce), accresce la sua energia e può salire dalla banda energetica di valenza a quella di conduzione**, entrando a far parte di un processo di conduzione elettrica. Ogni elettrone che sale nella banda di conduzione, inoltre, lascia “scoperta” una lacuna nella banda di valenza.

Nei fotorivelatori una giunzione pn, **inversamente polarizzata**, è percorsa, in **assenza di luce**, da una debole “**corrente di buio**” **inversa** (avente l'intensità di qualche **decina di nA**).

Quando la giunzione viene colpita da una radiazione luminosa di opportuna frequenza, la corrente inversa aumenta di intensità, in modo proporzionale all'intensità della luce incidente.

La frequenza di lavoro del fotorivelatore dipende dal materiale col quale è realizzato il componente, in quanto è legata all'entità del gap fra banda di conduzione e di valenza, gap che gli elettroni devono superare grazie all'azione di fotoni dotati di sufficiente energia.

Possiamo notare che il fenomeno su cui si basa la conversione ottico-elettrica che avviene nel ricevitore è esattamente l'inverso del processo sul quale si basa la conversione elettrico-ottica del trasmettitore.

Mentre la generazione del segnale luminoso a partire dal segnale elettrico si basa sulla ricombinazione di una coppia lacuna-elettrone (che “scompare” dal processo conduttivo, emettendo un fotone), la “riconversione” del segnale ottico in elettrico si fonda sulla **separazione di una coppia elettrone-lacuna, con conseguente salita dell'elettrone nella banda di conduzione a causa dell'assorbimento di un fotone** (cioè di luce) da parte dell'elettrone.

PARAMETRI DI UN FOTORIVELATORE

- **RESPONSIVITA' "R"**

E' il **rapporto** fra la **corrente generata dal componente** e la **potenza ottica incidente**.
Quindi, più è alta la responsività, maggiore è l'efficienza del fotorivelatore

$$R = \frac{I_u}{P. \text{ ottica}} \quad [\text{A/W}]$$

La responsività è direttamente proporzionale alla lunghezza d'onda secondo la relazione approssimata:

$$R = \frac{\lambda[nm]}{1240}$$

- **EFFICIENZA QUANTICA "η"**

$$\eta = \frac{\text{numero di coppie lacuna - elettrone prodotte}}{\text{numero di fotoni incidenti}}$$

- **SENSIBILITA' (SENSITIVITY)**

E' la **minima potenza ottica media**, espressa¹ in dBm, che risulta **necessaria per avere un certo BER** (tasso di errore) e tipicamente per avere:

BER = 10⁻⁹ per un intero collegamento

BER = 10⁻¹¹ per una tratta

Ricordiamo che il BER è definito come:

BER = numero di bit errati in ricezione / numero di bit totalmente trasmessi

$$P_{dBm} = 10 \cdot \log_{10} \left(\frac{P \text{ [mW]}}{1 \text{ mW}} \right)$$

- **CORRENTE DI USCITA**

E' la corrente prodotta dal fotorivelatore

$$I_u = R \cdot G \cdot P_{OTTICA}$$

dove:

R = responsività

G = guadagno di corrente

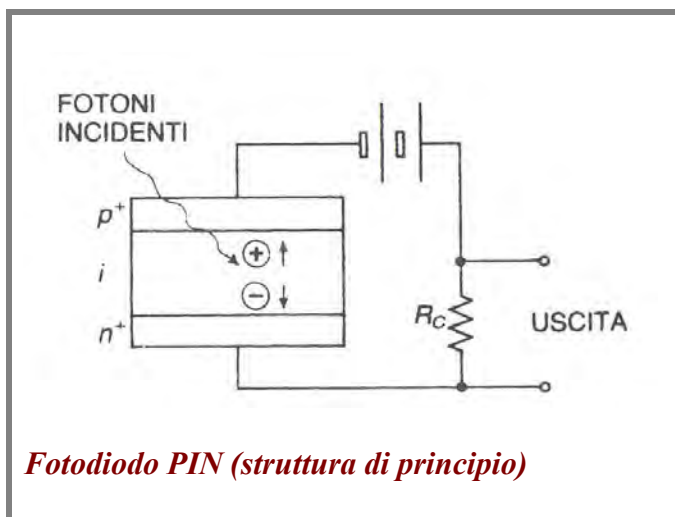
G = 1 per fotorivelatori PIN

G = 100-200 per fotorivelatori APD

FOTODIODO PIN (Positivo-Intrinseco-Negativo)

Come ogni fotorivelatore, lavora in **polarizzazione inversa** e presenta, sul circuito esterno, una **corrente proporzionale alla potenza ottica assorbita**.

La sua particolarità è di presentare, fra la zona p e la zona n, una zona intrinseca, cioè non drogata, oppure pochissimo drogata, in modo da avere un'area più grande in grado di assorbire fotoni e generare corrente.

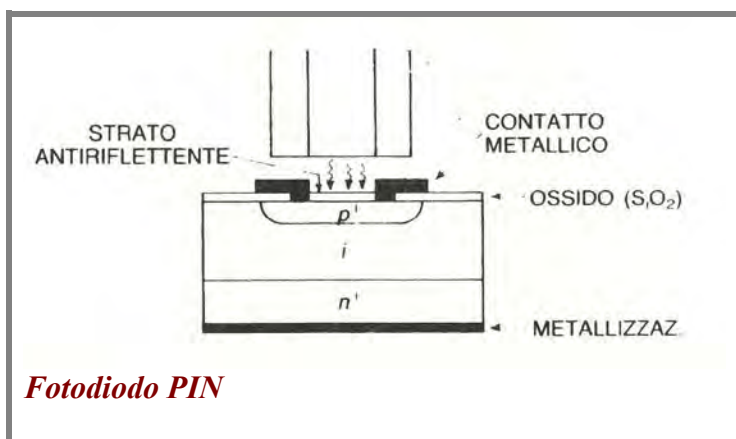


Nella zona intrinseca si crea un forte campo elettrico inverso in grado di favorire la separazione delle coppie elettrone-lacuna.

Il pregio dei PIN sono la semplicità costruttiva, il basso costo e le **basse tensioni inverse (20V-50V) che sono sufficienti al funzionamento del dispositivo**.

L'inconveniente fondamentale è la sua **bassa velocità di risposta**, dovuta al tempo necessario agli elettroni ad attraversare il dispositivo.

Per ridurre il tempo di risposta si dovrebbe rimpicciolire la zona intrinseca, ma ciò determinerebbe una diminuzione della sensibilità.



FOTODIODO APD (Avalanche Photo Diode o fotodiodo a valanga)

Come il precedente fotorivelatore, lavora in polarizzazione inversa e presenta, sul circuito esterno, una corrente proporzionale alla potenza ottica assorbita.

A differenza dei PIN però richiede l'applicazione di **forti tensioni inverse (100V-300V)** che lo portano a funzionare (analogamente ai diodi Zener) sfruttando l'effetto valanga, che determina una moltiplicazione delle cariche libere per ionizzazione da impatto. Conseguenza dell'effetto valanga è che l'APD presenta un **guadagno di corrente** (di valore tipicamente compreso fra 10 e 100), che consente al dispositivo di rispondere anche a segnali ottici di bassa potenza e quindi di avere sensibilità e velocità di risposta maggiori di quelle del PIN. Nell'uso dell'APD bisogna fare attenzione

- a non superare la tensione di breakdown (poco al di sotto della quale il dispositivo è polarizzato), il che porterebbe alla distruzione del dispositivo
- a non superare un certo valore del guadagno, perché la moltiplicazione avviene non solo per la corrente di segnale, ma anche per quella di rumore, per cui, aumentando il guadagno al di là di un certo limite, si corre il rischio di far diminuire eccessivamente il rapporto segnale/rumore.

Il pregio dell'APD è la capacità di rivelare, grazie al guadagno, segnali ottici molto deboli e quindi una **sensibilità maggiore** (di 10dB-20dB) di quella del PIN. L'inconveniente fondamentale è la complessità circuitale derivante essenzialmente dalle tensioni di polarizzazione più elevata.

PRESTAZIONI DEI FOTORIVELATORI

PARAMETRO	SIMBOLO e unità di misura	PIN	APD
RESPONSIVITA'	R [A/W]	0,5 ÷ 1	
SENSIBILITA'	[dBm]	-30 ÷ -40	-40 ÷ -50
TENSIONE di POLARIZZAZIONE	[V]	20 ÷ 50	>100
CORRENTE DI BUIO	[nA]	< 0,1	100 circa
COSTO		BASSO	ALTO

CAPITOLO 13

CODIFICA, ARCHIVIAZIONE ED ELABORAZIONE DELL'INFORMAZIONE DIGITALE

COMPLEMENTI

***RAPPRESENTAZIONE DEI NUMERI
NEI SISTEMI DIGITALI***

***NUMERI (ANCHE NON INTERI)
IN VIRGOLA MOBILE (FLOATING POINT)***

***NUMERI
IN VIRGOLA MOBILE (FLOATING POINT)
NORMALIZZATI***

***NUMERI NON INTERI IN VIRGOLA MOBILE NORMALIZZATI
SECONDO LO STANDARD IEEE 754***

***RAPPRESENTAZIONE “EFFETTIVA” PREVISTA DALLO STANDARD
IEEE 754***

***SISTEMA DI NUMERAZIONE OTTALE E CONVERSIONI
(CODICE OTTALE)***

***SISTEMA DI NUMERAZIONE ESADECIMALE (HEX)
(CODICE ESADECIMALE) E CONVERSIONI***

RAPPRESENTAZIONE DEI NUMERI NEI SISTEMI DIGITALI

NUMERI IN SEMPLICE PRECISIONE

Sono numeri:

- INTERI RELATIVI
- rappresentati IN COMPLEMENTO A DUE
- su DUE BYTE (16 bit, corrispondenti a 1word)

numero in semplice precisione

BYTE	BYTE
------	------

NUMERI IN DOPPIA PRECISIONE

Sono numeri:

- INTERI RELATIVI
- rappresentati IN COMPLEMENTO A DUE
- su QUATTRO BYTE (32 bit, corrispondenti a 2 word)

numero in doppia precisione

BYTE	BYTE	BYTE	BYTE
------	------	------	------

NUMERI (ANCHE *NON* INTERI) IN VIRGOLA MOBILE (FLOATING POINT)

Per VIRGOLA MOBILE o FLOATING POINT si intende un formato di rappresentazione dei numeri NON INTERI nel quale sia utilizzata la NOTAZIONE ESPONENZIALE. Cioè nel quale sia presente, come fattore, la base del sistema di numerazione adottato, elevata a un esponente intero (positivo o negativo).

Esistono pertanto infinite rappresentazioni in virgola mobile di un numero non intero.

Esempi

numero decimale NON intero	alcune possibili rappresentazioni IN VIRGOLA MOBILE
$27,8_{10}$	$0,278 \cdot 10^2$ $0,0278 \cdot 10^3$ $278 \cdot 10^{-1}$ $27,8 \cdot 10^0 \rightarrow 27,8$

Ribadiamo che sono “**in virgola mobile**” tutte le possibili **rappresentazioni** di un numero in base B nelle quali sia presente, come fattore, una potenza della base B.

Per uniformare l’elaborazione aritmetica su differenti macchine, l’ente normativo statunitense **IEEE** ha fissato lo standard “**754**” per il **calcolo in virgola mobile**¹.
“IEEE 754” è lo **standard più diffuso** nel campo del **calcolo automatico**.

¹ (IEEE Standard for Binary Floating-Point Arithmetic: **ANSI/IEEE Std 754-1985** o anche **IEC 60559:1989**, Binary floating-point arithmetic for microprocessor systems).

NUMERI IN VIRGOLA MOBILE (FLOATING POINT) NORMALIZZATI

Siccome sono possibili infinite rappresentazioni in virgola mobile di uno stesso numero, risulta opportuno sceglierne una come riferimento, in modo da avere una modalità standard di rappresentazione.

Per esempio, si potrebbe scegliere, come rappresentazione di riferimento in virgola mobile, una rappresentazione del tipo:

$-0,278... \cdot 10^{-3}$

caratterizzata da:

- parte intera (parte a sinistra dello zero) sempre nulla
- prima cifra a destra dello zero sempre diversa da zero

Oppure si può scegliere (come fa lo standard IEEE per la maggior parte dei numeri) una rappresentazione del tipo:

$-2,325... \cdot 10^{-2}$

caratterizzata da:

- **parte intera** (parte a **sinistra dello zero**) formata da **una sola cifra**, sempre diversa da zero
- prima cifra a destra dello zero che può essere indifferentemente nulla o non nulla

Il tipo di rappresentazione in virgola mobile che viene **scelto** (arbitrariamente) come **riferimento** è detto “rappresentazione **NORMALIZZATA** in virgola mobile”.

NUMERI NON INTERI IN VIRGOLA MOBILE (FLOATING POINT) NORMALIZZATI SECONDO LO STANDARD IEEE 754

Sempre con riferimento ai numeri anche non interi, lo standard prevede una rappresentazione BINARIA del tipo:

$$1, 000101... * 2^5$$

(dove il bit 1 a sinistra della virgola è la parte intera
e la stringa 000101, a destra della virgola, è la parte frazionaria)

La forma standard è quindi caratterizzata da:

- **parte intera** (parte a **sinistra dello zero**) sempre costituita da un **bit "1"**
- prima cifra a destra dello zero indifferentemente "0" o "1"

RAPPRESENTAZIONE "EFFETTIVA" PREVISTA DALLO STANDARD

La struttura prima illustrata non è ciò che, secondo lo standard, viene effettivamente rappresentato (e memorizzato nei dispositivi digitali).

L'IEEE 754 prevede infatti che:

- NON venga per nulla rappresentata la PARTE INTERA, in quanto vale sempre "1" nella normalizzazione adottata
- NON venga per nulla rappresentata la BASE "2", in quanto è implicito che la base di un numero binario sia "2"
- non venga rappresentato il vero esponente, bensì l'esponente al quale sia stato sommato il codice binario di 127 (si parla perciò di rappresentazione "in eccesso 127" o di rappresentazione "con offset").

Dato quindi il numero:

$34,5_{10} = 1,000101... \cdot 2^5$

ciò che viene effettivamente rappresentato è:



Riguardo all'**esponente**:

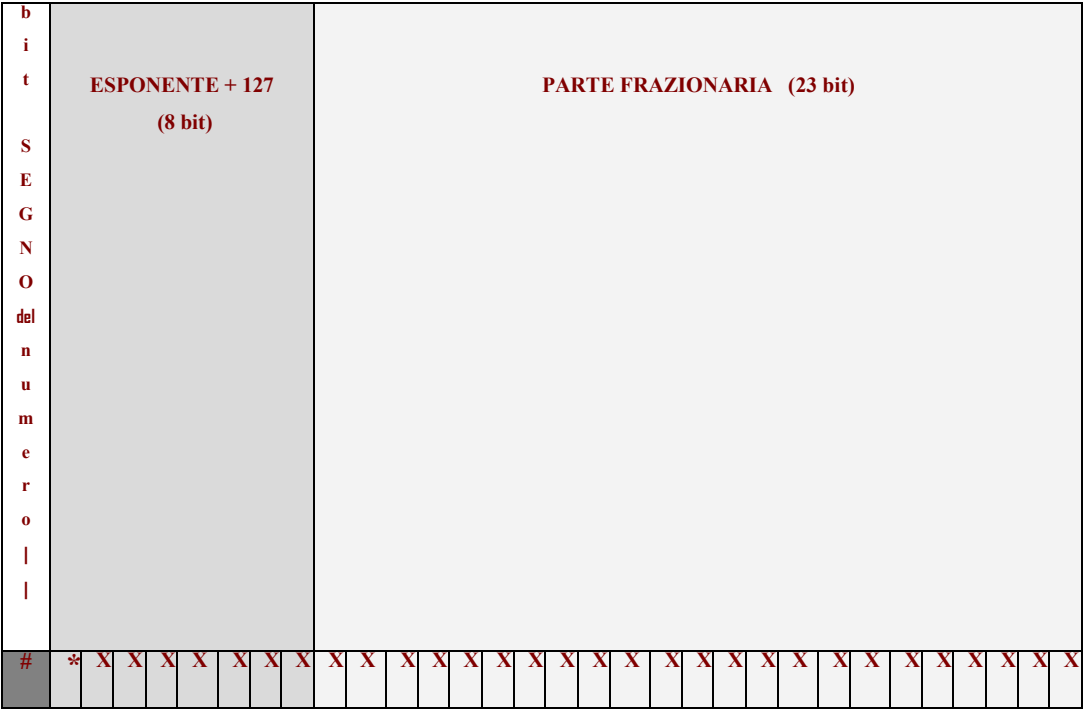
pesi	128	64	32	16	8	4	2	1
esponente di partenza (5)						1	0	1
offset 127 da sommare		1	1	1	1	1	1	1
esponente rappresentato "in eccesso 127" (5+127=132)	1	0	0	0	0	1	0	0

Se l'**esponente di partenza è positivo o nullo**, l'esponente effettivamente rappresentato vale almeno 128 e quindi il bit di peso 128 dell'esponente è **"1"**

Se invece ***l'esponente di partenza è negativo***, l'esponente effettivamente rappresentato è minore di 128 e quindi il bit di peso 128 dell'esponente è **"0"**.

In generale, quindi, ciò che viene effettivamente rappresentato è:

RAPPRESENTAZIONE “EFFETTIVA” PREVISTA DALLO STANDARD IEEE 754



LEGENDA

- #** → questo bit dà informazioni sul segno del numero:
0=positivo
1=negativo
- *** → questo bit dà informazioni sul segno dell’esponente:
1=positivo
0=negativo
- X** → bit con significato numerico

SISTEMA DI NUMERAZIONE OTTALE (CODICE OTTALE) E CONVERSIONI

Il sistema di numerazione ottale esprime tutti i numeri servendosi di otto simboli:

0, 1, 2, 3, 4, 5, 6, 7.

Siccome utilizza **otto cifre o simboli**, si dice è un sistema di numerazione “**in base 8**”.

Ovviamente un numero ottale non contiene le cifre decimali 8 e 9.

Analogamente al decimale e al binario, quello ottale è un sistema posizionale, nel senso che il valore (o peso) di una cifra dipende dalla posizione che la cifra ha all'interno del numero.

Per esempio, nel numero ottale 23, la cifra a destra (3) ha peso 8^0 , mentre la cifra a sinistra (2) ha peso 8^1 .

Rappresentiamo “graficamente” i pesi delle cifre di un numero ottale.

PESI delle cifre ottali, espresse come potenze della base 8		8^3	8^2	8^1	8^0
PESI delle cifre ottali		512	64	8	1
Esempio di numero ottale: 23				2	3

Per non creare ambiguità con gli altri sistemi di numerazione, si usa accompagnare i numeri ottali con il pedice “8”.

Esempi: 11111_8 , 100_8 , 343_8 .

Inoltre i numeri ottali non si leggono come quelli decimali.

Per esempio 100_8 si legge “uno-zero-zero” (e non “cento”).

CONVERSIONE DA DECIMALE A OTTALE

Vediamo ora come si converte un numero decimale (per esempio 74) in ottale, ossia come si determina il codice ottale di un numero decimale.

Si opera per successive “divisioni intere” per la base 8, conservando i resti che rappresenteranno le cifre della codifica ottale del numero.

La conversione prevede i seguenti passi:

- il numero da convertire (74) si divide per 8; il resto (2) rappresenterà “l’ultima cifra”, cioè la cifra più a destra del numero ottale
- il quoziente (9) si divide per 8; il resto (1) rappresenterà la “penultima cifra” del numero ottale
- il quoziente (1) è minore della base 8, per cui la conversione è terminata; la cifra “1” sarà la prima cifra (la terzultima) del numero ottale

				2
--	--	--	--	---

		1	2
--	--	---	---

	1	1	2
--	---	---	---

Il procedimento seguito può anche essere schematizzato mediante una tabella:

DIVIDENDO	DIVISORE	QUOZIENTE	RESTI (cifre del codice ottale del numero)
74	8	9	2
9	8	1 quoziente < 8 → FINE La cifra 1 rappresenterà l’ultimo resto e quindi la prima cifra del codice ottale	1
			1

La codifica ottale di 74_{10} è quindi 112_8 .

CONVERSIONE DA OTTALE A DECIMALE

Dato il numero ottale da convertire in decimale (per esempio: 203_8), lo incolonniamo in corrispondenza della sequenza dei pesi ottali:

Numero ottale da convertire			2	0	3
PESI delle cifre ottali), espresse come potenze della base 8	8^4	8^3	8^2	8^1	8^0
PESI delle cifre ottali	4096	512	64	8	1

Moltiplichiamo quindi ogni cifra ottale (diversa da zero) per il peso corrispondente e facciamo la somma dei prodotti:

$$X_{10} = 3 \cdot 1 + 2 \cdot 64 = 131_{10}$$

CONVERSIONI TRA SISTEMI OTTALE E BINARIO

OTTALE → BINARIO

Ad ogni cifra del numero ottale si associa la sua codifica binaria a tre bit.
L'insieme dei bit così ottenuti è il codice binario del numero ottale di partenza.

Esempio

Si trovi la codifica binaria di 527_8

Numero ottale da convertire		5	2	7
Codifica binaria delle singole cifre		101	010	111

Il codice binario di 527_8 è quindi 101**0**10111

BINARIO → OTTALE

Si suddivide, a partire da destra, il numero binario di partenza in gruppi di tre bit.
A ogni gruppo di tre bit si associa la corrispondente cifra ottale.
L'insieme delle cifre ottali così ottenute è il codice ottale del numero binario di partenza.

Esempio

Si trovi la codifica ottale di 11100101_2

Numero binario da convertire		11	100	101
Codifica ottale dei singoli gruppi di tre bit		3	4	5

Il codice ottale di 11100101_2 è quindi 345_8 .

SISTEMA DI NUMERAZIONE ESADECIMALE (HEX) (CODICE ESADECIMALE) E CONVERSIONI

Il sistema esadecimale è un sistema di numerazione “**in base 16**” e utilizza le sedici cifre (o simboli):

0, 1, 2, 3, 4, 5, 6, 7, 8, 9, A, B, C, D, E, F

Dove:

A=10₁₀

B=11₁₀

C=12₁₀

D=13₁₀

E=14₁₀

F=15₁₀

Analogamente al decimale, al binario e all’ottale, è un sistema posizionale nel senso che il valore (o peso) di una cifra dipende dalla posizione che la cifra ha all’interno del numero.

Per esempio, nel numero esadecimale 1A, la cifra a destra (A) ha peso 16⁰, mentre la cifra a sinistra (1) ha peso 16¹

Rappresentiamo “graficamente” i pesi delle cifre di un numero ottale.

PESI delle cifre esadecimali, espresse come potenze della base 16		16³	16²	16¹	16⁰
PESI delle cifre esadecimali		4096	256	16	1
Esempio di numero esadecimale:				1	A

Per non creare ambiguità con gli altri sistemi di numerazione, si usa contrassegnare i numeri esadecimali con il pedice “16” oppure con la lettera “h”.

Esempi:

Numeri esadecimali: 323₁₆, 323h, AF0₁₆, 1Ah.

CONVERSIONE DA DECIMALE A ESADECIMALE

Vediamo ora come si converte un numero decimale (per esempio 781) in esadecimale, ossia come si determina il codice hex di un numero decimale.

Si opera per successive “divisioni intere” per la base 16, conservando i resti, che rappresenteranno le cifre della codifica esadecimale del numero.

Se il resto è maggiore di 9, ad esso va sostituita la corrispondente cifra alfabetica.

Per esempio, al resto 11, va sostituita la lettera B.

La conversione prevede i seguenti passi:

- il numero da convertire (781) si divide per 16; il resto (13) corrisponde alla cifra D e rappresenterà “l’ultima cifra”, cioè la cifra più a destra del numero esadecimale
- il quoziente (48) si divide per 16; il resto (0) rappresenterà la “penultima cifra” del numero esadecimale
- il quoziente (3) è minore della base 16, per cui la conversione è terminata; la cifra “3” sarà la prima cifra (la terzultima) del numero esadecimale

				D
--	--	--	--	---

		0	D	
--	--	---	---	--

	3	0	D	
--	---	---	---	--

Il procedimento seguito può anche essere schematizzato mediante una tabella:

DIVIDENDO	DIVISORE	QUOZIENTE	RESTI (cifre del codice esadecimale del numero)
781	16	48	13 → D
48	16	3 quoziente < 16 → FINE La cifra 3 rappresenterà l’ultimo resto e quindi la prima cifra del codice esadecimale	0
			3

Il codice esadecimale del numero decimale 781 è 30D.

CONVERSIONE DA ESADECIMALE A DECIMALE

Dato il numero esadecimale da convertire in decimale (per esempio: 30D), lo incolonniamo in corrispondenza della sequenza dei pesi esadecimali.

Ad ogni cifra costituita da una lettera, sostituiamo l'equivalente decimale (per esempio, alla cifra D, sostituiamo il numero 13)

Numero esadecimale da convertire			3	0	D
Sostituzione dell'equivalente decimale alle lettere			3	0	13
PESI delle cifre esadecimali espresse come potenze della base 16	16^4	16^3	16^2	16^1	16^0
PESI delle cifre esadecimali	65536	4096	256	16	1

Moltiplichiamo quindi ogni cifra esadecimale (diversa da zero) per il peso corrispondente e facciamo la somma dei prodotti:

$$X_{10} = 13 \cdot 1 + 3 \cdot 256 = 781_{10}$$

CONVERSIONI TRA SISTEMI ESADECIMALE E BINARIO

ESADECIMALE → BINARIO

Ad ogni cifra del numero hex si associa la sua codifica binaria a quattro bit.
L'insieme dei bit così ottenuti è il codice binario del numero hex di partenza.

Esempio

Si trovi la codifica binaria di $A02_{16}$

Numero hex da convertire		A (10)	0	2
Codifica binaria delle singole cifre		1010	0000	0010

Il codice binario di $A02_{16}$ è quindi **101000000010**

BINARIO → ESADECIMALE

Si suddivide, a partire da destra, il numero binario di partenza in gruppi di quattro bit.
A ogni gruppo di quattro bit si associa la corrispondente cifra hex.
L'insieme delle cifre hex così ottenute è il codice hex del numero binario di partenza.

Esempio

Si trovi la codifica ottale di 111100101_2

Numero binario da convertire		1	1110	0101
Codifica hex dei singoli gruppi di tre bit		1	14 (E)	5

Il codice hex di 111100101_2 è quindi **1E5₁₆**.

CAPITOLO 16

MEMORIE DIGITALI

COMPLEMENTI

MEMORIE ROM A SELEZIONE UNIDIMENSIONALE: INCONVENIENTI

MEMORIE ROM A SELEZIONE BIDIMENSIONALE

MEMORIE ROM A SELEZIONE UNIDIMENSIONALE: INCONVENIENTI

Le memorie a indirizzamento unidimensionale si rivelano inadeguate quando la quantità di parole da memorizzare non è piccola.

Siccome infatti le righe devono essere selezionate da un indirizzo binario, esse vengono poste all'uscita di un decoder, che è costituito da tante porte AND quante sono le uscite (e quindi le righe).

E' evidente quindi che già con un numero di parole memorizzate non eccessivo, la quantità di porte richieste è inaccettabile.

MEMORIE ROM A SELEZIONE BIDIMENSIONALE

In una memoria, il **numero complessivo di bit di indirizzo necessari**¹, e quindi il numero di linee di selezione esterne, è determinato dal **numero di parole (locazioni)** che si vogliono in uscita. Se, per esempio, si vuole una memoria con **512=2⁹ parole (locazioni)**, saranno necessari **9 bit di indirizzo**, dove $9 = \log_2(512)$.

Di questi bit di indirizzo (nelle memorie a indirizzamento bidimensionale) solo un certo numero (per esempio 6 su 9) sarà utilizzato per selezionare le righe, mentre i rimanenti (3 nel caso in esame) costituiranno i bit di indirizzo dei multiplexer collegati alle colonne, bit di indirizzo che saranno comuni a tutti i multiplexer.

Per realizzare una memoria bidimensionale di **64 parole** ($n_{LOC} = 64$) di **4 bit ciascuna** ($n_{bit-loc} = 4$), per una **capacità totale 256 bit** ($n_{BIT} = 256$), risulteranno necessari **6 bit di indirizzo**, cioè un numero di bit di indirizzo dato da $\log_2(n_{LOC})$ e quindi da $\log_2(64)$. Per l'organizzazione di una tale memoria può procedere in questo modo:

- Si sceglie il **numero delle righe** (per esempio **$n_{RIGHE} = 8 = 2^3$**), le quali richiederanno un numero di linee di selezione di riga pari a **$\log_2(n_{RIGHE})$** . Nel nostro esempio le linee di selezione di riga sono 3.
- Di conseguenza il **numero delle colonne** n_{COL} sarà il **rapporto fra il numero totale di bit e il numero delle righe**:

$$n_{COL} = \frac{n_{BIT}}{n_{RIGHE}}$$

Nel nostro caso le colonne saranno:

$$n_{COL} = \frac{256}{8} = 32$$

- L'insieme delle colonne (ciascuna delle quali ha "larghezza" di un solo bit) viene suddiviso in un **numero n_{GR} di gruppi di colonne**, uguale al **numero di bit delle parole** (nel nostro caso $n_{GR} = 4$). Tutte le linee (8 nel nostro caso) relative a ciascun gruppo di colonne vengono applicate allo stesso multiplexer e solo ad esso.
- Il **numero di multiplexer** deve essere uguale al **numero di bit delle parole** o locazioni e quindi uguale al numero di gruppi di colonne. Nell'esempio in esame i MUX dovranno essere quattro e quindi ciascun multiplexer avrà 8 ingressi ($32/4$), ciascuno corrispondente a una delle colonne del gruppo. Per individuare ciascun ingresso, saranno necessari 3 bit.

¹ Per l'individuazione sia delle righe che delle colonne

Ciò implica che, a ogni parola o locazione di 4 bit corrisponde un indirizzo di 6 bit

$A_0 A_1 A_2 A_3 A_4 A_5$:

Tre di questi bit ($A_0 A_1 A_2$) formano l'indirizzo di riga.

Gli altri tre ($A_3 A_4 A_5$) formano l'indirizzo delle colonne all'interno di ciascun gruppo di 8 colonne, gruppo facente capo a un multiplexer.

Ogni singola combinazione dei tre bit di indirizzo di riga $A_0 A_1 A_2$ individua la riga di appartenenza della parola.

Ogni singola combinazione dei tre bit di indirizzo di colonna $A_3 A_4 A_5$ individua una colonna in ognuno dei 4 gruppi di 8 colonne.

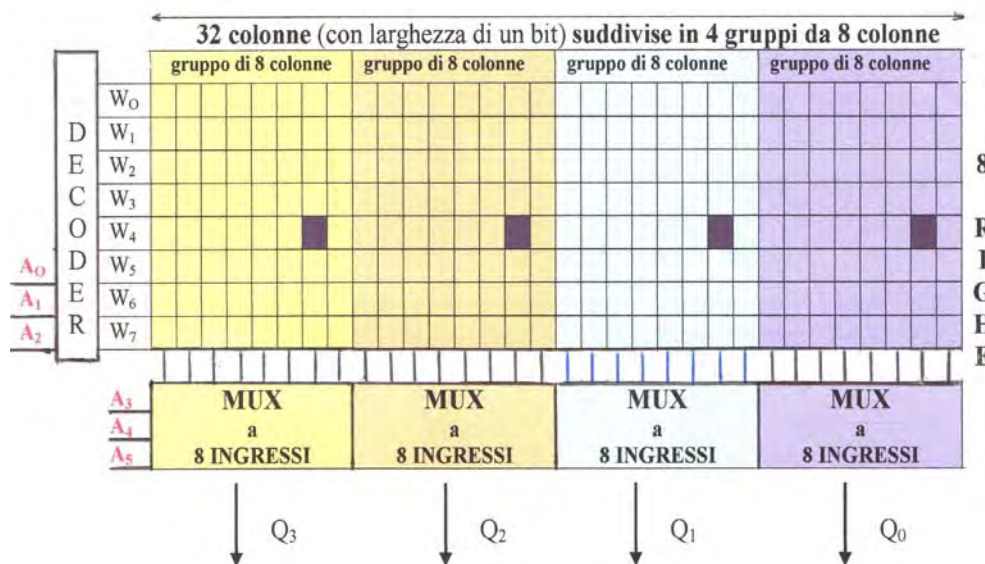
Per esempio $A_3 A_4 A_5 = 101$ individua la quinta colonna di ogni gruppo.

ROM

a 256 bit (con parole di 4 bit)

in 8 righe

in 32 colonne suddivise in gruppi di 8

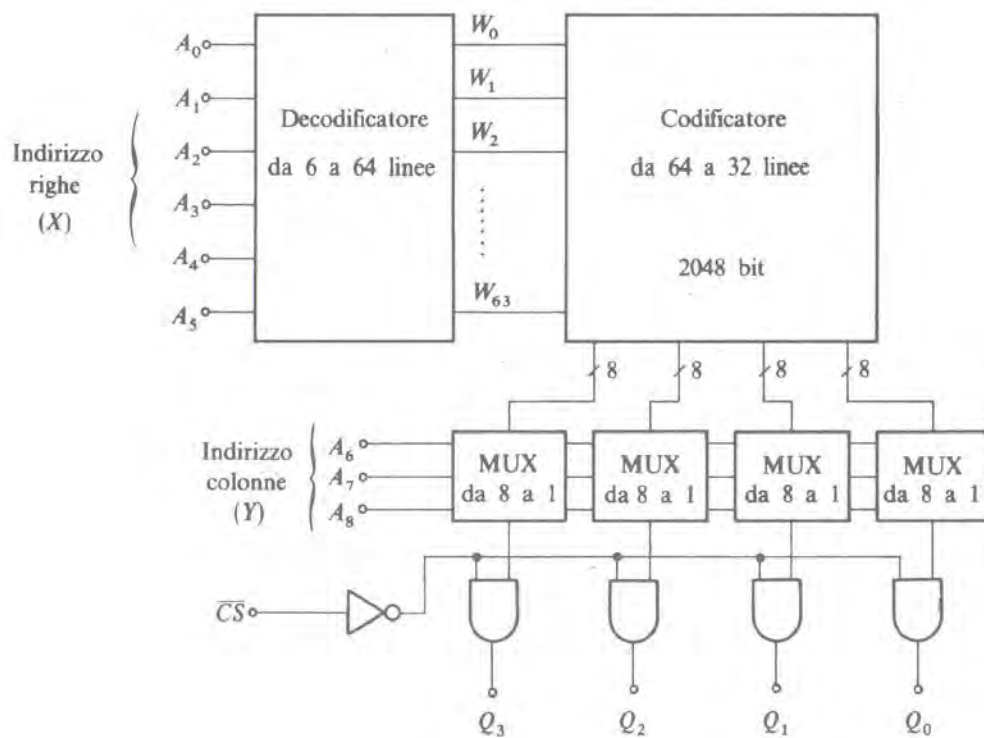


$A_0 A_1 A_2$ = indirizzi delle RIGHE $W_0, W_1, W_2, W_3, W_4, W_5, W_6, W_7$

$A_3 A_4 A_5$ = indirizzi di COLONNA all'interno di ciascun gruppo di 8 colonne

Nella figura le celle evidenziate in colore scuro costituiscono la parola (locazione) individuata dall'indirizzo di riga $A_0 A_1 A_2 = 100$ e dall'indirizzo di colonna $A_3 A_4 A_5 = 111$ e quindi dall'indirizzo complessivo

$A_0 A_1 A_2 A_3 A_4 A_5 = 100111$



ROM a 2048 bit (512x4) con indirizzamento bidimensionale

CAPITOLO 18

ANALISI IN FREQUENZA DEI SEGNALE

COMPLEMENTI

PROPRIETA' DI SIMMETRIA DEI SEGNALE PERIODICI E SVILUPPO IN SERIE DI FOURIER (SEGNALE PARI, DISPARI, EMISIMMETRICI)

SVILUPPI IN SERIE DI FOURIER DI ONDE QUADRE E RETTANGOLARI

ONDE RETTANGOLARI (NON QUADRE) CON DUTY-CYCLE DIVERSO DAL 50%

IL SEGNALE APERIODICO COME "LIMITE" DEL SEGNALE PERIODICO

SEGNALE APERIODICI SCOMPOSIZIONE NEL TEMPO E SPETTRO

*"SCOMPOSIZIONE NEL TEMPO" DEL SEGNALE APERIODICO:
INTEGRALE DI FOURIER O ANTITRASFORMATATA DI FOURIER*

*SPETTRO DEI SEGNALE APERIODICI: TRASFORMATATA DI FOURIER O F-
TRASFORMATATA*

*SPETTRO E TRASFORMATATA DI FOURIER DI UN SINGOLO IMPULSO
RETTANGOLARE*

PROPRIETA' DI SIMMETRIA DEI SEGNALE PERIODICI E SVILUPPO IN SERIE DI FOURIER

Se l'andamento temporale di un segnale periodico presenta determinate proprietà di simmetria, l'espressione dello sviluppo in serie di Fourier assume espressioni particolari o semplificate. Di conseguenza anche lo spettro può presentare caratteristiche specifiche.

Le proprietà che ci interessano sono:

- Simmetria **PARI**
- Simmetria **DISPARI**
- **EMISIMMETRIA**

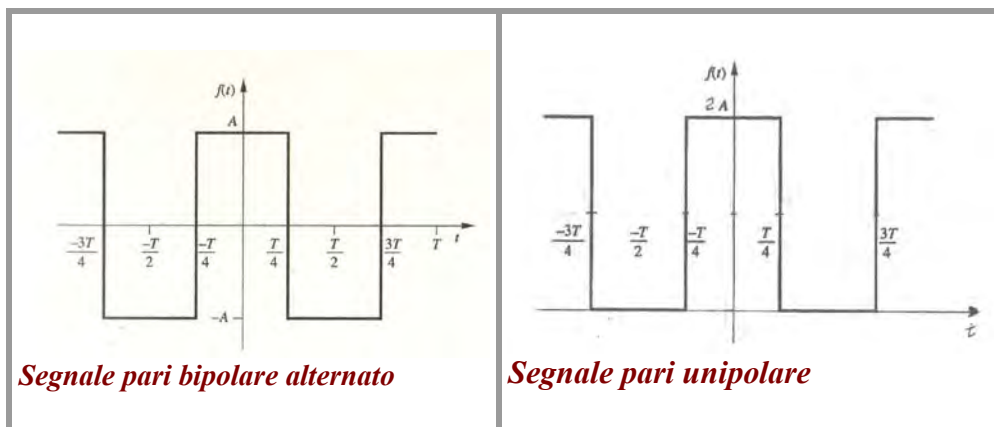
SEGNALE PARI

Un segnale $v(t)$ si dice “pari” quando è **simmetrico rispetto all'asse verticale**, cioè quando verifica la condizione:

$$v(t) = v(-t)$$

I segnali “pari” possono essere unipolari o bipolari, alternati o non alternati e il loro valore medio nel periodo può essere nullo o diverso da zero.

I segnali pari possono essere sviluppati in **serie di Fourier** di soli **coseni a fase nulla**.



SEGNALE DISPARI

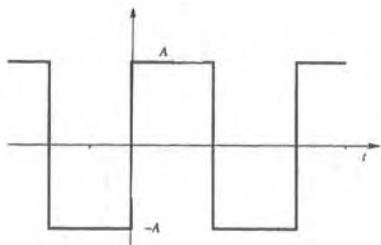
Un segnale $v(t)$ si dice “dispari” quando è **simmetrico rispetto all’origine del riferimento**, cioè quando verifica la condizione:

$$v(t) = -v(-t)$$

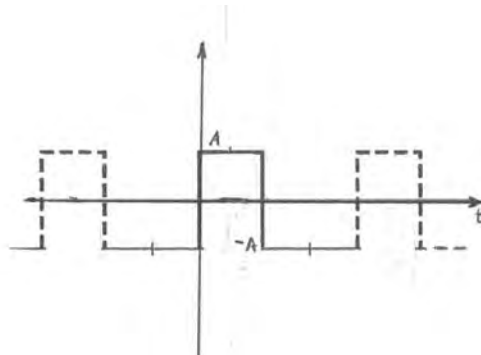
I segnali “dispari” devono necessariamente essere **bipolari**, ma possono essere alternati o non alternati e il loro valore medio nel periodo può essere nullo o diverso da zero.

I segnali dispari possono essere sviluppati in **serie di Fourier** di **solli seni a fase nulla**.

Un segnale dispari può essere reso pari (e viceversa) mediante una traslazione orizzontale lungo l’asse dei tempi.



*Segnale dispari alternato:
onda quadra (duty-cycle del 50%)*



*Segnale dispari non alternato:
onda rettangolare
(duty-cycle diverso dal 50%)*

SEGNALE EMISIMMETRICO

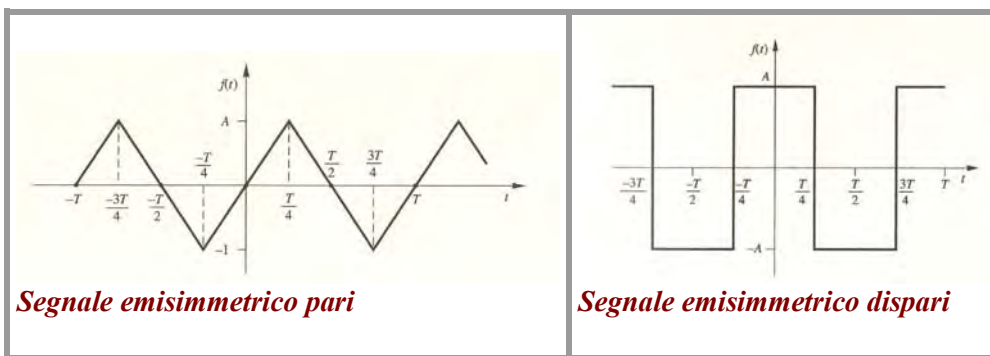
Un segnale $v(t)$ è “emisimmetrico” quando **la semionda positiva, ribaltata rispetto all’asse dei tempi e traslata in orizzontale, risulta sovrapponibile alla semionda positiva**, cioè quando verifica la condizione:

$$v(t) = -v\left(t + \frac{T}{2}\right)$$

I segnali “emisimmetrici” possono essere dispari o pari, ma devono necessariamente essere **bipolari alternati** e quindi con **valore medio nullo** nel periodo.

Lo **sviluppo in serie di Fourier** di un segnale emisimmetrico **presenta le sole armoniche dispari** (la prima, la terza, la quinta ecc.), mentre tutte le armoniche di ordine pari risultano nulle.

L’**emisimmetria** di un segnale **non** può essere ottenuta mediante **traslazione orizzontale** lungo l’asse dei tempi.

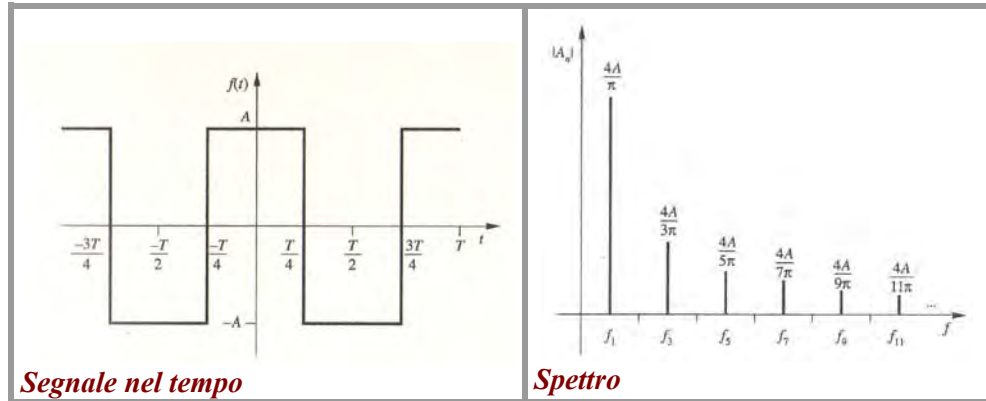


SVILUPPI IN SERIE DI FOURIER DI ONDE QUADRE E RETTANGOLARI

ONDA QUADRA con valore picco-picco 2A

BIPOLORE ALTERNATA

PARI



L'ONDA è QUADRA (duty- cycle del 50%), quindi è EMISIMMETRICA, per cui il suo sviluppo in serie di Fourier contiene sole armoniche dispari

L'ONDA è PARI, quindi presenta uno sviluppo con sole armoniche in coseno a fase nulla

$$v(t) = \frac{4A}{\pi} \cdot \cos(\omega_0 \cdot t) - \frac{4A}{3 \cdot \pi} \cdot \cos(3 \cdot \omega_0 \cdot t) + \frac{4A}{5 \cdot \pi} \cdot \cos(5 \cdot \omega_0 \cdot t) - \frac{4A}{7 \cdot \pi} \cdot \cos(7 \cdot \omega_0 \cdot t) + \dots$$

Si noti che le armoniche di ordine 3, 7, 11 sono col segno “-”

In forma più compatta:

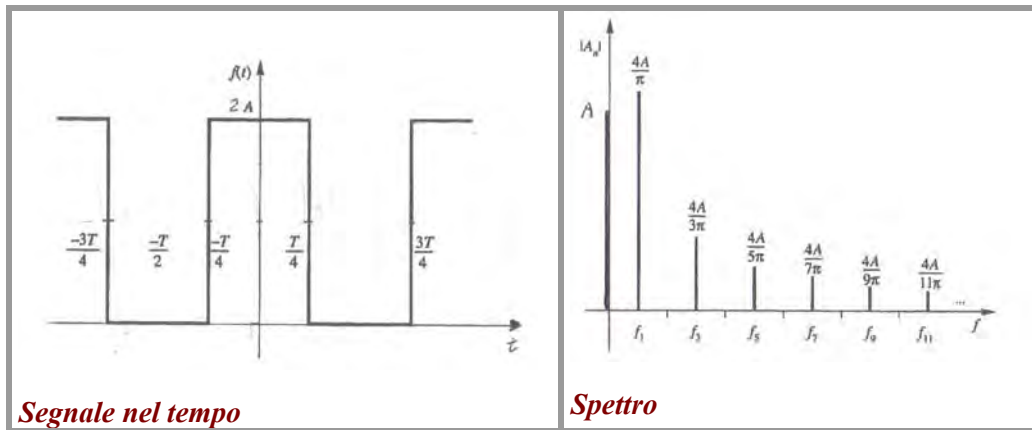
$$v(t) = \frac{4A}{\pi} \cdot \sum_{m=0}^{\infty} (-1)^m \cdot \frac{\cos[(2m+1) \cdot \omega_0 \cdot t]}{(2m+1)}$$

Dove:

- $2m+1 = n$ = ordine dell'armonica
- $m=0 \rightarrow$ prima armonica
- $m=1 \rightarrow$ terza armonica
- $m=2 \rightarrow$ quinta armonica

ONDA QUADRA con valore picco-picco $2A$ UNIPOLARE PARI

Si ottiene dalla precedente traslandola verso l'alto di A , cioè sommandole una componente continua (valor medio in T) di A



Lo sviluppo in serie di Fourier rimane **invariato (rispetto al caso precedente)**, tranne che per l'**aggiunta del valore medio**.

$$v(t) = V_{medio} + \frac{4A}{\pi} \cdot \cos(\omega_0 \cdot t) - \frac{4A}{3 \cdot \pi} \cdot \cos(3 \cdot \omega_0 \cdot t) + \frac{4A}{5 \cdot \pi} \cdot \cos(5 \cdot \omega_0 \cdot t) - \frac{4A}{7 \cdot \pi} \cdot \cos(7 \cdot \omega_0 \cdot t) + \dots$$

con:

$$V_{medio} = \frac{2A \cdot \frac{T}{2}}{T} = A$$

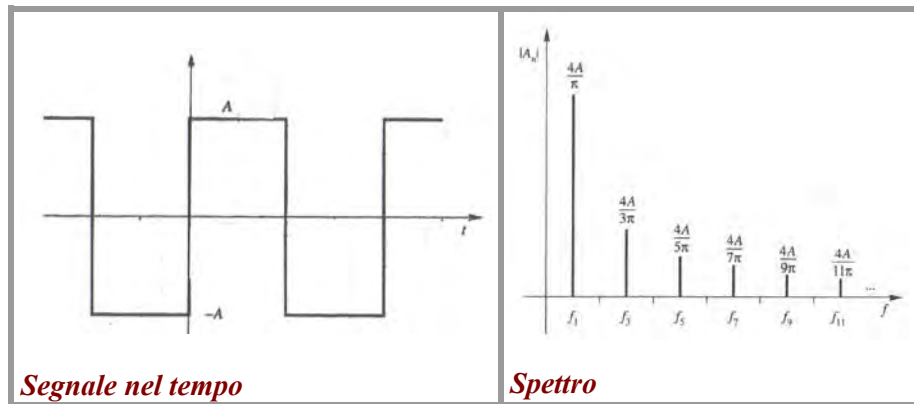
In forma più compatta:

$$v(t) = V_{medio} + \frac{4A}{\pi} \cdot \sum_{m=0}^{\infty} (-1)^m \cdot \frac{\cos[(2m+1) \cdot \omega_0 \cdot t]}{(2m+1)}$$

Dove:

- $2m+1 = n$ = ordine dell'armonica
- $m=0 \rightarrow$ prima armonica
- $m=1 \rightarrow$ terza armonica
- $m=2 \rightarrow$ quinta armonica

ONDA QUADRA con valore picco-picco 2A
BIPOLARE ALTERNATA
DISPARI



L'ONDA è QUADRA (duty- cycle del 50%) e quindi è EMISIMMETRICA, per cui il suo sviluppo in serie di Fourier presenta sole armoniche dispari

L'ONDA è DISPARI, per cui il suo sviluppo in serie di Fourier è con sole armoniche in seno a fase nulla

$$v(t) = \frac{4A}{\pi} \cdot \text{sen}(\omega_0 \cdot t) + \frac{4A}{3 \cdot \pi} \cdot \text{sen}(3 \cdot \omega_0 \cdot t) + \frac{4A}{5 \cdot \pi} \cdot \text{sen}(5 \cdot \omega_0 \cdot t) + \frac{4A}{7 \cdot \pi} \cdot \text{sen}(7 \cdot \omega_0 \cdot t) + \dots$$

Si noti che i termini sono tutti col segno “+”

In forma più compatta:

$$v(t) = \frac{4A}{\pi} \cdot \sum_{m=0}^{\infty} \frac{\text{sen}[(2m+1) \cdot \omega_0 \cdot t]}{(2m+1)}$$

Si noti che manca il termine $(-1)^m$

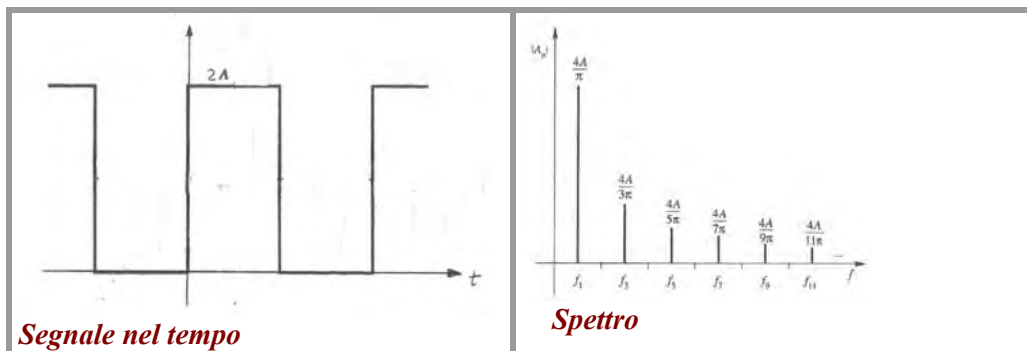
Nell'espressione è:

- $2m+1 = n =$ ordine dell'armonica
- $m=0 \rightarrow$ prima armonica
- $m=1 \rightarrow$ terza armonica
- $m=2 \rightarrow$ quinta armonica

ONDA QUADRA con valore picco-picco 2A

UNIPOLARE

Ottenuta per traslazione verticale di un'onda DISPARI



Si ottiene dalla precedente traslandola verso l'alto di A , cioè sommandole una componente continua (valor medio in T) di A

Lo sviluppo in serie di Fourier rimane **invariato (rispetto al caso precedente)**, tranne che per l'**aggiunta del valore medio**.

$$v(t) = V_{medio} + \frac{4A}{\pi} \cdot \text{sen}(\omega_0 \cdot t) + \frac{4A}{3 \cdot \pi} \cdot \text{sen}(3 \cdot \omega_0 \cdot t) + \frac{4A}{5 \cdot \pi} \cdot \text{sen}(5 \cdot \omega_0 \cdot t) + \frac{4A}{7 \cdot \pi} \cdot \text{sen}(7 \cdot \omega_0 \cdot t)$$

Si noti che i termini sono tutti col segno “+”

In forma più compatta:

$$v(t) = V_{medio} + \frac{4A}{\pi} \cdot \sum_{m=0}^{\infty} \frac{\text{sen}[(2m+1) \cdot \omega_0 \cdot t]}{(2m+1)}$$

Si noti che manca il termine $(-1)^m$

Osserviamo infine che:

$$V_{medio} = \frac{2A \cdot \frac{T}{2}}{T} = A$$

OSSERVAZIONE

Possiamo notare che per **tutte le onde quadre** (onde rettangolari con duty-cycle del 50%) con $V_{pp}=2A$ lo **spettro di ampiezza** (a prescindere dall'eventuale riga a frequenza zero) presenta coefficienti **del tipo**:

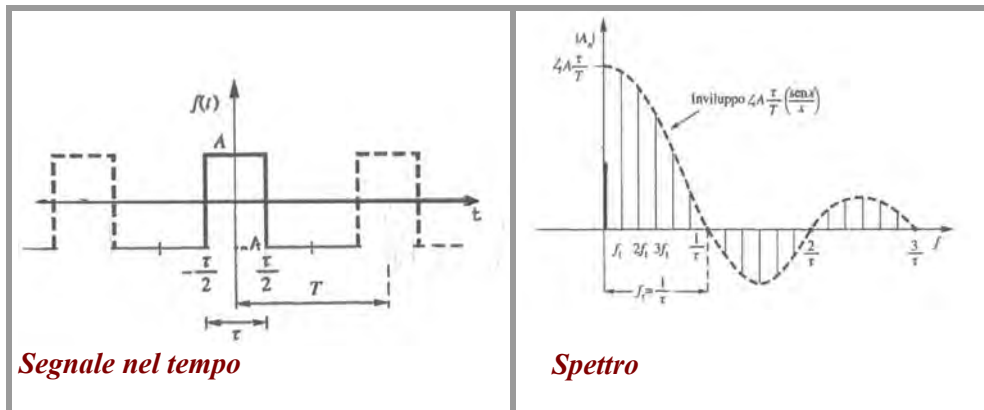
$$\frac{4A}{\pi}, \quad \frac{4A}{3 \cdot \pi}, \quad \frac{4A}{5 \cdot \pi}, \quad \frac{4A}{7 \cdot \pi}$$

La riga a frequenza zero (che rappresenta il valore medio e ha ampiezza A) si deve aggiungere quando l'onda è unipolare.

ONDE RETTANGOLARI (NON QUADRE) CON DUTY-CYCLE DIVERSO DAL 50%

NON sono emisimmetriche, quindi presentano anche le armoniche di ordine pari
Sono bipolari, ma NON alternate, perciò il loro valore medio NON è nullo.

ONDA RETTANGOLARE con valore picco-picco 2A CON DUTY-CYCLE DIVERSO DAL 50% PARI

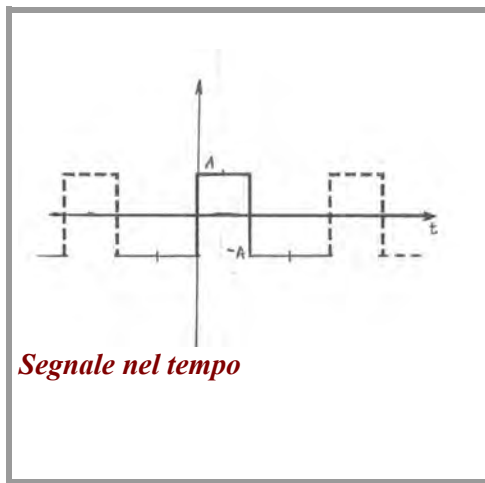


$$v(t) = V_{medio} + 4A \cdot \frac{\tau}{T} \cdot \sum_{m=0}^{\infty} \frac{\sin\left(n \cdot \omega_0 \cdot \frac{\tau}{2}\right)}{n \cdot \omega_0 \cdot \frac{\tau}{2}}$$

dove:

$$V_{medio} = \frac{A \cdot \tau - (T - \tau) \cdot A}{T} = \frac{A \cdot \tau - A \cdot T + A \tau}{T} = \frac{A \cdot (2\tau - T)}{T} = A \cdot \left(2 \cdot \frac{\tau}{T} - 1\right)$$

**ONDA RETTANGOLARE con valore picco-picco 2° CON DUTY-CYCLE DIVERSO DAL 50%
DISPARI**



$$v(t) = V_{medio} + \sum_{m=0}^{\infty} \frac{2A}{n \cdot \pi} \cdot [1 - \cos(n \cdot \omega_0 \cdot \tau)]$$

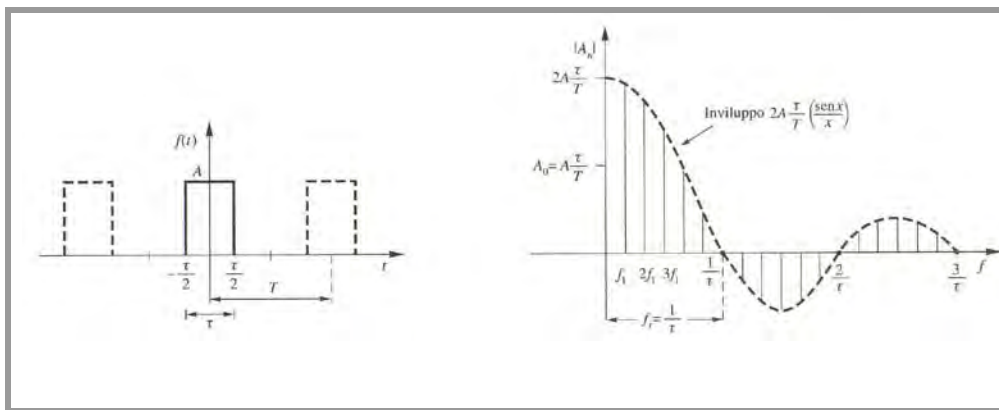
dove:

$$V_{medio} = \frac{A \cdot \tau - (T - \tau) \cdot A}{T} = \frac{A \cdot \tau - A \cdot T + A \tau}{T} = \frac{A \cdot (2\tau - T)}{T} = A \cdot \left(2 \cdot \frac{\tau}{T} - 1\right)$$

IL SEGNALE APERIODICO COME “LIMITE” DEL SEGNALE PERIODICO

Dato un segnale periodico, per esempio una sequenza di impulsi rettangolari di durata τ , di periodo T e di frequenza $f_0=1/T$, il suo spettro complesso è costituito da infinite righe che (essendo situate in $f_0, 2f_0, 3f_0, \dots, n f_0$) distano una dall'altra proprio di f_0 .

I vertici delle righe spettrali formano una curva ideale, chiamata “inviluppo”, che ha l'andamento di una oscillazione smorzata¹.



Se facciamo tendere all'infinito il periodo T della sequenza, possiamo pensare che l'impulso centrale resti fermo nella sua posizione, e che tutti gli altri “scompaiano” allontanandosi all'infinito. Per $T \rightarrow \infty$ quindi l'onda rettangolare si riduce a un unico impulso. In altri termini facendo tendere all'infinito il periodo T di un segnale periodico, tale segnale diventa aperiodico.

Che conseguenze ha sullo spettro del segnale il fatto di aver fatto diventare infinito il suo periodo? Se T tende all'infinito, la frequenza $f_0=1/T$ tende a zero, il che significa che le righe dello spettro tendono ad avvicinarsi sempre più una all'altra fino a diventare indistinguibili e a ricoprire completamente la banda dello spettro, mantenendo però inalterata la forma dell'inviluppo. Possiamo in definitiva dire che facendo tendere all'infinito il periodo di un segnale periodico, tale segnale diventa aperiodico e il suo spettro diventa continuo, ma conserva la forma dell'inviluppo.

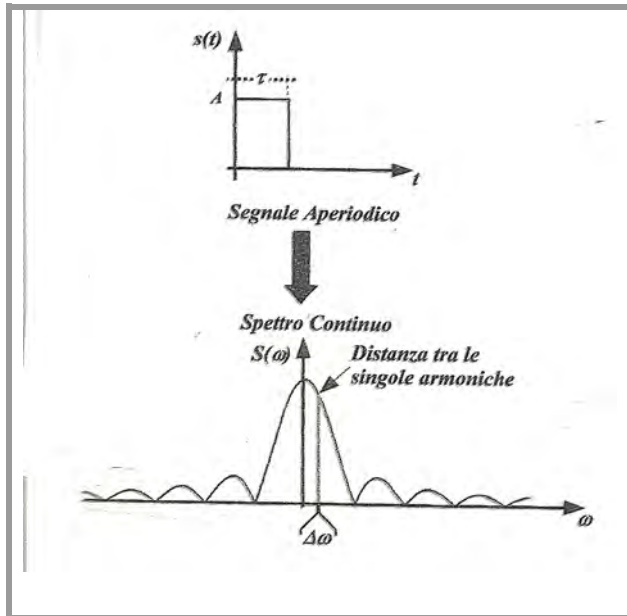
¹ In ogni lobo (“cupola”) dell'inviluppo sono contenute T/τ righe, l'ultima delle quali è nulla. Quindi ciascun lobo contiene $(T/\tau) - 1$ righe non nulle. Se per esempio è $T/\tau = 10$ (onda con duty-cycle del 10%), ogni lobo contiene 9 armoniche non nulle.

Le armoniche nulle si trovano in corrispondenza dei valori di frequenza $f = N \cdot (1/\tau)$, nei quali due lobi successivi si toccano.

Infine, per un'onda rettangolare con duty-cycle del 50%, si ha $T/\tau = 2$, per cui ciascun lobo è formato da una sola riga (due righe di cui una nulla)

SEGNALI APERIODICI SCOMPOSIZIONE NEL TEMPO E SPETTRO

Conseguenza del fatto che un segnale aperiodico può essere visto come un “particolare” segnale periodico, con periodo T tendente all’infinito e con frequenza f_0 tendente a zero, è che **le righe dello spettro non sono più separate una dall’altra, ma indistinguibili**.



Lo spettro del segnale aperiodico è uno spettro “**continuo**”, uno spettro cioè costituito da **infinite righe poste a distanza infinitesima una dall’altra**, righe i cui vertici formano una curva ideale detta “**curva di inviluppo $V(\omega)$** ”.

La curva di inviluppo $V(\omega)$ prende anche il nome di **trasformata di Fourier** o **F-trasformata** e “delimita” lo spettro costituendone il “contorno”.

$$V(\omega) = \int_{t=-\infty}^{t=+\infty} v(t) \cdot e^{-j\omega \cdot t} \cdot dt$$

Trasformata di Fourier

Riguardo allo **spettro di ampiezza** del segnale aperiodico, cioè riguardo al **modulo dello spettro** del segnale, esso è **delimitato** dal **modulo $|V(\omega)|$** dell’inviluppo, cioè dal modulo della trasformata di Fourier $V(\omega)$.

“SCOMPOSIZIONE NEL TEMPO” DEL SEGNALE APERIODICO: INTEGRALE DI FOURIER O ANTITRASFORMATATA DI FOURIER

Sappiamo che un segnale periodico può essere rappresentato, nel tempo, come la somma di un termine costante e di un numero infinito di armoniche, distinguibili l'una dall'altra.

Ciascuna armonica è un segnale sinusoidale.

Questo fatto è rappresentabile, da un punto di vista matematico, mediante una sommatoria di infiniti termini che prende il nome di sviluppo in serie di Fourier:

$$v(t) = A_0 + \sum_{n=1}^{\infty} C_n \sin(2\pi \cdot n \cdot f_0 \cdot t + \varphi_n)$$

E' possibile rappresentare in questo modo anche un segnale aperiodico? La risposta è negativa perché la sommatoria si può fare solo su entità discrete, cioè distinguibili l'una dall'altra, mentre, come abbiamo prima detto, lo spettro di un segnale aperiodico è continuo, cioè costituito da infinite righe a distanza infinitesima una dall'altra, righe che non sono distinguibili o separabili l'una dall'altra.

Se non è possibile fare la sommatoria, risulta però possibile fare l'integrale che può essere considerato come una “somma di infiniti termini infinitesimi”.

Di conseguenza, per i segnali aperiodici lo sviluppo in serie di Fourier è sostituito da un integrale, detto appunto **integrale di Fourier o antitrasformata di Fourier**, che rappresenta la “**scomposizione nel tempo**” del segnale aperiodico.

Possiamo anche dire che l'integrale di Fourier o antitrasformata di Fourier rappresenta l'**estensione, al caso dei segnali non periodici, dello sviluppo in serie di Fourier**.

L'espressione matematica dell'integrale di Fourier è:

$$v(t)_{\text{aperiodico}} = \frac{1}{2\pi} \int_{\omega=-\infty}^{\omega=+\infty} V(\omega) \cdot e^{+j\omega \cdot t} \cdot d\omega$$

Integrale di Fourier

La funzione continua $V(\omega)$, che compare all'interno dell'integrale, è la trasformata di Fourier del segnale $v(t)$ e, come già detto, rappresenta l'involuppo dello spettro del segnale aperiodico.

Osserviamo esplicitamente che, nel passaggio da un segnale periodico a un segnale non periodico, la rappresentazione temporale del segnale passa dallo sviluppo in serie di Fourier all'integrale di Fourier e l'operatore di sommatoria dello sviluppo in serie diventa un integrale.

Inoltre i coefficienti C_n (che rappresentano le ampiezze delle singole armoniche e quindi le “lunghezze” delle righe spettrali) vengono sostituiti dalla funzione $V(\omega)$ che rappresenta l'involuppo, ossia la linea di confine dello spettro continuo del segnale aperiodico:

$$v(t) = A_0 + \sum_{n=1}^{\infty} C_n \sin(2\pi \cdot n \cdot f_0 \cdot t + \varphi_n)$$

$$v(t)_{\text{aperiodico}} = \frac{1}{2\pi} \int_{\omega=-\infty}^{\omega=+\infty} V(\omega) \cdot e^{+j\omega \cdot t} \cdot d\omega$$

SPETTRO DEI SEGNALI APERIODICI: TRASFORMATA DI FOURIER O F-TRASFORMATATA

Come già accennato, la **trasformata di Fourier** o **F-trasformata** rappresenta la **curva di confine**, cioè l'**inviluppo**, dello **spettro** di un segnale aperiodico e sostituisce ciò che nello sviluppo in serie è rappresentato dai coefficienti C_n .

L'espressione matematica della trasformata di Fourier o F-trasformata è:

$$V(\omega) = \int_{t=-\infty}^{t=+\infty} v(t) \cdot e^{-j\omega t} \cdot dt$$

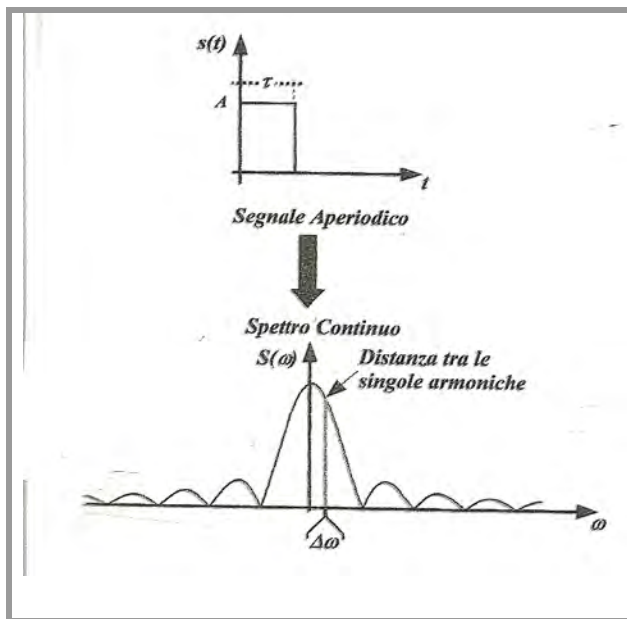
Trasformata di Fourier

NOTA

Anche la trasformata di Fourier è un integrale ma differisce dall'integrale di Fourier per questi motivi:

- il "risultato" della trasformata è una funzione della pulsazione e quindi della frequenza ($V(\omega)$), mentre il "risultato" dell'integrale di Fourier è una funzione del tempo ($v(t)$)
- l'integrazione avviene nel tempo e non nella pulsazione ω ;
infatti gli estremi di integrazione sono $t=-\infty$ e $t=+\infty$ (e non $\omega=-\infty$ e $\omega=+\infty$) e compare il prodotto per "dt" e non per "d ω "
- l'esponente del numero di Nepero "e" ha segno negativo

SPETTRO E TRASFORMATA DI FOURIER DI UN SINGOLO IMPULSO RETTANGOLARE



Lo spettro di un singolo impulso rettangolare di durata T e ampiezza A (che è un esempio di segnale aperiodico) è l'area contornata dalla curva di involuppo $V(\omega)$, cioè dalla trasformata di Fourier. Si può dimostrare che l'espressione matematica della trasformata di Fourier $V(\omega)$ dell'impulso rettangolare è:

$$V(\omega) = \frac{A \cdot T}{\pi} \cdot \text{sen}\left(\omega \cdot \frac{\tau}{2}\right)$$

CAPITOLO 19

FILTRI PASSIVI

COMPLEMENTI

FILTRI PASSIVI A INDUTTANZA E CAPACITA'

FILTRO LC PASSABASSO

CASCATA DI DUE CELLE ELEMENTARI LC PASSABASSO

FILTRO CL PASSAALTO

CASCATA DI DUE CELLE ELEMENTARI CL PASSAALTO

UN POSSIBILE FILTRO PASSABANDA LC: IL CIRCUITO LC RISONANTE PARALLELO

IL CIRCUITO LC PARALLELO COME FILTRO PASSABANDA

RISONANZA PARALLELO

FILTRI ELIMINABANDA O NOTCH

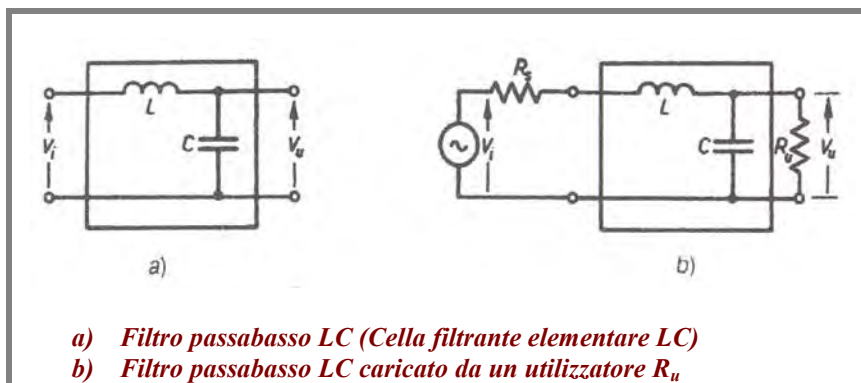
IL FILTRO NOTCH A PONTE DI WIEN

FILTRI A INDUTTANZA E CAPACITÀ

I circuiti selettivi che abbiamo esaminato nella loro realizzazione con resistenze e capacità, si possono anche ottenere utilizzando induttanze (al posto delle resistenze) e capacità, ottenendo il vantaggio di una **pendenza maggiore** nella curva del **modulo della risposta in frequenza**: la curva linearizzata del modulo della risposta in frequenza della cella elementare LC, cioè del filtro con un solo induttore e una sola capacità, ha una pendenza **40 dB/décade** contro i 20 dB/decade dell'analogo filtro a resistenza e capacità.

FILTRO LC PASSABASSO

Ha la stessa struttura del circuito RC passabasso, con la sola differenza dell'induttanza al posto della resistenza.



Anche nel circuito passabasso LC la tensione di ingresso si applica ai capi della serie e la tensione di **uscita** si preleva **ai capi del condensatore**.

Osserviamo che in entrambi i filtri passa basso (RC ed LC) il prelievo della tensione di uscita avviene su quell'impedenza del circuito il cui modulo decresce più velocemente all'aumentare della frequenza.

Nel circuito RC, al crescere della frequenza, il modulo ($Z=1/\omega C$) dell'impedenza del condensatore diminuisce, mentre il modulo R dell'impedenza del resistore è costante con la frequenza. Pertanto l'uscita è la tensione sul condensatore.

Nel circuito LC, al crescere della frequenza, il modulo ($Z=1/\omega C$) dell'impedenza del condensatore diminuisce, mentre il modulo ($Z = \omega L$) dell'impedenza aumenta. Perciò, anche in questo caso, l'uscita è la tensione sul condensatore.

La **frequenza di taglio**, cioè il valore di frequenza in corrispondenza del quale A_v vale $0,707A_{vMAX}$, è:

$$f_i = \frac{1}{2\pi \cdot \sqrt{L \cdot C}}$$

MASSIMO TRASFERIMENTO DI POTENZA (DAL GENERATORE ALL'UTILIZZATORE) E RESISTENZA CARATTERISTICA

Se il filtro è interposto fra un **generatore di resistenza interna R_G** e un dispositivo **utilizzatore**, schematizzabile mediante una **resistenza R_U** , affinché si abbia il **massimo trasferimento di potenza** dal generatore all'utilizzatore, deve essere verificata la condizione:

$$R_G = R_U = R_0$$

dove **R_0** è la **resistenza caratteristica** del filtro, che ha espressione:

$$R_0 = \sqrt{\frac{L}{C}}$$

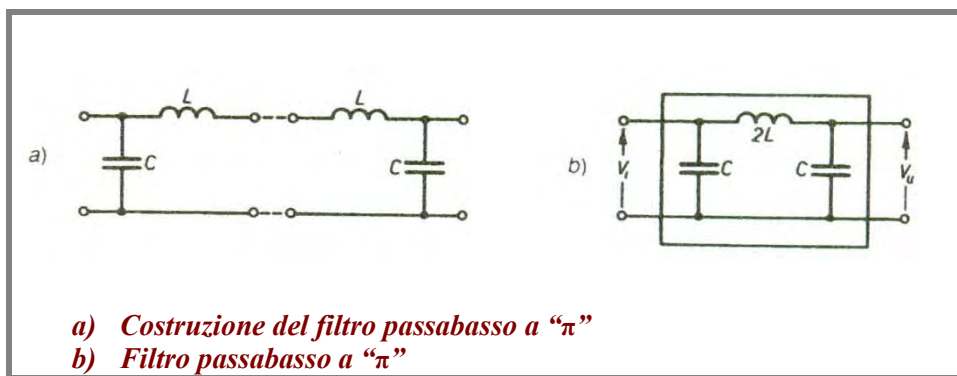
CASCATA DI DUE CELLE ELEMENTARI LC PASSABASSO

Il collegamento in cascata di due celle elementari LC determina un ulteriore aumento di pendenza (della curva del modulo della risposta in frequenza) rispetto alla cella elementare LC.

Esamineremo due diversi collegamenti di celle elementari LC: il collegamento a “ π ” e il collegamento a “T”

CASCATA DI DUE CELLE ELEMENTARI LC: FILTRO LC PASSABASSO A “ π ”

Collegando in cascata, in maniera “**speculare**”, con i rami delle induttanze L che convergono al centro (come nella figura seguente) due celle elementari LC e **sostituendo**, alle due induttanze L, **un'unica induttanza** di valore **2L**, si ottiene il **filtro passabasso a “ π ”**.



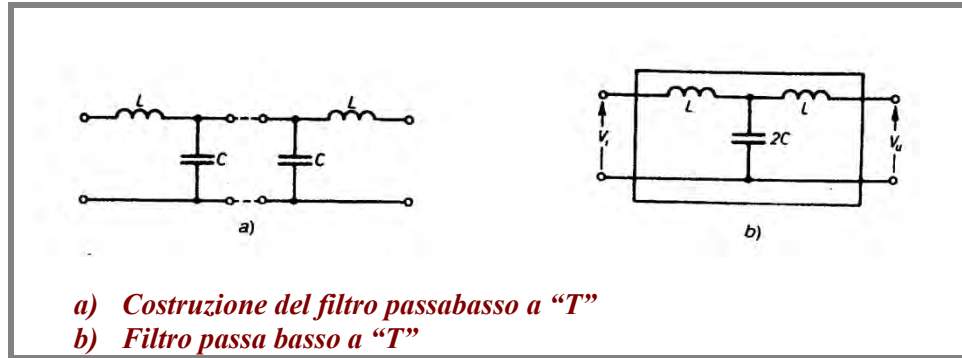
La frequenza di taglio e la resistenza caratteristica sono identiche a quelle della cella elementare passabasso LC:

$$f_c = \frac{1}{2\pi \cdot \sqrt{L \cdot C}}$$

$$R_0 = \sqrt{\frac{L}{C}}$$

CASCATA DI DUE CELLE ELEMENTARI LC: FILTRO LC PASSABASSO A “T”

Collegando in cascata, in maniera “**speculare**”, con i rami delle induttanze che divergono agli estremi (come nella figura seguente) due celle elementari LC e **sostituendo**, ai due condensatori, **un’unica capacità** di valore **2C**, si ottiene il **filtro passabasso a “T”**.



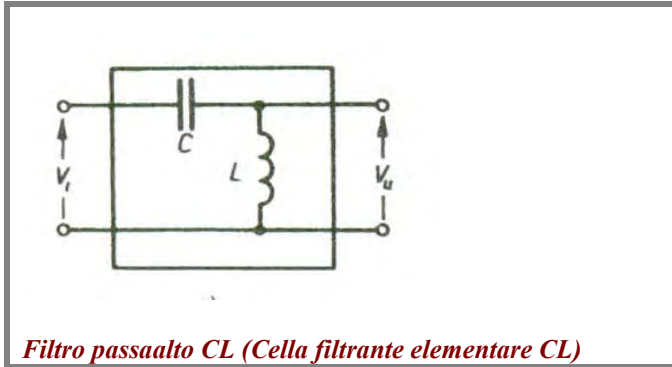
La frequenza di taglio e la resistenza caratteristica sono identiche a quelle della cella elementare passabasso LC:

$$f_c = \frac{1}{2\pi \cdot \sqrt{L \cdot C}}$$

$$R_0 = \sqrt{\frac{L}{C}}$$

FILTRO CL PASSAALTO

Ha la stessa struttura del circuito CR passaalto, con la sola differenza dell'induttanza al posto della resistenza.



Anche nel circuito passaalto LC la tensione di ingresso si applica ai capi della serie, ma la tensione di **uscita** si preleva **ai capi dell'induttore, che prende il posto della resistenza**.

Osserviamo che in entrambi i filtri passaalto (CR e CL) il prelievo della tensione di uscita avviene su quell'impedenza del circuito il cui modulo decresce più lentamente (o cresce più velocemente) all'aumentare della frequenza.

Nel circuito CR, al crescere della frequenza, il modulo ($Z=1/\omega C$) dell'impedenza del condensatore diminuisce, mentre il modulo R dell'impedenza del resistore è costante con la frequenza. Pertanto l'uscita è la tensione sulla resistenza.

Nel circuito CL, al crescere della frequenza, il modulo ($Z=1/\omega C$) dell'impedenza del condensatore diminuisce, mentre il modulo ($Z=\omega L$) dell'impedenza dell'induttore cresce con la frequenza.

Pertanto l'uscita è la tensione sull'induttore.

La **frequenza di taglio**, cioè il valore di frequenza in corrispondenza del quale A_V vale $0,707A_{VMAX}$, è:

$$f_i = \frac{1}{2\pi \cdot \sqrt{L \cdot C}}$$

MASSIMO TRASFERIMENTO DI POTENZA (DAL GENERATORE ALL'UTILIZZATORE) E RESISTENZA CARATTERISTICA

Anche in questo caso, se il filtro è interposto fra un **generatore di resistenza interna R_G** e un dispositivo **utilizzatore**, schematizzabile mediante una **resistenza R_U** , affinché si abbia il **massimo trasferimento di potenza** dal generatore all'utilizzatore, deve essere verificata la condizione:

$$R_G = R_U = R_0$$

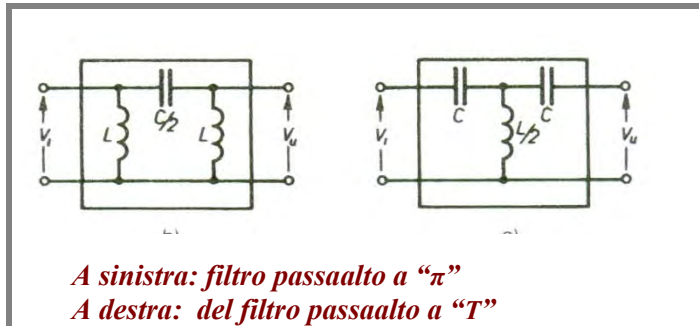
dove R_0 è la *resistenza caratteristica* del filtro, che ha espressione:

$$R_0 = \sqrt{\frac{L}{C}}$$

CASCATA DI DUE CELLE ELEMENTARI CL PASSAALTO

Il collegamento in cascata di due celle elementari CL determina un ulteriore aumento di pendenza (della curva del modulo della risposta in frequenza) rispetto alla cella elementare CL.

Anche in questi caso, esamineremo due diversi collegamenti di celle elementari LC: il collegamento a “ π ” e il collegamento a “T”



Per entrambi i filtri, la frequenza di taglio e la resistenza caratteristica sono identiche a quelle della cella elementare passaalto CL e passabasso LC:

$$f_c = \frac{1}{2\pi \cdot \sqrt{L \cdot C}}$$

$$R_0 = \sqrt{\frac{L}{C}}$$

CASCATA DI DUE CELLE ELEMENTARI LC: FILTRO LC PASSAALTO A “ π ”

Collegando in cascata, in maniera “speculare”, con i rami delle capacità C che convergono al centro (come nella figura seguente), due celle elementari CL e **sostituendo**, alle capacità, **un'unica capacità** di valore C/2, si ottiene il **filtro passaalto a “ π ”**.

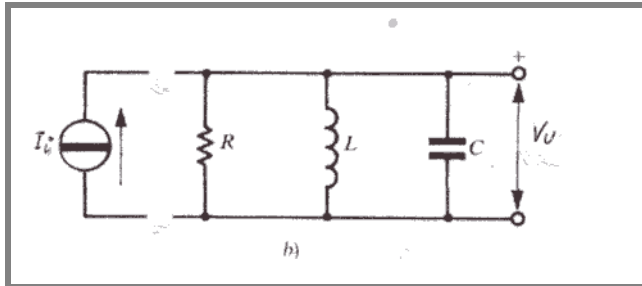
CASCATA DI DUE CELLE ELEMENTARI LC: FILTRO LC PASSAALTO A “T”

Collegando in cascata, in maniera “speculare”, con i rami delle capacità che divergono agli estremi (come nella figura precedente), due celle elementari CL e **sostituendo**, ai due induttori, **un'unica induttanza** di valore L/2, si ottiene il **filtro passaalto a “T”**.

UN POSSIBILE FILTRO PASSABANDA LC: IL CIRCUITO LC RISONANTE PARALLELO

Questo circuito è trattato approfonditamente nel capitolo relativo al ricevitore eterodina (in quanto è alla base dell'apparato di sintonia).

Pertanto qui se ne esporranno solo gli aspetti fondamentali.



Il circuito LC (o RLC) risonante è formato dal parallelo di una capacità e di un'induttanza. Nella rappresentazione circuitale viene aggiunta una resistenza R_p in parallelo, che schematizza le perdite del circuito, ossia il fatto che il condensatore e l'induttore non sono in realtà puramente capacitivi o induttivi, ma presentano una resistenza parassita che determina dissipazione di potenza in calore per effetto Joule.

La resistenza di perdita è schematizzata in parallelo, perciò:

- ❖ se il suo valore è basso significa che le perdite sono rilevanti
- ❖ se il suo valore è alto significa che le perdite sono trascurabili.

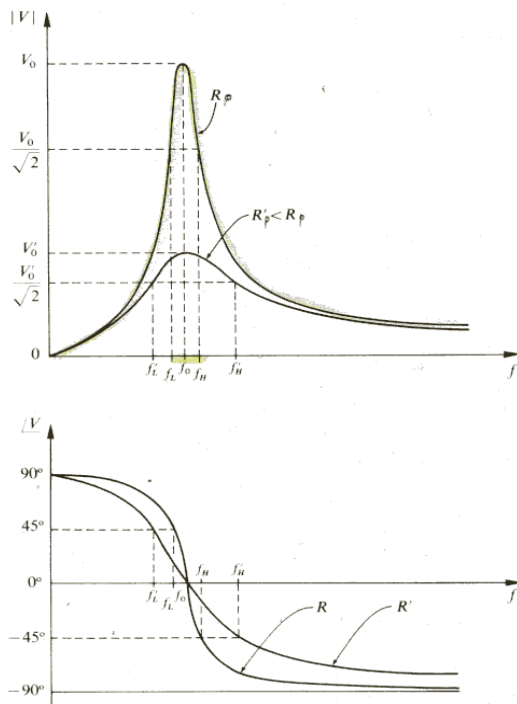
Il circuito è alimentato da un generatore ideale di corrente, che può essere schematizzato come un generatore ideale di tensione con resistenza interna R_G molto più grande del modulo dell'impedenza del carico.

Carico che in questo caso è proprio il parallelo $R_p//L//C$.

IL CIRCUITO LC PARALLELO COME FILTRO PASSABANDA

Il circuito LC parallelo si comporta come un filtro passabanda, che presenta il **massimo della tensione di uscita** in corrispondenza di una frequenza detta **frequenza di risonanza** (inversamente proporzionale ai valori di L e di C del filtro), la quale può anche essere vista (approssimativamente) come **frequenza di “centro-banda**.

La banda del filtro è tanto più stretta (e quindi la selettività è tanto maggiore) quanto **più elevato** è il **coefficiente Q di risonanza**.



Tensione di uscita (sul parallelo) in funzione della frequenza

$$\omega_0 = \frac{1}{\sqrt{L \cdot C}} \quad f_0 = \frac{1}{2\pi\sqrt{L \cdot C}} \quad \text{pulsazione e frequenza di risonanza}$$

$$Q = \frac{R_p}{\omega_0 \cdot L} = \omega_0 \cdot C \cdot R_p \quad \text{coefficiente di risonanza}$$

RISONANZA PARALLELO

Nel parallelo dell'induttore¹ e del condensatore le correnti in C e in L possono diventare notevolmente più grandi della corrente di ingresso (sovracorrenti)

Alla frequenza di risonanza

- ❖ l'**impedenza** diventa **massima** e puramente **resistiva**
- ❖ la **tensione di uscita** ai capi del parallelo, essendo data da $V_u = I_i |Z|$ risulta **massima**.

L'impedenza del circuito è:

$$Z(\omega) = \frac{1}{Y(\omega)} = \frac{1}{\frac{1}{R_p} + j\left(\omega \cdot C - \frac{1}{\omega \cdot L}\right)}$$

IL MODULO dell'impedenza è:

$$Z(\omega) = \frac{1}{Y(\omega)} = \frac{1}{\sqrt{\frac{1}{R_p^2} + \left(\omega \cdot C - \frac{1}{\omega \cdot L}\right)^2}}$$

L'**impedenza** risulta **massima** (in modulo) in corrispondenza del **valore di ω** , e quindi di **f**, per il quale la parentesi a denominatore diventa zero, valore che è detto "pulsazione di risonanza" o "**frequenza di risonanza**"

$$\omega = \frac{1}{\sqrt{L \cdot C}} = \omega_0 \quad \text{pulsazione di risonanza}$$

e quindi per

$$f = \frac{1}{2\pi\sqrt{L \cdot C}} = f_0 \quad \text{frequenza di risonanza}$$

In corrispondenza della **frequenza di risonanza** l'**impedenza** è **massima** e puramente **resistiva** e vale **R_p**.

Per tutte le frequenze diverse da quella di risonanza, l'impedenza è, in modulo, più bassa perché, al denominatore della sua espressione, si aggiunge il contenuto della parentesi, comunque positivo in quanto reale elevato al quadrato.

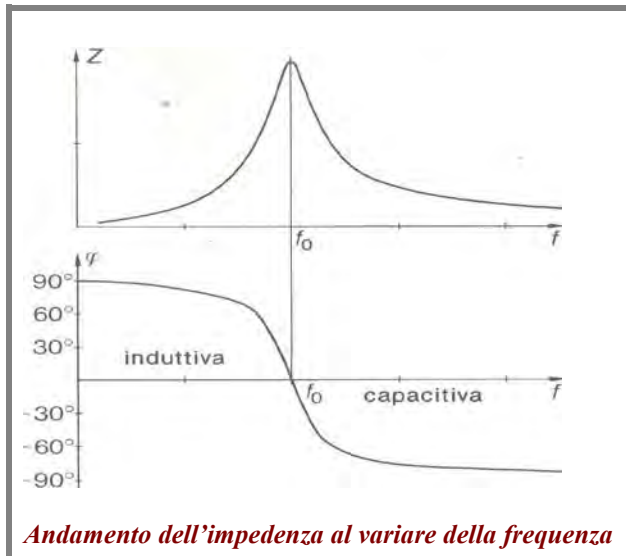
Di conseguenza la tensione di uscita ai capi del parallelo, essendo data da $V_u = I_i |Z|$, risulta massima.

¹ Ricordiamo che il valore L di induttanza dell'induttore è

- ❖ direttamente proporzionale al numero di spire e al diametro dell'avvolgimento
- ❖ inversamente proporzionale allo spessore del filo

In sintesi, alla frequenza di risonanza:

- ❖ il circuito si riduce alla sola resistenza parallelo R_p che rappresenta il massimo valore di impedenza
- ❖ la tensione di uscita sul parallelo è massima e vale $R_p \cdot I_i$, in quanto il generatore vede come carico la sola resistenza R_p



Le correnti sull'induttore e sul condensatore sono, in risonanza:

- ❖ uguali e opposte tra loro (di uguale ampiezza e sfasate di 180° una rispetto all'altra)
- ❖ Q volte più grandi della corrente I_i entrante nel parallelo, dove " Q " è il coefficiente di risonanza, definito di seguito.

Q è il coefficiente di risonanza del circuito e vale

$$Q = \frac{R_p}{Z_L} = \frac{R_p}{Z_c} = \frac{R_p}{\omega_0 \cdot L} = \omega_0 \cdot C \cdot R_p$$

Q ha un valore tanto più elevato quanto minori sono le perdite del circuito, schematizzate dalla R_p (cioè quanto più grande è la R_p)

Riguardo alle correnti del circuito oscillante parallelo è verificata, per ogni frequenza, la somma vettoriale:

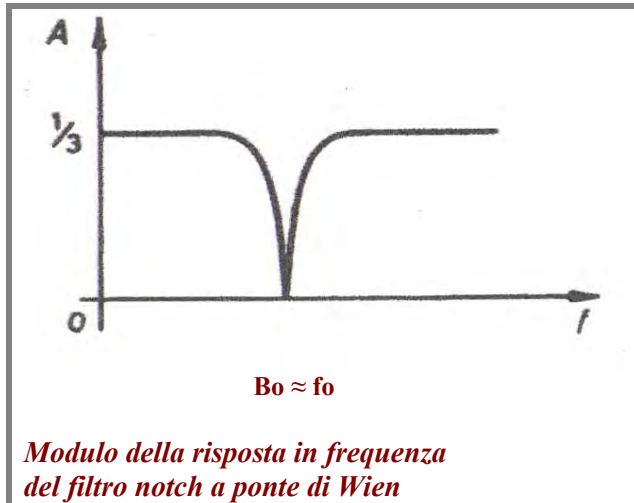
$$I_R + I_C + I_L = I_i$$

In RISONANZA, siccome I_C e I_L sono uguali e opposte, si ha $I_i = I_R$

FILTRI ELIMINA BANDA O NOTCH

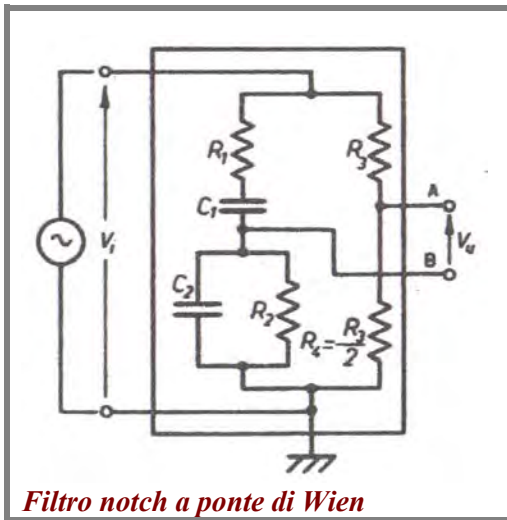
I filtri “eliminabanda” (o “a reiezione di banda” o “notch”) hanno la funzione di **eliminare** tutte le componenti di segnale **appartenenti a una certa banda B_0** più o meno stretta e di “lasciar passare tutte le componenti di segnale aventi frequenze esterne alla banda B_0 ”.

Fra i possibili circuiti notch accenneremo al filtro a **ponte di Wien**, caratterizzato da una **banda B_0 molto stretta**, come si vede dalla figura seguente:



IL FILTRO NOTCH A PONTE DI WIEN

Per questo tipo di filtro a reiezione di banda, essendo la **banda Bo, molto stretta**, si può ritenere che **tale banda si riduca a una sola frequenza**: la **frequenza centrale fo**.



$f_o = \frac{1}{2\pi \cdot R \cdot C}$	<i>Frequenza centrale del filtro</i> essendo: $R = R_1 = R_2$ $C = C_1 = C_2$ $R_4 = R_3/2$
--	---

$ \bar{A} = A = \frac{V_o}{V_i} = \frac{1}{3}$	<i>Valore di A, cioè di Vo/Vi, per frequenze diverse da quella centrale</i>
---	--

CAPITOLO 20

TRANSISTOR BJT E AMPLIFICATORI A BJT

COMPLEMENTI

RISPOSTA IN FREQUENZA DI UN AMPLIFICATORE RC

***DETERMINAZIONE DELLA FREQUENZA DI TAGLIO INFERIORE
RELATIVA A UNA SINGOLA CAPACITA'***

***DETERMINAZIONE DELLA FREQUENZA DI TAGLIO SUPERIORE
RELATIVA A UNA SINGOLA CAPACITA'***

FREQUENZA DI TAGLIO INFERIORE COMPLESSIVA

AMPLIFICATORI DI POTENZA E CLASSI DI AMPLIFICAZIONE

DISTORSIONE

DISTORSIONE LINEARE

DISTORSIONE ARMONICA O NON LINEARE

PARAMETRI DELLA DISTORSIONE ARMONICA O NON LINEARE

DISTORSIONE NON LINEARE DI INTERMODULAZIONE

RISPOSTA IN FREQUENZA DI UN AMPLIFICATORE RC

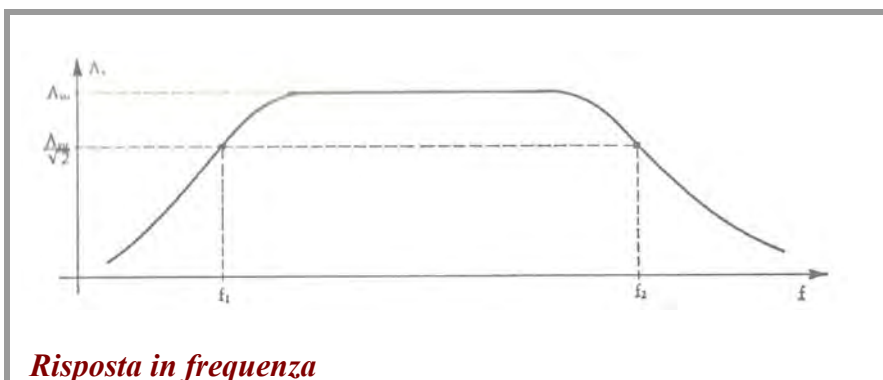
Gli amplificatori nei quali il collegamento del transistor (e della sua rete di polarizzazione) con il generatore e con il carico è realizzato mediante condensatori, detti appunto “di accoppiamento”, sono chiamati “ad accoppiamento RC” o “ad accoppiamento in alternata”.

I condensatori infatti bloccano le componenti continue del segnale, il che è necessario per questi motivi:

- le correnti continue derivanti dalla rete di polarizzazione non devono scorrere sul carico e sulla sorgente di segnale perché potrebbero danneggiarli.
- eventuali componenti continue presenti nella sorgente di segnale e nel carico (che potrebbero essere ulteriori stadi amplificatori) non devono “entrare” nel circuito di polarizzazione perché altererebbero i valori di progetto delle tensioni e le correnti di polarizzazione.

Un amplificatore RC che sia stato progettato per lavorare con segnali appartenenti a una banda centrata sulla frequenza f_0 presenta un valore di amplificazione che non è costante al variare della frequenza. L'amplificazione risulta:

- massima e costante in una banda centrata sulla frequenza di lavoro f_0
- decrescente al crescere della frequenza al di sopra della banda centrale
- crescente al crescere della frequenza al di sotto della banda centrale.



Possiamo descrivere questo comportamento dicendo che l'amplificatore RC ha una risposta in frequenza di tipo **passabanda**, cioè una risposta caratterizzata sia da una frequenza di taglio inferiore che da una frequenza di taglio superiore.

L'amplificatore RC alle **basse frequenze** si comporta come un filtro **passaalto** (nel senso che avviene un taglio delle componenti di segnale a frequenza più bassa) e alle **alte frequenze** come un filtro **passabasso** (nel senso che vengono tagliate le componenti di segnale a frequenza più alta). Sappiamo che il comportamento filtrante di un circuito è determinato dalla presenza di componenti reattivi (condensatori e/o induttori).

Ebbene, nel caso dell'amplificatore RC, il comportamento di tipo **passaalto** alle **basse frequenze** è determinato dai **condensatori di accoppiamento** (al generatore e al carico) e dal **condensatore di bypass o di fuga**.

Il comportamento filtrante **passabasso** alle **alte frequenze** è dovuto invece alle **capacità interne (parassite) del BJT** e cioè alla capacità della giunzione di emettitore e alla capacità della giunzione di collettore, capacità il cui valore non supera le centinaia di pF.

PERCHE' I CONDENSATORI DI ACCOPPIAMENTO E DI BYPASS DETERMINANO UN EFFETTO FILTRANTE PASSAALTO ALLE BASSE FREQUENZE

I condensatori di accoppiamento sono in serie al percorso del segnale.

Essendo caratterizzati da un'impedenza $Z_c = -j \cdot \frac{1}{\omega \cdot C} = -j \cdot \frac{1}{2\pi \cdot f \cdot C}$

il passaggio di corrente variabile su di essi determina una caduta di tensione, la cui entità si sottrae alla tensione utile che deve essere trasferita dal generatore al transistor oppure dall'amplificatore all'utilizzatore, cioè al carico.

Come si vede dalla sua espressione, l'impedenza è tanto più alta quanto più bassa è la frequenza. Mentre quindi alla frequenza di lavoro prevista, l'impedenza di questi condensatori deve essere trascurabile, al diminuire della frequenza l'impedenza diventa rilevante e la caduta su di essa sottrae segnale utile esercitando un filtraggio di tipo passaalto (taglio delle componenti di segnale a bassa frequenza).

Riguardo alla capacità di bypass C_E , la sua impedenza, al diminuire della frequenza, cresce e non cortocircuita più R_E .

Di conseguenza R_E esercita la sua azione limitatrice nei confronti della variazione di segnale, e l'amplificazione, al diminuire della frequenza, decresce (contribuendo così al taglio alle basse frequenze).

Si può anche dire che, siccome a bassa frequenza C_E non cortocircuita più R_E , si ha una caduta della tensione di segnale su R_E che determina una diminuzione di V_{CE} (2° principio di Kirchhoff alla maglia di uscita).

Osserviamo che, per la determinazione della frequenza di taglio inferiore, la frequenza "più importante" è la più alta.

Il motivo è che, alle basse frequenze, l'amplificatore si comporta come un circuito passaalto cioè come un filtro che esclude le frequenze più basse.

Se, per esempio, un segnale è sottoposto al filtraggio passaalto di tre filtri, che tagliano rispettivamente a 200Hz, a 100Hz e a 50Hz, risulta "decisivo" il taglio a 200Hz.

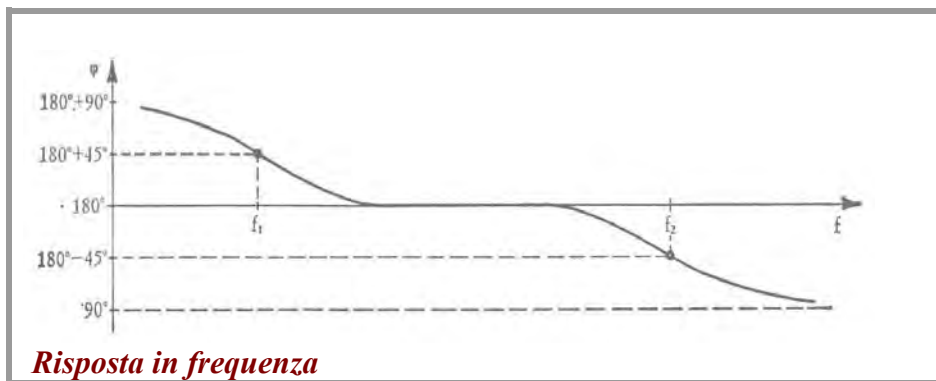
Se infatti percorriamo l'asse f dalle frequenze della banda passante (più alte) verso le più basse, la prima frequenza che incontriamo è quella più alta (200Hz). Essa taglia tutte le frequenze più basse di 200Hz e quindi risulta decisiva.

Le frequenze comprese fra 200 e 100 Hz risultano infatti già eliminate dal taglio a 200Hz, come pure le frequenze comprese fra 100Hz e 50 Hz e tutte quelle ancora più basse.

PERCHE' LE CAPACITA' PARASSITE DI GIUNZIONE DETERMINANO UN EFFETTO FILTRANTE PASSABASSO ALLE ALTE FREQUENZE

Ad alta frequenza le capacità di giunzione e dei collegamenti verso massa tendono a shuntare il segnale verso massa facendo diminuire il guadagno dell'amplificatore.

Osserviamo che, per la determinazione della frequenza di taglio superiore, la frequenza "più importante" è la più bassa (e non, come nel caso della frequenza di taglio inferiore, la più alta). Il motivo è che, alle alte frequenze, il comportamento dell'amplificatore è di tipo passabasso: se percorriamo quindi l'asse f dalle frequenze della banda passante verso le più alte, ci rendiamo conto che è la prima frequenza decisiva è la prima che incontriamo (la più bassa).



DETERMINAZIONE DELLA FREQUENZA DI TAGLIO INFERIORE RELATIVA A UNA SINGOLA CAPACITA' C_i

- Bisogna prima verificare se le capacità che concorrono alla frequenza di taglio complessiva sono interagenti o non interagenti: una capacità C_1 è **interagente** con una capacità C_2 , se C_2 **fa parte dell'impedenza che si vede** fra i terminali di C_1
Riguardo all'amplificatore a **emettitore comune** osserviamo che sono **interagenti** la **capacità di accoppiamento con il generatore** e la **capacità di bypass** di emettitore.
Riguardo all'amplificatore a collettore comune rileviamo che sono interagenti la capacità di accoppiamento con il generatore e la capacità di accoppiamento con il carico.
Il motivo è che **l'impedenza vista dalla base dipende dall'impedenza vista dall'emettitore e viceversa**. Non c'è invece dipendenza fra le impedenze di base e di collettore e fra le impedenze di emettitore e di collettore.
- Si determinano le singole frequenze di taglio secondo la relazione:

$$f_{tot} = \frac{1}{2\pi \cdot R_{eq} \cdot C} = \frac{1}{2\pi \cdot \tau_i}$$

avendo presente che R_{eq-i} è la resistenza equivalente vista ai capi del condensatore C_i dopo aver soppresso i generatori, e –nel caso di capacità interagenti- dopo aver sostituito con cortocircuiti le capacità che interagiscono con quella in esame. Vanno inoltre ignorate, considerandole come circuiti aperti, tutte le capacità parassite del transistor, cioè tutte quelle capacità che determinano la frequenza di taglio superiore.

DETERMINAZIONE DELLA FREQUENZA DI TAGLIO SUPERIORE RELATIVA A UNA SINGOLA CAPACITA' C_i

- Bisogna prima verificare se le capacità che concorrono alla frequenza di taglio complessiva sono interagenti o non interagenti: una capacità C_1 è interagente con una capacità C_2 , se C_2 fa parte dell'impedenza che si vede fra i terminali di C_1
- Si determinano le singole frequenze di taglio secondo la relazione:

$$f_{tot} = \frac{1}{2\pi \cdot R_{eq} \cdot C} = \frac{1}{2\pi \cdot \tau_i}$$

avendo presente che R_{eq-i} è la resistenza equivalente vista ai capi del condensatore C_i dopo aver soppresso i generatori, e –nel caso di capacità interagenti- dopo aver sostituito con circuiti aperti le capacità che interagiscono con quella in esame. Vanno inoltre ignorate, considerandole come cortocircuiti, le capacità di accoppiamento e di bypass, cioè tutte quelle capacità che determinano la frequenza di taglio inferiore.

FREQUENZA DI TAGLIO INFERIORE COMPLESSIVA

Risulta in ogni caso **maggiore o uguale** alla **più alta** delle singole frequenze di taglio

1. capacità **NON** interagenti

- capacità non interagenti con una **frequenza di taglio DOMINANTE**, cioè molto **più grande (almeno quattro volte)** di tutte le altre:
la **frequenza di taglio complessiva** è circa uguale alla **dominante**, cioè a quella di gran lunga più elevata:

$$f_{Ltot} = f_{dom}$$

- capacità non interagenti **SENZA una frequenza di taglio dominante**:
la **frequenza di taglio complessiva** è data dalla **media geometrica** delle singole frequenze di taglio e risulta **circa uguale alla più alta** delle singole frequenze di taglio:

$$f_{Ltot} = \sqrt{f_{L1}^2 + f_{L2}^2 + \dots + f_{LN}^2}$$

2. capacità **INTERAGENTI**

La frequenza di taglio complessiva si calcola come semplice somma:

$$f_{Ltot} = f_{L1} + f_{L2} + \dots + f_{LN}$$

e risulta **abbastanza più grande** della **più elevata** delle singole frequenze di taglio

AMPLIFICATORI DI POTENZA E CLASSI DI AMPLIFICAZIONE

Il problema delle classi di amplificazione riguarda gli amplificatori di potenza, in quanto gli amplificatori di segnale lavorano tutti in quella che chiameremo “classe A”.

Generalmente sono detti **amplificatori DI POTENZA** o amplificatori **FINALI** quelli che forniscono al carico potenze **SUPERIORI a 0,1W** (>100mW).

Gli amplificatori di potenza possono arrivare a fornire alcuni KW.

Gli amplificatori di potenza svolgono la funzione di fornire al carico la potenza richiesta e, nei loro componenti attivi, il segnale presenta ampie escursioni sia in corrente che in tensione, per cui tali componenti non possono più essere rappresentati col modello per piccoli segnali.

In relazione al loro campo di applicazione, gli amplificatori di potenza (il cui nome abbrevieremo con AdP) possono essere classificati in:

- **AdP AUDIO**
che trattano il segnale ad audiofrequenza, rendendolo adatto a pilotare gli altoparlanti, che sono carichi a bassa impedenza (4-8) Ω;
- **AdP a RF (RADIOFREQUENZA)**
che forniscono, ai segnali appartenenti alle gamme HF, VHF, UHF ecc, la potenza necessaria ad essere irradiati da un’antenna, che ne costituisce il carico; gli AdP a RF sono amplificatori selettivi impiegati nei trasmettitori e nei ricevitori radio e TV
- **AdP per APPLICAZIONI INDUSTRIALI**
che generalmente lavorano in bassa frequenza e sono adatti al pilotaggio di attuatori (convertitori di potenza elettrica in potenza meccanica), come servomotori, motori passo-passo ecc.

Per l’analisi delle potenze che coinvolgono l’amplificatore, i suoi componenti attivi (transistor) e il carico, si definiscono due parametri: il **rendimento di conversione** (o **efficienza**) e il **fattore di merito** (o **figura di merito**).

RENDIMENTO

$$\eta = \frac{\text{potenza utile che il segnale fornisce al carico}}{\text{potenza che l'amplificatore assorbe dall'alimentatore}^*}$$

$$\eta = \frac{P_L}{P_{ALIM}} = \frac{P_U}{P_{CC}}$$

* uguale in continua e in alternata

FIGURA DI MERITO

$$F = \frac{\text{massima potenza dissipata dal componente attivo dell'amplificatore}}{\text{massima potenza che l'amplificatore può fornire al carico}} = \frac{P_{DMAX}}{P_{Lmax}}$$

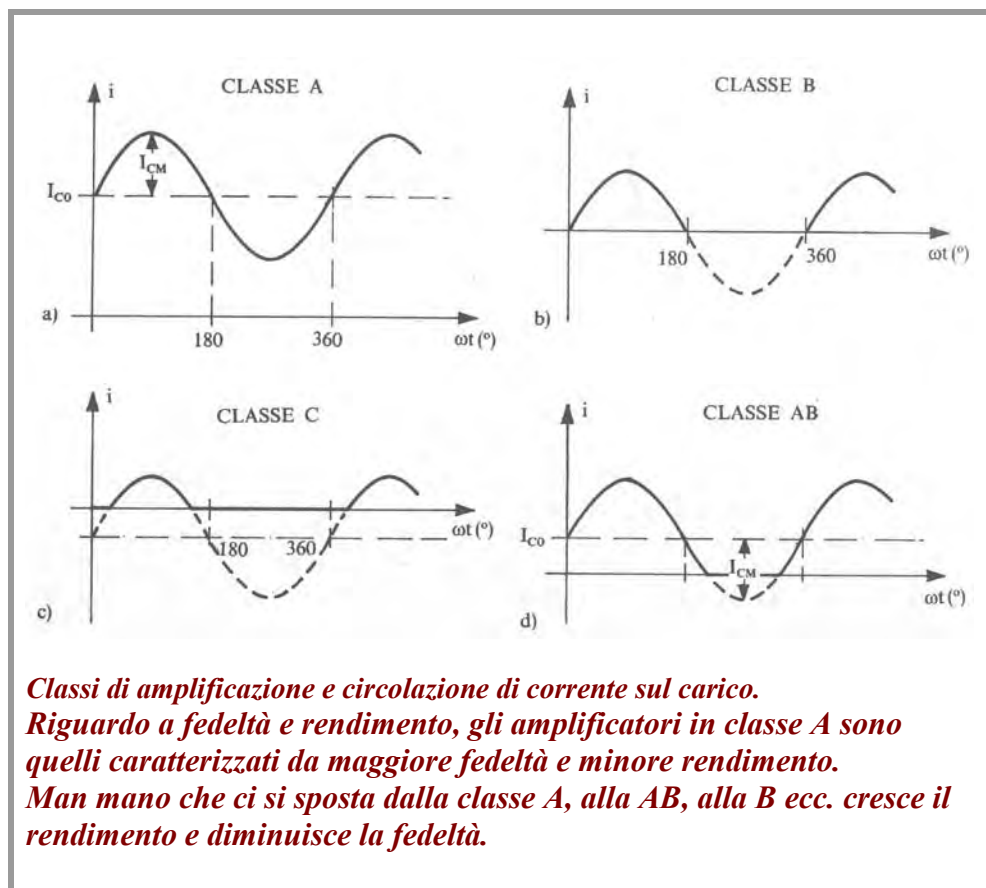
Vedremo ora come si determina la classe di funzionamento di un amplificatore.

Si suppone di applicare all'ingresso dell'amplificatore un segnale sinusoidale e si dice T il suo periodo (al quale corrisponde, su un asse ωt , un angolo di 2π), quindi si verifica se l'**intervallo di conduzione T'** , durante il quale il quale scorre **corrente nel circuito di uscita**, coincide con il periodo T del segnale di ingresso oppure è una frazione di T . Se, per i grafici dei segnali non facciamo riferimento all'asse t , ma all'asse ωt , parleremo non di intervalli temporali, ma di angoli di periodicità e di angoli di circolazione della corrente (angoli di conduzione). Chiameremo pertanto **angolo di circolazione** la **frazione di periodo angolare** durante la quale vi è **circolazione di corrente sul carico**.

Il segnale sinusoidale di ingresso, descrivendo un intero ciclo, ha un angolo di periodicità di 2π , mentre la corrente di uscita avrà, a seconda della classe di appartenenza dell'amplificatore, un angolo di circolazione α uguale a oppure minore di 2π . Premesso che il punto di lavoro a riposo o punto di riposo Q (Quiescent point) del transistor rappresenta i valori di tensione di uscita V_{CE} e corrente di uscita I_C del transistor in condizioni di riposo, cioè in assenza di segnale, possiamo dire che:

Il tempo (o l'angolo) di circolazione della corrente di uscita e, quindi, la classe di appartenenza dell'amplificatore, dipendono:

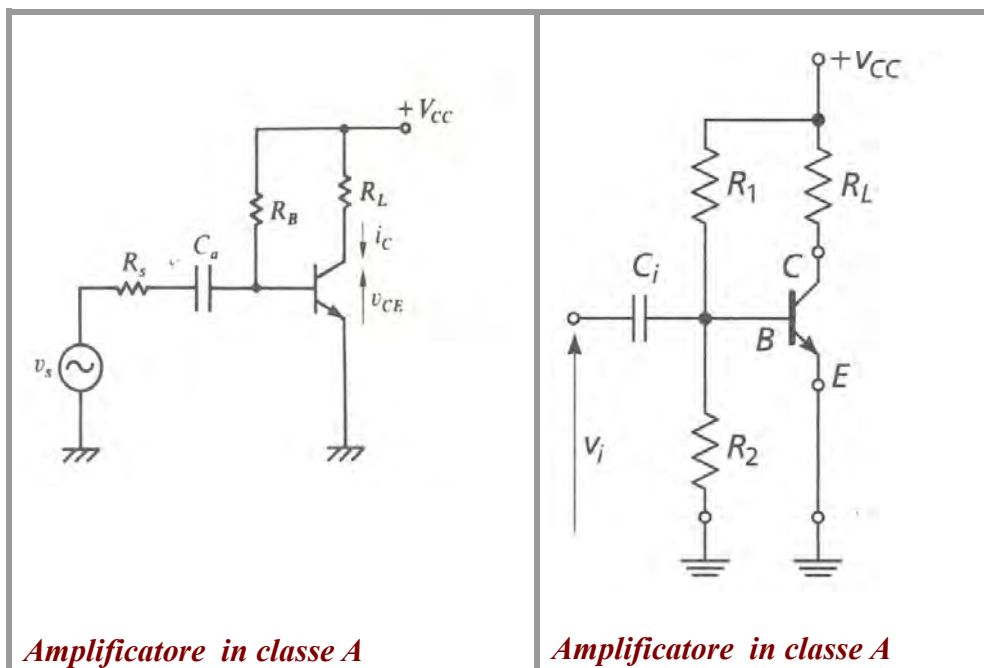
- dalla polarizzazione del componente attivo (per es. transistor) dell'amplificatore e quindi dalla posizione del punto di lavoro a riposo $Q = (I_{C0}, V_{CE0})$
- dall'ampiezza del segnale di ingresso.



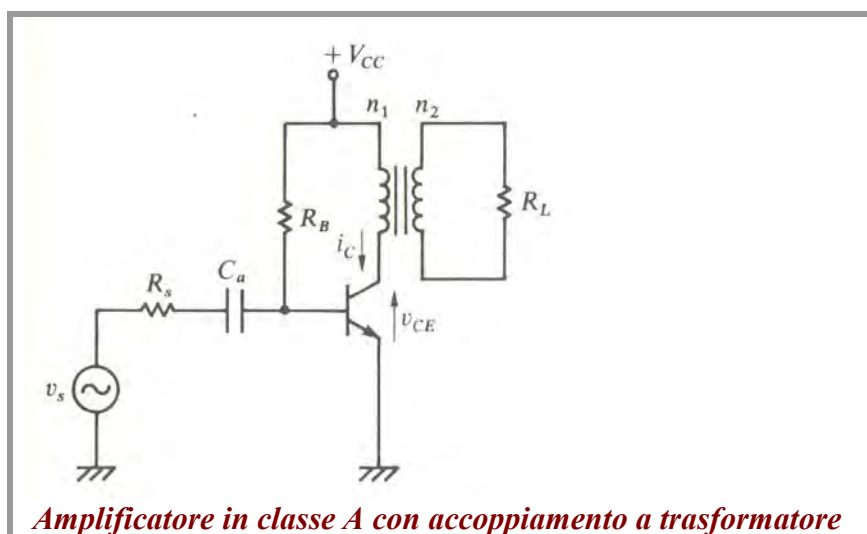
Come è illustrato nel grafico, l'amplificatore in **classe A** ha un intervallo T' di circolazione della corrente di uscita **uguale al periodo T** del segnale di ingresso ($T = T'$), il che significa che ha un **angolo di circolazione α** della corrente di uscita **pari a 2π** .

Ciò è dovuto al fatto che il punto di riposo Q_A e l'ampiezza del segnale di ingresso sono tali da far lavorare l'amplificatore in condizioni di linearità.

Il punto di riposo si trova infatti a un livello di tensione continua pari a circa la metà della tensione di alimentazione, il che significa che, in presenza di segnale (di opportuna ampiezza), questo punto potrà oscillare intorno alla posizione di riposo senza che la sua escursione venga bloccata.



Gli amplificatori per piccoli segnali (preamplificatori o amplificatori di segnale) lavorano in classe A. Sono realizzati in classe A anche alcuni amplificatori per alta fedeltà .

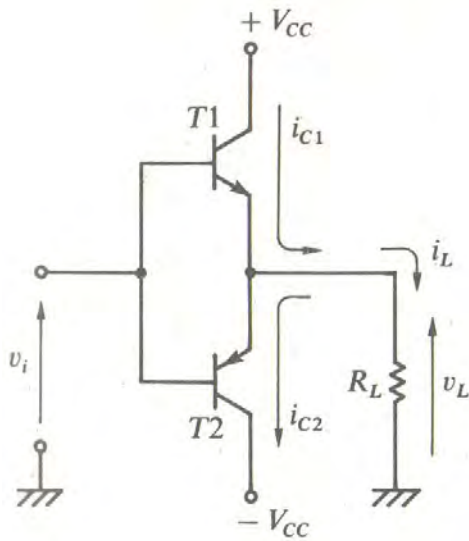


L'amplificatore in **classe B** ha un **intervallo T'** di circolazione della corrente di uscita **pari esattamente alla metà** del periodo T del segnale di ingresso (**$T' = T/2$**), il che significa che ha un **angolo α** di circolazione della corrente di uscita **pari a π** .

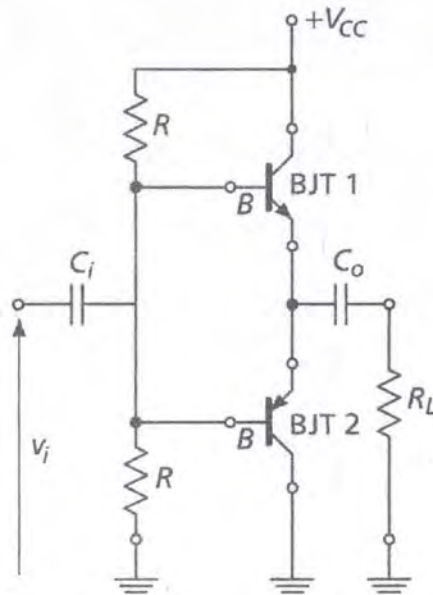
Ciò è dovuto al fatto che il componente attivo è polarizzato all'interdizione . Il punto di riposo Q_B infatti si trova al livello zero della corrente di collettore, il che significa che, in presenza di segnale, l'escursione del punto di lavoro viene bloccata a metà, ed esattamente che le semionde negative della corrente di uscita vengono tagliate.

Per fare in modo che la corrente sul carico circoli anche nel corso della semionda negativa del segnale di ingresso si utilizzano **due** transistor "complementari: uno di tipo **nnp** (che conduce durante la semionda positiva dell'ingresso) e uno di tipo **pnp** (che conduce durante la semionda negativa dell'ingresso).

Notiamo che il fatto che il **punto di riposo sia nella zona di interdizione** comporta che la corrente di riposo è praticamente nulla e quindi che la **potenza dissipata a riposo è trascurabile**.

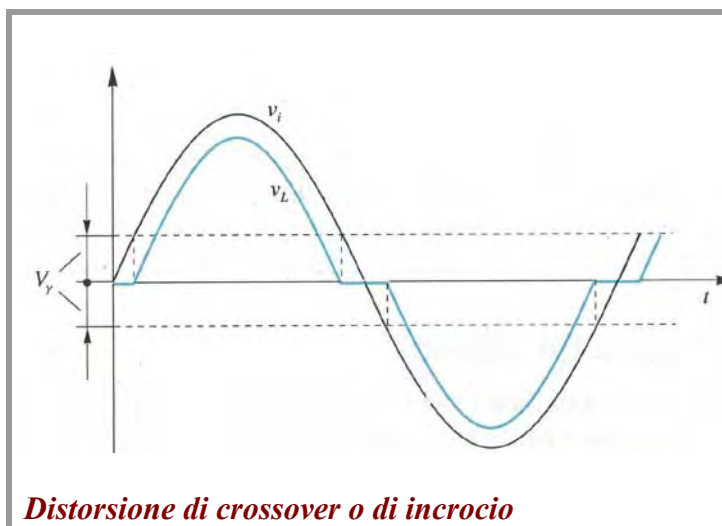


*Amplificatore in classe B con alimentazione duale.
Si noti la presenza di due componenti attivi*



*Amplificatore in classe B con alimentazione singola.
Si noti la presenza del condensatore in serie al carico, che svolge le funzioni dell'alimentazione negativa $-V_{CC}$*

Gli amplificatori in classe B sono caratterizzati da una **distorsione detta “di incrocio” o di “crossover”**. Questa distorsione è dovuta al fatto che la corrente comincia a scorrere sul carico solo quando la giunzione base-emettitore dei due transistor ha superato la tensione di soglia. Siccome in classe B i transistor sono polarizzati all’interdizione, deve trascorrere un piccolo intervallo di tempo (a partire dall’inizio della semionda del segnale di ingresso) perché i BJT entrino in conduzione e abbia inizio la circolazione di corrente sul carico. Ciò comporta che le semionde positive e negative sul carico non siano contigue, ma separate da un piccolo intervallo di tempo.



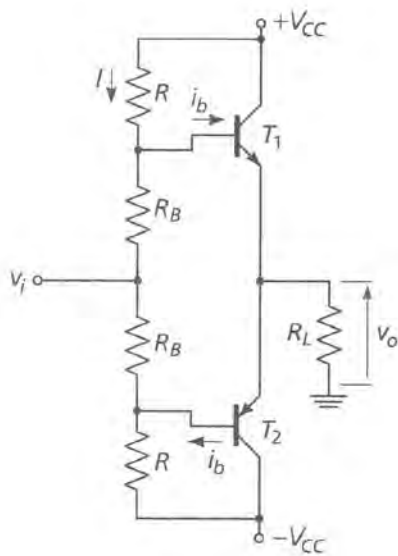
Per evitare la **distorsione detta di “crossover”** si evita di polarizzare l’amplificatore esattamente all’interdizione, cioè in classe B “pura”, e si preferisce fissare **il punto di riposo un po’ al di sopra dell’interdizione**, in modo che, appena hanno inizio le semionde del segnale di ingresso, i BJT siano già pronti alla conduzione. In pratica si fa in modo che, anche in assenza di segnale, le giunzioni base-emettitore siano già sul valore di tensione soglia, grazie all’interposizione di due resistenze, o meglio di due diodi, fra le due basi dei transistor.

Gli amplificatori polarizzati in questo modo vengono detti **“in classe AB”**.

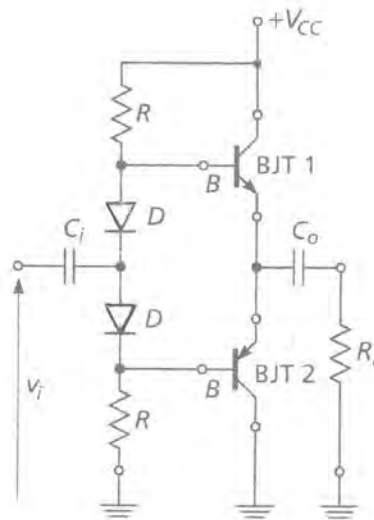
Essi hanno:

- un periodo T di circolazione della corrente di uscita **poco maggiore della metà del periodo** del segnale di ingresso ($T/2 < T' < T$)
- un angolo di circolazione α della corrente **poco maggiore di π** : ($\pi < \alpha < 2\pi$).

Sono realizzati in classe AB quasi tutti gli amplificatori di potenza audio.

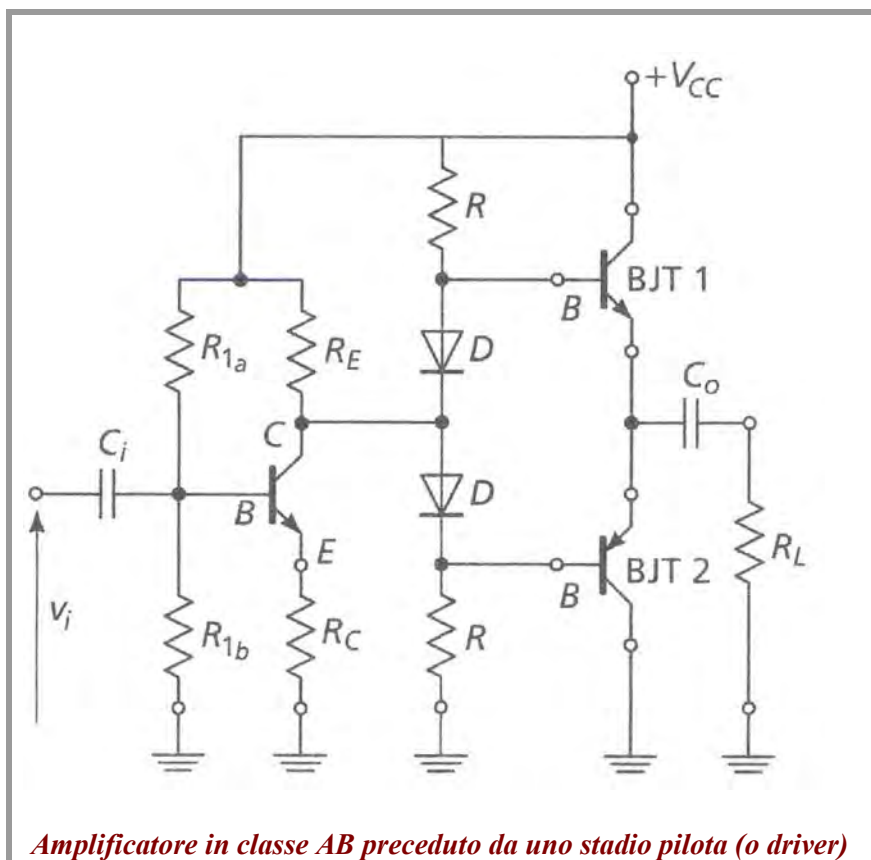


Amplificatore in classe AB con alimentazione duale. Si noti la presenza delle due resistenze R_B per l'eliminazione della distorsione di crossover



Amplificatore in classe AB con alimentazione singola. Si noti la presenza dei due diodi (per l'eliminazione della distorsione di crossover e del condensatore in serie al carico

AMPLIFICATORE IN CLASSE AB “COMPLETO”



Amplificatore in classe AB preceduto da uno stadio pilota (o driver)

L'amplificatore in **classe C** ha un intervallo T' di circolazione della corrente di uscita **minore della metà** del periodo T del segnale di ingresso ($T' < T/2$), il che significa che ha un angolo di circolazione della corrente di uscita **minore di π** .

Ciò è dovuto al fatto che il componente attivo è **polarizzato al di sotto dell'interdizione**.

Il punto di riposo Q_C infatti si trova al di sotto del livello zero della corrente di collettore, il che significa che, in presenza di segnale, vengono tagliate non solo le semionde negative della corrente di uscita, ma anche, parzialmente, le semionde positive.

La **corrente di uscita** ha quindi un **andamento impulsivo**.

Per ricostruire, a partire dall'uscita distorta, un'onda sinusoidale, viene normalmente utilizzato come carico (si vedano gli amplificatori a RF) un circuito risonante.

Sono realizzati in classe C gli **amplificatori RF finali** (di potenza) per **trasmettitori FM e CW** (Continuous Wave, per radioamatori).

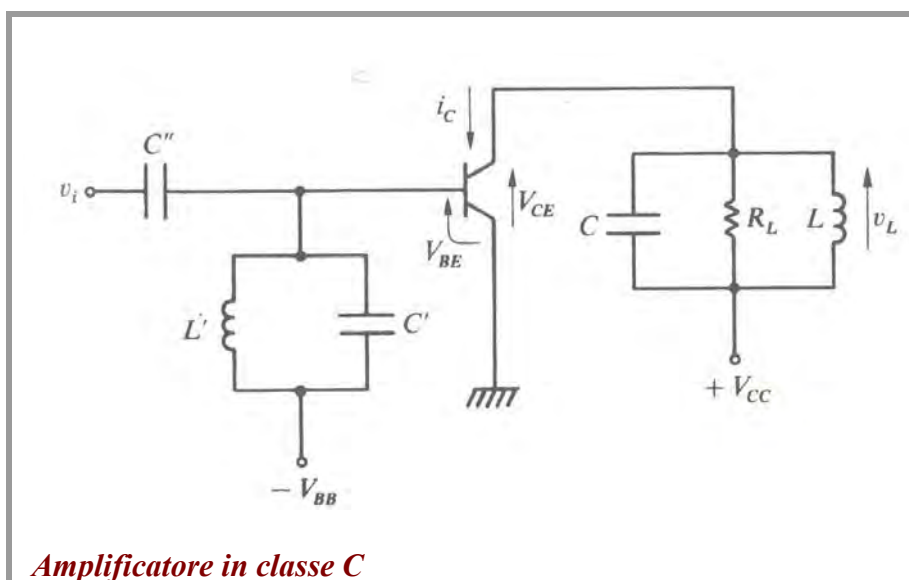
La classe C viene impiegata quasi esclusivamente per gli amplificatori a radiofrequenza perché questi dispositivi, essendo selettivi, non amplificano le armoniche successive presenti nella forma d'onda, ma amplificano solo la fondamentale (prima armonica), fornendo quindi in uscita un segnale sinusoidale, anche se il dispositivo attivo è polarizzato all'interdizione.

Il rendimento, tanto maggiore quanto più piccolo è l'angolo di circolazione della corrente di uscita, può avvicinarsi al 100%.

Per potenze inferiori alle centinaia di W gli amplificatori di potenza a RF sono realizzati con transistor.

Per potenze **superiori alle centinaia di W** gli amplificatori di potenza a RF sono realizzati **con tubi a vuoto** (valvole).

Non possono essere invece di classe C gli amplificatori audio in quanto essi utilizzano due componenti attivi funzionanti alternativamente i quali, in classe C, sarebbero entrambi interdetti per una frazione di periodo. Sarebbe pertanto impossibile ricostruire l'onda sinusoidale completa.



ALTRE CLASSI DI FUNZIONAMENTO

Sono state definite anche altre classi di funzionamento e, in particolare, la **classe D** (rendimento del 100% circa) che sfrutta tecniche di modulazione con portante impulsiva (e successiva demodulazione), e in particolare la **PWM (Pulse Width Modulation)** o “modulazione a durata di impulsi” o “modulazione a larghezza di impulsi”.

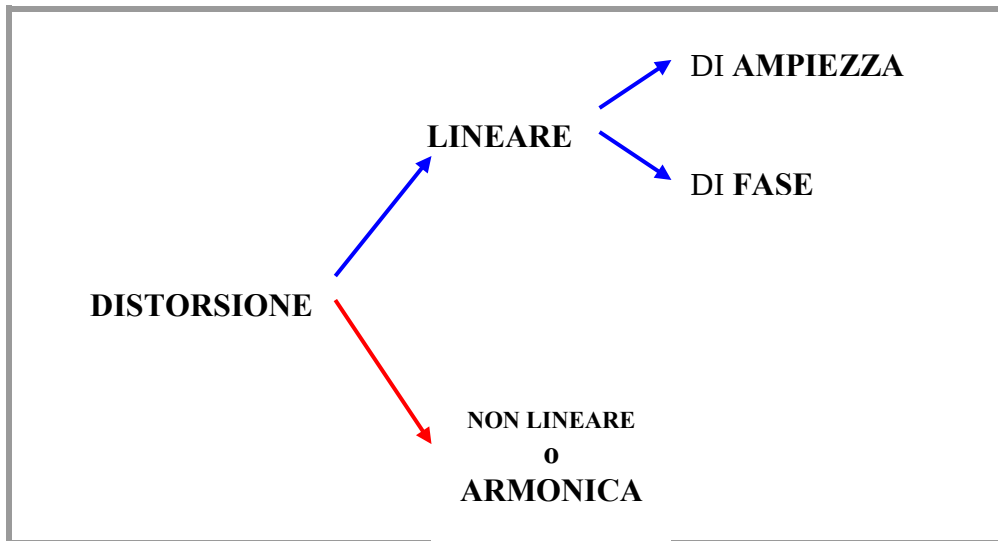
(Questa modulazione è anche nota come PDM o PLM).

L'uscita di un modulatore PWM è un treno di impulsi di ampiezza e frequenza costante, ma con duty-cycle dipendente dall'ampiezza del segnale modulante.

L'amplificatore PWM in classe D funziona nel modo seguente:

- 1) il segnale di ingresso, da amplificare, viene applicato, come modulante, a un modulatore PWM che fornisce in uscita un treno di impulsi con ampiezza e frequenza costanti e duty-cycle (e quindi durata degli impulsi alti) proporzionale all'ampiezza del segnale di ingresso, cioè del segnale modulante da amplificare;
- 2) il treno di impulsi viene amplificato da un amplificatore a BJT o a MOS, i quali vengono fatti funzionare in commutazione
- 3) il treno di impulsi, amplificato, viene demodulato e il risultato della demodulazione è proprio il segnale di partenza amplificato.

DISTORSIONE



La distorsione è **un'alterazione indesiderata** della **forma** di un segnale e può essere **lineare** o **non lineare**.

La distorsione **lineare** è dovuta a effetti **capacitivi o induttivi** e può interessare sia lo **spettro di ampiezza del segnale** (distorsione lineare di ampiezza), sia lo **spettro di fase** (distorsione lineare di fase).

La distorsione **non lineare** o distorsione armonica è determinata dalla presenza di componenti non lineari (diodi, transistor ecc.)

DISTORSIONE LINEARE

La distorsione lineare è determinata dalla presenza di **condensatori e/o induttori** nel circuito (componenti lineari passivi reattivi) e dalle **capacità parassite** dei componenti attivi.

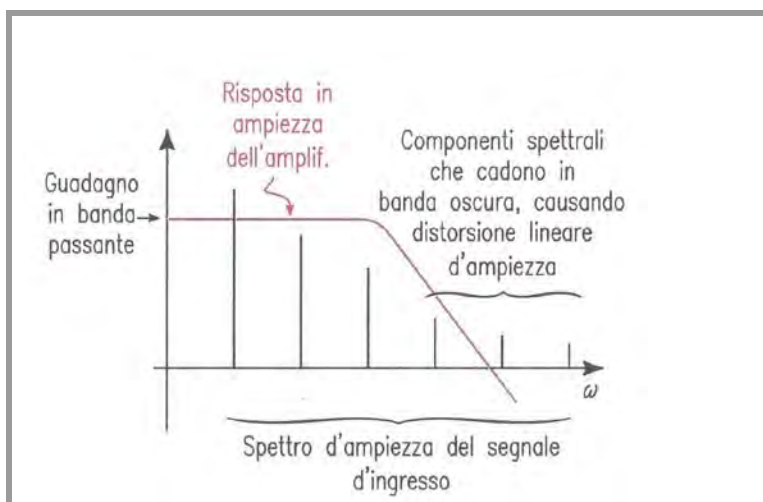
DISTORSIONE LINEARE DI AMPIEZZA

Si dice che un segnale ha subito una distorsione lineare di ampiezza quando **le armoniche**, e quindi le righe spettrali, **non sono state amplificate (o attenuate) nella stessa misura**.

Il segnale di uscita di un dispositivo che distorce linearmente presenta le stesse armoniche del segnale di ingresso, ma i rapporti fra le ampiezze delle diverse armoniche non sono più quelli originari. Di conseguenza si ha un'alterazione di forma del segnale.

Affinché un dispositivo non determini distorsione lineare di ampiezza, esso deve moltiplicare o dividere l'ampiezza di tutte le armoniche per uno stesso fattore costante k .

E' importante sottolineare che la distorsione lineare non genera nuove armoniche: l'uscita di un circuito che distorce linearmente non presenta ulteriori armoniche rispetto a quelle già presenti nel segnale di ingresso. Di conseguenza se il segnale di un ingresso è sinusoidale, e quindi se il suo spettro è formato da una sola riga, anche il segnale di uscita deve presentare una sola riga e risulta perciò ancora sinusoidale.



Come si contrasta la distorsione lineare di ampiezza?

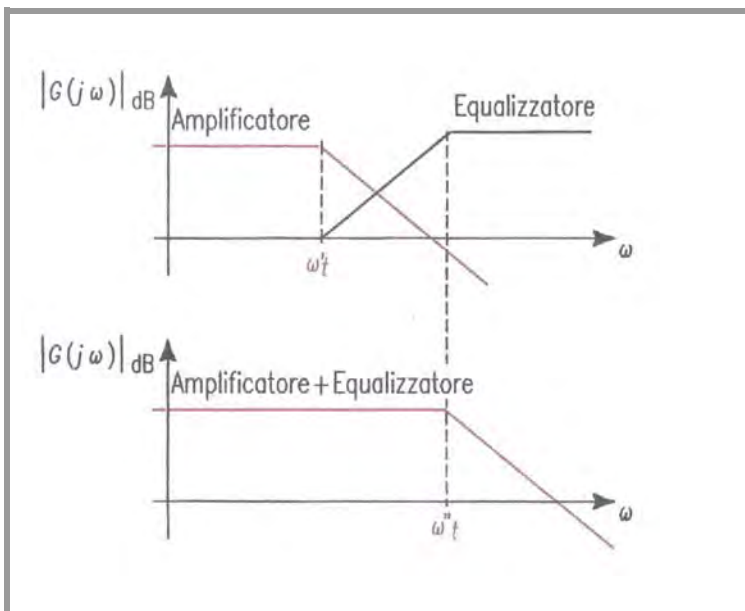
La risposta in frequenza di un circuito che contiene elementi capacitivi e/o induttivi non è piatta per ogni frequenza, ma presenta una banda passante (all'interno della quale la curva è piatta) e una diminuzione a sinistra e/o a destra della banda passante. La distorsione in esame si verifica proprio a quelle frequenze nelle quali la curva di risposta non è piatta, ma presenta una certa pendenza.

Pertanto, se vogliamo che il segnale di ingresso non subisca distorsione lineare di ampiezza, dobbiamo fare in modo che la sua frequenza venga a trovarsi all'interno della banda passante del circuito, cioè all'interno dell'intervallo di frequenze dove la risposta può essere considerata piatta.

La soluzione è quindi l'**allargamento della banda passante** del dispositivo, che può essere realizzato mediante un **filtro equalizzatore**, posto prima o dopo l'amplificatore.

L'equalizzatore effettua una "compensazione in frequenza", nel senso che ha la funzione di correggere la risposta del dispositivo distorto spostandone verso le alte frequenze il taglio superiore e verso le basse frequenze il taglio inferiore

Per esempio si può allargare la banda di un amplificatore con risposta di tipo passabasso, mediante il collegamento a un filtro passaalto.



DISTORSIONE LINEARE DI FASE

Anche la distorsione di fase, essendo lineare, non genera nuove armoniche: l'uscita di un circuito che distorce linearmente non presenta ulteriori armoniche rispetto a quelle del segnale di ingresso. **L'orecchio umano avverte la distorsione di fase in maniera estremamente ridotta rispetto a quella di ampiezza perché è più sensibile al contenuto spettrale del suono che alla sua forma d'onda.**

Perciò, per gli amplificatori audio la distorsione di fase è molto meno importante di quanto non lo sia in altre applicazioni, come, per esempio, nell'amplificazione video.

Se un amplificatore (con amplificazione **A**) sfasa di un angolo α un segnale sinusoidale

$$v_i(t) = V_i \cdot \text{sen}(\omega t)$$

il corrispondente segnale di uscita è

$$v_o(t) = A \cdot V_i \cdot \text{sen}(\omega t + \alpha)$$

che si può anche scrivere come:

$$v_o(t) = A \cdot V_i \cdot \text{sen} \left[\omega \left(t + \frac{\alpha}{\omega} \right) \right]$$

Questo segnale di uscita presenta un ritardo temporale

$$\Delta t = \frac{\alpha}{\omega}$$

cioè è traslato lungo l'asse dei tempi di Δt rispetto al segnale di ingresso.

Qualora il segnale applicato all'ingresso dell'amplificatore sia formato da più armoniche e qualora ciascuna armonica subisca, a causa dell'amplificatore, un differente ritardo temporale, allora la forma del segnale di uscita risulta alterata.

Per conservare la forma d'onda del segnale di ingresso è necessario che tutte le armoniche subiscano lo stesso ritardo temporale, cioè che il Δt risulti costante per ogni armonica e quindi per ogni valore di frequenza e di pulsazione. E' pertanto necessario che si abbia:

$$\Delta t = k \quad \forall \omega$$

$$\frac{\alpha}{\omega} = k \quad \forall \omega$$

$$\alpha = k \cdot \omega$$

L'ultima espressione ci dice che **(affinché il ritardo temporale sia costante con la frequenza, e quindi affinché non si abbia distorsione di fase) lo sfasamento angolare** introdotto dall'amplificatore deve **variare linearmente con la frequenza** oppure valere **0° o 180°**. Anche la distorsione di fase può essere contrastata mediante¹ un **filtro equalizzatore**.

¹ Osserviamo che, se si elimina con un equalizzatore la distorsione di ampiezza in una certa banda, rendendo piatto il modulo della risposta in frequenza del dispositivo distorcente, ciò determina anche l'eliminazione della distorsione di fase solo se, in quella banda, la fase della risposta in frequenza risulta uguale a 0° oppure a 180°, oppure se risulta linearmente variabile con la frequenza.

In conclusione, **affinché un amplificatore non produca distorsione lineare**, né di ampiezza né di fase, l'**amplificazione** deve avere un **modulo costante** e una **fase variabile linearmente con la frequenza**.

DISTORSIONE ARMONICA O NON LINEARE

La distorsione non lineare è la **presenza**, nel segnale di uscita, di **armoniche non contenute nello spettro del segnale di ingresso**.

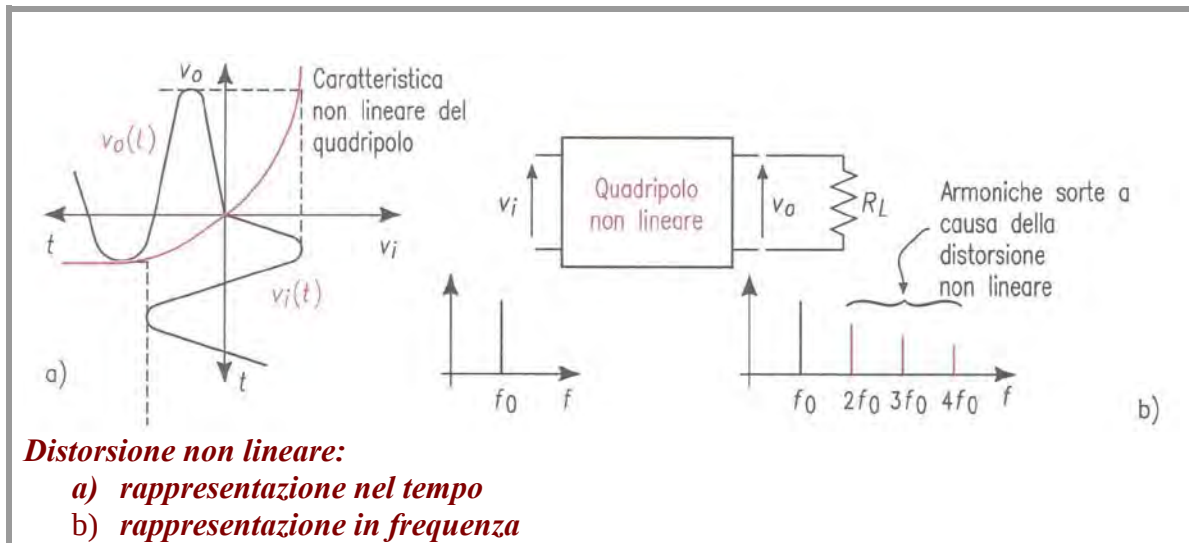
Un circuito produce una distorsione non lineare quando contiene elementi non lineari, come **diodi** o **transistor**, cioè componenti nei quali la relazione fra tensione e corrente non è graficamente rappresentabile mediante una retta.

Il funzionamento di un componente o di un dispositivo è descrivibile mediante una “curva caratteristica”, per esempio una curva che fornisce l’andamento della corrente in funzione della tensione. Se il componente è lineare, la sua curva caratteristica è una retta.

Se invece è non lineare, la sua curva caratteristica non è una retta, ma è per esempio di tipo quadratico o esponenziale. I diodi e i transistor hanno caratteristiche non lineari, se però ad essi sono applicati segnali di piccola ampiezza, questi vanno a interessare solo piccoli tratti delle curve caratteristiche, tratti che, per la loro piccolezza, possono essere approssimati a segmenti rettilinei: questo è ciò che accade negli amplificatori per piccoli segnali, dai quali è più agevole ottenere un comportamento lineare.

Diverso è invece il caso degli amplificatori di potenza perché al loro ingresso vengono applicati segnali di grande ampiezza. Di conseguenza i transistor sono costretti a lavorare in corrispondenza di tratti non lineari della loro caratteristica e si presenta il problema di eliminare o limitare la distorsione non lineare.

Sono proprio gli **amplificatori di potenza** i circuiti per i quali è fondamentale risolvere il problema della **distorsione armonica**. La distorsione di un amplificatore o di un filtro **si misura** ponendo in **ingresso** al quadripolo un **segnale puramente sinusoidale** (ossia un segnale il cui spettro è formato da una sola riga) e **analizzando lo spettro del segnale di uscita**.



PARAMETRI DELLA DISTORSIONE ARMONICA O NON LINEARE

DISTORSIONE DI ENNESIMA ARMONICA

$$D_n = \frac{V_n}{V_1} = \frac{\text{valore efficace dell'ennesima armonica}}{\text{valore efficace della prima armonica}}$$

Generalmente la distorsione si esprime in percentuale:

$$D_n = \frac{V_n}{V_1} \cdot 100$$

Particolarmente importante è la distorsione di terza armonica:

$$D_3 = \frac{V_3}{V_1} \cdot 100 = \frac{\text{valore efficace della terza armonica}}{\text{valore efficace della prima armonica}} \cdot 100$$

THD: DISTORSIONE ARMONICA TOTALE (TOTAL HARMONIC DISTORTION)

La distorsione armonica totale è il **valore quadratico medio delle distorsioni di ennesima armonica**, cioè delle distorsioni relative alle singole armoniche:

$$\text{THD} = D_{\text{TOT}} = \sqrt{(D_2^2 + D_3^2 + \dots + D_n^2)}$$

Oppure, in percentuale:

$$\text{THD}_{\%} = D_{\text{TOT}\%} = \sqrt{(D_2^2 + D_3^2 + \dots + D_n^2)} \cdot 100$$

Un ulteriore significato della THD si può ottenere sostituendo alle D_n le loro espressioni:

$$\text{THD} = \sqrt{\left(\frac{V_2^2}{V_1^2} + \frac{V_3^2}{V_1^2} + \dots + \frac{V_n^2}{V_1^2} \right)} \quad \text{da cui:}$$

$$\text{THD} = \sqrt{\frac{V_2^2 + V_3^2 + \dots + V_n^2}{V_1^2}} \quad \text{e quindi:}$$

$$\text{THD} = \frac{\sqrt{V_2^2 + V_3^2 + \dots + V_n^2}}{V_1} = \frac{\sqrt{\text{val. efficace complessivo delle armoniche superiori}}}{\text{valore efficace della prima armonica}}$$

La THD può essere scritta in termini di potenze su una resistenza di carico R_L dividendo per R_L numeratore e denominatore della precedente espressione:

$$THD = \sqrt{\frac{\frac{V_2^2}{R_L} + \frac{V_3^2}{R_L} + \dots + \frac{V_n^2}{R_L}}{\frac{V_1^2}{R_L}}}$$

$$THD = \sqrt{\frac{P_2 + P_3 + \dots + P_n}{P_1}} = \sqrt{\frac{\text{potenza delle armoniche superiori}}{\text{potenza della prima armonica}}}$$

Le **norme DIN²** impongono per gli **amplificatori audio** ad **alta fedeltà** una **distorsione armonica totale (THD) inferiore all'1%** in corrispondenza della potenza dichiarata dal costruttore.

I dispositivi commerciali ad alta fedeltà si mantengono ampiamente al di sotto di questo limite, presentando valori di distorsione molto più bassi.

Va inoltre notato che:

- L'orecchio umano è più sensibile alle distorsioni di ordine dispari (distorsione da terza armonica, da quinta armonica ecc.) perché queste determinano suoni aspri e sgradevoli
- purché la distorsione di terza armonica sia inferiore al 2%, **l'udito umano avverte con certezza solo distorsioni totali maggiori del 5%**
- l'ampiezza delle armoniche è rapidamente decrescente con la frequenza delle armoniche stesse, per cui, spesso, è sufficiente tener conto solo della seconda e della terza armonica.

² Il DIN (Deutsches Institut für Normung) è l'ente tedesco per la standardizzazione ed è membro dell'ISO. L'ISO (International Organization for Standardization) è l'organizzazione internazionale per la standardizzazione ed è l'ente più importante a livello mondiale per la definizione di norme tecniche. Fondata nel 1947, ha sede a Ginevra. Membri dell'ISO sono gli organismi nazionali di standardizzazione. In Italia le norme ISO vengono recepite dall'UNI. L'ISO coopera strettamente con l'IEC, responsabile per la standardizzazione degli equipaggiamenti elettrici.

DISTORSIONE NON LINEARE DI INTERMODULAZIONE

Riguardo agli amplificatori audio questa distorsione si rivela piuttosto fastidiosa nella riproduzione di brani eseguiti da molti strumenti. E' però utilmente sfruttata per la realizzazione di "mixer" o circuiti mescolatori, presenti, ad esempio, nei radioricevitori eterodina.

La distorsione da intermodulazione è un effetto della distorsione **non lineare** e si produce quando lo spettro del segnale applicato in ingresso al dispositivo contiene **almeno due componenti** a frequenza diversa.

Questa distorsione si manifesta con la presenza, nel segnale di uscita, di componenti spettrali con **frequenze multiple** delle frequenze di ingresso e anche di componenti spettrali con valori di frequenza dati dalla **somma**, dalla **differenza**, e, in generale dalla **combinazione lineare** delle frequenze presenti in ingresso.

Se, per esempio, il segnale di ingresso contiene due righe spettrali con frequenze f_a ed f_p , le componenti del segnale di uscita avranno valori di frequenza del tipo $\pm k_1 \cdot n \cdot f_a \pm k_2 \cdot n \cdot f_p$ (purché diano luogo a un valore positivo):

Il segnale di uscita potrà contenere quindi le combinazioni di frequenze sotto elencate:

$$\begin{aligned} & f_a, 2f_a, 3f_a, \dots, nf_a, \\ & f_p, 2f_p, 3f_p, \dots, nf_p, \\ & f_a + f_p, f_a + 2f_p, f_a + 3f_p, \dots, f_a + nf_p, \\ & f_a - f_p, f_a - 2f_p, f_a - 3f_p, \dots, f_a - nf_p, \\ & f_p + f_a, f_p + 2f_a, f_p + 3f_a, \dots, f_p + nf_a, \\ & f_p - f_a, f_p - 2f_a, f_p - 3f_a, \dots, f_p - nf_a, \\ & \pm k_1 \cdot f_a \pm k_2 \cdot f_p \\ & \pm k_1 \cdot n \cdot f_a \pm k_2 \cdot n \cdot f_p. \end{aligned}$$

Vediamo ora come può essere motivato questo tipo di distorsione.

Abbiamo detto che la caratteristica di trasferimento dei componenti attivi non è lineare.

Per esempio la **corrente di uscita** i_0 di un BJT, di un FET, di un tubo a vuoto (triolo o pentodo) è legata alla **tensione di ingresso** v da una legge del tipo:

$$i_0 = k_0 + k_1 \cdot v + k_2 \cdot v^2 + k_3 \cdot v^3 + \dots + k_n \cdot v^n \quad (*)$$

che può essere approssimata a una legge quadratica se tralasciamo i termini contenenti potenze superiori alla seconda, che risultano di entità trascurabile:

$$i_0 = k_0 + k_1 \cdot v + k_2 \cdot v^2$$

Ipotizziamo allora di applicare all'ingresso del dispositivo non lineare una tensione v formata dalla somma di due componenti sinusoidali, una a frequenza f_a e una a frequenza f_p :

$$v = V_a \cdot \sin(\omega_a \cdot t) + V_p \cdot \sin(\omega_p \cdot t)$$

la corrente di uscita risulterà:

$$v = k_0 + k_1 \cdot [V_a \cdot \sin(\omega_a \cdot t) + V_p \cdot \sin(\omega_p \cdot t)] + k_2 \cdot [V_a \cdot \sin(\omega_a \cdot t) + V_p \cdot \sin(\omega_p \cdot t)]^2$$

La presenza di termini sinusoidali **elevati al quadrato** (e a potenze più elevate) determina la nascita di componenti di segnale a **frequenze multiple** di quella originaria, come è intuibile dall'applicazione di relazioni trigonometriche del tipo:

$$\sin^2 \alpha = \frac{1 - \cos(2\alpha)}{2}$$

Inoltre, dallo sviluppo delle potenze dei binomi, si ottengono **prodotti di termini sinusoidali**, del tipo:

$$V_a \cdot \sin(\omega_a \cdot t) + V_p \cdot \sin(\omega_p \cdot t)$$

A loro volta questi prodotti danno luogo alla nascita di componenti spettrali con **frequenza somma e/o differenza** di quelle originaria, come si rileva dall'applicazione di relazioni trigonometriche come quelle di Werner:

$$\sin \alpha \cdot \sin \beta = \frac{1}{2} \cdot [\cos(\alpha - \beta) - \cos(\alpha + \beta)]$$

Ulteriori termini con frequenze multiple e combinazione lineare delle frequenze applicate in ingresso compariranno nell'espressione della corrente di uscita se assumiamo, per il componente elettronico, non una legge quadratica, ma una legge, come la (*), che contiene anche potenze superiori alla seconda.

CAPITOLO 21

TRANSISTOR A EFFETTO DI CAMPO (FET) E AMPLIFICATORI A FET

COMPLEMENTI

IL JFET: GENERALITÀ

IL FUNZIONAMENTO DEL JFET

CIRCUITO EQUIVALENTE DINAMICO DEL JFET

***CIRCUITO DI POLARIZZAZIONE “SEMPLICE” O DI
AUTOPOLARIZZAZIONE DEL JFET:
DIMENSIONAMENTO***

***RETE DI POLARIZZAZIONE E STABILIZZAZIONE A
PARTITORE DEL JFET: ANALISI***

IL JFET IN COMMUTAZIONE

MOSFET A SVUOTAMENTO (DEPLETION)

***POLARIZZAZIONE DEL MOSFET A SVUOTAMENTO
(DEPLETION)***

MOSFET AD ARRICCHIMENTO (ENHANCEMENT)

***POLARIZZAZIONE DEL MOSFET AD ARRICCHIMENTO
(ENHANCEMENT)***

I MOSFET IN COMMUTAZIONE

CIRCUITO EQUIVALENTE DINAMICO DEI MOSFET

IL JFET: GENERALITA'

Il JFET, come il BJT, è un componente allo stato solido a tre terminali.

I terminali sono:

- Il GATE (che è il terminale di controllo, come la base lo è per il BJT)
- Il SOURCE (che potremmo paragonare all'emettitore del BJT)
- Il DRAIN (che potremmo paragonare al collettore del BJT).

Come ogni transistor a effetto di campo, il JFET è simmetrico, nel senso che drain e source sono intercambiabili. Si usa però chiamare **source** il terminale **dal quale le cariche elettriche libere entrano nel componente** e **drain** il terminale **verso il quale esse si dirigono per poi uscirne**.

Le principali applicazioni del JFET sono:

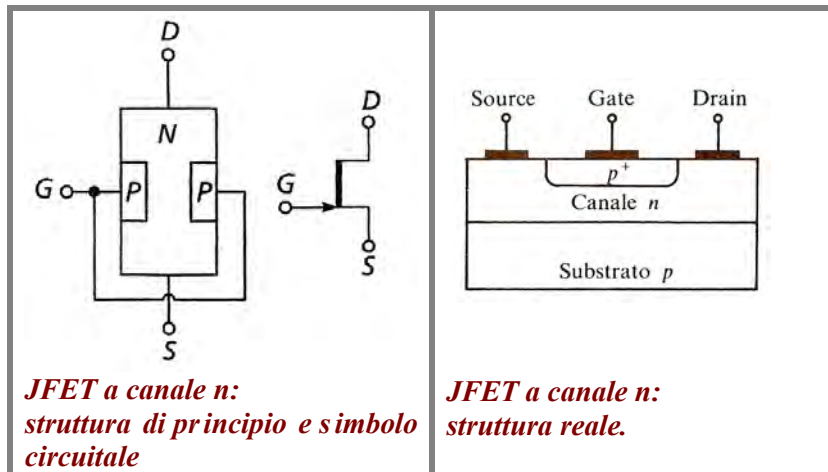
- **preamplificatori, microfoni amplificati, amplificatori** (per esempio: stadi iniziali differenziali di operazionali e di amplificatori audio);
- **oscillatori**;
- **interruttori analogici**, grazie alla mancanza di offset nello stato di massima conduzione, offset che invece è presente nel BJT sotto forma di $V_{CE(SAT)}$;
- **resistori controllati in tensione (VCR, Voltage Controlled Resistor)**; i VCR sono resistori variabili a tre terminali, nei quali la d.d.p. fra due piedini dipende dalla tensione applicata al terzo, cioè al terminale di controllo;
per una fissata V_{DS} , di basso valore, la resistenza del canale dipende dalla V_{GS} in maniera lineare, per cui sia il JFET che il MOSFET possono essere utilizzati come “resistenze controllate in tensione”, ossia come “VCR” (Voltage Controlled Resistor)
- **generatori di corrente costante**, grazie al fatto che le curve caratteristiche di uscita $I_D = I_D(V_{DS})$ sono orizzontali fra il punto di pinch-off e quello di rottura.

Il **JFET** può essere **di due tipi**:

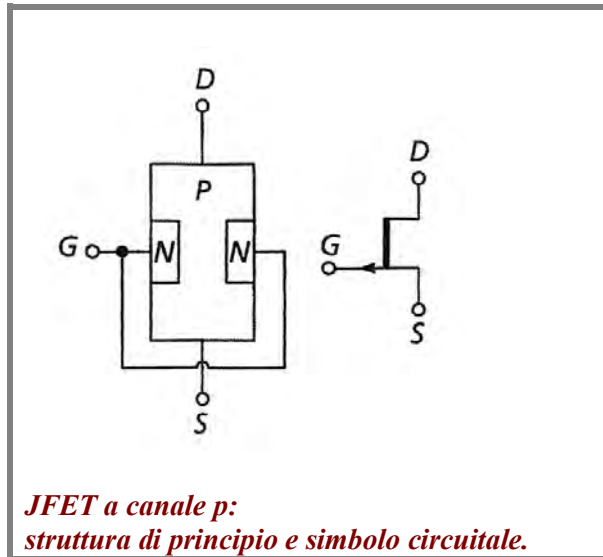
1. **JFET “a canale n”**: è una **barretta** di semiconduttore a **drogaggio n**, con **due zone laterali a forte drogaggio p** (che formano il **gate**).

La zona “n”, posta fra le due “isole p” è il **canale**.

Nel **simbolo circuitale**, la **freccia** che indica il **gate** è orientata **verso l'interno**.



2. **JFET “a canale p”**: è una barretta di semiconduttore a drogaggio p (il canale), con due isole laterali a forte drogaggio n (che costituiscono il gate). Nel **simbolo circuitale**, la **freccia** che indica il **gate** è orientata **verso l'esterno**.



Nel JFET a canale p il gate è formato da due “isole” a drogaggio n, mentre drain e source sono le estremità opposte del canale a drogaggio p.

IL FUNZIONAMENTO DEL JFET

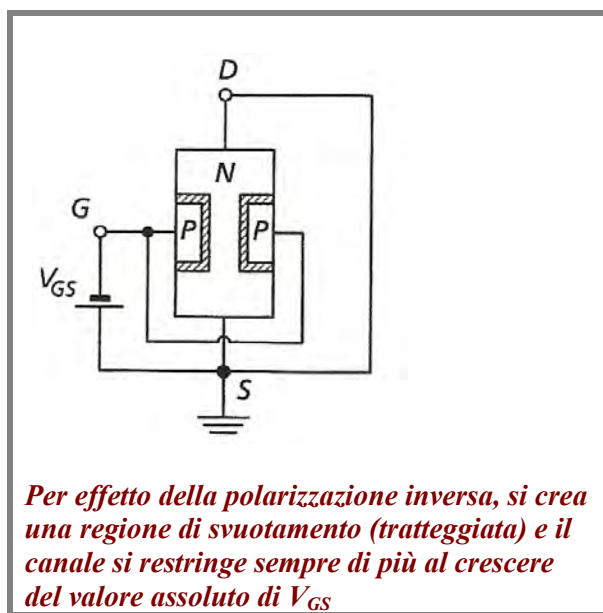
Nel normale funzionamento del JFET vengono simultaneamente applicate due tensioni di polarizzazione:

- la **tensione negativa** V_{GS} per la di **polarizzazione inversa** della **giunzione gate(p)/canale(n)**
- la **tensione positiva** V_{DS} applicata **ai capi del canale n** ossia **fra drain e source**. Essendo positiva, la V_{DS} porta il drain a un potenziale più alto di quello del source.

Semplicemente per comprendere meglio il funzionamento del JFET, esamineremo prima l'effetto della sola V_{GS} , poi l'effetto della sola V_{DS} e infine il funzionamento del componente in presenza di entrambe.

1. Comportamento con:

- **polarizzazione inversa del solo gate** rispetto al source, cioè con $V_{GS} < 0$
- **drain e source in cortocircuito**, ossia $V_{DS} = 0$

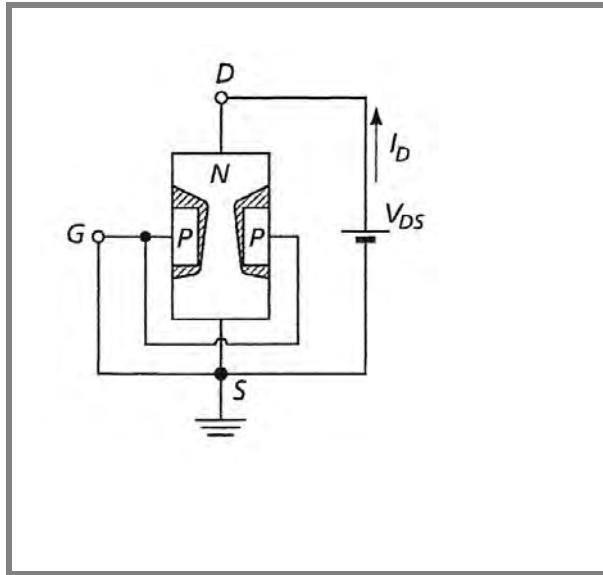


Per effetto della polarizzazione inversa, le cariche libere maggioritarie tendono ad abbandonare la giunzione. Gli elettroni della zona n, cioè del canale, si allontanano dalla superficie di giunzione, creando una zona priva di cariche libere (zona di svuotamento o depletion layer). Si determina quindi un **restringimento della sezione del canale** con **duttivo** originario e il suo **aumento di resistenza**.

Al crescere (in valore assoluto) della V_{GS} , la sezione utile del canale diminuisce sempre di più (e la resistenza cresce), fino a quando, **per un certo valore di V_{GS}** , detto “ V_p ” (o “**valore di pinch-off**” o ancora “**valore di strozzamento**”), il **canale** conduttivo di fatto **scompare**, rimpiazzato dalla regione di svuotamento, e presenta resistenza infinita. Nella situazione di pinch-off, anche se si applicasse una tensione fra D ed S, non potrebbe scorrere corrente.

2. Comportamento con:

- **polarizzazione del solo drain** rispetto al source, cioè con $V_{DS} > 0$
- **gate in cortocircuito** con il **source**, ossia $V_{GS} = 0$



La $V_{DS} > 0$ determina lo scorrimento di una corrente I_D , nel canale, con verso convenzionale uscente dal generatore.

Essendo una estremità del canale n (quella superiore nella figura) collegata al polo positivo del generatore ed essendo le isole p (costituenti il gate) a massa, la giunzione pn risulta inversamente polarizzata.

La polarizzazione inversa è più forte nel punto di applicazione della $V_{DS} > 0$ e si indebolisce man mano che ci si allontana da tale punto.

Si ha quindi un allontanamento di cariche libere maggioritarie, con la conseguente formazione di una regione di svuotamento, che sarà più pronunciata nella regione del canale più vicina al drain (regione superiore nella figura).

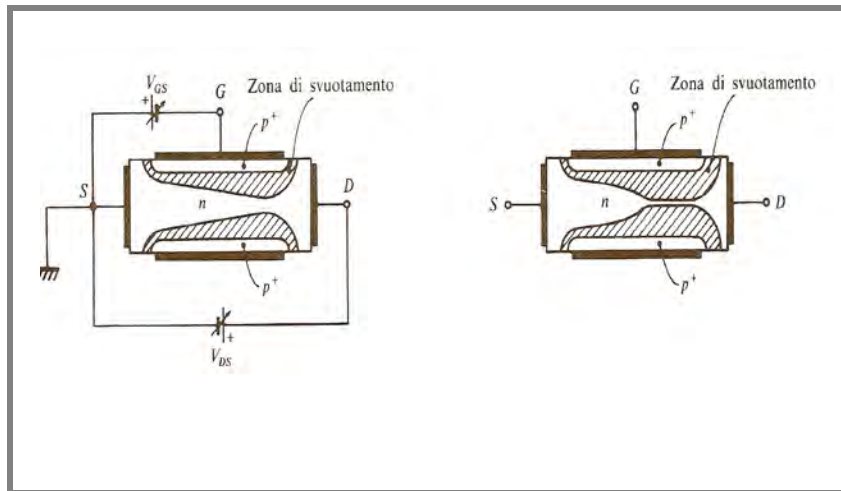
Al crescere della V_{DS} , aumentano sia la corrente I_D , sia l'entità della zona di svuotamento (la corrente aumenta anche se il canale conduttivo si restringe). Appena però la V_{DS} tocca il valore " V_p " ossia il "*valore di pinch-off*", cioè quando si ha $V_{DS} = |V_p|$, il canale si chiude completamente e la corrente I_D rimane ferma al valore assunto prima del manifestarsi del pinch-off.

(Osserviamo che, mentre il raggiungimento del valore di pinch-off da parte della V_{GS} azzerava completamente la corrente I_D , quando è la V_{DS} a uguagliare $|V_p|$, la I_D non si estingue, però non cresce più).

3. **Comportamento "normale"**, cioè con:

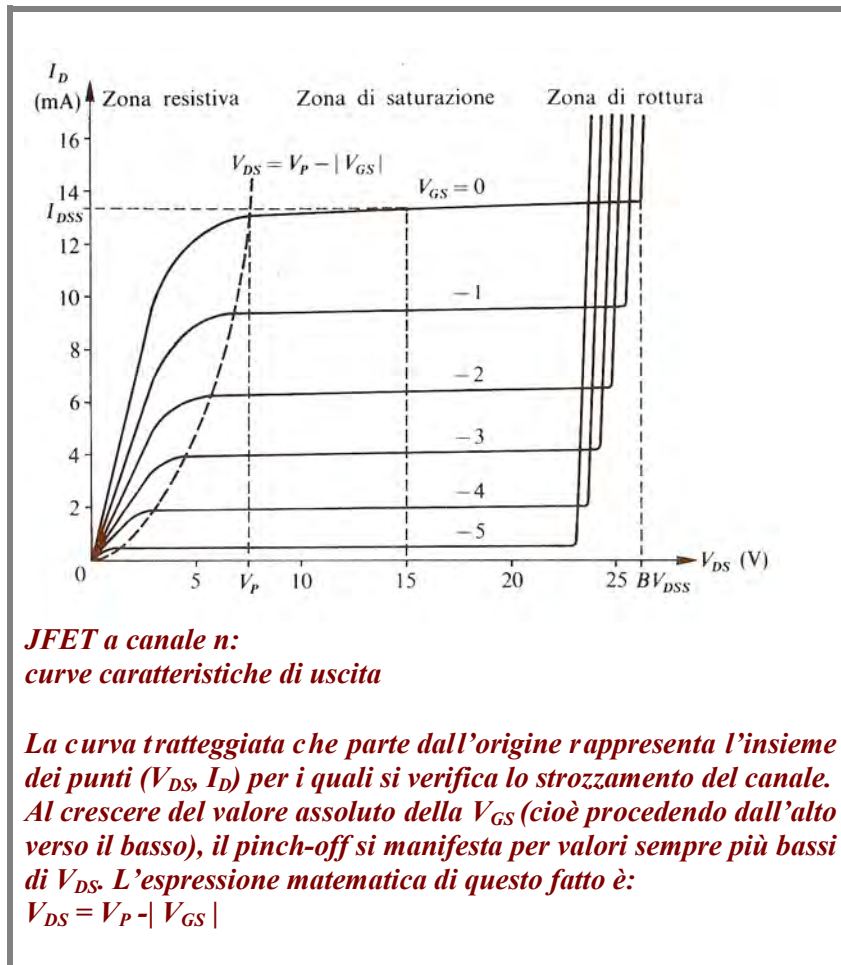
- **Polarizzazione inversa di gate** rispetto al source, cioè con $V_{GS} < 0$
- **polarizzazione di drain** rispetto al source, cioè con $V_{DS} > 0$

E' questa la situazione nella quale normalmente, viene fatto lavorare il JFET.



CURVE CARATTERISTICHE DI USCITA DEL JFET

Quella delle caratteristiche di uscita è una famiglia di curve nella quale **ciascuna curva** rappresenta l'andamento della corrente di drain I_D al crescere di V_{DS} , per un fissato valore della V_{GS} .



ANALISI DI UNA SINGOLA CURVA (PER $V_{GS} = \text{COSTANTE}$)

- Per bassi valori di V_{DS} il canale si comporta in maniera puramente resistiva e la corrente cresce linearmente con V_{DS}
- Quando V_{DS} , crescendo, si avvicina al valore di pinch-off la zona di svuotamento si estende, il canale si restringe e la velocità di crescita di I_D diminuisce, per effetto del restringimento (graficamente questo comportamento è descritto dalla “flessione” della curva)
- Una volta avvenuto lo strozzamento, la corrente conserva un valore pressoché costante, raggiungendo la saturazione.

ANALISI DELLE DIVERSE CURVE (AL VARIARE DI V_{GS})

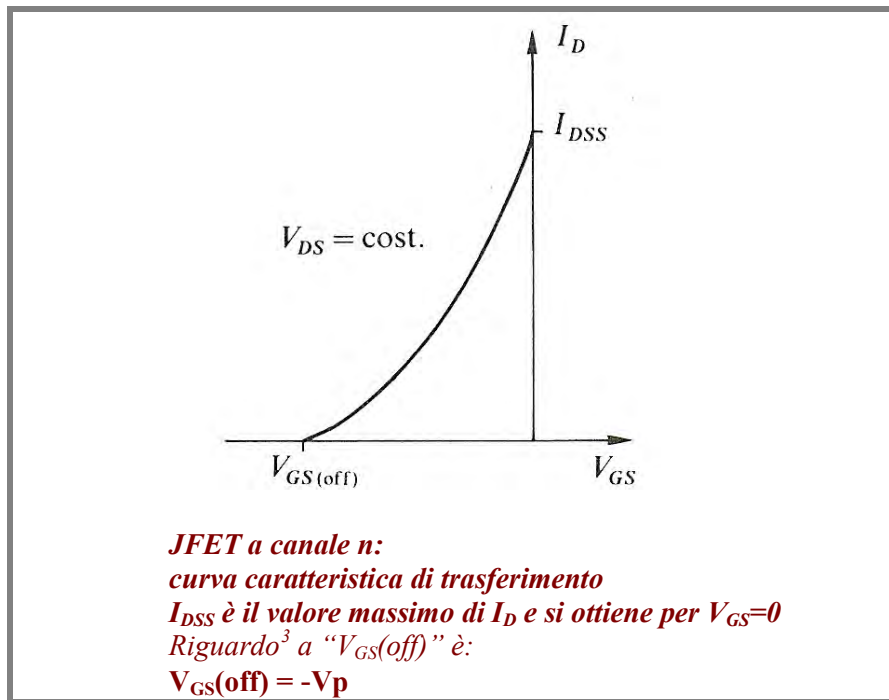
- Al crescere del valore assoluto della V_{GS} , cioè procedendo, di curva in curva, dall’alto verso il basso, notiamo che lo strozzamento si manifesta sempre più presto, cioè per valori di V_{DS} sempre più piccoli.
- Nella curva più in basso la regione di svuotamento comincia a formarsi¹ già con V_{DS} nulla. Di conseguenza il canale presenta una resistenza maggiore e, nella zona a comportamento resistivo, la crescita di I_D è più lenta.

¹ Attenzione: comincia a formarsi la zona di svuotamento, ma si è ancora lontani dallo strozzamento.

CARATTERISTICA DI TRASFERIMENTO O TRANSCARATTERISTICA

Questa curva ha un andamento quadratico e descrive il comportamento della corrente di drain I_D , in funzione di V_{GS} nella regione di saturazione, cioè nella regione nella quale le curve caratteristiche della figura precedente² sono orizzontali. La relazione espressa di questa curva è:

$$I_D = I_D(V_{GS}) \quad \text{per } V_{DS} = \text{costante}$$



Riportiamo l'equazione della transcaratteristica e la tabulazione di alcuni valori.

$$I_D = I_{DSS} \cdot \left(1 - \frac{V_{GS}}{V_p}\right)^2$$

V_{GS}	I_D
0	$\rightarrow I_{DSS}$ (val. max.)
$V_p/2$	$\rightarrow I_{DSS} \cdot (1-0,5)^2 = 0,25 \cdot I_{DSS}$
V_p	$\rightarrow 0$

Nella tabella che segue sono confrontati infine alcuni

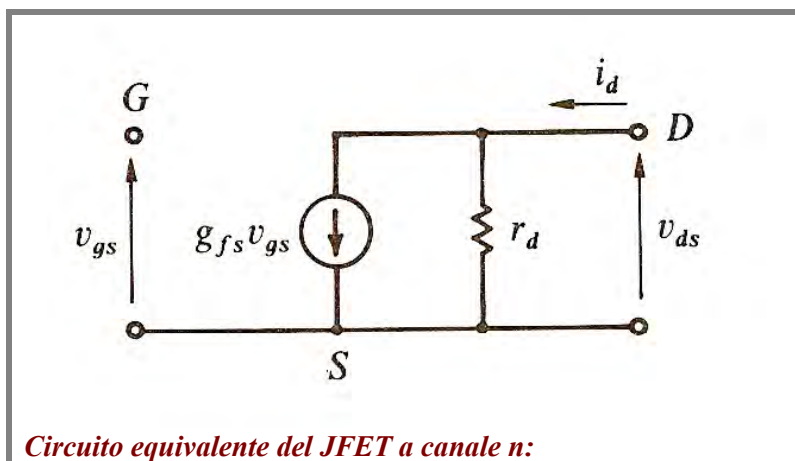
- ² La transcaratteristica può essere ottenuta dal grafico delle curve caratteristiche procedendo, dall'alto verso il basso, lungo una verticale passante per un punto V_{DS} della zona di saturazione.
- ³ Nel foglio-dati il costruttore fornisce generalmente $V_{GS(off)}$ e non V_p .
 $V_{GS(off)}$ è il valore negativo, tanto elevato in valore assoluto, da impedire il passaggio della I_D per ogni valore di V_{DS} , dal momento che, per $V_{GS} = V_{GS(off)}$, lo strozzamento del canale si ha già per $V_{DS}=0$. $V_{GS(off)}$, tipicamente ha valori interni a un intervallo che va da qualche decimo di volt agli 8V circa

parametri elettrici relativi a differenti JFET

Simbolo	2N3967		2N4977		2N5397		2N2498		Unità
	min	max	min	max	min	max	min	max	
I_{GSS}		— 0,1		— 0,5		— 0,1		10	nA
I_{DSS}	2,5	10	50		10	30	— 2	— 6	mA
$V_{GS(off)}$	— 2	— 5	— 4	— 10	— 1	— 6		6	V
$R_{DS(on)}$		400		15				800	Ω

Parametri elettrici di alcuni JFET

CIRCUITO EQUIVALENTE DINAMICO DEL JFET



Circuito equivalente del JFET a canale n:

Il circuito equivalente del JFET, con i suoi parametri, costituisce un **modello astratto di rappresentazione** del transistor, che ne descrive il comportamento attraverso componenti più semplici (lineari). In particolare descrive il funzionamento del JFET in bassa frequenza, per piccoli segnali, e quindi in condizioni di linearità. Il circuito equivalente di un transistor serve a determinare le amplificazioni e le resistenze di ingresso e di uscita e la frequenza di taglio inferiore dell'amplificatore basato su tale transistor.

Per svolgere questo procedimento si possono effettuare i seguenti passi:

- dato, per esempio, lo schema elettrico di un amplificatore a JFET, si disegna il **circuito dinamico dell'amplificatore**, cioè il circuito che ne descrive il funzionamento solo per quanto riguarda le componenti variabili nel tempo del segnale (il circuito dinamico si ottiene da quello originario sostituendo i generatori di tensione continua con dei cortocircuiti, o con le loro resistenze interne, e sostituendo i condensatori con dei circuiti aperti)
- nel circuito dinamico dell'intero amplificatore, **si sostituisce il circuito equivalente del JFET al posto del simbolo circuitale del JFET** stesso.
- basandosi sullo schema così ottenuto, si determinano le espressioni dei valori dei parametri desiderati.

PARAMETRI DEL CIRCUITO EQUIVALENTE

- **AMPLIFICAZIONE DI TENSIONE**

$$\mu = r_D \cdot g_{fs} = \left. \frac{\Delta v_{DS}}{\Delta v_{GS}} \right|_{i_D = COST}$$

- **RESISTENZA DI DRAIN**

$$r_D = r_{DS} = \left. \frac{\Delta v_{DS}}{\Delta i_D} \right|_{v_{GS} = COST}$$

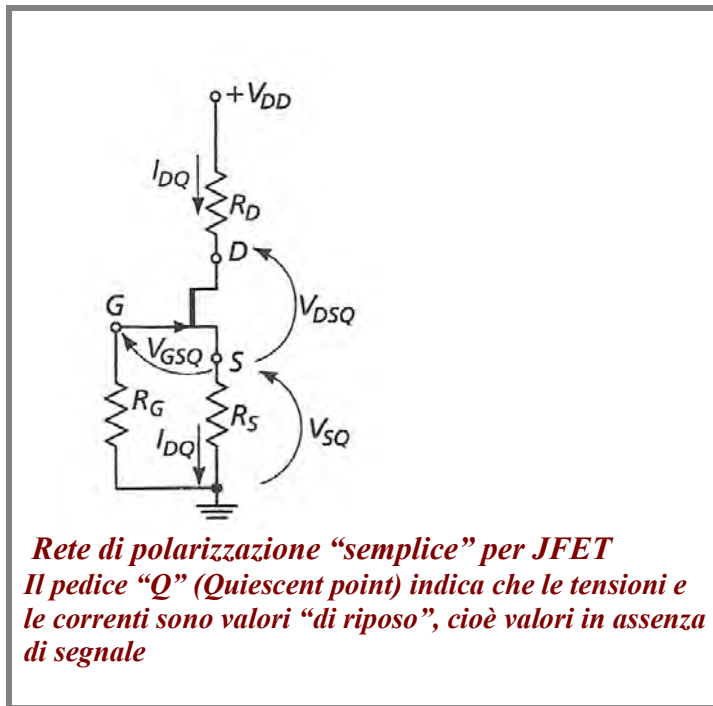
- **TRANSCONDUITTANZA**

$$g_m = g_{fs} = \left. \frac{\Delta i_D}{\Delta v_{GS}} \right|_{v_{DS} = COST}$$

DIMENSIONI E VALORI DEI PARAMETRI

PARAMETRO	UNITA' DI MISURA	CAMPO DI VALORI
μ	adimensionale	(50 ÷ 1000)
r_D	K Ω	(50 ÷ 1000)
g_m	mA/V	(0,1 ÷ 10)

CIRCUITO DI POLARIZZAZIONE “SEMPLICE” O AUTOPOLARIZZAZIONE DEL JFET: DIMENSIONAMENTO



PREMESSA

Il circuito ha lo scopo di fissare il **punto di riposo** nella **zona di saturazione** delle curve caratteristiche di uscita, dove queste sono **orizzontali**.

Il circuito deve **polarizzare inversamente** la giunzione “gate(p)/canale(n)” e lo fa in questo modo:

1. porta il **gate (G)** a potenziale di **massa** attraverso una resistenza R_G
2. porta a **potenziale positivo** il **source (S)** attraverso una resistenza R_S .

Riguardo al punto “1” osserviamo che, essendo la **giunzione gate-canale inversamente polarizzata**, la corrente I_G assume il valore I_{GSS} che è estremamente piccolo (circa $1 \text{ nA} = 10^{-9} \text{ A}$).

Di conseguenza, **il potenziale V_G del gate ha un valore praticamente nullo**, anche se la R_G , come normalmente accade, viene scelta di **elevato valore (qualche $M\Omega$)**.

Se, per esempio, si sceglie $R_G = 1 M\Omega$, il potenziale di gate risulta:

$$V_G = -R_G \cdot I_G = -R_G \cdot I_{GSS} = 10^6 \cdot 10^{-9} = 10^{-3} \text{ V} = 1 \text{ mV}$$

DIMENSIONAMENTO DEL CIRCUITO

DATI E INCOGNITE

Grandezze assegnate (si tratta di valori di riposo):

- V_{DD} = tensione di alimentazione
- I_D = corrente di drain
- V_{DS} = tensione drain-source
- V_{GS} = tensione gate-source

Grandezze da ricavare:

- R_G = resistenza di gate
- R_S = resistenza di source
- R_D = resistenza di drain

PROCEDIMENTO

- In base allo schema circuitale, considerando il transistor come un “nodo” possiamo scrivere:

$$I_S = I_D + I_G$$

Se **scegliamo** una **R_G** di **alcuni $M\Omega$** , possiamo ritenere che si abbia:

$$I_G \cong 0 \rightarrow I_S \cong I_D$$

- Determiniamo la R_S (ricordando che V_G è trascurabile rispetto a V_S):

$$R_S = \frac{V_S}{I_S} \cong \frac{V_S}{I_D} \cong \frac{V_{GS}}{I_D} \quad (\text{dove } V_{GS} \text{ e } I_D \text{ sono assegnate})$$

- Determiniamo la R_D applicando il secondo principio di Kirchhoff alla maglia di uscita:

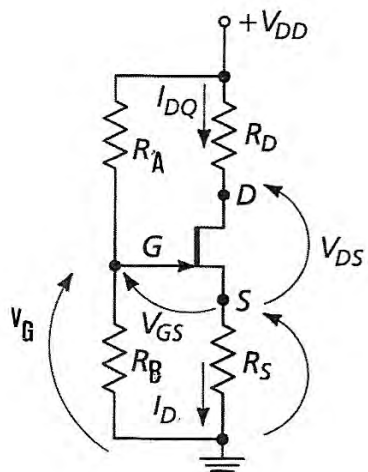
$$V_{DD} = R_D \cdot I_D + V_{DS} + R_S \cdot I_D \rightarrow R_D \cdot I_D = V_{DD} - V_{DS} - R_S \cdot I_D \rightarrow$$

$$\rightarrow R_D = \frac{V_{DD} - V_{DS} - R_S \cdot I_D}{I_D} \rightarrow$$

$$R_D = \frac{V_{DD} - V_{DS}}{I_D} - R_S$$

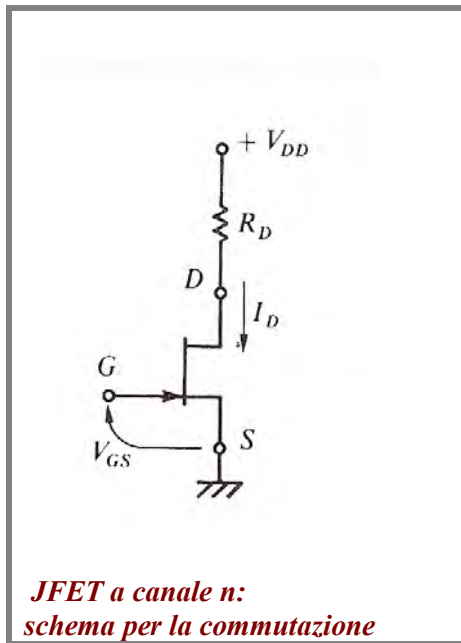
(dove tutti i valori al secondo membro sono assegnati o già calcolati)

RETE DI POLARIZZAZIONE E STABILIZZAZIONE A PARTITORE DEL JFET: ANALISI



rete di polarizzazione a partitore

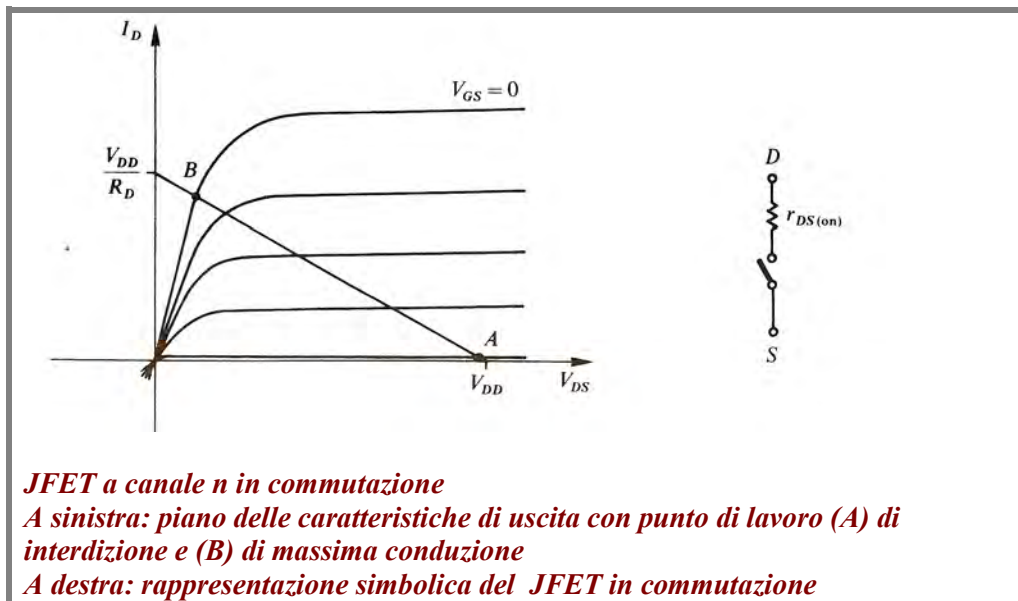
IL JFET IN COMMUTAZIONE



Il JFET può essere utilizzato come **interruttore elettronico**.

In tal caso il componente commuta velocemente fra due condizioni “estreme”:

- l'**interdizione**, nella quale si comporta come un **interruttore aperto**
- la **massima conduzione** (situazione che per il BJT è chiamata “**saturazione**”), nella quale si comporta come un **interruttore chiuso**, cioè come un **cortocircuito**.



INTERDIZIONE (OFF)

La condizione di interdizione è caratterizzata da:

- **corrente di drain** pressoché **nulla** (circa $I_D=0$)
- **tensione fra drain e source** “alta” e, in particolare, pressoché uguale alla tensione di alimentazione (circa $V_{DS}=V_{DD}$)
- **punto di funzionamento** $A=(V_{DD}, 0)$, in basso a destra nella figura.

Per ottenere l’interdizione deve essere: $V_{GS} < V_{GS(OFF)}$

MASSIMA CONDUZIONE (ON)

La condizione di massima conduzione è caratterizzata da:

- **corrente di drain** con il **valore massimo** consentito dal circuito: ($I_D=V_{DD}/R_D$)
- **tensione fra drain e source** “bassa”, pressoché nulla ($V_{DS}=0$)
- **punto di funzionamento** $B=(0, V_{DD}/R_D)$

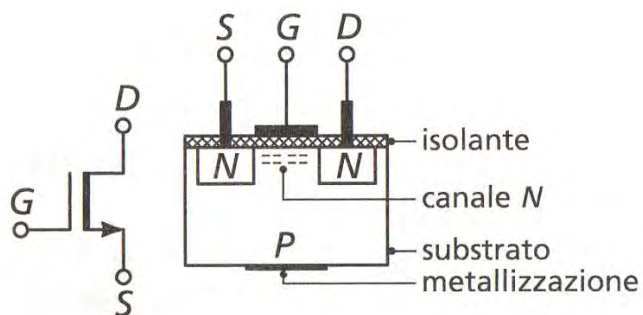
Per ottenere la massima conduzione deve essere: $V_{GS}=0$

MOSFET A SVUOTAMENTO o MOSFET DEPLETION (analogo al JFET)

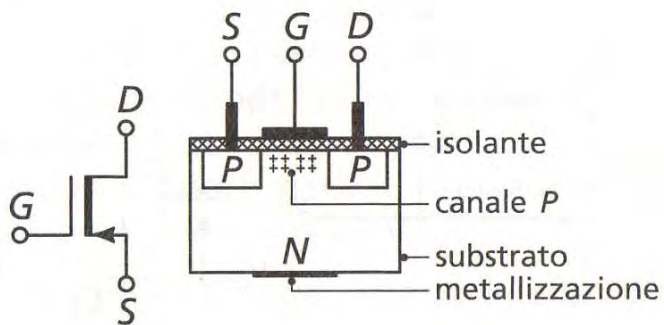
Il MOSFET, può essere, analogamente al JFET, sia a canale n che a canale p. Noi faremo sempre riferimento (salvo avviso contrario) al tipo a canale n, che è il più diffuso.

Come il JFET, il MOSFET a svuotamento ha un canale di conduzione pre-formato, cioè un canale che esiste anche in assenza di polarizzazione del gate. C'è però un differenza:

- in **assenza di polarizzazione** il **canale** del **JFET** ha la **dimensione massima consentita dalla struttura**, in quanto occupa tutta la sezione disponibile del componente; la dimensione del canale può essere quindi solo ridotta (e non aumentata); perciò al JFET si applicano tensioni V_{GS} soltanto negative;
- Nel MOSFET-depletion il **canale-preformato**, cioè il canale presente anche in assenza di polarizzazione, **non occupa tutta la sezione disponibile del componente**, occupa infatti **solo una parte** del substrato; il canale può essere quindi sia **ristretto**, rispetto alle dimensioni originarie, mediante una $V_{GS} < 0$ (così come si fa per il JFET) sia **allargato**, mediante una $V_{GS} > 0$, ottenendo un comportamento analogo a quello del MOSFET-enhancement, di cui si parlerà nelle pagine seguenti.



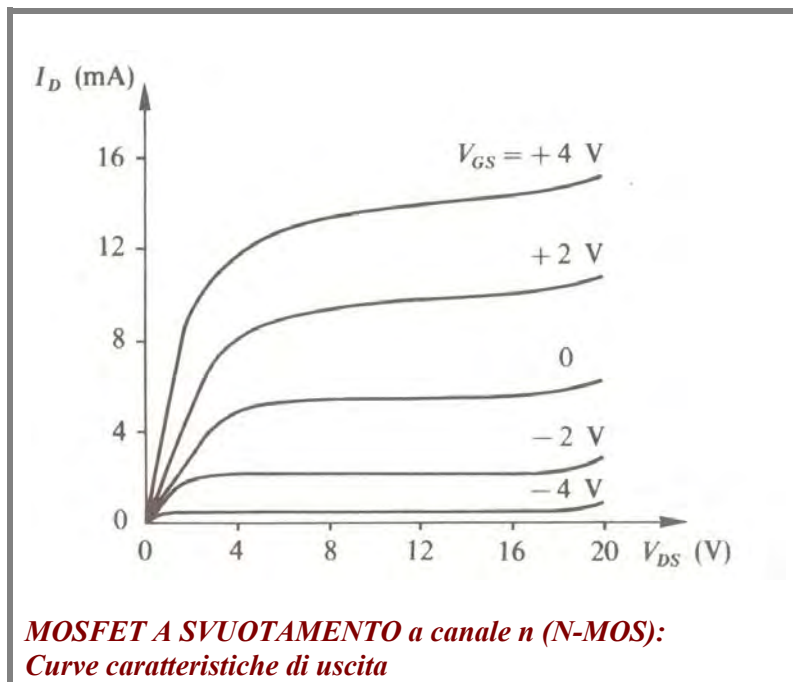
***MOSFET a canale n (N-MOS):
simbolo circuitale e struttura di principio.***



***MOSFET a canale p:
simbolo circuitale e struttura di principio.***

MOSFET DEPLETION A CANALE n
CURVE CARATTERISTICHE DI USCITA O DI DRAIN

Come per il JFET, le caratteristiche di uscita costituiscono una famiglia di curve. Ciascuna curva rappresenta l'andamento della corrente di drain I_D al crescere di V_{DS} , per un fissato valore della V_{GS} .



COMPORTAMENTO AL VARIARE DI V_{GS}

Se, per un fissato valore di V_{DS} sull'asse orizzontale, tracciamo idealmente una retta verticale e **ci muoviamo lungo questa verticale**, incontriamo curve caratterizzate da differenti valori di V_{GS} . Vedremo cosa succede quando la V_{GS} è nulla (curva posta nella zona centrale del grafico), quando la V_{GS} è positiva (zona superiore del grafico) e quando la V_{GS} è negativa (zona inferiore del grafico)

COMPORTAMENTO PER $V_{GS}=0$

Siccome il MOS-DEPLETION ha, come il JFET, il canale pre-formato, in esso si ha, come nel JFET, in presenza di una $V_{DS}>0$, una corrente drain che va da D ad S, anche quando è $V_{GS}=0$ (come si vede dalla curva $I_D = I_D(V_{DS})$ per $V_{GS}=0$)

COMPORTAMENTO PER $V_{GS}>0$

Quando la tensione applicata al gate è positiva (rispetto al source) si ha un **allargamento** del canale rispetto alle dimensioni che si hanno per $V_{GS}=0$

(e il componente si comporta come un MOSFET-ENHANCEMENT che analizzeremo in seguito).

Pertanto, a parità di V_{DS} , si ha una corrente I_D maggiore di quella che scorre quando è $V_{GS}=0$.

Il motivo di quanto accade è che, quando la tensione applicata al gate è positiva, le lacune del substrato p si allontanano dal fianco del canale n, per cui si forma, di fianco al canale, un fascia di cariche negative che determinano un allargamento del canale stesso. Vanno in direzione dell'allargamento del canale anche gli elettroni minoritari del substrato p, i quali si avvicinano al canale, attratti dal potenziale positivo applicato.

COMPORTAMENTO PER $V_{GS}<0$

Quando la tensione applicata al gate è negativa (rispetto al source) si ha un restringimento del canale rispetto alle dimensioni che si hanno per $V_{GS}=0$.

Pertanto, a parità di V_{DS} , si ha una corrente I_D minore di quella che scorre quando è $V_{GS}=0$ e, a maggior ragione, di quella scorre per $V_{GS}>0$.

Il motivo di quanto accade è che il potenziale negativo applicato al gate “respinge” parte degli elettroni che formano il canale, che così tende a restringersi.

Inoltre le lacune del substrato tendono ad avvicinarsi al fianco del canale n, ricombinandosi agli elettroni là presenti.

ANALISI DI UNA SINGOLA CURVA (PER $V_{GS}=\text{COSTANTE}$)

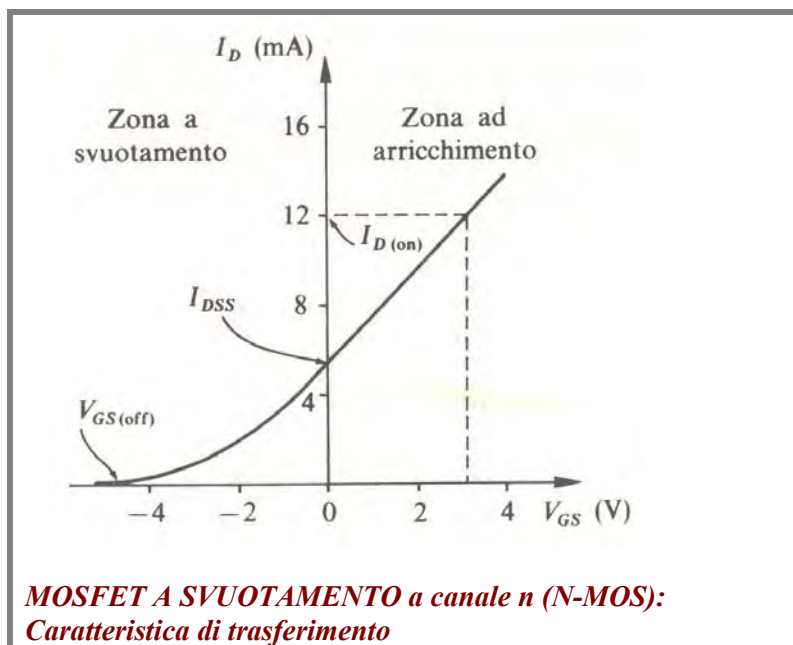
- Per bassi valori di V_{DS} il canale si comporta in maniera puramente resistiva e la corrente cresce linearmente con V_{DS} ;
- Quando V_{DS} , crescendo, si avvicina al valore di pinch-off, la zona di svuotamento si estende, il canale si restringe e la velocità di crescita di I_D diminuisce, per effetto del restringimento; (graficamente questo comportamento è descritto dalla “flessione” della curva);
- Una volta avvenuto lo strozzamento la corrente conserva un valore pressoché costante, raggiungendo la saturazione.

MOSFET DEPLETION A CANALE n

TRANSCARATTERISTICA (O CARATTERISTICA DI TRASFERIMENTO)

Questa curva ha un andamento analogo a quello della transcaratteristica del JFET. Però si sviluppa sia per valori di V_{GS} negativi, sia per valori di V_{GS} positivi. Per il JFET lo sviluppo della curva è limitata al semiasse negativo della V_{GS} , coerentemente col fatto che, per il FET a giunzione, deve essere sempre: $V_{GS} < 0$. La transcaratteristica descrive il comportamento della corrente di drain I_D in funzione di V_{GS} nella regione di saturazione, cioè nella regione nella quale le curve caratteristiche della figura precedente⁴ sono orizzontali. La relazione espressa da questa curva è:

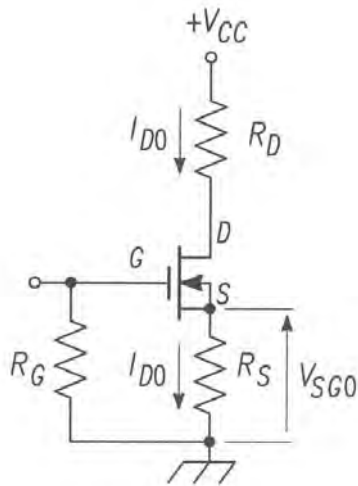
$I_D = I_D(V_{GS})$ per $V_{DS} = \text{costante}$



⁴ La transcaratteristica può essere ottenuta dalle curve caratteristiche procedendo Dal basso verso l'alto lungo una verticale passante per un punto V_{DS} della zona di saturazione.

POLARIZZAZIONE DEL MOSFET A SVUOTAMENTO o MOSFET DEPLETION (analogo al JFET)

Se il punto di riposo deve essere fissato in posizione diversa da $V_{GS}=0$, si può utilizzare un circuito di polarizzazione **analogo a quello del JFET**.



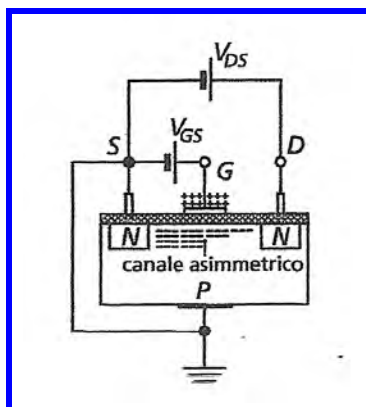
***Rete di polarizzazione per MOSFET a svuotamento
a canale n (N-MOS):***

Se invece è possibile fissare il punto di riposo proprio in corrispondenza di $V_{GS}=0$, non sono necessari i componenti per la polarizzazione della maglia di ingresso.

MOSFET AD ARRICCHIMENTO O ENHANCEMENT

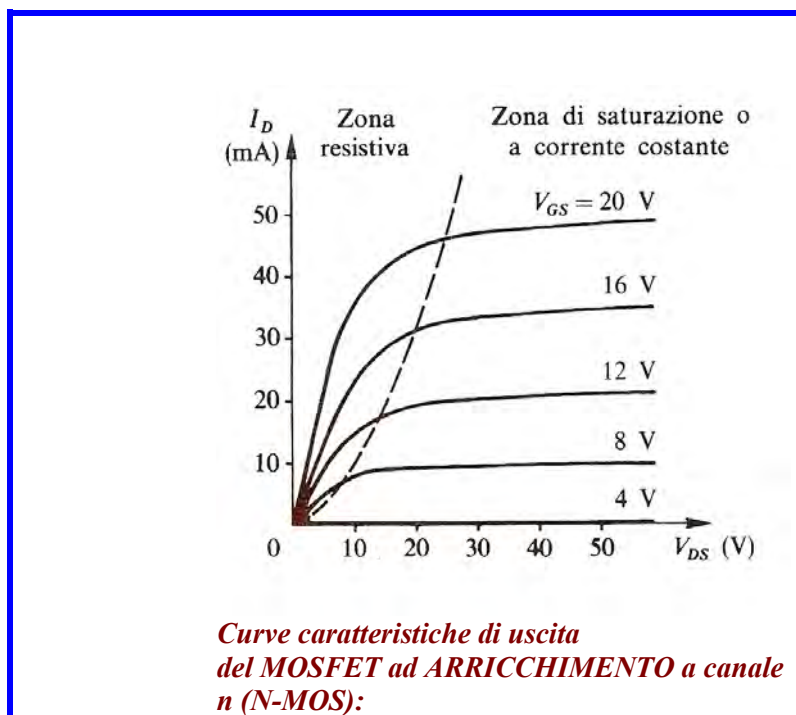
Questo FET, contrariamente al JFET e al MOS-DEPLETION è caratterizzato da un **canale conduttivo pre-formato**.

Il canale si forma **solo in seguito all'applicazione di un campo elettrico**, generato attraverso opportune tensioni di polarizzazione.



CARATTERISTICHE DI USCITA DEL MOSFET AD ARRICCHIMENTO (ENHANCEMENT)

Come per gli altri FET esaminati, le caratteristiche di uscita sono una famiglia di curve, ciascuna delle quali rappresenta l'andamento della corrente di drain I_D al crescere di V_{DS} , per un fissato valore di V_{GS} .



Una singola curva caratteristica di uscita descrive l'andamento della corrente I_D al crescere della tensione V_{DS} , per un fissato valore di V_{GS} .

Notiamo che la V_{GS} deve essere **sempre positiva** (mentre per il JFET deve essere sempre negativa e per il MOS-DEPLETION può essere sia negativa che positiva).

Se, a partire da $V_{DS} = 0$, facciamo crescere V_{DS} , **inizialmente il MOS si comporta in maniera resistiva** e il canale tende a restringersi dalla parte del drain.

Nella **zona resistiva** il **canale** si comporta come una resistenza di valore r_{ON} . Il valore di r_{ON} è funzione della profondità del canale e quindi della V_{GS} . Al crescere di V_{GS} , il canale diventa più profondo e la I_D aumenta.

Quando la V_{DS} raggiunge il valore $V_{DS} = V_{GS} - V_{GSth}$ (dove V_{GSth} è la tensione di soglia⁵) si ha lo strozzamento, o pinch-off, del canale e la corrente I_D non cresce più. Si entra quindi nella zona di saturazione. Come appare evidente anche nel caso del JFET e del MOS-DEPLETION, per i transistor a effetto di campo il termine “saturazione” è usato con un significato completamente diverso da quello che assume nel caso del BJT. La parola “saturazione” è usata, per i FET, solo con riferimento al fenomeno dello strozzamento e del conseguente blocco della crescita di I_D .

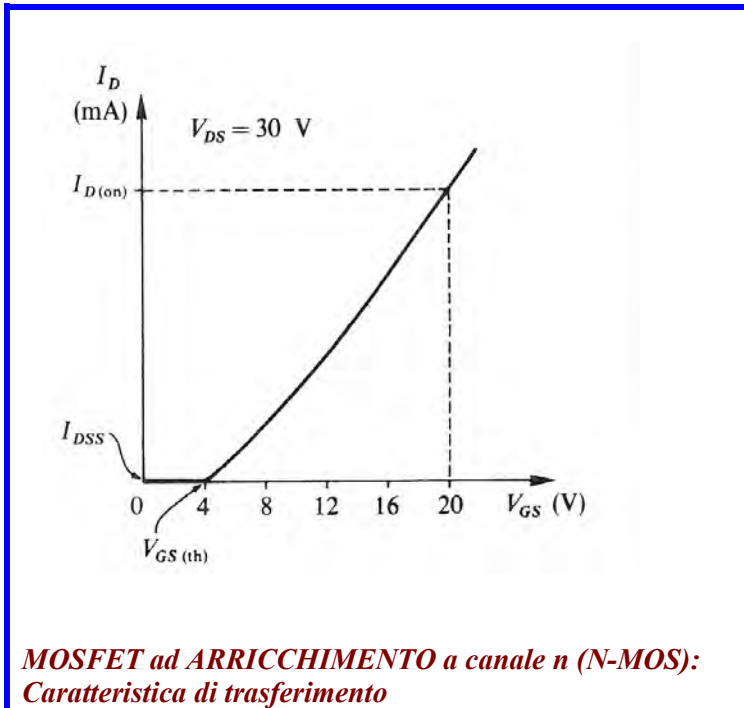
Per i FET la saturazione è la regione del piano delle caratteristiche di uscita nella quale il componente si comporta da amplificatore lineare, e corrisponde quindi alla “zona attiva” del BJT.

Questo significa che, nella zona di saturazione, la V_{GS} fa da tensione di controllo per la corrente di uscita I_D , così come nel BJT la corrente di base I_B controlla la corrente di uscita I_C .

⁵ La tensione di soglia è il valore di V_{GS} al di sotto del quale la I_D è trascurabile. Dipende dallo spessore dell'ossido e ha un valore compreso in un intervallo che va da meno di 1V (per i MOS più recenti) fino a qualche Volt

TRANSCARATTERISTICA DEL MOSFET AD ARRICCHIMENTO (ENHANCEMENT)

La transcaratteristica descrive graficamente l'andamento della corrente I_D in funzione della tensione di gate V_{GS} .



Osserviamo prima di tutto che:

- la V_{GS} deve essere **sempre positiva**, altrimenti il canale non si forma, mentre nel JFET è sempre negativa e nel MOS-DEPLETION può assumere entrambe le polarità;
- conseguentemente la curva interessa solo i valori positivi di V_{GS} , mentre per il JFET interessa solo i valori negativi e per il MOS-DEPLETION sia i valori negativi che quelli positivi

Dalla curva si vede che, per piccoli valori di V_{GS} , valori inferiori alla soglia V_{GSth} , la corrente di drain I_D è pressoché nulla, mentre per valori di V_{GS} superiori alla soglia V_{GSth} , cresce in maniera parabolica, cioè quadratica.

Dopo il superamento della soglia da parte di V_{GS} , accade che la corrente di uscita I_D risulta controllata dalla tensione V_{GS}

L'andamento della curva transcaratteristica è descritto dall'equazione:

$$I_D = K \cdot (V_{GS} - V_{GS,TH})^2$$

In genere il costruttore fornisce la coppia di valori

$(V_{GS(on)}, I_{D(on)})$, cioè le coordinate del punto di funzionamento di piena conduzione.

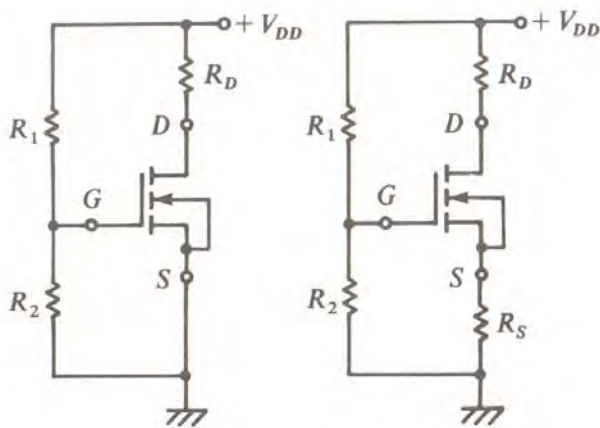
Il valore della costante K che compare nella precedente equazione non è fornito dal costruttore, ma può essere ricavato sostituendo, nell'equazione stessa, il valore di V_{GSth} e una coppia di valori (V_{GS}, I_D) appartenenti alla stessa curva caratteristica (possono essere usati anche gli stessi $V_{GS(on)}, I_{D(on)}$).

POLARIZZAZIONE DEL MOSFET AD ARRICCHIMENTO (ENHANCEMENT)

Per polarizzare un MOSFET ad arricchimento, o Mosfet enhancement, a canale n bisogna (contrariamente a quanto si fa per il MOSFET a svuotamento) portare il gate a un potenziale positivo rispetto al source: $V_{GS} > 0$.

Non è necessaria (se non per la stabilizzazione) una resistenza di source dal momento che l'altissima resistenza di ingresso del componente (un milione di $M\Omega$), impedisce lo scorrimento di corrente nella maglia di ingresso.

RETI DI POLARIZZAZIONE PER MOSFET AD ARRICCHIMENTO (ENHANCEMENT)

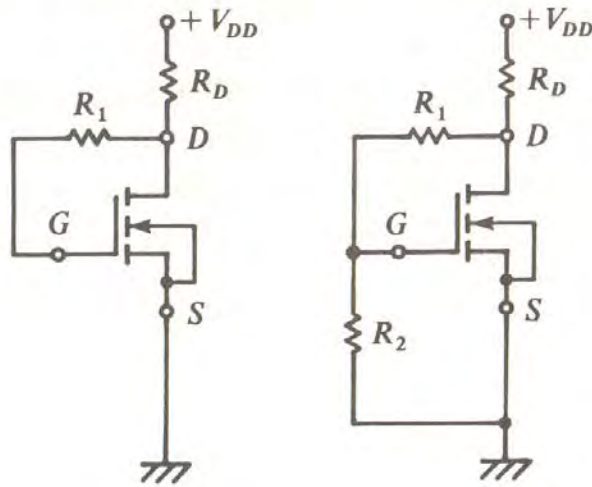


POLARIZZAZIONE

del MOSFET ad ARRICCHIMENTO (ENHANCEMENT)

A sinistra: semplice circuito di polarizzazione senza stabilizzazione. Il circuito porta il gate a un potenziale positivo rispetto al source

A destra: circuito di polarizzazione e stabilizzazione, effettuata mediante la resistenza di source. Anche questo circuito porta il gate a un potenziale positivo rispetto al source

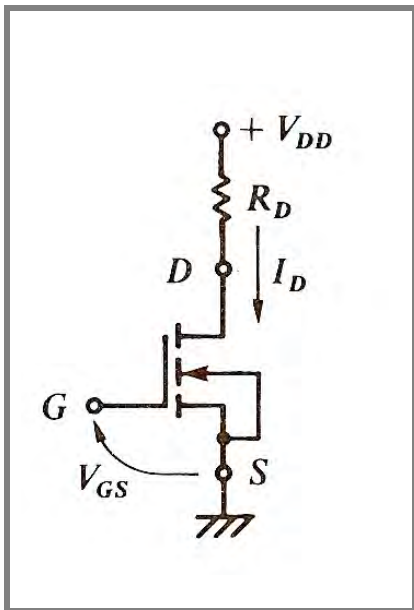


**POLARIZZAZIONE
del MOSFET ad ARRICCHIMENTO (ENHANCEMENT)**

A sinistra: circuito di polarizzazione e stabilizzazione termica, effettuata mediante la resistenza R_1 fra D e G. La presenza di R_1 agisce in questo modo: se I_D aumenta, cresce la caduta di tensione su R_D . Di conseguenza V_{DS} e V_{GS} devono diminuire.

A destra: circuito di polarizzazione e stabilizzazione termica, effettuata ancora attraverso la resistenza R_1 , che però appartiene al partitore R_1+R_2 , che consente di fissare la V_{GS}

I MOSFET IN COMMUTAZIONE



Il MOS, sia a svuotamento che ad arricchimento, può essere utilizzato come **interruttore elettronico**.

In tal caso il componente commuta velocemente fra due condizioni “estreme”:

- l'**interdizione**, nella quale si comporta come un **interruttore aperto**
- la **massima conduzione** (situazione che per il BJT è chiamata “*saturazione*”), nella quale si comporta come un **interruttore chiuso**, cioè come un **cortocircuito**.

INTERDIZIONE (OFF) DEGLI N-MOS

Per ottenere l'**interdizione** del **MOS-DEPLETION** (analogamente al JFET) deve essere:
 $V_{GS} < V_{GS(OFF)}$

Per ottenere l'**interdizione** del **MOS-ENHANCEMENT** deve essere: $V_{GS} < V_{GS(th)}$

(per i PMOS basta invertire i segni di disuguaglianza)

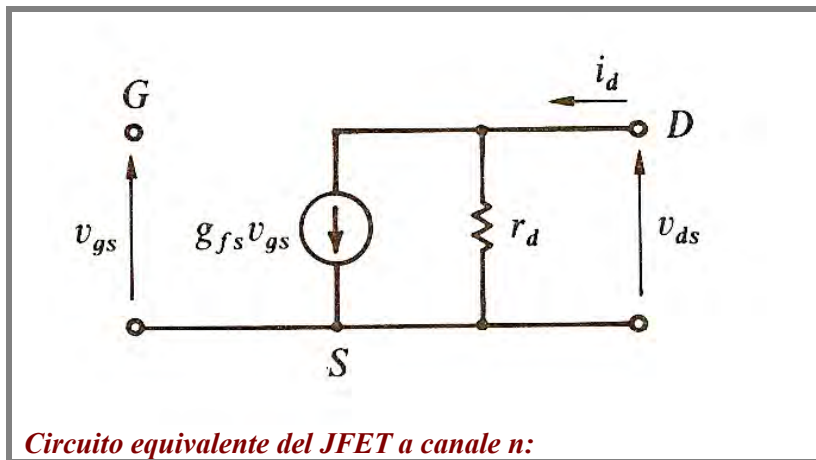
MASSIMA CONDUZIONE (ON) DEGLI N-MOS

Per ottenere la **massima conduzione** deve essere fornita una V_{GS} **sufficientemente alta**.

Nello stato di conduzione il MOS viene fatto lavorare in zona resistiva, perciò il costruttore fornisce il parametro:

$r_{DS(ON)}$ = **resistenza di canale** nella regione resistiva di funzionamento, per un fissato valore di V_{GS} .

CIRCUITO E QUIVALENTE DI NAMICO DE I M OSFET (sia DEPLETION, sia ENHANCEMENT)



Per tutti i tipi di MOSFET si può utilizzare il circuito equivalente del JFET, con le seguenti precisazioni:

- la transconduttanza g_m , o g_{fs} , può assumere anche valori maggiori di quelli del JFET (si veda la tabella)
- la resistenza di drain r_d ha valori molto più bassi di quelli del JFET (si veda la tabella)

PARAMETRI DEL CIRCUITO EQUIVALENTE

- AMPLIFICAZIONE DI TENSIONE

$$\mu = r_D \cdot g_{fs} = \left. \frac{\Delta v_{DS}}{\Delta v_{GS}} \right|_{i_D = \text{cost}}$$

- RESISTENZA DI DRAIN

$$r_D = r_{DS} = \left. \frac{\Delta v_{DS}}{\Delta i_D} \right|_{v_{GS} = \text{cost}}$$

- TRANSCONDUTTANZA

$$g_m = g_{fs} = \left. \frac{\Delta i_D}{\Delta v_{GS}} \right|_{v_{DS} = \text{cost}}$$

DIMENSIONI E VALORI DEI PARAMETRI

PARAMETRO	UNITA' DI MISURA	CAMPO DI VALORI
μ	adimensionale	(50 ÷ 1000)
r_d	K Ω	(1 ÷ 50)
g_m	mA/V	(0,1 ÷ 20)

CAPITOLO 22

AMPLIFICATORI OPERAZIONALI E LORO APPLICAZIONI

COMPLEMENTI

**CONVERTITORI
TENSIONE $\rightarrow$ CORRENTE**

**CONVERTITORE
CORRENTE $\rightarrow$ TENSIONE**

AMPLIFICATORE LOGARITMICO

AMPLIFICATORE LOGARITMICO “COMPENSATO”

AMPLIFICATORE ESPONENZIALE O ANTILOGARITMICO

**CIRCUITI CON OPERAZIONALE E DIODI PER LA MANIPOLAZIONE
DELLE FORME D'ONDA (RADDRIZZATORI)**

**USO DELL'ALIMENTAZIONE SINGOLA PER UN OPERAZIONALE CON
ALIMENTAZIONE DUALE**

L'OPERAZIONALE (IN CATENA APERTA) COME COMPARATORE

CLASSIFICAZIONE DEI COMPARATORI

TRIGGER DI SCHMITT (COMPARATORE CON ISTERESI)

CONVERTITORI TENSIONE → CORRENTE

A COSA SERVONO I CONVERTITORI V/I

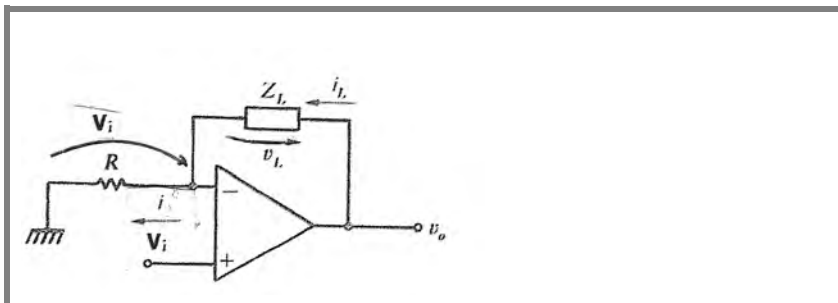
Un convertitore tensione → corrente è un circuito che fornisce in uscita, cioè sul carico, una corrente dipendente soltanto dalla tensione di ingresso, in maniera lineare (e non dal carico).

I convertitori tensione → corrente sono anche chiamati “**amplificatori a transconduttanza**” oppure “**generatori di corrente controllati in tensione**”

(VCIS Voltage-Controlled I-source)

Per ottenere un convertitore tensione-corrente basta inserire l'utilizzatore (il carico) nella rete di retroazione negativa di un amplificatore non invertente o di un amplificatore invertente.

Per entrambi i circuiti vale la stessa relazione I/O.



Convertitore $V \rightarrow I$, basato sull'amplificatore non invertente

Se, come corrente di uscita, consideriamo la i_L “entrante”, come nella figura, l'espressione della corrente di uscita risulta:

$$i_L = + \frac{V_i}{R}$$

ANALISI DEL CIRCUITO

Per il corto circuito virtuale si ha:

$$V^- = V^+$$

ossia:

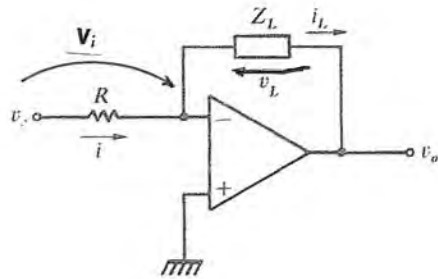
$$V_R = V_i$$

Essendo R_i elevatissima, possiamo ritenere che si abbia: $i_L = i$.

Perciò la legge di Ohm su R , cioè:

$$i = \frac{V_i}{R} \quad \text{si può scrivere come:} \quad i_L = \frac{V_i}{R} \quad (\text{dove } i_L \text{ è la corrente di uscita})$$

Osserviamo che potrebbe essere opportuno che R fosse una resistenza di precisione per poter effettuare in modo ottimale la regolazione della corrente.



Convertitore $V \rightarrow I$, basato sull'amplificatore invertente

Se, come corrente di uscita, consideriamo la i_L "uscente", come nella figura (e al contrario di come si è fatto per il circuito non invertente), l'espressione della corrente di uscita ha ancora il segno positivo e:

$$i_L = + \frac{v_i}{R}$$

ANALISI DEL CIRCUITO

Essendo R_i elevatissima, possiamo ritenere che si abbia: $i_L = i$.

Per il principio di massa virtuale il terminale invertente è a potenziale di massa, per cui:

$$v_R = v_i$$

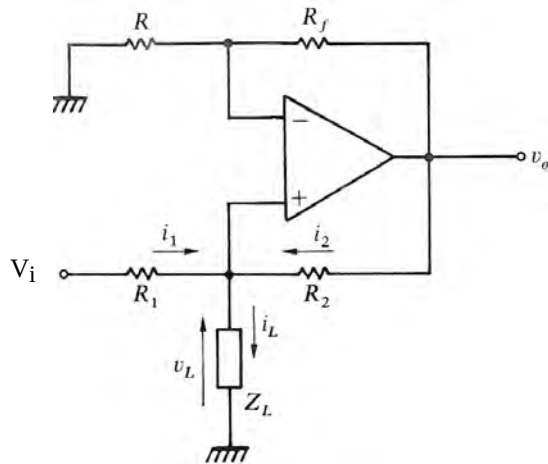
Perciò la legge di Ohm su R, cioè:

$$i = \frac{v_i}{R} \quad \text{si può scrivere come:} \quad i_L = \frac{v_i}{R} \quad (\text{dove } i_L \text{ è la corrente di uscita})$$

Osserviamo che potrebbe essere opportuno che R fosse una resistenza di precisione per poter effettuare in modo ottimale la regolazione della corrente.

CONVERTITORE $V \rightarrow I$ CON UTILIZZATORE (CARICO) RIFERITO A MASSA

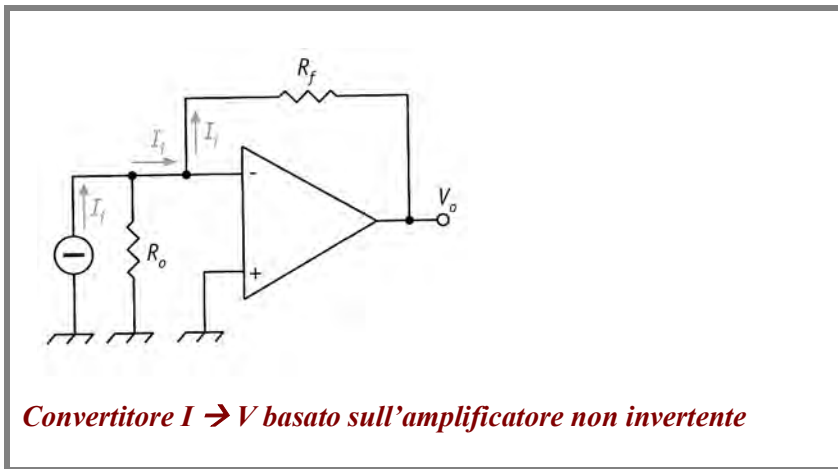
*Convertitore $V \rightarrow I$
con utilizzatore (carico) riferito a massa*



RELAZIONE INGRESSO-USCITA (I/O)

$$i_L = \frac{V_i}{R_1} \quad \text{se si è scelto:} \quad \frac{R_f}{R} = \frac{R_2}{R_1}$$

CONVERTITORE CORRENTE → TENSIONE



A COSA SERVE IL CIRCUITO

Un convertitore corrente → tensione trasforma un generatore reale di corrente in un generatore ideale di tensione. Deve cioè fornire una **tensione di uscita proporzionale** alla **corrente di ingresso** e **indipendente** sia dal **carico** che dalla **resistenza interna** del generatore della corrente di ingresso.

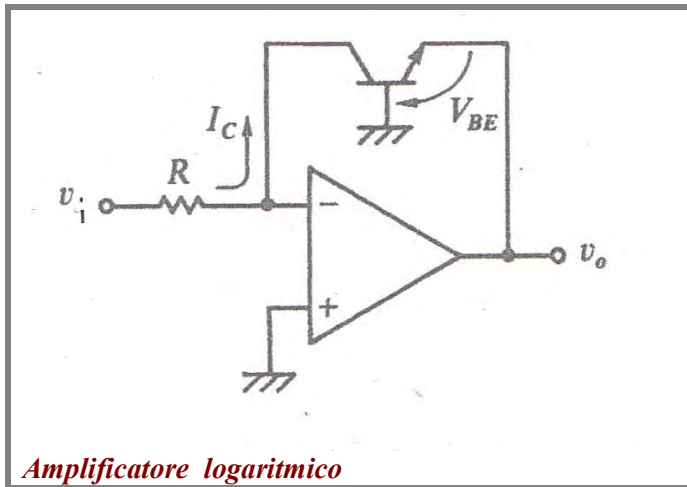
I convertitori corrente → tensione sono anche chiamati “**amplificatori a transresistenza**” oppure “generatori di tensione controllati in corrente” (ICVS, I-Controlled Voltage-source).

RELAZIONE INGRESSO-USCITA (I/O)

tensione di uscita = $-(\text{corrente di ingresso}) \cdot (\text{resistenza di retroazione})$

$$V_o = -R_f \cdot I_i$$

AMPLIFICATORE LOGARITMICO



$v_o = -\frac{KT}{q} \cdot \left[\ln(v_i) - \ln(R \cdot I_{BEsat}) \right]$	Relazione I/O generale¹
--	---

dove:

$k = 1,38 \cdot 10^{-23} \text{ [J/°K]} = \text{costante di Boltzmann}$
$T = \text{temperatura assoluta [°K]}$
$q = 1,6 \cdot 10^{-19} \text{ [C]} = \text{carica dell'elettrone}$
$KT/q = V_T = \text{"equivalente in Volt" della temperatura}$ (a temperatura ambiente è: $V_T = 26\text{mV}$)

¹ dove I_{BEsat} è il valore che nei data-sheet viene indicato come I_R e può valere:

$I_{BEsat} = (20\text{nA} \div 10\text{μA})$ a temperatura ambiente¹

$I_{BEsat} = (3\text{μA} \div 50\text{μA})$ alla temperatura di un centinaio di gradi

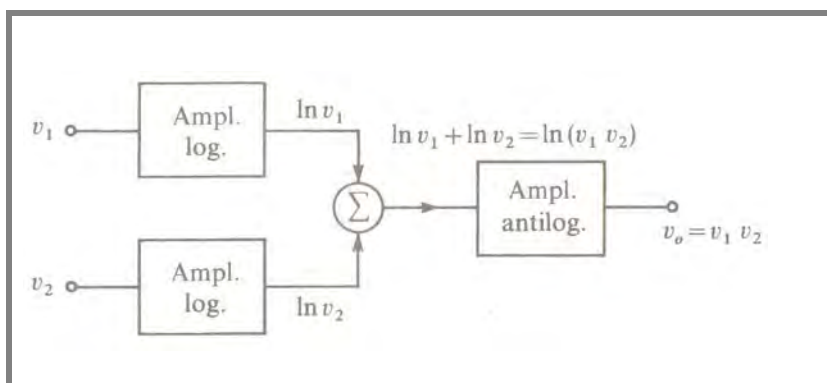
A COSA SERVE IL CIRCUITO

L'amplificatore logaritmo serve essenzialmente a:

- Amplificare con una dinamica di uscita compressa, cioè amplificare con una riduzione dell'escursione massima del segnale di uscita, rispetto a un normale amplificatore lineare. Ciò è utile quando l'escursione risulta troppo ampia.
Se, per esempio, nell'intervallo di tensioni di ingresso ($10\mu\text{V} \div 100\text{mV}$), la tensione di uscita di un normale amplificatore varia nel range ($10^{-3} \div 10$)V, l'uscita di un amplificatore logaritmico è compresa invece nell'intervallo ($\ln 10^{-3} \div \ln 10$)V, ossia nell'intervallo $(-6,9 \div 2,3)$ V, che è molto più ristretto.
- Effettuare l'elaborazione analogica di segnali, cioè ottenere tensioni di uscita che siano il logaritmo (naturale o decimale) delle tensioni di ingresso, oppure effettuare prodotti e rapporti attraverso i logaritmi.
Per esempio, il prodotto di due tensioni si può ottenere attraverso i logaritmi, come mostra l'espressione seguente:

$$V_1 V_2 = e^{\ln(V_1 V_2)} = e^{\ln(V_1) + \ln(V_2)}$$

Questa operazione può essere svolta mediante la struttura circuitale illustrata nello schema a blocchi:



Nello schema:

- I logaritmi $\ln(V_1)$ e $\ln(V_2)$ sono eseguiti da amplificatori logaritmici
- La somma $\ln(V_1) + \ln(V_2)$ è effettuata dal sommatore
- L'esponenziale $e^{\ln(V_1) + \ln(V_2)}$, equivalente al prodotto $V_1 V_2$, è eseguita da un amplificatore esponenziale (detto anche “amplificatore antilogaritmico”)

RELAZIONE INGRESSO USCITA DELL'AMPLIFICATORE LOGARITMICO

E' possibile esprimere questa relazione in diversi modi, come viene schematizzato di seguito:

$v_0 = -\frac{KT}{q} \cdot [\ln(v_i) - \ln(R \cdot I_{BEsat})]$	<i>Relazione I/O generale²</i>
---	---

Esprimendo la relazione mediante logaritmi decimali si ha:

$v_0 = -2,3 \cdot \frac{KT}{q} \cdot [\log_{10}(v_i) - \log_{10}(R \cdot I_{BEsat})]$	<i>Relazione I/O generale espressa mediante logaritmi decimali³</i>
---	--

A temperatura ambiente di 27°C, risultando **KT/q=26·10⁻³V**, la prima relazione si può scrivere come:

$v_0 = -0,026 \cdot [\ln(v_i) - \ln(R \cdot I_{BEsat})] \quad [V]$	<i>Relazione I/O a temperatura ambiente di 27°C</i>
--	---

La stessa relazione, espressa in mV; è:

$v_0 = -26 \cdot [\ln(v_i) - \ln(R \cdot I_{BEsat})] \quad [mV]$	<i>Relazione I/O a temp. Amb. di 27°C e espressa in mV</i>
--	--

La relazione espressa in mV, in base ai logaritmi decimali ,diventa:

$$v_0 = -2,3 \cdot 26 \cdot [\log_{10}(v_i) - \log_{10}(R \cdot I_{BEsat})] \quad [mV]$$

e quindi:

$v_0 = -59,8 \cdot [\log_{10}(v_i) - \log_{10}(R \cdot I_{BEsat})] \quad [mV]$	<i>Relazione I/O a temperatura ambiente di 27°C espressa in mV e mediante logaritmi decimali</i>
--	--

² dove I_{BEsat} può valere:

$I_{BEsat} = (20nA \div 10\mu A)$ a temperatura ambiente²

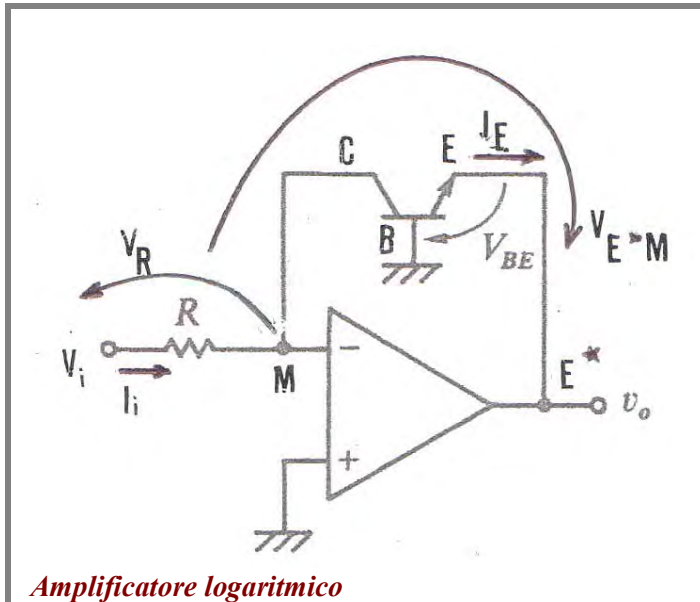
$I_{BEsat} = (3\mu A \div 50\mu A)$ alla temperatura di un centinaio di gradi

³ Utilizzando la relazione:

$$\ln(X) = \frac{\log_{10}(X)}{0,4343} \Leftrightarrow \ln(X) = 2,3 \cdot \log_{10}(X)$$

ANALISI DEL CIRCUITO

L'amplificatore logaritmico è essenzialmente un amplificatore invertente, con un BJT al silicio inserito fra il terminale di uscita e l'ingresso invertente. L'emettitore E del transistor è collegato al terminale di uscita dell'amplificatore operazionale.



Trattandosi di amplificatore invertente, valgono il corto circuito virtuale ($V^+ = V^-$) e il principio di massa virtuale ($V^+ = V^- = 0V$). Il terminale invertente è quindi a potenziale di massa e lo indichiamo con "M".

Il transistor, perciò, oltre ad avere la base collegata a massa, ha anche il collettore C a potenziale di massa, in quanto C è collegato al punto M.

Di conseguenza (osservando lo schema circuitale e avendo presente che, per il principio di massa virtuale, è: $V_{E^*M} = V_0$) possiamo affermare che:

$$V_0 = V_{E^*M} = V_{E^*} - V_M = V_E - V_B = V_{EB}$$

(dove abbiamo potuto sostituire V_B a V_M , dato che i punti B ed M sono entrambi a potenziale di massa)

Abbiamo quindi ottenuto:

$$V_0 = -V_{BE}$$

Per la giunzione base-emettitore di un transistor o di un diodo vale la relazione:

$$I_C = I_{BEsat} \left(e^{\frac{V_{BE}q}{kT}} - 1 \right)$$

da cui si ricava⁴:

$$I_C = I_{BEsat} \left(e^{\frac{V_{BE}q}{kT}} - 1 \right)$$

$$\frac{I_C}{I_{BEsat}} = e^{\frac{V_{BE}q}{kT}} - 1$$

$$V_{BE} = \frac{K \cdot T}{q} \cdot \ln \left(\frac{I_C}{I_{BEsat}} \right)$$

che, sostituita nell'espressione $V_0 = -V_{BE}$, fornisce:

$$V_0 = -\frac{kT}{q} \ln \left(\frac{I_C}{I_{BEsat}} \right) \quad (*)$$

A noi interessa che, nella relazione, figuri la tensione di ingresso V_i dell'amplificatore.

A questo proposito osserviamo dallo schema circuitale che, grazie al principio di massa virtuale, è:

$$V_i = V_R$$

e, per la legge di Ohm:

$$V_i = V_R I_i$$

Essendo poi infinita la resistenza di ingresso dell'operazionale, si ha: $I_i = I_C$, per cui: $V_i = V_R I_C$ da cui:

$$I_C = \frac{V_i}{R}$$

Sostituendo nell'espressione (*) di V_0 , si ha:

$$v_0 = -\frac{KT}{q} \cdot \ln \left(\frac{v_i}{R} \cdot \frac{1}{I_{BEsat}} \right)$$

$$v_0 = -\frac{KT}{q} \cdot \ln \left(\frac{v_i}{R \cdot I_{BEsat}} \right)$$

$$\frac{I_C}{I_{BEsat}} + 1 = e^{V_{BE} \frac{q}{kT}}$$

Essendo I_{BEsat} molto piccola, si può trascurare l'unità, ottenendo:

$$\frac{I_C}{I_{BEsat}} \cong e^{V_{BE} \frac{q}{kT}}$$

Applicando il logaritmo naturale a entrambi i membri si ha:

$$\ln \left(\frac{I_C}{I_{BEsat}} \right) \cong V_{BE} \frac{q}{kT}$$

$$V_{BE} \cong \frac{kT}{q} \ln \left(\frac{I_C}{I_{BEsat}} \right)$$

come volevamo appurare.

Applicando la proprietà dei logaritmi $\ln (X/Y) = \ln (X) - \ln (Y)$, si può scrivere:

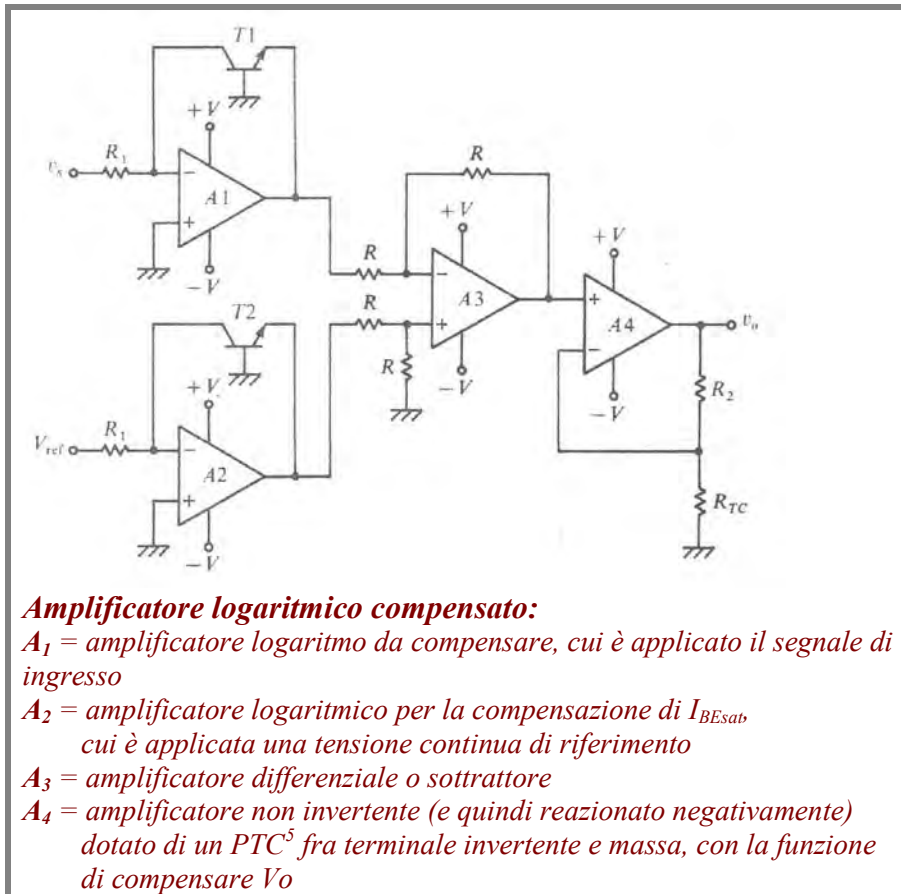
$$v_0 = -\frac{KT}{q} \cdot \left[\ln (v_i) - \ln (R \cdot I_{BEsat}) \right]$$

che è proprio la relazione che volevamo dimostrare.

AMPLIFICATORE LOGARITMICO “COMPENSATO”

Nella configurazione essenziale che abbiamo esaminato, l'amplificatore logaritmico presenta due inconvenienti:

- variazione di I_{BEsat} da transistor a transistor e con la temperatura
- dipendenza della tensione di uscita dalla temperatura (V_O risulta direttamente proporzionale alla temperatura). Una configurazione circuitale che risolve gli inconvenienti esposti è l'amplificatore logaritmico compensato, che viene mostrato nella figura seguente:

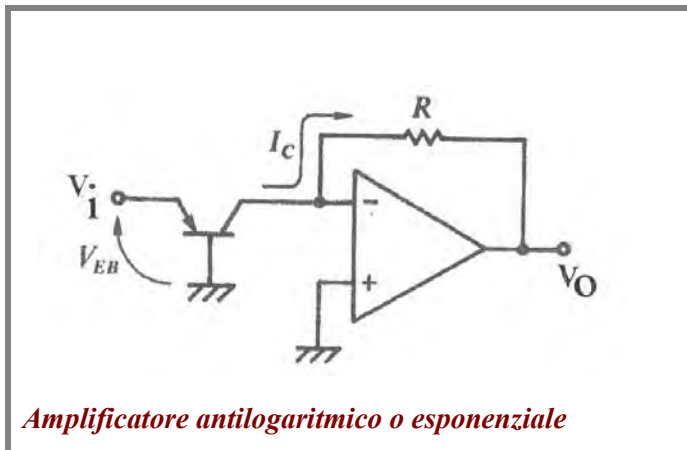


L'espressione della tensione di uscita dell'amplificatore logaritmico compensato è:

$$v_o = \left(1 + \frac{R_2}{R_{PTC}}\right) \cdot \frac{kT}{q} \cdot \ln \left(1 + \frac{v_i}{v_{RIF}}\right)$$

⁵ resistore variabile con la temperatura, a coefficiente di temperatura positivo (la resistenza elettrica del PTC cresce al crescere della temperatura)

AMPLIFICATORE ESPONENZIALE O ANTILOGARITMICO

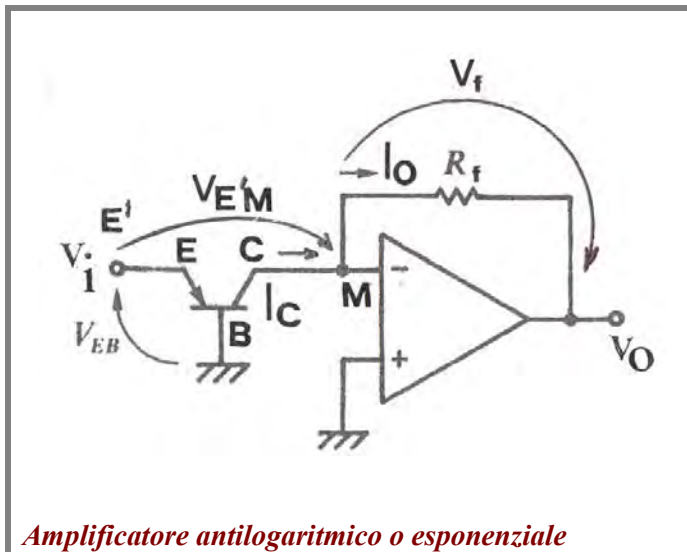


A COSA SERVE IL CIRCUITO

L'utilizzazione fondamentale dell'amplificatore antilogaritmico è nel campo dell'elaborazione analogica dei segnali, come già si è accennato nel precedente paragrafo sull'amplificatore logaritmico.

$v_O = R_f \cdot I_{BEsat} \cdot e^{v_i / kT}$	<i>Relazione I/O</i>
--	----------------------

ANALISI DEL CIRCUITO



L'amplificatore antilogaritmico è essenzialmente un amplificatore invertente, con un BJT di tipo **PNP**, al silicio, inserito fra l'ingresso invertente (M) dell'operazionale e il terminale (E') di applicazione della tensione di ingresso.

Trattandosi di amplificatore invertente, valgono il corto circuito virtuale ($V^+ = V^-$) e il principio di massa virtuale ($V^+ = V^- = 0V$). Il terminale invertente è quindi a potenziale di massa e lo indichiamo con "M".

Osservando lo schema circuitale, possiamo affermare che:

$$V_{E'M} = V_{E'} - V_M = V_E - V_M = -V_{BE} - V_M = -V_{BE} - 0V = -V_{BE}$$

ossia:

$$V_{E'M} = +V_{EB}$$

Per il principio di massa virtuale è $V_i = V_{E'M}$ e quindi:

$$V_i = +V_{EB} \quad (*)$$

Essendo infinita la resistenza di ingresso dell'operazionale, si ha: $I_C = I_o$

Ancora per il principio di massa virtuale si ha: $V_o = V_f$

e, per la legge di Ohm: $V_o = R_f I_o$ e quindi:

$$V_o = R_f I_C \quad (**)$$

Riguardo al BJT, che ha il collettore a massa, vale l'equazione:

$$I_C = I_{BEsat} \cdot \left(e^{V_{EB} \cdot q / kT} - 1 \right)$$

(equazione uguale a quella relativa al BJT di tipo NPN, ma con V_{EB} al posto di V_{BE})

Trascurando l'unità rispetto all'esponenziale, si ottiene:

$$I_C = I_{BEsat} \cdot e^{V_{EB} \cdot q / kT}$$

che, sostituita nella (**), fornisce:

$$v_o = R_f \cdot I_{BEsat} \cdot e^{v_i/q/kT}$$

che è proprio ciò che volevamo dimostrare.

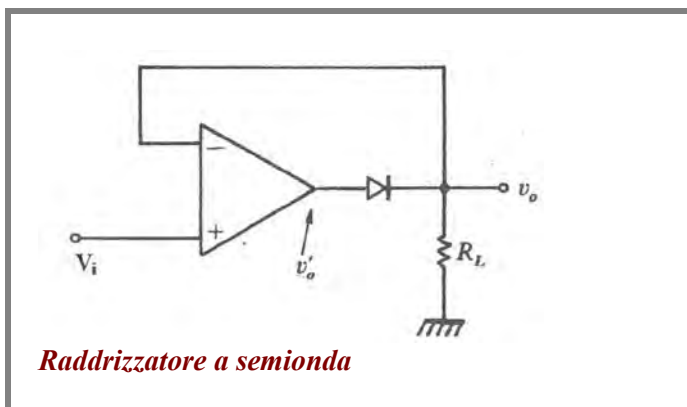
CIRCUITI CON OPERAZIONALE E DIODI PER LA MANIPOLAZIONE DELLE FORME D'ONDA

RADDRIZZATORI

I raddrizzatori ad operazionale presentano, rispetto a quelli realizzati con soli diodi, il vantaggio di **portare i diodi in conduzione a una tensione V_T/A_{OL} molto più bassa della tensione di soglia V_T .**

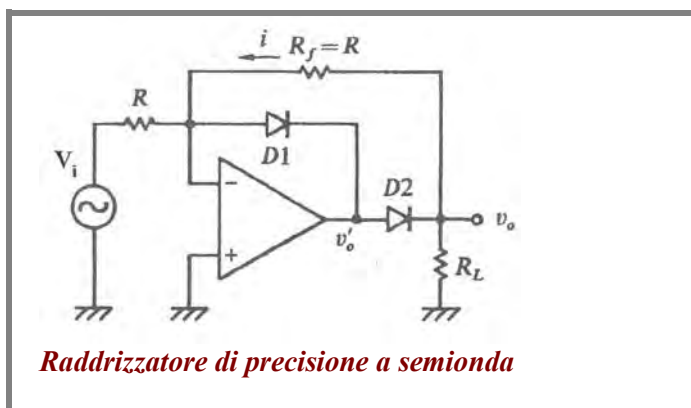
Offrono inoltre la possibilità di amplificare il segnale da raddrizzare.

RADDRIZZATORE A SEMIONDA



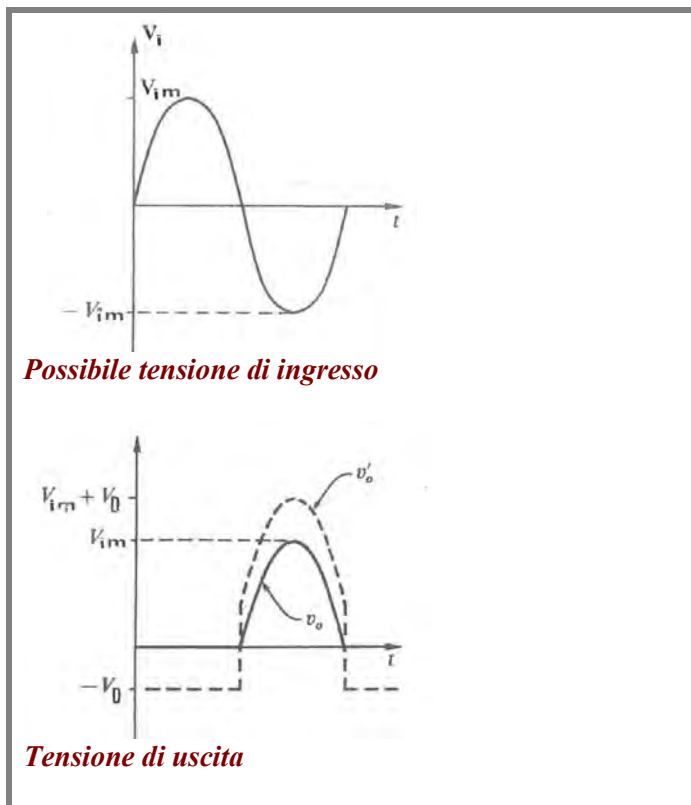
Questo raddrizzatore trasferisce in uscita le sole semionde positive della tensione di ingresso, facendo entrare in conduzione il diodo alla tensione V_T/A_{OL} molto più bassa della tensione di soglia. Però distorce i segnali in misura tanto maggiore quanto più alta è la loro frequenza.

RADDRIZZATORE DI PRECISIONE A SEMIONDA



Raddrizzatore di precisione a semionda

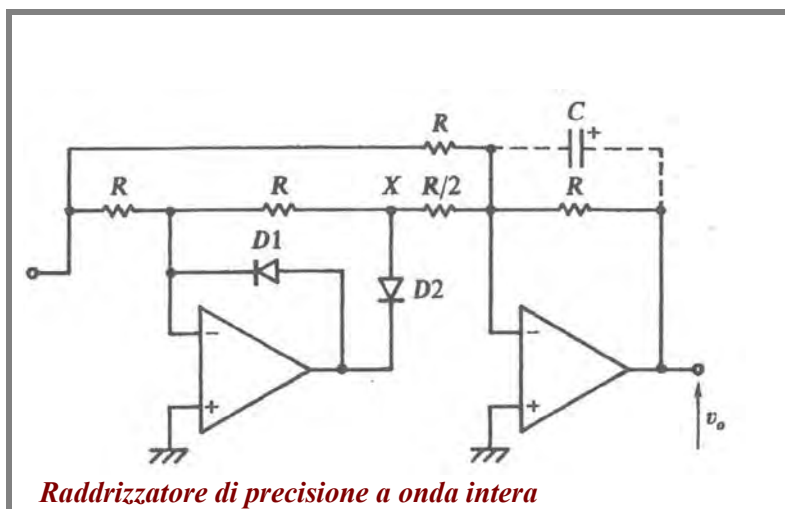
Questo circuito non determina distorsione nel segnale di uscita, che è formato dalle semionde negative di ingresso, raddrizzate ed eventualmente amplificate.



Possibile tensione di ingresso

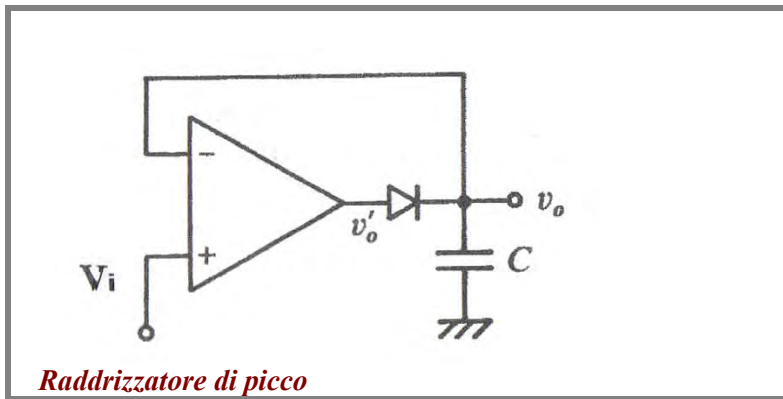
Tensione di uscita

RADDRIZZATORE DI PRECISIONE A ONDA INTERA



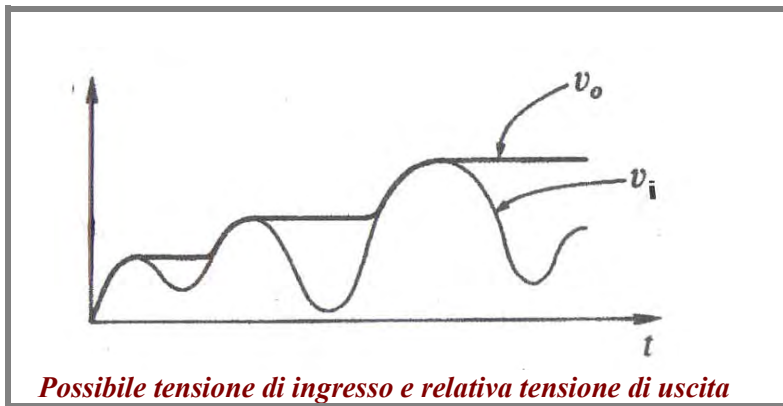
Il circuito è formato dalla **cascata** di un **raddrizzatore di precisione a semionda** e di un **sommatore invertente**.

RADDRIZZATORE DI PICCO

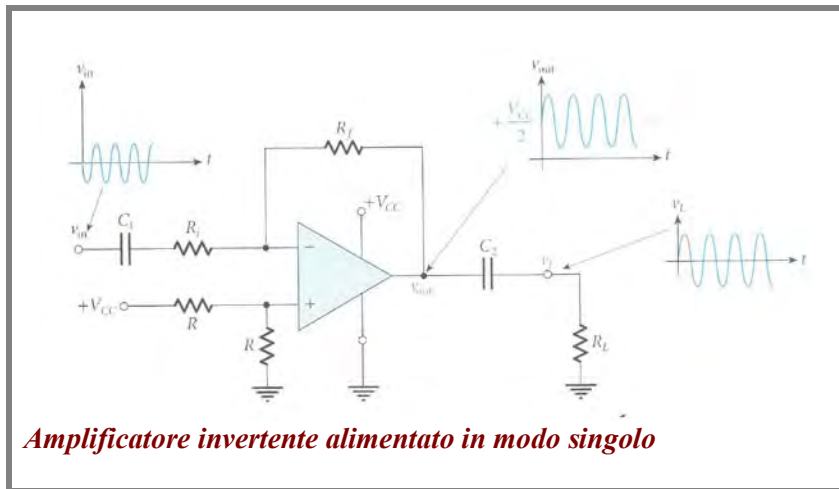


Il raddrizzatore di picco fornisce una tensione di uscita che segue quella di ingresso, finché la V_i non raggiunge il suo valore massimo.

Da questo istante in poi la V_o si mantiene costante sul valore V_{iMAX} .



USO DELL'ALIMENTAZIONE SINGOLA PER UN OPERAZIONALE CON ALIMENTAZIONE DUALE

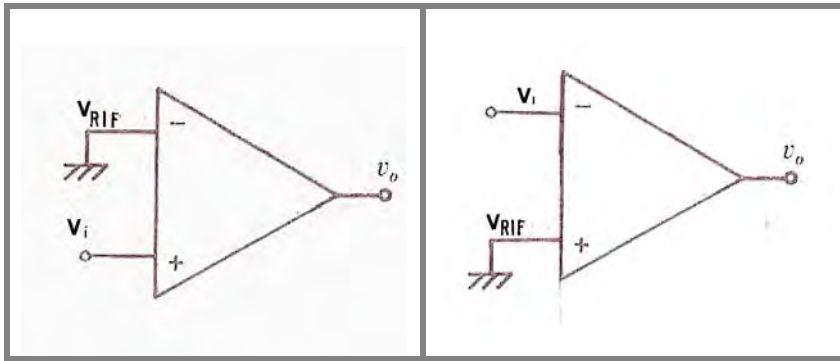


Ogni operazionale con alimentazione duale può essere alimentato in modo singolo, collegando a massa il terminale dell'alimentazione negativa.

Se non si modifica adeguatamente la rete esterna, l'alimentazione singola determina il dimezzamento dell'escursione della tensione di uscita, che si riduce al range $(0V \div V_{O(SAT)}^+)$. Se quindi non si prendono gli opportuni accorgimenti, le eventuali semionde negative del segnale di uscita vengono tagliate. Una possibile configurazione amplificatrice, in grado di amplificare, con alimentazione singola, anche segnali bipolari è quella della figura: si tratta di un **amplificatore invertente** nel quale, però il **potenziale del terminale non invertente non è a massa, ma viene portato al livello $V_{CC}/2$** (metà della tensione di alimentazione) mediante il partitore $R+R$ posto fra l'alimentazione $+V_{CC}$ e la massa. In questo modo, in assenza di segnale applicato al terminale invertente, la tensione di uscita non è nulla, ma assume il valore costante $V_{CC}/2$.

Quando poi viene applicato un segnale bipolare al terminale invertente, questo segnale si sovrappone alla componente continua $V_{CC}/2$ per cui la sua escursione può avvenire sia al di sopra che al di sotto di $V_{CC}/2$ e, in particolare le semionde negative del segnale di uscita possono svilupparsi fra il livello $V_{CC}/2$ e il livello di massa.

L'OPERAZIONALE (IN CATENA APERTA) COME COMPARATORE



Si è già precisato che l'operazionale può funzionare come amplificatore lineare (per segnali di ampiezza accettabile) soltanto se è in catena chiusa, e, precisamente, solo quando è collegato a una rete di retroazione negativa.

Quando invece non è collegato ad alcun componente esterno (cioè quando è ad anello aperto o "open loop"), l'amp. op. è, di fatto, sempre in saturazione: basta una tensione di ingresso che superi l'ampiezza di un centinaio di microvolt per portare il livello della tensione di uscita ai valori di saturazione.

Siccome quella di un centinaio di microvolt è un'ampiezza piccolissima, molto minore dell'ampiezza dei segnali che interessano le normali applicazioni elettroniche, si usa approssimare a zero questi centomicrovolt e ritenere che **l'operazionale in catena aperta sia portato in saturazione da qualunque segnale di ampiezza diversa da zero**, ossia che **l'operazionale open loop sia sempre in saturazione (positiva o negativa)**.

Il componente in anello aperto è utilizzato per realizzare **circuiti squadratori** e circuiti **comparatori**.

In ogni caso, qualunque sia il circuito realizzato (squadratore, comparatore invertente o non invertente), il comportamento dell'operazionale open loop può essere così sintetizzato:

- quando il **potenziale dell'ingresso non invertente supera** il potenziale dell'ingresso invertente, la **tensione di uscita** è uguale al valore di **saturazione positiva**
- quando il **potenziale dell'ingresso non invertente è inferiore** al potenziale dell'ingresso invertente, la **tensione di uscita** è uguale al valore di **saturazione negativa**.

$$V_i^+ > V_i^- \Rightarrow V_o = V_{O(SAT)}^+$$

$$V_i^+ < V_i^- \Rightarrow V_o = V_{O(SAT)}^-$$

CIRCUITO SQUADRATORE

Siccome la **tensione di uscita** dell'operazionale in catena aperta può assumere, qualunque sia il segnale di ingresso, solo due valori ($V_{O(SAT)}^+$ e $V_{O(SAT)}^-$), essa è sempre **un'onda rettangolare**, cioè un segnale a due soli livelli (segnale digitale binario).

Pertanto: **ogni operazionale in catena aperta si comporta come squadratore**, nel senso **che trasforma qualunque segnale di ingresso in un'onda rettangolare** che commuta fra i livelli $V_{O(SAT)}^+$ e $V_{O(SAT)}^-$.

CIRCUITI COMPARATORI

Sono operazionali in catena aperta (e quindi squadratori) ai quali sono applicati:

- una tensione v_i di **ingresso**,
- una tensione V_{RIF} di **referimento** (che può anche essere nulla).

CLASSIFICAZIONE DEI COMPARATORI

COMPARATORI INVERTENTI E NON INVERTENTI

- Se la **tensione v_i di ingresso** è **applicata** al **terminale non invertente** (e la tensione di riferimento è applicata al terminale invertente), il circuito viene chiamato “**comparatore non invertente**”
- Se la **tensione v_i di ingresso** è **applicata** al **terminale *invertente*** (e la tensione di riferimento è applicata al terminale non invertente), il circuito viene chiamato “**comparatore *invertente***”.

COMPARATORI DI ZERO E COMPARATORI CON RIFERIMENTO DIVERSO DA ZERO

- Se la **tensione di riferimento è nulla**, il **terminale** di ingresso destinato alla tensione di riferimento **viene posto a massa** e il circuito viene chiamato “**comparatore di zero**” o “**rivelatore di zero**” o “**zero crossing detector**”.

Il comparatore di zero **segnala**, con una **commutazione della tensione di uscita dal livello basso $V_{O(SAT)}^-$ al livello alto $V_{O(SAT)}^+$** (o viceversa) **il passaggio del segnale di ingresso per il livello zero**.

In particolare se il comparatore di zero è non invertente, il passaggio in salita per lo zero è segnalato da una commutazione dell'uscita dal livello basso al livello alto.

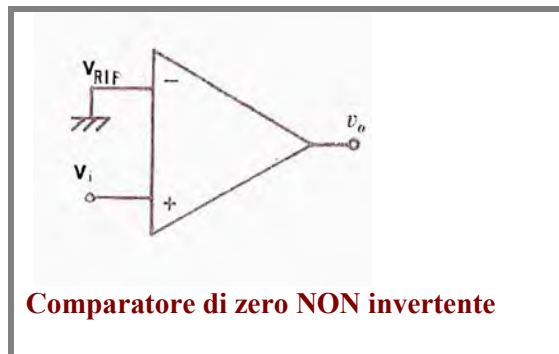
Se il comparatore di zero è *invertente*, il passaggio in salita per lo zero è segnalato da una commutazione dell'uscita *dal livello alto al livello basso*.

- Se la **tensione di riferimento è diversa da zero** (per esempio $V_{RIF} = 2V$), il circuito viene chiamato “**comparatore con riferimento diverso da zero**” o “**rivelatore di soglia**”.

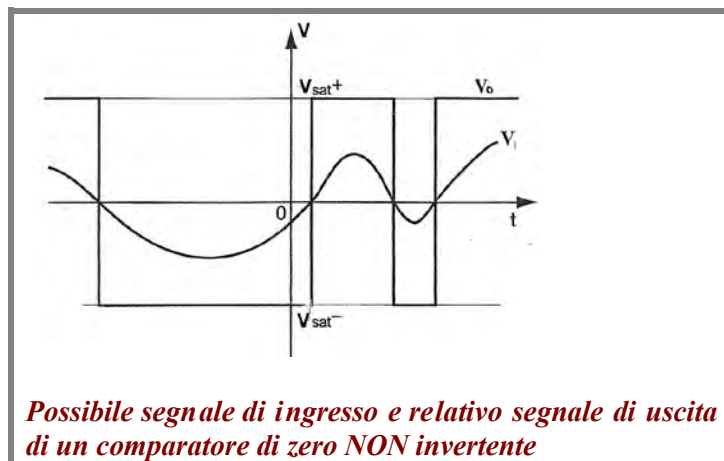
Al terminale di ingresso destinato alla tensione di riferimento **viene applicata una tensione continua di valore V_{RIF}** (per esempio $V_{RIF} = 2V$). Questo comparatore **segnala**, con una commutazione della tensione di uscita dal livello basso $V_{O(SAT)}^-$ al livello alto $V_{O(SAT)}^+$ (o viceversa) **il passaggio del segnale di ingresso per il livello V_{RIF}** , (per esempio per il livello $V_{RIF} = 2V$).

Osserviamo infine che la tensione di riferimento può anche essere negativa rispetto a massa e che, in qualche caso, non è continua come noi ora ipotizziamo, ma variabile anch'essa nel tempo.

COMPARATORE NON INVERTENTE CON RIFERIMENTO NULLO (RIVELATORE DI ZERO)



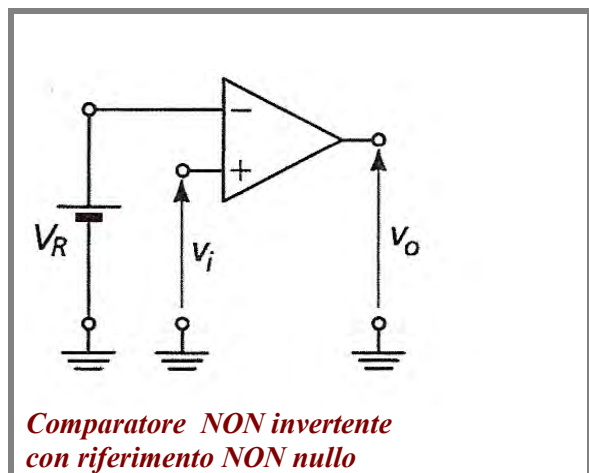
Essendo il comparatore non invertente, il segnale v_i di ingresso è applicato al terminale non invertente. Il segnale di riferimento sarebbe da applicare al terminale invertente. In questo caso il riferimento è zero e perciò il terminale invertente è collegato a massa.



FUNZIONAMENTO DEL CIRCUITO

Come si vede dalla figura, quando il **segnale di ingresso** (segnale **curvilineo** nella figura) **supera il livello zero** (cioè si trova al di sopra dell'asse orizzontale), la **tensione di uscita** è **alta** e vale $V_{O(SAT)}^+$. Quando invece il segnale di ingresso si trova **al di sotto del livello zero** (cioè si trova al di sotto dell'asse dei tempi), la **tensione di uscita** è **bassa** e vale $V_{O(SAT)}^-$.

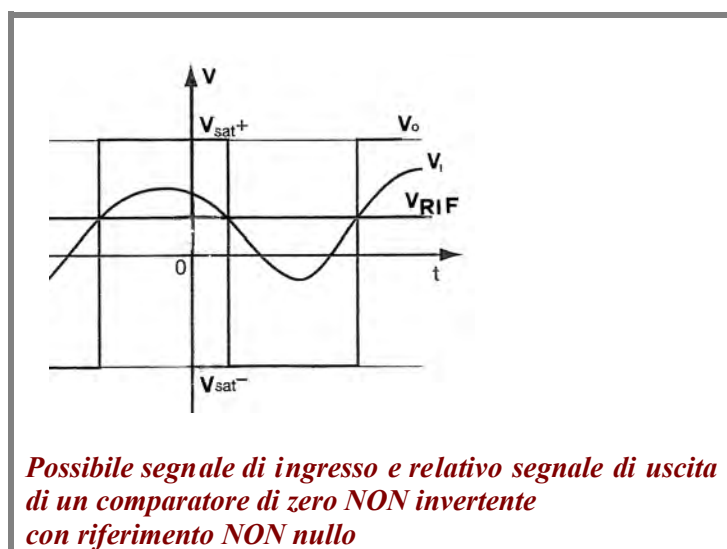
COMPARATORE NON INVERTENTE CON RIFERIMENTO NON NULO (RIVELATORE, NON INVERTENTE, DI SOGLIA)



Essendo il comparatore non invertente, il segnale v_i di ingresso è applicato al terminale non invertente.

Il segnale di riferimento deve essere applicato al terminale invertente.

Il segnale costante di riferimento è schematizzato con un generatore di tensione costante (batteria) applicato al terminale invertente dell'operazionale.

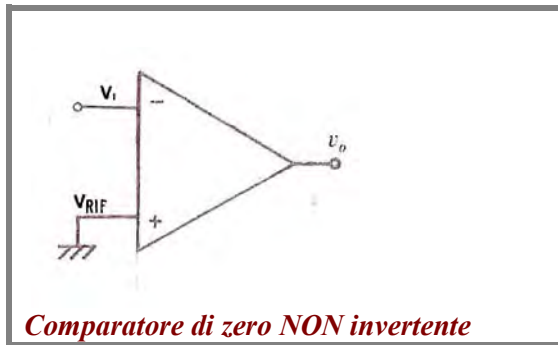


FUNZIONAMENTO DEL CIRCUITO

Quando il segnale di ingresso (segnale curvilineo nella figura) supera il livello V_{RIF} (cioè si trova al di sopra della semiretta orizzontale $V = V_{RIF}$), la tensione di uscita è alta e vale $V_{O(SAT)}^+$.

Quando invece il segnale di ingresso si trova al di sotto livello V_{RIF} (cioè si trova al di sotto della semiretta orizzontale $V = V_{RIF}$), la tensione di uscita è bassa e vale $V_{O(SAT)}^-$.

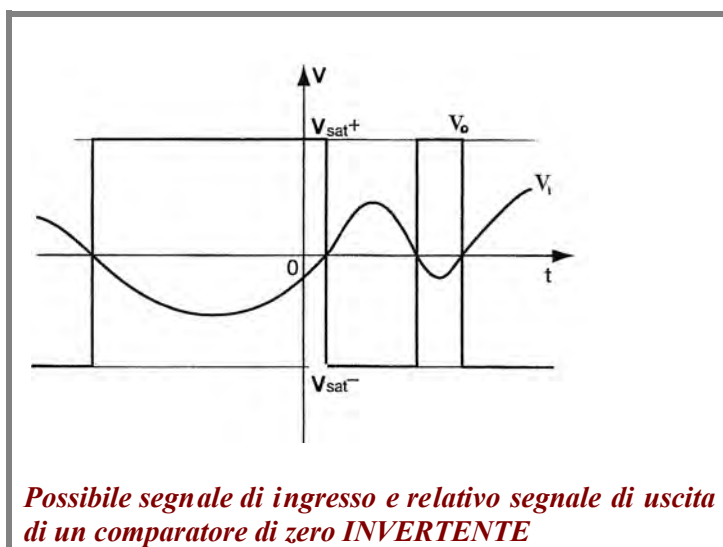
COMPARATORE INVERTENTE CON RIFERIMENTO NULLO (RIVELATORE DI ZERO)



Essendo il comparatore di tipo invertente, il segnale v_i di ingresso è applicato al terminale invertente.

Il segnale di riferimento dovrebbe essere applicato al terminale non invertente.

In questo caso il segnale di riferimento è zero, perciò il terminale non invertente dell'operazionale viene posto a massa.

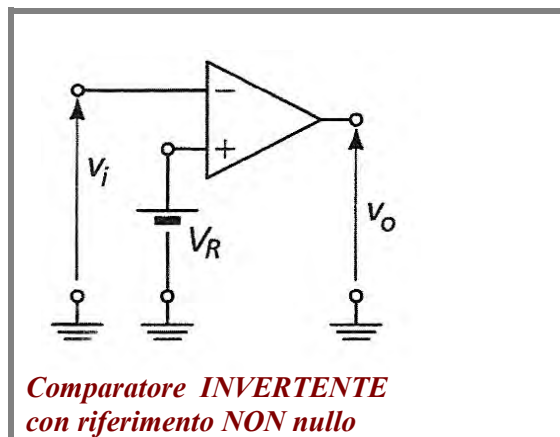


FUNZIONAMENTO DEL CIRCUITO

Quando il segnale di ingresso (segnale curvilineo nella figura) **supera** il livello **zero** (cioè si trova **al di sopra** dell'asse dei tempi), la tensione di uscita è **bassa** e vale $V_{O(SAT)}^-$.

Quando invece il segnale di ingresso si trova **al di sotto** del livello **zero** (cioè si trova **al di sotto** dell'asse dei tempi), la tensione di uscita è **alta** e vale $V_{O(SAT)}^+$.

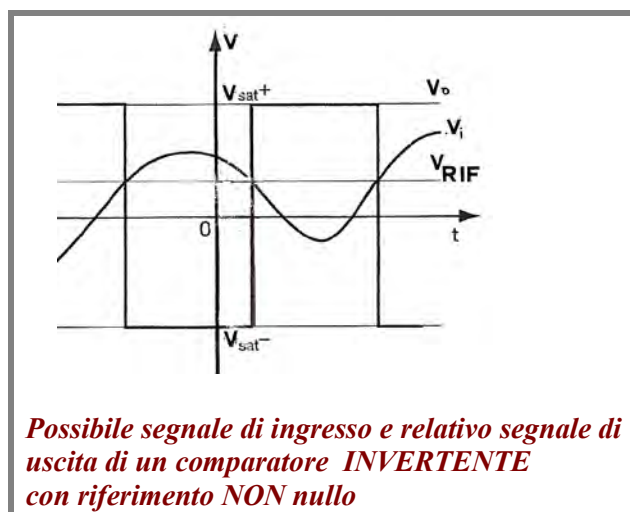
COMPARATORE INVERTENTE CON RIFERIMENTO NON NULLO (RIVELATORE INVERTENTE DI SOGLIA)



Essendo il comparatore di tipo invertente, il segnale v_i di ingresso è applicato al terminale invertente.

Il segnale di riferimento deve essere applicato al terminale non invertente.

Il segnale costante di riferimento è schematizzato con un generatore di tensione costante (batteria) applicato al terminale non invertente dell'operazionale.

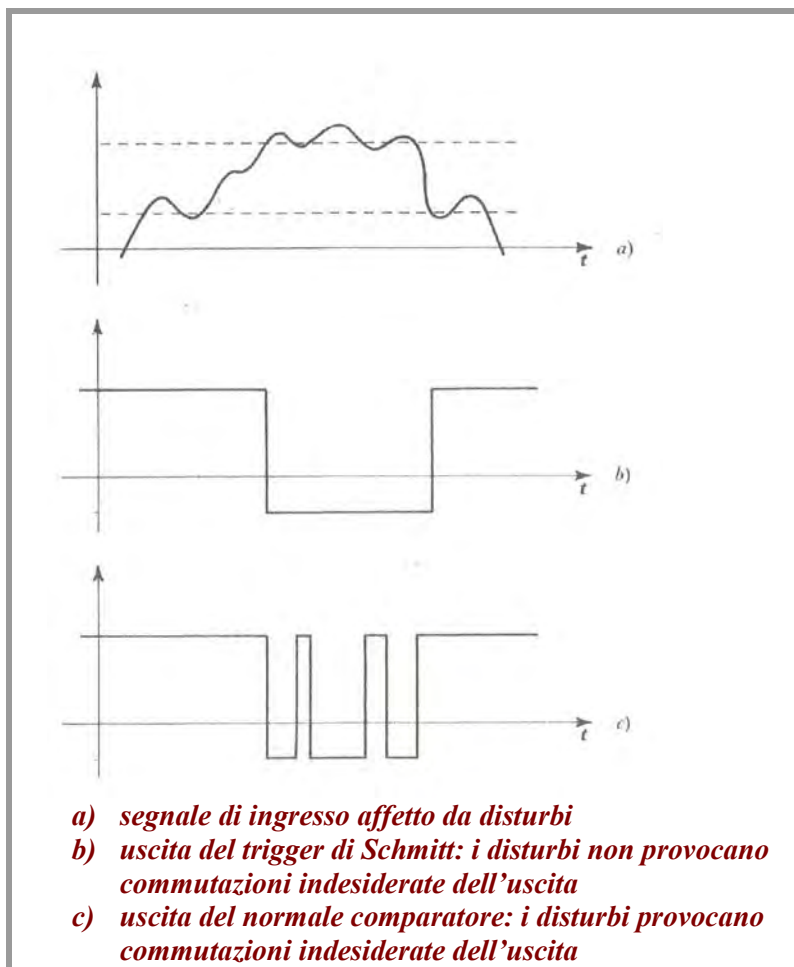


FUNZIONAMENTO DEL CIRCUITO

Quando il segnale di ingresso (segnale curvilineo nella figura) **supera** il livello V_{RIF} (cioè si trova **al di sopra** della semiretta orizzontale $V = V_{RIF}$), la tensione di uscita è **bassa** e vale $V_{O(SAT)}^-$.

Quando invece il segnale di ingresso si trova **al di sotto** del livello V_{RIF} (cioè si trova al di sotto della semiretta orizzontale $V = V_{RIF}$), la tensione di uscita è **alta** e vale $V_{O(SAT)}^+$.

TRIGGER DI SCHMITT (COMPARATORE CON ISTERESI)



A COSA SERVE IL CIRCUITO

Il trigger di Schmitt ha la stessa funzione del normale comparatore e cioè di “**segnalare**”, con una commutazione della tensione di uscita, **il passaggio della tensione di ingresso per un livello di riferimento V_{RIF}** , che può anche essere di 0V.

Il trigger di Schmitt presenta però il vantaggio di **evitare commutazioni indesiderate, determinate da piccole variazioni casuali del segnale di ingresso (dovute per esempio, al rumore).**

Perciò possiamo dire che il trigger è un’evoluzione del circuito comparatore.

Come nel comparatore “semplice”, la **tensione di uscita** del comparatore con isteresi è **un’onda rettangolare** che commuta fra due valori:

“ $+V_{SAT}$ ” e “ $-V_{SAT}$ ”.

Il comparatore con isteresi riesce a evitare le commutazioni indesiderate perché **la sua uscita commuta in corrispondenza non di uno, ma di due distinti valori di soglia**: uno posto al di sopra del livello di riferimento e uno posto al di sotto. I valori di soglia sono chiamati “**tensione di soglia superiore (V_T^+)**” e “**tensione di soglia inferiore (V_T^-)**”.

La **differenza** fra le due soglie è chiamata “**tensione di isteresi V_H** ” o semplicemente “**isteresi**”.

Il trigger evita tutte le commutazioni che avverrebbero per effetto di escursioni del segnale di ingresso di ampiezza minore della tensione di isteresi.

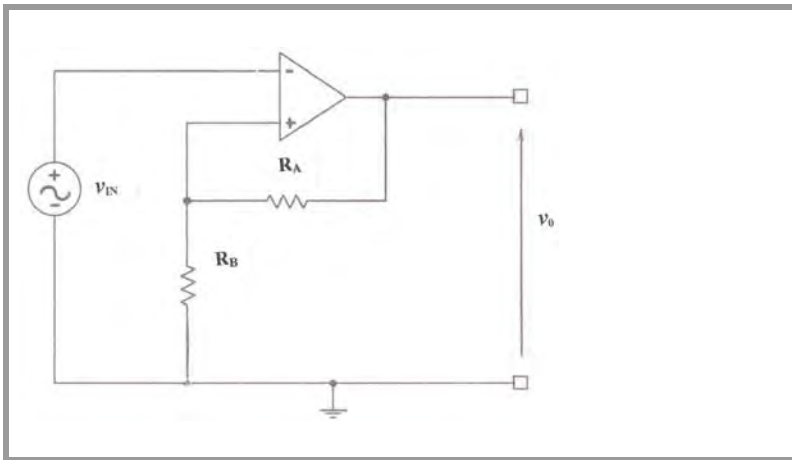
STRUTTURA DEL CIRCUITO

Il trigger di Schmitt può essere realizzato aggiungendo, al normale comparatore, **una rete di retroazione positiva (R_A , R_B)**, che colleghi il terminale di **uscita** al terminale di ingresso **non invertente**. La retroazione positiva riporta in ingresso una porzione del segnale di uscita in maniera tale che vada a sommarsi col segnale già presente all'ingresso e determina un aumento dell'amplificazione.

Nel comparatore con isteresi la **retroazione positiva** fa in modo che le **commutazioni** avvengano **istantaneamente**, anche in presenza di amplificazioni non infinite dell'operazionale. La velocità delle commutazioni risulta infatti dipendente solo dallo slew-rate dell'integrato.

La **retroazione** inoltre rende i **valori** delle tensioni di soglia **dipendenti** dal valore della tensione di **uscita**.

TRIGGER DI SCHMITT (COMPARATORE CON ISTERESI) INVERTENTE CON RIFERIMENTO ZERO



FUNZIONAMENTO DEL CIRCUITO

Il segnale di **ingresso** è applicato al **terminale invertente** dell'operazionale.

La tensione di **riferimento** è il potenziale di **massa** in quanto l'estremo inferiore della resistenza R_B di retroazione positiva è collegato a massa.

Il circuito **segnala**, mediante commutazioni dell'uscita, **il passaggio del segnale di ingresso per due valori di soglia, simmetrici rispetto al livello zero: " V_T^+ " e " V_T^- ".**

Quando il segnale di ingresso, crescendo, supera la soglia superiore V_T^+ , la tensione di uscita ha una commutazione in discesa, cioè passa dal livello alto $+V_{SAT}$ al livello basso $-V_{SAT}$; invece quando il segnale di ingresso, decrescendo, scende al di sotto della soglia inferiore V_T^- , la tensione di uscita ha una commutazione in salita, cioè passa dal livello basso al livello alto.

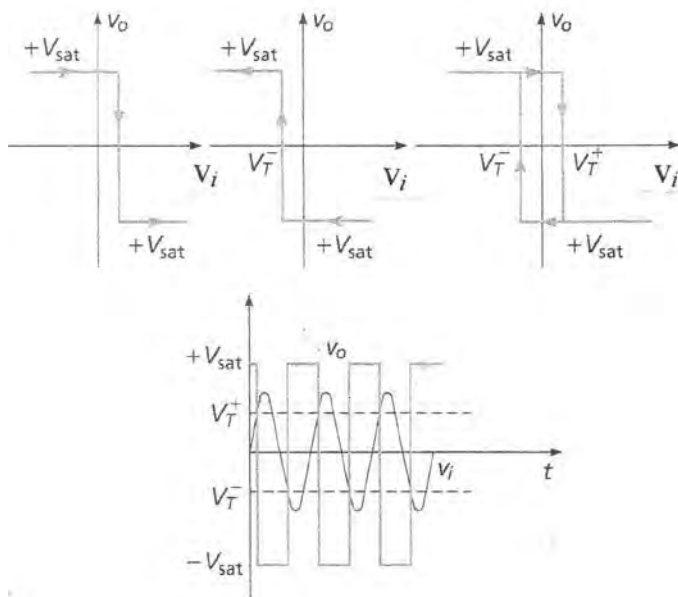
Si può dimostrare (vedi analisi del circuito) che i valori delle tensioni di soglia dipendono dai due valori (" $+V_{SAT}$ " e " $-V_{SAT}$ ") della tensione di uscita e dalle resistenze di retroazione positiva:

tensione di soglia inferiore
$$V_T^- = -\frac{R_B}{R_B + R_A} \cdot |-V_{SAT}|$$

tensione di soglia superiore
$$V_T^+ = +\frac{R_B}{R_B + R_A} \cdot |+V_{SAT}|$$

tensione di isteresi
$$V_H = |V_T^+| - |V_T^-| = 2 \cdot \frac{R_B}{R_B + R_A} \cdot |+V_{SAT}|$$

Osserviamo che l'**isteresi** è tanto più **ristretta** (e le due **soglie** sono tanto **più vicine fra loro**) quanto più la **resistenza R_B** (in basso nel circuito) è **piccola** rispetto ad R_A .



In alto: curve caratteristiche con V_u in funzione di V_i

In basso: andamento temporale dei segnali di ingresso e di uscita.

Quando il segnale di ingresso, crescendo, supera la soglia superiore V_T^+ , la tensione di uscita ha una commutazione in discesa, cioè passa livello alto al livello basso; invece quando il segnale di ingresso, decrescendo, va al di sotto della soglia inferiore V_T^- , la tensione di uscita ha una commutazione in salita, cioè passa livello basso al livello alto.

ANALISI DEL FUNZIONAMENTO

Prima di tutto ricordiamo che l'uscita di un operazionale in catena aperta o retroazionato positivamente può commutare solo negli istanti nei quali il potenziale dell'ingresso invertente e il potenziale dell'ingresso non invertente si uguagliano.

Supponiamo che inizialmente la tensione di uscita sia positiva e valga $+V_{SAT}$. Di conseguenza, essendo il circuito di tipo invertente, la tensione di ingresso (potenziale dell'ingresso invertente) deve risultare negativa.

Siccome l'uscita è positiva, risulta positivo (per effetto della retroazione) anche il potenziale dell'ingresso non invertente, che assume il valore:

$$V_T^+ = + \frac{R_B}{R_B + R_A} \cdot | +V_{SAT} |$$

Se la tensione di ingresso, applicata all'ingresso invertente, crescendo⁶, raggiunge e supera il potenziale V_T^+ del terminale non invertente (cioè la tensione di soglia superiore), si ha una commutazione della tensione di uscita, la quale, da positiva che era, diventa negativa e assume il valore $-V_{SAT}$.

Di conseguenza (grazie al collegamento di retroazione) il potenziale del terminale non invertente assume anch'esso un valore negativo e cioè il valore di soglia inferiore.

$$V_T^- = - \frac{R_B}{R_B + R_A} \cdot | -V_{SAT} |$$

A questo punto si avrà una nuova commutazione dell'uscita (dal valore negativo a un valore positivo) quando la tensione di ingresso, applicata all'ingresso invertente, toccherà, decrescendo⁷, il valore V_T^- del terminale non invertente.

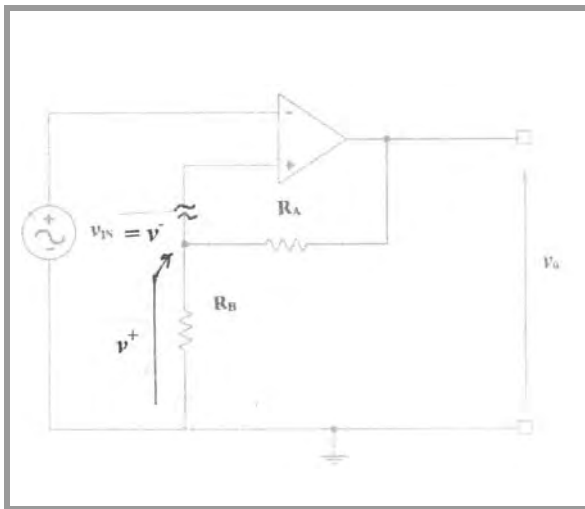
Pertanto, avvenuta la nuova commutazione, il funzionamento ora descritto si ripeterà indefinitamente.

⁶ Essendo, per ipotesi, inizialmente negativa, la tensione di ingresso può raggiungere il potenziale positivo V_T^+ solo crescendo.

⁷ Nell'istante in esame la tensione di ingresso, avendo superato la soglia superiore, è positiva, per cui può raggiungere il valore V_T^- solamente decrescendo.

ANALISI DEL CIRCUITO:

COME SI RICAVA L'ESPRESSIONE DELLE TENSIONI DI SOGLIA



Le soglie V_T^- e V_T^+ sono valori assunti dal potenziale V^+ dell'ingresso non invertente. Perciò dobbiamo trovare una relazione fra V^+ e la tensione di uscita V_0 .

Siccome l'impedenza di ingresso dell'operazionale è elevatissima, possiamo ritenere nulle le correnti entranti o uscenti dal terminale di ingresso e quindi ritenere "aperto" (interrotto) il collegamento fra il terminale "+" e il terminale comune ad R_A ed R_B .

Le resistenze R_A ed R_B sono allora in serie tra loro e la tensione di uscita risulta applicata ai capi della serie $R_A + R_B$.

Essendo poi V^+ la tensione ai capi di R_B , dalla regola del partitore di tensione si ha:

$$V^+ = V_0 \cdot \frac{R_B}{R_B + R_A}$$

Siccome V_0 può assumere solo due valori ($-V_{SAT}$ e $+V_{SAT}$), possiamo scrivere:

$$V^+ = \pm \frac{R_B}{R_B + R_A} \cdot |V_{SAT}|$$

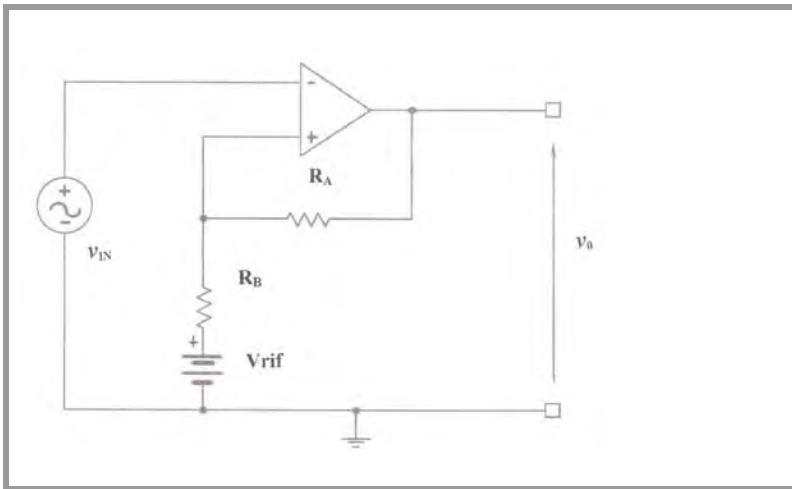
Da questa relazione si vede che anche V^+ può assumere solo due valori, i quali prendono il nome di “valori di soglia”:

$$-\frac{R_B}{R_B + R_A} \cdot |V_{SAT}| = V_T^-$$

e

$$+\frac{R_B}{R_B + R_A} \cdot |V_{SAT}| = V_T^+$$

TRIGGER DI SCHMITT (COMPARATORE CON ISTERESI) INVERTENTE CON RIFERIMENTO DIVERSO DA ZERO



FUNZIONAMENTO DEL CIRCUITO

Il segnale di **ingresso** è applicato al **terminale invertente** dell'operazionale.

La tensione di **riferimento** è schematizzata dal generatore di tensione continua V_{RIF} .

Il circuito **segnala**, mediante commutazioni dell'uscita, **il passaggio del segnale di ingresso per due valori di soglia** (“ $+V_{SAT}$ ” e “ $-V_{SAT}$ ”), posti uno al di sotto e uno al di sopra del livello V_{RIF} .

Quando il segnale di ingresso, crescendo, supera la soglia superiore V_T^+ , la tensione di uscita ha una commutazione in discesa, cioè passa dal livello alto $+V_{SAT}$ al livello basso $-V_{SAT}$; invece quando il segnale di ingresso, decrescendo, scende al di sotto della soglia inferiore V_T^- , la tensione di uscita ha una commutazione in salita, cioè passa dal livello basso al livello alto.

Si può dimostrare che i valori delle tensioni di soglia dipendono dai due valori (“ $+V_{SAT}$ ” e “ $-V_{SAT}$ ”) della tensione di uscita, dal valore di V_{RIF} e dalle resistenze di retroazione positiva:

tensione di soglia inferiore
$$V_T^- = -\frac{R_B}{R_B + R_A} \cdot |-V_{SAT}| + \frac{R_A}{R_B + R_A} \cdot V_{RIF}$$

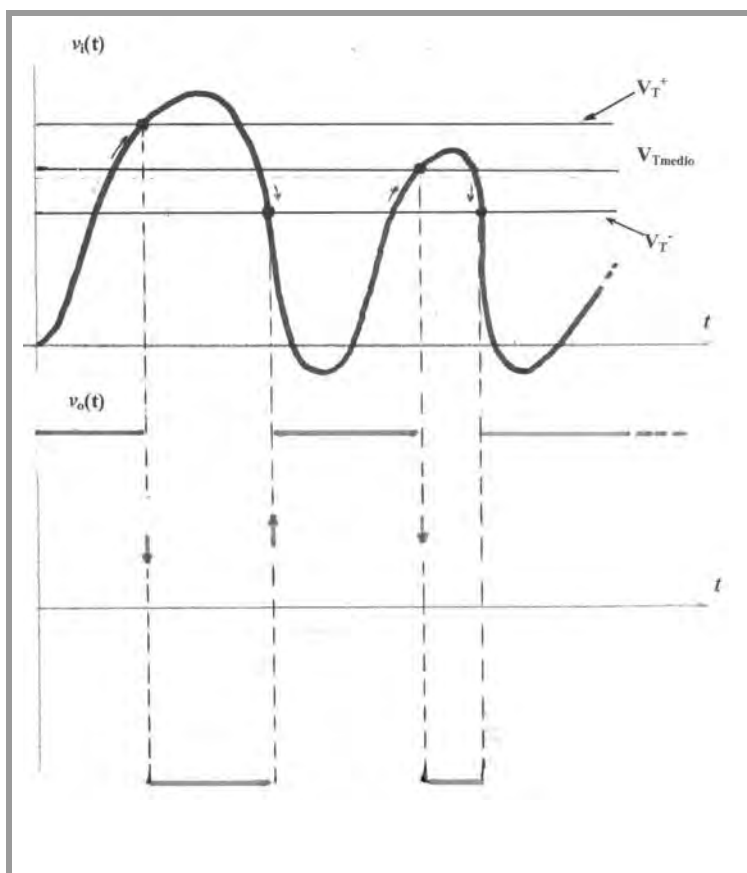
tensione di soglia superiore
$$V_T^+ = +\frac{R_B}{R_B + R_A} \cdot |+V_{SAT}| + \frac{R_A}{R_B + R_A} \cdot V_{RIF}$$

tensione di isteresi
$$V_H = |V_T^+| - |V_T^-| = 2 \cdot \frac{R_B}{R_B + R_A} \cdot |+V_{SAT}|$$

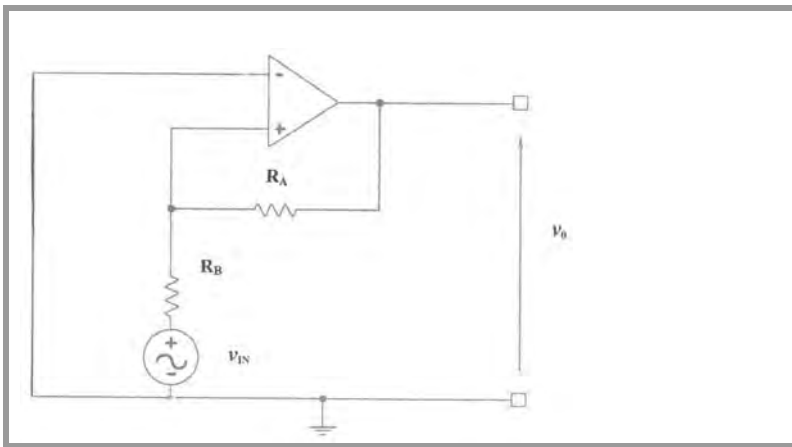
valore centrale di commutazione
$$V_{Tmedio} = \frac{|V_T^+| + |V_T^-|}{2} = \frac{R_A}{R_B + R_A} \cdot V_{RIF}$$

Osserviamo che:

- l'isteresi è tanto più **ristretta** (e le due **soglie** sono tanto **più vicine fra loro**) quanto più la **resistenza R_B** (in basso nel circuito) è **piccola** rispetto ad R_A .
- il valore centrale di commutazione (V_{Tmedio}) è (come si vede dall'ultima relazione) sempre al di sotto di V_{RIF} .



TRIGGER DI SCHMITT (COMPARATORE CON ISTERESI) NON INVERTENTE CON RIFERIMENTO ZERO



FUNZIONAMENTO DEL CIRCUITO

Il segnale di **ingresso** è applicato, attraverso la resistenza R_B , al **terminale non invertente** dell'operazionale.

La tensione di **riferimento** è il potenziale di **massa** in quanto il terminale invertente dell'operazionale è collegato a massa.

Il circuito **segnala**, mediante commutazioni dell'uscita, **il passaggio del segnale di ingresso per due valori di soglia, simmetrici rispetto al livello zero**: “ $+V_{SAT}$ ” e “ $-V_{SAT}$ ”.

Quando il segnale di ingresso, crescendo, supera la soglia superiore V_T^+ , la tensione di uscita ha una commutazione in salita, cioè passa dal livello basso $-V_{SAT}$ al livello alto $+V_{SAT}$; invece quando il segnale di ingresso, decrescendo, scende al di sotto della soglia inferiore V_T^- , la tensione di uscita ha una commutazione in discesa, cioè passa dal livello alto al livello basso.

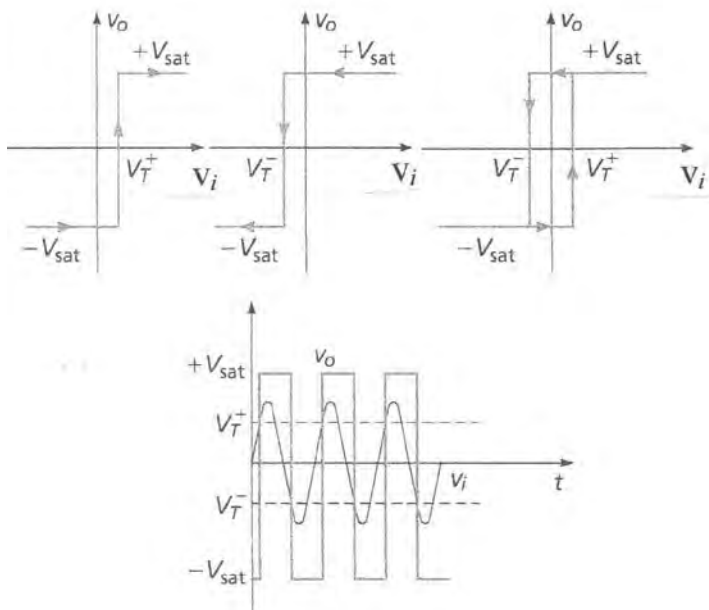
Si può dimostrare che i valori delle tensioni di soglia dipendono dai due valori (“ $+V_{SAT}$ ” e “ $-V_{SAT}$ ”) della tensione di uscita e dalle resistenze di retroazione positiva:

tensione di soglia inferiore
$$V_T^- = -\frac{R_B}{R_A} \cdot |-V_{SAT}|$$

tensione di soglia superiore
$$V_T^+ = +\frac{R_B}{R_A} \cdot |+V_{SAT}|$$

tensione di isteresi
$$V_H = |V_T^+| - |V_T^-| = 2 \cdot \frac{R_B}{R_A} \cdot |+V_{SAT}|$$

Osserviamo che l'**isteresi** è tanto più **ristretta** (e le due **soglie** sono tanto **più vicine fra loro**) quanto più la **resistenza R_B** (in basso nel circuito) è **piccola** rispetto ad R_A .

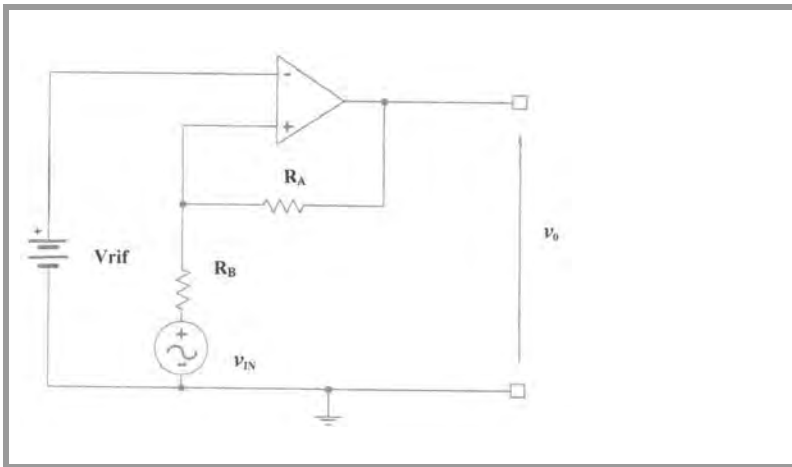


In alto: curve caratteristiche con V_u in funzione di V_i

In basso: andamento temporale dei segnali di ingresso e di uscita.

Quando il segnale di ingresso, crescendo, supera la soglia superiore V_T^+ , la tensione di uscita ha una commutazione in salita, cioè passa livello basso al livello alto; invece quando il segnale di ingresso, decrescendo, va al di sotto della soglia inferiore V_T^- , la tensione di uscita ha una commutazione in discesa, cioè passa livello alto al livello basso.

TRIGGER DI SCHMITT (COMPARATORE CON ISTERESI) NON INVERTENTE CON RIFERIMENTO DIVERSO DA ZERO



FUNZIONAMENTO DEL CIRCUITO

Il segnale di **ingresso** è applicato, attraverso la resistenza R_B , al terminale **non** invertente dell'operazionale.

La tensione di **riferimento** è schematizzata dal generatore di tensione continua V_{RIF} , posta fra il terminale invertente e la massa.

Il circuito **segnala**, mediante commutazioni dell'uscita, **il passaggio del segnale di ingresso per due valori di soglia** (“ $+V_{SAT}$ ” e “ $-V_{SAT}$ ”), posti uno al di sotto e uno al di sopra del livello V_{RIF} . Quando il segnale di ingresso, crescendo, supera la soglia superiore V_T^+ , la tensione di uscita ha una commutazione in salita, cioè passa dal livello basso $-V_{SAT}$ al livello alto $+V_{SAT}$; invece quando il segnale di ingresso, decrescendo, scende al di sotto della soglia inferiore V_T^- , la tensione di uscita ha una commutazione in discesa, cioè passa dal livello alto al livello basso.

Si può dimostrare che i valori delle tensioni di soglia dipendono dai due valori (“ $+V_{SAT}$ ” e “ $-V_{SAT}$ ”) della tensione di uscita, dal valore di V_{RIF} e dalle resistenze di retroazione positiva:

tensione di soglia inferiore
$$V_T^- = -\frac{R_B}{R_A} \cdot |-V_{SAT}| + \frac{R_B + R_A}{R_A} \cdot V_{RIF}$$

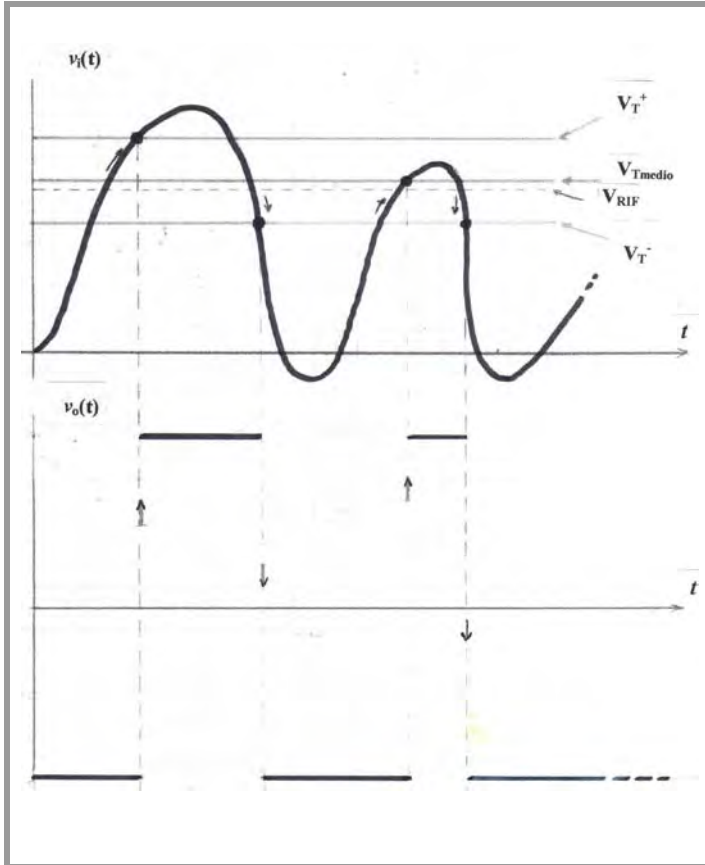
tensione di soglia superiore
$$V_T^+ = +\frac{R_B}{R_A} \cdot |+V_{SAT}| + \frac{R_B + R_A}{R_A} \cdot V_{RIF}$$

tensione di isteresi
$$V_H = |V_T^+| - |V_T^-| = 2 \cdot \frac{R_B}{R_A} \cdot |+V_{SAT}|$$

valore centrale di commutazione
$$V_{Tmedio} = \frac{|V_T^+| + |V_T^-|}{2} = \frac{R_B + R_A}{R_A} \cdot V_{RIF}$$

Osserviamo che:

- l'isteresi è tanto più **ristretta** (e le due **soglie** sono tanto **più vicine fra loro**) quanto più la **resistenza R_B** (in basso nel circuito) è **piccola** rispetto ad R_A .
- il valore centrale di commutazione (V_{Tmedio}) è (come si vede dall'ultima relazione) sempre al di sopra di V_{RIF} .



CAPITOLO 26

OSCILLATORI SINUSOIDALI

COMPLEMENTI

***TRATTAZIONE ANALITICA DELL'OSCILLATORE SINUSOIDALE:
LE CONDIZIONI DI MANTENIMENTO E DI INNESCO DELLE
OSCILLAZIONI***

STABILITA' IN FREQUENZA DEGLI OSCILLATORI

OSCILLATORI AL QUARZO (O A CRISTALLO)

TRATTAZIONE ANALITICA DELL'OSCILLATORE SINUSOIDALE: LE CONDIZIONI DI MANTENIMENTO E DI INNESCO DELLE OSCILLAZIONI

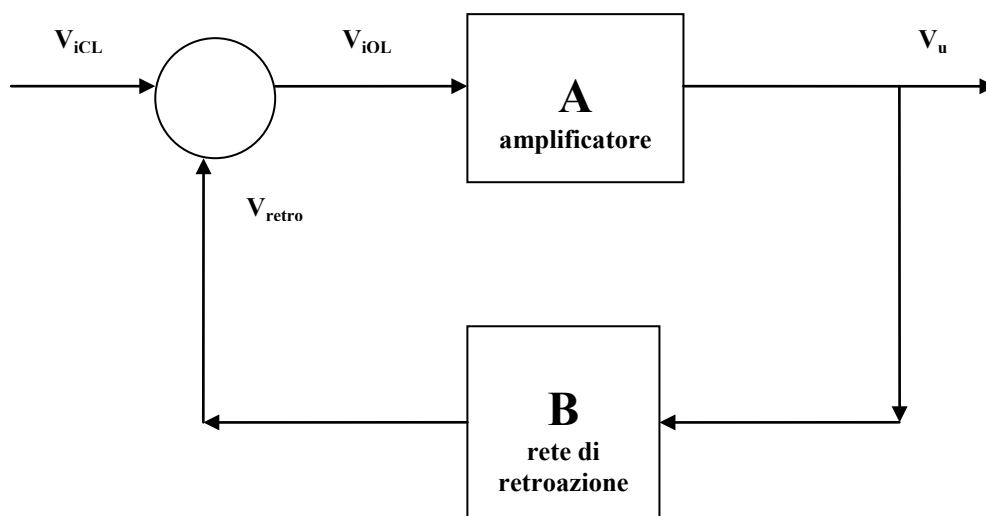
PREMESSA n.1

LA RETROAZIONE POSITIVA

La retroazione è la tecnica per la quale il segnale di uscita di un amplificatore (o una porzione di esso) viene riportato in ingresso.

Nella reazione positiva, il **segnale di retroazione**, cioè la porzione di segnale di uscita prelevata dalla rete di retroazione, risulta **in fase** con il **segnale già presente** all'ingresso dell'amplificatore e quindi **si somma** ad esso.

Di conseguenza l'**amplificazione** dell'intero circuito (cioè del circuito reazionato) **cresce** rispetto all'amplificazione del solo amplificatore non reazionato.



La crescita, per effetto della reazione positiva, del segnale presente all'ingresso dell'amplificatore può avere (in certe condizioni) un effetto analogo a quello di un vero e proprio segnale di ingresso. Nel senso che nasce, all'interno del circuito, un nuovo segnale, per cui (anche in assenza di un segnale effettivamente applicato dall'esterno) il sistema si comporta come se al suo ingresso fosse applicata una sollecitazione.

Se poi all'ingresso del sistema è già applicato un segnale dall'esterno, allora il “nuovo segnale” si sovrappone a quello originario determinando una tendenza del sistema all'instabilità, cioè ad assumere valori di uscita crescenti in maniera incontrollata.

L'effetto della reazione positiva è dannoso per un amplificatore a meno che l'amplificatore non venga utilizzato per realizzare un oscillatore, o, più in generale, un generatore di forme d'onda.

Nel caso di impiego in un oscillatore, l'amplificatore esercita la sua funzione sul rumore interno al circuito, mentre la rete di retroazione esegue il filtraggio del rumore stesso in modo da trasformarlo in segnale sinusoidale.

PREMESSA n.2

OSCILLATORI E RETROAZIONE POSITIVA

Come già accennato, gli oscillatori possono essere realizzati sfruttando la **retroazione positiva** degli **amplificatori**.

Ci proponiamo di **determinare le condizioni nelle quali un amplificatore e un circuito di retroazione sono in grado di innescare e mantenere un'oscillazione**.

Dalle condizioni di mantenimento e di innesco delle oscillazioni derivano le **relazioni di progetto degli oscillatori**, cioè le relazioni che **legano i valori dei componenti circuitali alla frequenza di oscillazione** del dispositivo.

PREMESSA n.3

RISPOSTE IN FREQUENZA DELL'AMPLIFICATORE E DELLA RETE DI RETROAZIONE

Con “A” indicheremo la funzione di risposta in frequenza (funzione di trasferimento nella variabile “f”) di un amplificatore e con “B” la risposta in frequenza della sua rete di retroazione.

Ricordiamo che come “**risposta in frequenza**” di un quadripolo può essere assunta l'espressione, in funzione della frequenza o della pulsazione, del **rapporto fra la tensione di uscita e la tensione di ingresso** del quadripolo.

Per esempio la risposta in frequenza di un filtro passivo RC passabasso è:

$$\overline{A(\omega)} = \frac{\overline{V_u(\omega)}}{\overline{V_i(\omega)}} = \frac{1}{1+j\omega RC}$$

Come si vede, la risposta in frequenza è una grandezza complessa, sarà quindi dotata di un modulo:

$$\left| \overline{A(\omega)} \right| = \frac{\left| \overline{V_u(\omega)} \right|}{\left| \overline{V_i(\omega)} \right|} = \frac{1}{\sqrt{1+(\omega RC)^2}}$$

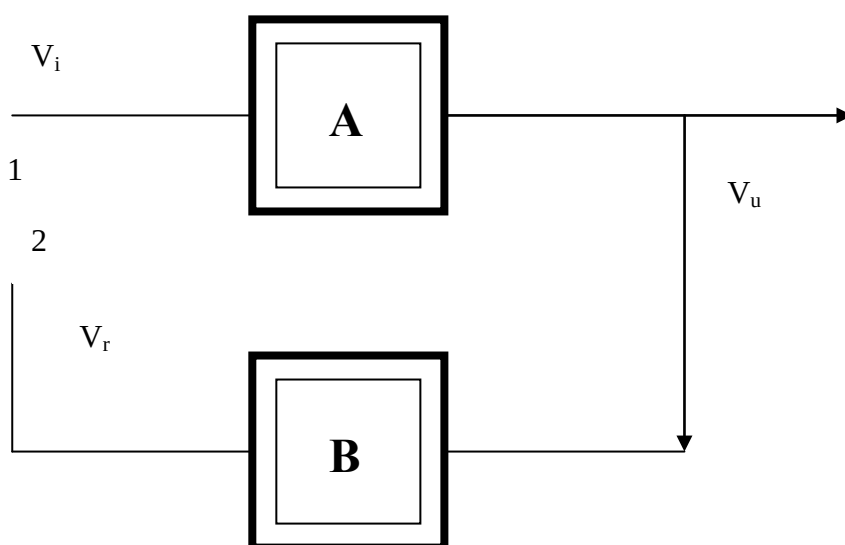
e di una fase:

$$\varphi(\omega) = -\arctg(\omega \cdot RC)$$

MANTENIMENTO E INNESCO DELLE OSCILLAZIONI

Consideriamo un dispositivo costituito da un amplificatore, indicato con la grandezza complessa $\bar{A}$, e da un blocco che funzioni da rete di retroazione (non specifichiamo ancora se positiva o negativa), indicato con la grandezza complessa $\bar{B}$.

Ipotizziamo che i segnali applicati siano sinusoidali e perciò usiamo numeri complessi.



$\bar{A}$ è il rapporto fra le tensioni di uscita e di ingresso dell'amplificatore e rappresenta la risposta in frequenza dell'amplificatore. $\bar{A}$ dipende dalla frequenza ed è anche detta "funzione di trasferimento":

$$\bar{A} = \bar{V}_u / \bar{V}_i$$

$\bar{B}$ è il rapporto fra le tensioni di uscita e di ingresso della rete di retroazione e rappresenta la risposta in frequenza della rete di retroazione.

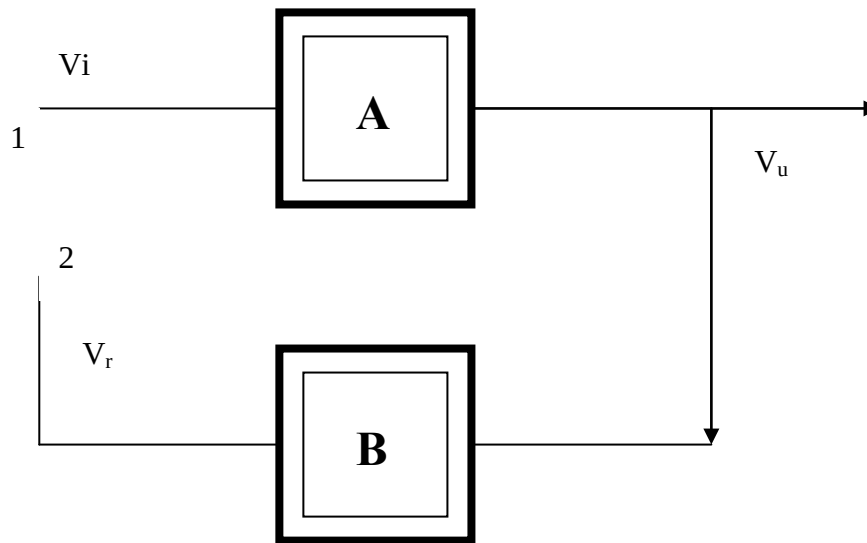
$\bar{B}$ dipende dalla frequenza ed è anche detta "funzione di trasferimento" della rete di retroazione.

Dallo schema del dispositivo in esame si vede che la tensione di ingresso del blocco di reazione $\bar{B}$ è la tensione di uscita dell'intero dispositivo, mentre la tensione di uscita di B è la tensione di retroazione V_r :

$$\bar{B} = \bar{V}_r / \bar{V}_u$$

Vogliamo appurare quali condizioni devono rispettare i blocchi $\bar{A}$ e $\bar{B}$ del sistema affinché questo sia in grado di MANTENERE un'oscillazione (cioè un segnale sinusoidale) a una determinata frequenza f_0 . Dire che il sistema è in grado di ***mantenere*** un'oscillazione a frequenza f_0 significa dire che ***se noi applichiamo dall'esterno un segnale sinusoidale di frequenza f_0 (il quale dà luogo a una uscita sinusoidale) e poi "stacchiamo" questo segnale, il sistema continua a produrre un'uscita sinusoidale anche in assenza della sollecitazione esterna.***

Noi vogliamo sapere a quali condizioni per $\bar{A}$ e per $\bar{B}$ è possibile questo mantenimento. Procediamo allora in questo modo:



Con l'anello di reazione aperto, tenendo cioè staccati i punti 1 e 2, applichiamo all'amplificatore un segnale sinusoidale V_i (generato dall'esterno) di frequenza f_0 . Questo segnale V_i sarà applicato all'amplificatore $\bar{A}$, all'uscita del quale darà luogo al segnale:

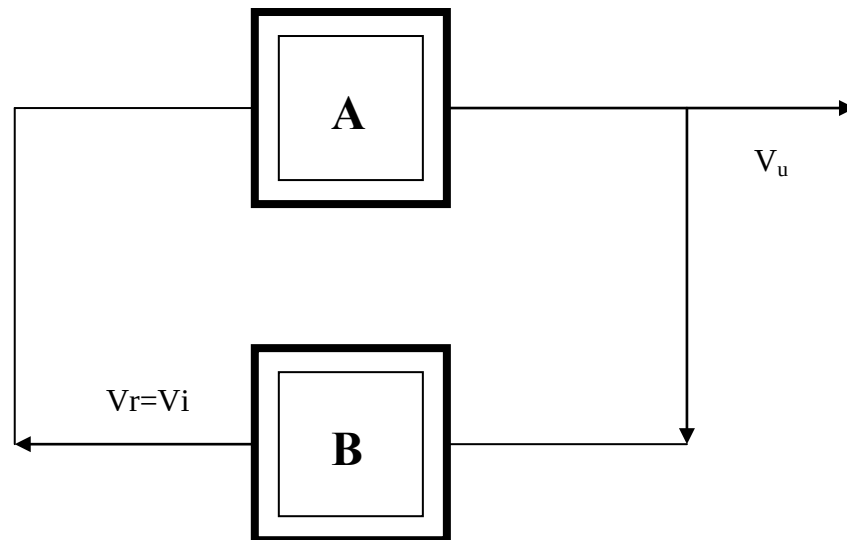
$$\bar{V}_u = \bar{A} \cdot \bar{V}_i$$

V_u poi risulterà applicato alla rete di reazione $\bar{B}$, all'uscita della quale darà luogo al segnale:

$$\bar{V}_r = \bar{B} \cdot \bar{V}_u = \bar{A} \cdot \bar{B} \cdot \bar{V}_i$$

Dunque l'applicazione del segnale sinusoidale V_i , di frequenza f_0 , al sistema ad anello aperto, produce, a valle del blocco B, il segnale $\mathbf{V_r = A \cdot B \cdot V_i}$ avente frequenza f_0 .

Chiediamoci ora cosa succederebbe se riuscissimo a ottenere, (all'uscita del blocco B) un segnale V_r uguale, in modulo e fase, a V_i e cosa succederebbe se (una volta rimosso V_i) applicassimo, all'ingresso del blocco A, proprio V_r , chiudendo l'anello di reazione.



Accade questo: il dispositivo continua a produrre un'uscita sinusoidale perché all'ingresso dell'amplificatore, pur non essendo più applicato il segnale esterno V_i , è ora applicato il segnale di retroazione V_r , uguale a V_i .

Il dispositivo cioè continua a produrre la stessa uscita sinusoidale a frequenza f_0 anche quando viene soppressa la sollecitazione esterna perché l'amplificatore "non può accorgersi" che ad esso non è più applicata la V_i , ma un segnale V_r proveniente dalla retroazione e non più dall'esterno.

Da quanto detto si vede che, per rendere possibile il mantenimento dell'oscillazione una volta cessata la sollecitazione esterna, è necessario che il segnale di reazione V_r sia uguale, in ampiezza (modulo) e in fase, al segnale V_i entrante nell'amplificatore.

Osservando la relazione:

$$\overline{V_r} = \overline{A} \cdot \overline{B} \cdot \overline{V_i}$$

Si vede che, affinché il mantenimento sia possibile (cioè affinché vi sia uguaglianza tra V_r e V_i) deve essere:

$\overline{A} \cdot \overline{B} = 1$	<i>condizione di BARKHAUSEN o di MANTENIMENTO delle oscillazioni</i> <i>(dove $A \cdot B$ = amplificazione di anello)</i>
---------------------------------------	--

La condizione impone che, in corrispondenza della frequenza f_0 desiderata, e solo di essa, il prodotto fra le funzioni di trasferimento (risposte in frequenza) del blocco amplificatore e del blocco di reazione deve essere pari a 1.

Tenendo presente che:

- due numeri complessi sono uguali quando hanno lo stesso modulo e la stessa fase (oppure la stessa parte reale e la stessa parte immaginaria)
- il numero 1 è un numero reale positivo e quindi ha fase nulla allora la condizione di Barkhausen

$$\overline{A} \cdot \overline{B} = 1$$

può essere sdoppiata in due condizioni, una sul **modulo** e una sulla *fase o argomento*:

$\left \overline{A} \cdot \overline{B} \right = 1$ <i>condizione di BARKHAUSEN applicata al modulo</i>	$\arg \overline{A} \cdot \overline{B} = 0^\circ$ <i>condizione di BARKHAUSEN applicata alla fase</i>
--	--

Ricordiamo che imporre:

$$\overline{A} \cdot \overline{B} = 1$$

equivale a imporre che il segnale presente all'uscita del blocco B di reazione abbia la stessa fase del segnale presente all'ingresso dell'amplificatore A e che i due segnali abbiano la stessa ampiezza.

Si può anche dire che la condizione di Barkhausen impone che il segnale presente all'ingresso dell'amplificatore ritorni allo stesso punto (dopo aver attraversato l'amplificatore stesso e la rete di reazione) con sfasamento nullo e senza modificazioni di ampiezza.

INNESCO DELLE OSCILLAZIONI

La condizione

$$\overline{A} \cdot \overline{B} = 1$$

è una condizione di mantenimento delle oscillazioni. Essa cioè garantisce, una volta soddisfatta, che un'oscillazione, già in qualche modo presente nel dispositivo, si conservi nel tempo. **Non** è però sufficiente a garantire l'**innescò** delle oscillazioni cioè **la nascita di oscillazioni in un sistema al quale non sia stata applicata nessuna sollecitazione dall'esterno se non la tensione continua di alimentazione.**

Per avere innescò di oscillazione è necessario che:

$\left \overline{A} \cdot \overline{B} \right > 1$ <i>condizione di BARKHAUSEN applicata al modulo</i>	$\arg \overline{A} \cdot \overline{B} = 0^\circ$ <i>condizione di BARKHAUSEN applicata alla fase</i>
---	---

Infatti una componente del rumore termico¹, di frequenza tale da rendere nullo lo sfasamento della maglia, si trova sempre presente all'ingresso dell'amplificatore.

Se il modulo dell'amplificazione di anello $\overline{A} \cdot \overline{B}$ è maggiore di 1 e la fase $\overline{A} \cdot \overline{B}$ è nulla, il segnale, cioè la componente di rumore, "ritorna" all'ingresso con ampiezza maggiore, viene nuovamente amplificato e quindi la sua ampiezza continua a crescere fino a quando non interviene uno dei seguenti fattori:

- la non linearità del dispositivo attivo (al limite la saturazione o l'interdizione o entrambe)
- elementi di controllo che riducano il modulo di $\overline{A} \cdot \overline{B}$ quando l'oscillazione si è innescata.

¹ Il rumore termico o rumore Johnson è provocato dal movimento casuale, di natura termica, degli elettroni all'interno di qualsiasi conduttore.

STABILITA' IN FREQUENZA DEGLI OSCILLATORI

Un requisito fondamentale dell'oscillatore è la **stabilità in frequenza**, cioè la sua **attitudine a mantenere costante**, sul valore prefissato f_0 , la **frequenza del segnale sinusoidale di uscita**. Il problema dell'instabilità in frequenza dell'oscillazione nasce dal fatto che, per vari motivi, la frequenza f_0 si allontana dal valore f_0 di progetto, per cui il circuito produce oscillazioni a una frequenza f' diversa da f_0 . (Ricordiamo che f_0 è anche la **frequenza in corrispondenza della quale lo sfasamento complessivo dell'anello deve risultare nullo**).

L'alterazione del valore di f_0 è dovuta al fatto che la **curva dello sfasamento dell'anello "amplificatore + rete di retroazione"**, in funzione della frequenza, subisce degli **spostamenti** che determinano lo slittamento della frequenza f_0 di oscillazione.

Gli spostamenti della curva di fase possono avvenire per effetto di fenomeni come variazioni di temperatura e di umidità, invecchiamento dei componenti, instabilità della tensione di alimentazione, influenza del carico sul circuito oscillatore.

Affinché l'oscillatore abbia una buona stabilità, deve accadere che, al variare della curva di fase, la variazione di valore della frequenza f_0 sia quanto più piccola possibile².

Se quindi si verifica una indesiderata **variazione di fase $\Delta\Phi$** , la conseguente **variazione Δf di frequenza**, attorno al valore originario f_0 , deve essere quanto **più piccola possibile**.

Pertanto, se definiamo una **variazione relativa di frequenza $\Delta f/f_0$** , questa variazione deve risultare **tanto più piccola**, rispetto a $\Delta\Phi$, **quanto più l'oscillatore è stabile**.

Data una **variazione di fase $\Delta\Phi$** , per una buona stabilità deve quindi essere:

$\Delta f/f_0 = \text{piccola}$

e, di conseguenza il valore del rapporto:

$$\frac{\Delta\Phi}{\left(\Delta f / f_0\right)}$$

deve risultare grande.

Il rapporto appena definito assume un valore che è tanto più elevato quanto maggiore è la stabilità in frequenza, perciò viene assunto come indice della stabilità del sistema. Questo rapporto è in genere indicato come " S_f " ed è chiamato **"coefficiente di stabilità in frequenza"**:

$$S_f = \frac{\Delta\Phi}{\left(\Delta f / f_0\right)} = \frac{\text{variazione di fase}}{\text{variazione relativa di frequenza}}$$

² Graficamente, se l'oscillatore è stabile (e quindi se la variazione di frequenza è piccola rispetto alla variazione di fase) la curva di sfasamento deve intersecare l'asse orizzontale delle frequenze in maniera molto verticale, così che a grosse variazioni di fase corrispondano piccole variazioni di frequenza.

I costruttori usano esprimere la stabilità anche in maniera diversa e cioè mediante i parametri:

- variazione relativa (espressa in parti per milione) della frequenza di oscillazione in rapporto al valore nominale f_0

$$S_{ppm} = 10^6 \cdot \frac{\Delta f}{f_0}$$

- variazione relativa (espressa in percentuale) della frequenza di oscillazione in rapporto al valore nominale f_0

$$S_{\%} = 100 \cdot \frac{\Delta f}{f_0}$$

Fra gli oscillatori RC, l'oscillatore a ponte di Wien con operazionale risulta più stabile di quello a sfasamento. Riguardo agli oscillatori LC a tre punti, la stabilità è tanto più alta, quanto maggiore è il fattore Q di qualità della rete risonante di retroazione. La stabilità in frequenza può essere notevolmente migliorata mediante l'introduzione di quarzi piezoelettrici nella rete di retroazione, al posto di un induttore, o al posto dell'intero circuito risonante di retroazione.

OSCILLATORI AL QUARZO (O A CRISTALLO)

I cosiddetti quarzi (Xtal) sono cristalli opportunamente tagliati, e inseriti fra due elettrodi metallici, che, grazie alla proprietà di piezoelettricità, sono utilizzati come componenti elettronici, con la funzione di rendere stabile nel tempo la frequenza di oscillatori sinusoidali e di generatori d'onda rettangolare (generatori di clock o multivibratori astabili).



Il cristallo di quarzo, da un punto di vista elettronico, si comporta come un circuito risonante, cioè come un filtro passabanda la cui frequenza centrale è chiamata appunto frequenza di risonanza. Ogni cristallo di quarzo è caratterizzato da una sua frequenza di risonanza (oltre che da altri parametri elettrici). La gamma di frequenze di risonanza dei quarzi è compresa orientativamente fra i 400 Hz e i 125 MHz.

Parametro	Frequenza			
	90 kHz	280 kHz	525 kHz	5 MHz
$R \quad [\Omega]$	15 k	1,35 k	220	150
$L \quad [H]$	137	27,7	7,8	0,785
$C \quad [pF]$	0,0235	0,0117	0,0115	0,00135
$C_0 \quad [pF]$	3,5	6,18	6,3	3,95

Parametri dei quarzi per diversi valori di frequenza

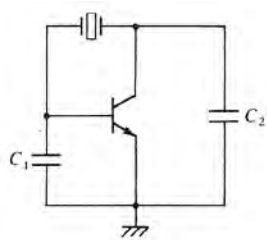
La frequenza di risonanza è tanto più alta quanto più sottile è la piastrina di quarzo. Siccome esiste un limite tecnologico all'assottigliamento della piastrina (che tende a spezzarsi) esiste anche un limite superiore alla frequenza di risonanza dei quarzi. **Tipicamente le frequenze di oscillazione dei quarzi sono minori di 20 MHz.** In quanto componente elettronico, il quarzo è caratterizzato da un'impedenza il cui valore e il cui carattere cambiano con la frequenza.

In corrispondenza della **frequenza di risonanza**, l'**impedenza** del quarzo è **puramente resistiva**. Per frequenze diverse da quella di risonanza, l'impedenza del quarzo è induttiva oppure capacitiva. Quando è inserito in un oscillatore, il quarzo viene fatto lavorare in bande di frequenza nelle quali il suo comportamento è induttivo (fermo restando che in risonanza l'impedenza del quarzo è puramente resistiva) oppure viene fatto lavorare, senza altri componenti reattivi, come "intero circuito risonante".

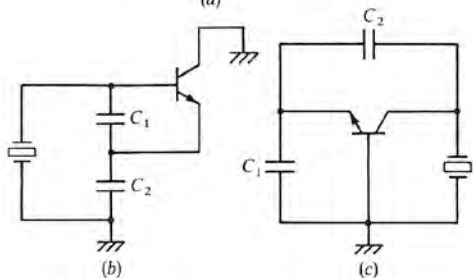
COME I QUARZI SONO INSERITI NEGLI OSCILLATORI

Ai fini della stabilità in frequenza i **quarzi** sono tipicamente inseriti nella rete di retroazione positiva di oscillatori LC "a tre punti", dove:

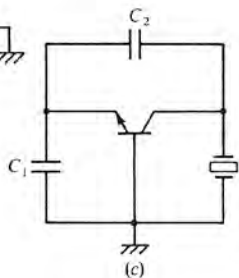
- o **sostituiscono un induttore** (oscillatore di Pierce/Colpitts)
- o sostituiscono una delle due induttanze, mentre al posto dell'altra vi è un intero circuito risonante (oscillatore di Miller/Hartley)
- o sostituiscono **l'intera rete di retroazione** (oscillatore di Meacham).



(a)



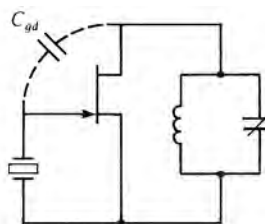
(b)



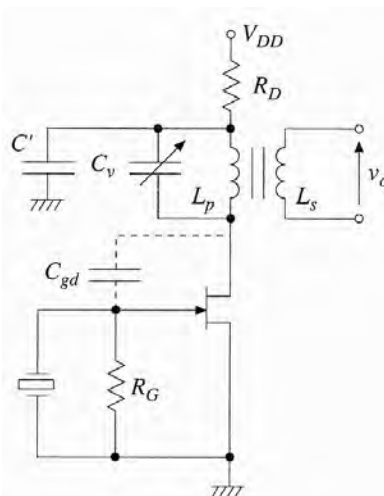
(c)

**Circuiti dinamici di oscillatori
Pierce/Colpitts:**

- a) con emettitore a massa
- b) con collettore a massa
- c) con base a massa



**Circuito dinamico di oscillatore
di Miller/Hartley**



**Oscillatore di Miller/Hartley:
circuitto completo**

FREQUENZA DI OSCILLAZIONE DELL'OSCILLATORE E FREQUENZA DI RISONANZA DEL QUARZO

Quando in un oscillatore viene inserito un quarzo, l'oscillatore deve produrre in **uscita** un segnale sinusoidale di **frequenza uguale** alla **frequenza di risonanza del quarzo**. Se il quarzo sostituisce l'intera rete LC di retroazione positiva dell'oscillatore, è evidente che la frequenza di oscillazione del circuito debba essere quella del quarzo. Se il quarzo è inserito al posto di un singolo induttore della rete di retroazione positiva dell'oscillatore, bisogna dimensionare gli altri componenti in modo che la frequenza di oscillazione del circuito coincida con quella di risonanza del quarzo. La scelta dei valori dei componenti si fa in base a relazioni derivanti dall'applicazione della condizione di Barkhausen all'oscillatore in questione.

Per esempio, per l'oscillatore di Colpitts/Pierce a BJT, la condizione di oscillazione è:

$$g_m = R_{XTAL} \cdot C_1 \cdot C_2 \cdot 4 \cdot \pi^2 \cdot f_R^2$$

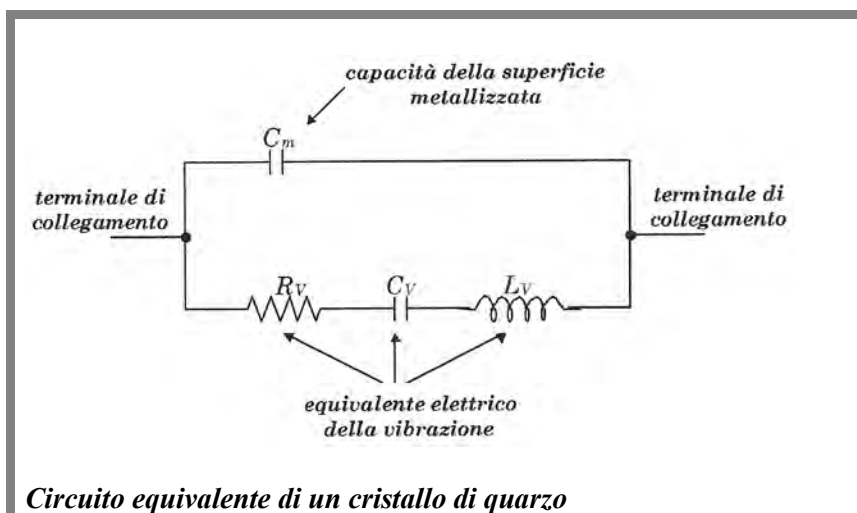
dove:

g_m = transconduttanza del BJT

C_1, C_2 = condensatori della rete di retroazione positiva

R_{XTAL} = resistenza di perdita del QUARZO

f_R = frequenza di risonanza

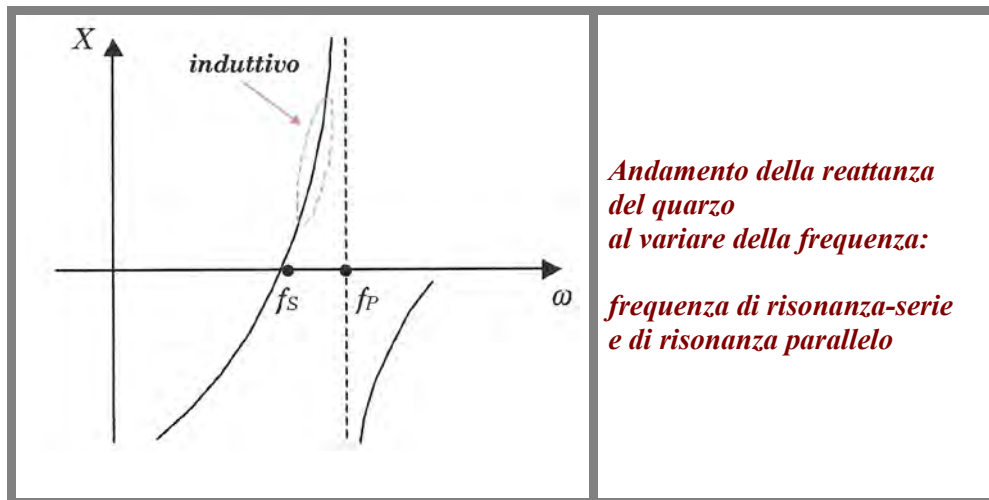


RISONANZA SERIE E RISONANZA PARALLELO

A rigore ciascun cristallo di quarzo ha due frequenze di risonanza: la frequenza f_s di risonanza-serie (più bassa) e la frequenza f_p di risonanza parallelo (leggermente più alta), ma essendo queste frequenze molto vicine fra loro, si può parlare di un'unica frequenza di risonanza.

L'esistenza di due frequenze di risonanza può essere spiegata analizzando il circuito equivalente del quarzo, nel quale sono presenti capacità e induttanze sia in serie che in parallelo.

Per frequenze inferiori alla f_s e per valori compresi fra f_s e f_p il quarzo si comporta da induttore. Per frequenze superiori a f_p si comporta da condensatore. Tipicamente viene fatto lavorare nella banda compresa fra f_s e f_p dove si comporta da induttore.



QUARZI OVERTONE

Sono in commercio quarzi, chiamati "overtone" che sono in grado di oscillare a frequenze multiple della frequenza fondamentale di risonanza.

Quando si inserisce un quarzo overtone in un oscillatore che si vuol far lavorare su una frequenza $f^* = Nf_R$ superiore alla fondamentale, bisogna impedire che l'oscillatore possa entrare in oscillazione per frequenze inferiori a f^* .

Ciò può essere realizzato dimensionando i componenti in modo che la rete di retroazione abbia, per $f < f^*$, comportamento induttivo e non capacitivo, così che venga a mancare, per $f < f^*$, la capacità la cui presenza è indispensabile (assieme all'induttanza realizzata dal quarzo) per il funzionamento di un oscillatore RC.

CAPITOLO 27

DERIVAZIONE E INTEGRAZIONE IN ELETTRONICA

OBIETTIVI

- Conoscere i campi di utilizzazione, soprattutto nell'ambito dei controlli automatici, dei circuiti derivatore e integratore
- Comprendere il funzionamento dei circuiti derivatore e integratore basati su amplificatore operazionale
- Saper analizzare il comportamento del derivatore e dell'integratore, nel tempo e in frequenza
- Saper dimensionare i circuiti derivatore e integratore

PREREQUISITI

- Significato delle operazioni di derivazione e integrazione
- Generalità sull'amplificatore operazionale
- Nozioni sulla risposta in frequenza
- Limitatamente ai paragrafi scritti in colore diverso: metodo della trasformata di Laplace (cap. 40)

L'OPERAZIONE DI DERIVAZIONE

Dato un segnale $v(t)$, la derivata temporale di $v(t)$ in un certo istante t^* rappresenta la velocità di variazione del segnale nell'istante t^* e si indica come:

$$\frac{d}{dt}v(t^*)$$

Esistono altri modi per esprimere la derivata temporale di $v(t)$ nell'istante t^* :

$$v'(t^*) = D[v(t^*)] = \frac{d}{dt}v(t^*) = \text{derivata di } v(t) \text{ in } t = t^*$$

L'andamento nel tempo della funzione derivata

$$\frac{d}{dt}v(t)$$

in un certo intervallo (t_1, t_2) ci mostra l'andamento della velocità di variazione del segnale $v(t)$ nell'intervallo.

Per esempio la velocità di variazione di un segnale costante $v(t)=K$ è zero perché un segnale costante non varia, ma conserva nel tempo sempre lo stesso valore. Quindi:

$$\frac{d}{dt}v(t) = \frac{d}{dt}K = 0$$

Invece la velocità di variazione di un segnale a rampa crescente $v(t)=5t$ è costante e vale 5, che è il valore del coefficiente angolare della retta che descrive l'andamento della rampa nel tempo. Quindi:

$$\frac{d}{dt}v(t) = \frac{d}{dt}[5 \cdot t] = 5 \cdot \frac{d}{dt}[t] = 5 \cdot 1 = 5$$

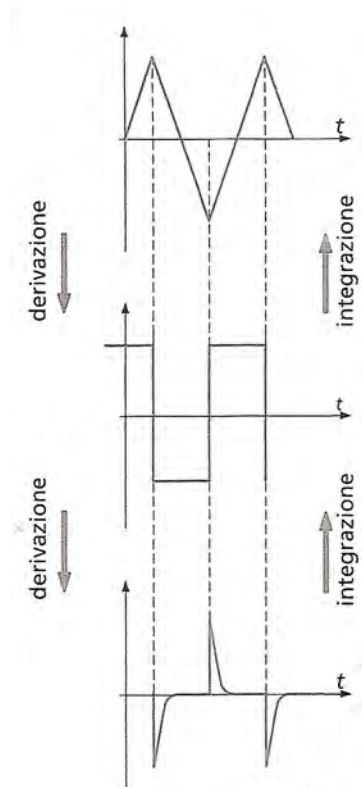
Una rampa crescente a velocità maggiore, per esempio $v(t)=7t$ avrà come derivata ancora una costante, ma di valore più elevato:

$$\frac{d}{dt}v(t) = \frac{d}{dt}[7 \cdot t] = 7 \cdot \frac{d}{dt}[t] = 7 \cdot 1 = 7$$

Invece una rampa decrescente, per esempio $v(t) = -9t$ avrà come derivata ancora una costante, ma di valore negativo:

$$\frac{d}{dt}v(t) = \frac{d}{dt}[-9 \cdot t] = -9 \cdot \frac{d}{dt}[t] = -9 \cdot 1 = -9$$

Funzione	Derivata	Funzione	Derivata
$y = \sqrt{x}$	$y' = \frac{1}{2\sqrt{x}}$	$y = \arccos x$	$y' = -\frac{1}{\sqrt{1-x^2}}$
$y = \frac{1}{x}$	$y' = -\frac{1}{x^2}$	$y = \arctan x$	$y' = \frac{1}{1+x^2}$
$y = \text{sen } x$ oppure $y = \sin x$	$y' = \cos x$	$y = \text{arccotan } x$	$y' = -\frac{1}{1+x^2}$
$y = \cos x$	$y' = -\text{sen } x$	$y = a^x$	$y' = a^x \ln a$
$y = \tan x$	$y' = \frac{1}{\cos^2 x}$	$y = e^x$	$y' = e^x$
$y = \cotan x$	$y' = -\frac{1}{\text{sen}^2 x}$	$y = \log_a x$	$y' = \frac{1}{x} \log_a e$
$y = \ln \text{sen } x $	$y' = \cotan x$	$y = \ln x$	$y' = \frac{1}{x}$
$y = \ln \left \tan \frac{x}{2} \right $	$y' = \frac{1}{\text{sen } x}$	$y = \ln \cos x $	$y' = -\tan x$
$y = \arcsen x$	$y' = \frac{1}{\sqrt{1-x^2}}$	$y = \ln \left \tan \left(\frac{x}{2} + \frac{\pi}{4} \right) \right $	$y' = \frac{1}{\cos x}$
<i>derivate di alcune funzioni matematiche</i>			

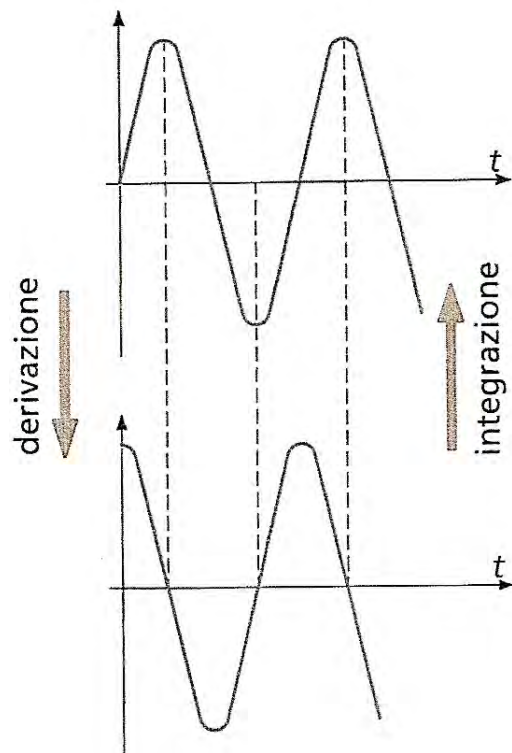


Dall'alto verso il basso:

- *La derivata dell'onda triangolare è l'onda rettangolare bipolare*
- *La derivata dell'onda rettangolare è una sequenza di impulsi: positivi in corrispondenza dei fronti di salita del segnale rettangolare, negativi in corrispondenza dei fronti di discesa*

Dal basso verso l'alto:

- *L'operazione inversa della derivazione è l'integrazione (che tratteremo nei prossimi paragrafi): se integriamo la sequenza di impulsi, riotteniamo l'onda rettangolare, se integriamo l'onda rettangolare, ritroviamo il segnale triangolare*



Seno e coseno

Dall'alto verso il basso:

- *La derivata del seno è il coseno*

Dal basso verso l'alto:

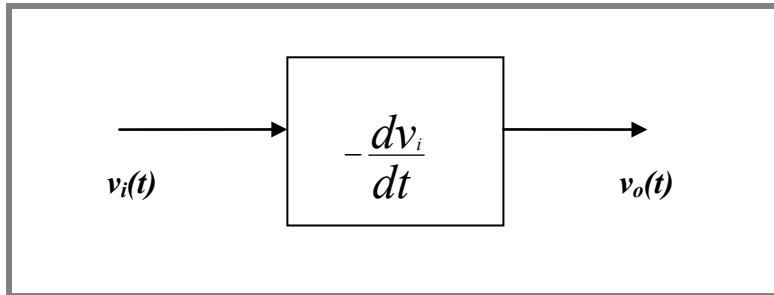
- *L'operazione inversa della derivazione è l'integrazione (che tratteremo nei prossimi paragrafi): se integriamo il coseno ritroviamo la funzione seno, a meno di uno sfasamento di 180°*

CIRCUITO DERIVATORE INVERTENTE CON OPERAZIONALE

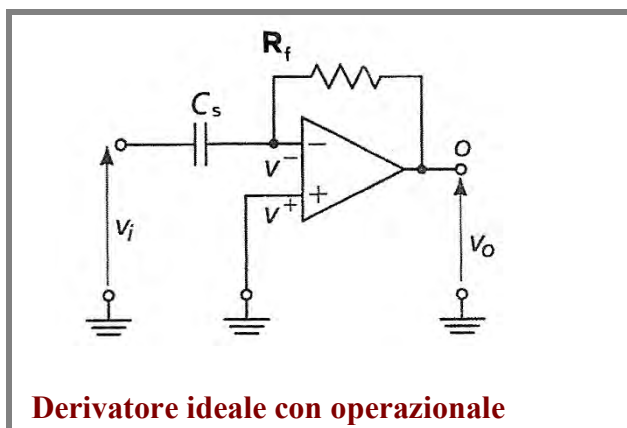
Il derivatore invertente è un dispositivo caratterizzato dal fatto che l'andamento temporale della tensione di uscita è proporzionale istante per istante alla velocità di variazione nel tempo (con il segno cambiato) della tensione di ingresso:

$$v_o = -K \cdot \frac{dv_i}{dt} \quad (\text{con } k = \text{costante positiva})$$

La tensione di uscita cioè è proporzionale alla derivata temporale della tensione di ingresso.



In pratica il derivatore è un circuito in grado di trasformare una tensione a rampa in tensione continua, una tensione triangolare in tensione ad onda rettangolare, una tensione ad onda rettangolare in una sequenza alternata di impulsi, una tensione sinusoidale in tensione cosinusoidale ecc.



A COSA SERVE IL DERIVATORE

Il derivatore ha numerosi campi di applicazione, ne ricordiamo qui alcuni:

- **Controllori PID:** il classico controllore dei sistemi di controllo industriali è di tipo “Proporzionale-Integrativo-Derivativo”, ed è essenzialmente formato dal parallelo di un amplificatore (azione proporzionale) di un integratore e di un derivatore che agiscono sul segnale di errore (di un sistema controllato in catena chiusa) in modo che questo segnale possa essere utilizzato efficacemente per correggere la grandezza controllata. Ora, il derivatore, essendo sensibile alla velocità di variazione del segnale di ingresso, risulta ottimale nel caso di segnali di errore rapidamente variabile nel tempo
- **Circuiti di trigger:** alcuni circuiti, come il multivibratore monostabile (generatore di singolo impulso rettangolare) hanno bisogno, per essere attivati, di un impulso molto stretto da applicare in ingresso. Il derivatore è in grado di fornire questo tipo di impulsi, grazie alla sua proprietà di trasformare in impulsi stretti i fronti di salita e di discesa di onde o impulsi rettangolari.

RELAZIONI INGRESSO-USCITA (I/O) DEL DERIVATORE IDEALE

$$v_o = -RC \cdot \frac{dv_i}{dt}$$

RELAZIONE INGRESSO-USCITA NEL TEMPO

$$\overline{V_o} = -j\omega RC \cdot \overline{V_i} \quad \text{con} \quad \omega = 2\pi f$$

RELAZIONE INGRESSO-USCITA IN FREQUENZA

$$\overline{V_o} = -s \cdot RC \cdot \overline{V_i}$$

RELAZIONE I/O NELLA VARIABILE “s” DI LAPLACE

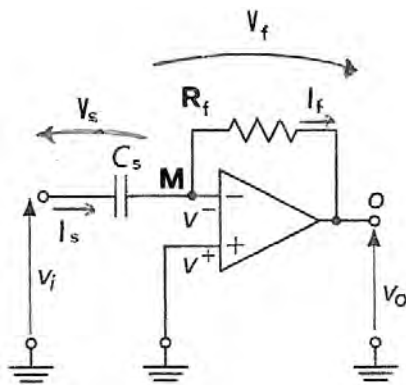
ANALISI NEL TEMPO DEL DERIVATORE IDEALE

Il derivatore è un circuito **lineare**, per cui vale il **corto circuito virtuale**:

i terminali di ingresso dell'operazionale hanno lo stesso potenziale:

$V^+ = V^-$. Inoltre vale il **principio di massa virtuale** perché il terminale di ingresso non invertente è collegato a massa ($V^+ = 0$). Di conseguenza anche il terminale invertente viene a trovarsi a potenziale nullo ($V^+ = V^- = 0$).

Il terminale “-” si trova quindi virtualmente a massa: chiamiamo “M” questo punto.



Di conseguenza la tensione V_o coincide con la caduta di potenziale V_f su R_f (in quanto entrambe ddp fra il terminale di uscita e un punto di massa), per cui possiamo scrivere, applicando la legge di Ohm su R_f :

$$v_o = -i_f \cdot R_f$$

Questa relazione, grazie al fatto che la resistenza di ingresso dell'operazionale è idealmente ∞ (e che quindi si ha: $i_f = i_s = i$) può essere scritta come:

$$v_o = -i \cdot R_f$$

Sempre per il principio di massa virtuale la tensione di ingresso v_i coincide con la tensione v_s ai capi del condensatore per cui possiamo scrivere:

$$i = c_s \cdot \frac{dv_i}{dt}$$

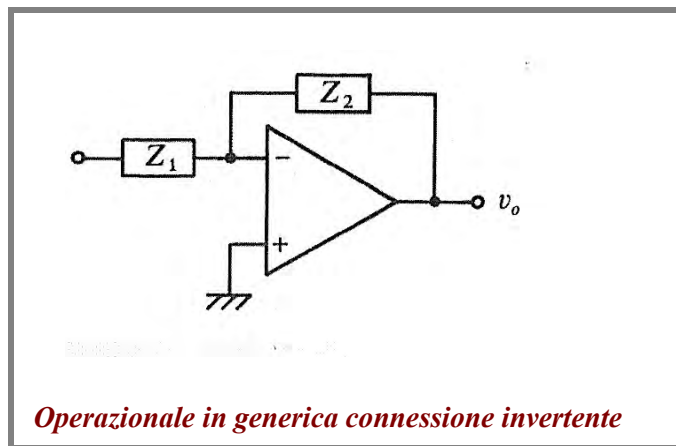
ove la tensione ai capi del condensatore C_s è proprio la v_i di ingresso.

Sostituendo questa relazione nella precedente abbiamo:

$$v_o = -R_f \cdot C_s \cdot \frac{dv_i}{dt}$$

che ci indica proprio che la tensione di uscita del derivatore è proporzionale, attraverso la costante -RC, alla derivata della tensione v_i di ingresso.

ANALISI IN FREQUENZA DEL DERIVATORE IDEALE



La struttura del derivatore presenta analogie con quella dell'amplificatore invertente. Ricordiamo che la funzione di risposta in frequenza, cioè l'amplificazione, dell'amplificatore invertente è:

$$\frac{\overline{V_o}}{\overline{V_i}} = \overline{A} = -\frac{R_2}{R_1}$$

Consideriamo ora il generico schema circuitale invertente della figura. La risposta in frequenza di questo circuito, considerando la sua analogia all'amplificatore invertente, si può scrivere avendo presente l'espressione dell'amplificazione dell'amplificatore invertente. Bisogna però tener conto che al posto di R_2 c'è la generica impedenza Z_2 e al posto di R_1 la generica impedenza Z_1 .

Procedendo in questo modo si ottiene:

$$\frac{\overline{V_o}}{\overline{V_i}} = \overline{A} = -\frac{\overline{Z_2}}{\overline{Z_1}}$$

Ora il derivatore è un caso particolare del circuito in figura con

$$\begin{cases} \bar{Z}_2 = R_f \\ \bar{Z}_1 = \frac{1}{j\omega C_s} = -j \frac{1}{\omega C_s} \end{cases}$$

Di conseguenza avremo:

$$\frac{\bar{V}_o}{\bar{V}_i} = \bar{A} = -\frac{\bar{Z}_2}{\bar{Z}_1} = -\frac{R_f}{\frac{1}{j\omega C_s}} = -j\omega \cdot R_f \cdot C_s$$

cioè:

$\frac{\bar{V}_o}{\bar{V}_i} = \bar{A} = -j\omega \cdot R_f \cdot C_s$	RISPOSTA IN FREQUENZA (ESPRESSIONE COMPLESSA)
--	--

Come si vede, la risposta in frequenza è una quantità complessa, anzi, in questo caso, puramente immaginaria. Il modulo e la fase della risposta saranno quindi:

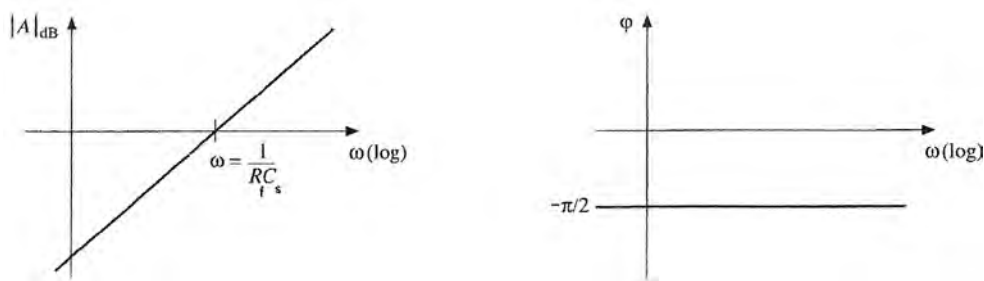
RISPOSTA IN FREQUENZA	
MODULO	$A = \omega \cdot R_f \cdot C_s = 2\pi \cdot f \cdot R_f \cdot C_s$
FASE	$\varphi = \arg(-j\omega R_f C_s) = \operatorname{arctg}\left(\frac{-\omega R_f C_s}{0}\right) = \operatorname{arctg}(-\infty) = -90^\circ$

Il derivatore ideale invertente impone un ritardo di fase costante negativo (sfasamento di -90°) per ogni valore di f a tutti i segnali applicati al suo ingresso.

Nel dominio della frequenza si può quindi scrivere: $\overline{V_o} = -j\omega R_f \cdot C_s \cdot \overline{V_i}$

Mentre nel dominio del tempo si ha: $v_o = -R_f \cdot C_s \cdot \frac{dv_i}{dt}$

Confrontando le due ultime equazioni possiamo notare come l'operatore di derivata temporale “ d/dt ” corrisponda, nel **dominio della frequenza**, a una **moltiplicazione per “ $j\omega$ ”**. Osserviamo infine che le relazioni nel dominio della frequenza sono più semplici (puramente algebriche) ma valgono solo in regime sinusoidale, mentre le relazioni del dominio del tempo sono matematicamente più complicate (contenendo derivate o integrali), ma sono valide per qualunque tipo di segnale.



Derivatore ideale: risposta in frequenza

- a sinistra: andamento **linearizzato** del **modulo** della risposta in frequenza in funzione della **pulsazione in scala logaritmica** (diagramma di **Bode**)
- A destra: andamento della **fase** (o **argomento**) della risposta in frequenza in funzione della **pulsazione in scala logaritmica** (diagramma di Bode)

COME SI OTTIENE LA RAPPRESENTAZIONE GRAFICA DELLE CURVE DI RISPOSTA IN FREQUENZA: I DIAGRAMMI DI BODE

PREMESSE

Abbiamo ricavato l'espressione del modulo della risposta in frequenza del derivatore che è:

$$A = \omega \cdot R_f \cdot C_S$$

e che è analoga all'equazione

$$y = m \cdot x$$

di una retta passante per l'origine del piano (x,y).

Riguardo all'equazione:

$$A = \omega \cdot R_f \cdot C_S$$

Puntualizziamo che:

- **A** è il modulo della risposta in frequenza, è una funzione della frequenza (e quindi della pulsazione $\omega=2\pi f$) e perciò potrebbe essere indicato come “A(f)” oppure come “A(ω)”
- **A** corrisponde alla **y** dell'equazione della retta ed è quindi la variabile **dipendente**, riportata sull'asse **verticale**
- **ω** corrisponde alla **x** dell'equazione della retta ed è quindi la variabile **indipendente**, riportata sull'asse **orizzontale**
- **$R_f C_S$** è una **costante** (che spesso è chiamata **τ**) e corrisponde al coefficiente angolare della retta.

LA RAPPRESENTAZIONE CHE UTILIZZIAMO: I DIAGRAMMI DI BODE

Pertanto se rappresentassimo, senza ulteriori accorgimenti, l'andamento del modulo A su un piano (ω, A) , questo andamento non sarebbe che una retta passante per l'origine, con coefficiente angolare m .

Normalmente però non si usa rappresentare le risposte in frequenza nel modo ora descritto, ma attraverso i cosiddetti diagrammi di Bode, nei quali:

- le **pulsazioni** vengono riportate sull'**asse orizzontale in scala logaritmica**, quindi la **variabile indipendente** non è ω , bensì **$\log \omega$**
- **A viene espresso in dB**, cioè come:
$$(A)_{dB} = 20 \cdot \log_{10} A$$
- **La costante RC viene anch'essa espressa in dB**, cioè come:
$$20 \cdot \log_{10} (R_f \cdot C_s) = \text{costante corrispondente a } q$$

Pertanto l'**equazione**

$$A = \omega \cdot R_f \cdot C_s$$

del **modulo della risposta in frequenza** diventa:

$$(A)_{dB} = 20 \cdot \log_{10} (\omega \cdot R_f \cdot C_s)$$

e per per la proprietà dei logaritmi

$$\log_{10} (A \cdot B) = \log_{10} (A) + \log_{10} (B)$$

si può scrivere come:

$$(A)_{dB} = 20 \cdot \log_{10} (\omega) + 20 \cdot \log_{10} (R_f \cdot C_s) \quad (*)$$

Questa equazione è quindi del tipo:

$$y = mx + q$$

e, in essa, si ha:

- $(A)_{dB}$ = variabile dipendente corrispondente alla y
- $20 \cdot \log_{10}(R_f \cdot C_s) = \text{costante} = \text{intersezione con l'asse verticale}$
- $\omega_t = \frac{1}{R_f C_s} = \text{intersezione con l'asse orizzontale (asse } \omega)$

L'intersezione con l'asse orizzontale si determina imponendo $(A)_{dB}=0$ nella (*), nel modo seguente:

$$0 = 20 \cdot \log_{10}(\omega) + 20 \cdot \log_{10}(R_f \cdot C_s)$$

$$20 \cdot \log_{10}(\omega) = -20 \cdot \log_{10}(R_f \cdot C_s)$$

$$\log_{10}(\omega) = -\log_{10}(R_f \cdot C_s)$$

e, per la proprietà dei logaritmi: $-\log(x) = \log(1/x)$:

$$\log_{10}(\omega) = +\log_{10}\left(\frac{1}{R_f \cdot C_s}\right)$$

$$\omega = \frac{1}{R_f \cdot C_s}$$

DIAGRAMMA DELLA FASE (O ARGOMENTO) DELLA RISPOSTA IN FREQUENZA

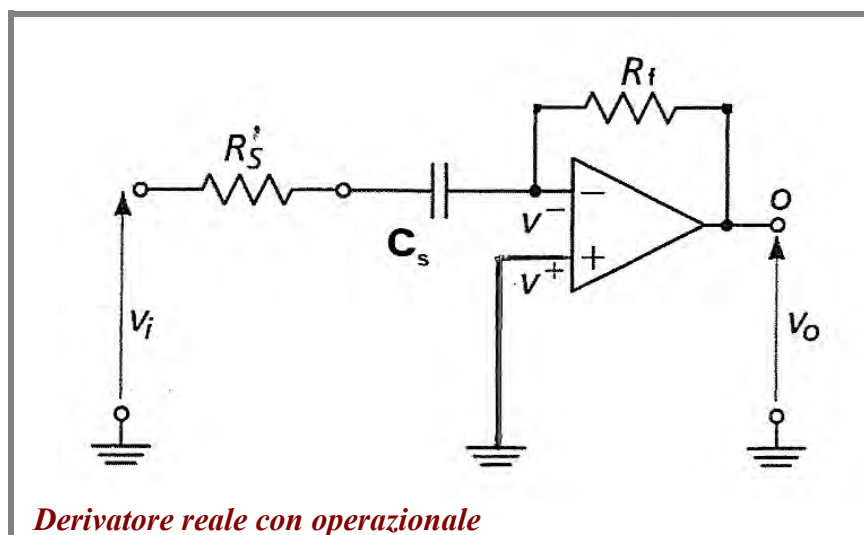
Nei diagrammi di Bode della fase della risposta in frequenza:

- Sull'asse orizzontale sono riportate le pulsazioni (o le frequenze) in scala logaritmica
- Sull'asse verticale è riportata la fase φ in gradi o in radianti

Per il derivatore reale invertente l'andamento della fase risulta costante e

quindi è comunque una retta orizzontale posta alla quota -90° .

DERIVATORE REALE CON OPERAZIONALE



$\overline{A} = \frac{-\frac{R_f}{R'_S}}{1 - j \frac{1}{\omega \cdot C \cdot R'_S}} \quad oppure \quad \overline{A(\omega)} = \frac{j\omega \cdot R_f \cdot C_s}{1 + j\omega R'_S \cdot C_s}$	
RISPOSTA IN FREQUENZA	
$ \overline{A} = \frac{\frac{R_f}{R'_S}}{\sqrt{1 + \frac{1}{(\omega \cdot C_s \cdot R'_S)^2}}} \quad oppure \quad \left \overline{A(\omega)} \right = \frac{\omega \cdot R_f \cdot C_s}{\sqrt{1 + (\omega \cdot R'_S \cdot C_s)^2}}$	
MODULO DELLA RISPOSTA IN FREQUENZA	
$\varphi = -180^\circ + \arctg \left[\frac{+1}{\omega \cdot R'_S \cdot C_s} \right]$ <i>oppure¹</i> $\varphi(\omega) = \arg \left(\overline{A(\omega)} \right) = -90^\circ - \arctg(\omega \cdot R'_S \cdot C_s)$	
FASE DELLA RISPOSTA IN FREQUENZA	

Il derivatore ideale, che è un circuito di tipo passaalto, e che quindi alle basse frequenze taglia il segnale, presenta alle alte frequenze un'amplificazione tendente all' ∞ per cui risulta troppo sensibile al rumore in alta frequenza. Il fatto che l'amplificazione per $\omega \rightarrow \infty$ tenda all'infinito lo si vede dalla espressione:

¹ Basandosi sulla relazione: $\arctg\left(\frac{1}{x}\right) = 90^\circ - \arctg(x)$

$$\bar{A} = -\frac{\bar{Z}_2}{\bar{Z}_1}$$

dell'amplificazione stessa dove il denominatore $\bar{Z}_1$ è una impedenza capacitiva, che al crescere della frequenza tende, in modulo, a zero. Sappiamo infatti che al crescere della frequenza il condensatore tende a comportarsi come un cortocircuito. Ad alta frequenza quindi l'impedenza a denominatore dell'amplificazione tende a zero e l'amplificazione tende ad ∞ .

Il fatto che per $\omega \rightarrow \infty$ l'amplificazione tende ad ∞ , lo si può vedere anche dalla espressione:

$$\bar{A} = -j\omega \cdot R_f \cdot C_s.$$

È evidente che al tendere di ω a ∞ , cresce indefinitamente anche il modulo A della risposta o amplificazione. Per ovviare a questo inconveniente che compromette il funzionamento del derivatore ideale, si aggiunge, **in serie al condensatore C_s** , una **resistenza R'_s** in modo che ad alta frequenza il modulo Z_1 dell'impedenza a denominatore della

$$\bar{A} = -\frac{\bar{Z}_2}{\bar{Z}_1}$$

non tenda più a zero.

Ora infatti l'impedenza al denominatore non è più

$$\frac{1}{j\omega C_s} \quad \text{bensì} \quad R'_s + \frac{1}{j\omega C_s}$$

È bene notare che la R'_s va messa in serie e non in parallelo a C_s perché se fosse messa in parallelo sarebbe cortocircuitata, alle alte frequenze, dall'impedenza capacitiva tendente a zero.

Il derivatore, con l'aggiunta della resistenza R'_s , prende il nome di “derivatore reale”.

COME SI RICAVANO IL MODULO E LA FASE DELLA RISPOSTA IN FREQUENZA DEL DERIVATORE REALE

$$\bar{A} = -\frac{\bar{Z}_2}{\bar{Z}_1}$$

$$\bar{A} = -\frac{R_f}{R'_s - j\frac{1}{\omega \cdot C_s}}$$

Dividendo numeratore e denominatore per R'_s si ha:

$\bar{A} = \frac{-\frac{R_f}{R'_s}}{1 - j\frac{1}{\omega \cdot C_s \cdot R'_s}}$	RISPOSTA IN FREQUENZA DEL DERIVATORE REALE (*)
--	---

In base alla definizione di modulo di un numero complesso, si ha:

$ \bar{A} = \frac{\frac{R_f}{R'_s}}{\sqrt{1 + \frac{1}{(\omega \cdot C_s \cdot R'_s)^2}}}$	MODULO DELLA RISPOSTA IN FREQUENZA DEL DERIVATORE REALE
---	--

Determiniamo ora, a partire dalla (*), la fase (o argomento) della risposta in frequenza del derivatore reale

$$\varphi = \arg \left[-\frac{R_f}{R_s'} \right] - \arg \left[1 - j \frac{1}{\omega \cdot R_s' \cdot C_s} \right]$$

Riguardo al secondo termine del secondo membro della precedente equazione, notiamo che è una quantità complessa del tipo “a+jb” (dove “a” è 1 e “b” è la frazione) e ricordiamo che la fase o argomento di un numero complesso a+jb, si ricava come: $\arg(a+jb) = \arctg(b/a)$

$$\varphi = \arg \left[-\frac{R_f}{R_s'} \right] - \arctg \left[\frac{-1}{\omega \cdot R_s' \cdot C_s} \right]$$

$$\varphi = \pm 180^\circ + \arctg \left[\frac{+1}{\omega \cdot R_s' \cdot C_s} \right]$$

$\varphi = -180^\circ + \arctg \left[\frac{+1}{\omega \cdot R_s' \cdot C_s} \right]$	FASE DELLA RISPOSTA IN FREQUENZA DEL DERIVATORE REALE (la curva assume valori compresi fra -90° e -180°) $\omega = 0 \Rightarrow \varphi = -90^\circ$ $\omega = \infty \Rightarrow \varphi = -180^\circ$
---	--

COME SI OTTENGONO I DIAGRAMMI DI BODE ASINTOTICI DEL MODULO DELLA RISPOSTA IN FREQUENZA

DETERMINAZIONE DEI POLI E DEGLI ZERI

Per determinare i poli e gli zeri ci serve la funzione di trasferimento in “s” del derivatore reale, che noi possiamo ottenere sostituendo “s” al posto di “jω” nell’espressione della risposta in frequenza che abbiamo già ricavato.

Partiamo quindi dalla risposta in frequenza:

$$\bar{A} = \frac{-\frac{R_f}{R'_S}}{1 - j \frac{1}{\omega \cdot C_S \cdot R'_S}} \quad \rightarrow \quad \bar{A} = \frac{-\frac{R_f}{R'_S}}{1 + \frac{1}{j \cdot \omega \cdot C_S \cdot R'_S}}$$

$$\bar{A}(s) = \frac{-\frac{R_f}{R'_S}}{1 + \frac{1}{s \cdot C_S \cdot R'_S}} = \frac{-\frac{R_f}{R'_S}}{\frac{1 + s \cdot C_S \cdot R'_S}{s \cdot C_S \cdot R'_S}} = -\frac{R_f}{R'_S} \cdot \frac{s \cdot C_S \cdot R'_S}{1 + s \cdot C_S \cdot R'_S} = -\frac{s \cdot C_S \cdot R_f}{1 + s \cdot C_S \cdot R'_S}$$

In definitiva la f. di t. in “s” è

$\bar{A}(s) = -\frac{s \cdot C_S \cdot R_f}{1 + s \cdot C_S \cdot R'_S}$	Funzione di trasferimento in “s” del derivatore reale
--	--

Da questa f. di t. in “s” individuiamo:

- **Zero nell’origine** $s = z = 0$
(valore di s che annulla in numeratore della f. di t. in s)
- **Polo** $s = p = -\frac{1}{C_S \cdot R_S}$
(valore di s che annulla in denominatore della f. di t. in s)
- **Pulsazione di rottura corrispondente al polo:**
 $\omega_i = |p| = +\frac{1}{C_S \cdot R_S}$

Il **numeratore** della funzione di trasferimento è il prodotto di un fattore “s” (che dà luogo a uno zero nell’origine: s=z=0) moltiplicato per la costante $R_f C_f$. Il solo zero nell’origine s=z=0 corrisponderebbe a una retta $(A)_{dB} = 20 \log \omega$, che interseca l’asse orizzontale (0dB) in $\omega = 1$ e sale con pendenza di +20 dB/decade. Se invece consideriamo, nel suo complesso, il **numeratore della A(s)** e cioè termine:

$$s \cdot C_S \cdot R_f \quad (\text{cui corrisponde, nel dominio della frequenza: } j\omega \cdot C_S \cdot R_f)$$

e se lo esprimiamo in dB, esso diventa:

$$(A)_{dB} = 20 \cdot \log_{10} (\omega \cdot C_S \cdot R_f)$$

e, per la proprietà $\log(x \cdot y) = \log(x) + \log(y)$, otteniamo l’equazione di una retta del piano ($\log \omega$, $(A)_{dB}$):

$$(A)_{dB} = 20 \cdot \log_{10} (\omega) + 20 \cdot \log_{10} (C_S \cdot R_f)$$

Questa retta vale 0 dB, e quindi **interseca l’asse orizzontale**, quando l’argomento del logaritmo vale 1 e cioè **quando si ha:**

$$\omega = + \frac{1}{C_S \cdot R_f}$$

(si veda, in proposito, la nota alla fine del paragrafo) inoltre ha una pendenza, in salita, di +20 dB/decade.

Valutiamo ora la retta in

$$\omega_t = + \frac{1}{C_S \cdot R_S}$$

Otteniamo:

$$[A(\omega_t)]_{dB} = 20 \cdot \log_{10} \left(\frac{1}{C_S \cdot R_S} \cdot C_S \cdot R_f \right) \rightarrow [A(\omega_t)]_{dB} = 20 \cdot \log_{10} \left(\frac{R_f}{R_S} \right)$$

e, in conclusione:

$[A(\omega_t)]_{dB} = \left(\frac{R_f}{R_S} \right)_{dB}$	<i>(valore in $\omega = \omega_t$ della retta rappresentativa del numeratore della risposta in frequenza)</i>
--	---

Riguardo al **denominatore** della A(s), esso è caratterizzato da un **polo semplice** in:

$$s = p = - \frac{1}{C_S \cdot R_S} \rightarrow$$

$\omega_t = + \frac{1}{C_S \cdot R_S}$
--

E quindi, graficamente è una **spezzata che, orizzontale fino a**

$$\omega = + \frac{1}{C_S \cdot R_S}$$

comincia a decrescere, in corrispondenza di tale pulsazione, con pendenza di -20dB/dec

Sommando i contributi del numeratore e del denominatore otteniamo il modulo della risposta in frequenza.

NOTA

L'intersezione con l'asse orizzontale si determina imponendo $(A)_{dB}=0$ nella (*), nel modo seguente:

$$(A)_{dB} = 20 \cdot \log_{10}(\omega \cdot C_S \cdot R_f)$$

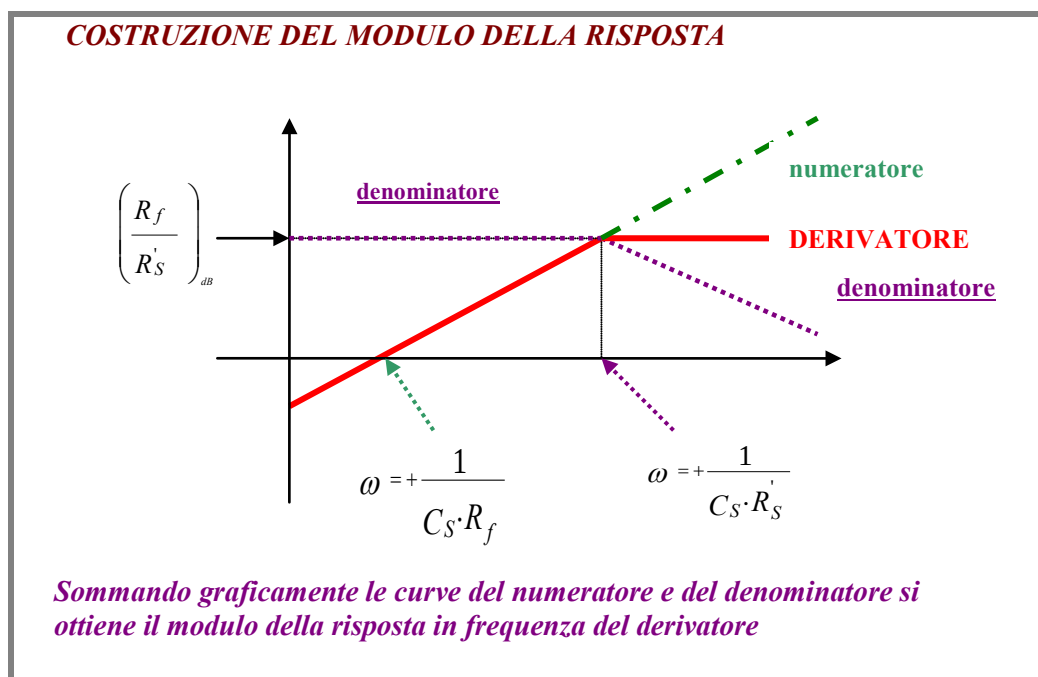
$$0 = 20 \cdot \log_{10}(\omega) + 20 \cdot \log_{10}(R_f \cdot C_S)$$

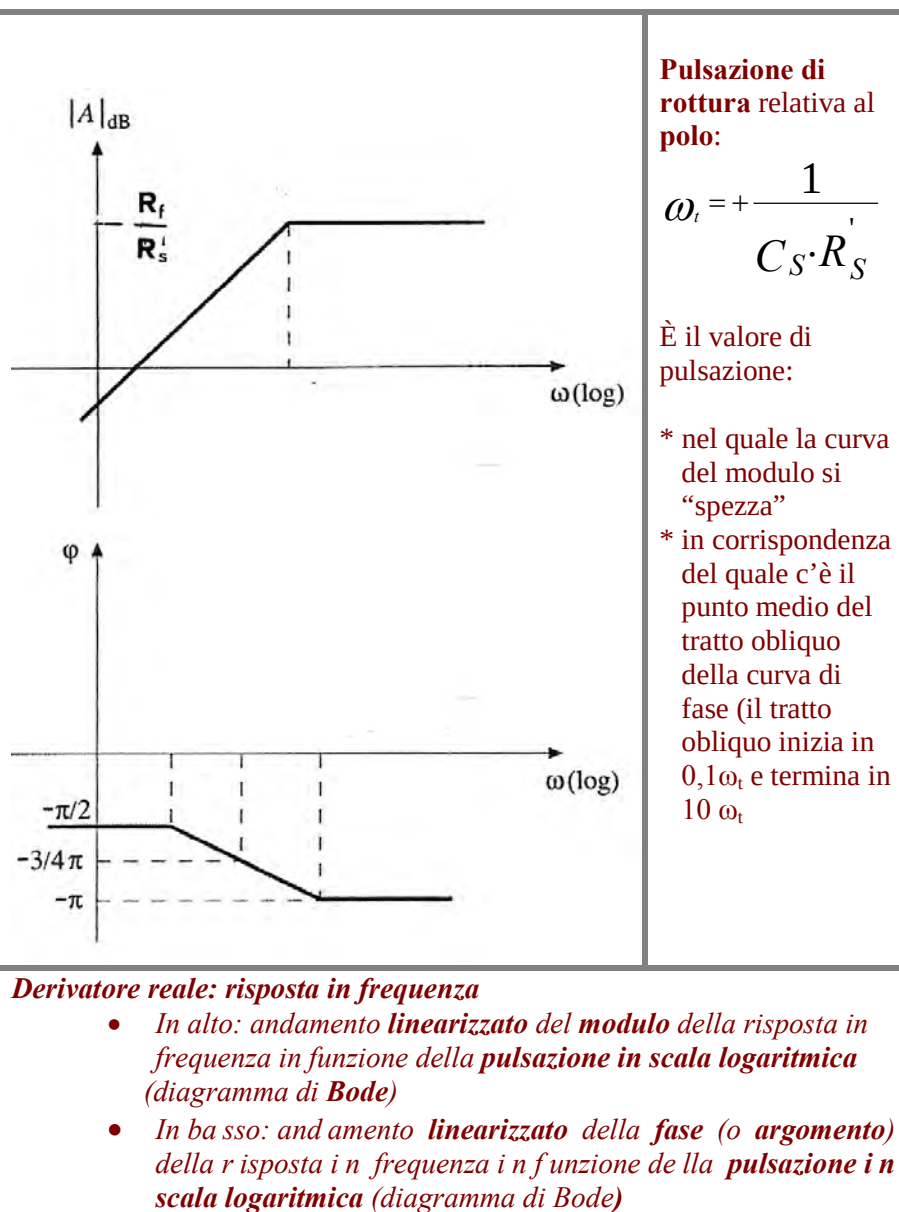
$$20 \cdot \log_{10}(\omega) = -20 \cdot \log_{10}(R_f \cdot C_S)$$

$$\log_{10}(\omega) = -\log_{10}(R_f \cdot C_S)$$

e, per la proprietà dei logaritmi: $-\log(x) = \log(1/x)$

$$\log_{10}(\omega) = +\log_{10}\left(\frac{1}{R_f \cdot C_S}\right) \quad \text{la cui soluzione è:} \quad \omega = \frac{1}{R_f \cdot C_S}$$





DIMENSIONAMENTO DEL DERIVATORE

Come si vede dai diagrammi di Bode asintotici, il derivatore reale **si comporta** effettivamente **da derivatore (tratto inclinato)** e non da amplificatore invertente (tratto orizzontale) **quando la pulsazione ω_s del segnale** applicato

in ingresso **è minore della pulsazione di rottura o di angolo** $\omega = \frac{1}{R_f C_s}$

relativa al polo (dove R'_f è la resistenza di limitazione dell'amplificazione posta in serie a C).

Quando si dimensiona un derivatore, bisogna quindi sempre imporre:

$$\omega_i \ll \frac{1}{R_f C_s}$$

$\uparrow$
Pulsazione del segnale

Qualora poi si voglia che il derivatore **si comporti** come il dispositivo **ideale** (che sfasa in ritardo il segnale di ingresso quasi esattamente di 90°), dovrà essere:

$$\omega_i = \frac{1}{100} \frac{1}{R_f C_f}$$

$\uparrow$
Pulsazione del segnale

Riguardo alla relazione fra la resistenza di limitazione R_s e la resistenza di retroazione R_f , se non vi sono particolari specifiche che ne impongono il valore, si può porre:

$$R_f = \frac{R_s}{10} \quad \leftrightarrow \quad R_s = R_f \cdot 10$$

Spesso però il valore della resistenza R_f di retroazione è vincolato da condizioni imposte sull'ampiezza del segnale di uscita.

INTEGRATORE

L'integratore è un dispositivo che effettua l'operazione inversa della derivazione, operazione detta integrazione. Perciò quando all'ingresso dell'integratore è applicata una tensione costante, all'uscita si avrà una tensione a rampa (la derivata della rampa è, infatti, una costante).

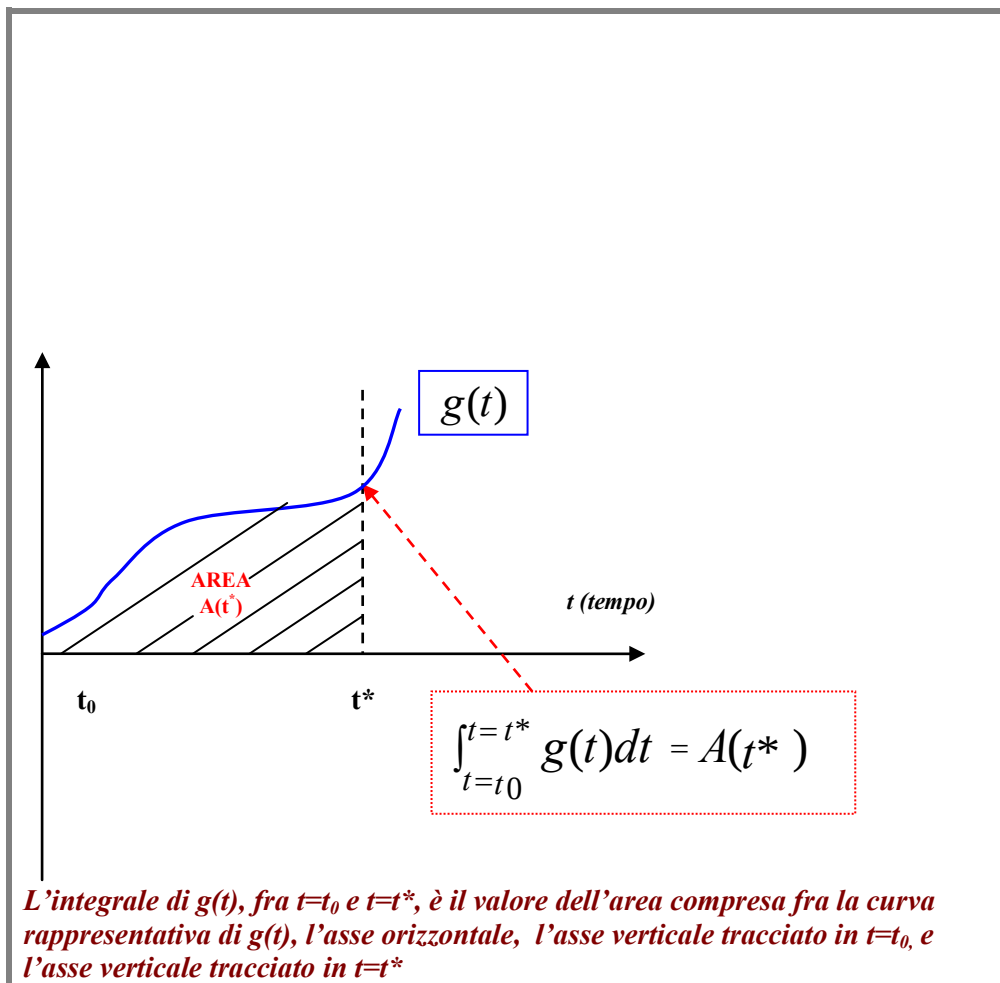
Se all'ingresso dell'integratore è applicata un'onda rettangolare, la tensione di uscita sarà triangolare.

Se all'ingresso dell'integratore è applicata una tensione di equazione $\sin(t)$, in uscita avremo una tensione di tipo $-\cos(t)$.

Per sapere quale sarà l'uscita dell'integratore in corrispondenza di un certo segnale di ingresso bisogna chiedersi qual è il segnale che, derivato, ci dà proprio la forma d'onda di ingresso.

Così, per esempio, se l'ingresso dell'integratore è una tensione nulla, l'uscita è una tensione costante in quanto la derivata di una costante è zero.

Data una funzione $g(t)$, la sua² **funzione integrale** $\int_0^t g(t)dt$ rappresenta in ogni istante il **valore dell'area compresa** fra la **curva rappresentativa della $g(t)$** e **l'asse orizzontale** nell'intervallo di tempo che va **dall'istante zero fino all'istante "t"**.



² La **funzione integrale $f(x)$** della funzione $\varphi(x)$ è la funzione che si ottiene integrando $\varphi(x)$ da un estremo inferiore fisso a un limite superiore variabile x :

$$f(x) = \int_a^x \varphi(x)dx$$

Per esempio, l'integrale della funzione costante $g(t)=5$, dall'istante $t=0$ all'istante $t=10$ è l'area di un rettangolo con base 10 e altezza 5. Per cui:

$$\int_{t=0}^{t=t^*} g(t)dt = 10 \cdot 5 = 50$$

Se poi calcoliamo i valori dell'integrale (e quindi delle aree associate) di una funzione $g(t)$ costante (a gradino) negli istanti t_1, t_2, t_3 ecc, e se diagrammiamo questi valori, otteniamo un andamento a rampa.

I punti della rampa sono i valori delle aree $A_1, A_1, A_3, \dots$.

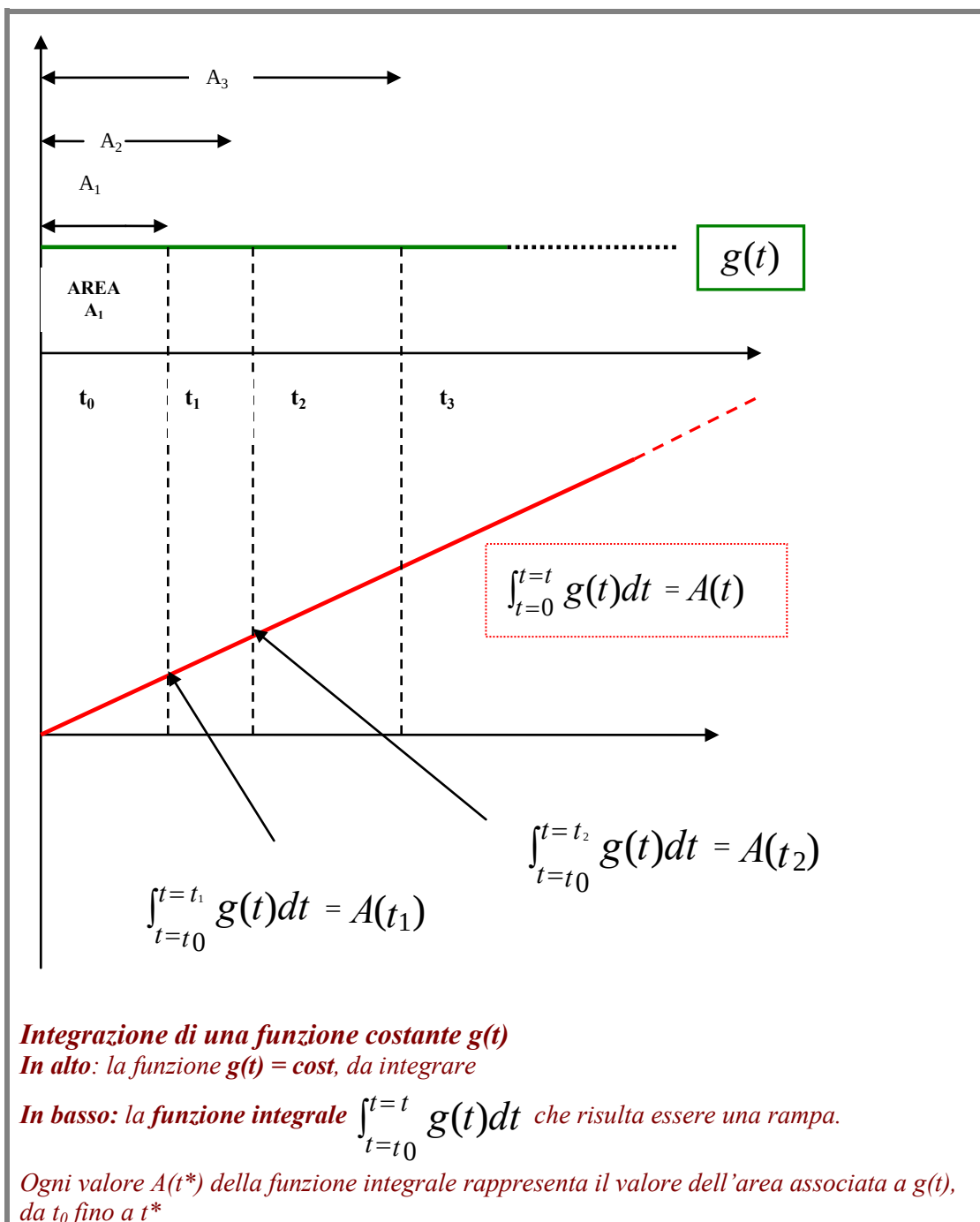
Questo significa che la funzione integrale di una funzione costante è una rampa (si vedano i diagrammi della figura seguente).

In sintesi, data una funzione $g(t)$ in un intervallo $(0, t^*)$, gli integrali, calcolati nei diversi istanti t_1, t_2, t_3 , ecc, sono numeri che rappresentano valori di aree ($A_1, A_1, A_3, \dots$). Se poi tracciamo, con questi valori, una curva, riferita all'asse dei tempi, essa individua una funzione, che è la "funzione integrale" di $g(t)$, cioè la funzione

$$\int_{t=0}^{t=t^*} g(t)dt$$

Abbiamo visto che questa funzione:

- è una costante se la $g(t)$ è identicamente nulla
- è una rampa se la $g(t)$ è costante
- è un'onda triangolare se la $g(t)$ è rettangolare.



INTEGRAZIONE MATEMATICA

Per eseguire matematicamente l'operazione di integrale sulle differenti funzioni, esistono specifiche tecniche di integrazione. Riportiamo una tabella con gli integrali di alcune funzioni elementari:

$\int x^\alpha dx = \frac{x^{\alpha+1}}{\alpha+1} + C \quad (\alpha \neq -1)$	$\int \frac{1}{\sin^2 x} dx = -\cotan(x) + C$
$\int \frac{1}{x} dx = \ln x + C$	$\int \tan x dx = -\ln \cos x + C$
$\int e^x dx = e^x + C$	$\int \cotan x dx = \ln \sin x + C$
$\int a^x dx = \frac{a^x}{\ln a} + C$	$\int \frac{1}{a^2 + x^2} dx = \frac{1}{a} \arctan \frac{x}{a} + C$
$\int \sin x dx = -\cos x + C$	$\int \frac{1}{a^2 - x^2} dx = \frac{1}{2a} \ln \left \frac{a+x}{a-x} \right + C$
$\int \cos x dx = \sin x + C$	$\int \frac{1}{\sqrt{a^2 - x^2}} dx = \arcsin \frac{x}{a} + C$
$\int \frac{1}{\cos^2 x} dx = \tan(x) + C$	$\int \frac{1}{\sqrt{a^2 \pm x^2}} dx = \ln \left x + \sqrt{a^2 \pm x^2} \right + C$
<i>Integrali di alcune funzioni matematiche</i>	

Osserviamo che gli integrali che compaiono in questa tabella non hanno estremi di integrazione, sono cioè “integrali indefiniti” (o “insiemi di primitive”).

In assenza di estremi di integrazione, data una funzione $g(t)$, esistono infinite funzioni che sono integrali di $g(t)$ e che differiscono una dall'altra per una costante additiva “C”.

Data una funzione $g(t)$, esiste quindi un numero infinito di integrali

$$\int g(t) dt$$

e, siccome tutte queste funzioni differiscono solo per una costante, questo fatto si esprime scrivendo:

$$\int g(t)dt = a(t) + C$$

Più esattamente **l'insieme di tutti gli integrali** di una funzione $g(t)$ si chiama **integrale indefinito** di $g(t)$, mentre la **singola funzione $a(t)$** dell'integrale indefinito è detta "**primitiva**" di $g(t)$.

Il fatto che, in assenza di estremi di integrazione esistono infiniti integrali di una funzione $g(t)$, lo si può verificare supponendo, per esempio, che:

$$g(t) = 5$$

Allora certamente la rampa

$$a(t) = 5 \cdot t$$

È un integrale di $g(t)$ perché:

$$\frac{d}{dt}[5 \cdot t] = 5$$

Però, siccome la derivata di una costante è zero, anche tutte le funzioni:

$$b(t) = 5 \cdot t + C \quad (\text{per ogni valore di } C)$$

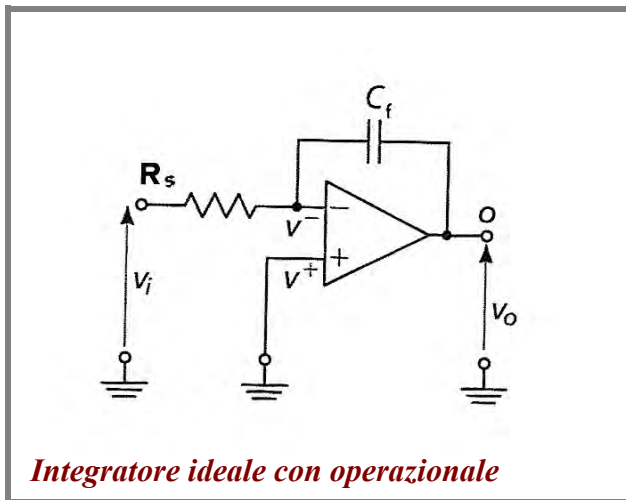
sono integrali di $g(t)$ in quanto la loro derivata vale 5.

Infatti:

$$\frac{d}{dt}[b(t)] = \frac{d}{dt}[5 \cdot t + C] = \frac{d}{dt}[5 \cdot t] + \frac{d}{dt}[C] = 5 + 0 = 5$$

INTEGRATORE IDEALE INVERTENTE CON OPERAZIONALE

L'integratore invertente è un dispositivo la cui tensione di uscita è proporzionale all'integrale (con il segno cambiato) della tensione d'ingresso.



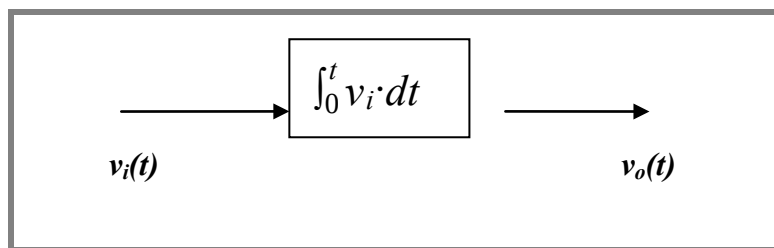
La relazione, nel tempo, fra i segnali di uscita e di ingresso dell'integratore ideale è data da:

$v_o(t^*) = -K \cdot \int_{t=0}^{t=t^*} v_i(t) dt$	RELAZIONE INGRESSO- USCITA (I/O)
--	---

A COSA SERVE L'INTEGRATORE

L'integratore è caratterizzato dal fatto che l'andamento temporale della tensione di uscita è proporzionale istante per istante all'integrale (con il segno invertito) della tensione di ingresso:

$$v_o(t^*) = -K \cdot \int_{t=0}^{t=t^*} v_i(t) dt \quad (\text{con } k = \text{costante positiva})$$



L'integratore ha un vastissimo campo di applicazioni in differenti settori dell'elettronica. Forniamo qui solo alcuni esempi significativi:

- **Convertitori analogico- digitale (ADC)**
- **Modulazione FM:** uno dei metodi di modulazione FM prevede che il segnale informativo modulante venga integrato prima di essere applicato a un modulatore PM (modulatore di fase), il che permette di ottenere la modulazione FM attraverso un modulatore PM (che presenta minori difficoltà realizzative rispetto al modulatore FM)
- **Generatori di segnale:** con un astabile e un integratore si può ottenere un generatore di onde rettangolari e triangolari
- **Controllori PID** o controllori di tipo “Proporzionale-Integrativo-Derivativo”: uno dei blocchi in parallelo che formano un PID è proprio l'integratore, che è particolarmente adatto al trattamento di segnali di errore prolungati nel tempo e di piccola ampiezza
- **Reti correttrici** di dispositivi controllori
- **Filtri attivi universali (UAF) o “a variabili di stato”:** sono basati su sommatore e integratori. Possono avere comportamento sia passabasso, sia passaalto sia passabanda. Tre terminali di uscita permettono di prelevare i segnali che hanno subito i tre diversi tipi di filtraggio
- **Sistemi di navigazione inerziale**

ANALISI DEL CONDENSATORE NEL TEMPO

Prima di passare all'analisi dell'integratore ricordiamo che, nel dominio del tempo, l'espressione della tensione ai capi del condensatore in funzione della corrente è:

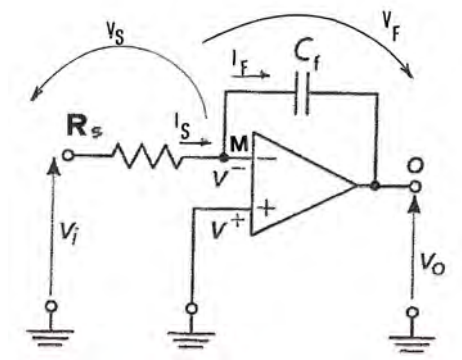
$$v(t^*) = \frac{1}{C} \int_{t=0}^{t^*} i(t) dt + v(o)$$

Dove:

- $v(t^*)$ = valore della tensione sul condensatore nell'istante t^*
- $\int_{t=0}^{t^*} i(t) dt$ = quantità di carica immagazzinata dall'istante $t=0$ fino all'istante $t=t^*$
- $v(0)$ = tensione ai capi del condensatore nell'istante iniziale $t=0$.
La $v(0)$ è nulla se supponiamo il condensatore inizialmente scarico.

INTEGRATORE IDEALE: ANALISI NEL DOMINIO DEL TEMPO

Il circuito è lineare, per cui vale il corto circuito virtuale ($V^+ = V^-$). Inoltre il terminale non invertente è a massa ($V^+ = 0$), e conseguentemente risulterà a potenziale di massa anche il terminale invertente ($V^- = V^+ = 0$).



Di conseguenza la tensione di ingresso v_i coincide con la caduta di tensione v_s su R_s :

$$v_i = R_s \cdot i_s \quad \text{da cui} \quad i_s = \frac{v_i}{R_s}$$

Questa relazione, grazie al fatto che la resistenza di ingresso dell'operazionale è idealmente ∞ (e che quindi si ha: $i_f = i_s = i$) può essere scritta come:

$$i = \frac{v_i}{R_s}$$

Per il principio di massa virtuale, la tensione di uscita v_o coincide con la tensione v_f sul condensatore (in quanto entrambe ddp fra il terminale di uscita e un punto di massa), per cui (utilizzando la relazione temporale corrente $\rightarrow$ tensione relativa al condensatore C_f) possiamo scrivere:

$$v_o(t) = -\frac{1}{C_f} \int_{t=0}^t i(t) \cdot dt + v_o(0)$$

Sostituendo in essa l'ultima espressione di “i” trovata, otteniamo:

$$v_o(t) = -\frac{1}{C_f} \int_{t=0}^t \frac{v_i(t)}{R_s} \cdot dt + v_o(o)$$

e, portando la costante moltiplicativa $1/R_s$ fuori dal segno di integrale:

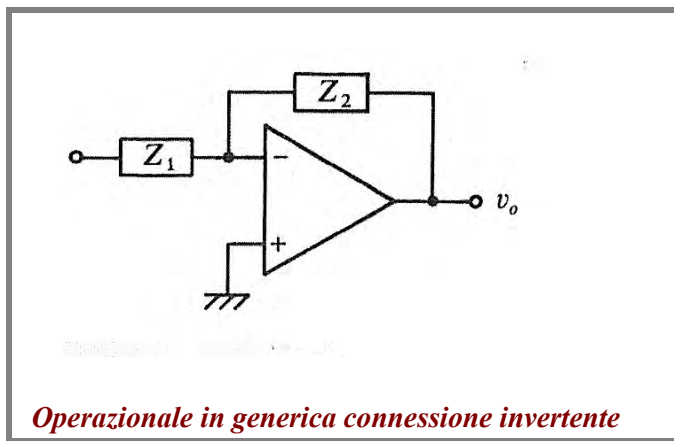
$$v_o(t) = -\frac{1}{R_s \cdot C_f} \cdot \int_{t=0}^t v_i(t) \cdot dt + v_o(o)$$

Che è appunto ciò che volevamo dimostrare.

INTEGRATORE IDEALE: ANALISI NEL DOMINIO DELLA FREQUENZA

Un operazionale in configurazione invertente, retroazionato negativamente mediante due generiche impedenze $\bar{Z}_2$ e $\bar{Z}_1$ ha una risposta in frequenza

esprimibile come: $\bar{A} = -\frac{\bar{Z}_2}{\bar{Z}_1}$.



L'integratore ideale si può realizzare proprio con la struttura ora descritta, con Z_2 impedenza capacitiva e Z_1 impedenza resistiva.

Z_2 è quindi l'impedenza del condensatore C_f e ha espressione:

$$\bar{Z}_2 = -j \frac{1}{\omega \cdot C_f}$$

Mentre Z_1 è la resistenza R_S e quindi ha espressione:

$$\bar{Z}_1 = R_S$$

La risposta in frequenza dell'integratore ideale risulta perciò:

$$\overline{A} = -\frac{-j\frac{1}{\omega \cdot C_f}}{R_s} = +j\frac{1}{\omega \cdot R_s \cdot C_f}.$$

La risposta è puramente immaginaria (nel senso che è una grandezza complessa, del tipo “a+jb” con a=0). Il modulo di una grandezza puramente immaginaria è la grandezza stessa, con segno sempre positivo e privata dell'unità immaginaria “j”.

La fase di una grandezza puramente immaginaria positiva è costante e vale +90°. Questo significa che il circuito integratore ideale sfasa in anticipo di un angolo di 90°, qualunque sia la frequenza, il segnale applicato al suo ingresso. Possiamo sintetizzare il tutto in una tabella:

RISPOSTA IN FREQUENZA DELL'INTEGRATORE IDEALE

ESPRESSIONE COMPLESSA	$\overline{A} = +j\frac{1}{\omega \cdot R_s \cdot C_f}$
MODULO	$ \overline{A} = \frac{1}{\omega \cdot R_s \cdot C_f}$
FASE	$\varphi = +90^\circ$

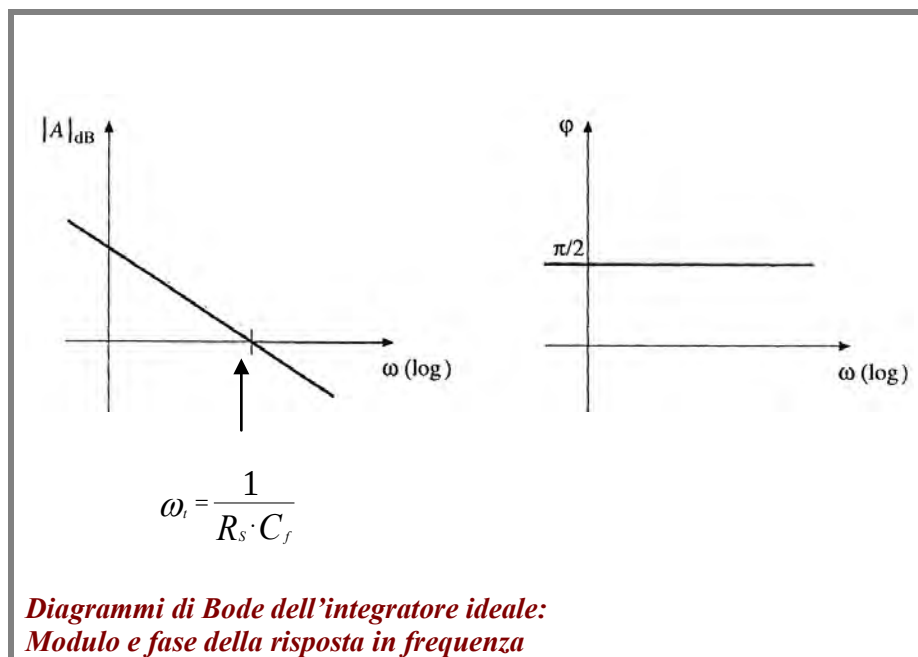
Analizzando l'espressione del **modulo** vediamo che il suo andamento è descritto da una equazione del tipo:

$$y = \frac{k}{x} \quad \rightarrow \quad k = x \cdot y$$

(con $k = \frac{1}{R_s \cdot C_f}$ e con $x = \omega$)

che è l'equazione di **un'iperbole equilatera**³.

Non si usa però rappresentare il modulo della risposta come funzione iperbolica, e si preferisce esprimerlo in dB e riferirlo a un'asse delle pulsazioni in scala logaritmica, così che si presenti come una retta inclinata. Si usa, di fatto, esprimere sia il modulo che la fase della risposta in frequenza come diagrammi di Bode (si veda la figura seguente).



³ Un'iperbole si dice "equilatera" quando ha i semiassi uguali. La lunghezza dei semiassi è $a = \sqrt{2k}$

COME SI RICAVALA IL DIAGRAMMA DI BODE DEL MODULO DELL'INTEGRATORE IDEALE

Dato il modulo della risposta in frequenza dell'integratore:

$$|A| = \frac{1}{\omega \cdot R_S \cdot C_f}$$

esprimiamolo in dB:

$$(A)_{dB} = 20 \cdot \log\left(\frac{1}{\omega \cdot R_S \cdot C_f}\right)$$

Per la proprietà dei logaritmi: $\log(1/x) = -\log(x)$ possiamo scrivere:

$$(A)_{dB} = -20 \cdot \log(\omega \cdot R_S \cdot C_f)$$

La proprietà: $\log(x \cdot y) = \log(x) + \log(y)$ ci permette di modificare l'equazione in questo modo:

$$(A)_{dB} = -20 \cdot \log(\omega) - 20 \cdot \log(R_S \cdot C_f)$$

Il diagramma di Bode si ottiene rappresentando questa equazione in un piano nel quale:

- le **pulsazioni** vengono riportate sull'**asse orizzontale** in **scala logaritmica**, quindi la **variabile indipendente** non è ω , bensì **log ω**
- **A** viene **espresso in dB**, cioè come:

$$(A)_{dB} = 20 \cdot \log_{10} A$$

- La costante " **$R_S C_f$** " viene anch'essa espressa in **dB**, cioè come:

$$20 \cdot \log_{10} (R_S \cdot C_f)$$

e rappresenta la costante q dell'equazione della retta, scritta come $y = mx + q$

Rappresentata in tale piano, l'equazione del modulo dell'integratore ideale, cioè:

$$(A)_{dB} = -20 \cdot \log(\omega) - 20 \cdot \log(R_S \cdot C_f) \quad (*)$$

è una **retta inclinata** caratterizzata da:

- **coefficiente angolare m = -20dB/décade (pendenza negativa)**
- **intersezione con l'asse orizzontale in:** $\omega = \frac{1}{R_S \cdot C_f}$

L'intersezione con l'asse orizzontale è stata determinata imponendo $(A)_{dB}=0$ nella (*), nel modo seguente:

$$0 = 20 \cdot \log_{10}(\omega) + 20 \cdot \log_{10}(R_S \cdot C_f)$$

$$20 \cdot \log_{10}(\omega) = -20 \cdot \log_{10}(R_S \cdot C_f)$$

$$\log_{10}(\omega) = -\log_{10}(R_S \cdot C_f)$$

Per la proprietà dei logaritmi: $-\log(x) = \log(1/x)$ si ha:

$$\log_{10}(\omega) = +\log_{10}\left(\frac{1}{R_S \cdot C_f}\right)$$

e quindi:

$$\omega = \frac{1}{R_S \cdot C_f}$$

INTEGRATORE REALE

Si osserva che, **per ω tendente a zero**, il condensatore C_f della rete di retroazione negativa dell'integratore ideale tende a comportarsi come circuito aperto, facendo tendere all'infinito l'impedenza Z_2 .

Di conseguenza anche il modulo della risposta in frequenza dell'integratore ideale:

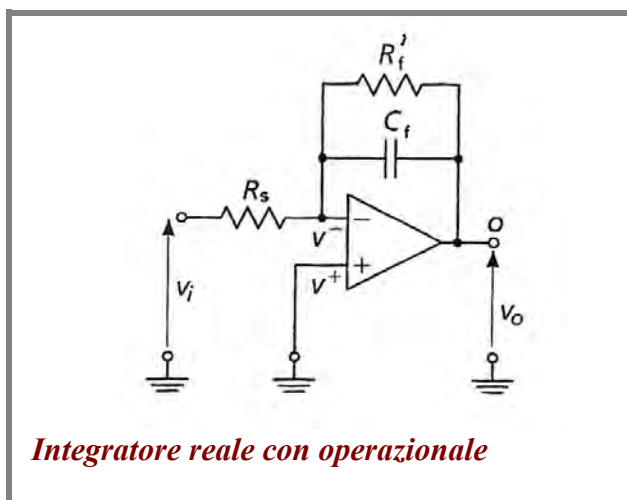
$$A = \frac{\bar{Z}_2}{Z_1} = \frac{\frac{1}{\omega C_f}}{R_s} = \frac{1}{\omega \cdot R_s \cdot C_f}$$

tende ad ∞ al tendere di ω a zero.

Ciò comporta che anche la sola presenza di un leggero offset (cioè di un indesiderato segnale a frequenza zero) può portare il dispositivo ideale in saturazione facendo venir meno la relazione

$$v_o(t) = -\frac{1}{R_s \cdot C_f} \cdot \int_{t=0}^t i \cdot dt + v_o(o) \cdot$$

Per ovviare a questo inconveniente, si pone **in parallelo** al condensatore una resistenza R'_f di **limitazione del guadagno dell'amplificatore-integratore alle basse frequenze**.



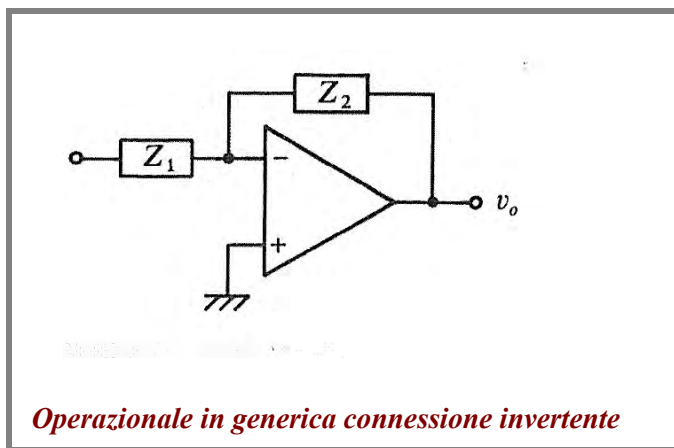
Che funzione svolge $\mathbf{R'_f}$?

Il problema nasce dall'impedenza del condensatore posto fra l'uscita e l'ingresso invertente, che diventa infinita (circuitto aperto) per $\omega=0$ e rende infinita l'amplificazione.

Il fatto che il condensatore, per $f=0$, si comporti come un circuito aperto significa che viene meno la retroazione negativa e che l'operazionale si ritrova ad essere in catena aperta.

La resistenza $\mathbf{R'_f}$ in parallelo alla capacità serve proprio a mantenere “finita” la Z_2 anche a frequenza nulla, e quindi anche a evitare che l'amplificazione dell'integratore diventi infinita.

INTEGRATORE REALE: ANALISI NEL DOMINIO DELLA FREQUENZA



In base alle considerazioni prima svolte, l'integratore reale si può realizzare aggiungendo una resistenza R'_f in parallelo al condensatore C_f dell'impedenza Z_2 .

Si può realizzare quindi un integratore reale, a partire dalla generica configurazione invertente della figura, ponendo:

- Z_2 = parallelo dell'impedenza di C_f e della resistenza aggiuntiva R'_f

$$\overline{Z_2} = \frac{R'_f}{1 + j\omega \cdot R'_f \cdot C_f}$$

- $Z_1 = R_s$, come nel circuito ideale

La risposta in frequenza dell'integratore reale risulta perciò:

$$\overline{A} = - \frac{\frac{R'_f}{1 + j\omega \cdot R'_f \cdot C_f}}{R_s} = - \frac{R'_f}{1 + j\omega \cdot R'_f \cdot C_f} \cdot \frac{1}{R_s} = - \frac{R'_f}{R_s} \cdot \frac{1}{1 + j\omega \cdot R'_f \cdot C_f}$$

e quindi:

$\overline{A} = \frac{-\frac{R'_f}{R_S}}{1 + j\omega \cdot R'_f \cdot C_f}$	Risposta in frequenza complessa dell'integratore reale (*)
---	---

Determinando il modulo del numeratore e del denominatore otteniamo:

$ \overline{A} = A = \frac{\frac{R'_f}{R_S}}{\sqrt{1 + (\omega \cdot R'_f \cdot C_f)^2}}$	Modulo della risposta in frequenza dell'integratore reale
--	--

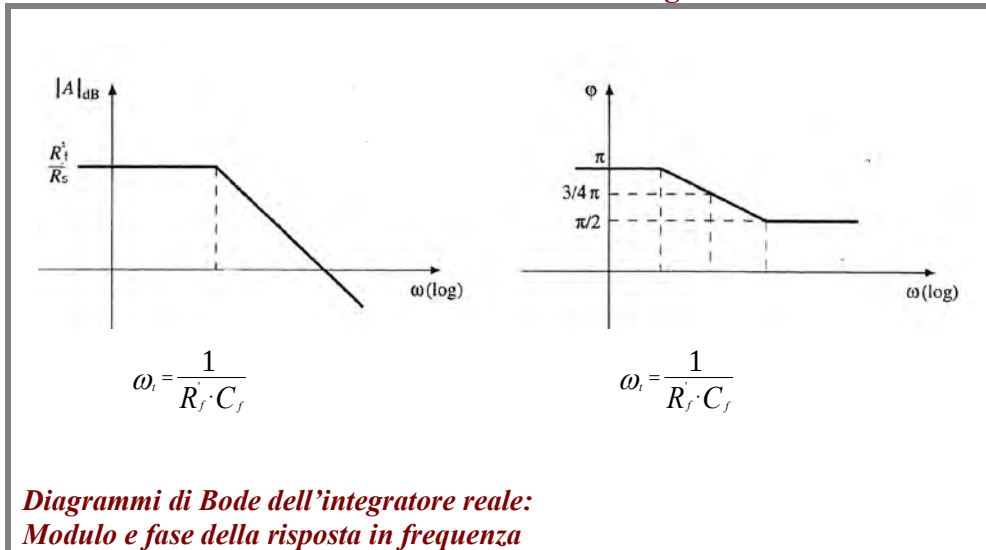
Determiniamo ora, dalla (*), la **fase** (o “**argomento**”) della risposta in frequenza complessa, avendo presente che la fase di un rapporto di grandezze complesse è la differenza fra la fase del numeratore e la fase del denominatore:

$$\varphi = \arg(\overline{A}) = \arg\left(-\frac{R'_f}{R_S}\right) - \arg(1 + j\omega \cdot R'_f \cdot C_f)$$

$$\varphi = \arg(\overline{A}) = \pm 180^\circ - \arctg\left(\frac{\omega \cdot R'_f \cdot C_f}{1}\right)$$

$\varphi = +180^\circ - \arctg(\omega \cdot R'_f \cdot C_f)$	fase della risposta in frequenza dell'integratore reale (assume valori decrescenti da +180° a +90°)
--	---

DIAGRAMMI ASINTOTICI di BODE dell' integratore REALE



INTEGRATORE REALE: RISPOSTA IN FREQUENZA

Nella figura precedente è rappresentato l'andamento (linearizzato e riferito a un asse delle pulsazioni in scala logaritmica) del modulo (espresso in dB) e della fase della risposta in frequenza dell'integratore reale.
Nelle pagine successive cercheremo di capire come si costruiscono i diagrammi.

RISPOSTA IN FREQUENZA DELL'INTEGRATORE REALE: ANALISI DELLA COSTRUZIONE DEL MODULO DEL DIAGRAMMA DI BODE

MEDIANTE GLI ZERI E POLI DELLA FUNZIONE DI TRASFERIMENTO IN “s”

La risposta in frequenza “complessa” dell’integratore reale è:

$$\overline{A} = \frac{-\frac{R'_f}{R_s}}{1+j\omega R'_f C_f} \quad \text{e il suo modulo è:} \quad |\overline{A}| = \frac{\frac{R'_f}{R_s}}{\sqrt{1+(\omega R'_f C_f)^2}}$$

Dalla risposta complessa, sostituendo “s” al posto di “jω” si ottiene la funzione di trasferimento in s:

$$\overline{A} = \frac{-\frac{R'_f}{R_s}}{1+s \cdot R'_f C_f}$$

La funzione di trasferimento presenta quindi:

- al numeratore un termine costante
- al denominatore un termine relativo a un polo semplice.

Uguagliando a zero il denominatore della funzione di trasferimento si ottiene il valore del polo:

$$s = p = -\frac{1}{R'_f \cdot C_f}$$

cui corrisponde una pulsazione di rottura:

$$\omega_t = |p| = \frac{1}{R'_f \cdot C_f}$$

così come si può osservare dalla figura precedente.

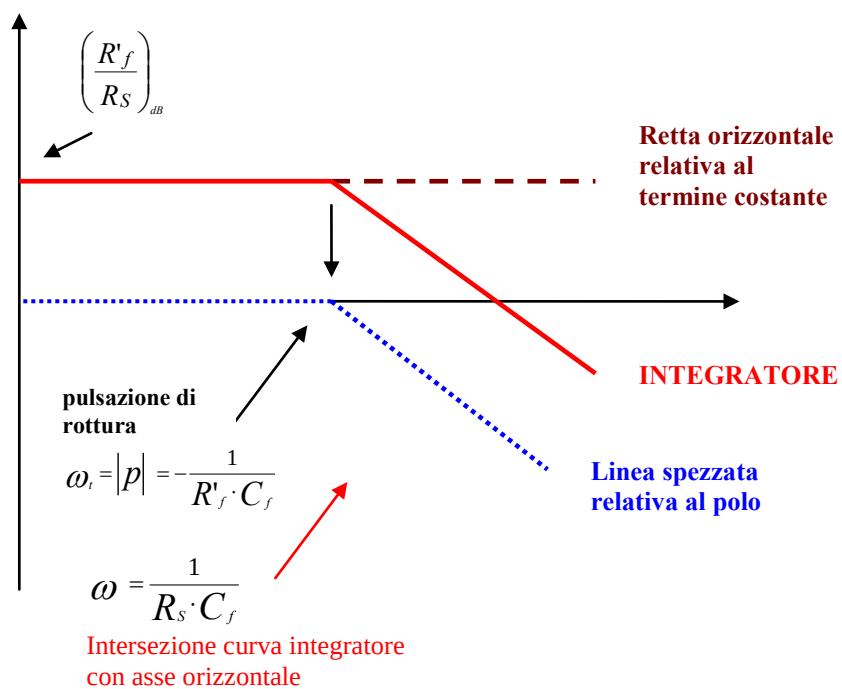
Nel diagramma di Bode del modulo della risposta in frequenza **un polo semplice** dà luogo a una linea “**spezzata**”, orizzontale fino alla pulsazione di rottura ω_t , e poi decrescente con pendenza di -20dB (si veda la figura seguente).

Invece il **termine costante** relativo al numeratore si presenta come una **retta orizzontale**:

$$(K)_{dB} = 20 \cdot \log \left(\frac{R'_f}{R_s} \right) = \left(\frac{R'_f}{R_s} \right)_{dB}$$

Dalla somma grafica della spezzata e della retta orizzontale si ottiene il diagramma complessivo del modulo della risposta, riportato nella figura che segue.

COSTRUZIONE DEL MODULO DELLA RISPOSTA



RISPOSTA IN FREQUENZA DELL'INTEGRATORE REALE: ANALISI DELLA COSTRUZIONE DEL MODULO DIAGRAMMA DI BODE

IN BASE ALL'APPROSSIMAZIONE ASINTOTICA DELLA RISPOSTA IN FREQUENZA

Il tracciamento del diagramma del modulo della risposta in frequenza può essere effettuato anche senza individuare i poli e gli zeri della funzione di trasferimento in s, bensì valutando il comportamento del modulo della risposta in frequenza (espresso in dB) per ω tendente a zero e per ω tendente ad infinito, valutando cioè il comportamento asintotico del modulo della risposta.

Dato il modulo della risposta in frequenza:

$$|A| = \frac{\frac{R'_f}{R_s}}{\sqrt{1 + (\omega R'_f C_f)^2}}$$

lo esprimiamo in dB

$$A_{dB} = 20 \cdot \log\left(\frac{R'_f}{R_s}\right) - 20 \cdot \log\left(\sqrt{1 + (\omega R'_f C_f)^2}\right) \quad (**)$$

Cominciamo ad esaminare l'andamento a bassa frequenza, cioè per

$$\omega \cdot R'_f \cdot C_f \ll 1 \quad \Leftrightarrow \quad \omega \ll \frac{1}{R'_f \cdot C_f}$$

L'ipotesi fatta sulla pulsazione ci permette di accettare l'approssimazione:

$$1 + (\omega R'_f C_f)^2 \cong 1$$

grazie alla quale l'equazione del modulo diventa:

$$A_{dB} = 20 \cdot \log\left(\frac{R'_f}{R_s}\right) - 20 \cdot \log(1)$$

ed essendo $\log(1)=0$:

$A_{dB} = 20 \cdot \log \left(\frac{R'_f}{R_S} \right)$	Espressione del modulo per $\omega \ll 1$
--	--

Questa equazione ci dice che a **bassa frequenza** il modulo è una **retta orizzontale** a quota:

$$\left(\frac{R'_f}{R_S} \right)_{dB}$$

Analizziamo ora il comportamento ad **alta frequenza** cioè per:

$$\omega \cdot R'_f \cdot C_f \gg 1 \quad \Leftrightarrow \quad \omega \gg \frac{1}{R'_f \cdot C_f}$$

Nel secondo termine del secondo membro dell'espressione (**) del modulo possiamo osservare che:

$$1 + (\omega \cdot R'_f \cdot C_f)^2 \cong (\omega \cdot R'_f \cdot C_f)^2$$

per cui l'equazione del modulo diventa:

$$A_{dB} = 20 \cdot \log \left(\frac{R'_f}{R_S} \right) - 20 \cdot \log \left(\sqrt{(\omega \cdot R'_f \cdot C_f)^2} \right)$$

e quindi:

$$A_{dB} = 20 \cdot \log \left(\frac{R'_f}{R_S} \right) - 20 \cdot \log (\omega \cdot R'_f \cdot C_f)$$

Siccome ci interessa il comportamento per valori di pulsazione molto alti, possiamo ritenere che, a tali valori, il termine costante:

$$20 \cdot \log \left(\frac{R'_f}{R_S} \right)$$

risultati trascurabile rispetto a

$$20 \cdot \log(\omega \cdot R'_f \cdot C_f)$$

per cui:

$$A_{dB} \cong -20 \cdot \log(\omega \cdot R'_f \cdot C_f)$$

Infine, applicando la proprietà $\log(x \cdot y) = \log(x) + \log(y)$,
il modulo si può esprimere come:

$A_{dB} \cong -20 \cdot \log(\omega) - 20 \cdot \log(R'_f \cdot C_f)$	Espressione del modulo per $\omega \gg 1$
---	--

Il diagramma di Bode si ottiene rappresentando questa equazione in un piano
nel quale:

- le **pulsazioni** vengono riportate sull'**asse orizzontale in scala
logaritmica**, per cui la **variabile indipendente** non è ω , bensì **log ω**
- **A** viene **espresso in dB**, cioè come:
 $(A)_{dB} = 20 \cdot \log_{10} A$
- La costante " **$R'_f C_f$** " viene anch'essa espressa in **dB**.

Rappresentata in tale piano, l'equazione

$$A_{dB} \cong -20 \cdot \log(\omega) - 20 \cdot \log(R'_f \cdot C_f) \quad (***)$$

è una **retta inclinata** caratterizzata da:

- **coefficiente angolare m = -20dB/décade (pendenza negativa)**
- **intersezione con l'asse orizzontale in:** $\omega = \frac{1}{R'_f \cdot C_f}$

L'intersezione con l'asse orizzontale è stata determinata imponendo $(A)_{dB}=0$ nella (***), nel modo seguente:

$$0 = 20 \cdot \log_{10}(\omega) + 20 \cdot \log_{10}(R'_f \cdot C_f)$$

$$20 \cdot \log_{10}(\omega) = -20 \cdot \log_{10}(R'_f \cdot C_f)$$

$$\log_{10}(\omega) = -\log_{10}(R'_f \cdot C_f)$$

Per la proprietà dei logaritmi: $-\log(x) = \log(1/x)$ si ha:

$$\log_{10}(\omega) = +\log_{10}\left(\frac{1}{R'_f \cdot C_f}\right)$$

$$\omega = \frac{1}{R'_f \cdot C_f}$$

DIMENSIONAMENTO DELL'INTEGRATORE

Come si vede dai diagrammi di Bode, l'integratore reale si comporta effettivamente da integratore solo per pulsazioni maggiori della pulsazione di rottura (tratto inclinato della spezzata).

Per pulsazioni minori funziona invece come amplificatore invertente (tratto orizzontale della spezzata).

Perciò quando dimensioniamo un integratore dobbiamo imporre

$$\omega_i > \omega_t \Leftrightarrow \omega > \frac{1}{R'_f \cdot C_f}$$

Con:

ω_i = pulsazione del segnale di ingresso

ω_t = pulsazione di rottura

In particolare **se imponiamo**:

$$\omega_i = 100 \cdot \frac{1}{R'_f \cdot C_f}$$

il dispositivo si **comporta** da integratore **pressoché ideale**, imponendo al segnale uno sfasamento in anticipo quasi esattamente di 90°.

L'ultima imposizione ci permette inoltre di **fare uso⁴ delle relazioni del dispositivo ideale**.

Riguardo ai valori di R'_f e di R_S , se non vi sono specifiche di progetto che li vincolino, si può porre:

$$R'_f = 10 \cdot R_S$$

⁴ per la determinazione della costante $\frac{1}{R C}$

ESEMPIO DI DIMENSIONAMENTO DI UN INTEGRATORE REALE

E' richiesto di progettare un integratore attivo che, ricevendo in ingresso un segnale sinusoidale con ampiezza picco-picco di 1V e frequenza 1 kHz, fornisca in uscita un segnale sinusoidale con ampiezza picco-picco di 10 V e sfasato di 90° in anticipo rispetto a v_i

OSSERVAZIONI PRELIMINARI

- 1) Per rispettare la specifica sull'**amplificazione** (che deve essere pari a 10) entra in gioco la costante di tempo $R_S \cdot C_f$ (e non la $R_f \cdot C_f$)
- 2) Non possiamo usare l'integratore ideale perché viene portato subito in saturazione dalla tensione di offset di ingresso. Useremo perciò l'integratore reale, ma nel campo di frequenza in cui esso si comporta da integratore ideale (campo di frequenze in cui la risposta in frequenza del dispositivo reale e quella del dispositivo ideale coincidono)
- 3) Se facciamo lavorare l'integratore reale ad una frequenza in cui si comporta da ideale, possiamo usare le relazioni dell'integratore ideale e cioè:

$$\bar{A} = -\frac{j \frac{1}{\omega \cdot C_f}}{R_s} = +j \frac{1}{\omega \cdot R_s \cdot C_f} \quad (\text{risposta in frequenza})$$

e la relazione ingresso-uscita

$$v_o(t) = -\frac{1}{R_s \cdot C_f} \cdot \int_{t=0}^t v_i \cdot dt + v_o(0)$$

che (in caso di ingresso sinusoidale e di condensatore inizialmente scarico) diventa⁵:

⁵ ricavata dalla precedente per $v_i(t) = V_{i(\max)} \sin(\omega \cdot t)$

$$v_o(t) = -\frac{1}{RC} \int_{t=0}^t V_{i(\max)} \cdot \sin(\omega t) dt$$

Applicando poi la regola matematica di integrazione:

$$\int \sin(\omega t) dt = -\frac{1}{\omega} \cdot \cos(\omega t)$$

otteniamo:

$$v_o(t) = \frac{+1}{R_S \cdot C_f \cdot \omega} \cdot V_{i(\max)} \cdot \cos(\omega t)$$

Da cui si può ottenere il **fattore di amplificazione**:

$A = \frac{V_{O(\max)}}{V_{i(\max)}} = \frac{1}{R_S \cdot C_f \cdot \omega}$	fattore di amplificazione (valido nel caso di ingresso sinusoidale)
--	--

- 4) Come si fa ad utilizzare l'integratore reale nel campo di frequenze in cui si comporta da ideale?

Come già accennato, la pulsazione del segnale di ingresso deve essere molto più elevata di quella di rottura::

$$\omega_i > \omega_t \quad \Leftrightarrow \quad \omega > \frac{1}{R'_f \cdot C_f}$$

Con:

ω_i = pulsazione del segnale di ingresso

ω_t = pulsazione di rottura, con espressione:

$$\omega_t = \frac{1}{R'_f \cdot C_f}$$

Nei casi (come quello in esame) in cui si vuole uno sfasamento, dell'uscita rispetto all'ingresso, pari effettivamente a $+90^\circ$ (comportamento da integratore ideale) bisogna imporre:

$$\omega_i = 100 \cdot \frac{1}{R'_f \cdot C_f} \quad \text{con } \omega_i \text{ pulsazione del segnale.}$$

Osserviamo esplicitamente che imporre la condizione su ω_i non significa scegliere la frequenza del segnale (dato che questa ci è imposta dal problema) bensì scegliere il valore del prodotto $\mathbf{R}'_f \mathbf{C}_f$

SVOLGIMENTO

Ricordiamo che è richiesto di progettare un integratore reale che, alla frequenza $f_i = 1 \text{ kHz}$, amplifichi di 10 e sfasi di $+90^\circ$.

Deve quindi essere:

$$f_s = 1 \text{ kHz} \leftrightarrow \omega_i = 6283 \left[\frac{\text{rad}}{\text{sec}} \right] \quad e \quad inoltre \quad \begin{cases} |\overline{A}(f_i)| = 10 \\ \varphi(f_i) = 90^\circ \end{cases}$$

La specifica sull'amplificazione va soddisfatta (facendo riferimento al fattore di amplificazione dell'integratore ideale nel caso di segnali sinusoidali) imponendo:

$$\frac{1}{R_s \cdot C_f \cdot \omega_i} = 10$$

Da

$$\frac{1}{R_s \cdot C_f \cdot \omega_i} = 10$$

una volta sostituito in essa il valore della pulsazione ω_i del segnale, otteniamo il valore di $\mathbf{R}'_f \mathbf{C}_f$

In questo caso si ottiene:

$$R_s \cdot C_f = 15,92 \cdot 10^{-6} \text{ [s]}$$

Si può quindi scegliere arbitrariamente il valore di C_f e questa scelta fisserà il valore di R_s .

Possiamo, per esempio, scegliere:

$$C_f = 15 \text{ [nF]} = 15 \cdot 10^{-9} \text{ [F]}$$

da cui:

$$R_S \cong \frac{15,92 \cdot 10^{-6}}{C_f}$$

$$R_S \cong \frac{15,92 \cdot 10^{-6}}{15 \cdot 10^{-9}} \cong 1061 \quad [\Omega]$$

La specifica sullo sfasamento va soddisfatta imponendo

$$\omega_i = 100 \cdot \frac{1}{R'_f \cdot C_f}$$

Siccome conosciamo ω_i e C_f , da questa relazione possiamo ottenere il valore dell'ultimo componente, cioè di R'_f .

$$R'_f = \frac{100}{C_f \cdot \omega_i}$$

$$R'_f = \frac{100}{15 \cdot 10^{-9} \cdot 6283} \cong 1061 \quad [K\Omega]$$

NOTA

Il valore di R'_f non è critico: se non interessa ottenere con precisione una fase di 90° si può scegliere anche un valore inferiore diminuendo così il guadagno R'_f/R_S valido per le componenti continue.

CAPITOLO 29

LA MODULAZIONE

COMPLEMENTI

I DIVERSI TIPI DI MODULAZIONE AM

SPETTRI DEI SEGNALE MODULATI IN AMPIEZZA SECONDO LE VARIE TECNICHE

CARATTERISTICHE E UTILIZZAZIONE DELLE DIVERSE MODULAZIONI DI AMPIEZZA

MODULAZIONE DSB O DSB-SC (Double Side Band – Suppressed Carrier)

MODULAZIONE SSB O SSB-SC (Single Side Band – Suppressed Carrier) o BLU (Banda Laterale Unica)

LA SOPPRESSIONE DELLA PORTANTE

MODULAZIONE VSB

DEMODULAZIONE DEI SEGNALE DSB

DEMODULAZIONE DEI SEGNALE A SINGOLA BANDA LATERALE (SSB)

MODULAZIONI CON SEGNALE INFORMATIVO DIGITALE E PORTANTE SINUSOIDALE

- *OOK o (ASK-OOK)*
- *ASK*
- *FSK*
- *PSK a DUE LIVELLI (2-PSK o B-PSK)*
- *2-PSK DIFFERENZIALE*
- *M-PSK (PSK A PIU' di DUE LIVELLI)*
- *MODULAZIONE QAM*
- *MODULAZIONE TCM (Trellis Coded Modulation) O MODULAZIONE QAM-TCM
O MODULAZIONE QAM CON CODIFICA TRELLIS*

DEMODULAZIONI

BANDA DI UN SEGNALE CON MODULANTE DIGITALE E PORTANTE SINUSOIDALE

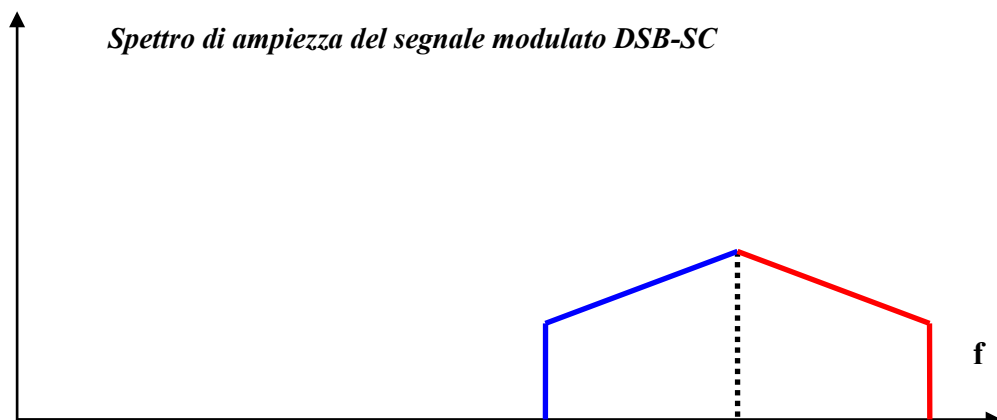
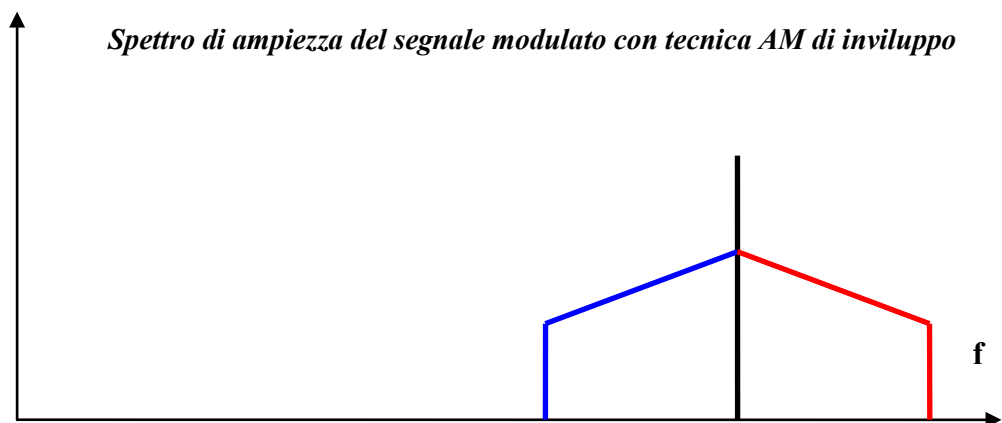
I DIVERSI TIPI DI MODULAZIONE AM

Oltre a quella di inviluppo, esistono altri tipi di modulazione di ampiezza. Tracciamo un quadro completo delle modulazioni di ampiezza:

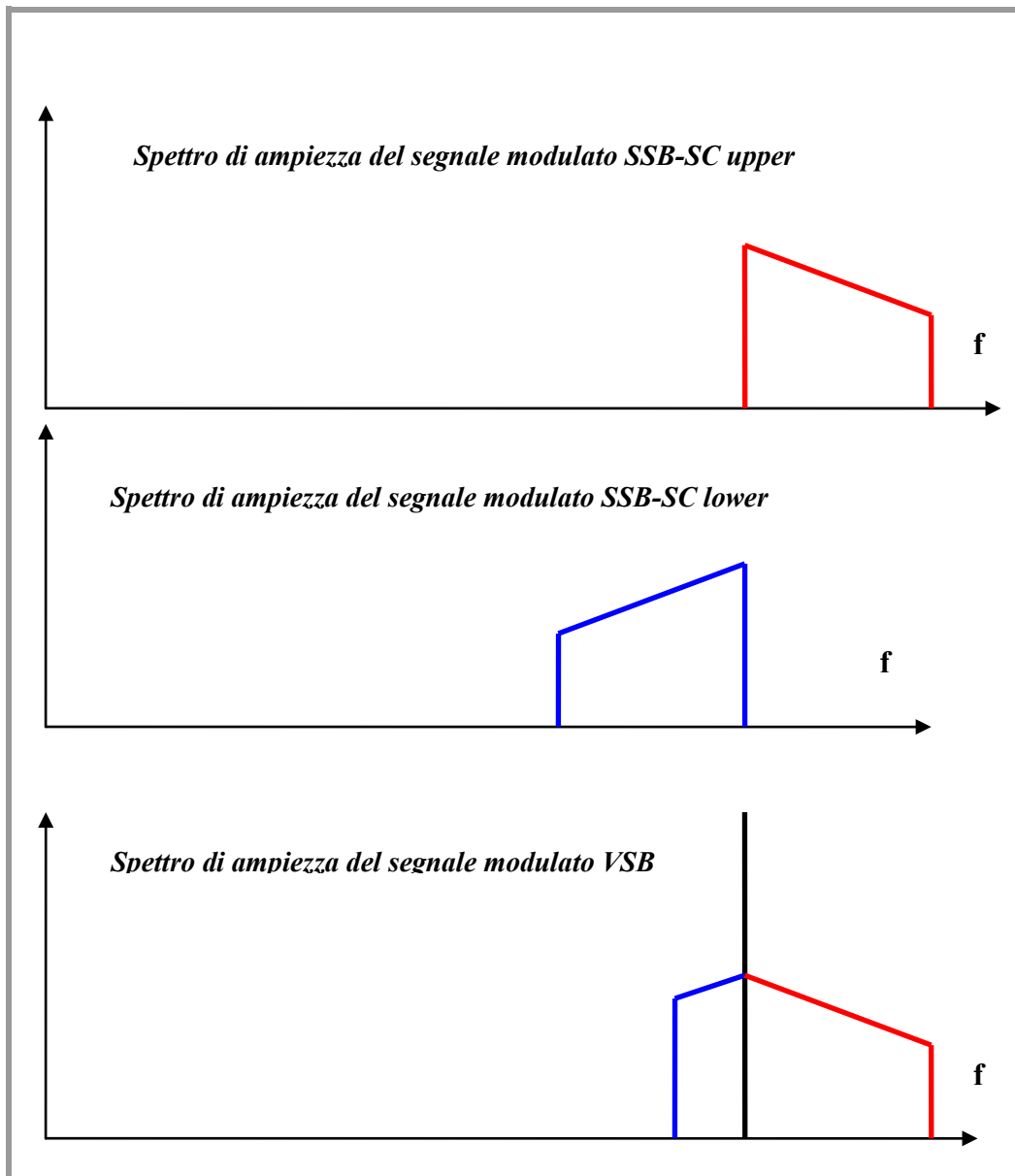
- Modulazione **AM di inviluppo**, usata per la radiodiffusione commerciale (broadcast)
- Modulazione **DSB-SC** (Double Side Band – Suppressed Carrier), o semplicemente DSB: modulazione a **doppia banda laterale con portante soppressa**.
- Modulazione **SSB-SC** (Single Side Band – Suppressed Carrier), o semplicemente SSB: modulazione a singola banda laterale con portante soppressa. E' anche chiamata "**BLU**", che sta per "**Banda Laterale Unica**".
- Modulazione **VSB**: modulazione di ampiezza a **banda vestigiale** o a banda laterale residua, usata per il segnale video televisivo.

Prima di esaminare le singole modulazioni, osserviamone gli spettri di ampiezza.

SPETTRI DEI SEGNALE MODULATI IN AMPIEZZA SECONDO LE VARIE TECNICHE



SPETTRI DEI SEGNALE MODULATI IN AMPIEZZA SECONDO LE VARIE TECNICHE



CARATTERISTICHE E UTILIZZAZIONE DELLE DIVERSE MODULAZIONI DI AMPIEZZA

MODULAZIONE	PREGI	INCONVENIENTI	UTILIZZAZIONE
AM DI INVILUPPO	Semplicità/economicità del ricevitore	Maggior potenza richiesta per la trasmissione della portante, che non contiene informazione	Radiodiffusione Broadcast
DSB-SC Doppia Banda Laterale con portante soppressa	Minor potenza richiesta al modulatore grazie alla soppressione della portante, che non contiene informazione	Maggior complessità/ maggior costo del ricevitore, che deve rigenerare la portante soppressa	E' una delle tecniche impiegate nell'elaborazione del segnale per la radiodiffusione FM stereo. Modulazione del segnale televisivo di cromaticità Trasmissioni punto-punto e bidirezionali
SSB-SC Singola Banda Laterale con portante soppressa	Riduzione dell'occupazione di banda. Minor potenza richiesta al modulatore grazie alla soppressione della portante e di una banda laterale	Maggior complessità/ maggior costo del ricevitore, che deve rigenerare la portante soppressa	Trasmissioni vocali (e non musicali) punto-punto e bidirezionali Comunicazioni marittime (GMDSS ¹) Trasmittenti militari Radioamatori Moltiplicazione a divisione di frequenza (FDM)
VSB Modulazione a Banda Vestigiale o residua	Riduzione dell'occupazione di banda.	Maggiore complessità	Modulazione del segnale televisivo di luminanza

¹ GLOBAL MARITIME DISTRESS AND SAFETY SYSTEM
sistema marittimo globale per il soccorso (distress) e la sicurezza (safety)

MODULAZIONE DSB O DSB-SC (Double Side Band – Suppressed Carrier)

Questo tipo di modulazione prevede la trasmissione di **entrambe le bande laterali** del segnale modulato, ma non della portante.

Spesso però la portante non viene eliminata del tutto, ma trasmessa con ampiezza molto ridotta per facilitare la sua ricostruzione in ricezione.

La DSB-SC consente un notevole risparmio di potenza senza danno per l'informazione.

La potenza di picco necessaria per trasmettere il segnale è la quarta parte di quella necessaria per l'AM di involuppo, mentre la tensione di picco è la metà.

L'inconveniente è che il demodulatore è più complesso e costoso di quello per l'AM di involuppo.

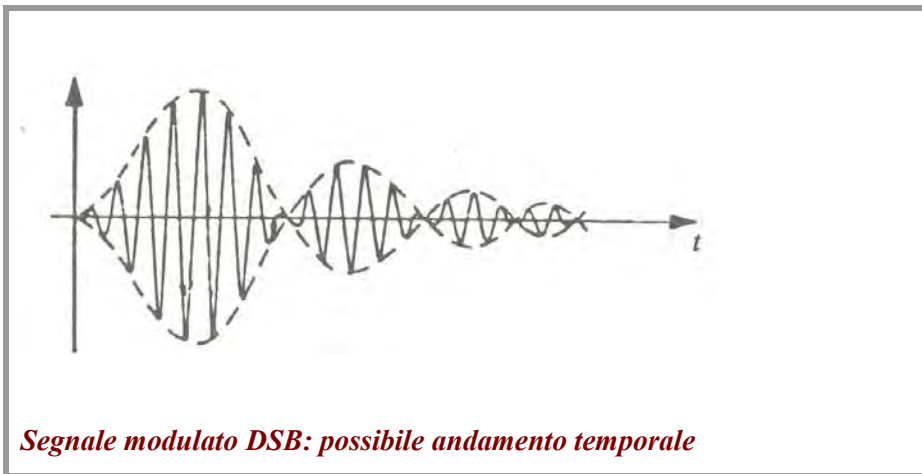
La DSB è utilizzata nella costruzione del segnale stereo e per trasmettere il segnale televisivo di cromaticanza.

La DSB non è usata nella radiodiffusione broadcast (cioè nella normale radiodiffusione con un unico trasmettitore e molti ricevitori) perché imporrebbe agli ascoltatori l'uso di ricevitori costosi.

Invece nelle trasmissioni punto-punto e bidirezionali (e in generale nelle situazioni in cui sono in gioco più trasmettitori) la DSB è preferibile perché richiede trasmettitori meno potenti.

L'inconveniente che i ricevitori siano più costosi è compensato dal fatto che ne servono pochi.

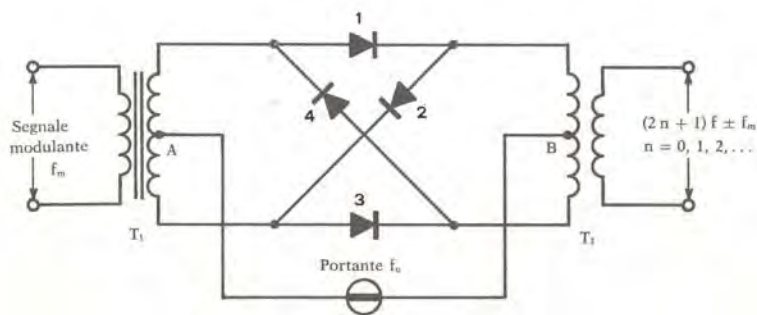
Riguardo alla AM di involuppo, precisiamo che è usata nelle trasmissioni broadcast perché il fatto che richieda più potenza in trasmissione, rispetto a DSB e a SSB, (e che necessiti di ricevitori semplici ed economici) la rende adatta proprio a una situazione nella quale l'informazione è distribuita da un unico trasmettitore a una molteplicità di ricevitori.



La DSB-SC si realizza mediante un **moltiplicatore**, come un **modulatore bilanciato ad anello**, seguito da un filtro passabasso.

Il modulatore bilanciato ad anello è basato su quattro diodi, disposti a traliccio, che conducono a coppie.

Modulatore bilanciato ad anello come moltiplicatore



Fra i terminali del trasformatore di uscita T_2 si preleva il “segnale prodotto”, che, dopo il filtraggio passabasso, diventerà il segnale DSB

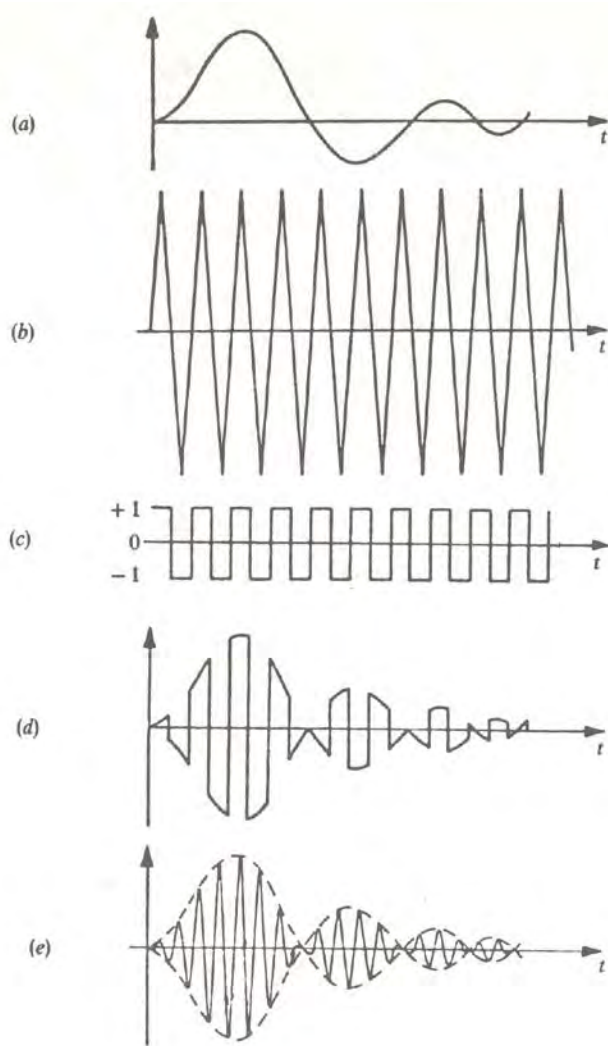
All'uscita del moltiplicatore, nel caso di modulante cosinusoidale, l'espressione del segnale è del tipo:

$$V_{DSB}(t) = K \cdot \cos(\omega_0 \cdot t) \cdot \cos(\omega_m \cdot t)$$

Nella DSB quindi viene quindi trasmesso soltanto il cosiddetto prodotto di modulazione. Applicando la formula trigonometrica di Werner, l'ultima espressione può essere scritta come:

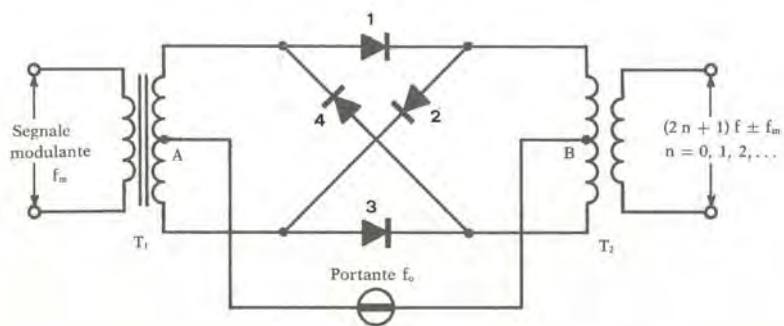
$$v_{DSB}(t) = \frac{A_0 \cdot A_M}{2} \cdot [\cos(\omega_0 - \omega_m) \cdot t + \cos(\omega_0 + \omega_m) \cdot t]$$

Dall'equazione del segnale si vede che lo **spettro** è formato soltanto da **due righe**, poste alle frequenze $f_0 - f_m$ ed $f_0 + f_m$, simmetriche rispetto alla frequenza f_0 della portante soppressa.



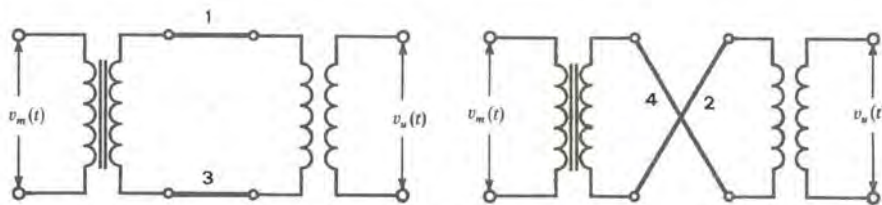
“Costruzione” del segnale DSB-SC con modulatore bilanciato ad anello:

- a) segnale modulante informativo*
- b) portante sinusoidale (è rappresentata come onda triangolare perché ipotizzata di grande ampiezza)*
- c) commutazioni dei diodi*
- d) uscita del modulatore (prima del filtraggio)*
- segnale modulato DSB-SC (dopo il filtraggio)*



Modulatore bilanciato ad anello

Fra i terminali del trasformatore di uscita T_2 , si preleva il “segnale prodotto”, che, dopo il filtraggio passabasso, diventerà il segnale DSB



A destra: schema equivalente del modulatore nel corso della semionda positiva del segnale modulante informativo.

A sinistra: schema equivalente del modulatore nel corso della semionda negativa del segnale modulante informativo.

MODULAZIONE SSB O SSB-SC (Single Side Band – Suppressed Carrier) O BLU (Banda Laterale Unica)

In questo tipo di modulazione viene trasmessa **una sola delle due bande laterali** (o l'inferiore o la superiore). La portante o viene eliminata del tutto o trasmessa con ampiezza molto ridotta per facilitare la demodulazione in ricezione.

La SSB-SC consente un notevole risparmio di potenza sia rispetto alla AM di inviluppo, sia rispetto alla DSB.

Inoltre è più immune dal fading selettivo, che, nella AM di inviluppo, può determinare distorsione (se colpisce la portante).

Come la DSB, presenta l'inconveniente di richiedere il reinserimento della portante in ricezione, per cui il demodulatore è più complesso e costoso di quello per l'AM di inviluppo.

E' utilizzata nelle trasmissioni per le quali è sufficiente l'intelligibilità del messaggio audio e non la qualità musicale.

Non usata nella radiodiffusione broadcast, la SSB è conveniente nelle comunicazioni (punto-punto e bidirezionali) fra stazioni mobili e/o fisse dove è presente da entrambi i lati un trasmettitore. Cioè quando risulta più importante ridurre costo, dimensioni e potenza del trasmettitore che non semplificare il ricevitore.

La SSB ha trovato utilizzazione nelle stazioni trasmittenti militari in onde corte, nei sistemi per radioamatori e nella multiplazione a divisione di frequenza (FDM).

Un altro esempio di applicazione della tecnica SSB riguarda le comunicazioni marittime, e in particolare le comunicazioni non satellitari, in radiotelefonica (RTF) a media e alta frequenza, nell'ambito del sistema marittimo globale per il soccorso e la sicurezza (GMDSS, Maritime Distress And Safety System).

La SSB si realizza, a partire da un segnale DSB-SC mediante **filtraggio** (eliminando una delle due bande del segnale DSB) o mediante **sfasamento**.

L'espressione del segnale modulato SSB (nel caso di modulante sinusoidale) è:

$$v_{SSB}(t) = \frac{A_0 \cdot A_M}{2} \cdot [\cos(\omega_0 - \omega_m) \cdot t] \quad \text{segnale SSB-L (Lower)}$$

oppure:

$$v_{SSB}(t) = \frac{A_0 \cdot A_M}{2} \cdot [\cos(\omega_0 + \omega_m) \cdot t] \quad \text{segnale SSB-U (Upper)}$$

Il segnale SSB è quindi costituito da una sola delle due componenti del segnale DSB e il suo **spettro** è formato da **una sola riga** posta alla frequenza $f_0 - f_m$ oppure alla frequenza $f_0 + f_m$.

LA SOPPRESSIONE DELLA PORTANTE

Si usa dire che la portante serve a trasportare l'informazione contenuta nella modulante che, da sola, incontrerebbe difficoltà o impossibilità nella propagazione o nella ricezione.

In realtà però ciò che è indispensabile non è la presenza fisica della portante nel corso della propagazione, ma le modificazioni che, grazie alla presenza –solo all'interno del trasmettitore- della portante, il segnale originario ha subito.

Dire che un sistema di modulazione prevede la soppressione della portante significa dire che il segnale modulato, negli intervalli di tempo in cui la modulante è nulla, risulta nullo (vedi DSB-SC) mentre nelle modulazioni in cui la portante non è soppressa accade che, negli intervalli nei quali la modulante è nulla, il segnale modulato non si annulla, ma è costituito dalla sola portante non modulata.

Nelle modulazioni che prevedono la soppressione della portante, anche se la portante viene soppressa, essa ha già determinato una modificazione del segnale cui è affidata l'informazione (segnale modulante informativo): prima della modulazione il segnale informativo era un segnale a bassa frequenza, poco adatto a propagarsi e a essere captato da un'antenna di dimensioni accettabili. Dopo la modulazione, l'informazione è affidata al segnale modulato che, sia in caso di presenza di portante, sia in caso di soppressione, è un segnale ad alta frequenza, adatto a propagarsi e a essere captato da antenne di dimensioni accettabili.

MODULAZIONE VSB

La modulazione di ampiezza a banda vestigiale o a banda laterale residua è usata per il segnale video televisivo di luminanza e prevede la trasmissione:

- della portante
- dell'intera banda laterale superiore
- di una parte (cioè di un vestigio o residuo) della banda laterale inferiore.

La VSB nasce dal compromesso fra l'esigenza di non ampliare troppo la larghezza di banda del segnale modulato televisivo (che risulterebbe troppo grande con l'adozione dell'AM di involuppo) e l'esigenza di salvaguardare la componente continua del segnale.

La componente continua (componente a frequenza zero) del segnale video di luminanza rappresenta il valore medio di luminosità del quadro e quindi non si può correre il rischio di tagliarla, come potrebbe avvenire se – al fine di restringere la banda- si utilizzasse la modulazione SSB.

Il segnale SSB infatti è ricavato da un segnale a doppia banda laterale mediante filtraggio e proprio l'azione di un filtro che elimini completamente una delle due bande laterali può compromettere la componente continua del segnale di luminanza.

DEMODULAZIONE DEI SEGNALE DSB

La demodulazione dei segnali a portante soppressa DSB ed SSB richiede la rigenerazione, nel ricevitore, di una portante caratterizzata dalla stessa frequenza e dalla stessa fase della portante originaria. Questa portante è detta: “portante di demodulazione” o “portante di riferimento”. La demodulazione del segnale DSB può essere effettuata in due modi:

- mediante **demodulazione di inviluppo**
- mediante **demodulazione a prodotto**.

DEMODULAZIONE DI INVILUPPO DI UN SEGNALE DSB

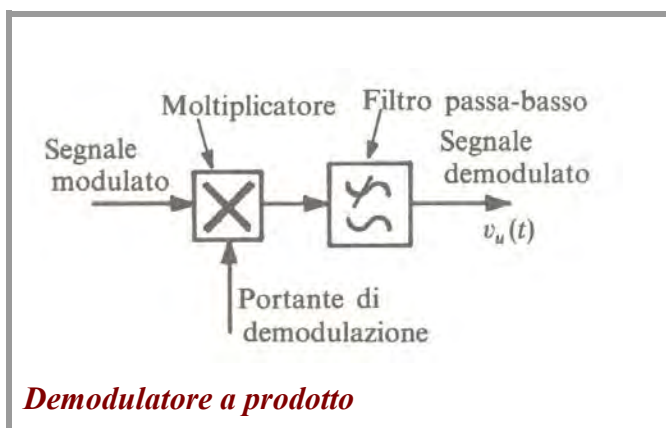
Si effettua la **somma** del segnale ricevuto (segnale DSB²) con la portante di demodulazione, cioè la portante rigenerata dal ricevitore stesso, e il segnale risultante si applica al rivelatore di inviluppo.

E' necessario che la portante abbia un'ampiezza molto maggiore di quella del segnale modulato.



DEMODULAZIONE A PRODOTTO DI UN SEGNALE DSB

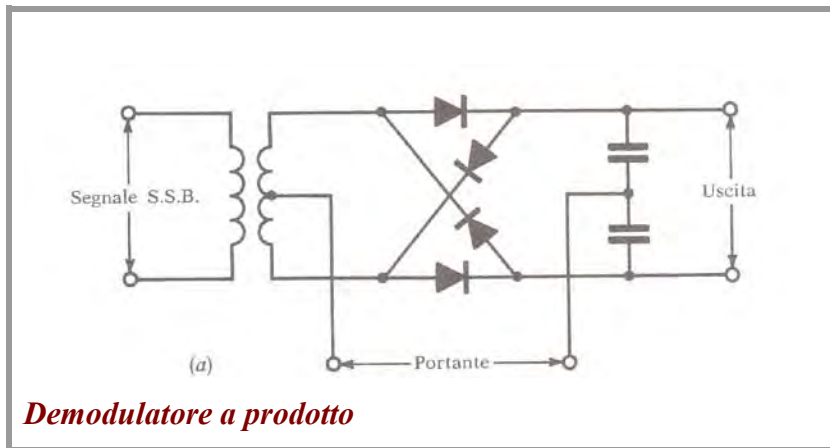
Si effettua il **prodotto** del segnale ricevuto (segnale DSB) con la portante di demodulazione, cioè la portante rigenerata dal ricevitore stesso, e il segnale prodotto viene sottoposto a filtraggio passabasso. Come **demodulatore a prodotto** si può usare un **modulatore bilanciato** analogo a quello utilizzato per la modulazione DSB. Il segnale da demodulare deve essere applicato al trasformatore di ingresso



² che è un prodotto di modulazione

DEMODULAZIONE DEI SEGNALE A SINGOLA BANDA LATERALE (SSB)

E' generalmente effettuata con un demodulatore a **prodotto** (realizzabile mediante un **modulatore bilanciato**).



MODULAZIONI CON SEGNALE INFORMATIVO DIGITALE E PORTANTE SINUSOIDALE

Sono le modulazioni che sono state usate, o sono tuttora utilizzate (QAM-Trellis) per i **modem fonici**, ma che trovano applicazione anche nelle **comunicazioni numeriche via satellite**, nel sistema **GPS** di posizionamento globale e nei **ponti radio**.

Tipicamente queste modulazioni sono utilizzate in tutti quei sistemi nei quali la banda passante del canale trasmissivo (o l'adozione della moltiplicazione di frequenza FDM) impongono restrizioni alla larghezza di banda utilizzabile.

Siccome i segnali digitali (come i messaggi generati da un PC) sono definibili "a banda larga", occorrono tecniche di modulazione in grado di trasformare uno spettro a banda larga in un spettro a banda stretta, affinché il segnale possa essere inviato su un canale a banda limitata (per esempio la linea telefonica con banda di $300\text{Hz} \div 3400\text{Hz}$).

Le fondamentali tipologie di modulazione con segnale informativo numerico e portante analogica sono:

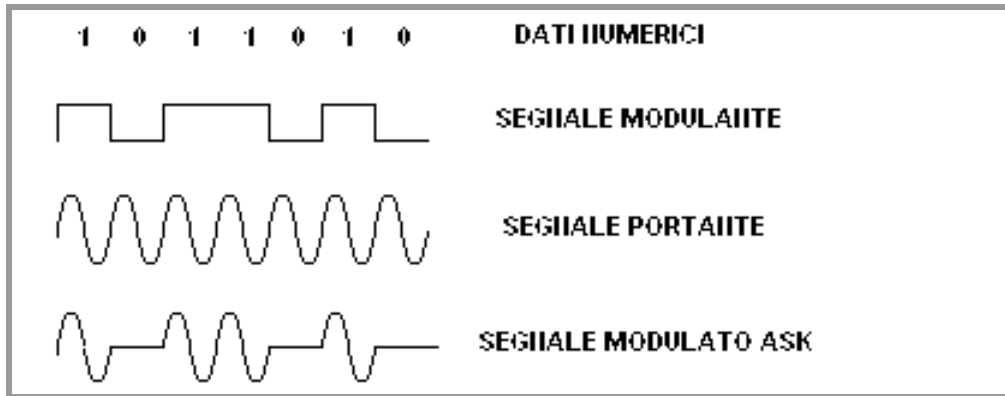
1. OOK (On-Off Keying): modulazione a spostamento (o salto) di ampiezza di tipo on/off
2. ASK (Amplitude Shift Keying): modulazione a spostamento (o salto) di ampiezza
3. FSK (Frequency Shift Keying): modulazione a spostamento (o salto) di frequenza
4. PSK (PHASE Shift Keying): modulazione a spostamento (o salto) di fase
5. QAM (Quadrature Amplitude Modulation): modulazione di ampiezza in quadratura; si tratta in realtà di una tecnica che modula sia in ampiezza che in fase.

Tipo di modulazione		Caratteristiche	Ambito tipico di impiego
Modulazioni di ampiezza	OOK	Poco resistente al rumore, bassa efficienza spettrale.	Obsoleta
	ASK	Poco resistente al rumore, bassa efficienza spettrale.	Obsoleta
Modulazioni di fase	M-PSK M-DPSK ($M \leq 8$)	Resistenza al rumore medio-alta; efficienza spettrale medio bassa.	WLAN, TV digitale satellitare UMTS, ecc.
Modulazioni miste ampiezza-fase	M-QAM; M-TCM ($M \geq 16$)	Alta o altissima efficienza spettrale; resistenza al rumore bassa.	Ponti radio, modem fonici e ADSL, TV digitale terrestre, ecc.
Modulazioni di frequenza	FSK; GFSK; MSK; GMSK	Alta resistenza al rumore; bassa efficienza spettrale.	GSM, DECT, Bluetooth, ecc.

Note: M = numero degli stati di modulazione; OOK = On Off Keying; ASK = Amplitude Shift Keying; PSK = Phase Shift Keying; D = Differential; QAM = Quadrature Amplitude Modulation; TCM = Trellis Coded Modulation; FSK = Frequency Shift Keying; MSK = Minimum Shift Keying; G = Gaussian.

OOK o (ASK-OOK)

E' una particolare tipologia di modulazione di ampiezza ASK, caratterizzata dal fatto che uno dei due livelli di ampiezza (che il segnale modulato può assumere) è nullo.



COME E' FATTO IL SEGNALE MODULATO

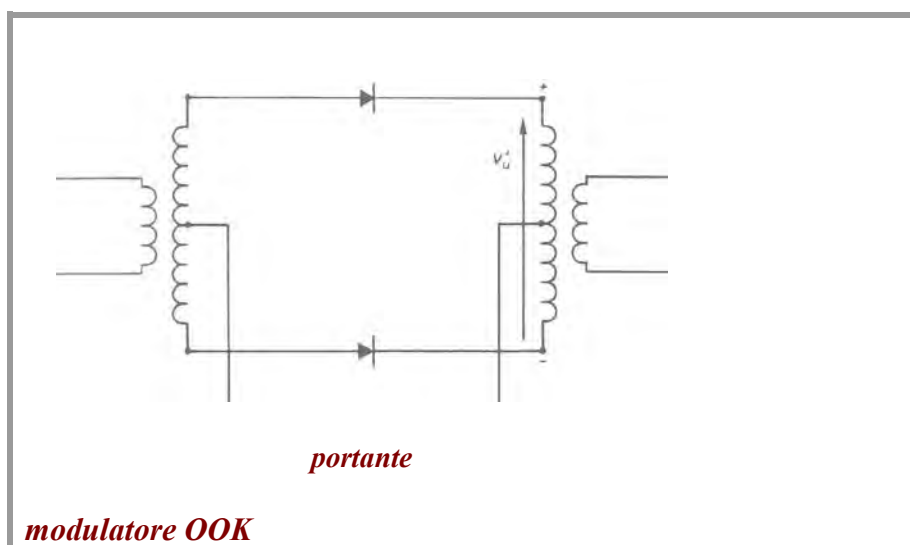
In corrispondenza dei bit "1" del segnale modulante informativo (**unipolare**), il segnale modulato ha un valore fisso di ampiezza diverso da zero (trasferimento della portante in uscita, con prefissata ampiezza).

Il segnale modulato OOK si presenta quindi come una sequenza di pacchetti di impulsi sinusoidali. In corrispondenza dei bit "0" del segnale modulante informativo, il segnale modulato ha valore nullo (assenza della portante).

IL MODULATORE OOK

Può essere realizzato con un modulatore bilanciato a due soli diodi.

- La portante è applicata ai terminali del trasformatore di ingresso
- Il segnale informativo modulante è applicato fra le prese centrali delle induttanze³ del circuito "centrale" e cioè ai capi dei diodi, che il segnale comanda
- Il segnale modulato si preleva fra i terminali del trasformatore di uscita.
-



³ Induttanze costituenti il secondario del trasformatore di ingresso e il primario del trasformatore di uscita

ASK

E' una modulazione di ampiezza che fa assumere al segnale modulato soltanto due valori di ampiezza, diversi da zero.

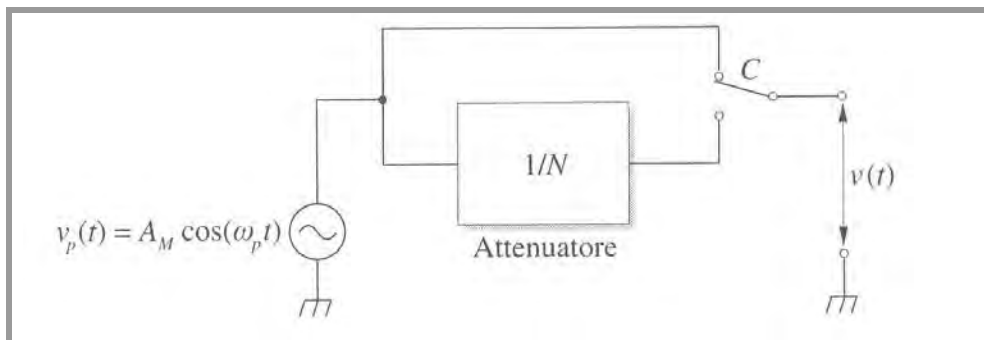
COME E' FATTO IL SEGNALE MODULATO

In corrispondenza dei bit "1" del segnale modulante informativo (**bipolare**), il segnale modulato riproduce la portante con un determinato valore di ampiezza V_{p1} . In corrispondenza del bit "0" del segnale modulante informativo, il segnale modulato riproduce ancora la portante, ma con un'ampiezza minore ($V_{p2}=V_{p1}/n$).

Il segnale modulato OOK si presenta quindi come una sequenza di pacchetti di impulsi sinusoidali, con ampiezza maggiore in corrispondenza del bit "1" del segnale modulante e un'ampiezza minore in corrispondenza del bit "0" del segnale modulante.

IL MODULATORE ASK

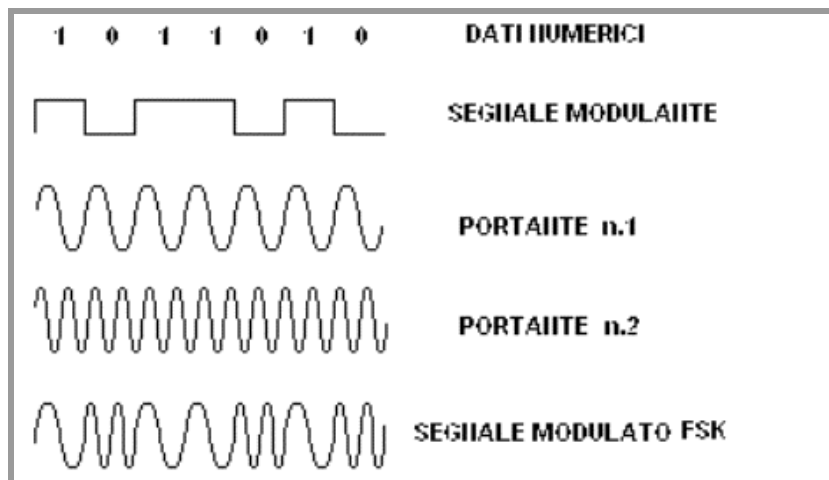
In linea di principio può essere pensato come un generatore di portante con un terminale a massa e l'altro terminale facente capo a due linee: una prima linea che riporta la portante al terminale di uscita così com'è, senza nessuna manipolazione, e una seconda linea collegata a un attenuatore che riduce di n volte l'ampiezza della portante. Un commutatore comandato dal segnale modulante informativo preleva in corrispondenza (per esempio) del bit "0" la portante attenuata e in presenza del bit uno la portante non attenuata.



FSK

E' una modulazione di frequenza che associa ai bit del segnale modulante informativo due soli valori di frequenza:

- la frequenza f_a , più alta, (tipicamente) associata al valore zero del bit
- La frequenza f_z , più bassa, (tipicamente) associata al valore uno del bit.



Esempio

	Standard V.21	Standard V.23
f_a (freq. più alta, associata al bit "0")	1180 Hz	2100 Hz
f_z (freq. più bassa, associata al bit "1")	980 Hz	1300 Hz
Velocità del segnale-dati (segnale modulante informativo)	300 bit/s	1200 bit/s

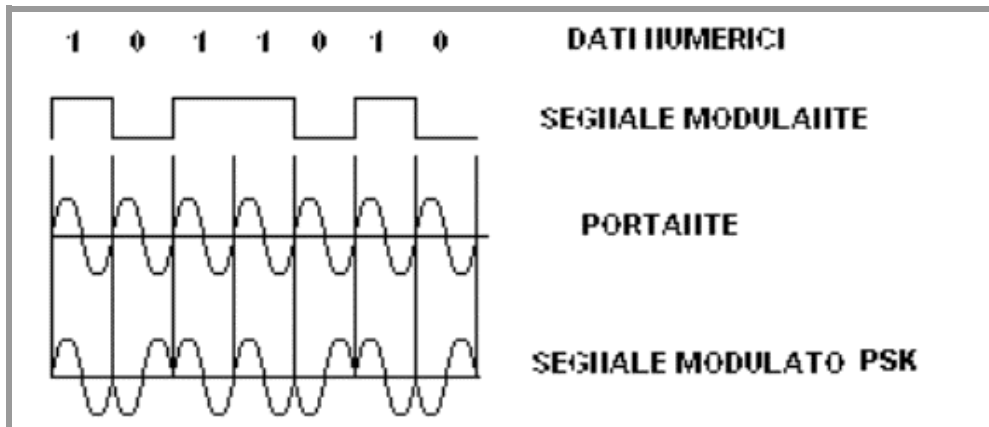
MODULATORE FSK

Può, in linea di principio, essere pensato come una coppia di oscillatori, che generano ciascuno una diversa frequenza, e un commutatore, comandato dal segnale modulante informativo, che porta in uscita il segnale del circuito a frequenza più alta in corrispondenza del bit zero e il segnale dell'oscillatore a frequenza più bassa in corrispondenza del bit 1.

PSK a DUE LIVELLI (2-PSK o B-PSK)

Questa modulazione impone alla portante due valori di sfasamento (per esempio 0° e 180°) in corrispondenza del valore “1” o “0” del bit del segnale informativo modulante.

Negli intervalli di tempo in cui il bit del segnale-dati vale “1”, il modulatore 2-PSK fornisce, come uscita, la portante non sfasata, mentre, negli istanti in cui il bit vale “0” fornisce come uscita la portante sfasata di 180° , cioè invertita rispetto all’asse dei tempi.



IL MODULATORE 2-PSK

Può essere realizzato con un **modulatore bilanciato ad anello**, con **quattro diodi** (due coppie contrapposte) o **moltiplicatore ad anello**, cioè con **lo stesso modulatore utilizzato per la modulazione DSB-SC**. Rispetto al caso DSB, però, **sono scambiate le coppie di terminali** di applicazione del segnale informativo e della portante.

La modulazione DSB è l’unico caso nel quale il segnale informativo è applicato al trasformatore di ingresso e la portante fra i punti centrali delle induttanze del circuito centrale.

Nella modulazione ASK e in tutte le altre modulazioni a spostamento, la **portante è sempre applicata al trasformatore di ingresso**, mentre la **modulante** comanda le commutazioni dei diodi e quindi è applicata fra i punti centrali delle induttanze del circuito “centrale”.

- La portante è applicata ai terminali del trasformatore di ingresso
- Il segnale informativo modulante è applicato fra le prese centrali delle induttanze⁴ del circuito “centrale” e cioè ai capi dei diodi, che il segnale comanda
- Il segnale modulato si preleva fra i terminali del trasformatore di uscita.

⁴ Induttanze costituenti il secondario del trasformatore di ingresso e il primario del trasformatore di uscita

2-PSK DIFFERENZIALE

In questa modulazione il valore logico del bit del segnale-dati provoca un salto di fase non rispetto alla portante, ma rispetto alla fase che aveva lo stesso segnale modulato nel tempo di bit precedente. In questo modo il rivelatore, in ricezione, avrà bisogno solo di un riferimento esatto di frequenza e non di fase.

La modulazione avviene nel modo seguente.

In base al segnale-dati (segnale modulante informativo) si costruisce un segnale dati “rielaborato” (realizzato per esempio mediante una porta EXOR). Il segnale-dati rielaborato modula la portante imponendole, ad ogni fronte di salita del segnale-dati rielaborato, una inversione di fase rispetto alla fase della stessa portante nel tempo di bit precedente.

In definitiva il segnale modulato è tale che:

- in presenza di ogni bit 1 del segnale-dati originario si ha uno sfasamento di 180° della portante (rispetto al tempo di bit precedente)
- in presenza del bit 0 del segnale-dati originario si ha uno sfasamento di 0° della portante (rispetto al tempo di bit precedente).

M-PSK (PSK A PIU' di DUE LIVELLI)

Questa modulazione impone alla portante più di due valori di sfasamento, per esempio quattro o otto.

Un singolo bit del segnale informativo può assumere solo due valori, per cui non può imporre più di due valori di sfasamento alla portante.

Perciò nelle modulazioni M-PSK i bit del segnale modulante informativo vengono accorpati in gruppetti, per esempio, di due bit (detti “dibit”) o in gruppetti di tre bit, detti “tribit”. I gruppetti vengono chiamati “simboli”.

Un modulatore che accorpi i bit del segnale informativo modulante in dibit, realizza la modulazione 4-PSK, perché una coppia di bit (un dibit) può assumere 2^2 cioè 4 differenti valori. Di conseguenza il modulatore può imporre quattro diverse fasi (per esempio 0° , 90° , 180° , 270°) al segnale modulato e realizzare quindi una modulazione 4-PSK.

Un modulatore che accorpi i bit del segnale informativo modulante in tribit, realizza la modulazione 8-PSK perché un tribit può assumere 2^3 cioè 8 differenti valori. Di conseguenza il modulatore può imporre otto diverse fasi (per esempio 0° , 45° , 90° , 135° , 180° , 225° , 270° , 315°) al segnale modulato e realizzare quindi una modulazione 8-PSK.

In generale una modulazione a M livelli può essere ottenuta creando simboli formati da un numero $n = \log_2 M$ di bit.

Ricordiamo che il logaritmo in base 2 di un numero M è l'esponente da dare a 2 per avere, come potenza M. Per esempio:

$\log_2(2)=1$, $\log_2(4)=2$, $\log_2(8)=3$, $\log_2(16)=4$, ecc.

Le modulazioni che **accorpano almeno due bit in un simbolo** (e il cui segnale **modulato** può assumere **più di due stati di modulazione**, per esempio più di due sfasamenti) sono chiamate “**multilivello**” e hanno il vantaggio di occupare una banda più stretta.

In generale, una modulazione con $M=2^n$ livelli, e con simboli ad n bit, occupa una banda n volte più stretta della corrispondente modulazione a due soli livelli

Ciò è giustificato dal “**criterio di Nyquist sulla relazione tra velocità e banda di un canale non rumoroso**”:

$$B_M \geq \frac{v_{MOD}}{n} = \frac{v_{MOD}}{\log_2 M} \quad (\text{dove } v_{MOD} = \text{velocità di modulazione} = \text{numero di simboli/sec})$$

Le modulazioni multilivello inoltre presentano una maggiore “**efficienza spettrale**”, efficienza definita come:

$$E = \frac{v}{B} \cdot \log_2 M \quad \left[\frac{\text{bit/s}}{\text{Hz}} \right]$$

MODULATORE 4-PSK o Q-PSK

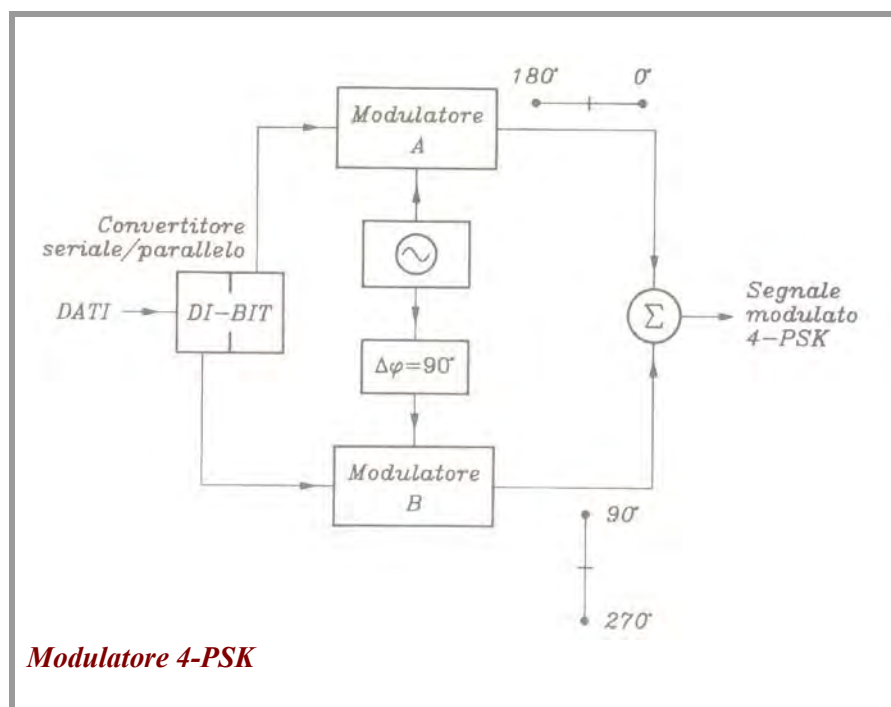
Un modulatore bilanciato ad anello (a quattro diodi) è un modulatore 2-PSK in quanto produce, a partire da una portante, un segnale caratterizzato da due livelli di fase.

Se pensiamo di utilizzare due modulatori bilanciati ad anello (moltiplicatori), con portanti sfasate fra loro di 90° e di sommare le loro uscite, otteniamo un modulatore 4-PSK.

Naturalmente il modulatore deve contenere un circuito in grado di aggregare in coppie i bit del segnale dati, cioè del segnale modulante informativo.

Tale circuito è un convertitore seriale-parallelo che, appunto, disponga in parallelo ogni due bit che riceve. Può essere realizzato con un registro SIPO a due bit. Un modulatore 4-PSK è quindi formato da:

- due modulatori 2-PSK, cioè due modulatori bilanciati ad anello a quattro diodi, ossia da due moltiplicatori
- un convertitore seriale-parallelo per generare i dibit
- un generatore di portante
- un circuito sfasatore di 90° per sfasare la seconda portante rispetto alla prima
- un circuito sommatore (in linea di principio un sommatore ad operazionale) che sommi i due segnali 2-PSK.



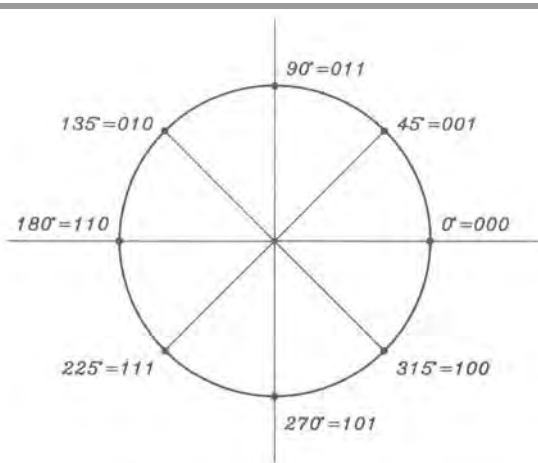
MODULATORE 8-PSK

Così come abbiamo realizzato un modulatore 4-PSK sommando le uscite di due modulatori 2-PSK (cioè di due moltiplicatori o modulatori bilanciati ad anello a quattro diodi), possiamo pensare di creare un modulatore 8-PSK sommando le uscite di due modulatori completi 4-PSK.

Nel modulatore 8-PSK lo sfasatore di portante dovrà essere da 45° (e non da 90°) e il convertitore seriale- parallelo dovrà essere a tre bit, e non a due (dovendo generare dei tribit e non dei dibit).

Un modulatore 4-PSK produce, a partire da una portante, un segnale caratterizzato da quattro livelli di fase

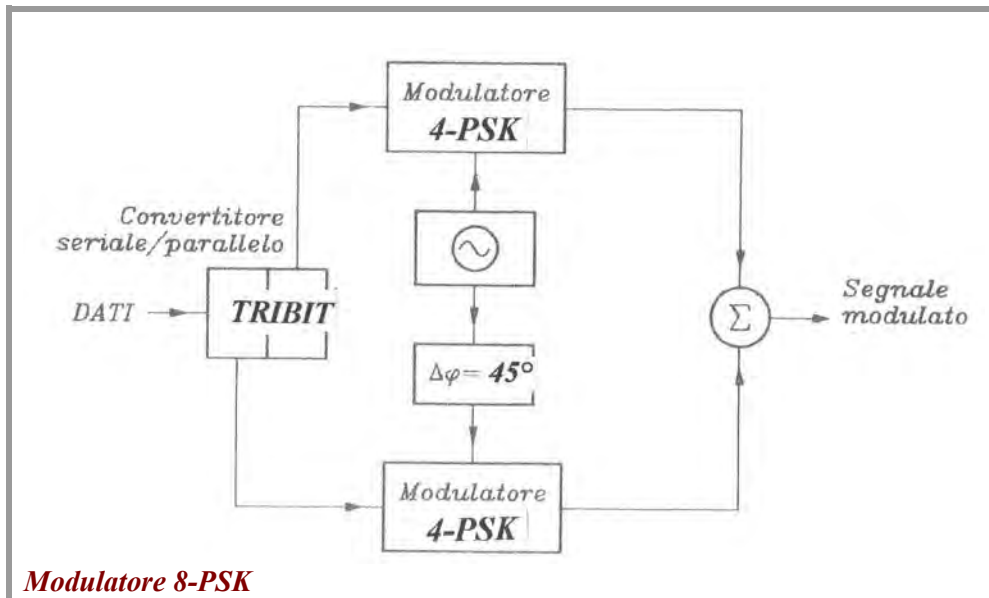
Se pensiamo di utilizzare due modulatori 4-PSK, con portanti sfasate fra loro di 45° e di sommare le loro uscite, otteniamo un modulatore 8-PSK.



Schema di distribuzione delle fasi nella modulazione 8-PSK

Un modulatore 8-PSK è quindi formato da:

- due modulatori 4-PSK
- un convertitore seriale-parallelo per generare i simboli detti tribit
- un generatore di portante
- un circuito sfasatore di 45° per sfasare la seconda portante rispetto alla prima
- un circuito sommatore (in linea di principio un sommatore ad operazionale) che sommi i due segnali 4-PSK.



MODULAZIONE QAM⁽⁵⁾

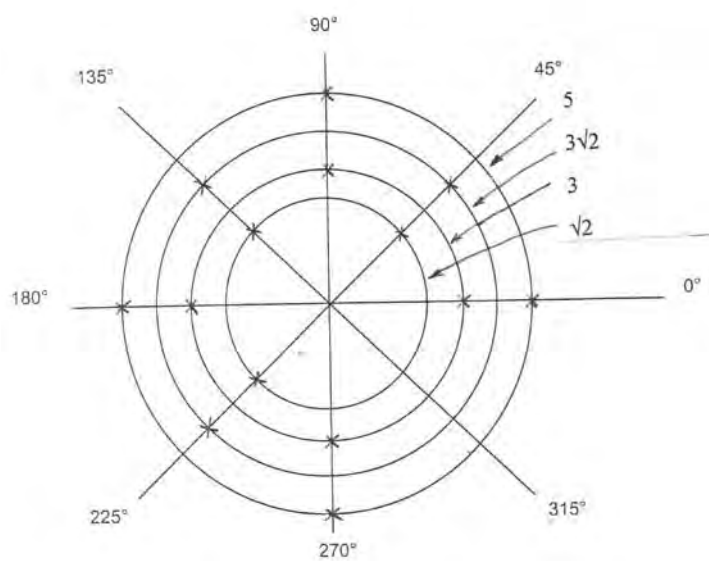
I punti di modulazione della modulazione PSK, essendo caratterizzati dalla stessa ampiezza e da fasi diverse, possono essere rappresentati lungo una circonferenza.

Al crescere del numero di punti di modulazione, e quindi al crescere del numero delle fasi, i punti di modulazione risultano sempre più vicini uno all'altro lungo la circonferenza, il che aumenta la probabilità di errore in fase di ricezione. Diventa perciò più facile confondere un punto con un altro. La modulazione QAM rimedia a questo inconveniente perché introduce, oltre a quella di fase, una modulazione di ampiezza, che distribuisce i punti di modulazione su più circonferenze, aumentando la distanza fra essi e quindi riducendo la probabilità di errore in ricezione.

La QAM quindi è una tecnica di modulazione “mista”: di fase e di ampiezza.

⁵ Detta da qualcuno “Q-PSK” (come la 4-PSK) oppure PSK-QAM

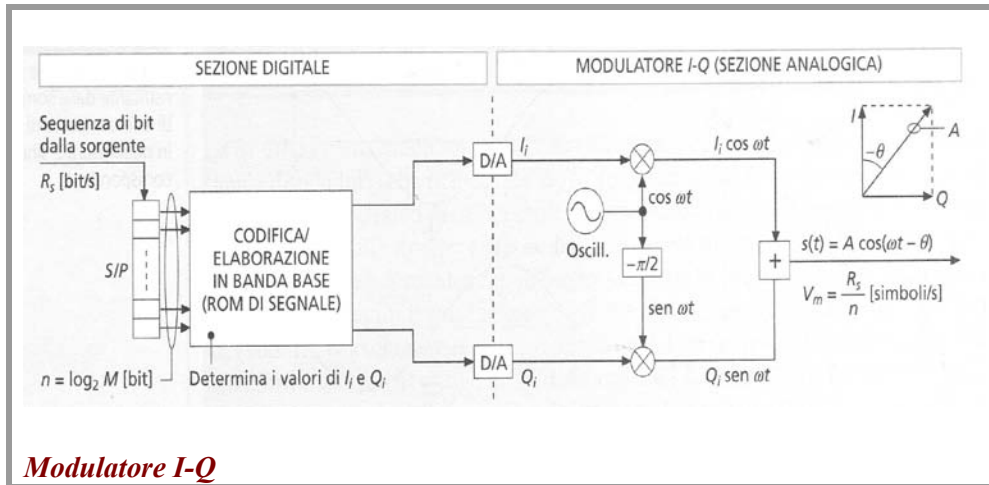
bit relativi alla fase			bit relativo all'ampiezza	ampiezza	fase
0	0	1	0	3	0°
0	0	1	1	5	
0	0	0	0	$\sqrt{2}$	45°
0	0	0	1	$3\sqrt{2}$	
0	1	0	0	3	90°
0	1	0	1	5	
0	1	1	0	$\sqrt{2}$	135°
0	1	1	1	$3\sqrt{2}$	
1	1	1	0	3	180°
1	1	1	1	5	
1	1	0	0	$\sqrt{2}$	225°
1	1	0	1	$3\sqrt{2}$	
1	0	0	0	3	270°
1	0	0	1	5	
1	0	1	0	$\sqrt{2}$	315°
1	0	1	1	$3\sqrt{2}$	



Modulazione 16-QAM, con, in basso, il diagramma a COSTELLAZIONE dei punti di modulazione

MODULATORE QAM

Analizzeremo il **MODULATORE I-Q** utilizzabile anche per le tecniche M-PSK, M-DPSK, M-QAM, M-TCM



Un segnale modulato che assuma M stati di modulazione, che possono essere M fasi, o M combinazioni fase-ampiezza, può essere ottenuto come somma vettoriale di un segnale seno e di un segnale coseno aventi ampiezze opportune. Un qualunque segnale modulato che si trovi in uno stato di modulazione caratterizzato da un'ampiezza A_i e da una fase Φ_i può essere espresso come:

$$v_{MOD}(t) = A_i \cdot \cos(\omega_p \cdot t - \Phi_i)$$

Dove A_i e Φ_i sono costanti perché sono riferite a un ben preciso stato di modulazione.

Per la formula trigonometrica di sottrazione:

$$\cos(\alpha - \beta) = \cos(\alpha) \cdot \cos(\beta) + \sin(\alpha) \cdot \sin(\beta)$$

possiamo scrivere:

$$v_{MOD}(t) = A_i \cdot \cos(\Phi_i) \cdot \cos(\omega_p \cdot t) + A_i \cdot \sin(\Phi_i) \cdot \sin(\omega_p \cdot t)$$

dove:

- $A_i \cdot \cos(\Phi_i)$ è **costante**, può essere considerata l'**ampiezza** di $\cos(\omega_p \cdot t)$ e viene chiamata "**I_i**"
- $A_i \cdot \sin(\Phi_i)$ è **costante**, può essere considerata l'**ampiezza** di $\sin(\omega_p \cdot t)$ e viene chiamata "**Q_i**".

In definitiva il segnale modulato è:

$$v_{MOD}(t) = I_i \cdot \cos(\omega_p \cdot t) + Q_i \cdot \sin(\omega_p \cdot t)$$

Questa espressione dimostra che è possibile ottenere un segnale modulato, con l'ampiezza e la fase desiderate, sommando un segnale coseno e un segnale seno (portanti in quadratura) di opportune ampiezza A_i e Q_i .

Ciascuno dei possibili stati di modulazione si ottiene scegliendo una coppia dei valori delle ampiezze I_i e Q_i ed effettuando la somma delle portanti in quadratura

$$I_i \cdot \cos(\omega_p \cdot t) \quad \text{e} \quad Q_i \cdot \sin(\omega_p \cdot t)$$

MODULAZIONE TCM (Trellis⁶ Coded Modulation) O MODULAZIONE QAM-TCM O MODULAZIONE QAM CON CODIFICA TRELLIS

Questa tecnica trae origine dal concetto di unificazione delle operazioni di codifica (numerica) e di modulazione (analogica).

La modulazione QAM-TCM utilizza una codifica detta “convoluzionale ”con memoria”, nel senso che la codifica dei bit dipende da quelli precedenti.

La codifica, aggiungendo informazione , rende ridondante quella originale, permettendo (in ricezione) la correzione diretta delle sequenze di simboli contenenti errori.

⁶ “traliccio” in inglese.

ESEMPI DI MODULAZIONI QAM E QAM TCM NEL CAMPO DEI MODEM

MODULAZIONE	NUMERO STATI O LIVELLI MODULAZIONE	STANDARD DI ADOZIONE	SIMBOLO
16-QAM	16	V.22BIS	Quadribit, formato da due dibit. Il dibit più significativo indica il quadrante, quello meno significativo il punto di modulazione all'interno del quadrante.
16-QAM	16	V.29	Quadribit, formato da tre bit che identificano la fase e un bit che indica l'ampiezza.
32-QAM Trellis o TCM	32	V.32	A cinque bit , di cui uno con funzioni di rilevazione dell'errore
128-QAM Trellis o TCM	128	V.32BIS	A sette bit: un sixbit più un bit di ridondanza con funzioni di rilevazione dell'errore

DEMODULAZIONI

SEGNALE	DEMODULATORE
ASK-OOK	Demodulatore di inviluppo + filtro passabasso Oppure: moltiplicatore (modulatore bilanciato ad anello a 4 diodi), con oscillatore che rigenera la portante + filtro passabasso + rivelatore di soglia
ASK	Demodulatore di inviluppo + filtro passabasso
2-FSK	PLL
2-PSK	moltiplicatore (modulatore bilanciato ad anello a 4 diodi), con oscillatore che rigenera la portante + filtro passabasso
2-DPSK	Non c'è bisogno della rigenerazione della portante. E' sufficiente moltiplicare il segnale modulato ricevuto, per una sua versione ritardata di un T_{bit} e filtrare con un passabasso. Il risultato rivela, con un segno positivo la presenza di un bit 1 e con un segno negativo la presenza di un bit 0. I dispositivi necessari sono quindi un circuito di ritardo, un moltiplicatore e un filtro passabasso.
QAM	Demodulatore IQ con: oscillatore di ricostruzione della portante, sfasatore di -45° di una delle due portanti, filtri passabasso, equalizzatore, estraattore di clock, convertitori A/D, decodifica ed elaborazione inversa in banda base, convertitore parallelo-seriale

BANDA DI UN SEGNALE CON MODULANTE DIGITALE E PORTANTE SINUSOIDALE

Da un punto di vista teorico, lo spettro dei segnali in esame è formato da un numero infinito di repliche (di ampiezza via via minore) di uno spettro base. Questo perché il segnale modulato è il risultato di un prodotto nel tempo fra una portante sinusoidale e l'onda modulante di tipo rettangolare. Tuttavia, tenendo conto del fatto che l'energia del segnale si distribuisce nella parte più bassa dello spettro, possiamo considerare non illimitata la banda di questi segnali e darle un valore finito.

La mera applicazione del “*criterio di Nyquist sulla relazione tra velocità e banda di un canale non rumoroso*” e cioè la relazione:

$$B_M \geq \frac{V_{MOD}}{n} = \frac{V_{MOD}}{\log_2 M}$$

non ci fornisce l'intera banda del segnale complessivo, ma ci dice soltanto che , al crescere del numero n di bit per simbolo (e quindi al crescere del numero $M=2^n$ dei livelli di modulazione), tale banda si restringe.

Per la valutazione della banda è necessario:

1. valutare quante repliche dello spettro del segnale sono necessarie per ricostruire, a partire da esse, il segnale di partenza nel tempo (le estensioni delle repliche necessarie, espresse in termini di velocità di trasmissione v_{TX} , o di frequenza di cifra F_c , e quindi in termini di bit/s) sono riportate nella tabella seguente:

SEGNALE	BANDA
2-ASK con modulante casuale (onda di tipo rettangolare non periodica)	$2 \cdot v_{TX} \rightarrow 2 \cdot F_C$
2-ASK con modulante a onda rettangolare periodica	$3 \cdot v_{TX} \rightarrow 3 \cdot F_C$
2-PSK con modulante casuale (onda di tipo rettangolare non periodica)	$2 \cdot v_{TX} \rightarrow 2 \cdot F_C$
2-PSK con modulante a onda rettangolare periodica	$3 \cdot v_{TX} \rightarrow 3 \cdot F_C$
2-FSK	$4 \cdot v_{TX} \rightarrow 4 \cdot F_C$
2-MSK	$(3/2) \cdot v_{TX}$

2. Se il segnale non è multilivello ($n=1$, $M=2$) la sua banda è quella fornita, come v_{TX} , o come F_c , dalla tabella. Se invece è multilivello, se, cioè, è caratterizzato da un numero n di bit per simbolo maggiore di 1 (e da un numero M di livelli di modulazione maggiore di 2) **si divide per “n”** il risultato ottenuto dalla tabella.

Se per esempio dobbiamo valutare la banda di un segnale 4-PSK modulato con modulante casuale (onda di tipo rettangolare non periodica), procediamo in questo modo.

Dalla tabella vediamo che l'estensione delle repliche spettrali da considerare, per la corrispondente modulazione non multilivello, cioè per la 2-PSK, è $2v_{TX}$.

Poi, considerando che la 4-PSK è una modulazione multilivello con $n=2$ ed $M=4$, dividiamo per $n=2$ il dato fornito dalla tabella ottenendo:

$$B_{4-PSK} = \frac{2v_{TX}}{n} = \frac{2v_{TX}}{2} = v_{TX}$$

CAPITOLO 30

RADIORICEVITORE ETERODINA

COMPLEMENTI

RISONANZA E SINTONIA. IL CIRCUITO LC PARALLELO

CIRCUITO OSCILLANTE LC PARALLELO

ANALISI DELLA CONDIZIONE DI RISONANZA

RILIEVO DELLA RISPOSTA IN FREQUENZA

IL CIRCUITO OSCILLANTE LC PARALLELO COME CIRCUITO DI SINTONIA

RISONANZA E SINTONIA. IL CIRCUITO LC PARALLELO

COS'E' LA RISONANZA IN GENERALE

Ogni corpo ha la sua **frequenza “propria” di oscillazione**, (o **frequenza di risonanza**) determinata dalle dimensioni del corpo, dalla distanza dei punti ai quali è vincolato, dai suoi parametri elettrici e magnetici ecc.

Quando sul corpo agisce una **sollecitazione periodica con frequenza uguale o multipla alla frequenza propria del corpo stesso**, allora il corpo risponde **con oscillazioni di ampiezza molto più grande** di quelle che gli sono state applicate.

La risonanza è quindi un **fenomeno di esaltazione delle oscillazioni** prodotte da un sistema quando questo è sottoposto a una sollecitazione periodica con frequenza uguale o multipla della frequenza propria del corpo.

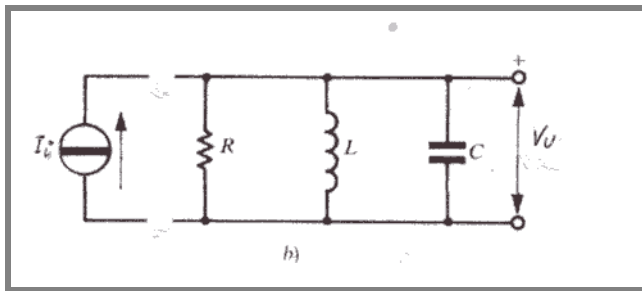
La risonanza è utilmente sfruttata negli ambiti elettrico e acustico, mentre ne l'ambito meccanico e delle costruzioni può risultare pericoloso.

In un circuito elettrico oscillante a induttanza e capacità l'energia elettrostatica viene continuamente convertita in energia magnetica e viceversa

NOTA

Osserviamo d'altra parte che in un circuito LC (e in generale in un sistema del secondo ordine con una coppia di poli complessi coniugati, sistema sottosmorzato, ossia con coefficiente di smorzamento minore di 1) si hanno oscillazioni anche nella risposta a un ingresso non sinusoidale, come l'ingresso a gradino.

CIRCUITO OSCILLANTE LC PARALLELO



E' formato dal parallelo di una capacità e di un'induttanza.

Nella rappresentazione circuitale viene aggiunta una resistenza R_p in parallelo che schematizza le perdite del circuito, ossia il fatto che il condensatore e l'induttore non sono in realtà puramente capacitivi o induttivi, ma presentano una resistenza parassita che determina dissipazione di potenza in calore per effetto Joule.

La resistenza di perdita è schematizzata in parallelo, perciò:

- ❖ se il suo valore è basso significa che le perdite sono rilevanti
- ❖ se il suo valore è alto significa che le perdite sono trascurabili.

RISONANZA PARALLELO

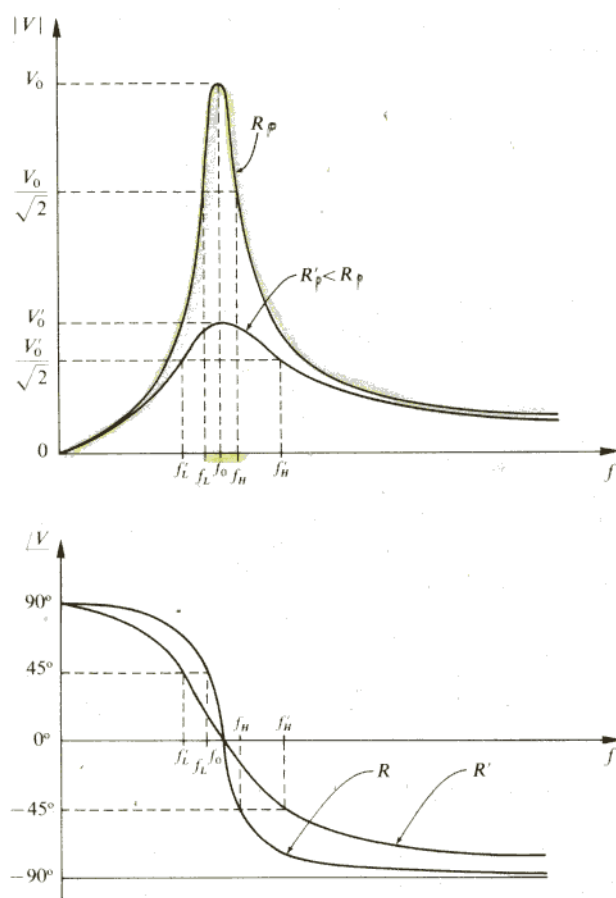
Nel parallelo di un induttore¹ e di un condensatore le correnti in C e in L possono diventare notevolmente più grandi della corrente di ingresso (sovracorrenti)

Alla frequenza di risonanza

- ❖ l'impedenza diventa **massima** e puramente **resistiva**
- ❖ la tensione di uscita ai capi del parallelo, essendo data da $V_u = I_i |Z|$ risulta **massima**.

¹ Ricordiamo che il valore L di induttanza dell'induttore è

- ❖ direttamente proporzionale al numero di spire e al diametro dell'avvolgimento
- ❖ inversamente proporzionale allo spessore del filo



Tensione di uscita (sul parallelo) in funzione della frequenza

ANALISI DEL CIRCUITO

L'impedenza del circuito, calcolata come inverso dell'ammettenza è:

$$Z(\omega) = \frac{1}{Y(\omega)} = \frac{1}{\frac{1}{R_p} + j\left(\omega \cdot C - \frac{1}{\omega \cdot L}\right)}$$

Al denominatore R_p è fissa e non può azzerarsi, mentre il contenuto della parentesi a denominatore dipende dalla pulsazione e quindi dalla frequenza, per cui può azzerarsi.

IL MODULO dell'impedenza è:

$$Z(\omega) = \frac{1}{Y(\omega)} = \frac{1}{\sqrt{\left(\frac{1}{R_p}\right)^2 + \left(\omega \cdot C - \frac{1}{\omega \cdot L}\right)^2}}$$

L'**impedenza** risulta **massima** (in modulo) per il valore di ω , e quindi di f , per il quale la parentesi a denominatore:

$$\left(\omega \cdot C - \frac{1}{\omega \cdot L}\right)^2$$

diventa zero.

Se imponiamo: $\left(\omega \cdot C - \frac{1}{\omega \cdot L}\right) = 0$

Il contenuto della parentesi risulta nullo per:

$$\omega = \frac{1}{\sqrt{L \cdot C}} = \omega_0 \quad \text{pulsazione di risonanza}$$

e quindi per

$$f = \frac{1}{2\pi\sqrt{L \cdot C}} = f_0 \quad \text{frequenza di risonanza}$$

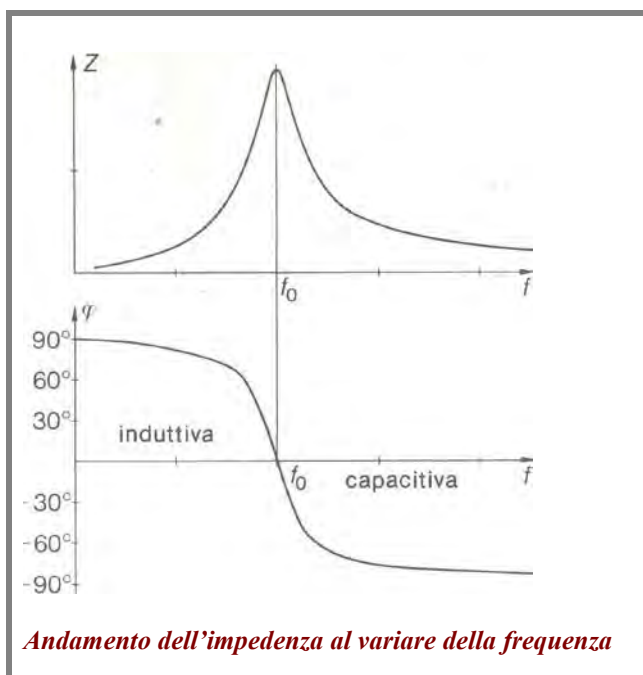
In corrispondenza della **frequenza di risonanza** l'impedenza è **massima e puramente resistiva** e vale R_p .

Per tutte le frequenze diverse da quella di risonanza l'impedenza è, in modulo, più bassa perché, al suo denominatore si aggiunge il contenuto della parentesi, comunque positivo in quanto reale elevato al quadrato.

Di conseguenza la tensione di uscita ai capi del parallelo, essendo data da $V_u = I_i \cdot |Z|$, risulta massima.

In sintesi, alla frequenza di risonanza:

- ❖ il circuito si riduce alla sola resistenza parallelo R_p che rappresenta il suo massimo valore di impedenza
- ❖ la tensione di uscita sul parallelo è massima e vale $R_p \cdot I_i$, in quanto il generatore vede come carico la sola resistenza R_p .



Andamento dell'impedenza al variare della frequenza

ANALISI DELLA CONDIZIONE DI RISONANZA (SOVRACORRENTI E COEFFICIENTE “Q” DI RISONANZA)

Abbiamo visto che, alla frequenza di risonanza, è:

$$\left(\omega \cdot C - \frac{1}{\omega \cdot L} \right) = 0$$

e cioè:

$$\omega \cdot C = \frac{1}{\omega \cdot L}$$

per cui si ha:

$$Z_L = j \omega_0 L = j \frac{1}{\omega_0 C} = -Z_C$$

ossia:

$$Z_L = -Z_C \quad (*)$$

che in modulo si scrive: $|Z_L| = |Z_C|$

Abbiamo anche appurato che, in risonanza,:

$$Z = R_p$$

$$V_u = R_p \cdot I_i$$

Di conseguenza possiamo scrivere :

$$I_L = \frac{V_u}{Z_{LO}} = \frac{R_p \cdot I_i}{Z_{LO}} = \frac{R_p}{Z_{LO}} \cdot I_i$$

$$I_C = \frac{V_u}{Z_{CO}} = \text{per la } (*) = -\frac{R_p \cdot I_i}{Z_{LO}} = \frac{R_p}{Z_{LO}} \cdot I_i = -I_L$$

dalle due ultime espressioni si vede che:

$$I_C = -I_L$$

e cioè che le correnti sull'induttore e sul condensatore sono, in risonanza:

- ❖ uguali e opposte tra loro (di uguale ampiezza e sfasate di 180° una rispetto all'altra)
- ❖ Q volte più grandi della corrente I_i entrante nel parallelo, dove “Q” è il coefficiente di risonanza, definito di seguito.

Q è il coefficiente di risonanza del circuito e vale

$$Q = \frac{R_p}{Z_L} = \frac{R_p}{Z_C} = \frac{R_p}{\omega_0 L} = \omega_0 \cdot C \cdot R$$

Q ha un valore tanto più elevato quanto minori sono le perdite del circuito, schematizzate dalla R_p (cioè quanto più grande è la R_p).

Riguardo alle correnti del circuito oscillante parallelo è verificata, per ogni frequenza, la somma vettoriale:

$$I_R + I_C + I_L = I_i$$

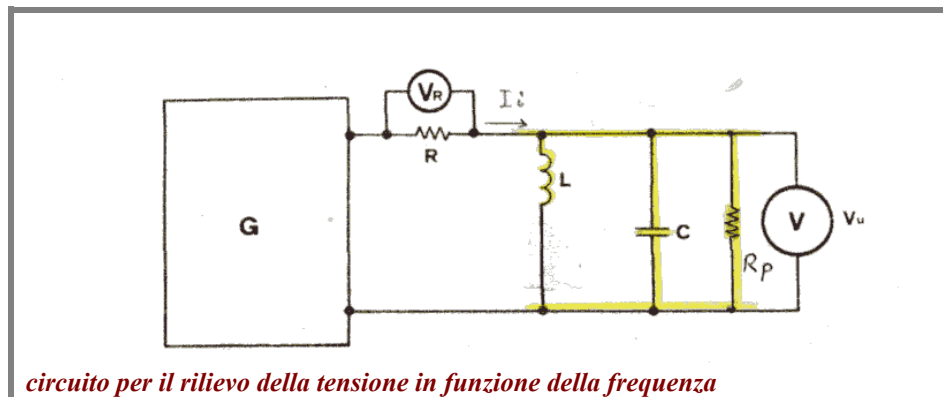
In RISONANZA, siccome I_C e I_L sono uguali e opposte, si ha $I_i = I_R$.

RILIEVO DELLA RISPOSTA IN FREQUENZA (TENSIONE DI USCITA IN FUNZIONE DELLA FREQUENZA)

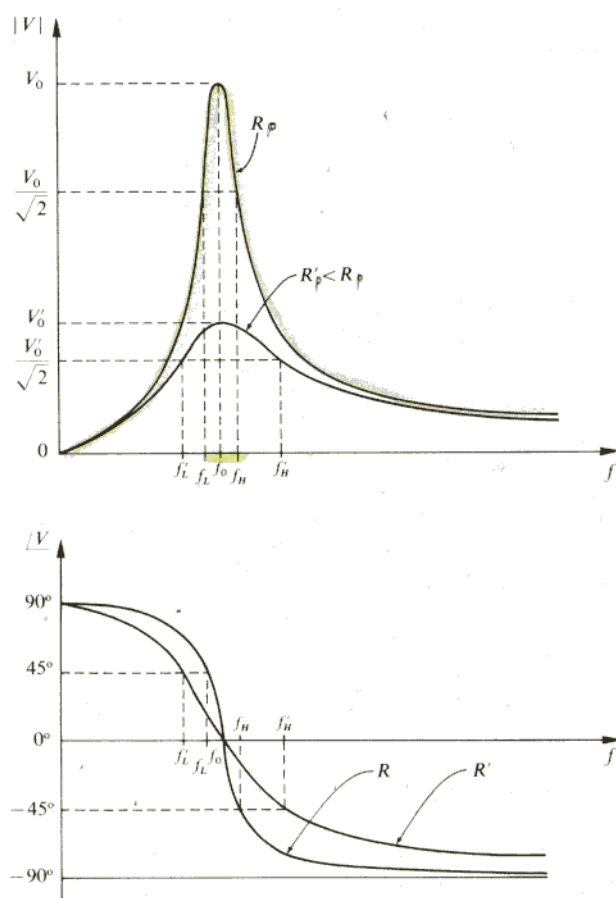
Per effettuare il rilievo della tensione di uscita in funzione della frequenza, il circuito risonante parallelo deve essere alimentato da un generatore ideale di corrente, cioè da un generatore di tensione con resistenza interna R_G molto più grande del modulo dell'impedenza del carico.

Carico che in questo caso è proprio il parallelo $R_p // L // C$.

La R_G viene rappresentata in serie al generatore di tensione.



Se facciamo variare la frequenza della corrente I_0 entrante nel parallelo, mantenendone costante l'ampiezza (e mantenendo quindi costante anche l'ampiezza della tensione V_R sulla resistenza interna R_G del generatore di corrente) otteniamo una risposta in frequenza del tipo illustrato nella figura seguente.



Tensione di uscita (sul parallelo) in funzione della frequenza

IL CIRCUITO OSCILLANTE LC PARALLELO COME CIRCUITO DI SINTONIA

Il fatto che a una particolare frequenza (frequenza di risonanza) la tensione ai capi del circuito risonante parallelo sia massima, mentre risulta più bassa per tutte le frequenze minori e maggiori di f_0 , rende adatto il **circuito oscillante parallelo** a selezionare lo spettro del segnale della stazione emittente desiderata da tutti gli altri spettri presenti all'antenna del ricevitore radio.

In altri termini l'azione filtrante **passabanda**² del circuito oscillante $R_p//L//C$ lo rende adatto a realizzare il **circuito di sintonia del radioricevitore**.

Il ricevitore però deve offrire la possibilità di scegliere la stazione da ascoltare, perciò affinché il circuito oscillante parallelo funzioni da circuito di sintonia, deve darci la possibilità di **variare la**

frequenza di risonanza $f_0 = \left(\frac{1}{2\pi \cdot \sqrt{LC}} \right)$ e quindi la **posizione della banda passante sull'asse delle frequenze**.

Questa variazione è resa possibile adottando, invece di un condensatore fisso, una **capacità variabile**, comandata dalla manopola o dal tasto di sintonia.

² Il circuito LC serie invece, presentando alla frequenza di risonanza il valore più basso di impedenza, è adatto a provocare un abbassamento di tensione nella banda di frequenze a cavallo della f_0 , per cui ha un comportamento eliminabanda

IL CIRCUITO OSCILLANTE LC PARALLELO COME CIRCUITO DI SINTONIA

Il fatto che a una particolare frequenza (frequenza di risonanza) la tensione ai capi del circuito risonante parallelo sia massima, mentre risulta più bassa per tutte le frequenze minori e maggiori di f_0 , rende adatto il **circuito oscillante parallelo** a selezionare lo spettro del segnale della stazione emittente desiderata da tutti gli altri spettri presenti all'antenna del ricevitore radio. In altri termini l'azione filtrante **passabanda**³ del circuito oscillante $R_p//L//C$ lo rende adatto a realizzare il **circuito di sintonia del radioricevitore**.

Il ricevitore però deve offrire la possibilità di scegliere la stazione da ascoltare, perciò affinché il circuito oscillante parallelo funzioni da circuito di sintonia, deve darci la possibilità di **variare la frequenza di risonanza**

$$f_0 = \left(\frac{1}{2\pi \cdot \sqrt{LC}} \right)$$

e quindi la **posizione della banda passante sull'asse delle frequenze**.

Questa variazione è resa possibile adottando, invece di un condensatore fisso, una **capacità variabile**, comandata dalla manopola o dal tasto di sintonia.

³ Il circuito LC serie invece, presentando alla frequenza di risonanza il valore più basso di impedenza, è adatto a provocare un abbassamento di tensione nella banda di frequenze a cavallo della f_0 , per cui ha un comportamento eliminabanda

CAPITOLO 31

RADAR

***DEFLESSIONE DEL PENNELLO ELETTRONICO
E PORTATA DI BASE (MASSIMA DISTANZA NON AMBIGUA)***

VARIAZIONE DEL PERIODO

***COMANDO DI “DURATA DELL’IMPULSO”
(REGOLAZIONE DELLA PORTATA DI BASE)***

SCALE DI DISTANZA

RADAR AERONAUTICI

RADAR DOPPLER

UTILIZZAZIONE DI UN RADAR IMBARCATO

RADAR AERONAUTICI DI TERRA

***SISTEMI RADAR PER LA MISURAZIONE DELLA QUOTA (RADAR ALTIMETRO E
RADIO ALTIMETRO)***

***PRINCIPIO DI FUNZIONAMENTO DEL RADIO ALTIMETRO E DEI RADAR DOPPLER
CW-FM***

RADAR ALTIMETRO: CONFRONTO CON IL RADIO ALTIMETRO

ALTRI TIPI DI RADAR

PREMESSA: L’EFFETTO DOPPLER

RADAR DOPPLER A ONDA CONTINUA

RADAR DOPPLE: ESPRESSIONE DELLO SCOSTAMENTO ΔF (BATTIMENTO)

RADAR DOPPLER: VELOCITA’ RELATIVA E VELOCITA’ RISPETTO AL SUOLO

RADAR FM-CW

RADAR IMPULSIVO CON INFORMAZIONE DOPPLER O RADAR MTI

RADAR METEO

DEFLESSIONE DEL PENNELLO ELETTRONICO E PORTATA DI BASE (MASSIMA DISTANZA NON AMBIGUA¹)

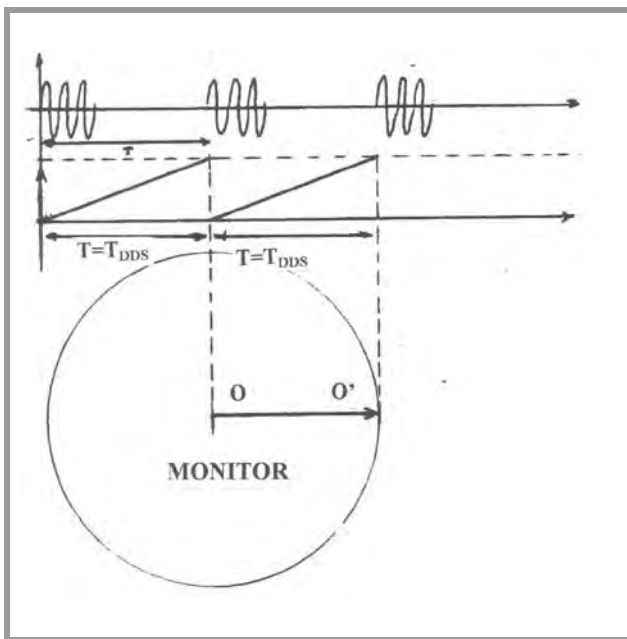
Nel tempo T (cioè nel **periodo PRI** di **ripetizione** degli impulsi) intercorrente fra la partenza di un pacchetto di impulsi e la partenza del pacchetto successivo, il pacchetto di impulsi **raggiunge il bersaglio e torna indietro**.

T è quindi il tempo di andata e ritorno dell'impulso.

Di conseguenza, essendo " c " la velocità delle onde e.m. nel vuoto,

il prodotto Tc è la **lunghezza del percorso di andata e ritorno del pacchetto**.

D'altra parte il tempo T è, nel tubo catodico, il **tempo impiegato dal fascio o pennello elettronico a percorrere l'intero raggio OO' del monitor** cioè per passare dal centro O del monitor al punto O' del bordo esterno.



Il tragitto OO' del pennello, descritto nel tempo T , corrisponde al percorso Tc di andata e ritorno del pacchetto di impulsi sinusoidali.

Se il raggio OO' del monitor (descritto dal pennello nel tempo T) corrisponde al percorso di andata e ritorno del pacchetto, ciò implica che, ad OO' e al tempo T , corrisponde un percorso di "sola andata" radar-bersaglio (cioè una distanza radar bersaglio) di lunghezza $\frac{1}{2} Tc$.

Se quindi stabiliamo una **corrispondenza** fra **percorsi di sola andata** e **raggi** sul monitor, possiamo dire che **ad OO' e al tempo T , corrisponde un percorso di "sola andata" radar-bersaglio (cioè una distanza radar-bersaglio) di lunghezza $\frac{1}{2} Tc$.**

Dunque un **punto luminoso** che si trovi sul **marginale esterno dello schermo**, cioè sulla circonferenza di raggio OO' , rappresenta un **bersaglio che si trovi alla distanza $\frac{1}{2} Tc$** dall'antenna radar.

Questa distanza $\frac{1}{2} Tc$ viene chiamata **portata di base del radar** (o "**massima distanza non ambigua**") e rappresenta quindi:

¹ Distanza al di là della quale i bersagli appaiono come echi di seconda traccia

- la distanza di un bersaglio la cui visualizzazione si trovi sul margine esterno del monitor, ossia “a fondo scala”
- la metà della distanza complessiva T_c di andata e ritorno percorsa da un pacchetto impulsivo prima che parta il successivo pacchetto.

VARIAZIONE DEL PERIODO

PREMESSA

la lunghezza del raggio sul monitor dipende solo dall'ampiezza V_{dds} e non dalla durata T_{dds} del dente di sega: se manteniamo fissa l'ampiezza del dente di sega a un valore di tensione (chiamiamolo $V_{MAX_{dds}}$) tale da consentire la descrizione del raggio massimo OO', allora il raggio descritto sarà quello massimo qualunque sia la durata T_{dds} del dente di sega. Ciò che cambierà sarà solo la velocità con la quale il pennello traccia il raggio.

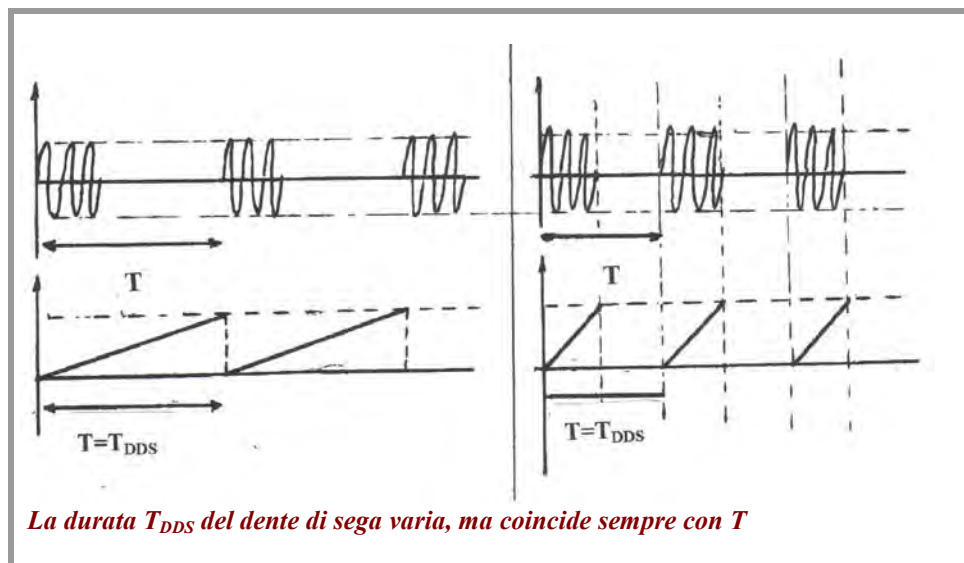
La variazione della durata T_{dds} del dente di sega può essere fatta in due modi:

1. facendo variare T_{dds} simultaneamente con il periodo T di ripetizione degli impulsi, cioè mantenendo T_{dds} sempre uguale a T
2. facendo variare T_{dds} indipendentemente da T e più esattamente scegliendo valori di T_{dds} minori di T

Esamineremo i due casi **supponendo, in entrambi, di mantenere costante sul valore massimo $V_{MAX_{dds}}$, l'ampiezza del dente di sega.** Così, in qualunque caso, **il raggio descritto dal pennello sul monitor sarà sempre quello massimo.**

Sarà cioè il raggio OO' che va dal centro alla circonferenza estrema del monitor.

1. VARIAZIONE SIMULTANEA DI T_{DDS} E DI T



Quando facciamo variare assieme la durata T_{DDS} dei denti di sega e il periodo T della sequenza di pacchetti, il raggio OO' coincide sempre con la massima distanza $\frac{1}{2}Tc$ esplorabile realmente dal singolo pacchetto prima che parta il successivo.

OO' coincide quindi con la portata reale.

In questa situazione, cambiare T equivale a cambiare simultaneamente sia la massima distanza esplorabile dal singolo pacchetto sia il valore, in Km o in miglia, del raggio OO' (portata di base). Cosa significa questo?

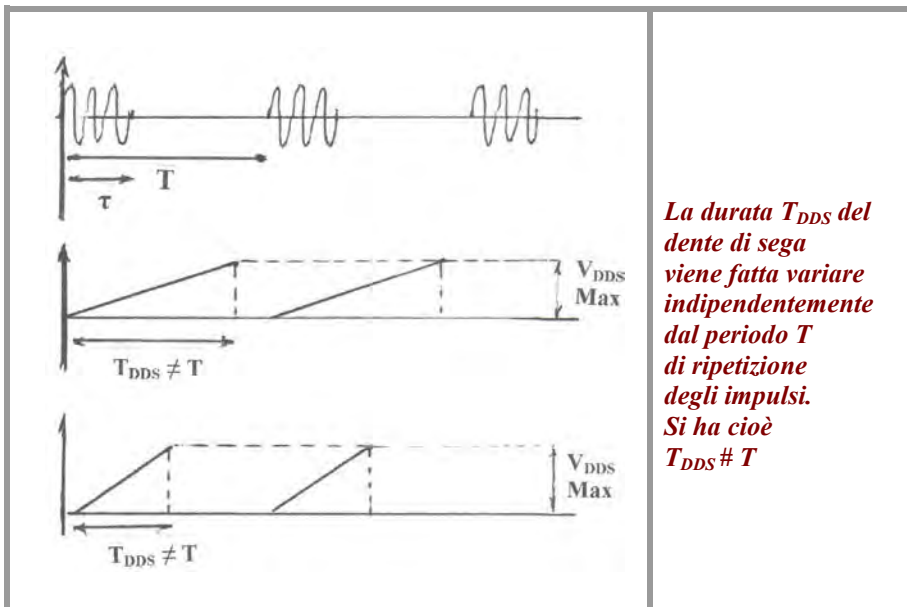
Il rapporto:

distanza realmente esplorata dal pacchetto / valore in Km della lunghezza del raggio OO'

rimane **costante** al variare di T in quanto la distanza realmente esplorata dal pacchetto e il valore (in Km o in miglia) della lunghezza del raggio variano nella stessa misura.

Il fatto che il rapporto rimane costante significa che **varia la portata** del radar, ma non la scala di rappresentazione che rimane costante.

2. VARIAZIONE DI T_{DDS} CON T FISSO



Se, lasciando inalterato il periodo T della sequenza di pacchetti, facciamo variare (diminuire) la sola durata T_{dds} del dente di sega, noi lasciamo inalterata la massima distanza esplorata realmente dal pacchetto. Di questa distanza però visualizziamo, sempre sull'intero raggio OO' del display, solo una porzione. Sfruttiamo cioè solo una parte dell'esplorazione del pacchetto.

Questo significa che in sostanza facciamo diminuire il valore in Km o in miglia del raggio OO'

Se per esempio abbiamo scelto $T_{DDS} = 0,5T$, sul display possiamo vedere solo i bersagli che sono compresi entro la metà della distanza realmente esplorata, cioè entro la metà di $\frac{1}{2} \cdot T_c$.

Cosa comporta tutto questo?

Siccome la massima distanza esplorata realmente dal pacchetto rimane fissa, mentre il valore in Km o in miglia del raggio OO' diminuisce, il rapporto

distanza realmente esplorata dal pacchetto / valore in Km della lunghezza del raggio OO' aumenta.

Ciò significa che (al diminuire di T_{dds}):

- la **portata** reale **diminuisce**
- la scala di rappresentazione diventa più grande, cioè si ha, sul display, la raffigurazione più dettagliata di una zona più ristretta.

Il fascio elettronico, nel corso del tempo T_{DDS} (minore di T) percorre velocemente l'intero raggio OO' del monitor, ma questo raggio rappresenta solo una porzione dalla massima distanza esplorata dal pacchetto.

Per la frazione di T eccedente T_{DDS} , il pennello è inattivo.

COMANDO DI “DURATA DELL’IMPULSO” (REGOLAZIONE DELLA PORTATA DI BASE)

Normalmente i radar sono dotati di un comando “durata dell’impulso”.

Questo comando permette di allungare o accorciare **simultaneamente** il periodo T di ripetizione dei pacchetti e la durata τ del singolo pacchetto e , di conseguenza, la portata di base².

Possono esistere più possibilità di regolazione:

1. solamente fra pacchetto (τ) corto e pacchetto lungo (e quindi, data la simultaneità di regolazione, fra periodo T corto e periodo T lungo)
2. fra pacchetto (τ) corto, medio o lungo (e quindi, data la simultaneità di regolazione, fra periodo T corto, medio o lungo)

Riguardo al caso 1 possiamo dire che:

- ❖ l’impulso τ lungo caratterizzato da grandi valori di τ e di potenza istantanea P_i , è associato a un intervallo T di ripetizione di 1ms e serve per la rilevazione dei bersagli lontani, oltre le 24 miglia (circa)
- ❖ L’impulso τ corto, caratterizzato da piccoli valori di τ e di potenza istantanea P_i presenta un intervallo T di ripetizione più piccolo e serve per la rilevazione dei bersagli vicini, dalle 6 alle 12 miglia (circa).

² (Infatti, aumentare T lasciando piccola la durata τ potrebbe portarci in una condizione in cui il pacchetto, pur avendo il tempo sufficiente per il raggiungimento di un bersaglio lontano non riesca a farlo perché la potenza media associata a un τ piccolo risulterebbe troppo bassa)

SCALE DI DISTANZA

La regolazione della durata del dente di sega (con T e τ fissi), ossia la variazione della SCALA di rappresentazione senza variazione della durata degli impulsi, viene effettuata da un comando (diverso da quello della portata di base) che è chiamato “comando di RANGE”.

Questo comando permette, per esempio, di selezionare le seguenti scale:

$3/8$, $3/4$, $3/2$, 3, 6, 12, 24, 48, 36

RADAR AERONAUTICI

RADAR DOPPLER

Vedi anche: altri tipi di radar

In generale il radar Doppler è un radar che:

- **Rileva** la presenza di **BERSAGLI IN MOVIMENTO**
- **DETERMINA la VELOCITA'** radiale di avvicinamento e di allontanamento **del bersaglio** (dall'antenna del trasmettitore radar).

Il radar Doppler deve essere provvisto di dispositivi in grado di **misurare** la **differenza Δf** tra la frequenza del segnale riflesso e la frequenza (f_{sorg}) del segnale emesso dall'antenna radar.

I radar Doppler sono usati come:

- radar per **difesa e combattimento aereo** (scoperta e inseguimento di bersagli mobili)
- radar per il **controllo del traffico aereo**
- radar **meteorologici**
- radar per la **rilevazione di velocità dei veicoli** (autovelox)
- rivelatori **antiintrusione** per sistemi di sicurezza ambientale

A bordo degli aerei i radar Doppler sono anche usati per **misurare la velocità dell'aeromobile** stesso.

DIFFERENZE CON IL “NORMALE” RADAR

A differenza del radar “normale” il radar Doppler :

- rileva la **velocità** del bersaglio
- è basato sulla misura di uno **scostamento di frequenza Δf** (e non di un intervallo di tempo)
- è dotato di un **mixer** per la determinazione di Δf e di un frequenzimetro per la sua misura

È caratterizzato da una **minore emissione di potenza**, il che lo rende meno individuabile in caso di operazioni belliche.

UTILIZZAZIONE DI UN RADAR IMBARCATO

In generale possiamo dire che il compito del radar è di scoprire un bersaglio e misurarne distanza, posizione e angolo e velocità relativa.

Volendo tracciare un quadro delle funzioni svolte dal radar, possiamo specificare le sue funzioni nel modo seguente.

AUSILIO ALLA NAVIGAZIONE

- **TERRAIN FOLLOWING** (“inseguimento del profilo del suolo”). Questa funzione del radar di bordo consiste nel rilevare il profilo del terreno muovendo il fascio d’antenna lungo il piano verticale. Il profilo ottenuto viene fornito a un elaboratore che ha il compito di pilotare il velivolo in modo che questo segua l’andamento del suolo a quota costante rispetto ad esso.
- **TERRAIN AVOIDANCE**³ (“aggiramento di ostacoli al suolo”). Questa funzione del radar di bordo consiste nel rilevare non soltanto il profilo altimetrico del terreno, ma anche eventuali ostacoli presenti sul suolo. Pertanto il fascio d’antenna deve essere mosso sia lungo il piano verticale che lungo il piano orizzontale. A prescindere da questa variante la funzione di terrain avoidance è simile alla precedente.
- **FIX DI POSIZIONE** (“correzione o aggiornamento della posizione del velivolo già calcolata dalla piattaforma inerziale o da altri dispositivi”). Il fascio d’antenna viene indirizzato verso un punto di coordinate note e il radar fornisce la distanza dal punto e l’angolo. L’elaboratore di navigazione determina la posizione del velivolo, la confronta con quella fornita dai sistemi di navigazione ed eventualmente la corregge.

GROUND MAPPING (“tracciamento di mappe del terreno”)

³ Fuga, scampo, annullamento, l’evitare

WEATHER AVOIDANCE (“aggiramento di perturbazioni atmosferiche”)

Con una opportuna scelta della frequenza di lavoro, un segnale radar può oltrepassare le nuvole e rilevare la presenza di gocce di pioggia e chicchi di grandine.

L’eco che perviene al ricevitore sarà tanto più intensa quanto più grandi saranno i bersagli.

Siccome poi, generalmente, la forza della precipitazione è proporzionale alla grandezza delle gocce, il radar fornirà un segnale di intensità proporzionale all’intensità della precipitazione.

RADAR AERONAUTICI DI TERRA

1) RADAR PRIMARI DI SORVEGLIANZA (PSR)

Sono normali radar a impulsi.

Non forniscono informazioni sull'**identità** degli aerei.

Sono formati da due componenti essenziali: il trasmettitore e il ricevitore.

Tutte le loro componenti sono installate nello **stesso sito**: per esempio in un aeroporto

- RADAR PRIMARI TERMINALI DI SORVEGLIANZA

Sono utilizzati per coprire le **zone attorno agli aeroporti**.

Possono avere una portata di **60** miglia

- RADAR PRIMARI DI SORVEGLIANZA DI ROTTA

Sono utilizzati per la copertura di **vaste aree** al fine di controllare gli aeromobili in rotta (cioè **in volo**).

Possono avere una portata di **100** miglia.

-

2) RADAR SECONDARI (SSR)

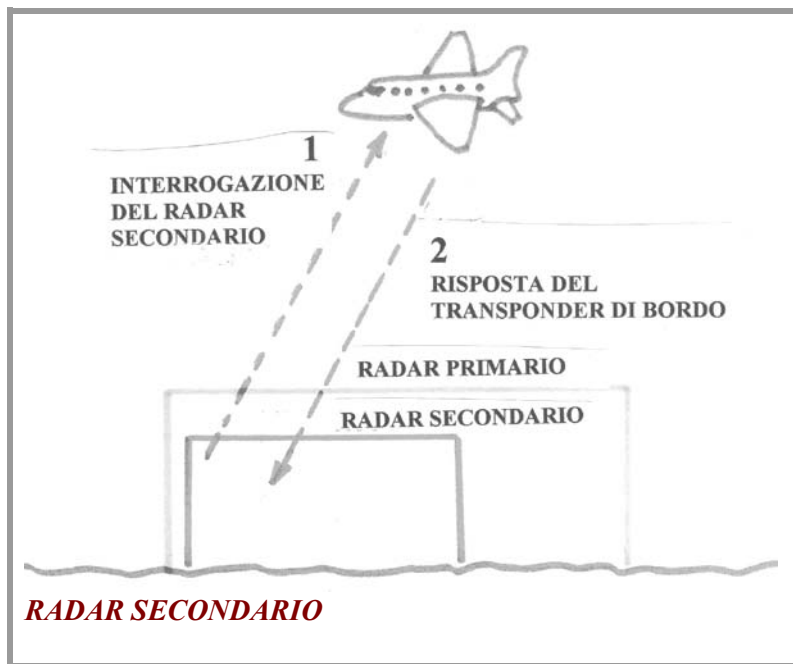
Forniscono **informazioni** sull'**identità** dell'aeromobile.

Sono formati da tre parti:

- INTERROGATORE (di terra)
- RISPONDITORE di BORDO o TRANSPONDER
- RICEVITORE O DECODIFICATORE di terra

Le componenti del radar secondario sono installate in due **siti diversi** (in aeroporto l'interrogatore e il ricevitore o decodificatore; a bordo dell'aeromobile il transponder).

Tipicamente **l'interrogatore di terra** e il **ricevitore di terra** sono **integrati** a un **radar primario**.



SISTEMI RADAR PER LA MISURAZIONE DELLA QUOTA RADAR ALTIMETRO E RADIO ALTIMETRO

A cosa servono il radar altimetro e il radio altimetro se già esiste l'altimetro barometrico?

L'altimetro barometrico è basato sulla misura della pressione atmosferica e perciò non può fornire un'esatta misura della distanza dal suolo.

E' perciò utilizzato per navigare lungo un'isobara (linea a pressione atmosferica costante), cioè per la cosiddetta "navigazione in aerovia".

Radio altimetro e radar altimetro invece sono in grado di fornire la quota effettiva rispetto al suolo e perciò vengono usati per il volo a bassa quota, nel corso dell'atterraggio, e per i missili intercontinentali.

Spesso i termini "RADIO altimetro" e "RADAR altimetro" sono usati come sinonimi.

A rigore però si parla di RADAR altimetro nel caso di un'apparecchiatura basata su segnali IMPULSIVI.

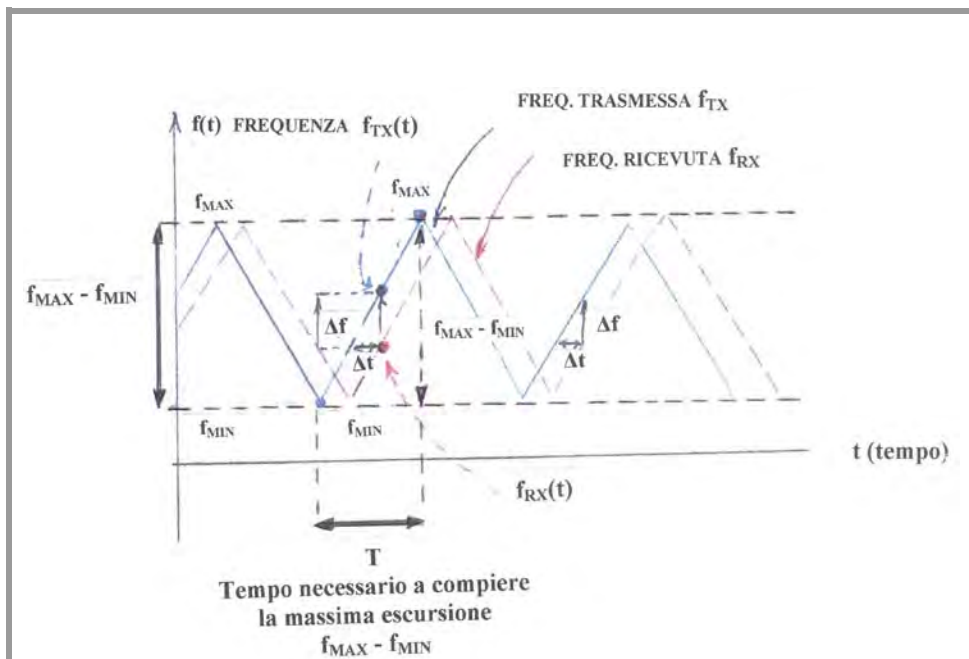
Si parla invece di **RADIO** altimetro nel caso di un'apparecchiatura basata su segnali **A ONDA CONTINUA**, in particolare modulati in frequenza.

I sistemi FM-CW (come il radar-altimetro) misurano la variazione istantanea Δf fra la frequenza emessa e quella ricevuta e, proprio perché sono basati sulla misura di una differenza di frequenze, sono intrinsecamente dei **radar Doppler**.

Questa variazione Δf è direttamente proporzionale all'intervallo di tempo di ritardo, Δt , nel quale il segnale radar colpisce l'obiettivo e poi ritorna.

Grazie alle tecniche di modulazione (per esempio la FM), è possibile misurare la distanza dal bersaglio anche con un radar CW. Il tempo di ritardo fra l'onda ricevuta e l'onda trasmessa può essere determinato attraverso una misura di differenza di frequenze.

Il modo più semplice per modulare l'onda è di variarne la frequenza con legge lineare.



Andamenti temporali della frequenza emessa e della frequenza ricevuta

NOTE

La differenza " $f_{max}-f_{min}$ " non è Δf e rappresenta la massima escursione possibile sia per la frequenza trasmessa che per la frequenza ricevuta.

$f_{TX}(t)$ e $f_{RX}(t)$ sono i valori della frequenza trasmessa e della frequenza ricevuta misurati entrambi nell'istante t .

PRINCIPIO DI FUNZIONAMENTO DEL RADIO ALTIMETRO E DEI RADAR DOPPLER CW-FM

Il segnale trasmesso verso il suolo è caratterizzato da una frequenza f_{TX} (variabile nel tempo) che viene fatta variare temporalmente con andamento triangolare.

Detto Δt il tempo che intercorre fra l'emissione del segnale trasmesso e la ricezione del segnale riflesso dal suolo, possiamo dire che il segnale ricevuto è caratterizzato da una frequenza f_{RX} che ha lo stesso andamento triangolare di f_{TX} , ma che differisce da f_{TX} della quantità Δf .

Da un punto di vista grafico l'andamento di f_{RX} è (in ogni istante t) traslato verticalmente (verso l'alto o verso il basso) della quantità Δf rispetto all'andamento di f_{TX} :

$$f_{TX}(t) - f_{RX}(t) = \Delta f \quad (1)$$

La *distanza* D è data da
$$D = \frac{1}{2} \cdot C \cdot \Delta t$$

dove⁴ il ritardo Δt ha espressione:
$$\Delta t = \frac{\Delta f}{f_{\max} - f_{\min}} T$$

con:

$f_{\max}$ = *frequenza massima*

$f_{\min}$ = *frequenza minima*

T = *periodo dell'escursione da $f_{\max}$ a $f_{\min}$*

Δf = *variazione tra la frequenza emessa e ricevuta: $f_{TX}(t) - f_{RX}(t)$*

Un'importante caratteristica dei radar CW-FM è che (contrariamente a quelli a impulsi) **non hanno una distanza minima d'impiego**.

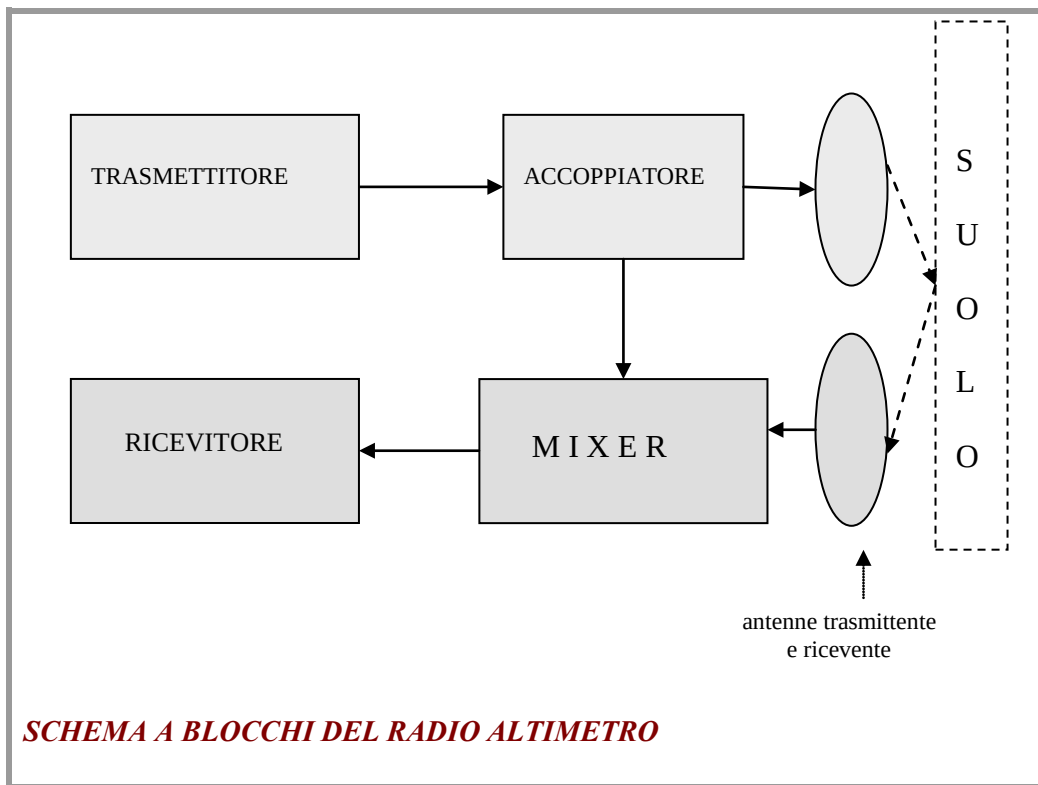
Non sono utilizzati per il rilevamento a lunga distanza, perché la potenza necessaria per la trasmissione sarebbe elevatissima. L'emissione di un'onda continua infatti richiede un maggior dispendio di potenza di quanta non ne richieda un segnale impulsivo (la cui potenza è diversa da zero solo per limitati intervalli di tempo).

Le utilizzazioni fondamentali sono:

- altimetri radar
- sistemi di inseguimento di missili

⁴ L'espressione del ritardo si può ricavare, guardando il grafico, dalla proporzione:

$$(f_{\max} - f_{\min}) : T = \Delta f : \Delta t$$



Nel radio altimetro, un'onda continua (CW, Continuous Wave) viene modulata in frequenza e dalla misura della frequenza "differenza" $f_{TX}(t) - f_{RX}(t)$, e cioè di Δf , si ottiene la quota h .

La frequenza differenza è determinata dal mixer, ai cui due ingressi sono applicate:

- la "frequenza ricevuta" f_{RX} proveniente dall'antenna ricevente
- la frequenza trasmessa f_{TX} proveniente dall'accoppiatore che applica al mixer un campione dell'onda trasmessa, onda caratterizzata appunto dalla frequenza f_{TX} .

RADAR ALTIMETRO: CONFRONTO CON IL RADIO ALTIMETRO

Il principio di funzionamento del radar altimetro è quello del normale radar a impulsi. Perciò non ci soffermeremo su questa apparecchiatura, ma faremo solo un confronto di prestazioni con il radio altimetro.

A quote **inferiori** al cosiddetto valore critico h_c , le **perdite di potenza** del radioaltimetro e del radar altimetro sono **uguali**: aumentano linearmente, al crescere della distanza dal suolo, di 6 dB/ottava (cioè di 6 deciBel per ogni raddoppio della distanza dal suolo).

A quote **superiori** al valore critico invece risulta più **conveniente** l'altimetro a **onda continua** (radio altimetro), in, quanto per $h > h_c$, le perdite di tale radar continuano a crescere di 6 dB/ottava, mentre le perdite dell'altimetro a impulsi, radar altimetro, crescono (sempre per $h > h_c$) di 9 dB/ottava.

Per bilanciare le maggiori perdite i radar altimetri (altimetri impulsivi) sono dotati di antenne con guadagno maggiore (e quindi con diagramma di irradiazione più stretto) rispetto ai sistemi a onda continua.

Comunque i sistemi impulsivi (radar altimetri) sono in grado di funzionare correttamente solo per angoli di pitch e di roll inferiori a 40°.

Precisiamo infine l'espressione della quota critica:

QUOTA CRITICA
$$h_c = \frac{c \cdot \delta \cdot G}{4}$$

con:

$c = 3 \cdot 10^8$ m/s = velocità delle onde e.m. nel vuoto

d = lunghezza dell'impulso

G = guadagno dell'antenna

ALTRI TIPI DI RADAR

RADAR DOPPLER

PREMESSA: L'EFFETTO DOPPLER

L'effetto Doppler è importante sia come strumento di indagine scientifica, sia per le sue numerose applicazioni tecnologiche, che spaziano dalla navigazione, alla difesa, dalla medicina, all'industria civile.

L'applicazione di maggiore interesse per chi opera nei settori nautico e aeronautico è il radar, o meglio alcune tipologie di radar.

I veri e propri radar Doppler sono usati come:

- radar per **difesa e combattimento aereo** (scoperta e inseguimento di bersagli mobili)

- radar per il **controllo del traffico aereo**
- radar **meteorologici**

Oltre che per la realizzazione di radar propriamente detti, l'effetto Doppler è utilizzato ad esempio nel campo dei

- dispositivi per la **rilevazione di velocità dei veicoli** (autovelox)
- rivelatori **antiintrusione** per sistemi di sicurezza ambientale
- misuratori **di velocità di navi e aeromobili**
- sistemi per la **navigazione aerea**
- **trasduttori**

Un esempio di effetto Doppler nel campo delle onde meccaniche lo abbiamo certamente sperimentato nella vita quotidiana: quando una sorgente sonora ci si avvicina, il suo suono appare più acuto (come se avesse una frequenza più alta). Quando invece si allontana, il suono appare più grave (frequenza più bassa).

L'effetto Doppler può essere descritto nel modo seguente.

Siano dati un osservatore e una sorgente d'onda, caratterizzata da un valore f_s di frequenza. Se c'è un **movimento relativo** fra osservatore e sorgente, il valore di frequenza f_{os} percepito dall'osservatore risulta diverso dal valore della frequenza emessa dalla sorgente, e questa **variazione di frequenza** dipende dalla **velocità** con la quale lo spostamento avviene.

In particolare, quando c'è **avvicinamento**, la frequenza percepita dall'osservatore appare più elevata di quella emessa dalla sorgente, per cui lo scostamento $\Delta f = f_{os} - f_s$ risulta **positivo** (e crescente al crescere della velocità).

Invece in caso di allontanamento, la frequenza percepita dall'osservatore appare più bassa di quella emessa dalla sorgente, per cui lo scostamento

$\Delta f = f_{os} - f_s$ risulta negativo (e crescente, in valore assoluto, al crescere della velocità).

Lo scostamento Δf è chiamato “**variazione Doppler**” o “**Doppler shift**”

L'effetto si produce se la sorgente si sposta e l'osservatore è fermo, se la sorgente è ferma e l'osservatore si sposta, se entrambi si spostano.

L'effetto Doppler consiste quindi in una variazione (o scostamento o slittamento) di frequenza che risulta proporzionale alla velocità relativa di avvicinamento o di allontanamento fra sorgente e osservatore.

Il fenomeno interessa sia le onde meccaniche, o onde di pressione (per esempio le onde sonore e le onde ultrasoniche), sia le onde elettromagnetiche (per esempio le onde radio e la luce, visibile o invisibile)

Ricordando che la radiazione luminosa a bassa frequenza tende al colore rosso, mentre quella ad alta frequenza tende al violetto, possiamo dire che, quando una sorgente luminosa si avvicina, a noi tende ad apparire viola (alta frequenza) mentre quando si allontana tende ad apparire rossa.

Perciò il fatto che la luce delle galassie appare rossa, significa che esse si allontanano e quindi che l'universo si espande.

EFFETTO DOPPLER IN MEDICINA

E' utilizzato in varie tecniche di **visualizzazione ecografica**.

ed è basato su onde meccaniche ultrasoniche (2,5 – 10) MHz.

Queste onde ultrasoniche⁵ attraversano facilmente i **tessuti molli** e i **liquidi**,

mentre non possono attraversare le ossa e i gas, pertanto possono essere usati per esaminare l'utero, il fegato e in generale gli organi pieni di liquido, ma non per analizzare il cervello e i polmoni.

Un emettitore (trasduttore piezoelettrico che trasforma correnti e tensioni elettriche in onde sonore) invia impulsi ultrasonici a una determinata frequenza.

Quando questi impulsi rimbalzano contro un oggetto **IN MOVIMENTO**

(per es. il sangue in un vaso o il cuore pulsante di un feto nell'utero),

la frequenza degli echi, proprio per'effetto doppler, è diversa dalla frequenza degli impulsi inviati e la variazione di frequenza è funzione –tra le altre cose- della velocità dell'oggetto in movimento.

Un sensore rileva la differenza di frequenza, e da questo dato vengono ricavate informazioni sulla velocità del sangue o sulla frequenza del battito cardiaco del feto.

⁵ Il campo di frequenze qui indicato comprende valori molto più elevati di quelli propriamente indicati come ultrasuoni

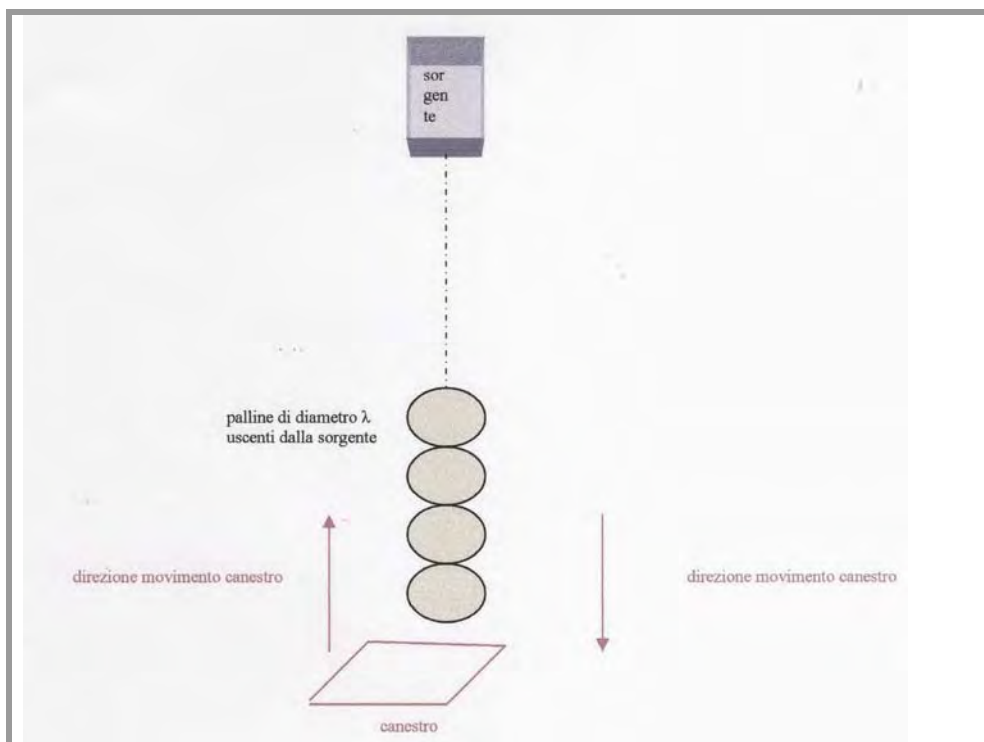
NOTA

Per rendere più evidente l'incremento di frequenza che si produce per effetto doppler quando si riduce la distanza fra la sorgente e la superficie ricevente (e la diminuzione di frequenza che si produce in seguito a un aumento della distanza), possiamo ricorrere al paragone che segue.

Possiamo immaginare che la SORGENTE sia un TUBO che emette, in direzione assiale, PALLINE di diametro λ , in numero di f palline al secondo
dove:

- f può rappresentare la frequenza iniziale dell'onda, quando la distanza sorgente-superficie non varia
- λ può rappresentare la lunghezza d'onda del segnale

Possiamo immaginare poi che la **superficie ricevente** sia superficie di un **canestro** che, o è fermo, o si muove lungo l'asse della sorgente



La superficie ricevente (canestro), se è fissa, viene attraversata da f palline al secondo (per esempio da 1000 palline/sec se $f=1000$)

Se la superficie si avvicina alla sorgente, viene attraversata da un numero di palline al secondo tanto più elevato quanto maggiore è la velocità di avvicinamento e quanto minore è il diametro delle palline.

Se la superficie si allontana dalla sorgente, viene attraversata da un numero di palline al secondo tanto più basso quanto maggiore è la velocità di avvicinamento e quanto minore è il diametro delle palline.

ESPRESSIONE ANALITICA DELL'EFFETTO DOPPLER

Se una sorgente in movimento sta emettendo onde con una frequenza f' , allora un osservatore percepirà le onde con una frequenza f'' data da:

$$f'' = f' \frac{v_p \pm v_{oss}}{v_p \mp v_{sorg}}$$

v_p = velocità di propagazione dell'onda

v_{sorg} = velocità della sorgente

v_{oss} = velocità dell'osservatore

Il segno è positivo al numeratore e negativo al denominatore se sorgente ed osservatore si avvicinano; in questo modo infatti se aumentano la velocità v_{sorg} della sorgente o la velocità v_{oss} dell'osservatore o entrambe, aumenta anche f'' (rispetto ad f')

$$f'' = f' \frac{v_p + v_{oss}}{v_p - v_{sorg}} \quad (\text{avvicinamento})$$

Il segno sarà invece negativo al numeratore e positivo al denominatore se sorgente e osservatore si allontanano

$$f'' = f' \frac{v_p - v_{oss}}{v_p + v_{sorg}} \quad (\text{allontanamento})$$

Se applichiamo le relazioni precedenti a un'onda elettromagnetica che si propaga nel vuoto, possiamo sostituire $c=3 \cdot 10^8 \text{ m/s}$ al posto di v_p , ottenendo:

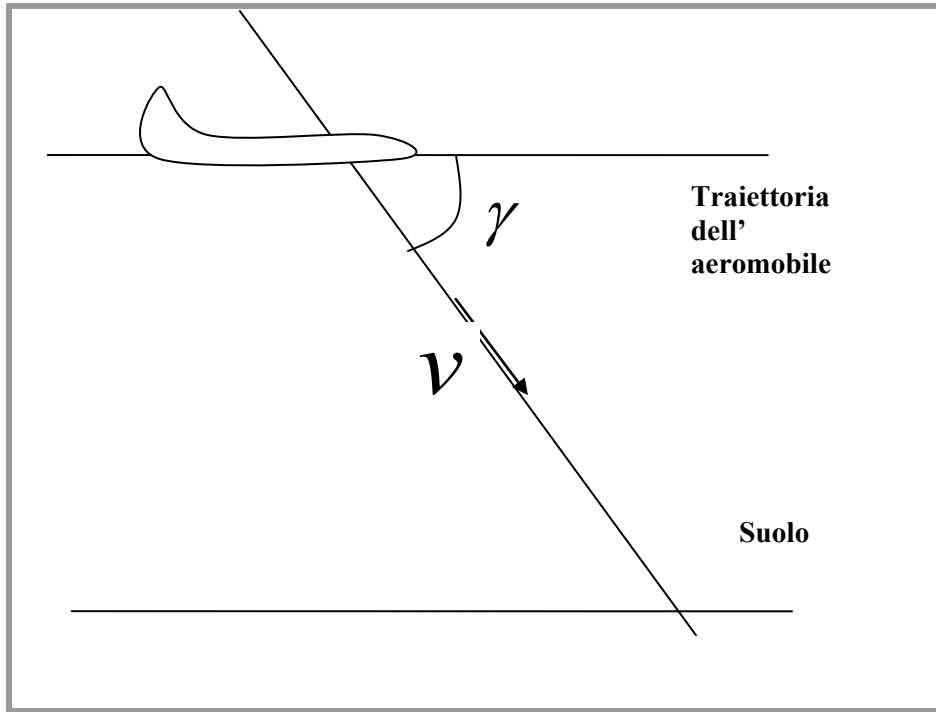
$$f'' = f' \frac{c + v_{oss}}{c - v_{sorg}} \quad (\text{avvicinamento})$$

che, nel caso di allontanamento relativo fra sorgente e osservatore, diventa:

$$f'' = f' \frac{c - v_{sorg}}{c + v_{oss}} \quad (\text{allontanamento})$$

NOTA

Le velocità che compaiono nelle precedenti relazioni sono valutate lungo la direzione radar-bersaglio, mentre a noi interessa la velocità rispetto al suolo; perciò calcoleremo la componente della velocità rispetto al suolo.



Dal disegno vediamo che

γ = angolo fra il suolo e la congiungente radar-bersaglio

La componente di v rispetto al suolo è: $v_{suolo} = v \cos \gamma$

Sostituendo $v_{suolo} = v \cos \gamma$ al posto delle v nell'ultima espressione di f'' (e considerando $v_{sorg} = v_{oss}$), otteniamo:

$$f'' = f^I \frac{c + v \cos \gamma}{c - v \cos \gamma} \quad (\text{avvicinamento})$$

A questo punto possiamo determinare lo scostamento doppler di frequenza ($f_D = \Delta f = f'' - f^I$) come:

$$f_D = \Delta f = f^I \frac{c + v \cos \gamma}{c - v \cos \gamma} - f^I$$

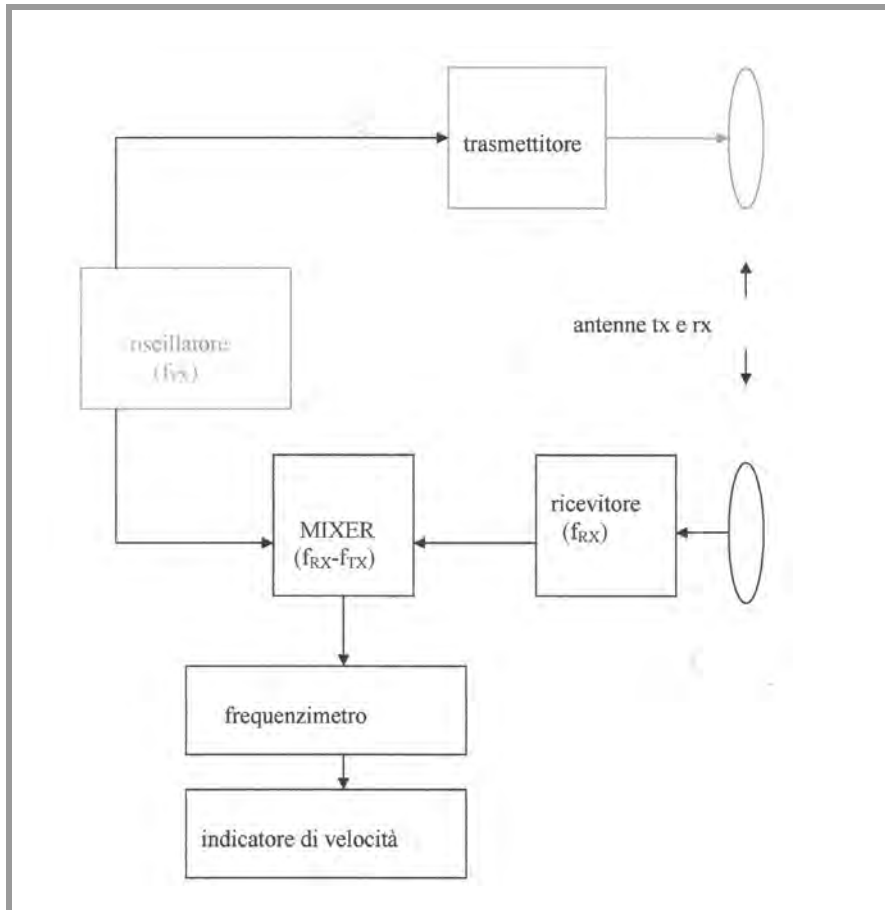
Il risultato è:

$$f_D = \Delta f = \frac{1}{\lambda} 2v \cos \gamma$$

RADAR DOPPLER A ONDA CONTINUA

I radar Doppler sono usati come:

- radar per **difesa e combattimento aereo** (scoperta e inseguimento di bersagli mobili)
- radar per il **controllo del traffico aereo**
- radar **meteorologici**



L'apparato caratterizzante è un mixer o mescolatore che fornisce, come segnale di uscita, la differenza Δf (battimento) tra la frequenza f_{RX} del segnale ricevuto e la frequenza f_{TX} del segnale trasmesso (frequenza dell'oscillatore).

RADAR DOPPLER: ESPRESSIONE DELLO SCOSTAMENTO Δf (BATTIMENTO)

Il radar Doppler è quindi basato sulla misura dello **scostamento di frequenza Δf** , che è legato alla lunghezza d'onda della sorgente e alla velocità del bersaglio dalla relazione:

$$\Delta f = \pm 2 \cdot \frac{v_{rel}}{\lambda_{sorg}}$$

dove:

λ_{sorg} = lunghezza d'onda della sorgente radar

f_{sorg} = frequenza della sorgente radar

v_{rel} = velocità relativa fra bersaglio (aereo) e sorgente radar

Essendo poi $\lambda = \frac{c}{f}$ si può scrivere:

$$\Delta f = \pm 2 \cdot \frac{v_{rel}}{\frac{c}{f_{sorg}}}$$

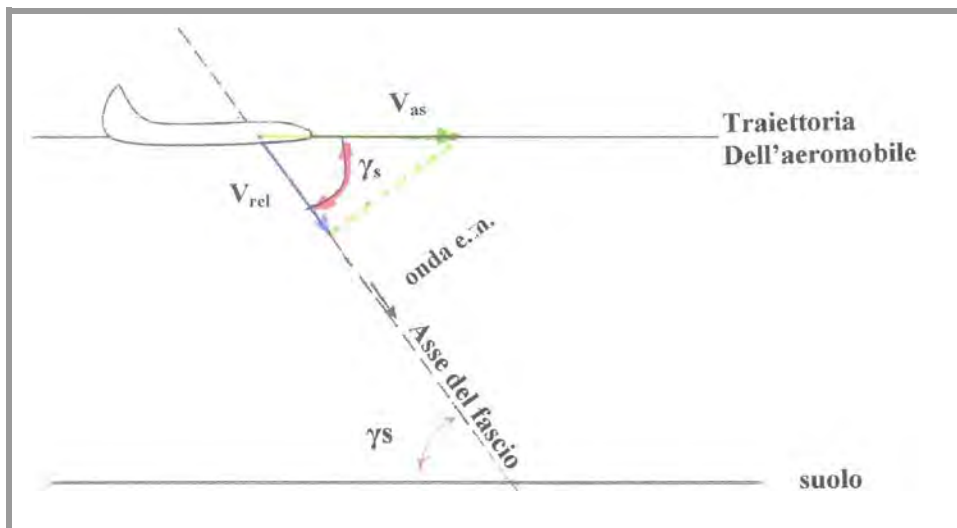
ossia:

$$\Delta f = \pm 2 \cdot f_{sorg} \cdot \frac{v_{rel}}{c}$$

RADAR DOPPLER: VELOCITA' RELATIVA E VELOCITA' RISPETTO AL SUOLO

Mediante il radar doppler posto a bordo si vuole determinare la velocità v_{as} del velivolo, nell'ipotesi che la traiettoria sia parallela al suolo.

Indichiamo con γ_s l'**angolo** che la direzione dell'aeromobile (e quindi il suolo) forma con l'asse del fascio d'onda emesso dall'antenna radar.



Il radar Doppler non misura la velocità del velivolo rispetto al suolo, ma la velocità relativa V_{rel} dell'aeromobile, cioè la componente di v_{as} (velocità rispetto al suolo) lungo l'asse del fascio d'onda irradiato dall'antenna radar. Quindi la grandezza che compare nelle espressioni dello scostamento Doppler:

$$\Delta f = \pm 2 \cdot \frac{v_{rel}}{\lambda_{sorg}} \quad (*)$$

$$\Delta f = \pm 2 \cdot f_{sorg} \cdot \frac{v_{rel}}{c} \quad (**)$$

non è la velocità rispetto al suolo, ma la velocità relativa.

Come si vede dal disegno, la velocità relativa può essere espressa come:

$$v_{rel} = v_{as} \cdot \cos \gamma_s$$

Perciò, se vogliamo che nelle espressioni (*) e (**) di Δf compaia la velocità assoluta, dobbiamo sostituire in esse espressione di v_{rel} .

Effettuando la sostituzione otteniamo:

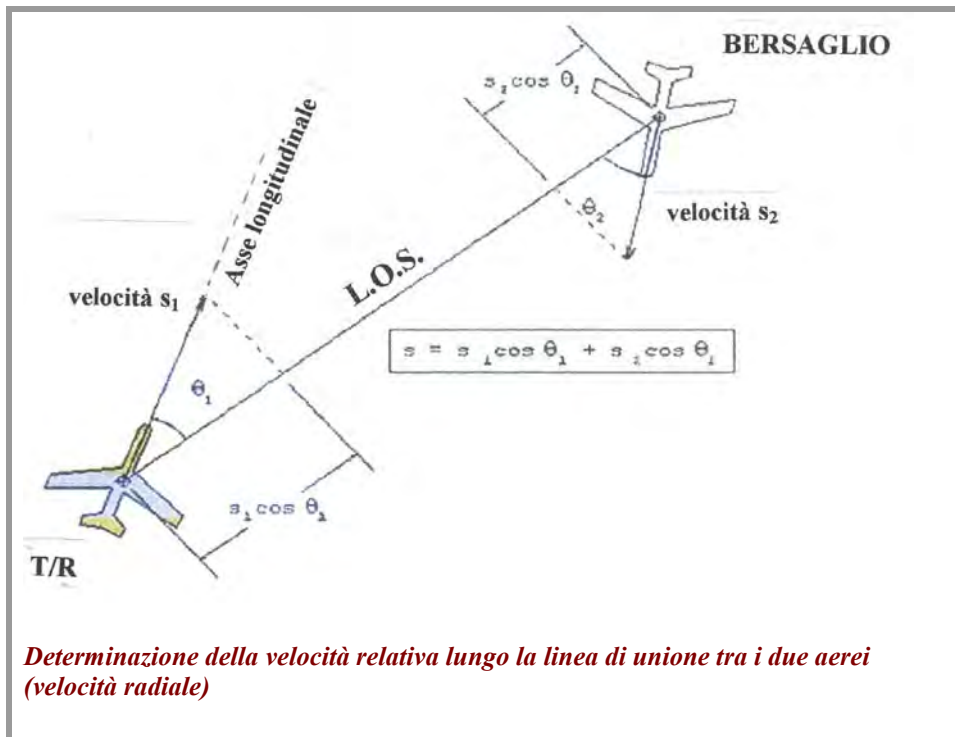
$$\Delta f = \pm 2 \cdot \frac{v_{as}}{\lambda_{sorg}} \cdot \cos \gamma_s$$

$$\Delta f = \pm 2 \cdot f_{sorg} \cdot \frac{v_{as}}{c} \cdot \cos \gamma_s$$

Il caso più generale è che la traiettoria dell'aeromobile non sia parallela al suolo, ma che la sua velocità v_{as} sia un **vettore** che forma un **angolo γ** con l'asse del fascio d'onda emesso dall'antenna radar.

NOTA

Se il velivolo trasmettitore/ricevitore e il suo bersaglio sono in moto l'uno verso l'altro, la frequenza percepita dal ricevitore è maggiore di quella emessa se invece si allontanano è minore. Il valore risultante della variazione dipenderà dalla velocità relativa tra il T/R e l'obiettivo lungo una linea, nota come *line-of-sight* (LOS), che è la linea di unione tra i due oggetti (raggio).



La variazione Doppler può essere calcolata conoscendo le **velocità effettiva** s_1 del T/R la velocità effettiva s_2 dell'obiettivo (quella diretta lungo il suo asse longitudinale) e gli angoli compresi tra la loro direzione e la linea di unione, denominati Θ_1 e Θ_2 .

La somma delle due velocità lungo la linea di unione (velocità radiale) è data dalla somma delle singole velocità:

$$s = s_1 \cos(\Theta_1) + s_2 \cos(\Theta_2)$$

Tale valore può essere interpretato anche come la velocità istantanea con la quale i due aerei, percorrendo la LOS, tendono alla collisione.

L'angolo Θ_1 rappresenta il **rilevamento relativo** con l'obiettivo (spesso indicato come **RIPLO** o **rilevamento polare**, ossia quello che utilizza le ore dell'orologio per indicare la posizione del bersaglio). Oppure, la differenza tra la *prua* (angolo compreso tra il Nord e la direzione dell'asse longitudinale dell'aereo) del T/R e il **rilevamento vero** dell'obiettivo (angolo compreso tra la direzione del Nord e la direzione dell'aereo che vogliamo indicare).

RADAR FM-CW, *Frequency Modulated Continuous Wave* (RADAR A ONDA CONTINUA MODULATA IN FREQUENZA)

Il radar CW-FM è intrinsecamente un radar Doppler.

Per il principio di funzionamento, si veda, in “**radar aeronautici**”, il **radio-altimetro**, che è proprio un’applicazione del radar in questione

RADAR IMPULSIVO CON INFORMAZIONE DOPPLER O RADAR MTI (MOVING TARGET INDICATION)

Lo schema a blocchi di un radar impulsivo con informazione Doppler contiene sia dispositivi che sono caratteristici del normale radar a impulsi (cioè del radar impulsivo senza informazione Doppler), sia dispositivi che sono invece caratteristici del radar Doppler a onda continua (CW).

Si vedano in proposito gli schemi semplificati:

- del radar Doppler a onda continua (nel paragrafo ad esso dedicato)
- del “normale radar a impulsi

Il fondamentale dispositivo che è presente nel *radar impulsivo con informazione Doppler* (e che è caratteristico del normale **radar a impulsi**) è il **modulatore** o **generatore di impulsi**, il quale ha la funzione di generare gli impulsi che abilitano e disabilitano periodicamente il trasmettitore in modo che il segnale irradiato risulti impulsivo e non continuo.

Il fondamentale dispositivo che è presente nel *radar impulsivo con informazione Doppler* e che è caratteristico del *radar Doppler a onda continua (CW)* è il *mixer* che fornisce lo scostamento Doppler Δf , dal quale si determina la velocità del bersaglio.

Osserviamo inoltre che, tipicamente, nel radar impulsivo con informazione Doppler non è presente un magnetron (come nel normale radar impulsivo), ma un **klystron**, che è sempre un tubo a vuoto (valvola termoionica), ma ha caratteristiche diverse.

RADAR METEO (RADAR DOPPLER METEOROLOGICO)

Il RADAR meteorologico misura la quantità di pioggia, neve e ghiaccio (idrometeore) presenti nell'atmosfera. Può acquisire dati in tre dimensioni e monitorare aree estese, ad esempio un volume fino a 200 km di distanza e 10 km di altezza dal suolo in pochi minuti.

Fino a poco tempo non esistevano radar meteorologici con la funzione di acquisire dati finalizzati all'analisi operativa delle perturbazioni atmosferiche. I soli radar meteorologici funzionanti erano nati per esclusivi scopi di ricerca e non erano gestiti tramite calcolatori.

Attualmente i sistemi di nuova generazione permettono un controllo automatico del sensore, l'ottimizzazione e la facilità d'impostazione delle modalità operative, la gestione e la memorizzazione di grandi quantità di dati.

Il radar meteo **emette brevi impulsi di onde elettromagnetiche** di elevata potenza nell'atmosfera lungo la direzione di puntamento dell'antenna, che può variare sia in azimut che in elevazione. **I pacchetti di onde sono assorbiti dalle idrometeore presenti nell'atmosfera e reirradiati in tutte le direzioni** tra cui quella del RADAR.

L'analisi del **segnale di ritorno**, che prende il nome di **riflettività**, è effettuata nell'apparato ricevente del RADAR e permette di determinare **l'intensità della precipitazione**.

La direzione di puntamento dell'antenna e il **tempo impiegato dal segnale** nel percorso andata-ritorno consentono di localizzare le idrometeore in termini di **direzione e distanza**.

Inoltre piccole **variazioni nella frequenza** dell'eco di ritorno permettono, attraverso l'effetto Doppler, di misurare la **velocità radiale** e quindi di stimare la **direzione di spostamento** dell'evento meteorologico.

I RADAR meteorologici operano nell'intervallo di frequenze delle microonde perché la lunghezza d'onda delle microonde è confrontabile con la dimensione delle idrometeore stesse, di solito si utilizzano le **bande S e C**

	BANDA S	BANDA C
Lunghezza d'onda	15-7.5 cm	7.5-3.75 cm
Frequenza	2-4 GHz	4-8 GHz

L'attenuazione aumenta al diminuire della lunghezza d'onda (e quindi al crescere della frequenza) ed è per questo motivo che la **banda S** è la più **indicata** per le regioni tropicali e per quelle aree dove **uragani, tornado e cicloni sono più probabili**.

I radar in banda S tuttavia, comportano notevoli **inconvenienti in termini di dimensioni e stabilità strutturale**: siccome le dimensioni dell'antenna sono proporzionali alla lunghezza d'onda, per un fascio di 1 grado a 10 cm di lunghezza d'onda è necessario un riflettore di 7,3 m di diametro, sconsigliabile sia per l'alto costo che per la notevole instabilità meccanica.

I radar in banda C realizzano il **miglior compromesso** tra problematiche realizzative e prestazioni meteorologiche e perciò, attualmente sono **utilizzati dalla maggior parte dei nuovi impianti** installati in regioni **non tropicali**. Tornando all'esempio precedente, a 5 cm di lunghezza d'onda un fascio della stessa apertura è ottenibile con un'antenna di soli 3,7 m;

Esistono, inoltre, radar che utilizzano valori inferiori di lunghezza d'onda (frequenze più elevate) e precisamente:

	BANDA X	BANDA K
Lunghezza d'onda	3.75-2.5 cm	2.5-0.75 cm
Frequenza	8-12 GHz	12-40 GHz

Si tratta, però, di impianti di uso è problematico, poiché a queste lunghezze d'onda il fascio radar è soggetto ad un notevole assorbimento atmosferico.

Con lunghezze d'onda superiori a 3 cm e nell'ipotesi che le particelle abbiano un diametro inferiore ai 5 mm, la teoria permette di legare, in un'unica relazione nota come *equazione del radar*, la riflettività di un volume di atmosfera con l'intensità dell'eco, le caratteristiche del radar e quelle del mezzo in cui avviene la propagazione dell'onda elettromagnetica. In particolare, risulta che il fascio può subire una notevole attenuazione per la presenza di precipitazioni poste tra l'antenna ed il volume in osservazione.

Altri fattori possono poi concorrere a modificare il fascio radar durante la sua propagazione; essi sono di natura sia geometrica che fisica (legati a variazioni dell'indice di rifrazione atmosferico) o dovuti alla presenza di ostacoli naturali o artificiali. In assenza di contromisure, possono portare un'errata valutazione della situazione meteorologica.

Una volta ottenuto un valore di riflettività il più possibile corretto, esso può essere convertito in quantità di pioggia tramite una relazione di questo tipo:

$$Z = a R^b$$

dove

- *R è la precipitazione in mm/h,*
- *Z è la riflettività*
- *a e b sono costanti il cui valore dipende dal tipo di precipitazione.*

Il problema della determinazione delle costanti è uno dei più discussi in radar-meteorologia; in letteratura esistono moltissime coppie di coefficienti calcolate per diversi tipi di precipitazione e verificate sperimentalmente mediante confronto con una rete di sensori pluviometrici posti nella zona spazzata dal radar. E' evidente, tuttavia, come questa relazione sia puramente empirica e soggetta ad incertezze dipendenti soprattutto dalla distribuzione delle dimensioni delle gocce nel volume in esame e dalla durata temporale della misura. Per migliorare il grado di precisione vengono usati i cosiddetti sistemi a doppia polarizzazione, dove la riflettività viene valutata su due piani tra loro ortogonali. Infatti, tanto più le gocce sono grandi, tanto più assumono, durante la caduta, una forma schiacciata e tanto più la riflettività orizzontale differisce da quella verticale, fornendo così un'indicazione della distribuzione delle dimensioni. In questo caso esistono delle relazioni diverse da quella citata dato che entrano in gioco altri parametri come *la riflettività differenziale, la fase differenziale etc.* Inoltre, la polarizzazione duale permette di distinguere le gocce di pioggia dalle particelle di ghiaccio, che durante la caduta conservano una forma pressoché sferica.

Infine, si deve sottolineare che **(per quanto possa essere impreciso, rispetto al satellite meteorologico, che vede essenzialmente la cima delle nubi , ed i sensori di punto, che rilevano dati su di un punto dello spazio), un radar fornisce una misura delle grandezze osservabili su tutto lo spazio e rileva in maniera completa la struttura dei fenomeni meteorologici.**

CAPITOLO 32

I TRASDUTTORI

COMPLEMENTI

TIPOLOGIE DI TRASDUTTORI

TRASDUTTORI DI TEMPERATURA

INTERRUTTORI BIMETALLICI

TERMOCOPPIE

PROBLEMI NELLA LINEARIZZAZIONE DI UNO STRUMENTO A TERMOCOPPIA

SENSORI DI TEMPERATURA INTEGRATI

SENSORE INTEGRATO AD 590

SENSORE INTEGRATO LM 35

TERMISTORI NTC E PTC

TERMORESISTENZE O RTD

TRASDUTTORI DI FORZA (CELLE DI CARICO)

TRASDUTTORI DI FORZA PIEZOELETTRICI

TRASDUTTORI DI PRESSIONE

PREMESSA: LA MISURA DELLA PRESSIONE

TRASDUTTORI DI PRESSIONE BASATI SU “PURI ELEMENTI ELASTICI”

TRASDUTTORI DI PRESSIONE CAPACITIVI

TRASDUTTORI DI PRESSIONE INDUTTIVI

TRASDUTTORI ESTENSIMETRICI (ESTENSIMETRI O STRAIN-GAUGE) *

TRASDUTTORI DI POSIZIONE

ENCODER ASSOLUTO

POTENZIOMETRI

TRASFORMATORI DIFFERENZIALI O LVDT

SINCRO

RESOLVER O RISOLUTORI

TRASDUTTORI DI SPOSTAMENTO CAPACITIVI

TRASDUTTORI DI SPOSTAMENTO OTTICI A RIFLESSIONE

TRASDUTTORI DI SPOSTAMENTO A ULTRASUONI

PIEZOELETTRICITÀ E MAGNETOSTRIZIONE

TRASDUTTORI DI POSIZIONE A EFFETTO HALL

INTERRUTTORI DI PROSSIMITÀ OTTICI O FOTOELETTRICI

TRASDUTTORI DI VELOCITÀ

ENCODER INCREMENTALE

ACCELEROMETRO (TRASDUTTORE DI ACCELERAZIONE)

TRASDUTTORI DI LIVELLO
MISURATORI DI PRESSIONE DIFFERENZIALE
MISURATORI DI CAPACITÀ ELETTRICA

TRASDUTTORI DI ENERGIA RADIANTE
UNA POSSIBILE CLASSIFICAZIONE
CELLULA FOTOELETTRICA
FOTOMOLTIPLICATORI
TUBI INTENSIFICATORI DI IMMAGINE
FOTORESISTENZE
FOTODIODO
FOTOTRANSISTOR
DISPOSITIVI FOTOVOLTAICI (PILE O CELLE SOLARI)

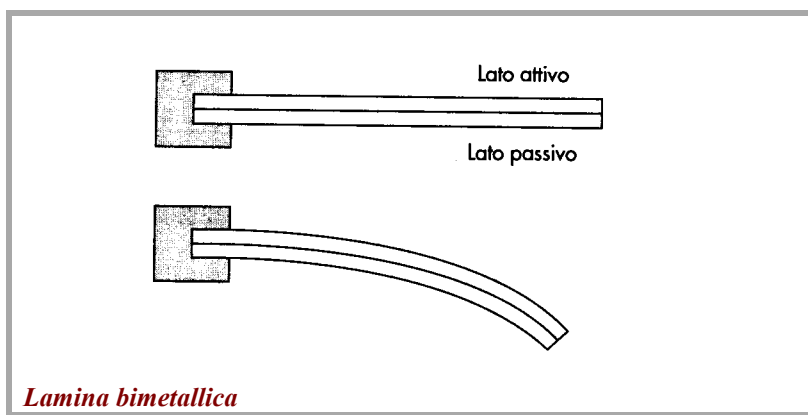
RIVELATORI PIROELETTRICI
SENSORI CHIMICI E DI UMIDITA'
SENSORE DI GAS DI SCARICO PER AUTO (SONDA LAMBDA)
SENSORI DI IONI (ISFET O CHEMFET)

TIPOLOGIE DI TRASDUTTORI

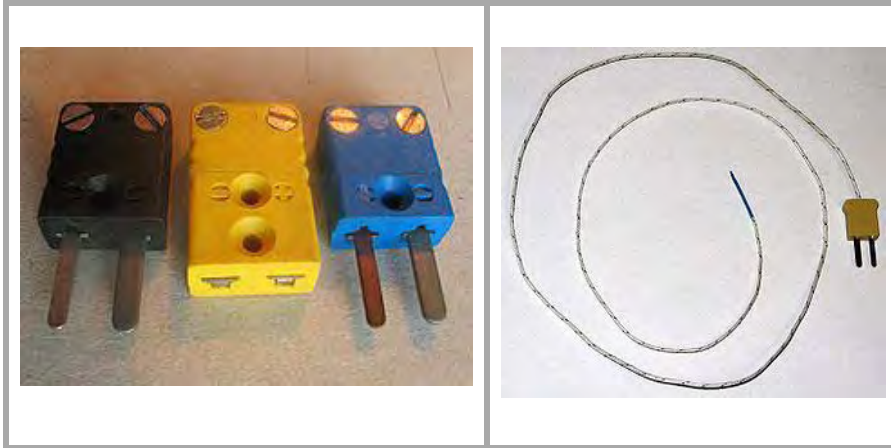
TRASDUTTORI DI TEMPERATURA

1) INTERRUTTORI BIMETALLICI

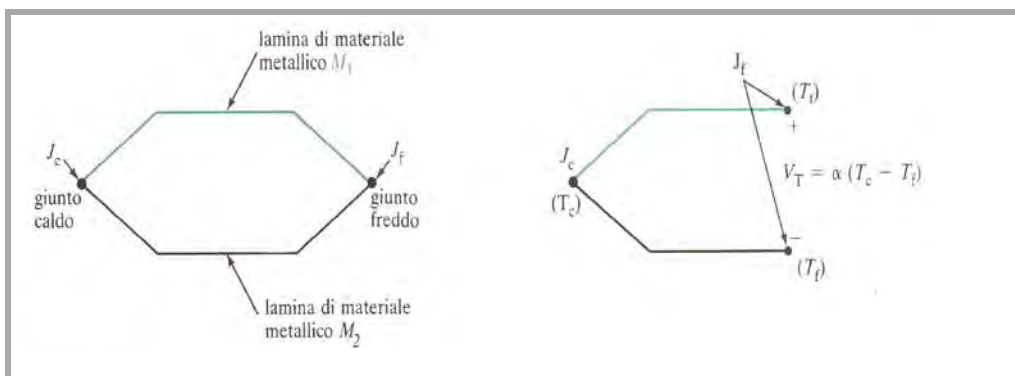
Sono costituiti da due lamine metalliche, con differente coefficiente di dilatazione, e fissate l'una all'altra. Al crescere della temperatura dell'ambiente, una delle due lamine si dilata più dell'altra e ciò provoca una flessione della struttura, la quale andrà a toccare un interruttore determinandone la commutazione.



2) TERMOCOPPIE



Sono basate sull'effetto Seebeck, che si produce quando si uniscono tra loro, mediante saldature, due metalli differenti, dando luogo a due punti di contatto, dette giunzioni.



Se i punti di giunzione vengono a trovarsi a temperature diverse, si crea una corrente elettrica (termocorrente) di intensità proporzionale alla differenza $(T_c - T_f)$ fra la temperatura della giunzione calda e la temperatura della giunzione fredda. Se, come normalmente si usa fare, la struttura non viene chiusa da entrambe le parti, ma viene lasciata aperta in corrispondenza della giunzione fredda, non vi è circolazione di corrente, ma fra i terminali del giunto freddo si può prelevare una tensione proporzionale alla differenza di temperatura fra i due giunti:

$$V_T = \alpha \cdot (T_c - T_f)$$

Le coppie di materiali maggiormente usate e i relativi campi di temperatura di impiego sono:

- cromo-costantana ($100 \div 1000$)°C
- ferro- costantana ($0 \div 800$)°C
- rame-costantana ($-200 \div 350$)°C

Ricordiamo che la **costantana** è una lega di **rame** (Cu 60%) e **nickel** (Cr 40%).

Gli **inconvenienti** delle termocoppie sono:

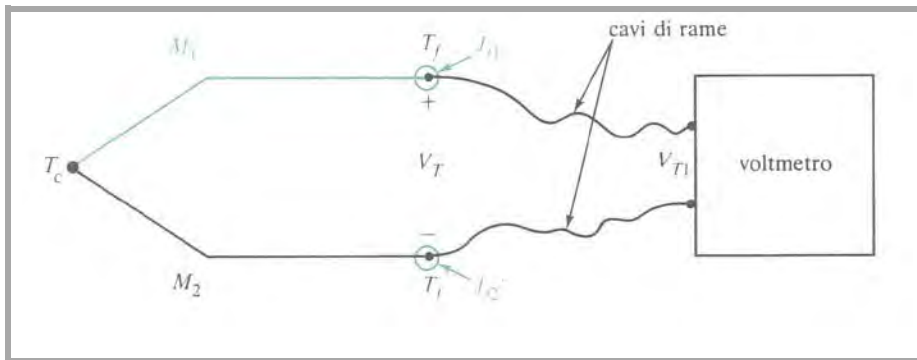
- scarsa linearità della caratteristica di trasferimento
- basso livello della tensione di uscita (che è dell'ordine della decina di mV).

I **pregi** sono:

- costi contenuti
- campo di temperature di lavoro molto ampio
- buona ripetibilità e buona resistenza alle sollecitazioni.

PROBLEMI NELLA LINEARIZZAZIONE DI UNO STRUMENTO A TERMOCOPPIA

Per misurare la tensione della termocoppia, bisogna collegarla a un voltmetro, ma nei punti di contatto fra i terminali del voltmetro e i due metalli della termocoppia, si creano due nuovi giunti termici “parassiti”. Quindi il voltmetro misurerebbe una tensione che dipende anche dalle due nuove giunzioni.



Questo problema si può risolvere inserendo i due giunti “parassiti” in un contenitore isotermico, che li mantenga alla stessa temperatura del contenitore stesso, in modo che non forniscano alcuna tensione indesiderata.

Così il voltmetro fornirebbe la tensione:

$$V_{TM} = \alpha \cdot (T_c - T_{ISO})$$

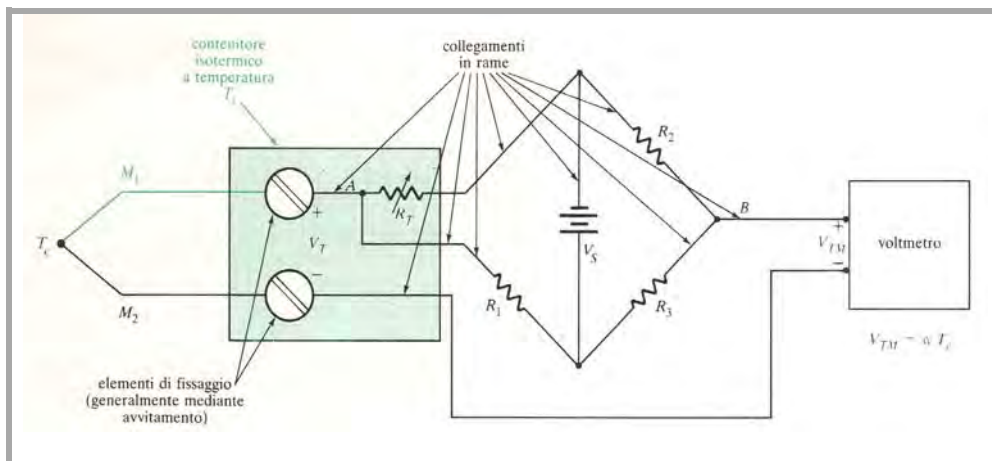
C'è però ancora il problema che la tensione dipende anche dalla temperatura T_{ISO} del contenitore isotermico, mentre noi vorremmo che il voltmetro ci fornisse direttamente:

$$V = \alpha \cdot T_c$$

Si potrebbe pensare di mantenere il contenitore alla temperatura di 0°C , ma ciò sarebbe scomodissimo.

Si adotta allora questa tecnica:

- si collega uno dei due terminali della termocoppia a una termoresistenza R_T il cui valore deve essere crescente con la temperatura del contenitore isotermico: $R_T = \alpha \cdot T_{ISO}$
- R_T deve costituire un ramo di un ponte di Wheatstone
- si inserisce la termoresistenza nel contenitore isotermico
- detto R_{T0} il valore della termoresistenza R_T alla temperatura di 0°C , si equilibra il ponte in corrispondenza di 0°C , cioè si fa in modo che, a 0°C , la tensione di uscita del ponte risulti nulla ($V_{AB}=0$); questo risultato si ottiene scegliendo i valori di R_1 , R_2 ed R_3 in modo che risulti:
$$R_{T0} \cdot R_3 = R_1 \cdot R_2$$



Dall'osservazione del circuito, si vede che la tensione misurata dal voltmetro è esprimibile come:

$$V_{TM} = V_T + V_{AB}$$

dove:

V_T = tensione di uscita della termocoppia

V_{AB} = tensione di uscita del ponte

Se il ponte è equilibrato a 0°C , cioè se si ha $V_{AB}=0$ per $T_{ISO}=0^\circ\text{C}$, allora la V_{AB} risulterà proporzionale a T_{ISO} , cioè si avrà $V_{AB} = \alpha \cdot T_{ISO}$

Il voltmetro pertanto ci fornirà:

$$V_{TM} = V_T + V_{AB} \quad \text{con} \quad V_T = \alpha \cdot (T_c - T_{ISO}) \quad \text{e} \quad V_{AB} = \alpha \cdot T_{ISO}$$

e quindi:

$$V_{TM} = \alpha \cdot T_c - \alpha \cdot T_{ISO} + \alpha \cdot T_{ISO}$$

$$V_{TM} = \alpha \cdot T_c$$

così come si voleva.

3) SENSORI DI TEMPERATURA INTEGRATI

Sono realizzati con semiconduttori e basati sulla **diminuzione di tensione** che, a parità di corrente, si produce ai capi di una giunzione pn **all'aumentare della temperatura**.

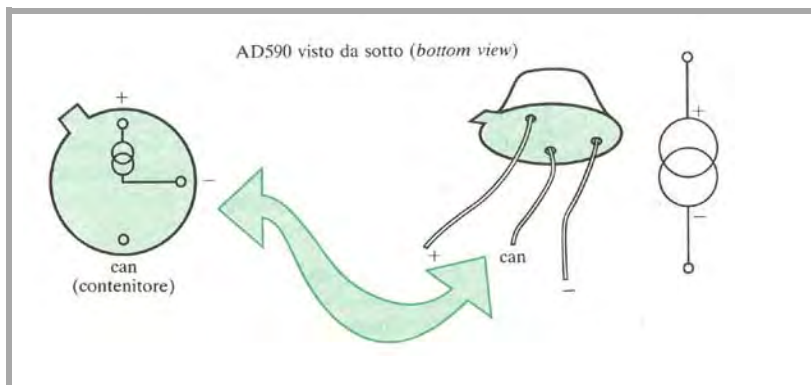
Tale diminuzione è dell'ordine di 2 mV/°K.

Naturalmente, dire che, a corrente costante, la tensione ai capi della giunzione diminuisce al crescere della temperatura, equivale a dire che, a parità di tensione, **la corrente cresce** al crescere della temperatura.

SENSORE INTEGRATO AD 590

Si comporta come un **generatore di corrente ad alta impedenza interna**.

Ciò permette di rendere il sistema di misura indipendente da cadute di tensione e di porre il trasduttore lontano dal punto di misura.



L'AD 590 fornisce in uscita una corrente proporzionale alla temperatura:

$$I = k \cdot T = k'(T_c + 273)$$

Dove: $k = \text{cost} = 1 \text{ mA/}^\circ\text{K}$

T = temperatura in gradi **Kelvin**

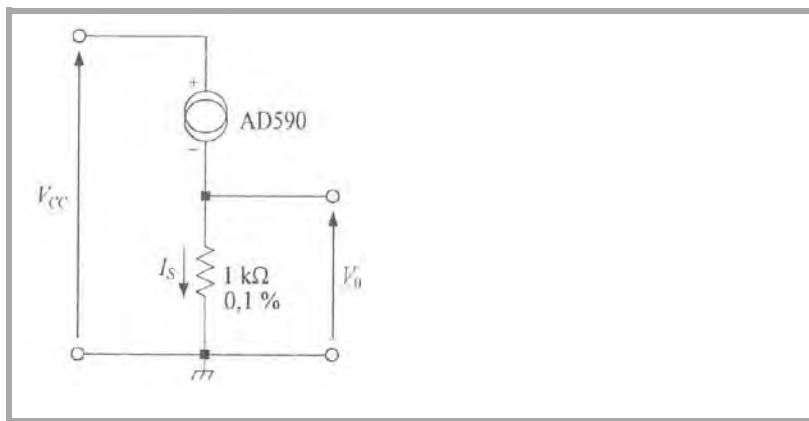
T_c = temperatura in gradi centigradi.

Il campo delle temperature di lavoro è (-55-150)°C

Il campo delle tensioni di alimentazione è (4-30)V

Si dice che l'AD 590 ha "l'uscita in corrente", ossia che fornisce una corrente (e non una tensione) proporzionale alla temperatura perché, per effetto dell'elevata impedenza interna, si comporta come un generatore di corrente, cioè come un generatore il cui valore di corrente non è influenzato dal valore dell'impedenza di carico. Quindi il segnale di uscita dell'AD590, essendo una corrente, non dipende dal carico e pertanto è indipendente dalla lunghezza della linea che collega il trasduttore al punto di misura: più lunga è la linea, più grande è la sua resistenza, ma, per quanto detto precedentemente, tale resistenza (che è una resistenza di carico) non altera il valore della corrente.

E' possibile convertire l'uscita in corrente in una uscita in tensione, facendo scorrere la corrente di uscita su un resistore o applicandola a un convertitore corrente-tensione a operazionale.



SENSORE INTEGRATO LM 35

Fornisce in uscita una tensione proporzionale alla temperatura:

$$V = k \cdot T_c = k \cdot (T - 273)$$

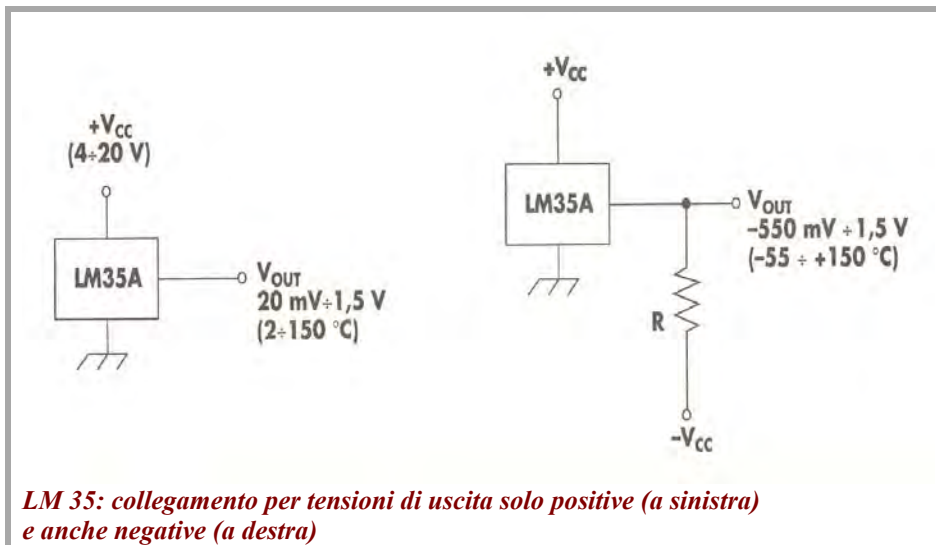
dove: $k = \text{cost} = 10 \text{ mV}/^\circ\text{C}$

T_c = temperatura in gradi **centigradi**

T = temperatura in gradi Kelvin

Il campo delle temperature di lavoro è lo stesso del precedente sensore
(-55÷150)°C

Il campo delle tensioni di alimentazione è (4÷20)V



Si dice che l'LM 35 ha "l'uscita in tensione", ossia che fornisce una tensione (e non una corrente) proporzionale alla temperatura perché, avendo una bassa impedenza interna, si comporta come un generatore di tensione, cioè come un generatore che eroga una tensione il cui valore non è influenzato dal valore dell'impedenza di carico.

4) TERMISTORI NTC E PTC

Sono realizzati con materiali **semiconduttori** o dal comportamento analogo (per esempio **miscele di ossidi metallici sottoposti a drogaggio**).

L'intensità del drogaggio è **bassa** per gli **NTC** (e perciò il coefficiente di temperatura è **negativo**), alta per i PTC.

Il loro funzionamento si basa sulla proprietà di variare la resistenza elettrica al variare della temperatura.

Pregi: rapidità di risposta, sensibilità

Inconvenienti: ridotto campo di linearità.

Gli NTC sono termistori a Coefficiente di Temperatura Negativo, cioè il loro valore di resistenza diminuisce all'aumentare della temperatura (carbone, grafite e leghe derivate).

Gli NTC sono utilizzati nei controlli di temperatura, dove la linearità non è fondamentale e in ambiti nei quali il range di temperature di lavoro è ristretto, come nel caso dei termometri clinici.

Esistono NTC con campo di temperatura $(-100 \div 450)^{\circ}\text{C}$.

I PTC sono termistori a Coefficiente di Temperatura Positivo, cioè il loro valore di resistenza aumenta con la temperatura. Ciò però avviene solo nel tratto centrale della caratteristica di funzionamento.

Il campo di temperature dei PTC è più ristretto di quello degli NTC:.

5) TERMORESISTENZE o RTD (Resistance Temperature Detector)

Sono realizzate in **metallo** (platino, nichel e talvolta rame) e si basano sulla proprietà dei metalli di aumentare la loro resistenza elettrica all'aumentare della temperatura, secondo la legge:

$$R_T = R_0(1 + \alpha T)$$

dove:

R_T = Resistenza alla temperatura T in $^{\circ}\text{C}$

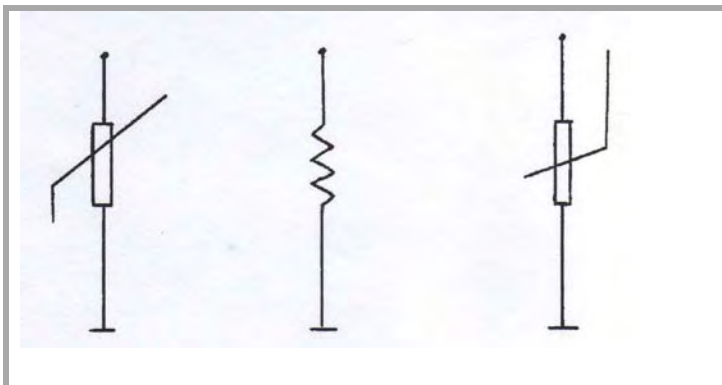
R_0 = Resistenza alla temperatura $T = 0^{\circ}\text{C}$

α = coefficiente di temperatura espresso in $(^{\circ}\text{C})^{-1}$
(per il platino vale $3,85 \cdot 10^{-3} (^{\circ}\text{C})^{-1}$)

Questa legge però è approssimata, per cui il comportamento reale non è lineare.

Le termoresistenze presentano, tipicamente, a temperatura ambiente, valori di
($20 \div 20000$) Ohm.

Il campo di temperature delle migliori termoresistenze (platino)
è $(-200 \div +800)^{\circ}\text{C}$.



TRASDUTTORE	RANGE DI FUNZIONAMENTO	LINEARITA'	SENSIBILITA'	ACCORGIMENTI DA ADOTTARE
TERMOCOPPIE	Molto ampio (-200 ÷ 1900)°C	Lineare solo per piccoli tratti del range	scarsa	Amplificazione Inserimento del giunto freddo e dei contatti in un contenitore isotermico Inserimento in un ponte di Wheatstone
TERMO RESISTENZE METALLICHE RTD	ampio (-200 ÷ 800)°C	Non lineare (lineare soltanto in un intervallo, non ampio, attorno a 0°C)	limitata	Amplificazione linearizzazione
TERMISTORI NTC	(-100 ÷ 400)°C	Non lineare	alta	Linearizzazione Misurazione di una tensione e non di un valore resistivo
TERMISTORI PTC	Più ristretto del range degli NTC	Non lineare	alta	Linearizzazione

TRASDUTTORI DI FORZA (CELLE DI CARICO)

La forza è misurata:

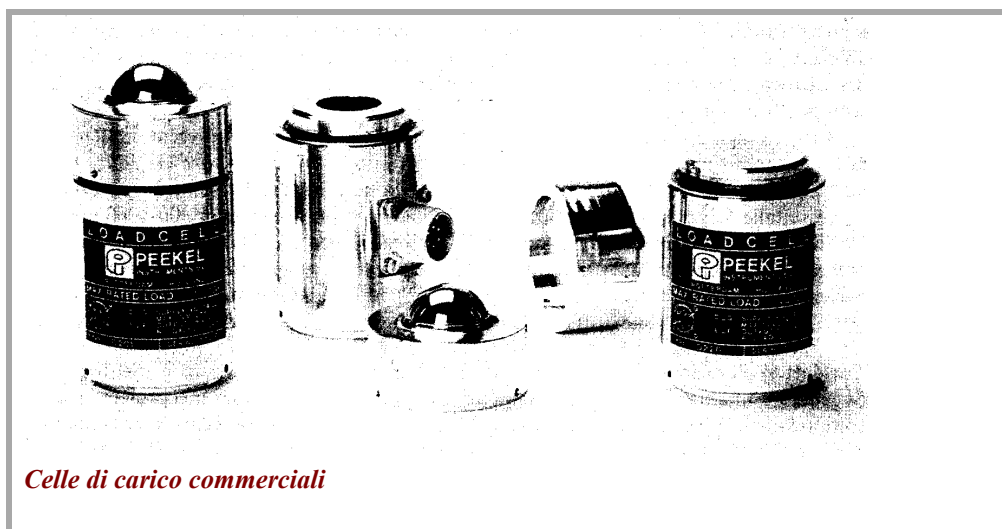
- in chilogrammi-peso (Kg_p) nel sistema metrico decimale
 - in Newton (N) nel Sistema Internazionale (SI).
- Il Newton è l'intensità di quella forza che, applicata alla massa di 1 Kg, le imprime l'accelerazione di 1 m/s^2 . Osserviamo esplicitamente che il Newton è circa 10 volte più piccolo del Kg_p : $1 \text{ N} = 1 \text{ Kg}_p / 9,81$

I trasduttori di forza, detti anche “celle di carico” si basano su una duplice conversione:

- conversione della forza in deformazione o in deflessione di un elemento elastico
- conversione della deformazione in un segnale elettrico (mediante un trasformatore differenziale o un estensimetro, o un trasduttore piezoelettrico¹).

forza $\rightarrow$ deformazione ----- deformazione $\rightarrow$ segnale elettrico

Le celle di carico sono usate per la misura di forze e di masse (pesatura).



¹ elencati in ordine di forze crescenti

TRASDUTTORI DI FORZA PIEZOELETTRICI

Sono basati sull'effetto piezoelettrico, che si manifesta in alcuni materiali, come il quarzo. Grazie a questo effetto, il materiale genera una tensione elettrica quando è sottoposto a sollecitazioni meccaniche quali compressione, flessione o stiramento.

I vantaggi di questi trasduttori sono:

- grande sensibilità
- ampio intervallo di frequenze di funzionamento ($20 \div 20000$) Hz.

Si veda anche il paragrafo sugli estensimetri.

TRASDUTTORI DI PRESSIONE

PREMESSA: LA MISURA DELLA PRESSIONE

Da un punto di vista dimensionale, la pressione è un rapporto forza/superficie. Le unità di misura più diffuse sono riportate nella tabella seguente.

Premettiamo qualche considerazione.

La più “piccola” delle unità di misura che consideriamo è il “Pascal” (Pa).

$1 \text{ Pascal} = 1 \text{ Newton} / 1 \text{ metroquadro}$.

Siccome $1 \text{ Newton} = 1 \text{ Kg}_{\text{peso}}/9,81$
possiamo dire che:

il Pascal è la forza di circa un ettogrammo distribuita sulla superficie di un metroquadro.

Il “Torricelli” (Torr) o “millimetro di mercurio” (mmHg) vale circa 133 Pascal:

$1 \text{ Torr} = 1 \text{ mmHg} = 133,32 \text{ Pa}$.

La “atmosfera” (atm) è quasi uguale al “bar” ed entrambe le unità valgono circa 100'000 Pascal:

$1 \text{ atm} = 1,01325 \cdot 10^5 \text{ Pa}$

$1 \text{ bar} = 10^5 \text{ Pa}$.

Quindi sia l'atmosfera che il bar equivalgono alla forza di circa 10'000 Kg_{peso}
(100'000 ettogrammi_{peso}) distribuiti sulla superficie di un m^2 .

E' anche utilizzato l'ettoPascal”, un multiplo del Pascal costituito da 100 Pa ed equivalente a 1 millibar.

Esiste inoltre il “Pound force per Square Inch” (psi), che equivale a circa 7000 Pascal:

$1 \text{ psi} = 6895 \text{ Pa}$.

TIPI DI MISURA DELLA PRESSIONE

1. Pressione **DIFFERENZIALE**: è la differenza di pressione fra due ambienti, uno dei quali è preso come **RIFERIMENTO**.
2. Pressione **RELATIVA**: è la pressione misurata rispetto alla pressione dell'**AMBIENTE**.
3. Pressione **assoluta**: è la pressione misurata rispetto **al vuoto**.

TRASDUTTORI DI PRESSIONE BASATI SU “PURI” ELEMENTI ELASTICI

In questi trasduttori la pressione viene rilevata utilizzando ELEMENTI ELASTICI come DIAFRAMMI, CAPSULE, SOFFIETTI, TUBI DI BOURDON.

Il DIAFRAMMA è una membrana circolare elastica, con i bordi fissi, che, sotto l'azione di una pressione, si flette in misura proporzionale alla pressione stessa. E' tipicamente realizzata in acciaio e può essere liscia o corrugata. Potendo subire piccole deformazioni, è utilizzata solo in caso di piccole pressioni.

La CAPSULA o ANEROIDE può sopportare una pressione pressoché doppia di quella del diaframma in quanto è realizzata con due diaframmi corrugati uniti lungo il bordo.

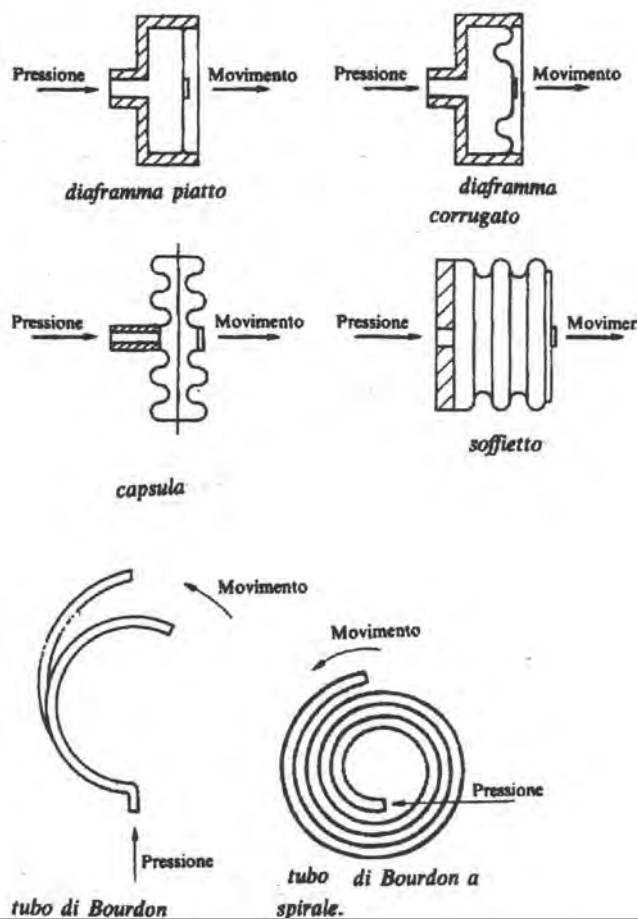
Il SOFFIETTO è un tubo dalle pareti dotate di convoluzioni ed è adatto a pressioni basse.

Il TUBO DI BOURDON è un tubo flessibile, incurvato in forma di spirale o di “C”, chiuso a una estremità, il quale tende a raddrizzarsi quando l'estremità aperta è sottoposta a una pressione.

Può misurare pressioni elevate, fino a 35 MPa.

Tabella di equivalenza fra alcune unità di pressione e pascal.

Altre unità di pressione	pascal [Pa] ($= 1 \text{ N/m}^2$)
atmosfera [atm]	$1,01325 \times 10^5 \text{ Pa}$
millimetri di Hg [mmHg]	133,32 Pa
torr ($= 1 \text{ mmHg}$)	133,32 Pa
bar	10^5 Pa
pound force per square inch [psi]	$6,895 \times 10^3 \text{ Pa}$

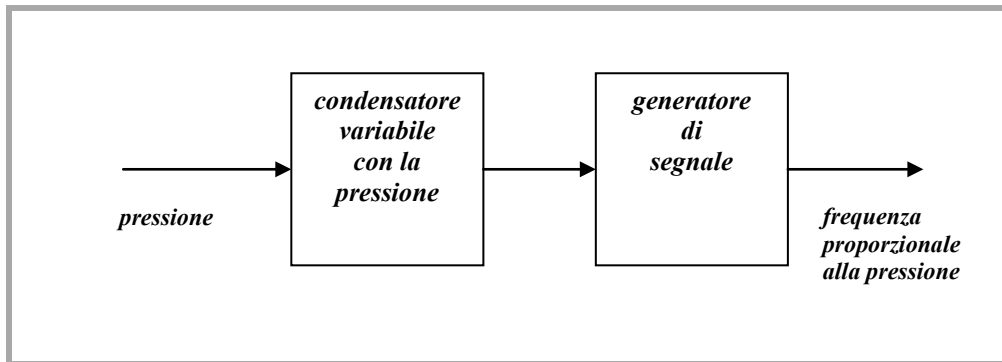


TRASDUTTORI DI PRESSIONE CAPACITIVI

Il condensatore utilizzato in questi trasduttori ha una delle due armature costituita da un diaframma mobile.

La pressione applicata al diaframma ne determina lo spostamento e quindi produce una variazione di capacità del condensatore.

Il condensatore variabile così realizzato può essere inserito in un generatore di segnale. Si ottiene così una frequenza proporzionale alla tensione.



Trasduttori di questo tipo sono tipicamente usati per misure di pressione assolute, essendo dotati di una cavità sotto vuoto.

Un esempio di utilizzazione è la misurazione della pressione delle onde acustiche.

I vantaggi sono:

- agevole interfacciabilità con dispositivi elettronici
- dimensioni contenute
- precisione elevata.

TRASDUTTORI DI PRESSIONE INDUTTIVI

La pressione viene applicata a un diaframma metallico mobile.

Il diaframma è posto in prossimità di una bobina percorsa da corrente alternata.

Si determina così, nella bobina, una variazione del valore di induttanza dipendente dalla distanza bobina-diaframma e quindi dalla pressione.

TRASDUTTORI ESTENSIMETRICI

Si veda il paragrafo seguente

TRASDUTTORI ESTENSIMETRICI O STRAIN-GAUGE

Detti anche STRAIN-GAGE, possono essere realizzati con materiale **conduttore** o **semiconduttore**.

Se realizzati con conduttore, si basano sulla variazione di resistenza di un corpo sottoposto a deformazione geometrica.

Sappiamo infatti che la resistenza di un corpo aumenta all'aumentare della sua lunghezza e diminuisce all'aumentare della sezione:

$$R = \rho * (L / S) \quad (\text{dove } \rho \text{ è la resistività del materiale}).$$

Un diffuso tipo di strain-gauge è quello detto "a francobollo". E' costituito infatti da un filamento resistivo (per es. di costantana) applicato a un foglio di carta o di resina o di plastica, foglio che viene a sua volta fissato all'oggetto del quale si vuole misurare la deformazione.

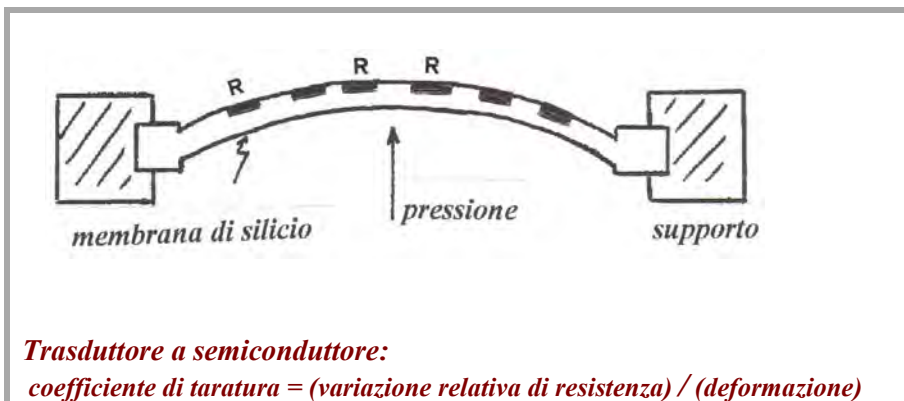
Se invece il trasduttore è a **semiconduttore** si basa sull'effetto **piezoresistivo** del semiconduttore che si manifesta in una **variazione della resistenza a causa di una deformazione meccanica**, *deformazione che, alterando la struttura molecolare del materiale, ne varia la resistività.*

In seguito a una **trazione**, la **resistenza** complessiva del wafer di semiconduttore **aumenta**.

In seguito a una compressione la resistenza complessiva del wafer diminuisce.

Per esempio i sensori di pressione al silicio sono realizzati con un chip sul quale viene incisa una membrana. Sulla membrana vengono create, per impiantazione ionica², le piste resistive.

La membrana di silicio (che, se è quadrata, può misurare qualche mm di lato) è fissata su un supporto che, pur impedendone lo spostamento, ne consente la deformazione in seguito all'effetto della pressione.



Pregi fondamentali degli strain-gauge sono:

- robustezza e resistenza alla fatica
- sensibilità e linearità buone
- basso tempo di risposta

² L'**impiantazione ionica** è una tecnica di **drogaggio** del semiconduttore che ne prevede il bombardamento (a temperatura non elevata) con ioni di materiale drogante. La profondità di penetrazione degli ioni droganti nel semiconduttore non supera la decina di μm (micron) e dipende dall'energia e dalla concentrazione degli ioni. Gli ioni, ottenuti da vapori di materiale drogante, sono di boro per il drogaggio di tipo p e di **fosforo** per il drogaggio di tipo n.

- costi non elevati

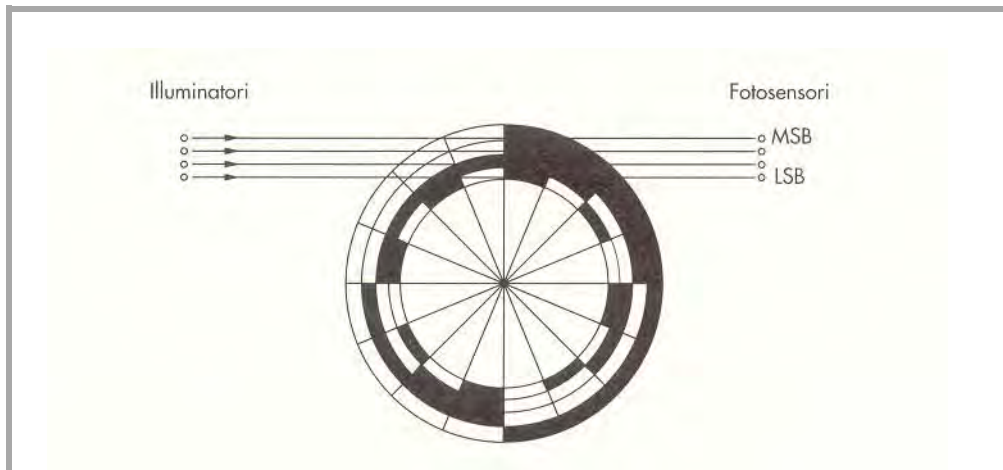
Il difetto fondamentale è la sensibilità alla temperatura con conseguente alterazione dei parametri elettrici.

Si pone rimedio a ciò mediante circuiti di *compensazione termica*, cioè *circuiti* che compensano le variazioni di resistenza indesiderate (cioè quelle dovute a effetti termici) in modo che *la variazione di resistenza complessiva dipenda solo da fattori meccanici e non termici*.

TRASDUTTORI DI POSIZIONE

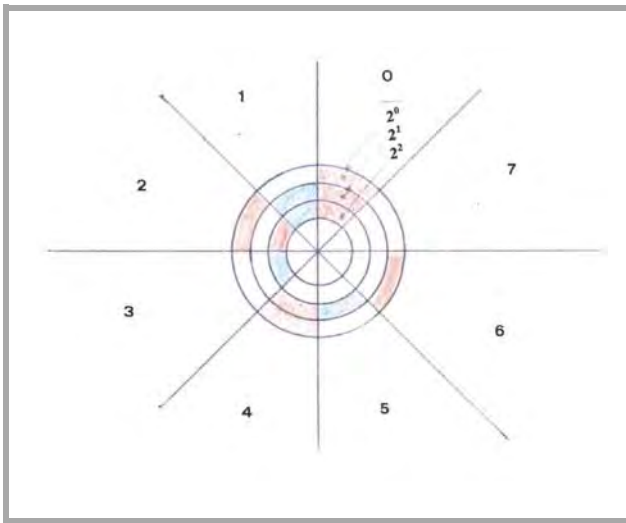
ENCODER ASSOLUTO

E' un trasduttore digitale di posizione che ci fornisce, sotto forma di segnale digitale binario, l'informazione relativa alla posizione assunta da un organo rotante (per esempio un albero motore).



E' costituito da un disco trasparente (solidale all'albero) suddiviso in un certo numero di corone circolari e settori circolari ("spicchi").

Per esempio 3 corone circolari e 8 settori, con ciascun settore corrispondente a un angolo di $(360/8)^\circ$ cioè di 45° .



In ogni settore (o “spicchio”) avremo 3 elementi, alcuni dei quali possono essere resi opachi (anneriti). In questo modo possiamo associare una informazione agli elementi trasparenti o opachi. Possiamo per es. attribuire il significato di “1” agli elementi trasparenti e il significato di “0” agli elementi **opachi**.

Così a ogni “spicchio” viene assegnato un numero espresso mediante 3 elementi trasparenti o opachi e quindi mediante 3 bit.

Detto N il numero delle corone (nell’esempio dell’ultima figura è $N=3$) e quindi dei bit disponibili, si possono codificare 2^N (in questo caso $2^3=8$) posizioni.

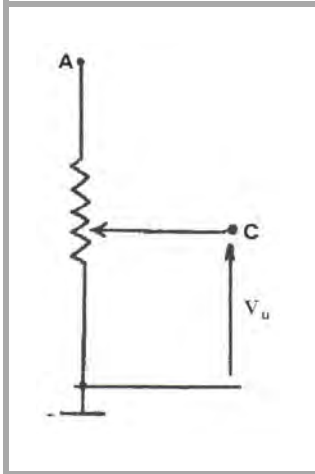
Da una parte del disco vi sono 3 diodi emettitori di luce (LED). Dall’altra parte sono posti dei ricevitori di luce (fotodiodi o fototransistor).

Avendo indicato con N il numero dei bit e quindi delle corone circolari, il numero dei settori, e quindi delle posizioni codificabili sarà 2^N .

POTENZIOMETRI COME TRASDUTTORI DI POSIZIONE

Un trasduttore di posizione analogico può essere semplicemente costituito da un potenziometro, cioè da un resistore variabile a 3 terminali.

Il potenziometro fornisce una tensione V_u proporzionale allo spostamento rettilineo di un corpo vincolato al cursore C.



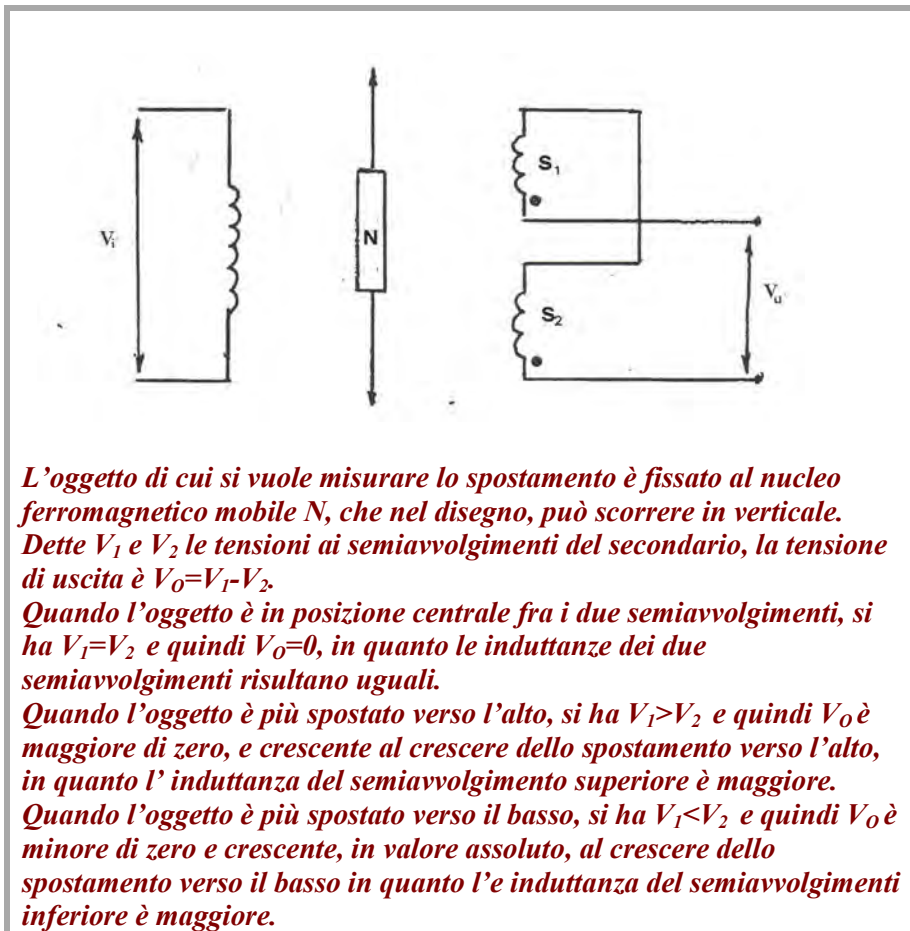
Esistono anche potenziometri in grado di fornire tensioni proporzionali a uno spostamento angolare α . In essi, se il corpo è vincolato al cursore, si preleva una tensione di uscita proporzionale allo spostamento del corpo:

$$V_u = V_i \cdot (\alpha / \alpha_{\max})$$

dove:

- $\alpha_{\max}$ = angolo relativo alla resistenza totale e quindi alla massima escursione possibile del corpo in movimento
- α = angolo relativo all'escursione del cursore determinata dal corpo in movimento

TRASFORMATORI DIFFERENZIALI o LVDT (Trasformatori Differenziali Variabili Linearmente)



Sono trasduttori di posizione che forniscono al circuito secondario una tensione proporzionale allo spostamento di un organo meccanico solidale a un nucleo mobile di materiale ferromagnetico.

Sono costituiti da:

- un **circuito primario** alimentato da una **tensione alternata**
- un **secondario** composto da due **sezioni collegate in serie**
- un nucleo mobile che varia l'accoppiamento fra il primario e il secondario, determinando una variazione del coefficiente di accoppiamento K (il cui valore³ è compreso fra 0 e 1) e quindi della tensione al secondario. Infatti la tensione V_0 ai terminali di uscita dell'avvolgimento secondario di una mutua induttanza ha un'espressione del tipo:

$$\overline{V}_0 = j \cdot \omega \cdot M \cdot \overline{I}_1$$

³
 $k = \frac{M}{L_1 \cdot L_2}$

con:

I_1 = *corrente del circuito primario*

e con :

$$M = K \sqrt{L_1 \cdot L_2}$$

Nel **caso ideale** in cui tutte le linee di flusso magnetico del primario si concatenano al secondario si ha:

$$K = 1$$

ed

$$M = \sqrt{L_1 \cdot L_2} = M_{MAX}$$

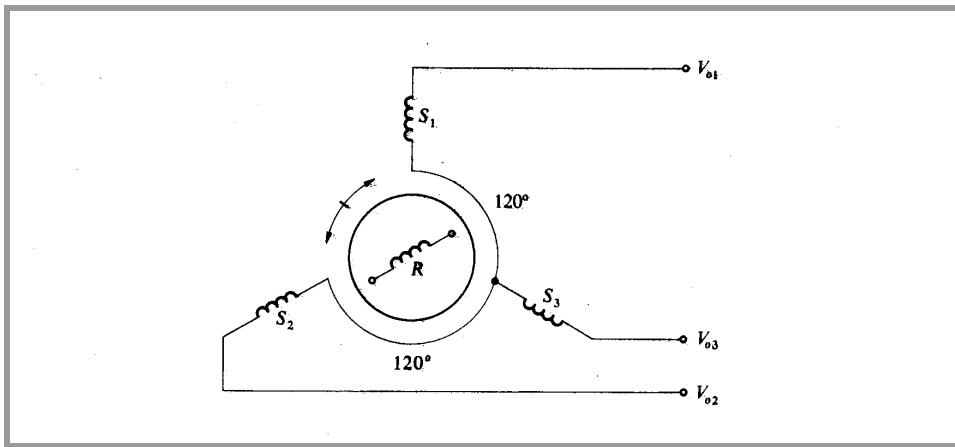
Nelle **situazioni reali** invece si ha:

$$M = K \sqrt{L_1 \cdot L_2} < M_{MAX}$$

in quanto **K è minore di 1** e dipende dall'efficienza dell'accoppiamento induttivo.

SINCRO

COME E' FATTO UN SINCRO



E' una piccola **macchina elettrica rotante** che funziona **in alternata**, utilizzata come **trasduttore di spostamento o di posizione**.

Il sincro (o syncro) può anche essere considerato come un particolare **tipo di trasformatore** ad avvolgimenti mobili che, attraverso la modifica del coefficiente M di mutua induzione, converte in segnale elettrico la posizione angolare di un albero meccanico.

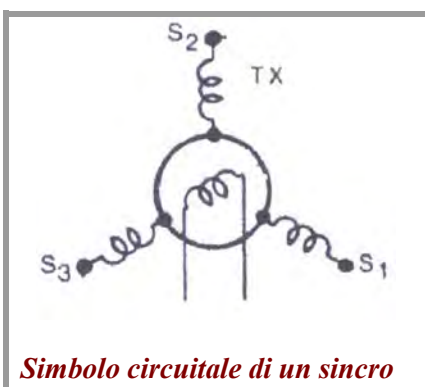
L'aspetto esterno è simile a quello di un generatore sincrono trifase.

Tuttavia il sincro differisce dal generatore sincrono trifase per diversi motivi, tra i quali:

- il rotore del sincro è alimentato in alternata, mentre quello del gen. sincrono è alimentato in c.c.
- nel sincro si hanno tensioni statoriche non nulle anche a rotore fermo, il che non avviene nel gen. sincrono
- le tensioni del sincro non sono sfasate di 120° , ma sono in fase o in opposizione di fase

Lo **statore** ha **tre avvolgimenti** disposti spazialmente a 120° uno dall'altro, alimentati in alternata a 50Hz o 60Hz o 400Hz

Il **rotore** ha un **unico avvolgimento**, alimentato (anch'esso in alternata) mediante **anelli e spazzole**.



A COSA SERVE

Un singolo sincro è un **trasduttore di posizione** che **trasforma uno spostamento angolare in segnale elettrico**.

Spesso –anche in campo navale- i sincro sono **usati in coppia per trasmettere a distanza l'informazione relativa a uno spostamento angolare α** al fine di consentire una telemisura dell'angolo o per imporre una rotazione di α gradi a un meccanismo remoto (servomeccanismo).

Se, per esempio, si vuole far ruotare un meccanismo (a pilotaggio elettrico), situato in un locale diverso da quello da cui si impartisce il comando, si usa una coppia di sincro.

Nel locale di comando un sincro-trasmettitore trasforma in segnale elettrico l'angolo impostato dall'operatore.

Il segnale elettrico viene inviato (su un'apposita linea che collega i rotori dei due sincro: trasmettitore e ricevitore), al sincro-ricevitore.

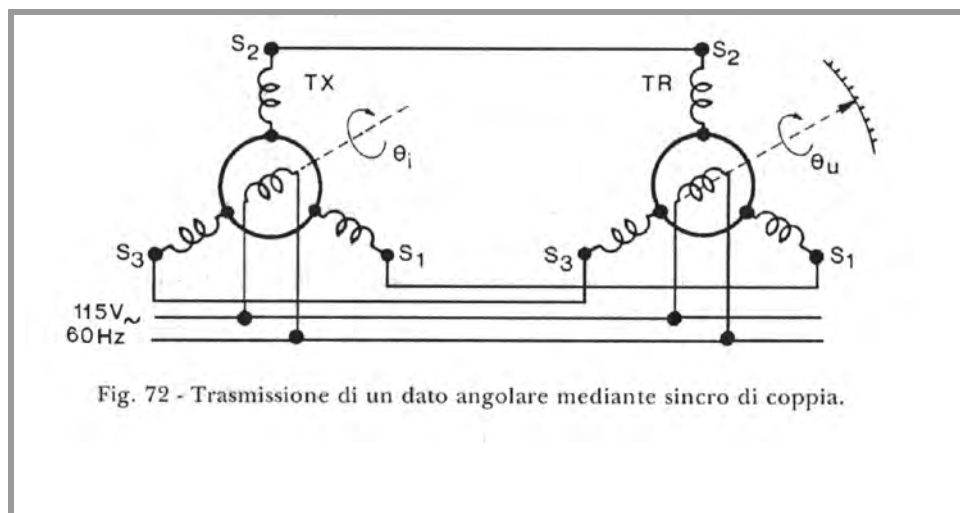
Il sincro-ricevitore infine fornisce un segnale elettrico di pilotaggio al meccanismo che si vuol far ruotare dell'angolo α .

In altri termini, se nel sincro-trasmettitore spostiamo il rotore di un certo angolo, si determina una variazione delle tensioni trifasi fra gli avvolgimenti dello statore del trasmettitore stesso.

Siccome ciascuno dei tre avvolgimenti dello statore del trasmettitore è collegato al corrispondente avvolgimento statorico del ricevitore, la variazione delle tensioni trifasi statoriche del trasmettitore si trasmette agli avvolgimenti di statore del ricevitore.

L'applicazione di questo segnale al sincro-ricevitore determina uno spostamento del rotore dello stesso angolo impostato sul trasmettitore. Se il sincro-ricevitore serve a **spostare l'indice** (lancetta) di un **indicatore per telemisure**, fornendo una coppia meccanica allo strumento, si parla di “**sincro di coppia**”

Se il sincro-ricevitore deve solo a **fornire un segnale elettrico di pilotaggio** per un servomeccanismo si parla di “**sincro di controllo**”



In campo navale i sincro sono utilizzati:

- per i “**telegrafi di macchina**”, ossia per effettuare la trasmissione di ordini relativi all'andatura della nave dalla plancia (o dalla centrale di propulsione) alla macchine
- come **ripetitori di angoli di barra e di timone**
- come **ripetitori di girobussola**
- come **ripetitori di solcometro**

RESOLVER O RISOLUTORI

Simili ai syncro, convertono la rotazione angolare di un rotore in un segnale elettrico corrispondente al seno o al coseno dell'angolo descritto dal rotore.

Esistono anche risolutori che eseguono, a partire dall'angolo di rotazione, elaborazioni più complesse, come, per esempio, la conversione delle coordinate cartesiane in coordinate polari.

TRASDUTTORI DI SPOSTAMENTO CAPACITIVI

Sono basati sulla variazione di capacità di un condensatore, a un'armatura del quale viene fissato l'oggetto di cui si vuole rilevare la posizione.

Quando l'oggetto si muove, varia la distanza fra le armature e di conseguenza varia il valore di capacità.

La variazione del valore di capacità può essere ottenuta anche facendo in modo che l'oggetto mobile modifichi la superficie di un'armatura o il valore della costante dielettrica.

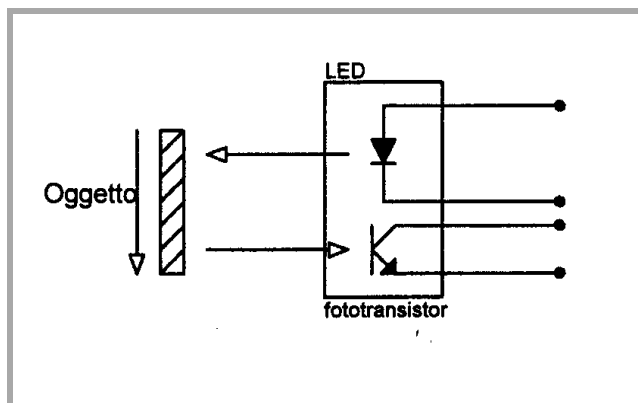
Il condensatore variabile può essere inserito in un circuito che converta, per esempio, la variazione di capacità in variazione di tensione.

Un inconveniente dei trasduttori capacitivi è la notevole sensibilità ai disturbi dell'ambiente.

TRASDUTTORI DI SPOSTAMENTO OTTICI A RIFLESSIONE

Sono basati su un componente che emetta luce (un led) e da un fotorivelatore (cioè un componente, come il fotodiodo o il fototransistor, il cui comportamento dipende dalla luce che lo colpisce).

L'oggetto di cui si vuole rilevare lo spostamento deve essere riflettente o deve essergli applicata una superficie riflettente.



Quando l'oggetto si trova di fronte al trasduttore riflette la luce del led.

La luce colpisce il fotorivelatore che produce una corrente non nulla.

Man mano che l'oggetto si allontana, riflette, verso il rivelatore, una quantità sempre minore di luce. Di conseguenza diminuisce la corrente prodotta dal rivelatore stesso.

Quando l'oggetto si allontana tanto da non essere più colpito dalla luce del led, non vi è riflessione e il fotorivelatore è interessato da corrente nulla.

TRASDUTTORI DI SPOSTAMENTO A ULTRASUONI

Sono costituiti da un dispositivo che emette onde elettromagnetiche di frequenza ultrasonica, compresa cioè nell'intervallo (40÷400) KHz e un dispositivo in grado di rilevare le onde eventualmente riflesse da un ostacolo, ostacolo che è proprio l'oggetto del quale si vuole determinare la distanza.

La rilevazione della distanza dell'oggetto avviene attraverso la misura del tempo impiegato dall'onda e.m. a colpire l'oggetto e a tornare indietro verso il sensore: più lontano è l'oggetto, maggiore è il tempo necessario all'onda per il percorso di andata e ritorno.

Siccome:

$$velocità = distanza / tempo$$

da cui:

$$distanza = velocità * tempo$$

conoscendo la velocità delle onde e.m. nel vuoto che è $c = 3 \cdot 10^8$ m/s

e misurando il tempo T di andata e ritorno si ha:

$$d = \frac{1}{2} v \cdot T$$

dove c'è la divisione per 2 perché il tempo misurato è quello di andata e ritorno e quindi è un tempo doppio del tempo proporzionale alla distanza.

Il principio di funzionamento di questi trasduttori è analogo a quello del radar.

Un aspetto positivo di questi trasduttori è (grazie alle elevate frequenze di lavoro) l'immunità dai "normali" rumori industriali.

PIEZOELETTRICITA' E MAGNETOSTRIZIONE

La piezoelettricità è una **conversione** fra **energia elettrica** ed **energia meccanica** e consiste in **due fenomeni, uno inverso dell'altro**:

- quando il cristallo è sottoposto a una deformazione meccanica, genera una tensione elettrica
- quando il cristallo è sottoposto a una tensione elettrica, subisce una deformazione meccanica.

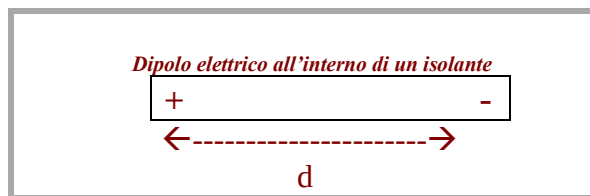
La piezoelettricità si manifesta in **alcuni materiali cristallini**, detti **cristalli piezoelettrici** e nelle cosiddette **ceramiche piezoelettriche**, che sono **insiemi di microcristalli**.

Le ceramiche piezoelettriche possono essere basate sui seguenti materiali:

- **soluzioni solide** di **Titanato-Zirconato di Piombo**⁴, che rappresentano la scelta attualmente più diffusa e che sono indicate con la sigla **PZT**
- **Titanato di Bario**⁵ (materiale di base della prima ceramica piezoelettrica che sia stata realizzata).

I cristalli piezoelettrici:

- non devono presentare un centro di simmetria nella loro struttura
- devono essere in grado, sotto l'effetto di un campo elettrico esterno, di **allineare** in una stessa direzione gran parte dei **momenti di dipolo elettrico**⁶ presenti al loro interno, cioè gran parte delle **coppie di cariche positive e negative** che si trovano al loro interno.



⁴ $\text{Pb}(\text{Ti}, \text{Zr})\text{O}_3$

⁵ BaTiO_3

⁶ Il **momento di dipolo elettrico** o **asse polare** può essere definito come l'**effetto** di una **coppia di cariche elettriche** uguali (una positiva e l'altra negativa) poste a distanza d una dall'altra:

MOMENTO “p” di DIPOLO elettrico:

vettore caratterizzato da

- modulo $|p| = Q \cdot d$ (“carica” per “distanza fra le cariche”)
- direzione: direzione coincidente con quella di d , che va dalla carica negativa verso la positiva

Per **dominio** si intende una piccola zona di materiale nella quale tutti i momenti di dipolo elettrico sono allineati nella stessa direzione

REALIZZAZIONE DELLE CERAMICHE PIEZOELETTRICHE

La ceramica piezoelettrica, per essere “**attiva**”, cioè per manifestare effettivamente la proprietà di piezoelettricità, deve essere sottoposta a un processo di lavorazione, che in parte si svolge ad **alte temperature**⁷ e in parte prevede l’uso di **campi elettrici**.

I passi principali della realizzazione di una ceramica piezoelettrica sono:

- MISCELAZIONE
- ESSICCAZIONE
- **CALCINAZIONE** (introduzione in un **forno** a 800÷900°C) per far avvenire delle reazioni chimiche che rendono omogeneo il materiale
- MACINAZIONE (riduzione in granuli di 5 micron)
- SECONDA ESSICCAZIONE
- GRANULAZIONE con aggiunta di leganti
- STAMPAGGIO (mediante **pressa** per dare la forma voluta e migliorare le caratteristiche elettriche
- **SINTERIZZAZIONE** a **1100°C** (è un trattamento **termico** delle polveri, che serve, fra l’altro, per rimuovere le porosità, conferire compattezza e durezza)
- **METALLIZZAZIONE** di due superfici del pezzo (realizzazione degli elettrodi)
- **POLARIZZAZIONE** (mediante **campo elettrico** esterno di **alcuni KV/mm** per rendere la ceramica **piezoelettricamente attiva**): viene effettuata alla temperatura di 100°C mediante **l’applicazione di un campo elettrico** che determina un’espansione permanente del materiale nella direzione del campo e una contrazione nelle due direzioni opposte.

⁷ Temperature che devono comunque essere inferiori al “punto di Curie” del materiale, in corrispondenza del quale il materiale perde irreversibilmente le proprietà piezoelettriche

L'EFFETTO PIEZOELETTRICO

Quando gli viene applicata una tensione **sinusoidale**, il cristallo piezoelettrico **vibra meccanicamente** in regime di **oscillazioni forzate**.

Quando gli viene applicata una tensione continua, il cristallo subisce simultaneamente un'espansione in una direzione e una compressione nella direzione perpendicolare, deformazioni che permangono fino a che la tensione è applicata.

In particolare, quando gli viene applicata una tensione **continua** con la **stessa polarità** (ma ampiezza minore) della **tensione di polarizzazione** subita in fase di **lavorazione**, il cristallo subisce una temporanea **espansione** lungo i **piani paralleli** a quelli degli elettrodi e una temporanea contrazione nella direzione opposta.

Quando gli viene applicata una tensione **continua** con polarità opposta (ma ampiezza minore) a quella della tensione di polarizzazione subita in fase di lavoro, il cristallo subisce una temporanea contrazione lungo i piani paralleli a quelli degli elettrodi e una temporanea espansione nella direzione opposta.

Le deformazioni, a meno che non ci si trovi alla frequenza di risonanza meccanica, sono molto piccole (decimi di micron).

APPLICAZIONI DELL' EFFETTO PIEZOELETTRICO

Le principali applicazioni dell'effetto piezoelettrico sono:

- Il **SONAR** (SOund NAvigation and Ranging), basato sulla trasmissione di onde di pressione, soniche o ultrasoniche, nell'acqua, e sulla riflessione, da parte di un ostacolo, di tali onde
- **microfoni** e capsule microfoniche
- avvisatori acustici (**buzzer**)
- **ecografia medica**
- **ultrasuonoterapia** (per artriti, reumatismi, pulizia del viso e del corpo)
- **accelerometri** (**pace-maker**, **airbag**)
- **sensori antidetonazione**
- **lavaggi industriali** di componenti elettronici, ferri chirurgici ecc.
- prove di **integrità dei materiali** (per esempio **scafi di imbarcazioni in vetroresina**)

MAGNETOSTRIZIONE

In qualche caso (quando, come nel lavaggio a ultrasuoni, è richiesta una notevole potenza) la generazione di ultrasuoni non è realizzata mediante un trasduttore piezoelettrico ma attraverso la magnetostrizione.

La magnetostrizione è un fenomeno che, come la piezoelettricità, comporta la **deformazione meccanica** di un materiale. Questa deformazione però è dovuta al campo **magnetico** e non (come avviene per la piezoelettricità) al campo elettrico.

La magnetostrizione si manifesta solo in materiali **ferromagnetici**, come l'**oro**, il **nichel** e il **cobalto** e le loro **leghe**.

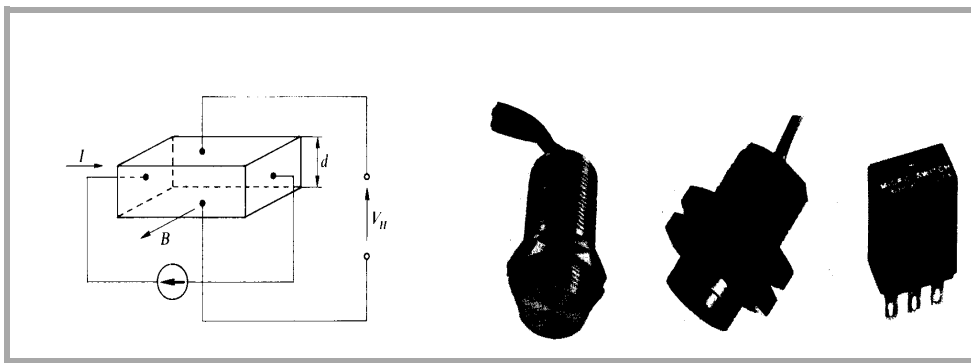
TRASDUTTORI DI POSIZIONE A EFFETTO HALL

Si basano sul principio seguente:

Se una barretta di semiconduttore⁸, attraversata da una corrente I , viene immersa in un campo magnetico B di induzione, allora, sulle cariche costituenti la corrente I , si manifesta una forza (forza di Lorenz) che tende a spingere le cariche lungo una direzione perpendicolare al piano individuato dalla corrente I e da B .

In presenza di un campo magnetico si manifesta quindi, tra le superfici della barretta, una ddp, sempre lungo una direzione perpendicolare al piano di I e di B .

Se all'oggetto da rilevare viene applicato un elettromagnete⁹, ossia una sorgente di campo, l'avvicinarsi dell'oggetto determina l'aumento della ddp.



(Nella figura il piano di I e di B è perpendicolare al foglio, quindi le cariche risultano spinte, a seconda del segno, in alto o in basso lungo il piano del foglio).

⁸ la barretta potrebbe anche essere metallica, ma l'effetto Hall sarebbe molto meno rilevante

⁹ il campo B potrebbe anche essere generato da un circuito solidale al trasduttore, in tal caso l'oggetto (di natura ferrosa) provocherebbe una variazione della riluttanza del circuito.

L'espressione della ddp di Hall è:

$$V_H = k_H \cdot \frac{B \cdot I}{d}$$

dove:

d = spessore della barretta

K_H = coefficiente di Hall, basso nei metalli ed elevato nei semiconduttori, in particolare nell'arseniuro di gallio.

Esistono trasduttori di Hall,

- “lineari” cioè in grado di fornire un'uscita proporzionale all'entità del campo magnetico (e quindi variabile con continuità al variare della posizione dell'oggetto)
- “non lineari” ossia “on-off”, in grado cioè di rilevare solo la presenza o l'assenza dell'oggetto, utilizzati come interruttori di prossimità.

Gli aspetti positivi dei trasduttori di Hall sono:

- alta frequenza di funzionamento (fino a 100 KHz)
- ampio range di temperature di lavoro (-40, +150)°C
- maggiore economia e robustezza rispetto ad altri trasduttori di posizione, come gli encoder (la risoluzione dei trasduttori di Hall è però minore).

INTERRUTTORI DI PROSSIMITA'

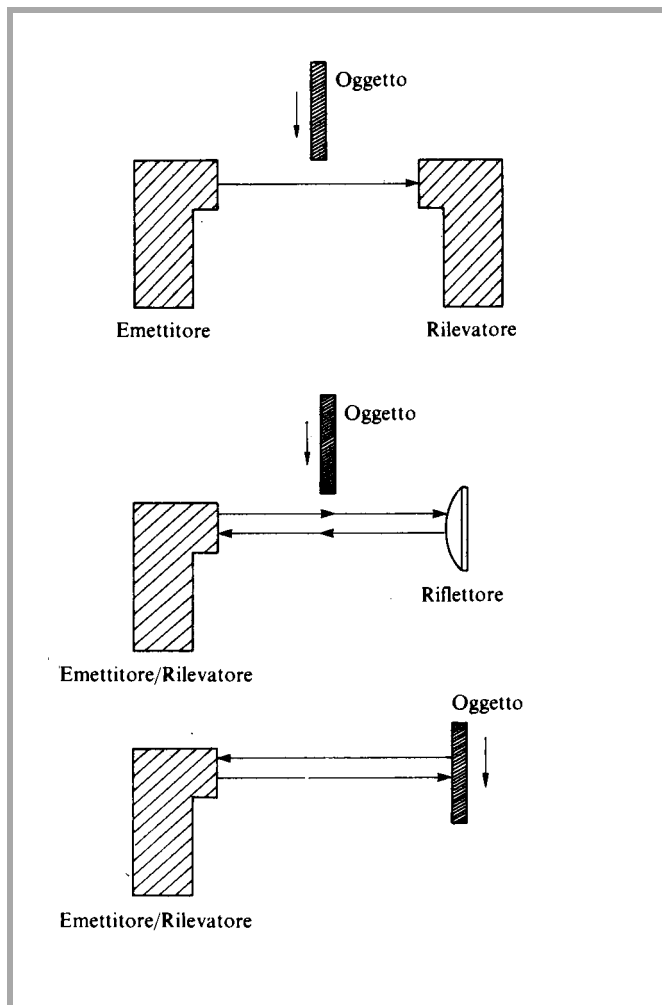
Sono sensori on-off che rivelano la vicinanza o la non vicinanza di un oggetto, e perciò sono impiegati come rivelatori di presenza/assenza, rivelatori di passaggio, finecorsa ecc.

Analogamente ai trasduttori “lineari” o “proporzionali” di posizione già analizzati possono essere:

- o elettromeccanici (richiedono il contatto fisico con l'oggetto)
- o pneumatici (richiedono il contatto fisico con l'oggetto)
- o induttivi
- o capacitivi
- o a effetto hall
- o ottici

Accenneremo qui solo agli interruttori di prossimità ottici.

INTERRUTTORI DI PROSSIMITA' OTTICI O FOTOELETTRICI



Nella figura più in alto è illustrato un interruttore ottico

A SBARRAMENTO.

Una sorgente ottica (emettitore) emette un raggio luminoso verso un rilevatore.

Quando fra emettitore e rilevatore non è presente alcun oggetto, il raggio luminoso non è interrotto e il rilevatore non emette alcun segnale.

Quando fra emettitore e rilevatore si interpone un oggetto opaco, il raggio luminoso è interrotto e il rilevatore emette un segnale.

La distanza di lavoro è generalmente di (2÷30) m.

Nella figura centrale è illustrato un interruttore ottico A RIFLESSIONE (REFLEX).

Il sensore è costituito, da una parte, da un dispositivo emettitore/rilevatore e, dalla parte opposta, da un riflettore.

Quando fra emettitore/rilevatore e riflettore non è presente alcun oggetto, il raggio luminoso viene riflesso dal riflettore e torna verso l'emettitore/rilevatore.

Quando invece fra emettitore/rilevatore e riflettore si interpone un oggetto (che deve essere NON riflettente), il raggio luminoso NON torna verso l'emettitore/rilevatore.

La distanza di lavoro è generalmente di (2÷10) m.

Nella figura inferiore è illustrato un interruttore ottico A DIFFUSIONE.

Analogamente al reflex, sensore è costituito, da una parte, da un dispositivo emettitore/rilevatore, ma dalla parte opposta, non vi è alcun riflettore.

Quando fra emettitore/rilevatore e riflettore non è presente alcun oggetto, il raggio luminoso NON torna verso l'emettitore/rilevatore.

Quando invece fra emettitore/rilevatore e riflettore si interpone un oggetto (che deve essere RIFLETTENTE), il raggio luminoso torna verso l'emettitore/rilevatore.

La distanza di lavoro è generalmente di (20cm ÷ 1,5m).

TRASDUTTORI DI VELOCITA'

Nelle prime pagine del capitolo sono stati trattati:

- un trasduttore di velocità analogico (la **dinamo tachimetrica**)
- un trasduttore di velocità digitale (l'**encoder tachimetrico**).

Un ulteriore trasduttore *digitale* di velocità è l'*encoder incrementale* che fornisce un segnale che dà conto anche del *verso* del movimento meccanico (rotazione).

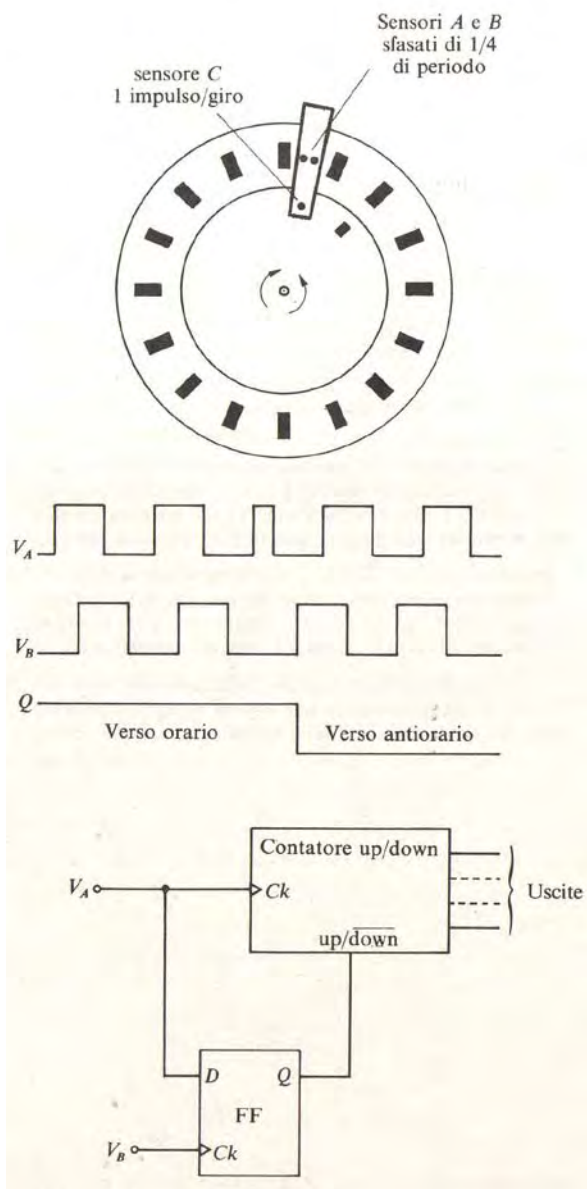
ENCODER INCREMENTALE

Ha la stessa struttura dell'encoder tachimetrico "normale", però non possiede un solo sensore, ossia una sola coppia led-fotodiodo, ma tre.

I sensori A e B servono per il rilievo della posizione, il sensore C per la velocità.

Se il disco dell'encoder gira in senso orario, il sensore A (coppia led-fotodiodo) produce un'onda rettangolare con impulsi alti aventi un anticipo di un quarto di periodo rispetto a quelli dell'onda quadra prodotta dal sensore B (altra coppia led-fotodiodo) in quanto i fori del disco incontrano, ruotando, prima A e poi B.

Se invece il disco ruota in senso antiorario i fori incontrano prima il sensore B e poi A, per cui la tensione di uscita v_a (onda rettangolare) prodotta da A, risulta in ritardo rispetto alla v_b prodotta da B.



Encoder incrementale

Dallo schema circuitale si vede che

- il segnale v_a del sensore A viene applicato simultaneamente all'ingresso-dati di un flip-flop (FF) di tipo D e all'ingresso di clock di un contatore bidirezionale
- il segnale v_b del sensore B viene applicato all'ingresso di clock del flip-flop, clock che abilita gli ingressi in corrispondenza dei fronti di salita di v_b .

L'uscita Q del FF è poi collegata al terminale di commutazione UP/DOWN del contatore bidirezionale che ci indicherà il verso di rotazione del movimento meccanico contando in avanti o all'indietro.

Dai simboli circuitali si vede che il FF è sincronizzato sui fronti di salita del clock e che quindi l'uscita può commutare solo negli istanti relativi a tali fronti.

Il FF è di tipo D e quindi funziona in questo modo:

- quando il FF è abilitato, la sua uscita assume il valore dell'ingresso applicato in quell'istante
- se invece è disabilitato, l'uscita conserva il valore precedente.

L'onda quadra v_a (proveniente dal sensore A) è applicata all'ingresso D del FF, mentre l'onda quadra v_b è utilizzata come segnale di clock per il FF stesso.

Se il segnale v_b di clock del FF (a causa della rotazione oraria del disco) è sempre in ritardo di $T/4$ rispetto al segnale di ingresso v_a del FF, allora i fronti di salita del clock vengono a prodursi sempre in istanti in cui il segnale v_a è già alto, quindi il dispositivo acquisisce (in ognuno di questi istanti) un livello alto che si propaga in uscita.

D'altra parte, se consideriamo tutti gli istanti successivi a quelli dei fronti di salita (in cui l'uscita rimane alta), in tali istanti l'uscita non può cambiare, qualunque sia stato l'ingresso, e quindi rimane alta.

In conclusione, il fatto che i fronti di salita del clock capitano sempre quando l'ingresso D è alto, fa sì che l'uscita sia sempre alta.

L'uscita alta del FF, applicata al terminale UP/DOWN del contatore, fa in modo che il conteggio proceda in avanti.

Se invece il segnale v_b di clock del FF (a causa della rotazione antioraria del disco) è sempre in anticipo di $T/4$ rispetto al segnale di ingresso v_a del FF, allora, in corrispondenza dei fronti di salita del clock, si ha ingresso basso e quindi uscita bassa.

Negli istanti successivi a quelli dei fronti di salita l'uscita mantiene il valore precedente, cioè quello basso, per cui il segnale è sempre basso.

In sintesi:

- Rotazione oraria $\rightarrow$ onda v_b (clock del FF) in ritardo $\rightarrow$ uscita Q (del FF) alta $\rightarrow$ conteggio in avanti (UP)
- Rotazione antioraria $\rightarrow$ onda v_b (clock del FF) in anticipo $\rightarrow$ uscita Q (del FF) bassa $\rightarrow$ conteggio all'indietro (DOWN).

ACCELEROMETRO (TRASDUTTORE DI ACCELERAZIONE)

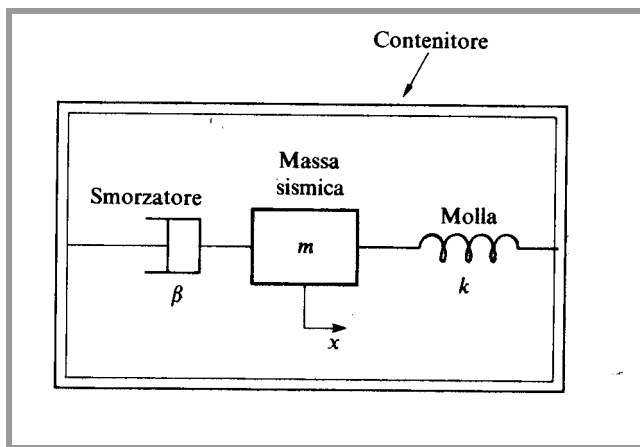
E' un trasduttore di accelerazione e quindi anche di decelerazione e di vibrazioni. Osserviamo che la decelerazione e/o le vibrazioni possono precedere la collisione di un veicolo con un ostacolo.

Gli accelerometri possono essere basati su differenti fenomeni fisici. Esistono, per esempio, accelerometri **meccanici** (come quello di seguito descritto) o piezoelettrici.

La misura dell'accelerazione avviene attraverso una misura di spostamento.

In particolare viene misurato lo spostamento di un corpo, detto "massa sismica" che è sospeso mediante una molla e uno smorzatore (pistone) all'interno di un contenitore solidale al veicolo o all'oggetto di cui si vuole rilevare l'accelerazione.

Per esempio il contenitore può essere fissato a bordo di un'auto e può essere utilizzato per l'azionamento di un air bag.



Come è possibile misurare l'accelerazione attraverso una misura di spostamento?

La massa sismica (sulla quale l'accelerometro è basato) è fissata al contenitore attraverso una molla elastica K. Perciò lo spostamento della

massa viene a coincidere con lo spostamento x (allungamento o compressione) dell'estremo libero della molla.

Riguardo alla molla ricordiamo che vale la legge di Hooke:

$$F = K \cdot x$$

dove:

F = sollecitazione (forza) cui la molla è sottoposta

K = costante elastica

x = spostamento dell'estremo libero della molla.

Tenendo conto della 2^a legge della dinamica: $F = m \cdot a$

e sostituendo nella legge di Hooke otteniamo:

$$m \cdot a = K \cdot x$$

da cui:

$$a = \frac{K}{m} \cdot x$$

dove:

m = massa sismica

K = costante elastica della molla

a = accelerazione

Da quest'ultima relazione si vede che:

l'accelerazione è proporzionale allo spostamento x attraverso la costante k/m e quindi può essere rilevata attraverso una misura di spostamento.

La misura di spostamento può essere effettuata:

- con un trasformatore differenziale (LVDT) in caso di basse frequenze, ossia fino a 100 Hz
- con un estensimetro, per frequenze fino a 10 KHz
- con un trasduttore piezoelettrico per frequenze fino a 10 KHz. Ricordiamo che la piezoelettricità è la proprietà (che alcuni materiali, come il quarzo, hanno) di generare una tensione elettrica se sottoposti a sollecitazioni meccaniche quali compressione, flessione o stiramento.

TRASDUTTORI DI LIVELLO

Misurano il livello di un liquido contenuto in un recipiente e possono essere essenzialmente di due tipi:

1. misuratori di pressione differenziale
2. misuratori di variazione di capacità elettrica.

MISURATORI DI PRESSIONE DIFFERENZIALE

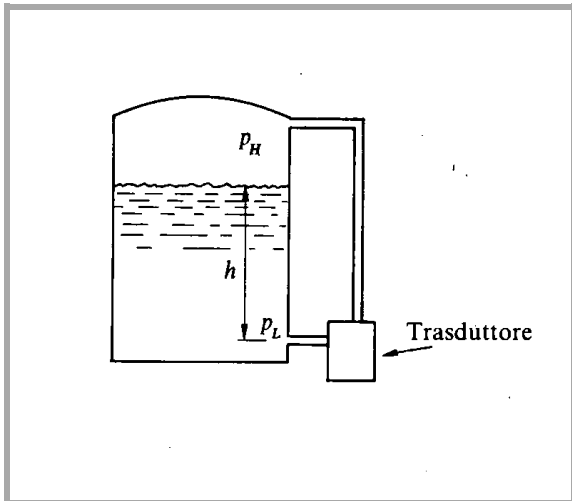
Conoscendo il peso specifico w del liquido, il livello h viene misurato rilevando due pressioni:

p_H = pressione al di sopra del livello del liquido

p_L = pressione in corrispondenza del fondo del recipiente

e realizzando il calcolo seguente:

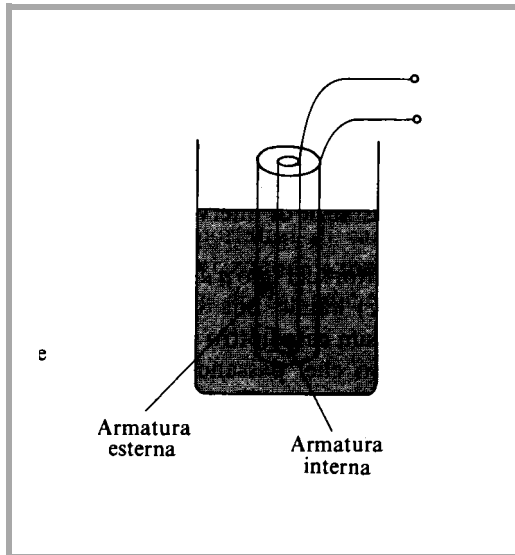
$$h = \frac{p_H - p_L}{w}$$



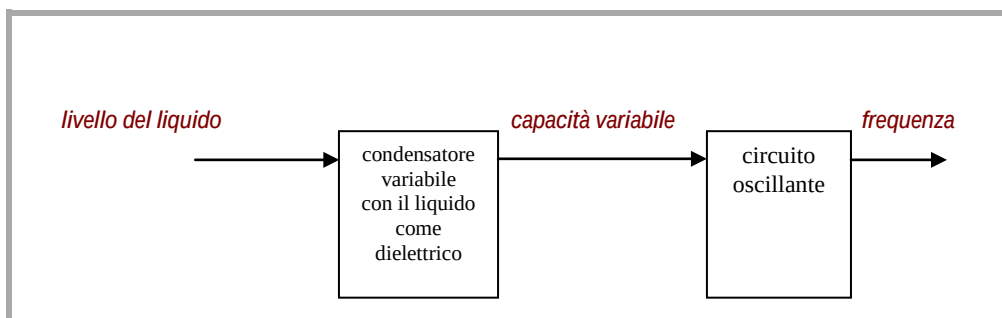
MISURATORI DI CAPACITÀ ELETTRICA

All'interno del serbatoio viene realizzato un condensatore il cui dielettrico è proprio il liquido del quale si vuole misurare il livello.

Al variare del livello del liquido varia la capacità del condensatore.



Il condensatore così realizzato può poi essere inserito in un circuito oscillante, ottenendo così una frequenza proporzionale al livello del liquido.



TRASDUTTORI DI ENERGIA RADIANTE. DISPOSITIVI FOTOELETTRICI E PIROELETTRICI

I trasduttori di energia radiante o fotorivelatori sono quei dispositivi che convertono una radiazione visibile o infrarossa (ossia una radiazione elettromagnetica appartenente alla regione ottica dello spettro) in un segnale elettrico.

Sono detti “**rivelatori di fotoni**” i trasduttori di energia radiante che appartengono alla “famiglia” dei dispositivi fotoelettrici cioè dei dispositivi elettronici basati su uno dei fenomeni seguenti:

- effetto fotoemittente o fotoelettrico
- effetto fotoconduttivo
- effetti fotoconduttivo di giunzione
- effetto fotovoltaico.

Sono invece detti “**rivelatori piroelettrici**” quei fotorivelatori¹⁰ che sono basati su un particolare fenomeno chiamato “piroelettricità”.

¹⁰ Detti anche “termici”

UNA POSSIBILE CLASSIFICAZIONE

- Dispositivi fotoelettrici (rivelatori di fotoni) basati sull' effetto **FOTOEMITTENTE O FOTOELETTRICO**
 - **Cellula fotoelettrica** o cella fotoelettrica o fotocellula
 - **Fotomoltiplicatore**
 - **Intensificatore di immagine**
- Dispositivi fotoelettrici (rivelatori di fotoni) basati sull' effetto **FOTOCONDUTTIVO**
Fotoresistenza al solfuro di cadmio (CdS)
- Dispositivi fotoelettrici (rivelatori di fotoni) basati sull' effetto **FOTOELETTRICO DI GIUNZIONE E SULL'EFFETTO FOTOVOLTAICO**
 - **Fotodiodo**
 - **Fototransistor**
 - **Dispositivo fotovoltaico (Cella o pila solare)**
- Dispositivi fotoelettrici (rivelatori di fotoni) basati **SULL'EFFETTO FOTOVOLTAICO**
 - **Cella o pila solare**
- **RIVELATORI PIROELETTRICI**

EFFETTO FOTOEMITTENTE O FOTOELETTRICO

Un raggio luminoso, colpendo la superficie di un materiale, produce una emissione di elettroni.

EFFETTO FOTOCONDUTTIVO

Si ha nei semiconduttori e consiste in un aumento di conducibilità del semiconduttore quando è investito da una irradiazione di luce da parte di una sorgente esterna

EFFETTO FOTOELETTRICO DI GIUNZIONE

Si ha nei *semiconduttori drogati* e consiste nel fatto che, in una giunzione pn **inversamente** polarizzata, si ha un ***aumento della corrente inversa*** quando la ***giunzione è colpita da luce*** proveniente da una sorgente esterna

EFFETTO PIROELETTRICO

Si ha in particolari materiali ferromagnetici e consiste nell'insorgere di una ddp fra gli elettrodi di un cristallo che abbia subito una polarizzazione elettrica interna per effetto della radiazione incidente.

EFFETTO FOTOVOLTAICO

L'effetto fotovoltaico è la trasformazione diretta (senza l'ausilio di sorgenti esterne, ossia di alimentazione del componente) della radiazione luminosa in energia elettrica.

L'effetto si produce quando viene illuminato uno dei due elettrodi di un pezzo di semiconduttore opportunamente trattato.

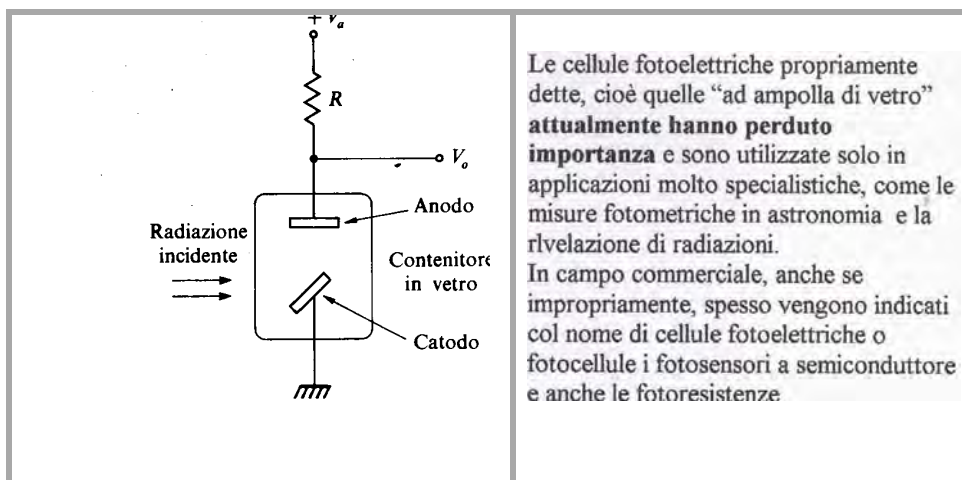
DISPOSITIVI FOTOELETTRICI (RIVELATORI DI FOTONI) BASATI SULL' EFFETTO FOTOEMITTENTE O FOTOELETTRICO

CELLULA FOTOELETTRICA O CELLA FOTOELETTRICA O FOTOCELLULA

E' un'ampolla di vetro, riempita di gas, oppure nella quale è stato realizzato il vuoto.

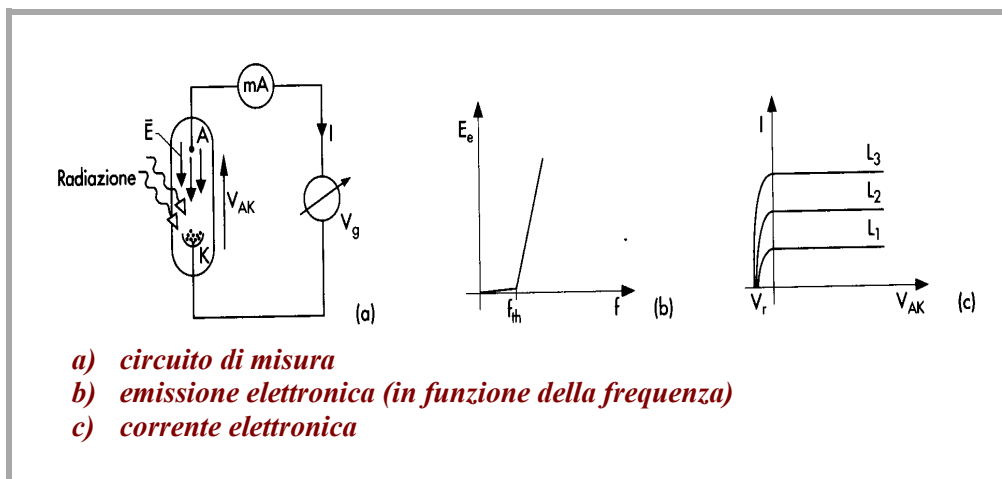
L'ampolla contiene un catodo di materiale fotoemittente che, quando viene colpito dalla luce (radiazione e.m.), emette elettroni.

Gli elettroni vengono raccolti dall'anodo e danno luogo a una corrente elettrica nel circuito.



La corrente che si produce :

- è chiamata fotocorrente
- aumenta all'aumentare dell'intensità luminosa incidente
- non dipende dal valore della tensione V_{AK} fra anodo e catodo.



Un parametro caratterizzante della fotocellula è la **sensibilità luminosa** definita come:

$$s = \frac{\text{corrente} [\mu\text{A}]}{\text{intensità di luce incidente} [\text{lumen}]}$$

La sensibilità è massima in corrispondenza di una certa lunghezza d'onda (e quindi di una certa frequenza) della luce incidente, lunghezza d'onda che dipende dal materiale di realizzazione del catodo.

La lunghezza d'onda λ per la quale la sensibilità è massima è detta “ λ di picco” della cellula.

MATERIALE CATODO	LUNGHEZZA D'ONDA DI PICCO DELLA SENSIBILITA' [nanometri, nm]	SENSIBILITA' $\left[\frac{\mu\text{A}}{\text{lumen}} \right]$
ARGENTO-OSSIDO-CESIO	800 (ir ¹¹)	25
ARGENTO-OSSIDO-RUBIDIO	420 (ir)	6,5
ANTIMONIO-CESIO	400 (ir)	40

Le fotocellule a gas offrono la possibilità di aumentare la sensibilità aumentando la tensione, sono però più “lente” di quelle a vuoto.

Le fotocellule a vuoto sono le sole adatte a essere usate come trasduttori di luce perché la loro corrente dipende linearmente dall'intensità della radiazione elettromagnetica che le colpisce.

Le fotocellule a vuoto possono essere usate anche quando la frequenza della radiazione da rilevare supera i 10 KHz.

Gli **inconvenienti** delle fotocellule sono:

- l'ingombro
- la fragilità
- l'elevata tensione di polarizzazione (circa 100 V).

¹¹ InfraRosso: le lunghezze d'onda elencate (e le corrispondenti frequenze) appartengono a quella parte dello spettro elettromagnetico detto “regione dell'infrarosso”.

FOTOMOLTIPLICATORI

Differiscono dalle fotocellule per la presenza, oltre al catodo e all'anodo, di più anodi supplementari (dinodi) disposti sequenzialmente e posti a tensioni progressivamente crescenti.

Quando gli elettroni vengono emessi dal catodo, vengono attratti, in successione, dai dinodi. Su ciascun dinodo si produce un'emissione secondaria di altri elettroni, col risultato di moltiplicare il numero di elettroni e di amplificare la corrente rispetto al valore originario. I fotomoltiplicatori hanno un rapporto segnale/rumore (S/N) molto alto e trovano applicazione

- come sensori di intensità luminose molto deboli
- come sensori di radiazioni nucleari.

TUBI INTENSIFICATORI DI IMMAGINE

Sfruttano la pochissima luce presente quando l'occhio non è in grado di vedere, moltiplicandone l'effetto fino a centomila volte.

Rendono quindi possibile la visione di oggetti invisibili a occhio nudo.

Sono usati:

- in campo militare
- nella navigazione subacquea
- nei microscopi elettronici
- nelle apparecchiature a raggi X.

DISPOSITIVI FOTOELETTRICI BASATI SULL' EFFETTO FOTOCONDUTTIVO

FOTORESISTENZA

Tipicamente le fotoresistenze sono realizzate in solfuro di cadmio e presentano un valore di resistenza che diminuisce linearmente al crescere dell'intensità luminosa.

Sono utilizzate come rivelatori di intensità luminosa quando le variazioni di intensità sono lente (il tempo di risposta è di 100 ns).

Il semiconduttore utilizzato per rilevare radiazioni luminose dello spettro VISIBILE è il SOLFURO DI CADMIO (CdS)¹² per il quale il massimo effetto fotoconduttivo si ha per lunghezze d'onda corrispondenti a radiazioni luminose con gamma di colori compresa fra il verde e il rosso.

La **lunghezza d'onda critica** λ_c della fotoresistenza è quella lunghezza d'onda della radiazione incidente che produce il **massimo effetto fotoconduttivo**:

$$\lambda_c = \frac{h \cdot c}{E_g} = \frac{1,24}{E_g}$$

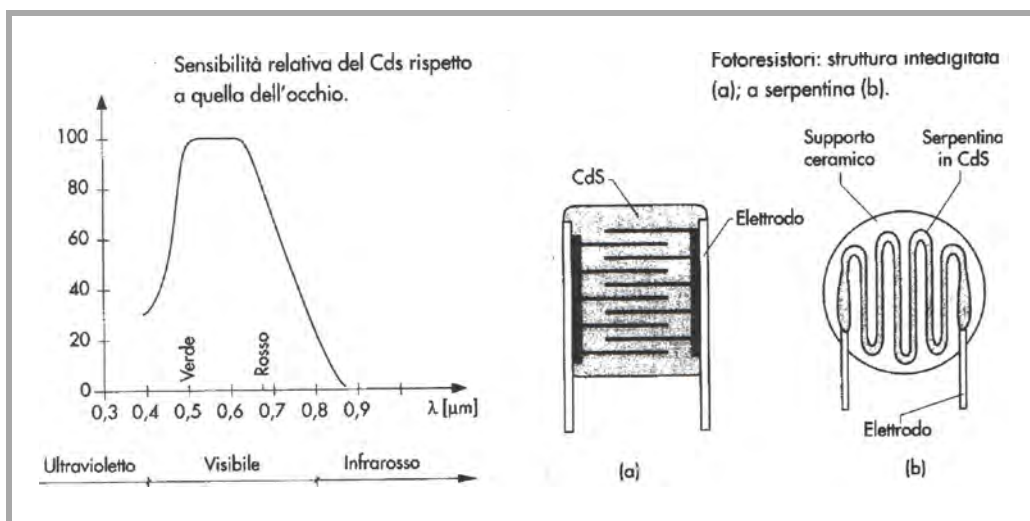
dove:

h = costante di Planck = $4,14 \cdot 10^{-15}$ [eVs]

E_g = gap energetico = energia necessaria all'elettrone per passare dalla banda di valenza alla banda di conduzione.

Le fotoresistenze sono usate, oltre che come trasduttori di intensità luminosa, come:

- controllo della fiamma negli scaldacqua a gas
- rilevatori di intrusione (antifurto)
- misura a distanza della temperatura.



¹² Per la rivelazione della radiazione VISIBILE è anche usato il seleniuro di cadmio (CdSe). Per la rivelazione dell'infrarosso sono invece usati il solfuro di piombo (PbS) e il seleniuro di piombo (PbSe).

DISPOSITIVI FOTOELETTRICI BASATI SULL' EFFETTO FOTOELETTRICO DI GIUNZIONE E SULL'EFFETTO FOTOVOLTAICO

FOTODIODO

La corrente inversa è proporzionale all'intensità luminosa incidente ed è indipendente dalla tensione di polarizzazione inversa.

Siccome il fotodiodo risponde alla radiazione con un aumento di corrente, se si vuole un segnale in tensione bisogna convertire la corrente in tensione facendola passare su una resistenza.

Come parametro caratteristico viene fornita anche la sensibilità spettrale a $0,8 \mu\text{m}$.

Rispetto alla fotoresistenza, il fotodiodo presenta:

- migliore risposta in frequenza, grazie ai brevi tempi di risposta (100 ps)
- migliore linearità nella relazione fra intensità luminosa e corrente
- corrente di buio influenzata dalla temperatura (aspetto negativo)
- superficie attiva minore (aspetto negativo)

I fotodiodi sono usati come:

- interruttori comandati dalla luce
- per poter contare il numero di interruzioni di un fascio luminoso in dispositivi di conteggio.

FOTOTRANSISTOR

La corrente di saturazione inversa è proporzionale all'intensità luminosa incidente.

La corrente di collettore è:

$$I_C = (\beta + 1) \cdot (I_{lum} + I_{term})$$

dove:

I_{lum} = componente della corrente dovuta all'effetto luminoso

I_{term} = componente della corrente dovuta all'effetto termico

Rispetto al fotodiodo, il fototransistor presenta:

- scarsa linearità
- minore larghezza di banda e quindi tempi di risposta più lunghi
- maggiore sensibilità

Per queste ragioni i fototransistor

- non sono adatti ad applicazioni di misura
- sono utilizzati come INTERRUITORI e nei casi in cui servono correnti di uscita elevate.

DISPOSITIVI FOTOELETTRICI BASATI SULL' EFFETTO FOTOVOLTAICO

DISPOSITIVI FOTOVOLTAICI (PILE O CELLE SOLARI)

Quando l'effetto fotoelettrico di giunzione si manifesta in *assenza di tensione di polarizzazione* si parla di effetto *fotovoltaico* e il componente si comporta come un “*generatore fotovoltaico*”.

Le celle fotovoltaiche o celle solari (o pile fotovoltaiche) si basano sull'effetto fotovoltaico. La pila solare può essere realizzata mediante due elettrodi separati da uno strato semiconduttore.

L'effetto fotovoltaico è la trasformazione diretta (senza l'ausilio di sorgenti esterne, ossia di alimentazione del componente) della radiazione luminosa in energia elettrica.

L'effetto si produce quando viene illuminato uno dei due elettrodi di un pezzo di semiconduttore opportunamente trattato.

Il trattamento consiste in

- drogaggio molto elevato (rispetto al normale fotodiodo)
- area superficiale molto estesa
- riduzione della riflessione.

Il principio di funzionamento è la generazione di coppie lacuna-elettrone libere e la loro separazione mediante il campo elettrico della zona di svuotamento.

L'illuminazione di uno degli elettrodi determina la nascita di una f.e.m.

La fem fotovoltaica è la differenza di potenziale che si manifesta ai capi del componente fotovoltaico non polarizzato quando viene colpito dalla luce.

L'esistenza di questa fem consente di collegare direttamente il dispositivo fotovoltaico al carico, senza bisogno di alimentazione.

La fem fotovoltaica è dell'ordine di $(0,4 \div 0,5)V$ per le celle al silicio e di $0,1 V$ per le celle al germanio

Il componente fotovoltaico fornirà al carico una corrente inversa che va dal catodo verso l'anodo e che quindi risulta uscente dall'anodo.

La tensione a vuoto ai capi della cella dipende dal logaritmo dell'illuminazione del componente.

La corrente di corto circuito varia linearmente con l'illuminazione.

I rivelatori fotovoltaici sono particolari fotodiodi che differiscono dai normali fotodiodi perché producono una corrente negativa (corrente inversa, che va dall'anodo verso il catodo) in corrispondenza di una tensione V_{AK} positiva.

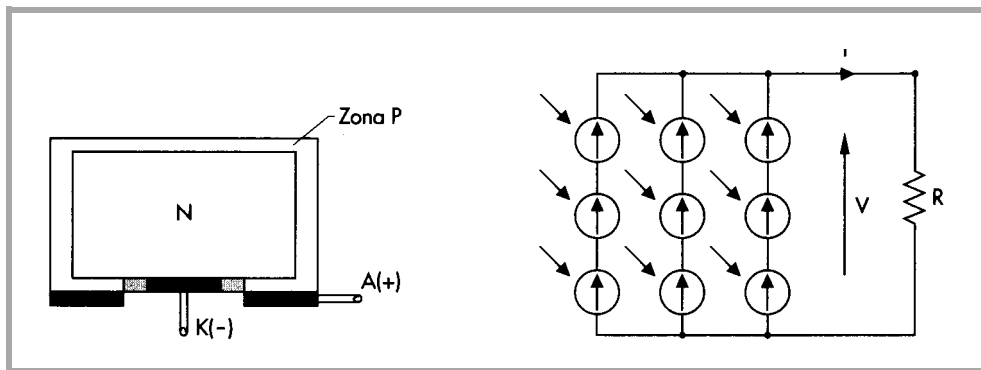
Pertanto il punto di funzionamento dei dispositivi fotovoltaici si trova nel quarto quadrante (essendo caratterizzato da tensione positiva e corrente negativa), il che significa che questi dispositivi cedono energia al carico senza bisogno di alcuna sorgente esterna di alimentazione o polarizzazione. Si comportano quindi come generatori elettrici.

Le celle solari sono realizzate in silicio o in arseniuro di gallio

La cella è costituita da un pezzo di silicio con drogaggio N, ricoperto da un sottile strato con drogaggio P.

Le correnti e le tensioni erogabili dalla singola cella sono basse.

Perciò, per aumentare la tensione, più celle vengono collegate in serie (array di celle).
Più array sono collegati in parallelo per aumentare la corrente.



RIVELATORI PIROELETTRICI

I rivelatori piroelettrici sono costituiti da un particolare cristallo¹³ di materiale ferromagnetico posto fra due elettrodi.

Per effetto della variazione di temperatura dovuta alla radiazione incidente, il cristallo subisce una **polarizzazione elettrica interna**. Questa polarizzazione si manifesta come **ddp fra gli elettrodi**. Questi dispositivi sono caratterizzati da :

- prezzo non elevato
- buona velocità di risposta
- elevata sensibilità.

Il loro inconveniente principale è la perdita delle proprietà piroelettriche alla temperatura di Curie. Vengono impiegati:

- in impianti anti-intrusione (con raggio di azione di qualche metro)
- nella misura di temperatura dei forni.

¹³ Per esempio: tantalato di litio (LiTaO_3)

SENSORI CHIMICI E DI UMIDITA'

I sensori chimici sono dispositivi che rivelano la presenza di una determinata sostanza attraverso la variazione di una proprietà elettrica del sensore stesso. Per esempio la conducibilità, la capacità, la ddp fra due elettrodi.

UNA POSSIBILE CLASSIFICAZIONE

- sensori di GAS
- sensori di IONI (presenti in una soluzione)
- sensori di umidità

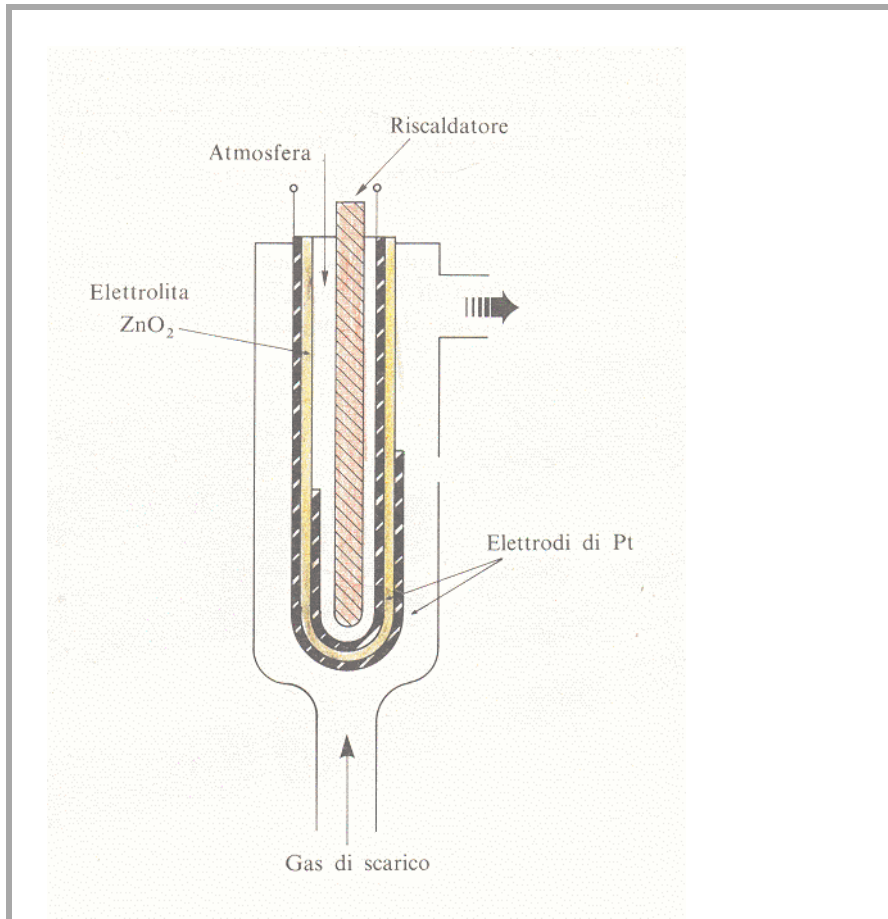
SENSORI DI GAS A OSSIDO DI STAGNO

L'ossido di stagno è un semiconduttore la cui conducibilità varia in presenza di gas combustibili.

I sensori in esame, realizzati su substrato di silicio, sfruttano proprio la variazione di conducibilità per rilevare la presenza di gas di scarico.

SENSORE DI GAS DI SCARICO PER AUTO (SONDA LAMBDA)

Le due superfici di un materiale **ceramico** (elettrolita¹⁴ a ossido di zirconio) sono esposte una all'atmosfera e l'altra ai gas di scarico.



La superficie esposta all'atmosfera presenta concentrazione costante di ossigeno, mentre la concentrazione dell'altra superficie è variabile per effetto dei gas di scarico.

La differenza di concentrazione di ossigeno fra le due superfici determina la nascita di una f.e.m. fra le due superfici.

Questa f.e.m. risulta proporzionale al logaritmo del rapporto delle due concentrazioni.

In sintesi il sensore rileva una f.e.m. proporzionale al rapporto fra la concentrazione di ossigeno nell'atmosfera e la concentrazione di ossigeno dei gas di scarico.

Sulle superfici dell'elettrolita sono presenti contatti elettrici in platino.

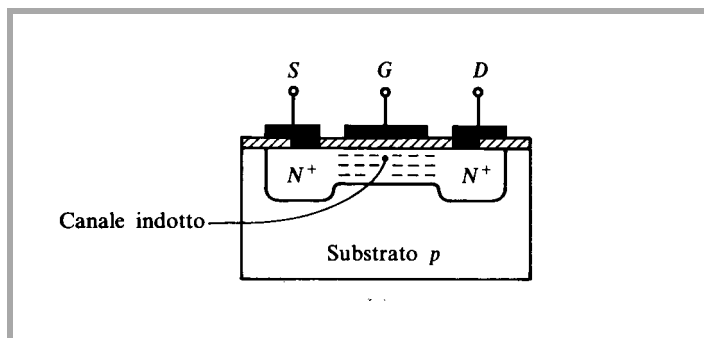
Il sensore lavora a una temperatura di circa 650° che viene ottenuta mediante un riscaldatore.

¹⁴ Composto le cui molecole (quando è allo stato fuso o in soluzione) si dissociano in ioni positivi e negativi

SENSORE DI IONI

(chiamato ISFET, Ion Sensitive FET oppure CHEMFET, CHEMical FET)

Il funzionamento dei normali mosfet ad arricchimento è basato sulla formazione di un canale conduttivo fra source e drain, in corrispondenza dell'applicazione di una tensione al gate.



Anche il CHEMFET si basa sulla formazione di un canale conduttivo fra drain e source.

La formazione del canale non è però dovuta alla applicazione di una tensione, ma alla presenza di una membrana polimerica sensibile agli ioni, membrana che sostituisce l'elettrodo di gate.

SENSORI DI UMIDITA'

Sono molto diffusi i sensori di umidità capacitivi, basati sulla variazione che l'umidità impone alla costante dielettrica dei materiali plastici e ceramici con i quali si riempie lo spazio fra le armature del condensatore.

CAPITOLO 35

ANTENNE

COMPLEMENTI

TIPI DI ANTENNE

DIPOLO ELEMENTARE O DIPOLO CORTO

DIPOLO A $\lambda/2$ O ANTENNA A MEZZ'ONDA

ANTENNE RETTILINEE RISONANTI (O ANTENNE LUNGHE UN NUMERO INTERO DI MEZZE LUNGHEZZE D'ONDA)

DIPOLO RIPIEGATO

ANTENNE VERTICALI POSTE AL SUOLO (DIPOLO MARCONIANO, ANTENNA GROUND PLANE, ANTENNE DI LUNGHEZZA COMPRESA FRA $\lambda/4$ E $\lambda/2$ E ANTENNE DI LUNGHEZZA SUPERIORE A $\lambda/2$)

ANTENNA YAGI

ANTENNE TRASMITTENTI TELEVISIVE

ANTENNE A LARGA BANDA (LOG-PERIODICHE E CON RIFLETTORE)

ANTENNE PARABOLICHE O ANTENNE A SUPERFICIE

PARABOLA CASSEGRAIN

ANTENNA A TROMBA

ANTENNA A TELAIO O ANTENNA LOOP

ANTENNA A FERRITE

ARRAY O ALLINEAMENTI O SCHIERE DI ANTENNE

ALLINEAMENTO BROADSIDE

ALLINEAMENTO ENDFIRE

ANTENNE RICEVENTI PER TELEFONIA CELLULARE

CAMPI DI UTILIZZAZIONE DELLE ANTENNE RICEVENTI

TIPI DI ANTENNE

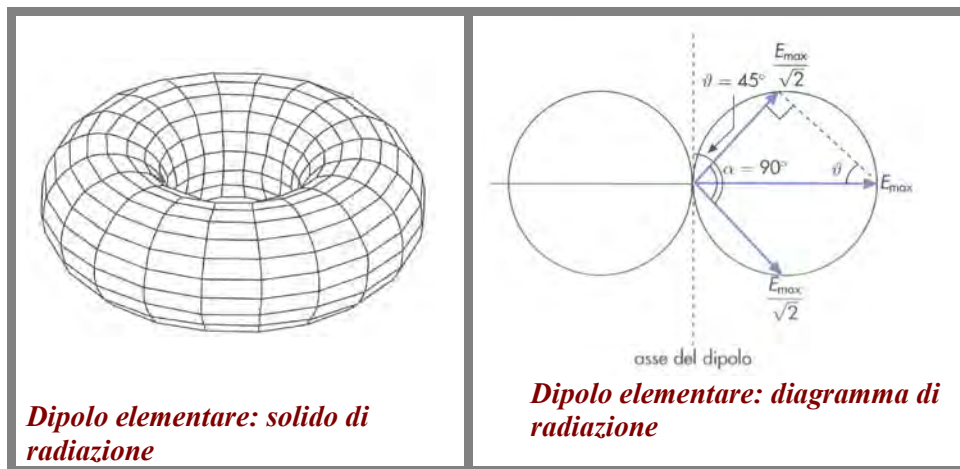
DIPOLO ELEMENTARE O DIPOLO CORTO

E' un'antenna *rettilinea omnidirezionale* caratterizzata da

- lunghezza L molto minore della lunghezza d'onda ($L \ll \lambda$)
- **corrente costante** lungo tutto il conduttore.

Il dipolo elementare è utilizzato anche come **sensore di campo elettrico** e presenta:

- GUADAGNO di direttività $G = 1,5$ cui corrisponde $G_{dB} = 1,76$
- RESISTENZA di irradiazione $R_i = 80\pi^2(L/\lambda)^2$ Ohm.



La **potenza irradiata** si può calcolare come:

$$P_i = \frac{40 \cdot \pi^2 \cdot I^2 \cdot L^2}{\lambda^2} = \frac{\pi \cdot Z_0 \cdot I^2 \cdot L^2}{3 \cdot \lambda^2}$$

L'**area efficace** è data da:

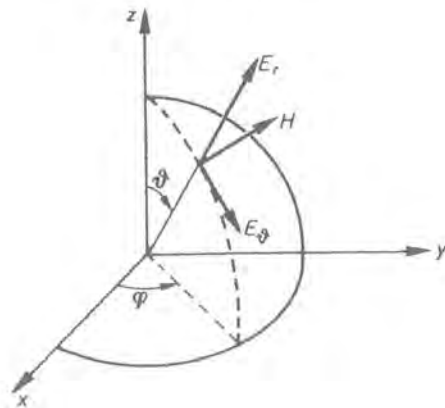
$$A_{EFF} = \frac{1,5 \cdot \lambda^2}{4\pi}$$

Le espressioni (ricavate dalle equazioni di Maxwell) dei **campi elettrico e magnetico** del **dipolo elementare**, disposto lungo l'asse z del riferimento, sono riportate nella figura seguente.

$$E_r = \frac{60 IL}{r^2} \cos \theta \left(1 - j \frac{\lambda}{2\pi r} \right) e^{-j2\pi r/\lambda}$$

$$E_\theta = j \frac{60 \pi IL}{\lambda r} \sin \theta \left(1 - j \frac{\lambda}{2\pi r} - \frac{\lambda^2}{4 \pi^2 r^2} \right) e^{-j2\pi r/\lambda}$$

$$H_\phi = j \frac{IL}{2\lambda r} \left(1 - j \frac{\lambda}{2\pi r} \right) \sin \theta e^{-j2\pi r/\lambda}$$



Campo prodotto da un dipolo elementare

Le tre precedenti espressioni, valutate però a grande distanza dall'antenna, cioè per $r \gg \lambda/\pi$, si possono scrivere come:

$$E_r \approx 0$$

$$E_\theta = j \frac{60 \pi IL}{\lambda r} \sin \theta e^{-j2\pi r/\lambda}$$

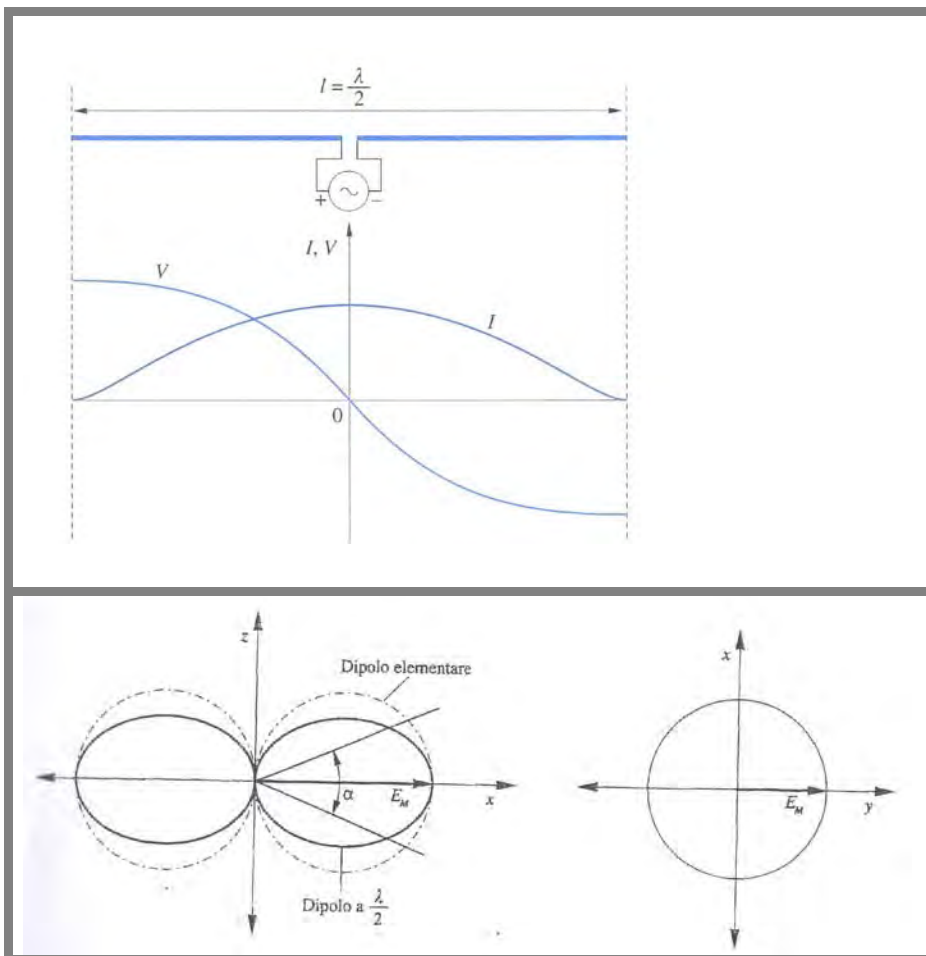
$$H_\phi = j \frac{IL}{2\lambda r} \sin \theta e^{-j2\pi r/\lambda}$$

Il campo a grande distanza dall'antenna è detto campo di radiazione, mentre il campo di induzione è quello prevalente nelle vicinanze dell'antenna

DIPOLO A $\lambda/2$ O ANTENNA A MEZZ'ONDA

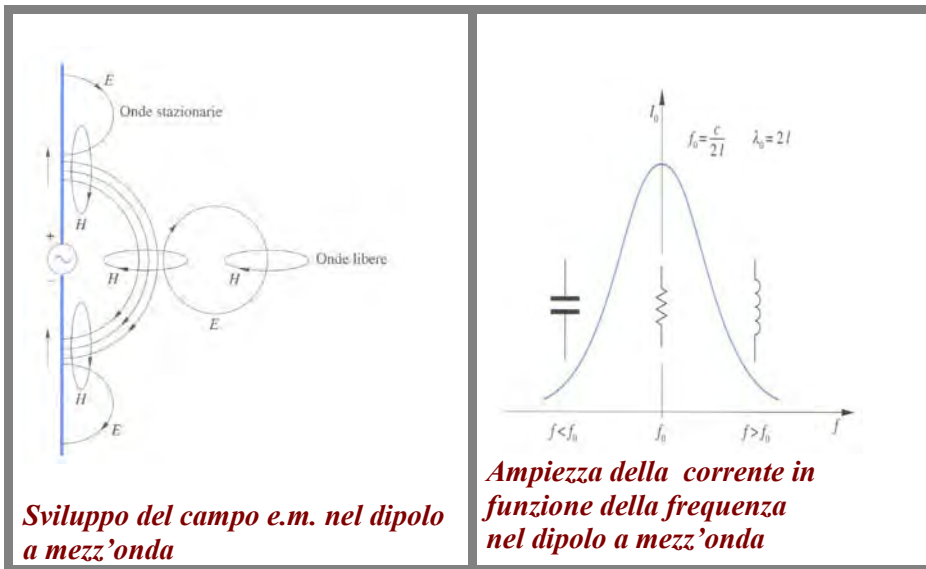
E' un'antenna costituita da un conduttore **filiforme, rettilineo**, lungo la metà della lunghezza d'onda del segnale da ricevere, **interrotto a metà della lunghezza** per l'applicazione dell'**alimentazione**, la quale deve essere posta in corrispondenza di un **massimo di corrente**.

Si tratta di un'antenna **risonante** che (contrariamente al dipolo marconiano a $\lambda/4$) non è posta al suolo, ma nello spazio libero, ossia a una **significativa distanza dal suolo** (e da altre strutture metalliche), distanza che deve essere pari **a diverse lunghezze d'onda**. E' impiegato **in VHF ed in UHF** ed è anche utilizzato, come **dispositivo di confronto**, per determinare il guadagno di altre antenne. Può essere posto in posizione **orizzontale o verticale**.



Il dipolo a $\lambda/2$ presenta:

- **GUADAGNO** di direttività **$G = 1,65$**
- **RESISTENZA** di irradiazione **$R_i = 73 \text{ Ohm}$**
- **RESISTENZA** di dissipazione **$R_d = 42 \text{ Ohm}$**
- **ANGOLO DI RADIAZIONE** **$\alpha = 78^\circ$**



La **potenza irradiata** si può calcolare come:

$$P_i = \frac{2,435}{8 \cdot \pi} \cdot Z_0 \cdot I^2$$

Il campo elettrico IRRADIATO, cioè il campo A GRANDE DISTANZA, (ossia a una distanza superiore ad alcune lunghezze d'onda) del dipolo a $\lambda/2$ è dato dall'espressione:

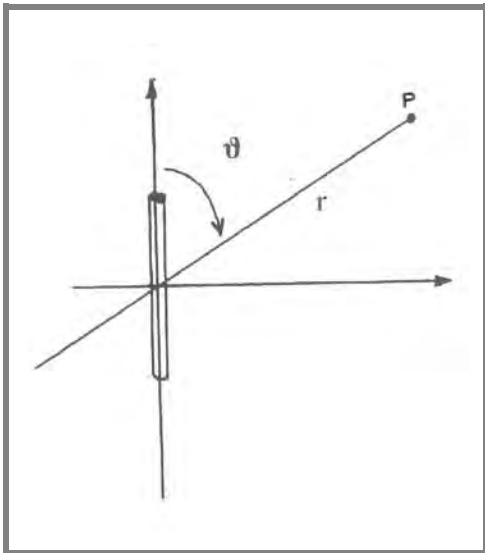
dove:

$$E_\vartheta = \frac{60 I_0}{r \cdot \sin \vartheta} \cdot \cos\left(\frac{\pi}{2} \cdot \cos \vartheta\right)$$

I_0 = *valore massimo della corrente lungo l'ANTENNA*

r = *distanza del punto P, nel quale viene calcolato il campo, dal centro dell'antenna*

θ = *angolo formato, dal conduttore d'antenna, con la congiungente il centro dell'antenna con il punto P*

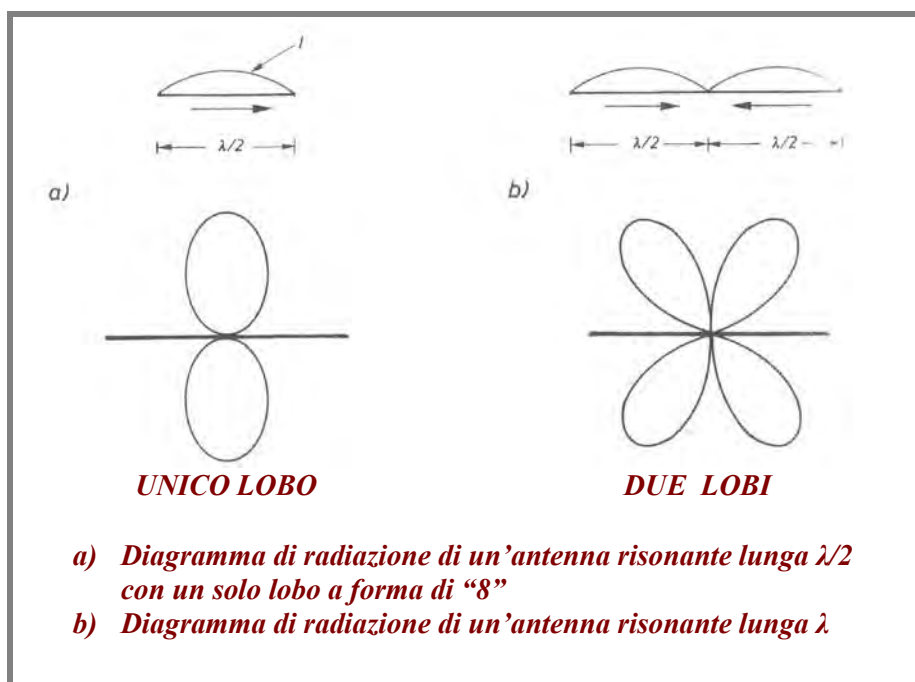


La precedente espressione di E_θ rappresenta un caso particolare (per $L = \text{lunghezza dell'antenna} = \lambda/2$) dell'espressione del campo di un'antenna rettilinea risonante con lunghezza uguale a un numero DISPARI di mezza lunghezze d'onda, di cui ora parleremo.

ANTENNE RETTILINEE RISONANTI (O ANTENNE LUNGHE UN NUMERO INTERO DI MEZZE LUNGHEZZE D'ONDA)

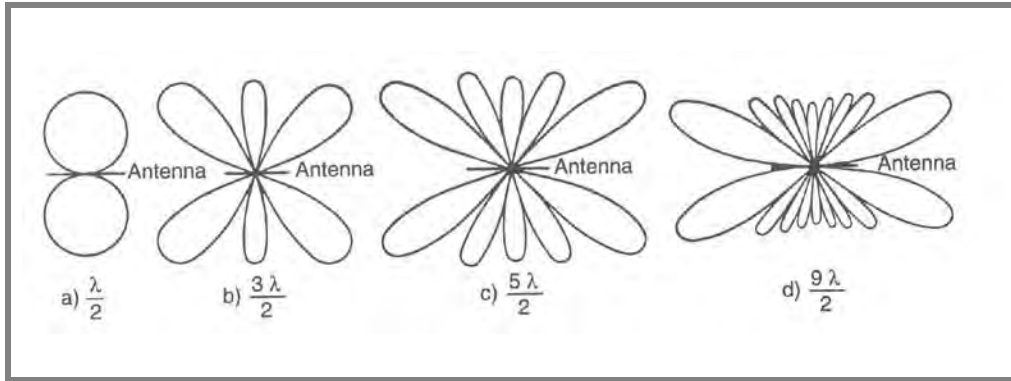
Un'antenna rettilinea si dice “risonante” se la sua lunghezza L è un **numero intero** di “**mezze lunghezze d'onda**”, cioè se verifica la condizione:

$$L = K \cdot \frac{\lambda}{2} \quad \text{con: } K = \text{numero INTERO}$$



Il diagramma di radiazione di un'antenna rettilinea risonante è caratterizzato da un **numero N di lobi** che risulta:

- pari al **numero K di mezze lunghezze d'onda** costituenti la lunghezza fisica:
 $N_{\text{LOBI}} = K$
- pari circa alla **metà del numero** di **intere lunghezza d'onda** costituenti la lunghezza fisica.



Le espressioni del campo elettrico delle antenne di cui stiamo parlando sono:

1. nel caso che la lunghezza L dell'antenna sia un multiplo **DISPARI** di **MEZZE** lunghezze

d'onda:
$$E_{\vartheta} = \frac{60 I_0}{r \cdot \sin \vartheta} \cdot \cos \left(\frac{\pi \cdot L}{\lambda} \cdot \cos \vartheta \right)$$

dalla quale abbiamo ricavato l'espressione del campo per il dipolo a $\lambda/2$

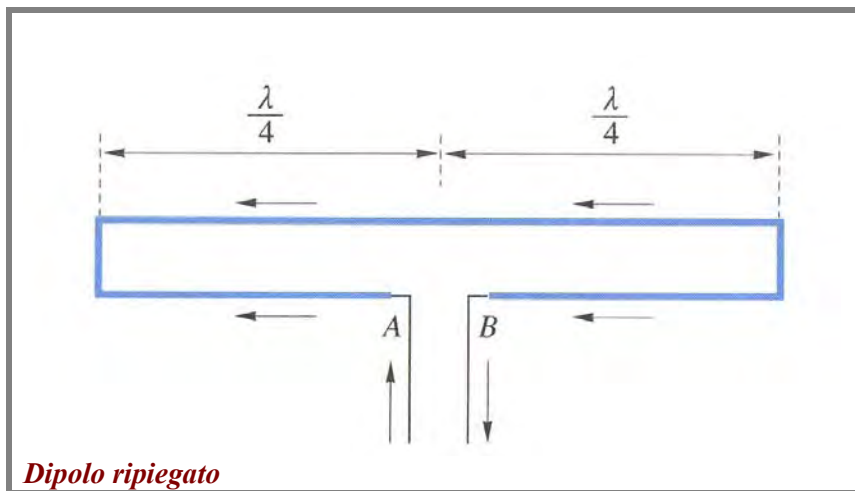
2. nel caso che la lunghezza L dell'antenna sia un multiplo **PARI** di **MEZZE** lunghezze

d'onda:
$$E_{\vartheta} = \frac{60 I_0}{r \cdot \sin \vartheta} \cdot \sin \left(\frac{\pi \cdot L}{\lambda} \cdot \cos \vartheta \right)$$

DIPOLO RIPIEGATO (FOLDED DIPOLE)

Anche se può essere usato da solo, questo dipolo è impiegato spesso come elemento radiante di antenne più complesse, come l'antenna Yagi.

Può essere visto come una variante del dipolo a $\lambda/2$, caratterizzata però da un maggiore diametro e quindi da una maggiore larghezza di banda. Perciò questo dipolo è un'antenna **a banda larga**.



Il dipolo ripiegato presenta:

- **GUADAGNO di direttività** ($G = 1,65 \rightarrow G_{DB} = 2,15$), UGUALE al dipolo a $\lambda/2$
- **RESISTENZA di irradiazione** circa QUATTRO volte maggiore:
 $R_i = 73 \cdot 4 = 300 \text{ Ohm}$
- **POTENZA IRRADIATA:**

$$P_i = \frac{9,74}{8 \cdot \pi} \cdot Z_0 \cdot I^2$$

ANTENNE VERTICALI POSTE AL SUOLO:

- **DIPOLO MARCONIANO**
- **ANTENNA GROUND PLANE**
- **ANTENNE DI LUNGHEZZA COMPRESA
FRA $\lambda/4$ e $\lambda/2$**
- **ANTENNE DI LUNGHEZZA SUPERIORE a $\lambda/2$**

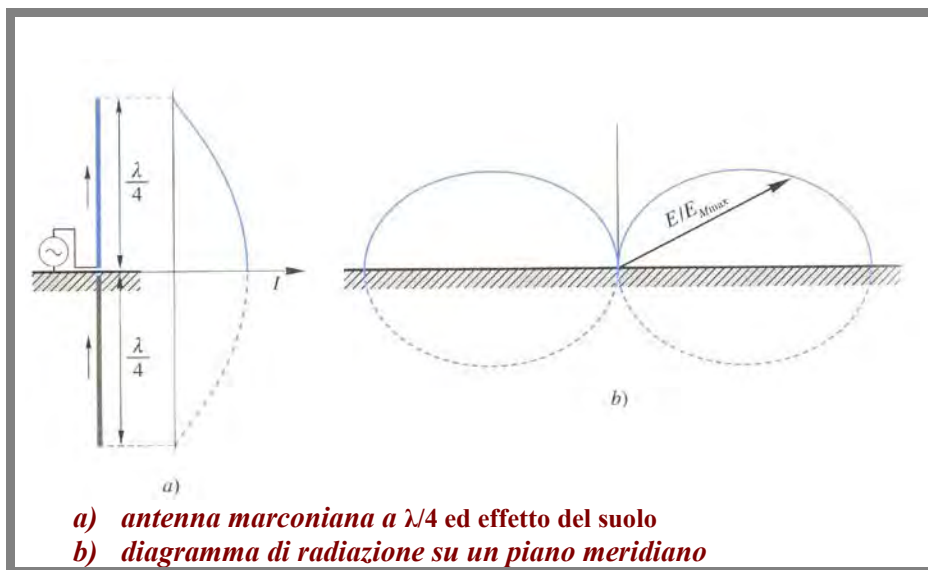
DIPOLO MARCONIANO O ANTENNA A STILO O SEMIDIPOLO (di lunghezza $\lambda/4$)

E' un'antenna utilizzata nelle **radiodiffusioni** in onde **medie e lunghe** e in particolare nella realizzazione di **torri autoirradianti per MF** e nella gamma di frequenze riservata alle comunicazioni **CB** (Banda Cittadina) a **27 MHz**, cui corrispondono lunghezze d'onda sui **(15-20)m**.

E' anche utilizzata come **antenna veicolare**, nel qual caso deve essere dotata di "contrappeso" che può essere costituito dallo stesso veicolo.

A questo proposito ricordiamo che, in relazione alle antenne poste in prossimità del suolo, per **contrappeso** (counterpoise) si intende una maglia di fili intrecciati che viene posta a breve distanza dal suolo per aumentarne la conducibilità, nel caso fosse scarsa.

Quella "a stilo" è un'antenna **verticale di lunghezza $\lambda/4$** , però **si comporta come un dipolo a $\lambda/2$** perché, essendo alimentata alla base da un generatore che ha l'altro morsetto collegato a terra, **sfrutta l'effetto del suolo**. Il suolo riflette l'energia che lo colpisce, comportandosi come un prolungamento ideale dell'antenna stessa, prolungamento che è chiamato "**antenna immagine** del dipolo reale."



Rispetto al dipolo a $\lambda/2$, il dipolo marconiano presenta:

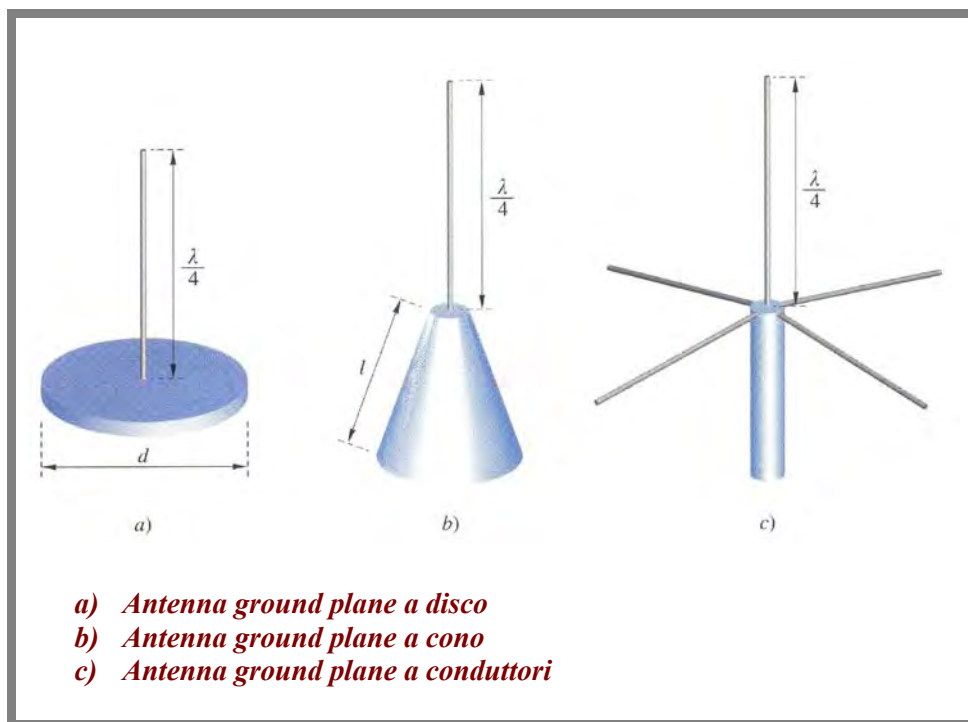
- **GUADAGNO** di direttività **DOPPIO** ($G = 3,33$, corrispondente a $G_{dB} = 5,18$)
- **RESISTENZA** di irradiazione **DIMEZZATA** ($R_i = 36,5 \text{ Ohm}$, contro 73 Ohm)
- **ANGOLO DI RADIAZIONE DIMEZZATO** ($\alpha = 39^\circ$, contro 78°)

ANTENNA GROUND PLANE

Una variante del dipolo marconiano è l'antenna ground plane dotata di un piano conduttore artificiale realizzato con raggi metallici.

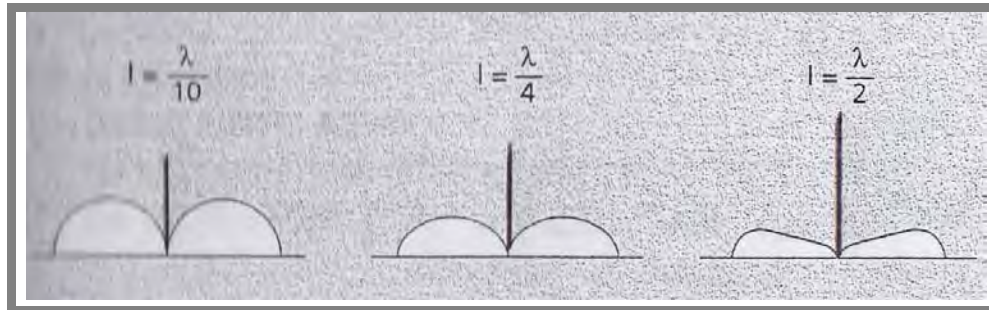
La resistenza di radiazione varia con l'inclinazione dei raggi fra 36,5 Ohm (con conduttori radiali orizzontali) e 73 Ohm (con conduttori radiali verticali).

Si può quindi ottenere anche una resistenza di 52 Ohm che permette di utilizzare come feeder di antenna un cavo coassiale con la stessa impedenza (RG 58).



ANTENNE DI LUNGHEZZA COMPRESA FRA $\lambda/4$ e $\lambda/2$

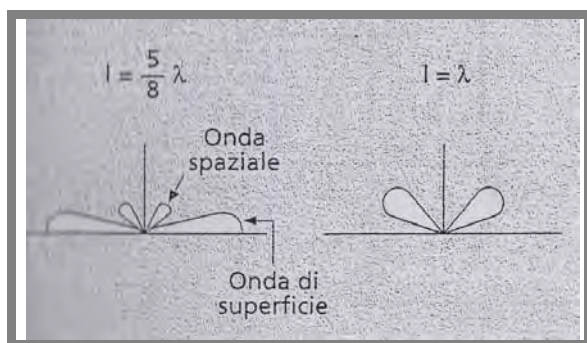
Il diagramma di radiazione è analogo a quello del dipolo marconiano. Al crescere dell'altezza però, l'energia tende a concentrarsi lungo la superficie del suolo.



ANTENNE DI LUNGHEZZA SUPERIORE a $\lambda/2$

Il diagramma di radiazione comincia a presentare lobi secondari.

Per altezze superiori a $(3/4)\lambda$ prevale l'energia emessa secondo angolazioni diverse da quella orizzontale.



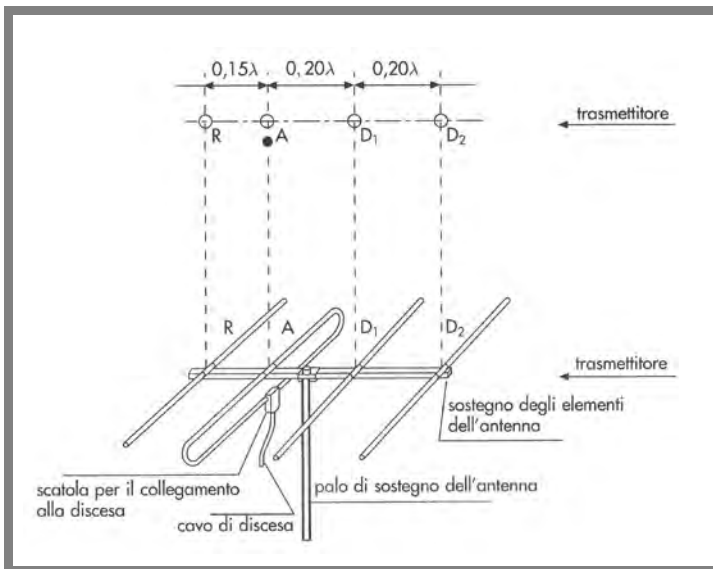
ANTENNA YAGI

L'antenna Yagi è un'antenna tipicamente usata per la ricezione **televisiva** in **VHF** e **UHF** a eccezione dei casi (UHF, bande IV e V) nei quali è necessaria un'antenna a banda larga. In questi ultimi casi la Yagi può essere insoddisfacente perché ha una **banda di ricezione piuttosto stretta** intorno alla frequenza per la quale è stata progettata.

L'antenna Yagi è costituita

- da un dipolo **attivo**, che può essere un **dipolo a $\lambda/2$** o un **dipolo ripiegato**
- da un certo numero di dipoli passivi detti direttori
- da un dipolo passivo anch'esso, detto riflettore (che si trova “alle spalle” del dipolo attivo rispetto ai direttori).

L'**aggiunta** di elementi **passivi** fa **aumentare il guadagno direzionale** rispetto al dipolo unico.

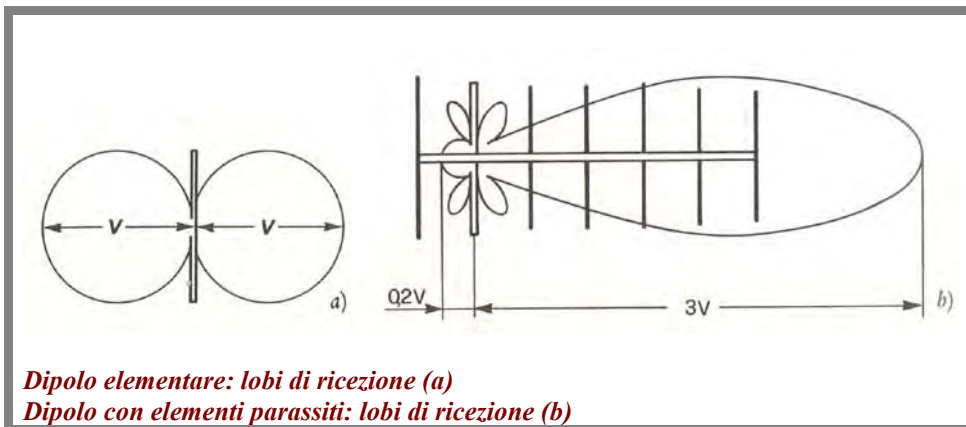
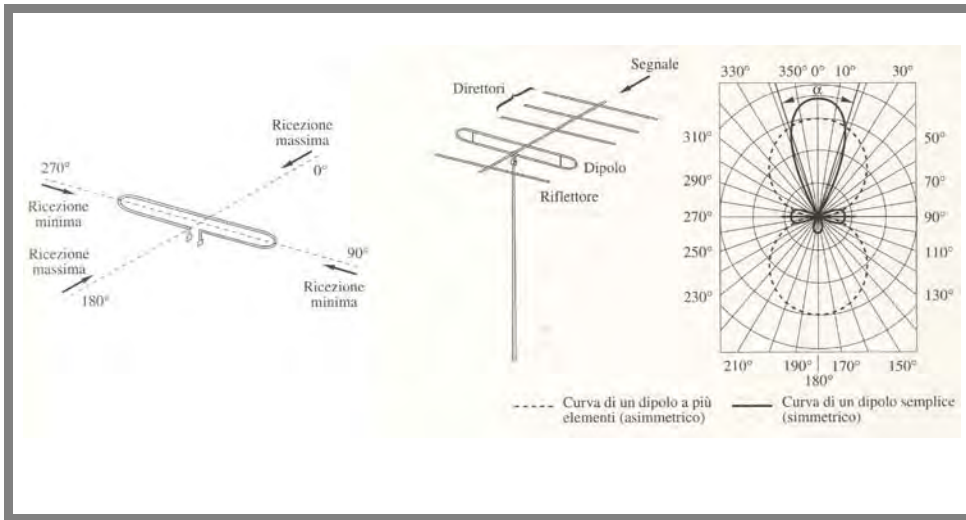


Per la ricezione della gamma VHF le antenne Yagi sono costituite da un numero di elementi compreso fra 4 e 8 e in particolare

- da un dipolo attivo
- da un numero di direttori compreso fra 2 e 6
- da un riflettore.

DIRETTORI

I direttori devono il loro nome al fatto che **dirigono dalla loro parte il lobo principale** del diagramma di irradiazione complessivo dell'antenna Yagi.



Dipolo elementare: lobi di ricezione (a)

Dipolo con elementi parassiti: lobi di ricezione (b)

In presenza di direttori, l'irradiazione è **massima** nella **direzione** che va **dal dipolo alimentato ai direttori stessi**.

EFFETTI DELLA DISTANZA FRA I DIRETTORI

All'AUMENTARE della distanza fra i direttori:

- diminuisce il guadagno
- AUMENTANO la resistenza di irradiazione e la larghezza di banda

Al *diminuire* della distanza fra i direttori:

- aumenta il GUADAGNO
- *diminuiscono* la resistenza di irradiazione e la larghezza di banda.

RIFLETTORE

Il riflettore “respinge” il diagramma di radiazione dalla parte opposta rispetto all'antenna alimentata. Dalla parte del riflettore il diagramma di irradiazione è più corto, mentre si allunga dalla parte dei direttori.

ANTENNE TRASMITTENTI TELEVISIVE

Le antenne trasmettenti televisive sono formate da **gruppi di dipoli elementari ripiegati, collegati in parallelo** tra loro ed eccitati con **correnti di uguale fase ed ampiezza**.

Si ottengono così **schiere di quattro elementi allineati**, riuniti in modo opportuno sui quattro lati di un parallelepipedo verticale a sezione quadrata.

Posteriormente alle schiere dei dipoli sono sistemati dei **riflettori metallici piani** con la funzione di **evitare l'irradiazione posteriore** e quindi di accrescere le direttività delle antenne.

La **direzione di irradiazione è perpendicolare all'asse di allineamento degli elementi**. L'insieme è racchiuso in un **tubo verticale**.



ANTENNE A LARGA BANDA

Sono tipicamente utilizzate in UHF (bande IV e V), cioè nei casi nei quali è appunto necessaria un'antenna a banda larga.

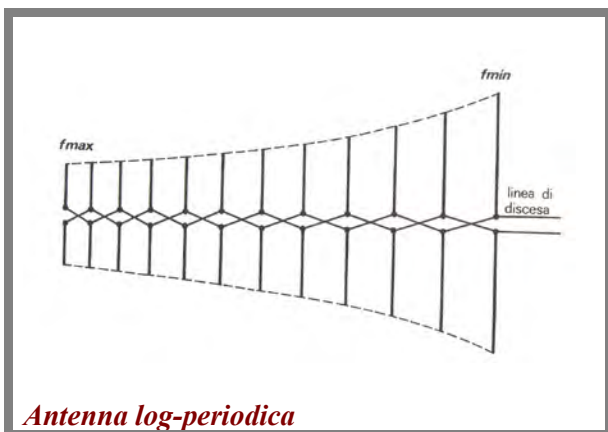
Fra le antenne a larga banda annoveriamo:

1. l'antenna **LOG-PERIODICA**
2. l'antenna a “**DOPPIO V**” con **RIFLETTORE a GRIGLIA**
3. l'antenna a “**X**” con **RIFLETTORE a GRIGLIA**
4. l'antenna con **RIFLETTORE a PANNELLO**

ANTENNE LOG-PERIODICHE

Sono antenne ad accordo variabile che hanno la stessa struttura delle Yagi, ma sono costituite da una serie di dipoli, **tutti alimentati**, di lunghezza decrescente.

Per ogni frequenza ricevuta, solo un certo numero dei dipoli si accorda su tale frequenza, mentre gli altri fanno da direttori e riflettori.



ANTENNE A “V” CON RIFLETTORE A GRIGLIA

Sono antenne **risonanti a onda intera**: ogni elemento è costituito da due dipoli sagomati a “V”, con i vertici convergenti al centro e isolati tra loro e dal supporto. Il riflettore (di dimensioni superiori alle lunghezze dei dipoli) è sagomato opportunamente ed è costituito da una rete di maglie fitte.

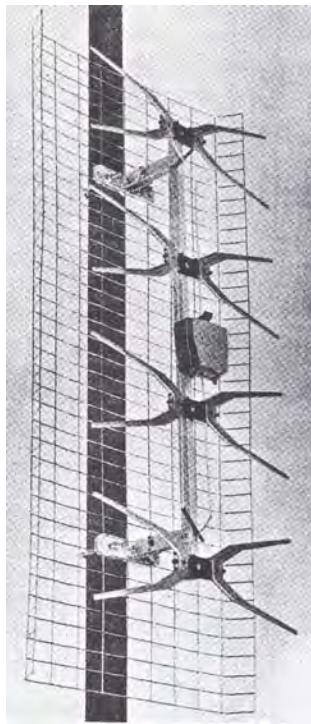
Il riflettore a griglia ha le funzioni di:

- agire da schermo per le onde che provengono dalla parte posteriore dell'antenna
- concentrare sull'antenna una parte dell'energia che proviene dal trasmettitore.

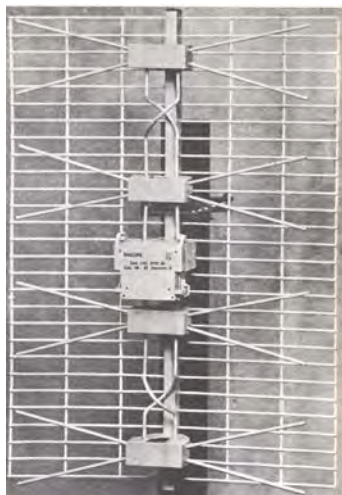


ANTENNE CON RIFLETTORE A PANNELLO

Sono caratterizzate da un discreto angolo di apertura che le rende adatte alle situazioni in cui i segnali delle emittenti (nelle bande IV e V UHF) non provengono dalla stessa direzione.



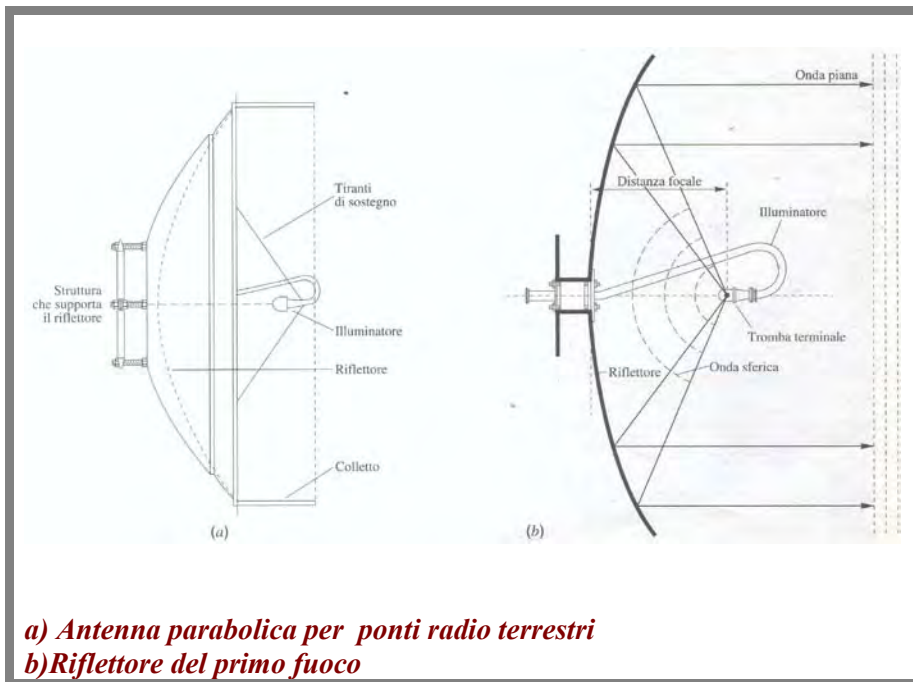
***Antenna a pannello a larga banda
per VHF (canali 21-68)***



***Antenna a pannello per la banda V
(canali 38÷72)***

ANTENNE PARABOLICHE (ANTENNE A SUPERFICIE)

Le antenne a parabola (e, più in generale, le cosiddette antenne a superficie) si utilizzano nel campo delle **microonde** o comunque a frequenze per le quali la propagazione delle onde radio avviene con **modalità ottiche** e, di conseguenza, è necessario un **elevato guadagno di direttività**. Ricordiamo che la propagazione ottica avviene in linea retta e in **condizione di visibilità** fra trasmettitore e ricevitore.

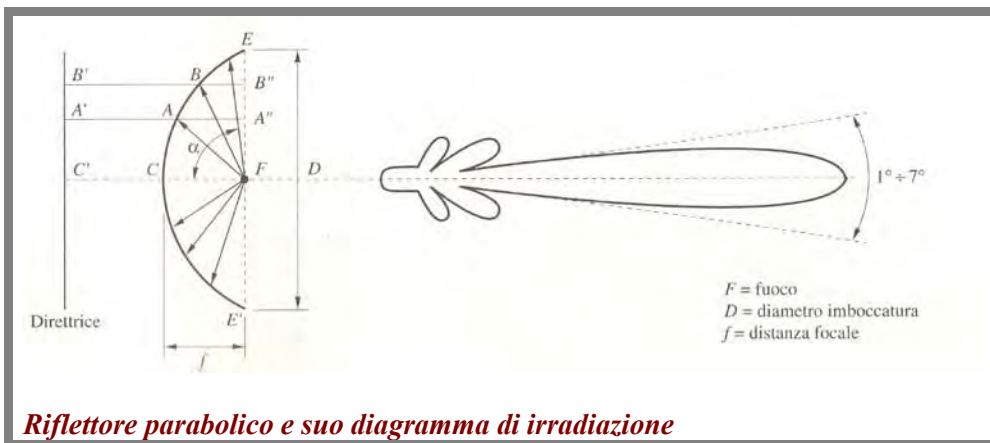


Analogamente a quanto avviene in una torcia o nel faro di un'automobile (dove la luce della lampadina viene riflessa e focalizzata da un paraboloide riflettente) l'antenna parabolica riflette e focalizza, mediante la sua superficie metallica interna, le onde e.m. provenienti da un dispositivo detto "illuminatore". L'illuminatore è posto nel punto focale della parabola stessa ed è una guida d'onda tronca o con terminazione a tromba.

Quanto detto riguardo all'irradiazione di onde e.m. attraverso un'antenna parabolica vale anche (in base al principio di reciprocità) quando si utilizza l'antenna come ricevente.

In quest'ultimo caso le onde e.m. provenienti dallo spazio circostante vengono raccolte dalla parabola che le convoglia verso l'illuminatore il quale le guida agli apparati di ricezione.

Il diagramma di radiazione di un'antenna parabolica è quello tipico di un'antenna direttiva a elevato guadagno: lobi secondari molto piccoli e **lobo principale stretto e allungato**.



DETERMINAZIONE DEI PARAMETRI DI UN'ANTENNA PARABOLICA

EFFICIENZA η

$$\eta = A_{eq} / A_g$$

dove

A_{eq} = area di un'antenna equivalente ideale illuminata uniformemente

A_g = area geometrica dell'antenna in esame

Valore tipico: $\eta = 0,55$

ANGOLO DI APERTURA A 3dB

$$\alpha = 70 * (\lambda / D_a)$$

dove: D_a = diametro dell'antenna

GUADAGNO DIRETTIVO (relazione esatta)

$$G = \eta \cdot \left(\frac{\pi \cdot D_a}{\lambda} \right)^2$$

GUADAGNO DIRETTIVO

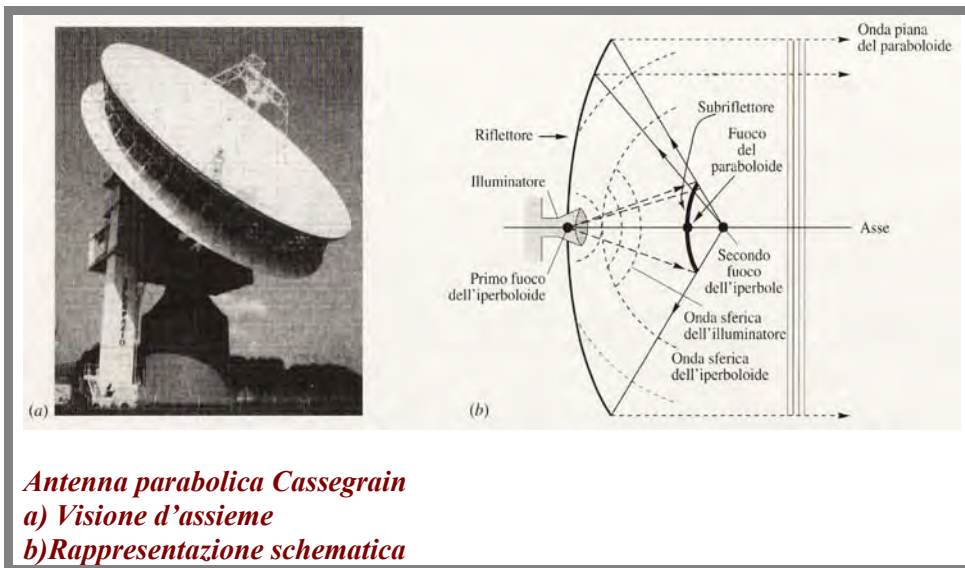
(relazione **approssimata**, avendo assunto, per l'efficienza, il valore tipico $\eta = 0,55$)

$$G = 5,5 \cdot \left(\frac{D_a}{\lambda} \right)^2$$

PARABOLA CASSEGRAIN

E' un'antenna di grandi dimensioni che differisce dalla comune antenna a parabola per due caratteristiche:

- l'illuminatore, invece che nel punto focale del riflettore parabolico, termina proprio sul riflettore interrompendone la superficie
- presenza di un secondo riflettore (subriflettore).

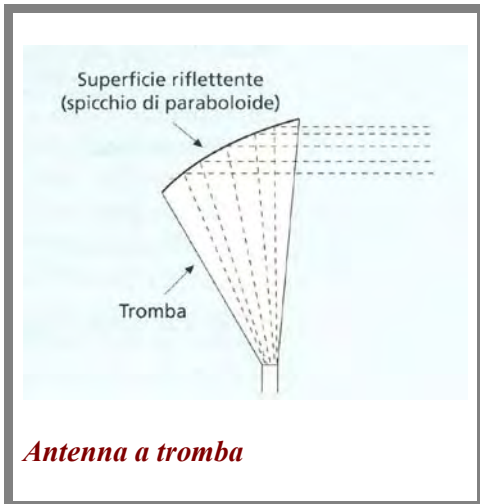


Questi accorgimenti determinano la riduzione della lunghezza del feeder e una **maggiore uniformità nell'illuminazione** della parabola.

ANTENNA A TROMBA (HORN REFLECTOR)

E' un'antenna di grandi dimensioni costituita da:

- un illuminatore a forma di tromba conica
- un riflettore a specchio di parabola

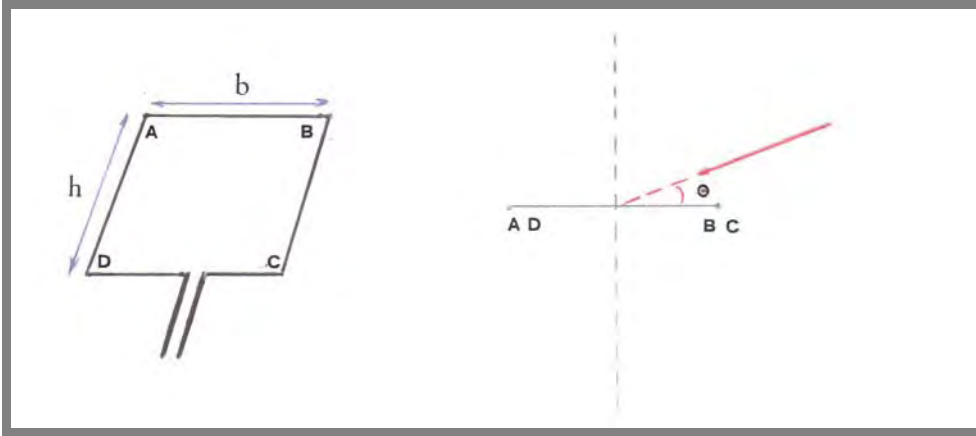


Questi accorgimenti permettono alla radiazione riflessa di non incontrare alcun ostacolo.

ANTENNA A TELAIO (ANTENNA LOOP)

E' utilizzata nei **radiogoniometri**, cioè negli apparati per la rilevazione della provenienza dei segnali radio e come sensore di campo magnetico.

Per le sue caratteristiche di direttività può essere usata in presenza di disturbi di provenienza localizzata.



Per un'antenna a telaio, di forma rettangolare e di **area $A = b \cdot h$** (e nell'ipotesi che l'antenna sia di **piccole dimensioni** rispetto alla lunghezza d'onda del segnale: **$b/\lambda \ll 1$**) l'espressione della **tensione $v(t)$ indotta** è data da:

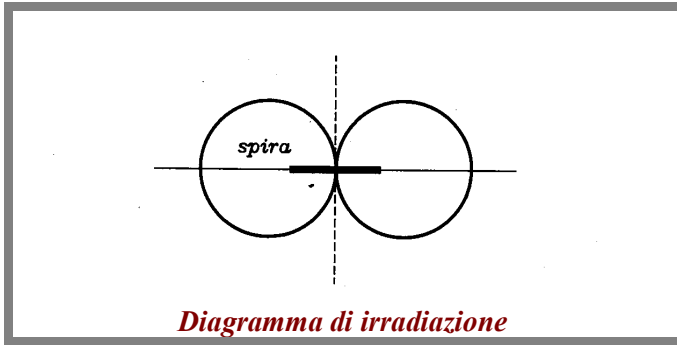
$$v(t) = \frac{2 \cdot \pi \cdot A \cdot E}{\lambda} \cdot \cos \theta \cdot \cos(\omega t) \quad \text{tensione indotta}$$

dove:

- $\frac{2 \cdot \pi \cdot A \cdot E}{\lambda} \cdot \cos \theta$ è l' **AMPIEZZA** “V” della tensione indotta
- A è l'area del telaio
- θ è l'angolo che il campo elettrico forma con il piano della spira del telaio

DIAGRAMMA DI RADIAZIONE

Il diagramma di irradiazione nel piano perpendicolare al piano della spira si può costruire proprio valutando l'espressione della tensione indotta $v(t)$ al variare dell'angolo ϑ che il vettore campo E forma con il piano della spira.
Il massimo si ha per $\vartheta = 0$, cioè nel piano della spira



TELAIO AD N SPIRE

Se l'antenna è formata da N spire, l'ampiezza della tensione va moltiplicata per N :

$$v(t) = N \cdot \frac{2 \cdot \pi \cdot A \cdot E}{\lambda} \cdot \cos \theta \cdot \cos(\omega t)$$

Il valore massimo V della tensione (cioè il valore per $\cos(\omega t)=1$) è:

$$V = N \cdot \frac{2 \cdot \pi \cdot A \cdot E}{\lambda} \cdot \cos \theta$$

e il valore massimo nella direzione di massima irradiazione ($\vartheta=0$) è:

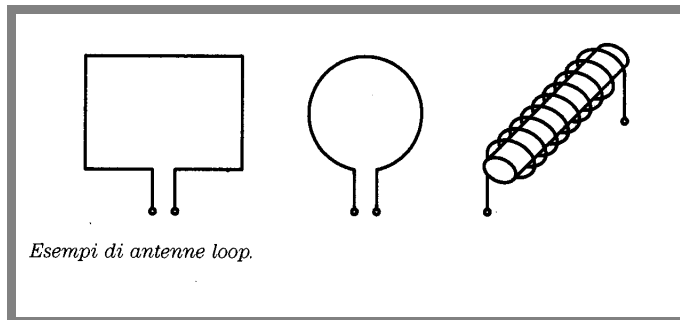
$$V_M = N \cdot \frac{2 \cdot \pi \cdot A \cdot E}{\lambda} \quad (****)$$

ALTEZZA EFFICACE

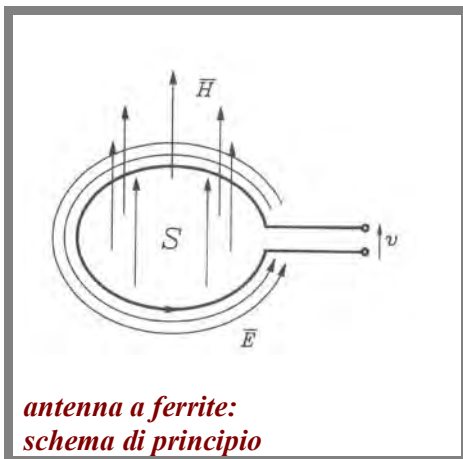
Ricordando che, in generale, l'altezza efficace è il fattore di proporzionalità tra la f.e.m dell'antenna e il campo elettrico:

$V = h_{EFF} E$ e confrontando con la (****) si ha:
$$h_{EFF} = N \cdot \frac{2 \cdot \pi \cdot A}{\lambda}$$

Notiamo che l'altezza efficace è inversamente proporzionale a λ cioè che diminuisce al crescere di λ e quindi al diminuire della frequenza.



ANTENNA A FERRITE

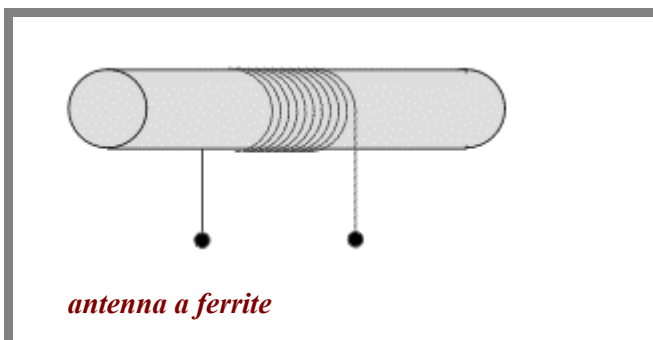


Può essere considerata “antenna a telaio” anche l’antenna a ferrite.

la tensione che si induce sull’antenna a ferrite è dovuta al fenomeno dell’induzione elettromagnetica, descritto dalla legge di Faraday. Perciò ai capi della spira si rileva una tensione:

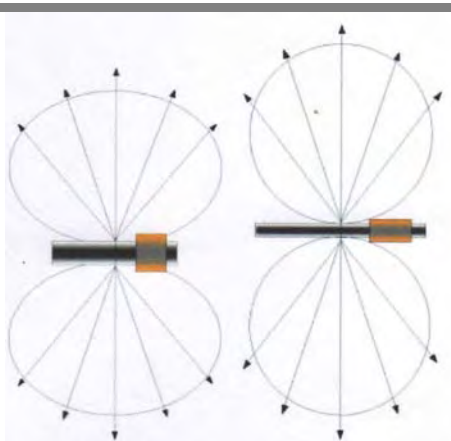
$$v = -\mu_0 \cdot \frac{dH}{dt}$$

Si può anche dire che la tensione che si induce sull’antenna nasce dal fatto che, associato al campo magnetico, è presente anche un campo elettrico posto ortogonalmente al campo magnetico stesso e quindi parallelo al conduttore, su cui induce una corrente, e quindi una tensione.



Sfruttando questo principio si realizza una **tipica antenna** per la **banda AM commerciale a Onde medie** (520-1605 KHz), utilizzando una bacchetta di ferrite lunga circa 10 cm, avvolta da spire di conduttore. Essendo piccola, l’antenna a ferrite può essere posta all’interno delle radioline funzionanti in questa banda (pensiamo, per confronto, a quali dimensioni dovrebbe avere invece un’antenna verticale, a dipolo oppure a $\lambda/4$, visto che la lunghezza d’onda è, in questo caso specifico, di circa 300 m).

Quest’antenna può essere definita “di tipo magnetico” perché la componente magnetica del flusso elettromagnetico incidente, concatenandosi con l’avvolgimento, induce su questo una tensione proporzionale al numero di spire. Il nucleo di ferrite ha la funzione di convogliare e concentrare il flusso magnetico all’interno dell’avvolgimento in modo da indurre tensioni sufficientemente elevate. Per questo motivo le sue dimensioni non sono piccole (per esempio 1÷10 cm).



Diagrammi di irradiazione di antenne a ferrite: le spire (che formano un cilindro avvolto attorno alla bacchetta di ferrite) sono viste in sezione e i piani nei quali giacciono sono perpendicolari al foglio

Il diagramma di irradiazione dell'antenna a ferrite è toroidale, con il massimo lungo l'asse delle spire e uno zero nella direzione perpendicolare e per questo motivo, nelle radioline commerciali in onde medie, l'antenna a ferrite è collocata orizzontalmente (cioè parallelamente alla direzione del campo magnetico e **ortogonalmente a quella del campo elettrico**).

Questo diagramma consente di utilizzare le antenne a telaio (antenne loop) nel campo della *radiogoniometria*, cioè per la localizzazione delle direzioni di provenienza dei segnali radio; in questo caso i telai sono applicati su supporti rotanti. Per effetto della bidirezionalità del diagramma, è possibile individuare la direzione di propagazione ma non il verso, per cui in genere a queste antenne vengono accoppiate antenne verticali al fine di rendere il diagramma di irradiazione unidirezionale.

ARRAY O ALLINEAMENTI O SCHIERE DI ANTENNE

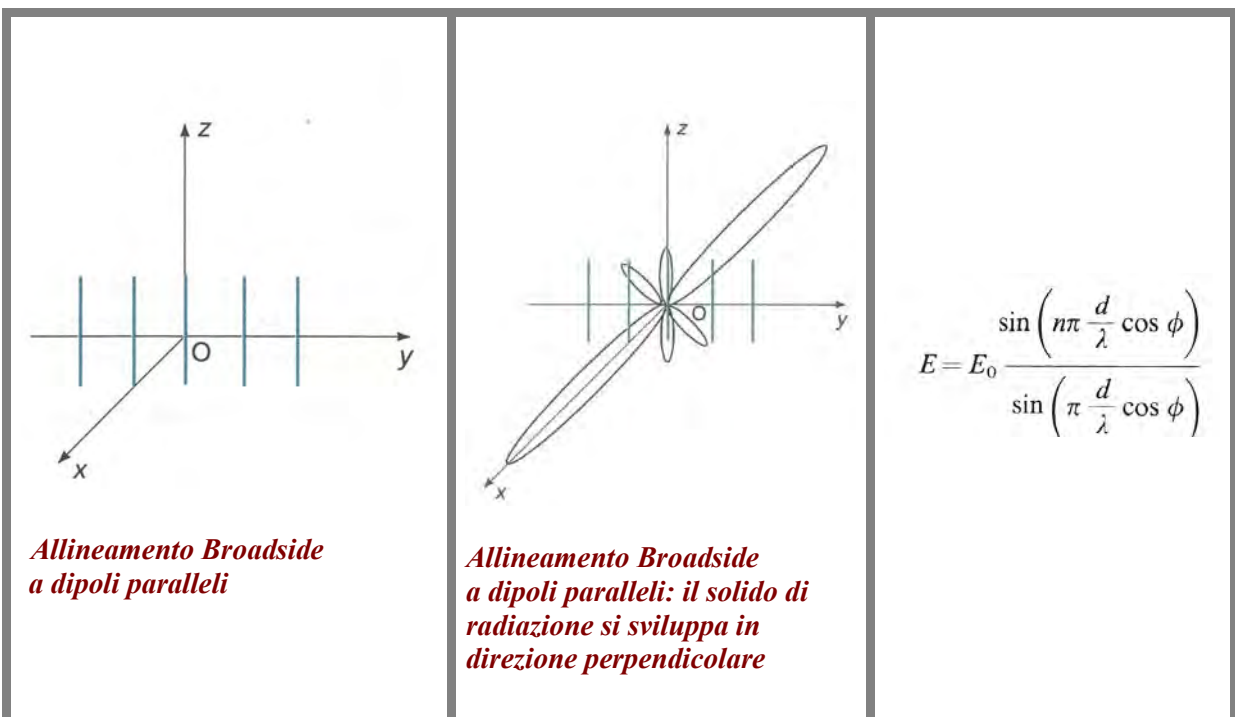
Sono raggruppamenti di antenne che hanno la funzione di produrre un **campo molto più direttivo** rispetto al campo ottenibile da una singola antenna.

ALLINEAMENTO BROADSIDE

Le antenne, che possiamo supporre isotropiche (ma che possono anche essere dipoli), sono disposte con i centri lungo una retta, cioè lungo una direzione assiale.

Il **CAMPO COMPLESSIVO** che viene prodotto risulta:

- molto **piccolo in direzione assiale**, in quanto i contributi dei singoli radiatori presentano fasi diverse (al crescere del numero delle antenne il campo tende ad annullarsi)
- **MASSIMO in direzione PERPENDICOLARE** all'allineamento (in quanto i contributi dei singoli radiatori si sommano).



L'allineamento Broadside è formato da più antenne:

- con **correnti in fase**
- **equidistanti** fra loro (per esempio di $\lambda/2$)

Se le antenne non sono isotropiche, ma sono dipoli, l'allineamento può essere:

- a dipoli paralleli (con i dipoli disposti come un filare di alberi)
- a dipoli allineati (con i dipoli disposti come i tratti della linea di mezzeria di una strada).

ALLINEAMENTO ENDFIRE

L'allineamento Endfire è formato da più antenne:

- con **correnti aventi la stessa ampiezza, ma progressivamente sfasate**
- **equidistanti** fra loro.

Lo sfasamento $\Delta\Phi_i$ fra le correnti di antenne adiacenti viene scelto uguale allo sfasamento

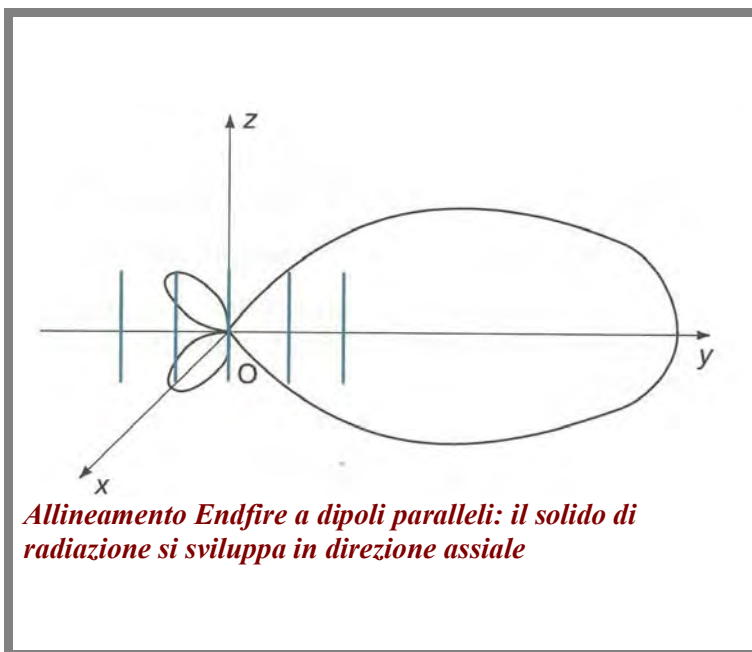
$$\Delta\Phi = 2\pi \cdot \frac{d}{\lambda}$$

dovuto alla distanza d fra le antenne.

Per esempio, se la distanza fra le antenne è $\lambda/2$, lo sfasamento progressivo fra le correnti è:

$$\Delta\Phi_i = 2\pi \cdot \frac{\lambda/2}{\lambda} = \frac{2\pi}{2} = \pi \quad \text{cioè} \quad \Delta\Phi_i = \pi$$

Le antenne che possiamo supporre isotropiche (ma che possono anche essere dipoli) sono disposte lungo una retta, cioè lungo una direzione assiale.



Il **CAMPO COMPLESSIVO** che viene prodotto **si sviluppa** essenzialmente in **direzione ASSIALE** (cioè nella direzione nella quale si compone l'allineamento) mentre risulta molto **piccolo** direzione perpendicolare.

CORTINA DI ANTENNE

E' l'insieme di due o più array di radiatori che ha l'effetto di aumentare ulteriormente la direttività dei singoli allineamenti.

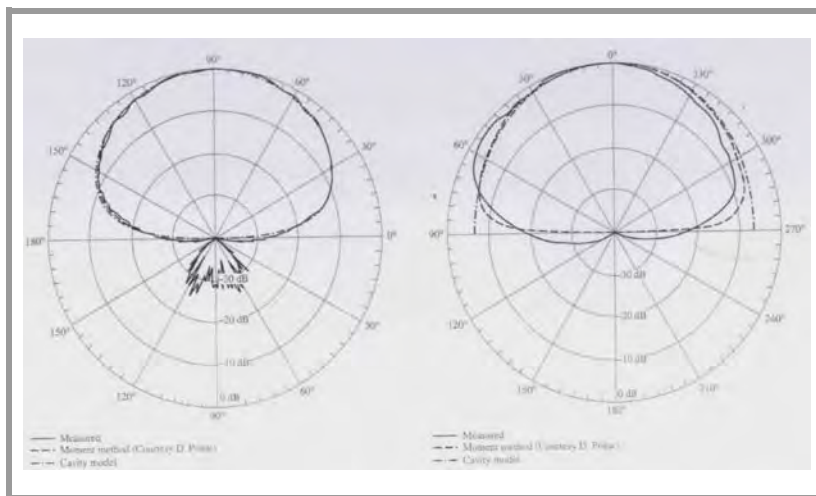
La cortina consente inoltre di modellare il solido di radiazione in base a particolari esigenze e può essere dotata di uno schermo riflettente qualora si voglia irradiare in un solo verso.

ANTENNE RICEVENTI PER TELEFONIA CELLULARE

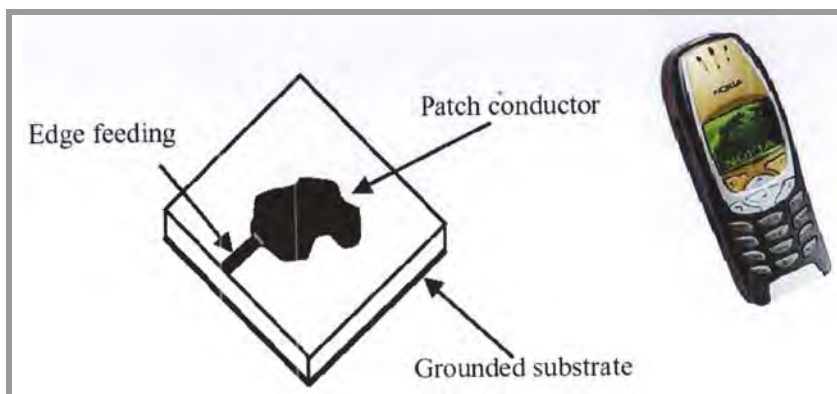
1) ANTENNE “PATCH” O “A MICROSTRISCIA

Sono realizzate mediante una **striscia (patch)** di **materiale conduttore**, stampata su un **supporto dielettrico** (per es: vetronite) che presenta una **metallizzazione sulla faccia opposta**. Le dimensioni sono dell'ordine di alcune decine di millimetri.

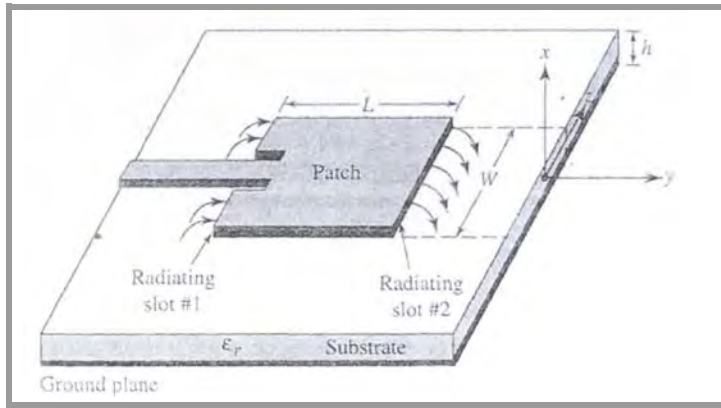
Fra la striscia e la metallizzazione resta intrappolato un campo elettromagnetico. L'irradiazione è resa possibile dai bordi del patch, in corrispondenza dei quali si creano delle fenditure che si comportano come “aperture”.



Il **solido di radiazione** si sviluppa, in maniera uniforme, nel **semispazio superiore** alla striscia.



Il patch è, come il dipolo a $\lambda/2$, una **struttura risonante**, per cui **la banda è stretta**. Si può allargare la banda sovrapponendo altri patch – strutture “stacked patch- che risuonino a frequenza leggermente diversa.



Modificando la geometria della striscia è possibile variare la frequenza di risonanza, la polarizzazione, il diagramma di radiazione (beam pattern) e l'impedenza dell'antenna. L'**alimentazione** del *patch* può avvenire tramite **cavo coassiale** o **linea a microstriscia** (variando la posizione del punto di alimentazione è possibile modificare l'impedenza d'ingresso dell'antenna e le caratteristiche del campo irradiato).

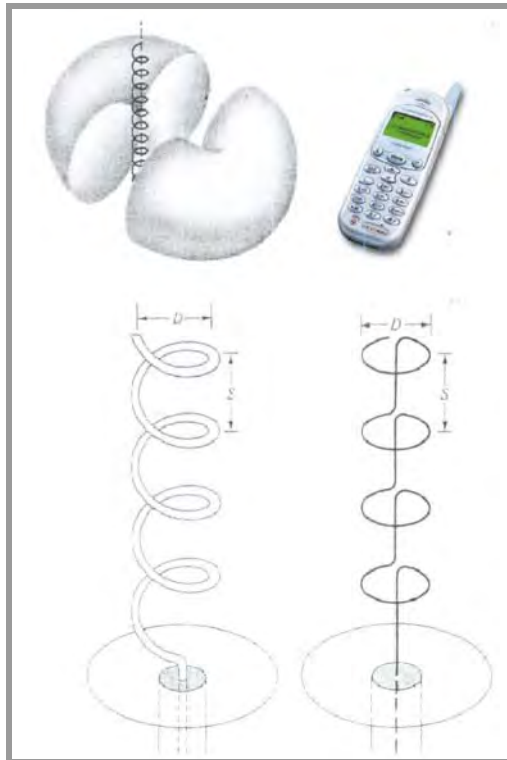
Gli inconvenienti di questo tipo di antenna sono la **bassa efficienza** e la **scarsa larghezza di banda**. Per aumentare la larghezza di banda è possibile incrementare lo spessore del substrato dielettrico; questa soluzione può però causare una diminuzione dell'efficienza globale dell'antenna per effetto dell'innesco di onde superficiali.

2) ANTENNE A ELICA

Hanno due modalità di funzionamento: il modo "normale" e il modo assiale".

- **Modo NORMALE**

Detta "L" la lunghezza complessiva dell'antenna, questa modalità si ha per $L \ll \lambda$. Nel loro funzionamento "normale" le antenne ad elica presentano **irradiazione massima** in direzione **perpendicolare** al loro **asse** (e minima lungo l'asse), descrivendo un solido di radiazione simile a quello del dipolo elementare.



- **Modo ASSIALE**

Questa modalità si ha quando le **dimensioni** dell'antenna (passo e diametro) sono **confrontabili** con la **lunghezza d'onda λ** .

Nel loro funzionamento “assiale” le antenne ad elica presentano **irradiazione massima lungo l'asse**.



CAMPI DI UTILIZZAZIONE DELLE ANTENNE RICEVENTI

A) ANTENNE PER ONDE LUNGHE E MEDIE (30KHz-3MHz)

SERVIZI:

- **RADIODIFFUSIONE** (su distanze di qualche migliaio di Km)
- **SERVIZIO MARITTIMO**

Presentano dimensioni comprese fra un decimo della lunghezza d'onda del segnale da captare e la lunghezza d'onda stessa.

Possono essere classificate in:

1. Antenne INTERNE ALL'EDIFICIO

- 1.1 antenna **FILARE**,
- 1.2 antenna **A FERRITE**

2. Antenne ESTERNE ALL'EDIFICIO

- 2.1 antenne **CARICATE**: antenna a "L rovesciata" e antenna a "T"

TIPOLOGIE

ANTENNA FILARE

E' ancora utilizzata negli apparecchi radio tradizionali.

E' un filo di rame (ricoperto, intrecciato o semplice), lungo circa 50 cm, con una estremità libera.

Le piccole dimensioni sono consentite dalla forte sensibilità dei moderni ricevitori.

ANTENNA A FERRITE

L'antenna a ferrite è generalmente interna al ricevitore ed è una bobina di filo isolato avvolta intorno a una bacchetta di materiale ferroso. I terminali di ingresso del radioricevitore vengono collegati ai capi del filo costituente la bobina di antenna.

ANTENNE ESTERNE CARICATE

Sono particolari dipoli a $\lambda/4$ e, in quanto dipoli, equivalgono a circuiti risonanti.

Il nome di queste antenne è dovuto al fatto che sono state caricate con un componente reattivo (induttore o condensatore) al fine di ottenere la stessa frequenza di risonanza di un'antenna più lunga.

Hanno una lunghezza compresa, orientativamente, fra i 3 e i 15 metri.

Sono realizzate con una treccia di rame disposta orizzontalmente rispetto al suolo e perpendicolare alla direzione del segnale da captare.

B) ANTENNE PER ONDE CORTE
(3-30) MHz

SERVIZI:

- **RADIODIFFUSIONE**

TIPOLOGIE

1. **DIPOLO ORIZZONTALE A $\lambda/2$**
2. **DIPOLO VERTICALE A $\lambda/4$ (Dip. Marconiano o “a STILO”)**
3. **DIPOLI CARICATI**
4. **ANTENNE YAGI**

C) ANTENNE PER SISTEMI:

- **RADIO FM**
- **TELEVISIVI NELLE BANDE I e II (VHF)**

TIPOLOGIE

1. **DIPOLO ORIZZONTALE A $\lambda/2$**
2. **DIPOLO VERTICALE A $\lambda/4$ (Dip. Marconiano o “a STILO”) RIPIEGATO**
3. **ANTENNE YAGI**

D) ANTENNE RICEVENTI PER SISTEMI TELEVISIVI IN UHF

TIPOLOGIE

- **ANTENNE A LARGA BANDA**

E) ANTENNE PER MICROONDE

SERVIZI:

- **RADAR**
- **TELEVISIVI NELLE BANDE I e II (VHF)**

TIPOLOGIE

- **ANTENNE A SUPERFICIE (PARABOLE)**

CAPITOLO 36

PLL

OBIETTIVI

- **Conoscere lo schema a blocchi del PLL**
- **Comprendere il funzionamento del PLL**
- **Conoscere le applicazioni fondamentali del PLL**
- **Comprendere e saper dimensionare i circuiti modulatori e demodulatori di frequenza basati su PLL**

PREREQUISITI

- **Parametri fondamentali dei segnali sinusoidali**
- **La modulazione e la demodulazione di frequenza**
- **Schema a blocchi di un sistema di controllo in catena chiusa**
- **Filtri**

PLL (PHASE LOCKED LOOP) O “ANELLO AD AGGANCIO DI FASE”

Il PLL è un integrato di basso costo che può svolgere diverse funzioni tra le quali quella di demodulatore di segnali FM (dispositivo che può funzionare anche da demodulatore FSK).

Applicando in ingresso al PLL il segnale modulato FM, si ottiene in uscita un segnale proporzionale al segnale modulante informativo.

Oltre che come demodulatore FM, il PLL può essere usato come:

- modulatore FM
- demodulatore PSK
- dispositivo per la ricostruzione delle portanti DSB ed SSB
- dispositivo per il controllo automatico di frequenza (CAF) nei radioricevitori
- moltiplicatore e sintetizzatore di frequenza
- filtro attivo.

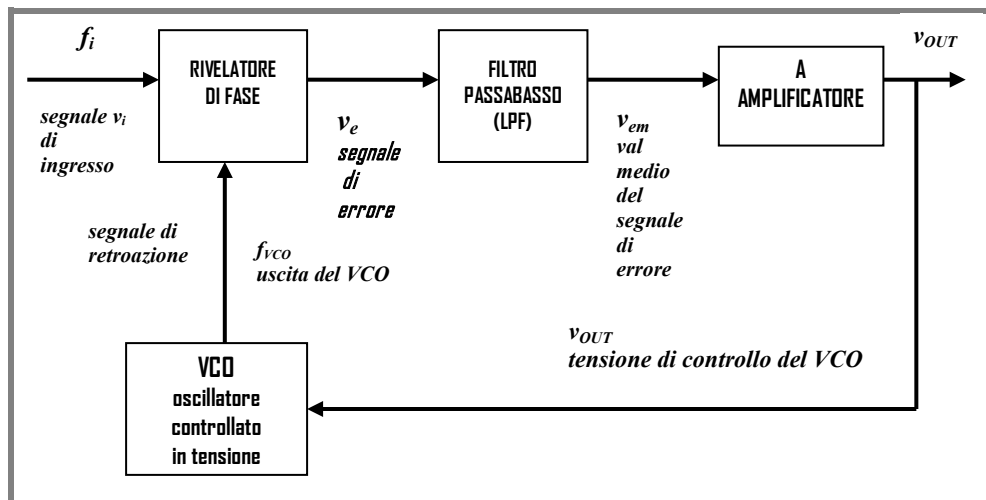
Il PLL è un sistema retroazionato negativamente e quindi può essere visto come sistema di controllo. Esso è in grado **di rendere uguali tra loro due frequenze**.

Le due frequenze che vengono rese uguali sono:

- la frequenza f_i del segnale applicato in ingresso
- la frequenza f_{VCO} del segnale di retroazione prodotto dall'oscillatore controllato in tensione (VCO), presente all'interno del PLL.

Osserviamo che i segnali le cui frequenze vengono rese uguali sono i due ingressi del rivelatore o comparatore di fase.

A seconda del tipo di PLL la differenza di fase tra i due segnali può essere resa nulla o costante.



Il PLL è costituito da:

- un circuito rivelatore (o comparatore) di fase che rappresenta il “nodo di confronto” del sistema e fornisce il segnale di errore
- un filtro passabasso (anche se necessario al funzionamento di diversi dispositivi basati sul PLL, il filtro può non essere interno all’integrato, ma deve essere realizzato esternamente, come accade per l’integrato digitale 4046 e per l’integrato analogico 565)
- un oscillatore controllato in tensione (VCO), che fornisce il segnale di retroazione
- un amplificatore.

**PREMESSA AL FUNZIONAMENTO DEL PLL:
LA FASE ISTANTANEA DI UN SEGNALE SINUSOIDALE**

Dato un segnale sinusoidale di **pulsazione** ω_0 e di **fase iniziale** Φ_0 :

$$v(t) = V \cdot \sin(\omega_0 \cdot t + \phi_0)$$

Si definisce **fase istantanea** del segnale la quantità:

$$\phi(t) = \omega_0 \cdot t + \phi_0$$

Dati due segnali con fasi istantanee $\phi_1(t)$ e $\phi_2(t)$, si definisce **sfasamento istantaneo** (fra i due segnali) la differenza:

$$\Delta\phi(t) = \phi_1(t) - \phi_2(t)$$

Osserviamo inoltre che **la pulsazione** $\omega_0=2\pi f_0$ di un segnale sinusoidale **può essere vista come derivata temporale della fase istantanea**:

$$\frac{d}{dt}\phi(t) = \frac{d}{dt}(\omega_0 \cdot t + \phi_0) = \omega_0$$

e che quindi la fase istantanea può essere vista come integrale della pulsazione:

$$\phi(t) = \int_0^t \omega_0 \cdot dt + k \quad (\text{con } k = \text{costante})$$

FUNZIONAMENTO DEL PLL

Il comparatore o rivelatore di fase fornisce in uscita una tensione di errore $v_e(t)$ il cui valore medio $v_{em}(t)$, rivelato dal filtro passabasso, è proporzionale allo sfasamento istantaneo:

$\phi_i(t) - \phi_{vco}(t)$ che si può anche esprimere come¹:

$$(\omega_i - \omega_{vco}) \cdot t + (\phi_i - \phi_{vco})$$

Il valore medio della tensione di errore (uscita del filtro passabasso) viene prima amplificato e poi retroazionato dal VCO, verso il quale si comporta come segnale di controllo, imponendo alla frequenza f_{vco} di uscita del VCO una variazione che tende a ridurre la differenza $f_i - f_{vco}$ tra la frequenza prodotta in uscita dal VCO e la frequenza di ingresso del PLL.

Quando i valori di f_{vco} e di f_i sono sufficientemente vicini, la retroazione negativa del PLL impone al VCO di “agganciarsi” al segnale di ingresso (cioè rende la frequenza di uscita f_{vco} del VCO uguale alla frequenza f_i del segnale di ingresso).

Il VCO quindi varia la sua frequenza in modo da ridurre le variazioni della differenza fra le fasi istantanee.

¹ Infatti:

$$\phi_i(t) = \omega_i \cdot t + \phi_i$$

$$\phi_{vco}(t) = \omega_{vco} \cdot t + \phi_{vco}$$

$$\phi_i(t) - \phi_{vco}(t) = \omega_i \cdot t + \phi_i - (\omega_{vco} \cdot t + \phi_{vco}) = (\omega_i - \omega_{vco}) \cdot t + (\phi_i - \phi_{vco})$$

FREE RUN

Si dice “frequenza di free run”, e si indica con f_{FR} , la **frequenza che viene mantenuta “spontaneamente” all’uscita del VCO quando la tensione di ingresso o di controllo del VCO è nulla.**

Viene detto “stato di free run” o anche “stato di non aggancio” lo stato di attesa nel quale il sistema non risulta sensibile alla frequenza del segnale di ingresso e la frequenza presente all’uscita del PLL è quella di free run del VCO.

STATO DI CATTURA

E’ la **condizione transitoria** fra lo stato di free run e quello di aggancio, nella quale il VCO, pilotato dal valor medio della tensione di errore (cioè dall’uscita del filtro), **varia con continuità la sua frequenza di uscita (f_{VCO}) finché questa non diventa uguale alla frequenza di ingresso f_i del comparatore e dell’intero sistema.**

STATO DI AGGANCIO

Nello stato di aggancio, che è lo stato di regime del PLL, e che è anche l’unica condizione di equilibrio possibile, le **frequenze f_{VCO} e f_i , applicate al comparatore di fase, sono uguali**, mentre, **a seconda del comparatore usato, le fasi possono essere anch’esse uguali ($\Phi_i = \Phi_{VCO}$) o differire di una piccola quantità costante ($\Phi_i - \Phi_{VCO} = \Phi_d = \text{costante}$).**

PLL ANALOGICI E PLL DIGITALI

Si parla di PLL “digitali” e di PLL “analogici”.

Un PLL è di tipo analogico quando il rivelatore di fase è un mixer (è questo il caso di alcuni PLL integrati come il 565).

Un PLL è di tipo digitale quando il rivelatore di fase è un circuito EXOR oppure un circuito “edge triggered”, cioè sincronizzato su un fronte d’onda, come un flip-flop (è questo il caso del PLL integrato 4046).

Il segnale di errore $v_e(t)$, se il rivelatore è analogico, è un segnale analogico che contiene componenti proporzionali alla differenza $(\omega_i - \omega_{VCO})$, alla differenza $(\Phi_i - \Phi_{VCO})$, ma anche componenti che devono essere eliminate mediante filtraggio.

Se il rivelatore è digitale, il segnale di errore è una tensione a due soli livelli, derivante, istante per istante, dall’operazione di exor fra il segnale di ingresso e il segnale di retroazione, oppure dall’azione di un flip-flop.

Un **rivelatore di fase digitale** è un circuito che fornisce in uscita un segnale il cui **valore medio**, ottenuto mediante filtraggio passabasso, è una **tensione proporzionale alla differenza tra le fasi istantanee (sfasamento istantaneo)** dei due ingressi, ossia proporzionale alla quantità:

$$\phi_i(t) - \phi_{VCO}(t)$$

che si può anche esprimere come:

$$(\omega_i - \omega_{VCO}) \cdot t + (\phi_i - \phi_{VCO})$$

Il valore medio dell’uscita del rivelatore di fase risulta quindi dipendente anche dalle quantità:

$$\omega_i - \omega_{VCO} \quad (\text{differenza di pulsazioni e quindi di frequenze})$$

$$\phi_i - \phi_{VCO} \quad (\text{differenza di fasi iniziali})$$

IL PLL DIGITALE 4046

L'integrato 4046 è realizzato in tecnologia CMOS, a 16 piedini, e può essere alimentato in modo singolo o duale.

Esso contiene:

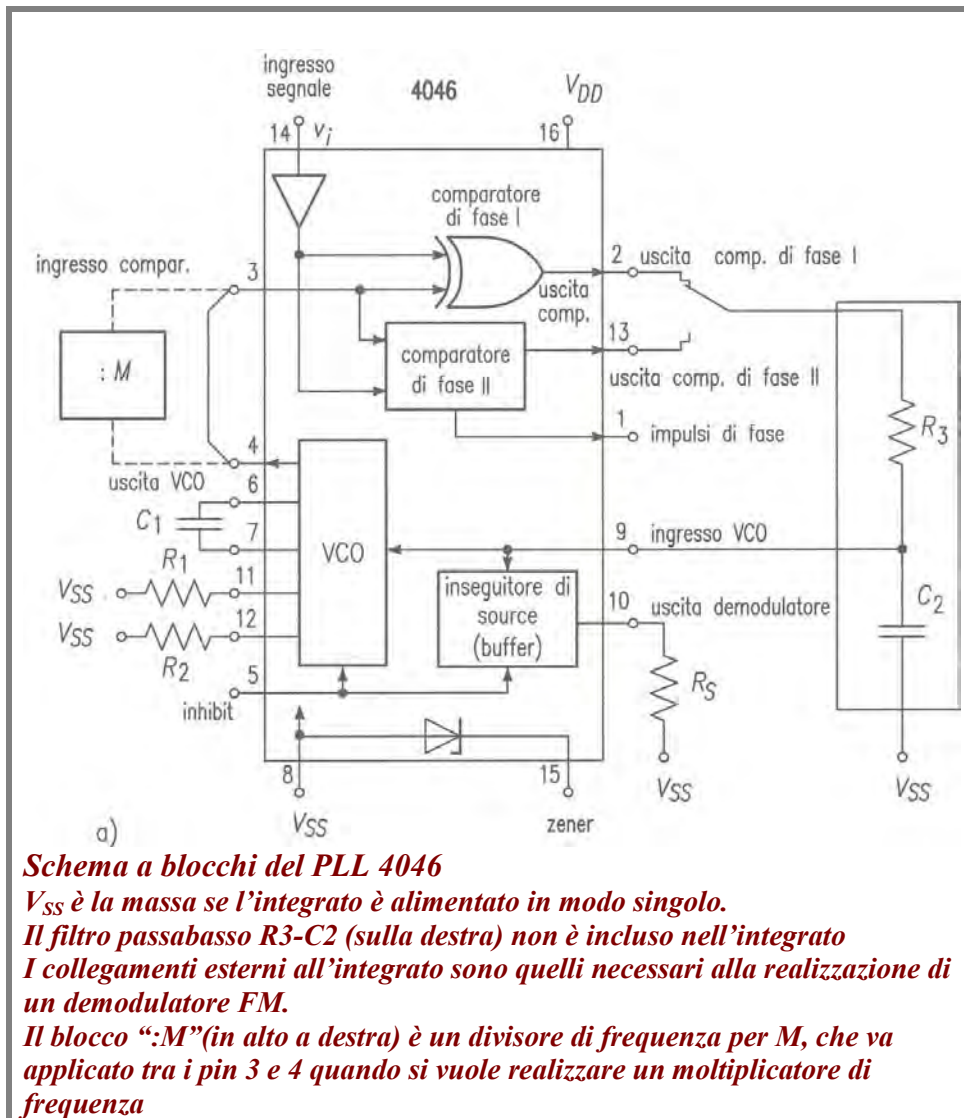
- **due circuiti rivelatori (o comparatori) di fase** che rappresentano il “nodo di confronto” del sistema e forniscono il segnale di errore
- un **oscillatore controllato in tensione (VCO)**, che fornisce il segnale di retroazione
- un **amplificatore**.

Il filtro passabasso invece non è presente all'interno dell'integrato, ma deve essere realizzato esternamente con un condensatore e una resistenza.

Il condensatore va posto fra l'ingresso del VCO (pin 9) e la massa (o l'alimentazione negativa se l'integrato è alimentato in modo duale).

La resistenza va posta fra l'ingresso del VCO (pin 9) e l'uscita di uno dei due comparatori di fase (pin 13 oppure 2).

La presenza del filtro è essenziale perché esso, in base alla tensione di ingresso applicatagli dal comparatore di fase, fornisce un diverso livello del segnale di uscita, che fa da segnale di controllo per il VCO.



IL VCO

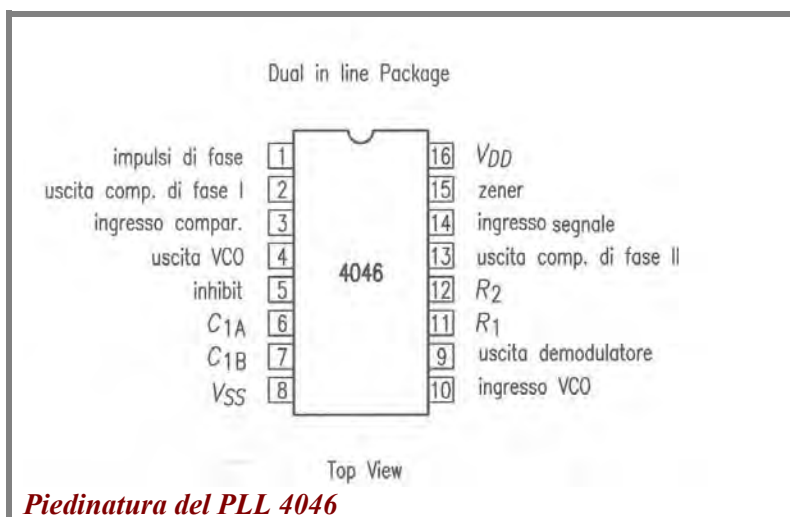
La massima frequenza di oscillazione del VCO dipende dalla tensione di alimentazione:

TENSIONE DI ALIMENTAZIONE	MASSIMA FREQUENZA DI OSCILLAZIONE
5V	600KHz
12V	1,2MHz

I COMPARATORI O RIVELATORI DI FASE

Nel 4046 sono presenti due comparatori che hanno l'ingresso in comune (pin 3) e uscite distinte:

- il “**comparatore 1**”, detto anche “comparatore a basso rumore” basato su una porta **EXOR**; può lavorare nell'intervallo di sfasamento **$(0 \div \pi)$** e richiede un segnale di **ingresso digitale**, con duty-cycle del 50%; la sua **uscita** è sul **pin 2**;
- il “**comparatore 2**”, detto anche “comparatore a larga banda” basato su un dispositivo sincronizzato sul fronte d'onda, come un **flip-flop SR**; può lavorare **nell'intervallo di sfasamento $(0 \div 2\pi)$ più ampio** di quello del “comparatore 1” (il che garantisce campi di cattura e di aggancio più estesi) e **funziona anche con un segnale di ingresso analogico**; risulta quindi più versatile del comparatore 1; la sua **uscita** è sul **pin 13**.



SEGNALE DI ERRORE DEL COMPARATORE DI TIPO 2

Nel PLL 4046 il segnale di errore non è la differenza analogica fra la tensione di ingresso del PLL e la tensione di uscita del VCO, ma è un segnale a due soli livelli, che è applicato all'ingresso del filtro passabasso e che può configurarsi come:

- livello **basso** (massa o tensione di alimentazione negativa) quando si ha $f_i < f_{VCO}$, cioè quando la frequenza del segnale di uscita del VCO è maggiore di quella di ingresso del PLL (questo livello basso, fa **scaricare il condensatore** del filtro e l'uscita del filtro, applicata al VCO come segnale di controllo, fa **diminuire** f_{VCO} e quindi tende a ridurre la differenza fra f_i e f_{VCO})
- livello **alto** quando si ha $f_i > f_{VCO}$, cioè quando la frequenza del segnale di uscita del VCO è minore di quella di ingresso del PLL (questo livello alto, fa **caricare il condensatore** del filtro e l'uscita del filtro, applicata al VCO come segnale di controllo, fa **aumentare** f_{VCO} e quindi, anche in questo caso, tende a ridurre la differenza fra f_i e f_{VCO})
- **treno di impulsi** (che tende ad annullare la differenza fra le fasi istantanee) se $f_i = f_{VCO}$ e in particolare:
- se $\Phi_i < \Phi_{VCO}$ (se cioè la fase iniziale dell'ingresso del PLL è minore della fase iniziale dell'uscita del VCO, e quindi se **il segnale di ingresso è in ritardo** rispetto all'uscita del VCO), si ha un treno di impulsi **negativi** con duty-cycle proporzionale alla deviazione di fase istantanea $\Delta\Phi = \Phi_i - \Phi_{VCO}$, che fa diminuire il valore di fase del segnale del VCO fino all'annullamento della differenza $\Phi_i - \Phi_{VCO}$ fra le fasi iniziali
- se $\Phi_i > \Phi_{VCO}$ (se cioè la fase iniziale del segnale di ingresso è maggiore della fase iniziale dell'uscita del VCO e quindi se **il segnale di ingresso è in anticipo** rispetto all'uscita del VCO) si ha un treno di impulsi **positivi** con duty-cycle proporzionale alla deviazione di fase istantanea $\Delta\Phi = \Phi_i - \Phi_{VCO}$, che fa aumentare il valore di fase del segnale del VCO fino all'annullamento della differenza $\Phi_i - \Phi_{VCO}$ fra le fasi iniziali- se infine $\Phi_i = \Phi_{VCO}$ (se cioè la fase iniziale del segnale di ingresso è uguale alla fase iniziale dell'uscita del VCO) l'uscita del "comparatore 2" va in stato di alta impedenza, per cui la tensione sul condensatore del filtro rimane invariata, lasciando invariata anche la tensione di controllo del VCO. Si mantiene così uno sfasamento nullo

fra le tensioni di ingresso del PLL e di uscita del VCO (situazione che si presenta sempre quando è verificata la condizione di aggancio).

INTERVALLO DI FREQUENZE DI LAVORO

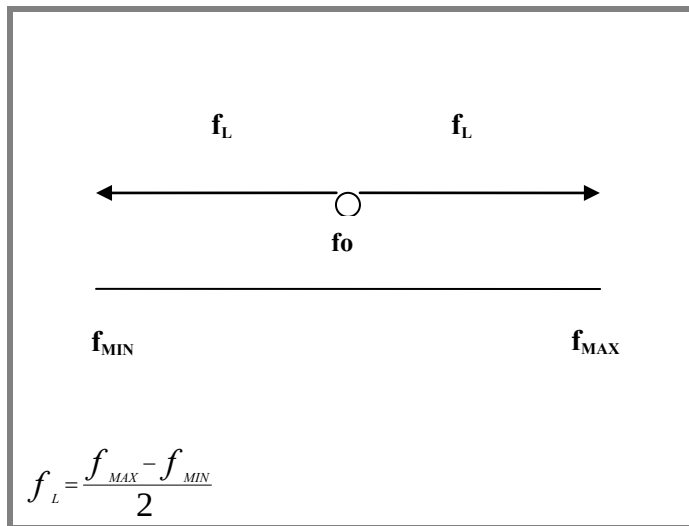
Le frequenze $f_{\text{MIN}} \div f_{\text{MAX}}$ sono i valori-limite di frequenza del VCO interno al PLL.

L'intervallo di frequenze di lavoro ($f_{\text{MIN}} \div f_{\text{MAX}}$), desiderate per il PLL (e la frequenza centrale f_0 di tale intervallo) possono essere scelti attraverso il dimensionamento:

- della resistenza R_1 posta fra il **pin 11** del VCO interno al PLL e la **massa** (o l'alimentazione negativa se l'integrato è alimentato in maniera duale)
- della resistenza R_2 posta fra il **pin 12** del VCO interno al PLL e la **massa** (o l'alimentazione negativa se l'integrato è alimentato in maniera duale)
- del condensatore C_1 posto fra i **pin 6 e 7** del VCO interno al PLL.

In particolare osserviamo che:

- il valore **centrale** f_0 del range di frequenza dipende **inversamente** da R_1 e C_1
- la presenza di una f_{MIN} diversa da zero, cioè di un **offset**, dipende da R_2 ;
in assenza di R_2 cioè con il pin 12 non collegato ($R_2=\infty$), si ha $f_{\text{MIN}}=0$;
in presenza di R_2 , la f_{MIN} dipende **inversamente** da R_2 e C_1 .



Per le due resistenze R_1 ed R_2 sono consigliati valori compresi fra $10 \text{ K}\Omega$ e $1 \text{ M}\Omega$.

Per C_1 valori a partire da 50 pF .

DIMENSIONAMENTO DEI CIRCUITI BASATI SUL PLL 4046

Distingueremo più casi a seconda dell'intervallo di frequenze di funzionamento richiesto e delle grandezze assegnate:

Intervallo ($f_{\text{MIN}} \div f_{\text{MAX}}$), con f_{MIN} diverso da zero (cioè con offset)

CASO 1

Sono assegnati:

f_0 = frequenza centrale

f_L = semiampiezza dell'intervallo ($f_{\text{MIN}} \div f_{\text{MAX}}$)

PROCEDIMENTO

- Si calcolano $f_{\text{MIN}} = f_0 - f_L$ ed $f_{\text{MAX}} = f_0 + f_L$
- dal diagramma “B” (che fornisce f_{MIN} in funzione di R_2 e C_1), in base anche alla tensione di alimentazione scelta, si ricavano
 - la capacità di accordo C_1 del VCO
 - la resistenza R_2
- dal diagramma “C” (che fornisce R_2/R_1 in funzione di $f_{\text{MAX}}/f_{\text{MIN}}$), in base anche alla tensione di alimentazione scelta, si determina R_2/R_1 dal quale (avendo già calcolato R_2) si ricava R_1

CASO 2

Sono assegnati

f_{MIN} = frequenza minima (offset)

f_{MAX} = frequenza massima

PROCEDIMENTO

- dal diagramma “B” (che fornisce f_{MIN} in funzione di R_2 e C_1), in base anche alla tensione di alimentazione scelta, si ricavano
 - la capacità di accordo C_1 del VCO
 - la resistenza R_2
- dal diagramma “C” (che fornisce R_2/R_1 in funzione di $f_{\text{MAX}}/f_{\text{MIN}}$), in base anche alla tensione di alimentazione scelta, si determina R_2/R_1 dal quale (avendo già calcolato R_2) si ricava R_1

Intervallo ($0 \div f_{MAX}$), con $f_{MIN} = 0$ (cioè senza offset)

E' sempre $R_2 = \infty$ e quindi il pin 12 è non collegato

CASO 1

E' assegnata f_0 = frequenza centrale = $f_{max}/2$

PROCEDIMENTO

- dal diagramma "A" (che fornisce f_0 in funzione di R_1 e C_1), in base anche alla tensione di alimentazione scelta, si ricavano
 - la capacità di accordo C_1 del VCO
 - la resistenza R_1

CASO 2

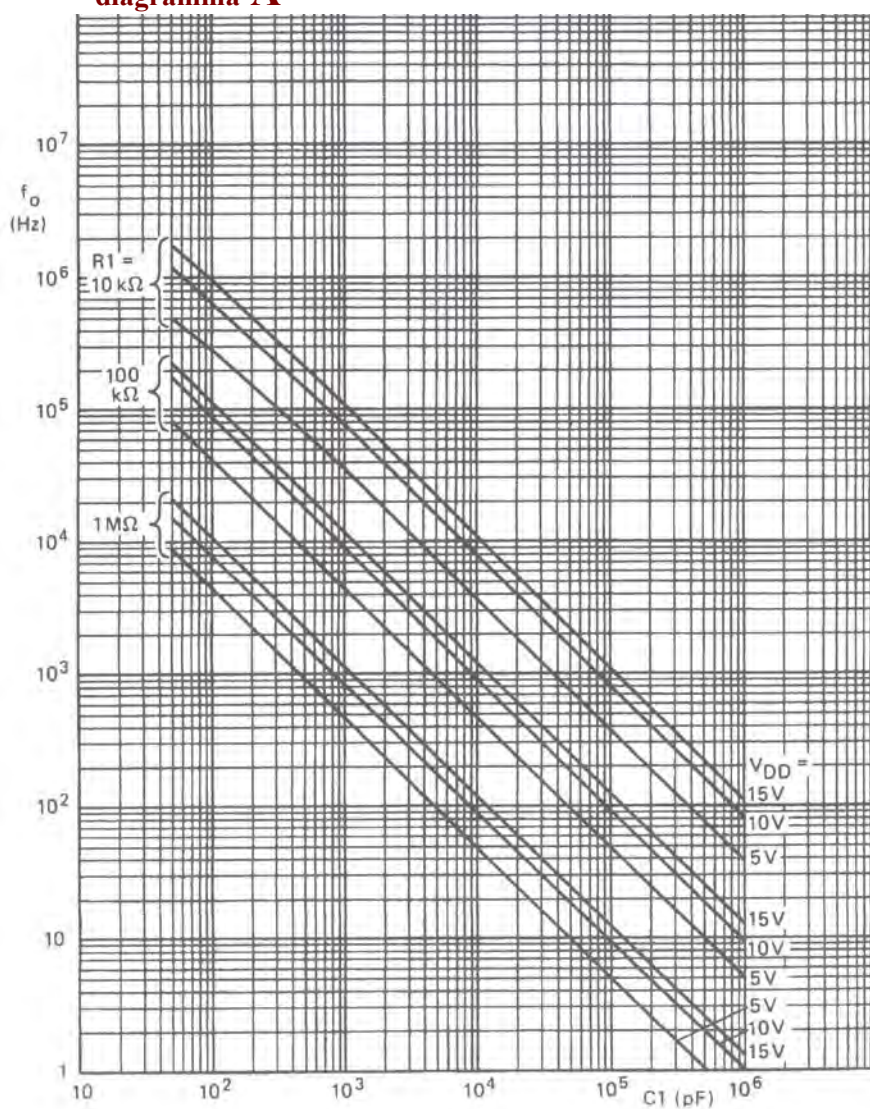
E' assegnata f_{MAX} = frequenza massima

PROCEDIMENTO

- si calcola $f_0 = \frac{f_{MAX}}{2}$
- come nel caso precedente, dal diagramma "A" (f_0 in funzione di R_1 e C_1), in base anche alla tensione di alimentazione scelta, si ricavano
 - la capacità di accordo C_1 del VCO
 - la resistenza R_1

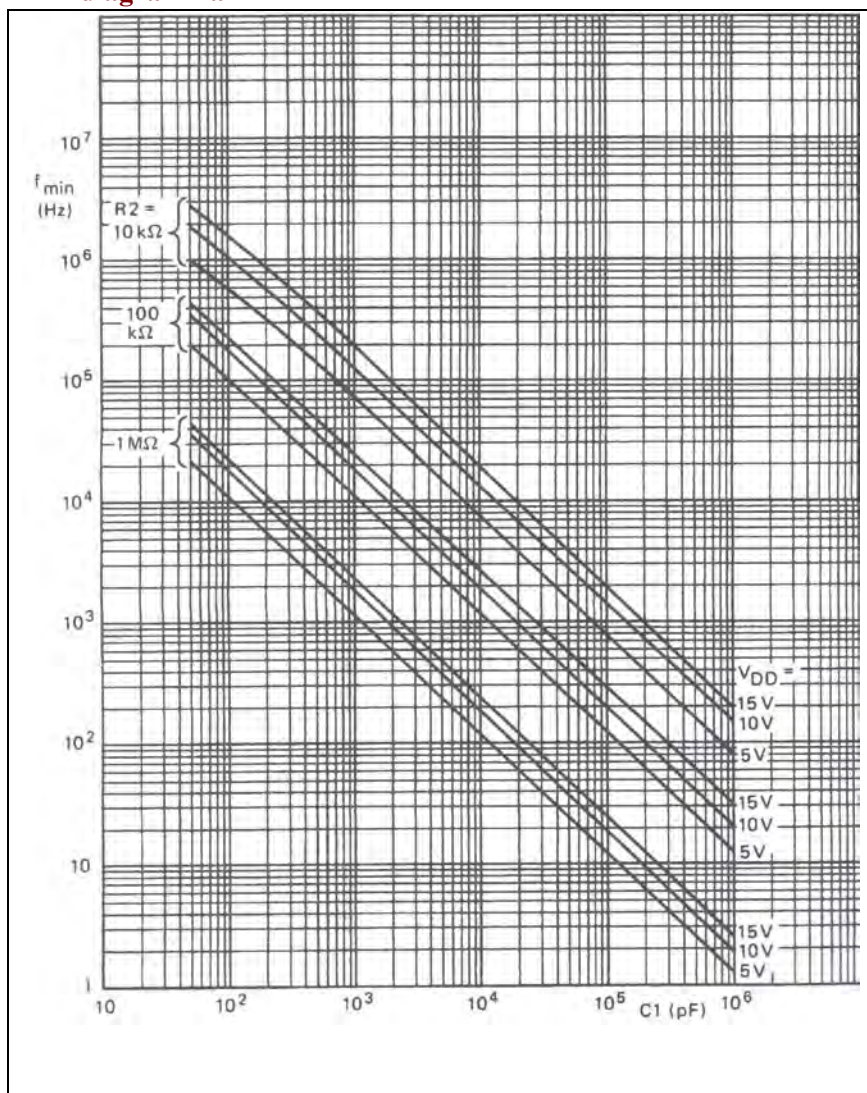
DIAGRAMMI PER IL DIMENSIONAMENTO DEI CIRCUITI BASATI SUL PLL 4046

diagramma A



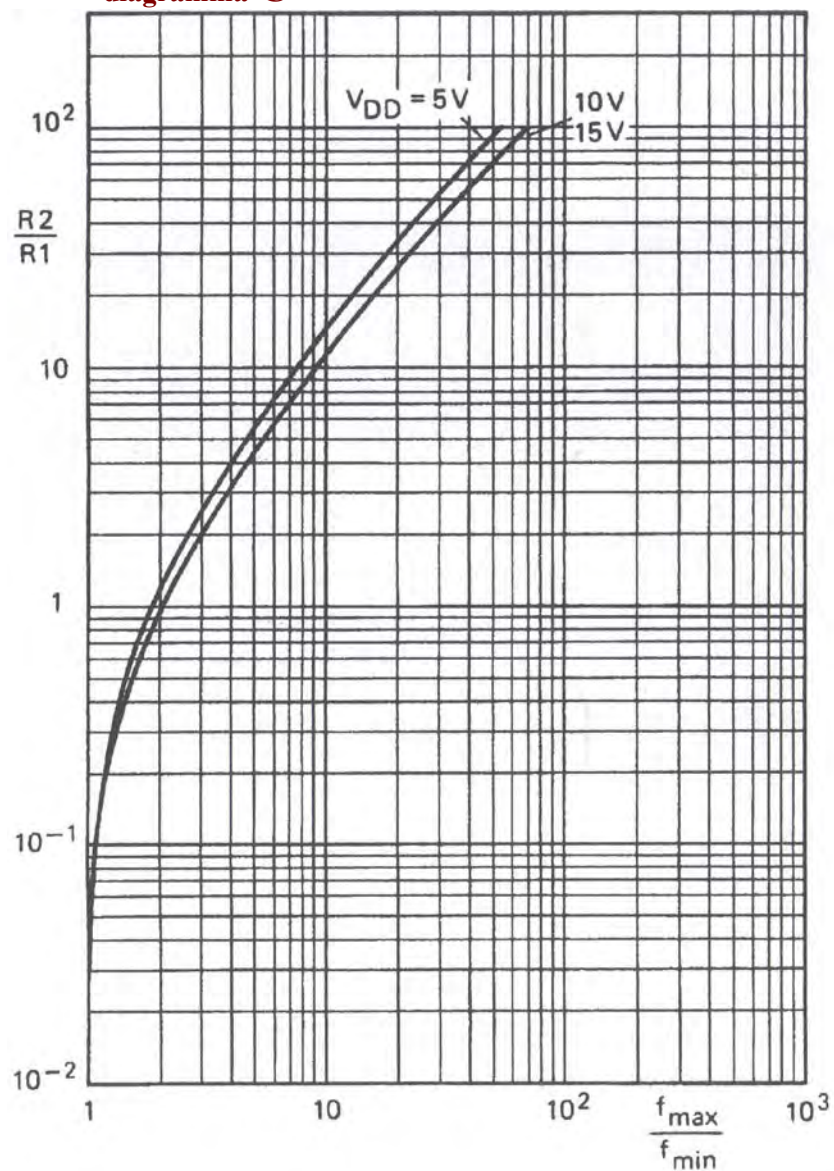
Diagrammi per il dimensionamento dei circuiti basati sul PLL 4046
Frequenza centrale f_o in funzione di R_1 e C_1

diagramma B



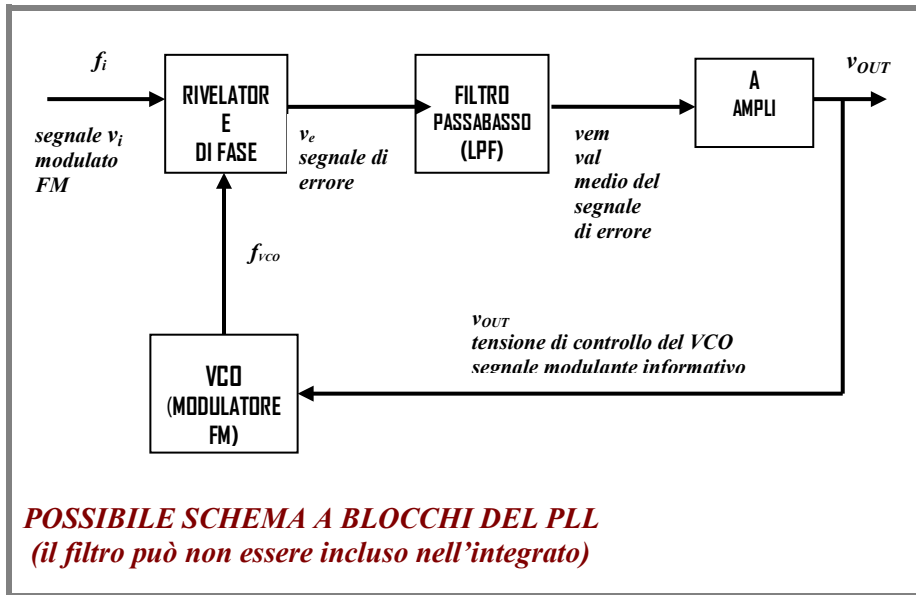
Diagrammi per il dimensionamento dei circuiti basati sul PLL 4046
Frequenza minima $f_{\min}$ (offset) in funzione di R_2 e C_1

diagramma C



Diagrammi per il dimensionamento dei circuiti basati sul PLL 4046
 R_2/R_1 in funzione di $f_{\max}/f_{\min}$

DEMODULAZIONE FM MEDIANTE PLL



IL PLL COME DEMODULATORE DI FREQUENZA

Il PLL è un sistema di controllo e quindi tende a rendere uguali il segnale di ingresso e il segnale di retroazione, che è l'uscita di un VCO (contenuto nel PLL stesso).

Come segnale di ingresso possiamo applicare al PLL il segnale FM da demodulare.

Il VCO è un modulatore di frequenza: fornisce in uscita un segnale che risulta modulato in frequenza dalla tensione di controllo (v_{OUT} nello schema) applicata al suo ingresso.

Questo significa che, nel momento in cui il segnale di uscita del VCO sia un segnale modulato FM, il segnale di controllo è il segnale modulante informativo.

Ora il PLL, in condizioni di equilibrio, rende l'uscita del VCO uguale al segnale di ingresso del PLL, cioè al segnale modulato in frequenza. Di conseguenza il segnale di controllo del VCO rappresenta proprio il segnale modulante informativo che si desidera.

FUNZIONAMENTO DEL PLL COME DEMODULATORE DI FREQUENZA

Il segnale FM modulato in frequenza, cioè il segnale da demodulare, è il segnale $v_i(t)$ che viene applicato all'ingresso del comparatore di fase e quindi dell'intero PLL.

La frequenza di tale segnale FM è la f_i di ingresso, cioè la frequenza del segnale modulato FM.

Quando il PLL è in condizioni di aggancio, la frequenza del segnale di ingresso (cioè la frequenza f_i , variabile, del segnale di ingresso) coincide con la frequenza f_{VCO} , anch'essa variabile, di uscita del VCO, il che significa che l'uscita del VCO ha una frequenza che (in condizioni di aggancio) varia come quella del segnale modulato in frequenza.

Ora, siccome l'oscillatore VCO è un modulatore di frequenza, il fatto che alla sua uscita vi sia la frequenza del segnale modulato FM significa che la sua tensione di controllo $v_{OUT}(t)$ rappresenta il segnale informativo modulante, cioè il segnale che noi desideriamo ottenere dal PLL .

Prelevando dunque la tensione di uscita $v_{OUT}(t)$ del PLL, la quale funge anche da tensione di controllo per il VCO, si ottiene il segnale modulante informativo, ossia si demodula il segnale $v_i(t)$ di ingresso modulato in frequenza (segnale FM)

IL PLL COME MOLTIPLICATORE E COME SINTETIZZATORE DI FREQUENZA

Il PLL può funzionare anche come **sintetizzatore di frequenze**, cioè come un **dispositivo che, a partire da una frequenza di riferimento f_i , è in grado di fornire, all'uscita del VCO, una frequenza che sia, con f_i , in una relazione stabilita dall'utente.**

Per comprendere come il PLL possa funzionare da sintetizzatore, vediamo prima come può funzionare da semplice **moltiplicatore di frequenza**.

MOLTIPLICATORE DI FREQUENZA

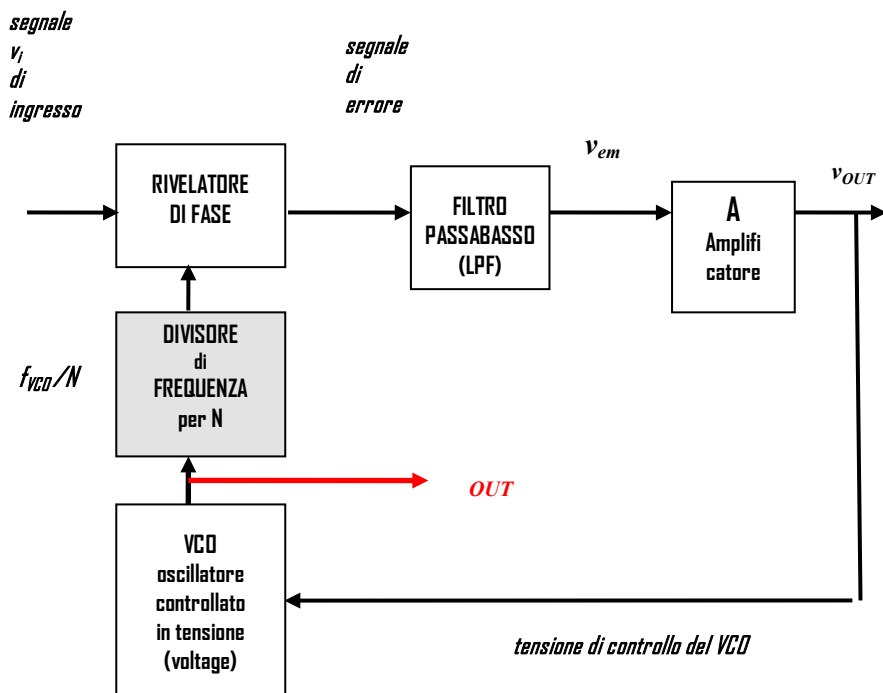
Basta **aggiungere**, fra il comparatore di fase e il VCO, un “**divisore di frequenza per N**” (per esempio un “contatore modulo N”, adoperabile una volta rese compatibili con esso le altre grandezze del circuito).

Siccome, nella condizione di aggancio, devono risultare uguali tra loro le frequenze presenti ai due ingressi del comparatore di fase, dovrà essere:

$$f_i = \frac{f_{VCO}}{N}$$

da cui:

$$f_{VCO} = f_i \cdot N$$



MOLTIPLICATORE DI FREQUENZA
 Fornisce, all'uscita del VCO, la frequenza:
 $f_{VCO} = N f_i$

SINTETIZZATORE DI FREQUENZA

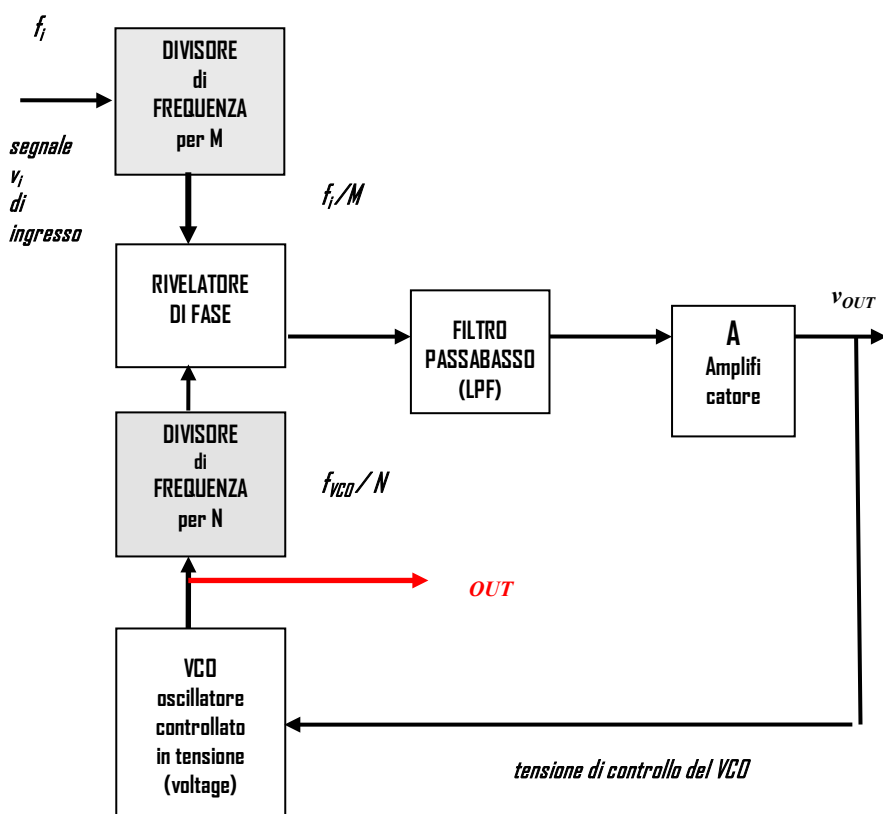
Con lo schema riportato di seguito si ottiene, se N ed M sono numeri interi fissati dall'utente, e se f_i è una frequenza campione, un sintetizzatore di frequenza.

Nella condizione di aggancio le due frequenze di ingresso del comparatore di fase devono essere uguali e quindi deve risultare:

$$\frac{f_i}{M} = \frac{f_{VCO}}{N}$$

da cui:

$$f_{VCO} = f_i \cdot \frac{N}{M}$$



SINTETIZZATORE DI FREQUENZA

Fornisce, all'uscita del VCO, la frequenza:

$$f_{VCO} = f_i \cdot \frac{N}{M}$$

CAPITOLO 40

PRINCIPI E TECNICHE DI AUTOMAZIONE

COMPLEMENTI

SPECIFICHE DI PROGETTO DI UN SISTEMA DI CONTROLLO E RETI CORRETRICI O DI COMPENSAZIONE

STABILITÀ

PRECISIONE

VELOCITÀ O RAPIDITÀ DI RISPOSTA (PRONTEZZA)

SODDISFACIMENTO DELLE SPECIFICHE IN FASE DI PROGETTO DEL SISTEMA DI CONTROLLO

RETI CORRETRICI O RETI DI COMPENSAZIONE

CONTROLLORE: SPECIFICHE

ERRORE A REGIME (O ERRORE STATICO), SEGNALE DI RIFERIMENTO E SEGNALE DI RETROAZIONE

ESPRESSIONE DELL'ERRORE A REGIME NELLA VARIABILE "S" DI LAPLACE

ERRORE A REGIME, TIPO DI SISTEMA E GRADO DEL SEGNALE DI INGRESSO

SPECIFICHE NEL TEMPO SULLA RAPIDITÀ DI RISPOSTA DI UN SISTEMA CONTROLLATO (PARAMETRI DELLA RISPOSTA AL GRADINO)

SPECIFICHE IN FREQUENZA SULLA RAPIDITÀ DI RISPOSTA DI UN SISTEMA CONTROLLATO (PARAMETRI DELLA RISPOSTA ARMONICA O RISPOSTA IN FREQUENZA)

METODI CLASSICI DI SINTESI DEL CONTROLLORE

STABILITÀ DEL SISTEMA DA CONTROLLARE E DEL SISTEMA IN CATENA CHIUSA

INFORMAZIONI SULLA STABILITÀ DEL SISTEMA RETROAZIONATO ATTRAVERSO L'ANALISI DELLA FUNZIONE $G(\omega)H(\omega)$

ENUNCIATO DEL CRITERIO DI BODE

MARGINI DI AMPIEZZA E DI FASE

RETI CORRETRICI O DI COMPENSAZIONE

SPECIFICHE DI PROGETTO DI UN SISTEMA DI CONTROLLO E RETI CORRETTRICI O DI COMPENSAZIONE

Le **specifiche** (ossia i requisiti, le qualità) di **stabilità, precisione e risposta in frequenza**, di cui tratteremo, si riferiscono al **SISTEMA DI CONTROLLO**.

Per “sistema di controllo” dobbiamo intendere un sistema di controllo in catena chiusa e cioè il **sistema retroazionato formato da**

- **campo (sistema da controllare, chiamato anche “processo” o “impianto”)**
- **controllore**
- **blocco di retroazione.**

In sostanza la funzione di un sistema di controllo è fare in modo che la grandezza controllata o grandezza di uscita si comporti (**una volta trascorso il transitorio iniziale**) conformemente all’andamento del segnale di riferimento.

*In maniera equivalente **il segnale di retroazione** o di misura (proporzionale all’andamento effettivo della grandezza controllata, ossia della grandezza di uscita) deve **riprodurre** il più fedelmente possibile **il segnale di riferimento** (detto anche **segnale di ingresso o set point**).*

Quando per esempio si vuole che la temperatura di un ambiente rimanga costante, il segnale di riferimento è una costante (graficamente una retta orizzontale) e il sistema di controllo dovrebbe fare in modo che, a regime, anche l’andamento temporale della grandezza controllata sia una retta orizzontale posta alla stessa quota.

Se invece si vuole che la temperatura aumenti linearmente nel tempo dobbiamo applicare come segnale di riferimento o di ingresso una rampa crescente e il sistema di controllo dovrebbe far sì che la temperatura (grandezza controllata) abbia nel tempo, a regime, un andamento a rampa crescente con la stessa pendenza.

Il sistema di controllo deve fare in modo:

- ❖ che la grandezza di uscita sia in grado di seguire l’andamento del segnale di ingresso o riferimento (requisito di stabilità)
- ❖ che la grandezza di uscita si adegui al segnale di riferimento con la velocità richiesta dall’efficacia del controllo (requisito di prontezza o velocità di risposta)
- ❖ che, **terminato il transitorio**, lo scostamento fra uscita e riferimento sia più piccolo possibile (requisito di precisione o requisito di errore a regime).

Le specifiche possono essere fissate nel **dominio del tempo** quando si impongono, per esempio, il tempo di salita o il tempo di assestamento

all'**andamento temporale della risposta** del sistema controllato ai segnali canonici (per esempio al gradino).

Oppure le specifiche possono riferirsi al *dominio della frequenza* quando impongono determinate caratteristiche (banda passante, picco di risonanza ecc.) alla *risposta in frequenza o risposta armonica*, del sistema controllato.

STABILITA'

Riguardo alla specifica di stabilità¹, dobbiamo precisare che essa è indispensabile anche riguardo al soddisfacimento delle altre due, perché se un sistema non è stabile, la sua risposta a una sollecitazione, anche di breve durata, tende ad assumere ampiezza infinita, il che comporta che l'uscita del sistema non potrà mai essere un grado di seguire l'andamento del segnale di riferimento.

La stabilità di un sistema viene definita facendo riferimento alla risposta impulsiva del sistema, ossia alla risposta all'impulso di Dirac e può essere analizzata studiando la funzione di trasferimento $G(s)$ del sistema e in particolare studiando i valori della variabile s per i quali il denominatore della $G(s)$ si annulla e la $G(s)$ tende all'infinito, studiando cioè i poli della $G(s)$.

PRECISIONE

Il segnale di retroazione o di misura (proporzionale all'andamento effettivo della grandezza controllata, ossia della grandezza di uscita) deve riprodurre il più fedelmente possibile il segnale di riferimento (detto anche segnale di ingresso o set point).

Lo scostamento a regime (per $t \rightarrow \infty$, cioè a transitorio finito) fra il segnale di retroazione e il segnale di riferimento si chiama “errore a regime” e deve essere o nullo o finito. Non deve cioè superare un valore massimo prefissato.

VELOCITA' O RAPIDITA' DI RISPOSTA (PRONTEZZA)

E' la velocità con la quale il sistema “reagisce” alle variazioni del segnale di riferimento o di ingresso.

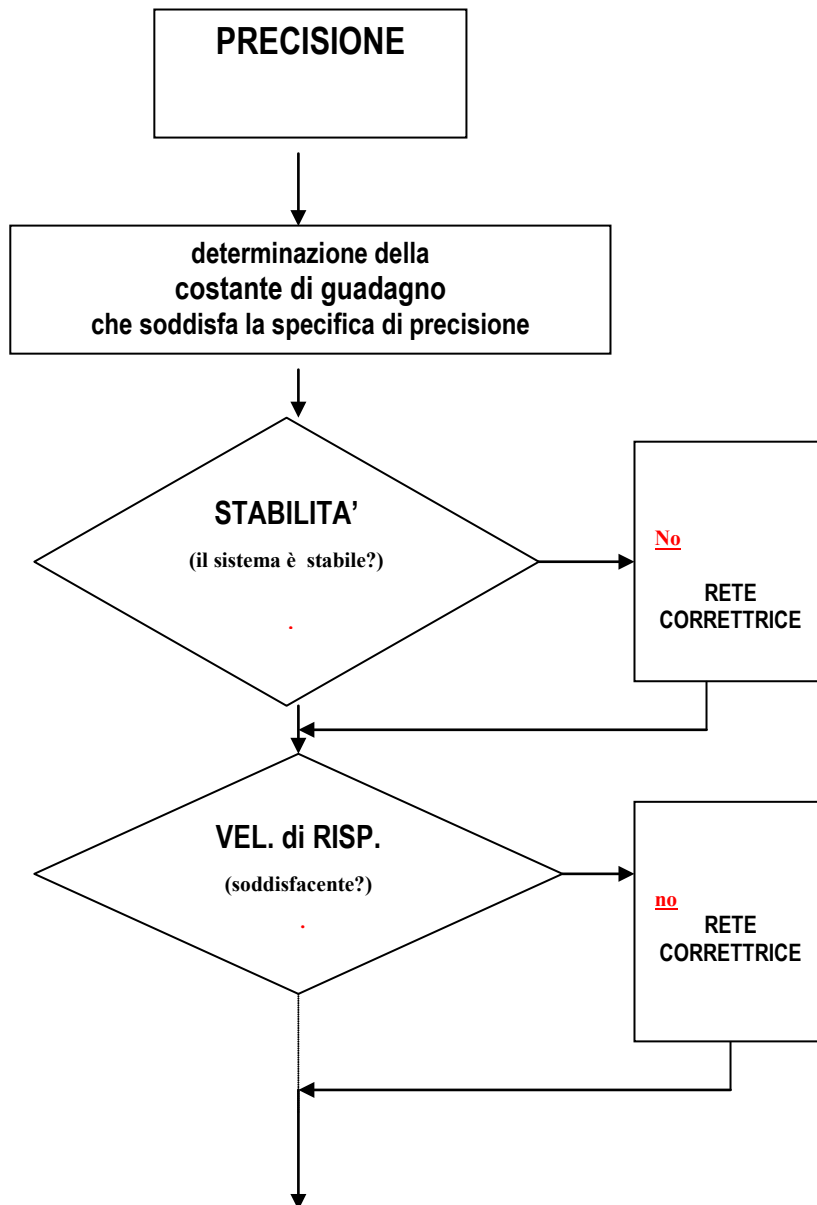
E' legata al tempo di estinzione dei transitori e al tempo di salita nella risposta al gradino.

¹ Le specifiche sulla stabilità riguardano in particolare il cosiddetto **marginale di ampiezza** e il cosiddetto **marginale di fase** (parametri che ci indicano di quanto il modulo e la fase del sistema possano variare senza che il sistema entri in uno stato di instabilità).

Le specifiche sulla stabilità possono anche riferirsi all'ampiezza del picco di risonanza della risposta armonica o alla sovraelongazione della risposta temporale a un segnale canonico.

SODDISFACIMENTO DELLE SPECIFICHE IN FASE DI PROGETTO DEL SISTEMA DI CONTROLLO

Tipicamente si procede secondo l'ordine seguente:



RETI CORRETTRICI O RETI DI COMPENSAZIONE

Nei sistemi di controllo ad amplificazione elettronica, le reti correttici sono circuiti elettronici che vengono inseriti del sistema di controllo per rimediare a **eventuali inconvenienti** (come l'**instabilità**) che possono crearsi quando il sistema da controllare **viene chiuso in retroazione**.

Queste reti possono avere anche la funzione di soddisfare alcune specifiche riguardanti le **prestazioni in transitorio** (e quindi la **velocità di risposta**) come:

- un determinato tempo di assestamento
- un determinato tempo di salita
- la sovraelongazione o overshoot
- un determinato valore della frequenza di taglio superiore nella risposta in frequenza.

Benché il sistema di controllo ad anello chiuso sia basato sulla retroazione negativa (il che farebbe pensare a un aumento della stabilità), può accadere che il sistema da controllare, di per sé stabile, diventi instabile in seguito al collegamento in retroazione. Questo perché in certe condizioni, e a certe frequenze, la retroazione negativa può diventare positiva (con guadagno di anello maggiore di 1) e produrre instabilità.

Le reti correttrici, o reti di compensazione, sono circuiti **a resistenza e capacità (passivi o attivi)**.

Possono essere reti **integratrici**, reti **derivatrici**, reti di **anticipo e/o ritardo**.

Nel sistema di controllo possono essere inseriti nella linea di azione diretta (cioè nel controllore, “in serie” al sistema da controllare) oppure come blocco di reazione di uno dei blocchi del sistema di controllo.

CONTROLLORE: SPECIFICHE

Torniamo ora sulle specifiche del sistema di controllo, approfondendone alcuni aspetti.

STABILITA'

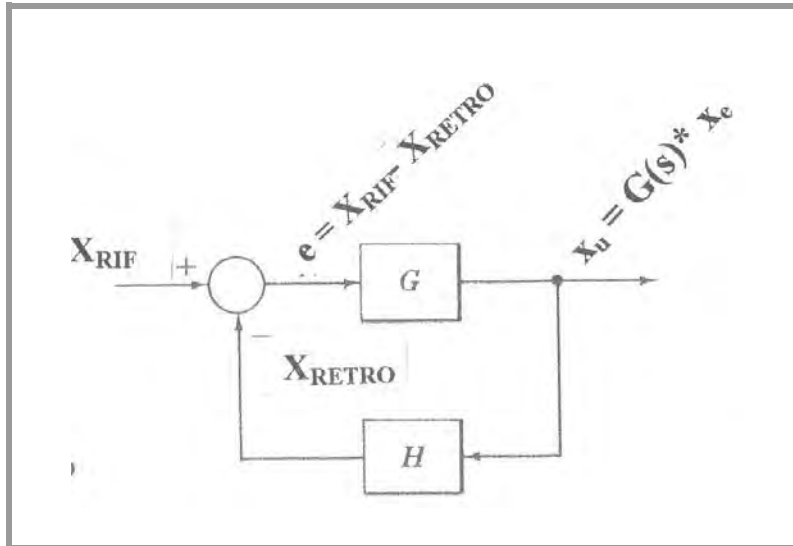
Sono definite una stabilità “semplice” e una stabilità “asintotica”.

Un sistema è asintoticamente stabile se la sua risposta, a un impulso di ampiezza limitata, tende ad annullarsi (per t che tende all'infinito). E' invece semplicemente stabile se la risposta non si annulla, ma rimane limitata in ampiezza.

Inoltre bisogna distinguere fra stabilità del sistema controllato (campo) “in se” e stabilità del sistema complessivo in catena chiusa.

La stabilità di un sistema viene definita facendo riferimento alla risposta impulsiva del sistema, ossia alla risposta all'impulso di Dirac e può essere analizzata studiando la funzione di trasferimento $G(s)$ del sistema.

ERRORE A REGIME (O ERRORE PERMANENTE O ERRORE STATICO), SEGNALE DI RIFERIMENTO E SEGNALE DI RETROAZIONE



Nella definizione di errore non viene direttamente coinvolta la vera e propria uscita del sistema.

Per definizione l'errore a regime è lo **scostamento** (per t che tende **all'infinito**²) fra il segnale X_{RIF} di **riferimento** o di ingresso e il segnale X_{retro} di **retroazione** o di misura:

$$e(t) = \lim_{t \rightarrow \infty} [x_{RIF} - x_{RETRO}]$$

Se l'errore è nullo, allora (per $t \rightarrow \infty$) si ha:

$$[x_{RIF}]_{t \rightarrow \infty} = [x_{RETRO}]_{t \rightarrow \infty}$$

² cioè dopo l'assestamento del sistema, a transitorio esaurito

Invece se l'errore a regime è costante e vale K_e , deve essere

$$K_e = \lim_{t \rightarrow \infty} [x_{RIF} - x_{RETRO}]$$

$$[x_{RIF}]_{t \rightarrow \infty} - [x_{RETRO}]_{t \rightarrow \infty} = K_e$$

$$[x_{RIF}]_{t \rightarrow \infty} = [x_{RETRO}]_{t \rightarrow \infty} + K_e$$

ESPRESSIONE DELL'ERRORE A REGIME NELLA VARIABILE "S" DI LAPLACE

L'espressione dell'errore nel tempo è:

$$e(t) = x_{RIF}(t) - x_{RETRO}(t)$$

Facendone la L-trasformata si ha:

$$e(s) = x_{RIF}(s) - x_{RETRO}(s)$$

Essendo poi:

$$x_{RETRO}(s) = H_0 \cdot X_U(s) = H_0 \cdot G(s) \cdot e(s)$$

si può scrivere:

$$e(s) = x_{RIF}(s) - H_0 \cdot G(s) \cdot e(s)$$

da cui:

$$e(s) \cdot [1 + H_0 \cdot G(s)] = x_{RIF}(s)$$

e quindi:

$$e(s) = x_{RIF}(s) \cdot \frac{1}{1 + H_0 \cdot G(s)}$$

L'errore a regime, nel tempo, è definito come limite:

$$e_\infty = \lim_{t \rightarrow \infty} [e(t)]$$

Per il teorema sulle L-trasformate detto "del valore finale", si ha:

$$\lim_{t \rightarrow \infty} [e(t)] = \lim_{s \rightarrow 0} [e(s)] \cdot s$$

Per cui si può scrivere:

$$e_\infty = \lim_{s \rightarrow 0} \left[s \cdot x_{RIF}(s) \cdot \frac{1}{1 + H_0 \cdot G(s)} \right]$$

APPROFONDIMENTO

PERCHE' LA PRESENZA DI POLI NELL'ORIGINE IN $G(s)$ TENDE A RIDURRE L'ERRORE A REGIME

A regime, l'errore è definito, nella variabile “s”, come:

$$e_{\infty} = \lim_{s \rightarrow 0} \left[s \cdot x_{RIF}(s) \cdot \frac{1}{1 + H_0 G(s)} \right] \quad (*)$$

Dove $x_{RIF}(s)$ è formato, se il segnale di riferimento è almeno di grado 1, cioè se è almeno una rampa, da fattori $1/s$.

Se $G(s)$ ha poli nell'origine, cioè se contiene dei fattori $1/s$, tali fattori inseriti nella funzione

$$\frac{1}{1 + H_0 G(s)} \quad (**)$$

determinano la presenza di fattori “s” al numeratore della funzione stessa, che quindi si presenterà nella forma³:

$$\frac{s^n}{s^n + k \cdot H_0}$$

Insieme al fattore “s” già presente nell'espressione (*) di e_{∞} , questi fattori “s” presenti al numeratore della (**) condurrebbero (per s tendente a 0) ad azzerare l'errore e_{∞} . Bisogna tener presente però che i fattori $1/s$ eventualmente presenti in $x_{RIF}(s)$ possono “semplificarsi” con i fattori “s” del numeratore della (**), o con alcuni di essi, per cui la loro azione sarebbe vanificata.

Se quindi, dopo la semplificazione, fra i fattori “s” del numeratore di e_{∞} e i fattori “s” presenti al denominatore, non rimane alcun fattore “s”, l'errore risulta finito.

Se invece, dopo la semplificazione, rimangono fattori s solo al numeratore di e_{∞} , l'errore, essendo definito per t tendente a zero, risulta nullo.

Se infine, dopo la semplificazione, rimangono fattori s solo al denominatore di e_{∞} l'errore risulta infinito.

$$\frac{1}{1 + k \cdot H_0 \cdot \frac{1}{s^n}} \rightarrow \frac{1}{\frac{s^n + k \cdot H_0}{s^n}} \rightarrow \frac{s^n}{s^n + k \cdot H_0}$$

ERRORE A REGIME, TIPO DI SISTEMA E GRADO DEL SEGNALE DI INGRESSO

Esiste, come si è detto, una relazione fra l'entità dell'errore a regime, il "tipo" del sistema W ad anello chiuso e il "grado" del segnale applicato al suo ingresso.

Prima di esplicitare questa relazione richiamiamo i concetti di "tipo" di sistema e "grado" del segnale.

COSA SI INTENDE PER TIPO DI SISTEMA

Un sistema si dice di tipo N se la sua fdt in s ha N poli nell'origine.

Dire che la fdt di un sistema ha un polo nell'origine significa dire che il suo denominatore si annulla per $s=0$ e che quindi, se viene scomposto in fattori, presenta un fattore s^N .

Per esempio la fdt

$$W(s) = k / [s * (1 + \tau s)]$$

presenta un polo s semplice (ossia elevato alla prima potenza) per cui il sistema è di tipo 1.

Invece la fdt

$$W(s) = k / [s^2 * (1 + \tau s)]$$

presenta un polo doppio nell'origine per cui il sistema è di tipo 2.

COSA SI INTENDE PER GRADO DI UN SEGNALE DI INGRESSO

Il grado di un segnale $x(t)$ è l'esponente che la variabile tempo (t) ha nell'espressione di $x(t)$.

Per esempio un ingresso a parabola $x(t) = kt^2$ è di grado 2.

Un ingresso a rampa $x(t) = kt = k t^1$ è di grado 1.

Il gradino invece è di grado 0 in quanto la variabile t non compare nella sua espressione (o meglio compare come t^0).

Il gradino infatti si esprime come

$$X(t) = \begin{cases} 0 & \text{per } t < 0 \\ 1 & \text{per } t > 0 \end{cases}$$

ENTITA' DELL'ERRORE A REGIME IN FUNZIONE DEL TIPO DI SISTEMA E DEL GRADO DELL'INGRESSO

L'errore a regime permanente nella risposta di un sistema di tipo N_T a un ingresso di grado N_G ha un valore:

NULLO se il tipo del sistema è maggiore del grado dell'ingresso

FINITO se il tipo del sistema è uguale al grado dell'ingresso

INFINITO se il tipo del sistema è minore del grado dell'ingresso

Pertanto:

se il sistema è di tipo ZERO, l'errore non è mai nullo, ma è finito

nel caso di ingresso a gradino, cioè di ingresso di grado zero, mentre è infinito nel caso di ingressi di grado maggiore (rampa e parabola);

se il sistema è di tipo UNO, l'errore è nullo nel caso di ingresso a gradino, cioè di ingresso di grado zero (grado<tipo), è finito nel caso di ingresso a rampa, cioè di ingresso di grado 1 (grado=tipo) mentre è infinito nel caso di ingresso a parabola (grado>tipo);

se il sistema è di tipo DUE, l'errore non è mai nullo, ma è finito nel caso di ingresso a parabola, cioè di ingresso di grado 2 (grado=tipo), mentre è nullo nel caso di ingressi a rampa e a gradino (grado<tipo).

TABELLA DELL'ERRORE A REGIME

Tipo di sistema	Ingresso		
	gradino	rampa	parabola
0	$\frac{1}{1 + H_0 G_{st}}$	∞	∞
1	0	$\frac{1}{H_0 G_{st}}$	∞
2	0	0	$\frac{1}{H_0 G_{st}}$

NOTA: per “ G_{st} ” si intende “ K_{BODE} ”

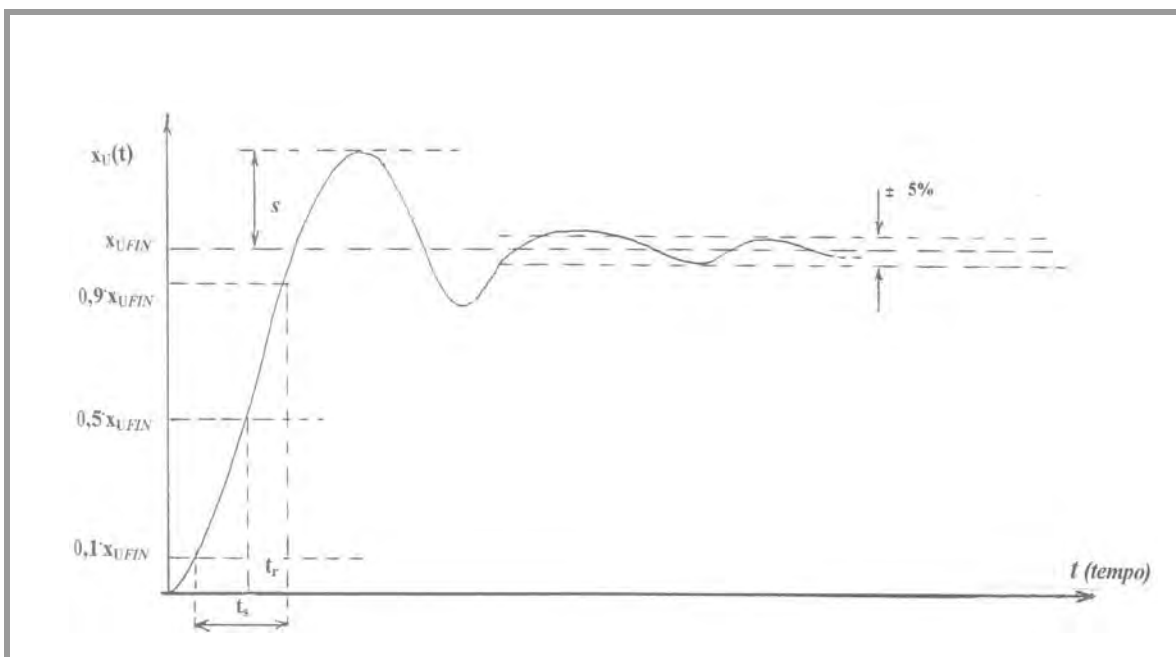
SPECIFICHE NEL TEMPO SULLA RAPIDITA' O VELOCITA' DI RISPOSTA DI UN SISTEMA CONTROLLATO (PARAMETRI DELLA RISPOSTA AL GRADINO)

Molto spesso i sistemi retroazionati (come i sistemi di controllo che ci interessano) hanno, anche se sono di ordine elevato, una risposta al gradino come quella dei sistemi del secondo ordine, cioè dei sistemi con funzione di trasferimento del tipo:

$$G(s) = k / [(s-p_1)(s-p_2)] = \text{per esempio} = 90 / (s^2 + 6s + 18)$$

E, in particolare⁴, hanno la risposta tipica dei sistemi del 2° ordine caratterizzati da una coppia di poli complessi coniugati e da un coefficiente di smorzamento ζ minore di 0,707.

La risposta al gradino in questione ha l'andamento temporale in figura:



⁴ a causa della presenza di una coppia di poli dominanti

I parametri caratteristici della risposta al gradino (o risposta **indiciale**)

di un sistema **in retroazione** evidenziati nel grafico sono:

t_r = TEMPO di RITARDO = tempo impiegato dalla risposta al gradino per raggiungere il 50% del valore finale (valore a regime X_{uFIN})

t_s = TEMPO di SALITA = tempo necessario al sistema per portare la risposta al gradino dal 10% del valore finale al 90% del valore finale X_{uFIN}

s = SOVRAELONGAZIONE = massimo scostamento (spesso espresso come % di X_{uFIN}) del valore istantaneo della risposta rispetto al valore a regime X_{uFIN}

t_{ass} = TEMPO di ASSESTAMENTO = SETTLING TIME = tempo necessario perché la risposta del sistema non si superi più un valore di scostamento del 5% dal valore a regime della risposta X_{uFIN}

GUADAGNO = rapporto fra il valore X_{uFIN} di regime della risposta al gradino e l'ampiezza X_i del gradino di ingresso.

NOTA

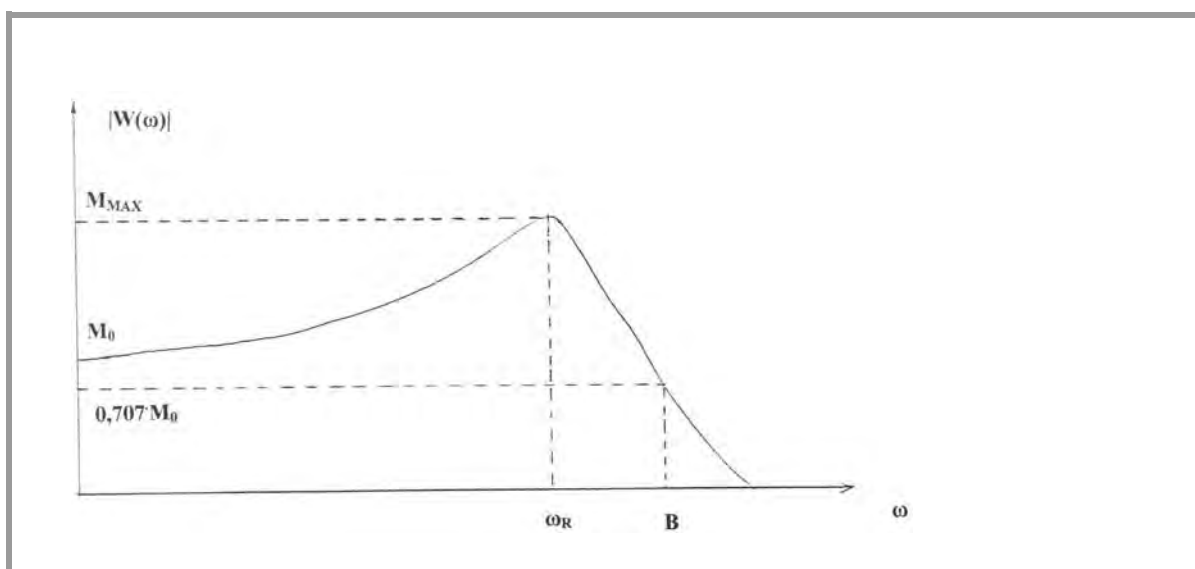
Lo smorzamento si può determinare confrontando la funzione di trasferimento, per esempio:

$$G(s) = 90 / (s^2 + 6s + 18)$$

con la sua forma standard.

SPECIFICHE IN FREQUENZA SULLA RAPIDITA' O VELOCITA' DI RISPOSTA DI UN SISTEMA CONTROLLATO (PARAMETRI DELLA RISPOSTA ARMONICA O RISPOSTA IN FREQUENZA)

Nella figura seguente è rappresentato il modulo della risposta armonica (diagramma di Bode in modulo) di un sistema in retroazione. Osserviamo che la presenza, anche in sistemi di ordine superiore al secondo, di poli dominanti (cioè di poli vicini all'asse immaginario) fa sì che il diagramma di Bode della risposta in frequenza di un sistema in retroazione abbia grossomodo sempre la forma del grafico che segue



I parametri tipici per l'imposizione delle specifiche sono:

M_0 = MODULO della risposta armonica IN CONTINUA cioè alla frequenza ZERO (i valori tipici sono compresi fra 1,1 e 1,3)

M_{MAX} = Valore MASSIMO del MODULO della risposta armonica

ω_R = PULSAZIONE di RISONANZA = pulsazione alla quale il
MODULO della risposta in frequenza assume il valore massimo cioè
per la quale si ha $|W(\omega_R)| = M_{MAX}$

B = BANDA PASSANTE = intervallo di frequenze compreso fra il valore
 $f=0$ e il valore di frequenza in corrispondenza del quale il modulo
della risposta armonica vale $0,707 \cdot M_0$, cioè 0,707 volte il valore in
continua.

M_R = MODULO DI RISONANZA = M_{MAX} / M_0 =
valore massimo / val. in continua

RELAZIONI SPERIMENTALI FRA SPECIFICHE NEL TEMPO E SPECIFICHE IN FREQUENZA

$$t_s = \text{cost} / B$$

dove:

t_s = tempo di salita o di risposta

B = banda passante del sistema

cost = costante di valore compreso fra 0,3 e 0,45

Significato: quanto più alte sono le frequenze alle quali il sistema riesce a rispondere (cioè quanto più larga è la banda) tanto più breve sarà il tempo di risposta.

METODI CLASSICI DI SINTESI DEL CONTROLLORE

Si attuano di solito nel dominio della frequenza **e quindi prevedono che le specifiche formulate nel tempo siano convertite in frequenza mediante le apposite relazioni sperimentali.**

Si basano inoltre sull'ipotesi che il **comportamento del sistema ad anello chiuso** sia determinato esclusivamente da una **coppia di poli dominanti** (dominanti sono i poli più vicini all'asse immaginario).

La sintesi avviene **per tentativi**: prima si aggiusta il guadagno del sistema, poi si soddisfano le specifiche di stabilità e prontezza mediante **reti correttrici**.

Le reti corretttrici sono reti che, aggiunte al sistema da controllare, **modificano la funzione di trasferimento (e quindi la risposta armonica) del sistema complessivo** in modo da poter soddisfare le specifiche di **stabilità e prontezza**.

Esempi di reti corretttrici sono:

la rete ritardatrice (phase lag)

la rete anticipatrice (phase lead)

la rete a sella o rete ad anticipo e ritardo (phase lead-lag).

STABILITÀ DEL SISTEMA DA CONTROLLARE E DEL SISTEMA IN CATENA CHIUSA

La funzione di trasferimento in s di un sistema è un modello matematico importante anche perché da essa possiamo trarre informazioni sulla stabilità del sistema, cioè sull'attitudine del sistema a ritornare nello stato originario dopo l'applicazione, in ingresso, di una sollecitazione di tipo impulsivo.

Ricordiamo che i poli della funzione di trasferimento sono i valori di s che annullano il denominatore della $G(s)$ e in corrispondenza dei quali la $G(s)$ tende all'infinito.

Gli zeri della funzione di trasferimento sono i valori di s che annullano il numeratore della $G(s)$ e in corrispondenza dei quali la $G(s)$ assume valore zero.

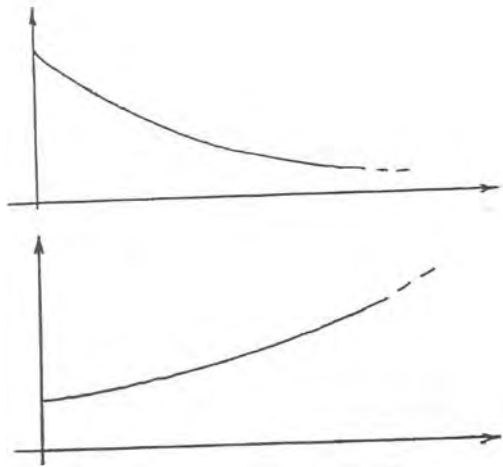
Per i sistemi che prendiamo in considerazione (sistemi continui, lineari, stazionari, a costanti concentrate, cioè con funzione di trasferimento razionale fratta) la stabilità è legata alla posizione dei poli della funzione di trasferimento $G(s)$.

Esiste un legame fra i poli della fdt $G(s)$ di un sistema e l'andamento temporale delle risposte del sistema alle sollecitazioni.

La presenza di un polo p della fdt indica la presenza di un termine esponenziale e^{pt} (esprimibile anche come $e^{a+j\omega t}$) nella risposta all'impulso del sistema.

Se il polo ha una parte reale negativa l'esponenziale è decrescente, quindi non diverge e non determina instabilità.

Se invece il polo ha una parte reale positiva l'esponenziale è crescente, quindi diverge e rappresenta un fattore di instabilità.



*In alto: termine della
risposta temporale
all'impulso dovuto a un:*

***polo semplice reale
negativo***

$e^{\alpha.t}$ con α minore di zero

Esempio:

$$\frac{1}{s+3} = \frac{1}{s-(-3)} \rightarrow \text{polo} = -3 \rightarrow L^{-1}\left(\frac{1}{s+3}\right) = e^{-3.t}$$

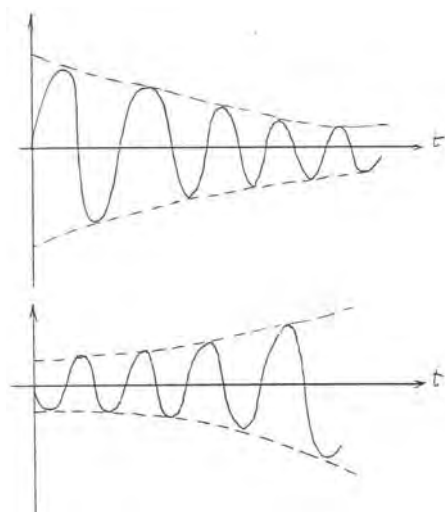
*In basso: termine della
risposta temporale
all'impulso dovuto a un:*

polo semplice reale positivo

$e^{\alpha.t}$ con α maggiore di zero

Esempio:

$$\frac{1}{s-3} = \frac{1}{s-(+3)} \rightarrow \text{polo} = +3 \rightarrow L^{-1}\left(\frac{1}{s-3}\right) = e^{+3.t}$$



In alto: termine della risposta temporale all'impulso dovuto a una coppia di poli complessi coniugati a parte reale negativa

$$k \cdot e^{\alpha t} \cdot \sin(\omega t) \quad \text{con } \alpha \text{ minore di zero}$$

Esempio:

$$\frac{3}{s^2 + 4s + 13} = \frac{3}{[s - (-2 + j3)][s - (-2 - j3)]}$$

$$\rightarrow \text{poli: } -2 + j3; -2 - j3$$

$$L^{-1}\left(\frac{3}{[s - (-2 + j3)][s - (-2 - j3)]}\right) = k \cdot e^{-2t} \cdot \sin(3t)$$

In basso: termine della risposta temporale all'impulso dovuto a una coppia di poli complessi coniugati a parte reale positiva

$$k e^{\alpha t} \cdot \sin(\omega t) \quad \text{con } \alpha \text{ maggiore di zero}$$

Esempio:

$$\frac{3}{s^2 - 4s + 13} = \frac{3}{[s - (+2 + j3)][s - (+2 - j3)]}$$

$$\rightarrow \text{poli: } +2 + j3; +2 - j3$$

$$\rightarrow L^{-1}\left(\frac{3}{[s - (+2 + j3)][s - (+2 - j3)]}\right) = k \cdot e^{+2t} \cdot \sin(3t)$$

E' fondamentale quindi il ruolo dei poli nella stabilità del sistema $G(s)$.

Se il sistema $G(s)$ ha tutti i poli a parte reale negativa, allora $G(s)$ è stabile.

INFORMAZIONI SULLA STABILITA' DEL SISTEMA RETROAZIONATO ATTRAVERSO L'ANALISI DELLA FUNZIONE $\overline{G(\omega)} \cdot \overline{H(\omega)}$

Quando si progetta il sistema di controllo in catena chiusa $W(s)$ di un impianto $G(s)$ (sistema da controllare) ci interessa la stabilità del sistema di controllo $W(s)$, stabilità sulla quale si impongono delle specifiche (requisiti).

In particolare noi vogliamo avere la possibilità di agire sulla stabilità del sistema ad anello chiuso W avendo informazioni solo sul sistema in catena aperta G^*H .

Precisiamo che GH è il prodotto delle risposte dei blocchi G ed H quando l'anello di retroazione è aperto, questo prodotto viene detto “funzione guadagno di anello” o “**funzione di trasferimento ad anello aperto**”.

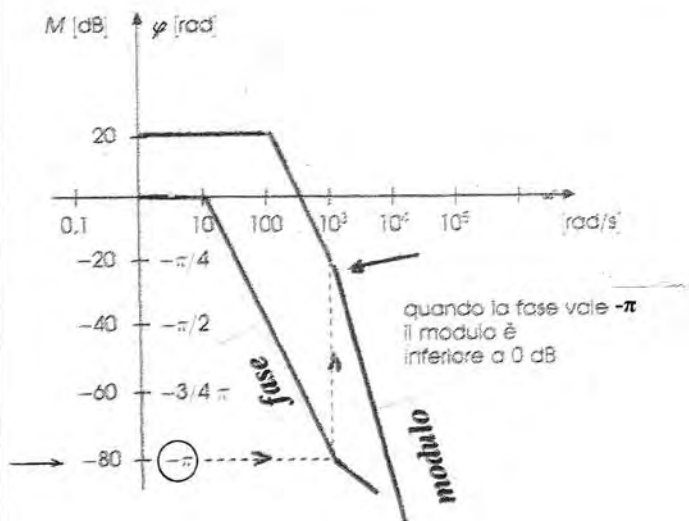
A questa esigenza risponde il **Criterio di BODE**⁵ che ci permette di stabilire **se un sistema retroazionato $W = G / (1 + G^*H)$ sia stabile o meno, conoscendo soltanto la risposta armonica del sistema ad anello aperto G^*H .**

ENUNCIATO DEL CRITERIO DI BODE

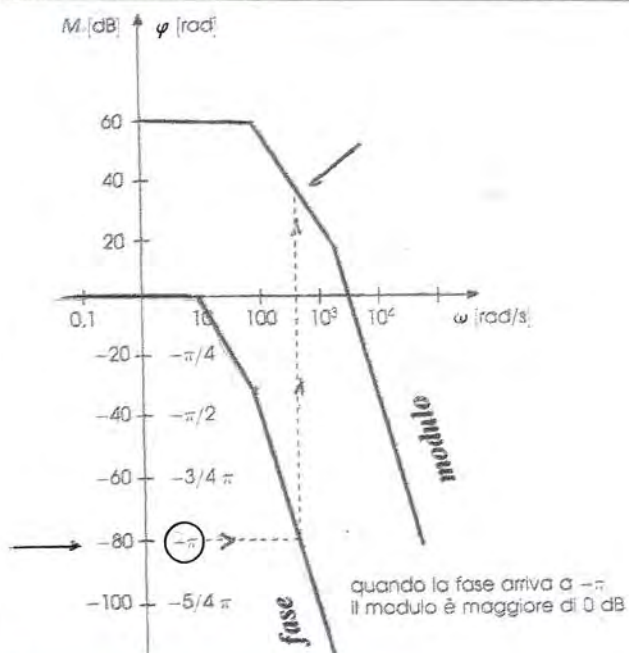
Il sistema in catena chiusa $W(\omega)$, cioè il sistema retroazionato, è instabile se esiste un valore di pulsazione ω in corrispondenza del quale la funzione di trasferimento ad anello aperto $G(\omega)H(\omega)$ ha fase uguale a $\pm \pi$ e modulo maggiore di 1

Di conseguenza, affinché un sistema retroazionato sia stabile, non deve esistere nessun valore di pulsazione per il quale la funzione di trasferimento ad anello aperto abbia simultaneamente fase uguale a $\pm \pi$ e modulo maggiore di 1.

⁵ Analogo al Criterio di Nyquist



Sistema di controllo stabile



Sistema di controllo instabile

MARGINI DI AMPIEZZA E DI FASE

Tra le specifiche più spesso richieste sulla stabilità di un sistema di controllo in catena chiusa vi sono il margine di ampiezza M_a e il margine di fase M_f che danno una misura di quanto il sistema sia stabile, ossia di quanto sia lontano dalla instabilità.

I valori tipicamente richiesti per M_a sono compresi *tra 12dB e 16dB*.

MARGINE DI AMPIEZZA O DI GUADAGNO M_a

Indichiamo con ω_π (e lo chiamiamo “pulsazione di inversione di fase”) il

valore di pulsazione in corrispondenza del quale la fase di $\overline{G(\omega)} \cdot \overline{H(\omega)}$ vale $+\pi$ o $-\pi$.

Il margine di guadagno $M_a|_{dB}$ di un sistema W , in catena chiusa, è il modulo, espresso in dB, della funzione di trasferimento di anello

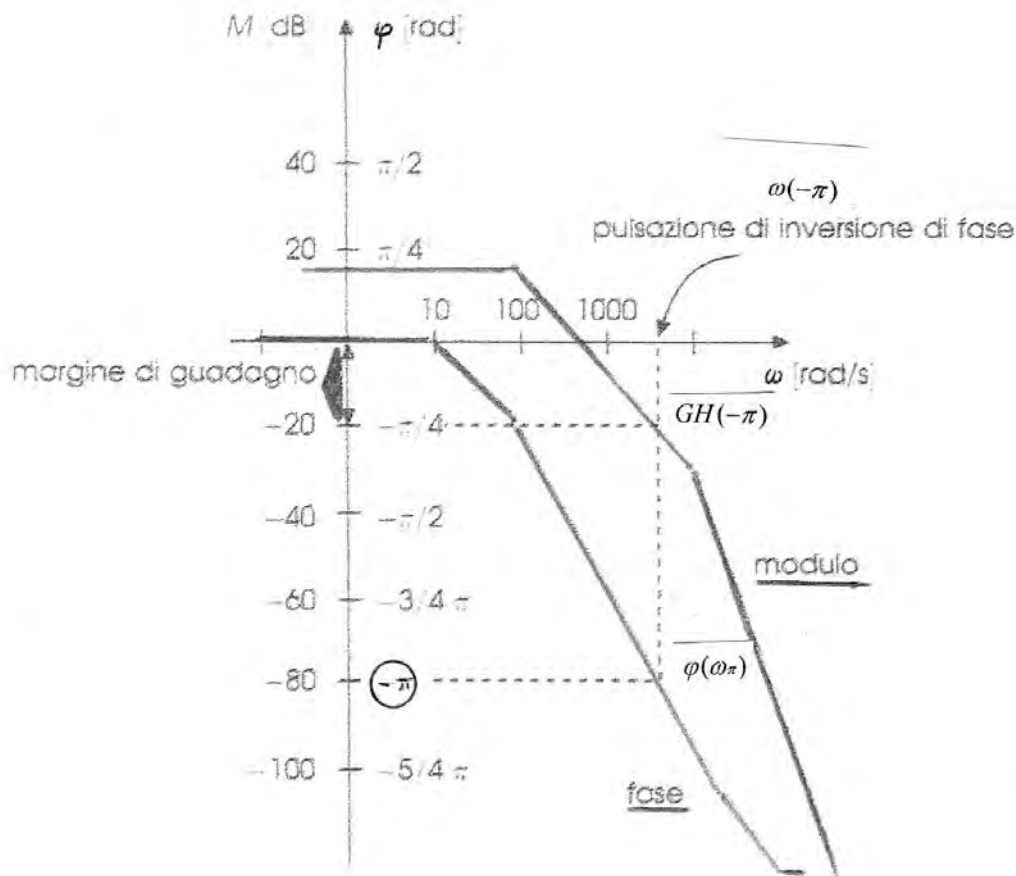
$\overline{G(\omega)} \cdot \overline{H(\omega)}$ in corrispondenza della **pulsazione ω_π alla quale la fase del sistema vale $+\pi$ o $-\pi$, cioè in corrispondenza della pulsazione di inversione di fase**. Il margine di fase va considerato convenzionalmente sempre positivo⁶.

⁶ Notiamo che il margine di fase, essendo definito per sistemi stabili, è rappresentabile come un segmento che si trova sempre al di sotto dell'asse orizzontale. Per cui, in dB, è sempre negativo. Se infatti non lo fosse, questo significherebbe

che, in corrispondenza della pulsazione ω_π in cui la fase vale $-\pi$, il modulo della $\overline{G(\omega)} \cdot \overline{H(\omega)}$ sarebbe maggiore di 1 (maggiore di zero se espresso in dB) e quindi che il sistema sarebbe instabile. Comunque noi, dal momento che definiamo il margine di guadagno solo per sistemi stabili, possiamo considerarlo positivo senza incorrere in equivoci

Per l'individuazione del margine di ampiezza M_a sui diagramma di Bode si procede in questo modo:

- si tracciano, uno sotto l'altro i diagrammi di Bode del modulo (in dB) e della fase della funzione guadagno di anello GH
- sul diagramma della fase si individua la pulsazione ω_π per la quale la fase vale $-\pi$, e quindi la fase $\varphi(\omega_\pi)$
- a partire da $\varphi(\omega_\pi)$, si traccia una verticale verso il diagramma delle ampiezze (cioè dei moduli). Il punto del diagramma dei moduli corrispondente a ω_π è $|GH|_\pi$. Il segmento compreso fra il punto $|GH|_\pi$ e l'asse orizzontale è il margine di ampiezza M_a .



MARGINE DI FASE M_f

Indichiamo con ω_1 (e lo chiamiamo “pulsazione di attraversamento dell’asse orizzontale”) il valore di pulsazione in corrispondenza del quale si ha:

$$\left| \overline{G(\omega)} \cdot \overline{H(\omega)} \right|_{dB} = 0_{dB}$$

oppure, in maniera equivalente:

$$\left| \overline{G(\omega)} \cdot \overline{H(\omega)} \right| = 1$$

Per quanto precedentemente detto, si può affermare che **il margine di fase M_f**

è lo **scostamento della fase di $\overline{G(\omega)} \cdot \overline{H(\omega)}$** , in corrispondenza della pulsazione ω_1 di attraversamento, rispetto a $+\pi$ o $-\pi$.

Dette quindi:

- ϕ_G la fase di $\overline{G(\omega)}$
- ϕ_H la fase di $\overline{H(\omega)}$
- ϕ_{GH} la fase di $\overline{G(\omega)} \cdot \overline{H(\omega)}$,

il margine di fase può essere definito come:

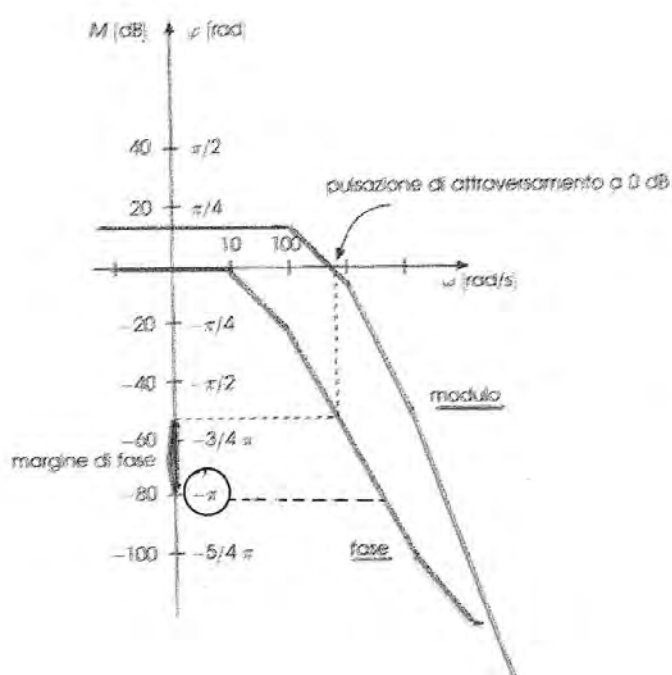
$$M_\phi = \phi_{GH}(\omega_1) - (-\pi)$$

Aggiungiamo infine che il margine di fase può anche essere definito come **l'ulteriore sfasamento, che in ω_1 , il sistema può sopportare prima di diventare instabile, cioè prima di scendere al di sotto della quota $-\pi$, o prima di salire al di sopra della quota $+\pi$.**

Per l'individuazione del margine di fase sui diagramma di Bode si procede in questo modo:

- si tracciano, uno sotto l'altro i diagrammi di Bode del modulo (in dB) e della fase della funzione guadagno di anello GH
- sul diagramma del modulo si individua la pulsazione ω_1 per la quale il modulo vale 1, ossia la pulsazione alla quale il diagramma del modulo (che è in dB) interseca l'asse orizzontale.
- a partire da questo punto di intersezione si abbassa una verticale fino a intercettare il sottostante diagramma della fase.

Il segmento verticale compreso fra il diagramma della fase e la retta a -180° rappresenta il margine di fase M_f . Se questo **segmento** si trova **al di sopra della retta a -180°** , allora il sistema è **stabile**.



RETI CORRETRRICI O DI COMPENSAZIONE

Le reti correttrici o reti di compensazione sono dispositivi, generalmente elettronici, che vengono introdotti in un sistema di controllo retroazionato per correggere la risposta del sistema alle sollecitazioni, ossia per soddisfare le specifiche di comportamento del sistema.

La rete correttrice fa parte del blocco regolatore o controllore⁷ del sistema.

La rete correttrice viene spesso introdotta per risolvere il conflitto fra precisione e stabilità del sistema: siccome la precisione è tanto migliore quanto maggiore è il numero di poli nell'origine della fdt dei blocchi (non facenti parte della linea di retroazione) e quanto più elevato è il valore della costante di guadagno della funzione di anello GH, ecco che per conseguire una buona precisione si compromette la stabilità. Si fa intervenire a questo punto la rete correttrice, che, con l'introduzione dei suoi poli e/o zeri nella fdt complessiva del sistema, può ripristinare la stabilità.

Anche l'esigenza di avere un determinato tempo di salita o di assestamento o una determinata banda (parametri questi relativi al comportamento in transitorio del sistema), può essere soddisfatta grazie a una rete correttrice.

Le reti correttrici fondamentali sono la rete integratrice (filtro RC passabasso), la rete derivatrice (filtro CR passaalto), la rete ritardatrice, la rete anticipatrice, la rete a ritardo e anticipo o rete "a sella".

Possono essere realizzate come circuiti passivi (con soli resistori e condensatori) o come circuiti attivi, con l'aggiunta di amplificatori operazionali (op.amp.).

L'azione svolta dalle reti correttrici è detta **compensazione**.

⁷ Il regolatore comprende generalmente un dispositivo che esegue la differenza fra il segnale di riferimento e quello di retroazione (generando un segnale di errore), una o più reti correttrici che elaborano il segnale di errore, un amplificatore di segnale, un amplificatore di potenza e un attuatore avente la funzione di agire sul sistema da controllare. I regolatori possono essere realizzati con dispositivi meccanici (molle e ammortizzatori), pneumatici, idraulici (strozzature e serbatoi) e, soprattutto, elettrici ed elettronici. Ci occuperemo solo di questi ultimi perché, avendo prestazioni migliori, tendono a soppiantare tutti gli altri.

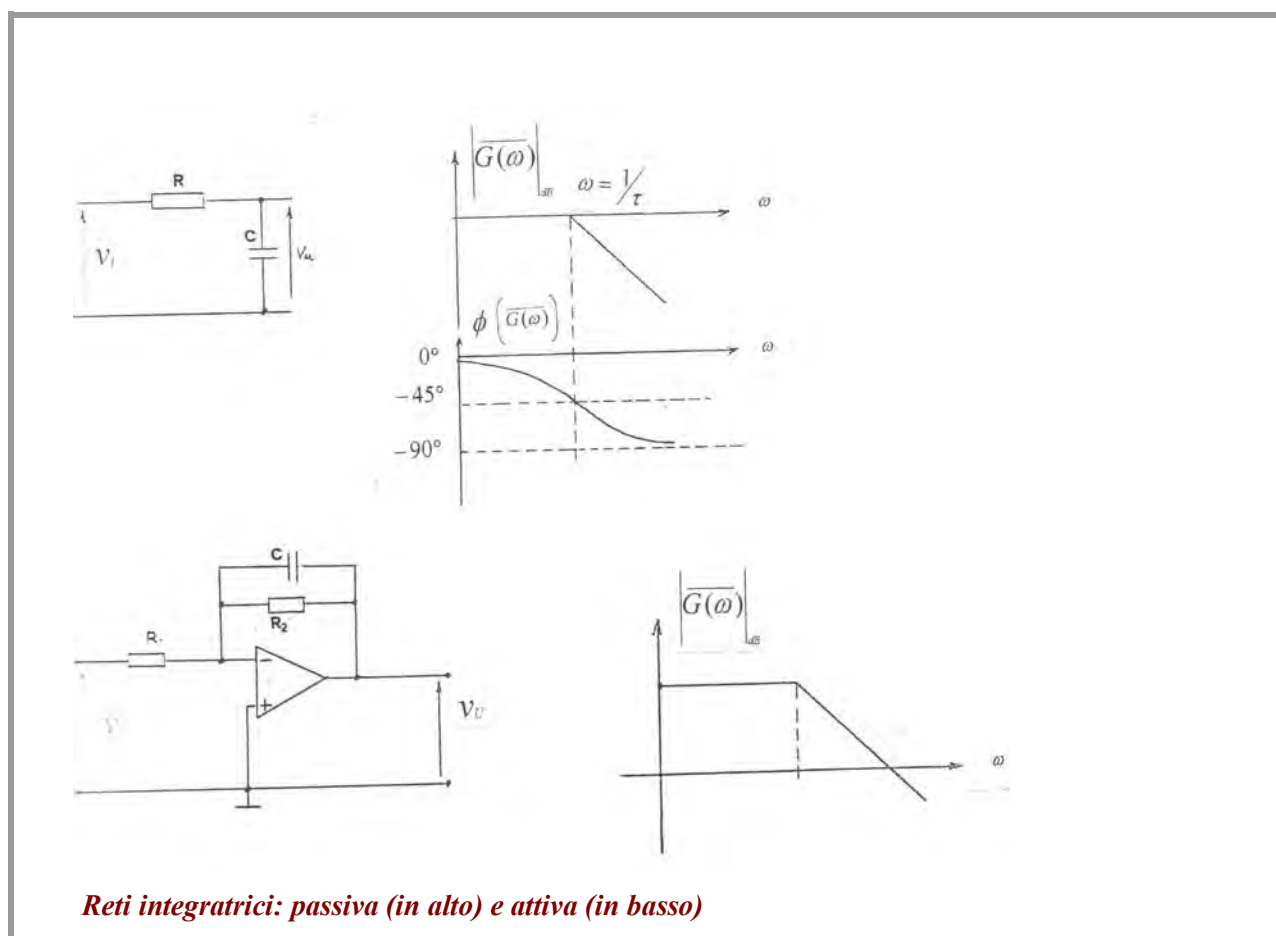
RETE INTEGRATRICE

Può essere realizzata con un circuito attivo o passivo.

Nel caso passivo è un semplice filtro RC passabasso con funzione di trasferimento:

$$G(s) = 1 / (1 + \tau s)$$

La presenza, nella fdt, del polo $p = -1/\tau$ ne consente l'utilizzazione come rete correttiva, operante la cosiddetta **compensazione con polo dominante**.



Scegliendo τ in modo che il polo della rete risulti molto inferiore al valore del primo polo della funzione guadagno di anello GH, si abbassa la frequenza di taglio superiore del sistema, **migliorandone la stabilità**.

Il polo della rete integratrice fa sì che il diagramma del modulo di GH intersechi l'asse orizzontale ad un valore di pulsazione minore di quello originario, valore in corrispondenza del quale la fase di GH si trova ancora al di sopra della linea a -180° .

RETE DERIVATRICE

Può essere realizzata con un circuito attivo o passivo.

Nel caso passivo è un semplice filtro CR passaalto con funzione di trasferimento:

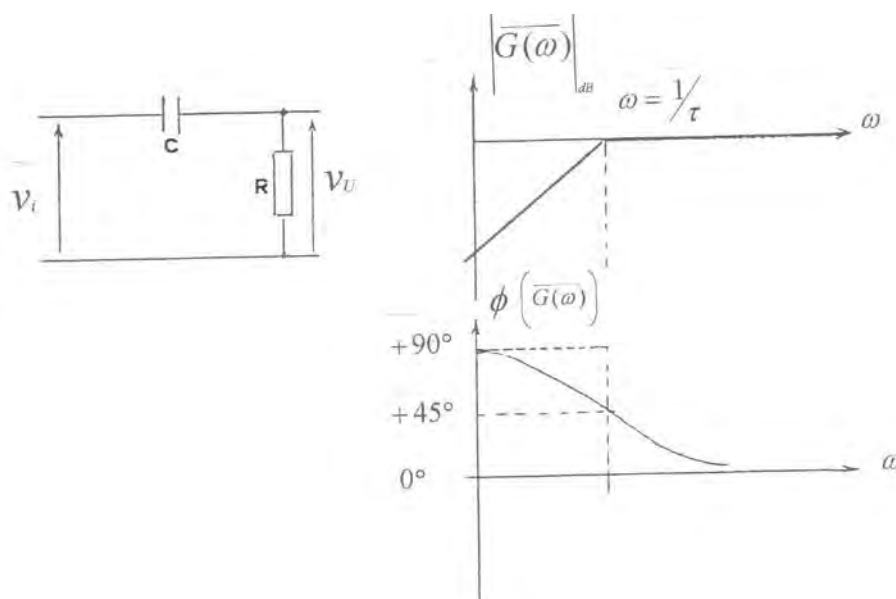
$$G(s) = \tau s / (1 + \tau)$$

con $\tau = RC$

La $G(s)$ è dotata di un **polo $p = -1/\tau$** e di uno **zero nell'origine**.

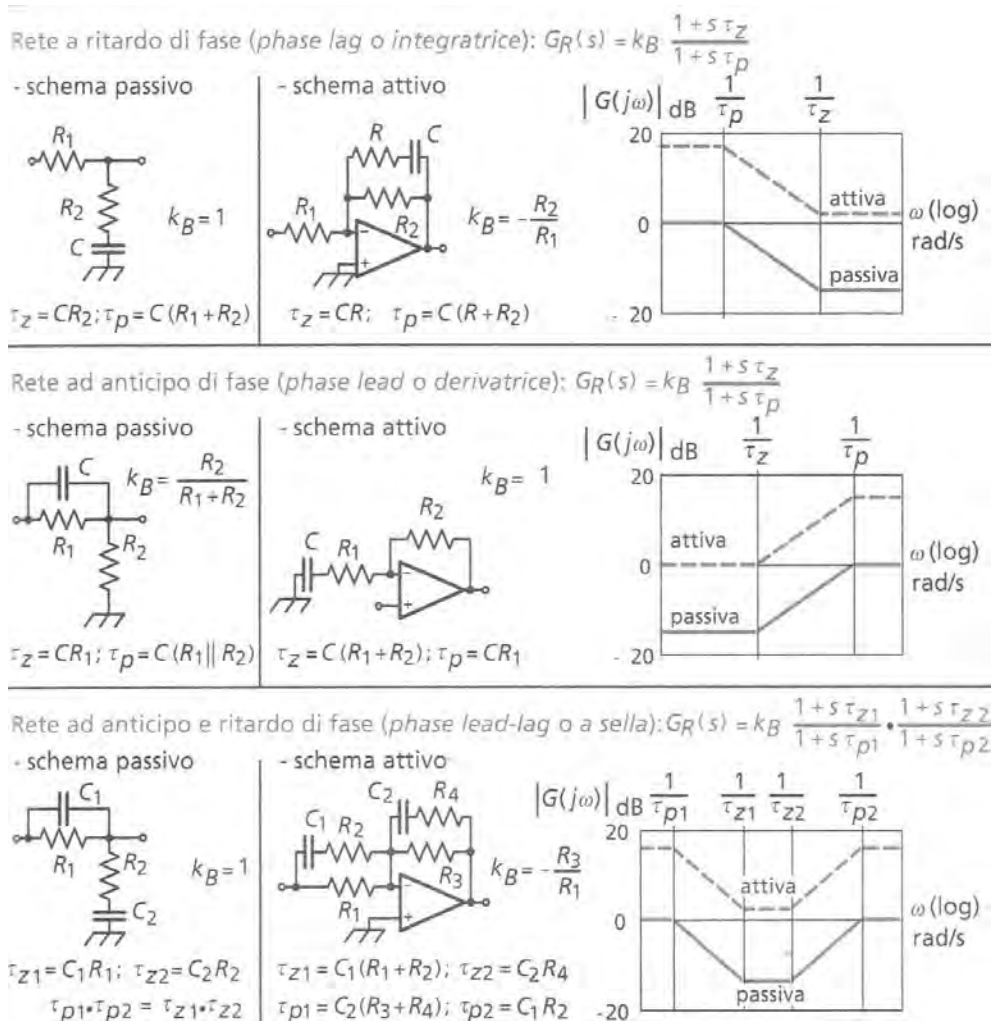
La rete derivatrice produce un anticipo di fase per tutte le pulsazioni finite ed è impiegata di rado per la stabilizzazione dei sistemi di controllo.

Questa rete non può essere utilizzata da sola in un sistema di controllo retroazionato perché blocca la componente continua del segnale, mentre è proprio in continua che il sistema di controllo deve avere un elevato guadagno di anello.



Rete derivatrice

ALTRE RETI CORRETRICI O DI COMPENSAZIONE



ALTRE RETI CORRETTRICI O DI COMPENSAZIONE

RETE A RITARDO DI FASE (PHASE LAG)

- **POSIZIONAMENTO POLI E ZERI**

La pulsazione dello **zero della rete corretttrice** va posizionata in corrispondenza del **polo reale del sistema**.

- La pulsazione del **polo della rete** va posizionata *una decade prima dello zero del sistema*

***Effetti:** migliora la stabilità, rallenta la risposta del sistema, rende la frequenza di taglio complessiva 10 volte più piccola, ottimizza le prestazioni a bassa frequenza, ritarda la fase fra polo e zero.*

RETE AD ANTICIPO DI FASE (PHASE LEAD)

- **POSIZIONAMENTO POLI E ZERI**

La pulsazione dello **zero della rete corretttrice** va posizionata in corrispondenza del **secondo polo reale del sistema**.

- La pulsazione del **polo della rete** va posizionata *una decade dopo della pulsazione dello zero del sistema*

***Effetti:** migliora la stabilità, accelera la risposta del sistema, rende la frequenza di taglio complessiva 10 volte più grande, ottimizza le prestazioni ad alta frequenza, riduce il guadagno statico, esercita azione attenuatrice, anticipa la fase fra zero e polo.*

RETE A SELLA O RETE AD ANTICIPO E RITARDO DI FASE (PHASE LEAD-LAG)

- **POSIZIONAMENTO POLI E ZERI**

La pulsazione dello **zero della rete corretttrice** va posizionata in corrispondenza del **primo polo reale del sistema**.

- La pulsazione del **polo della rete** va posizionata *una decade dopo della pulsazione del polo intermedio del sistema*

***Effetti:** accelera la risposta del sistema, attenua nell'intervallo compreso fra i due poli della rete, in un primo intervallo ritarda la fase, in un secondo la anticipa.*

CAPITOLO 41

COMPONENTI ELETTRONICI PER IL CONTROLLO DELLA POTENZA ELETTRICA

OBIETTIVI

- **Comprendere la specificità dell'SCR e degli altri tiristori rispetto agli altri componenti**
- **Comprendere il principio di funzionamento di SCR, TriAC e DiAC e le differenze fra tali componenti**
- **Conoscere e saper valutare i principali parametri dei tiristori**
- **Comprendere la tecnica del “controllo in fase”**
- **Conoscere e saper analizzare alcuni circuiti applicativi dei tiristori**

PREREQUISITI

- **La giunzione pn**
- **Diodi e transistor**

I COMPONENTI ELETTRONICI PER IL CONTROLLO DI POTENZA

I componenti per l'elettronica di potenza come l'SCR, il TriAC e altri hanno come utilizzazione principale il controllo della potenza erogata a un carico, per esempio a un motore elettrico. Il campo di applicazione più specifico è quello delle tensioni e correnti "altissime", anche se non è esclusa l'utilizzazione in ambiti differenti.

Questo controllo può essere di tipo on/off, cioè può consistere nelle sole funzioni di accensione e spegnimento, oppure può essere un controllo "continuo" nel senso che la potenza fornita al carico può essere regolata con continuità e non per salti.

Alcuni dei componenti che tratteremo consentono il passaggio della corrente in un solo verso, e sono perciò definiti "unidirezionali". Altri invece sono bidirezionali, nel senso che la corrente li può percorrere in entrambi i versi.

I principali componenti per l'elettronica di potenza sono:

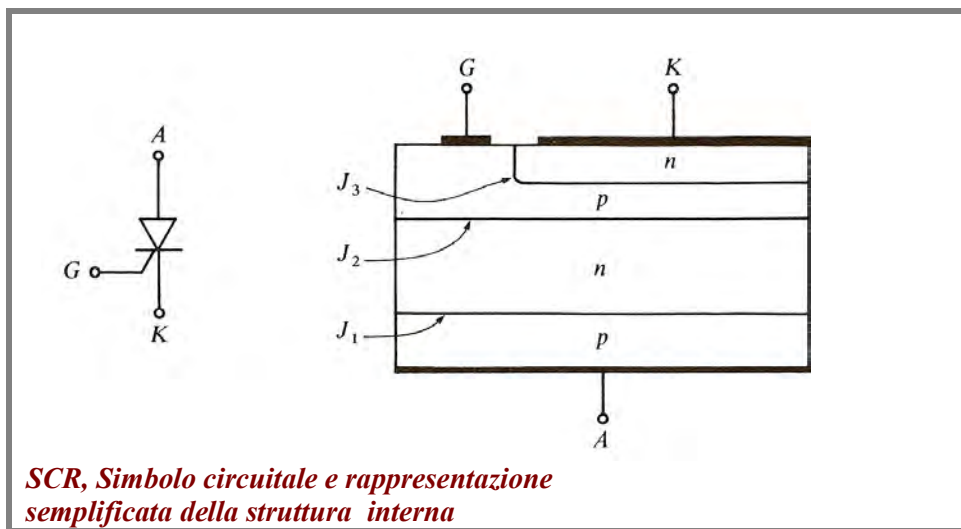
- 1) l' **SCR**, Raddrizzatore Controllato al Silicio o diodo controllato, che è un componente unidirezionale a **tre terminali** e che viene anche detto "**tiristore**" (spesso però con questo termine si indicano tutti i componenti per il controllo della potenza, come si farà anche in questo testo); l'SCR può **controllare correnti** di **alcune migliaia di ampère** e tensioni di **alcune migliaia di volt**;
- 2) il **TriAC** (TRIodo AC, triodo per alternata), che è un componente **a tre terminali** analogo all'SCR, ma **bidirezionale**;
- 3) il **DiAC** (Diodo AC, diodo per alternata), che è un componente bidirezionale a **due terminali**.

Tratteremo sinteticamente anche il **GTO** (Gate Turn Off, tiristore spegnibile dal gate) e l'**SCS** (Silicon Controlled Switch, interruttore controllato al silicio) che sono "evoluzioni" dell' SCR.

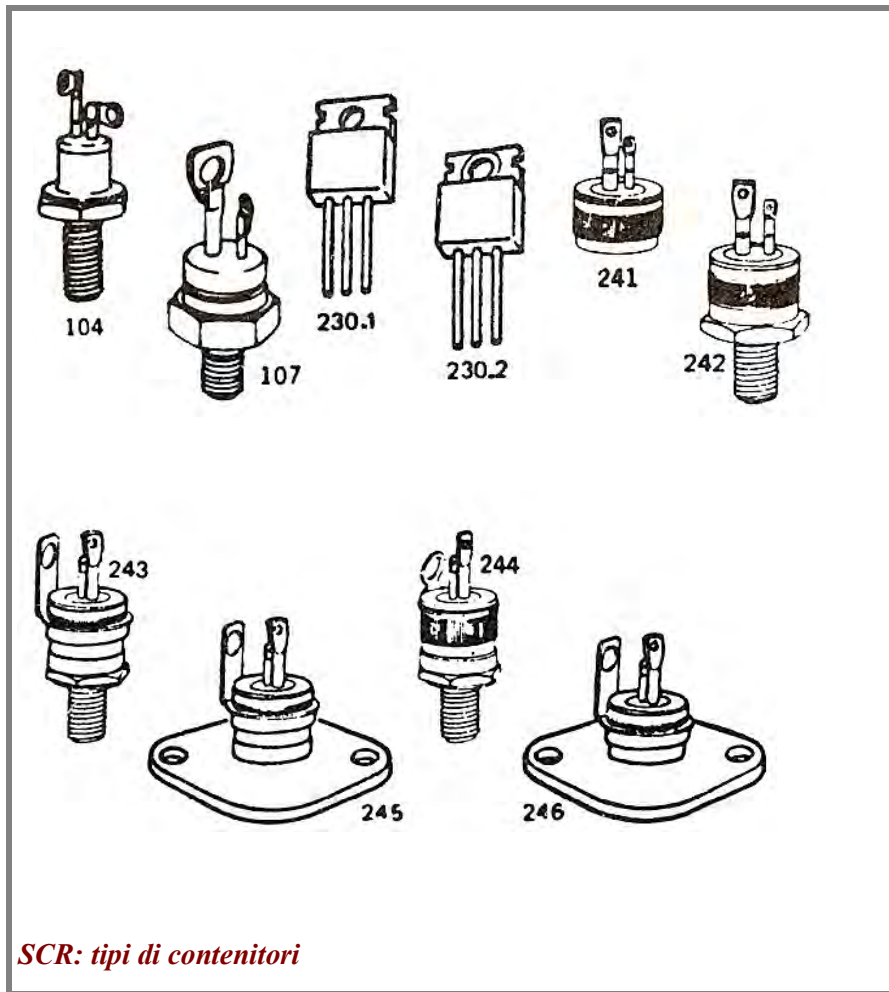
Questi componenti si sono rivelati particolarmente importanti nel settore **navale** e nel settore **ferroviario**, ove sono utilizzati per il **controllo dei motori elettrici** in alternata.

SCR (RADDRIZZATORE CONTROLLATO AL SILICIO O DIODO CONTROLLATO)

Vedremo ora, soltanto a grandi linee e in maniera molto schematica come è fatto e come funziona un SCR.



La struttura dell'SCR presenta tre giunzioni pn e tre terminali (anodo, catodo e gate). Anodo e catodo non sono intercambiabili (uno è collegato alla zona p, l'altro alla zona n), quindi il componente non è "simmetrico".



COME FUNZIONA L'SCR IN ASSENZA DI SEGNALI DI CONTROLLO AL TERMINALE DI GATE

Analizzeremo il comportamento dell' SCR in due situazioni: quando l'anodo è potenziale più basso del catodo ($V_{AK} < 0$) e quando l'anodo è, invece, a potenziale più alto del catodo ($V_{AK} > 0$)

SITUAZIONE $V_{AK} < 0$

Quando la tensione dell'anodo (terminale “positivo”) è minore di quella del catodo, sono inversamente polarizzate due giunzioni su tre (quelle “esterne”), mentre una sola giunzione è direttamente polarizzata.

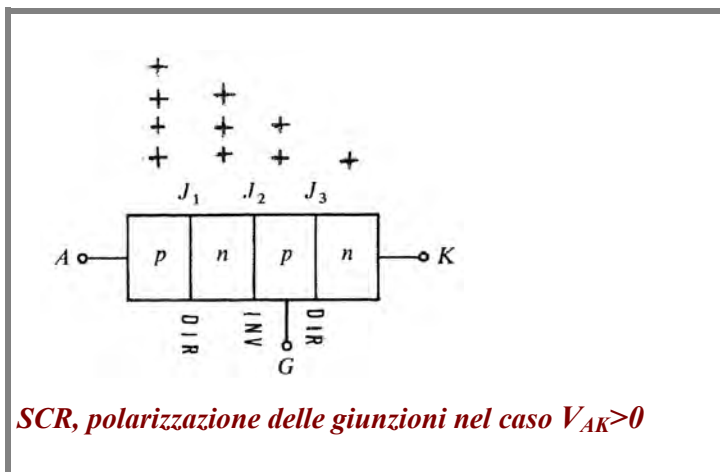
Per $V_{AK} < 0$, l' SCR si comporta come un diodo inversamente polarizzato.

Trascurando le piccole correnti inverse, possiamo dire che il componente non conduce fino a quando, per tensioni inverse di valore molto elevato, non avviene il breakdown, situazione nella quale l' SCR non può funzionare.



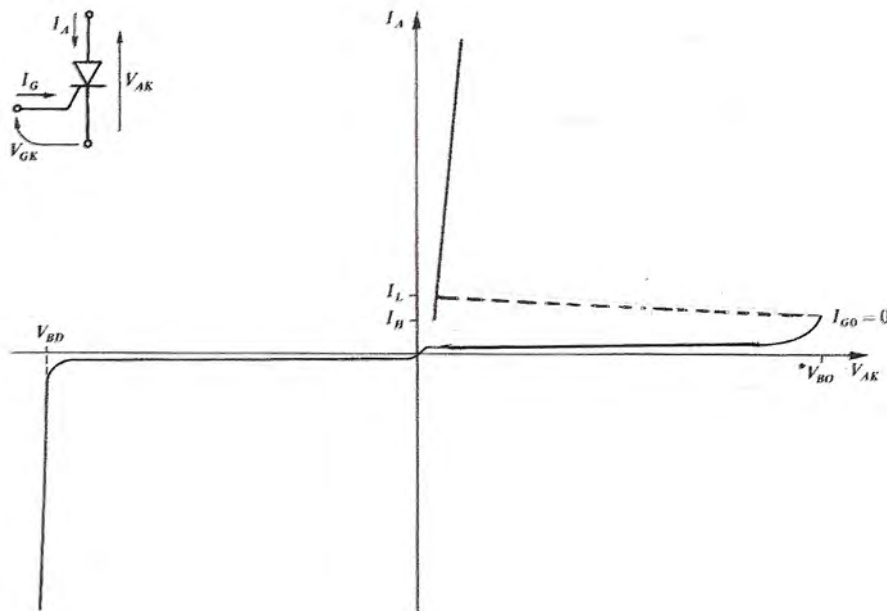
SITUAZIONE $V_{AK} > 0$

In questa situazione l' SCR ha un comportamento molto diverso da quello del "normale" diodo, per effetto del verificarsi di un fenomeno particolare noto come **rottura diretta** o **breakover**.



Il componente funziona nel modo di seguito descritto.

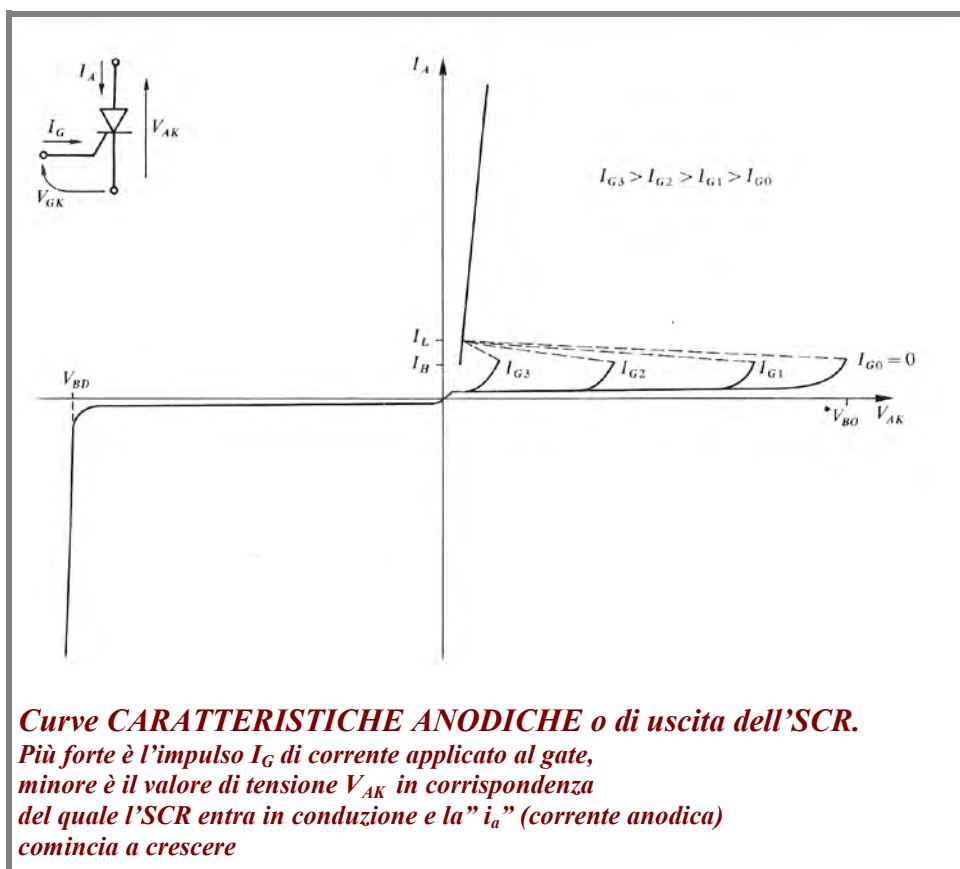
Se (a partire da 0V) facciamo crescere la tensione V_{AK} , la corrente cresce molto lentamente (in pratica l'SCR non conduce) fino a quando non viene raggiunto un ben preciso valore V_{BO} di tensione, specifico per ogni "sigla" di SCR, detto **tensione di breakover o di "rottura di retta"**. Il valore di V_{BO} è generalmente dell'ordine delle **centinaia di V**, ma può anche **superare i 1000V**. Appena viene raggiunta la tensione di breakover, fra anodo e catodo si sviluppa una rilevante corrente che innesca un **processo a valanga** per il quale **la barriera di potenziale dell'unica giunzione (J_2) in versamente polarizzata viene neutralizzata**. Di conseguenza la tensione V_{AK} fra anodo e catodo precipita al valore di circa 1,5 V (pari alla somma delle cadute di potenziale sulle due giunzioni in conduzione) e la corrente comincia a crescere come in un normale diodo. Si hanno cioè grandi aumenti di corrente per piccoli incrementi di tensione.

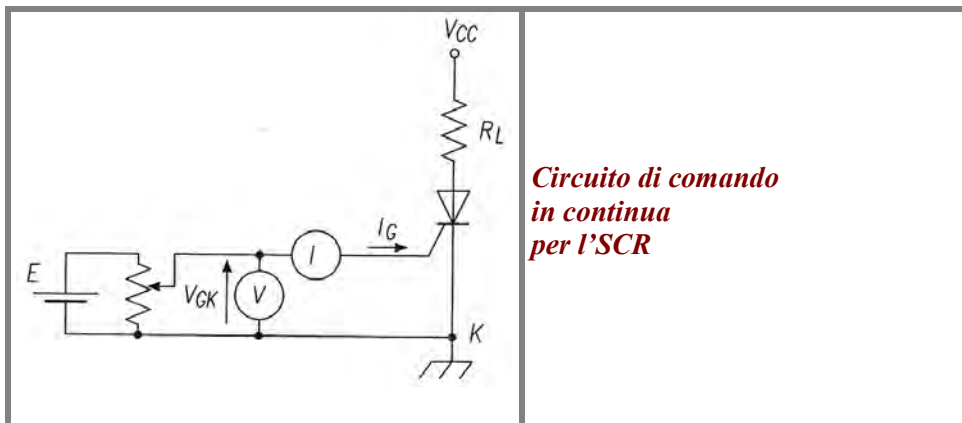


Il valore di corrente corrispondente alla tensione di **breakover**, cioè il valore in corrispondenza del quale **comincia la rapida crescita di corrente** è detto **corrente di latching o di aggan cio** e lo indichiamo con I_L . Il valore della corrente di latching può essere dell'ordine delle **decine di mA**.

COME FUNZIONA L'SCR IN PRESENZA DI SEGNALI DI CONTROLLO AL TERMINALE DI GATE

Ciò che abbiamo finora descritto è il funzionamento in assenza di correnti (e tensioni) di gate. Se noi **applichiamo al gate** un impulso di **corrente**, l'SCR entra in conduzione diretta a un valore di tensione V_{AK} minore di quello di breakover. Più forte è l'impulso di corrente (tensione) applicato al gate, minore è il valore di tensione V_{AK} in corrispondenza del quale l'SCR entra in conduzione. Esiste un valore della corrente (tensione), da applicare al gate, in corrispondenza del quale l'SCR entra subito in conduzione. Le ampiezze tipiche delle correnti di controllo o comando applicata al gate sono comprese fra **qualche decimo di mA** e **qualche decina di mA**. Il comportamento dell'SCR ora descritto è illustrato graficamente mediante curve chiamate “caratteristiche anodiche” o “caratteristiche di uscita” dell'SCR.





DISINNESCO O SPEGNIMENTO DELL'SCR

Nel funzionamento dei sistemi di comando o di controllo si presenta abitualmente la necessità di disinnescare (spegnere) l'SCR, cioè di farlo uscire dallo stato di conduzione nel quale sia venuto a trovarsi. Per disinnescare l'SCR bisogna fare in modo che la corrente anodica scenda al di sotto di un particolare valore di corrente, detto valore di **mantenimento** o di **holding** I_H . Il valore della corrente di holding, poco minore di quello della corrente di latching, può essere dell'ordine delle **decine di mA**.

In pratica il disinnescamento si ottiene mediante **l'inversione della polarità** fra anodo e catodo o mediante **l'apertura del circuito**. A proposito del disinnescamento di legga il paragrafo “circuiti di spegnimento dell'SCR”.

IL FOGLIO-DATI DELL'SCR

Osserviamo che, normalmente, nel foglio-dati dell'SCR, il costruttore non indica la tensione V_{BO} di breakover, ma alcuni valori di tensione poco minori di V_{BO} , che però non portano il componente al breakover.

Questi valori, in ordine decrescente di ampiezza, sono:

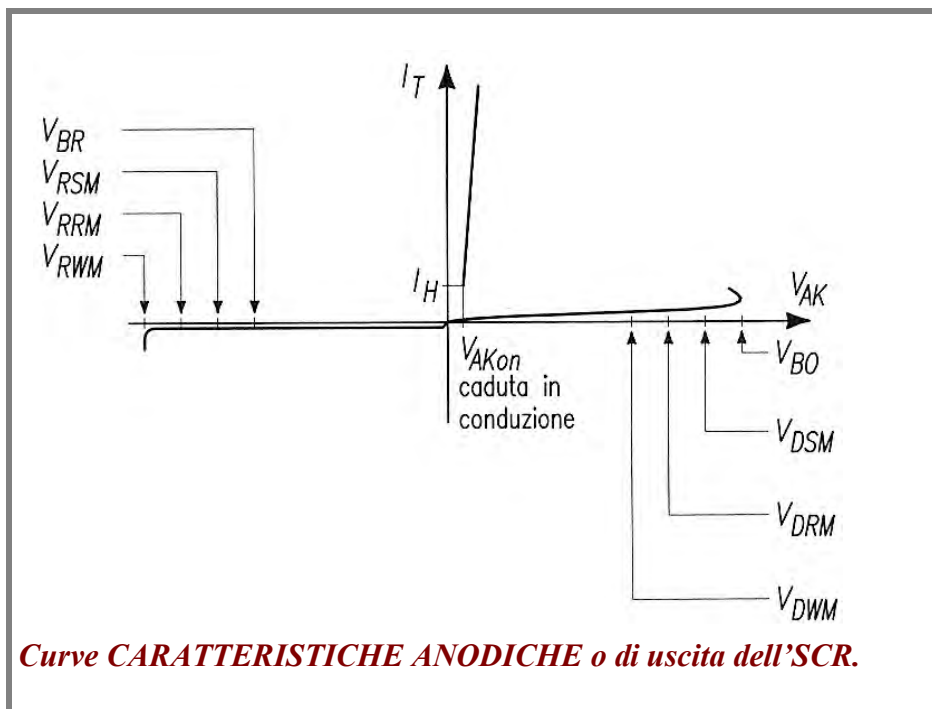
- V_{DSM} = tensione **limite** diretta di picco **non** ripetitiva (Direct **Single** Max Voltage)
- V_{DRM} = tensione **limite** diretta di picco **ripetitiva** (Direct **Ripetitive peak** Max Voltage)
- V_{DWM} = **massima** tensione diretta di picco **applicabile in maniera costante** (**massima tensione diretta di lavoro**, Direct **Working** Max Voltage)

Nel data-sheet compaiono anche alcuni valori caratteristici di tensioni inverse:

- V_{RSM} = tensione **limite** inversa di picco **non** ripetitiva (Reverse **Single** Max Voltage)
- V_{RRM} = tensione **limite** inversa di picco **ripetitiva** (Reverse **Ripetitive peak** Max Voltage)
- V_{DWM} = **massima** tensione inversa di picco **applicabile in maniera costante** (**massima tensione inversa di lavoro**, **Reverse Working** Max Voltage)

Ricordiamo infine i due diversi valori limite di **corrente** anodica forniti dal costruttore:

- I_{Tav} = massima corrente media del componente in conduzione = massimo valore di corrente che può scorrere con continuità nel componente (“av” sta per *average*, media)
- I_{Trms} = massima corrente efficace del componente in conduzione = massima entità del valore efficace della corrente che può scorrere con continuità nel componente (“rms” sta per *root mean square*, valore quadratico medio o valore efficace)



Può essere utile confrontare i parametri caratteristici di alcuni SCR, come nella tabella che segue.

sigla componente	I_H corrente di HOLDING	I_L corrente di LATCHING	I_G corrente di TRIGGER (da applicare al gate)	V_{GT} tensione di TRIGGER (da applicare fra gate e anodo)	V_{DDR} valore di picco (val. max) della tensione ripetitiva applicabile fra A e K
<i>MCR 100-3</i>		<i>10 mA</i>	<i>0,2 mA</i>	<i>0,8 V</i>	<i>100 V</i>
<i>MCR 100-4</i>		<i>10 mA</i>	<i>0,2 mA</i>	<i>0,8 V</i>	<i>200 V</i>
<i>MCR 100-6</i>		<i>10 mA</i>	<i>0,2 mA</i>	<i>0,8 V</i>	<i>400 V</i>
<i>MCR 100-8</i>		<i>10 mA</i>	<i>0,2 mA</i>	<i>0,8 V</i>	<i>600 V</i>
<i>CR 12 M</i>	<i>25 mA</i>	<i>42 mA</i>	<i>35 mA</i>	<i>1,2 V</i>	<i>(50÷1300) V</i>
<i>TIC 116</i> <i>TIC 126</i>	<i>(40÷70) mA</i>		<i>(5÷20) mA</i>	<i>(0,2÷2,5) V</i>	<i>(50÷800) V</i>

L'SCR IN CONTINUA E L'SCR IN ALTERNATA

L'SCR può lavorare sia in circuiti alimentati in continua, sia in circuiti in alternata. Bisogna però aver presente che:

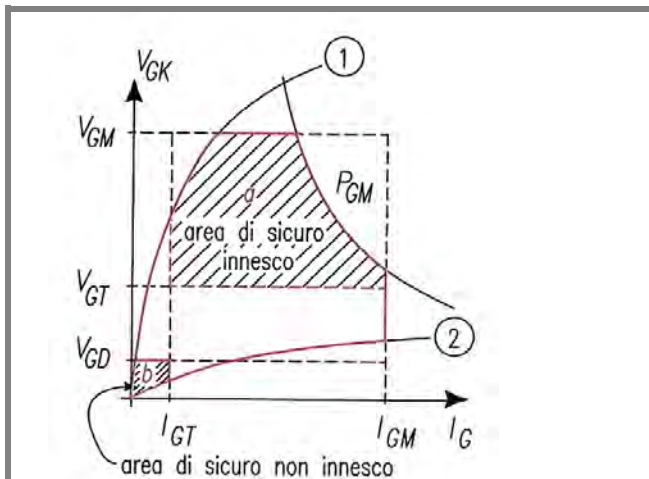
- Nei circuiti **in continua** gli impulsi di **comando** (impulsi di trigger applicati al gate) possono agire **in qualunque istante**, mentre per il **disinnesco** del componente è necessario un **apposito circuito**.
- Nei circuiti **in alternata** gli impulsi di **comando** (impulsi di trigger applicati al gate) possono agire *solo* negli *intervalli di tempo* corrispondenti alle *semionde positive* della tensione di alimentazione, mentre per il *disinnesco* del componente *non* occorrono circuiti ad hoc: *lo spegnimento avviene automaticamente* a ogni *inversione della tensione di alimentazione*, ossia in corrispondenza *dei passaggi per lo zero* di tale tensione.

CARATTERISTICHE DI CONTROLLO O DI GATE DELL'SCR

Dette anche caratteristiche di ingresso o di innesco, sono le curve più frequentemente fornite dai costruttori e sono analoghe alle caratteristiche del diodo (fra gate e catodo vi è infatti una giunzione). Ciascuna curva raffigura il legame che sussiste fra la corrente di gate I_G e la tensione di gate (chiamata V_{GK} oppure V_G).

Data la forte variabilità dei parametri, anche fra SCR con la stessa sigla¹, il costruttore fornisce le due curve limite (indicate come “1” e “2” nella figura). Nel grafico sono anche tracciate:

- l'**iperbole di massima dissipazione di potenza sul gate**, cioè l'insieme dei punti di funzionamento (I_G , V_{GK}) cui corrisponde la massima potenza dissipata sul gate (P_{GM})
- l'**area di innesco sicuro** che è delimitata dalle due caratteristiche di gate, dall'iperbole P_{GM} , dalle rette tratteggiate verticali $I_G = I_{GT}$ e $I_G = I_{GM}$ e infine dalle rette orizzontali tratteggiate $V_{GK} = V_{GT}$ e $V_{GK} = V_{GM}$; quest'area è l'insieme delle coppie di valori (I_G , V_{GK}) che garantiscono l'innesco dell' SCR
- la **zona di innesco incerto** (in basso a sinistra), corrispondente a bassi valori di tensioni e correnti di gate; questa regione si restringe al crescere della temperatura.

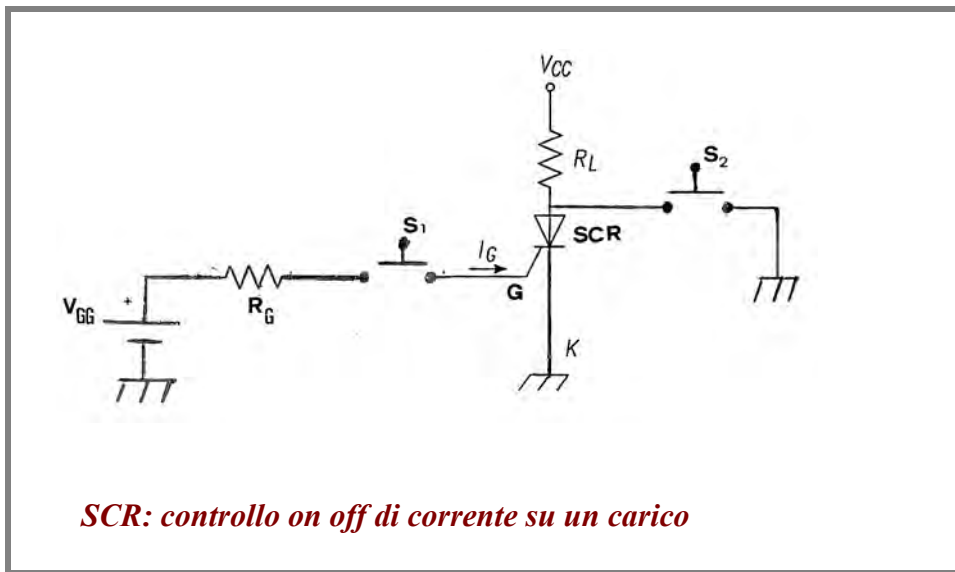


Curve caratteristiche di gate dell'SCR:
con “1” e “2” sono indicate le due curve limite,
di andamento analogo alla caratteristica del diodo

¹ (dispersione delle caratteristiche)

L'SCR PER IL CONTROLLO ON/OFF

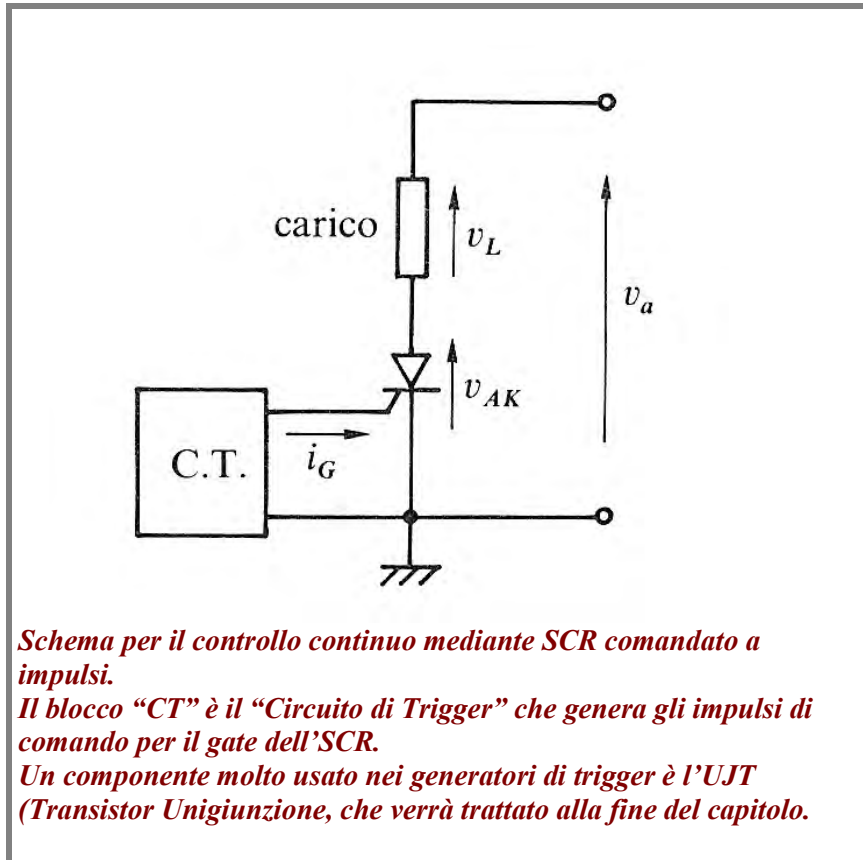
Analizziamo ora uno schema di principio dell'utilizzazione dell'SCR per il controllo on/off della corrente su un carico.



La chiusura momentanea di S_1 consente l'innesco dell'SCR e quindi il passaggio di corrente sull'utilizzatore R_L .

La momentanea chiusura di S_2 determina la deviazione della corrente dall'SCR. Con S_2 chiuso, la corrente di anodo diminuisce fino a scendere al di sotto del valore della corrente di holding I_H , con il conseguente "spegnimento" dell'SCR e il blocco del passaggio di corrente sul carico

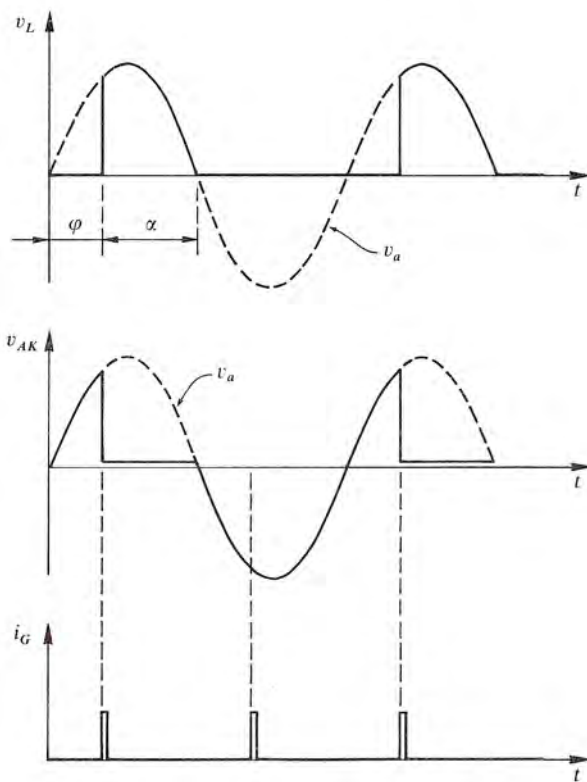
**L'SCR PER IL CONTROLLO CONTINUO DI POTENZA:
SCHEMA DI PRINCIPIO**



Questo tipo di controllo continuo viene chiamato "**controllo in fase**" o "**parzializzazione di fase**". Una possibile realizzazione circuitale prevede infatti la presenza di un generatore di impulsi di ampiezza sufficiente all'innesco dell'SCR e di frequenza regolabile.

In questo modo, variando la frequenza degli impulsi, si modifica il periodo di conduzione e quindi (come risulta evidente facendo riferimento a un asse ωt invece che a un asse t) l'angolo di conduzione.

Variando l'angolo di conduzione varia il valor medio della tensione fornita al carico e dunque la potenza fornita al carico.



Forme d'onda nel controllo in fase mediante SCR comandato a impulsi.

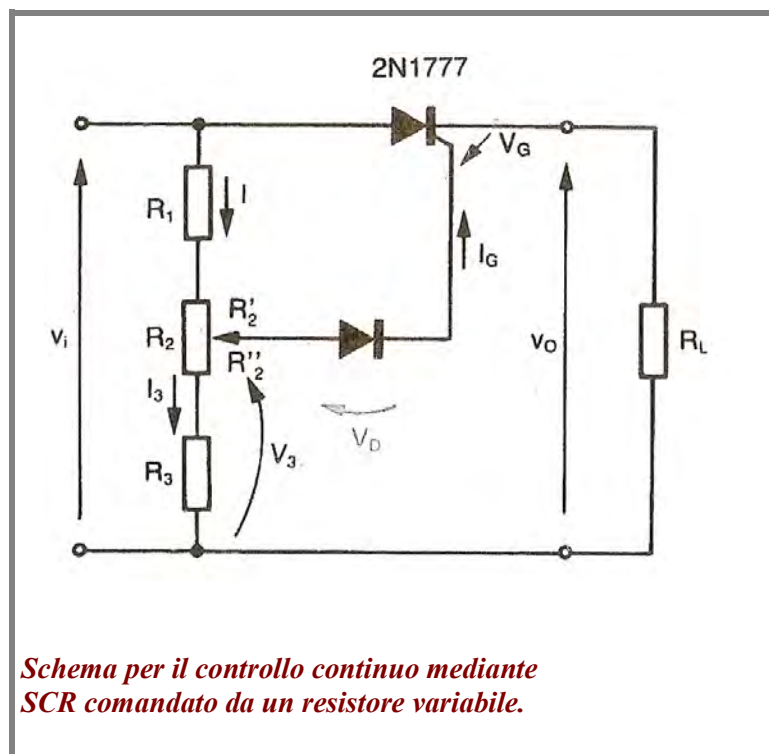
A partire dall'alto:

- *tensione v_a di alimentazione (tratteggiata), tensione v_L sul carico e angolo α di conduzione (solo in corrispondenza dell'angolo di conduzione la tensione di alimentazione risulta applicata al carico)*
- *tensione v_{AK} fra anodo e catodo dell'SCR; nel corso dell'angolo di conduzione l'SCR si comporta come un c.c. e la sua v_{AK} è pressoché nulla*
- *impulsi di corrente generati dal circuito di trigger e applicati al gate dell'SCR*

L'angolo di conduzione, e quindi la potenza fornita al carico, possono essere variati anche senza fare ricorso a un generatore di impulsi, ma servendosi di un semplice resistore variabile.

Spostando (verso l'alto nella figura) il cursore, si aumenta la tensione applicata al gate dell'SCR e quindi si rende più bassa la tensione V_{AK} in corrispondenza della quale l'SCR entra in conduzione.

Sappiamo infatti che più alta è la tensione applicata al gate, maggiore è l'angolo di conduzione.

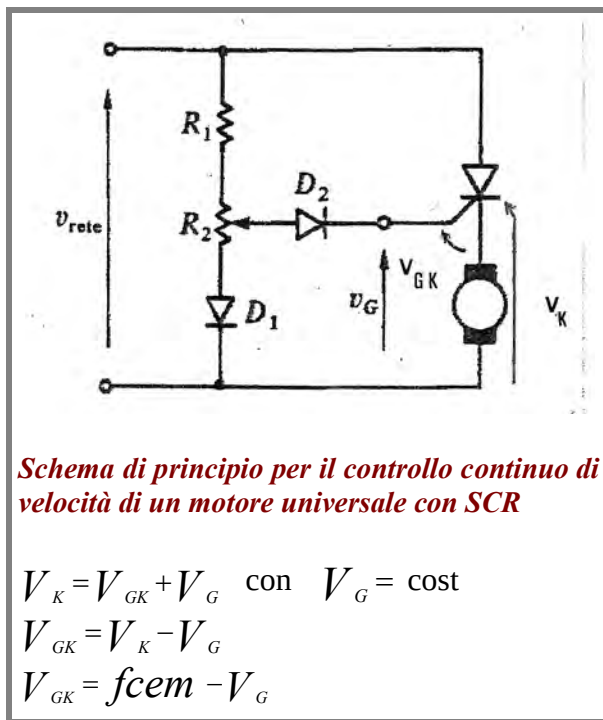


CONTROLLO CONTINUO DI VELOCITA' CON SCR

Un esempio di controllo continuo (o controllo in fase) mediante SCR è il controllo di velocità di un motore universale, per esempio il motore di un trapano.

Il motore universale è un motore che può funzionare sia in continua che in alternata.

La sua struttura è quella di un motore c.c. (ossia in continua), ma il campo magnetico di statore è prodotto non da un magnete permanente, bensì da un avvolgimento collegato in serie con l'avvolgimento di rotore o armatura. Per effettuare questo controllo ci si basa sulla presenza della forza controelettromotrice ai capi del motore: questa tensione è proporzionale al numero dei giri del motore stesso.



Nello schema di principio del circuito la forza controelettromotrice (f.c.e.m.) corrisponde alla tensione V_K fra catodo e massa.

Il circuito (per una fissata posizione del trimmer, cioè per un fissato valore di V_G) tende a mantenere costante la velocità del motore, che supponiamo essere di un trapano, nel modo ora descritto.

Se per qualche motivo (per esempio: parete da perforare meno dura) la velocità di rotazione aumenta, aumenta anche la f.c.e.m. V_K che è proporzionale al numero di giri. Di conseguenza deve diminuire la tensione di controllo v_{GK} dell'SCR, come è verificabile mediante il 2° principio di Kirchhoff. La diminuzione della tensione di controllo determina l'aumento del valore di tensione anodica V_{AK} per il quale l'SCR va in conduzione e, di conseguenza, l'angolo di conduzione diminuisce.

Ne segue che il motore riceve minore potenza e la velocità decresce, tendendo a riportarsi al valore originario.

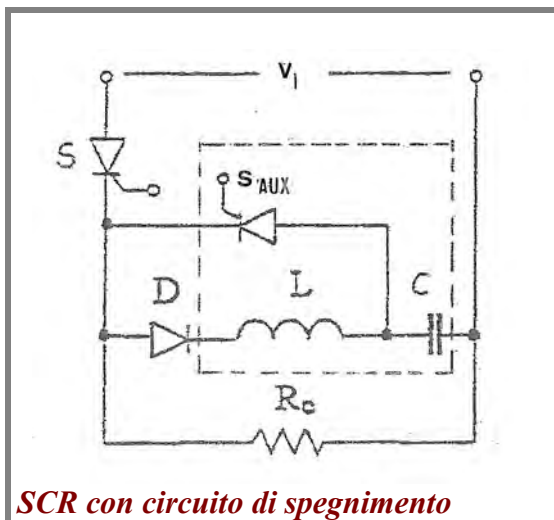
Analogamente si può vedere che, se la velocità del trapano diminuisce, il circuito di controllo tende a farla aumentare ripristinando il valore originario. Se infatti (per effetto per esempio dell' aumento di durezza della parete da perforare) la velocità di rotazione diminuisce, diminuisce la f.c.e.m. V_K che è proporzionale al numero di giri. Di conseguenza deve aumentare la tensione di controllo V_{GK} dell'SCR, come è verificabile mediante il 2° principio di Kirchhoff. L'aumento della tensione di controllo determina la diminuzione del valore di tensione anodica V_{AK} per il quale l'SCR va in conduzione e, di conseguenza, l'angolo di conduzione aumenta, con conseguente aumento di potenza e di velocità.

CIRCUITO DI SPEGNIMENTO CON SCR

Un inconveniente dell'SCR è che, per avere certezza del suo disinnesco, la corrente deve scendere al di sotto del valore di mantenimento almeno per un breve intervallo di tempo (20-100 μ s), il che accade certamente se rendiamo negativa la tensione V_{AK} , cioè se applichiamo al catodo una tensione maggiore di quella anodica.

Per questo motivo i dispositivi basati sull'SCR possono presentare la **necessità** di **circuiti di disinnesco** (per esempio in presenza di **carichi ohmico-induttivi** e nel funzionamento **in continua**).

In base a quanto detto il circuito di spegnimento deve essere in grado di portare nell'istante desiderato il catodo dell' SCR a un potenziale più alto di quello dell'anodo.



Nella figura è rappresentato un circuito ad SCR (S₁) dotato del circuito di spegnimento racchiuso nel tratteggio e basato essenzialmente su:

- un SCR ausiliario S_{AUX}
- un condensatore
- un induttore.

Il circuito di spegnimento funziona nel modo descritto di seguito.

Quando l'SCR principale conduce (e l' SCR ausiliario S_{AUX} è disinnesco), il condensatore si carica.

La tensione alla quale il condensatore si carica è maggiore del valore massimo di quella di ingresso v_i per la presenza dell'induttanza.

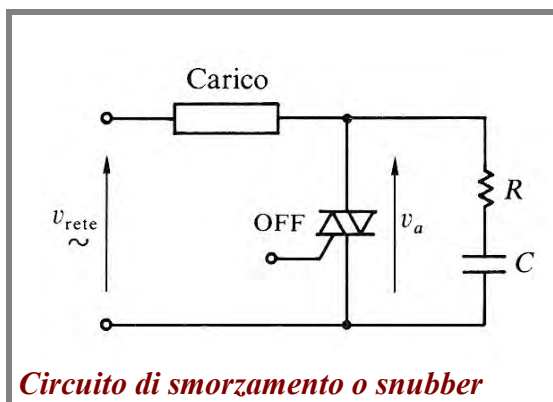
(L'induttore continua a fornire corrente al condensatore, e a caricarlo fino a un valore di tensione pari quasi a $2 V_{imax}$, anche dopo che il condensatore ha raggiunto il valore V_{imax}).

Quando si vuole disinnescare l'SCR principale basta innescare, comandandolo con un impulso, S_{AUX} . L'SCR ausiliario, diventando in pratica un c.c. porterà il suo catodo a un potenziale (V^*) maggiore di quello (V_{imax}) dell'anodo, determinando il disinnescamento dell'SCR principale.

SNUBBER (RETE SMORZATRICE) PER TIRISTORI

E' una serie RC che viene posta in parallelo al tiristore per evitare inneschi indesiderati (inneschi in assenza di comando di gate) dovuti a variazioni troppo rapide della tensione fra i terminali principali, cioè alle situazioni nelle quali la velocità di variazione della tensione v_a (cioè la derivata dv_a/dt , espressa per esempio in $v/\mu s$) ha un il valore troppo elevato. La situazione critica può verificarsi quando:

- il tiristore è applicato a un carico induttivo (corrente in ritardo di 90° sulla tensione)
- l'angolo di conduzione prescelto è 180°
- il tiristore si spegne dopo l'ultimo impulso, quando la corrente scende al di sotto del valore di mantenimento.



Nella situazione descritta, nel momento in cui il tiristore si spegne, la tensione anodica si porta, in maniera pressoché istantanea, al massimo valore della tensione di alimentazione con una velocità di variazione (dv_a/dt) molto elevata. La presenza di una serie RC in parallelo al tiristore fa sì che, nel momento in cui il componente si spegne comportandosi come un circuito aperto, si crei (con il carico induttivo) un circuito RLC il quale fa in modo che la v_a presenti, invece di un istantaneo innalzamento, delle oscillazioni smorzate.

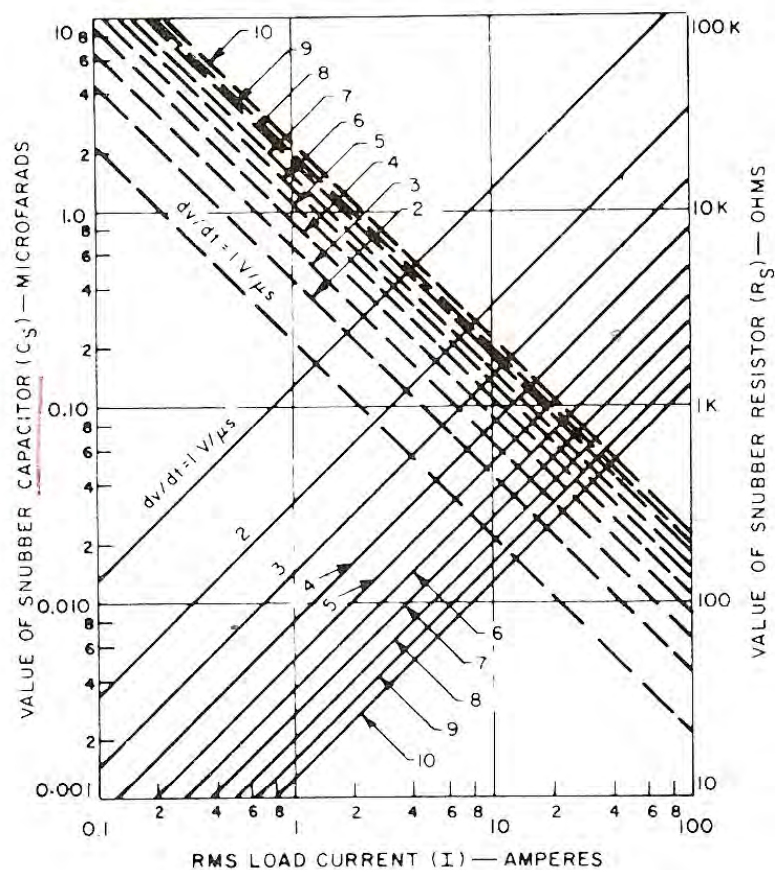
Per facilitare il dimensionamento di R, L e C le aziende produttrici di tiristori forniscono un diagramma nel quale sono riportate due famiglie di curve a $dv_a/dt = \text{costante}$: una famiglia tratteggiata (rappresentativa di R) e una famiglia a tratto intero (rappresentativa di C).

L'asse orizzontale è quotato in termini di valori efficaci di corrente anodica in corrispondenza di angoli di conduzione di 180° .

L'asse verticale sinistro è quotato in termini di C.

L'asse verticale destro è quotato in termini di R.

Noto il valore della corrente (per es. 40 A) e stabilito (in base al valore massimo indicato dal data-sheet) il valore desiderato di dv_a/dt , si traccia, a partire dal valore di corrente, una verticale che intersechi le due curve caratterizzate dal valore desiderato di dv_a/dt .



Diagrammi per il dimensionamento dello snubber

Dal punto di intersezione della verticale con la curva tratteggiata, procedendo in orizzontale verso destra, si ottiene il valore di R.

Dal punto di intersezione della verticale con la curva a tratto intero, procedendo in orizzontale verso sinistra, si ottiene il valore di C.

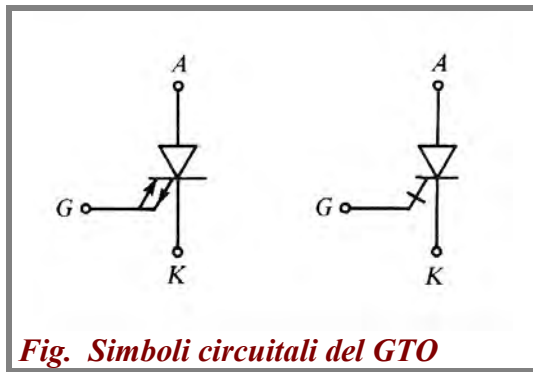
ALTRI COMPONENTI UNIPOLARI

GTO

Gate Turn Off (tiristore spegnibile dal gate)

E' simile all'SCR, ma (a differenza di questo - e come i transistor) può essere riportato all'interdizione mediante un impulso negativo applicato fra gate e catodo.

Lo spegnimento può comunque avvenire, come per l'SCR, abbassando la corrente principale al di sotto del valore I_H di mantenimento o di holding.



I vantaggi del GTO sono:

- **bassa corrente di innesco** sul gate
- può pilotare correnti di parecchi Ampère, commutare fino a 800 A e bloccare fino a 400 V
- **basso tempo di intervento**: il componente va in interdizione in pochi microsecondi, il che lo rende adatto a essere impiegato in applicazioni ad **alta velocità**
- semplicità di inserzione nel circuito
- tensione di interdizione molto piccola (pochi V negativi)

SCS

Silicon Controlled Switch (interruttore controllato al silicio)

Unidirezionale come l'SCR da cui deriva, differisce da esso per la presenza di **due terminali (gate)** di controllo: il **gate anodico** e il **gate catodico**.

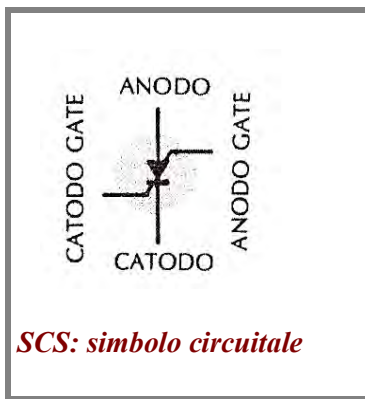
L'**innescò (turn on)** può avvenire in uno dei modi seguenti:

- mediante un impulso **negativo** sul **gate di anodo G_A**
- mediante un impulso **positivo** sul **gate di catodo G_K**

Il disinnescò (turn off) può avvenire in uno dei modi seguenti:

- mediante un impulso positivo sul gate di anodo G_A
- mediante un impulso negativo sul gate di catodo G_K

-



Rispetto all'SCR ha l'inconveniente di sopportare tensioni e correnti meno elevate, ma ha il pregio di avere:

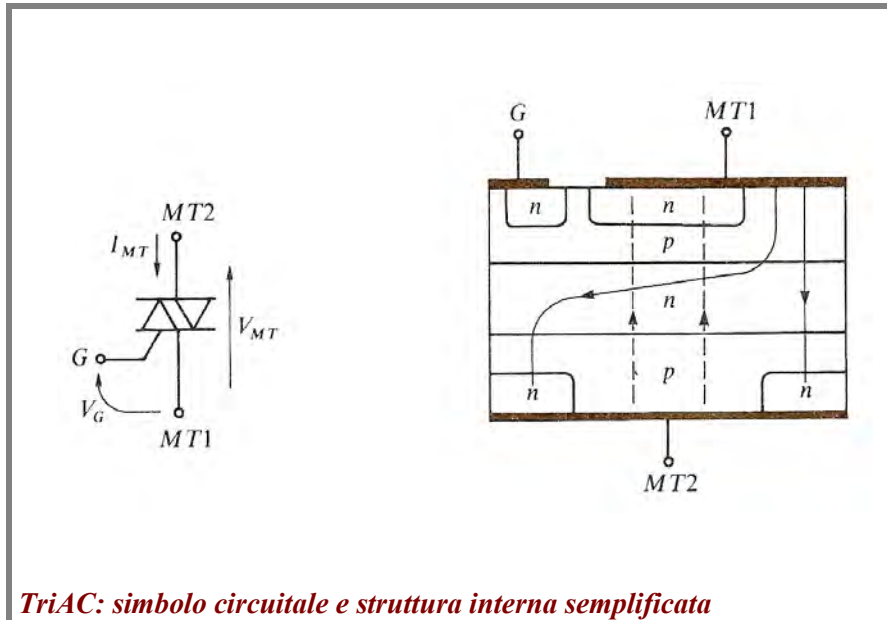
- tempi di disinnescò più brevi
- disinnescò ottenibile con un impulso applicato a uno dei due gate
- possibilità di essere impiegato in applicazioni di tipo digitale, come circuiti di temporizzazione, registri e contatori

COMPONENTI BIPOLARI TriAC, DiAC E APPLICAZIONI

TriAC

TRIodo AC (Triodo per alternata)

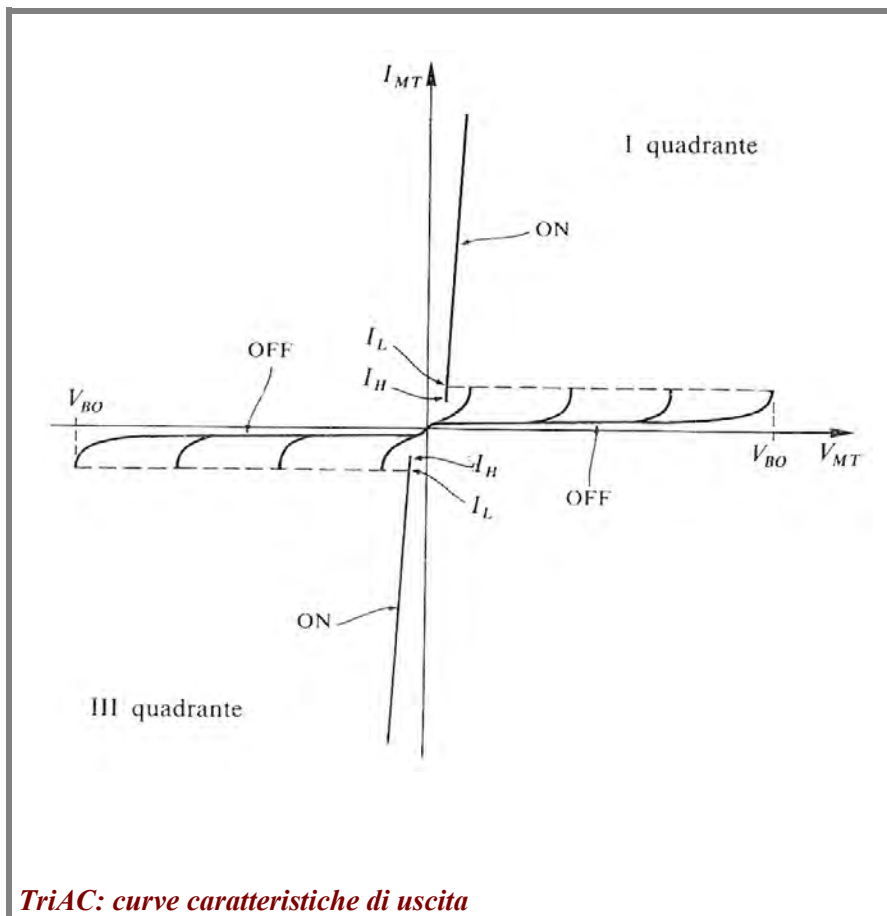
Il simbolo circuitale è simile a quello dell'SCR, ma la presenza di due triangolini orientati in direzioni opposte indica che, contrariamente all'SCR, il TriAC è bidirezionale, cioè può lasciar scorrere la corrente in entrambi i versi.



Perciò non si parla più di anodo e di catodo, risultando questi terminali interscambiabili, ma genericamente di “terminali principali” (Main Terminal) MT_1 e MT_2 .

Da quanto detto, si può intuire che le caratteristiche del TriAC sono analoghe a quelle dell'SCR, però si sviluppano sia nel primo quadrante sia nel terzo (con correnti e tensioni di segno opposto).

Come si vede dalla figura, le caratteristiche risultano simmetriche rispetto all'origine del riferimento.



Nel primo quadrante è

$$V_{MT} > 0$$

$$I_{MT} > 0$$

e la tensione di comando V_G può essere di polarità uguale e opposta a $V_{MT} > 0$. Cioè può essere $V_G > 0$ oppure $V_G < 0$. Però la sensibilità di innesco è lievemente maggiore con la V_G positiva. Ciò è dovuto alla struttura non perfettamente simmetrica del dispositivo.

Nel terzo quadrante è

$$V_{MT} < 0$$

$$I_{MT} < 0$$

e, anche stavolta, la tensione di comando V_G può essere di polarità uguale e opposta a V_{MT} . Cioè può essere $V_G > 0$ oppure $V_G < 0$. Però la sensibilità di innesco è molto maggiore con la V_G negativa.

NOTE

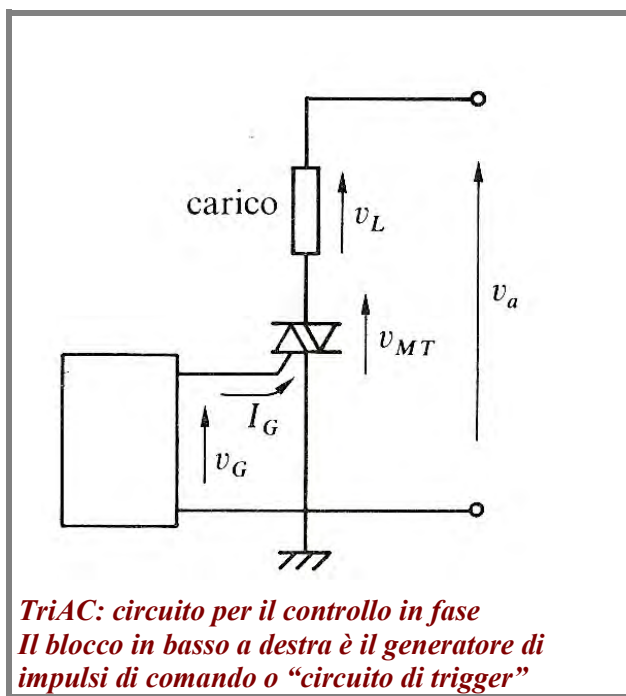
Per V_{MT} si intende $V_{MT2} - V_{MT1}$, come si evince dal simbolo circuitale. Sempre dal simbolo si vede che V_G è il potenziale del gate rispetto a MT_1 e che la corrente I_{MT} è ritenuta positiva quando va da MT_2 a MT_1 .

L'impulso di trigger in ingresso è sempre riferito a MT_1 , così come, nel caso dell'SCR, è sempre relativo al catodo.

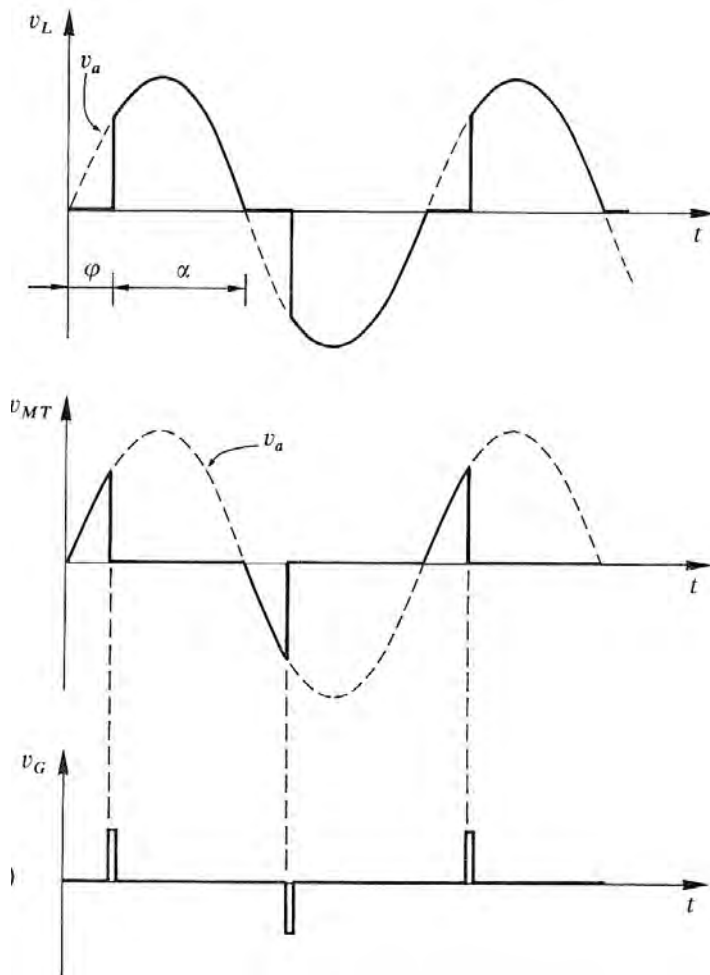
CONTROLLO IN FASE CON TriAC

Anche il TriAC, come l'SCR, può essere utilizzato (oltre che come interruttore, che attivi o disattivi “completamente” un dispositivo) per il controllo in fase (controllo continuo), inserendolo in un circuito come quello mostrato nella figura seguente.

Si parlerà però di **controllo in fase a onda intera** -e non a semionda- dal momento che il TriAC è bidirezionale.



Vediamo ora come avviene il controllo in fase con TriAC:



**Forme d'onda nel controllo in fase mediante TriAC
comandato a impulsi.**

A partire dall'alto:

- tensione v_a di alimentazione (tratteggiata), tensione v_L sul carico e angolo α di conduzione (solo in corrispondenza dell'angolo di conduzione la tensione di alimentazione risulta applicata al carico)
- tensione v_{MT} fra i terminali principali del TriAC; nel corso dell'angolo di conduzione il TriAC si comporta come un c.c. e la sua v_{MT} è pressoché nulla
- impulsi di corrente, alternativamente positivi e negativi generati dal circuito di trigger e applicati al gate del TriAC

All'arrivo dell'impulso di trigger (impulso di controllo), il TriAC passa in conduzione, sicché la tensione di alimentazione v_a viene a cadere sul carico ($v_L=v_a$).

Il TriAC si spegne quando la sua corrente I_{MT} scende sotto il valore di mantenimento I_H . In pratica si può ritenere che il TriAC si spenga quando la tensione di alimentazione si inverte, cioè in corrispondenza del passaggio dell'onda per lo zero.

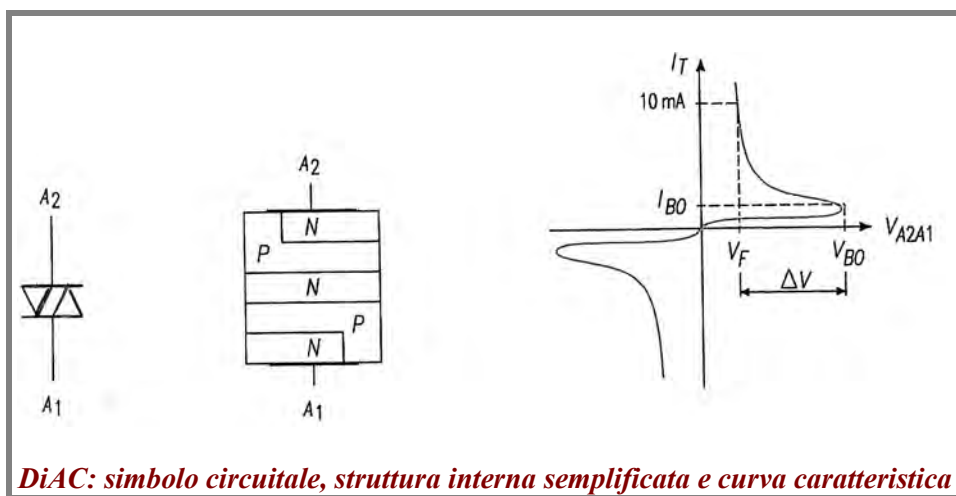
Per riaccendere il TriAC nel corso della semionda negativa occorre un altro impulso di gate, il quale, per avere una buona sensibilità all'innesco, è opportuno che sia negativo.

DiAC

DIodo AC (Diodo controllato per alternata)

Il DiAC è un dispositivo bidirezionale, simile in ciò al TriAC, ma sprovvisto del terminale di gate ossia del terminale di controllo.

Ovviamente quindi non può essere innescato mediante impulsi applicati all'elettrodo di controllo, che non esiste.



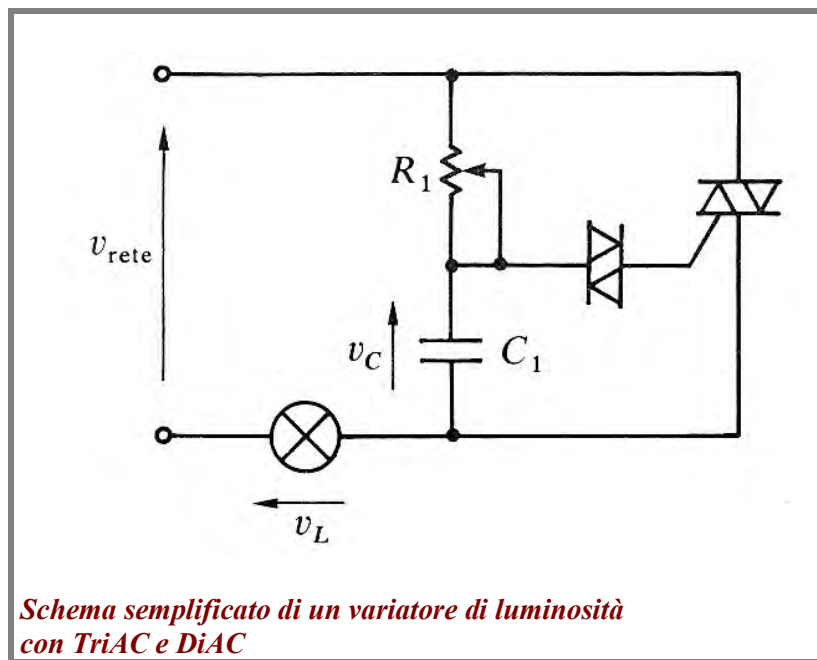
Il DiAC entra in conduzione quando la tensione applicata ai due morsetti, secondo una qualsiasi delle due possibili polarità, raggiunge il valore di breakover. Una volta che ai capi del dispositivo sia stato raggiunto il valore di breakover, la corrente fluisce attraverso il DiAC nel verso imposto dalla polarità della tensione presente ai due terminali. A questo punto ai capi del DiAC si stabilisce una tensione minore del valore di breakover. Il DiAC esce di conduzione, cioè si disinnesca, quando la corrente scende al di sotto del valore di mantenimento. Possiamo dire che il DiAC è un tiristore di piccola potenza, utilizzato quasi esclusivamente per l'innescò del TriAC e di altri componenti affini. Un tipico valore della tensione di breakover è 30 V. Comunque **la V_{BO} è compresa nel range (28÷39) V** (contro i 30-1200V dell'SCR). La corrente di innescò è dell'ordine delle centinaia di μA .

Sigla	V_{BO}		
	min	typ	max
DB 3	28	32	36
DB 4	35	40	45
DC 34	30	34	38
DC 38	35	38	42
DC 42	39	42	45

*DiAC: valori della tensione di breakover
per alcune sigle di DiAC*

UN' APPLICAZIONE DEL TriAC E DEL DiAC: VARIATORE DI LUMINOSITA' PER LAMPADE A INCANDESCENZA (DIMMER)

E' un circuito, chiamato talvolta "dimmer" che permette di variare, mediante un trimmer, la frazione di periodo (della tensione di alimentazione) durante la quale la lampada è alimentata, facendo variare così la luminosità. Osserviamo prima di tutto che la lampada si accende per due volte in ciascun periodo della tensione di alimentazione e, in particolare, si accende in tutti gli istanti in cui la tensione v_c ai capi del condensatore uguaglia la tensione di breakover (positiva o negativa) del DiAC.



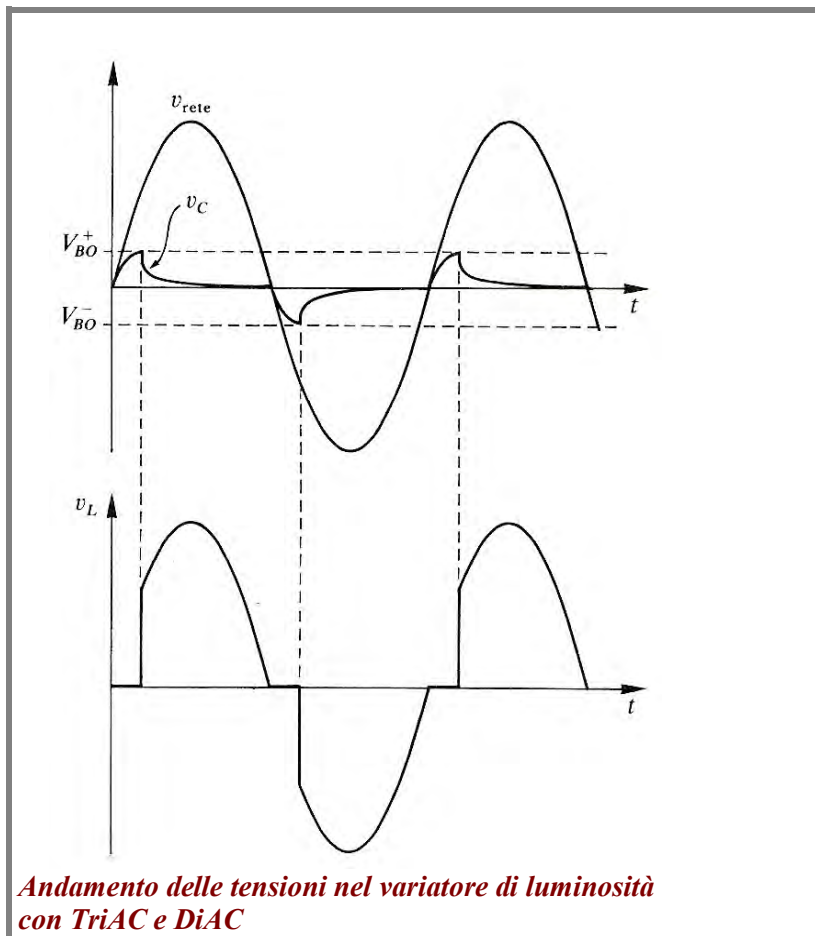
Appena infatti la v_c uguaglia la tensione di breakover del DiAC, il DiAC si innescando fornendo al TriAC un impulso di corrente sufficiente a mandarlo in conduzione e quindi a determinare l'accensione della lampada.

Una volta che DiAC e TriAC siano innescati, il condensatore si scarica attraverso il resistore variabile e il TriAC.

Il TriAC, poi, si disinnescano quando la tensione di alimentazione si azzerà (passaggio dalla semionda positiva alla semionda negativa).

Un nuovo innescamento (caratterizzato però da tensione sul condensatore negativa) avviene nel corso della semionda negativa della tensione di alimentazione: il

condensatore si carica (con polarità opposta rispetto a quanto avveniva durante la semionda positiva) fino a quando la tensione ai suoi capi non uguaglia il valore di breakover negativo del DiAC.

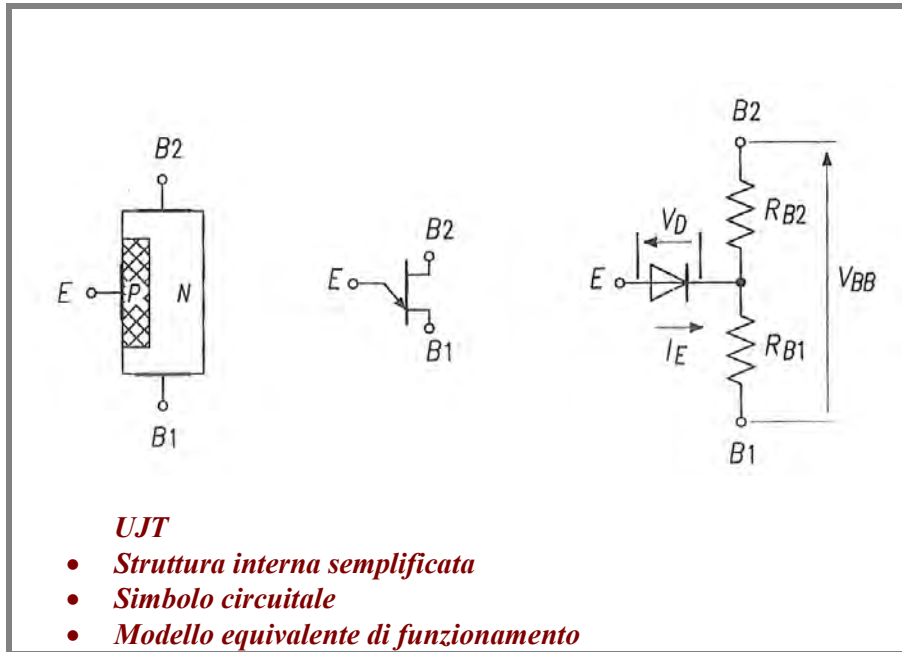


REGOLAZIONE

Agendo sul resistore variabile è possibile variare l'intervallo temporale (e quindi l'angolo di conduzione) e, di conseguenza, la luminosità.

Diminuendo il valore di resistenza (cursore verso l'alto nel disegno) la costante di tempo di carica del condensatore diminuisce, per cui la tensione sul condensatore raggiunge più rapidamente il valore di breakover del DiAC.

COMPONENTI PER L'INNESCO DELL'SCR E DEL TRIAC: L'UJT O TRANSISTORE UNIGIUNZIONE



Il transistor unigiunzione UJT è un componente a tre terminale detti:

- emettitore
- base 1 (B1)
- base 2 (B2)

L'UJT è formato da una barretta di silicio a drogaggio "n", nella quale è realizzata una giunzione "pn", con una zona "p" fortemente drogata, facente capo all'emettitore.

Per spiegare il funzionamento dell'UJT si è concepito un circuito equivalente (un modello teorico) formato da un diodo e da due resistenze R_{B1} ed R_{B2} .

Il diodo rappresenta la giunzione pn, mentre R_{B1} ed R_{B2} rappresentano rispettivamente la resistenza interna equivalente della zona di silicio compresa fra l'emettitore e la base B1 e quella della zona di silicio compresa fra l'emettitore e la base B2.

La **somma** $R_{B1}+R_{B2}$ rappresenta la resistenza complessiva fra i terminali delle due basi. **Il valore che questa resistenza assume quando I_E è nulla** (cioè quando R_{B1} assume il suo massimo valore) è detto “**resistenza di interbase**”, è indicato come **$R_{BB}=R_{B1}+R_{B2}$** e ha un valore compreso fra i 5 K Ω e i 10 K Ω .

Il valore di **R_{B1}** varia in maniera inversamente proporzionale al valore della corrente di emettitore I_E . Per questo la **R_{B1}** è rappresentata talvolta come una **resistenza variabile**. (Al crescere di I_E il valore di **R_{B1}** passa da diversi K Ω a poche decine di Ω).

Si definisce “**rapporto intrinseco**” o “**rapporto di stand-off**” la grandezza (di valore normalmente compreso fra 0,5 e 0,8):

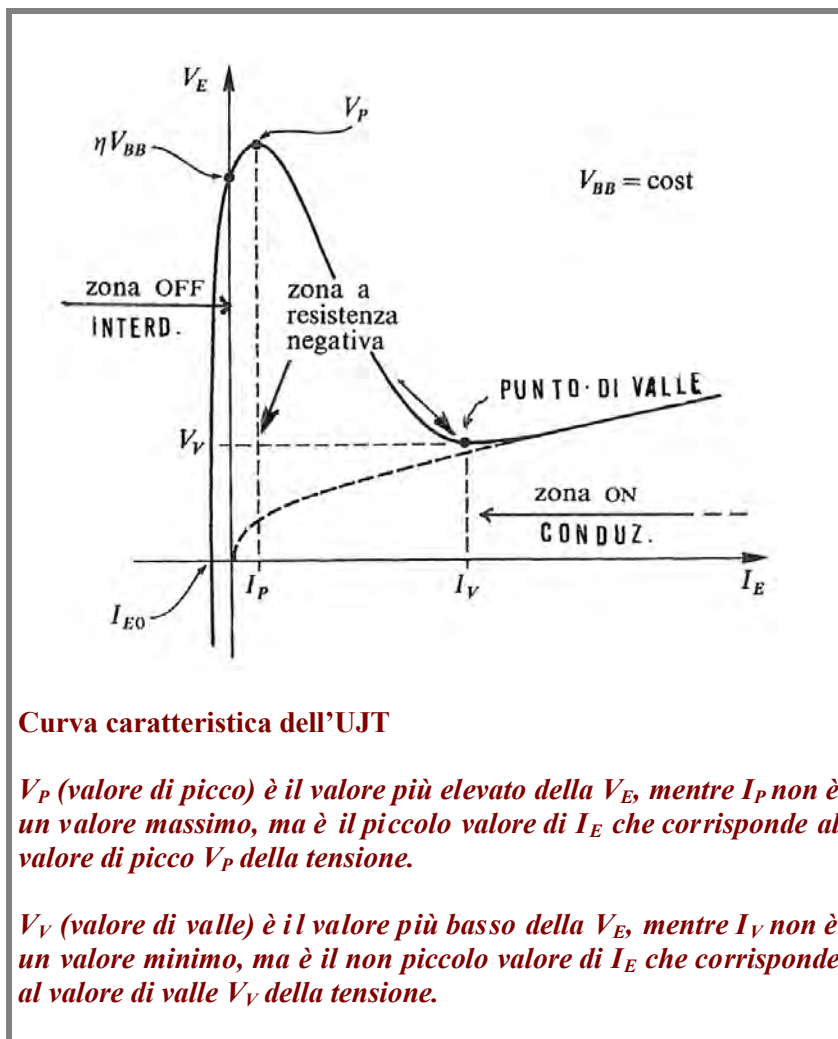
$$\eta = \frac{R_{B1}}{R_{B1} + R_{B2}} = \frac{R_{B1}}{R_{BB}}$$

Quando all’UJT è applicata una tensione V_{BB} , che abbracci entrambe le basi, la **tensione ai capi della resistenza R_{B1}** può essere espressa, mediante la regola del partitore di tensione come:

$$V_{R_{B1}} = \frac{R_{B1}}{R_{B1} + R_{B2}} \cdot V_{BB} = \eta \cdot V_{BB}$$

Analizziamo ora il comportamento dell'UJT e la sua **curva caratteristica** $V_E = V_E(I_E)$, che descrive l'andamento della tensione V_E di emettitore in funzione della corrente I_E di emettitore.

La curva caratteristica che rappresentiamo si ottiene per un fissato valore² costante della tensione V_{BB} .



² Per valori di V_{BB} più bassi, avremmo un abbassamento della curva perché, al diminuire di V_{BB} , diminuisce V_P .

Detta V_γ la tensione di soglia (potenziale di barriera) del diodo rappresentativo della giunzione pn dell'UJT, possiamo descrivere il funzionamento dell'UJT suddividendolo in “zone di funzionamento”.

ZONA DI INTERDIZIONE O ZONA OFF $(I < I_P)$

Fino a quando la tensione V_E applicata all'emettitore risulta inferiore al valore di picco, cioè al valore

$$V_P = \eta \cdot V_{BB} + V_\gamma$$

La giunzione non è polarizzata direttamente e la corrente di emettitore è assente o trascurabile.

La regione di funzionamento ora descritta è detta “regione di interdizione” o “zona off”

ZONA A “RESISTENZA NEGATIVA” $(I_P \leq I < I_V)$

Appena la tensione di emettitore V_E raggiunge il valore di picco, la giunzione risulta direttamente polarizzata, per cui il diodo inizia a condurre e I_E comincia ad essere apprezzabile. Nella barretta di silicio a drogaggio n viene iniettato un consistente flusso di cariche che innalza fortemente la concentrazione dei portatori liberi di carica, abbassando, di conseguenza, drasticamente il valore di R_{B1} . Diminuisce allora anche la tensione V_E , finché R_{B1} non si stabilizza sul suo valore minimo, che è di qualche decina di Ohm. In corrispondenza del valore minimo di R_{B1} , anche la tensione V_E scende al livello minimo, detto “livello di valle”. Osserviamo che dal momento in cui, avendo la V_E raggiunto il valore di picco, l'UJT si è innescato, esso ha cominciato a funzionare in una zona a resistenza “negativa” (V_E diminuisce all'aumentare di I_E), nella quale rimane fino al raggiungimento del punto di valle.

ZONA DI CONDUZIONE O DI SATURAZIONE $(I \geq I_V)$

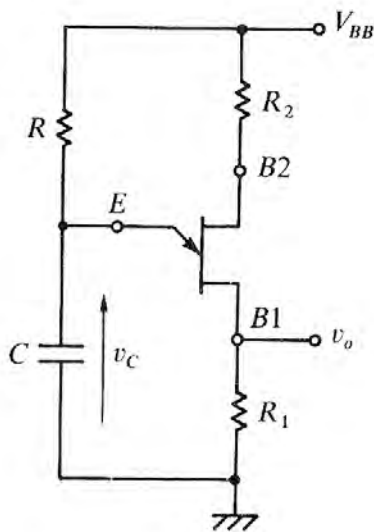
Chiamiamo I_V la corrente di valle. A partire dal punto di valle, la V_E riprende, lentissimamente a crescere al crescere di I_E , come in un normale diodo. Questa terza zona di funzionamento è detta “di conduzione” (o “di saturazione”)

UTILIZZAZIONI DELL'UJT

I transistor unigiunzione, per le loro caratteristiche, sono impiegati:

- per generare tensioni a dente di sega
- per generare tensioni impulsive
- per circuiti di temporizzazione
- per **comandare** componenti di **commutazione di p otenza** come **SCR e TriAC**

GENERATORE DI IMPULSI AD UJT (OSCILLATORE A RILASSAMENTO) PER IL PILOTAGGIO DI SCR



*Generatore di impulsi
(oscillatore a rilassamento)*

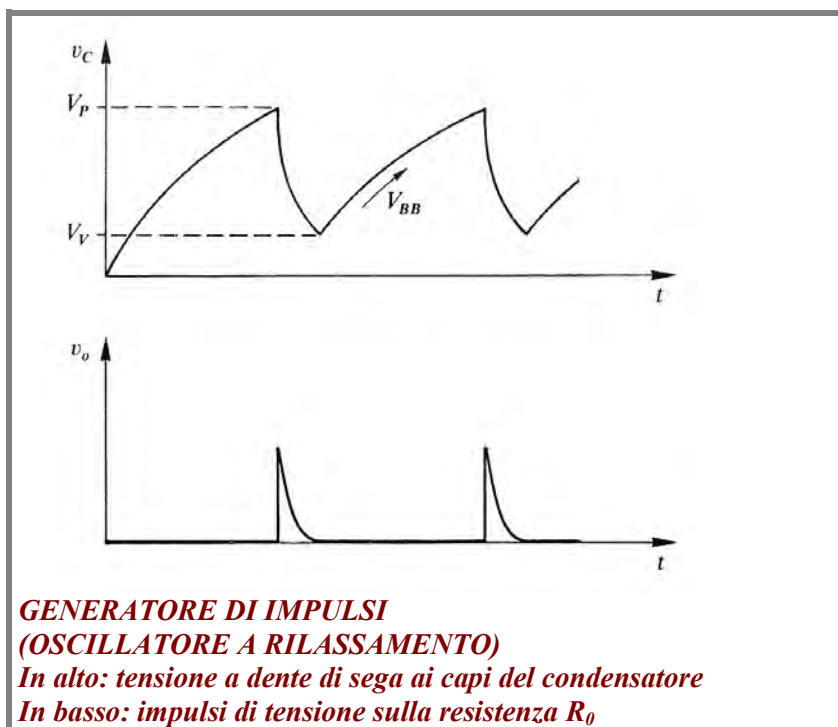
VALORI TIPICI

$R = \text{alcuni } K\Omega \quad (R \gg R_1)$

$R_1 = (10 \div 100) \Omega \quad (\text{resistore di basso valore sul quale vengono prelevati gli impulsi di uscita})$

$R_2 = (100 \div 1000) \Omega \quad (\text{resistore di stabilizzazione termica})$

Il circuito fornisce, **ai capi di R_1** , una sequenza periodica v_o di **impulsi** che possono essere utilizzati **per innescare SCR anche di grossa potenza**. L'ampiezza degli impulsi è di **alcuni volt**.



Il periodo della sequenza di impulsi è dato da:

$$T = R \cdot C \cdot \ln \left(\frac{V_{BB} - V_V}{V_{BB} - V_P} \right)$$

Nelle ipotesi semplificative: $V_V = 0$ e $V_P = \eta \cdot V_{BB}$

il periodo si può esprimere come:

$$T \cong R \cdot C \cdot \ln \left(\frac{1}{1 - \eta} \right)$$

CAPITOLO 42

RETI DI COMPUTER CABLATE E WIRELESS

COMPLEMENTI

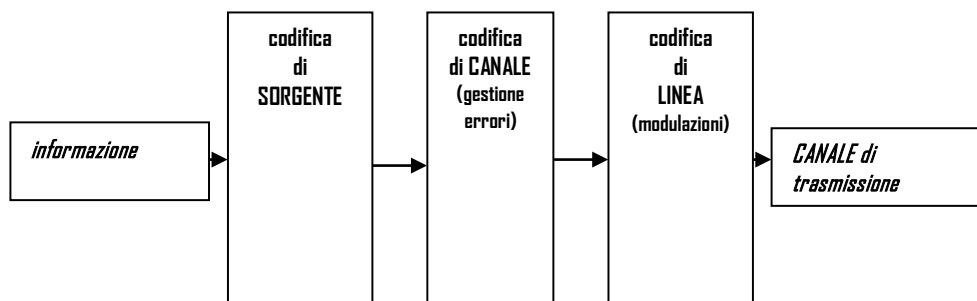
TECNICHE E TECNOLOGIE PER RETI CABLATE E WIRELESS
TERMINOLOGIA DEI SISTEMI DI COMUNICAZIONE NUMERICA
CODIFICA DI LINEA
CODIFICA DI CANALE: METODI ARQ E FEC PER IL CONTROLLO DEGLI ERRORI
CLASSIFICAZIONE DELLE METODOLOGIE DI TRATTAMENTO DEGLI ERRORI
CRC (CONTROLLO CICLICO DI RIDONDANZA)
ALTRI METODI ARQ PER IL TRATTAMENTO DEGLI ERRORI:
CONTROLLO DI PARITÀ, CHECKSUM
PROTEZIONE A RIDONDANZA DI BLOCCO: VRC,LRC
METODI FEC (FORWARD ERROR CORRECTION) PER LA CORREZIONE DIRETTA DEGLI ERRORI
MODULAZIONE PSK
MODULAZIONE QAM-TCM
SCM (SIGNAL CODE MODULATION)
MODULAZIONE GFSK (FSK GAUSSIANA)
MODULAZIONE GMSK (MSK GAUSSIANA)
SPREAD SPECTRUM O “SPETTRO ESPANSO”
DSSS (SPETTRO ESPANSO MEDIANTE SEQUENZA DIRETTA)
FHSS (SPETTRO ESPANSO CON SALTO DI FREQUENZA PORTANTE)
CDMA (ACCESSO MULTIPO A DIVISIONE DI CODICE)
TDMA (ACCESSO MULTIPO A DIVISIONE DI TEMPO)
CSMA/CD
CSMA/CA
MDMA

ONDE CONVOGLIATE (POWERLINE)

TERMINOLOGIA DEI SITI WEB

TECNICHE E TECNOLOGIE PER RETI CABLATE E WIRELESS CODIFICHE E MODULAZIONI PROTOCOLLI DI COMUNICAZIONE

TERMINOLOGIA DEI SISTEMI DI COMUNICAZIONE NUMERICA



CODIFICA

Per “codifica”, in generale, si intende il complesso di elaborazioni alle quali l’informazione deve essere sottoposta per poter essere trasmessa.

Il segnale subisce, nell’ordine, le codifiche di sorgente, di canale, di linea.

CODIFICA DI SORGENTE

All'informazione viene **associato un alfabeto**, cioè un **insieme discreto di simboli**, per esempio i **bit**.

Esempio di codici sorgente sono i codici ASCII, Morse, Baudot.

Una codifica di sorgente è quindi quella che effettua un computer trasmittente, prima di inviare un messaggio.

CODIFICA DI CANALE

Sono detti “codici di canale” i codici per la **gestione degli errori**.

Sono di questo tipo il “codice a maggioranza” il “codice di parità”, il CRC (Codice Ciclico a Ridondanza).

CODIFICA DI LINEA

Per “codifica di linea” si intende quel complesso di operazioni, come, per esempio, la **modulazione** in banda base, e la modulazione in banda traslata, che servono a **rendere il segnale** (contenente l'informazione) **adatto al transito in un determinato canale** di comunicazione.

I codici di linea che realizzano la modulazione in banda base sono il Bifase, il Bifase Differenziale, il Manchester, l'HDB3, il Miller.

La codifica di linea inoltre **trasforma** un messaggio digitale, costituito da una sequenza di “1” e di “0” **logici** in un segnale “fisico” **elettrico o ottico**.

CODIFICA DI LINEA

Dopo essere stato “generato” mediante una codifica di sorgente e dopo aver – eventualmente- subito la codifica di canale per la protezione dagli errori, un messaggio digitale viene sottoposto alla codifica di linea.

La **codifica di linea** è l’operazione che **trasforma** un messaggio digitale, costituito da una sequenza di “1” e di “0” **logici** in un segnale costituito da una **sequenza di impulsi elettrici o ottici**.

Si può anche dire che la codifica di linea **definisce il formato degli impulsi** elettrici o ottici ai quali il messaggio viene affidato.

Se sul mezzo trasmissivo si realizza un **canale passabasso** (canale che occupa una banda che include la frequenza zero), la codifica di canale ha la funzione di **minimizzare la distorsione**.

Se invece sul mezzo trasmissivo si realizza un *canale passabanda* (canale che occupa una banda che esclude la frequenza nulla e le frequenze molto basse), allora la codifica di linea deve effettuare una *traslazione dello spettro* e quindi tale codifica è una *modulazione in banda traslata* (per esempio le modulazioni numeriche su portanti analogiche utilizzate per la trasmissione-dati su rete telefonica commutata: *FSK, PSK, QAM. ecc.*).

L’implementazione della codifica di linea, ossia la trasformazione del segnale digitale in una sequenza di impulsi di determinato formato, è effettuata dal trasmettitore che (grazie a un codificatore e a un generatore di clock) trasmette il segnale digitale sul canale alla cadenza imposta dal clock di trasmissione. Gli impulsi possono poi essere eventualmente sagomati con un filtro.

Siccome anche il ricevitore lavora al ritmo di un clock, in quanto deve effettuare la lettura dei bit ricevuti in determinati istanti (al “centro” del tempo di bit), è opportuno che il clock di ricezione sia sincronizzato con il clock di trasmissione.

Spesso l’informazione di sincronizzazione viene estratta dallo stesso segnale trasmesso: in tal caso il segnale deve essere tale da consentire agevolmente l’estrazione dell’informazione di sincronizzazione (deve contenere una riga spettrale alla frequenza di clock).

In definitiva il segnale elettrico associato al codice di linea deve avere le seguenti caratteristiche:

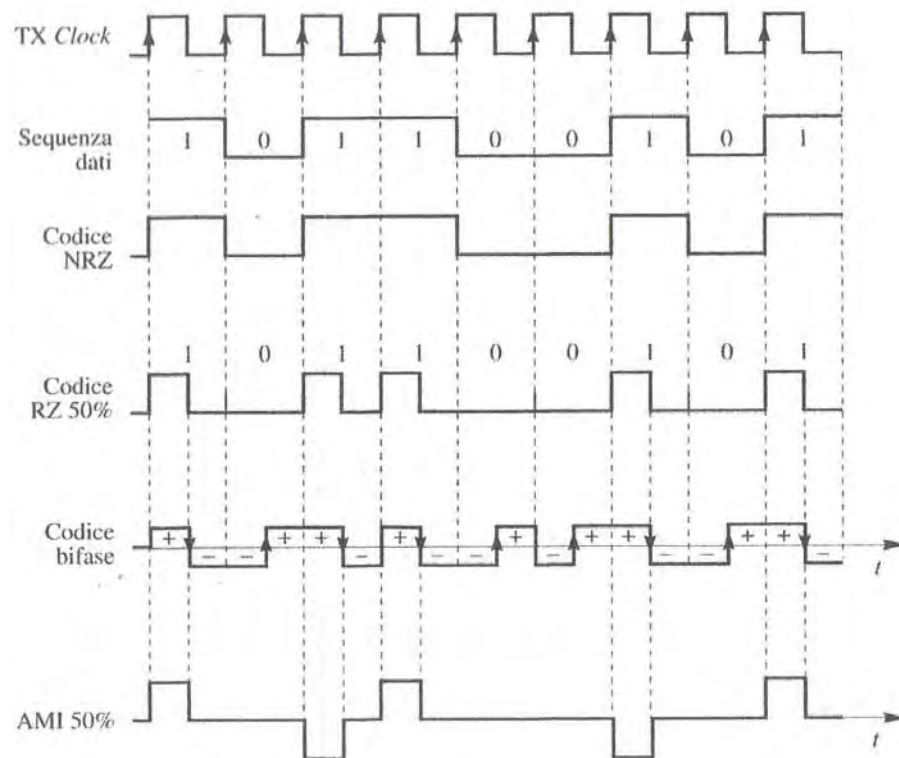
- spettro che sia contenuto all’interno della banda del canale, per evitare distorsione

- deve consentire, in maniera semplice, l'aggancio fra il clock di ricezione e il clock di trasmissione.

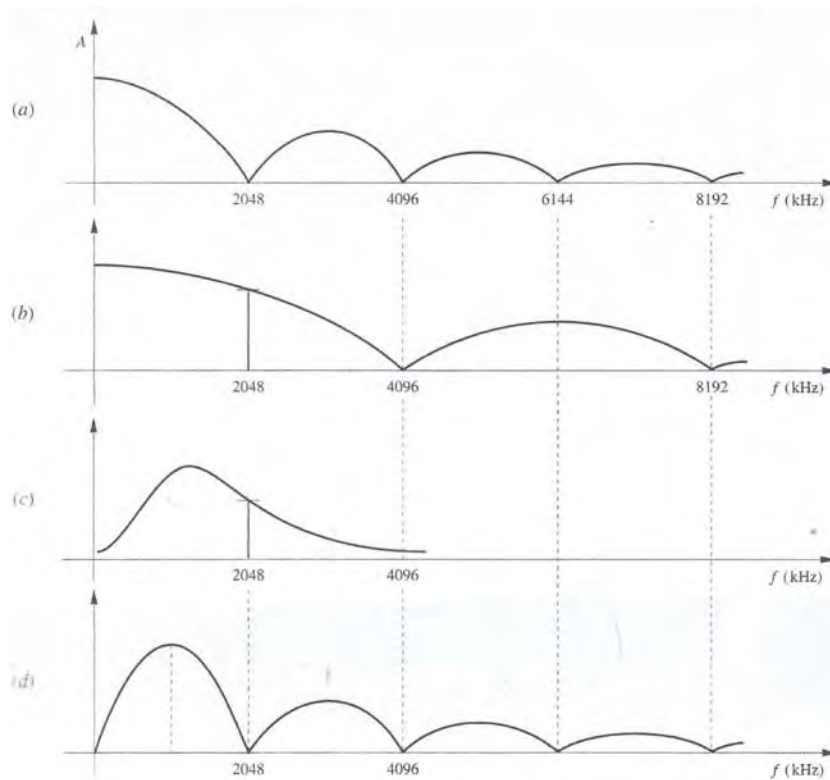
Inoltre nel caso di trasmissione su cavi metallici deve presentare valore medio nullo, dal momento che, generalmente, si richiede l'isolamento galvanico (accoppiamento in AC) fra trasmettitore e linea.

Possiamo individuare tre tipi di codici di linea:

- Codici BINARI (NRZ, RZ, Manchester o Bifase, Bifase differenziale, mB-nB per F.O.)
- Codici PSEUDOTERNARI (AMI, HDB-3)
- Codici MULTILIVELLO (2B-1Q)



Codici di linea: andamenti temporali



Codici di linea: SPETTRI

- a. spettro del codice NRZ
- b. spettro del codice RZ al 50%
- c. spettro del codice bifase
- d. spettro del codice AMI al 50%

CODICE di LINEA (1)	TIPO	UTILIZZAZIONE	COMPONENTE CONTINUA (VAL. MEDIO)	COMPONENTE SPETTRALE a FREQUENZA di CLOCK	SEGNALE ELETTRICO CORRISPONDENTE
NRZ Non Return Zero = Not Return to Zero = Senza passaggio per lo zero	BINARIO senza "passaggio per lo zero"	non risponde ai requisiti di un codice di linea perciò è utilizzato per collegare apparato ad apparato e su F.D.	DIVERSA DA ZERO (il che è negativo)	NO (il che è negativo)	UNIPOLARE (solo nullo e positivo)
RZ Return Zero = Return to Zero = Con passaggio per lo zero	BINARIO con "passaggio per lo zero"	non utilizzato su mezzo metallico, può essere ricavato, per conversione, dal codice utilizzato in linea	DIVERSA DA ZERO (il che è negativo)	SI (il che è positivo)	UNIPOLARE (solo nullo e positivo)
BIFASE o MANCHESTER	BINARIO con "passaggio per lo zero"	LAN (reti locali di computer)	NULLA (il che è positivo)	SI (il che è positivo)	BIPOLARE (valori positivi e negativi: $+V_0$; $-V_0$)
BIFASE DIFFERENZIALE	BINARIO con "passaggio per lo zero"	trasmissione dati: modem in banda base	NULLA (il che è positivo)	SI (il che è positivo)	BIPOLARE (valori positivi e negativi: $+V_0$; $-V_0$)
AMI 50% Alternate Mark Inversion cioè con inversione alternata di polarità degli impulsi associati al bit "1"	PSEUDOTERNARIO con "passaggio per lo zero" (assume 3 livelli, due dei quali però sono associati allo stesso simbolo)	collegamenti numerici dedicati trasmissione segnali PCM (USA)	NULLA (il che è positivo)		BIPOLARE a TRE livelli (valori nullo, positivo negativo: 0; $+V_0$; $-V_0$)
HDB-3 High Density Bipolar = Bipolare ad alta densità- Pseudoternario	PSEUDOTERNARIO con "passaggio per lo zero" (assume 3 livelli, due dei quali però sono associati allo stesso simbolo)	trasmissione segnali PCM (Europa)	NULLA (il che è positivo)		
MULTILIVELLO 2B-1Q = 2Binary-1Quaternary = bit del segnale-dati raggruppati a 2 a 2; segnale associato a quattro livelli	QUATERNARIO (a 4 livelli) 00 $\rightarrow$ $-V_0$ 01 $\rightarrow$ $-V_0/3$ 11 $\rightarrow$ $+V_0/3$ 10 $\rightarrow$ $+V_0$ (non c'è 0V)	trasmissione ISDN: collegamento utente centrale			

<u>CODICE</u> di <u>LINEA</u> <u>(2)</u>	TIPO	UTILIZZAZIONE	COMPONENTE CONTINUA (VAL. MEDIO)	COMPONENTE SPETTRALE a FREQUENZA di CLOCK	SEGNALE ELETTRICO CORRISPONDENTE
codici mB-nB esempi: 1B-2B 6B-8B	a un blocco di m bit in ingresso viene fatto corrispondere un blocco di n bit in uscita, in modo che si abbia un numero uguale di "1" e di "0" anche in sequenze molto lunghe	CODICI di linea per trasmissioni su FIBRA OTTICA (F.O.)			
codici mB-(n+1)N esempio 5B-6B	evitano l'inconveniente degli mB-nB (incremento della velocità di trasmissione richiesta e della ridondanza, conseguenti all'aumento del numero di bit di uscita rispetto a quelli di ingresso)	CODICI di linea per trasmissioni su FIBRA OTTICA (F.O.)			

Forniamo ora ulteriori informazioni sui codici precedentemente elencati.

<u><i>CODICE di LINEA</i></u>		
<u><i>APPROFONDIMENTO</i></u>	<u><i>Costruzione del CODICE</i></u>	<u><i>Caratteristiche</i></u>
<u><i>(1)</i></u>		
NRZ Non Return Zero = Not Return to Zero = Senza passaggio per lo zero	Il segnale associato al codice NRZ si mantiene costante (cioè si mantiene allo stesso livello alto o basso) per tutta la durata del tempo di bit. Il segnale può essere bipolare e in tal caso: al bit "0" si associa la tensione $-V_0[V]$ al bit "1" si associa la tensione $+V_0[V]$ Se invece il segnale è unipolare: al bit "0" si associa la tensione $0[V]$ al bit "1" si associa la tensione $+V_0[V]$	
RZ Return Zero = Return to Zero = Con passaggio per lo zero	Il codice RZ viene ricavato dall'NRZ: - riducendo al 50% la durata dell'impulso relativo al bit "1" - lasciando al livello 0 il segnale nel tempo relativo al bit "0"	Avendo una componente a frequenza di clock il segnale RZ permette la sincronizzazione fra trasmettitore e ricevitore. La presenza della componente continua e la larghezza della banda occupata (doppia rispetto al corrispondente RZ) ne escludono, di norma, l'impiego su cavo metallico.
BIFASE o MANCHESTER	Il codice si implementa mediante un segnale bipolare caratterizzato da una transizione di livello al centro del tempo di bit. Il segnale si "costruisce" in questo modo: il bit "0" corrisponde a un periodo del segnale di clock inalterato ; il bit "1" corrisponde a un periodo del segnale di clock <u>sfasato di 180°</u>	Il segnale relativo al codice Manchester presenta sempre una transizione al centro del tempo di bit, il che facilita la rivelazione in ricezione. Alla perfetta alternanza di bit "0" e "1" (sequenza 101010...) corrisponde un segnale con frequenza metà di quella di clock.
BIFASE DIFFERENZIALE	Analogamente al Manchester, il codice si implementa mediante un segnale bipolare caratterizzato da una transizione di livello al centro del tempo di bit, stavolta però l'inversione di fase non si riferisce al clock, ma al segnale stesso in corrispondenza del precedente tempo di bit. Il segnale si "costruisce" in questo modo: in corrispondenza del bit "0" si trasmette la stessa "porzione" di segnale del precedente tempo di bit , senza sfasamenti; in corrispondenza del bit "1" si trasmette la <u>"porzione" di segnale del precedente tempo di bit, sfasata di 180°</u>	

<u>CODICE</u> <u>di</u> <u>LINEA</u> <u>APPROFONDIMENTO</u> <u>(2)</u>	<u>COSTRUZIONE del CODICE</u>	<u>CARATTERISTICHE</u>
AMI 50% Alternate Mark Inversion = con inversione alternata di polarità degli impulsi associati al bit "1"	Il codice si implementa mediante un segnale bipolare a tre livelli (valori nullo, positivo, negativo: 0; $+V_0$; $-V_0$) che vengono fissati nel modo seguente: bit "0" → assenza di impulso (0V) bit "1" → impulsi alternativamente positivi ($+V_0$) e negativi ($-V_0$)	Si parla di AMI al 50% quando, come nel caso che stiamo trattando, gli impulsi hanno durata pari al 50% del tempo di bit, in quanto il codice è ricavato a partire dall'RZ. Dall'AMI 50% è possibile ricavare, mediante raddrizzamento a onda intera, il corrispondente RZ, che permette di sincronizzare il clock di ricezione a quello di trasmissione. Vantaggio del codice AMI: possibilità di rivelare errori attraverso la verifica della regola dell'alternanza. Inconveniente: siccome in corrispondenza di uno zero non si trasmette nulla in linea, una lunga sequenza di bit "0" può determinare la perdita del sincronismo. Possibili rimedi sono l'inserimento di uno scrambler ("mescolatore di bit") a monte del codificatore AMI oppure l'adozione del codice HDB-3 che impedisce il formarsi di sequenze di più di tre zeri consecutivi Esiste anche un AMI al 100%, caratterizzato da impulsi la cui durata copre l'intero tempo di bit, in quanto il codice è ricavato a partire dall'RZ.

<u>CODICE</u> <u>di</u> <u>LINEA</u> <u>APPROFONDIMENTO</u> <u>(3)</u>	<u>COSTRUZIONE del CODICE</u>	<u>CARATTERISTICHE</u>
HDB-3 High Density Bipolar = Bipolare ad alta densità- Pseudoternario	Premessa: Indichiamo con "V" un "impulso di violazione", cioè un impulso che viola la regola dell'alternanza di polarità fra gli impulsi che codificano i bit "1" (gli zeri sono codificati da assenza di impulsi) Indichiamo invece con "B" un impulso che <i>rispetta</i> la regola dell'alternanza di polarità fra gli impulsi che codificano i bit "1" (gli zeri sono codificati da assenza di impulsi). Pertanto B ha segno opposto a V. Il segnale HDB-3 viene costruito con le stesse leggi dell'AMI, con la differenza che tutte le sequenze di quattro zeri consecutivi vengono sostituite con: la sequenza BOOV se, a partire dall'ultima violazione V, vi è un numero pari di bit "1" la sequenza OOOV se, a partire dall'ultima violazione V, vi è un numero dispari di bit "1"	Altri codici analoghi all'HDB-3 sono: Il BGZS (Bipolare con Sostituzione di 6 zeri) che sostituisce a ogni gruppo di 6 zeri consecutivi la sequenza: BOVB0V Il CHDB-4 (HDB-Compatibile) che sostituisce, a ogni gruppo di 5 zeri consecutivi, quella sequenza (scelta fra OOOOV e OOB0V) che mantiene nullo il valor medio.
MULTILIVELLO 2B-1Q = 2Binary-1Quaternary = bit del segnale-dati raggruppati a 2 a 2; segnale associato a quattro livelli	I bit del segnale-dati vengono raggruppati a 2 a 2 e si costruisce un segnale a quattro livelli in base alla seguente regola: 00 → -V₀ → -2,5V (valori tipici) 01 → -V₀/3 → -0,833 11 → +V₀/3 → +0,833 10 → +V₀ → +2,5V (non c'è 0V)	
codici mB-nB esempi: 1B-2B 6B-8B	a un blocco di m bit in ingresso viene fatto corrispondere un blocco di n bit in uscita, in modo che si abbia un numero uguale di "1" e di "0" anche in sequenze molto lunghe	
codici mB-(n+1)N esempio 5B-6B	Evitano l'inconveniente degli mB-nB (incremento della velocità di trasmissione richiesta e della ridondanza, conseguenti all'aumento del numero di bit di uscita rispetto a quelli di ingresso)	

CODIFICA DI CANALE: METODI ARQ E FEC PER IL CONTROLLO DEGLI ERRORI

Nella sequenza di elaborazioni cui viene sottoposto il messaggio binario nel corso della trasmissione-dati, la codifica di canale provvede all'introduzione di una **ridondanza** nei codici, finalizzata alla **rivelazione**, e, se possibile, alla **correzione** degli **errori**.

Ogni sistema di telecomunicazione è sottoposto a distorsioni e a rumore.

Nei sistemi digitali il rumore, sommandosi al segnale utile, ne modifica la forma.

Se il livello del rumore è rilevante, l'alterazione del segnale può essere tale da determinare un'interpretazione errata dell'informazione trasmessa: un bit "1" può essere interpretato come "0" o viceversa.

La presenza di errori è descritta dal "rapporto segnale/rumore" (indicato come S/N o come SNR).

Risulta dipendere da S/N anche il "**tasso di errore**", ossia il **BER (Bit Error Rate)**, definito come:

$$\text{BER} = \frac{\text{numero di bit ERRATI}}{\text{numero di bit totalmente TRASMESSI}}$$

EFFETTO DEGLI ERRORI

L'effetto degli errori sul segnale digitale è diverso a seconda della natura del messaggio binario. Infatti:

- il segnale-dati, che è intrinsecamente digitale, tollera solo una piccolissima quantità di errori, in quanto l'errore altera "direttamente" il contenuto informativo;
in trasmissione-dati, per esempio, il tasso di errore –generalmente- non deve superare il valore $\text{BER}=10^{-6}$ (un bit errato per ogni milione di bit trasmessi)
- i segnali digitalizzati (cioè i segnali che, come la voce e l'immagine, sono originariamente analogici) tollerano una maggiore quantità di errori, in quanto gli errori (entro un determinato tasso) comportano solo un accettabile degrado di qualità; in telefonia digitale, per esempio, è tollerato un tasso di errore massimo $\text{BER}=10^{-3}$ (un bit errato per ogni mille bit trasmessi).

CLASSIFICAZIONE DELLE METODOLOGIE DI TRATTAMENTO DEGLI ERRORI

Si individuano due metodologie:

- **ARQ:** metodi a **richiesta (REquest) Automatica di Ripetizione**
- **FEC:** metodi a **Correzione diretta (Forward) degli Errori**

CRC (Controllo Ciclico di Ridondanza) (Cyclic Redundancy Check)

Nell'ambito della **codifica di canale**, il CRC è un metodo di tipo **ARQ** per il **controllo** (e la successiva correzione) **degli errori** prodottisi nel corso di una trasmissione-dati, a causa del rumore.

Sono detti ARQ i metodi a **Richiesta (reQuest) Automatica di Ripetizione**.

Con i metodi ARQ in un primo momento il ricevitore si limita a rivelare (senza correggerli) gli errori. Successivamente chiede la ritrasmissione dei blocchi di bit che contengono errori: gli errori vengono quindi **corretti per ri-trasmissione**. Nella tecnica ARQ la rivelazione degli errori presuppone l'adozione di un **codice a blocco**, cioè di un codice per il quale il messaggio (segnale-dati binario) è suddiviso in blocchi di determinata lunghezza.

In coda a ciascun blocco viene aggiunto un prefissato numero di bit di controllo (o di ridondanza) che il codificatore determina in base a un certo algoritmo¹, che è noto anche al ricevitore.

<i>B l o c c o</i>		<i>B l o c c o</i>	
<i>gruppo di bit informativi</i>	<i>gruppo di bit di ridondanza o di controllo</i>	<i>gruppo di bit informativi</i>	<i>gruppo di bit di ridondanza o di controllo</i>

Il **ricevitore**, che **conosce l'algoritmo**, è in grado, analizzando i bit ricevuti, di **rivelare la presenza di errori**.

Se infatti un blocco contiene errori, la sequenza di bit di controllo, ricalcolata dal ricevitore, risulterà diversa da quella trasmessa.

In pratica però il ricevitore, invece di ricalcolare realmente i bit di ridondanza, può effettuare delle elaborazioni matematiche sui bit ricevuti (bit di ridondanza e bit informativi del singolo blocco) e il tipo di risultato ottenuto farà capire se il blocco contiene, o meno, errori.

Il metodo CRC è in grado di rivelare più errori anche con commutazione simultanea.

Oltre al CRC, sono metodi di tipo **ARQ**:

- il **CONTROLLO DI PARITA'**
- il **CHECKSUM (SOMMA DI CONTROLLO)**

¹ Per "algoritmo" si intende un qualunque procedimento che conduce a un risultato in un numero finito di passi

PREMESSE ALLA DETERMINAZIONE DEI BIT DEL CRC

1. Polinomio associato a una stringa (o blocco) di bit:
dato un blocco di n bit, ad esso si può associare un polinomio, nella variabile x,
di grado n-1, i cui coefficienti sono proprio i bit del blocco.

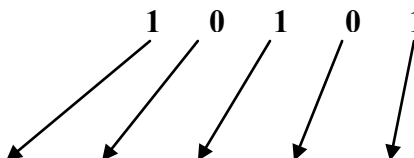
Esempio:

dato il blocco 10101 (costituito da n=5 bit),

ad esso può essere associato il polinomio:

$$x^4 + x^2 + 1$$

che si ricava nel modo seguente:


$$1 \cdot x^4 + 0 \cdot x^3 + 1 \cdot x^2 + 0 \cdot x^1 + 1 \cdot x^0 \quad \rightarrow \quad x^4 + x^2 + 1$$

2. L'aggiunta di m bit in coda a un blocco di bit equivale a elevare di m unità il grado del polinomio (e quindi equivale a moltiplicare il polinomio per x^m)

Data, per esempio, la stringa di bit 10101

alla quale corrisponde il polinomio $x^4 + x^2 + 1$

se in coda aggiungiamo due bit "0" otteniamo la stringa:

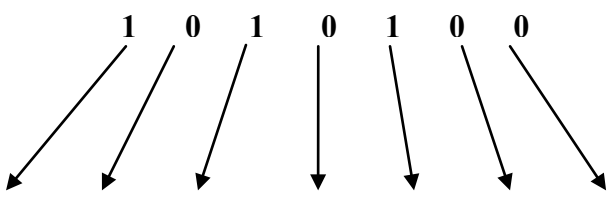
1010100

Alla quale corrisponde il polinomio binario

$$x^6 + x^4 + x^2$$

che ha un grado maggiore di m=2 unità rispetto al polinomio associato alla stringa precedente.

Verifichiamo la corrispondenza fra la nuova stringa e il "suo" polinomio:


$$1 \cdot x^6 + 0 \cdot x^5 + 1 \cdot x^4 + 0 \cdot x^3 + 1 \cdot x^2 + 0 \cdot x^1 + 0 \cdot x^0 \quad \rightarrow$$

$$\rightarrow x^6 + x^4 + x^2$$

DETERMINAZIONE DEL CRC

Con il termine “**CRC**” indicheremo ora l’insieme dei bit di controllo posti al termine di ogni blocco.

La sequenza delle operazioni per ricavarlo è la seguente:

1. si **suddivide** la sequenza in **gruppi di n bit**
2. **a ciascun gruppo si associa un polinomio $p(x)$ di coefficienti binari**, detto “**polinomio messaggio**”, di **grado $n-1$** , avente come **coefficienti i bit** del gruppo.

Un bit “0” determina l’assenza del corrispondente termine del polinomio.

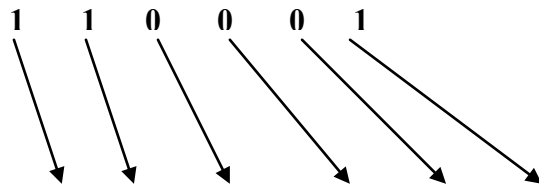
Se, a puro titolo di esempio, consideriamo il gruppo di $n=6$ bit:

110001,

ad esso corrisponde il polinomio di grado 5:

$$p(x) = x^5 + x^4 + 1$$

L’associazione del polinomio al gruppo di bit avviene in questo modo:



$$p(x) = 1 \cdot x^5 + 1 \cdot x^4 + 0 \cdot x^3 + 0 \cdot x^2 + 0 \cdot x^1 + 1 \rightarrow$$

$$p(x) = x^5 + x^4 + 1$$

3. Si sceglie (fra quelli proposti dall'ente normativo ITU-T), un *polinomio* $g(x)$ di *grado "g"*, detto "*polinomio generatore*" che useremo come *divisore* del polinomio $p(x)$.

Il polinomio $g(x)$ è noto sia al trasmettitore che al ricevitore e non dipende dal messaggio.

Se, per semplicità², si sceglie un grado $g=3$, il polinomio generatore sarà di grado 3, e potrà essere, per esempio:

$$g(x) = x^3 + 1$$

² Nella pratica il polinomio generatore ha un grado molto più elevato di quello scelto come esempio.

Ecco alcuni polinomi generatori proposti dall'ITU-T

$$g(x) = x^{16} + x^{12} + x^5 + 1 \quad (\text{CRC - CCITT})$$

$$g(x) = x^{16} + x^{15} + x^2 + 1 \quad (\text{CRC - 16})$$

$$g(x) = x^{12} + x^{11} + x^3 + x^2 + x + 1 \quad (\text{CRC - 12})$$

$$g(x) = x^{16} + x^{11} + x^4 + 1 \quad (\text{CRC - CCITT inverso})$$

$$g(x) = x^{16} + x^{14} + x + 1 \quad (\text{CRC - 16 inverso})$$

4. Si aggiunge, in coda al blocco di bit da trasmettere, un numero “g” di bit “0” di ridondanza (il che equivale a moltiplicare per x^g il polinomio messaggio $p(x)$, oppure a elevare di g unità il grado di $p(x)$).

Se, per esempio, scegliamo $g=3$, il gruppo di bit

110001

Diventa:

110001**000**

Riguardo al polinomio messaggio ($p(x) = x^5 + x^4 + 1$ nell'esempio),

lo modifichiamo nel modo seguente:

$$p'(x) = x^g \cdot p(x) = x^3 \cdot p(x) = x^8 + x^7 + x^3$$

(notiamo che il grado di $p'(x)$ è aumentato di 3 rispetto al grado di $p(x)$ e che i suoi coefficienti sono: 110001000)

5. Si esegue, in “aritmetica modulo due”³, il rapporto:

$$\frac{p'(x)}{g(x)} \quad \text{cioè:} \quad \frac{p(x) \cdot x^g}{g(x)}$$

Di questo rapporto non ci interessa il quoziente $q(x)$, bensì il **polinomio-resto $r(x)$** .

Infatti, con i **coefficienti del polinomio-resto $r(x)$** , costruiamo il **codice CRC** da inserire in coda al blocco al posto degli zeri di ridondanza aggiunti al punto 4.

Nell'esempio in esame, eseguendo la divisione “modulo due” si ottiene un “**resto**”, cioè un **CRC** pari a **111**.

Il “rapporto modulo due” si ottiene effettuando l'operazione logica di exor fra i coefficienti del polinomio presente al numeratore e i coefficienti del polinomio presente al denominatore.

³ L'aritmetica dei polinomi binari si effettua in “modulo due”, per cui: nella somma non si considerano i riporti, nella differenza non si considerano i prestiti e le operazioni di somma e di differenza si equivalgono e vengono eseguite mediante exor bit per bit.

6. Come avviene il controllo in ricezione?

Il trasmettitore ha aggiunto a ciascun blocco di n bit un CRC di g bit (ottenuti come coefficienti del polinomio-resto r(x)).

Il ricevitore divide il “polinomio messaggio modificato”, al quale è stato aggiunto (al posto degli ultimi g zeri) il CRC, per il polinomio generatore, che gli è noto.

Esegue cioè il rapporto:

$$\frac{p(x) \cdot x^g + CRC}{g(x)}$$

Se il **resto** di tale rapporto⁴, calcolato in “aritmetica modulo due” risulta **nullo**, allora il blocco **non** contiene **errori**.

Se invece il resto è diverso da zero, il polinomio contiene errori, che verranno corretti mediante la ritrasmissione del blocco, che sarà richiesta dal ricevitore.

⁴ Rapporto nel quale il CRC è il resto r(x), pari a 111 nell'esempio e dove, relativamente all'esempio, è:

$$p(x) \cdot x^g + CRC = 110001000111$$

ESEMPIO DI DIVISIONE MODULO DUE

Vedremo ora, nella tabella che segue, come si esegue la divisione modulo due fra:

$p'(x)$, avente coefficienti: 110001000

e

$g(x)$, avente coefficienti: 1001

Precisiamo che di questo rapporto ci interessa il resto, e non il quoziente, che potremmo anche non calcolare affatto.

Sotto il dividendo $p'(x)$ riportiamo (seconda riga della tabella) il divisore $g(x)$ e facciamo l'exor bit per bit, cioè la divisione modulo due.

Sotto la stringa di bit ottenuta come risultato dell'exor, si riporta (quinta riga della tabella) nuovamente il divisore $g(x)$ allineandolo, a sinistra, con il primo bit "1" della precedente stringa.

Prima di eseguire il successivo exor, si "abbassano" (quarta riga della tabella) dal dividendo, tante cifre quante ne servono per "coprire" completamente $g(x)$.

Per esempio, se per completare tale allineamento serve un solo bit, si "abbassa" dal dividendo un solo bit; se invece ne occorrono tre, se ne abbassano tre.

Riguardo al quoziente, si mette un 1 quando, per la prima volta, si riporta il divisore sotto il dividendo e ogni volta che, "abbassando" una cifra del dividendo, si ottiene un resto costituito da un numero di bit sufficiente ad allineare il quoziente sotto di esso, a partire dal primo 1 del resto.

Quando invece, "abbassando" una cifra del dividendo, il resto non raggiunge un numero di bit tale da realizzare l'allineamento di cui sopra, si mette uno zero al quoziente e si abbassa un'ulteriore cifra.

Se anche questa cifra è insufficiente, si aggiunge un altro zero al quoziente, si "abbassa" ancora un'altra cifra e così via.

[illegible]

CRC IN SINTESI

Riepiloghiamo in maniera sintetica ciò che si è scritto a proposito della tecnica CRC.

Si suddivide la sequenza di bit da trasmettere in gruppi da n bit
(supponiamo che uno di questi gruppi di bit sia 110001)

Cos'è il CRC

Il CRC è un gruppo di bit “supplementari (cioè “di ridondanza”) che viene aggiunto in coda al gruppo di bit da trasmettere

1 1 0 0 0 1	Bit del CRC
-------------	--------------------

Come si costruisce il CRC

- Si aggiunge un numero g di zeri in coda al gruppo di bit da trasmettere
(per esempio $g=3$ zeri: 000)

1 1 0 0 0 1	0 0 0
-------------	--------------

- I **bit che formeranno il CRC** (e che saranno inseriti al posto degli zeri) sono il **resto** di una particolare operazione detta “**divisione modulo due**” fra la stringa originaria completata con gli zeri e un gruppo di g bit (polinomio divisore o generatore) noti sia al trasmettitore che al ricevitore:

$$\text{Bit del CRC} = \text{resto della divisione modulo 2} = \frac{\text{stringa originaria completata con g zeri}}{\text{bit del polinomio divisore}}$$

(ipotizziamo $g=3$ e supponiamo che il resto della divisione sia “111”.
Il CRC sarà allora proprio “111”)

- I **bit che formano il CRC** vengono **inseriti in coda** alla sequenza originaria, **al posto degli zeri aggiuntivi**

Sequenza originaria di bit	Bit del CRC
1 1 0 0 0 1	1 1 1

Come avviene il controllo in ricezione

Il ricevitore calcola il resto della divisione modulo due:

$$\frac{\text{stringa originaria dei bit del messaggio} + g \text{ bit del CRC}}{\text{bit del polinomio divisore}}$$

Se il resto è nullo, non ci sono errori

Se il *resto* è *diverso da zero*, questo significa che *ci sono stati errori in trasmissione* e il *ricevitore chiede* automaticamente la *ritrasmissione del gruppo di bit*.

ALTRI METODI ARQ PER IL TRATTAMENTO DEGLI ERRORI: CONTROLLO DI PARITA', CHECKSUM

PROTEZIONE A RIDONDANZA DI BLOCCO

Oltre al CRC, che è un metodo “ciclico”, sono metodi di trattamento degli errori “a richiesta automatica di ripetizione” (ARQ):

- il controllo di parità
- il checksum
- la protezione a ridondanza di blocco (VRC, LRC)

CONTROLLO DI PARITA'

Questa tecnica consiste nell'**aggiungere**, alla stringa di bit che codifica un carattere, **un bit “0” o un bit “1”** in modo tale che **il numero di bit “1”** presenti nella sequenza “carattere + bit di ridondanza” **risulti sempre pari** (“parità pari” o “even parity”) oppure risulti sempre dispari (“parità dispari” o “odd parity”).

Il ricevitore controlla il carattere ricevuto e, se la regola prestabilita di parità non risulta rispettata, ciò significa che almeno un bit è errato.

Un limite di questo metodo è l'incapacità di rivelare errori multipli, come, per esempio, la presenza simultanea di due errori nello stesso carattere.

Il controllo di parità è normalmente associato al codice sorgente ASCII, con l'aggiunta di un ottavo bit ai sette del codice sorgente.

CHECKSUM O SOMMA DI CONTROLLO

E' una famiglia di tecniche di controllo degli errori che prevede:

- la **suddivisione del segnale-dati in blocchi** costituiti da un determinato numero di parole di otto o sedici bit (più eventuali bit di riempimento)
- **l'esecuzione di una particolare operazione algebrica** sul blocco di dati e **l'inserimento del risultato in coda al blocco.**

Tra le possibili operazioni ricordiamo

- la somma “bit per bit” delle parole del blocco
- la somma “modulo due” (exor) delle parole del blocco
- il complemento a uno della somma dei complementi di tutte le parole
- la somma delle parole del blocco, la divisione della somma per 255, il prelievo del resto (a otto bit) della divisione e la sua utilizzazione come checksum.

PROTEZIONE A RIDONDANZA DI BLOCCO: VRC (Controllo di Ridondanza Verticale) LRC (Controllo di Ridondanza Longitudinale o orizzontale)

Le tecniche VRC ed LRC sono un'estensione del semplice controllo di parità. Dati m caratteri, formati ciascuno da n bit, essi possono essere rappresentati (o realmente memorizzati in memorie digitali) in una matrice di n righe ed m colonne. Facendo l'esempio di 3 caratteri di 5 bit ciascuno, si ha la matrice:

1	1	0	0	0	<i>carattere o parola</i>
1	0	0	1	1	<i>carattere o parola</i>
1	1	1	1	0	<i>carattere o parola</i>

VRC

La tecnica VRC consiste nel fare un controllo di parità su ciascuna colonna, determinando, per ogni colonna, un bit di parità verticale, e formando, con tali bit, un carattere di controllo che può essere inviato in linea dal trasmettitore e confrontato con quello che il ricevitore determina (sul messaggio ricevuto) utilizzando gli stessi criteri adottati dal trasmettitore.

1	1	0	0	0
1	0	0	1	1
1	1	1	1	0
1	0	1	0	1

LRC

Analogamente, la tecnica VRC consiste nel fare un controllo di parità su ciascuna riga, determinando, per ogni riga, un bit di parità longitudinale, e formando, con tali bit, un carattere di controllo, che può essere inviato in linea dal trasmettitore e confrontato con quello che il ricevitore determina (sul messaggio ricevuto) utilizzando gli stessi criteri adottati dal trasmettitore.

1	1	0	0	0	0
1	0	0	1	1	1
1	1	1	1	0	0

CONTROLLO DI PARITA' INCROCIATO

Considerando sia i bit di parità verticale che quelli di parità longitudinale, si attua il controllo di parità incrociato, che permette anche di individuare un bit globale di controllo.

1	1	0	0	0	0
1	0	0	1	1	1
1	1	1	1	0	0
1	0	1	0	1	<u>1</u>

METODI FEC (Forward Error Correction) PER LA CORREZIONE DIRETTA DEGLI ERRORI

Nell'ambito della codifica di canale, un metodo di controllo degli errori completamente diverso dall'ARQ è il metodo FEC.

Proprio l'aggettivo "forward", cioè "diretto", esplicita la differenza fra le due metodologie: mentre la tecnica ARQ può effettuare la correzione degli errori solo mediante la ri-trasmissione dei blocchi di bit affetti da errori, la tecnica **FEC** consente al ricevitore di correggere gli errori **senza che il messaggio debba essere ritrasmesso**.

Gli aspetti caratterizzanti del metodo di correzione diretta degli errori sono i seguenti.

- In trasmissione i dati sono sottoposti a una codifica "convoluzionale" caratterizzata dal fatto che ciascun bit generato in un certo istante dal codificatore (e quindi presente all'uscita del codificatore stesso) dipende non solo dal bit presente all'ingresso del dispositivo in quello stesso istante, ma anche da un certo numero di bit precedenti;
- In ricezione si utilizzano procedimenti di decodifica che non si basano solo sul controllo del singolo bit, ma anche sull'analisi della sequenza nella quale il bit è inserito⁵; il ricevitore cioè verifica che la sequenza di bit ricevuti rispetti le regole in base alle quali è stata costruita in trasmissione, regole che sono conosciute anche dal ricevitore.

Una delle più note tecniche per il trattamento degli errori in fase di decodifica è l'"**algoritmo di Viterbi**" che adotta il criterio della "**stima della sequenza a massima verosimiglianza**" (MLSE, Maximum Likelihood Sequence Estimation); l'algoritmo cioè è in grado di scegliere (una volta ricevuta una sequenza contenente errori) quella, fra le sequenze possibili, che, con maggiore probabilità, è stata effettivamente emessa dal trasmettitore;

La tecnica FEC può essere attuata anche su un **canale unidirezionale** perché **il ricevitore non deve comunicare nulla al trasmettitore**.

⁵ Si potrebbe dire che il metodo FEC effettua una correzione che tiene conto del contesto nel quale il bit si trova, così come una persona che legge è in grado di rilevare e correggere direttamente una parola sbagliata presente in un brano, cioè in un "contesto", proprio tenendo conto del contesto. Per esempio nella frase "voglio invitarti a cane", la parola "cane", pur essendo di per sé corretta grammaticalmente, risulta errata nel contesto della frase e la correzione dell'errore con il vocabolo giusto (cena) è immediata.

I metodi ARQ richiedono invece la bidirezionalità del canale, in quanto il ricevitore deve poter comunicare al trasmettitore se i blocchi ricevuti contengono errori, per avere la possibilità di chiederne la ritrasmissione.

MODULAZIONE PSK

Usata, tra gli altri sistemi, da **Zigbee** e dal **GPS**, è una modulazione a spostamento di fase (**Phase Shift Keying**) con portante analogica e segnale informativo modulante di tipo digitale.

COME SI INQUADRA LA PSK

La PSK appartiene alla tipologia di modulazione con segnale informativo modulante digitale e portante analogica. Appartiene cioè a una tipologia di modulazione adatta alla trasmissione di dati numerici su un canale analogico. Non a caso la PSK è una delle modulazioni che sono state impiegate, o che sono tuttora utilizzate (da sole o congiuntamente ad altre) dai modem in banda fonica. Questi modem permettono la comunicazione fra dispositivi digitali (computer ecc.) attraverso la “normale” rete telefonica, ossia attraverso la rete telefonica commutata PSTN, caratterizzata dalla limitazione di banda (300-3400) KHz. Le altre modulazioni appartenenti alla stessa tipologia, e quindi utilizzabili o utilizzate dai “voice modem”, sono:

- FSK (Modulazione a Spostamento di Frequenza, Frequency Shift Keying)
- ASK (Modulazione a Spostamento di Ampiezza, Amplitude Shift Keying)
- QAM (Modulazione di ampiezza in quadratura: modulazione “mista” di ampiezza e fase, la quale fa uso di due portanti in quadratura)
- TCM (Modulazione con codifica a traliccio, Trellis Code Modulation): è una versione più sofisticata di QAM, con l'utilizzo di codifiche ridondanti per il trattamento dell'errore ed è la tecnica utilizzata dagli attuali modem

Qual è lo scopo di queste modulazioni?

Il segnale-dati è un segnale digitale ed è pertanto caratterizzato da uno spettro che ha una banda molto larga e dal fatto che concentra gran parte del suo contenuto energetico alla basse frequenze.

E' evidente quindi che segnali di questo tipo non possono viaggiare “così come sono” su un canale trasmissivo (come per esempio la rete telefonica commutata) che taglia le frequenze molto basse e che ha una banda piuttosto ristretta.

Le modulazioni in esame risolvono il problema perché restringono la banda del segnale (rendendolo analogico) e modificano la distribuzione spettrale di energia.

TIPI DI PSK

A seconda del numero di possibili sfasamenti che si vogliono imporre (da parte del segnale informativo modulante) alla portante sinusoidale, e a seconda del modo in cui lo sfasamento viene imposto, si possono individuare differenti “sottotipi” di modulazioni PSK

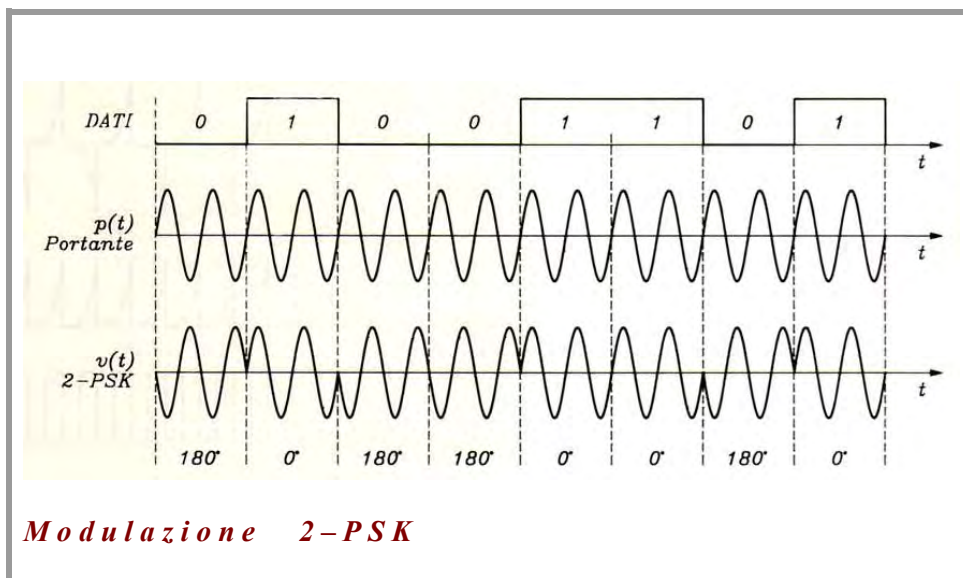
- 2PSK o BPSK (PSK binaria)
- 2PSK differenziale
- 4PSK o Q-PSK
- 8PSK

I differenti sfasamenti vengono chiamati “stati di modulazione” o “livelli di modulazione”, così il segnale modulato 2PSK è caratterizzato da due stati di modulazione, il segnale modulato 4PSK è caratterizzato da 4 stati di modulazione ecc.

LA MODULAZIONE 2PSK

La più semplice modulazione a salto di fase è la 2PSK che (come si vede dalla figura) agisce nel modo seguente:

- in corrispondenza del bit “0” viene imposto alla portante uno sfasamento di 0° ossia la portante non viene sfasata
- in corrispondenza del bit “1” viene imposto alla portante uno sfasamento di 180° ossia la portante viene “invertita”.



Naturalmente è anche possibile fare il contrario: invertire la portante (sfasarla di 180°) nel corso della durata del bit “0” e lasciarla invece inalterata nel corso della durata del bit “1”

La modulazione PSK può essere descritta anche mediante una rappresentazione vettoriale sul piano delle fasi.

LA MODULAZIONE 4PSK o QPSK COME MODULAZIONE “MULTILIVELLO”

Prima di tutto dobbiamo chiederci come è possibile, a partire da un segnale digitale, costituito da cifre (bit) che possono assumere solo due valori, ottenere un segnale modulato caratterizzato da più di due stati o livelli di modulazione, caratterizzato cioè, in questo caso, da più di due valori di sfasamento.

Questo è reso possibile dall'accorpamento della sequenza originaria di bit in gruppetti di più bit che sono applicati simultaneamente al modulatore.

I gruppetti di bit sono chiamati “simboli” o “codici” di trasmissione.

Il “simbolo” o “codice” di trasmissione di una modulazione a quattro stati o livelli, come la 4-PSK, è il dibit (coppia di due bit).

Il “simbolo” o “codice” di trasmissione di una modulazione a otto stati o livelli, come la 8-PSK, è il tribit (accorpamento di tre bit).

In generale il “simbolo” o “codice” di trasmissione di una modulazione a M stati o livelli è un accorpamento di un numero di bit pari a $\log_2 M$.

I vantaggi delle modulazioni multilivello sono:

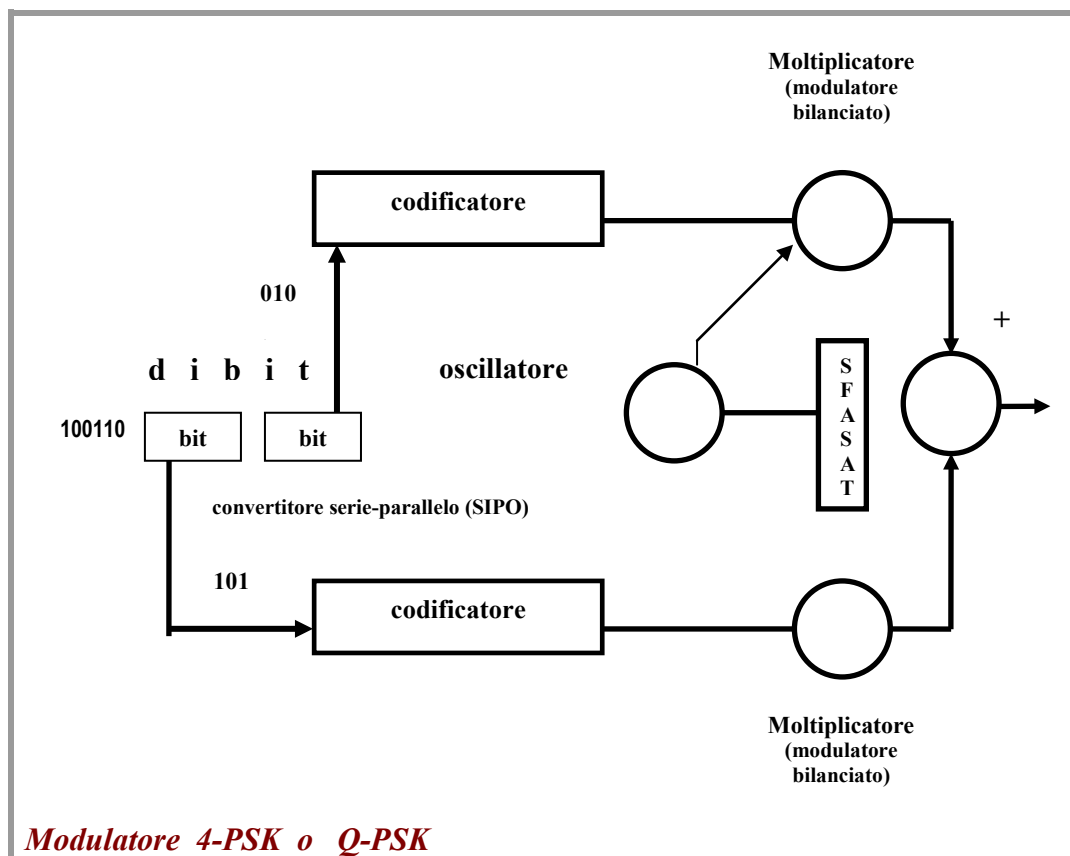
- a parità di bit trasmessi nell'unità di tempo, una minore occupazione di banda (e quindi, a parità di banda, una maggiore velocità)
- una maggiore efficienza di trasmissione, dove, per **efficienza**, si intende il parametro:

$$E = (v/B) * \log_2 M$$

LA MODULAZIONE 4PSK o QPSK

La 4-PSK o Q-PSK è una modulazione a spostamento di fase che produce un segnale modulato caratterizzato da quattro possibili valori di fase (per esempio 45°, 135°, 225°, 315°): da qui il nome di 4-PSK.

Questa tecnica è detta anche QPSK perché, come si vede nella figura seguente, il segnale informativo modula simultaneamente due portanti in quadratura, cioè due portanti sinusoidali sfasate fra loro di 90°.



Modulatore 4-PSK o Q-PSK

La modulazione 4PSK (di cui, tipicamente, si serve la tecnica CDMA, della quale parleremo tra poco) si svolge nel modo di seguito descritto.

Il segnale informativo originario, il quale è un segnale digitale, viene suddiviso in coppie di bit dette dibit.

I due bit del dibit vengono applicati simultaneamente al modulatore: ciascun bit va all'ingresso di uno dei due codificatori presenti nel modulatore.

Tutti i primi bit di ciascun dibit vengono applicati, uno dopo l'altro, a un codificatore. Il segnale di uscita del codificatore modula una portante sinusoidale

Tutti i secondi bit di ciascun dibit sono applicati, uno dopo l'altro, all'altro codificatore che modula una seconda portante sinusoidale, sfasata di 90° rispetto alla prima.

Una volta modulate, le due portanti in quadratura vengono sommate dando luogo a un segnale modulato in fase a 4 stati, ossia a un segnale modulato che può assumere quattro differenti valori di fase.

MODULAZIONE QAM-TCM

Le modulazioni QAM-TCM sono utilizzate nei **modem** per trasmissione dati su linea telefonica commutata, nella **Tv via satellite** e nei **ponti radio**.

Il numero di stati di modulazione ottenibile è di 256 e oltre.

TIPI DI MODULAZIONE QAM TCM

16 QAM

I bit del segnale dati vengono raggruppati in QUADRIBIT e i livelli di modulazione sono $2^4=16$.

Se la velocità di trasmissione è 9600 bit/s, la velocità di modulazione è di $9600/4=2400$ baud.

Riguardo al significato dei bit del quadribit vi sono più possibilità:

1. Il quadribit è suddiviso in due dibit e il diagramma a costellazione degli stati di modulazione è a reticolo. Il dibit più significativo codifica il quadrante del reticolo, il bit meno significativo indica la posizione del punto all'interno del quadrante; questo è lo standard V32 per modem a 9600 bit/s ;
2. Il bit più significativo indica l'ampiezza, gli altri tre la fase (a ciascuna delle 8 possibili fasi sono associati due valori di ampiezza); questo tipo di modulazione è detta anche QPSK (PSK-QAM) ed è prevista dallo standard V.29; il diagramma di costellazione si sviluppa su circonferenze concentriche

32 TCM o 32 QAM TRELLIS

I bit del segnale dati vengono raggruppati in QUADRIBIT per cui, se la velocità di trasmissione è 9600 bit/s, la velocità di modulazione è di $9600/4=2400$ baud.

Però, ad ogni quadribit, viene aggiunto un quinto bit (determinato a partire dai primi due bit del quadribit stesso). Si ottengono così $2^5=32$ livelli di modulazione, grazie a questo quinto bit che viene detto “di ridondanza” e che ha lo scopo di ridurre gli errori.

Con l'aumento degli stati di modulazione infatti si limita la possibilità di avere errori nella transizione fra due stati successivi.

64 TCM o 64 QAM TRELLIS

Non usata in nessuno standard per modem

128 TCM o 128 QAM TRELLIS

I bit del segnale dati vengono raggruppati in SIXBIT (gruppetti di 6 bit) per cui, se la velocità di trasmissione è 14400 bit/s, la velocità di modulazione è di $14400/6=2400$ baud.

Però, ad ogni sixbit, viene aggiunto un settimo bit (determinato a partire dai primi due bit del sixbit stesso) si ottengono così $2^7=128$ livelli di modulazione.

Questa modulazione è usata nello standard V.32bis.

NOTA

Velocità normalizzate , in **bit/s**, per modem fonici e in banda base:

300, 600, 1200, 2400, 3600, 4800

14400, 19200, 28800, 48000, 56000, 64000, 7200, 96000

112000, 128000, 144000.

SCM (Signal Code Modulation)

CARATTERISTICHE

- Concepita per la BANDA MILLIMETRICA MMW, MilliMeter Wave: (20÷42) GHz
- Adatta per l'integrazione fra reti cablate (wired) e reti senza fili (wireless)
- È una tecnica mista analogico-digitale che RIMODULA i segnali QPSK (segnali analogici adatti a viaggiare su cavo) in modo da renderli adatti alla trasmissione via radio (wireless) mediante portanti allocate nella banda delle onde millimetriche

PREGI

- Unisce i vantaggi della modulazione analogica (maggiore facilità di trasmissione) a quelli della modulazione DIGITALE (maggiore possibilità di RIDURRE gli ERRORI)

MODULAZIONE GFSK (FSK gaussiana)

Utilizzata da Bluetooth, e dai cordless a standard DECT, la **GFSK (FSK gaussiana)** è una variante della modulazione **FSK** a spostamento di **frequenza** (Frequency Shift Keying), usata per anche per alcuni modem fonici. La **GFSK** è anche simile alla GMSK adottata nel GSM.

MODULAZIONE MSK

MSK sta per **Minimum Shift Keying** cioè “**Minimo Spostamento**” ed è utilizzata al posto della FSK quando si vuole evitare l’allargamento dello spettro, tipico appunto della FSK.

In realtà la MSK può essere vista:

- o come caso particolare della modulazione a spostamento di frequenza FSK
- o come caso particolare della modulazione a spostamento di fase PSK

Se vediamo la MSK come variante della PSK (ed è ciò che faremo), possiamo dire che il segnale modulato in MSK viene creato in base alla seguente legge:

- negli intervalli di tempo (t_{bit} cioè tempi di bit) nei quali il segnale modulante informativo vale **ZERO**, il segnale modulato PSK presenta uno **sfasamento negativo**, cioè un **ritardo**, di **-90°** rispetto alla **portante sinusoidale**; il segnale modulato PSK coincide quindi con la portante sinusoidale, ma è sfasato di -90° rispetto ad essa
- negli intervalli di tempo nei quali il segnale modulante informativo vale **UNO**, il segnale modulato PSK presenta uno sfasamento positivo, cioè un anticipo, di +90°; il segnale modulato PSK coincide quindi con la portante sinusoidale, ma è sfasato di +90° rispetto ad essa.

Questa legge è valida a patto che la deviazione di frequenza Δf (del segnale modulato rispetto alla portante) sia legata al tempo di bit dalla seguente relazione:

$$\Delta f = \frac{1}{4 \cdot T_{bit}}$$

MODULAZIONE GMSK (MSK GAUSSIANA)

Modulazione a Minimo Spostamento di Frequenza con filtraggio Gaussiano

Questa modulazione prevede che la **sequenza di bit della sorgente** (cioè il segnale modulante informativo) venga filtrato **prima** di essere applicato al modulatore MSK. La “G” che compare nell’acronimo GMSK sta per “gaussiana”, in quanto la risposta in frequenza del filtro utilizzato ha un andamento gaussiano, cioè a campana. Questo accorgimento riduce la probabilità di interferenza fra canali adiacenti perché elimina i lobi secondari dello spettro del segnale modulante informativo.

SPREAD SPECTRUM O “SPETTRO ESPANSO”

E' un **insieme di tecniche** di comunicazione (comprendente DSSS e FHSS) che prevedono l'**espansione dello spettro** di un segnale **su una banda W molto più larga di quella originaria**, ottenendo così la distribuzione della potenza del segnale di partenza su un intervallo di frequenza più ampio e quindi una **densità spettrale di potenza molto più bassa**, paragonabile a quella del **rumore bianco**.

I vantaggi offerti dalle tecniche “spread spectrum” sono:

1. possibilità di **utilizzo della stessa banda** da parte di **più utenti**, purché questi utilizzino codici di spreading uno diverso dall'altro.
Con “spread spectrum”, infatti, l'ingresso di un nuovo utente nella banda non interferisce con gli altri utenti, ma contribuisce solo ad aumentare il rumore di fondo, inconveniente al quale si rimedia con le tecniche di correzione diretta degli errori (FEC)
2. possibilità di **operare in bande non soggette a licenza** e in **bande già utilizzate da sistemi tradizionali** di telecomunicazione, **senza interferire** con altri utenti (ma aumentando solamente il rumore di fondo). Per esempio, le reti a corto raggio Wi-Fi e Bluetooth lavorano nella banda senza licenza ISM (Industriale-Scientifica-Medica), mentre il GPS (Sistema di Posizionamento Globale) opera in bande utilizzate da sistemi tradizionali.
3. **Ridurre la probabilità di essere intercettati**, grazie al fatto che il **segnale trasmesso** a spettro espanso è **confondibile con il rumore**. Con l'espansione di spettro infatti, la densità spettrale di potenza del segnale ha le stesse caratteristiche di quella del rumore bianco. Ricordiamo che il rumore bianco (che nel tempo è assimilabile a un impulso ideale) ha una densità di potenza costante per ogni frequenza.
4. **Diminuzione della probabilità di errore** a parità di bit rate
5. **Aumento del bit rate** a parità di probabilità di errore.

DSSS - Direct Sequence Spread Spectrum - (Spettro espanso mediante sequenza diretta)

Usata da Zigbee, DSSS è, come FHSS, una tecnica di comunicazione a spettro espanso. In particolare DSSS è la tecnica per la quale **lo spettro di un segnale digitale** viene **allargato** mediante la **moltiplicazione (nel tempo)** dei bit del segnale stesso **con i simboli (chip) di una sequenza digitale pseudocasuale**, detta “sequenza PN” (PseudoNoise sequence).

L’effetto di questa moltiplicazione nel tempo è l’espansione dello spettro su una banda più larga e quindi la distribuzione della potenza su un intervallo di frequenza più ampio (intervallo che è proprio la banda della sequenza PN)

La durata **T_c** dei singoli **chip** deve essere **molto minore** della durata **T_b** del singolo bit

s e g n a l e

bit di durata T_b	bit	bit	bit	bit	bit
---	-----	-----	-----	-----	-----

PN o (sequenza di codice pseudocasuale)

c	c	c	c								c
h	h	h	h								h
i	i	i	i								i
p	p	p	p								p

Ogni chip ha durata T_c

Una volta effettuata l’espansione di spettro, la tecnica DSSS richiede che il **segnale espanso**, per poter essere trasmesso via radio, **moduli una portante analogica**, tipicamente con tecnica a spostamento di fase **PSK** (o comunque con una opportuna modulazione con portante analogica e con segnale informativo modulante digitale)

I fondamentali **parametri** dell’espansione di spettro per sequenza diretta sono:

- **S_f** = spreading factor = **fattore di espansione** =
= **numero di chip per ciascun bit**
- **R_c** = chip rate = **numero di chip al secondo** = $\frac{1}{T_c}$

FHSS (Spettro espanso con salto di frequenza portante)

La tecnica “**Frequency Hopping Spread spectrum**” (spettro espanso con salto di frequenza portante) è usata da **Bluetooth**, da **Home RF** e dal **GPS**.

FHSS è, come DSSS, una tecnica di comunicazione a spettro espanso ed è la tecnica per la quale:

- **la banda W di espansione** di un sistema (cioè la banda sulla quale è stata espansa -mediante spreading- la banda originaria di un segnale) viene **suddivisa** in un certo numero di **canali radio**, ciascuno dei quali ha la larghezza di banda di una “normale” modulazione (per es. PSK)
- per **ciascun canale** viene creata una **frequenza portante**
- il sistema cambia, **a intervalli regolari di tempo Δt** , la frequenza portante utilizzata, effettuando così un **salto (hop) di frequenza**
- **per ogni intervallo di tempo Δt** il sistema **trasmette un burst** (una raffica) di n bit, dopodiché, **effettuato il salto successivo**, trasmette **un altro burst su una frequenza portante diversa**.

In sintesi, FHSS è una tecnica che prevede l'utilizzazione, mediante salti di frequenza, di un certo numero di frequenze portanti. Ciascuna portante è utilizzata per un limitato intervallo di tempo e si trova in una porzione della banda espansa W sufficiente a contenere lo spettro di un segnale modulato.

I fondamentali **parametri** dell'espansione di spettro a salto di frequenza sono:

- x (hop/s) = **frequenza di hopping** =
= **numero di salti di frequenza al secondo**
- $\Delta t = \frac{1}{x}$ = tempo di burst = tempo di utilizzazione continuativa di una frequenza portante
- **SEQUENZA di hopping** = **successione delle frequenze portanti utilizzate dal sistema**.

La tecnica FHSS **migliora l'immunità alle interferenze** perché **un eventuale segnale interferente a banda stretta può disturbare la comunicazione solo per un intervallo di tempo molto breve**, cioè per il tempo di utilizzazione della portante che si trova nella banda occupata dal disturbo.

La tecnica di frequency hopping è anche usata dal GSM, ma con un basso numero di portanti (15 al massimo), una bassa frequenza di hopping (217 hop/s) e un

$$\Delta t = \frac{1}{217} = 4,6m$$

(corrispondente a un cambio di frequenza ogni 4,6 ms).

La tecnica del salto di frequenza prevede dunque l'effettuazione, per un certo numero di volte al secondo (**1600 in Bluetooth**), di un “salto” di frequenza “portante” e l'utilizzazione di un determinato numero di **valori di frequenza (79 in Bluetooth)**. Le **prerogative** del frequency hopping sono:

1. quando su una certa frequenza sono presenti **interferenze**, **questa frequenza non viene assegnata** a nessuna portante
2. **maggiore protezione dalle intercettazioni** e dagli accessi fraudolenti nella rete
3. **eliminazione dei disturbi a radiofrequenza** (per esempio i disturbi dovuti ai forni a microonde)
4. possibilità di utilizzare **tecnica di correzione diretta degli errori** (FEC, Forward Error Correction).

CDMA (Accesso Multiplo a Divisione di Codice)

CDMA è una tecnica di accesso multiplo basata sull'espansione dello spettro, cioè su "spread spectrum", per la quale più utenti (riconoscibili ciascuno in modo univoco da un proprio codice) possono utilizzare la stessa portante nello stesso intervallo di tempo (timeslot).

ORTOGONALITA' DEI CODICI PN

Per la realizzazione dell'accesso multiplo a divisione di codice è necessario che tutte le sequenze pseudocasuali PN utilizzate per l'espansione dello spettro siano ORTOGONALI tra loro.

Ad ogni utente di una portante è assegnato un differente codice PN, ortogonale ai codici assegnati a tutti gli altri. Ciò rende possibile la condivisione della stessa portante da parte di molti utenti senza che i segnali relativi ai vari utenti possano interferire in alcun modo fra loro.

Codici ortogonali fra loro:

- **non interferiscono** uno con l'altro
- **non possono essere sostituiti** uno con l'altro

e perciò **identificano univocamente l'utente cui sono assegnati.**

L'ortogonalità dei codici consente, alla tecnica di accesso **CDMA**, di **far utilizzare a più utenti la stessa portante** e lo **stesso timeslot** (intervallo di tempo) **senza reciproche interferenze.**

Possiamo dire che due segnali codificati ortogonalmente, anche se occupano la stessa banda, non interferiscono uno con l'altro, si "ignorano" reciprocamente senza disturbarsi fra loro, analogamente a quanto avviene per due onde elettromagnetiche che abbiano la stessa frequenza, ma che siano polarizzate⁶ una verticalmente e l'altra orizzontalmente.

Le due onde non interferiscono fra loro e possono essere captate singolarmente da due antenne riceventi con la giusta polarizzazione.

L'ortogonalità dei codici è indispensabile per la demodulazione.

La **ricezione** di un segnale a spettro espanso (che è simile al rumore bianco), relativo a un certo utente, richiede un'operazione inversa all'espansione: il **despreading**. Nel

⁶ Un'onda è polarizzata verticalmente quando il vettore campo ELETTRICO è PERPENDICOLARE rispetto al suolo. L'onda polarizzata verticalmente deve essere emessa e ricevuta da antenne VERTICALI rispetto al suolo.

despreading il segnale espanso viene moltiplicato con una sequenza PN esattamente uguale a quella usata in trasmissione.

Se la sequenza PN non è la stessa usata in trasmissione, il segnale estratto mediante despreading non sarà quello originario, ma solo rumore.

CARATTERISTICHE di CDMA

- E' una tecnica di accesso multiplo in base alla quale molti segnali vengono trasmessi CONTEMPORANEAMENTE, in parallelo, a più utenti
- I segnali emessi sono SEPARATI in FASE e in CORRELAZIONE, usando DIFFERENTI CODIFICHE per i simboli.

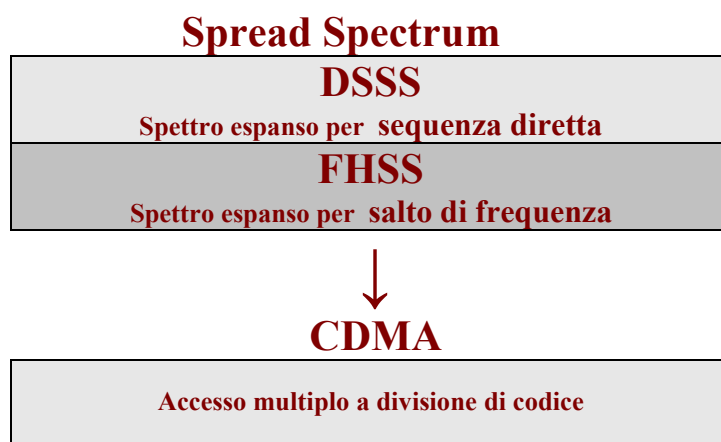
PREGI

I segnali risultano:

- più IMMUNI al RUMORE e all'INTERFERENZA INTERSIMBOLICA⁷ (ISI, InterSymbol Interference)

INCONVENIENTI

- alto TASSO di ERRORE (BER, Bit Error Rate)



⁷ L'ISI consiste nell'allargamento, durante la trasmissione, della durata dei singoli impulsi, i quali tendono a sovrapporsi. La forma originaria del segnale subisce pertanto un'alterazione che può comportare l'erronea interpretazione del messaggio associato alla sequenza di impulsi.

TDMA (Accesso Multiplo a Divisione di Tempo)⁸

CARATTERISTICHE

- E' una tecnica di multiplazione (accesso multiplo) in base alla quale molti segnali vengono trasmessi IN DIFFERENTI INTERVALLI DI TEMPO (TIME SLOT) nella stessa banda di un canale trasmissivo.
- In ciascun intervallo di tempo è trasmesso UN SOLO SEGNALE, che può occupare anche l'INTERA BANDA del CANALE

PREGI

- Piccolo TASSO di ERRORE (BER, Bit Error Rate)

INCONVENIENTI

- forte INTERFERENZA INTERSIMBOLICA (ISI, InterSymbol Interference) per elevate velocità di trasmissione e per lunghe distanze.

⁸ Esiste anche un “Accesso Multiplo a Divisione di Frequenza ” FDMA. In questa tecnica la banda a disposizione del sistema viene suddivisa in n canali radio che rappresentano le “risorse radio” Ciascun canale è centrato su una frequenza portante ed è unidirezionale. Per ottenere una comunicazione bidirezionale bisogna quindi utilizzare due canali

CSMA/CD - Carrier Sense Multiple Access with Collision Detection – (Accesso Multiplo con rivelazione di portante e rilevamento di collisione)

CSMA è una tecnica di accesso multiplo, cioè una tecnica che permette a più utenti (per esempio a più PC collegati in rete) di utilizzare lo stesso mezzo trasmissivo (per esempio lo stesso bus), cioè di “condividere” il mezzo trasmissivo in base a determinate regole.

La tecnica di accesso multiplo utilizzata per le reti Ethernet (le quali adottano la topologia a stella) è **CSMA/CD**

CSMA/CD è una tecnica:

- “**non** deterministica”, dal momento che non permette di calcolare a priori i tempi di accesso al bus
- “**a contesa**” perché più stazioni entrano in competizione per l'utilizzazione del bus.

Il computer, che vuole accedere all'unico cavo della rete, invia i suoi dati solo dopo aver accertato che sulla rete non siano già presenti altri segnali.

In presenza di altri segnali, il computer, per evitare collisioni di dati, inizia (o ricomincia) a trasmettere soltanto dopo un tempo di durata casuale.

Se la collisione si verifica, vengono **persi i dati di entrambi i computer** che hanno trasmesso simultaneamente.

CSMA/CD rende possibile la realizzazione di diversi “sottotipi” di rete Ethernet, caratterizzati da differenti portanti fisici (cavi metallici ecc.) e quindi da differenti velocità di trasmissione e differenti distanze massime fra stazioni

CSMA/CD risulta non necessario se i ripetitori di rete (hub) sono sostituiti da apparati più sofisticati (switch) che evitano le collisioni di dati.

CSMA/CA - Carrier Sense Multiple Access with Collision Avoidance⁹- (Accesso Multiplo con rivelazione di portante e prevenzione delle collisioni)

Nelle reti **wireless**, in cui le stazioni trasmettono e ricevono sullo stesso canale, **non è possibile rilevare le collisioni come nel CSMA/CD** (utilizzata da Ethernet), dato che è difficile e costoso costruire un apparato che possa contemporaneamente trasmettere ed ascoltare, pertanto le collisioni non possono essere rilevate e devono essere evitate.

Pertanto, nel momento in cui una stazione vuole tentare una trasmissione, testa il canale. Se il canale è occupato, la stazione attiva un conteggio di durata casuale (durata detta “tempo di *back off*”), che viene effettuato solo durante i periodi di inattività del canale.

Quando il conteggio è terminato, la stazione ricontrolla il canale. Se il canale risulta libero lo “prenota” ed attende per un certo intervallo di tempo. Se il canale continua ad essere libero (cioè non ci sono state altre prenotazioni) la stazione trasmette.

CSMA/CA, usata anche dal bus CAN (Computer Area Network) differisce inoltre da CSMA/CD, per il fatto che, **in caso di collisione** (cioè nel caso in cui due stazioni, nonostante tutti gli accorgimenti, trasmettano simultaneamente il loro messaggio), una delle due stazioni **smette di trasmettere** (in base a un meccanismo di priorità) **senza che il messaggio dell’altra stazione venga danneggiato**.

In CSMA/CD invece, in caso di trasmissione concomitante, la collisione viene rilevata da entrambe le stazioni, ma i pacchetti di dati vanno perduti.

⁹ Da “to avoid” = evitare

MDMA (Accesso Multiplo Multi Dimensionale) (tappa successiva a TDMA e a CDMA)

CARATTERISTICHE

- molti segnali trasmessi
- per ogni velocità di trasmissione si impiega la massima banda disponibile
- il segnale è modulato in tutti i modi possibili (ampiezza, frequenza, fase)

PREGI

- maggiore immunità al rumore rispetto a TDMA e a CDMA
- minor TASSO di ERRORE (BER, Bit Error Rate)

UTILIZZAZIONI

- LAN Wireless per Bluetooth
- Reti terrestri via cavo

UWB

La tecnica UWB (Ultra Wide Band, cioè banda ultralarga), è utilizzata dai sistemi **Wi-Max** su rete senza fili, per la realizzazione di **reti a corto raggio (LAN)**. E' anche usata in sistemi di **radiolocalizzazione**.

UWB è basata sulla **ricetrasmisione di impulsi**, formati, ciascuno, **da pochi periodi di una portante a radiofrequenza**, di **brevissima durata temporale** (decine di ps ÷ alcuni ns) e quindi con spettro che occupa una banda molto larga.

PREGI

- Trasmissioni **difficilmente intercettabili** perché simili al rumore di fondo per effetto dalla bassa intensità
- **Non interferisce** con sistemi già esistenti
- **Basso consumo.**

ONDE CONVOGLIATE (powerline)

La comunicazione per onde convogliate (Power Line Communication, PLC) o **modulazione su rete elettrica** è una tecnica che permette di **trasferire segnali** (che contengono informazione) **attraverso la rete elettrica**, la cui utilizzazione naturale è il trasporto di potenza elettrica sotto forma di tensioni e correnti.

La comunicazione per onde convogliate è attualmente usata anche per la **trasmissione-dati** e la **domotica**, ma non è una tecnica nuovissima. E' stata infatti utilizzata, prima dell'avvento della telefonia mobile, per comunicare con i treni in marcia, attraverso le linee di alimentazione elettrica dei convogli.

TERMINOLOGIA DEI SITI WEB

(alcune note terminologiche in ordine alfabetico)

A

APPLET

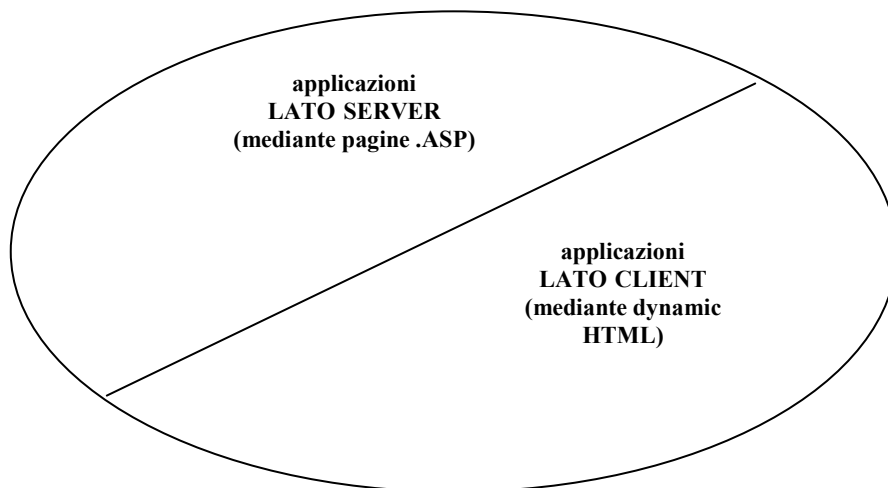
Programma compilato, scritto in linguaggio Java, che ha estensione **.class**

Gli applet possono essere utilizzati per aggiungere funzionalità avanzate a un sito web, attraverso la realizzazione dei cosiddetti “componenti di Frontpage”.

APPLICAZIONE WEB

Un'applicazione web è un **insieme di programmi** che consente di **effettuare l'input**, **l'elaborazione** e **l'output dei risultati** in **maniera distribuita** su **più computer** di una rete.

L'applicazione web consente, in sostanza, all'utente (client) di interagire con i server web (inviando dati e non soltanto ricevendone).



ASP

PAGINA ASP (Pagina Attiva lato-Server)

La **pagina ASP** è una pagina di testo che contiene, oltre alle righe di testo e ai tag (marcatori o etichette) HTML, anche **istruzioni di programma**¹⁰ che vengono **eseguite dal server web prima** che la pagina sia inviata al browser (cioè a Internet Explorer o altro) e quindi prima che sia visualizzata dall'utente.

Le pagine ASP permettono dunque di realizzare una “**applicazione web distribuita**”, cioè un **insieme di programmi** in grado di **operare su più computer di internet** (cioè in grado di eseguire, su più computer di internet, operazioni di input, elaborazione e output di dati).

In altri termini la pagina ASP fa in modo che venga **presentata all'utente l'elaborazione**, da parte del **server**, delle righe di programma presenti in essa (elaborazione fornita sotto forma di **documento HTML**).

Proprio da questo fatto deriva il nome di “**pagina attiva lato-server**” o “Active Server Page”

Possiamo anche dire che la pagina ASP definisce la **risposta di un server web a un browser**, in quanto contiene le istruzioni che il server dovrà eseguire.

Le **istruzioni di programma** sono **racchiuse** fra i **tag** `<%` e `%>` e sono chiamate **comandi di scripting lato-server**.

I linguaggi di scripting (cioè i linguaggi per la programmazione di script) più diffusi sono

- Javascript
- VBScript (Visual Basic Script).

La tecnica ASP si serve dei protocolli TCP/IP e in particolare di http.

La stessa funzione delle pagine ASP può essere svolta dall'inserimento, nelle pagine HTML, di codice (istruzioni di programma) in **linguaggio PHP** (che, attualmente sta per “**Hypertext PreProcessor**”, mentre in passato l'acronimo si scioglieva in “Personal Home Page”)

Anche le PAGINE HTML DINAMICHE contengono righe di programma che però, in questo caso (contrariamente a quanto avviene per le pagine ASP), sono eseguite dal browser e non preventivamente dal server.

¹⁰ (Possiamo perciò dire che la pagina .ASP è un documento di testo in linguaggio script)

PAGINE ASPX O PAGINE ASP.NET

La **pagine ASPX**, dette anche **ASP.NET** o ancora **WEB FORM** sono una nuova generazione di ACTIVE SERVER PAGE, cioè una **nuova generazione di pagine .ASP**.

La finalità delle pagine ASP.NET sono:

- la creazione di siti web veloci, sicuri e versatili
- la comunicazione fra differenti applicazioni, mediante uno **standard di programmazione unificato**¹¹ per la creazione di software che possa essere distribuito **su qualunque supporto elettronico** (PC, PDA, smartphones)

Gli ambienti web realizzati con ASP.NET, e quindi con pagine ASPX, presentano, rispetto a quelli realizzati con pagine ASP3.0 (e che presentano pagine attive lato server di tipo ASP) le caratteristiche seguenti:

- uso di linguaggi orientati agli oggetti
- compilazione (anziché interpretazione) del codice
- DEBUG e TRACE per un maggior controllo del codice generato
- LOGICA SERVER: tutti i controlli che prima erano effettuati lato client (es: validazione dei dati in un form) sono effettuati lato server con maggiore flessibilità e sicurezza
- interfaccia ADO.NET invece di ADO: la nuova interfaccia ADO.NET di collegamento con le basi di dati permette di **legare le basi di dati direttamente al web** in maniera versatile e robusta, il che è molto utile per la creazione di motori di ricerca e report (Ricordiamo che con ADO l'accesso ai dati è solo sequenziale)

¹¹ Indicato come MICROSOFT.NET FRAMEWORK o anche come DOTNET FRAMEWORK

B

BOOKMARK

E' il cosiddetto “segnalibro” degli ipertesti. E' anche detto “**àncora**” o “ancoraggio” ed è la parola del testo o l'elemento grafico (immagine) che rimanda al link.

C

CLIENT FTP

Software per realizzare una connessione FTP.

Per esempio:

- FileZilla
- Smart FTP Client

COOKIE

Il cookie è un file di testo nel quale sono memorizzati i dati che l'operatore (cioè un utente connesso a internet con il suo computer client) digita nel browser, per esempio in Internet Explorer.

I cookies vengono memorizzati nell'hard disk del client, cioè nell'hard disk del computer dell'utente connesso a internet.

D

D-HTML (documenti Dynamic-HTML)

I documenti Dynamic-HTML sono **documenti** di tipo **client-side** nel senso che le **istruzioni** sono **eseguite** non dal server, ma **direttamente dal browser**, dopo che la pagina è stata inviata dal computer server al computer client.

Il documento DHTML, serve quindi a realizzare applicazioni web distribuite, come le pagine .ASP, queste ultime, però, sono di tipo server-side.

DNS (Domain Name Server)

Con il termine “DNS” viene **comunemente** indicato un **nome a dominio**.

Precisamente per DNS si intende un **sistema di codifica** oppure un **server** (name server) che **associa** un **nome a dominio** a un **indirizzo IP numerico**.

E

EDITOR VISUALI o SISTEMI-AUTORE o WEB EDITOR

Per “editor” intendiamo i software mediante i quali si possono scrivere le pagine che costituiscono un sito web e si può creare la struttura di collegamenti fra queste pagine. Gli editor visuali sono pacchetti software anche indicati con l’acronimo WISIWIG (What You See Is What You Get), sono cioè editor del tipo “ciò che vedi è ciò che ottieni”, nel senso che ogni modifica effettuata dal programmatore produce immediatamente un risultato identico a quello che apparirebbe visualizzando la pagina con un browser (come Netscape o Internet Explorer)

Esempi di editor visuali sono:

- *Microsoft Expression Web Designer* (che è il nuovo nome di *Frontpage*)
- *Visual Studio 2005*
- *Flash*
- *Web Site Builder*

ESTENSIONI DEL SERVER DI FRONTPAGE

Le estensioni del server di Frontpage¹² sono un **insieme di programmi** aggiuntivi **installati nel server web**, cioè nel server di distribuzione.

Esse servono:

- per l’**esecuzione** dei cosiddetti “**componenti di Frontpage**”
- per l’**esecuzione** delle **pagine .ASP**
- per la **pubblicazione di un sito** mediante il protocollo **HTTP**
- per la **gestione remota dei siti**

I programmi indicati come “estensioni del server di Frontpage” utilizzano:

- il protocollo **HTTP**
- le interfacce di programmazione standard del WEB e cioè CGI (Common Gateway Interface) e ISAPI (Internet Server Application Programming Interface).

¹² Attualmente esiste la versione “rinnovata” di Frontpage che si chiama **Microsoft Web Expression**

F

FOGLI DI STILE O CCS (Cascade Style Sheets)

Il CCS è un **documento di testo senza formattazione**, con estensione “.CCS” che contiene **definizioni di formato** (per esempio colore dei caratteri, allineamenti del testo ecc.), le quali possono essere usate per la **formattazione “simultanea”** di un **insieme** di pagine web, o, in generale, di uno o più elementi HTML.

H

HOSTING

E’ il complesso dei servizi informatici (accesso al server, spazio di memoria sull’hard disk del server, caselle di posta, eventuali software ecc.) che il maintainer o il provider offrono all’utente

HYPERLINK

(vedi il sinonimo “**LINK**”)

HTML

(HyperText Markup Language)

E’ un linguaggio di markup (annotazione) basato su tag (marcatori) e utilizzato per la **creazione, l’elaborazione e la formattazione dei documenti da diffondere sul WWW** (World Wide Web).

I **documenti HTML** infatti possono essere **interpretati e riprodotti** dai **browser**¹³ di Internet.

La versione **più recente** è **HTML4** che supporta ***fogli di stile (CCS)*** e linguaggi di **scripting**.

HTML è più limitato del più moderno ***XML*** ed è **un’applicazione** dell’ancora più generale **linguaggio SGML**.

¹³ Internet Explorer, Netscape ecc.

XML (eXtensible Markup Language) è uno standard W3C e consente:

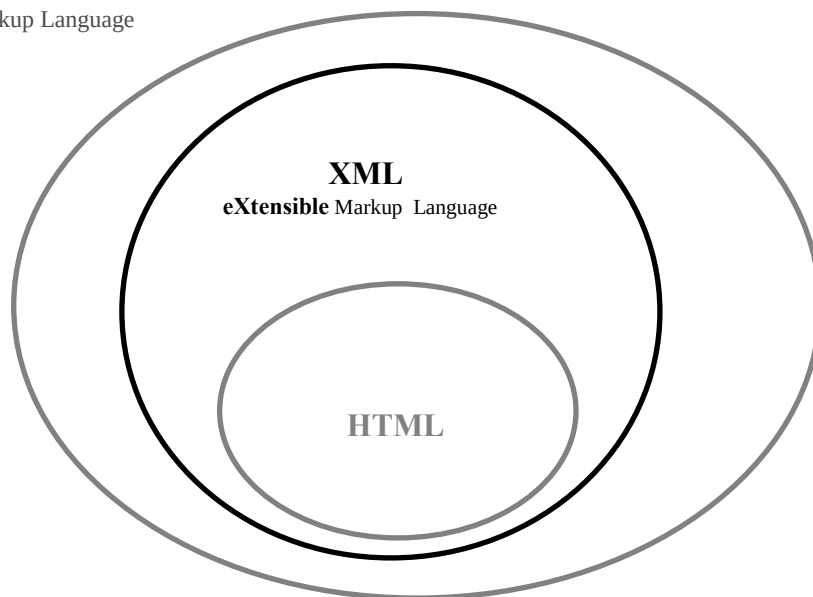
- di creare tag personalizzati
- di organizzare e presentare le informazioni con maggiore flessibilità

SGML (Standard **Generalized** Markup Language) utilizza tag per contrassegnare, nei documenti, non solo testi, ma anche grafica, così da indicare ai browser come visualizzare questi elementi e come rispondere alle azioni dell'utente (mediante il mouse e l'attivazione di collegamenti).

SGML è stato standardizzato nel 1986 dall'ISO al fine di:

- permettere la creazione di documenti indipendenti dalla piattaforma e dalle applicazioni
- includere nel documento informazioni relative al formato, alle informazioni collegate, e a un indice analitico
- creare tag personalizzati

SGML
Standard **Generalized**
Markup Language



I

IIS (Internet Information Service) E WEB SERVER

“Microsoft IIS” è un **software “server web”**, cioè un’**applicazione software** (ossia un **programma**), presente nei **computer server web** e ha la funzione di **far comunicare il server web con i browser**¹⁴ **mediante il protocollo http**.

IIS è quindi un programma “**web server**” ed è un’applicazione presente nelle versioni “server” e “professional” dei sistemi operativi:

“Microsoft **Windows NT 2000 server**”.

“Microsoft **Windows XP**”

Il software “web server” analogo a IIS, ma relativo ai sistemi operativi *UNIX e LINUX*, è **APACHE**.

(Si veda anche alla voce “web server”)

L

LINK (O HYPERLINK)

Un link è un collegamento ipertestuale, cioè il collegamento (in un determinato punto) di una pagina a un altro oggetto, come un’immagine, un filmato, un file audio o un file contenente un’altra pagina web o un indirizzo email.

L’oggetto che viene collegato alla pagina (per es. l’immagine) deve avere un suo indirizzo.

Il punto di partenza del link, per esempio una porzione di testo, viene chiamato “**àncora**” o “**ancoraggio**”.

LISTE DI DISTRIBUZIONE MULTIPLE

Sono **gruppi di indirizzi** ai quali si possano **inviare lettere in una sola volta** (per esempio tutti gli insegnanti di una classe).

¹⁴ I browser, come Internet Explorer e Netscape, possono essere definiti “software client”

N

NEWSLETTER MANAGER

Servizio di un provider (es: Register) per inviare email

P

PAGINA ASP

(Pagina Attiva lato-Server)

Vedi ASP.

R

RENDERING

Per “rendering” si intende la “**resa**” di una pagina **nei vari browser**, ossia **il modo** in cui la pagina viene visualizzata in Internet Explorer in Netscape ecc.

S

SCRIPT

Codice sorgente (cioè righe di programma scritte dal programmatore) inserito in un documento destinato al web (cioè alla pubblicazione su internet).

SGML

(Standard Generalized Markup Language)

SGML (Standard **Generalized** Markup Language) può essere definito come una “estensione” di HTML, in quanto include HTML (e anche XML).

SGML è quindi un linguaggio di markup (annotazione) basato su tag (marcatori) e utilizzato per la **creazione, l’elaborazione e la formattazione dei documenti da diffondere sul WWW** (World Wide Web).

SGML utilizza tag per contrassegnare, nei documenti, non solo testi, ma anche grafica, così da indicare ai browser come visualizzare questi elementi e come rispondere alle azioni dell’utente (mediante il mouse e l’attivazione di collegamenti).

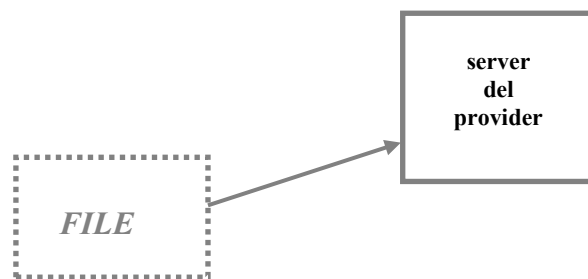
SGML è stato standardizzato nel 1986 dall’ISO al fine di:

- permettere la creazione di documenti indipendenti dalla piattaforma e dalle applicazioni
- includere nel documento informazioni relative al formato, alle informazioni collegate, e a un indice analitico
- creare tag personalizzati

(vedi anche HTML e XML).

STREAMING

E' lo **scaricamento di un file** (per esempio di un file video), fatto in modo che il **video** possa essere visto *prima* che il **caricamento sia completato**.



UPLOAD:
il file viene CARICATO nel SITO

SCRIPT

Codice sorgente (cioè righe di programma scritte dal programmatore) inserito in un documento destinati al web (cioè alla pubblicazione su internet).

T

TAG (Marcatore o Etichetta)

I tag sono delle “parole-chiave” del linguaggio HTML.

Possono essere costituiti anche da una sola lettera e sono sempre racchiusi¹⁵ fra i simboli “<” e “>”.

I tag servono a specificare qualcosa o a fornire indicazioni al browser che presenterà la pagina web all’utente.

Per esempio possono specificare il formato (grassetto, corsivo, maiuscolo ecc.) che deve avere il testo.

Oppure possono indicare l’inserimento di un’immagine (tag “IMG”).

I tag sono quindi quegli elementi del linguaggio HTML che precedono o seguono le istruzioni e gli elementi (es. un testo) del linguaggio.

Esistono tag “singoli” (es: META e IMG) e tag “doppi” (es:<BODY>....</BODY>).

¹⁵ da soli (per esempio: <HEAD>) o insieme ad altri elementi (per esempio: <IMG SRC “nomefoto.jpeg”)

TAG: ALCUNI ESEMPI

NOME del TAG	Significato	Singolo o doppio	Attributi
<P>...</P>	Paragraph (Paragrafo)	Doppio	ALIGN (allineamento)

	BReak (Ritorno a capo o interruzione di linea)	Singolo	
<H1>...</H1> Il pedice indica che il carattere è più grande	Heading (Titolo o intestazione)	Doppio	
...	Font	Doppio	face, size, color
<blockquote>...</blockquote>	evidenziazione del testo	doppio	
<code>...</code>	testo con carattere "macchina per scrivere"	Doppio	
<tt>...</tt>			
...	Bold/ Grassetto	Doppio	
<I>...</I>	Italic/ corsivo	Doppio	
<U>...</U>	underlined/ Sottolineato	Doppio	
...	emphasize	Doppio	
 ...	emphasize	Doppio	
...	image / inserimento immagine	Singolo	SRC (Source=indirizzo del file immagine)
			WIDTH (lunghezza immagine)
			HEIGHT (altezza immagine)
			ALT (ALternative Text)
			BORDER (bordo immagine)
			ALIGN (allineamento)
			HSPACE e VSPACE (spazio vuoto attorno all'immagine)
<A>...	Anchor (àncora cioè punto di partenza d'un link)	Doppio	HREF / reference cioè punto di arrivo del link

TAG: ALCUNI ESEMPI

NOME del TAG	Significato		Attributi
<TABLE>...</TABLE>	Tabella	Doppio	Border BG color Background color, Bordercolor.
<TR>...</TR>	Table Row (singola riga di una tabella)	Doppio	Align Width
<TD>...</TD>	Tabel Date (singola colonna di una tabella)	Doppio	ALIGN
			COLSPAN (unione in orizzontale)
			ROWSPAN (unione in verticale)
			WIDTH
<TH>...</TH>	Tabel Heading (intestazione tabella)	Doppio	Align Width

TEMA (THEME)

E' un formato “uniforme”, comune a tutte le pagine di un sito web, che può essere definito attraverso un foglio di stile o CCS.



URL (Uniform Resource Locator)

E' l'**indirizzo** di un **file** (che contiene una pagina web, un'immagine o altro) o di un'**altra risorsa del web**.

Un esempio di URL è il seguente:

<http://www.azienda.com/servizi/market.htm>

dove:

<http://www.azienda.com/servizi/market.htm> è l'URL **assoluto**

[servizi/market.htm](#) è l'URL **relativo**.

W

WEB SERVER

Per SERVER WEB si può intendere sia un **software**, sia un **computer** nel quale sia installato il software omonimo.

In ogni caso è il computer che prende il nome dal software che vi è installato

SOFTWARE “WEB SERVER”

E’ un software che serve a fornire le pagine web (o meglio le informazioni che servono a visualizzarle) ai browser che ne facciano richiesta.

I programmi Web Server utilizzano il protocollo http.

In altri termini, il software server web **attende le richieste di pagine web** (da parte di browser come Mozilla, Opera, Internet Explorer), elabora queste richieste e **fornisce i dati in base ai quali i browser possono presentare le pagine web** agli utenti. Siccome serve a fornire i dati necessari alla visualizzazione di pagine web, il software server web è installato in computer sui quali risiedono siti web e quindi, tipicamente, nei computer “server web” dei provider e dei maintainer. Qualche volta il software server web è installato sui computer ove si trova la copia locale di un sito web (residente nello spazio web di un provider) per permettere di visualizzare l’anteprima del sito.

I software “server web” più diffusi sono:

- **APACHE**

- **MICROSOFT INTERNET INFORMATION SERVICE (IIS).**

In sintesi, il software “web server” è un programma che tipicamente si trova nei computer “server web” dei provider e dei maintainer.

Tuttavia, nei computer di persone o enti che non sono provider né maintainer, si trovano programmi web server (inclusi nel sistema operativo) con la sola funzione di testare i siti là presenti in copia locale.

Il software server web IIS è incluso nei seguenti sistemi operativi di Microsoft:

Windows 2000 Professional, Windows 2000 Server, Windows XP Professional, Windows XP Server

L’IIS presente nelle versioni “server” è un software completo, mentre quello incluso nelle versioni “professional” è incompleto e inadatto alla realizzazione di un computer “server web” professionale: è adatto solo per visualizzare e testare i siti in via di realizzazione (vedi alla voce IIS).

COMPUTER “WEB SERVER”

Un web server è un **file server**, cioè un **computer** che **dispone**:

- del **servizio www** per la **distribuzione di file** in una rete TCP/IP (rete internet o intranet) mediante il protocollo HTTP (Hyper Text Transfer Protocol)
- del **servizio FTP** (File Transfer Protocol) per il caricamento e lo scaricamento di file organizzati in cartelle.

Ogni server web *deve essere dotato*:

- del **protocollo TCP/IP**
- di un **indirizzo IP statico** (pubblico se opera in Internet, privato se lavora in Intranet)
- di un **nome di dominio** (pubblico se opera in Internet, privato se lavora in Intranet)
- di **dischi formattati con file system** (es: NTFS di Windows 2000) in grado di limitare l’accesso agli utenti non autorizzati.

WISIWIG (What You See Is What You Get)

Vedi EDITOR VISUALI.



XML **(HyperText Markup Language)**

E' un linguaggio di markup (annotazione) basato su tag (marcatori) e utilizzato per la **creazione, l'elaborazione e la formattazione dei documenti da diffondere sul WWW** (World Wide Web).

I **documenti HTML** infatti possono essere **interpretati e riprodotti dai browser**¹⁶ di Internet.

XML (eXtensible Markup Language) è più vasto e moderno di HTML ed è uno standard W3C che consente:

- di creare tag personalizzati
- di organizzare e presentare le informazioni con maggiore flessibilità rispetto ad HTML.

XML è “incluso” nel più generale SGML.

SGML (Standard **Generalized** Markup Language) utilizza tag per contrassegnare, nei documenti, non solo testi, ma anche grafica, così da indicare ai browser come visualizzare questi elementi e come rispondere alle azioni dell'utente (mediante il mouse e l'attivazione di collegamenti).

SGML è stato standardizzato nel 1986 dall'ISO al fine di

- permettere la creazione di documenti indipendenti dalla piattaforma e dalle applicazioni
- includere nel documento informazioni relative al formato, alle informazioni collegate, e a un indice analitico
- creare tag personalizzati.

(vedi anche HTML)

¹⁶ Internet Explorer, Netscape ecc.

CAPITOLO 43

INTRODUZIONE A DSL, ALLA TV DIGITALE TERRESTRE E ALLA RADIO DIGITALE

OBIETTIVI

- **Conoscere, per grandi linee, i principi di funzionamento dei sistemi in oggetto**
- **Conoscere, per grandi linee, le tecniche che sono alla base dei sistemi in oggetto**
- **Conoscere i principali standard**

PREREQUISITI

- **Basi di elettronica**
- **Analisi spettrale**
- **Modulazioni**

DSL

Con l'acronimo DSL (Digital Subscriber Line, Linea Digitale di utente), oppure con xDSL si indica un insieme di tecniche aventi lo scopo di permettere e ottimizzare l'**esercizio simultaneo della trasmissione dati** e della **fonìa** su doppino di rame.

Ogni terminale DSL è **perennemente collegato in rete**, quindi non avvengono più né la connessione (con la relativa operazione di dial-up, ossia di selezione), né la sconnessione.

Nella terminologia, la lettera "x" viene di volta in volta sostituita con le iniziali della specifica tecnica della famiglia DSL che interessa. Per esempio "ADSL" sta per "DSL Asimmetrica", in quanto questa particolare tecnica DSL (che è proprio quella della quale ci occuperemo in questa sede) prevede che la **velocità** dei dati che l'**utente riceve**, velocità di **downstream**, sia **maggiore** della velocità dei dati che l'utente trasmette (upstream).

xDSL è un acronimo generico per definire una famiglia di servizi dedicati, la cui lettera "x" sta a significare:

- *ADSL Asymmetric Digital Subscriber Line: 1,5 Mbps-384 kbps/384-128 kbps*
- *HDSL High-bit-rate Digital Subscriber Line: 1.5 Mbps/1.5 Mbps (4 fili)*
- *SDSL Single-line Digital Subscriber Line: 1.5 Mbps/1.5 Mbps (2 fili)*
- *VDSL Very high Digital Subscriber Line: 13 Mbps-52 Mbps/1.5 Mbps- 2.3 Mbps.*
- *IDSL ISDN Digital Subscriber Line: 128 Kbps/128 Kbps.*
- *RADSL Rate Adaptive Digital Subscriber Line: 384 Kbps/128 Kbps*
- *UDSL Universal Digital Subscriber Line: 1.0 Mbps-384 Kbps/384 Kbps-128 Kbps chiamata anche "splitterless" DSL o DSL-Lite, dato che non richiede alcun separatore di segnale (splitter).*

dove Xbps/Ybps equivale a X=velocità in ricezione, Y=velocità in trasmissione.

ADSL

ADSL è una tecnica di comunicazione a larga banda che permette l'esercizio simultaneo della fonìa e della trasmissione di dati anche multimediali (file audio e video), utilizzando la tradizionale linea telefonica, cioè il doppino in rame, con notevoli vantaggi rispetto a:

- modem fonico su rete telefonica commutata
- modem in banda-base su linea dedicata
- ISDN¹.

La tecnica ADSL consente quindi l'**utilizzo simultaneo voce-dati**, cioè permette, con un'unica linea telefonica tradizionale, di effettuare simultaneamente una trasmissione telefonica e la trasmissione-dati.

(Anche l'accesso-base ISDN permette la stessa cosa, però richiede la completa digitalizzazione del segnale fonico da utente a utente, il che non è richiesto da DSL. D'altra parte il tradizionale modem fonico su normale linea telefonica commutata non permette invece la simultaneità fra conversazioni telefoniche e trasmissione-dati: una esclude l'altra.)

Le tecniche **DSL** prevedono l'**eliminazione** del **filtraggio** passabasso a 4 KHz che si effettua sulle linee telefoniche standard. Ricordiamo che la restrizione dalla banda telefonica fra i 300 e i 3400 KHz non è dovuta al doppino (che ha una banda passante di alcuni MHz) bensì, appunto, a dispositivi di rete come i filtri di canale, situati nelle centrali telefoniche per limitare le interferenze.

DSL può lavorare in più d'una modalità, cui corrispondono differenti tecniche:

- **CAP** (Carrierless Amplitude and Phase modulation o modulazione di ampiezza e fase senza portante), che è un'evoluzione della modulazione a singola portante QAM
- **MVL** (Multiple Virtual Line o Linea Multipla Virtuale)
- **DMT** (Discrete MultiTone Technology o Tecnica Discreta Multitono), detta anche "modulazione multiportante", che si serve di una molteplicità di portanti modulate ciascuna con tecnica **QAM**.

¹ "Integrated Services Digital Network", o ISDN, è una rete digitale che dà supporto a molti servizi di **voce** e **dati**.

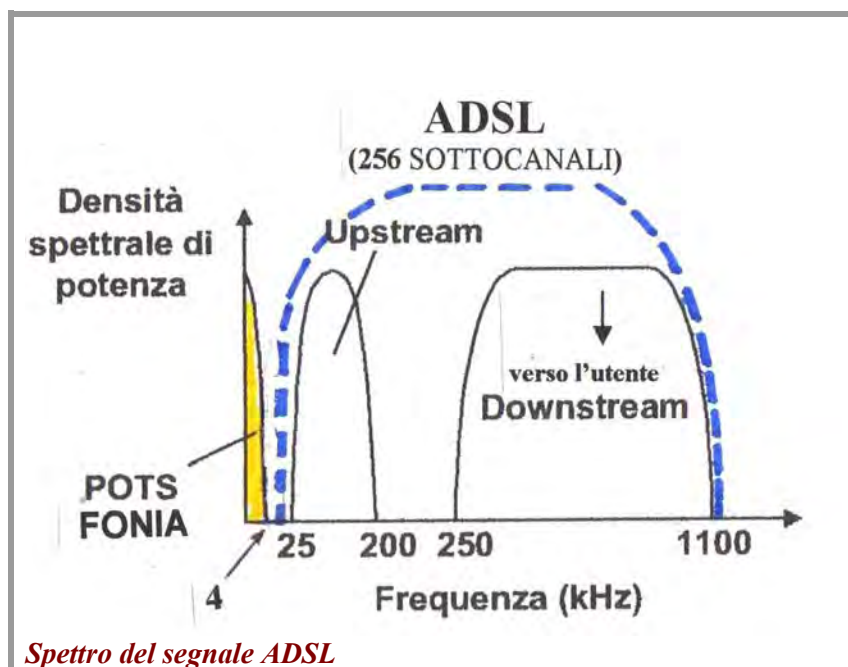
ADSL CON MODULAZIONE DMT

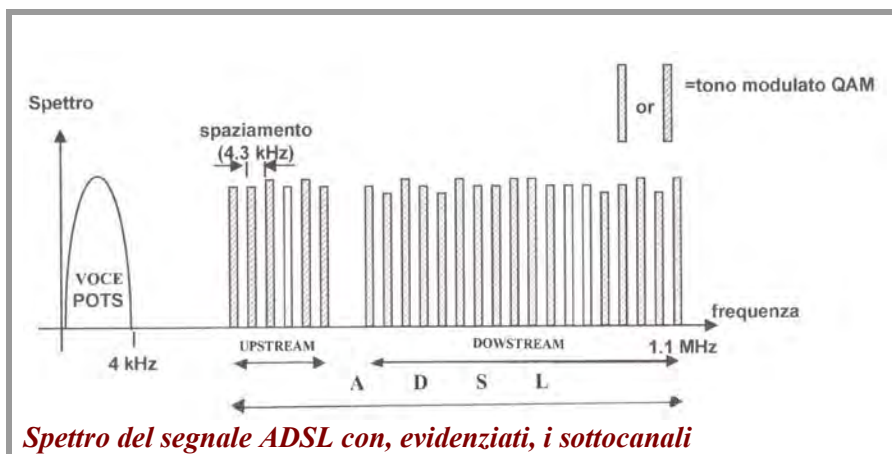
Noi faremo riferimento esclusivamente **alla modalità DMT** che è entrata a far parte degli standard ANSI (T1.413), ETSI e ISU.

Il segnale ADSL utilizza, per la trasmissione-dati, un intervallo di frequenze posto al di là dei 3400 Hz della banda telefonica, rendendo così possibile lo svolgimento contemporaneo delle telefonate e della trasmissione-dati.

Rispetto alla tradizionale trasmissione-dati su linea commutata PSTN (che richiede, da parte del flusso informativo digitale, la modulazione di un'unica portante), la tecnica DSL prevede la **suddivisione del flusso informativo digitale** in una **molteplicità di sottoflussi**, ciascuno dei quali modula con **tecnica QAM** una **diversa portante** posta al di là della banda telefonica.

Con la tecnica DMT quindi il flusso digitale di informazione viene suddiviso in 256 sottoflussi informativi e ciascun sottoflusso informativo, cioè ciascun sottocanale, modula, con tecnica QAM, una portante sinusoidale. Si può dire che la DMT crea, sulla stessa linea, 256 “modem virtuali” funzionanti simultaneamente.

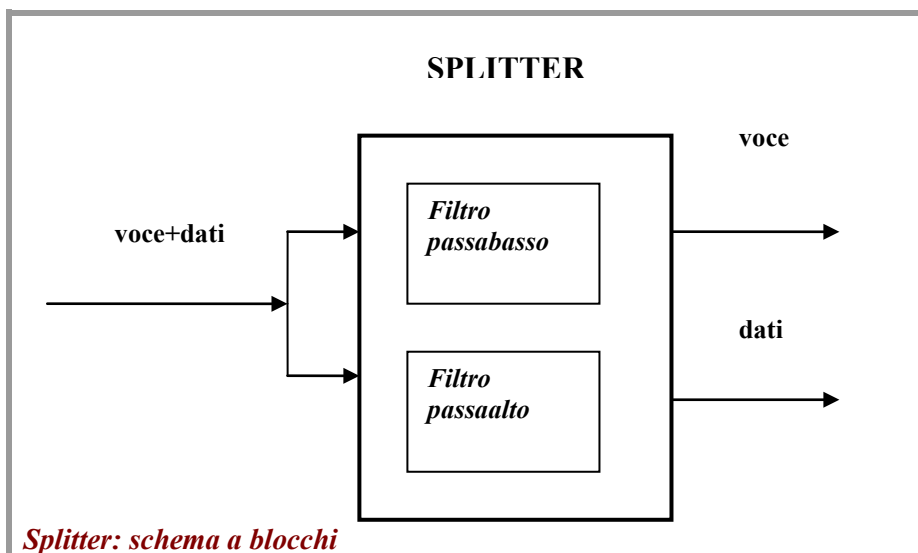




IMPIANTI DOMESTICI ADSL

LO SPLITTER

Lo splitter (o filtro DLS) è formato da **due filtri**: un passabasso e un passaalto e ha la funzione di **smistare** il segnale contenente la componente fonica (voce) e la componente-dati (DSL) **su due differenti linee**.



Nello splitter entra il segnale complessivo con le componenti in banda base (voce) e in banda traslata (dati) e in uscita non sono eliminate né la componente a frequenza più alta né quella a frequenza più bassa, ma avviene solo una separazione. Pertanto il filtro non è un semplice passabasso, ma un “insieme passabasso+passaalto”, cioè un dispositivo che separa le due componenti smistandole su linee diverse, senza sopprimerne nessuna delle due. Quindi lo splitter:

- Impedisce al segnale-dati (banda superiore) di entrare nel telefono
- Impedisce al segnale fonico (banda inferiore, voce) di entrare nel modem ADSL.

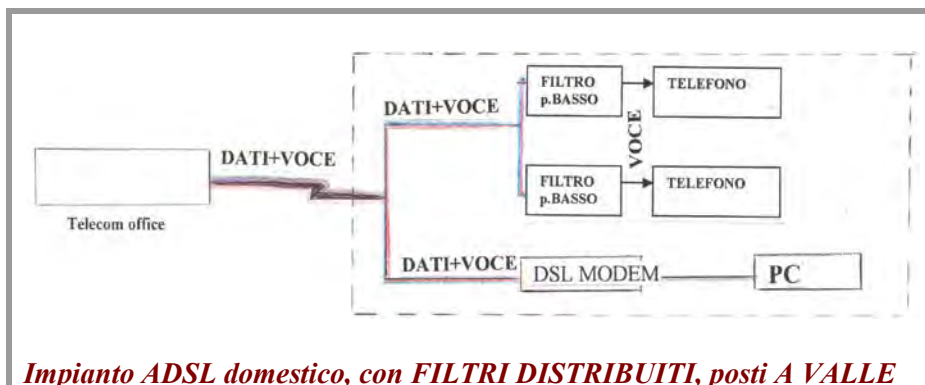
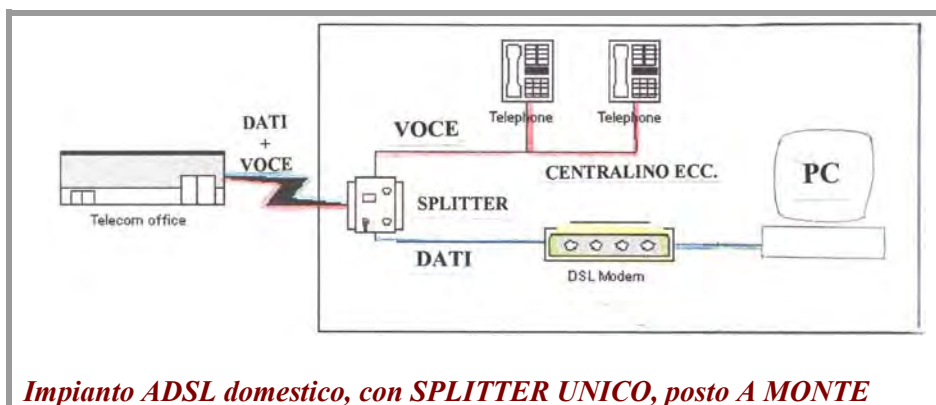


CONFIGURAZIONI DI IMPIANTI DOMESTICI ADSL

Per la realizzazione dell'ADSL in un'abitazione o in un ufficio non è necessario modificare l'impianto, ma basta aggiungere un modem ADSL (modem a banda larga) e uno splitter (o normali filtri passabasso).

Esistono due configurazioni tipiche di impianti domestici ADSL:

- L'impianto con splitter a monte, utilizzato quando sono presenti dispositivi come centralini, sistemi di allarme, telesoccorso ecc.
- L'impianto con filtri distribuiti posti a valle, utilizzato quando non sono presenti i dispositivi particolari di cui sopra. In questo caso i filtri impiegati non sono splitter, ma semplici passabasso.



REQUISITI DELLA RETE “ESTERNA”

Nelle centrali telefoniche di zona deve essere presente un particolare “**multiplexer di accesso**”, detto “**DSLAM**” (**DSL Access Multiplexer**), dotato di modem compatibile con quello dell’utente, il quale **concentra più linee ADSL in un’unica linea ATM**, cioè in una di quelle linee ad alta velocità (ordine dei Gbit/s) che collegano le reti telefoniche fra loro. Attraverso una linea ATM quindi il DSLAM connette l’utente al commutatore-voce e al provider (ISP).

LA MODULAZIONE QAM

La QAM, modulazione di ampiezza in quadratura è, insieme a ASK, FSK e PSK, una **modulazione utilizzata tipicamente per i modem fonici** (modem “tradizionali”), i quali hanno l’esigenza di trasferire l’informazione contenuta in un segnale a banda larga (segnale digitale, emesso, per esempio, da un computer) su un segnale a banda stretta, **in grado di transitare su un canale anch’esso a banda stretta**, come la rete telefonica commutata.

La QAM, e le altre modulazioni per modem fonici, sono caratterizzate da un **segnale informativo** modulante **digitale** e una **portante sinusoidale**. Il segnale informativo digitale modula la portante variandone o l’ampiezza (ASK), o la frequenza (FSK), o la fase (PSK), oppure, come nel caso QAM, l’ampiezza e la fase insieme.

Il segnale modulato QAM può assumere un certo numero (per esempio 16) di stati di modulazione, corrispondenti, ciascuno, a un gruppo di quattro bit del segnale informativo, e caratterizzati, ciascuno, da un valore di fase e da un valore di ampiezza.

Per esempio un segnale QAM che assuma 16 stati di modulazione (segnale 16-QAM) potrebbe assumere 8 valori di fase (0° , 45° , 90° , 135° , 180° , 225° , 270° e 315°) e, per ciascun valore di fase, due differenti valori di ampiezza (per esempio 3 e 5 oppure $\sqrt{2}$ e $3\sqrt{2}$).

Uno stato di modulazione sarà per esempio (0° , 3V), un altro sarà (0° , 5V), e così via, secondo quanto riportato nella tabella seguente.

STATO DI MODULAZIONE	FASE	AMPIEZZA [V]	BIT RELATIVI ALLA FASE	BIT RELATIVO ALL'AMPIEZZA
1°	0°	3	001	0
2°	0°	5	001	1
<i>3°</i>	<i>45°</i>	$\sqrt{2}$	<i>000</i>	<i>0</i>
<i>4°</i>	<i>45°</i>	$3\sqrt{2}$	<i>000</i>	<i>1</i>
5°	90°	3	010	0
6°	90°	5	010	1
<i>7°</i>	<i>135°</i>	$\sqrt{2}$	<i>011</i>	<i>0</i>
<i>8°</i>	<i>135°</i>	$3\sqrt{2}$	<i>011</i>	<i>1</i>
9°	180°	3	111	0
10°	180°	5	111	1
<i>11°</i>	<i>225°</i>	$\sqrt{2}$	<i>110</i>	<i>0</i>
<i>12°</i>	<i>225°</i>	$3\sqrt{2}$	<i>110</i>	<i>1</i>
13°	270°	3	100	0
14°	270°	5	100	1
<i>15°</i>	<i>315°</i>	$\sqrt{2}$	<i>101</i>	<i>0</i>
<i>16°</i>	<i>315°</i>	$3\sqrt{2}$	<i>101</i>	<i>1</i>

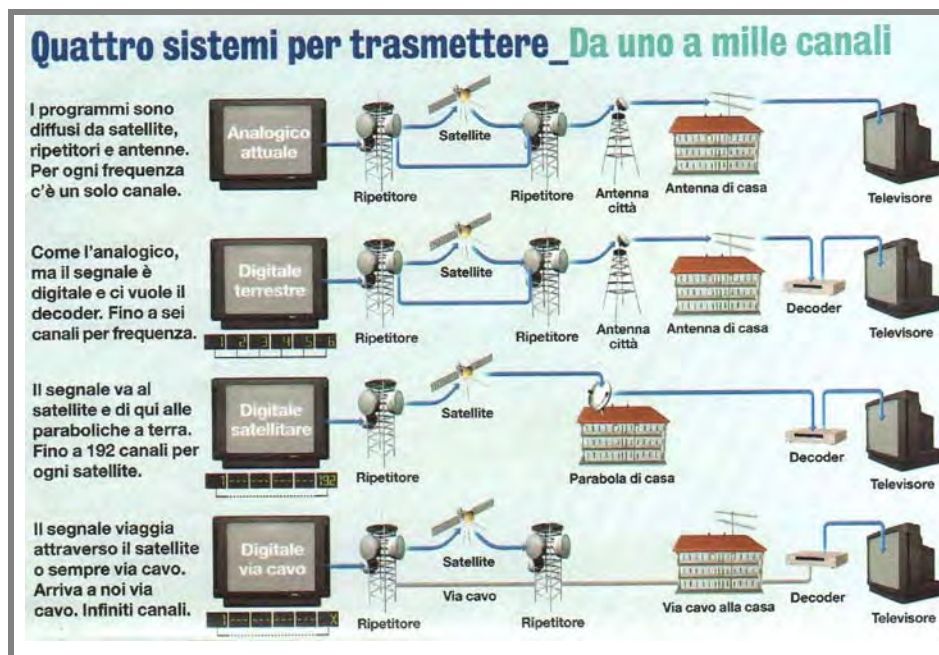
LA TELEVISIONE DIGITALE TERRESTRE (DTT)

Sono state realizzate le televisioni digitali via satellite, via cavo e via etere.

Gli standard su cui si fondano sono:

- **DVB-S** (Digital Video Broadcasting) per il satellite
- **DVB-C** (Digital Video Broadcasting) per il cavo
- **DVB-T** (Digital Video Broadcasting) per la TV digitale **terrestre** via etere.

Ci occuperemo di quest'ultima.



L'OCCUPAZIONE DI BANDA

E' prevista la realizzazione di 4÷5 programmi digitali per ogni canale analogico a radiofrequenza dell'attuale TV analogica. Siccome la banda di ciascun canale analogico è di 8 MHz, ciascun programma digitale occuperà una banda di $1,6 \text{ MHz} \div 2 \text{ MHz}$.

IL RICEVITORE-DECODER (STP o IRD)

La ricezione del segnale televisivo digitale terrestre comporta per l'utente l'adozione di **un solo dispositivo in più** (rispetto alla tv analogica) e cioè il **ricevitore-decoder**.

Il ricevitore-decoder è comunemente indicato con i seguenti acronimi:

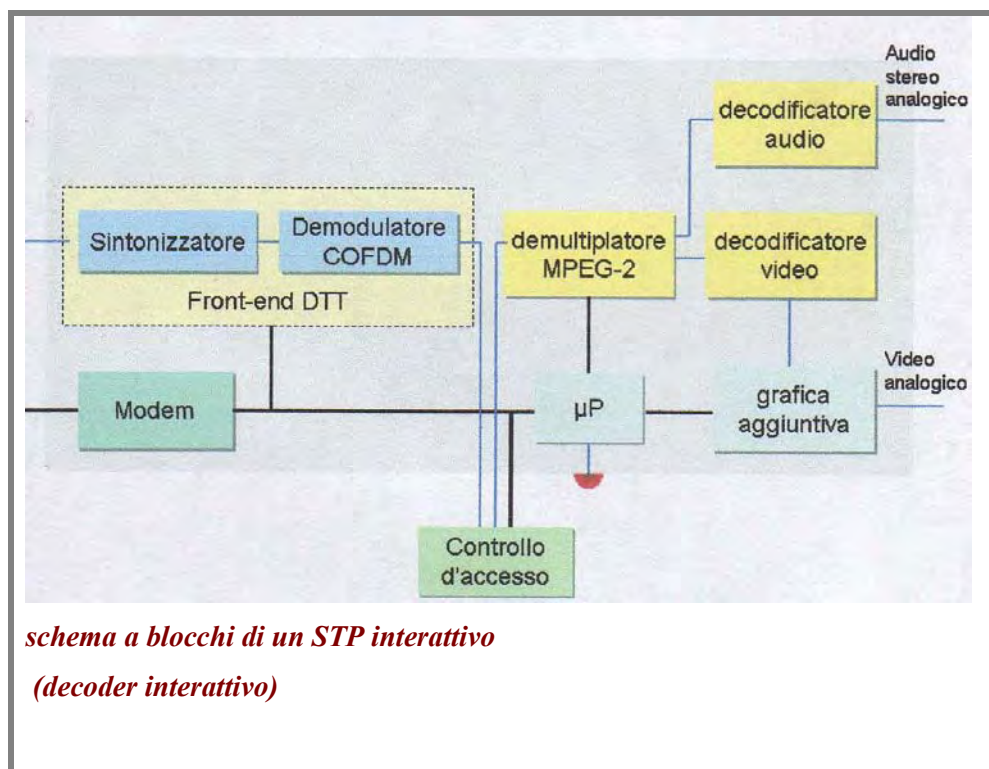
➤ **STP, Set Top Box**

(il nome deriva dal fatto che l'apparecchio può essere posto al di sopra del televisore)

➤ **IRD, Ricevitore-Decodificatore Integrato**

L'apparecchio televisivo domestico e le antenne sono gli stessi della TV analogica non satellitare (la parabola non funzionerebbe).

Già da tempo sono in commercio ricevitori televisivi “digitali” che incorporano le funzioni del decoder.



*schema a blocchi di un STP interattivo
(decoder interattivo)*

L'INTERATTIVITA'

Il cosiddetto “canale di ritorno”, cioè l'invio di **informazioni** all'emittente TV **da parte dell'utente** mediante il telecomando del ricevitore-decodificatore è costituito in realtà dalla **rete telefonica**: il decoder di tipo interattivo contiene un modem che trasmette sulla linea telefonica -dopo averlo adattato ad essa- il segnale del telecomando.

LA TRASMISSIONE E LA RICEZIONE

La trasmissione e la ricezione avvengono secondo i seguenti passi:

- * Il segnale emesso dalla stazione trasmittente, cioè il **segnale proveniente dalle telecamere** e dai microfoni, deve essere digitale. Questo comporta che il segnale proveniente dalle telecamere e dai microfoni sia stato sottoposto a **conversione analogico-digitale** oppure che telecamera e microfono forniscono segnali già digitalizzati.
- * I segnali digitali **modulano** un certo numero di **portanti analogiche**
- * I segnali **modulati e multiplati** sono irradiati dalle antenne trasmittenti
- * I segnali **modulati e multiplati** sono ricevuti dall'antenna ricevente dell'utente e **inviati al decoder** interposto fra l'antenna e il ricevitore
- * Il decoder riconverte il segnale ricevuto in un normale segnale televisivo analogico e lo applica al ricevitore domestico, cui è collegato attraverso un connettore *scart*.

GLI STANDARD

La televisione digitale terrestre è basata sullo standard **DVB-T** che a sua volta si basa sullo standard di compressione di immagini in movimento **MPEG-2** (Motion Picture Expert Group, gruppo di esperti per le immagini in movimento).

Ricordiamo che lo standard MPEG-2 è un'evoluzione di JPEG (Joint Picture Expert Group) e di MPEG-1.

LE MODULAZIONI E LE MULTIPLAZIONI ADOTTATE

Lo standard DVB-T prevede l'adozione della modulazione- multiplazione **COFDM (Multiplazione Codificata ORTOGONALE a Divisione di Frequenza²)**.

La COFDM è una modulazione-multiplazione multiportante con:

- segnale informativo modulante digitale
- portanti analogiche

Come indica anche il nome, viene adottata una **multiplazione in frequenza** (ed è grazie a questa multiplazione che, nella banda occupata da un canale analogico, possono coesistere più programmi digitali, per esempio cinque)

Il gruppo di programmi televisivi digitali relativi a un canale (cioè a una frequenza di portante assegnata) modula un gruppo di portanti **ortogonali tra loro**, e quindi **non interferenti**.

Esiste quindi una **molteplicità di portanti** per ogni **singolo canale**, cioè per ogni gruppo di programmi digitali relativi a un canale.

Cosa significa **ortogonalità** della modulazione?

L'ortogonalità si riferisce alle portanti.

Le portanti sono ortogonali tra loro quando **in corrispondenza del massimo dello spettro di una di esse**, gli spettri di tutte le altre risultano nulli.

Una portante cioè è ortogonale alle altre se, quando il suo spettro ha ampiezza massima, le ampiezze degli spettri di tutte le altre portanti assumono valore zero.

La COFDM è una tecnica di trasmissione caratterizzata dalla suddivisione del segnale di informazione ad alta velocità trasmissiva in **molti flussi paralleli a bassa velocità**, multiplati a **divisione di frequenza**.

² La tecnica **OFDM** è anche utilizzata per **ADSL** (trasmissione dati su doppino di rame), **DAB** (Digital Audio Broadcasting, diffusione radio digitale), **WiFi** (modem-router wireless), **WiMAX** (Worldwide Interoperability for Microwave ACCess), cioè in presenza di canali caratterizzati da selettività in frequenza

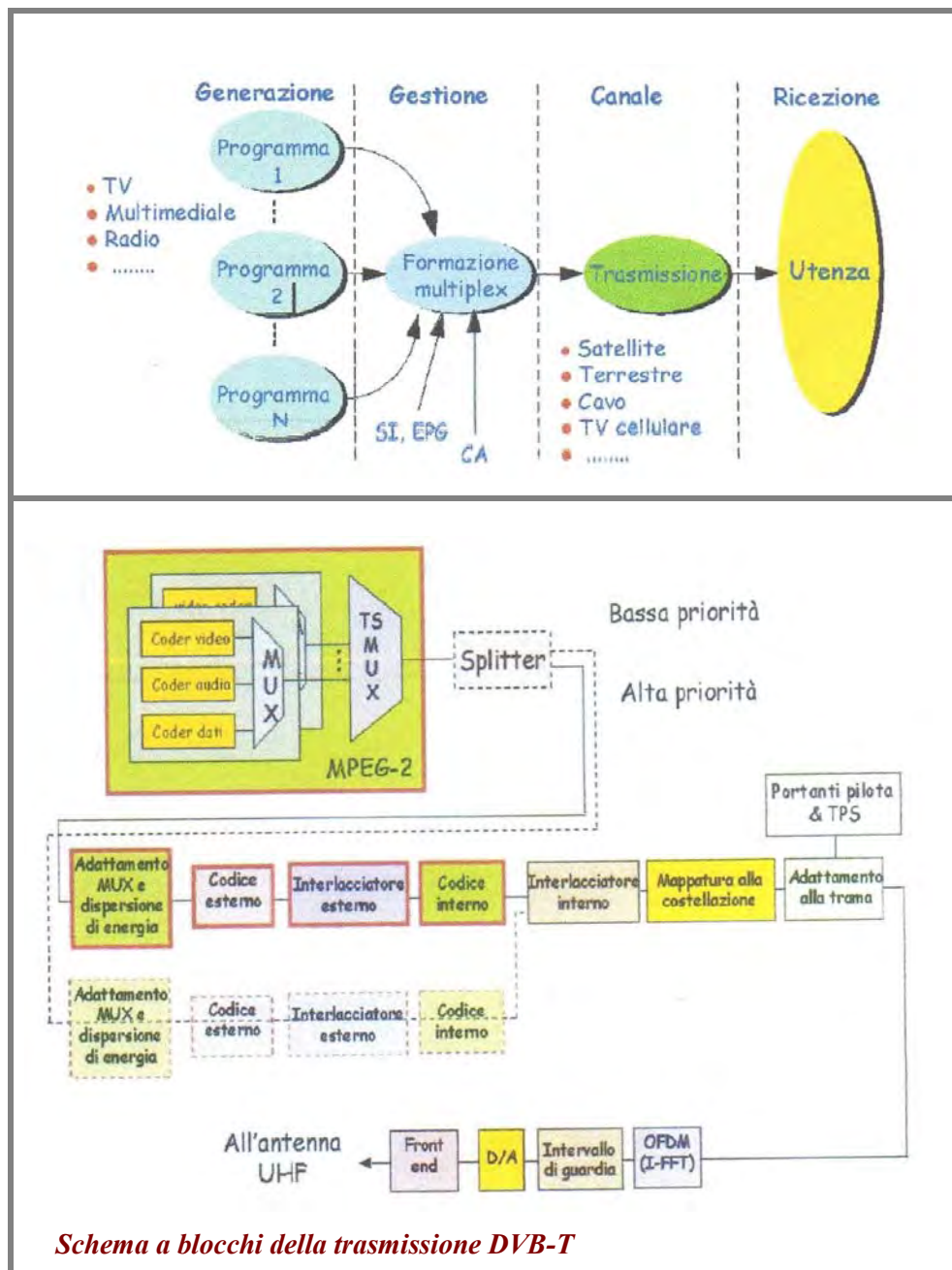
Il flusso-dati totale viene pertanto distribuito tra moltissime portanti (a banda stretta e quindi a bassa velocità di trasmissione) all'interno della banda del canale di diffusione.

Il segnale OFDM è la somma risultante dalla moltiplicazione a divisione di frequenza di N segnali con portanti a frequenze:

$0, f_0, 2f_0, \dots, (N-1)f_0$

A ciascuna delle portanti viene poi applicata una modulazione digitale del tipo **Q-PSK, M-QAM** ecc.

Il processo OFDM è attuato per mezzo di una trasformata INVERSA di Fourier I-FFT (Inverse Fast Fourier Transform).



Schema a blocchi della trasmissione DVB-T

IL RUOLO FONDAMENTALE DELLA COMPRESSIONE

Se il **segnale televisivo digitale** venisse trasmesso dall'emittente così com'è, cioè **senza essere sottoposto a compressione**, occuperebbe una **banda molto più larga di quella del segnale analogico**, il che renderebbe improponibile la televisione digitale.

Solo con il processo di compressione si riesce a restringere la banda occupata dal segnale televisivo digitale in modo da renderla “interessante” per l'applicazione pratica.

Comprimere un segnale digitale significa ridurre il numero di bit ad esso associato eliminando le “ridondanze” cioè i bit relativi a informazioni non indispensabili o esprimibili in forma più sintetica.

COMPRESSIONE VIDEO SPAZIALE

Supponiamo di dover trasmettere il segnale video relativo a due strisce colorate: una striscia rossa alta 2 cm e una striscia blu alta 5cm.

Invece di trasmettere, per ogni pixel (puntino) delle strisce, i bit relativi alle informazioni di colore, luminosità ecc, si può effettuare una compressione spaziale trasmettendo le informazioni di colore luminosità ecc una sola volta per ciascuna striscia, con l'aggiunta dell'indicazione dell'altezza della striscia.

COMPRESSIONE VIDEO TEMPORALE

Supponiamo di dover trasmettere il segnale video relativo a un quadro (cioè a una scena) in movimento, nel quale un oggetto si sposta, assumendo più posizioni successive, rispetto a uno sfondo fisso (automobile o persona che procede lungo una strada alberata).

Invece di trasmettere i bit relativi a tutte le scene successive si può -avendo memorizzato le immagini- trasmettere, una sola volta, l'informazione relativa alla scena completa (soggetto e sfondo) e poi soltanto le variazioni del quadro successivo rispetto al precedente.

COMPRESSIONE VIDEO BASATA SULL'IMPERFETTA PERCEZIONE DELL'IMMAGINE

Esiste un livello minimo di luminanza al di sotto del quale l'occhio umano non riesce a distinguere i dettagli.

Perciò, in fase di compressione, si elimina l'informazione relativa ad elementi di immagine caratterizzate da luminanza inferiore a una soglia prefissata.

COMPRESSSIONE AUDIO

La compressione del segnale audio è effettuata in base alle leggi della psicoacustica, cioè in base alle particolari modalità di percezione del suono da parte del l'orecchio e del cervello.

In fase di compressione si tiene conto, in particolare, dei seguenti fatti:

- l'udito umano è sensibile ai suoni in misura diversa a seconda della loro frequenza; esiste, per ogni frequenza, una diversa soglia di percettibilità; la maggiore sensibilità si ha nell'intervallo di frequenze $(1\div5)$ KHz; dire che, nell'intervallo $(1\div5)$ KHz, la **sensibilità** dell'udito umano è **maggiore** equivale a dire che in questa banda la soglia di percettibilità del suono è più bassa; pertanto, in fase di compressione si elimina l'informazione relativa a tutti i suoni di intensità minore della soglia di percettibilità
- mascheramento temporale: un suono molto forte rende meno memorizzabili, cioè “maschera” i suoni che temporalmente lo precedono e lo seguono immediatamente nel tempo.

COMPRESSIONE CON MPEG-2

La **velocità necessaria per trasmettere un segnale televisivo numerizzato non compresso** è di circa **120÷140 Mbit/s**.

Se il segnale è ad alta definizione è necessaria una velocità di 1000 Mbit/s.

In termini di velocità di trasmissione MPEG-2 permette di passare dai 140 Mbit/s (che sarebbero necessari per la trasmissione di un segnale TV digitale non compresso)

a una velocità di:

- ➔ **8 Mbit/s**, sufficienti per codificare **scene in movimento** (realizzando così un rapporto di compressione di $140/8 = 17,5$ circa)
- ➔ **2 Mbit/s**, sufficienti per codificare **scene poco dinamiche**, per esempio nel campo della teledidattica, (realizzando così un rapporto di compressione di $140/2 = 70$).

PREGI DELLA TV DIGITALE TERRESTRE

RISPETTO ALLA TV ANALOGICA

- Incremento del numero di canali: più canali digitali possono occupare la stessa banda di un singolo segnale analogico
- Riduzione della potenza richiesta al trasmettitore: il segnale digitale richiede una potenza minore
- Realizzazione di effetti speciali impossibili in formato analogico
- Possibilità di realizzare molte copie di un filmato senza alcun degrado qualitativo
- Semplificazione dello scambio internazionale, il quale può avvenire indipendentemente dallo standard (NTSC, PAL, SECAM MPEG, D2 MAC).

PREGI DELLA TV DIGITALE TERRESTRE

ANCHE RISPETTO ALLE TV DIGITALI SATELLITARE E VIA CAVO

- Copertura capillare del territorio (sfruttando le infrastrutture della TV analogica)
- PORTABILITA' del servizio, grazie ad antenne poco costose e semplici da installare
- REGIONALITA', cioè possibilità di trasmettere solo in zone geograficamente limitate, il che non sarebbe possibile né con la TV satellitare (che irradia su territori molto vasti) né con la TV via cavo, per la quale sarebbe difficoltoso e costoso coprire un territorio, che pur essendo geograficamente limitato, risulterebbe comunque vasto.

COFDM

Con la TV digitale a ognuna delle attuali portanti televisive analogiche viene associato un gruppo di programmi digitali detto bouquet o multiplex o ensemble.

Questo multiplex va, nel suo complesso, a modulare (con tecnica di tipo M-QAM o Q-PSK) un gruppo di specifiche portanti dette “sotto frequenze”, le quali sono ortogonali in frequenza.

Il complesso di segnali digitali è spezzettato in molteplici frammenti, ognuno dei quali modula una delle numerosi portanti ortogonali.

Vediamo come avviene, con maggior dettaglio, il processo COFDM:

- *il flusso dei dati digitali ad alta velocità³ (cioè il multiplex o bouquet di un canale, associato a una delle attuali portanti analogiche) viene SUDDIVISO in N sottoflussi, paralleli e più lenti⁴*
- *gli N sottoflussi di dati⁵ modulano (con tecnica di tipo M-QAM o Q-PSK) N sottoportanti⁶ ortogonali fra loro, formando così N segnali analogici modulati*
- *con le N sottoportanti modulate⁷ si costruisce un segnale⁸ $g(t)$*
- *da $g(t)$ si ottengono due serie di coefficienti numerici con i quali si modula la portante di canale f_0 (f_0 è una delle portanti dell'attuale sistema TV analogico)*

OSSERVAZIONI

Il processo COFDM non avviene fisicamente, nel senso che NON richiede OSCILLATORI e integratori, ma è una elaborazione matematica realizzata mediante integrati digitali o via software.

3 Velocità o bit-rate: R [bit/s]

4 cioè con velocità R/N (bit/s)

5 La suddivisione **dell'informazione su più portanti** serve a ridurre l'**interferenza intersimbolica** ed è utile su canali molto **distorcenti** o su canali affetti da propagazione **multipath** (multi-cammino), caratterizzate da riflessioni del segnale contro ostacoli

6 Il fatto che i sottoflussi, derivanti dalla suddivisione del flusso principale, vadano a modulare una MOLTEPLICITA' di sottoportanti realizza la MULTIPLAZIONE a DIVISIONE di FREQUENZA cui fa riferimento la sigla COFDM

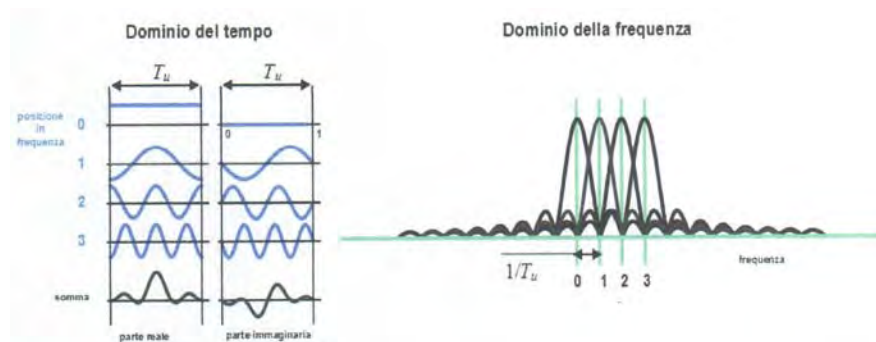
7 Le sottoportanti che trasportano dati utili sono:

1512 nella modalità 2k

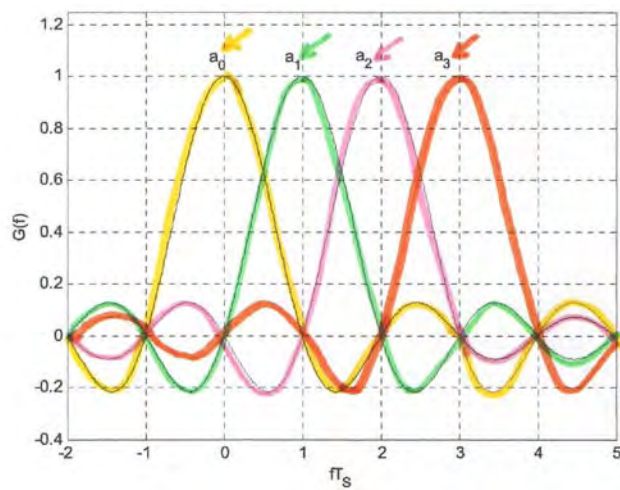
6048 nella modalità 8k

1536 nel DAB, Digital Audio Broadcasting (radio digitale)

8 $g(t)$ viene elaborato con l'algoritmo matematico noto come “trasformata veloce di Fourier INVERSA” (Inverse FFT)



representazione nel tempo e in frequenza della modulazione OFDM



ortogonalità in frequenza delle sottoportanti OFDM

LA RADIO DIGITALE

RDS (Radio Data System) e DAB (Digital Audio Broadcasting): DI COSA SI TRATTA

Il Data Radio System RDS è un sistema radio per l'**invio**, agli ascoltatori, di **informazioni ausiliarie** di natura **digitale**, contemporaneamente al segnale audio e senza interferire con esso.

Le specifiche di RDS, indicato anche come RADIODATA o TELEMATICA RADIODIFFUSA² sono state fissate dall'EBU (European Broadcasting Union). L'RDS, che funziona sia con trasmissioni monofoniche che stereo, richiede la generazione, da parte dell'emittente, di una **sottoportante aggiuntiva** di **57 KHz**, in una posizione tale, quindi, da **non interferire con lo spettro del segnale radiofonico analogico**.

La velocità di trasmissione dei dati RDS è di circa **1,8 Kbit/s** (1187,5 bit/s)

Mentre per "RDS" si intende la sola trasmissione aggiuntiva di dati da parte di emittenti FM analogiche, per "DAB" si intende l'intero sistema per la diffusione di programmi audio digitalizzati e di dati aggiuntivi.

Il sistema DAB è stato sviluppato negli anni '90 nell'ambito del progetto EUREKA 147.

Quindi RDS è solo un servizio digitale aggiunto alla radiofonia analogica, mentre DAB è la radiofonia interamente digitale.

COSA SERVE PER UTILIZZARE RDS

E' necessario un **ricevitore predisposto allo scopo**, dotato, fra l'altro, di **display a 8 caratteri**.

Le **autoradio** sono generalmente predisposte per l'RDS

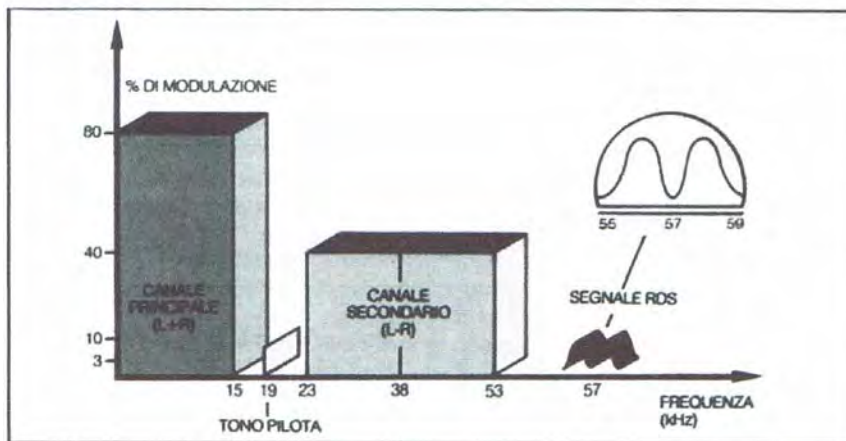
COSA SERVE PER UTILIZZARE DAB

Sono necessari:

1. un **DECODIFICATORE** digitale (che può avere le dimensioni di un cambia-CD per auto)
2. un normale **ricevitore AM/FM**, predisposto ad essere interfacciato al decodificatore.

Il principale settore di sviluppo del DAB sembra essere quello del car-stereo, ma dal 1999 gli apparati ricevitori DAB hanno cominciato ad apparire anche nei cataloghi di sistemi HiFi non automobilistici.

L'antenna ricevente è una semplice antenna omnidirezionale o fissa o mobile o portatile (simile a quella della tradizionale radio analogica).



SEGNale MODULANTE DI UN'EMITTENTE FM CON CODIFICA RDS

COSA FA RDS

Le funzioni che potenzialmente l'RDS può mettere a disposizione sono numerose. Non tutte però risultano in pratica disponibili. Tra le funzioni principali ricordiamo:

- **PROGRAM SERVICE (PS o PSN)**

E' la funzione che permette di **VISUALIZZARE** il **NOME** dell' **EMITTENTE** o qualunque altro testo. Ciascun messaggio può essere composto da un massimo di 50 scritte (eventualmente scorrevoli) di 8 caratteri.

- **ALTERNATIVE FREQUENCIES (AF)**

Nel caso che l'emittente irradia su più frequenze, il ricevitore RDS può **CAMBIARE AUTOMATICAMENTE FREQUENZA**, sintonizzandosi sulla portante (dell'emittente in esame) il cui segnale risulta più forte.

Il ricevitore può sintonizzarsi su una delle 25 frequenze (al massimo) che l'emittente ha comunicato.

- **TRAFFIC PROGRAM (TP)**

Il ricevitore RDS può essere impostato in modo che, quando la stazione emittente comincia a trasmettere **INFORMAZIONI** sul **TRAFFICO AUTOMOBILISTICO**, esso stesso (il ricevitore) **INTERROMPA OGNI** altra funzione (per esempio la lettura di un cd o di una cassetta) **PER** proporre l'ascolto **DEL BOLLETTINO**.

- Riguardo alla RAI, essa inserisce all'inizio e alla fine dei notiziari “onda verde” un jingle che viene riconosciuto dal ricevitore, il quale attiva la funzione RDS-TP determinando o l'aumento del volume o la commutazione dalla lettura (di CD o cassette) all'ascolto del notiziario.

- **PROGRAM TYPE**

Permette di VISUALIZZARE sul display la TIPOLOGIA (a scelta fra 32 disponibili) di PROGRAMMA (NEWS, SPORT, EDUCATE....) o di GENERE MUSICALE (POP, ROCK....)

- **PROGRAM IDENTIFICATION (PI)**

E' un codice di 16 bit che viene trasmesso dal network e lo identifica.

Il codice PI viene continuamente controllato dal ricevitore per assicurarsi di essere sintonizzato sull'emittente desiderata dall'utente.

I 16 bit sono suddivisi in tre aggruppamenti di 4 bit con i seguenti significati

- codice di nazione: 4 bit
- identificazione dell'area coperta: 4 bit
- numero di riferimento della stazione: 8 bit.

CARATTERISTICHE DI DAB

Il sistema DAB si basa essenzialmente sulle seguenti tecniche:

- codifica digitale dell'informazione
- modulazione con segnale modulante informativo digitale e portante analogica
- compressione
- correzione d'errore.

Le principali caratteristiche dello standard DAB sono:

- possibilità di trasmettere un **gruppo di 6÷8 programmi stereo** (e di dati addizionali) occupando una **banda di soli 1,5 KHz**. Ciascun programma stereo, di elevata qualità, ha una velocità di 256 Kbit/s
- i programmi del gruppo o pacchetto sono multiplati in frequenza e utilizzano una stessa portante radio
- codifica del segnale audio stereo secondo lo standard MUSICAM⁹ (Masking-pattern adapted Universal Sub-band Integrated Coding And Multiplexing)
- possibilità (grazie alla multiplazione) di utilizzare la stessa frequenza di diffusione per coprire aree geografiche adiacenti (cosa non possibile nella radiodiffusione tradizionale)
- potenzialità di invio all'utente di stringhe di dati multimediali (testi, disegni, foto) e di informazioni dettagliate in tempo reale sulla viabilità.

⁹ standard presente anche in MPEG e in mp3

DAB E TECNICA COFDM

DI MODULAZIONE E MULTIPLAZIONE

In maniera analoga alla TV digitale, la radio digitale DAB¹⁰ utilizza la tecnica COFDM (Coded ORTHOGONAL Frequency Division Multiplexing), in base alla quale il complesso dei dati digitali emessi da una stazione modula un gruppo di circa 1500 sottoportanti¹¹ ortogonali in frequenza (si veda TV digitale).

¹⁰ Attenzione: la tecnica COFDM è usata da DAB e non da RDS

¹¹ Non confondere il gruppo di portanti ortogonali con la sottoportante RDS di 57 KHz

CAPITOLO 44

APPLICAZIONI NAUTICHE ED AERONAUTICHE DELL'ELETTROTECNICA E DELL'ELETTRONICA

COMPLEMENTI

IL SISTEMA DI CARTOGRAFIA ELETTRONICA ECDIS (Electronic Chart Display and Information System)

L'INTERFACCIA UTENTE

MODALITÀ DI VISUALIZZAZIONE

GIROSCOPIO, PIATTAFORMA INERZIALE E NAVIGAZIONE INERZIALE

- ***IL GIROSCOPIO***
- ***ROTORE DEL GIROSCOPIO***
- ***APPLICAZIONI DEL GIROSCOPIO***
- ***LA PRECESSIONE GIROSCOPICA***
- ***PICKOFF***
- ***GIROSCOPI A DUE GRADI DI LIBERTÀ***
- ***PRECESSIONE DI UN GIROSCOPIO A DUE GRADI DI LIBERTÀ***
- ***NAVIGAZIONE INERZIALE***
- ***LA PIATTAFORMA INERZIALE***
- ***TIPOLOGIE DI PIATTAFORME INERZIALI***
- ***CALCOLATORE DI NAVIGAZIONE***
- ***ACCELEROMETRI***
- ***TORQUERS (GENERATORI DI COPPIA)***
- ***FUNZIONAMENTO DI UNA PIATTAFORMA INERZIALE DI TIPO "GIMBALLED"***
- ***GIROSCOPIO DI RATEO***
- ***PIATTAFORMA INERZIALE "STRAPDOWN" O "ANALITICA"***

IL SISTEMA DI CARTOGRAFIA ELETTRONICA E CDIS (Electronic Chart Display and Information System)

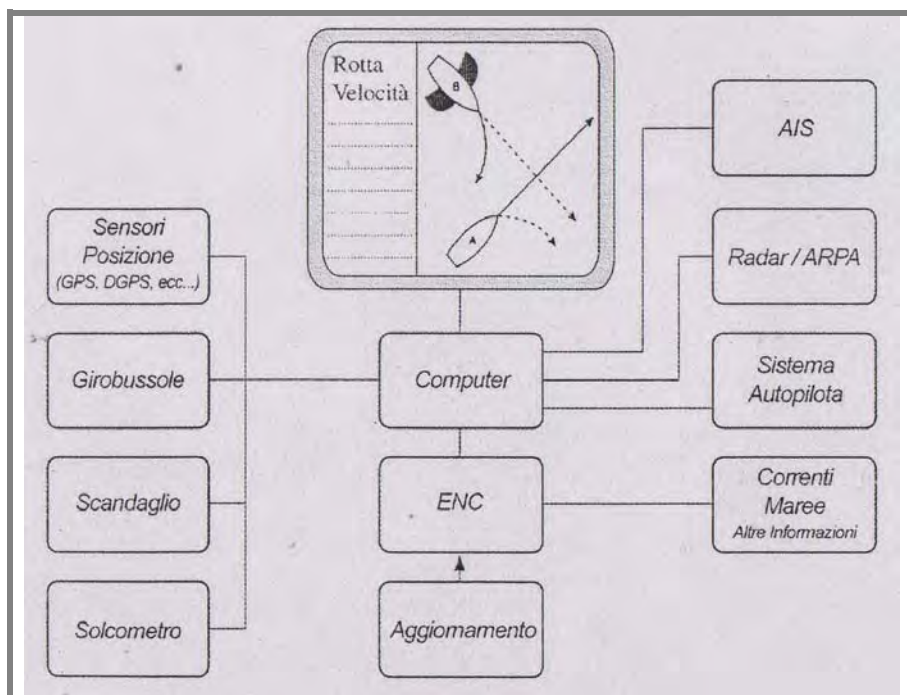
L'ECDIS è il complesso hardware e software di cartografia elettronica, atto a fornire un supporto alla navigazione, interfacciabile con reti di comunicazione e strumenti di navigazione.

- L'hardware include un processore, un monitor ad alta risoluzione, una tastiera, un trackball (per la traslazione della carta e il richiamo di funzioni e informazioni) un hard-disk, interfacce per lo scambio di dati con le reti GSM, Internet, IMMARSAT e con strumenti di navigazione quali radar-ARPA, GPS e altri.
- Il software controlla la gestione dei dati e la loro presentazione sul monitor. Consente inoltre l'interazione con gli strumenti di bordo.
- I dati sono costituiti da:
 - carte nautiche elettroniche (ENC, Electronic Navigational Chart), ossia dall'insieme delle informazioni che, elaborate e fornite al monitor, visualizzano le carte nautiche corredate di informazioni supplementari
 - il database SENC (System Electronic Navigational Chart), al quale si accede attraverso l'apparecchiatura ECDIS, database risultante dalla trasformazione e dall'aggiornamento delle ENC da parte del sistema.

Osserviamo esplicitamente che ciò che viene visualizzato e direttamente utilizzato, per le funzioni di navigazione, dal sistema ECDIS, è il database SENC , mentre le ENC rimangono inalterate nel sistema e archiviate in esso.

Il sistema ECDIS è stato standardizzato dall'IMO (Organizzazione Marittima internazionale, dall'IHO (Organizzazione Idrografica internazionale) e dall'IEC (Commissione Elettrotecnica internazionale).

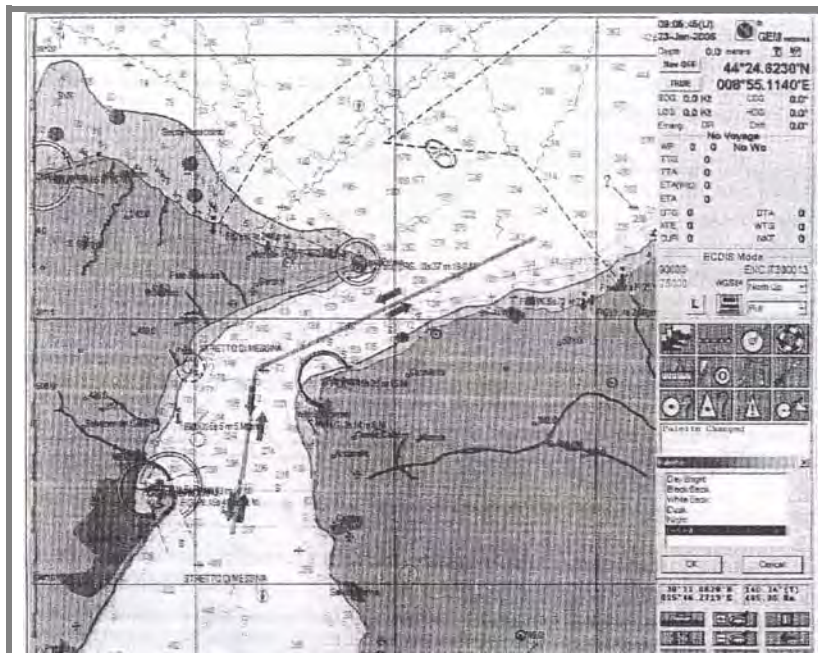
La convenzione IMO-SOLAS del 2002 stabilisce infine che l'ECDIS (dotato di opportuni sistemi di backup) può sostituire le tradizionali carte nautiche.



L'INTERFACCIA UTENTE

L'interfaccia utente presenta sul monitor la carta e le altre informazioni e funzioni:

- l'immagine della vera e propria carta occupa gran parte dello schermo
- su una barra laterale o superiore o inferiore sono visualizzati:
 - i dati della nave provenienti dal GPS (coordinate della nave, velocità e angolo di rotta), dalla girobussola (direzione della rotta) e dal solcometro (velocità)
 - i dati della navigazione in corso (data waypoints, tempi di arrivo, tempi stimati, distanza dall'arrivo)
 - i dati relativi alla carta (nome, scala, orientamento, ecc.)
 - pulsanti con differenti funzioni (misure di angoli e distanze, informazioni sui bersagli ARPA, variazioni della scala di visualizzazione, uomo in mare)
 - uno o più menu contenenti funzioni di pianificazione, monitoraggio e registrazione della rotta, dati della nave, la selezione degli allarmi, la configurazione dei sensori, l'aggiornamento manuale e la scelta delle unità di misura delle carte.



MODALITA' DI VISUALIZZAZIONE

Con il sistema ECDIS l'utente può scegliere il livello di dettaglio delle informazioni che vuole siano visualizzate. Esistono in particolare tre modalità di visualizzazione: "display base", "display standard" e "all other information".

- **Categoria di visualizzazione DISPLAY BASE**

E' la modalità di visualizzazione meno dettagliata ed è considerata insufficiente per una navigazione sicura. Deve contenere almeno le indicazioni seguenti:

- linea di costa
- isobata di sicurezza della propria nave, scelta dal navigante
- pericoli isolati sommersi a profondità inferiore all'isobata di sicurezza e posti nelle acque definite sicure dall'isobata di sicurezza

- pericoli isolati emersi posti entro l'area definita sicura dall'isobata di sicurezza (ponti cavi sospesi ecc)

- **Categoria di visualizzazione DISPLAY STANDARD**

E' la modalità che si attiva per default all'accensione del sistema. Oltre a tutte le informazioni fornite dal display base, contiene ulteriori informazioni per il monitoraggio e la pianificazione della rotta, informazioni il cui livello può essere modificato dal navigante e che sono:

- isobata zero o linea di riferimento scandaglio
- indicazione degli aiuti alla navigazione
- passaggi (bordi di fairway), canali ecc
- oggetti visibili di una certa rilevanza e radar
- aree proibite e regolamentate
- bordi della scala della carta
- note di avvertenze

- **Categoria di visualizzazione ALL OTHER INFORMATION**

Consente la visualizzazione di tutti gli altri tipi di informazione che possono essere richiamate dal navigante, in funzione delle proprie esigenze. Queste informazioni sono:

- fondali
- condotte e cavi sottomarini
- rotte delle navi
- dettagli di tutti i pericoli isolati
- dettagli degli aiuti alla navigazione
- testo di note di avvertenza
- data di edizione della ENC
- datum geodetico
- deviazione magnetica
- reticolato
- nomi dei luoghi

VISUALIZZAZIONE DELLA PROPRIA NAVE

La nave può essere rappresentata:

- con un piccolo cerchio (nella navigazione in mare aperto)
- con la propria forma, su carte a scala maggiore (in caso di avvicinamento alla costa e di entrata nei porti)
- con la propria forma in pianta e in scala vera (nel caso, per esempio, di manovre in acque ristrette)

MODI DI PRESENTAZIONE

- modalità "movimento vero o assoluto": la nave si muove sulla carta ferma
- modalità "movimento relativo": la nave è ferma al centro del monitor, mentre la carta scorre al di sotto di essa in direzione opposta.

Inoltre le informazioni della carta possono essere visualizzate:

- con il nord rivolto verso l'alto, come nelle rappresentazioni cartacee
- in ogni altra direzione
- con la direzione della nave rivolta sempre verso l'alto (conformemente a ciò che si vede dal ponte di comando)

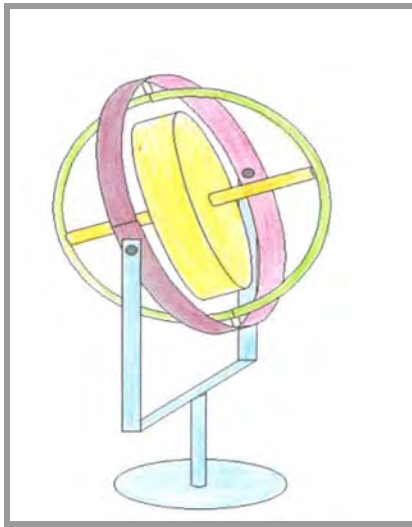
FUNZIONI DI NAVIGAZIONE

Fra le funzioni di navigazione offerte dal sistema ENCIS di cartografia elettronica ricordiamo::

- la pianificazione della rotta
- il controllo automatico della rotta pianificata
- segnalazione di allarmi in caso di pericoli, malfunzionamenti e discrepanze fra i dati
- sovrapposizione radar: l'immagine radar può essere sovrapposta all'immagine ECDIS per individuare errori, che ,altrimenti, sarebbero di difficile rilevazione
- registrazione del viaggio.

GIROSCOPIO, PIATTAFORMA INERZIALE E NAVIGAZIONE INERZIALE

IL GIROSCOPIO



Il giroscopio¹ è un dispositivo basato sulla proprietà, (che² hanno i **corpi posti in veloce moto rotatorio**), di **mantenere invariata**³ la **direzione del proprio asse di rotazione**⁴, rispetto a un sistema **inerziale**.

.Funzionano in base al principio del giroscopio la **bussola giroscopica** e le **piattaforme inerziali**, che consentono il rilievo della posizione (e di altri parametri) senza alcun ausilio esterno e senza emettere alcun segnale.

*Oltre al giroscopio meccanico, di cui ci stiamo occupando, esistono anche giroscopi “ad anello **LASER**” (**RLG**, Ring Laser Gyro) e giroscopi **MEMS** (Sistemi **Micro-Elettro-Meccanici**), costituiti da componenti elettronici e meccanici miniaturizzati.*

¹ In inglese “**gyroscope**” o semplicemente “**gyro**”

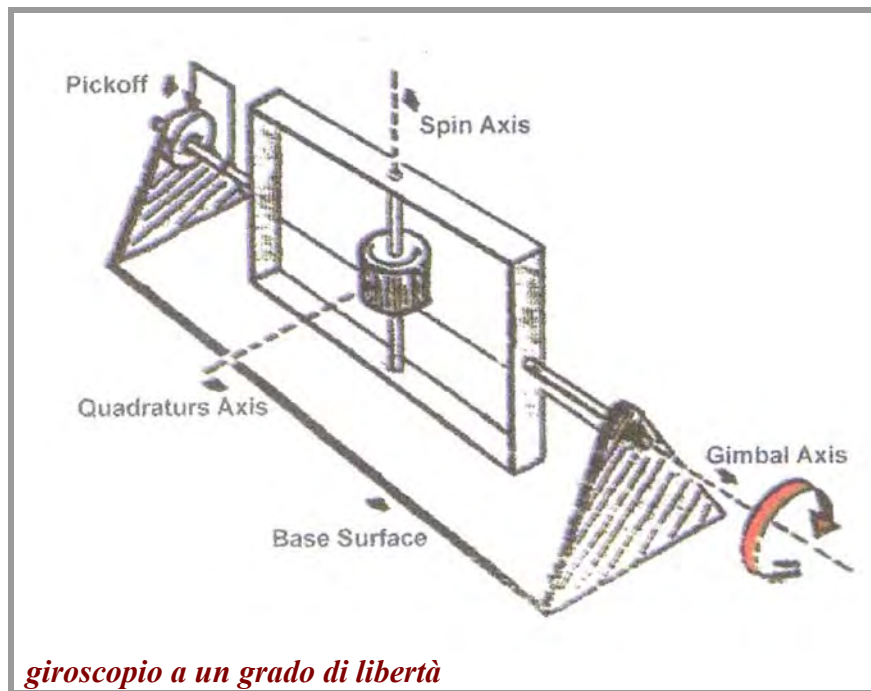
² in determinate condizioni

³ come il cielo delle stelle fisse

⁴ **asse di spin** o **asse di figura**

Il giroscopio meccanico è formato da:

- un **rotore** (o volàno o ruota), tipicamente metallico
- uno o più **anelli** di sospensione detti **gimbal**⁵ (il numero di anelli dipende dai “gradi di libertà”) di cui il giroscopio è dotato cioè dal numero di piani sui quali l’asse del rotore può spostarsi
- **trasduttori di velocità angolare**, per esempio pickoff
- **una base o piattaforma**



⁵ Dall'inglese “gimbals” = sospensione cardanica

ROTORE DEL GIROSCOPIO

E' un corpo, di massa preponderante rispetto a quella degli anelli, in grado di ruotare a velocità elevata grazie a un motore (generalmente elettrico).

Nella girobussola⁶, per esempio, il rotore è posto in movimento da un motore elettrico asincrono che ruota alla velocità di (10000-20000) giri/min.

La **rotazione del rotore** avviene **attorno** a un asse chiamato **asse di spin**.

*La capacità **del rotore** di contrastare gli spostamenti del suo asse è tanto maggiore quanto più elevata è la **velocità** di rotazione.*

Questo fenomeno è rilevabile, per esempio, in una trottola, che mantiene verticale il suo asse di spin finché la velocità di rotazione rimane abbastanza elevata.

*In assenza di sollecitazioni esterne, l'asse di spin si mantiene immobile rispetto a un sistema **inerziale** di riferimento e quindi rispetto al **cielo delle stelle fisse**.*

Rispetto alla terra⁷, che è un sistema **non inerziale**, l'**asse di spin** risulta **ruotare in verso opposto** a quello della terra stessa.

*Se si orienta l'asse di spin in **direzione di una stella**, esso conserva **quella posizione** e segue il moto apparente della stella, mentre cambia la sua posizione rispetto alle pareti e agli oggetti dell'ambiente in cui si trova.*

⁶ Una **girobussola** è un particolare tipo di bussola (chiamata anche **direzionale**), cioè un sistema di navigazione per trovare una direzione fissata, non basato sul campo magnetico terrestre, ma sulle proprietà del giroscopio

⁷ la terra non è un sistema inerziale perché è in movimento

APPLICAZIONI DEL GIROSCOPIO

1. GIROBUSSOLA

E' usata sia negli aerei che nelle navi

VANTAGGI

- indica il **Nord geografico** (invece del polo Nord magnetico)
- è **insensibile** ai disturbi magnetici
- la forza direttiva dell'elemento sensibile è anche 150 volte maggiore di quella dell'ago magnetico

INCONVENIENTI

- richiede un motore, generalmente elettrico, per mettere in funzionamento il giroscopio
- necessita di un riallineamento periodico (mediante bussola magnetica) per compensare gli effetti della precessione apparente dovuta alla rotazione terrestre
- ha un costo elevato: perciò, in campo nautico, è usata solo su navi (natanti di dimensioni non piccole) e su imbarcazioni da pesca e da diporto di un certo tonnellaggio.

2. PIATTAFORMA (NAVIGATORE) INERZIALE

3. E' utilizzata per la navigazione aerea e missilistica

INCONVENIENTI

- La precisione fornita è inversamente proporzionale alla durata della navigazione per effetto di un fenomeno di deriva nel tempo.

LA PRECESSIONE GIROSCOPICA

E' il **movimento** che l'**asse di spin**, cioè l'asse di rotazione, compie **quando al rotore viene applicata una coppia di forze**.

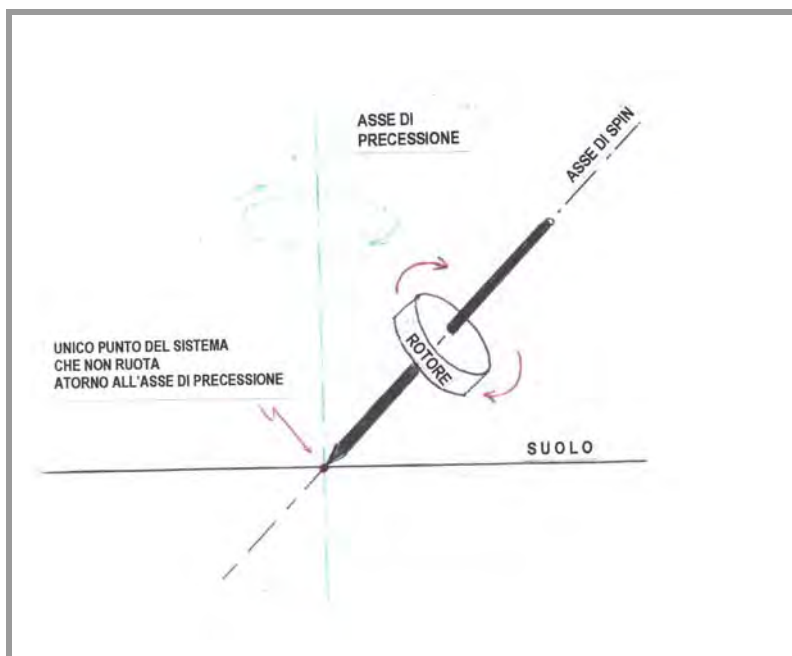
Per effetto della precessione, l'asse di spin abbandona la posizione originaria e si porta su un piano differente.

Quindi la precessione comporta che l'asse di spin ruoti a sua volta intorno a un altro asse, chiamato asse di precessione.

Nel caso della trottola, la precessione fa sì che, mentre l'asse di spin è obliquo rispetto al suolo, si produca anche una rotazione dell'asse di spin stesso (e del baricentro della trottola) attorno all'asse perpendicolare al suolo.

Possiamo dire che:

- *la precessione è una reazione del giroscopio all'applicazione di una coppia di forze*
- *se, d'altra parte, al giroscopio viene imposto dall'esterno un moto di precessione, il giroscopio reagisce con una coppia di forze.*



PICKOFF

In generale per “pickoff” intenderemo un **dispositivo** in grado di **rilevare i movimenti rotatori dei giroscopi**.

Un pickoff può essere, per esempio, realizzato con un particolare **trasformatore** dotato di una parte mobile o con **condensatori a capacità variabile**.

PICKOFF A TRASFORMATORE

In questo tipo di pickoff il **trasformatore** viene utilizzato come **trasduttore di velocità angolare**. Il pickoff è basato su un trasformatore che è in grado di **variare il coefficiente di accoppiamento** e quindi di **variare la tensione al secondario** in funzione di una **velocità angolare**.

Nel pickoff giroscopico la variazione della tensione al secondario è lineare e in particolare è direttamente proporzionale alla VELOCITA' giroscopica della rotazione (dell'anello) applicata in ingresso al pickoff stesso.

La variazione della tensione di uscita è dovuta alla variazione del valore del coefficiente di accoppiamento, il quale dipende dall'angolo di rotazione dell'anello o gimbal del giroscopio.

Quando la velocità giroscopica è nulla, o il giroscopio è in posizione di riposo, si ha in uscita una tensione o nulla o trascurabile.

Nel pickoff, una parte fissa del trasformatore, che contiene gli avvolgimenti primario e secondario, è collegata al contenitore del giroscopio, mentre una **parte mobile**, contenente **ulteriore materiale ferromagnetico**, è collegata **all'anello o gimbal**.

Quando l'anello ruota per effetto di una velocità non nulla di rotazione del giroscopio in movimento, questa rotazione relativa viene rilevata dalla parte fissa e della parte mobile del trasformatore del pickoff, il che determina la variazione dell'accoppiamento e.m, e quindi del coefficiente di accoppiamento.

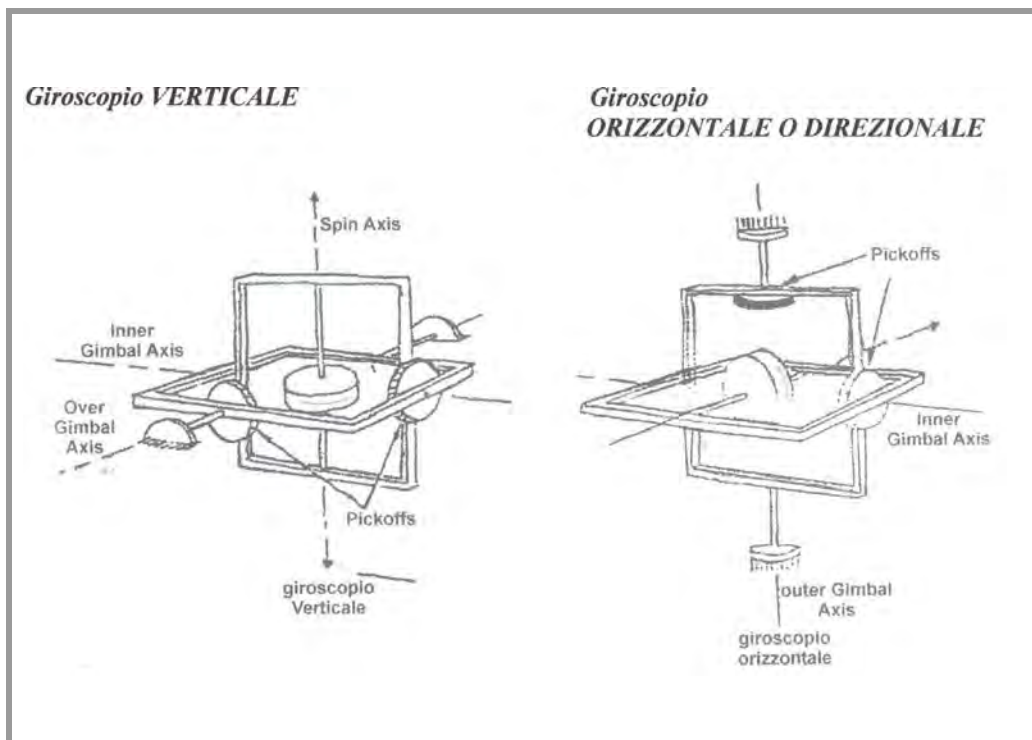
La parte fissa⁸ e la parte mobile del trasformatore del pickoff permettono di ottenere, in maniera continua, un segnale di uscita dal giroscopio.

⁸ non a contatto

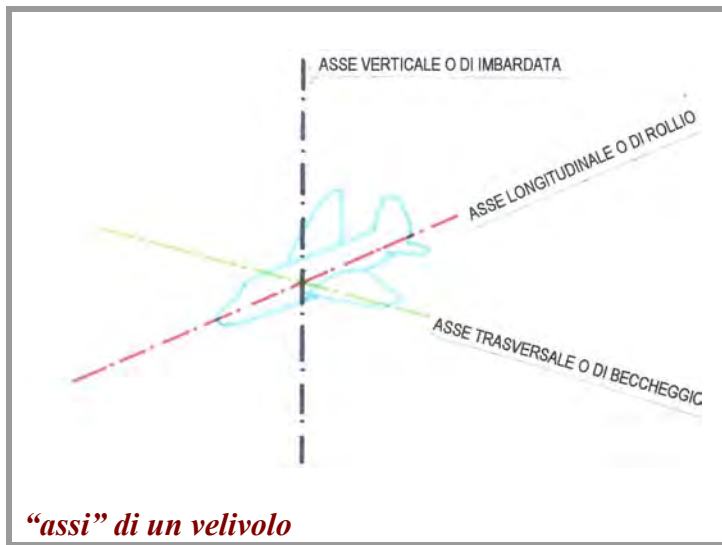
Il pickoff, essendo basato su un trasformatore, richiede una tensione elettrica alternata per funzionare, e **alternata** è anche la **tensione di uscita**.

Quando le vicende (i movimenti) del giroscopio danno luogo a una velocità di rotazione diversa da zero la tensione alternata di uscita del pickoff cresce in ampiezza.

I pickoff, nei giroscopi a due gradi di libertà, rilevano il movimento di rotazione (dell'asse di rotazione del rotore, o asse di spin) sia intorno all'asse del gimbal interno, sia intorno all'asse del gimbal esterno.



GIROSCOPI A DUE GRADI DI LIBERTA'



In navigazione aerea i giroscopi a due gradi di libertà si distinguono in:

- **GIROSCOPI VERTICALI:** in ogni luogo l'**asse di spin** è orientato secondo la **verticale**, cioè come il **raggio terrestre** in quel luogo; forniscono informazioni su **rollio e beccheggio** di un velivolo
- **GIROSCOPI ORIZZONTALI O DIREZIONALI:** l'asse di spin rimane costantemente orientato in una determinata direzione, parallela alla superficie terrestre, svolgono quindi le stesse funzioni della bussola magnetica.

PRECESSIONE DI UN GIROSCOPIO A DUE GRADI DI LIBERTA'

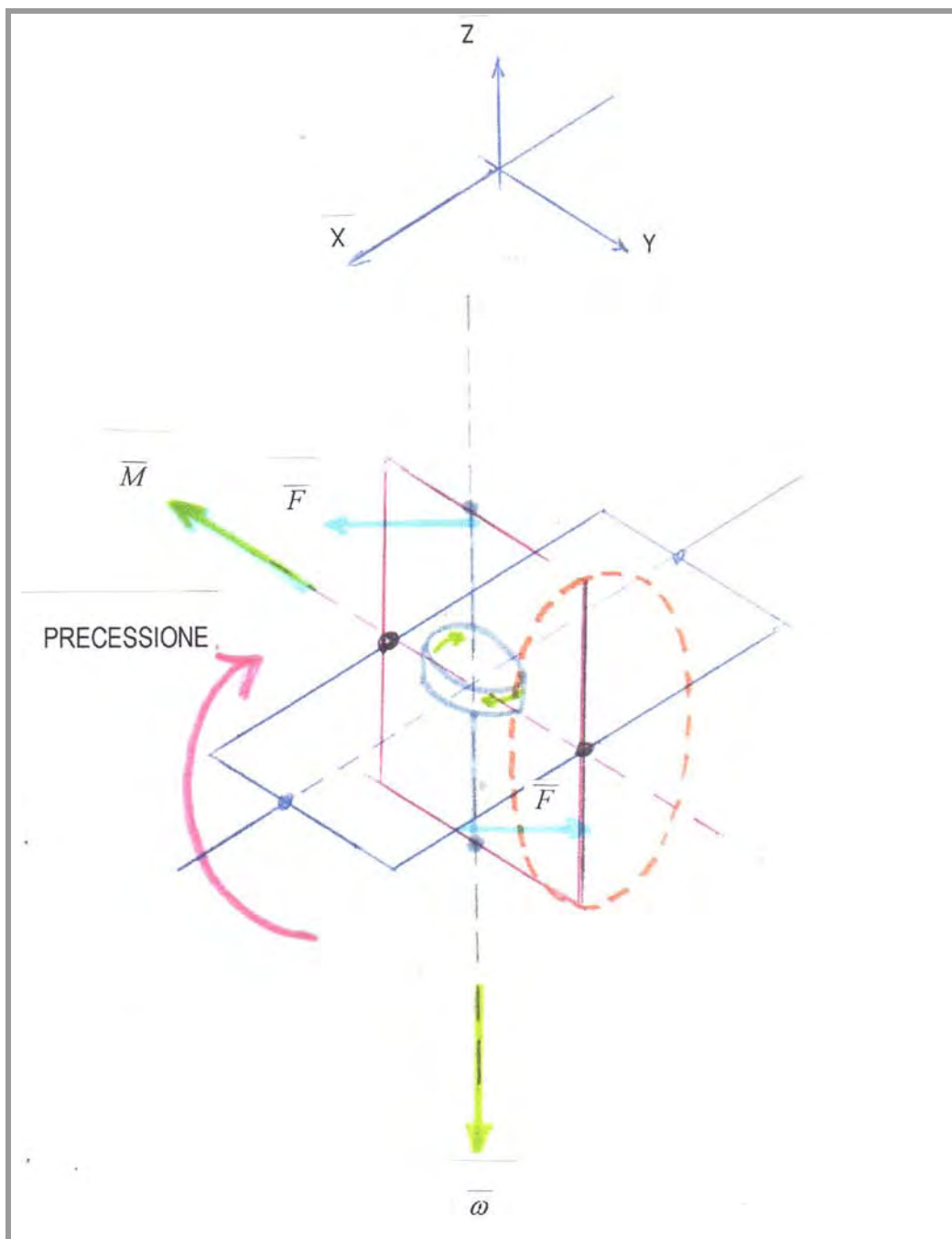
Dopo aver definito:

1. ω = **vettore velocità angolare** del rotore = vettore che è **perpendicolare al piano di rotazione** e che vede (con la punta della sua freccia rappresentativa) la rotazione avvenire in senso antiorario
2. M = **vettore MOMENTO di una coppia di forze** F = vettore che è **perpendicolare al piano nel quale giacciono le forze** e che vede (con la punta della sua freccia rappresentativa) la rotazione provocata dalle forze avvenire in senso antiorario,

possiamo dire che:

se si applica al giroscopio una coppia di forze di momento M , esso **reagisce** con un **moto di precessione**:

- la cui **orientazione** è tale da **spingere il vettore velocità angolare ω a sovrapporsi al vettore momento M**
- tale che l'asse di spin si porti su un piano perpendicolare al piano in cui giacciono le forze



NAVIGAZIONE INERZIALE (INS, Inertial Navigation System)

La navigazione inerziale è una tecnica, **autonoma e non radiante**⁹, che consente di determinare:

- la **posizione** (latitudine e longitudine) di un velivolo -o di un sottomarino o di una nave o di un'**astronave** o di un veicolo terrestre- misurando la sua accelerazione e sottoponendola due volte all'operazione matematica di integrazione
- l'**assetto** (prora, rollio, beccheggio)
- la **velocità**, effettuando l'integrazione una volta sola.

LA PIATTAFORMA INERZIALE

La piattaforma inerziale è un sistema che determina la **posizione** del velivolo (ed eventualmente l'assetto e la velocità) **in ogni istante** a partire dalla conoscenza:

- delle coordinate geografiche del punto di partenza
- delle accelerazioni
- della distanza percorsa.

La piattaforma inerziale è formata da:

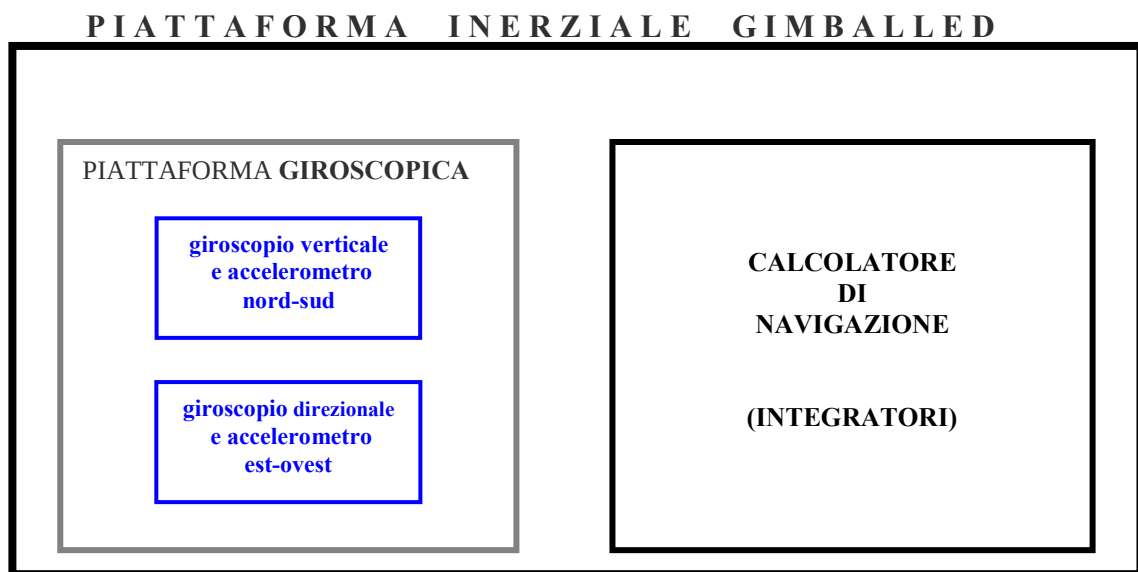
- una **piattaforma giroscopica, dotata di giroscopi e di accelerometri**, cioè di trasduttori di accelerazione
- un **calcolatore di navigazione**.

⁹ che non emette, cioè, segnali di alcun tipo

TIPOLOGIE DI PIATTAFORME INERZIALI

Le piattaforme possono essere distinte essenzialmente in:

- piattaforme **GIMBALLED** o “a sospensioni cardaniche” o “asservite”, che sono quelle di tipo “tradizionale”, basate sulla meccanica di precisione, nelle quali i giroscopi sono applicati mediante **sospensioni cardaniche**
- piattaforme **STRAPDOWN** (o “analitiche”), di concezione più moderna, nelle quali i giroscopi e gli accelerometri sono **vincolati rigidamente**¹⁰ e orientati secondo i tre assi del velivolo.



¹⁰ senza sospensioni cardaniche che ne consentano spostamenti

PIATTAFORMA GIROSCOPICA

E' un **supporto** che, in ogni luogo, si mantiene **sempre perpendicolare alla verticale del luogo**, grazie alla presenza di giroscopi e all'effettuazione di correzioni.

Possono per esempio esservi:

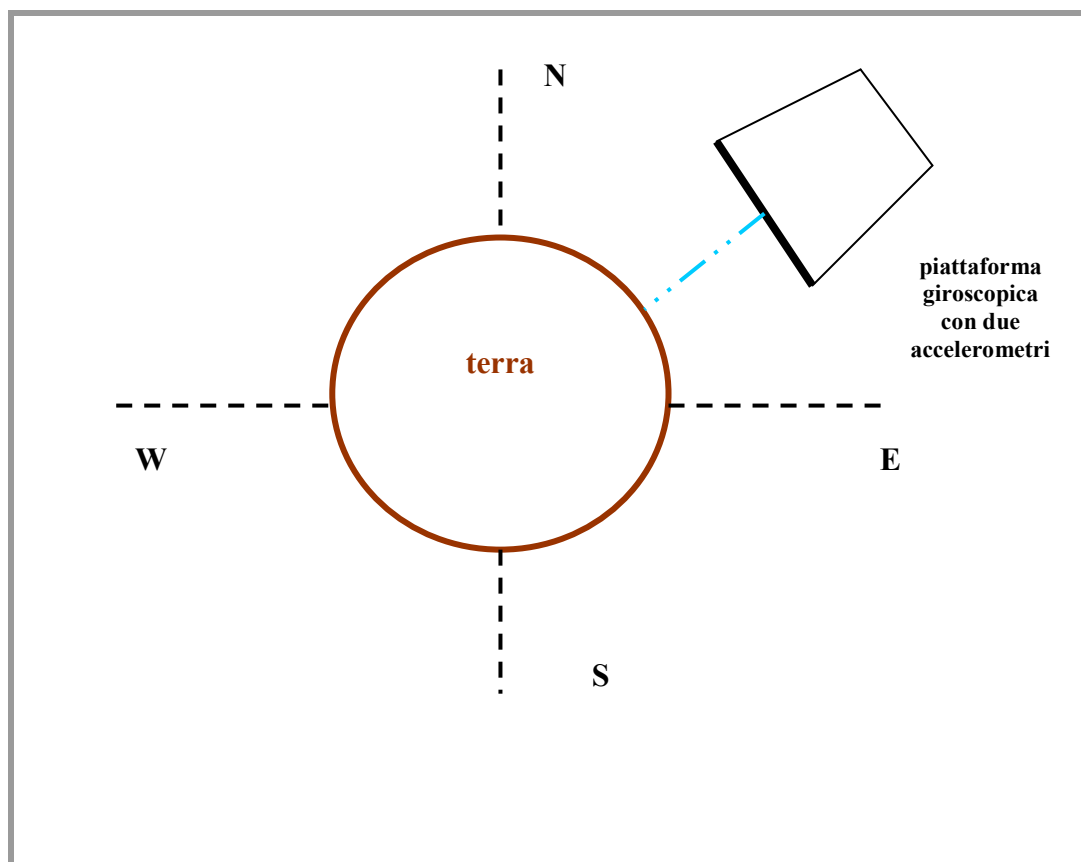
- due giroscopi a due gradi di libertà
- tre giroscopi a un grado di libertà.

PIATTAFORMA CON DUE GIROSCOPI A DUE GRADI DI LIBERTA'

Sono presenti un **giroscopio** verticale e un giroscopio orizzontale.

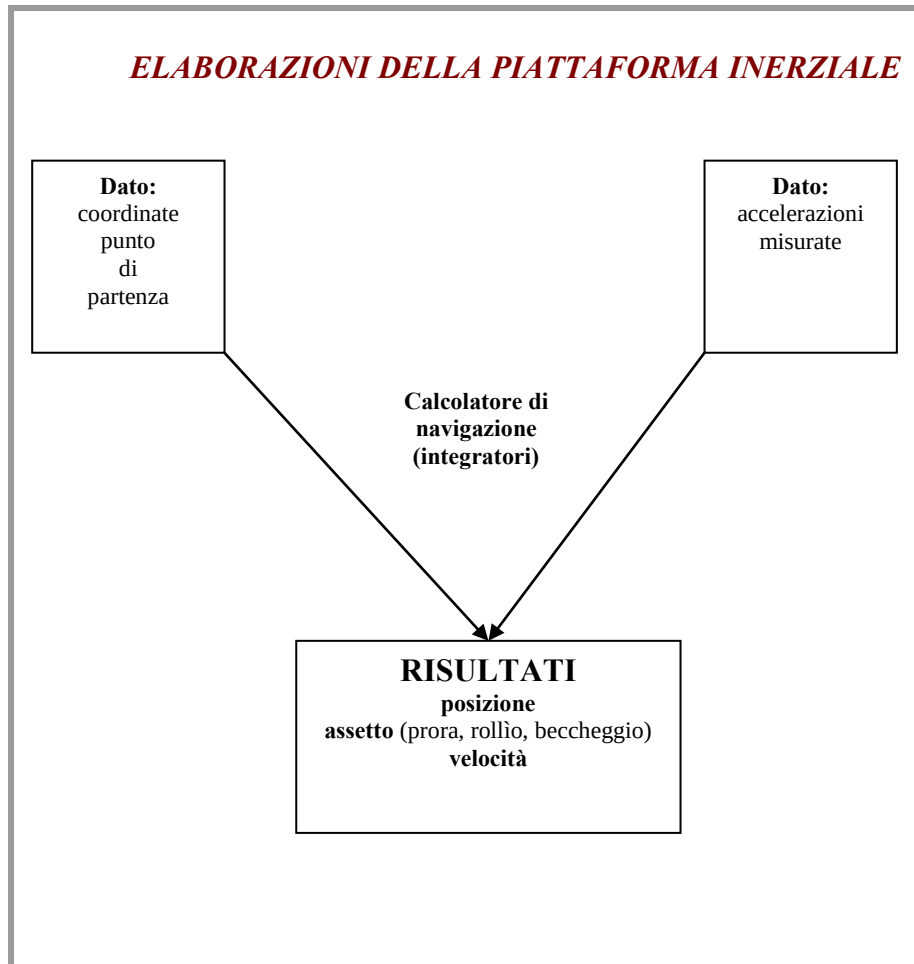
La piattaforma contiene inoltre **due accelerometri**:

- uno con **asse sensibile** nella direzione **Nord-Sud**
- uno con **asse sensibile** nella direzione **Est-Ovest**.



CALCOLATORE DI NAVIGAZIONE

Determina la velocità del velivolo e la distanza percorsa a partire dall'istante iniziale (e quindi la posizione) in base, essenzialmente, operazioni matematiche di integrazione effettuate sui segnali di uscita degli accelerometri.



ACCELEROMETRI

L'ACCELEROMETRO IN GENERALE

E' un trasduttore di accelerazione e quindi anche di decelerazione e di vibrazioni. Osserviamo, tra l'altro, che la decelerazione e/o le vibrazioni possono precedere la collisione di un veicolo con un ostacolo.

Gli accelerometri possono essere basati su differenti fenomeni fisici. Esistono, per esempio, accelerometri **meccanici** o piezoelettrici.

Cosa vorremmo misurare con l'accelerometro?

In navigazione inerziale noi vorremmo che l'accelerometro misurasse l'**accelerazione relativa** dv/dt dovuta alle **variazioni di moto del velivolo**, o, in generale, del veicolo.

accelerometro ideale $\rightarrow$ *accelerazione relativa* $\frac{dV}{dt}$

Cosa misura effettivamente, in generale, un accelerometro?

In generale un accelerometro è in grado di misurare la differenza fra l'accelerazione inerziale (dovuta alle **variazioni di moto del velivolo**) e l'accelerazione gravitazionale.

Misura cioè:

$$a_{inerz} - G$$

Cosa misura un accelerometro se il suo asse di ingresso (asse lungo il quale può scorrere la sua massa interna) è perpendicolare al campo gravitazionale?

Se l'asse di ingresso è perpendicolare al campo gravitazionale, l'accelerometro misura la componente, lungo l'asse di ingresso, dell'accelerazione inerziale (dovuta alle **variazioni di moto del velivolo**).

Misura cioè¹¹:

$$a_{inerz} = \frac{dV}{dt} + a_{trascinamento} + a_{Coriolis}$$

TERNA DI ACCELEROMETRI

Un singolo accelerometro misura l'accelerazione lungo una direzione.

Perciò, per poter individuare anche la direzione nello spazio tridimensionale del vettore accelerazione, sono necessari **tre accelerometri**.

Ciascuno dei tre accelerometri ha l'**asse di ingresso** disposto **lungo uno degli assi di una terna** cartesiana e quindi rileva la componente dell'accelerazione lungo tale asse.

La somma delle componenti a_x , a_y , a_z rappresenta il vettore accelerazione.

¹¹ dove:

$$a_{trascinamento} = \omega_{terra}^2 \cdot R \cdot \cos \theta \quad \text{con} \quad R \cdot \cos \theta = \text{raggio del parallelo terrestre (nel luogo)}$$

$$a_{coriolis} = 2 \cdot \omega_s \wedge v \quad \text{con} \quad \begin{aligned} \omega_s &= \omega + \omega_{aerom} \\ \wedge &= \text{prodotto vettoriale} \\ v &= \text{velocità del velivolo} \end{aligned}$$

Quest'ultima accelerazione (Accelerazione Di Coriolis)

è dovuta al movimento, con velocità angolare ω , della terra rispetto allo spazio inerziale.

ACCELEROMETRO PENDOLARE

E' uno degli accelerometri meccanici più diffusi ed è formato da una **massa m** che **si sposta dalla sua posizione di equilibrio** quando viene **sollecitata da un'accelerazione non gravitazionale**.

L'**accelerazione** determina una **coppia** il cui **momento**

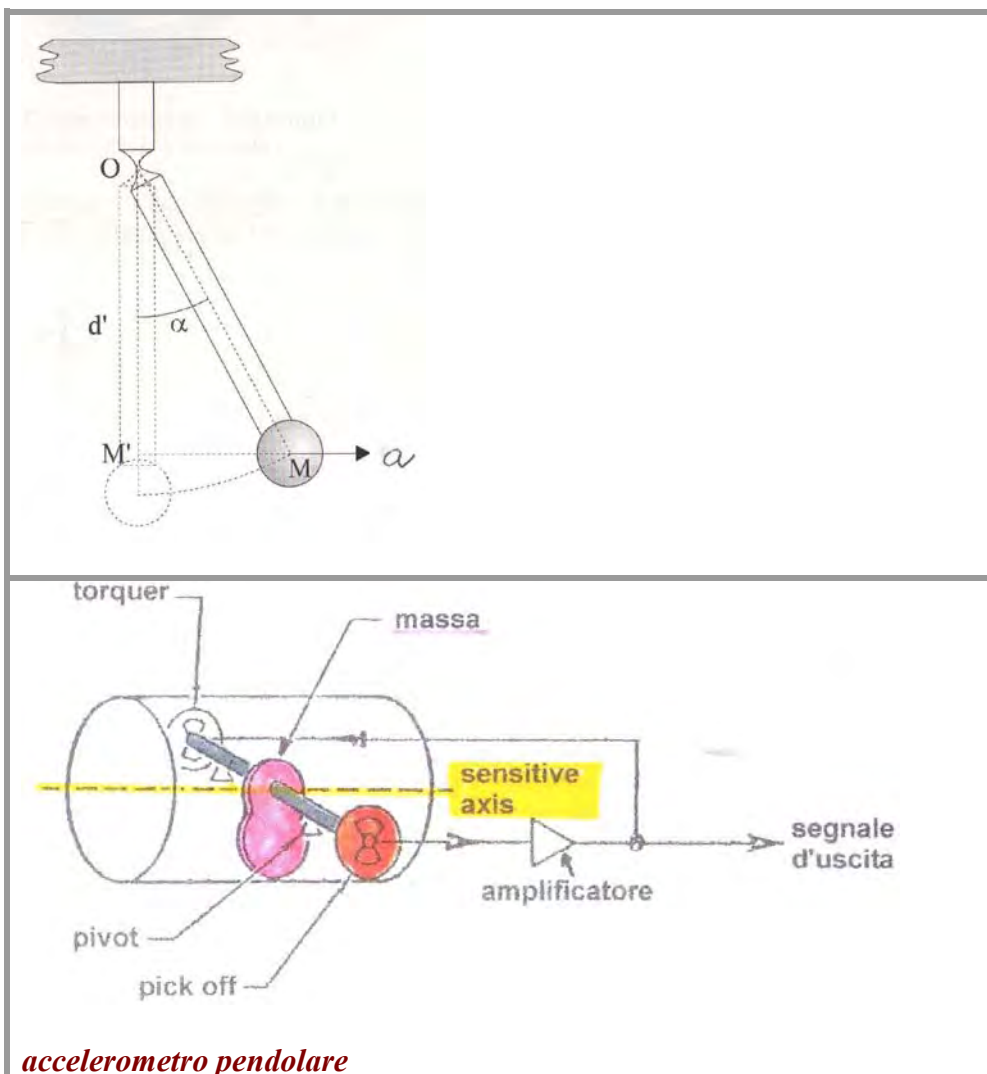
$$M = m \cdot a_{\text{inerz}} \cdot d$$

$$(d' = \text{braccio della coppia} = d \cdot \cos \alpha)$$

risulta **proporzionale all'accelerazione stessa**.

Un trasduttore (pickoff) produce poi una corrente proporzionale alla coppia e quindi all'accelerazione.

Per mantenere costante la sensibilità dell'accelerometro si fa in modo che il pendolo ritorni sempre alla posizione di equilibrio mediante un torquer, cioè un generatore di coppia.



ACCELEROMETRI MINIATURIZZATI “Q-flex”

Negli accelerometri noti come “Q-flex” la massa è una sottile lamina di metallo, inserita fra le armature di un condensatore.

Quando la lamina subisce uno spostamento, la capacità del condensatore varia. Ciò determina la nascita di una corrente che percorre una bobina e l’insorgere di un campo magnetico che riporta la lamina nella posizione originaria.

La corrente risulta proporzionale all’accelerazione che la ha prodotta e dà luogo a un’uscita digitale. Il condensatore svolge il ruolo di pickoff.



TORQUERS (GENERATORI DI COPPIA)

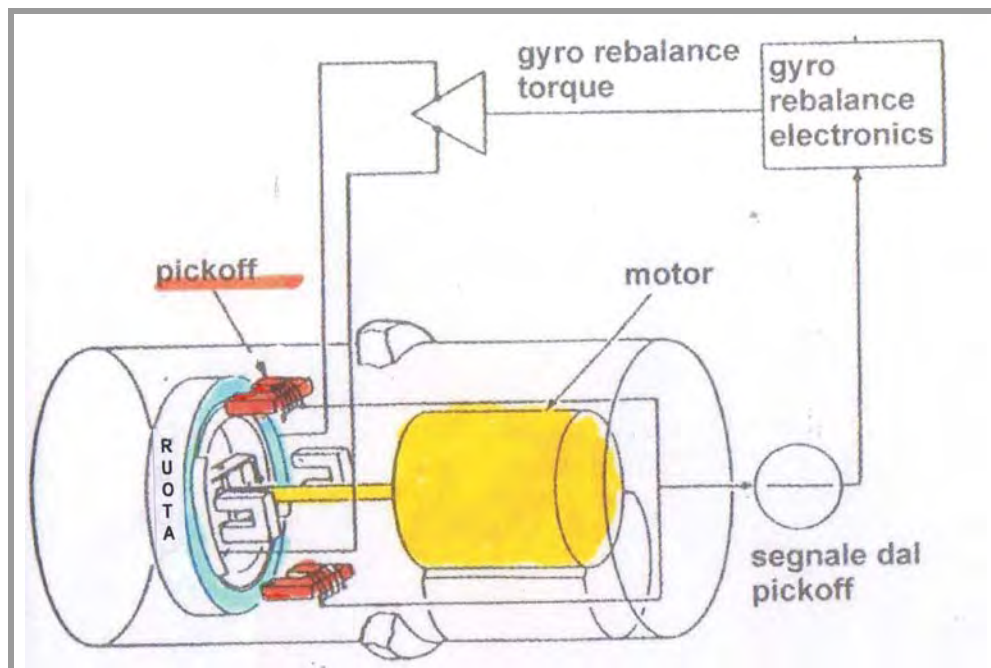
I torquers sono dispositivi che **generano una coppia di forze** le quali agiscono sulla ruota (rotore) del giroscopio.

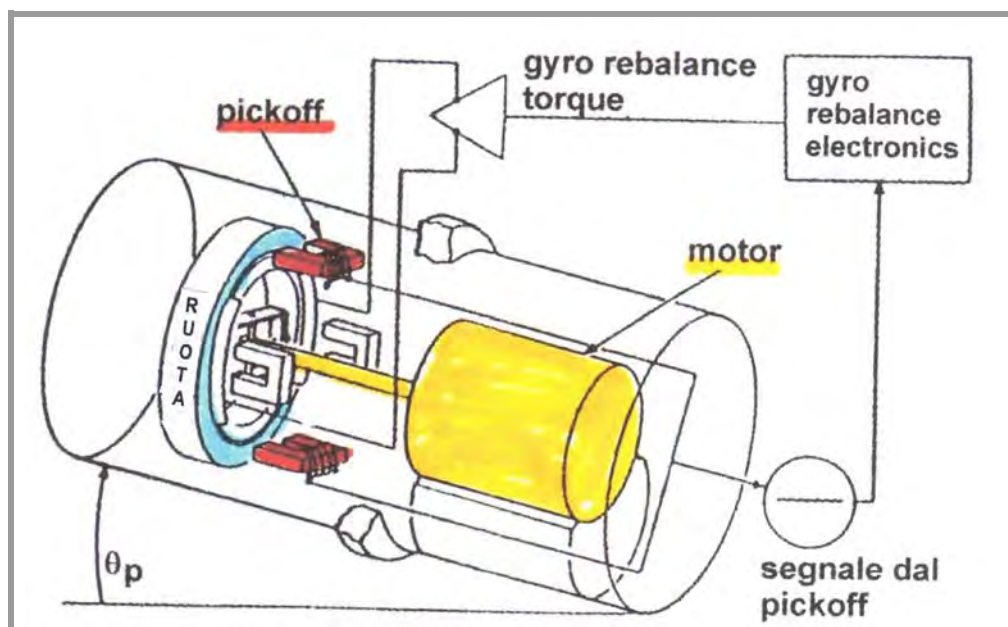
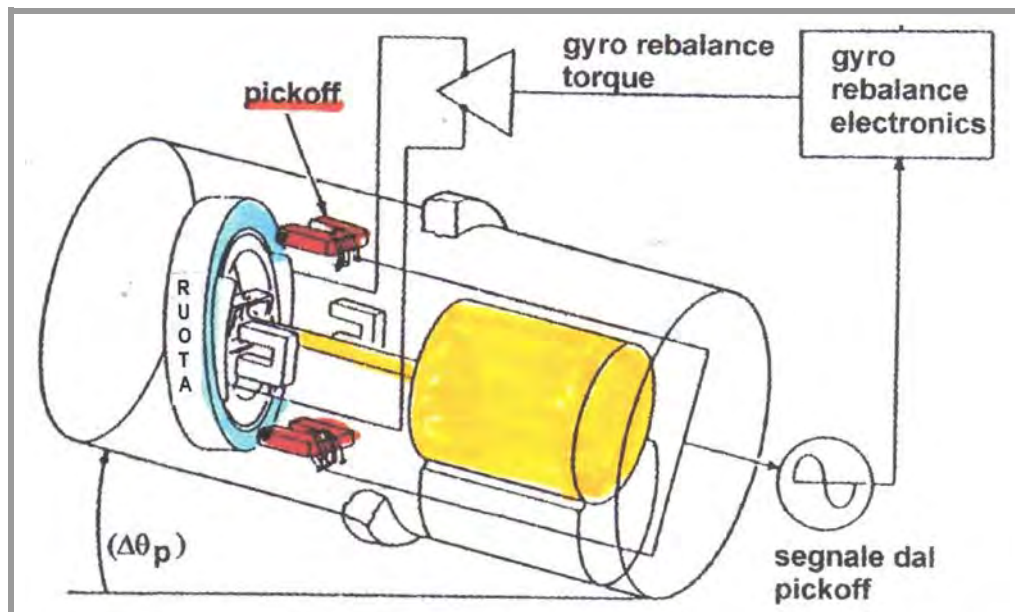
I torquers sono pilotati dai segnali di uscita dei pickoff (cioè dei trasduttori di spostamento angolare) e consentono il ri-orientamento della ruota (rotore) del giroscopio e quindi della piattaforma.

I pickoff rilevano la posizione della ruota del giroscopio e, quando c'è uno spostamento relativo fra essi e la ruota, generano segnali elettrici che attivano i torquers.

I pickoff sono quindi dei trasduttori, mentre i torquers sono attuatori.

Nelle figure che seguono si vede che, quando i pickoff rilevano uno spostamento della ruota giroscopica rispetto ad essi (per cui i pickoff vengono a trovarsi a differenti distanze dalla ruota), inviano ai circuiti di controllo dei torquers un segnale che sollecita i torquers a ripristinare la posizione originaria della ruota.





FUNZIONAMENTO DI UNA PIATTAFORMA INERZIALE DI TIPO “GIMBALLED” O “A SOSPENSIONE CARDANICA” O “ASSERVITA”

I giroscopi della piattaforma tenderebbero a mantenere la stessa orientazione che avevano **nell'istante iniziale** del viaggio, cioè a conservare l'orientazione rispetto a un **riferimento inerziale**, come quello delle stelle fisse.

Il giroscopio verticale tenderebbe a mantenere l'asse di spin nella direzione della verticale del luogo iniziale del viaggio, mentre il giroscopio direzionale (orizzontale) tenderebbe a mantenere l'asse di spin su un piano parallelo alla superficie terrestre (sempre del luogo iniziale del viaggio), in una direzione prescelta, per esempio il nord.

Nella navigazione inerziale invece è necessario che i giroscopi indichino **istante per istante** la verticale **del luogo nel quale ci si trova** (per quanto riguarda il giroscopio verticale) e la direzione prescelta, per esempio il nord (per quanto riguarda il giroscopio direzionale).

Perché è necessario che i giroscopi siano orientati in funzione del luogo in cui ci si trova e non in base al luogo di partenza (e quindi non in base al riferimento inerziale)?

Perché gli accelerometri (dai quali, mediante l'operazione matematica di integrazione, si ottengono la latitudine e la longitudine del luogo in cui ci si trova) devono avere i loro assi sensibili orientati perpendicolarmente alla superficie terrestre del luogo in cui ci si trova.

in questo modo l'accelerometro risulta sollecitato solo dalle **accelerazioni inerziali**, dovute **al moto del velivolo** e il suo asse sensibile risulta perpendicolare al campo gravitazionale.

Siccome quindi i giroscopi tendono a mantenere l'orientazione (che noi abbiamo imposto nell'istante iniziale del viaggio) rispetto al riferimento inerziale, mentre ci interessa che assumano, istante per istante, l'orientazione relativa al luogo in cui ci si trova, allora **è necessario che il sistema di navigazione inerziale apporti delle correzioni alle posizioni dei giroscopi** in modo che essi siano riferiti al luogo in cui ci si trova.

E' proprio in base all'entità delle correzioni da effettuare istante per istante alla posizione della piattaforma per adeguarla alla verticale del luogo (e alla direzione nord nel piano della verticale), che vengono calcolate la velocità del velivolo e le sue coordinate.

Chi apporta le correzioni all'orientazione della piattaforma?

La correzione di posizione della piattaforma in funzione del luogo in cui ci si trova è effettuata dai torquers che, pilotati dagli accelerometri, fanno ruotare i giroscopi, e quindi la piattaforma.

Lo scostamento fra l'orientazione desiderata istante per istante per i giroscopi (che è quella relativa, istante per istante, al luogo in cui ci si trova) e la posizione che essi hanno assunto all'istante iniziale del viaggio (e che tenderebbero a conservare) varia istante per istante ed è determinata dai seguenti fattori:

- **rotazione terrestre**
- **movimento del velivolo**
- **accelerazione centripeta**
- **accelerazione di Coriolis.**

GIROSCOPIO DI RATEO

Se il segnale di uscita del giroscopio è un segnale elettrico, possiamo dire che il **giroscopio di rateo** (Rate-Giroscope) è un **trasduttore di velocità di rotazione**, cioè un trasduttore di **velocità angolare**.

Il **segnale di uscita** del Rate-Giroscope:

- è una **tensione elettrica** proporzionale al numero di gradi angolari al secondo, relativi a un fenomeno di rotazione
- è calibrato in mV per gradi al secondo, nel senso che **varia di un certo numero di mV** (tipicamente 50 mV) **per ogni grado di rotazione** descritto in un secondo
- è generato da un dispositivo, inserito nel giroscopio stesso, che è detto “pickoff”.

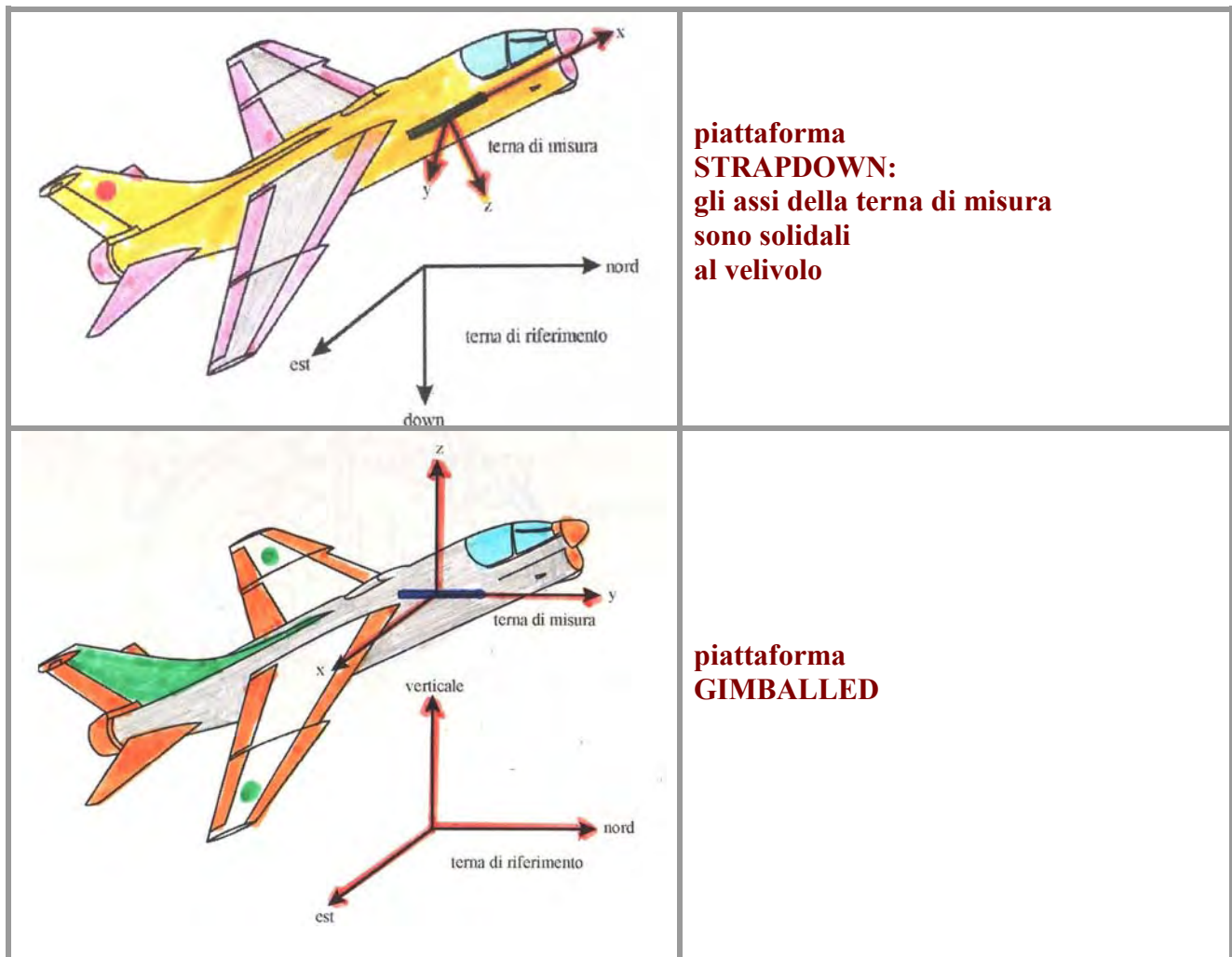
Se, per esempio, un'automobile compie un percorso circolare di 360° in 15 secondi, allora la velocità angolare dell'auto è di $360/15$ di gradi al secondo, cioè di 24 gradi/secondo.

A questa rotazione (assumendo una variazione unitaria di tensione di 50mV) corrisponderà una variazione di tensione di $50 \cdot 24 = 1200$ mV

PIATTAFORMA INERZIALE “STRAPDOWN” O “ANALITICA”

Le piattaforme inerziali strapdown sono di concezione più moderna delle tradizionali “gimballed” e sono in uso dal 1980 circa.

Mentre nelle piattaforme gimballed i giroscopi sono applicati alla struttura del velivolo attraverso sospensioni cardaniche¹² (che consentono almeno un grado di libertà), nelle strapdown **i giroscopi e gli accelerometri sono vincolati rigidamente** alla struttura del **velivolo** e fanno riferimento, per le misurazioni, a una **terna di assi solidale al velivolo** stesso (**terna XYZ: X = avanti, Y = destra, Z = basso**).



Nelle piattaforme strapdown quindi le casse (involucri) dei giroscopi sono solidali all'aeromobile e il **rotore** ruota rispetto alla cassa, **sollecitato dalla velocità angolare dell'aeromobile**, che determina un moto di **precessione forzata**.

In queste condizioni gli accelerometri misurano le componenti dell'accelerazione non gravitazionale rispetto a un riferimento solidale al velivolo e perciò forniscono direttamente le accelerazioni lineari dell'aeromobile.

¹² e quindi sono isolati dal moto angolare dell'aeromobile in quanto la piattaforma è stabilizzata secondo gli assi di una terna di riferimento.

Il movimento dei rotori è –anche in questo caso- rilevato dai pickoff.

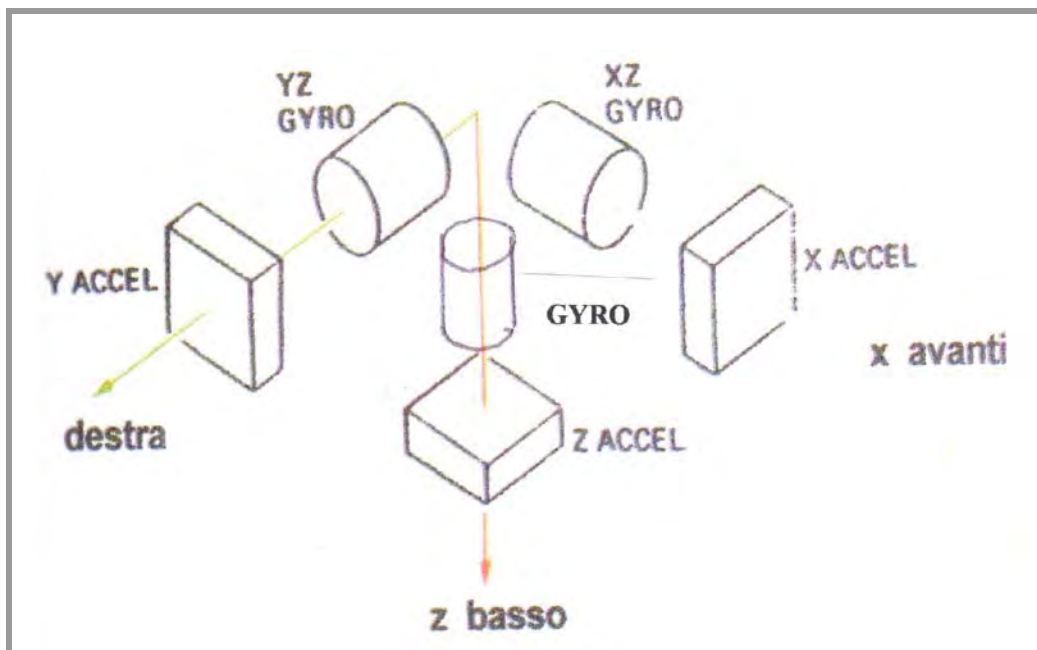
L'informazione sul moto passa quindi dal pickoff al **generatore di coppia (torquer)** che **riporta il rotore nella posizione originaria**.

Ai fini della **determinazione della posizione**, è **necessario** conoscere le componenti dell'accelerazione non gravitazionale rispetto alla terna di riferimento “esterna” **NORD-EST-DOWN (NED)**, mentre le componenti di accelerazione che realmente sono state misurate sono quelle relative alla terna XYZ solidale al veicolo.

Si rende quindi necessaria una TRASFORMAZIONE DI COORDINATE, cioè una PROIEZIONE, **sugli assi della terna “esterna” di riferimento**

NORD-EST-DOWN, delle accelerazioni misurate rispetto agli assi del velivolo.

La trasformazione è effettuata dal **CALCOLATORE DI NAVIGAZIONE** che riceve, con stretta periodicità (per esempio ogni dieci millisecondi) dai pickoff le informazioni sui giroscopi.



LABORATORIO

INDICE

LEGGI FONDAMENTALI E COMPONENTI ELETTRONICI DI BASE

RILIEVO DELLA CARATTERISTICA TENSIONE-CORRENTE DI UN RESISTORE (LEGGE DI OHM) E DI UNA LAMPADINA A FILAMENTO

VERIFICA SPERIMENTALE DEL PRINCIPIO DI THEVENIN

RILIEVO STATICO DELLA CARATTERISTICA DIRETTA DI DIODI RADDRIZZATORI O SCHOTTKY

RILIEVO STATICO DELLA CARATTERISTICA INVERSA DI DIODI RADDRIZZATORI O SCHOTTKY

AMPLIFICATORI A TRANSISTOR

DIMENSIONAMENTO E VERIFICA DEL FUNZIONAMENTO DELLA RETE POLARIZZAZIONE A *DUE* ALIMENTATORI PER BJT, FUNZIONANTE IN ZONA LINEARE

DIMENSIONAMENTO E VERIFICA DEL FUNZIONAMENTO DI UNA RETE DI POLARIZZAZIONE AUTOMATICA A PARTITORE PER BJT

RILIEVO DELLA RISPOSTA IN FREQUENZA DI UN AMPLIFICATORE A BJT A EMTTITORE COMUNE

ELETTRONICA DIGITALE

REALIZZAZIONE E VERIFICA DEL FUNZIONAMENTO DI UN CONTATORE ASINCRONO MODULO 4 BASATO SU FLIP-FLOP “D”

REALIZZAZIONE E VERIFICA DEL FUNZIONAMENTO DI UN CONTATORE ASINCRONO MODULO 9

CONFIGURAZIONI AMPLIFICATRICI DELL'AMPLIFICATORE OPERAZIONALE

VERIFICA DEL FUNZIONAMENTO IN CONTINUA DI UN SOTTRATTORE (AMPLIFICATORE DIFFERENZIALE) CON OPERAZIONALE

GENERATORI DI SEGNALE

REALIZZAZIONE E VERIFICA DEL FUNZIONAMENTO DI UN GENERATORE D'ONDA RETTANGOLARE (BASATO SU TIMER 555) A FREQUENZA REGOLABILE

**VCO
(OSCILLATORE CONTROLLATO IN TENSIONE) BASATO SU OPERAZIONALI LM 358 AD ALIMENTAZIONE SINGOLA**

ESPERIENZE DI TELECOMUNICAZIONI

MISURA DELL'INDICE DI MODULAZIONE AM CON IL METODO DEL TRAPEZIO

RADIORICEVITORE A CIRCUITO RIFLESSO

MODULATORE E DEMODULATORE FM BASATI SU DUE PLL DIGITALI INTEGRATI 4046

CIRCUITI BASATI SUL CONTROLLO DELLA POTENZA ELETTRICA

**REGOLAZIONE DI VELOCITA' DI UN MOTORINO CC CON TECNICA PWM
(Pulse Width Modulation, Modulazione di durata degli impulsi)**

**LAMPADA COMANDATA DA UN INTERRUTTORE CREPUSCOLARE
CON FOTORESISTENZA**

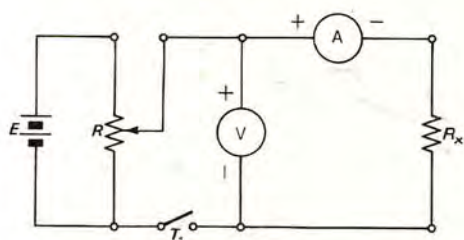
**REALIZZAZIONE E VERIFICA DEL FUNZIONAMENTO DI UN INVERTER CON DUE
BJT A SEMIPONTE**

SISTEMA DI CONTROLLO ON/OFF DELLA TEMPERATURA DI UN CONTENITORE

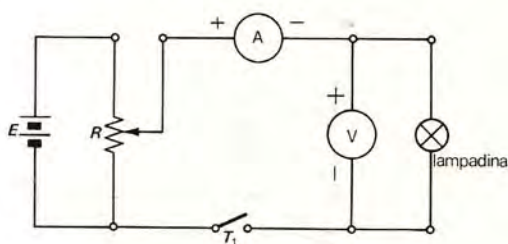
**VARIATORE DI LUMINOSITA' (DIMMER) PER LAMPADA A INCANDESCENZA
CON DIAC E TRIAC**

LEGGI FONDAMENTALI E COMPONENTI ELETTRONICI DI BASE

RILIEVO DELLA CARATTERISTICA TENSIONE-CORRENTE DI UN RESISTORE (LEGGE DI OHM) E DI UNA LAMPADINA A FILAMENTO



a) Circuito per i resistori.



b) Circuito per le lampade.

STRUMENTI

- ALIMENTATORE
- MULTIMETRO DIGITALE
- TASTO

COMPONENTI

- RESISTENZE: $1 \times 1\text{K}\Omega$, TRIMMER $10\text{K}\Omega$
- LAMPADINA DA 25 V A FILAMENTO DI TUNGSTENO

SCOPO DELLA MISURA

Lo scopo dell'esperienza è di tracciare le curve dell'andamento della corrente in funzione della tensione applicata sia nel caso di una resistenza di valore fisso che di una lampadina da 25V a filamento.

Ciò che ci aspettiamo è che:

- la curva caratteristica $I=I(V)$ della resistenza sia lineare, cioè sia una retta (conformemente alla legge di Ohm: $I=V/R$)
- la curva caratteristica $I=I(V)$ della lampadina a filamento non sia lineare, ma sia una curva con pendenza decrescente al crescere della tensione applicata, dal momento che le lampade

a filamento di tungsteno presentano un valore di resistenza interna crescente con la tensione applicata (e con la temperatura).

I CIRCUITI DI MISURA

- Il circuito A per la misura relativa alla resistenza è nella configurazione con voltmetro a monte, adatto alla misura voltamperometrica di grosse resistenze.
- Il circuito B per la misura relativa alla lampadina è nella configurazione con voltmetro a valle, adatto alla misura voltamperometrica di piccole resistenze (la resistenza equivalente di tali lampadine è bassa).

PROCEDIMENTO

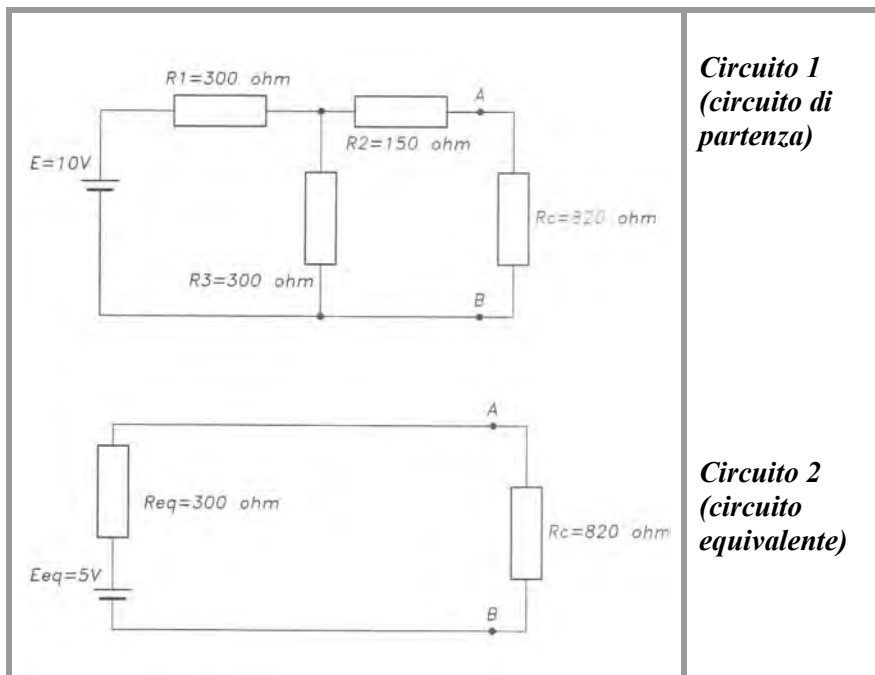
Il procedimento di misura è uguale per i due casi (resistore e lampadina) ed è il seguente:

- Si realizzano i circuiti 1 e 2
- Partendo da 0V, si aumentano gradualmente, mediante il trimmer, i valori della tensione applicata al circuito
- Per ogni valore di tensione applicata al circuito si misurano:
 - la tensione ai capi della resistenza (o della lampadina) mediante il voltmetro
 - la corrente che scorre nella resistenza (lampadina) mediante l'amperometro
- Si trascrivono i valori ottenuti in una tabella nella quale si aggiunge il valore della resistenza equivalente (della resistenza o della lampadina) calcolato come V/I
- Si riportano le coppie di valori tensione-corrente presenti nella tabella su un grafico avente sull'asse orizzontale le tensioni e sull'asse verticale le correnti.

RESISTORE (1KΩ)		
TENSIONE [V]	CORRENTE [mA]	RESISTENZA EQUIVALENTE $R=V/I$ [Ω]
5		
10		
15		
20		
25		

LAMPADINA DA 25 V A FILAMENTO DI TUNGSTENO		
TENSIONE [V]	CORRENTE [mA]	RESISTENZA EQUIVALENTE $R=V/I$ [Ω]
5		
10		
15		
20		
25		

VERIFICA SPERIMENTALE DEL PRINCIPIO DI THEVENIN



STRUMENTI

- ALIMENTATORE
- MULTIMETRO DIGITALE

COMPONENTI

- RESISTENZE: 3x300Ω, 1x150Ω, 2x820Ω, TRIMMER 10K Ω

PREMESSA

Il principio di Thevenin applicato fra i punti A e B del circuito di partenza (circuito 1) afferma che la parte di circuito a destra della retta AB equivale alla serie di un generatore equivalente E_{eq} e di una resistenza equivalente R_{eq} , come rappresentato nel circuito 2.

Il valore del generatore equivalente è la tensione V_{AB0} a vuoto, cioè la tensione che si misura fra A e B una volta staccata la resistenza R_C . La resistenza equivalente è quella che si misura fra A e B, staccata R_C e cortocircuitato il generatore E. In base al principio di Thevenin si ha quindi:

$$\bullet \quad E_{eq} = V_{AB0} = E \cdot \frac{R_3}{R_1 + R_3} = 10 \cdot \frac{300}{300 + 300} = 5V$$

(Applicando il partitore di tensione e considerando che su R_2 non scorre corrente)

$$\bullet \quad R_{eq} = \frac{R_1 \cdot R_3}{R_1 + R_3} + R_2 = 150 + 150 = 300\Omega$$

La verifica del principio di Thevenin consiste quindi nell'appurare che il circuito 1 e il circuito 2 risultano equivalenti riguardo al regime di tensioni e di correnti che si instaura a destra della retta AB, e quindi nel verificare che:

- la tensione V_{AB} del circuito 1 (con la R_C staccata) è uguale alla V_{AB} del circuito 2
- la corrente che scorre nella resistenza di carico R_C è la stessa nei due circuiti.

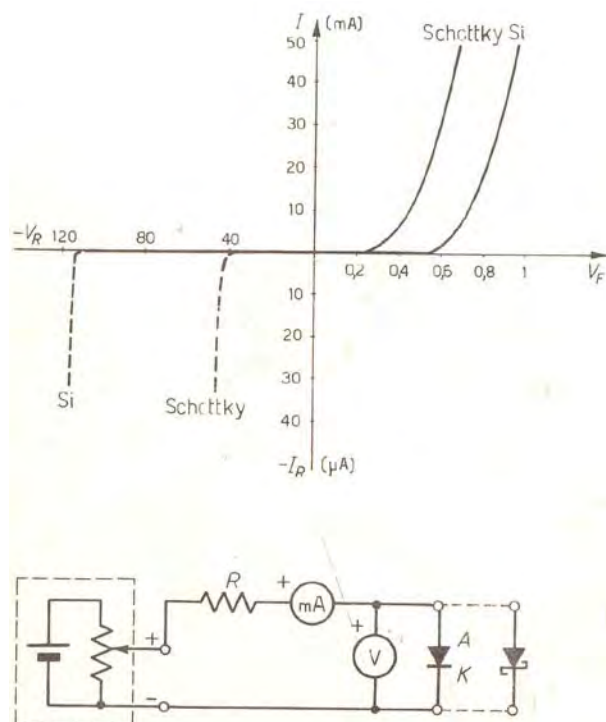
PROCEDIMENTO

- realizzare il circuito 1
- realizzare il circuito 2 utilizzando un alimentatore da 5V o un alimentatore di tensione superiore insieme con un trimmer in modo che forniscano 5V
- nel circuito 1 misurare:
 - la corrente I su R_C (con un amperometro in serie a R_C)
 - la tensione V_{AB} fra A e B (con un voltmetro in parallelo a R_C)
 - la tensione V_{AB0} a vuoto, cioè la tensione che si misura fra A e B una volta staccata la resistenza R_C
- nel circuito 2 misurare:
 - la corrente I' su R_C (con un amperometro in serie a R_C)
 - la tensione V'_{AB} fra A e B (con un voltmetro in parallelo a R_C)
 - la tensione E_{eq} del generatore equivalente

Deve risultare:

- $V_{AB0} = E_{eq}$
- $V_{AB} = V'_{AB}$
- $I = I'$

RILIEVO STATICO DELLA CARATTERISTICA DIRETTA DI DIODI RADDRIZZATORI O SCHOTTKY



In alto, nel semipiano destro del grafico, la curva caratteristica diretta del diodo

Questa esperienza si propone di ricavare un certo numero di punti (V_F , I), che ci permettano di tracciare la curva caratteristica diretta di un diodo raddrizzatore oppure di un diodo Schottky.

La curva caratteristica diretta di un diodo (rappresentata nel semipiano destro della figura) descrive **l'andamento della corrente diretta** (detta I oppure I_F , o I_D) del diodo in funzione di valori crescenti e positivi della tensione diretta

$$V_F = V_{AK} = V_A - V_K,$$

ai capi del diodo, con V_A maggiore o uguale a V_K .

Si parla invece di curva caratteristica inversa (rappresentata nel semipiano sinistro della figura) quando i valori della corrente che scorre nel diodo sono ricavati in corrispondenza di valori negativi della

$$V_{AK} = V_A - V_K$$

cioè per V_A minore o uguale a V_K .

STRUMENTI

- **ALIMENTATORE REGOLABILE da 0V A 2V**
- **MILLIAMPEROMETRO**
- **VOLTMETRO**

COMPONENTI

- **DIODO AL SILICIO (1N4148 o 1N914 o 1N4007)
oppure DIODO SCHOTTKY BAT86 O BYS21**
- **R= 100Ω**

PREMESSA: LA TENSIONE DI SOGLIA

Fra le possibili definizioni di tensione di soglia scegliamo (per l'esecuzione della misura) la seguente: **la tensione di soglia di un diodo è il valore della tensione V_{AK} fra anodo e catodo in corrispondenza del quale la corrente diretta tocca il valore di 1mA.**

IL CIRCUITO DI MISURA

- lo schema di misura voltamperometrica è “con voltmetro a valle”, così come è opportuno fare nel caso di misura di piccole resistenze. Stiamo lavorando infatti con un diodo in polarizzazione diretta, la cui resistenza interna, una volta superata la tensione di soglia, è molto bassa
- La resistenza R (100Ω) inserita nel circuito serve per la limitazione della corrente che scorre nel diodo
 - se non si dispone di un alimentatore regolabile si può usare un alimentatore fisso (p. es. 12V) con in parallelo un trimmer (p. es. 10KΩ), dal cui cursore si preleva la tensione da applicare al circuito.

PROCEDIMENTO (RIFERITO AL DIODO AL SILICIO)

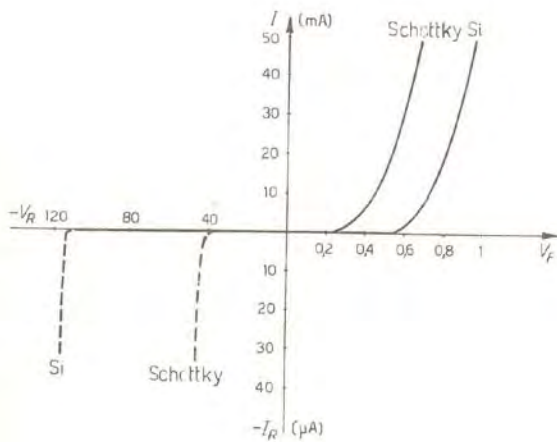
Collegati gli strumenti come indicato si procede nel modo seguente:

- si applica al circuito, mediante un generatore, una tensione **continua (perciò si parla di rilievo “statico”)**, inizialmente nulla e poi di **ampiezza** positiva **crescente** (per esempio: 0V, 0,1V, 0,2V.....) fino a raggiungere o superare lievemente le tensioni di soglia e cioè (0,5÷0,6)V per il diodo al silicio e (0,2÷0,3)V per lo Schottky.
- possono essere scelti valori di V_{AK} a **intervalli di 0,1V** fino all'approssimarsi della tensione di soglia, in prossimità della quale è opportuno, per una maggior precisione, ridurre gli intervalli, per esempio a **0,02V (20mV)**
- per ogni valore prescelto di V_{AK} (rilevato sul voltmetro), si legge sul milliamperometro il corrispondente valore di corrente I_F .
- i valori di V_{AK} e I_F vengono riportati su una tabella
- i valori di V_{AK} e I_F vanno riportati su un diagramma avente in ascissa i valori di V_{AK} (in V) e in ordinata i valori di I_F (in mA).
- volendo, possiamo anche misurare la resistenza diretta “statica” del diodo ($R_F = V_{AK}/I_F$) in corrispondenza dei diversi valori di V_{AK} e I_F .

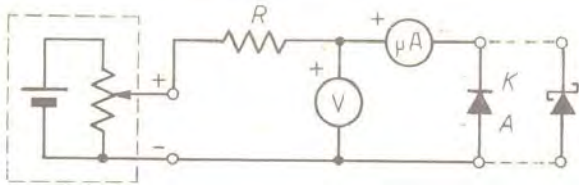
TENSIONE DIRETTA V_{AK} [V]	CORRENTE DIRETTA I_F [mA]	RESISTENZA DIRETTA R_F=V_{AK}/I_F [Ω]
0,10		
0,20		
0,30		
0,40		
0,50		
0,52		
0,54		
0,56		
0,58		
0,60		

TENSIONE DI SOGLIA:

RILIEVO STATICO DELLA CARATTERISTICA INVERSA DI DIODI RADDRIZZATORI O SCHOTTKY



Nel semipiano sinistro del grafico è raffigurata la curva caratteristica inversa del diodo. Con “ $-V_R$ ” è indicata la tensione inversa $V_{AK} = V_A - V_K$, con il segno cambiato.



Circuito di misura

Ricaveremo un certo numero di punti (V_{AK} , I_R), al fine di tracciare la curva caratteristica inversa di un diodo raddrizzatore oppure di un diodo Schottky.

La curva caratteristica inversa di un diodo (rappresentata nel semipiano sinistro della figura) descrive l'andamento della **corrente inversa (I_R) del diodo in funzione di valori negativi e crescenti in valore assoluto della tensione:**

$V_{AK} = V_A - V_K$,
ai capi del diodo, con V_A minore di V_K .

STRUMENTI

- **ALIMENTATORE (O ALIMENTATORI IN SERIE) in grado di erogare la tensione di breakdown del diodo**
- **MILLIAMPEROMETRO**
- **VOLTMETRO**

COMPONENTI

- **DIODO AL SILICIO**
- **1N4148 (con $V_{Rmax} = 75V$) o 1N914 (con $V_{Rmax} = 75V$)
oppure
DIODO SCHOTTKY
BAT86 (con $V_{Rmax} = 50V$) o BYS21 (con $V_{Rmax} = 45V$)**
- **$R = 47\text{ K}\Omega$**

PREMESSA: IL BREAKDOWN

Un diodo raddrizzatore, inversamente polarizzato non conduce (o, più esattamente, è percorso da una piccolissima corrente inversa) fino a quando non viene raggiunta la tensione di breakdown. Tipicamente questa tensione misura, in valore assoluto, dai 50V in su. E' quindi probabile che, con la strumentazione di cui dispone un laboratorio scolastico, questo valore non possa essere raggiunto, per cui ci si dovrà limitare a misurare i valori, pressoché nulli, della corrente prima del breakdown. Se invece si è in grado di toccare la tensione di breakdown, tale valore può essere individuato dalla deviazione che si produce nel microamperometro quando, aumentando lentamente la tensione di polarizzazione inversa, si raggiunge il valore di breakdown.

IL CIRCUITO DI MISURA

- lo schema di misura voltamperometrica è “con voltmetro a monte”, così come è opportuno fare nel caso di misura di grosse resistenze. Stiamo lavorando infatti con un diodo in polarizzazione inversa, la cui resistenza interna, prima del verificarsi del breakdown, è molto alta.
- La resistenza R inserita nel circuito serve per la limitazione della corrente che scorre nel diodo.

PROCEDIMENTO (RIFERITO AL DIODO AL SILICIO)

Collegati gli strumenti come indicato si procede nel modo seguente

- si applica al generatore una tensione **continua (perciò si parla di rilievo “statico”)**, inizialmente nulla e poi di **ampiezza** negativa **crescente** in valore assoluto(per esempio: -10V, -20V, -30V.....) fino a raggiungere, se possibile, il breakdown
- per ogni valore prescelto di V_{AK} (rilevato sul voltmetro), si legge sul microamperometro il corrispondente valore di corrente .
- i valori di V_{AK} e I_R vengono riportati su una tabella
- i valori di V_{AK} e I_R vanno riportati su un diagramma avente in ascissa i valori di V_{AK} (in V) e in ordinata i valori di I_R (in μA).

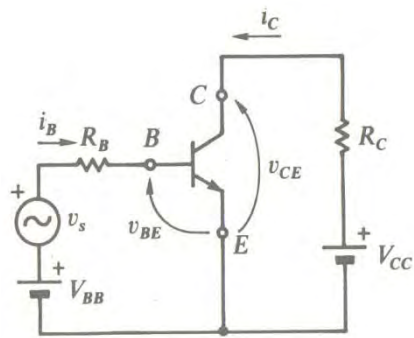
TENSIONE INVERSA V_{AK} [V]	CORRENTE INVERSA I_R [μA]
-10	
-20	
-30	
-40	
-50	
.....	

TENSIONE DI BREAKDOWN:

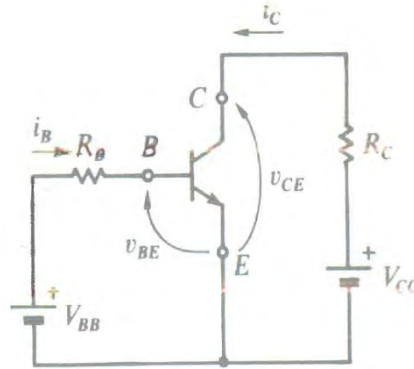
AMPLIFICATORI A TRANSISTOR

DIMENSIONAMENTO E VERIFICA DEL FUNZIONAMENTO DELLA RETE POLARIZZAZIONE A DUE ALIMENTATORI PER BJT, FUNZIONANTE IN ZONA LINEARE

(in modo che sia: $V_{CE} = V_{CC}/2$)



*Amplificatore a BJT a
emettitore comune:
a sinistra il circuito di ingresso,
a destra il circuito di uscita.
Il segnale di uscita è sfasato di
 180° rispetto a quello di
ingresso.*



*BJT a emettitore comune: la sola
rete di polarizzazione*

STRUMENTI:

- ALIMENTATORI ($V_{BB}=5V$; $V_{CC}=12V$)
- MULTIMETRO DIGITALE

COMPONENTI

- BJT 2N3903
- $R_C=120\ \Omega$
- $R_B=5180\ \Omega$ (utilizzare il valore commerciale più vicino: $5100\ \Omega$)

IL DIMENSIONAMENTO**Grandezze ASSEGNATE o SCELTE**

(o che possono essere ritenute note):

- I_C
- V_{CC}
- V_{BB}
- h_{FE}
- $V_{BE}\approx 0,7V$ (Si può ritenere che in zona lineare la tensione ai capi della giunzione base-emettitore non si discosti da questo valore)

Grandezze DA DETERMINARE:

- R_C
- I_B
- R_B

PROCEDIMENTO DI CALCOLO

1. DETERMINAZIONE DI V_{CE} :

si impone che si abbia $V_{CE} = \frac{V_{CC}}{2}$

2. CALCOLO DI R_C

Fissata la V_{CE} , si applica il 2° principio di Kirchhoff alla maglia di uscita, e si esplicita rispetto a R_C

$$V_{CE} + R_C \cdot I_C = V_{CC}$$

$$R_C \cdot I_C = V_{CC} - V_{CE}$$

$$R_C = \frac{V_{CC} - V_{CE}}{I_C}$$

3. CALCOLO DI I_B

Dalla definizione dell'amplificazione di corrente $h_{FE} = \frac{I_C}{I_B}$ si ottiene:

$$I_B = \frac{I_C}{h_{FE}}$$

4. CALCOLO DI R_B

Si applica il 2° principio di Kirchhoff alla maglia di ingresso e si esplicita rispetto a R_B :

$$V_{BB} - R_B \cdot I_B - V_{BE} = 0$$

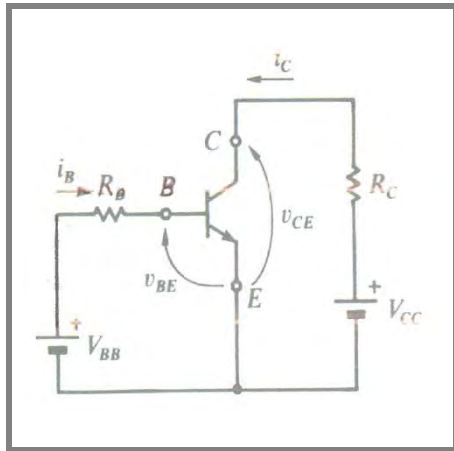
$$R_B \cdot I_B = V_{BB} - V_{BE}$$

$$R_B = \frac{V_{BB} - V_{BE}}{I_B}$$

$$R_B = \frac{V_{BB} - 0,7}{I_B}$$

APPLICAZIONE DEL PROCEDIMENTO AL NOSTRO PROGETTO

Determinare i valori di R_C , I_B ed R_B della rete di polarizzazione in figura, con BJT 2N3903, con $V_{BB}=5V$ e con $V_{CC}=12V$, in modo che la corrente di collettore risulti di 50 mA.



Dal foglio-dati, sappiamo che per il BJT 2N3903 è:

- $h_{FE} = 30$, con $I_C = 100$ mA
- $h_{FE} = 60$, con $I_C = 50$ mA

Possiamo inoltre ipotizzare: $V_{BE} \approx 0,7V$.

1. DETERMINAZIONE DI V_{CE} :

si impone che si abbia $V_{CE} = \frac{V_{CC}}{2} = \frac{12}{2} = 6V$

2. CALCOLO DI R_C

Fissata la V_{CE} , si applica il 2° principio di Kirchhoff alla maglia di uscita,

$$V_{CE} + R_C \cdot I_C = V_{CC} \quad \Rightarrow \quad R_C \cdot I_C = V_{CC} - V_{CE}$$

$$R_C = \frac{V_{CC} - V_{CE}}{I_C} = \frac{12 - 6}{50 \cdot 10^{-3}} = 120\Omega$$

3. CALCOLO DI I_B

Dalla definizione dell'amplificazione di corrente $h_{FE} = \frac{I_C}{I_B}$ si ottiene:

$$I_B = \frac{I_C}{h_{FE}} = \frac{50mA}{60} = 0,83mA$$

CALCOLO DI R_B

Si applica il 2° principio di Kirchhoff alla maglia di ingresso:

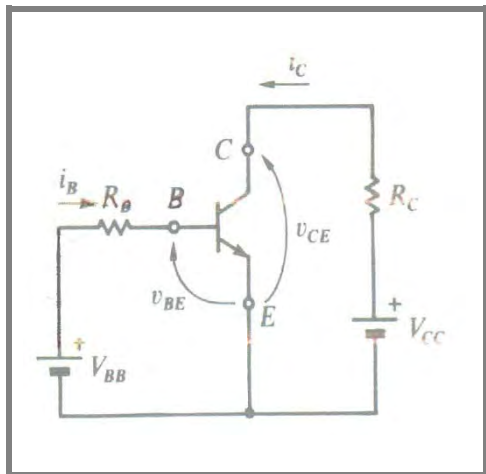
$$V_{BB} - R_B \cdot I_B - V_{BE} = 0 \rightarrow R_B \cdot I_B = V_{BB} - V_{BE} \rightarrow$$

$$R_B = \frac{V_{BB} - V_{BE}}{I_B}$$

$$R_B = \frac{5 - 0,7}{0,83 \cdot 10^{-3}} = \frac{4,3 \cdot 1000}{0,83} = \frac{4300}{0,83} = 5180\Omega$$

ESPERIENZA

Realizzare su breadboard la rete di polarizzazione relativa all'esempio precedente, con i valori ricavati (di R_C , I_B , ed R_B) e verificare, **misurandole col multimetro**, che le tensioni e le correnti di polarizzazione (V_{BE} , V_{CE} , I_B , I_C) siano corrispondenti a quelle scelte o assegnate.



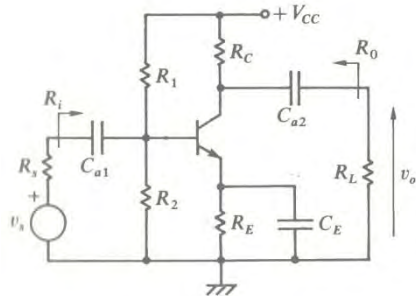
MISURE

V_{BE}		V_{CE}		I_B		I_C	
Valore teorico	Valore misurato	Valore teorico	Valore misurato	Valore teorico	Valore misurato	Valore teorico	Valore misurato
0,7 V		6V		0,83mA		50mA	

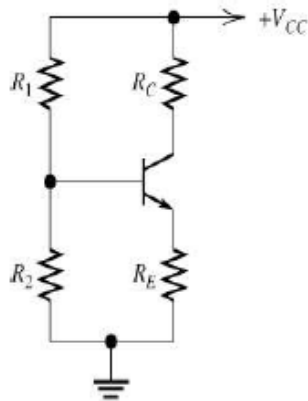
DIMENSIONAMENTO E VERIFICA DEL FUNZIONAMENTO DI UNA RETE DI POLARIZZAZIONE AUTOMATICA A PARTITORE PER BJT

Seguendo il procedimento indicato nel capitolo sul BJT, si progetti una rete di polarizzazione automatica a partitore con:

- BJT BC 107 B, caratterizzato da $h_{FEtyp}=290$
- corrente di collettore: $I_C = 2 \text{ mA}$
- tensione unica di alimentazione: $V_{CC}=12\text{V}$



Amplificatore a BJT a emettitore comune:
a sinistra il circuito di ingresso,
a destra il circuito di uscita.
Il segnale di uscita è sfasato di 180° rispetto a quello di ingresso.



BJT a emettitore comune: la sola rete di polarizzazione

STRUMENTI:

- ALIMENTATORI
- MULTIMETRO DIGITALE

COMPONENTI

- TRANSISTOR: BJT BC 107 B
- $R_C = 2,2 \text{ K}\Omega$
- $R_E = 560 \text{ }\Omega$
- $R_1 = 68 \text{ K}\Omega$
- $R_2 = 15 \text{ K}\Omega$
- $C_E = 560 \text{ nF}$

In base al procedimento suddetto, otteniamo i valori teorici:

NOME DELLA GRANDEZZA	GRANDEZZA	VALORE
V_{CE}	tensione fra collettore ed emettitore	6 V
V_{RE}	tensione sulla resistenza R_E	1,2 V
R_C	resistenza di collettore	2,4 K Ω
R_E	resistenza di stabilizzazione termica di emettitore	600 Ω
I_B	corrente di base	6,8 μA
$I_{\text{partitore}}$	corrente che scorre su R_1 ed R_2	0,136 mA
V_B	potenziale del collettore rispetto a massa	1,9 V
R_1	resistenza “superiore” del partitore	74 K Ω
R_2	resistenza “inferiore” del partitore	14 K Ω
C_E	condensatore di bypass	$\geq 530 \text{ nF}$

REALIZZAZIONE DEL CIRCUITO E MISURE

Realizziamo su breadboard il circuito in esame.

Riguardo ai valori di resistenza calcolati, useremo i valori commerciali più vicini.

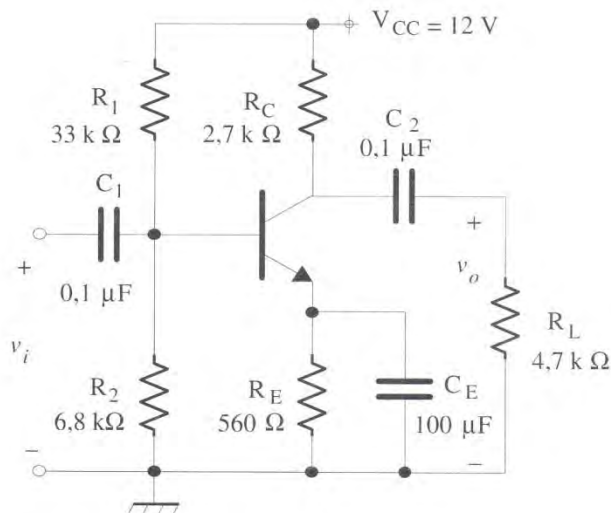
Le misure vanno eseguite con il multimetro digitale.

TABELLA DELLE MISURE

V_{BE}		V_{RE}		V_{CE}		I_B		I_C	
teorico	misurato	teorico	misur	teorico	misur	teor	misur	teor	misur
0,7 V		1,2 V		6 V		6,8 μA		2 mA	

RILIEVO DELLA RISPOSTA IN FREQUENZA DI UN AMPLIFICATORE A BJT A EMETTITORE COMUNE

Il circuito



STRUMENTI

- ALIMENTATORE
- GENERATORE DI SEGNALE (tra il terminale libero di C_1 e massa)
- OSCILLOSCOPIO (ai capi della resistenza di carico R_L)

COMPONENTI

- Si veda lo schema circuitale

Si vuole rilevare per punti la **curva del modulo della risposta in frequenza** dell'amplificatore della figura, con i valori dei componenti riportati sulla figura stessa.

Il **modulo della risposta in frequenza** è la **curva** che descrive, **al crescere della frequenza, l'andamento dell'amplificazione di tensione** $A_V = V_U/V_i$, dove V_U e V_i sono le ampiezze (moduli) della tensione di uscita e della tensione di ingresso.

Collegati gli strumenti come precedentemente indicato si procede nel modo seguente

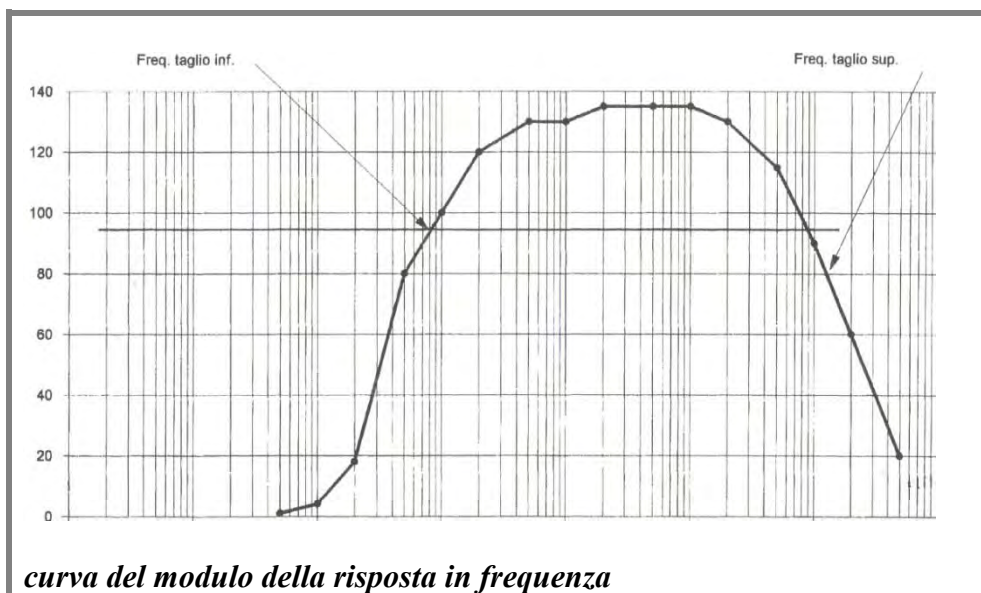
- si applica al circuito, mediante il generatore, una tensione **sinusoidale**, avente **sempre la stessa ampiezza** (per esempio **10mV**), ma **frequenza crescente** (a partire per esempio da 100 Hz)
- per ogni frequenza prescelta si leggono sull'oscilloscopio le ampiezze V_U e V_i , e se ne fa il rapporto, che rappresenta un punto della risposta A_V (l'andamento di A_V dovrà essere prima crescente, poi quasi costante, poi decrescente)
- i valori di V_U , V_i e A_V vengono riportati su una tabella
- i valori di A_V vanno riportati su un diagramma "semilogaritmico" avente in ascissa le frequenze in scala logaritmica e in ordinata A_V in scala lineare

- (a partire dal grafico o dalla tabella) si possono infine determinare le frequenze di taglio inferiore e la frequenza di taglio superiore come quei valori di frequenza in corrispondenza dei quali si ha $A_v = 0,707 \cdot A_{vMAX}$

FREQUENZA [KHz]	V_i [mV]	V_U [mV]	$A_v = V_U/V_i$
0,1	10		
0,2	10		
0,5	10		
1	10		
2	10		
5	10		
10	10		
20	10		
50	10		
100	10		
200	10		
500	10		
1000	10		
2000	10		
5000	10		
6000	10		

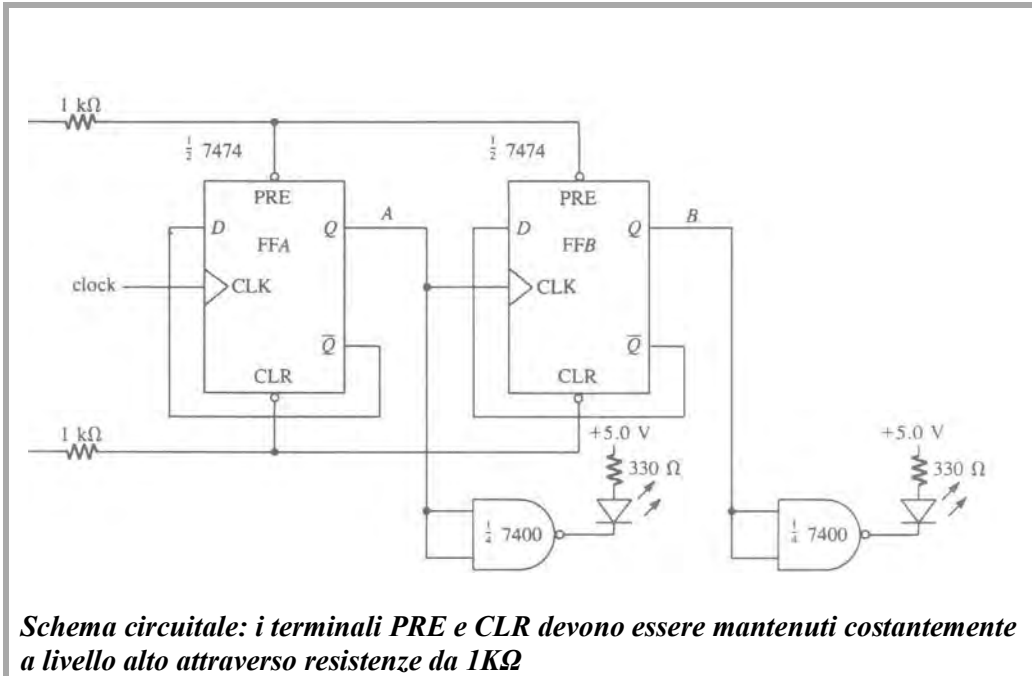
La curva di risposta in frequenza dovrebbe risultare come quella della figura seguente con

- **frequenza di taglio inferiore di circa 800 Hz**
- **frequenza di taglio superiore di circa 900 KHz**



ELETTRONICA DIGITALE

REALIZZAZIONE E VERIFICA DEL FUNZIONAMENTO DI UN CONTATORE ASINCRONO MODULO 4 BASATO SU FLIP-FLOP “D”

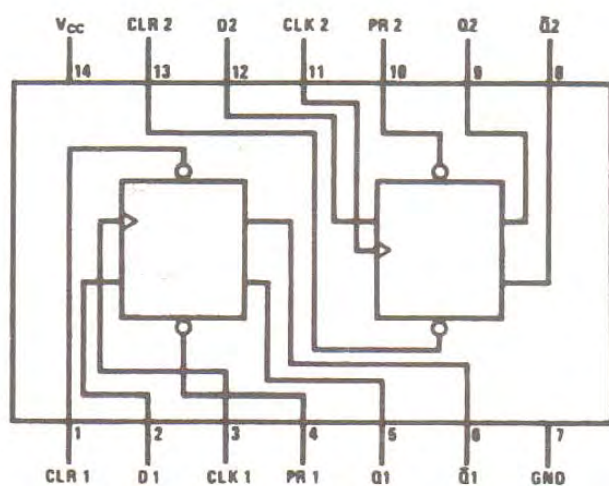


Un contatore digitale è asincrono quando l'ingresso di clock è applicato solo al primo dei FF della cascata e non a tutti. Il contatore che vogliamo realizzare è "modulo 4", cioè è in grado di contare da 0 a 3. Vengono contati gli impulsi applicati all'ingresso di clock del primo FF. All'arrivo del quarto impulso il contatore si azzerà.

Il conteggio viene visualizzato attraverso l'accensione e lo spegnimento dei due led collegati alle uscite dei due FF.

I FF sono di tipo "D" (Data flip-flop) e sono collegati in configurazione "toggle o di commutazione" con l'uscita negata permanentemente collegata all'ingresso D dei dati. I due FF utilizzati sono inclusi in un unico integrato: il '7474 (FF doppio, triggerato sul fronte di salita del clock) di cui mostriamo la piedinatura. Siccome gli ingressi asincroni di "set diretto" (PRE ossia PRESET) e "reset diretto" (CLEAR), attivi a livello basso, non saranno utilizzati, essi sono stabilmente collegati all'alimentazione di +5V.

Dual-In-Line Package



TL/F/6526-1

DM5474 (J) DM7474 (N)

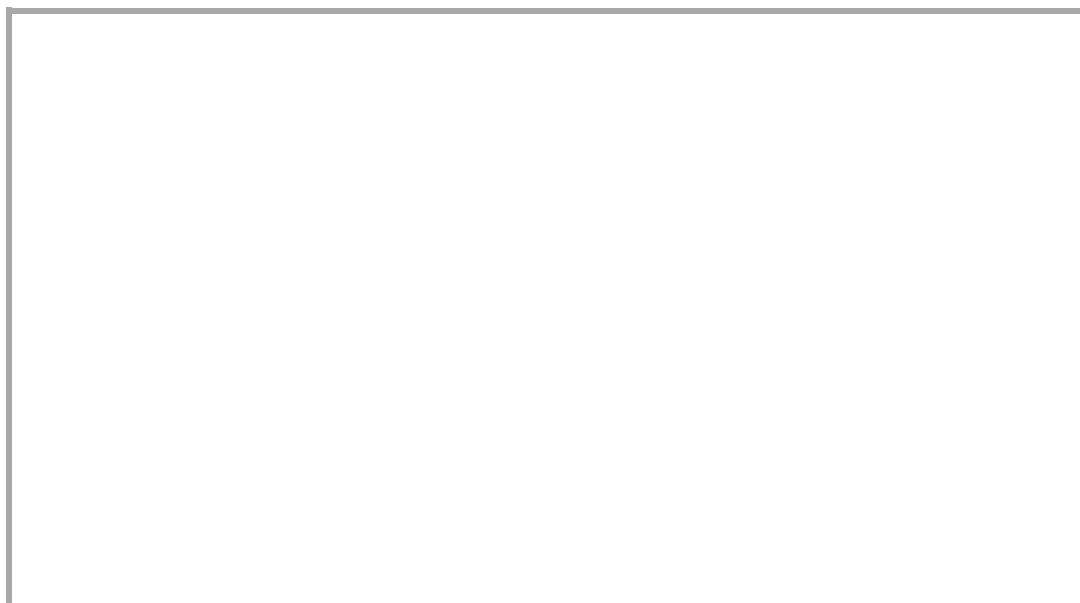
MISURA IN BASSA FREQUENZA

- tenendo conto delle piedinature si collegano i componenti conformemente allo schema circuitale
- si fornisce la tensione di alimentazione (+5V)
- gli impulsi da contare saranno forniti da un generatore onda rettangolare, impostato, per esempio, sulla frequenza di 1 Hz; l'uscita TTL-CMOS del generatore deve essere quindi collegata all'ingresso di clock del 1° FF
- osservare la sequenza di accensione/spegnimento dei due LED,, verificando che la coppia di led dia luogo a un conteggio binario (in avanti o all'indietro?)

IMPULSI CONTATI	LED "B"	LED "A"	CORRISPONDENTE DODICE BINARIO	CORRISPONDENTE CODICE DECIMALE
0	ON/OFF	ON/OFF		
1	ON/OFF	ON/OFF		
2	ON/OFF	ON/OFF		
3	ON/OFF	ON/OFF		
4	ON/OFF	ON/OFF		

MISURA A 1 KHz

- tenendo conto delle piedinature si collegano i componenti conformemente allo schema circuitale
- si fornisce la tensione di alimentazione (+5V)
- gli impulsi da contare saranno forniti da un generatore d'onda rettangolare, impostato sulla frequenza di 1 KHz; l'uscita TTL-CMOS del generatore deve essere quindi collegata all'ingresso di clock del 1° FF
- osservare sui due canali di un oscilloscopio le forme d'onda presenti sulle uscite B ed A dei due FF, disegnarle e interpretarle in base alla numerazione binaria





DM5474/DM7474 Dual Positive-Edge-Triggered D Flip-Flops with Preset, Clear and Complementary Outputs

General Description

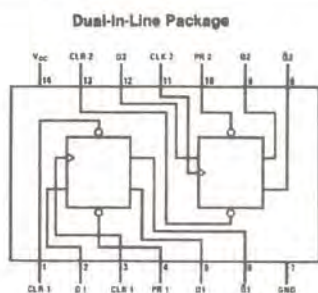
This device contains two independent positive-edge-triggered D flip-flops with complementary outputs. The information on the D input is accepted by the flip-flops on the positive going edge of the clock pulse. The triggering occurs at a voltage level and is not directly related to the transition time of the rising edge of the clock. The data on the D input may be changed while the clock is low or high without affecting the outputs as long as the data setup and hold times are not violated. A low logic level on the preset or clear inputs will set or reset the outputs regardless of the logic levels of the other inputs.

Absolute Maximum Ratings (Note 1)

Supply Voltage	7V
Input Voltage	5.5V
Storage Temperature Range	- 65°C to 150°C

Note 1: The "Absolute Maximum Ratings" are those values beyond which the safety of the device can not be guaranteed. The device should not be operated at these limits. The parametric values defined in the "Electrical Characteristics" table are not guaranteed at the absolute maximum ratings. The "Recommended Operating Conditions" table will define the conditions for actual device operation.

Connection Diagram



TL/F/0526-1

DM5474 (J) DM7474 (N)

Function Table

Inputs				Outputs	
PR	CLR	CLK	D	Q	$\bar{Q}$
L	H	X	X	H	L
H	L	X	X	L	H
L	L	X	X	H*	H*
H	H	↑	H	H	L
H	H	↑	L	L	H
H	H	L	X	Q _O	$\bar{Q}_O$

H = High Logic Level

X = Either Low or High Logic Level

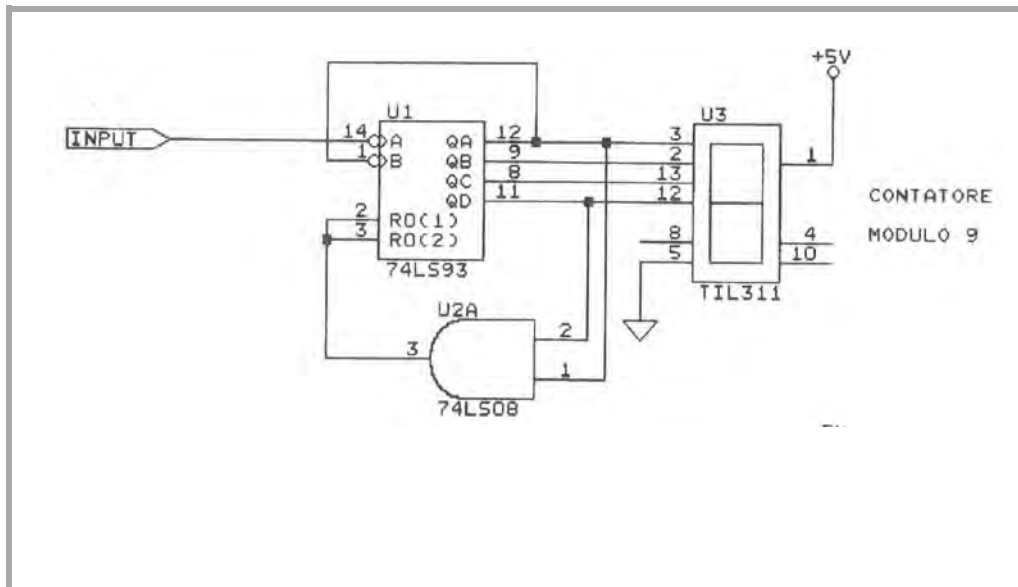
L = Low Logic Level

↑ = Positive-going transition of the clock.

* = This configuration is nonstable; that is, it will not persist when either the preset and/or clear inputs return to their inactive (high) level.

Q_O = The output logic level of Q before the indicated input conditions were established.

REALIZZAZIONE E VERIFICA DEL FUNZIONAMENTO DI UN CONTATORE ASINCRONO MODULO 9



COMPONENTI

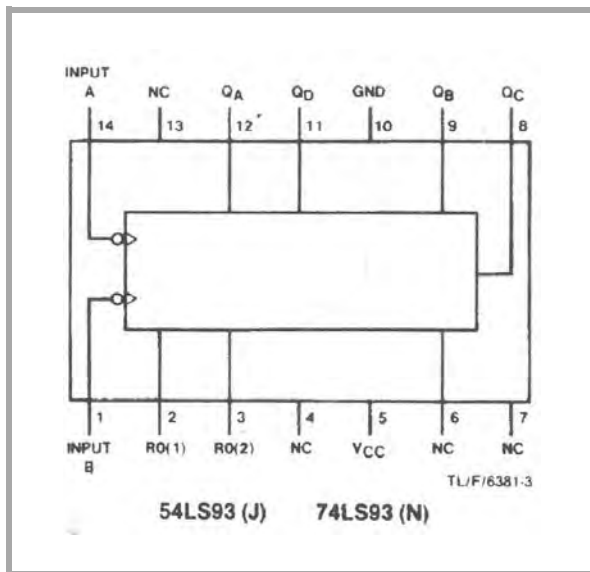
- contatore modulo 16 '7493
- display esadecimale TL311
- Porta logica AND 74LS08

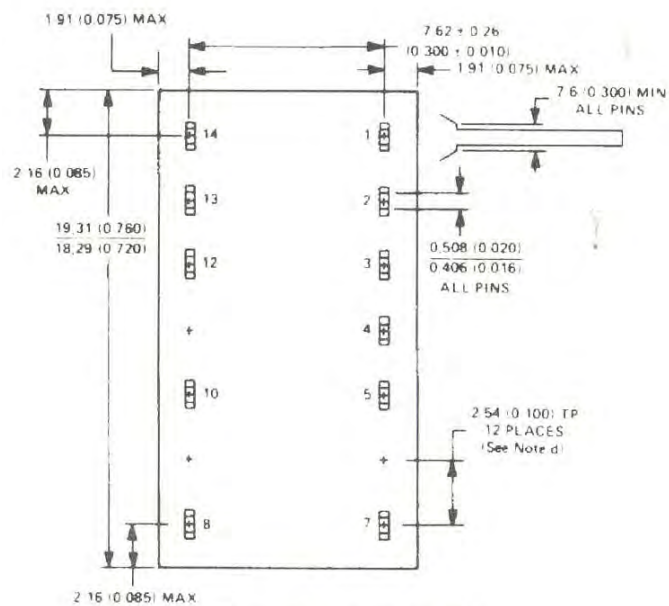
STRUMENTI

- generatore di segnali
- alimentatore
- oscilloscopio

Dato un contatore integrato commerciale, per esempio modulo 16, si può ottenere un contatore di modulo minore (per esempio modulo 9) collegando alcune uscite del contatore all'ingresso di una porta AND e collegando l'uscita della AND al terminale di "reset diretto" del contatore.

Le uscite da collegare sono Q_A e Q_D , ossia quelle che diventerebbero alte in corrispondenza della formazione del codice binario di 9 (e cioè 1001). In questo modo la formazione del codice 1001 viene bloccata e tutte le uscite del contatore si portano sullo zero. Useremo il contatore modulo 16 '7493. In questo modo il dispositivo conta fino ad 8, comportandosi da contatore "modulo 9".



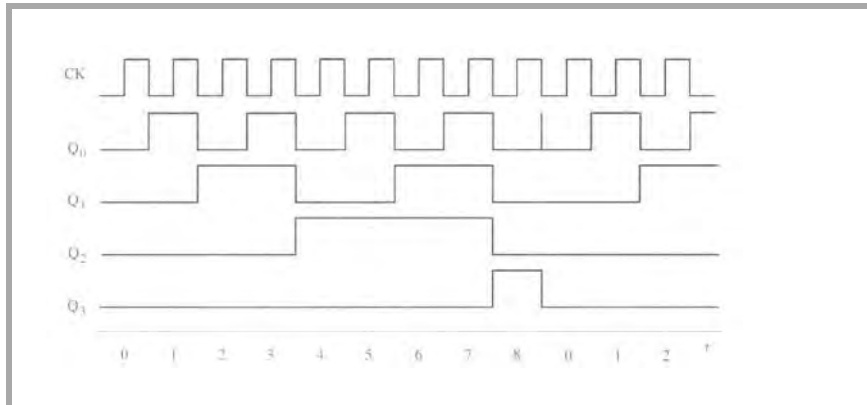


- PIN 1 LED SUPPLY VOLTAGE
- PIN 2 LATCH DATA INPUT B
- PIN 3 LATCH DATA INPUT A
- PIN 4 LEFT DECIMAL POINT CATHODE
- PIN 5 LATCH STROBE INPUT
- PIN 6 OMITTED
- PIN 7 COMMON GROUND
- PIN 8 BLANKING INPUT
- PIN 9 OMITTED
- PIN 10 RIGHT DECIMAL POINT CATHODE
- PIN 11 OMITTED
- PIN 12 LATCH DATA INPUT D
- PIN 13 LATCH DATA INPUT C
- PIN 14 LOGIC SUPPLY VOLTAGE, V_{CC}

DISPLAY ESADECIMALE TIL 311

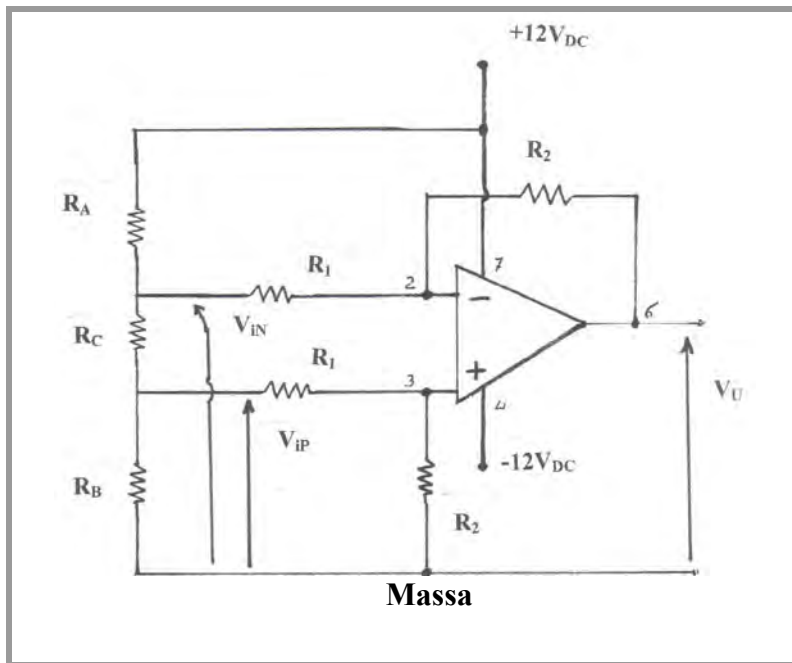
MISURA

- Si collegano il contatore la porta e il display secondo lo schema e in base alle piedinature fornite.
- Gli impulsi da conteggiare possono essere forniti dall'uscita TTL-CMOS di un generatore d'onda rettangolare e applicati a uno dei pin di ingresso del contatore
- Il corretto funzionamento del contatore può essere verificato attraverso il display o visualizzando le quattro uscite del contatore mediante un oscilloscopio



CONFIGURAZIONI AMPLIFICATRICI DELL'AMPLIFICATORE OPERAZIONALE

VERIFICA DEL FUNZIONAMENTO IN CONTINUA DI UN SOTTRATTORE (AMPLIFICATORE DIFFERENZIALE) CON OPERAZIONALE



STRUMENTI
ALIMENTATORE
MULTIMETRO DIGITALE

COMPONENTI
 $R_A = R_B = R_C = 3 \times 470 \text{ Ohm}$
 $R_2 = 2 \times 56 \text{ K}\Omega$
 $R_1 = 2 \times 47 \text{ K}\Omega$

CALCOLI

RELAZIONE I/O:

$$V_U = \frac{R_2}{R_1} \cdot (V_{ip} - V_{in})$$

- $R_A = R_B = R_C$: partitore per l'applicazione di due tensioni continue di ingresso
- **VALORI TEORICI DELLE TENSIONI DI INGRESSO E DI USCITA**

$$V_{ip} = \frac{R_B}{R_B + R_C + R_A} \cdot 12V = 4V$$

$$V_{in} = \frac{R_B + R_C}{R_B + R_C + R_A} \cdot 12V = 8V$$

$$V_U = \frac{56}{47} \cdot (4 - 8) = 1,2 \cdot (-4) = -4,8V$$

IL CIRCUITO

Osserviamo prima di tutto che le tre resistenze R_A , R_B , R_C non fanno parte del circuito sottrattore, ma costituiscono un partitore che ci consente di ottenere (a partire da un unico alimentatore) due differenti tensioni continue (V_{ip} e V_{in}) da applicare al sottrattore come segnali di ingresso. In questo modo l'amplificatore differenziale, in base alla relazione I/O eseguirà l'operazione:

$$V_U = \frac{56}{47} \cdot (4 - 8) \cong -1,2 \cdot 4 \cong -4,8V$$

Il sottrattore quindi esegue la differenza dei segnali di ingresso e la amplifica di 1,2 volte.

Si tenga presente che la relazione I/O qui riportata è valida solo se le due resistenze R_2 sono uguali tra loro e le due resistenze R_1 sono uguali tra loro.

PROCEDIMENTO DI MISURA

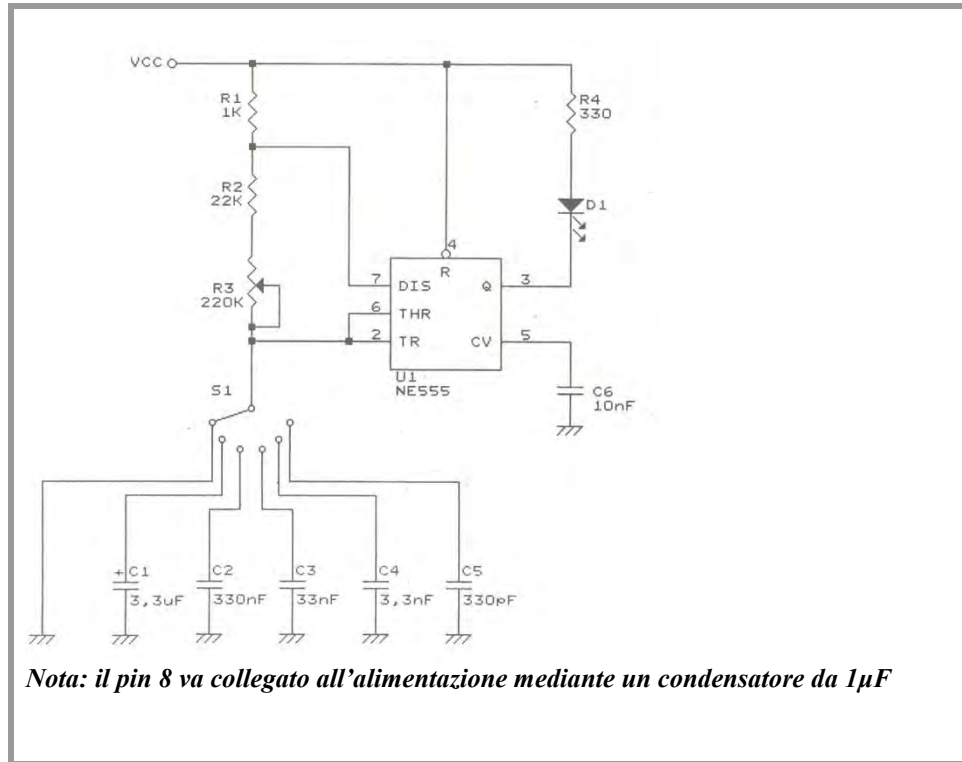
- Si realizza il circuito come nella figura, tenendo conto della piedinatura dell'operazionale
- si eseguono i calcoli teorici delle due tensioni di ingresso e della tensione di uscita
- con il multimetro digitale si misurano le tensioni di ingresso e di uscita (e se ne valuta eventualmente lo scostamento dai valori teorici)
- si ripete la misura con differenti valori delle resistenze del partitore.

MISURE DA ESEGUIRE

partitore di ingresso R_A, R_B, R_C (Ohm)	tensione di ingresso V_{ip}		tensione di ingresso V_{in}		tensione V_u di uscita	
	Valore teorico	Valore misurato	Valore teorico	Valore misurato	Valore teorico	Valore misurato
470, 470, 470						
<i>cambiare i valori</i>						

GENERATORI DI SEGNALE

REALIZZAZIONE E VERIFICA DEL FUNZIONAMENTO DI UN GENERATORE D'ONDA RETTANGOLARE (BASATO SU TIMER 555) A FREQUENZA REGOLABILE



STRUMENTI

- ALIMENTATORE
- MULTIMETRO DIGITALE
- OSCILLOSCOPIO

COMPONENTI

Quantità	Sigla di riferimento	Valore	Caratteristiche
1	C_1	3,3 μ F	Elettrolitico 16 V
1	C_2	330 nF	Plastico
1	C_3	33 nF	Plastico
1	C_4	3,3 nF	Plastico
1	C_5	330 pF	Ceramico
1	C_6	10 nF	Plastico
1	C_7	1 μ F	Elettrolitico 16 V
1	D_1		LED rosso
1	R_1	1 k Ω	1/4 W
1	R_2	22 k Ω	1/4 W
1	R_3	220 k Ω	Potenziometro lineare
1	R_4	330 Ω	1/4 W
1	S_1		Commutatore 6 posizioni
1	U1	555	Timer

RELAZIONI DI PROGETTO

$$T_0 = T_{carica} + T_{scarica} \cong 0,7 \cdot C \cdot (R_1 + 2 \cdot (R_2 + R_3))$$

Periodo del segnale di uscita
*in funzione dei valori del condensatore
e delle resistenze della rete esterna*

$$f_0 = \frac{1}{T_{carica} + T_{scarica}} \cong \frac{1}{0,7 \cdot C \cdot (R_1 + 2 \cdot (R_2 + R_3))}$$

Frequenza del segnale di uscita
*in funzione dei valori del condensatore
e delle resistenze della rete esterna*

IL CIRCUITO

Il generatore di onda rettangolare o multivibratore astabile o generatore di clock è un dispositivo che fornisce in uscita un'onda rettangolare a una frequenza determinata, che dipende dai componenti usati (in base alle relazioni di progetto).

Il circuito che si intende realizzare è un **generatore di clock** che ha la particolarità di **poter far variare la frequenza del segnale di uscita**.

La frequenza può essere regolata:

- **a scatti** (selezionando cioè un **intervallo di frequenze**) operando sul **commutatore S₁**, che offre la possibilità di **connettersi a sei diversi condensatori C₁.....C₅** (come si vede dalla relazione di progetto la frequenza dipende inversamente da C)
- **con continuità**, agendo sul **trimmer R₃** (come si vede dalla relazione di progetto, la frequenza dipende inversamente da R₃)

C'è quindi la possibilità di effettuare prima una regolazione "grossolana", mediante il commutatore e i condensatori, e poi (all'interno dell'intervallo selezionato con il commutatore) una regolazione "fine" mediante il trimmer.

Sostituendo nelle relazioni di progetto i valori dei componenti, si possono calcolare gli intervalli di frequenza relativi a ciascuno dei sei condensatori selezionabili mediante il commutatore.

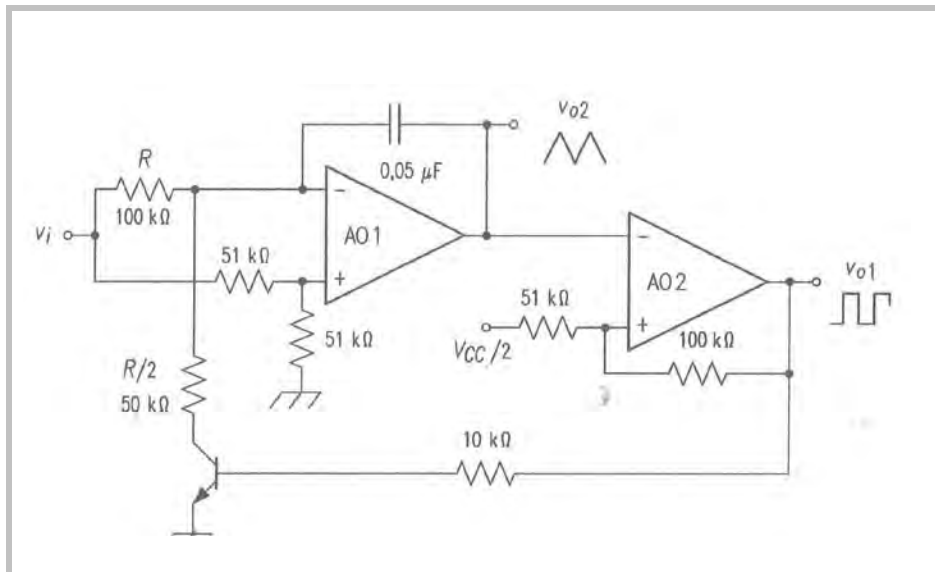
CONDENSATORE SELEZIONATO	C ₁ 3,3μF	C ₂ 330 nF	C ₃ 33 nF	C ₄ 3,3 nF	C ₅ 330 pF
RANGE DI FREQUENZE TEORICHE (Hz)	0,89-9,6	8,9-96	89-960	8900-96000	89000-96000

PROCEDIMENTO DI MISURA

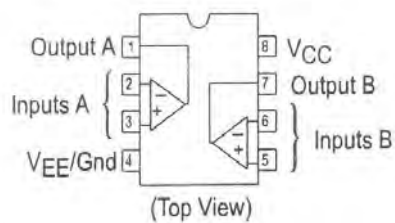
- Si realizza lo schema della figura, tenendo conto della piedinatura del timer 555, e si collega l'oscilloscopio tra il piedino 3 di uscita e la massa
- Si posiziona il commutatore S_1 sul condensatore di capacità maggiore (C_1), il quale darà luogo alle frequenze più basse. Si ruota il trimmer da un estremo all'altro: il numero di accensioni al secondo del LED ci fornirà i valori di frequenza
- Si posiziona il commutatore S_1 sul condensatore di capacità C_2 , il quale darà luogo a frequenze comprese fra gli 8, 9 e i 96 Hz. Si ruota il trimmer da un estremo all'altro: le frequenze più basse (entro i 20 Hz circa) potranno essere misurate attraverso il numero di accensioni al secondo del LED, le frequenze più alte dovranno essere misurate sull'onda rettangolare visualizzata dall'oscilloscopio
- Si posiziona il commutatore S_1 sul condensatore di capacità C_3 , il quale darà luogo a frequenze comprese fra gli 89 e i 960 Hz. Si ruota il trimmer da un estremo all'altro: il led rimarrà sempre acceso, per cui tutte le frequenze dovranno essere misurate sull'onda rettangolare visualizzata dall'oscilloscopio
- Si posiziona il commutatore S_1 sul condensatore di capacità C_4 e si ripetono le operazioni del passo precedente
- Si posiziona sul condensatore di capacità C_5 e si ripetono le operazioni del passo precedente
- I valori misurati si riportano sulla tabella seguente:

CONDENSATORE SELEZIONATO	C_1 3,3μF	C_2 330 nF	C_3 33 nF	C_4 3,3 nF	C_5 330 pF
RANGE TEORICO DI FREQUENZE (Hz)	0,89-9,6	8,9-96	89-960	8900-96000	89000-96000
RANGE MISURATO DI FREQUENZE (Hz)					

VCO (OSCILLATORE CONTROLLATO IN TENSIONE) BASATO SU OPERAZIONALI LM 358 AD ALIMENTAZIONE SINGOLA



PIN CONNECTIONS



LM 358

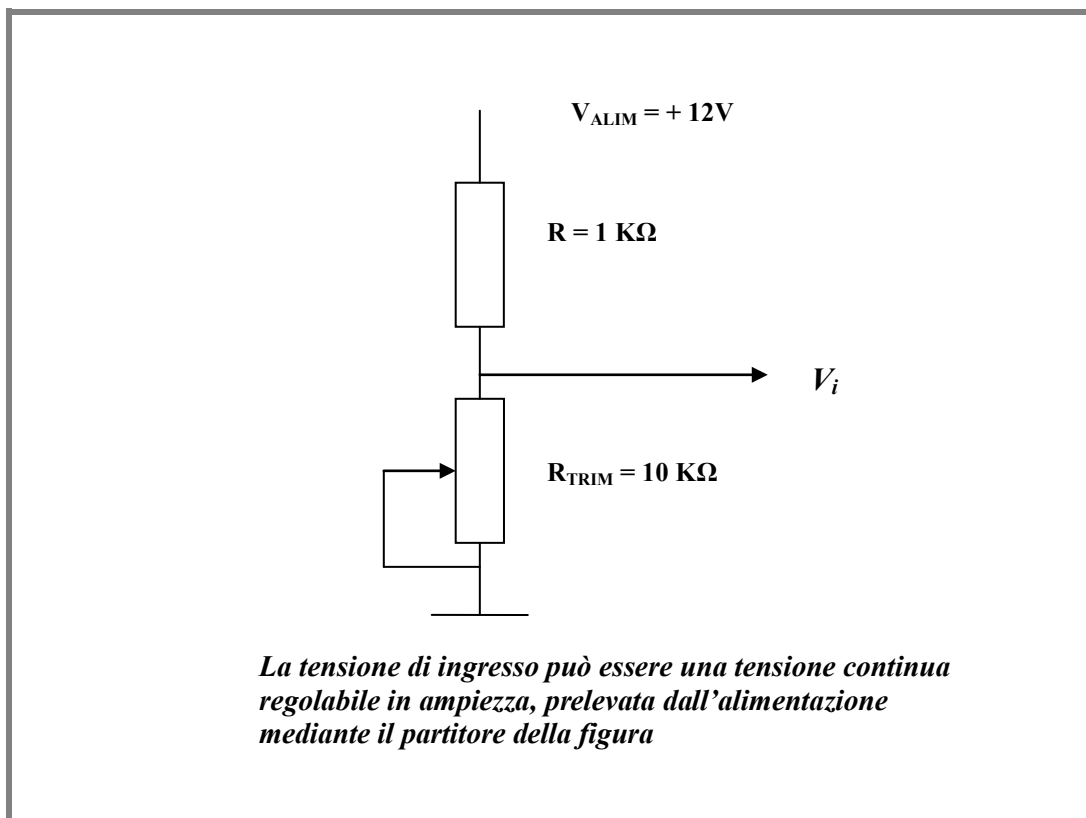
STRUMENTI
MULTIMETRO DIGITALE
OSCILLOSCOPIO
COMPONENTI
SI VEDA LO SCHEMA
Il BJT può essere un 2N2222

IL CIRCUITO

L'oscillatore controllato in tensione (VCO) effettua una conversione tensione-frequenza ed è anche un modulatore di frequenza perché la **variazione della frequenza** del segnale di **uscita** (rispetto al valore di frequenza in assenza di segnale di ingresso) è **proporzionale all'ampiezza del segnale applicato in ingresso**.

Il circuito da realizzare è basato su due operazionali ad alimentazione singola LM 358. L'operazionale AO1 realizza, con la sua rete di retroazione negativa, un integratore, mentre l'operazionale AO2 realizza, con la sua rete di retroazione positiva un comparatore con isteresi o trigger di Schmitt.

Sia il segnale di uscita dell'integratore che il segnale di uscita del trigger presentano quindi una frequenza proporzionale all'ampiezza del segnale di ingresso, cosa che ci proponiamo di verificare con questa misura.



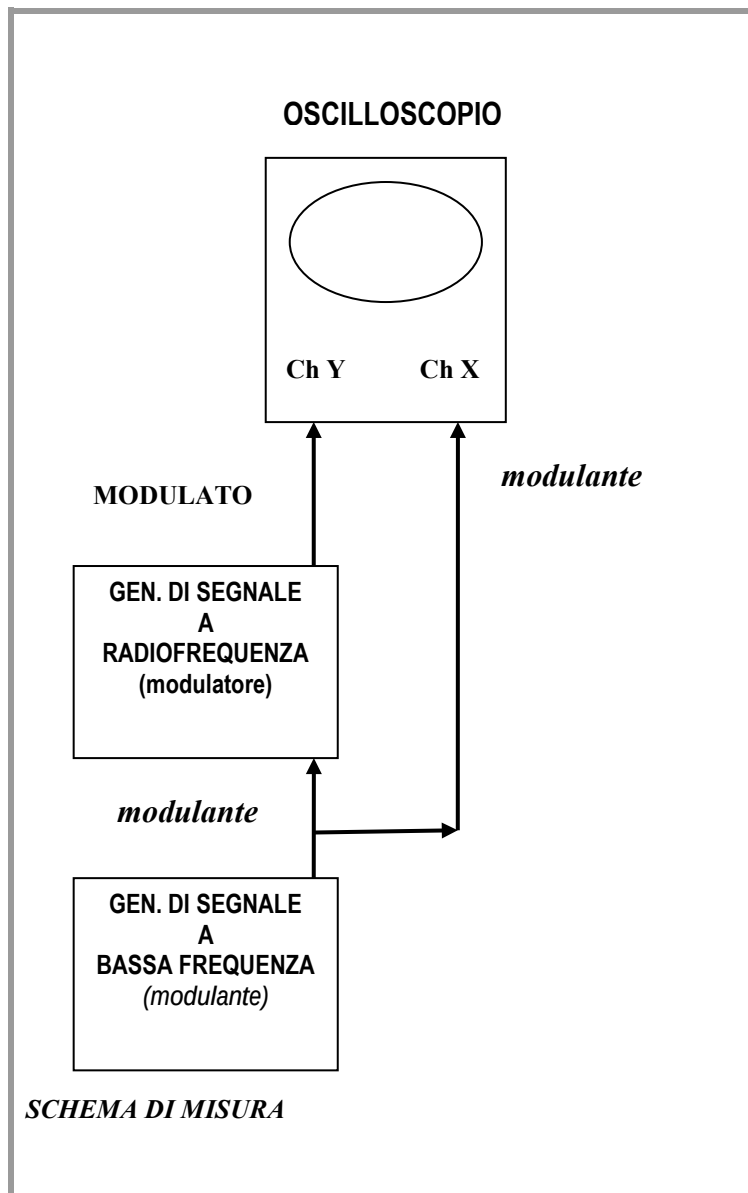
PROCEDIMENTO DI MISURA

- Si misurino (con il multimetro) differenti valori del segnale di ingresso e (con l'oscilloscopio) le corrispondenti frequenze di uno dei segnali di uscita (per esempio la tensione rettangolare di uscita del trigger)
- Si tracci un grafico che riporti i valori di frequenza dell'uscita in funzione dei valori di ampiezza dell'ingresso.

[illegible]

ESPERIENZE DI TELECOMUNICAZIONI

MISURA DELL'INDICE DI MODULAZIONE AM CON IL METODO DEL TRAPEZIO



STRUMENTI

OSCILLOSCOPIO

GENERATORE DI SEGNALE

GENERATORE DI SEGNALE (utilizzabile come modulatore AM)

COMPONENTI

NESSUNO

LO SCHEMA DI MISURA

Questa esperienza consiste nella visualizzazione, sul monitor dell'oscilloscopio, di un trapezio, le cui due basi "a" e "b" ci permettono di determinare l'indice di modulazione di un segnale modulato AM.

Per visualizzare il trapezio è necessario che:

- al canale orizzontale (ch **X**) dell'oscilloscopio venga applicato il segnale **modulante informativo**
- al canale verticale (ch **Y**) venga applicato il segnale modulato in ampiezza, cioè il **segnale AM**.

Per fare ciò dobbiamo eseguire una modulazione di ampiezza generando (mediante un generatore di segnale) un segnale informativo a bassa frequenza e applicando tale segnale a un secondo generatore (stavolta ad alta frequenza), il quale funzionerà da modulatore, fornendo in uscita un segnale modulato AM.

Il segnale modulante informativo viene applicato al canale X, mentre il segnale modulato AM viene applicato al canale Y.

Sul monitor risulterà visualizzato un trapezio nel quale:

- la base maggiore "a" è proporzionale alla massima ampiezza del segnale modulato AM

$$a = k \cdot 2 \cdot V_{AM(MAX)} = k \cdot 2 \cdot A_0 \cdot (1 + m)$$

(dove A_0 è l'ampiezza della portante non modulata)

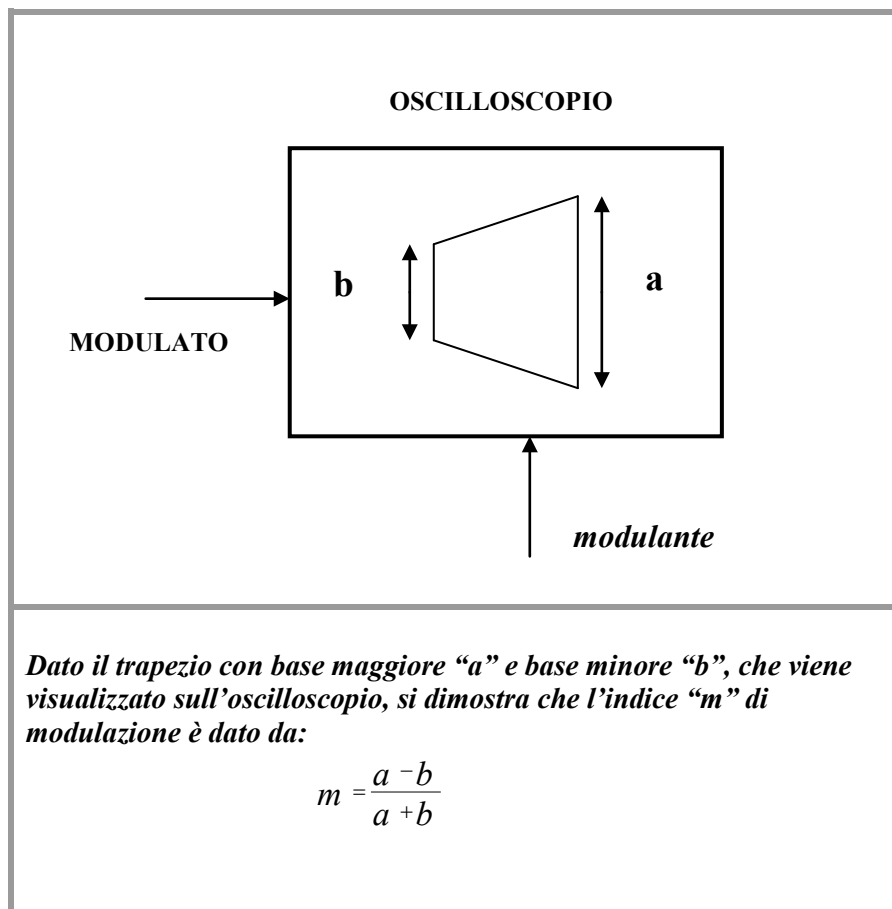
- la base minore "b" è proporzionale alla minima ampiezza dello stesso segnale modulato AM

- $b = k \cdot 2 \cdot V_{AM(min)} = k \cdot 2 \cdot A_0 \cdot (1 - m)$

Eseguendo il rapporto a/b, ed esplicitandolo rispetto a "m", si ottiene l'espressione:

$$m = \frac{a - b}{a + b}$$

dalla quale si vede che l'indice di modulazione si può determinare conoscendo soltanto le basi "a" e "b" del trapezio visualizzato.

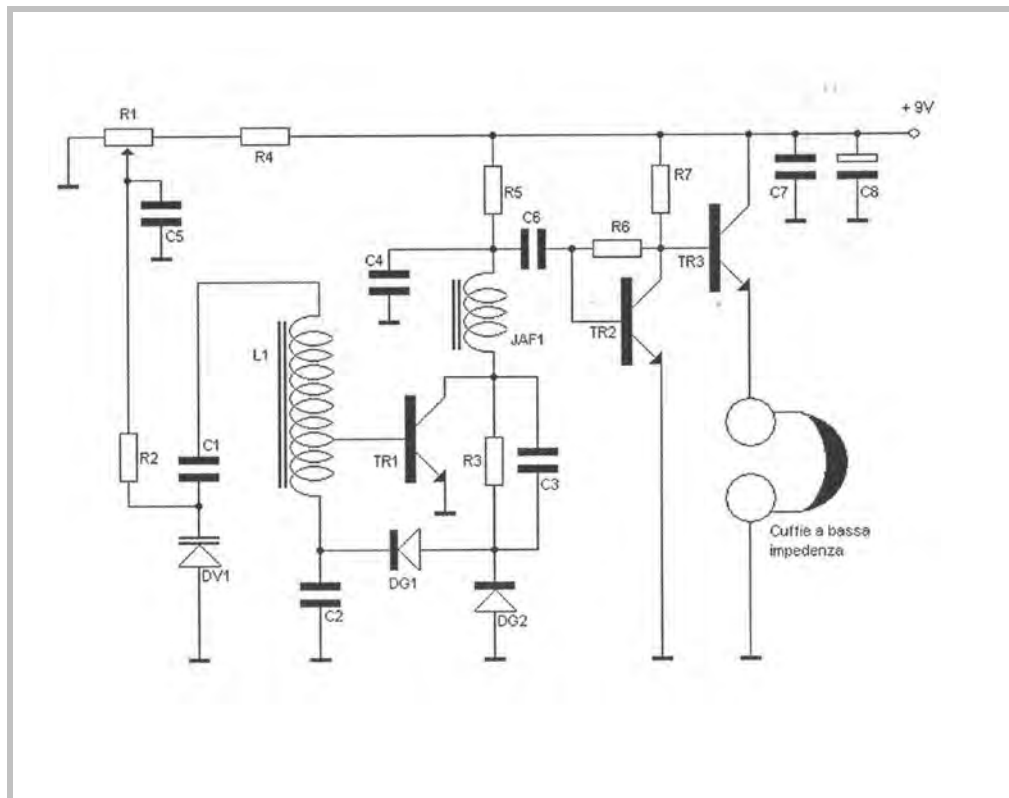


ESECUZIONE DELLA MISURA

- Si realizza lo schema indicato nella prima figura
- Si visualizza il trapezio sull’oscilloscopio e si misurano le basi “a” e “b” del trapezio
- Si calcola l’indice di modulazione sostituendo i valori delle basi “a” e “b” del trapezio nell’espressione:

$$m = \frac{a - b}{a + b}$$

RADIORICEVITORE A CIRCUITO RIFLESSO (REALIZZAZIONE SU BREADBOARD)



IL RICEVITORE

Il ricevitore in questione non è un eterodina, ma un ricevitore a circuito riflesso, caratterizzato dalla presenza di un unico stadio, che lavora sia da amplificatore a radiofrequenza che da amplificatore ad audiofrequenza.

I radioricevitori di questo tipo, basati su tubi a vuoto invece che su transistor, furono utilizzati nei primi anni del novecento.

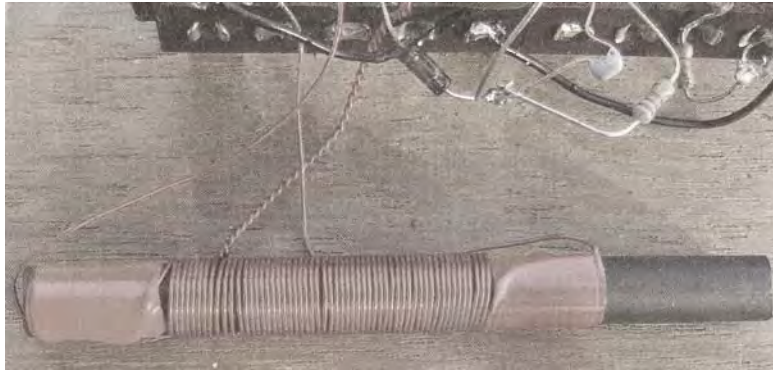
LISTA COMPONENTI	
R1	100 kohm pot. Lin.
R2	39 kohm ¼ w
R3	220 kohm ¼ w
R4	100 ohm ¼ w
R5	2,2 kohm ¼ w
R6	220 kohm ¼ w
R7	2,2 kohm ¼ w
C1	2,2 nf cer. 50 v.
C2	47 nf cer. 50 v.
C3	10 nf cer. 50 v.
C4	4,7 nf cer. 50 v.
C5	100 nf cer. 50 v.
C6	100 nf cer. 50 v.
C7	100 nf cer. 50 v.
C8	22 uf 25 v. Elettr.
TR1,2,3	BC505 o equiv.
DG1,2	AA117 o equiv.
JAF1	2,7 Mhenry
DV1	BB113 o equiv.
L1	Vedi testo

FUNZIONAMENTO DEL RICEVITORE

La sintonia del segnale captato dall'antenna è effettuata dall'induttanza L1 e dal diodo a capacità variabile (varicap o varactor) DV1. Il segnale selezionato viene applicato alla base del BJT TR1 che lo amplifica.

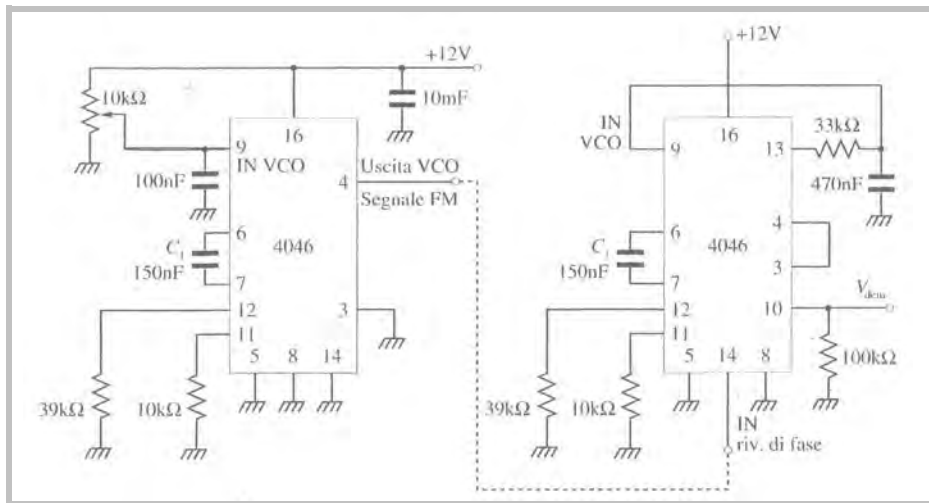
Sul collettore del transistor il segnale a radiofrequenza (che non può attraversare l'induttanza JAF1) converge verso il demodulatore formato dai due diodi al germanio DG1 e DG2.

Il segnale demodulato (segnale a frequenza audio) si presenta quindi ai capi del condensatore C2 e, attraverso la bobina L1, viene nuovamente applicato a TR1, dopodiché (potendo attraversare, grazie alla bassa frequenza, l'induttore JAF1) viene amplificato dai BJT TR2 e TR3 ed è pronto per poter essere ascoltato in cuffia (per esempio una cuffia da computer).



La bobina di sintonia L1 può essere costruita avvolgendo 60 spire di rame smaltato da 0,3 mm su un barretta di ferrite con diametro di 0,6-0,7 cm e lunghezza di 5-6 cm. La presa del collegamento alla base del BJT TR1 va fatta alla 12° spira a partire dal lato collegato a massa

MODULATORE E DEMODULATORE FM BASATI SU DUE PLL DIGITALI INTEGRATI 4046



STRUMENTI

MULTIMETRO DIGITALE

OSCILLOSCOPIO O FREQUENZIMETRO

COMPONENTI

SI VEDA LO SCHEMA

IL SISTEMA

Il sistema proposto è formato da un trasmettitore e da un ricevitore (basati ciascuno su un integrato PLL). Il primo dispositivo modula in frequenza la portante (generata internamente) con un **segnale modulante generato da noi regolando il trimmer da 10 KΩ**, collegato all'alimentazione e applicato al terminale (pin 9) di ingresso del VCO interno al PLL.

All'uscita di questo primo PLL (che funziona da modulatore) viene quindi ottenuto un segnale modulato in frequenza, che è applicato, attraverso un cavetto, al pin 14 (rivelatore di fase) del PLL demodulatore.

Il demodulatore, cioè il secondo PLL, deve “restituire” (sul pin n.10) un segnale demodulato proporzionale al segnale modulante che abbiamo fornito al modulatore attraverso il trimmer.

LA MISURA

Per verificare il funzionamento del **sistema nel suo complesso** si può verificare che esiste un intervallo di frequenze nel quale l'ampiezza del segnale di uscita del demodulatore è proporzionale all'ampiezza del segnale modulante. Un intervallo cioè nel quale il segnale modulante informativo viene ricostruito dal demodulatore. La misura è eseguita con **due multimetri** che vanno posti:

- uno fra il cursore del trimmer e massa
- l'altro fra il terminale di uscita (pin 10, segnale di controllo del VCO interno) del secondo PLL e massa.

Preliminarmente si può verificare la funzionalità del ***solo modulatore*** controllando che ***i valori della frequenza di uscita del modulatore siano proporzionali ai valori di ampiezza del segnale modulante impostato col trimmer.***

La verifica del funzionamento del modulatore può essere fatta fissando, con il trimmer, dei valori di tensione all'ingresso del VCO del modulatore, riportando su una tabella i valori della frequenza dell'onda rettangolare di uscita del VCO (pin 9 del PLL) e facendone il rapporto.

La misura può essere eseguita mediante:

- un ***multimetro*** (voltmetro) inserito fra il ***cursore del trimmer e massa***
- un ***oscilloscopio o un frequenzimetro*** posto fra ***l'uscita del primo PLL e massa.***

VERIFICA DEL FUNZIONAMENTO DEL MODULATORE

V_I $V_{\text{modulante (trimmer)}}$ [V]	f_{VCO} (FREQUENZA dell'uscita del VCO) [Hz]	f_{VCO}/V_I
0		
0,5		
1		
1,5		
2		
2,5		
3		
3,5		
4		
4,5		
5		
5,5		
6		
6,5		
7		
7,5		
8		
8,5		
9		
9,5		
10		

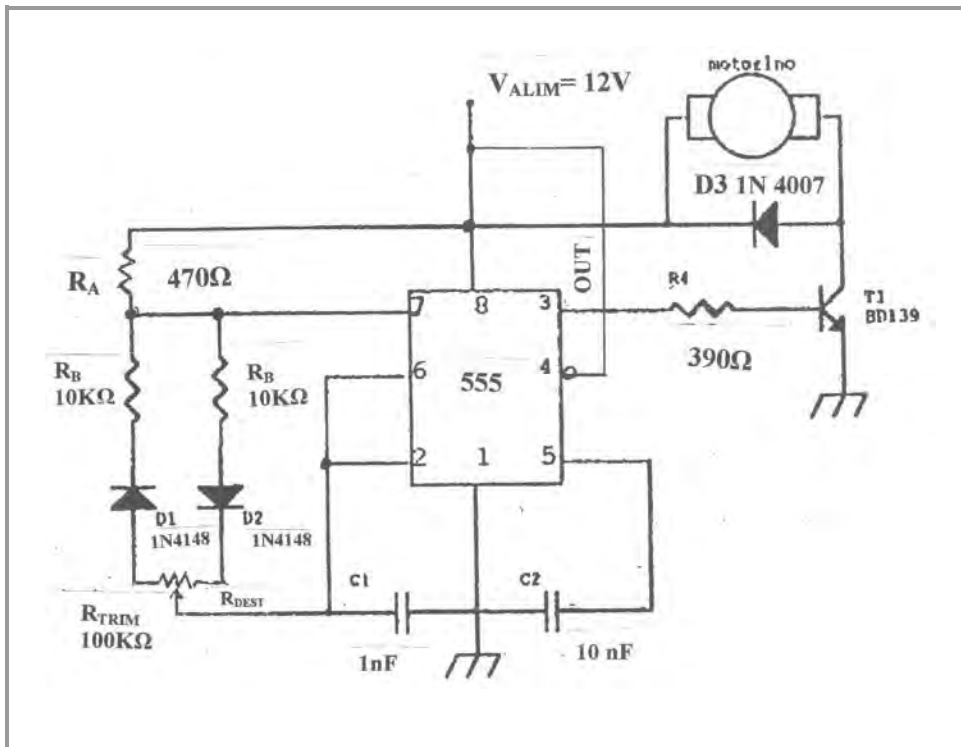
VERIFICA DEL FUNZIONAMENTO DEL SISTEMA COMPLESSIVO

V_I AMPIEZZA con V_I=V_{modulante} (trimmer) [V]	V_{II} AMPIEZZA dell'uscita del VCO del 2° PLL (pin 10)	V_{II}/ V_I	f [Hz]
0			
0,5			
1			
1,5			
2			
2,5			
3			
3,5			
4			
4,5			
5			
5,5			
6			
6,5			
7			
7,5			
8			
8,5			
9			
9,5			
10			

CIRCUITI BASATI SUL CONTROLLO DELLA POTENZA ELETTRICA

REGOLAZIONE DI VELOCITA' DI UN MOTORINO CC CON TECNICA PWM (Pulse Width Modulation, Modulazione di durata degli impulsi)

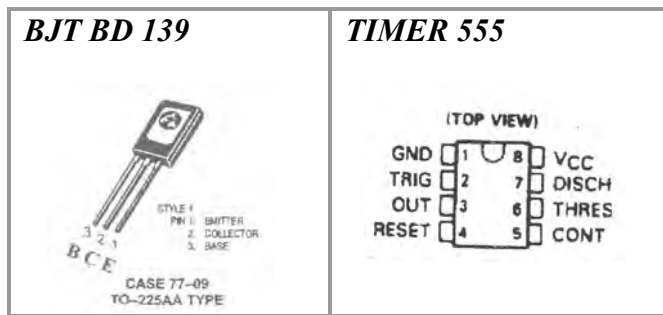
Un motore CC non deve necessariamente essere alimentato con una vera e propria tensione continua, ma anche mediante una tensione unipolare e variabile nel tempo, tipicamente **un'onda rettangolare unipolare** con frequenza e ampiezza costante e **duty-cycle variabile**.



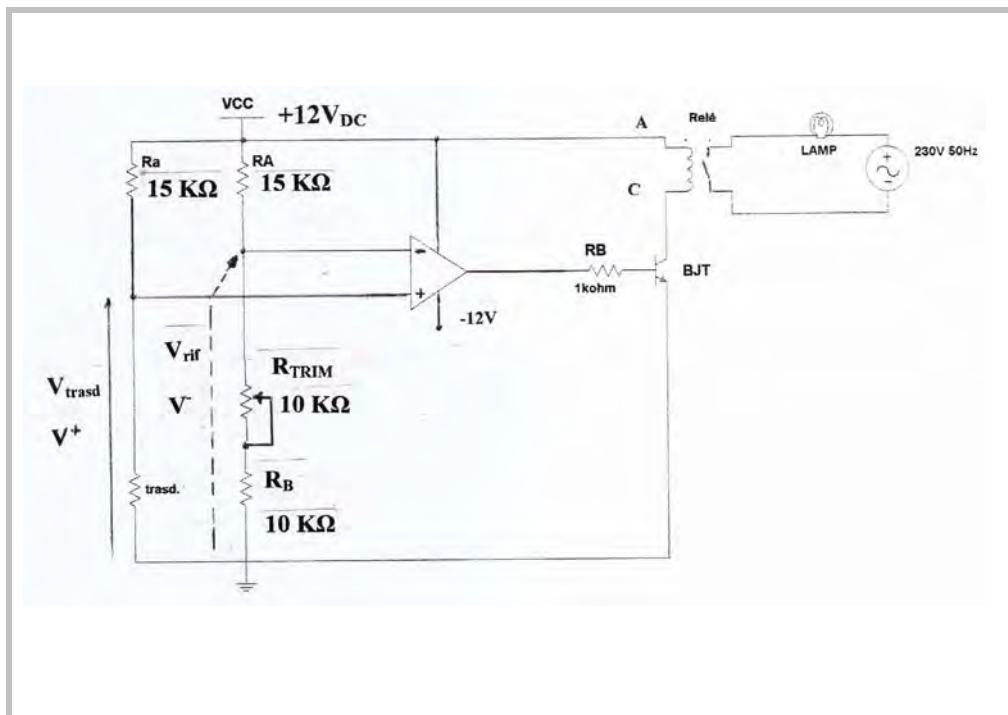
Siccome la velocità del motore viene a dipendere in sostanza dal valore medio della tensione nel periodo, valore medio che è proporzionale al duty-cycle, ecco che variando il duty-cycle dell'onda rettangolare di alimentazione, se ne varia il valor medio, il che fa variare, proporzionalmente, la velocità del motore.

Il circuito che si intende realizzare è basato appunto su un generatore d'onda rettangolare realizzato mediante il temporizzatore integrato 555. In questo momento non ci interessa il funzionamento del 555: ci basta sapere che dal suo terminale di uscita (pin n.3) è prelevabile un'onda rettangolare a **duty-cycle regolabile**, la quale pilota un transistor che interrompe o consente l'alimentazione del motore. Ampiezza e frequenza dell'onda rettangolare sono fisse. La regolazione del duty-cycle e quindi della velocità avviene mediante il trimmer, posto, nello schema, in basso a sinistra.

Se il circuito è realizzato correttamente la variazione di velocità in funzione del duty-cycle è valutabile "ad occhio". Sono comunque possibili misure più accurate.



LAMPADA COMANDATA DA UN INTERRUTTORE CREPUSCOLARE CON FOTORESISTENZA



ATTENZIONE

Questa esperienza, prevedendo l'uso della tensione di rete, va condotta sotto attenta vigilanza del docente

La parte di circuito a sinistra dei punti AC (circuito di controllo) può essere realizzata su breadboard. La parte a destra (circuito di potenza, con relè, lampadina, spina di rete) mediante un normale cavo elettrico con adeguata portata.

STRUMENTI
MULTIMETRO DIGITALE
COMPONENTI
SI VEDA LO SCHEMA
Inoltre
BJT 2N 2222 oppure BD 243
RELE' 3A 250V_{AC} 3A 30V_{DC}
FOTORESISTENZA
Lampadina 230V 50W
Portalampada
Spina

FINALITA' DEL SISTEMA

Il sistema ha la funzione di far accendere una lampadina quando la luminosità dell'ambiente scende al di sotto di un certo livello.

La luminosità dell'ambiente è rilevata da una fotoresistenza, un trasduttore la cui resistenza interna varia inversamente con la luce che lo colpisce.

Il trasduttore è collegato all'ingresso non invertente di un comparatore ad operazionale, la cui uscita risulta alta quando la tensione V^+ sul trasduttore supera la V^- , il che si verifica in condizioni di buio.

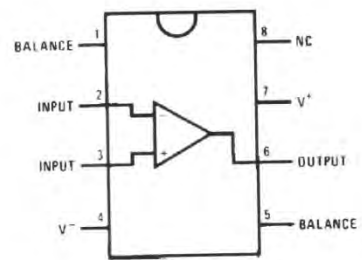
L'uscita alta dell'operazionale pone in conduzione un BJT (con funzione di amplificatore di corrente) che a sua volta determina la chiusura di un relé e quindi l'accensione della lampadina.

FUNZIONAMENTO DEL SISTEMA

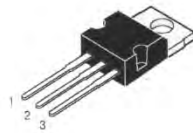
Al diminuire della luminosità dell'ambiente, il valore di resistenza del trasduttore cresce. Cresce quindi anche la tensione V_{trasd} ai suoi capi che coincide con il potenziale V^+ dell'ingresso non invertente dell'operazionale. Appena V^+ , crescendo, supera il potenziale V^- del terminale invertente, l'uscita dell'operazionale (che funziona da comparatore) diventa alta, facendo scattare il relé e determinando l'accensione della lampadina.

All'aumentare della luminosità dell'ambiente, il valore di resistenza del trasduttore diminuisce. Diminuisce quindi anche la tensione V_{trasd} ai suoi capi, che coincide con il potenziale V^+ dell'ingresso non invertente dell'operazionale. Appena V^+ , decrescendo, scende al di sotto del potenziale V^- del terminale invertente, l'uscita dell'operazionale diventa bassa, facendo scattare il relé, che si apre determinando lo spegnimento della lampadina.

Integrato TL081



Transistor BJT BD243



1base - 2collettore- 3emettitore

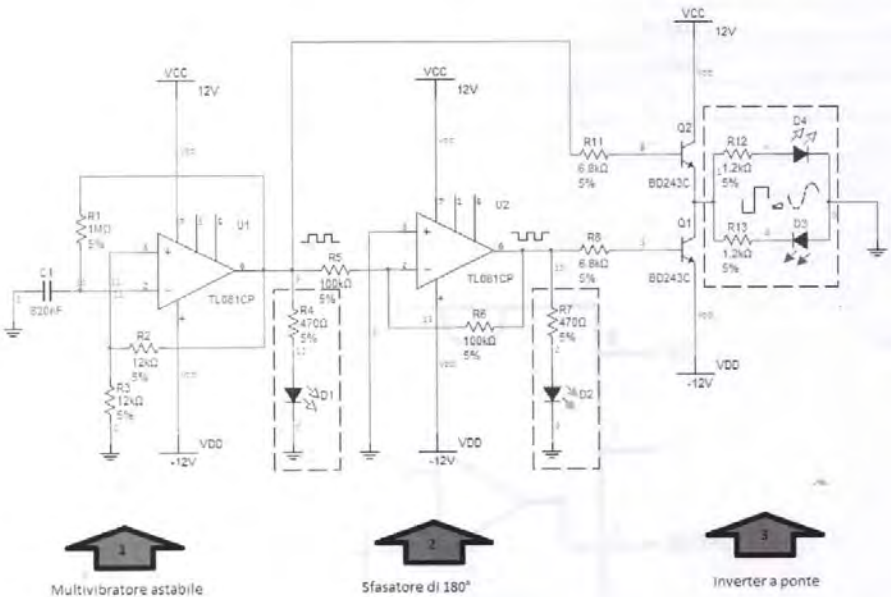
PROCEDIMENTO DI MISURA

- Posizionare la lampadina in modo tale che, quando sarà accesa, la sua luce non possa colpire la fotoresistenza
- Porre il cursore del trimmer a metà corsa e poi regolarlo in modo che, in piena luce, la lampada sia spenta
- Abbassare la luminosità dell'ambiente, o fare ombra sulla fotoresistenza, fino a quando si accende la lampadina. Se ciò non accade riprovare ruotando il trimmer (o, se necessario sostituendolo con uno di valore più elevato).
- Si eseguano le misure indicate nella tabella seguente

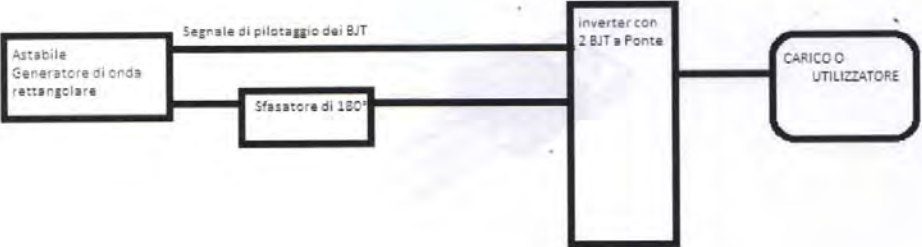
	AL BUIO	ALLA LUCE
STATO DELLA LAMPADINA (accesa/spenta)		
VALORE RESISTIVO DELLA FOTORESISTENZA		
V_{trasd}=V⁺ (tensione sulla fotoresistenza, coincidente col potenziale del terminale non invertente dell'operazionale)		
V⁻ (potenziale del terminale invertente dell'operazionale, coincidente con la tensione sulla serie R _B +R _{TRIM})		
V_{U40} (tensione di uscita dell'operazionale)		
V_C Potenziale del collettore del BJT rispetto a massa		
V_{AC} (tensione ai capi della bobina del relé)		

REALIZZAZIONE E VERIFICA DEL FUNZIONAMENTO DI UN INVERTER CON DUE BJT A SEMIPONTE

Schema circuitale:



Schema a blocchi:



STRUMENTI

- ALIMENTATORE
- MULTIMETRO DIGITALE

COMPONENTI

- C1= 820 nF
- R₁= 1 MΩ
- R₂= R₃= 12 KΩ
- R₅= R₆= 100 KΩ
- R₄= R₇= 100 KΩ
- R₈= R₁₁= 6,8 KΩ
- R₁₂= R₁₃= 1,2 KΩ
- AMPLIFICATORI OPERAZIONALI: 2 x TL081
- TRANSISTOR BJT: 2 x BD 243
- LED ROSSI (2)
- LED VERDI (2)
- MOTORINO CC

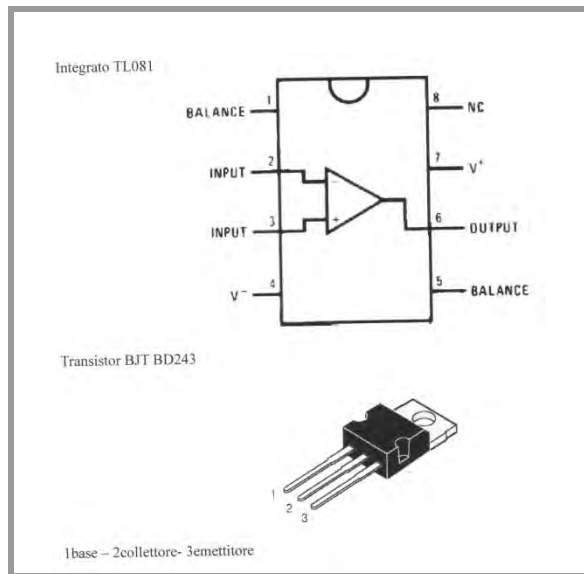
COSA FA IL CIRCUITO

L'inverter, in generale, è un dispositivo che converte tensioni e correnti continue in tensioni e correnti bipolari. Svolge quindi la funzione inversa dell'alimentatore.

Quello che noi vogliamo realizzare è un rudimentale inverter, il quale converte una tensione continua di 12 V in una tensione bipolare che viene applicata al carico.

Nella prima fase del collaudo il carico è una coppia di LED in antiparallelo: se il dispositivo funziona correttamente, i LED dovranno accendersi alternativamente a periodi di tempo regolari.

Una volta appurata la funzionalità del circuito con i led, si può sostituire ad essi un motorino cc che, per effetto della bipolarità della corrente, dovrà periodicamente invertire il senso di rotazione.



FUNZIONAMENTO E STRUTTURA DEL CIRCUITO

Il nucleo del dispositivo è la coppia di BJT cui è collegato il carico.

I BJT sono pilotati ciascuno da un'onda rettangolare. Le due onde sono sfasate fra loro di 180° (quando una è alta l'altra è bassa), in modo che quando un transistor è in conduzione (per effetto dell'impulso alto dell'onda rettangolare) l'altro è interdetto.

Siccome i due transistor fanno scorrere la corrente in verso opposto, l'alternarsi dei BJT in conduzione determina lo scambio periodico del verso della corrente sul carico e quindi la presenza di tensione e corrente bipolare sul carico (cioè sull'antiparallelo di LED prima, e sul motorino poi).

Chi fornisce le onde rettangolari di pilotaggio ai BJT?

Il blocco 1 dello schema è un multivibratore astabile ad operazionale (cioè un generatore d'onda rettangolare) la cui uscita pilota la base del BJT posto in alto nel disegno.

Il segnale di uscita del blocco 1 viene anche prelevato dal blocco 2, che lo sfasa di 180° e lo applica alla base del secondo transistor (in basso nello schema). Il blocco 2 è un amplificatore invertente ad operazionale.

LE FASI DELL'ESPERIENZA

E' consigliabile realizzare i circuiti uno alla volta e verificarne il funzionamento prima di realizzare il blocco successivo.

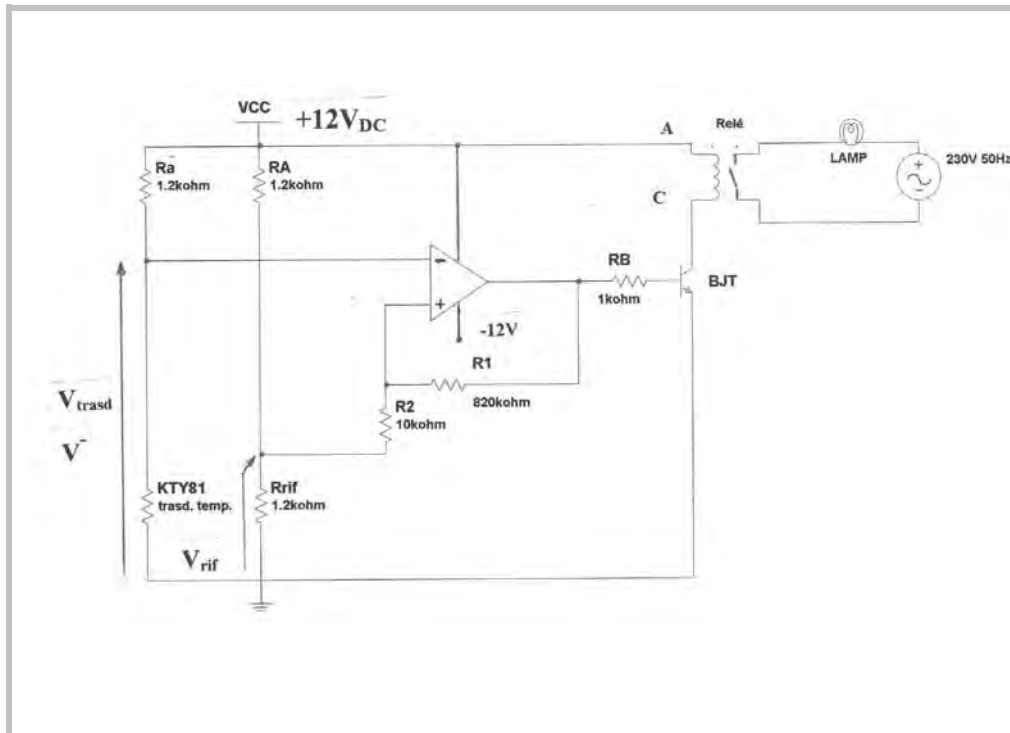
Si può procedere nel modo seguente:

- Si realizza il blocco 1 (generatore di onda rettangolare) e se ne verifica il funzionamento attraverso l'accensione intermittente del LED 1
- Si realizza il blocco 2 (sfasatore di 180°) e se ne verifica il funzionamento attraverso l'accensione intermittente del LED 2: il LED 2 deve risultare acceso quando il LED 1 è spento e viceversa
- Si realizza il semiponte di BJT, collegato al parallelo dei LED di prova, verificando attraverso il comportamento dei LED il corretto funzionamento del dispositivo. Cronometrando i tempi di accensione dei LED si possono misurare (in maniera grossolana) il periodo e la frequenza di funzionamento dell'inverter
- Se il circuito funziona correttamente con il parallelo di LED di carico, si può sostituire ai LED un motorino cc (in funzione del quale l'inverter è stato dimensionato). Anche stavolta, cronometrando i periodi di rotazione, in un verso e nell'altro, del motorino si possono misurare (in maniera grossolana) il periodo e la frequenza di funzionamento dell'inverter.

La frequenza teorica del generatore d'onda rettangolare, e quindi dell'intero dispositivo, si può determinare con l'espressione approssimata:

$$f = \frac{1}{2,2 \cdot R_1 \cdot C_1}$$

SISTEMA DI CONTROLLO ON/OFF DELLA TEMPERATURA DI UN CONTENITORE



ATTENZIONE

Questa esperienza, prevedendo l'uso della tensione di rete, va condotta sotto attenta vigilanza del docente.

Lampadina e trasduttore devono essere posti all'interno del contenitore di cui si vuole controllare la temperatura.

La parte di circuito a sinistra dei punti AC (circuito di controllo) può essere realizzata su breadboard. La parte a destra (circuito di potenza) mediante un normale cavo elettrico con adeguata portata.

STRUMENTI
MULTIMETRO DIGITALE
TERMOMETRO

COMPONENTI
SI VEDA LO SCHEMA

Inoltre

BJT BD 243

RELE' 3A 250V_{AC} 3A 30V_{DC}

TRASDUTTORE KTY 81-1

Lampadina 230V 150W

Porta lampada

Spina

FINALITA' DEL SISTEMA

Si vuole mantenere costante, per esempio a 50°C, la temperatura di un contenitore (una scatola di cartone o di polistirolo) utilizzando come elemento riscaldante una lampadina alimentata dalla tensione di rete e come trasduttore un PTC KTY 81-1, la cui resistenza elettrica aumenta all'aumentare della temperatura, come indicato nella tabella fornita dal costruttore e riportata nelle pagine seguenti.

La temperatura desiderata di 50° viene imposta attraverso la resistenza di riferimento R_{rif} il cui valore (1200 Ohm) corrisponde a 50°C circa secondo la legge di funzionamento $R=R(T)$ del trasduttore.

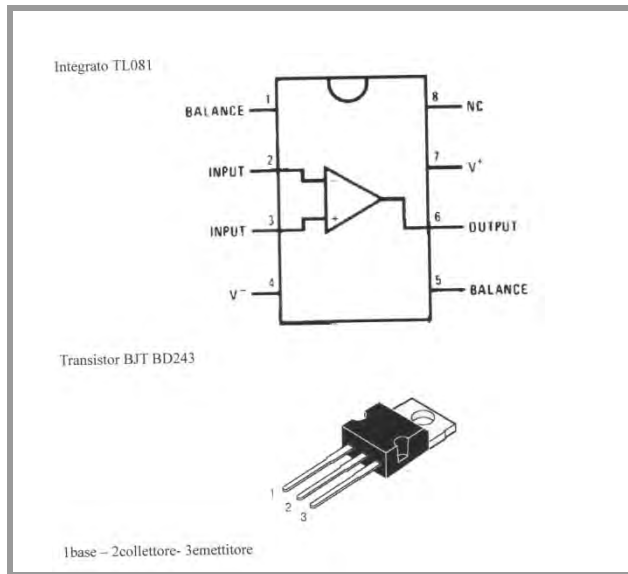
Ai capi di R_{rif} sarà presente una tensione che costituirà la tensione di riferimento del trigger di Schmitt invertente.

Il confronto fra la tensione proporzionale alla temperatura effettiva del contenitore (V_{trasd} coincidente con V^- , cioè con il potenziale dell'ingresso invertente dell'operazionale) e la tensione V_{rif} , proporzionale alla temperatura desiderata, è effettuato da un trigger di Schmitt invertente ad operazionale che pilota, attraverso un transistor, il relé che consente l'accensione e lo spegnimento dell'elemento riscaldante.

FUNZIONAMENTO DEL SISTEMA

Quando, al crescere della temperatura del contenitore, la tensione V_{trasd} sul trasduttore cresce, superando la tensione di soglia superiore del trigger, la tensione di uscita dell'amp op. diventa bassa, il relè si apre e la lampadina si spegne.

Quando invece, al diminuire della temperatura del contenitore, la tensione V_{trasd} sul trasduttore decresce, scendendo al di sotto della tensione di soglia inferiore del trigger, la tensione di uscita dell'amp op. diventa alta, il relè si chiude e la lampadina si accende.



PROCEDIMENTO DI MISURA

- Appena data tensione al sistema, la temperatura del contenitore è certamente minore di quella desiderata (50°C), per cui la lampadina si accende. **Si misuri V_{rif}** (che rimarrà costante durante tutta l'esperienza). Una volta che la lampadina sia accesa, il contenitore si riscalda progressivamente, finché, raggiunta e superata la temperatura di 50°C , la lampadina dovrebbe spegnersi.
- Nell'istante in cui la lampadina **si spegne**, misurare:
 - la **temperatura** nel contenitore
 - V_{trasd} coincidente con V^-
 - V^+ (che in questo istante coincide con la tensione di soglia superiore del trigger)
 - V_{rif}
- Una volta che la lampadina sia spenta, il contenitore (partendo da una temperatura superiore ai 50°) si raffredda progressivamente, fino a scendere al di sotto di 50°C . A questo punto la lampadina dovrebbe accendersi. **Nell'istante di accensione** si misurino:
 - la **temperatura** nel contenitore
 - V_{trasd} coincidente con V^-
 - V^+ (che in questo istante coincide con la tensione di soglia inferiore del trigger)
 - V_{rif}

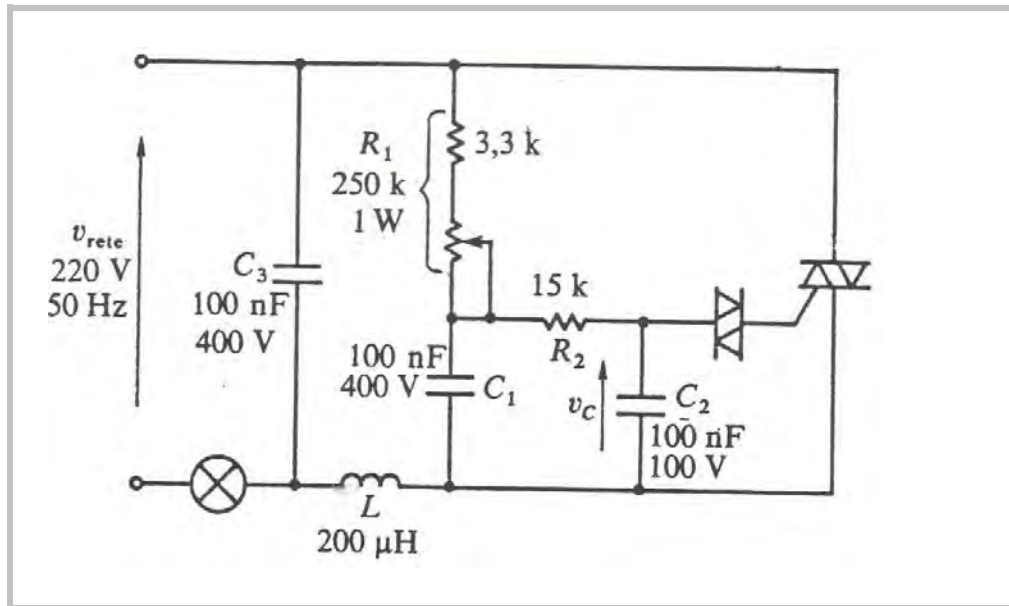
Silicon temperature sensors

KTY81-1 series

Table 1 Ambient temperature, corresponding resistance, temperature coefficient and maximum expected temperature error for KTY81-110 and KTY81-120 $I_{\text{cont}} = 1 \text{ mA}$

AMBIENT TEMPERATURE		TEMP. COEFF.	KTY81-110				KTY81-120				
(°C)	(°F)	(%/K)	RESISTANCE (Ω)			TEMP. ERROR (K)	RESISTANCE (Ω)			TEMP. ERROR (K)	
			MIN.	TYP.	MAX.		MIN.	TYP.	MAX.		
-55	-67	0.99	475	490	505	±3.02	470	490	510	±4.02	
-50	-58	0.98	500	515	530	±2.92	495	515	535	±3.94	
-40	-40	0.96	552	567	582	±2.74	547	567	588	±3.78	
-30	-22	0.93	609	624	638	±2.55	603	624	645	±3.62	
-20	-4	0.91	669	684	698	±2.35	662	684	705	±3.45	
-10	14	0.88	733	747	761	±2.14	726	747	769	±3.27	
0	32	0.85	802	815	828	±1.91	793	815	836	±3.08	
10	50	0.83	874	886	898	±1.67	865	886	907	±2.88	
20	68	0.80	950	961	972	±1.41	941	961	982	±2.66	
25	77	0.79	990	1000	1010	±1.27	980	1000	1020	±2.54	
30	86	0.78	1029	1040	1051	±1.39	1018	1040	1061	±2.68	
40	104	0.75	1108	1122	1136	±1.64	1097	1122	1147	±2.97	
50	122	0.73	1192	1209	1225	±1.91	1180	1209	1237	±3.28	
60	140	0.71	1278	1299	1319	±2.19	1265	1299	1332	±3.61	
70	158	0.69	1369	1392	1416	±2.49	1355	1392	1430	±3.94	
80	176	0.67	1462	1490	1518	±2.8	1447	1490	1532	±4.3	
90	194	0.65	1559	1591	1623	±3.12	1543	1591	1639	±4.66	
100	212	0.63	1659	1696	1733	±3.46	1642	1696	1750	±5.05	
110	230	0.61	1762	1805	1847	±3.83	1744	1805	1865	±5.48	
120	248	0.58	1867	1915	1963	±4.33	1848	1915	1982	±6.07	
125	257	0.55	1919	1970	2020	±4.66	1899	1970	2040	±6.47	
130	266	0.52	1970	2023	2077	±5.07	1950	2023	2097	±6.98	
140	284	0.45	2065	2124	2184	±6.28	2043	2124	2205	±8.51	
150	302	0.35	2145	2211	2277	±8.55	2123	2211	2299	±11.43	

VARIATORE DI LUMINOSITA' (DIMMER) PER LAMPADA A INCANDESCENZA CON DIAC E TRIAC



ATTENZIONE

Questa esperienza, prevedendo l'uso della tensione di rete, va condotta sotto attenta vigilanza del docente.

Condensatori, induttori, resistori, trimmer e DIAC possono essere collegati su breadboard, mentre la parte di circuito comprendente il TRIAC, la lampada e la spina per l'alimentazione di rete può essere realizzata su normale filo da elettricista di sufficiente portata.

Al Triac (che non deve essere toccato quando il circuito è sotto tensione) può essere applicato uno zoccolo che permetta di effettuare le saldature.

Sarebbe opportuno racchiudere il TRIAC in un contenitore di plastica.)

STRUMENTI
MULTIMETRO DIGITALE
OSCILLOSCOPIO
COMPONENTI
SI VEDA LO SCHEMA
Inoltre
TRIAC TIC 216 M
Lampadina 230V
Portalampada
Spina

FUNZIONAMENTO DEL SISTEMA

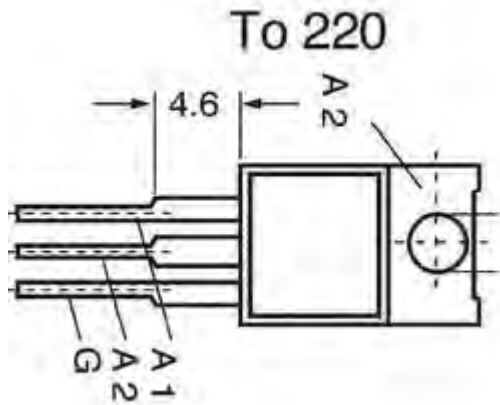
La variazione di luminosità si effettua attraverso il trimmer facente parte di R_1 . Il trimmer agisce sull'innesco del DIAC che a sua volta pilota il TRIAC.

Quanto più piccola è la resistenza selezionata dal trimmer, tanto più veloce è la carica capacitiva, tanto maggiore è l'angolo di conduzione del TRIAC e tanto più intensa è la luminosità della lampada.

Il gruppo R_2C_2 serve a rendere uniformemente graduale l'accensione della lampadina, che altrimenti presenterebbe un "salto".

Il gruppo LC_3 costituisce un filtro contro i disturbi ad alta frequenza che si generano in corrispondenza dei fronti di salita della corrente.

TRIAC TIC 216 M= BTB 06-600 D



PROCEDIMENTO DI MISURA

Una volta realizzato il circuito, la sua funzionalità è verificabile ad occhio. E' comunque possibile visualizzare le forme d'onda mediante oscilloscopio.

CAPITOLO 24

ALIMENTATORI LINEARI E SWITCHING

OBIETTIVI

- **Comprendere la funzione svolta dagli alimentatori e sapere in quali apparecchiature sono applicati**
- **Conoscere la differenza fra alimentatori lineari e alimentatori switching o a commutazione**
- **Saper rappresentare l'alimentatore lineare come schema a blocchi e comprendere la funzione di ciascun blocco**
- **Conoscere le tipologie di alimentatore lineare (a semionda, a onda intera, duale)**
- **Conoscere e saper valutare i parametri e le forme d'onda di un alimentatore lineare**
- **Saper dimensionare – per grandi linee – un alimentatore lineare**
- **Conoscere le prestazioni e almeno una delle tipologie di alimentatori switching**

PREREQUISITI

- **Basi di elettrotecnica**
- **Nozioni fondamentali su: filtri, raddrizzatori a diodo, trasformatore**

ALIMENTATORI (POWER SUPPLY o PSU, Power Supply Unit)

La maggior parte dei dispositivi e dei circuiti elettronici richiede, per funzionare, un'alimentazione, cioè una “fornitura” di energia elettrica sotto forma di tensioni e correnti **continue**.

L'alimentazione può essere singola, come ad esempio per le porte logiche TTL, che richiedono +5V, o duale, come per gli amplificatori operazionali, che possono richiedere +12V e -12V.

Le **tensioni** fornite dagli alimentatori sono normalmente comprese **fra qualche volt e qualche decina di volt**, con valori di **corrente** che possono andare **da qualche mA alle decine di A**.

Per alimentare i dispositivi elettronici, vengono comunemente impiegati alimentatori e batterie.

Le batterie però sono essenzialmente impiegate nella apparecchiature portatili a causa della loro portata limitata e del costo piuttosto elevato.

Più frequentemente sono impiegati gli **alimentatori**, che sono dispositivi in grado di **convertire la tensione alternata di rete in una tensione continua dell'ampiezza desiderata**.

Gli alimentatori possono essere “**non stabilizzati**” oppure “**stabilizzati**”.

Un alimentatore si dice “stabilizzato” quando la **tensione continua di uscita è insensibile alle variazioni della tensione di rete e del carico**.

Gli alimentatori possono inoltre essere di tipo “**lineare**” o di tipo “**switching**” (o “a commutazione”).

Gli alimentatori **switching** prevedono la **conversione della tensione di ingresso in onda rettangolare** e la **successiva riconversione in una tensione continua del valore desiderato**.

Essi possono anche essere “**a risonanza**” (SRPS) nel caso in cui **le onde rettangolari vengono rese sinusoidali** per mezzo di un **filtro risonante LC**. Questa conversione riduce l’emissione dei disturbi elettromagnetici dovuti alle armoniche degli impulsi rettangolari.

Una delle possibili classificazioni degli alimentatori è:

- ❖ Alimentatori “**lineari**” (alimentatori “**tradizionali**”)
- ❖ Alimentatori “**switching**” o “**a commutazione**” o “**SMPS**” (Switching Mode Power Supply)
- ❖ Alimentatori **switching “a risonanza”** o “**SRPS**”.



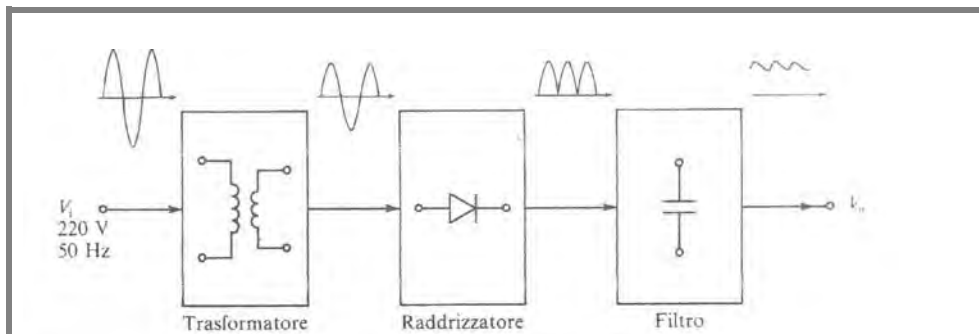
Alimentatore lineare: si notano il trasformatore e il condensatore di livellamento



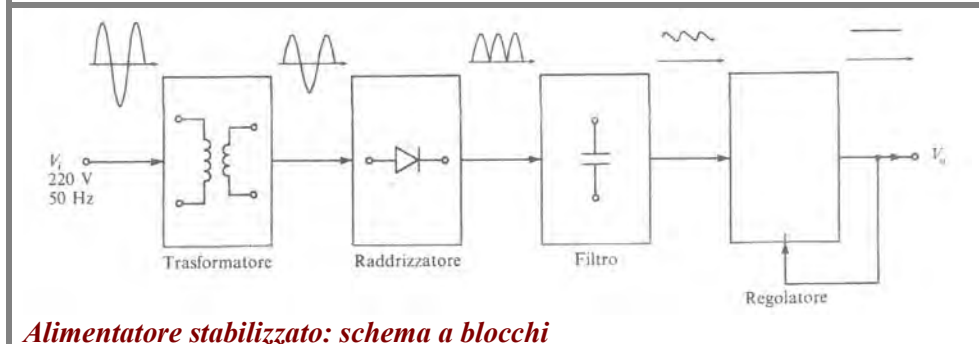
Alimentatore switching o a commutazione

Gli alimentatori **tradizionali**, basati sulla cascata “**trasformatore-raddrizzatore-filtro-stabilizzatore**” sono **lineari**.

Gli amplificatori lineari hanno una **struttura circuitale più semplice** rispetto agli switching, ma **non risultano convenienti quando la potenza richiesta supera i 50W**. Tratteremo per primi gli alimentatori non stabilizzati.



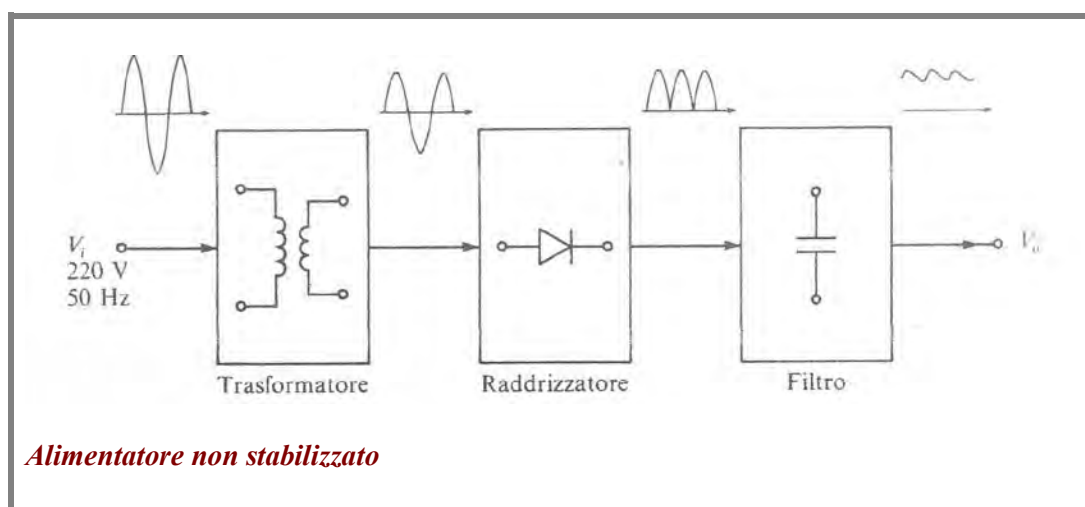
Alimentatore non stabilizzato: schema a blocchi



Alimentatore stabilizzato: schema a blocchi

ALIMENTATORI LINEARI NON STABILIZZATI

Un alimentatore lineare può essere rappresentato mediante uno schema a blocchi che comprende un trasformatore, un raddrizzatore e un filtro passabasso.



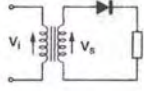
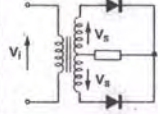
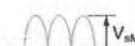
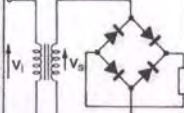
SCHEMA A BLOCCHI

La funzione dei singoli blocchi è illustrata di seguito.

- **TRASFORMATORE:** converte la tensione alternata di rete in una tensione ancora alternata, ma di ampiezza diversa.
Inoltre isola galvanicamente i componenti dell'alimentatore dalla rete.
- **RADDRIZZATORE:** è un sistema di diodi che converte la tensione di rete, caratterizzata da valor medio nullo, in una tensione a valor medio diverso da zero (tensione unipolare o unidirezionale pulsante); i raddrizzatori possono essere “a semionda” oppure a onda intera¹;
per la realizzazione di un raddrizzatore a semionda è sufficiente un solo diodo, un raddrizzatore a onda intera può essere realizzato con quattro diodi (raddrizzatore a ponte di Graetz) oppure con due soli diodi, nel qual caso il trasformatore, costituente il blocco a monte, deve avere il secondario a presa centrale.
- **FILTRO PASSABASSO o FILTRO DI LIVELLAMENTO:**
la tensione unipolare pulsante presente all'uscita del raddrizzatore è formata da una componente continua (che rappresenta il valore medio nel periodo) cui è sovrapposta una componente periodica; il filtro ha la funzione di eliminare o ridurre

¹ I raddrizzatori a onda intera sono anche detti “a doppia semionda”

la componente periodica; il filtro può essere realizzato, come nello schema a blocchi, da un solo condensatore posto in parallelo al carico, o da un solo induttore posto in serie al carico o, più in generale, da uno o più circuiti RC oppure LC (questi ultimi hanno il vantaggio di ridurre la dissipazione di potenza per effetto Joule).

Tipo	Schema elettrico	Forma d'onda in uscita	Valore medio della tensione di uscita	Massima tensione inversa per ogni diodo
A semionda			$\frac{V_{sM}}{\pi}$	V_{sM}
A doppia semionda			$\frac{2 V_{sM}}{\pi}$	$2 V_{sM}$
A ponte			$\frac{2 V_{sM}}{\pi}$	V_{sM}

Tipi di raddrizzatore e loro caratteristiche.

La “massima tensione inversa per ogni diodo” è spesso indicata come “PIV”, Peak Inverse Voltage

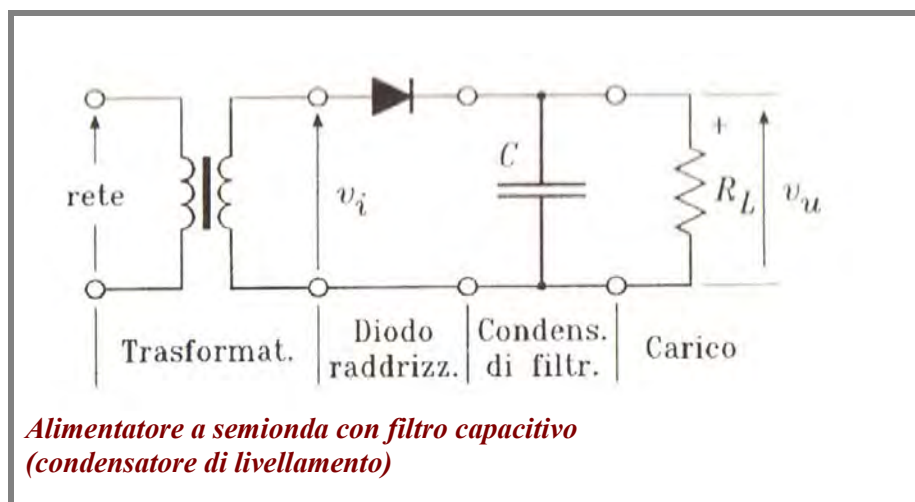
RIPPLE O ONDULAZIONE RESIDUA

La tensione di uscita di un alimentatore **non** è in realtà continua, ma contiene **un'ondulazione residua (ripple)**, che converrà sia il più possibile piccola.

Il ripple è un segnale indesiderato che, se il carico è un altoparlante, dà luogo a un suono cupo di fondo, il cosiddetto “ronzio”.

Spesso si usano indifferentemente i termini “ondulazione” “ripple” o “ronzio”.

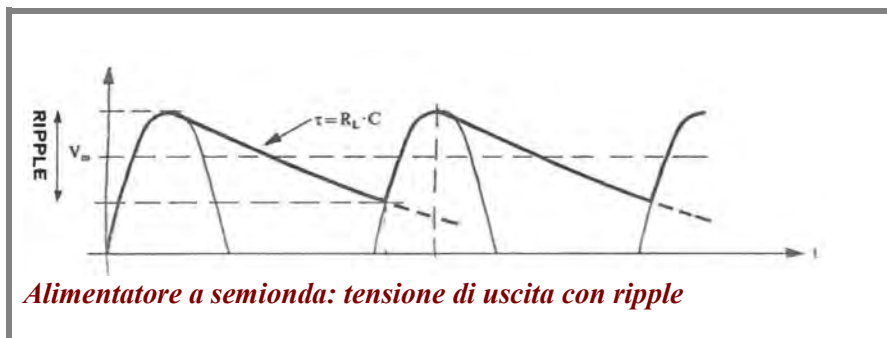
ALIMENTATORE A SEMIONDA CON FILTRO CAPACITIVO



Quello dell'alimentatore non stabilizzato con filtro capacitivo è lo schema più semplice di alimentatore.

COME FUNZIONA L'ALIMENTATORE A SEMIONDA

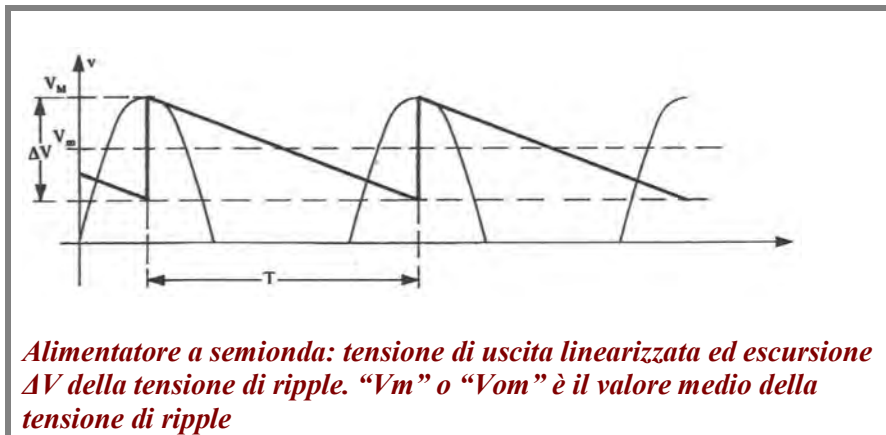
L'alimentatore a semionda fornisce in uscita, cioè sul condensatore e sul carico, una tensione unipolare, ma non costante (e quindi **non** continua), caratterizzata da un'ondulazione chiamata ripple, che costituisce un disturbo.



La tensione di uscita è quindi formata da una componente continua V_{Om} a cui è sovrapposta una componente periodica di disturbo, con frequenza f , che è appunto il ripple. Vedremo che il valore efficace della tensione di ripple è:

$V_{rEFF} = \frac{V_{Om}}{2 \cdot \sqrt{3} \cdot f \cdot R_L \cdot C}$	<i>tensione di ripple: valore efficace</i>
--	--

Il ripple è un disturbo che deve essere ridotto il più possibile.



Per misurare l’entità del ripple rispetto alla componente continua utile del segnale di uscita si definisce il fattore di ripple:

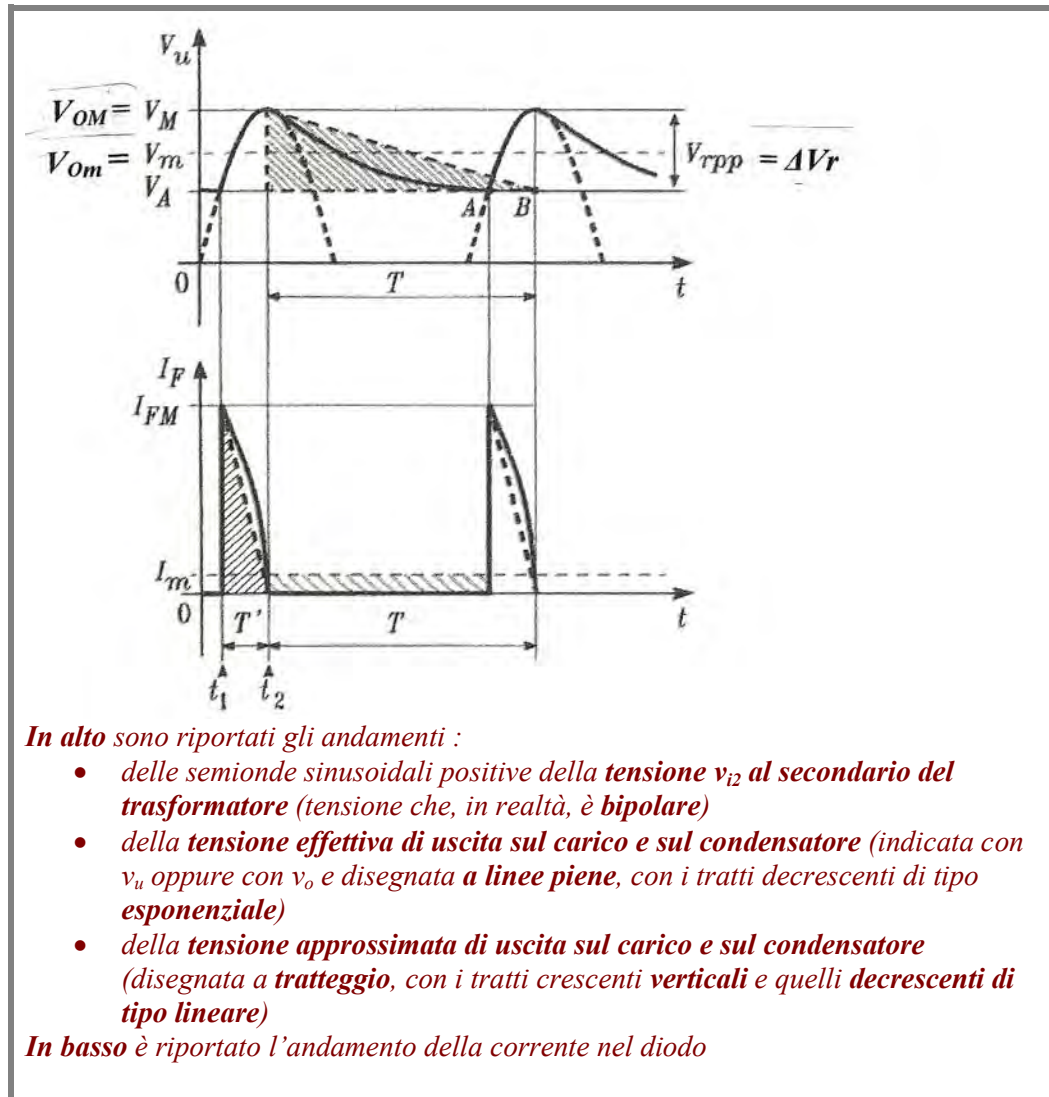
$r_{\%} = \frac{V_{reff}}{V_{Om}} = \frac{\text{valore efficace della tensione di ripple}}{\text{valore medio della tensione di uscita}} \cdot 100$	<i>fattore di ripple:</i>
---	---------------------------

che per l’alimentatore in esame assume l’espressione:

$r_{\%} = \frac{1}{2 \cdot \sqrt{3} \cdot f \cdot R_L \cdot C} \cdot 100$	<i>fattore di ripple:</i>
---	---------------------------

ANALISI DELL'ALIMENTATORE A SEMIONDA CON FILTRO CAPACITIVO

ANDAMENTI DELLE TENSIONI, RIPPLE E SUA RIDUZIONE



Durante la prima metà della semionda positiva della tensione di alimentazione (e quindi della tensione al secondario del trasformatore), il diodo, ponendosi in conduzione, carica il

condensatore al valore massimo V_{OM} (o semplicemente V_M) della tensione v_{i2} al secondario del trasformatore.

Durante la seconda metà della semionda positiva, e fino alla successiva semionda positiva, il diodo rimane interdetto e **il condensatore si scarica** sulla resistenza di carico R_L con **legge esponenziale**, per cui il valore istantaneo della tensione di uscita (che indicheremo con V_u oppure con V_o) è:

$$v_o(t) \equiv v_u(t) = V_{OM} \cdot e^{\frac{-t}{R \cdot C}}$$

Nella prima parte del secondo periodo, quando la tensione al secondario del trasformatore supera il valore residuo della tensione di uscita (v_o) ai capi del condensatore, il diodo si pone nuovamente in conduzione e riporta la tensione di uscita (ai capi del condensatore e del carico) sul valore massimo V_{OM} della tensione v_{i2} al secondario del trasformatore.

Poi, a partire dalla seconda metà della semionda del secondo periodo, si ha una nuova scarica del condensatore e l'andamento torna a ripetersi sempre nello stesso modo. Quello che abbiamo ora descritto è l'andamento effettivo della tensione di uscita, andamento che nella figura precedente è disegnato a tratto pieno.

Possiamo notare che la **tensione effettiva di uscita sul carico e sul condensatore** (indicata con v_u oppure con v_o e disegnata a **linee piene**, con i tratti decrescenti di tipo **esponenziale**) è unipolare, ma non è costante: infatti varia con periodicità, senza invertirsi, intorno a un livello costante (V_{Om} o V_m), dando luogo all'ondulazione indesiderata che è chiamata "ripple".

Riguardo alla tensione di **ripple**, ne indicheremo con V_{rpp} , oppure con ΔV_r , il **valore picco-picco**, cioè l'**escursione**, mentre il **valore efficace** lo indichiamo con $V_{r_{eff}}$.

Detto V_{Om} (o V_m) il valore medio della tensione di uscita, si definisce il "**fattore di ripple**" come il **rapporto** fra il valore efficace della tensione di ripple e il valore medio V_{Om} della tensione di uscita (tensione sul carico e sul condensatore):

$r = \frac{V_{\text{reff}}}{V_{\text{Om}}} = \frac{\text{valore efficace della tensione di ripple}}{\text{valore medio della tensione di uscita}}$	<i>fattore di ripple:</i>
--	---------------------------

Il fattore di ripple può anche essere espresso in forma percentuale:

$r_{\%} = \frac{V_{\text{reff}}}{V_{\text{Om}}} \cdot 100$	<i>fattore di ripple:</i>
--	---------------------------

L'ondulazione residua presente nella tensione raddrizzata e filtrata dà luogo a disturbi (ronzii) negli apparati alimentati e quindi deve essere mantenuta al valore più basso possibile. Affinché ciò accada occorre che la costante di tempo $\tau = R_L \cdot C$ sia più grande possibile.

La capacità deve essere tanto più elevata quanto maggiore è la corrente assorbita dal carico, cioè quanto più piccola è la R_L .

(Analoghi risultati, a parità di costante di tempo $\tau = R_L \cdot C$, si ottengono con un raddrizzatore a onda intera, come vedremo in seguito).

Se si vuole ridurre il ripple senza ricorrere a capacità C di valore troppo elevato, si interpone, fra il condensatore di livellamento e il carico, un filtro $R'C'$ passabasso che riduca ulteriormente l'ondulazione sovrapposta alla tensione continua utile presente a capi di C .

Si possono collegare anche più filtri in cascata, ottenendo una notevole attenuazione del ripple. Quest'attenuazione è massima quando i componenti delle diverse sezioni hanno gli stessi valori.

Quando la **corrente nel carico** è piuttosto **elevata**, si usa un **filtro LC**, evitando così la caduta di tensione sulla resistenza, dovuta alla componente continua della corrente.

ANALISI DELL' ALIMENTATORE A SEMIONDA CON FILTRO CAPACITIVO

ANDAMENTO APPROSSIMATO DELLA TENSIONE DI USCITA

Per conoscere il legame esistente fra il valore del ripple e i componenti dell'alimentatore, è opportuno approssimare l'andamento della tensione di uscita (che, come già detto, è la tensione sia sul carico che sul condensatore) con un andamento linearizzato a onda triangolare.

L'approssimazione si effettua nel modo seguente:

- il tratto sinusoidale di ricarica del condensatore viene sostituito con un tratto verticale di ampiezza $V_{r_{pp}}$, ossia ΔV_r , il che equivale a supporre che la carica del condensatore avvenga istantaneamente
- il tratto esponenziale di scarica del condensatore sul carico R_L viene linearizzato e portato fino all'asse della successiva semionda di ingresso (a una quota inferiore di ΔV_r rispetto al valore massimo V_{OM} della tensione di uscita); ciò equivale a supporre che la tensione sul condensatore (tensione di uscita) vari linearmente durante la scarica.
- Se la tensione sul condensatore decresce in maniera lineare, la **corrente sul condensatore** dovrà ritenersi **costante**, per la relazione tensione-corrente del condensatore:

$$i = C \cdot \frac{dV}{dt}$$

(la derivata di una funzione lineare, ossia di una funzione rappresentabile con una retta inclinata, è infatti una costante esprimibile con una retta orizzontale).

Diciamo **I** la **corrente sul condensatore**. Questa **corrente I**, essendo **costante**, coincide con il suo **valore medio I_m**.

Notiamo inoltre che, quando cercheremo la relazione fra la tensione di ripple e i componenti del circuito, esamineremo l'intervallo di tempo in cui il condensatore si scarica. In questo intervallo il diodo è interdetto, per cui il carico riceve corrente dal condensatore e non dalla rete, e tutta la corrente del condensatore va sul carico. Cerchiamo ora la relazione fra tensione di ripple e componenti del circuito.

Ricordiamo che la capacità C di un condensatore è definita come rapporto fra la quantità di carica Q e la differenza di potenziale V fra le armature del condensatore:

$$C = \frac{Q}{V}$$

Di conseguenza, se con ΔQ indichiamo la carica perduta dal condensatore nel periodo T, e con ΔV_r il valore picco-picco della tensione di ripple, possiamo scrivere:

$$C = \frac{\Delta Q}{\Delta V_r}$$

e quindi:

$\Delta V_r = \frac{\Delta Q}{C}$	<i>tensione di ripple: valore picco-picco o escursione</i>
-----------------------------------	--

Siccome ΔQ è la carica perduta nell'intervallo T dal condensatore (per fornire al carico la corrente I), essa è data dal prodotto:

$$\Delta Q = I \cdot T$$

Di conseguenza (sostituendo l'espressione di ΔQ nell'equazione precedente) l'escursione o valore picco-picco della tensione di ripple diventa:

$\Delta V_r = \frac{I \cdot T}{C}$	<i>tensione di ripple: valore picco-picco o escursione</i>
------------------------------------	--

da cui:

$V_{r \max} = \frac{I \cdot T}{2 \cdot C} = \frac{I}{2 \cdot f \cdot C}$	<i>tensione di ripple: valore massimo</i>
--	---

Siccome I è la corrente sul carico, che supponiamo continua, detta V_{Om} la tensione media sul carico R_L , possiamo applicare la legge di Ohm ad R_L e scrivere:

$$I = \frac{V_{Om}}{R_L}$$

che, sostituita, al posto di I, nell'ultima espressione di ΔV_r fornisce:

$\Delta V_r = \frac{V_{Om}}{R_L} \cdot \frac{T}{C}$	<i>tensione di ripple: valore picco-picco o escursione</i> *
---	--

Determiniamo ora il valore massimo V_{rMAX} della tensione di ripple, che è la metà del valore picco-picco:

$V_{rMAX} = \frac{\Delta V_r}{2} = \frac{1}{2} \frac{V_{Om}}{R_L} \cdot \frac{T}{C} = \frac{1}{2} \frac{V_{Om}}{f \cdot R_L \cdot C}$	<i>tensione di ripple: valore massimo</i>
---	---

In particolare ci interessa il valore efficace della tensione di ripple. Siccome la tensione di ripple è stata approssimata a un'onda triangolare, il suo valore efficace sarà $\sqrt{3}$ volte più piccolo del valore massimo:

$$V_{rEFF} = \frac{V_{rMAX}}{\sqrt{3}} = \frac{\Delta V_r}{2 \cdot \sqrt{3}} = \frac{1}{2 \cdot \sqrt{3}} \cdot \frac{V_{Om}}{R_L} \cdot \frac{T}{C}$$

Abbiamo quindi ottenuto:

$V_{reff} = \frac{V_{Om} \cdot T}{2 \cdot \sqrt{3} \cdot R_L \cdot C}$	<i>valore efficace della tensione di ripple per l'alimentatore a semionda</i>
--	---

Volendo far comparire la frequenza $f=1/T$ dell'onda sinusoidale di ingresso, si ottiene:

$V_{rEFF} = \frac{V_{Om}}{2 \cdot \sqrt{3} \cdot f \cdot R_L \cdot C}$	<i>valore efficace della tensione di ripple per l'alimentatore a semionda</i>
--	---

Riguardo al **fattore di ripple**:

$$r_{\%} = \frac{V_{reff}}{V_{Om}} \cdot 100$$

sostituendo in esso l'espressione del valore efficace della tensione di ripple otteniamo:

$$r_{\%} = \frac{\frac{V_{Om}}{2 \cdot \sqrt{3} \cdot f \cdot R_L \cdot C}}{V_{Om}} \cdot 100$$

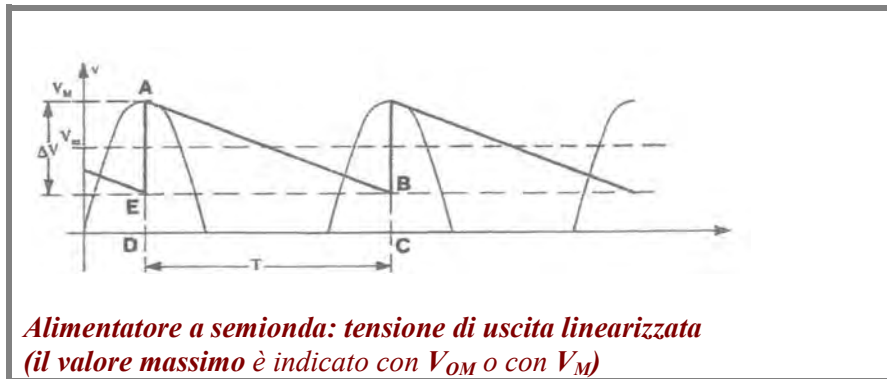
e quindi:

$r_{\%} = \frac{1}{2 \cdot \sqrt{3} \cdot f \cdot R_L \cdot C} \cdot 100$	<i>fattore di ripple per l'alimentatore a semionda</i>
---	--

Dalle relazioni ottenute si vede che la **tensione di ripple** e il **fattore di ripple** sono **inversamente proporzionali** al prodotto $f R_L C$, cioè $f \tau$.

Per una buona stabilizzazione, cioè per **ridurre il ripple**, occorre non soltanto una **capacità elevata**, ma anche una **resistenza di carico elevata**.

VALORE MEDIO DELLA TENSIONE DI USCITA DELL'ALIMENTATORE A SEMIONDA



Da un punto di vista matematico, il **valore medio**, nel periodo T , di un segnale $v(t)$ è definito come **rapporto** fra l'**integrale definito** del segnale (esteso al periodo) e il **periodo stesso**, cioè come:

$$V_{Om} = \frac{1}{T} \cdot \int_t^{t+T} v_o(\tau) \cdot d\tau$$

L'**integrale definito** del segnale, esteso al periodo, rappresenta però il valore dell'**area A_T** , compresa fra $v(t)$ e l'asse orizzontale, limitatamente alla lunghezza del periodo.

Perciò (siccome le aree in questo caso sono quelle di figure geometriche elementari), possiamo calcolare il **valore medio** come **rapporto** fra l'**area A_T** e il **periodo T** , determinando le aree "graficamente":

$$V_{Om} = \frac{A_T}{T}$$

Dalla figura si vede che tale area A_T può essere vista come la somma :

- dell'area A_{TRI} del triangolo ABE di base T e altezza ΔV_r
- dell'area A_{RET} del rettangolo EBCD di base T e altezza ED , pari a $V_{OM} - \Delta V_r$

Pertanto si ha:

$$V_{Om} = \frac{A_T}{T} = \frac{A_{TRI} + A_{RET}}{T} = \frac{(\Delta V_r \cdot T) / 2 + (V_{OM} - \Delta V_r) \cdot T}{T} = \frac{\frac{\Delta V_r}{2} + (V_{OM} - \Delta V_r)}{1}$$

e quindi:

$$V_{Om} = \frac{\Delta V_r}{2} + (V_{OM} - \Delta V_r) = V_{OM} - \Delta V_r + \frac{\Delta V_r}{2} = V_{OM} - \frac{\Delta V_r}{2}$$

In definitiva il valore medio della tensione di uscita è dato da:

$V_{Om} = V_{OM} - \frac{\Delta V_r}{2}$	Il valore medio della tensione di uscita dell'alimentatore a semionda è dato dal valore massimo dell'uscita meno il valore massimo del ripple
--	--

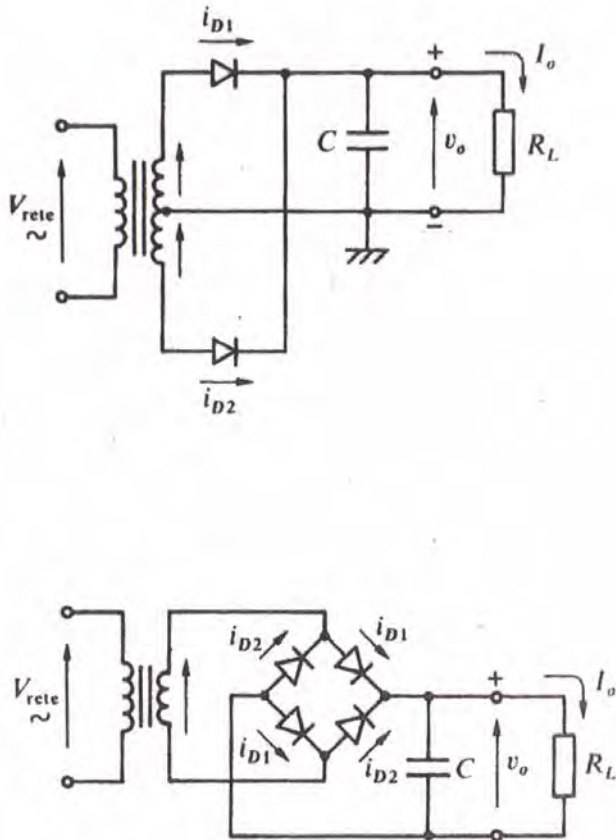
e, ricordando che:

$$\Delta V_r = \frac{I \cdot T}{C} = \frac{I}{f \cdot C}$$

si ha:

$V_{Om} = V_{OM} - \frac{I}{2 \cdot f \cdot C}$	valore medio della tensione di uscita dell'alimentatore a semionda
---	---

ALIMENTATORI A ONDA INTERA



Alimentatori a onda intera

- *In alto: alimentatore con trasformatore a presa centrale (la tensione su ciascun semiavvolgimento del secondario verrà chiamata v_{i2})*
- *In basso: alimentatore a ponte di Graetz (verrà chiamata v_{i2} la tensione sull'unico avvolgimento del secondario)*

CONFRONTO FRA ALIMENTATORI A ONDA INTERA E ALIMENTATORI A SEMIONDA

L'alimentatore **a semionda** è usato quando sono richieste basse correnti di carico, con *potenze normalmente inferiori al Watt*.

Nel circuito a semionda il diodo, in ciascun periodo della tensione di ingresso, conduce solo nell'unico breve intervallo di tempo nel quale la tensione v_{i2} al secondario del trasformatore supera la tensione v_o di uscita sul condensatore.

Per la rimanente parte del periodo, il diodo è interdetto e la corrente è quella fornita al carico dalla scarica condensatore stesso. Questa corrente attraversa il carico, determinando, nella tensione ai suoi capi, un'accentuata ondulazione (ripple).

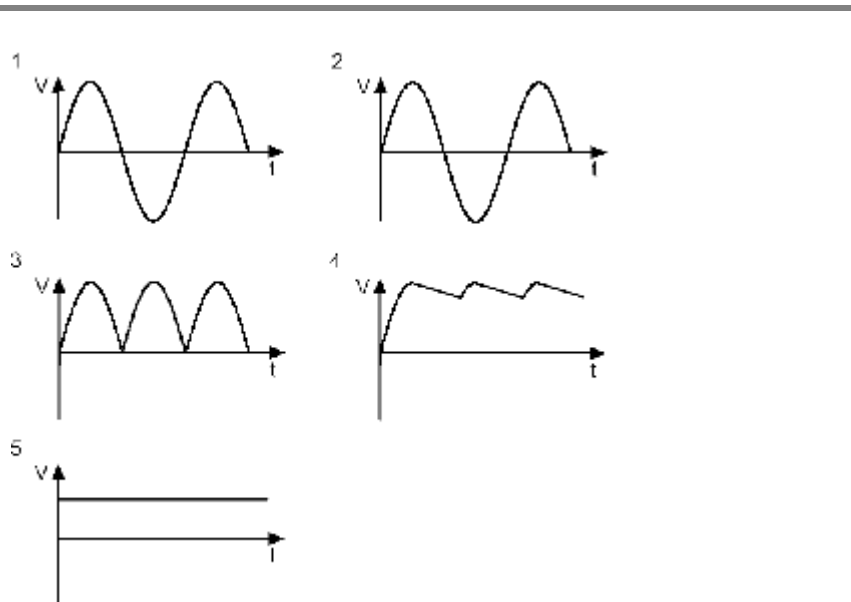
L'alimentatore a semionda inoltre presenta *un forte picco di corrente* sul diodo, dovuto al fatto che questo "impulso" di corrente *deve provvedere da solo*, e in un tempo molto breve, *alla ricarica del condensatore*.

Nei raddrizzatori **a onda intera** invece (sia a ponte che con trasformatore a presa centrale) **la ricarica del condensatore** avviene **ad ogni semiperiodo** della tensione di rete e quindi **due volte in ciascun periodo** (ciò accade perché vi è almeno un diodo che conduce in ogni semiperiodo della tensione di ingresso).

In sintesi, negli alimentatori a onda intera:

- **l'ampiezza del ripple** è più contenuta
- **i picchi di corrente** sui diodi sono **meno elevati**.

In particolare, a parità di resistenza di carico e di capacità del condensatore di livellamento, **l'ampiezza della tensione di ripple**, il **fattore di ripple** e la **corrente di picco** dei diodi risultano **dimezzati**.



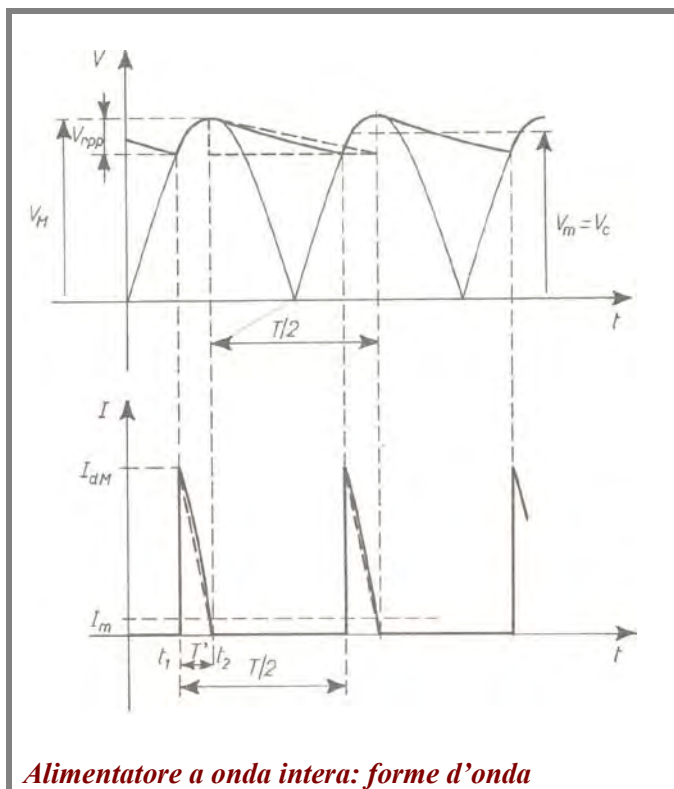
Alimentatore a onda intera: forme d'onda

- 1. tensione di ingresso (tensione di rete)***
- 2. tensione al secondario del trasformatore***
- 3. tensione unipolare pulsante che sarebbe presente all'uscita del raddrizzatore in assenza del filtro***
- 4. tensione di uscita caratterizzata da ripple***
- 5. tensione di uscita senza ripple (la riduzione del ripple si effettua mediante regolatori o stabilizzatori di tensione, presenti solo negli alimentatori stabilizzati, che tratteremo in seguito)***

ANALISI DEGLI ALIMENTATORI A ONDA INTERA (A PONTE DI GRAETZ E CON TRASFORMATORE A PRESA CENTRALE)

Anche nel caso degli alimentatori a onda intera con filtro capacitivo, per conoscere il legame esistente fra il valore del ripple e i componenti dell'alimentatore, si segue lo stesso procedimento adottato per il circuito a semionda. Si approssima cioè l'andamento della tensione di uscita (che, come già detto, è la tensione sia sul carico che sul condensatore) con un andamento linearizzato a onda triangolare.

Siccome all'uscita del raddrizzatore sono presenti anche le semionde che derivano dal raddrizzamento delle semionde negative della tensione al secondario del trasformatore, la scarica del condensatore dura $T/2$ (e non T).



Di conseguenza (in base a quanto visto per l'alimentatore a semionda) si ha:

- $\Delta Q = I \cdot \frac{T}{2}$
- il valore picco-picco della tensione di ripple **si dimezza**, diventando:

$$\Delta V_r = \frac{I \cdot T}{2 \cdot C} \rightarrow \Delta V_r = \frac{I}{2 \cdot f \cdot C}$$

- il valore massimo della tensione di ripple, definito sempre come:

$$V_{rMAX} = \frac{\Delta V_r}{2} \quad \text{si dimezza, diventando:} \quad V_{rMAX} = \frac{I}{4 \cdot f \cdot C}$$

- anche il valore efficace della tensione di ripple subirà una divisione per due:

$$V_{rEFF} = \frac{V_{rMAX}}{\sqrt{3}} = \frac{\Delta V_r}{2 \cdot \sqrt{3}} = \frac{1}{4 \cdot \sqrt{3}} \cdot \frac{V_{Om}}{R_L} \cdot \frac{T}{C}$$

In definitiva per gli alimentatori a onda intera si ha:

$V_{rEFF} = \frac{V_{Om} \cdot T}{4 \cdot \sqrt{3} \cdot R_L \cdot C}$	valore efficace della tensione di ripple per gli alimentatori a onda intera
--	--

e facendo comparire la frequenza $f=1/T$ dell'onda sinusoidale di ingresso:

$V_{rEFF} = \frac{V_{Om}}{4 \cdot \sqrt{3} \cdot f \cdot R_L \cdot C}$	valore efficace della tensione di ripple per gli alimentatori a onda intera
--	--

e inoltre:

$r_{\%} = \frac{1}{4 \cdot \sqrt{3} \cdot f \cdot R_L \cdot C} \cdot 100$	fattore di ripple per l'alimentatore a onda intera
---	---

VALORE MEDIO DELLA TENSIONE DI USCITA PER L'ALIMENTATORE A ONDA INTERA

Dalla figura che illustra le forme d'onda, si può ricavare un'espressione del valore medio della tensione di uscita identica a quella trovata per il circuito a semionda, e cioè un'espressione del tipo:

$$V_m = V_{OM} - \frac{\Delta V_r}{2}$$

Dove però l'ampiezza ΔV_r del ripple è la metà di quella del circuito a semionda, dal momento che, ora, in un periodo della tensione di rete vi sono due semionde positive della tensione raddrizzata e la tensione di uscita dell'alimentatore (tensione sul condensatore e sul carico) ha a disposizione la metà di tempo per decrescere.

Essendo poi:

$$\Delta V_r = \frac{I \cdot T}{2 \cdot C} = \frac{I}{2 \cdot f \cdot C}$$

si ha:

$V_{Om} = V_{OM} - \frac{I}{4 \cdot f \cdot C}$	valore medio della tensione di uscita per gli alimentatori a onda intera
---	---

Dove al denominatore è presente un 4 invece del 2 che figura nel circuito a semionda.

DIMENSIONAMENTO DI UN ALIMENTATORE A ONDA INTERA (A PONTE o CON TRASFORMATORE A PRESA CENTRALE)

SPECIFICHE DI PROGETTO

Le specifiche di progetto, cioè i dati assegnati, sono, in genere:

- la **massima corrente costante di uscita** I_O che l'alimentatore deve essere in grado di erogare (il "più grande valor medio" ottenibile per I_O)
- la **tensione continua di uscita** (cioè il valor medio V_{Om} , nel periodo, della v_O)

SVOLGIMENTO

- Se non è richiesto uno specifico valore del fattore "r" di ripple, si può imporre che sia:

$$r_{\%} \leq 10 \%$$

che corrisponde a²:

$$\frac{\Delta V_r}{V_{Om}} \leq 35 \%$$

² Infatti:

$$\Delta V_r = 2 \cdot V_{rMAX} = 2 \cdot V_{rEFF} \cdot \sqrt{3}$$
$$\frac{\Delta V_r}{V_{Om}} \cdot 100 = \frac{2 \cdot V_{rEFF} \cdot \sqrt{3}}{V_{Om}} \cdot 100 = 2 \cdot \sqrt{3} \cdot \frac{V_{rEFF}}{V_{Om}} \cdot 100 = 2 \cdot \sqrt{3} \cdot r_{\%} = 2 \cdot \sqrt{3} \cdot 10 \cong 35\%$$

- **Dimensionamento della CAPACITA' del filtro (primo modo)**

Le formule di dimensionamento della capacità del filtro derivano dall'espressione:

$$C = \frac{\Delta Q}{\Delta V_r}$$

Questa espressione di "C", tenendo conto del fatto che:

$$\Delta Q = I_{Om} \cdot \frac{T}{2}$$

si può scrivere come:

$$C = \frac{I_{Om} \cdot (T / 2)}{\Delta V_r}$$

ed essendo:

$$T = \frac{1}{f} \rightarrow \frac{T}{2} = \frac{1}{2 \cdot f}$$

si può anche esprimere come:

$$C = \frac{I_{Om}}{2 \cdot f \cdot \Delta V_r}$$

Moltiplicando e dividendo il denominatore per $2\sqrt{3}$ e ricordando che:

$$V_{rEFF} = \frac{\Delta V_r}{2\sqrt{3}}$$

si ha:

$$C = \frac{I_{Om}}{2 \cdot f \cdot \Delta V_r \cdot \frac{2\sqrt{3}}{2\sqrt{3}}}$$

$$C = \frac{I_{Om}}{4 \cdot f \cdot \sqrt{3} \cdot \frac{\Delta V_r}{2\sqrt{3}}}$$

e quindi:

$C = \frac{I_{Om}}{4 \cdot \sqrt{3} \cdot f \cdot V_{rEFF}}$	Formula di dimensionamento della CAPACITA' del filtro per un alimentatore a onda intera In funzione di V_{rEFF}
--	---

Se poi applichiamo la definizione di fattore di ripple:

$$r = \frac{V_{rEFF}}{V_{Om}} \quad \rightarrow \quad V_{rEFF} = r \cdot V_{Om}$$

e sostituiamo l'espressione di V_{rEFF} trovata in quella di C, possiamo scrivere:

$C = \frac{I_{Om}}{4 \cdot \sqrt{3} \cdot f \cdot r \cdot V_{Om}}$	Formula di dimensionamento della CAPACITA' del filtro, per un alimentatore a onda intera, in funzione del fattore di ripple "r"
--	--

Essendo infine:

$$R_L = \frac{V_{Om}}{I_{Om}}$$

possiamo anche scrivere:

$C = \frac{1}{4 \cdot \sqrt{3} \cdot f \cdot r \cdot R_L}$	Formula di dimensionamento della CAPACITA' del filtro, per un alimentatore a onda intera, in funzione della resistenza di carico e del fattore di ripple "r"
--	---

- **Dimensionamento della CAPACITA' del filtro (secondo modo)**

In maniera alternativa, un valore della capacità del filtro (tale da assicurare un fattore di ripple minore del 10%) può essere assegnato, senza effettuare calcoli, ma tenendo conto dei seguenti risultati sperimentali:

$C=(2\div4)$ mF per ogni ampère di I_{Om} , nel caso in cui $V_{Om}=(5\div10)$ V

$C=(1\div2)$ mF per ogni ampère di I_{Om} , nel caso in cui $V_{Om} > 10$ V

- **Dimensionamento del TRASFORMATORE attraverso la tensione (V_{i2}) al suo avvolgimento secondario**

Dimensionare il trasformatore significa determinare il suo rapporto-spire N_2/ N_1 , che è dato da:

$$\frac{N_2}{N_1} = \frac{V_{i2EFF}}{V_{reteEFF}} = \frac{V_{i2EFF}}{230V}$$

E' quindi necessario determinare il valore efficace V_{i2EFF} della tensione sul secondario:

$$V_{i2EFF} = \frac{V_{Om} + \frac{\Delta V_r}{2} + k \cdot V_D}{0,92 \cdot \sqrt{2}}$$

dove:

V_{i2EFF}	Tensione al secondario del trasformatore (nel caso di trasformatore a presa centrale si riferisce a un semiavvolgimento)
V_{Om}	Tensione continua di uscita (valore medio)
ΔV_r	Valore picco-picco della tensione di ripple
k	Numero di diodi contemporaneamente in conduzione: k=2 per alimentatore a ponte k=1 in tutti gli altri casi
V_D	Caduta di tensione sul singolo diodo
0,92	Coefficiente pratico che tiene conto delle perdite nel trasformatore

Per il valore efficace I_{I2EFF} della corrente sul secondario si possono usare i valori pratici:

$I_{I2EFF} =$	$1,8 \cdot I_{Om}$ per alimentatore a ponte
	$1,2 \cdot I_{Om}$ per alimentatore a presa centrale
	$2,2 \cdot I_{Om}$ per alimentatore a semionda

- **OSSERVAZIONI sul dimensionamento**

Quanto più la tensione di uscita è bassa, tanto più piccolo deve essere il ripple, il quale pesa maggiormente su una tensione bassa.

Per avere un ripple piccolo, la costante di tempo $\tau = RC$ del filtro deve essere grande, in modo che il transitorio di scarica sia lento e la tensione di uscita scenda poco in profondità (ricordiamo, a proposito, che la durata di un transitorio è circa $T^* = 5\tau = 5RC$).

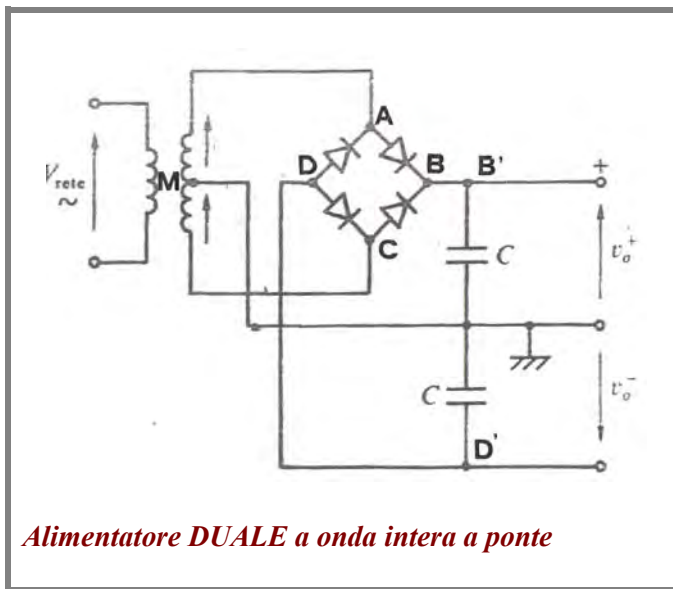
Dunque, **se la tensione di uscita è bassa la capacità del condensatore deve essere grande.**

Conseguentemente, al crescere della tensione di uscita, la capacità del filtro deve risultare più bassa.

Si può infine osservare che, al crescere della capacità, il ripple diminuisce.

ALIMENTATORI DUALI

L'alimentatore duale è un alimentatore in grado di fornire una **coppia** di *tensioni continue simmetriche rispetto alla massa* (per esempio $V^+ = +12V$ e $V^- = -12V$) cioè due tensioni di **uguale ampiezza** e di **polarità opposta**.



ESEMPIO

Supponiamo che:

- $V_{rete}=230V$
- Tensione totale sull'avvolgimento secondario: 24V
- Tensione su ciascun semiavvolgimento secondario: 12V

Trascurando le cadute di tensione sui diodi si ha:

- $V_{B'} = V_B = +12V$
- $V_M = 0V$
- $V_{D'} = V_D = -12V$
- $V_0^+ = V_{B'M} = V_{B'} - V_M = +12 - 0 = +12V$
- $V_0^- = V_{D'M} = V_{D'} - V_M = -12 - 0 = -12V$
- **I potenziali** in corrispondenza delle “**punte**” delle frecce sono **+12V e -12V**.

ALIMENTATORE DUALE A PONTE

L'alimentatore duale a ponte può essere realizzato, a partire dall'alimentatore singolo a ponte, nel modo seguente:

- si collega a massa il punto centrale M dell'avvolgimento secondario del trasformatore
- fra la massa e il punto B della diagonale BD del ponte (diagonale di “prelievo”) si inserisce il primo condensatore di livellamento (in parallelo al quale si può applicare un primo carico su cui ricadrà una tensione “ V_0^+ ”, positiva rispetto a massa)
- fra la massa e l'altro punto (D) della diagonale BD del ponte (diagonale di “prelievo”) si inserisce un secondo condensatore di livellamento (in parallelo al quale si può collegare un secondo carico sul quale si avrà una tensione “ V_0^- ”, negativa rispetto a massa e con ampiezza uguale a quella di V_0^+).

ALIMENTATORI STABILIZZATI LINEARI

REGOLATORI O STABILIZZATORI DI TENSIONE

L'alimentatore formato da trasformatore, raddrizzatore e filtro è detto "alimentatore non stabilizzato" e fornisce, all'uscita del filtro, una tensione che è unipolare, ma non costante, una tensione, cioè, che varia al variare della tensione di ingresso e al variare del carico. Questo inconveniente non è tollerabile per la maggior parte dei circuiti elettronici.

Si ricorre allora allo **stabilizzatore o regolatore di tensione**, che viene inserito **a valle del filtro di livellamento**.

Uno stabilizzatore di tensione può essere definito come un circuito che permette di convertire una tensione unipolare lentamente variabile nel tempo in una tensione praticamente costante.

L'alimentatore dotato di tale dispositivo è detto "alimentatore stabilizzato".

Tratteremo per ora i regolatori o stabilizzatori **lineari**, così detti perché **il transistor**, che costituisce l'elemento di controllo, **lavora in zona lineare** (non in saturazione né in interdizione).

COME FUNZIONA UN REGOLATORE O STABILIZZATORE DI TENSIONE

Un regolatore di tensione funziona in questo modo: un **amplificatore di errore** confronta una porzione della **tensione di uscita V_O** dell'alimentatore con una **tensione di riferimento V_Z** , precisa e stabile, e **produce un "segnale di errore"**. Il segnale di errore **comanda un elemento di controllo** che ha il compito di **contrastare le eventuali variazioni della tensione di uscita**, dovute a variazioni della tensione di ingresso oppure al carico.

L'amplificatore di errore può essere realizzato per esempio con un **amplificatore differenziale** basato su **operazionale**.

L'elemento di controllo può essere un transistor BJT o MOS.

TIPOLOGIE DI REGOLATORI DI TENSIONE

I regolatori o stabilizzatori lineari possono essere:

- di tipo “*serie*”, quando l’**elemento di controllo**, cioè il **BJT**, è posto **in serie al carico** e **si oppone** alle variazioni della tensione di uscita, **facendo variare la tensione ai suoi capi**; questa tipologia è di gran lunga la **più utilizzata**
- di tipo “*parallelo* o “*shunt*”, quando l’**elemento di controllo**, cioè il BJT, è posto direttamente **in parallelo al carico** e **si oppone** alle variazioni della tensione di uscita, **mantenendo la tensione ai suoi capi sostanzialmente costante** e variando la corrente assorbita.

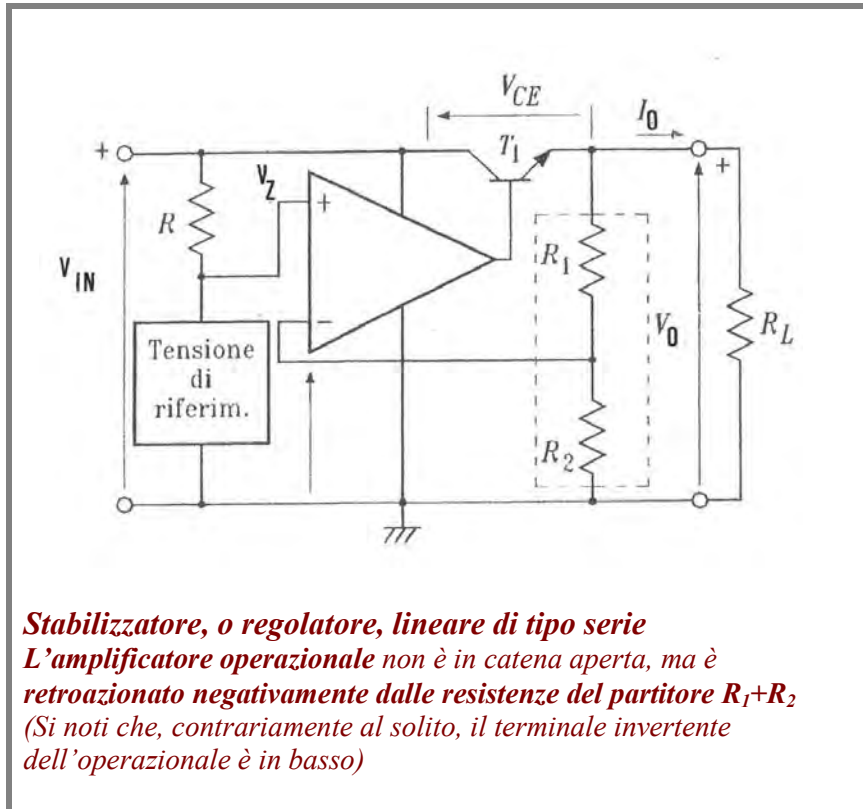
IL BJT COME ELEMENTO DI CONTROLLO (PASS ELEMENT) DEL REGOLATORE

L’elemento di controllo dei regolatori, e cioè il BJT, agisce sfruttando la **possibilità di modificare la resistenza fra emettitore e collettore** agendo sulla corrente di base I_B o sulla tensione base-emettitore V_{BE} .

Se nella base del transistor **si inietta una forte corrente** (oppure se la V_{BE} è sufficientemente elevata) il **BJT** si porta in uno stato di **forte conduzione**, nel senso che il suo punto di funzionamento sarà caratterizzato da **elevata corrente** e bassa tensione.

Se invece la *corrente di base* del transistor è *debole* (oppure se la V_{BE} è bassa) il BJT si porta in uno stato di bassa conduzione, nel senso che il suo punto di funzionamento sarà caratterizzato da corrente di piccola ampiezza e tensione elevata.

REGOLATORE O STABILIZZATORE DI TIPO “SERIE”



$$V_o = V_z \cdot \left(1 + \frac{R_1}{R_2}\right)$$

Tensione di uscita del regolatore serie:
essa dipende solo dalla tensione di riferimento V_z e dalle resistenze del partitore, e quindi risulta costante

Si dimostra che, grazie all'azione del regolatore, la tensione V_o sul carico viene a dipendere dalla sola tensione di riferimento V_z (e non da V_{IN} e da R_L), per cui la tensione sul carico è insensibile alle variazioni della tensione ai capi del filtro dell'alimentatore e ad eventuali variazioni del carico stesso.

ANALISI DEL REGOLATORE-SERIE

Nell'ipotesi che l'amplificazione di tensione possa considerarsi infinita, per l'operazionale vale il corto circuito virtuale, per cui: $V^+ = V^-$.

Determiniamo l'espressione di V^- mediante la regola del partitore di tensione:

$$V^- = V_o \cdot \frac{R_2}{R_2 + R_1}$$

Essendo poi $V^- = V^+ = V_Z$ (come si vede dallo schema circuitale), si ha:

$$V_Z = V_o \cdot \frac{R_2}{R_2 + R_1}$$

da cui possiamo ricavare la tensione di uscita:

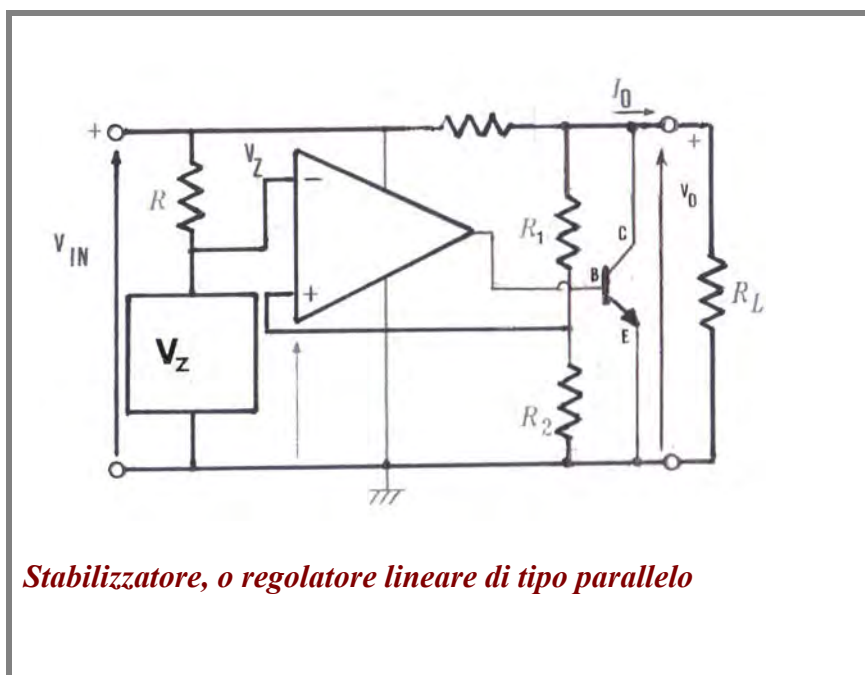
$$V_o = \frac{V_Z}{\frac{R_2}{R_2 + R_1}} \rightarrow V_o = V_Z \cdot \frac{R_2 + R_1}{R_2}$$

$V_o = V_Z \cdot \left(1 + \frac{R_1}{R_2}\right)$	Tensione di uscita del regolatore serie: essa dipende solo dalla tensione di riferimento V_Z e dalle resistenze del partitore, e quindi è resa costante dal regolatore
--	--

Osserviamo infine che, nello schema del regolatore serie, il BJT realizza un inseguitore di tensione (emitter follower) perché è collegato a collettore comune. Infatti:

- nel regolatore la tensione V_{IN} è applicata al collettore del BJT e ricopre un ruolo analogo a quello della V_{CC} dello schema generale dell'inseguitore
- il carico R_L del regolatore è in serie all'emettitore
- la tensione di base del transistor rappresenta l'ingresso della sezione a collettore comune del regolatore.

REGOLATORE O STABILIZZATORE DI TIPO “PARALLELO” o “SHUNT”

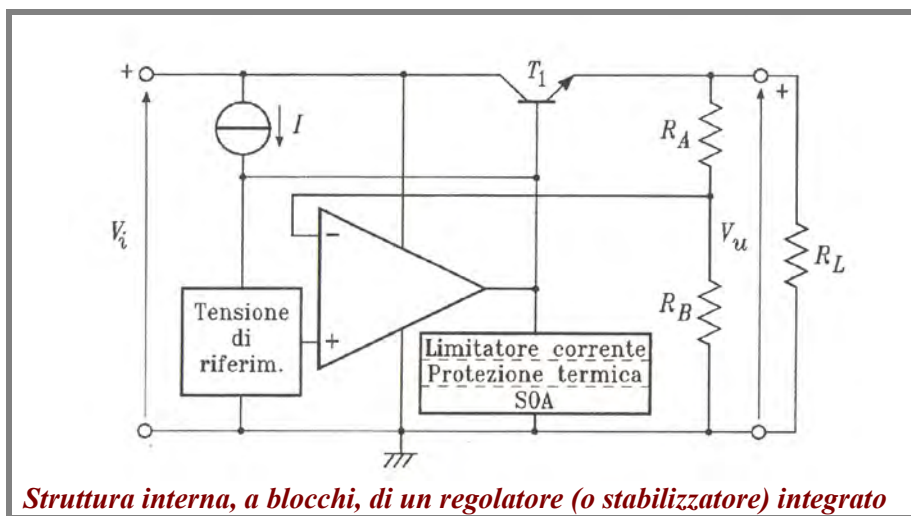


Stabilizzatore, o regolatore lineare di tipo parallelo

REGOLATORI INTEGRATI LINEARI A TRE TERMINALI

Sono disponibili in commercio regolatori di tensione in forma integrata, cioè con tutti i componenti realizzati su un unico chip di silicio.

Gli stabilizzatori integrati di tipo serie funzionano secondo gli stessi principi su cui si basano gli analoghi dispositivi discreti.



La loro struttura semplificata (come si vede dalla figura) prevede, oltre al circuito di stabilizzazione, alcuni circuiti di protezione, totalmente o parzialmente interni.

Le protezioni consistono in:

- limitazione della corrente di uscita
- protezione termica che interdice i componenti di controllo (uno o più BJT, eventualmente in connessione Darlington) quando la temperatura del chip supera un livello di guardia
- protezione SOA (Safe Operating Area, area di funzionamento sicuro) impedisce al punto di funzionamento dell'elemento-serie di uscire dalla zona di funzionamento sicuro, riducendo automaticamente la corrente che attraversa il dispositivo.

Il valore della tensione d uscita è determinato da due resistenze che in alcuni integrati sono interne (regolatori della serie 78XX e 79XX) e in altri sono esterne (regolatori con uscita variabile).

REGOLATORI INTEGRATI 78XX (PER TENSIONI CONTINUE POSITIVE) E 79XX (PER TENSIONI CONTINUE NEGATIVE)

La 78XX e la 79XX sono due serie molto diffuse di stabilizzatori integrati a tre terminali. Le cifre “XX” indicano la **tensione nominale dell’uscita stabilizzata**, la quale può assumere valori fissi compresi fra 5V e 24V.

Per esempio il regolatore per +5V nominali è indicato con la sigla 7805.

La struttura interna di questi stabilizzatori rispecchia lo schema a blocchi generale dei regolatori-serie illustrata in precedenza.

Il costruttore ottiene i singoli valori fissi della tensione di uscita scegliendo i valori di una coppia di resistenze interne.

Oltre alla tensione nominale di uscita e alla corrente massima di uscita, sono particolarmente importanti altri due parametri:

- il “**dropout**”, cioè la differenza minima che deve esserci fra la tensione di ingresso (più elevata) e quella di uscita (più bassa), per ottenere un funzionamento corretto
- la **corrente di riposo o di polarizzazione I_Q** , cioè la corrente assorbita dall’integrato, la quale fuoriesce dal terminale GND e non è erogata al carico.

Nella sigla possono anche essere presenti delle lettere. Esse indicano il valore della corrente di uscita, secondo la tabella seguente:

<i>lettera</i>	I_u [A]	<i>osservazioni</i>
<i>L</i>	0,1	
<i>C</i>	0,5	con load regul. $SV_L = 0,5\%$
<i>M</i>	0,5	con load regul. $SV_L = 1\%$
<i>-</i>	1	
<i>S</i>	2	
<i>CB</i>	2	con $V_u = 13,8$ V
<i>T</i>	3	
<i>H</i>	5	
<i>P</i>	10	

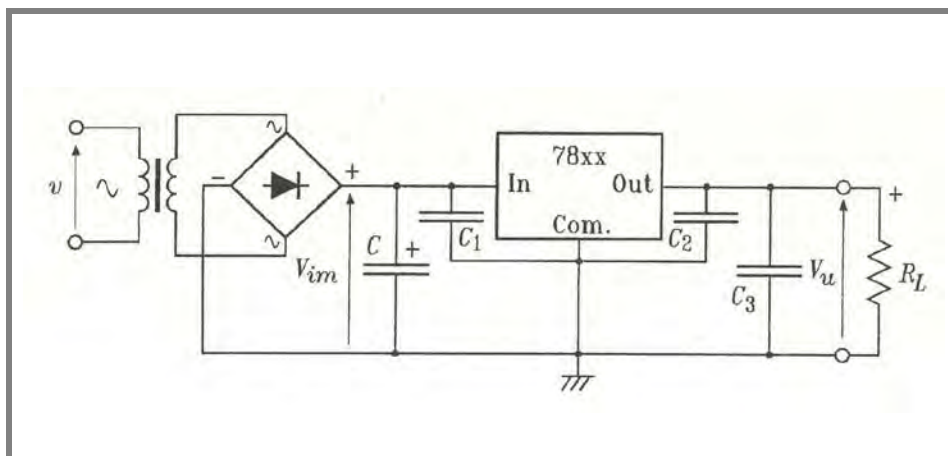
Componente	funzione e polarità	V_i (V)	V_o (V)	I_o max. (A)	I_{sc} max. (A)	Line reg. %	Load reg. %	I_o (mA)	R_r (dB) Ripple Rejection	k_{f_c} ($^{\circ}\text{C}/\text{W}$)	k_{f_d} ($^{\circ}\text{C}/\text{W}$)	Pack.	Drop-out Voltage (V)
$\mu\text{A} 7805$	Fisso pos.	$7 \div 35$	$4,8 \div 5,2$	1	1,2	1	1	8	62	5	65	TO-220	2,5
$\mu\text{A} 7806$	Fisso pos.	$8 \div 35$	$5,75 \div 6,25$	1	1,2	1	1	8	59	5,5	45	TO-3	
$\mu\text{A} 7808$	Fisso pos.	$10 \div 35$	$7,7 \div 8,3$	1	1,2	1	1	8	56	5	65	TO-220	2,5
$\mu\text{A} 7812$	Fisso pos.	$14 \div 35$	$11,5 \div 12,5$	1	1,2	1	1	8	55	5,5	45	TO-3	
$\mu\text{A} 7815$	Fisso pos.	$17 \div 35$	$14,4 \div 15,6$	1	1,2	1	1	8	54	5	65	TO-220	2,5
$\mu\text{A} 7818$	Fisso pos.	$20 \div 35$	$17,3 \div 18,7$	1	1,2	1	1	8	53	5,5	45	TO-3	
$\mu\text{A} 7824$	Fisso pos.	$26 \div 40$	$23 \div 25$	1	1,2	1	1	8	50	5	65	TO-220	2,5
$\mu\text{A} 78H05$	Fisso pos.	$8,5 \div 20$	$4,8 \div 5,2$	5	7	1	2	10	60	5,5	45	TO-3	
$\mu\text{A} 78G$	Reg. pos.	$7,5 \div 40$	$5 \div 35$	1	1,2	0,75	1	5	62	2,5	38	TO-3	3,5
										8	80	Pow. W	3
										6	47	TO-3	

Parametri di alcuni integrati della serie 78XX

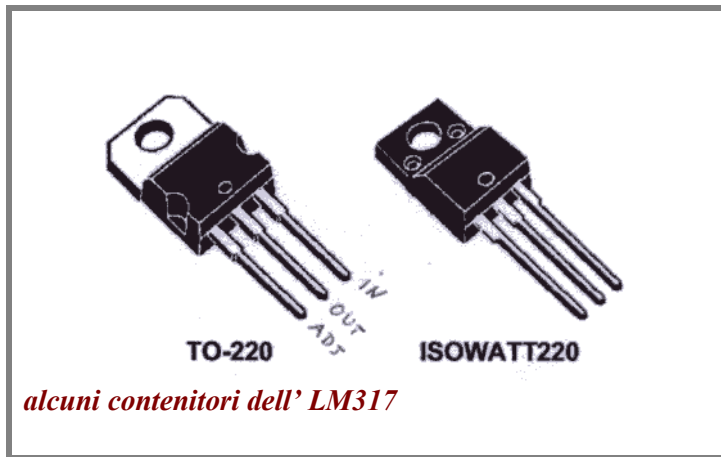
Nella figura seguente illustriamo un possibile schema realizzativo di alimentatore stabilizzato con regolatore integrato a tre pin.

Nella figura sono stati indicati dei condensatori che non sono strettamente indispensabili al funzionamento: la loro funzione è di evitare l'eventuale innesco di oscillazioni.

I valori di capacità da utilizzare sono generalmente indicati dal costruttore nel foglio-dati dello stabilizzatore.



REGOLATORI INTEGRATI AD AMPIEZZA VARIABILE LM 317



COSA PERMETTE DI FARE L' LM317

E' un regolatore di **tensione continua** di **ampiezza variabile**.

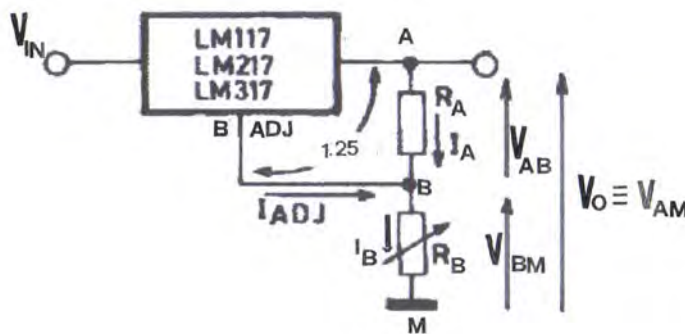
La sua configurazione più semplice richiede di collegare il componente a **due sole** resistenze esterne.

L'LM317 permette di:

1. Ottenere una **tensione rigorosamente costante** fra il terminale OUT di uscita (terminale A) e il terminale ADJ (terminale B) pari a **1,25V circa**. Questa tensione è indipendente dall'eventuale carico applicato fra OUT e ADJ (cioè fra A e B). E' necessario inserire comunque una resistenza fra OUT e ADJ.
2. **Ottenere una corrente I_A costante sulla resistenza R_A** collegata fra OUT e ADJ. Il valore di questa corrente è:

$$I_A = \frac{1,25}{R_A}$$

3. **Ottenere una tensione V_{AM} del valore desiderato**, rigorosamente costante e indipendente dal carico, fra il terminale OUT di uscita (terminale A) e la massa, ossia ai capi della "colonna" di resistenze R_A e R_B .



*Realizzazione "di base"
del regolatore variabile di tensione*

$$V_O = V_{AB} + V_{BM} = 1,25 + I_A \cdot R_B$$

$$V_O = 1,25 + \frac{1,25}{R_A} \cdot R_B$$

$$V_O = 1,25 \cdot \left(1 + \frac{R_B}{R_A}\right) \quad [V]$$

La tensione V_{AM} , cioè la V_O , è indipendente dal carico in quanto sarà il regolatore (grazie alle sue caratteristiche funzionali) a fornire l'eventuale sovrappiù di corrente che dovesse essere richiesto dal carico, mentre I_A e I_B rimarranno costanti.

CAMPO DI VALORI DI R_B

Una volta fissate la tensione V_{IN} di ingresso e la resistenza R_A fra OUT (punto A) e ADJ, il regolatore riesce a mantenere una **tensione di riferimento costante di 1,25V** (fra A e ADJ) e una **tensione di uscita V_O costante e indipendente dal carico**, ma ciò accade solo se la **R_B è minore di un certo valore**.

Per esempio:

$$V_{IN}=12V \rightarrow R_B \leq 800\Omega \div 900\Omega$$

$$V_{IN}=15V \rightarrow R_B \leq 1200\Omega$$

La **massima R_B possibile** è tanto **più grande** quanto **più elevata** è la **tensione di ingresso**.

VALORE DI I_A

Per poter imporre in maniera semplice il valore della tensione desiderata fra OUT e massa bisogna fare in modo che su R_B scorra la stessa corrente che passa su R_A .

Deve cioè essere: $I_B = I_A$, il che si verifica se I_{ADJ} risulta trascurabile rispetto a I_A .

Per rendere I_{ADJ} trascurabile rispetto a I_A , dobbiamo fare in modo che la I_A sia molto più grande, per esempio di 50 volte, rispetto alla I_{ADJ} il cui valore massimo è $I_{ADJ} = 0,1 \text{ mA}$.

Possiamo perciò imporre:

$$I_A = 50 \cdot I_{ADJ}$$

$$I_A = 50 \cdot 0,1 \text{ mA}$$

$$I_A = 5 \text{ mA}$$

Con I_A compreso fra **5 mA** e **10 mA** è certamente possibile trascurare I_{ADJ} .

IMPOSIZIONE DEL VALORE DI TENSIONE DI USCITA V_{AM} DESIDERATO FRA IL TERMINALE OUT DI USCITA E LA MASSA E DIMENSIONAMENTO DELLE RESISTENZE R_A E R_B

$$R_A = \frac{1,25V}{I_A} = \frac{1,25V}{5 \cdot 10^{-3}} = \frac{1250V}{5} = 250\Omega$$

I valori consigliati dal costruttore sono compresi fra 120Ω e 240Ω .

Supponiamo che il valore di tensione desiderato sia $V_{AM} = 9V$

Questo significa che su R_B dovrà esservi la tensione:

$$V_{BM} = V_{AM} - 1,25V$$

$$V_{BM} = 9V - 1,25V$$

$$V_{BM} = 7,75V$$

Di conseguenza dovrà essere:

$$R_B = \frac{V_{BM}}{I_B} = \frac{V_{BM}}{I_A} = \frac{7,75}{5 \cdot 10^{-3}} = \frac{7750}{5} = 1550\Omega$$

LM 317: ALCUNI DATI

$$1,25V \leq V_{AM} = V_U \leq 27V$$

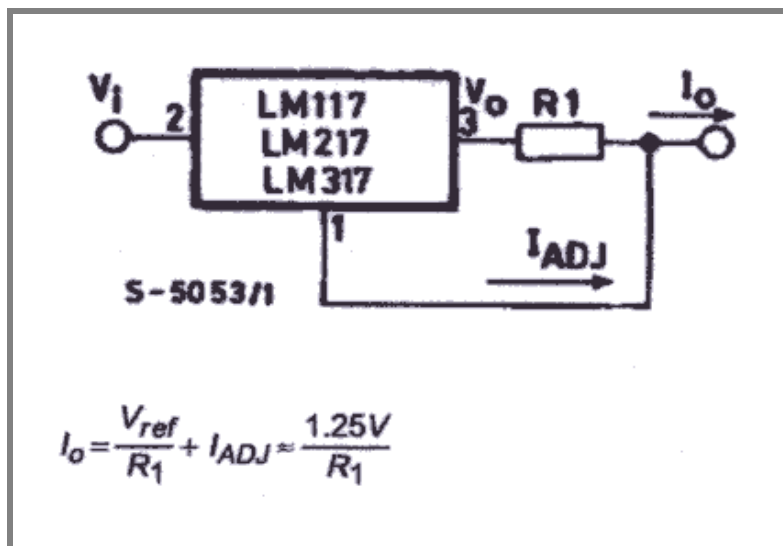
15 W = massima potenza dissipabile con contenitore TO-220

range di temperature dell'LM 337: $(0 \div +125)^{\circ}C$

range di temperature dell'LM 317: $(-55 \div +150)^{\circ}C$

Per avere una determinata tensione di uscita $V_U = V_{AM}$ (fra OUT e massa) è necessario applicare una tensione di ingresso maggiore di almeno 3V rispetto all'uscita (tensione di dropout: vedi "reference voltage" nel data sheet)

REGOLATORE DI CORRENTE



PARAMETRI DEI REGOLATORI DI TENSIONE

LINE REGULATION: REGOLAZIONE DI LINEA (O DI INGRESSO)

Esprime la **stabilità della tensione di uscita** nei confronti della **variazione della tensione di ingresso**, facendo riferimento alla tensione nominale di uscita:

$$\text{Line-reg}_{\%} = \frac{\frac{\text{variazione tensione di uscita}}{\text{variazione tensione di ingresso}}}{\text{valore nominale tensione uscita}} \cdot 100$$

$$\text{Line-reg}_{\%} = \frac{\frac{\Delta V_o}{\Delta V_{IN}}}{V_o} \cdot 100$$

dove:

$$\frac{\Delta V_o}{\Delta V_{IN}} = \text{variaz. della tens. di uscita per un volt di variaz. della tens. di ingresso}$$

Se, per esempio, la variazione della tensione di uscita di un regolatore (con V_o nominale di 5V) è di 1 mV in corrispondenza della variazione di 1V della tensione di ingresso, la sua regolazione percentuale di linea si determina in questo modo:

$$\text{Line-reg}_{\%} = \frac{\frac{10^{-3}}{1}}{5} \cdot 100 = \frac{10^{-3}}{5} \cdot 100 = \frac{10^{-1}}{5} = \frac{0,1}{5} = 0,02 \%$$

LOAD REGULATION: REGOLAZIONE DI CARICO (O DI USCITA)

Esprime la **stabilità della tensione di uscita** nei confronti della **variazione della corrente I_O sul carico**, facendo sempre riferimento alla tensione nominale di uscita:

$$\text{Load-reg}_{\%} = \frac{\frac{\text{variazione tensione di uscita}}{\text{variazione corrente sul carico}}}{\text{valore nominale tensione uscita}} \cdot 100$$

$$\text{Load-reg}_{\%} = \frac{\frac{\Delta V_o}{\Delta I_o}}{V_o} \cdot 100$$

dove:

$$\frac{\Delta V_o}{\Delta I_o} = \text{variaz. della tens. di uscita per un ampère di variaz. della corrente sul carico}$$

Se per esempio la variazione della tensione di uscita di un regolatore (con V_O nominale di 5V) è di 25 mV in corrispondenza della variazione di 1A della corrente sul carico, la sua regolazione percentuale di carico si determina in questo modo:

$$\text{Line-reg}_{\%} = \frac{\frac{25 \cdot 10^{-3}}{1}}{5} \cdot 100 = \frac{25 \cdot 10^{-3}}{5} \cdot 100 = 5 \cdot 10^{-1} = 5 \cdot 0,1 = 0,5 \%$$

Un'altra possibile definizione è:

$$\text{Load-reg}_{\%} = \frac{\text{tensione a pieno carico} - \text{tensione a carico minimo}}{\text{tensione nominale}} \cdot 100$$

ossia:

$$\frac{\text{tens. con carico che assorbe la max corrente} - \text{tens. con carico che assorbe la min corr.}}{\text{tensione con carico nominale}} \cdot 100$$

che, in formula, è:

$$\text{Load-reg}_{\%} = \frac{(V_{O-full} - V_{O-min})}{V_{O-nom}} \cdot 100$$

E' possibile trovare, nei data-sheet e nei manuali, un'ulteriore definizione che differisce dalla precedente soltanto per la presenza, a denominatore, della tensione a pieno carico al posto della tensione nominale:

$$\text{Load-reg}_{\%} = \frac{(V_{O-full} - V_{O-min})}{V_{O-full}} \cdot 100$$

REIEZIONE DEL RIPPLE

Esprime la stabilità della tensione di uscita rispetto al ripple e può essere definita come una "regolazione di linea applicata al ripple" (l'ingresso del regolatore è una tensione unipolare affetta da ripple):

$$(R_r)_{dB} = 20 \cdot \log_{10} \left(\frac{\Delta V_{IN}}{\Delta V_O} \right)$$

ALIMENTATORI STABILIZZATI SWITCHING o A COMMUTAZIONE (o SMPS, Switching Mode Power Supply)

Gli alimentatori stabilizzati lineari sono affidabili e circuitalmente semplici, però presentano alcuni inconvenienti:

- sul componente di controllo-serie del regolatore (tipicamente un BJT) si viene a dissipare troppa potenza, il che abbassa il rendimento dell'alimentatore e obbliga a ricorrere a dissipatori di calore voluminosi
- il trasformatore di ingresso dell'alimentatore deve avere dimensioni elevate perché lavora a una frequenza bassa come quella di rete.
Perché i trasformatori per bassa frequenza sono più voluminosi?
L'impedenza induttiva è

$$\bar{Z}_L = j \cdot \omega \cdot L$$

(e cresce al crescere della frequenza). Quindi, per ottenere un fissato valore di Z_L , è necessario un valore di induttanza L tanto più elevato quanto più bassa è la frequenza. Siccome l'entità di L aumenta all'aumentare delle dimensioni (diametro) e del numero di spire dell'induttore, questo dovrà essere più ingombrante se la frequenza è più bassa.

Viceversa, a frequenze alte, gli induttori possono avere dimensioni minori (pochi cm^3 se la frequenza è dell'ordine di qualche KHz).

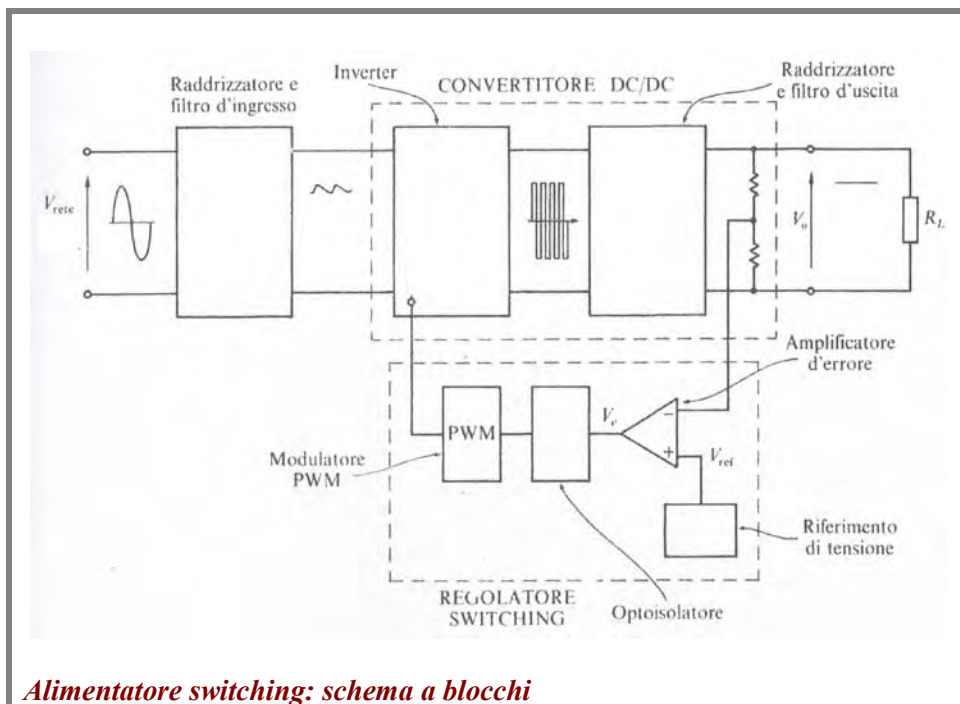
Questi inconvenienti limitano l'impiego degli **alimentatori lineari** al campo delle **basse potenze**, e, in particolare, alle potenze **minori di 50W** (potenze alle quali il costo di un alimentatore lineare è normalmente minore di quello di uno switching).

Per **potenze più elevate** (orientativamente **da 50W in su**) si usano gli **alimentatori switching**, presenti nei computer e nei televisori e diffusissimi nelle applicazioni industriali. Sono anche in commercio piccoli switching da usare come alimentatori esterni per dispositivi vari.

Una particolarità di questi alimentatori è che la **tensione unipolare di ingresso** passa attraverso una **conversione in onda rettangolare** (per poi essere riconvertita in tensione continua). Il dispositivo che effettua questa conversione è detto **inverter**.

Gli alimentatori switching sono caratterizzati dal fatto che **la stabilizzazione della tensione di uscita su un determinato livello avviene nel corso di un processo di conversione della tensione raddrizzata unipolare (prodotta nel primo blocco) in un'onda rettangolare** a frequenza fissa e duty-cycle variabile, processo effettuato da un **inverter** e regolato da un

sistema di controllo in catena chiusa basato sulla **modulazione PWM** (Pulse Width Modulation, modulazione di durata egli impulsi).
 Successivamente l'onda rettangolare viene riconvertita in tensione continua.



Alimentatore switching: schema a blocchi

PREGI DEGLI ALIMENTATORI SWITCHING

Gli alimentatori a commutazione consentono di ovviare agli inconvenienti di quelli lineari. Negli switching infatti, gli elementi di controllo (BJT, MOSFET, SCR dell'inverter) lavorano in commutazione, cioè passano, continuamente e velocemente, dall'interdizione (stato OFF) alla saturazione (stato ON) e viceversa, il che comporta una **dissipazione di potenza molto minore** da parte degli elementi di controllo e un conseguente **aumento del rendimento**, che può **superare l'80%**. Invece, negli alimentatori lineari, i componenti di controllo lavorano nella zona lineare delle caratteristiche (funzionamento in classe A), con maggior consumo.

Inoltre gli alimentatori switching funzionano a frequenze piuttosto elevate (20÷200) KHz e quindi consentono di impiegare **componenti induttivi meno ingombranti**.

INCONVENIENTI DEGLI ALIMENTATORI SWITCHING

In generale, il **ripple** degli alimentatori switching è **maggiore** di quello degli alimentatori lineari.

Perciò, nel caso in cui siano richiesti valori di ripple molto contenuti, anche nelle alimentazioni a commutazione è opportuno ricorrere a regolatori di tipo lineare.

Inoltre gli SMPS sono caratterizzati da **maggiore complessità circuitale** e da una notevole **emissione di disturbi elettromagnetici** (EMI, Electro Magnetic Interference).

SWITCHING SENZA TRASFORMATORE E CON TRASFORMATORE

Gli alimentatori switching possono essere “con” o “senza” *trasformatore*. Quest'ultimo, se presente, è ***contenuto nell'inverter del convertitore DC-DC***.

PREMESSA AL FUNZIONAMENTO DELL'ALIMENTATORE SWITCHING: LA MODULAZIONE PWM

Per comprendere il funzionamento degli alimentatori switching è indispensabile conoscere il principio di funzionamento della **modulazione PWM**.

La PWM è una tecnica che **permette di variare la durata**, ossia la “**larghezza temporale**” degli impulsi alti un’onda rettangolare con frequenza fissa, in funzione dell’ampiezza di un altro segnale.

Variare la durata degli impulsi alti di un’onda rettangolare a frequenza fissa significa anche variare gli intervalli di tempo nei quali l’impulso è basso e dunque significa modificare il duty-cycle ($\delta\%$) dell’onda.

Pertanto possiamo anche dire che la PWM è una tecnica che **permette di variare il duty-cycle** di un’onda rettangolare con frequenza fissa, **in funzione** dell’ampiezza di un **altro segnale**.

L’onda rettangolare, della quale si fa variare il $\delta\%$, è chiamata “**segnale modulato PWM**”.

Il segnale in funzione del quale viene fatto variare il $\delta\%$ dell’onda rettangolare è detto “**segnale modulante**” e lo indichiamo con “ v_e ”.

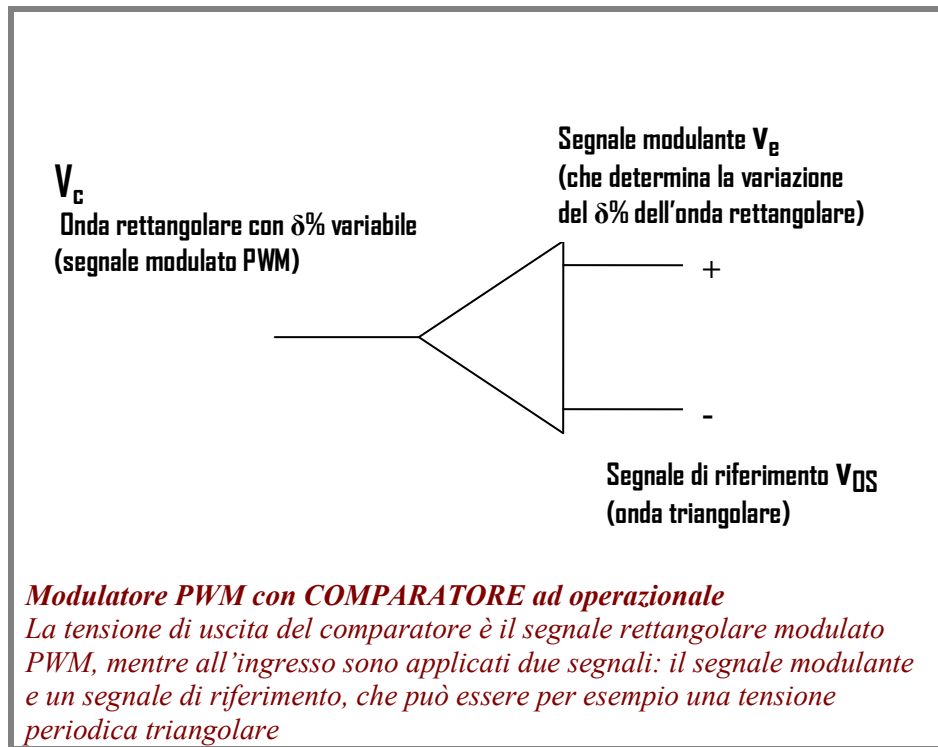
Il segnale modulante “ v_e ” potrebbe essere, per esempio, la tensione ai capi di un potenziometro o la **tensione di uscita di un amplificatore o di un comparatore**.

Nel processo PWM quindi **le variazioni di ampiezza di un segnale modulante v_e** producono le variazioni di duty-cycle di un’onda rettangolare, per esempio secondo questa legge:

- al **crescere** dell’**ampiezza** del **segnale modulante v_e** , **cresce** il duty-cycle dell’onda rettangolare
- al **diminuire** dell’ampiezza del segnale modulante v_e , il duty-cycle dell’onda rettangolare **decresce**
- quando l’ampiezza del segnale modulante v_e è costante, si mantiene costante anche il duty-cycle dell’onda rettangolare.

Il dispositivo che realizza la tecnica PWM è chiamato **modulatore PWM**.

Tipicamente negli alimentatori switching è presente un **modulatore PWM basato su un comparatore ad operazionale**.

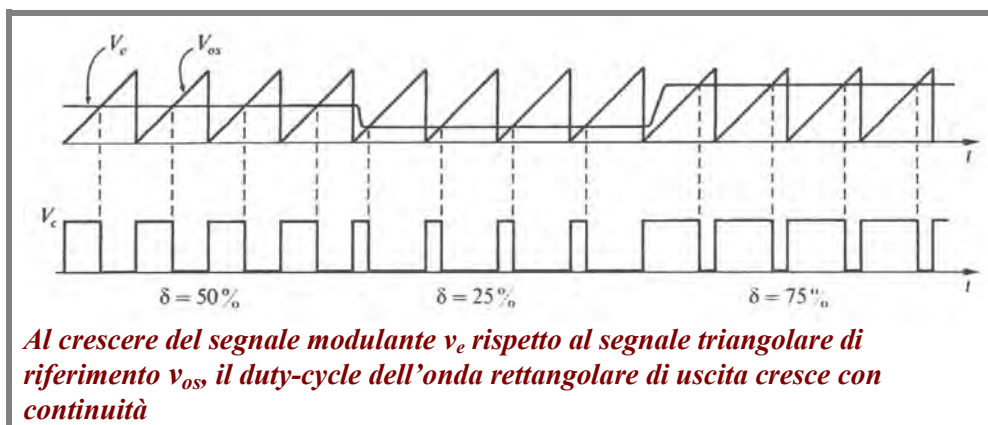


Il modulatore basato sul comparatore prevede l'applicazione **in ingresso**, oltre che del segnale modulante, anche di un **segnale di riferimento v_{OS}** (per esempio **un'onda triangolare**):

dal confronto fra il segnale modulante v_e e v_{OS} , risulta determinato il duty-cycle dell'onda rettangolare di uscita.

Se, come nell'esempio della figura, il segnale modulante è applicato all'ingresso non invertente dell'operazionale, il modulatore funziona in questo modo:

- negli intervalli di tempo nei quali il segnale **modulante v_e** è **maggiore** del segnale triangolare di riferimento v_{OS} , **l'onda rettangolare** di uscita è **alta**
- negli intervalli di tempo nei quali il segnale **modulante v_e** è **minore** del segnale triangolare di riferimento v_{OS} , **l'onda rettangolare** di uscita è **bassa**.



Come si vede dalla figura, **quanto più è alta la modulante v_e rispetto alla tensione di riferimento v_{os} , tanto più larga è la porzione di periodo nella quale l'onda rettangolare di uscita è alta**, e tanto maggiore è il duty-cycle.
 (infatti al crescere di v_e rispetto alla tensione di riferimento v_{os} , si allarga la porzione di periodo nella quale $v_e > v_{os}$ e nella quale l'onda rettangolare di uscita è alta, per cui aumenta il duty-cycle).

SCHEMA A BLOCCHI DELL'ALIMENTATORE SWITCHING

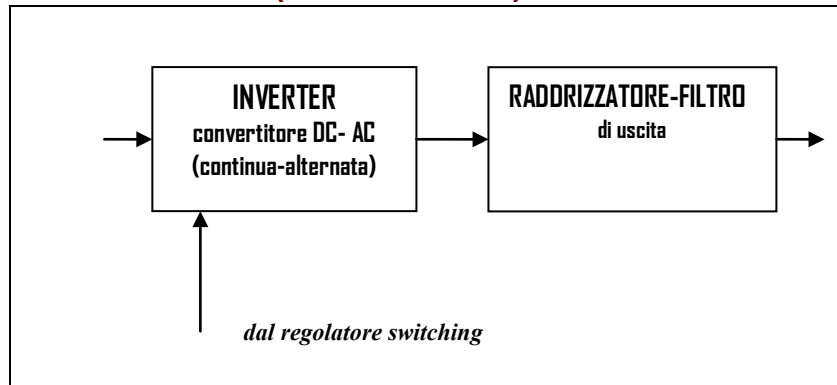
Confrontando lo schema a blocchi di un alimentatore switching con quello dell'alimentatore lineare si nota subito che:

- Nell'alimentatore switching il **primo blocco** è un **raddrizzatore con filtro capacitivo** e non un trasformatore; il trasformatore, se è presente, si trova all'interno dell'inverter o convertitore DC/AC;
- Il **secondo blocco** è il **convertitore DC/DC** o **convertitore “continua- continua”**, cioè un dispositivo che **converte la tensione unipolare applicata al suo ingresso in una tensione continua di ampiezza diversa**; il convertitore DC-DC si compone, a sua volta, di due dispositivi: un inverter e un raddrizzatore-filtro di uscita:
 - l'**inverter** converte la **tensione raddrizzata e filtrata** (presente all'uscita del primo blocco dell'alimentatore) in **un'onda rettangolare** ad alta frequenza **il cui duty-cycle è determinato dal regolatore switching**
 - il **raddrizzatore-filtro di uscita** ha come **ingresso proprio l'onda rettangolare** generata dall'inverter e la trasforma in una **tensione continua di ampiezza prefissata**.

Nel suo complesso quindi **il convertitore DC/DC converte una tensione unipolare (che è quella del primo raddrizzatore, posto all'ingresso dell'alimentatore) in una tensione unipolare costante (tensione continua stabilizzata), che è la V_O di uscita dell'intero alimentatore**; i convertitori DC/DC (e quindi gli alimentatori nei quali sono contenuti) possono essere di due tipi: “**con** trasformatore interno” e “**senza** trasformatore interno”;

ribadiamo infine che è proprio nel corso della conversione continua-continua che la tensione manipolata dall'alimentatore diventa un'onda rettangolare, e che è proprio in questa fase che avviene la stabilizzazione, grazie al blocco indicato come “regolatore switching”;

CONVERTITORE DC-DC (CONTINUA-CONTINUA)



- Il **terzo blocco** dell'alimentatore a commutazione è il “**regolatore switching**”, blocco indicato anche come “**circuiti controllo e modulazione PWM**”; il regolatore determina la **stabilizzazione della tensione di uscita** e, come vedremo fra poco, è essenzialmente formato da due circuiti:

- un amplificatore di errore (che confronta una porzione della tensione di uscita con una tensione di riferimento)
- un modulatore PWM (che modifica il duty-cycle dell'onda rettangolare prodotta dall'inverter, in funzione del segnale di errore)

Il regolatore switching **mantiene costante la tensione di uscita** dell'alimentatore **confrontando** una frazione di tale tensione con una tensione di riferimento e **contrastando le variazioni della V_O** mediante alterazione del duty-cycle dell'onda rettangolare.

La variazione del duty-cycle è effettuata dal **modulatore PWM**, che è comandato dal segnale di errore, cioè dal segnale prodotto dall'amplificatore di errore. L'uscita del modulatore è collegata all'inverter, ove si genera l'onda rettangolare.

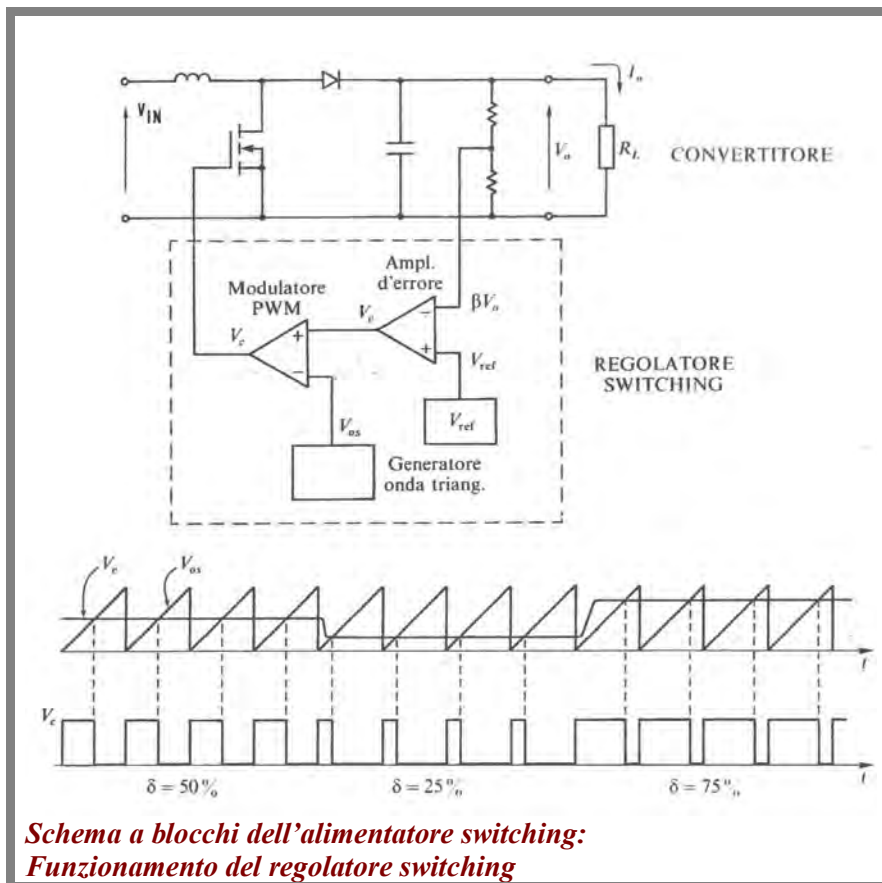
Alterando il duty-cycle ($\delta\%$) dell'onda rettangolare, si altera il valore delle aree sottese dall'onda e quindi se ne modifica il valor medio nel periodo; resta così determinata l'ampiezza della tensione raddrizzata e filtrata, cioè dell'uscita dell'alimentatore.

La stabilizzazione della tensione di uscita avviene in questo modo:

se il valore della tensione di uscita tende ad aumentare, il regolatore fa diminuire il valore del duty-cycle dell'onda rettangolare, determinando così una diminuzione della tensione continua di uscita che è ottenuta dall'onda rettangolare mediante raddrizzamento e filtraggio;

questa diminuzione dell'uscita contrasta l'aumento che tendeva a prodursi, mantenendo stabile l'ampiezza dell'uscita.

Analogamente, se la tensione di uscita subisse una indesiderata diminuzione, i circuiti di controllo e modulazione PWM (cioè i circuiti del regolatore switching) produrrebbero un aumento del $\delta\%$ dell'onda rettangolare (mantenendone costanti il periodo e quindi la frequenza) determinando così un aumento del valor medio dell'onda e dell'ampiezza della tensione continua di uscita dell'alimentatore.



TIPI DI CONVERTITORI DC-DC

Come si è già detto, i convertitori continua-continua possono essere o “con trasformatore interno” o “senza trasformatore interno”

CONVERTITORI DC-DC CON TRASFORMATORE

Sono utilizzati per alimentare **computer, monitor e televisori**.

La presenza del trasformatore interno al convertitore (trasformatore di piccolo ingombro con nucleo di ferrite) consente di ottenere l'**isolamento galvanico del carico dalla rete**. Il fatto che la parte ad alta tensione sia separata dall'uscita a bassa tensione, rende questi convertitori **adatti alle situazioni** nelle quali la **tensione di uscita** sia **molto più bassa o molto più alta** di quella di ingresso, alle situazioni, cioè, caratterizzate da una **grande differenza fra le tensioni di ingresso e di uscita**.

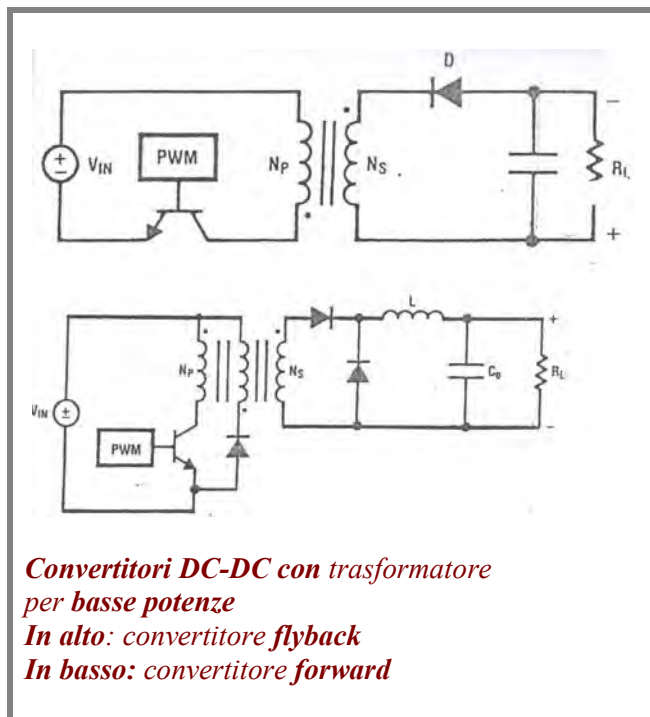
Si osservi che la presenza di un trasformatore nell'inverter rende necessario un isolamento galvanico anche nell'anello di reazione, isolamento che può essere realizzato mediante un ulteriore trasformatore o mediante un optoisolatore.

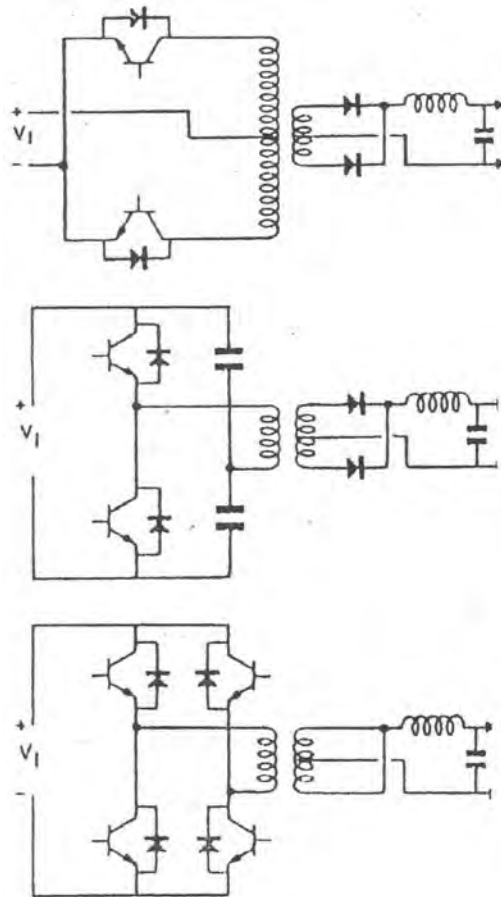
L'optoisolatore o fotoaccoppiatore è il complesso di un diodo LED a raggi infrarossi (IRED) e di un fototransistor (o altro fotorivelatore o fotosensore, come fotodiodo, fotodarlington o fototriac) separati da un dielettrico trasparente che garantisca l'isolamento elettrico.

Il LED a infrarossi, direttamente polarizzato, emette una radiazione che, incidendo sulla giunzione base collettore (inversamente polarizzata) del fototransistor, dà luogo a una corrente.

Rileviamo infine che la presenza del trasformatore nel convertitore DC/DC svincola la tensione di uscita dalla stretta dipendenza dal duty-cycle e la lega al rapporto di trasformazione. Si riportano di seguito le varie tipologie di convertitori “continua-continua” **con trasformatore**:

- convertitore **FLYBACK** (o a **interdizione**), per **basse** potenze
- convertitore **FORWARD** (o a **conduzione**), per **basse** potenze
- convertitore **PUSH-PULL** (o in **controfase**), per potenze **elevate**
- convertitore a **SEMIPONTE**, per potenze **elevate**
- convertitore a **PONTE**, per potenze **elevate**





**Convertitori DC-DC con trasformatore
per alte potenze**

**Dall'alto: convertitore *push-pull*, convertitore a
mezzo ponte, convertitore a ponte**

CONVERTITORI DC-DC SENZA TRASFORMATORE

I convertitori **senza** trasformatore sono **largamente usati** nell'ambito dei dispositivi elettronici a **bassa tensione**, cioè per **convertire** una **tensione di piccola ampiezza** in una **seconda tensione sempre di piccola ampiezza**.

Si possono per esempio attuare:

- la conversione in discesa da 24V a 12V
- la conversione in salita da 3V a 12V
- l'inversione di polarità da +12V a -12V

Essi sono quindi in grado di convertire una bassa tensione continua V_{IN} di ingresso in una bassa tensione V_O di uscita, di ampiezza minore o maggiore di V_{IN} ed eventualmente invertita rispetto a V_{IN} .

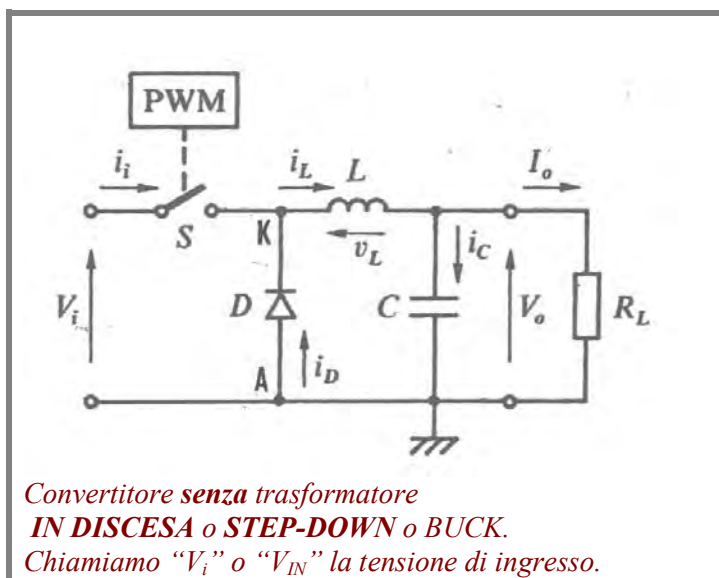
L'assenza del trasformatore interno al convertitore, e quindi la mancanza di isolamento galvanico del carico dalla rete, **non consente** l'utilizzazione di questi convertitori per passare direttamente dai 230V della tensione di rete al valore desiderato di tensione continua.

Si riportano di seguito le varie tipologie di convertitori continua-continua:

- convertitore **IN DISCESA** (o **STEP-DOWN** o BUCK, probabilmente **il più diffuso**)
- convertitore **IN SALITA** (o **STEP-UP** o **BOOST**)
- convertitore **INVERTENTE** (o BUCK-BOOST o STEP-UP-DOWN).

CONVERTITORE DC-DC SENZA TRASFORMATORE DI TIPO “STEP-DOWN” O “IN DISCESA”

E' il convertitore DC-DC **più utilizzato** negli **alimentatori switching** delle **apparecchiature elettroniche di larga diffusione**.



Il dispositivo è essenzialmente formato da:

- un interruttore elettronico (transistor BJT o MOS) che è alla base dell’inverter e che, essendo pilotato dal modulatore PWM, lavora a frequenza fissa e duty-cycle variabile
- un filtro passabasso LC
- un diodo di libera circolazione che permette lo scorrimento della corrente anche quando l’interruttore elettronico è aperto (OFF)

$$V_o = V_{IN} \cdot \frac{T_{ON}}{T_{ON} + T_{OFF}}$$

dove:

T_{ON} = parte del periodo T nella quale l’interruttore è chiuso (ON)

T_{OFF} = parte del periodo T nella quale l’interruttore è aperto (OFF)

Essendo $T_{ON} < T_{ON} + T_{OFF}$, si ha $V_o < V_{IN}$ (la tensione di uscita è minore di quella di ingresso)

**RELAZIONE
I/O**

ANALISI DEL CONVERTITORE SENZA TRASFORMATORE “STEP-DOWN” O “IN DISCESA”

PREMESSE

1. Il punto “A”, anodo del diodo, è a massa, per cui: $V_A=0$
2. Dallo schema circuitale si vede che la tensione sull'induttore è esprimibile come:
 $V_L = V_K - V_O$
3. La tensione V_{IN} (o V_i) è l'ingresso dell'inverter e del convertitore DC/DC e coincide con l'uscita del raddrizzatore-filtro di ingresso dell'alimentatore. La V_{IN} e la V_O sono positive e possono ritenersi sostanzialmente costanti: la V_{IN} perché è l'uscita del primo blocco dell'alimentatore, blocco costituito da un raddrizzatore-filtro, la V_O perché è la tensione sul carico, stabilizzata dal regolatore switching.

SITUAZIONE A INTERRUTTORE CHIUSO (ON)

Chiamiamo T_{ON} la durata dell'intervallo di tempo nel quale l'interruttore è chiuso.

Se l'interruttore è chiuso, si ha $V_{IN} = V_K$

Di conseguenza (per la premessa “2”) è: $V_L = V_{IN} - V_O = \text{cost}$

Riguardo alla corrente nell'induttore, essa è legata alla tensione da una relazione di tipo integrale:

$$i_L = \frac{1}{L} \cdot \int v_L(t) \cdot dt$$

Abbiamo appurato che v_L è una costante e sappiamo che l'integrale di una costante è una rampa (crescente se la costante è positiva, decrescente se la costante è negativa), di conseguenza la corrente sull'induttore ha un andamento lineare che (eseguendo l'integrale):

$$i_L(t) = \frac{1}{L} \cdot \int v_L(t) \cdot dt = \frac{1}{L} \cdot \int (V_{IN} - V_O) \cdot dt = \frac{V_{IN} - V_O}{L} \cdot \int 1 \cdot dt = \frac{V_{IN} - V_O}{L} \cdot t$$

risulta:

$$i_L(t) = \frac{V_{IN} - V_O}{L} \cdot t \quad \text{che è proprio l'andamento di una rampa crescente.}$$

Se valutiamo la i_L negli istanti iniziale e finale dell'intervallo (0, T_{ON}), otteniamo:

$$i_L(0) = \frac{V_{IN} - V_O}{L} \cdot 0 = 0 \quad i_L(T_{ON}) = \frac{V_{IN} - V_O}{L} \cdot T_{ON}$$

Pertanto l'incremento di i_L nell'intervallo (0, T_{ON}), cioè la grandezza:

$$\Delta i_L^+ = i_L(T_{ON}) - i_L(0)$$

risulta:

$$\Delta I_L^+ = \frac{V_{IN} - V_o}{L} \cdot T_{ON}$$

SITUAZIONE A INTERRUITTORE APERTO (OFF)

Chiamiamo T_{OFF} la durata dell'intervallo di tempo nel quale l'interruttore è aperto.

L'induttore tende a conservare il valore di corrente presente nell'istante finale dell'intervallo precedente (nel quale l'interruttore era chiuso).

Questa corrente è tale da polarizzare direttamente il diodo.

Ai capi del diodo ci sarà quindi la tensione di conduzione diretta, e cioè circa:

$$V_{AK} = 0,6V$$

avremo quindi

$$V_A - V_K = 0,6V$$

$$0 - V_K = 0,6V$$

$$V_K = -0,6V$$

Ricordando la premessa 2, possiamo scrivere $V_L = V_K - V_O$

da cui: $V_L = -0,6 - V_O$

E, trascurando 0,6V rispetto a V_L e a V_O :

$$V_L = -V_O \quad (\text{con } V_O = \text{cost})$$

L'andamento della corrente sull'induttore sarà quindi dato da:

$$i_L(t) = \frac{1}{L} \cdot \int v_L(t) \cdot dt = -\frac{1}{L} \cdot \int V_o \cdot dt = -\frac{V_o}{L} \cdot \int 1 \cdot dt = -\frac{V_o}{L} \cdot t$$

In definitiva:

$$i_L(t) = -\frac{V_o}{L} \cdot t$$

Se valutiamo la i_L negli istanti iniziale e finale dell'intervallo (0, T_{OFF}), otteniamo:

$$i_L(0) = 0 \quad i_L(T_{OFF}) = -\frac{V_o}{L} \cdot T_{OFF}$$

Pertanto l'incremento di I_L nell'intervallo (0, T_{ON}), cioè la grandezza:

$$\Delta I_L^- = i_L(T_{OFF}) - i_L(0)$$

risulta:

$$\Delta I_L^- = -\frac{V_o}{L} \cdot T_{OFF}$$

A regime incremento e il decremento di corrente che abbiamo calcolato devono risultare uguali, cioè deve essere:

$$\Delta I_L^+ = \Delta I_L^-$$

e quindi:

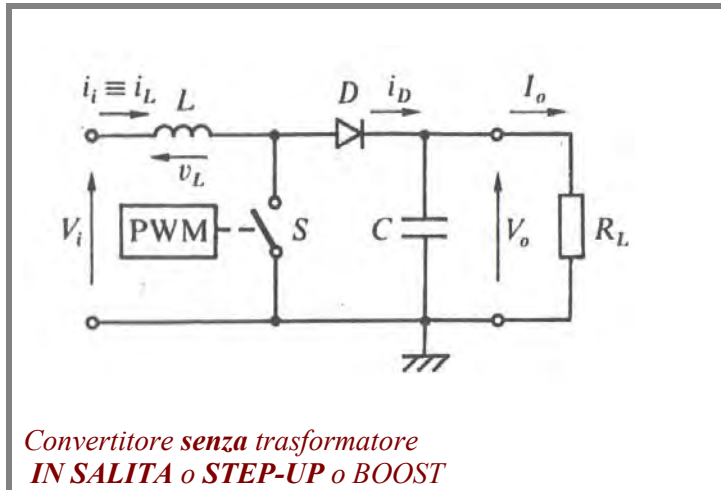
$$\frac{V_{IN} - V_O}{L} \cdot T_{ON} = \frac{V_O}{L} \cdot T_{OFF} \rightarrow (V_{IN} - V_O) \cdot T_{ON} = V_O \cdot T_{OFF} \rightarrow$$
$$(T_{ON} + T_{OFF}) \cdot V_O = T_{ON} \cdot V_{IN}$$

da cui:

$$V_O = V_{IN} \cdot \frac{T_{ON}}{T_{ON} + T_{OFF}}$$

che è proprio la relazione che volevamo dimostrare.

CONVERTITORE DC-DC SENZA TRASFORMATORE DI TIPO "STEP-UP" O "IN SALITA" O "BOOST"

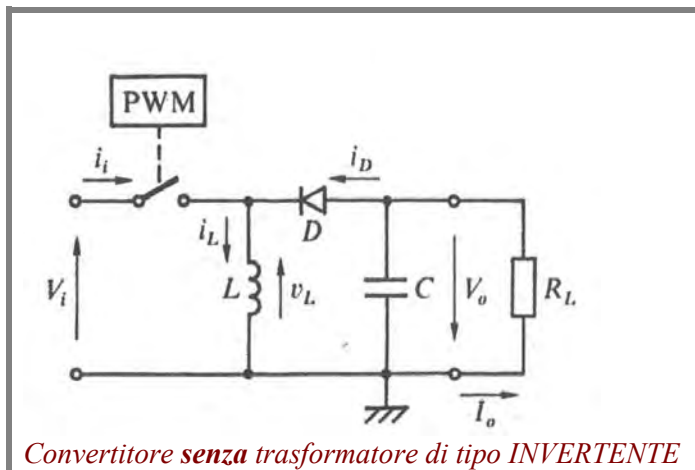


L'interruttore elettronico, a BJT o a MOS, realizza l'inverter e, essendo pilotato dal modulatore PWM, lavora a frequenza fissa e duty-cycle variabile. L'interruttore è posizionato in modo che:

- quando è chiuso (ON), permetta all'induttore di accumulare energia
- quando è aperto (OFF), permetta alla corrente dell'induttore di attraversare il filtro e il carico

$V_o = V_{IN} \cdot \frac{T_{ON} + T_{OFF}}{T_{OFF}}$ dove: T_{ON} = parte del periodo T nella quale l'interruttore è chiuso (ON) T_{OFF} = parte del periodo T nella quale l'interruttore è aperto (OFF) Essendo: T_{OFF} (denominatore) < $T_{ON} + T_{OFF}$ (numeratore), si ha: $V_o > V_{IN}$ (la tensione di uscita è maggiore di quella di ingresso)	RELAZIONE I/O
--	---------------------------------

CONVERTITORE DC-DC SENZA TRASFORMATORE DI TIPO “INVERTENTE”



L'interruttore elettronico, a BJT o a MOS, realizza l'inverter e, essendo pilotato dal modulatore PWM, lavora a frequenza fissa e duty-cycle variabile. L'interruttore è posizionato in modo che:

- quando è chiuso (ON), permetta all'induttore di accumulare energia
- quando è aperto (OFF), permetta alla corrente dell'induttore di attraversare il carico in verso tale da produrre su di esso una tensione di polarità negativa rispetto a massa

$V_o = -V_{IN} \cdot \frac{T_{ON}}{T_{OFF}}$ dove: T_{ON} = parte del periodo T nella quale l'interruttore è chiuso (ON) T_{OFF} = parte del periodo T nella quale l'interruttore è aperto (OFF) Siccome T_{ON} può essere maggiore o minore di T_{OFF}, anche la tensione di uscita V_o può essere maggiore o minore della tensione di ingresso V_{IN} Inoltre la V_o ha sempre polarità opposta rispetto a V_{IN}	RELAZIONE I/O
---	---------------------------------