



Claudio Di Girolamo (jordan20)

LE ONDE ELETTROMAGNETICHE COME MEZZO TRASMISSIVO: PONTI RADIO TERRESTRI (4)

17 December 2013

1. Continuando a... trasmettere sul ponte radio numerico

Riprendiamo la nostra trattazione sulla trasmissione numerica implementata su un ponte radio a microonde.

1.1 Codifica di linea su ponte radio numerico

Come anticipato nelle conclusioni alla [terza parte](#), è il momento di analizzare la maniera con cui l'informazione viene opportunamente manipolata e resa idonea ad attraversare l'etere e giungere al destinatario.

La codifica utilizzata sui ponti radio numerici in esercizio in Telecom Italia è di tipo **non lineare**, ovvero fa uso di codici *caratterizzati da leggi che fanno dipendere il generico simbolo*, che chiamo a_n , *sia alla sequenza completa pronta per essere trasmessa*, che indico con a_m , *sia ai simboli precedentemente codificati* $a_n - k$.

Nella pratica, questi codici sono specificati con delle sigle mnemoniche che richiamano le condizioni imposte sulle sequenza codificata (ad esempio *non più di k zeri consecutivi*). Essi possono essere classificati in:

- **Codici alfabetici.** Effettuano una *trasformazione* di un blocco di Z simboli a_m in un blocco di W simboli a_n , presupponendo che la corrispondenza fra le m^Z parole possibili sulla sequenza a_m sia *univoca* con tutte o parte delle n^W parole possibili sulla sequenza a_n ; ciò non esclude comunque che una parola della sequenza a_m corrisponda a più di una parola di a_n (*codifica non univoca*).

Facciamo un esempio. Supponiamo che il codificatore di linea alfabetico, con sequenza binaria a_m in ingresso, operi la codifica considerando blocchi di $Z = 3$ bit ($m = 2$) fornendo in uscita blocchi di $W = 2$ impulsi a 3 livelli ($n = 3$, ovvero 0, + e -). Si avrà pertanto la corrispondenza fra le $n^W = 3^2 = 9$ possibili parole della sequenza a_n codificata con le $m^Z = 2^3 = 8$ possibili

parole della sequenza a_m da codificare. La schematizzazione dell'esempio è di seguito mostrata:

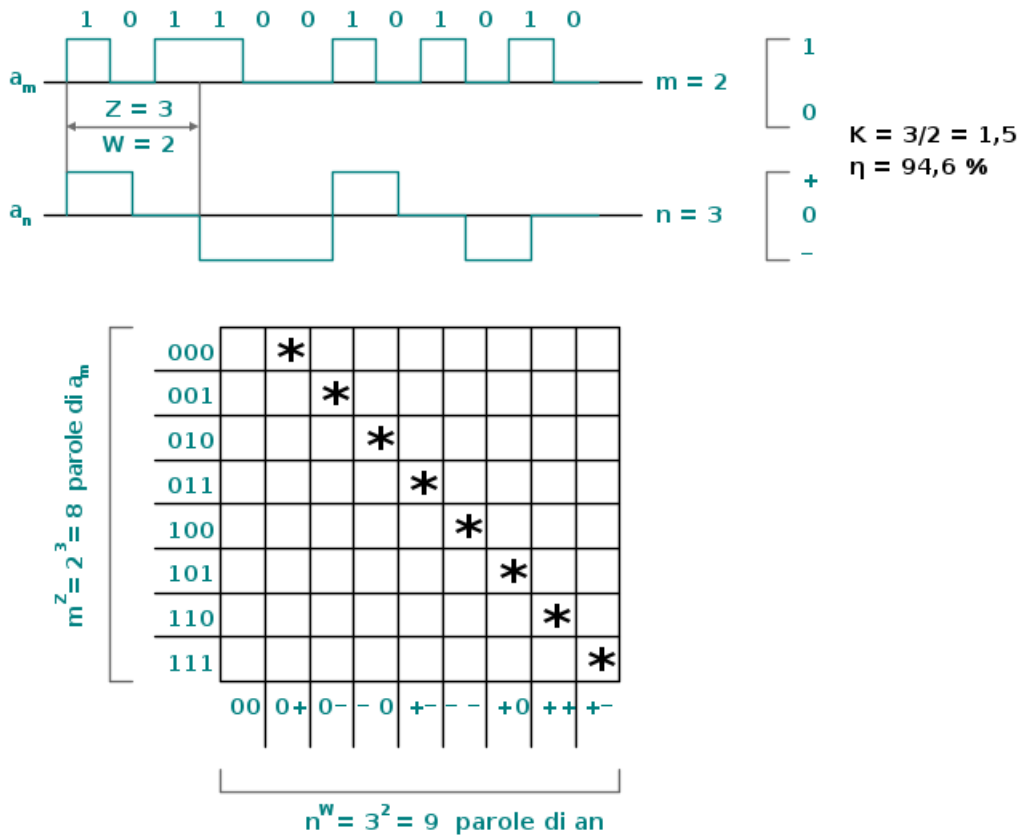


Fig.1

Anticipando qualche nozione di rilevazione d'errore in ricezione, osserviamo che la matrice di corrispondenza fra i blocchi delle due sequenze è il punto su cui agire per ottenere la sagomatura spettrale del canale desiderata e l'energia per l'estrazione del segnale di temporizzazione, mentre la sequenza di uscita **00**, non prevista, costituisce la base per un controllo di qualità in servizio effettuando la verifica della sua eventuale presenza nel codice di linea ricevuto.

- **Codici non alfabetici.** Hanno lo scopo di *eliminare le lunghe sequenze di zeri* sulla sequenza di ingresso al fine di garantire l'autosincronizzazione degli apparati di linea. A loro volta tali codici sono classificabili in *non modali* e *modali* a seconda che la sequenza di riempimento da inserire al posto della sequenza di k zeri entranti sia sempre la stessa o altrimenti dipendente dai dati precedenti. Per esempio se la sequenza a_m è 011010000101 $k = 4$, la sequenza a_n , adottando per esempio una codifica *bipolare alternata* (a tal

proposito rimando al paragrafo 1.1. di [questo](#) mio precedente articolo sulla *tecnica PCM*) diventa:

$$0 + -0 - S_1 S_2 S_3 S_4 - 0 + k = 4$$

dove $S_1 \div S_4$ rappresenta la sequenza di riempimento.

Se il codice è non modale la sequenza $S_1 \div S_k$ è sempre la stessa, mentre se è modale dipende anche dai dati trasmessi.

1.2 Gli scrambler

Uno degli obiettivi della codifica di linea è quello di agevolare il recupero del segnale di temporizzazione per tutte le operazioni di riconoscimento e ricostruzione della sequenza numerica. Questo in pratica significa fare in modo di arricchire quanto più possibile le *transizioni* del segnale emesso sul mezzo trasmissivo, in quanto normalmente la temporizzazione viene proprio estratta dai dati ricevuti.

Un altro obiettivo è quello di sagomare (o per lo meno *ridistribuire*) lo *spettro di potenza* del segnale in modo da renderlo adatto alle caratteristiche del portante fisico.

Per raggiungere tali obiettivi, esiste un dispositivo numerico che, posto all'uscita della sorgente di informazione, trasforma la sequenza numerica originaria in una nuova sequenza numerica (solitamente entrambe binarie) con diverse caratteristiche, nella fattispecie:

- se la sequenza di sorgente è periodica, quella di uscita è ancora periodica ma di periodo molto maggiore tale che possa essere ritenuta **pseudocasuale**;
- la sequenza d'uscita possiede molte transizioni tra livelli.

Con queste peculiarità, questo dispositivo denominato **scrambler** (letteralmente *mescolatore*) può essere convenientemente adoperato per:

- *agevolare il recupero della temporizzazione* nella trasmissione numerica;
- *distribuire in modo pressoché uniforme lo spettro di potenza del segnale d'uscita* in modo da ridurre eventuali "righe" elevate dello spettro che potrebbero interferire, per il loro alto livello o per eventuali non linearità in gioco, con altri canali adiacenti, specie nei ponti radio numerici, dove l'invio ai modulatori di segnali con spettri privi di righe ad alto livello è una tecnica regolarmente applicata;
- *ottenere segretezza delle comunicazioni* in quanto la sequenza di uscita dello scrambler è ottenuta con l'uso di una "chiave" senza la quale non è possibile rendere intellegibile la sequenza stessa.

Per quanto detto, lo scrambler è pertanto da ritenersi un dispositivo che *produce una codifica di linea* ed in genere viene posto subito all'ingresso della sequenza proveniente dalla sorgente (rivedere le Fig.15 e Fig.16 della terza parte).

Dal punto di vista del principio di funzionamento, gli scrambler possono suddividersi in due famiglie:

1. scrambler a *sequenze pseudocasuali*;
2. scrambler *autosincronizzanti*.

I primi sono utilizzati piuttosto nella trasmissione dati via cavo; i secondi nei ponti radio numerici, per cui analizziamo solo questi.

1.2.1 Scrambler autosincronizzanti

L'idea di base di questi dispositivi consiste nel fatto che la sequenza d'uscita è in grado di influenzare lo stato delle celle di un registro a scorrimento e simultaneamente dipendente dai dati d'ingresso. Questo fatto permette al dispositivo duale, il **descrambler**, di essere in grado di recuperare automaticamente eventuali perdite di sincronismo senza avere bisogno di procedure di riallineamento, poiché le informazioni per ottenere quest'ultimo sono ricavate decodificando la sequenza ricevuta, quando ritorna corretta.

Di contro a questo enorme vantaggio vi è l'inconveniente di una moltiplicazione degli errori dovuti al sistema trasmissivo.

Il principio di funzionamento di uno scrambler/descrambler autosincronizzante è analizzabile partendo dall'analisi della seguente figura:

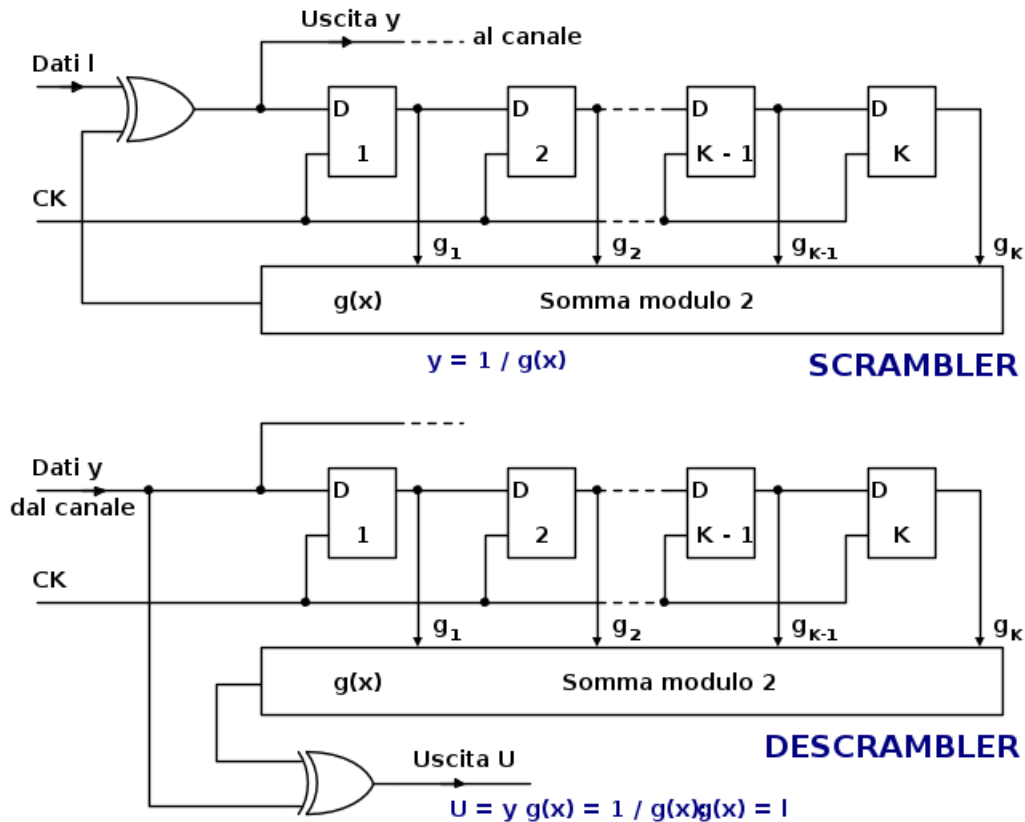


Fig.2

In pratica:

- dato un registro a scorrimento di k celle, esso viene predisposto per realizzare un generatore pseudocasuale secondo un dato polinomio primitivo $g(x)$;
- se si realizza una struttura di un circuito che effettui la *divisione* dei dati di ingresso per $g(x)$, si ottiene uno scrambler;
- se si realizza una struttura di un circuito che esegua la *moltiplicazione* dei dati forniti dal canale per $g(x)$, si ottiene un descrambler;
- se la sequenza di ingresso allo scrambler ha un periodo lungo "S", la sequenza di uscita ha un periodo dipendente dallo stato dello scrambler (e cioè dallo stato dei singoli flip-flop). E' dimostrabile che *esiste un solo stato* dello scrambler, per ogni fase della sequenza di ingresso, per cui il periodo della sequenza d'uscita è ancora lungo "S"; ad ogni modo, eccetto questo caso, la sequenza di uscita ha un periodo che risulta dal *minimo comune multiplo* tra "S" e $2^k - 1$. Pertanto si dimostra che lo scrambler moltiplica per $2^k - 1$ la sequenza "S" di ingresso rendendola pseudocasuale con le

medesime caratteristiche statistiche del generatore pseudocasuale dovuto al polinomio primitivo $g(x)$.

1.3 Modulazione numerica di fase

Nella modulazione di fase numerica, il segnale-dati modulante con livello binario agisce su un modulatore in grado di variare la fase della portante entro prefissati valori. Consideriamo la seguente figura:

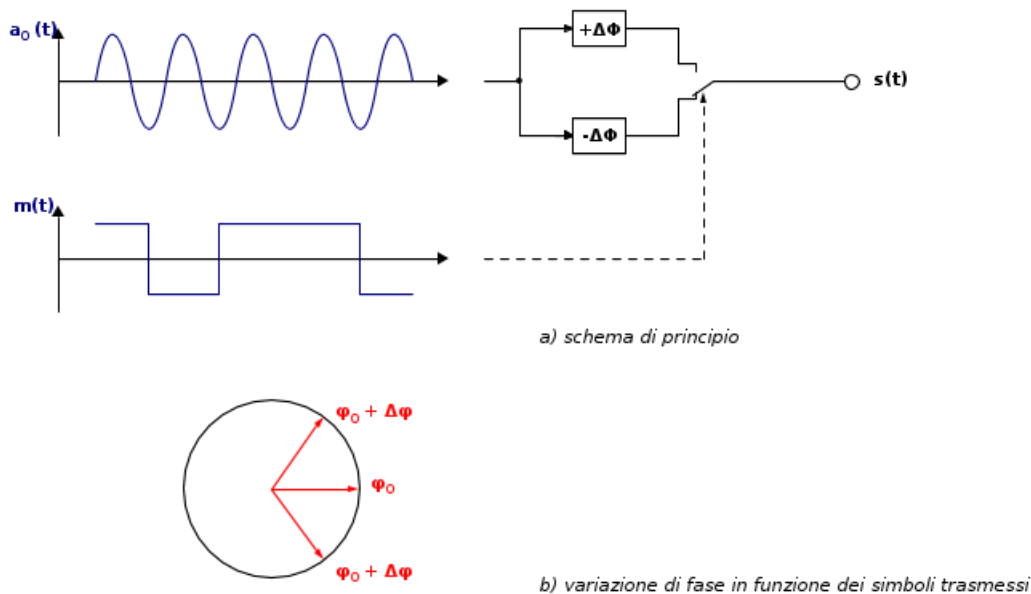


Fig.3

Osserviamo che la funzione modulante $m(t)$ è rappresentata dall'azione del commutatore (deviatore), alle cui posizioni sono associati univocamente i due stati logici della sequenza da trasmettere; l'effetto ottenuto sulla portante è di modificare l'argomento (Fig.3.b)) secondo la legge:

$$\varphi(t) = \varphi_0 + \Delta\Phi \cdot m(t) \quad (1)$$

Come si nota, il valore da assegnare alla deviazione di fase $\Delta\Phi$ può essere arbitrariamente scelto. Ne risulta una tipica modulazione angolare che automaticamente la esclude dalle modulazioni lineari; di fatto, questa modulazione è raramente impiegata (ad eccezione di alcuni collegamenti in ponte radio satellitare) e comunque di scarsa rilevanza pratica.

Esiste però un *caso particolare della modulazione di fase* che, al contrario, assume grande importanza per il larghissimo impiego che incontra nella modulazione numerica; si tratta della **modulazione PSK** (*Phase Shift Keying*), la cui caratteristica, che ne fa una *modulazione lineare*, è quella di *produrre salti di fase*

pari a $\pi / 2$.

La classe delle modulazioni numeriche è oggi considerata con grande interesse, in quanto offre soluzioni notevoli in termini di *efficienza spettrale* e *immunità al rumore*. Classifichiamo alcune peculiarità:

- le funzioni di modulazione e di demodulazione sono ottenute mediante *operazioni lineari di moltiplicazione tra il segnale modulante e la portante*;
- ciò è assimilabile ad una conversione di frequenza, o meglio ad una *traslazione dello spettro di banda base alla frequenza della portante (in modulazione) e viceversa (in demodulazione)*;
- la configurazione dello spettro dell'onda modulata è pertanto *una esatta replica dello spettro in banda base*, per cui il comportamento del canale trasmissivo può essere analizzato basandosi sulle caratteristiche del segnale in banda base.

Il segnale $s(t)$ così ottenuto viene genericamente descritto dalla seguente forma d'onda:

$$s(t) = m_I(t) \cdot A_I \sin \omega_0 t + m_Q(t) \cdot A_Q \cos \omega_0 t \quad (2)$$

dove A_I ed A_Q sono le ampiezze di due sottoportanti in quadratura (cioè con sfasamento reciproco pari a $\pi / 2$), ciascuna modulata in ampiezza da una distinta funzione modulante numerica $m_I(t)$ ed $m_Q(t)$, a cui possono essere assegnati più livelli. L'impiego della *modulazione su due assi ortogonali* è un espediente davvero molto valido al fine di ridurre l'occupazione di banda del segnale modulato.

Poiché la portante trasmessa è la somma vettoriale delle due sottoportanti, l'effetto che se ne ottiene è quello di una *modulazione contemporanea di ampiezza e di fase*. Ciò significa che, utilizzando la consueta forma di rappresentazione vettoriale della portante in rotazione con velocità ω_0 su un piano di riferimento, il segnale modulato può essere rappresentato mediante un dato numero di punti, costituenti la cosiddetta **costellazione**, sul piano suddetto, corrispondenti alle posizioni assunte dall'estremo mobile del vettore stesso, negli istanti in cui esso è sincrono con l'orologio della sequenza da trasmettere.

In realtà, solo la modulazione *2PSK* consente di associare idealmente ad ognuno dei due stati che le sono tipici, uno dei due possibili livelli della sequenza binaria; modulazioni più complesse (quali la *4PSK*, la *8PSK*, la *16QAM* ed altre) possono essere ottenute *associando ad ognuno degli N_S stati della portante, più bit della sequenza-dati da trasmettere*.

Poiché il numero di combinazioni possibili di k bit è pari a 2^k , anche il numero n degli stati discreti che la portante modulata dovrà esibire corrisponderà precisamente ad $n = 2^k$.

In queste modulazioni, definite *multilivello*, i bit della trama, anziché venire trasmessi

ad uno ad uno, *vengono trattati a pacchetti di m bit* attraverso l'azione di un convertitore serie-parallelo, che li emette su vie separate verso il modulatore, con una cadenza di invio pari a $1 / k$ volte la frequenza di cifra originaria.

In quest'ultima considerazione risiede l'interesse per la modulazione multilivello; infatti la larghezza di banda richiesta al canale trasmissivo è proporzionale proprio alla frequenza di cifra del flusso numerico da trasmettere, sicché: *riducendo di k volte la velocità d'invio dei simboli, si riduce dello stesso rapporto la banda necessaria per la trasmissione dello spettro in una determinata trama numerica.*

Un aspetto negativo è invece rappresentato dal *peggioramento della robustezza*, cioè della *resistenza agli effetti di disturbo*, tanto più rilevante al crescere della complessità della costellazione. La ragione risiede sostanzialmente nella riduzione delle **regioni di decisione**, termine con il quale definiamo *le porzioni di piano entro le quali il vettore rappresentativo della portante può discostarsi dalla posizione teorica, coincidente con un punto della costellazione, senza che venga originato un errore*. Queste aree di decisione (che saranno meglio discusse nelle prossime parti), sono tanto più ridotte quanto maggiore è il frazionamento del piano in conseguenza della complessità della costellazione.

Si comprende pertanto che gli stati discreti da assegnare alla portante saranno tanto più facilmente discriminabili, quanto maggiore sarà attorno a ciascuno di questi l'area di decisione.

1.3.1 Modulazione PSK binaria

Le modulazioni lineari sono basate sul **metodo PSK**, cioè sulla "manipolazione a salto di fase"; tuttavia è interessante notare che se si trattasse realmente di modulazione di fase esse non sarebbero formalmente lineari. Di fatto, la linearità della modulazione PSK risiede nel principio per cui *i salti di fase sono sempre e soltanto di Π* .

L'espressione analitica della modulazione PSK binaria si ricava considerando il seguente schema di principio:

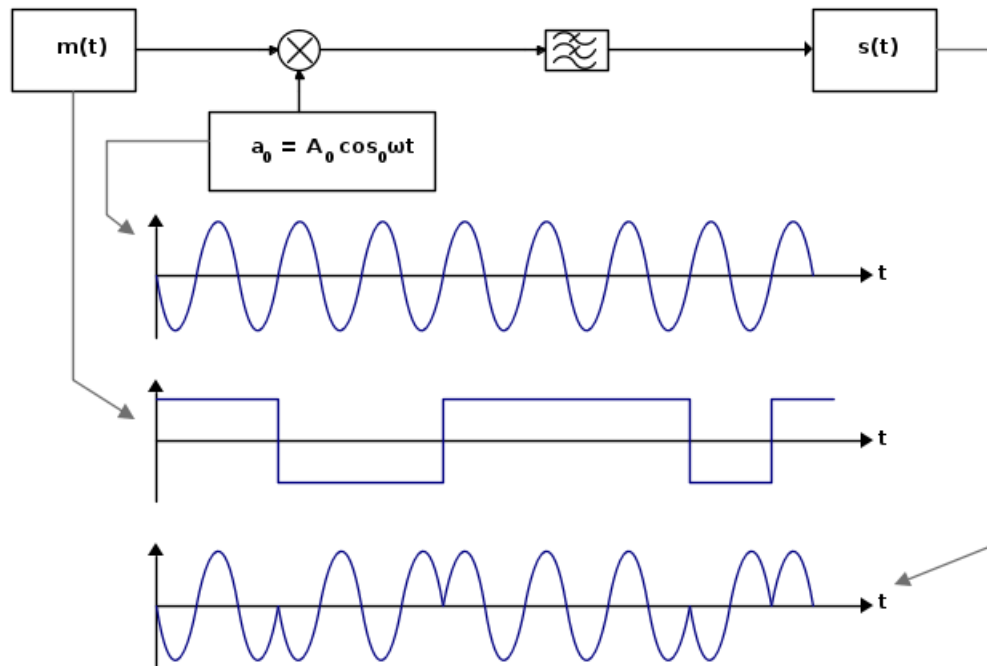


Fig.5

Esso mette in evidenza il carattere di *modulazione a prodotto*:

$$s(t) = m(t) \cdot A \cos \omega_0 t \quad (3)$$

Possiamo notare che tale espressione risulta essere un caso particolare della (2), da cui è praticamente esclusa la sottoportante in quadratura.

Il segnale modulante è di tipo binario (simboli "1" e "0") e nulla vieta di assegnare alla funzione modulante $m(t)$ due livelli logici in forma bipolare, ovvero "+1" e "-1". Pertanto:

- per $m(t) = 1 \Rightarrow s(t) = +A \cos \omega_0 t$;
- per $m(t) = -1 \Rightarrow s(t) = -A \cos(\omega_0 t + \pi)$.

Il risultato del prodotto fornisce quindi una portante in cui *ad ogni transizione del segnale dati, la fase della portante si rovescia istantaneamente*.

1.3.2 Modulazione QPSK

La **QPSK** (*Quadrature Phase Shift Keying*) rappresenta il caso più semplice di *modulazione su due portanti tra loro ortogonali*; il principio di funzionamento è di seguito schematizzato:

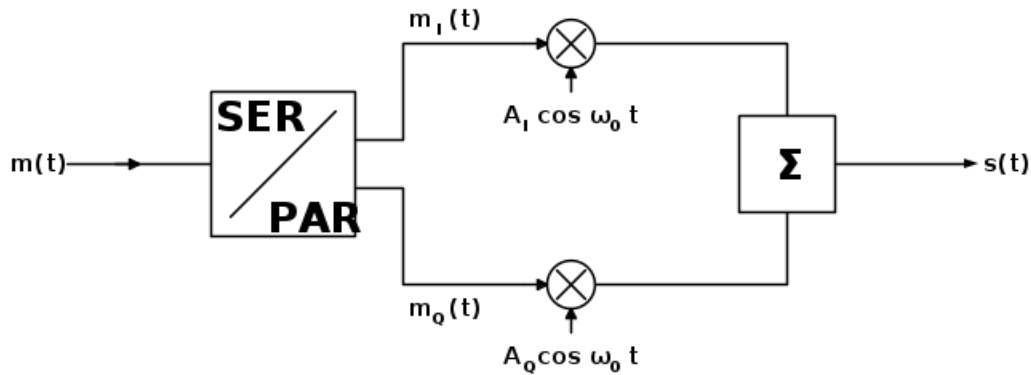


Fig.6

Le due portanti vengono modulate ciascuna secondo il principio della modulazione PSK binaria da due sequenze di dati ottenute dal segnale $m(t)$.

Il segnale modulato $s(t)$ è poi semplicemente generato *sommando linearmente le due portanti*; il che permette ancora comunque di considerare la modulazione 4PSK come l'insieme di due canali binari indipendenti, aventi la stessa frequenza f_0 ; il fatto di presentare relazione reciproca di fase a $\pi / 2$ permette ai due canali di *coesistere senza disturbarsi reciprocamente*.

Il principale vantaggio del metodo 4PSK sta nel *risparmio della banda*, dovuto alla possibilità che esso offre di dimezzare la velocità di simbolo, come anticipato. Consideriamo adesso la seguente figura:

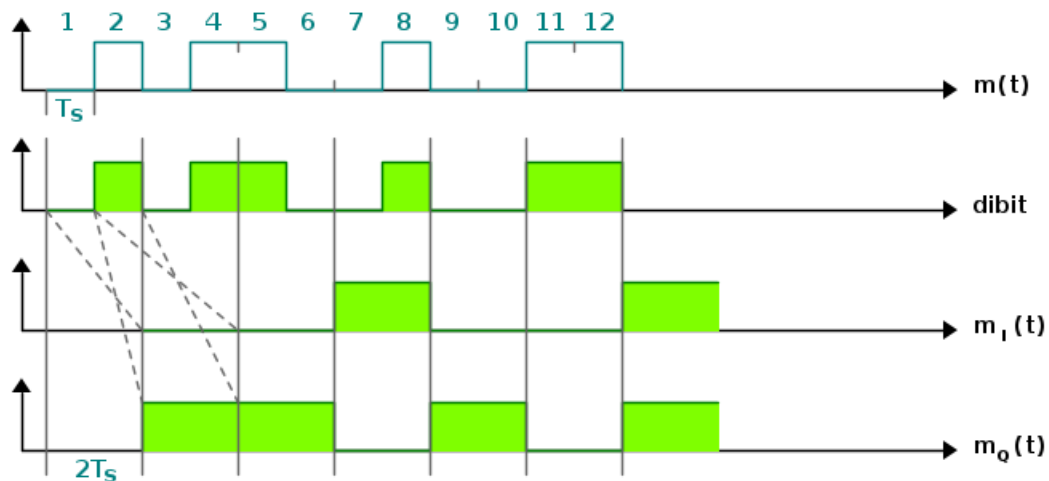
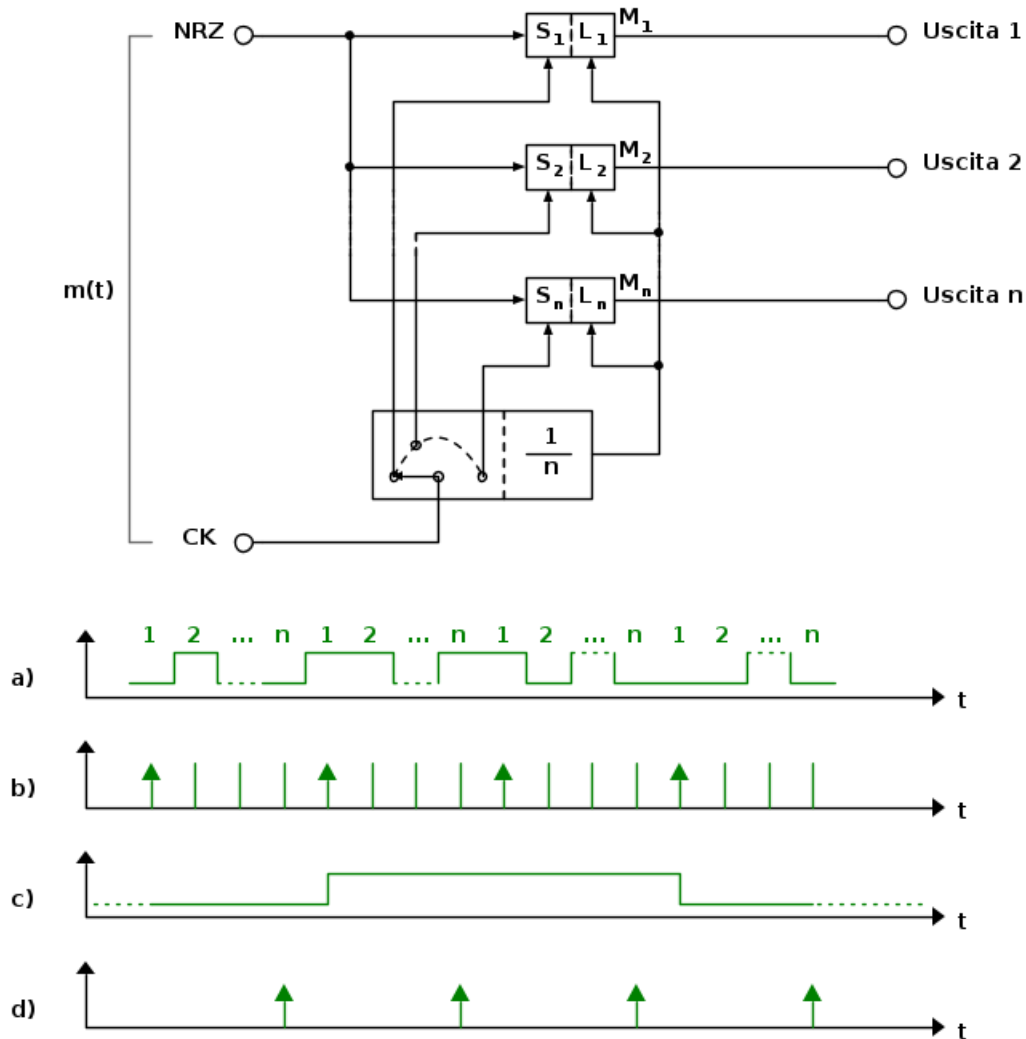


Fig.7

Il flusso di bit $m(t)$ viene suddiviso su due vie, ciascuna operante con un ritmo pari alla metà della velocità di simbolo. Le due sequenze generano ciascuna uno spettro che è largo la metà di quello della sequenza originaria; la somma delle due modulazioni poi, fa sì che tali spettri si sovrappongano, senza che ciò produca un

allargamento dello spettro risultante!

L'operazione di suddividere in due vie i bit della trama numerica è affidato al *convertitore serie-parallelo* di seguito schematizzato (nel caso esteso a n vie):



- a) segnale $m(t)$ seriale formato NRZ + CK
 b) distribuzione temporale del CK
 c) contenuto della cella di memoria S1
 d) istanti di lettura (a frequenza CK/n)

Fig.8

Il circuito riceve separatamente dati, presentati mediante *codifica NRZ*, e segnale di sincronismo CK dai circuiti che lo precedono, solitamente convertitori di codice. Ogni via dispone di una memoria, che legge la trama nell'istante t di sua competenza. Dopo n tempi di clock, ogni cella di memoria ha incamerato un bit della sequenza $m(t)$:

- il primo bit nella cella S_1 ;
- il secondo bit nella cella S_2 ;
- l'n-esimo bit nella cella S_n .

A questo punto si ha la fase di lettura delle celle di memoria comandata da un impulso di clock ottenuto per divisione ennesima del clock d'ingresso. Pertanto, mentre la scrittura sulle memorie è sequenziale, la lettura è parallela su n vie. Concluso pertanto un ciclo all'n-esimo bit, inizia un nuovo ciclo reiterando il meccanismo.

Nel caso 4PSK, il convertitore serie-parallelo alimenta semplicemente due vie, il che equivale a dire che sulla via $m_I(t)$ si troveranno tutti i *bit dispari* (1,3,5,...) e sulla via $m_Q(t)$ tutti i *bit pari* (2,4,6,...), con velocità dimezzata rispetto a quella di $m(t)$. Le due vie I e Q emettono pertanto una coppia di bit paralleli (il cosiddetto *dibit*) che agisce simultaneamente sui due moltiplicatori lineari. L'azione dei due flussi modulanti in quadratura è di seguito schematizzata:

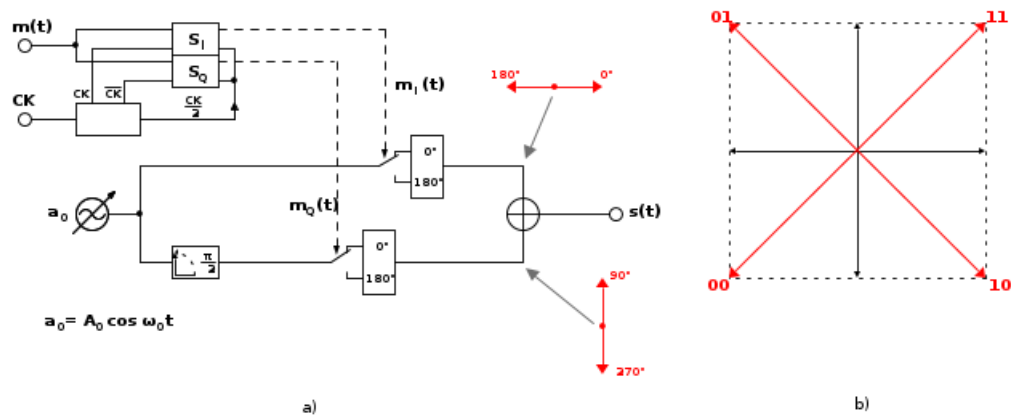


Fig.9

Tali flussi pilotano i commutatori, ciascuno dei quali produce una modulazione $0^\circ - 180^\circ$ esattamente del tipo 2PSK precedentemente esaminata. Notiamo tuttavia che, pur essendo entrambi i rami alimentati dalla stessa sorgente a_0 , una sottoportante viene sfasata di $\pi / 2$, per cui lo sfasamento reciproco tra due sottoportanti vale $\pm \pi / 2$.

Infine, dopo la modulazione 2PSK, le due sottoportanti entrano in un circuito sommatore che ne esegue la somma vettoriale. E' immediato osservare che *il vettore somma potrà assumere una fra quattro possibili posizioni di fase* (si faccia riferimento al diagramma vettoriale di Fig.9b)), *a cui corrisponde una combinazione univoca del dibit presente sui moltiplicatori*.

1.3.3 Codifica differenziale

Nella modulazione PSK, grazie all'azione del circuito che rende coerente l'oscillazione locale necessaria per la demodulazione, la fase di questo può trovarsi allineata ed agganciata in modo del tutto casuale ad una qualunque delle posizioni discrete di fase della portante. Inoltre, per una qualsivoglia discontinuità nel collegamento, il vettore del segnale locale può sganciarsi da una posizione di fase e riagganciarsi così su di un'altra.

Se ad ogni fase della portante venisse rigidamente assegnato un determinato simbolo, *si creerebbe una sostanziale ambiguità nel corretto riconoscimento, in ricezione, degli stessi simboli.*

Un possibile metodo per eliminarla è quello di trasmettere una sequenza di contenuto noto. E' il caso, ad esempio dei ponti radio satellitari TDMA che, operando su flussi discontinui, all'inizio della trasmissione inviano il cosiddetto *preambolo*, ovvero un segnale noto, riconosciuto come esatto solo se la fase del segnale coerente è corretta. Nel caso invece dei ponti radio terrestri, l'ambiguità viene eliminata impiegando la **codifica differenziale** *che permette di associare ad ogni simbolo trasmesso, non più un valore assoluto, bensì un valore relativo (cioè una variazione) della fase della portante.* Poiché una variazione di fase viene sempre riconosciuta in modo non ambiguo, qualunque sia la posizione assoluta assunta precedentemente, questo artificio risolve il problema.

Per comprendere il meccanismo di tale codifica rifacciamoci al caso della modulazione 2PSK e consideriamo la seguente figura:

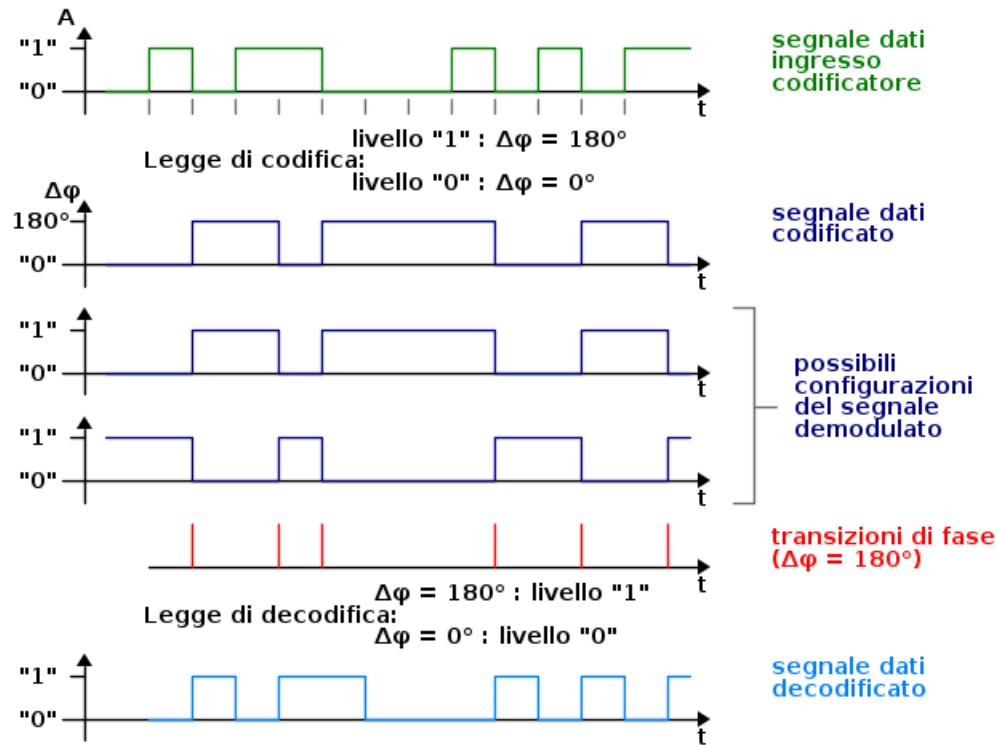


Fig.10

La legge di codifica provoca sulla portante una transizione di fase ($0 \Rightarrow \pi$) ogni volta che la sequenza da trasmettere presenta un simbolo "1" e, viceversa, non provocando nessuna transizione quando il simbolo da trasmettere è uno "0". Indicando con $a(k)$ la sequenza originaria dei simboli da trasmettere, il codificatore invia al modulatore una nuova sequenza $a'(k)$ tale per cui si abbia:

$$a'(k) = a(k) \oplus a'(k-1) \quad (4)$$

dove il segno $\oplus$ indica la *somma modulo 2* realizzata con la funzione logica dell'OR esclusivo (XOR), eseguita tra il segnale $a(k)$ e il segnale di uscita $a'(k)$ *ritardato* di un tempo di simbolo.

1.3.4 Modulazioni PSK di grado superiore

Estendendo il concetto di modulazioni ad inversione di fase, sono stati realizzati, nel corso degli anni, modulatori capaci di esibire 8 e 16 livelli discreti di fase.

Utilizzando il solito circuito moltiplicatore, possono venire generate separatamente modulazioni 2PSK, che vengono successivamente sommate previa opportuna rotazione di fase. Lo schema di un modulatore 8PSK così concepito e la relativa costellazione sono di seguito illustrati:

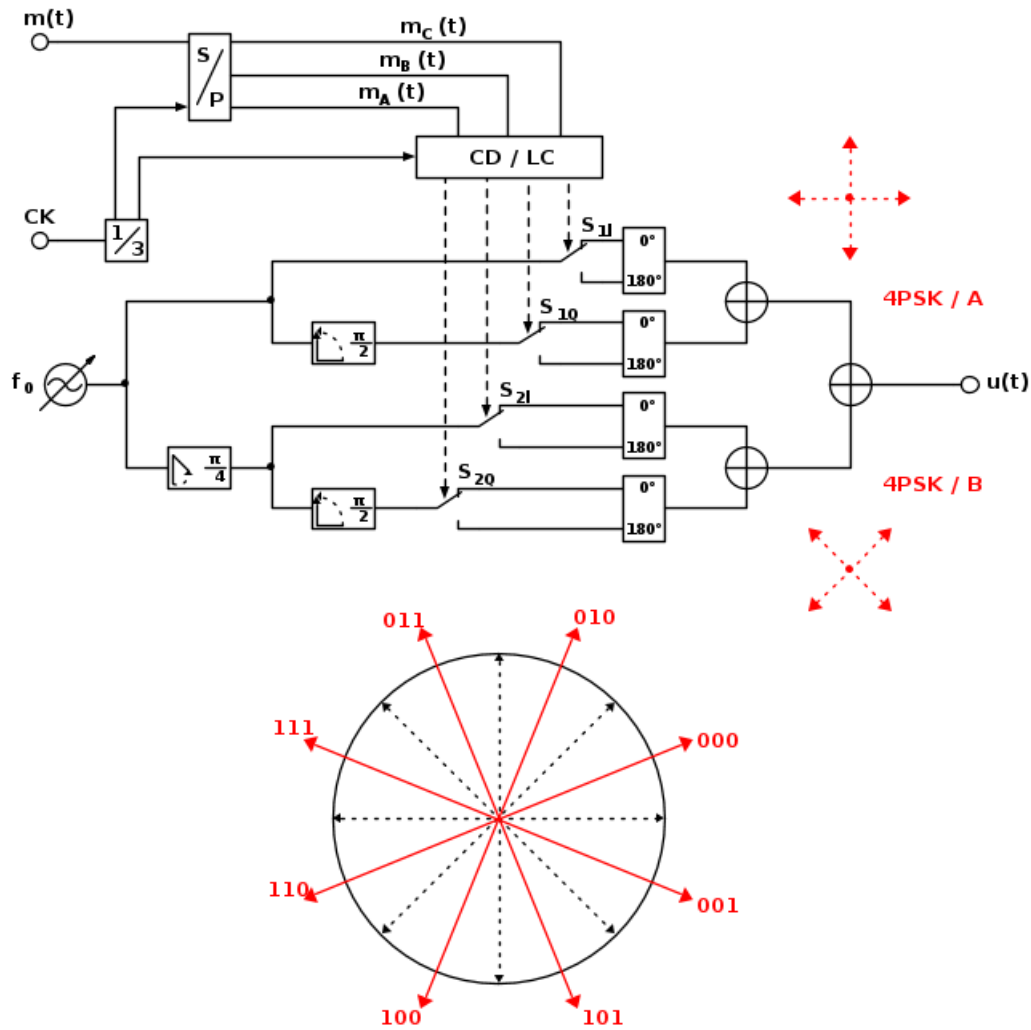


Fig.11

Le 8 fasi vengono ottenute sommando due modulazioni quadrifase sfasate tra loro di $\pi / 4$, a loro volta ottenute da modulazioni binarie, col metodo PSK. La costellazione ottenuta presenta così 8 punti, che dividono la circonferenza in archi pari a $\pi / 4$. A ciascun punto è associata una parola binaria composta da tre bit. Per tale ragione, il flusso dei dati d'ingresso $m(t)$ viene ripartito su tre vie separate $m_A(t)$, $m_B(t)$ e $m_C(t)$, la cui velocità di simbolo è $1 / 3$ del bitrate di $m(t)$. Ciò consente quindi di ridurre di tre volte la larghezza dello spettro occupato dalla portante modulata.

I tre flussi forniti dal convertitore serie-parallelo (S/P) vengono poi elaborati dal circuito per la codifica differenziale (CD) e da una logica di comando (LC) che fornisce i criteri di pilotaggio per i quattro invertitori comandati.

Sia nella 4PSK che nella 8PSK, le terne di bit vengono codificate secondo il codice GRAY per ottenere il vantaggio di dimezzare gli errori sulla trama numerica, in conseguenza di errato riconoscimento del simbolo sul processo di demodulazione.

Si intuisce infine che aumentando il numero dei modulatori 4PSK e sommandoli con opportune relazioni di fase, è teoricamente possibile aumentare il numero di livelli di modulazione, anche se poi in pratica tale possibilità non viene sfruttata. Infatti, quando il numero di livelli supera quello di 16, la modulazione PSK può perdere di interesse a causa della notevole complessità degli apparati di mo-demodulazione e, nel caso delle applicazioni su ponte radio, della riduzione dei margini nei confronti del rumore termico e dei disturbi nella propagazione; in tal caso, pertanto, è conveniente passare ad altri metodi di modulazione.

1.3.5 Modulazione QAM

La **modulazione QAM** (*Quadrature Amplitude Modulation*) fa parte della classe delle *modulazioni lineari miste*, ossia quella categoria di modulazioni dove *intervengono a definire il simbolo trasmesso sia variazioni di ampiezza della portante, sia variazioni d'angolo*.

La QAM è fondamentalmente una modulazione multilivello quadrifase che esibisce la tipica costellazione in cui i livelli modulanti sono distribuiti nei punti d'incrocio di una rete ortogonale. Di seguito sono illustrati lo schema a blocchi di un modulatore *16QAM*, la costellazione associata e il cronogramma relativo alla conversione serie-parallelo e all'andamento dei flussi modulanti:

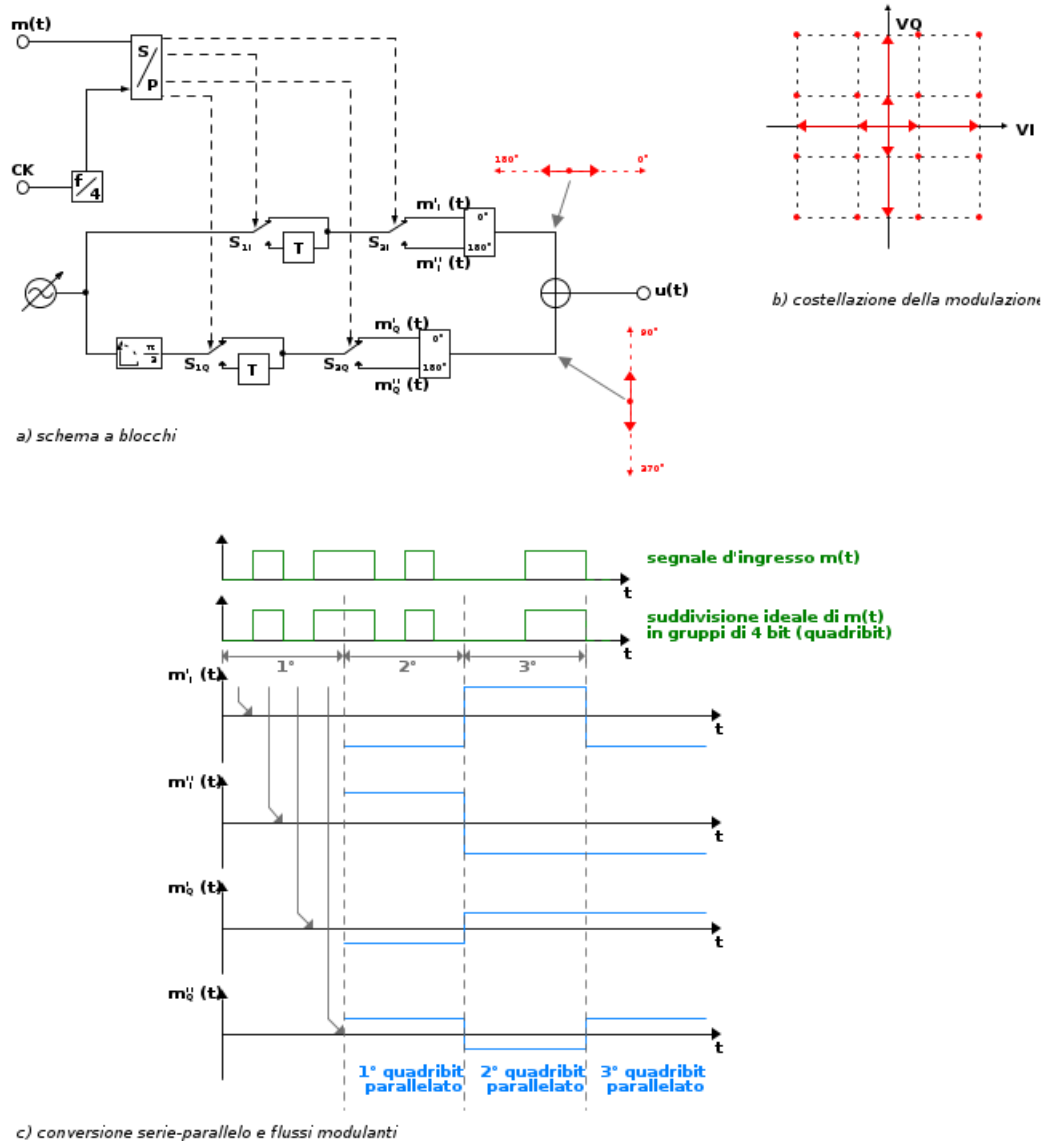


Fig.12

In particolare notiamo che la diversa distanza dei 16 punti al centro del diagramma denota una modulazione d'ampiezza del vettore portante, la cui posizione angolare dipende anche dalla modulazione PSK impartita sulle due sottoportanti in quadratura.

Sempre osservando la Fig.12, la conversione serie-parallelo trasforma il segnale-dati $m(t)$ in 4 segnali sincroni di durata quadrupla; con ciò l'ingombro di banda viene ridotto di 4 volte rispetto ad un sistema bifase di pari capacità. Si nota che i comandi agenti sulle sottoportanti in quadratura, schematizzati come deviatori, sono quattro, in particolare:

- S_{1I} ed S_{1Q} agiscono sull'ampiezza delle sottoportanti applicando una modulazione d'ampiezza numerica ASK, o meglio un suo caso particolare denominato *modulazione OOK (On-Off Keying)* la cui azione è brevemente schematizzata nella seguente figura:

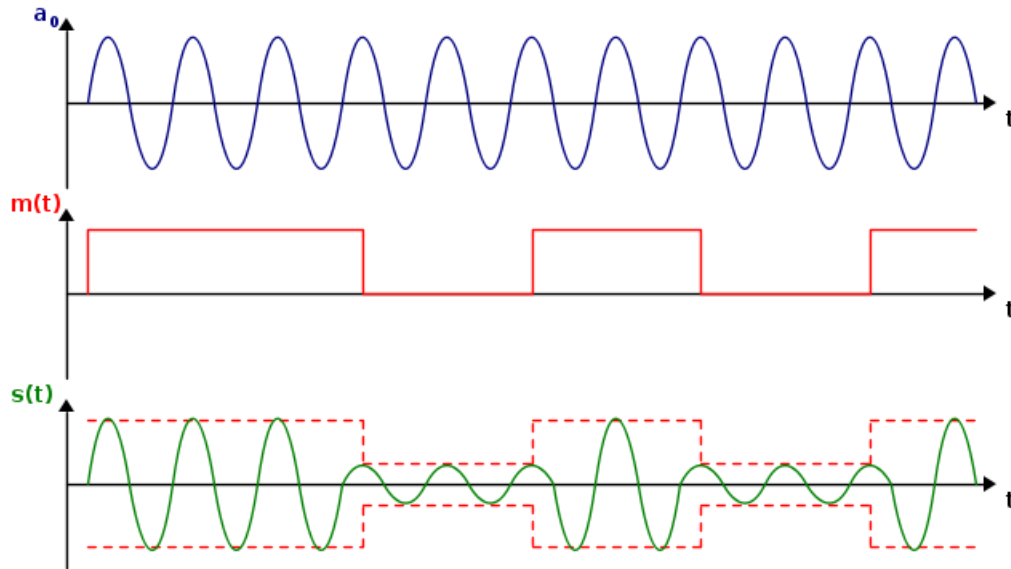


Fig.13

- S_{2I} ed S_{2Q} agiscono sulla fase delle sottoportanti, secondo il principio già visto a proposito della modulazione 4PSK.

Esaminando separatamente le due vie I e Q, è immediato osservare che il vettore uscente dal moltiplicatore presenta 4 possibili situazioni di ampiezza e fase, corrispondenti alle 4 combinazioni dei commutatori S_1 e S_2 ; l'operazione di somma dei due vettori così modulati, equivale ad esaminarne la composizione vettoriale quando essi sono riportati sul piano cartesiano mutuato da entrambi.

Le coppie di coordinate definite da tutte le possibili configurazioni dei due vettori, determinano la costellazione della modulazione, che indica gli stati discreti che la portante può assumere in un arbitrario tempo di clock.

Per quel che riguarda i ponti radio numerici, la QAM rappresenta il processo di modulazione ottimale in quanto realizza un ottimo compromesso tra efficienza spettrale e complessità circuitale.

Dal punto di vista analitico, l'equazione dell'onda modulata QAM presenta la somma di due termini, rispettivamente in seno e coseno, indicatore dell'ortogonalità delle due sottoportanti, ciascuno modulato in ampiezza da un segnale logico, rispettivamente $m_I(t)$ e $m_Q(t)$:

$$s(t) = A_I \cdot m_I(t) \cdot \sin \omega_0 t + A_Q \cdot m_Q(t) \cdot \cos \omega_0 t \quad (5)$$

Supponendo che le sequenze modulanti in quadratura presentino ciascuna 2^n livelli discreti, il numero dei livelli modulanti sarà quindi pari a 2^{2n} . Aumentando il numero dei livelli d'ampiezza delle due sequenze modulanti, aumenta di conseguenza il numero dei possibili stati della portante e la complessità della costellazione. Così, ad esempio, considerando 4 livelli d'ampiezza, si ottengono 8 diverse possibilità di modulazioni sui singoli vettori I e Q e un totale di 64 possibili stati discreti sulla portante modulata in quadratura. E' il caso della **modulazione 64 QAM**, implementata nei ponti radio numerici a grande capacità attualmente in esercizio.

2. Funzionamento dei modulatori numerici reali

Gli apparati trasmissivi implementati su un ponte radio numerico, possiedono modulatori in grado di modulare una portante a radiofrequenza in due modi:

1. *elaborando direttamente la portante* generata dall'oscillatore locale, mediante dispositivi capaci di operare nel campo delle microonde;
2. *modulando una portante a frequenza intermedia (FI)*, la quale viene successivamente convertita alla frequenza del canale radio da trasmettere, mediante un convertitore in salita (*up-converter*).

Vengono di seguito esaminati entrambi i metodi, con la premessa che, pure essendo il primo assai meno oneroso, modulazioni più sofisticate vengono realizzate sempre col secondo metodo, che offre anche il vantaggio di operare a frequenze portanti relativamente basse 70 MHz o 140 MHz.

Entrambi i metodi hanno comunque in comune tutta la parte trattata riguardante *l'elaborazione del segnale in banda base*, che precede il vero e proprio processo di modulazione.

2.1 Circuiti in banda base

Il segnale numerico perviene al ponte radio da apparati diversi, che possono essere un sistema multiplex a 2 Mb/s o, nei casi di ordini gerarchici superiori, da un multiplatore. La sistemistica della rete prevede che tra il ponte radio e gli apparati che lo precedono, possa esistere una linea fisica di una certa lunghezza; pertanto è necessario che il segnale in banda base venga trattato attraverso una serie di operazioni che ne eseguono la *rigenerazione*, come di seguito illustrato:

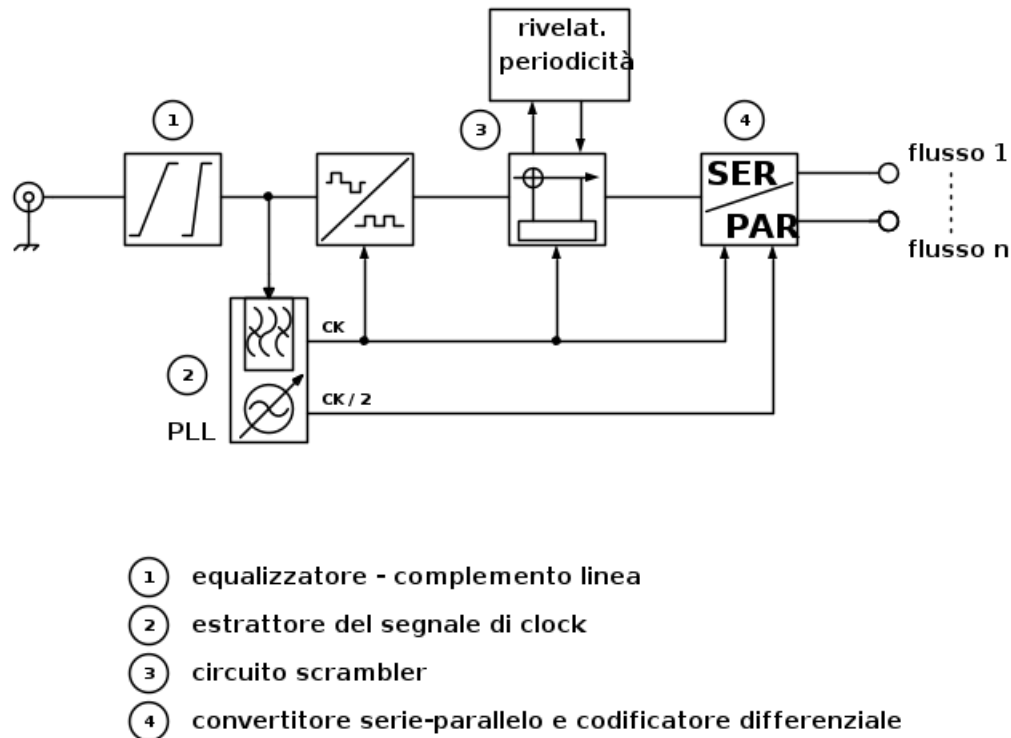


Fig.14

L'involuppo del segnale, infatti, può aver subito distorsioni d'ampiezza lungo la linea e quindi il primo stadio di un generico rigeneratore è un **equalizzatore a più celle**, che viene opportunamente predisposto in sede d'impianto a seconda della lunghezza della linea.

Successivamente, il segnale subisce una *conversione di codice e di formato*. Infatti, le esigenze di trattamento del segnale in fase di modulazione, richiedono che esso si presenti in formato NRZ (*No Return to Zero*), mentre in linea il segnale richiede di essere codificato con criteri particolari, che tengono conto delle caratteristiche fisiche della linea, quali l'esigenza di mantenere nulla la componente continua (argomento ampiamente esposto nei precedenti miei articoli sulla tecnica PCM). Pertanto viene cambiato il formato, che da bipolare viene reso unipolare mediante raddrizzamento, nonché il codice, che da formato HDB3 o simile, viene ricondotto al codice binario originale.

Il segnale di sincronizzazione (clock) viene estratto dal flusso numerico entrante mediante un PLL che impiega il VCO con oscillatore quarzato, la cui frequenza e fase vengono tenute agganciate alla ricorrenza delle transizioni del segnale d'ingresso. Lo stesso CK va poi a ritemporizzare i dati NRZ, che quindi si presentano all'uscita del rigeneratore nelle condizioni richieste per il pilotaggio degli stadi successivi.

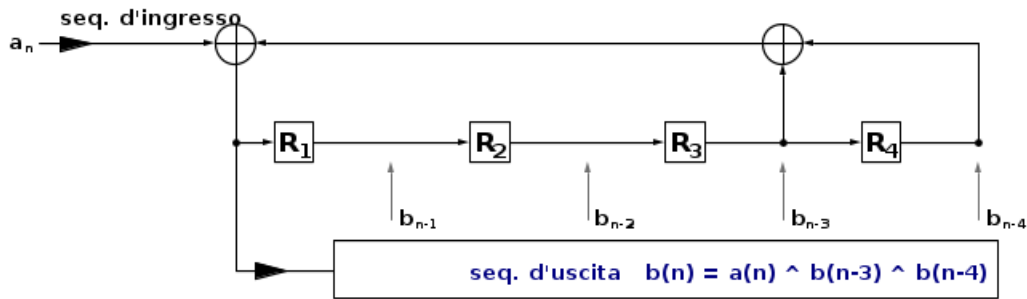
Segue lo scambler, che esegue la funzione di modificare continuamente, secondo una legge matematica nota, il contenuto del segnale numerico, al fine di "pseudocasualizzare" la sequenza modulante, anche in caso di ripetitività delle

configurazioni di trama. Questo garantisce la caratterizzazione statistica sullo spettro del segnale modulato; vengono così eliminate *eventuali parole periodiche* e le corrispondenti righe fisse sullo spettro, le quali potrebbero creare notevoli problemi (interferenze su altri sistemi radio, difficoltà nell'aggancio della portante del modulatore, ecc.).

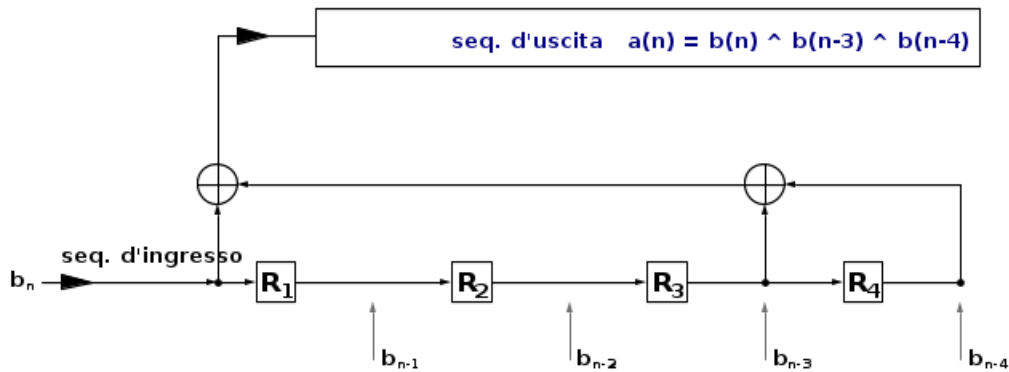
Gli scrambler implementati sono di due tipi:

1. *autosincronizzanti*;
2. *scrambler set-reset*.

Una schematizzazione del sistema scrambler/descrambler autosincronizzante implementato è di seguito riportata:



a) circuito scrambler



a) circuito descrambler (per entrambi gli schemi, il simbolo $\wedge$ è equivalente all'OR ESCLUSIVO)

Fig.15

Essendo a_n la funzione d'ingresso, la funzione d'uscita b_n vale:

$$b_n = a_n \oplus b_{n-3} \oplus b_{n-4} \quad (6)$$

Si può dimostrare che la *periodicità P della funzione in uscita*, è data dalla relazione:

$$P = k(2^N - 1) \quad (7)$$

in cui k è la periodicità dei simboli in ingresso. Si nota che *tale periodicità sarà tanto maggiore, quanto più alto è il numero N delle celle dello scrambler*. Esiste comunque sempre una condizione per la quale $P = k$, cioè la periodicità in uscita è uguale a quella d'ingresso, tanto meno ricorrente quanto naturalmente è maggiore N . Nei casi pratici, è comune utilizzare un valore di $N = 11$.

Il fatto che il circuito sia autosincronizzante, significa che non richiede siano mantenuti in passo, in qualche modo, i rispettivi circuiti in ricezione e trasmissione, per risalire alla corretta sequenza di simboli, in quanto ciò è garantito dalla legge matematica che governa il sistema; solo nello stato iniziale di avvio e stabilizzazione del circuito, esso può compromettere un numero massimo di errori pari al numero delle celle, dopodiché il sistema va a regime senza più commettere errori.

Lo scrambler set-reset, pur essendo molto simile all'autosincronizzante, richiede l'impostazione di uno stato iniziale correlato tra i due circuiti, per risalire alla sequenza originaria. A fronte di questa complicazione, esso ha però il vantaggio di essere del tutto indifferente alla periodicità d'entrata e di risultare meno sensibile alla propagazione di eventuali errori.

Il convertitore serie-parallelo che segue lo scrambler ha la funzione di *smistare N bit consecutivi su altrettante vie parallele d'uscita*, dirette ad alimentare altrettanti modulatori elementari PSK. N vale:

2 nel caso di modulazioni in quadratura 4PSK;

3 nel caso 8PSK;

4 nel caso 16PSK e 16QAM.

e così via per modulazioni di crescente complessità. Il relativo circuito è fondamentalmente un registro in cui la scrittura avviene consecutivamente nelle N celle su comando del CK di cifra, mentre la lettura avviene contemporaneamente a ritmo pari al tempo di simbolo. L'insieme degli N bit d'uscita è infatti definito come il **simbolo**, che va a modificare lo stato della portante con cadenza di periodo pari a $T_s = N \cdot T_b$ (dove T_b è il periodo di ricorrenza di bit del segnale numerico originario).

Il codificatore differenziale consente di superare il problema *dell'ambiguità di fase sull'aggancio della portante ricostruita* in ricezione. La legge di codifica associa ad ognuno degli N possibili simboli, come esposto nel primo paragrafo, una variazione prestabilita della fase della portante. Nel caso 4PSK, la corrispondenza tra simboli trasmessi e variazioni di fase è la seguente:

simbolo salto di fase

00	0°
01	$+90^\circ$

10	-90°
11	180°

2.2 Tecniche di filtraggio numerico in banda base

L'esigenza di limitare quanto più possibile l'estensione dello spettro del segnale numerico modulato e, nel contempo, di controllare e minimizzare l'effetto dell'interferenza intersimbolica (di cui discuteremo estesamente nel prossimo articolo) impone di definire con opportuni criteri la sagomatura del canale di trasmissione. Come noto, il filtraggio che realizza il migliore compromesso tra le suddette esigenze è quello che presenta una *simmetria complementare attorno alla frequenza di taglio* (-6 dB) e la cui curva di risposta ha la forma a coseno rialzato. Di fatto, le sagomature che rispondono a tale legge sono infinite, differenziandosi (come argomentato nella terza parte) per i valori di roll-off compresi tra 0 (filtro ideale di Nyquist) ed 1. Nella pratica si usano sagomature di canale con valore di roll-off molto prossimo a **0,5**.

Nel caso di ponte radio, *il filtraggio è ripartito in vari filtri lungo il processo di elaborazione del segnale modulato attraverso stadi in banda base, frequenza intermedia e radiofrequenza*. Il filtraggio in banda base è distribuito fra trasmissione e ricezione, dove assolve a funzioni supplementari quali: attenuazione del rumore termico (in ricezione), contenimento dell'ingombro spettrale (in trasmissione).

Nei sistemi multilivello ad alta velocità di cifra *sono richieste sagomature di banda molto precise, con fase lineare e soprattutto simmetriche rispetto alla portante*; ottenere tali caratteristiche è una questione davvero problematica con filtri convenzionali LC a costanti concentrate o distribuite. Al filtraggio realizzato a frequenza intermedia o a RF *si preferisce quello ottenuto in banda base*, che garantisce una stabilità e precisione altrimenti impossibili ed inoltre consente di ricorrere vantaggiosamente all'impiego delle *tecniche di filtraggio digitale*.

Si impiegano a tal proposito **filtri numerici o trasversali BTF** (*Bit Transversal Filters*), essenzialmente costituiti da registri a scorrimento associati ad una rete di pesatura resistiva, come di seguito illustrato:

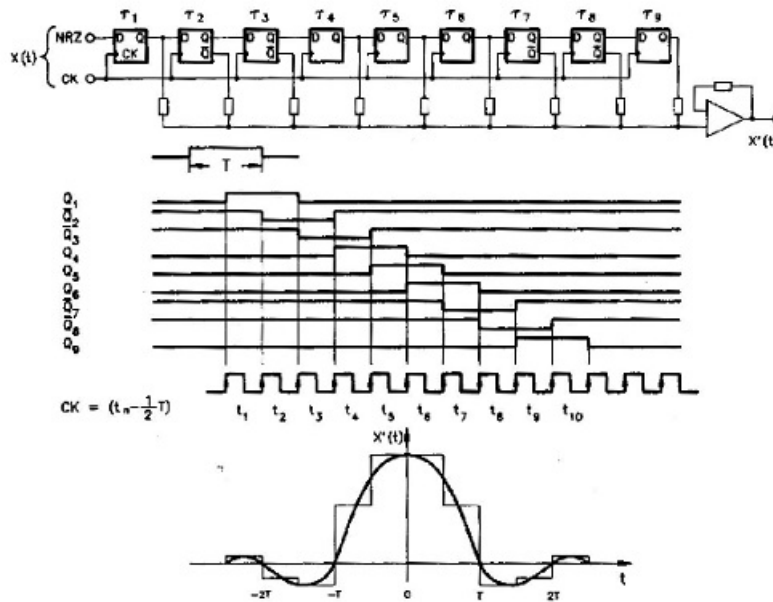


Fig.16.jpg

Il segnale d'ingresso $x(t)$ attraversa gli n elementi di ritardo (τ), che vengono pilotati da un clock a frequenza multipla di quella di simbolo. Si può vedere dai cronogrammi che le code del simbolo vengono forzate a passare per "0" in corrispondenza dei tempi del simbolo precedenti e successivi.

Il numero di gradini per ogni impulso è determinato dal numero di registri adoperati; poiché l'altezza dei gradini è proporzionale al valore delle resistenze, si comprende intuitivamente che, ammessa la simmetria della rete di pesatura, il segnale di uscita è simmetrico e quindi il filtro *introduce un ritardo di gruppo costante*.

Dal punto di vista spettrale, il contenuto armonico dell'onda a gradini è inferiore a quello dell'onda quadra; lo spettro trasmesso è pertanto meno esteso rispetto a quello dell'onda quadra d'ingresso e lo è tanto meno quanto maggiore è il numero dei gradini. Tuttavia, non è necessario usare lunghe catene di registri, poiché comunque lo spettro viene limitato anche dal filtro passa basso posto a valle del BTF, avente lo scopo di eliminare la periodicità del segnale filtrato.

Notiamo infine che il BTF offre elevata versatilità, in quanto attraverso un *opportuno dimensionamento della rete di pesatura e della cadenza di scorrimento dei registri*, possono essere realizzati filtri a qualunque frequenza con qualsivoglia sagomatura.

2.3 Circuiti modulatori a RF

Per le modulazioni più semplici (2PSK o 4PSK) vengono spesso utilizzati **modulatori a radiofrequenza**, per i vantaggi di semplicità implementativa offerti. La modulazione viene ottenuta mediante *variazioni del percorso dell'onda portante*; gli elementi che costituiscono un modulatore RF sono pertanto:

- una linea di lunghezza tale da creare, in base alla velocità di propagazione del segnale, lo sfasamento desiderato;
- un dispositivo munito di una porta di ingresso/uscita dove possa aver luogo l'allungamento del percorso del segnale;
- un interruttore elettronico comandato dal segnale-dati che, per l'alta frequenza, è ottenuto mediante l'impiego di un *diodo PIN* (tipo *p-Intrinseco-tipo n*).

Molto brevemente, questo diodo possiede una zona p ed una zona n fortemente drogate, separate da una zona a drogaggio bassissimo, tale che la concentrazione delle cariche intrinseche sia quella prevalente. Qualora lo si polarizzi con tensione inversa, la zona intrinseca viene svuotata dalle cariche mobili e il diodo oppone un'elevatissima resistenza (conducibilità nulla). In condizioni di polarizzazione diretta, nella zona intrinseca vengono iniettate cariche mobili dalle regioni adiacenti e la sua resistenza si abbassa (l'aumento delle cariche minoritarie incrementa la conducibilità del semiconduttore) presentando valori dipendenti con funzione quasi lineare dalla corrente diretta. Si comprende pertanto che il diodo PIN può venire usato come *resistenza variabile su comando elettronico* e, in condizioni estreme come il caso presente, può comportarsi come elemento di commutazione.

La seguente figura mostra un modulatore 2PSK a *variazione di percorso* con *circolatore a ferrite RF* (elemento che sarà discusso più avanti):

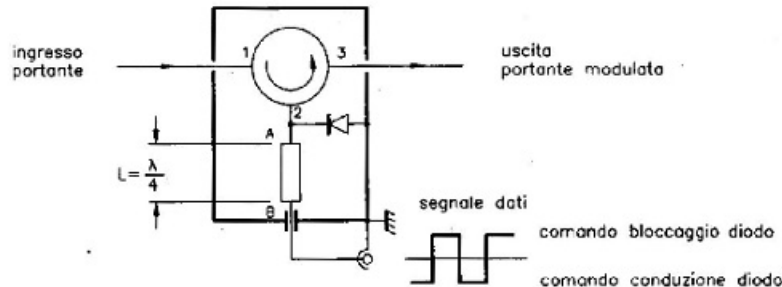


Fig.17.jpg

In corrispondenza del diodo (punto A), qualora questo sia in saturazione, si ha riflessione totale del segnale uscente dalla porta 2. Se invece il diodo è interdetto, il segnale viene riflesso dal cortocircuito che si trova in B, in corrispondenza della terminazione della linea. Fra le due condizioni *la differenza di percorso vale esattamente $2L$* , dove L è la distanza tra A e B. Appare quindi evidente che ponendo $L = \lambda / 4$, il salto di fase ottenibile è pari a 180° .

Un modulatore quadrifase si può ottenere ponendo in cascata due distinti modulatori a circolatore, come di seguito illustrato:

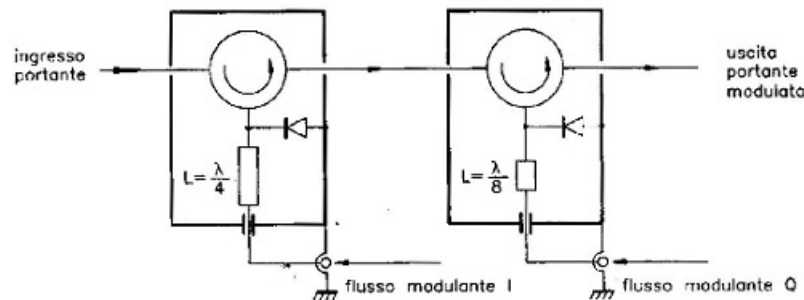


Fig.18.jpg

Il secondo modulatore è caratterizzato da una linea lunga $L = \lambda / 8$ ed è in grado di produrre salti di fase pari a 90° . Inviando ai diodi simboli binari paralleli trattati da un'opportuna logica di comando, si ottengono le quattro tipiche combinazioni di fase richieste dalla modulazione 4PSK.

Un modulatore a variazione di percorso è realizzabile anche mediante impiego di un **ibrido a 3 dB**; l'ibrido è sostanzialmente un dispositivo a quattro porte che può essere realizzato con struttura in guida d'onda o su microstrip con la tecnologia del film sottile:

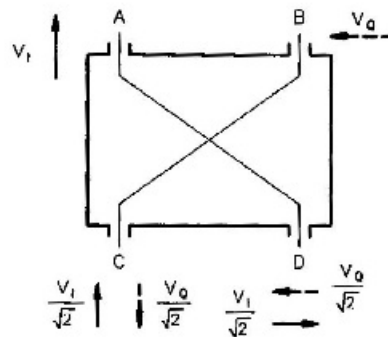


Fig.19.jpg

È costituito da due tratti di linea a RF mutuamente accoppiati, capaci di scambiarsi energia: la tensione V presente sulla porta d'ingresso A, viene suddivisa in due tensioni d'ampiezza $V/\sqrt{2}$ (quindi inferiore a 3 dB rispetto a V) su ciascuna delle due porte d'uscita C e D, con la particolarità che la fase in D è ruotata di 90° rispetto ai vettori in A e C.

Questo effetto si ha in ogni caso sulle tensioni trasferite alle uscite opposte alla porta d'ingresso, la quale può essere indifferentemente una delle quattro. Nessuna tensione invece si presenta sulla porta adiacente a quella d'ingresso (B nel caso di Fig.19), purché siano nulle le eventuali riflessioni sulle porte d'uscita.

Facciamo un esempio considerando la seguente illustrazione:

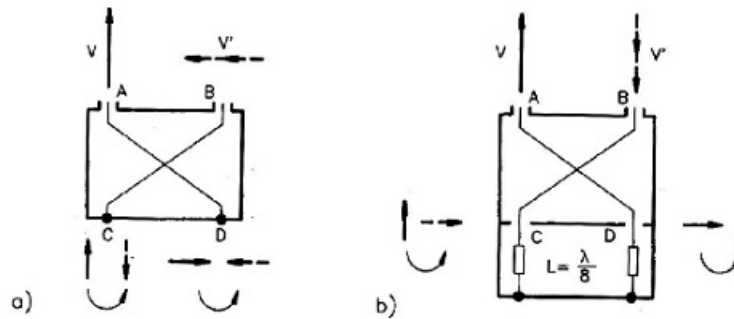


Fig.20.jpg

Supponiamo di creare sulle uscite C e D una riflessione totale (Fig.20 a)); poiché nel punto di riflessione i vettori vengono sfasati di 180° , si comprende facilmente che su B si troverà un vettore V' , somma dei contributi delle due riflessioni, che risultano in fase tra loro. V' si presenta ruotato di 90° rispetto a V , con modulo che, assumendo nulle le perdite all'interno dell'ibrido e nei punti di riflessione, è uguale a quello di V . Analogamente al caso di modulatore a circolatore, *se la riflessione viene spostata a valle di una linea di lunghezza L* (si faccia ancora riferimento alla Fig.20), *la fase in uscita subirà un ritardo corrispondente ad un percorso lungo $2L$* . In particolare, la fase di una uscita subirà una rotazione di 90° su una linea lunga $\lambda/8$, e di 180° su una lunga $\lambda/4$. Si può quindi ottenere un modulatore 4PSK impiegando due modulatori bifase capaci di realizzare rispettivamente salti di 90° e 180° , impieganti ciascuno come commutatori veloci una coppia di diodi PIN, come di seguito illustrato:

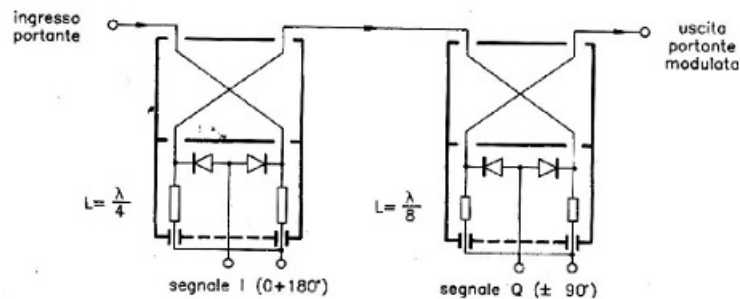


Fig.21.jpg

Questa configurazione di modulatore a RF fornisce prestazioni migliori rispetto a quello a circolatore, sia per la struttura bilanciata del primo, sia per le migliori caratteristiche di stabilità in temperatura dell'ibrido nei confronti del circolatore.

2.4 Modulatori QAM a frequenza intermedia

Le modulazioni multilivello di maggiore complessità vengono realizzate su circuiti a frequenza intermedia per ovvi motivi di ripetibilità e di minore criticità dei circuiti. Consideriamo il seguente schema esplicativo:

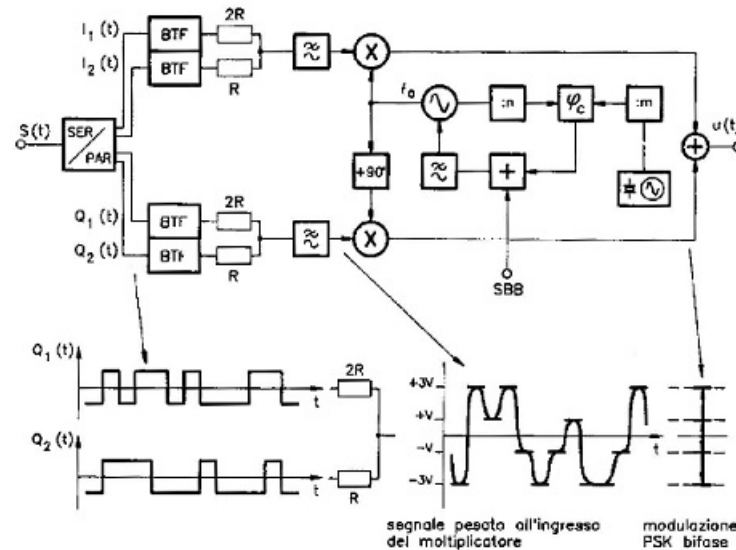


Fig.22.jpg

Il cuore del circuito è l'invertitore comandato, che realizza il modulatore elementare PSK; nei modulatori con configurazione in quadratura, questi sono due, alimentati ciascuno con fase appropriata da un unico oscillatore locale. La precisione e la stabilità della frequenza generata (dell'ordine di 1 parte su 10^6), vengono garantite da un oscillatore stabilizzato a quarzo o, più spesso, da un VCO stabilizzato con circuito PLL. In tal caso (riferendoci sempre alla Fig.22) esiste un secondo oscillatore controllato a quarzo, che genera una frequenza di riferimento a cui il VCO s'aggancia attraverso l'anello di controllo, in frequenza e fase. Il circuito così realizzato consente anche di *modulare il VCO in FM con segnali analogici di servizio*; la coesistenza dello spettro di tale modulazione con quella numerica è resa possibile dal fatto che in tale porzione di banda la densità spettrale del segnale numerico è bassa, il che limita in misura accettabile (per una sorta di compensazione) i rispettivi degradi delle due modulazioni.

Il segnale generato dall'oscillatore locale a FI possiede livello sufficiente per comandare la saturazione o l'interdizione dei diodi del modulatore, indipendentemente dai segnali modulanti $I(t)$ e $Q(t)$, che presentano livelli di tensione abbastanza piccoli da poter ritenere trascurabile il loro effetto sui diodi. Le quattro sequenze di bit $I_1(t)$, $I_2(t)$, $Q_1(t)$ e $Q_2(t)$ fornite dai circuiti in banda base, vengono filtrate e sagomate dai filtri digitali BTF, e con ciò perdono l'originario

formato ad onda rettangolare, presentandosi invece con una sagomatura più arrotondata. Su ciascun ramo, le due rispettive sequenze accedono ad un circuito costituito dalle resistenze R e $2R$, che *assegnano ai due segnali un peso rispettivamente di V e $3V$* ; quindi vengono sommate per ottenere il segnale a quattro livelli che va a modulare in fase la sottoportante relativa.

Sul circuito moltiplicatore, avviene poi l'interazione tra le due grandezze, che corrisponde a una modulazione PSK bifase tale da *mantenere inalterati i due livelli d'ampiezza su ciascuna delle due possibili fasi generate*. Sul secondo ramo del modulatore, tutto avviene con perfetta analogia, però sull'asse in quadratura; ed è solo sul sommatore d'uscita che avviene la composizione vettoriale dei due segnali in questione, da cui si genera la nota costellazione QAM.

Il circuito è completato da un'opportuna dotazione di radici d'allarme, che segnalano situazioni anomale quali mancanza dati all'ingresso o lungo il percorso della banda base, mancanza di potenza FI in uscita o irregolarità di funzionamento del PLL.

2.5 Conversione di frequenza e amplificazione a RF

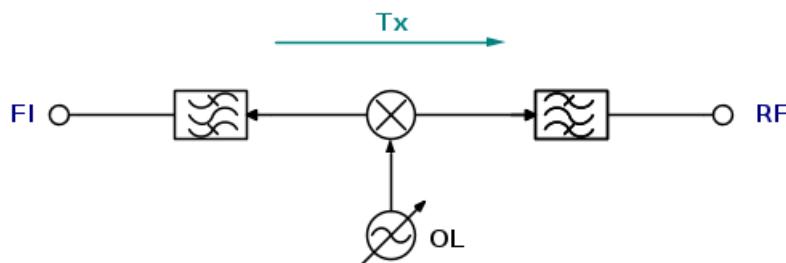
Nel caso in cui, come mostrato nel sottoparagrafo precedente, la modulazione sia effettuata su una portante a FI, occorre poi realizzare un ulteriore trattamento del segnale modulato, per *trasporto sul canale a RF* prescelto per la trasmissione.

Il circuito che adempie a tale funzione è il **convertitore IF/RF**, detto anche **up-converter**, in quanto traspone verso l'alto (in termini di banda) un segnale originariamente più basso di frequenza.

Esiste anche il processo perfettamente opposto, usato in ricezione, che effettua la trasposizione del segnale modulato a RF in un analogo segnale a FI; in tal caso il relativo circuito è il **convertitore RF/IF**, chiamato **down-converter**. Per adesso consideriamo soltanto il primo convertitore.

2.5.1 Up-converter

Il circuito convertitore è così schematizzabile:



Fondamentalmente esso è composto da un generatore a RF, un circuito moltiplicatore e un opportuno filtro di selezione. La trasposizione in frequenza avviene mediante *operazione lineare di moltiplicazione tra due frequenze, assimilabile ad una modulazione d'ampiezza*. Come detto più volte, il risultato finale vede la comparsa nello spettro di due componenti laterali, simmetricamente disposte attorno alla frequenza dell'oscillatore locale OL, corrispondenti alla somma e alla differenza tra le due frequenze che hanno partecipato all'operazione.

Nel caso dell'up-converter, il prodotto tra il segnale a FI (modulato) e quello fornito dall'OL, produce due bande laterali aventi rispettivamente frequenza $OL+FI$ (banda laterale superiore) e $OL-FI$ (banda laterale inferiore) e con contenuto spettrale pari alla somma degli spettri di FI e OL. Quest'ultima, ovviamente, non deve introdurre nulla che possa degradare il contenuto informativo portato da FI e perciò si pone particolare cura nel progettare l'oscillatore OL per ottenere la *massima purezza spettrale*. Consideriamo la seguente figura:

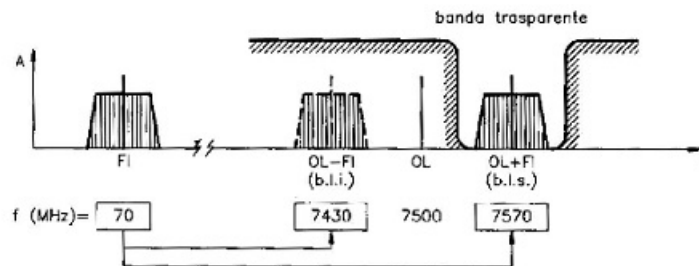


Fig.24.jpg

In questo esempio osserviamo che, volendo trasmettere a 7570 MHz un segnale modulato disponibile a FI (70 MHz), occorre che l'OL generi una frequenza pari a 7500 MHz e che il filtro a valle del moltiplicatore sia centrato a 7570 MHz, con selettività sufficiente per bloccare sia la banda laterale indesiderata (nel caso presente, quella inferiore centrata a 7430 MHz), sia la frequenza dell'OL che, avendo esaurito la sua funzione, non deve essere trasmessa dal trasmettitore (per ovvi motivi di inutile dispendio energetico). Notiamo inoltre che nulla cambierebbe qualora si decidesse di adottare la banda laterale inferiore, eliminando quella superiore. Si comprende che scegliendo opportunamente la f_{OL} e di conseguenza il relativo filtraggio, è possibile traslare lo spettro del segnale a FI in qualunque frequenza della gamma a RF.

Per quel che riguarda la *tecnica di conversione*, nel caso dei ponti radio a microonde, viene usata la *tecnologia a film sottile* per realizzare su un substrato dielettrico (*allumina, duroid*) un unico microcircuito che prende il nome di **convertitore bilanciato** di seguito schematizzato:

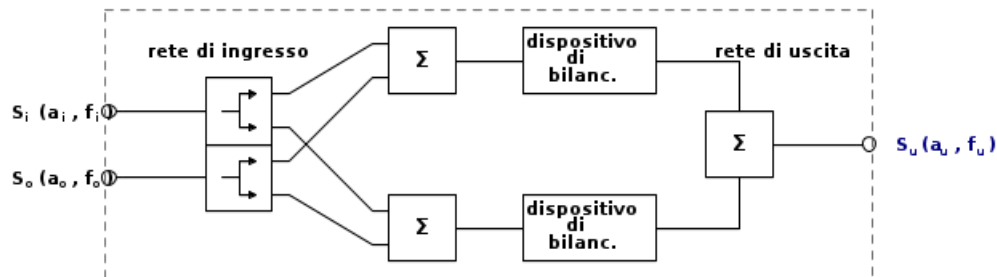


Fig.25

Questo circuito permette di migliorare il disaccoppiamento tra gli ingressi dei segnali a frequenza diversa e, soprattutto, consente la soppressione in uscita (tanto più quanto meglio è bilanciato il circuito) della riga spettrale dell'OL e i contributi di rumore che esso inevitabilmente porta con sé.

Nella conversione di trasmissione il *rendimento energetico* assume importanza rilevante in quanto la potenza dei segnali a microonde è una risorsa costosa che è opportuno non sprecare. Nel caso della conversione ad alto livello, ad esempio, l'uscita del convertitore coincide con l'uscita verso le antenne e, poiché in tale punto è richiesta una potenza di alcuni watt, il livello dei segnali all'ingresso del convertitore deve risultare *proporzionalmente più alto*, sicché sono alcuni watt che vengono persi nel processo di conversione. Nei primi ponti radio numerici, venivano impiegati dei convertitori parametrici impieganti dei varactor capaci di introdurre delle non linearità al fine di ridurre le perdite o addirittura introdurre un guadagno di potenza. La soluzione odierna consiste invece nell'effettuare la *conversione IF/RF a basso livello*, ovvero adottando circuiti e livelli di potenza analoghi a quelli di ricezione. Il segnale prodotto dalla conversione viene successivamente amplificato, come illustrato nel prossimo sottoparagrafo.

2.5.2 Amplificatori di potenza e tecniche di predistorsione

La potenza disponibile all'uscita di un modulatore a RF o di un convertitore IF/RF è dell'ordine di pochi milliwatt; occorre pertanto una successiva amplificazione per elevare la potenza al livello di trasmissione richiesto. La tecnica di **amplificazione diretta a RF** è ormai sufficientemente consolidata: impiegando *transistor MESFET* progettati per lavorare nel campo delle microonde, prodotti con la *tecnologia dell'arseniuro di gallio (GaAs)* (oggi sono in fase di sperimentazione degli amplificatori RF all'erbio, il cui impiego è già consolidato nei collegamenti in fibra ottica), è possibile ottenere le potenze necessarie alla trasmissione sui ponti radio.

Il dimensionamento di un amplificatore a RF, tuttavia, non è un problema banale. Richiede una particolare attenzione anzi, in quanto incide sulle caratteristiche di qualità e affidabilità del ponte stesso. Per quanto riguarda l'affidabilità, è intuitivo che il limitato rendimento della conversione DC/RF (corrente continua/

radiofrequenza) impone attenzione alla temperatura di funzionamento del MESFET. Infatti, solo una piccola parte della potenza assorbita dall'amplificatore viene trasformata in potenza utile a RF, tipicamente dell'ordine di qualche unità percentuale; il resto della potenza di alimentazione si trasforma in calore all'interno del dispositivo elettronico, raggiungendo valori molto elevati (anche maggiori ai 150°C). Occorre pertanto dissipare correttamente all'esterno tale eccesso di calore, che facilmente insidia la vita del dispositivo. Ciò si ottiene calcolando accuratamente la resistenza termica tra quest'ultimo e l'ambiente esterno, nonché dimensionando adeguatamente un heatsink verso l'aria ambiente, compatibilmente alle specifiche indicate nei datasheet dal costruttore.

Normalmente il parametro di riferimento per i FET di potenza è quello identificato con $P_{1\text{dB}}$ indicante la minima potenza d'uscita alla quale il guadagno dello stadio viene compromesso a 1 dB. Viene cioè intercettato un punto della caratteristica di trasferimento del dispositivo prossimo alla saturazione dove, per livelli crescenti di potenza in uscita, *il dispositivo inizia a perdere in linearità*. Da tale punto a salire è associata una distorsione d'ampiezza rapidamente crescente che, soprattutto per segnali contenenti modulazione d'ampiezza significativa, quali le modulazioni QAM, comporterebbero un consistente peggioramento della qualità di trasmissione in termini di interferenza intersimbolica e conseguente aumento del tasso d'errore.

Per garantire quella linearità tale da offrire bassa distorsione, anche in presenza di picchi di ampiezza di modulazione, occorre dimensionare la potenza d'uscita P_u ad un valore inferiore a $P_{1\text{dB}}$. La loro differenza $P_{bo} = P_{1\text{dB}} - P_u$ prende il nome di **back-off** dell'amplificatore.

Tuttavia è comprensibile come, all'aumentare del back-off, cresce il costo dell'amplificatore e se ne riduca il rendimento, dal momento che occorre scegliere componenti sovradimensionati e, di conseguenza, generosamente alimentati.

Si tollera pertanto un certo tasso di distorsione, *correggendolo* però per altra via, ovvero impiegando delle **tecniche di predistorsione**: il loro principio (siano esse in BB, a FI o RF) è quello di *introdurre a monte dell'amplificatore delle non linearità tali da compensare quelle dell'amplificatore stesso* (analogamente a quanto viene fatto negli OP AMP per compensare, ad esempio, le tensioni di offset). Queste non linearità si presentano tipicamente con due effetti:

1. non linearità tra potenza d'ingresso e potenza d'uscita, detta *conversione AM/FM*;
2. variazione di fase in uscita, non lineare in funzione della potenza di ingresso, detta *conversione AM/PM*.

La rete di predistorsione offre pertanto la capacità di generare caratteristiche ampiezza-potenza e fase-potenza tali da compensare quelle dell'amplificatore. Lo schema di principio tipico di un amplificatore a 7 GHz è quello di seguito riportato:

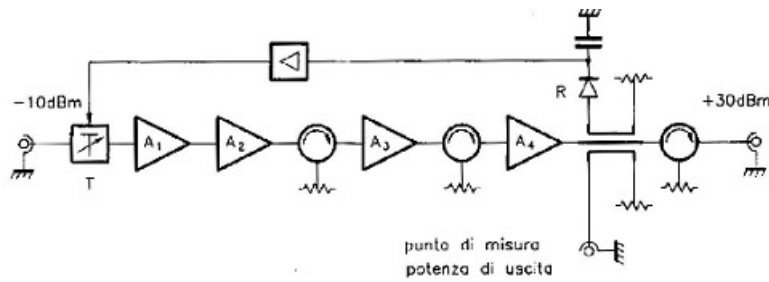


Fig.26.jpg

Il guadagno complessivo (40 dB) è ottenuto con quattro stadi attivi a GaAs, separati da isolatori in ferrite. All'uscita dell'ultimo stadio, un doppio accoppiatore direzionale provvede a fornire un punto di misura per il controllo della potenza d'uscita, nonché una tensione rivelata che controlla un attenuatore a diodo PIN presente all'ingresso del primo stadio. Con ciò si realizza un *circuito AGC (Automatic Gain Control)*, che permette di *mantenere costante la potenza d'uscita*, garantendo peraltro il necessario back-off in qualunque condizione operativa.

Il bilancio energetico di massima di un circuito del genere mostra che, a fronte di una potenza d'uscita di circa 1 mW, la potenza assorbita dall'alimentazione è di una trentina di watt; il rendimento dell'intero amplificatore è pertanto stimabile sul 3%.

3. Affasciamento delle portanti a RF

Il complesso di diramazione o **branching**, rappresenta un elemento importante ed interessante nella tecnica dei ponti radio a microonde: si tratta di un *dispositivo affasciatore* che in modo molto efficiente convoglia più frequenze su un solo sistema radiante (feeder + antenna), consentendo da un lato di sfruttarne vantaggiosamente le caratteristiche di grande larghezza di banda, dall'altro di evitare il moltiplicarsi di antenne ingombranti e costose. Un esempio di branching per un sistema 1+1 è quello mostrato nella seguente figura:

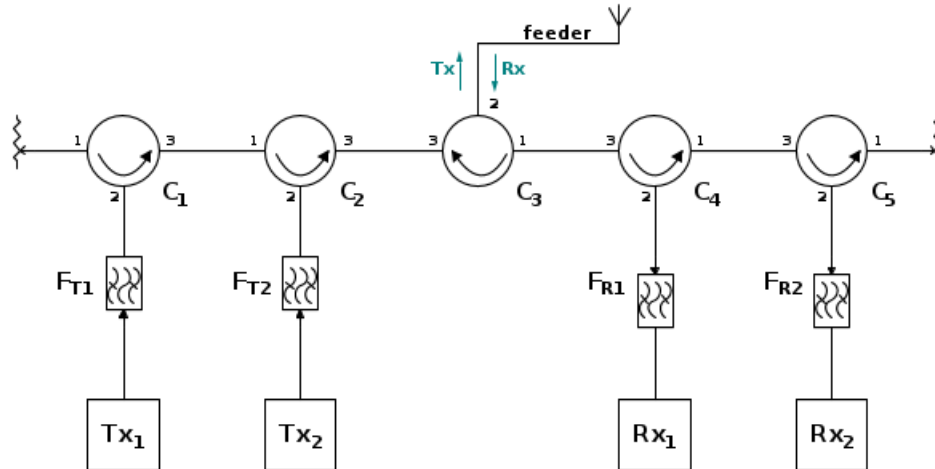


Fig.27

Gli elementi che lo costituiscono sono *filtri di canale* e circolatori a ferrite.

3.1 I filtri di canale

I **filtri di canale** schematizzati in Fig.27 sono di tipo passa banda; risultano sintonizzati sulle frequenze portanti dei rispettivi apparati trasmettenti e riceventi, e svolgono alcune tipiche funzioni, quali:

- per i ricevitori e trasmettitori a conversione, abbattano la frequenza immagine conseguente al processo di conversione in frequenza e in ogni caso provvedono ad eliminare segnali indesiderati fuori banda;
- per i trasmettitori numerici, partecipano alla sagomatura dello spettro trasmesso;
- possiedono inoltre la prerogativa di presentare, in corrispondenza della loro bocca d'accesso, un'impedenza teoricamente nulla per le frequenze esterne alla banda passante, le quali vengono pertanto *integralmente riflesse all'indietro*.

3.2 Il circolatore

Il **circolatore** è un dispositivo che sfrutta la proprietà giromagnetica di particolari materiali ceramici anisotropi (ferriti). Il campo magnetico alternato prodotto dall'onda incidente sulla ferrite, provoca negli elettroni di tale materiale un moto di precessione, simile a quello di una trottola. Sottoponendo la ferrite ad un campo magnetico continuo di entità tale da provocare la saturazione magnetica, si ottiene un *effetto di non reciprocità* nei riguardi della direzione di propagazione del campo elettromagnetico, in funzione della concordanza o discordanza dei campi magnetici che interagiscono nella ferrite.

Su tale principio si basano gli **isolatori direzionali per microonde**, realizzati disponendo all'interno di un tronchetto di guida d'onda una piastrina di ferrite che viene sottoposta al forte campo magnetico prodotto da un magnete permanente esterno alla guida. Il dispositivo così ottenuto (isolatore *uniline*) presenta *attenuazioni differenziate nei due sensi di propagazione dell'onda e.m.*: dell'ordine di pochi decimi di decibel nel senso favorevole e nettamente più elevati (> 25 dB) nell'altro senso. L'isolatore direzionale viene vantaggiosamente usato per proteggere le sorgenti dei segnali da possibili disadattamenti del carico.

Sullo stesso principio si basa il circolatore, con la differenza che esso dispone di tre porte, ciascuna delle quali presenta *bassa attenuazione di transito verso una delle due porte adiacenti e alta attenuazione nei confronti dell'altra*. Il circolatore è definito da un determinato *senso di circolazione del segnale* che dipende dalla polarizzazione del campo e.m. prodotto dalle espansioni polari, come mostrato di seguito:

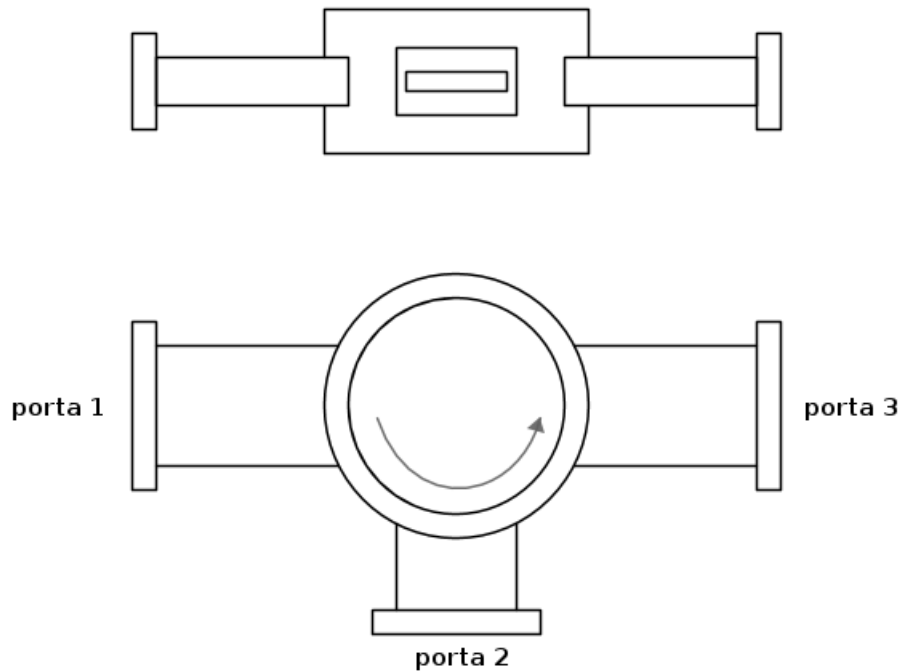


Fig.28

e indica la successione delle porte adiacenti nel senso di minore attenuazione. Osserviamo che un segnale entrante dalla porta 2 trova minima attenuazione procedendo nel senso della freccia ed esce dalla porta 3 attenuato di pochi decimi di decibel. Lo stesso segnale, misurato sulla porta 1, presenterà un'attenuazione molto elevata. Analogamente si comportano i segnali posti all'ingresso delle porte 2 e 3.

3.3 Il complesso di diramazione

Le proprietà del circolatore, di concerto a quelle dei filtri di canale, consentono di ottenere sui segnali radio il desiderato effetto di affasciamento, con minime perdite. Facendo riferimento alla Fig.26, si può seguire il percorso della portante emessa dal trasmettitore Tx2, attraverso C2 e C3, fino alla porta 2 di C3, a cui è connesso il feeder.

La portante emessa da Tx1, attraversato C1, si ritrova sulla porta 2 di C2; ma qui si trova il filtro FT2 che, offrendo a questa frequenza la sua zona oscura, la riflette integralmente verso la stessa porta 2 di C2. Da qui, il tragitto fino all'antenna è lo stesso seguito da Tx2. Identico ragionamento vale per le portanti di ricezione.

L'ultima porta di ciascuna serie di circolatori (1 di C1 e 1 di C5), è terminata su carico resistivo per dissipare segnali residui non utili al sistema radio.

Si comprende che la catena di coppie circolatore-filtro di canale si può estendere indefinitamente, con l'avvertenza che l'attenuazione di percorso cresce all'aumentare del numero di circolatori interposti tra l'ultimo della serie e quello centrale (detto *circolatore d'antenna*).

4. Conclusioni alla quarta parte

Con la quarta parte si chiude l'analisi delle tecniche e dei sistemi atti alla manipolazione dell'informazione da inviare sul ponte radio. La [quinta](#) e sesta parte verteranno pertanto su tutto quello che riguarda esclusivamente la ricezione e l'estrazione del messaggio a contenuto informativo. L'analisi sarà particolarmente dettagliata per quel che riguarda lo studio dell'interferenza intersimbolica, la sagomatura del canale e la codifica di linea *partial response*, utilizzata massicciamente dai ponti radio numerici, passando per il dimensionamento e la scelta del ricevitore ottimo (per la riduzione del tasso e della probabilità di errore), nonché cenni sulla rigenerazione del segnale numerico e la relativa discussione sui demodulatori adatti allo scopo.

Bibliografia

1. *Formazione specialistica e training on the job presso Telecom Italia S.p.A. (2007 - 2013);*
2. *Appunti, dispense e materiale didattico messo a disposizione nei corsi di Fondamenti Di Comunicazioni Elettriche e Campi Elettromagnetici tenuti presso la Facoltà di Ingegneria Elettronica dell'Università Degli Studi Di Palermo (2011 - 2013).*

3. G. CORAZZA, A. MANINDIETRI, C. MONTEBELLO: <<Circuiti a microonde>>. Ed. Patron - Bologna;

4. DAGLISH: <<Amplificatori per microonde a basso rumore>>. Calderini, Bologna.

Estratto da "<http://www.electroyou.it/mediawiki/index.php?title=UsersPages:Jordan20:le-onde-elettromagnetiche-come-mezzo-trasmissivo-ponti-radio-terrestri-4>"