

Il rumore nei circuiti e dispositivi per la radiofrequenza

Il rumore è un segnale casuale che si comporta a segnali deterministici, il cui valore, istante per istante, fluttua. Gli esperimenti sono casuali, ossia non è possibile dare una descrizione deterministica.

I disturbi non sono indotti solo dall'esterno di un dispositivo, ma possono anche essere introdotti internamente a esso: si può dimostrare che il rumore è una caratteristica fisica non eliminabile da un segnale. In altre parole, all'interno di un dispositivo c'è un qualcosa che genera rumore.

Il rumore è un processo a media nulla, ma non per questo possiamo ignorarlo: il problema in effetti è che il valore quadratico medio del processo non è nullo; dato un segnale:

$$x(t) = x_0(t) + \underbrace{S(t)}_{\text{questo è la parte di rumore!}}$$

Se la media $\langle S(t) \rangle$ è nulla, si ha che:

$$\langle S^2(t) \rangle \neq 0$$

Dunque, la potenza media, che è proporzionale al quadrato del processo, non è nulla; essa si dice "potenza di rumore".

Ovviamente si tratta di un valore molto piccolo, ma ciò non cambia il seguente fatto: per un amplificatore radio, a bassa potenza, una certa componente della potenza di uscita dipende dalla componente rumorosa; se l'ingresso è nullo, la potenza di uscita non sarà nulla. Ciò comporta anche limitazioni alla dinamica dell'amplificatore: se l'ampiezza del segnale di ingresso è costante

talile con il rumore, allora non si può discriminare in uscita cosa sia rumore e cosa segnale. Con segnali troppo piccoli, la potenza di uscita è dovuta al rumore, e ciò danneggia fortemente gli stadi venitori in un sistema di telecomunicazioni.

Al fine di minimizzare questo rumore, è necessario introdurre una caratterizzazione dei dispositivi sotto il punto di vista del rumore.

La caratterizzazione è ovviamente in termini statistici, dal momento che si parla di processi stocastici (processi casuali variabili nel tempo). Una potenza è la caratterizzazione in termini di valor medio (nel tempo o di insieme):

- Per media nel tempo si intende, dato uno specifico dei segnali, farne la media per ogni istante di tempo;
- Per media di insieme si intende, fissato un certo t_0 la media di tutte le realizzazioni (valore atteso, E)

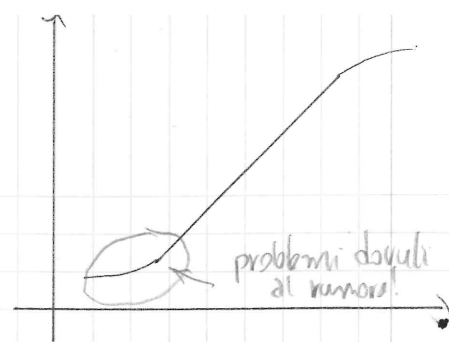
Quando le due medie non coincidono, si parla di "segnali ergodici". Considereremo sempre segnali di questo tipo.

Quello di più interessante è la media del processo per sé stesso: la funzione di autocorrelazione del processo:

$$R_{x_m x_n}(t, \tau) = \langle x_m(t) x_n^*(t+\tau) \rangle \quad m, n = 1, 2 \quad [o \quad m, n = 1]$$

Se due processi son diversi ($m, n = 1, 2$), si parla di "funzione di mutua correlazione" tra due segnali.

Si definiscono queste funzioni tra due tempi diversi, ma se il processo è stazionario si ha dipendenza dal solo τ , ossia dalla sola differenza dei due tempi.



Spesso si parla degli spettri di correlazione: $\{F\} \{R\}$

Se $\tau=0$, si considera in sostanza il valor medio quadratico del processo (ovvero dunque la potenza media).

Per i processi si approssimano a spettri costanti, ossia "rumori bianchi"; nella realtà ciò non esiste, poiché se no si avrebbe potenza infinita.

Nel rumore bianco, solo per $\tau=0$ il processo è correlato (si parla di "processo senza memoria"); esso sarà dunque variabile con grande rapidità.

Rappresentazione del rumore.

Dato un bipolo, il suo modello equivalente rumoroso si può ricavare introducendo un generatore di rumore equivalente.

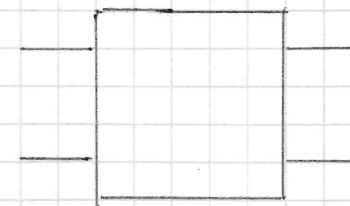
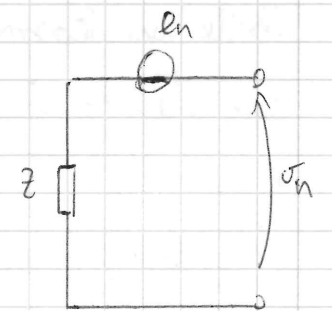
Lo spettro di questo generatore sarà una cosa di questo tipo: dato che:

$$\sigma = Z_i$$

$$\rightarrow S_{\sigma} = |Z|^2 S_i$$

Esiste ovviamente anche un equivalente Norton per la rappresentazione del circuito rumoroso.

Questo discorso può essere esteso a un 2-porte



Anche in questo caso i due generatori vanno rappresentati in termini di spettri: nella fattispecie, si avranno due spettri di autocorrelazione e uno di mutua correlazione (sia che si vogliono cedere termini su aperto sia che si vogliono cedere correnti in corto); al fine di ricavare i valori, però, serve conoscere le cause fisiche e i modelli del rumore, dunque si dovrebbe ripartire dalla fisica del dispositivo, cosa che non faremo.

Rumore termico

Esiste un teorema noto detto "teorema di fluttuazione-dissipazione": esso afferma che se un sistema è dissipativo, allora esso fluttua, e viceversa un sistema può fluttuare poiché scambia energia con l'ambiente esterno, e dunque dissipa.

Si può dire che ogni sistema è in equilibrio termico con l'ambiente, e ciò è possibile poiché esso è in grado di dissipare. Per esempio, un'ipotetica induttanza ideale, non potendo dissipare energia, non può fluttuare: solo ciò che è passivo può.

Per quanto riguarda il rumore termico,

$$S_{in} = 4k_B T R \quad [\text{spettro dovuto al rumore termico}]$$

Purtroppo ciò vale per i soli bipoli: nei FET o nei diodi, il rumore termico è solo una parte del rumore!

In pratica si hanno dei modelli, giustificabili fisicamente, dapprima ricavati per FET a canale lungo, ma poi "corretti", in modo da ottenere modelli "ad-hoc" per FET a microonde.

Modello PRC - modello a 2 e a 1 temperatura

Il modello PRC è il più elaborato, o contiene i sette parametri

intrinseci più quelli estrinseci; oltre al modello "tradizionale", si aggiunge un generatore di rumore in parallelo all'ingresso e uno in parallelo all'uscita (ign e ion). Von der Ziel dimostrò che:

$$S_{ind} = 4k_B T_0 g_m P \quad [\text{in realtà } P \text{ è un termine correttivo}]$$

ciò che capita è che la carica nel canale fluttua, come anche fluttua la mobilità; ciò darà luogo a un'equivalente fluttuazione totale della corrente al drain.

Come la corrente di rumore di ^{drain}gate fluttua, così anche quella di gate, seguendo la C_{gs} :

$$S_{ing} = 4k_B T_0 \frac{\omega^2 C_{gs}^2}{g_m} R \quad [R \text{ è un altro coefficiente di correzione}]$$

Secondo Von der Ziel esiste inoltre una correlazione "unitaria" tra le due correnti, ma nel modello più raffinato si ha:

$$S_{ing,ind} = jC \sqrt{S_{ind} S_{ing}} \quad [C \text{ parametro}]$$

I coefficienti P , R , C , che danno il nome al modello PRC, sono coefficienti correttivi per tener conto degli effetti di canale corto, che nei modelli di Von der Ziel non erano in conto.

Se le cose sono fatte bene, i tre parametri sono indipendenti dalla frequenza (e possono essere ottenuti o da una simulazione fisica o mediante misura); del momento che le difficoltà tecniche sono molte, questo modello viene considerato empirico.

Sperimentalmente si è visto che le formule si possono tranquillamente esprimere, in modo del tutto equivalente, usando come parametri le temperature di gate e di drain, T_g e T_d .

In questo caso, si parla di "modello a 2 temperature"; spesso, però, $T_G = T_0$ (T_0 è la temperatura ambiente), dunque si diminuisce il contributo di T_G , ottenendo il "modello a 1 temperatura".

Spesso, nei documenti, si trova direttamente il modello PRC.

Altra di rumore

Spesso si usa, come parametro, la "cifra di rumore"; essa, in realtà, è legata agli spettri, dunque ce ne occuperemo.

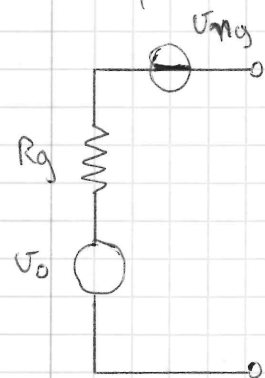
Definiamo in questo modo la cifra di rumore (noise figure):

$$NF = \frac{P_{nd,L}}{P_{nd,L}'}$$

Ossia, come il rapporto della potenza su di un carico Z_L e della potenza (o meglio, della densità di potenza) supposta che il 2-porte (il FET) non sia rumoroso; ciò indica la qualità del nostro circuito.



Una nota: se il generatore vuole erogare all'esterno, deve avere un'impedenza interna dotata di parte reale non nulla; questa, però, a sua volta, introdurrà del rumore termico.



Il FET dunque lavorerà avendo già in ingresso del rumore; poi il FET ne aggiungerà altro di suo, avendo sul carico una situazione ancora peggiore.

Tutto il fatto che il generatore introduce rumore, vogliamo capire

quanto il FET ne aggiunge. Data la potenza di rumore sul carico, vedo quanta era della sorgente, e così ricavo quella del FET.

Si può dire che:

$$NF = 1 + \frac{P_{nd,L}}{P_{nd,L}'}$$

Dove P' è la potenza che va sul carico (di rumore) dovuta solo al generatore, e P solo al 2-porte; si può vedere che, nel caso che il 2-porte è ideale (non introduce rumore), $NF=1$: meno di così non si riesce a fare. A volte si esprime in dB.

Ciò che però può essere importante negli ingegni queste potenze, che sono disponibili, moltiplicando per il guadagno disponibile G_{disp} le potenze di ingresso:

$$\frac{P_{nd,L}}{P_{nd,L}'} = \frac{G_{disp} P_{nd,in}}{G_{disp} P_{nd,in}'} = \frac{P_{nd,in}}{\frac{4k_B T R_g}{L R_g}} = \frac{P_{nd,in}}{k_B T} \quad \left[\begin{array}{l} \text{La potenza di rumore dovuta} \\ \text{alla sorgente è rumore termico!} \end{array} \right]$$

$$\rightarrow NF = 1 + \frac{P_{nd,in}}{k_B T}$$

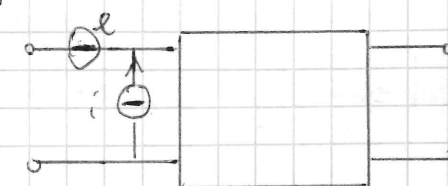
Vogliamo ora capire quanto il FET vi ed aggiunge. A volte ciò si esprime in termini di "temperatura equivalente": si tratta di una temperatura fittizia che il FET dovrebbe avere per produrre un rumore equivalente:

$$T_n = T_0 (NF - 1)$$

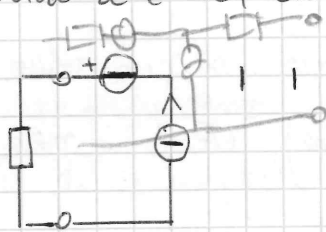
Ciò non ha significato fisico.

Si può chiedere, a partire dal modello PRC, P'' ?

Per prima cosa, a partire dal modello prima visto con un rumore in ingresso e



uno di usata, si deve riportare tutto in ingresso, dal momento che vogliamo calcolare p_{in} ; si trovano, nel circuito equivalente "tutto in ingresso", un generatore di corrente in parallelo e uno di tensione in serie. "e" e "i" sono tra loro correlati; ciò che si può dimostrare è che esiste una rappresentazione in cui i generatori sono indipendenti, a patto di introdurre due impedenze z_c e $-z_c$, che possono per l'appunto scouplare i generatori.



Come posso calcolare la cifra di rumore?

Considero, per ipotesi, che:

$$S_{enc} = 4 k_B T r_n$$

$$S_{inc} = 4 k_B T g_n$$

Voglio calcolare $\bar{v}$; la cifra di rumore sarà lo spettro S_{ov} ma, poiché z_c e $-z_c$ eliminano i propri contributi, rimane solo la R_g .

Ho che:

$$v = e_{ng} - e_{nc} + i(z_g + z_c)$$

$$\rightarrow S_{ov} = \langle v v^* \rangle = [e_{ng} - e_{nc} + (z_g + z_c)i] [e_{ng}^* - e_{nc}^* + (z_g + z_c)^* i^*] =$$

$$= S_{eng} + S_{enc} + S_{eng} + S_{enc} + S_{eng} + S_{enc} + |z_g + z_c|^2 S_i$$

$$= 4 k_B T R_g + 4 k_B T R_n + |z_g + z_c|^2 4 k_B T g_n$$

In fatti, FET e generatore sono correlati, come anche i_{nc} e e_{nc} (dato le impedenze z_c e $-z_c$).

Se dunque divido tutto per S_{eng} , il quale vale $4 k_B T R_g$:

$$NF = \frac{4 k_B T R_g + 4 k_B T R_n + |z_c + z_g|^2 4 k_B T g_n}{4 k_B T R_g} =$$

$$= 1 + \frac{r_n}{R_g} + \frac{g_n}{R_g} |z_c + z_g|^2$$

Servono dunque 4 parametri per poter calcolare la cifra di rumore.

Variando il generatore, la cifra di rumore a un certo punto assume un valore minimo: NF_{min} . Esso si ha quando:

$$R_{g0} = \sqrt{\frac{r_n}{g_n} + R_c^2} \quad \& \quad X_{g0} = -X_c$$

Quando questa condizione è verificata,

$$NF_{min} = 1 + 2 g_n R_c + 2 g_n \sqrt{\frac{r_n}{g_n} + R_c^2}$$

Dunque, il valore di NF_{min} dipende sostanzialmente dei parametri del FET. Data questa definizione, si può esprimere NF come:

$$NF = NF_{min} + \frac{4 g_n}{R_g} \frac{|z_g - z_{g0}|^2}{|z_{g0}|^2}$$

Questa formula ci dice, al variare del generatore, di quanto ci spostiamo dal valore ottimo.

Del punto di vista sperimentale di solito si usa questa formula, che ancora una volta contiene 4 parametri da calcolare/minimizzare.

Le sono formule che permettono di passare dal modello PRC a queste grandezze, ricadendoci dunque al rumore. Il modello è più utile alla finca, queste formule alle misure (ottenute variando z_g e minimando).

Questa caratterizzazione è interessante perché se vedo a rappresentare i Z_g per cui la cifra di rumore assume sempre lo stesso valore (ad esempio 3 dB), nel dominio delle riflettanze, il luogo dei punti per cui la cifra di rumore è costante è un cerchio. Anche in questo caso, si potrà identificare un punto per cui la cifra di rumore è quella minima. Si noti che nel piano Γ_a è anche possibile disegnare dei cerchi a guadagno costante (quello che ci interessa è il guadagno disponibile, del momento che si parla di Γ_a).

Si può vedere che, disegnando assieme i cerchi a guadagno costante e quelli a NF costante, posso decidere allo stesso tempo quale deve essere la cifra di rumore e quanto il guadagno.

Si deve a questo punto affrontare una scelta: i punti a guadagno massimo e a cifra di rumore minima non sono assolutamente coincidenti!

Volendo minimizzare NF, andrò da scegliere il suddetto punto; il guadagno per questo punto è detto "guadagno associato", ed è di solito molto piccolo, ma ciò non ci piace. Tattenden di solito di "primi studi", si vorrebbe un'amplificazione elevata.

Si dovrà richiedere un trade-off: non si progetterà nei punti ideali, ma si dovranno accettare 0,5 dB in più o simili, per ottenere il guadagno.

Si noti che "guadagno" e "adattamento" vanno di pari passo: progettare a minimo rumore implica sopportare l'ingresso, ottenendo un'elevata potenza riflessa: se il ROS è molto elevato, può essere che io cerchi troppo una linea, dal momento che tanta riflessione

troppa potenza.

Si ricordi che:

$$NF_{min} \approx 1 + 2K_e \frac{f}{f_T} \sqrt{g_m(R_g + R_s) + K_e}$$

Più si sale con la frequenza, più il rumore aumenta.

Questa formula, ricavata dopo vari "giri" dal modello PRC, è anche pensabile come un'evoluzione della "storica" formula di Fiksel: essa fu sviluppata prima della nascita del modello PRC, sulla base del modello fino del FET.

K_e vale circa 0,1 per gli HEMT, e circa 0,3 per i MESFET; ciò suggerisce che gli HEMT siano, per il rumore, meglio dei MESFET, a parità degli altri parametri.

Allo stato attuale gli HEMT van anche bene per le medie potenze, e sono sempre più usati.

Pensando alla linea generata dai FET, possiamo ricordare che g_m dipende dal punto di lavoro, da I_D : più bassa è I_D , più lo sarà g_m , ma, per questo il rumore migliora, il guadagno purtroppo peggiora.

Anche qui l'HEMT è meglio: la massima transconduttanza non richiede correnti elevatissime, dunque si può stare più lontani dalla zona ad alto rumore.

NR2

Amplificatori lineari

In questa trattazione considereremo solo due tipi di amplificatori:

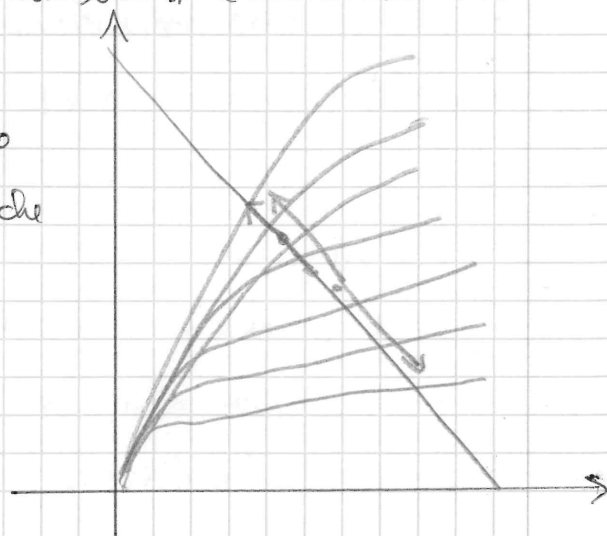
- ad alto guadagno
- a minimo rumore.

Al fine di realizzare il progetto, si dovrà far fronte a un certo numero di specifiche:

- Sul guadagno: si vuole che un amplificatore amplifichi almeno un tot. Ciò ci porta a scegliere topologie a singolo stadio o con più stadi in cascata.
- Sulla banda: a seconda della larghezza di banda richiesta, 2%, 5%, un'ottava, o molto largo, potrà usare amplificatori ad anello aperto o chiuso (in questo caso, prestando attenzione alla stabilità).
- Sulla tecnologia (ibrida/mondicica).
- Sulle tecniche di progetto: amplificatori ibridati (sempre adattati energeticamente), o distribuiti (a microonde, sono amplificatori meno soggetti ai vincoli tra banda e guadagno).

Prima di tutto, si deve scegliere il dispositivo attivo da usare (FET o HBT), che deve essere in grado di soddisfare abbondantemente le specifiche.

Dato il FET, sarà scegliere un punto di lavoro; si deve tener conto, oltre che del guadagno, della dinamica: se si sceglie una gm troppo elevata si alza troppo la corrente, la "ordinata", dunque

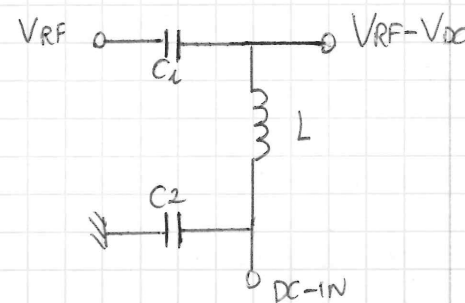


con un ingombro molto minore si arriva molto prima in saturazione, mandando tutta la potenza alle armoniche.

Si deve costruire dunque la rete di stabilizzazione, mediante procedimenti precedentemente introdotti.

Al fine di impostare Γ_g e Γ_L poi, si usano delle reti di adattamento, in modo da "far vedere" al FET le impedenze che vogliamo noi.

Per quanto riguarda le reti di polarizzazione, spenderemo qualche parola in più: si hanno reti note come "bias T". Si ha qualcosa di questo genere: a parametri concentrati, si fa in modo da disaccoppiare parte in DC e parte a RF: i due segnali in uscita devono essere sommati, ma un segnale non deve modulare l'altro! Devono essere isolati tra loro! E, soprattutto, la DC non deve assolutamente concorre al circuito RF.



C_1 fa disaccoppiamento tra la parte DC e quella RF; idealmente si vuole un corto circuito tra la DC e dove deve andare, e ciò è fatto con l'induttanza; C_2 garantisce ulteriormente il fatto che la RF veda un corto, "rafforzando" il filtro.

Nonostante ciò, si uscirà di avere disturbi della DC a quello decade di kHz, e bisogna starne attenti, perché "toccano" la RF.

Esiste una versione "distribuita" del circuito, che permette di evitare l'induttore, lasciando però i condensatori, aggiungendo i tratti di lunghezza elettrica $\lambda/4$.

Questa tecnica funziona a banda stretta, ma qui il tratto agirà molto meglio da "aperto" per la RF, poiché non si hanno le perdite tipiche degli induttori reali.

Di solito nel circuito ibrido si usa la rete distribuita, nel monolitico il circuito concentrato, ma se nel bias-T si deve avere tanta potenza, allora la soluzione distribuita può essere usata anche nel caso monolitico.

Per portare l'alimentazione si userà un filo; esso può essere dimensionato in modo da stabilire l'induttanza "wire-bond", ottenendo dunque anche il filtro, assieme alla DC.

Esistono componenti che realizzano le bias-T, molto precise ed ampio spettro, ma, essendo ingombranti, si usano solo in laboratorio, non nei layout.

Stabilizzazione

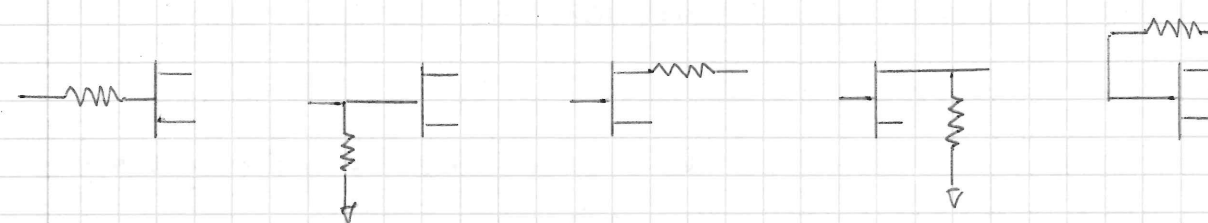
Bisogna verificare se il sistema è stabile sia alla frequenza di lavoro, sia a tutte le altre frequenze: da 0 Hz a qualche decina di GHz.

Si procede in questa maniera:

- fino a poco prima della frequenza di lavoro si cerca di rendere il dispositivo incondizionatamente stabile, ossia stabile per ogni carico o generatore.
- in banda si può o lavorare imponendo dei carichi tali da essere di certo nella zona stabile, o stabilizzare il dispositivo anche alla frequenza di lavoro, in modo che progettare con questo sia più "sicuro". Da un lato si perde il guadagno, ma dall'altro si può garantire la realizzabilità dell'adattamento coniugato, ottenendo

riflessioni molto basse. Una volta si tendeva a fare progetti senza stabilizzatore in banda, oggi è invece il trend che va per la maggiore: si stabilizza, guardando K o uno dei g_1 .

Per stabilizzare occorre introdurre dissipazioni nell'anello; ciò si fa introducendo una resistenza, dove passa la RF:



Se il circuito è ad anello aperto si tende a escludere la soluzione "a ponte", poiché introduce un loop. Questa si usa soprattutto ad anello chiuso.

Dal momento che si deve dissipare tanto in bassa frequenza e poco a frequenza più alta, si deve introdurre o una capacità in parallelo alla resistenza, o un'induttanza in serie ad essa.

Si preferisce mettere blocchi in serie, in modo da evitare di fare dei via-hole; questo blocco RC si mette di solito al gate, poiché nel drain si ha molta più corrente, dunque più dissipazione, più calore. Se però devo realizzare un LNA, dal momento che un resistore genera del rumore, si evita di metterlo sull'ingresso, usando per questi casi il drain.

Nel caso retroazionato, si usa di solito l'induttanza, e pure una capacità in serie, in modo da realizzare un disaccoppiamento tra le continue nel gate e nel drain.

Di solito non si gioca nel source, perché esso viene fornito già col-

legato a mano; mettere elementi di induttanza introduce reazioni, per questo è evitabile.

Per il progetto si userà il modello, chiedendo che il circuito finale sia molto stabile, in modo da sentirsi "sicuri" anche nel layout.

Adattamento

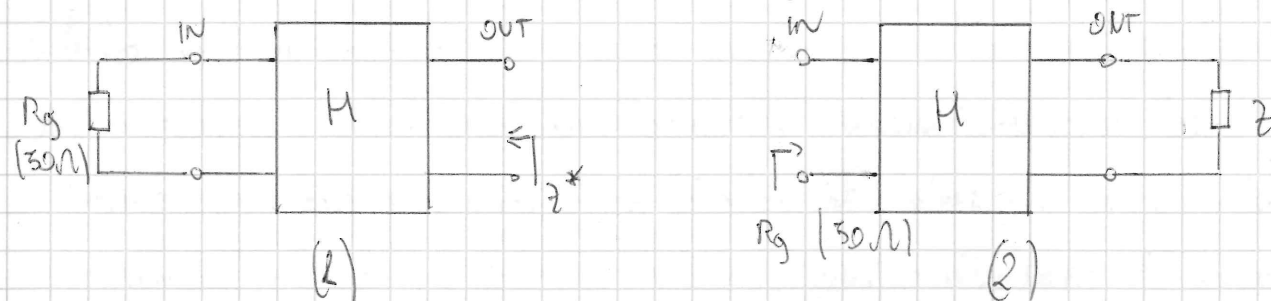
Per realizzare l'adattamento è necessario realizzare trasformazioni di impedenza: sia il generatore sia il carico devono essere dimensionati o comunque "visti", in modo tale da massimizzare la potenza erogata dal FET. Volendo per esempio un amplificatore a guadagno massimo, Z_G' deve essere uguale all'impedenza ottima di guadagno, e idem il carico:

$$Z_G' = Z_{Gopt} \quad ; \quad Z_L' = Z_{Lopt}$$

Volendo progettare a banda stretta, ci sono:

- reti puramente reattive (RMA): Reactive Matched Amplifier
- reti con elementi dissipativi (LMA): Lossy Matched Amplifier [a banda più larga ma rumorosi]

Si hanno in sostanza due possibilità di agire:



In entrambi i casi, IN è la porta alla quale si collega il generatore, OUT quella a cui si collega il dispositivo attivo, il FET; supponiamo che il nostro generatore abbia resistenza interna pari a 50Ω, e il FET resistenza pari a Z.

Lo si rappresentano queste due figure?

- 1) Nel caso 1 si chiede di far vedere, in uscita dalla rete di adattamento, Z^* , a partire dal 50Ω: tutta la potenza in OUT deve essere inghiottita dal FET, quando in ingresso ci sono 50Ω.
- 2) Nel caso 2, conicendo "a destra", "OUT", ora attaccando il FET a OUT, dunque a Z, si chiede di vedere in "IN" il 50Ω, in modo da adattare il generatore, dunque che la potenza del generatore venga interamente inghiottita dalla rete di adattamento.

Entrambe le condizioni andranno soddisfatte: io vorrei massimizzare il trasporto di potenza sia tra generatore e rete, sia tra rete e FET; si ha però un caso particolare: se la rete è puramente reattiva, una volta che il generatore ha erogato tutta la potenza disponibile, essendo la rete reattiva, essa non può assorbire potenza, dunque la darà al FET. Automaticamente dunque è soddisfatta anche l'altra condizione.

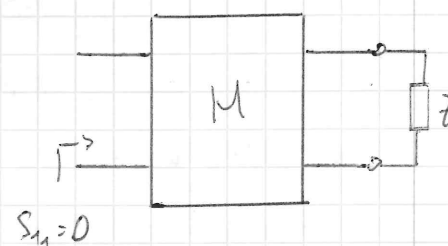
Se la rete è con perdite, resistiva, allora può anche succedere che la rete sia adattata da una parte e disadattata dall'altra, perché si può avere della potenza riflessa, che però viene dissipata dalle resistenze, facendo vedere adattamento dalle altre parti.

Per adattare dunque posso soddisfare una delle condizioni, e automaticamente, usando reti reattive, l'adattamento sarà realizzato.

Se per qualche motivo si può tollerare un disadattamento tra OUT e FET, si può adattare anche il solo generatore, mediante una rete resistiva, che "mangia" le riflessioni. Si dimensiona la rete in modo da "mangiare" queste riflessioni.

ciò che si realizza di solito con reti reattive, dunque a banda stretta, è la condizione 2, ma vedere il 50Ω a partire da Z , perché è più semplice realizzare il 50Ω che l'impedenza complessa.

Il problema sarà questo: in uscita della rete M , vogliamo vedere $S_{11}=0$, dunque vicino a -20 dB o -30 dB (quindi dei buoni adattamenti).



Z nel caso dell'amplificatore a guadagno massimo è l'impedenza ottima del FET; può essere, a seconda che si stia adattando al generatore o al carico, Z_{opt}^* o Z_{opt} .

Reti di adattamento

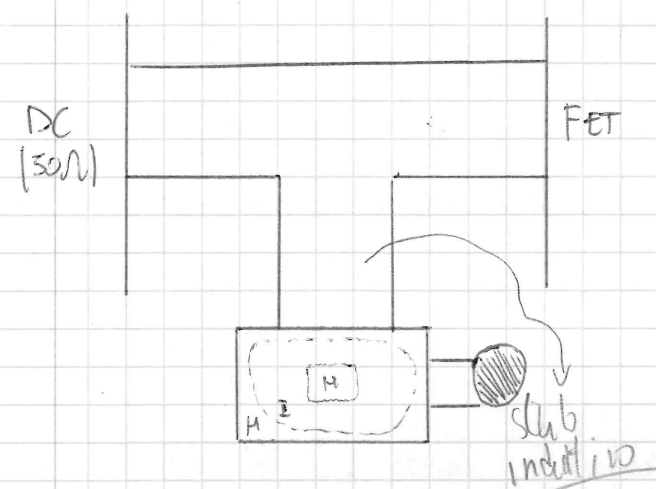
Al solito esistono due tipi di soluzioni: a parametri concentrati (soprattutto usate nei sistemi monolitici) o a parametri distribuiti. Ciò che serve, in sostanza, è una rete che sia in grado di far vedere una certa impedenza/admittanza. A seconda della condizione ($G_{in} < 1$ o $R_{in} < 1$) si userà una rete o l'altra (L diretto o rovesciato). Esistono formule, ma spesso si finisce per usare l'ottimizzatore.

La soluzione a parametri distribuiti è simile: una rete con una linea e uno stub parallelo. Il fatto di avere due linee di trasmissione permette di avere quattro gradi di libertà: impedenze caratteristiche e lunghezze delle linee. Ciò che si fa è fissare le impedenze caratteristiche, per giocare solo sulle lunghezze delle linee. Volendo si potrebbe usare una singola linea, e giocare anche sulla Z_{01} , ma è un po' infelice del momento che impos-

tere un valore preciso di Z_{01} è difficile, e sbagliando anche di poco il progetto non funziona. Usando due linee è possibile fare un po' di "trimming" e mettere a posto il circuito.

Usare stub chiusi in aperto è molto usato, per questo si dice uno svantaggio: essi tendono a invecchiare.

Si hanno soluzioni "ottimizzate", del tipo: al posto di mettere lo



stub in aperto, lo mette in corto, ma ora che ci sono molto una piastrina MM, una in ceram-
sare Metello Isolante Metello, dunque collego una piastrina alla pista in modo da mettere un condensatore "in serie" e il buco

verso massa. Se collego al condensatore un filo con la DC, posso realizzare con la stessa rete adattamento e alimentazione. Lo stub stesso farà da induttanza (servirebbe poi un disaccoppiamento per la DC dei 50Ω).

A questo punto si hanno tutti gli "ingredienti", e il progetto del layout sembra facile, ma non è proprio così: quando si tende a ottimizzare il layout si possono formare angoli, "T", "croci", o altri elementi. Ogni spigolo inedito, introduce discontinuità alla propagazione, dunque vanno tagliati, mediante un "bend chamfered": smussature. Altrimenti potrebbe permettere di fare l'angolo "tondo", che è ancora meglio.

Stessa cosa vale per le giunzioni a T tra una linea e lo stub, o un incrocio.

Altra non idealità è il cambio

di lunghezza delle linee,

che si "gradualizza" mediante un tapering.

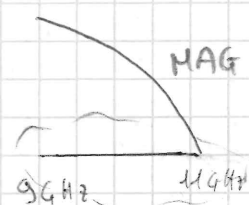
Altro accorgimento spesso usato è l'aggiunta di linee $\lambda/2$, per distanziare due stub ed evitare per esempio che si accoppino.

Queste cose non possono essere progettate a priori, ma valte solo dopo aver visto un primo layout, che deve essere modificato, per esser corretto.

Amplificatori a banda larga

Quando si vogliono progettare amplificatori a banda un po' più larga (anche non larghissima però), per esempio da 9 a 11 GHz con frequenza centrale 10 GHz, si deve fare qualcosa di diverso: la prima idea che potrebbe venire sarebbe quella di chiedere un MAG per tutta la banda, ma ciò non ha senso: il MAG (o G_{max} se sia) non è costante in tutta la banda; servirebbe inoltre un adattatore che trasformi il $SO R$ di carico in Z_{opt} , e il $SO R$ di generatore in Z_{opt} , su tutta la banda; per il limite di Fano, ciò non è realizzabile.

Ciò che dovrai fare è usare un gran numero di reti LC o stub, dunque molti elementi in cascata. Il problema principale continua però a essere il fatto che, volendo usare il MAG, esso non è costante.



Si può fare ciò: farci in modo che la rete risultante guadagni quanto il MAG (circa), nel punto più avanti in frequenza, e sia circa quanto il MAG minimo del nostro range di frequenze.

Questo fatto comporta necessariamente la presenza di un disadattamento nella rete, al fine di equalizzare il guadagno: a banda larga, un po' di disadattamento, almeno con questa tecnica, va accettato.

Ciò dunque introduce della potenza che viene riflessa dalle reti di matching, anche se solo verso il FET.



Una nota: se la rete di matching è reattiva la potenza riflessa si dumpa o sul carico o sul generatore; se le reti sono resistive, in alternativa, la potenza può venir dissipata nella rete e si può chiedere che, il generatore (o il carico), non accetti più potenza riflessa.

Noi useremo reti reattive, anche se capita di usarne anche di resistive.

Ciò che si fa di solito è introdurre disadattamento solo a una porta del FET: o la rete di ingresso o quella di uscita faranno l'equalizzazione, e una "appiattiranno" il guadagno, mentre l'altra trasferirà tutta la potenza disponibile, già equalizzata (o che sarà equalizzata).

Queste reti saranno un po' più complicate di quelle già viste!

ci vorranno almeno una linea e due stub, in modo da avere almeno una doppia risonanza.

Amplificatori reazionati

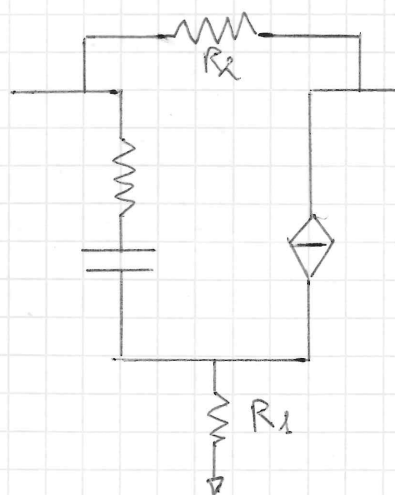
Questa tecnica allarga la banda, ma non tantissimo: volendo se un circuito amplifichi su bande di 10 o 20 GHz, è per forza necessario usare una controreazione. Ciò è brutto, dal momento che la reazione riduce il guadagno (per quanto aumenti la stabilità), dunque si cerca di entrare l'uso.

Si deve partire da dispositivi con un S_{22} molto elevato, in modo che, anche se la reazione riduce il guadagno, qualcosa rimanga.

Dati schemi circuitali noti, si progetta con gli ottimizzatori, facendo ciò:

- 1) Si guarda il limite a bassa frequenza;
- 2) si realizza la reazione prima con elementi resistivi per avere un bel guadagno in bassa frequenza, poi reattivi per allargare la banda.
- 3) di solito, si usa mettere (per il punto 2) delle induttanze: essendo i poli capacitivi, le induttanze li compensano (inductive peaking)

Si consideri il seguente modello, con due reazioni: R_1 e R_2 .
Si tenga a mente che, in serie a R_2 va messa una capacità (di solito piccola, e per avere una grossa, che fa parte bene, la si mette nella rete di



alimentazione) di disaccoppiamento per la continua tra gate e drain.

A partire da questo modello si deve avere esattamente energetico in larga banda, senza aggiungere reti di adattamento (che taglierebbero la banda); poiché in principio si considerano i limiti a bassa frequenza, si possono trascurare molte capacità. Si darebbe a questo punto vedere i limiti per bassa frequenza dei parametri S della rete; si potrebbe verificare che S_{11} e S_{22} sono uguali. Volendo adattare a ingresso e uscita, si trova un vincolo tra R_1 , R_2 , e R_0 (impedenza di riferimento, di solito pari a 50Ω):

$$R_1 = \frac{R_0^2}{R_2} - \frac{1}{g_m}$$

Se questa ipotesi è verificata:

$$S_{11}^P = S_{22}^P = 0; \quad S_{12}^P = \frac{R_0}{R_0 + R_2}; \quad S_{21}^P = \frac{R_0 - R_2}{R_0}$$

Si noti che noi vogliamo $R_2 > 0$; in tal caso, vogliamo che g_m sia grossa, in modo da ridurre il termine sottrattivo, rendendolo molto piccolo.

Se tutto ciò è rispettato, ossia se $S_{11}^P = S_{22}^P = 0$, ho che il FET, grazie alla reazione (l'apice "P" si riferisce proprio all'intero blocco reazionato), lavora bene (adattato) direttamente con carichi e generatori da 50Ω ! Se dunque tutto è adattato, il guadagno del FET reazionato sarà pari a $|S_{21}^P|^2$!

Ricapitolando: lavorando su R_1 , abbiamo adattato; abbiamo poi

visto che $|S_{21}^f|$ vale:

$$S_{21}^f = \frac{R_0 - R_2}{R_0} \quad [R_0 = 50\Omega]$$

Imponendo R_2 , imposto il guadagno! Una volta stabilito il valore del guadagno, dunque di $|S_{21}^f|$, scelgo una R_2 idonea.

Invertendo la relazione di prima,

$$R_2 = R_0 [1 + |S_{21}^f|] \quad [\text{si noti che } S_{21}^f \text{ è } S_{21}^f(0 \text{ Hz})]$$

Poi, però, S_{21}^f è legato a S_{21} del FET. Che si può fare?

Beh, dal vincolo di R_1 si ha (per $R_1 > 0$):

$$g_m \geq \frac{R_2}{R_0^2} \quad ; \quad [\text{moltiplico ambo i membri per } 2R_0]$$

$$\rightarrow 2R_0 g_m \geq 2R_0 \frac{R_2}{R_0^2}$$

Ricordando che, per basse frequenze,

$$S_{21}(0) \approx -2g_m R_0$$

Si ha che:

$$|S_{21}| \geq 2 [1 + |S_{21}^f(0)|] \approx 2 |S_{21}^f(0)|$$

Ossia, il S_{21} del FET deve essere almeno il doppio di quello dello stadio finale controreazionato.

Qualche accorgimento: ora, il posto di R_1 (che sta sul source, posto un po' anticipato da terra), potrà essere come grado di libertà la g_m : essa, in fondo, dipende dalla lunghezza di gate W . Giocando su W , dunque, la transconduttanza può

diventare un parametro di progetto.

Se chiediamo che $R_1 = 0$, otteniamo che:

$$\frac{R_0^2}{R_2} = \frac{1}{g_m}$$

Ciò si può ottenere, risalendo la periferia del FET.

A questo punto, un'ulteriore nota: un parametro paramita molto critico è R_{os} : se era è confrontabile con i 50Ω , tutto ciò che abbiamo detto salta: le approssimazioni diventano cattive.

Ciò può essere importante quando tocchiamo W : se si usala il FET con una W appropriata, aumenta il valore della conduttanza, dunque diminuisce quello della resistenza.

Una correzione può essere fatta mettendo, al posto di R_0 (nelle formule), $R_0 \oplus R_{os}$, ma è solo una correzione empirica.

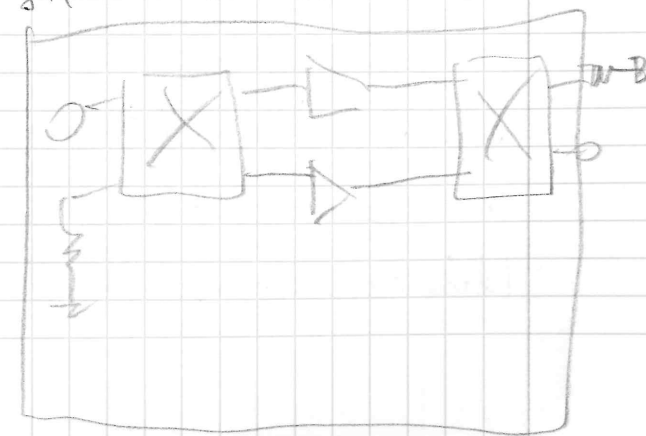
L'effetto principale di R_{os} sarà di sedattare l'usata.

Amplificatore bilanciato

Una tecnica per il progetto di amplificatori di potenza, ma anche a lunga banda, è data dagli amplificatori bilanciati.

Dati due amplificatori A e B, progettati con le loro reti di equalizzazione (edattamento), tali da avere un certo grado di disadattamento, avendo a banda larga, uno schema in cui ci si può basare è il seguente:

Dati A e B identici (come guadagno e disadattamento), li si collega a



degli accoppiatori direzionali a 90° e -3dB (magari Lange, essendo i migliori con banda passante). Alla porta 1 si mette il segnale di ingresso, alla 3 l'uscita, e le altre sono chiuse in corto a sottratto.

Si consideri l'accoppiatore di ingresso: le sue porte 3 e 4 sono caricate su delle Z_L che sono le impedenze di ingresso degli stadi; se applico un'onda progressiva alla porta 1, a_1 , se essa è adattata, non devo avere alcuna b_1 ; b_1 ha contributi da a_3 e a_4 , essendo 2 isolata; avendo che:

$$a_3 = \Gamma_L b_3 = \Gamma_L - \frac{j}{\sqrt{2}} a_1 \quad \text{e} \quad a_4 = \Gamma_L b_4 = \Gamma_L \frac{j}{\sqrt{2}} a_1$$

Essendo i due contributi in controfase, essi si sommano ed eliminiamo; infatti:

$$b_1 = -\frac{j}{\sqrt{2}} a_3 + \frac{j}{\sqrt{2}} a_4$$

Ma l'eliminazione si avrà se e solo se Γ_L sarà uguale sulle due porte.

Al contrario, nella porta 2, la potenza uscente si va a sommare, e si dissipa tutta su R_0 , evitando di girare nel circuito.

Per quanto concerne il guadagno, quando gli amplificatori "guardano" verso l'ingresso, essi "vedono" l'impedenza di normalizzazione, essendo tutto chiuso su R_0 ; gli amplificatori lavorano dunque in adattamento, e i FET (con le reti annessi) guadagnano $|S_{21}|^2$.

Si spera poi che esso sia vicino al MAG. Lo stadio risultante (bilanciato) guadagnerà $|S_{21}|^2$!

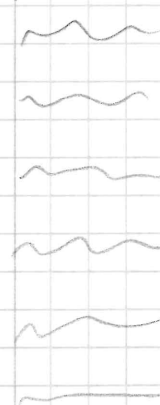
Ricordando le formule:

$$\Gamma_{in} = S_{11} + \frac{S_{21} S_{12} \Gamma_L}{1 - S_{22} \Gamma_L} \quad \left[\begin{array}{l} \text{poiché gli amplificatori "vedono"} \\ R_0 \end{array} \right]$$

Essendo $\Gamma_L = 0$, in ingresso al circuito si vede il S_{11} del sistema "amplificatore + matching"; la stessa cosa vale per la porta di uscita.

Questi sono i concetti dei FET, ciò che si vede dall'ingresso dell'amplificatore matchato. Tutto ciò che però viene utile finisce sulla R_0 della porta 2 (la porta che agli amplificatori facciamo vedere i carichi Z_{Leq} adattati).

È dunque un amplificatore a banda larga presenta un disadattamento, mediante questo schema si riesce a eliminarlo, dirottandolo su una resistenza dissipativa.



Descrivo l'impedenza! slide

(38) ! :-P

1h 35 ma fatto ! :-P