



RADIOCOMMUNICATIONS

Principes & circuits

Joël REDOUTEY

Novembre 2010

SOMMAIRE

CIRCUITS ELECTRONIQUES UTILISES EN RADIOCOMMUNICATIONS	6
Chapitre I LA MODULATION	7
1 Modulation d'amplitude.....	8
1.1 Spectre d'un signal modulé en amplitude.....	8
1.2 Modulation d'amplitude à bande latérale unique (BLU).....	10
2 Modulations angulaires.....	11
3 Modulation de fréquence	11
3.1 Modulation de fréquence à bande étroite.....	12
3.2 Modulation de fréquence à bande large (FM).....	12
3.3 Largeur de bande d'un signal FM.....	13
4 Modulation de phase.....	14
CHAPITRE II composants ACTIFS ET PASSIFS en hautes Fréquences	15
1 Effet de peau	15
2 Inductance en haute fréquence	16
3 Condensateur en haute fréquence.	17
4 DIODES.....	18
4.1 Diodes Schottky	19
4.2 Diodes PIN	20
5 TRANSISTOR BIPOLAIRE.....	21
5.1 MODELE HAUTE FREQUENCE DU TRANSISTOR BIPOLAIRE.....	21
5.2 Evolution du gain en courant β (ou h_{21}) en fonction de la fréquence.	22
5.3 Fréquence de transition	23
5.4 Bande passante d'un étage amplificateur.	24
5.5 Influence du boîtier et des connexions.....	25

5.6	Impédance d'entrée.....	26
5.7	Impédance de sortie	27
5.8	Modèle général du transistor.....	28
6	TRANSISTOR A EFFET DE CHAMP MOS.....	28
6.1	frequence de transition	29
6.2	Effet MILLER.....	30
6.3	Réponse en fréquence d'un amplificateur à MOSFET.....	30
CHAPITRE III LES CIRCUITS RESONANTS		33
1	Circuit résonant série	33
1.1	Coefficient de surtension Q.....	33
2	Circuit résonant parallèle.....	35
2.1	Coefficient de surtension Q.....	35
3	Récapitulation.....	36
4	Facteur de qualité en charge (loaded Q).....	37
Exemple		38
5	Transformation d'impédance.....	38
5.1	Prise intermédiaire capacitive	39
Exemple		41
CHAPITRE IV ADAPTATION D'IMPEDANCE.....		43
1	Puissance maximale fournie par un générateur relié à une charge quelconque.	43
2	Circuits d'adaptation d'impédance	44
3	Circuit d'adaptation en L.....	44
3.1	Transformation d'une résistance pure	44
3.2	Facteur de qualité du circuit.....	47
Exemple		47
3.3	Cas où les resistances terminales ne sont pas des résistances pures	49

Exemple	49
3.4 Bande passante et facteur de qualité	50
4 Réseaux d'adaptation en Té ou en Pi.....	51
4.1 Circuit d'adaptation d'impédance en PI.....	51
Exemple	52
Calcul de la valeur des éléments (à $F = 100$ MHz):.....	53
CHAPITRE V Changement de fréquence	54
1 MELANGEUR MULTIPLICATIF	54
2 MELANGEUR A COMMUTATION	56
2.1 MELANGEUR A DIODES	57
2.2 MELANGEUR EN ANNEAU	59
2.3 MELANGEUR ACTIF	61
2.4 MELANGEURS A REJECTION DE LA FREQUENCE IMAGE.....	62
Généralités	65
1.1 Définitions.....	65
1.2 Performances d'un oscillateur.....	65
Principe de fonctionnement d'un oscillateur	66
3 Oscillateurs hautes fréquences.....	68
3.1 Oscillateur COLPITTS à transistor bipolaire.....	69
4 Oscillateurs à quartz	71
3 Oscillateurs commandés en tension.....	75
4 Boucle à verrouillage de phase (PLL)	76
5 Synthétiseur de fréquence.....	83
CHAPITRE VII La Démodulation	87
1 Démodulation d'amplitude.....	87

2	Démodulation de fréquence.....	88
	ANNEXES	90
1	Conversion parallèle → série	91
2	Conversion série → parallèle.....	92
3	Facteur de qualité.....	93
	THEOREME DE MILLER.....	94
	LE BRUIT.....	95
	LA DISTORSION	101

CIRCUITS ELECTRONIQUES UTILISES EN RADIOCOMMUNICATIONS

Le 21^{ème} siècle sera pour l'électronique le siècle du *sans fil* et du *sans contact*.

Chacun peut le constater dans la vie de tous les jours: télécommande de nos appareils multimédia, téléphone mobile ou sans fil, télépéage autoroutier, étiquette électronique (RFID), ouverture automatique de portes, boucle locale radio (WLAN), etc.

La plupart de ces applications font appel aux transmissions par ondes électromagnétiques, que nous pouvons qualifier de *radiocommunications* au sens large.

Ce document présente les techniques de base et les principales fonctions utilisées dans les systèmes de radiocommunication. Son objectif est de fournir les informations nécessaires à la compréhension des circuits électroniques hautes fréquences (1 à 1000 MHz) utilisés dans les systèmes de transmission qu'ils soient analogiques ou numériques.

La propagation des ondes électromagnétiques sera supposée connue.

Les principaux thèmes abordés sont:

- La *modulation*, opération qui consiste à faire transporter une information par une onde porteuse à haute fréquence, et la *démodulation* qui est l'opération inverse.
- *L'amplification à haute fréquence*, notamment le comportement des composants semi conducteurs
- Le *filtrage* permettant de sélectionner un signal dans un spectre radioélectrique de plus en plus encombré
- *L'adaptation d'impédance* nécessaire pour optimiser les transferts d'énergie
- Le *changement de fréquence* et les circuits associés,...

On étudiera également quelques fonctions classiques et les circuits qui permettent de les réaliser:

- Oscillateurs et générateurs de signaux haute fréquence
- Modulateurs et démodulateurs
- Amplificateurs,...

CHAPITRE I

LA MODULATION

On sait qu'une onde électromagnétique se propage dans l'espace libre, à condition qu'elle soit émise par une antenne dont les dimensions doivent être en rapport avec la longueur d'onde $\lambda = C/f$.

(C: célérité de la lumière, f: fréquence). Plus une antenne est courte devant la longueur d'onde, moins elle est efficace.

Un signal dans la bande audio par exemple de fréquence $f=3\text{kHz}$ correspond à une longueur d'onde électrique de $300\,000 / 3\,000 = 100\text{ km}$. On ne pourra donc pas émettre ce signal directement.

On a alors recours à une onde électromagnétique haute fréquence, par exemple de fréquence $f=100\text{ MHz}$, ce qui correspond à une longueur d'onde de 3m , pour transporter ce signal.

Cette haute fréquence est appelée **porteuse** (carrier) et l'opération qui consiste à incorporer le signal utile à la porteuse s'appelle la **modulation**.

Considérons un signal sinusoïdal haute fréquence $s(t) = A \sin(2\pi f_c t + \phi)$

En observant son expression, on voit que l'on peut faire varier trois paramètres:

- L'amplitude A
- La fréquence f_c
- La phase Φ

En faisant varier un de ces paramètres au rythme du signal à transmettre, on obtient un signal haute fréquence modulé. Nous aurons donc trois types de modulation:

- La modulation d'amplitude (AM)
- La modulation de fréquence (FM)
- La modulation de phase (PM)

Le signal utile est généralement variable dans le temps. Son spectre, que l'on appelle **bande de base**, s'étend entre les fréquences $f_{\min}$ (voisine de 0) et $f_{\max}$. On notera que la fréquence $f_{\max}$ doit être faible devant la fréquence de la porteuse.

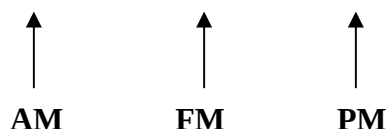
Dans un but de simplification, nous considérerons par la suite que le signal à transmettre est sinusoïdal de la forme $e(t) = E(t) \sin \omega_m t$

$f_m = \omega_m / 2\pi$ est la fréquence de modulation

$f_c = \omega_c / 2\pi$ est la fréquence porteuse

$f_c \gg f_m$

$$s(t) = S(t) \cdot \sin [2\pi f_c(t) t + \phi(t)]$$



$$e(t) = E(t) \sin[2\pi f_m(t)]$$

1 MODULATION D'AMPLITUDE ANALOGIQUE

Le signal de modulation peut être analogique (audio, video,...) ou impulsionnel (télégraphie, radar, télécommande,...).

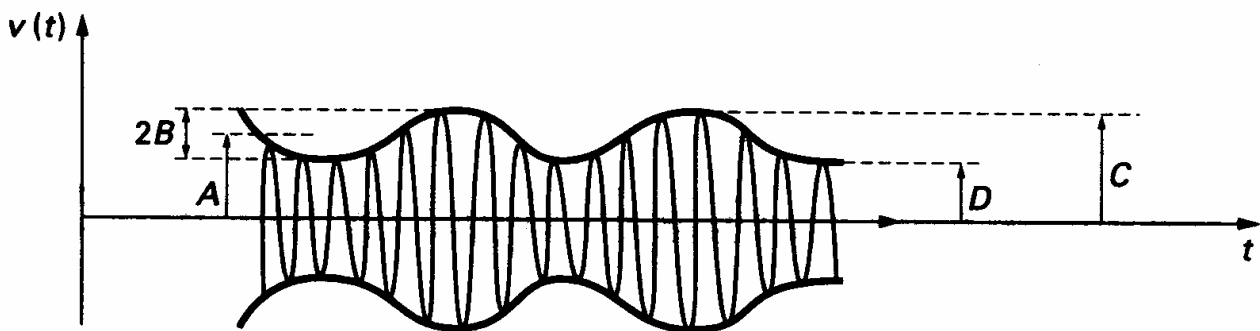
L'expression générale d'un signal modulé en amplitude s'écrit:

$$e(t) = A(1 + m \cos \omega_m t) \cos \omega_c t$$

$\omega_m = 2\pi f_m$ pulsation de modulation

$\omega_c = 2\pi f_c$ pulsation de la porteuse

m est appelé **taux de modulation** ($0 < m < 1$)



1.1 SPECTRE D'UN SIGNAL MODULE EN AMPLITUDE

$$e(t) = A(1 + m \cos \omega_m t) \cos \omega_c t$$

$$e(t) = A \cos(\omega_c t) + mA \cos(\omega_m t) \cos(\omega_c t)$$

transformons le produit de cosinus en somme

$$2 \cos a \cos b = \cos(a+b) + \cos(a-b)$$

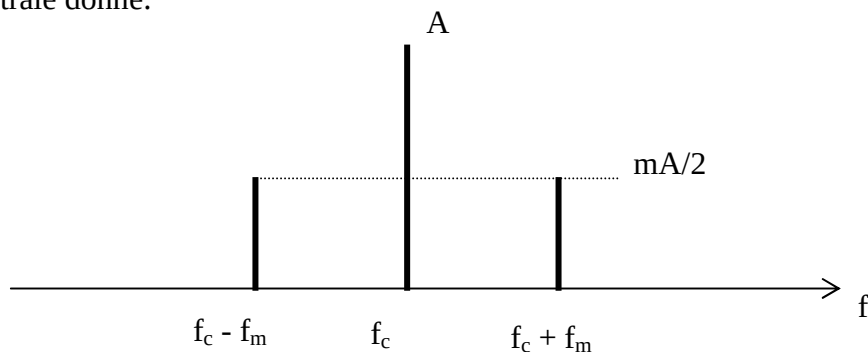
il vient:

$$e(t) = A \cos(\omega_c t) + (mA/2) \cos(\omega_c - \omega_m)t + (mA/2) \cos(\omega_c + \omega_m)t$$

On voit apparaître trois termes:

- Une composante à la fréquence de la porteuse f_c
- Une composante à la fréquence $f_c - f_m$
- Une composante à la fréquence $f_c + f_m$

L'analyse spectrale donne:



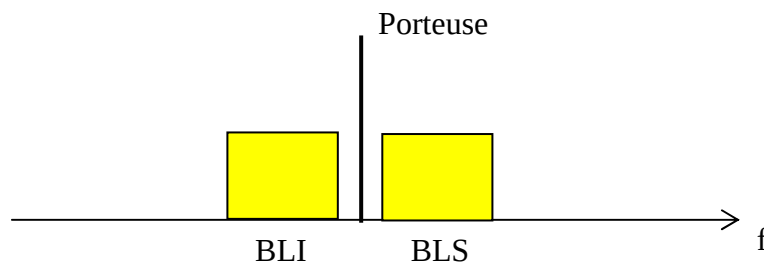
On peut remarquer que:

- L'onde porteuse ne sert qu'au transport de l'information
- Si le taux de modulation est de 100%, la moitié de l'énergie est concentrée dans la porteuse
- Lorsque le signal de modulation est variable dans le temps (en amplitude et en fréquence), ce qui est le cas le plus général, on n'obtient plus deux raies latérales mais deux **bandes latérales** (sidebands) dont la largeur est égale à la bande passante du signal de modulation et qui s'étendent de chaque côté de la porteuse.

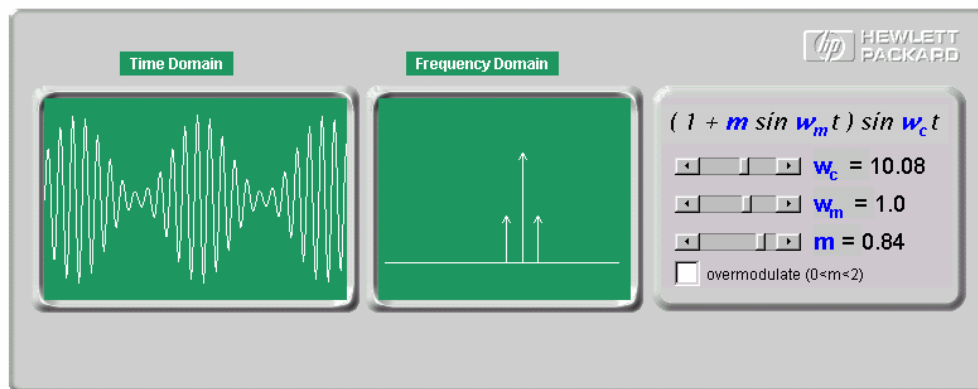
Par exemple, la bande passante d'un signal téléphonique est 300 Hz - 3000 Hz.

Le spectre d'un signal modulé en amplitude par une voie téléphonique s'étendra 3 kHz de chaque côté de la porteuse, soit un encombrement spectral total de 6 kHz.

D'une manière générale la bande passante du signal à transmettre est appelée bande de base.



- L'information utile est contenue dans les bandes latérales, et elle est redondante. Une seule bande latérale suffit pour reconstituer le signal.

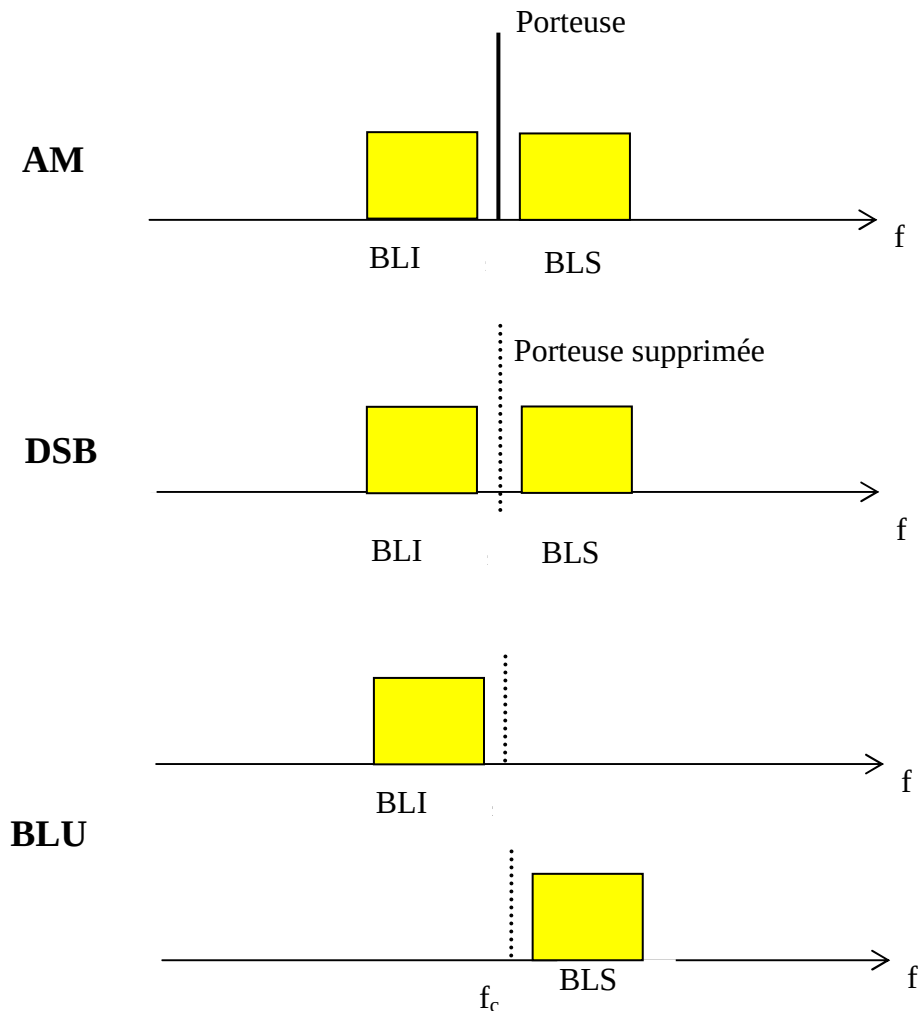


La modulation d'amplitude est très utilisée. On peut citer notamment:

- La radiodiffusion en petites et grandes ondes (PO et GO) ainsi qu'en ondes courtes (OC) avec une bande passante de la modulation limitée à 9 kHz
- La télévision analogique diffusée par voie Hertzienne (VHF et UHF)
- Le trafic aéronautique à courte distance (tour de contrôle, ...) en VHF

1.2 MODULATION D'AMPLITUDE A BANDE LATÉRALE UNIQUE (BLU)

Dans le but d'augmenter l'efficacité des systèmes à modulation d'amplitude et d'en diminuer l'encombrement spectral, on peut effectuer un traitement du signal permettant de supprimer la porteuse et une des bandes latérales. On obtient ainsi un signal à **bande latérale unique** (Single Side Band ou SSB).



L'utilisation d'un **modulateur équilibré** permet de moduler et de supprimer la porteuse. On obtient un signal DSB (Double Side Band). Par filtrage, on peut supprimer une des bandes latérales.

L'efficacité est considérablement augmentée et l'encombrement spectral est divisé par deux, ce qui permet de mettre plus de canaux de communication dans le même espace.

La BLU est utilisée pour les communications lointaines en bandes décadiques (3 - 30 MHz): communications maritimes, ambassades, ...

Par convention, la bande latérale inférieure (BLI) est utilisée au dessous de 10 MHz et la bande latérale supérieure au dessus de 10 MHz.

2 MODULATIONS ANGULAIRES

Soit un signal porteur $e(t) = A \cos(\omega_c t + \varphi_0)$

En posant $\Phi(t) = \omega_c t + \varphi_0$ on peut définir une pulsation instantanée $\omega_i = d\Phi / dt$
d'où $\Phi(t) = \int \omega_i dt$

Si $\Phi(t)$ est rendue variable en fonction d'un signal de modulation $f(t)$, on obtient une modulation angulaire.

3 MODULATION DE FREQUENCE

On fait varier la pulsation instantanée de la porteuse linéairement en fonction de la modulation.

$$\omega_i(t) = \omega_c + k_2 f(t)$$

On a donc:

$$\Phi(t) = \int \omega_i dt = \omega_c t + \varphi_0 + k_2 \int f(t) dt$$

On voit que dans ce cas la phase de la porteuse varie linéairement avec l'intégrale du signal de modulation.

Supposons, dans un but de simplification, que le signal de modulation soit sinusoïdal de la forme:

$$f(t) = a \cos(\omega_m t)$$

La pulsation instantanée peut s'écrire : $\omega_i(t) = \omega_c + \Delta\omega_{\text{crête}} \cdot \cos(\omega_m t)$ où $\Delta\omega_{\text{crête}}$ est une constante qui dépend de l'amplitude a du signal de modulation et des propriétés du circuit de modulation. On a toujours $\Delta\omega_{\text{crête}} \ll \omega_c$

$\Phi(t)$ peut se mettre sous la forme:

$$\Phi(t) = \int \omega_i dt = \omega_c t + \varphi_0 + (\Delta\omega_{\text{crête}} / \omega_m) \sin \omega_m t$$

Pour simplifier l'expression on peut choisir la phase à l'origine nulle: $\varphi_0 = 0$ et poser :

$$m = \Delta\omega_{\text{crête}} / \omega_m = \Delta f_{\text{crête}} / f_m$$

Le signal modulé s'écrit donc:

$$e(t) = A \cdot \cos(\omega_c t + m \cdot \sin \omega_m t)$$

- m est l'**indice de modulation**. Il représente l'écart maximal de phase de la porteuse.
- $\Delta f_{\text{crête}}$ est la **déviatiion maximale** (shift) de la fréquence de la porteuse (par rapport à sa fréquence en l'absence de modulation)

D'après la relation trigonométrique $\cos(a+b) = \cos a \cdot \cos b - \sin a \cdot \sin b$, on peut écrire:

$$e(t) = A[\cos(\omega_c t) \cdot \cos(m \cdot \sin \omega_m t) - \sin(\omega_c t) \cdot \sin(m \cdot \sin \omega_m t)]$$

3.1 MODULATION DE FREQUENCE A BANDE ETROITE

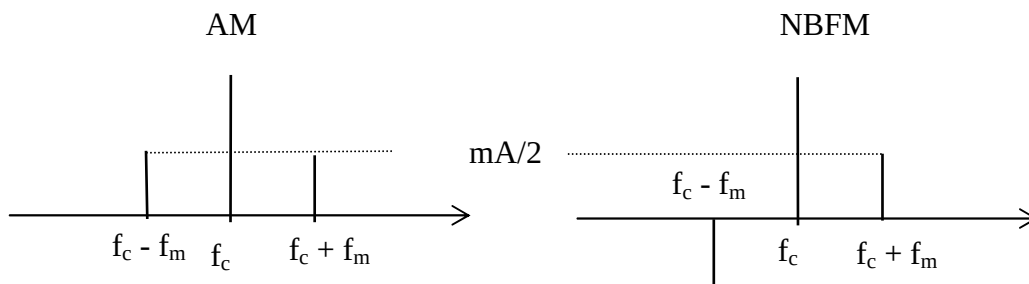
On se place dans le cas où $m \ll \pi/2$ ($m < 0,2$ dans la pratique)

si $m \ll \pi/2$ on a : $\cos(m.\sin \omega_m t) \approx 1$ et $\sin(m.\sin \omega_m t) \approx m.\sin \omega_m t$
d'où $e(t) = A[\cos(\omega_c t) - m.\sin(\omega_m t).\sin(\omega_c t)]$

d'après la relation $\sin a.\sin b = 1/2 [\cos(a-b) - \cos(a+b)]$ on peut écrire:

$$e(t) = A\cos(\omega_c t) + (mA/2).\cos(\omega_c + \omega_m)t - (mA/2).\cos(\omega_c - \omega_m)t$$

Cette relation est très similaire à celle de la modulation d'amplitude à l'exception de la phase de la bande latérale inférieure qui est inversée.



L'encombrement spectral est identique à celui de la modulation d'amplitude.

La modulation de fréquence à bande étroite (Narrow Band Frequency Modulation ou NBFM) est largement utilisée dans les réseaux mobiles ou portatifs (Taxis, ambulances, pompier, police, trafic portuaire, VHF marine...)

3.2 MODULATION DE FREQUENCE A BANDE LARGE (FM)

Nous sommes dans le cas où m n'est pas petit. Le signal modulé s'écrit:

$$e(t) = A[\cos(\omega_c t).\cos(m.\sin \omega_m t) - \sin(\omega_c t).\sin(m.\sin \omega_m t)]$$

On peut développer les termes $\cos(m.\sin \omega_m t)$ et $\sin(m.\sin \omega_m t)$ à l'aide des fonctions de Bessel:

$$\cos(m.\sin \omega_m t) = J_0(m) + 2J_2(m)\cos(2\omega_m t) + J_4(m)\cos(4\omega_m t) + \dots$$

$$\sin(m.\sin \omega_m t) = 2J_1(m).\sin\omega_m t + 2J_3(m).\sin 3\omega_m t + \dots$$

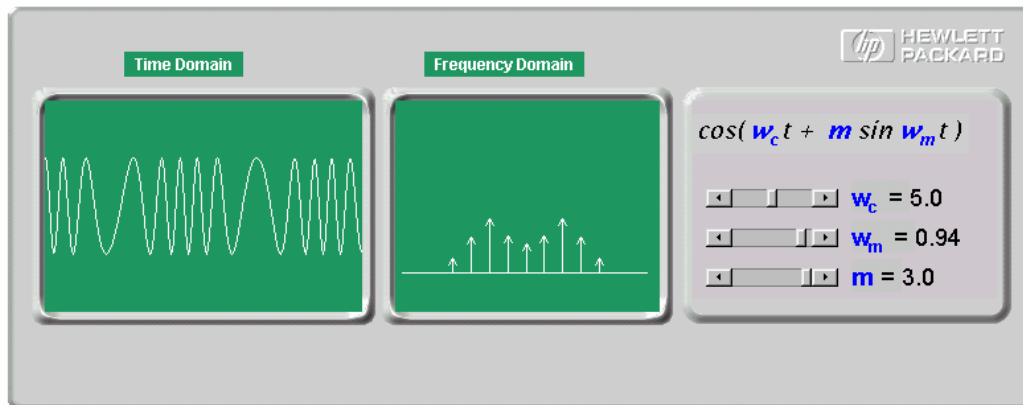
où J_n représente la fonction de Bessel du $n^{\text{ème}}$ ordre.

On a finalement:

$$e(t) = J_0(m) - J_1(m)[\cos(\omega_c - \omega_m)t - \cos(\omega_c + \omega_m)t] + J_2(m)[\cos(\omega_c - 2\omega_m)t + \cos(\omega_c + 2\omega_m)t] - J_3(m)[\cos(\omega_c - 3\omega_m)t - \cos(\omega_c + 3\omega_m)t] + J_4(m)[\cos(\omega_c - 4\omega_m)t + \cos(\omega_c + 4\omega_m)t] + \dots$$

On obtient un spectre constitué de la fréquence porteuse et d'une infinité de raies latérales dont l'amplitude est proportionnelle à $J_n(m)$.

Plus m augmente, plus le spectre est large et donc plus la bande passante doit être importante.



La modulation de fréquence à bande large est utilisée en radiodiffusion (bande FM 88 à 108 MHz), en télévision par satellite (analogique), dans les faisceaux hertziens , ...

3.3 LARGEUR DE BANDE D'UN SIGNAL FM

Dans la pratique, le spectre d'un signal FM ne s'étend pas à l'infini. L'amplitude des raies latérales devient négligeable à partir d'un certain écart par rapport à la porteuse (plus ou moins grand selon l'indice de modulation).

La bande passante nécessaire pour une réception sans distorsion peut être déterminée en comptabilisant le nombre de raies significatives (par exemple celles dont l'amplitude est supérieure à 1% de l'amplitude de la porteuse non modulée).

L'étendue spectrale dépend de l'indice de modulation m, qui dépend lui même de l'amplitude et de la fréquence du signal de modulation:

$$m = \Delta f_{\text{crête}} / f_m$$

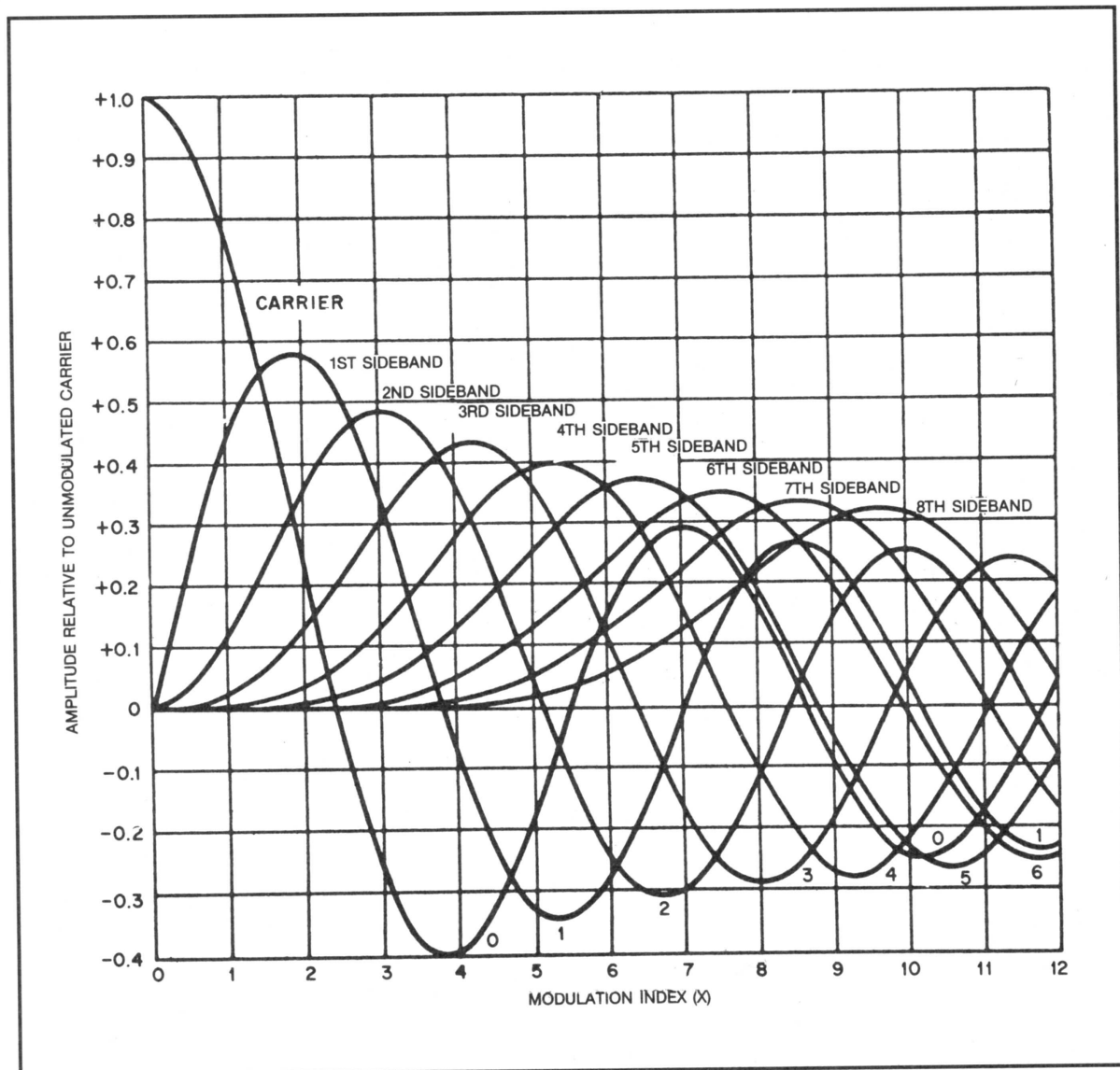
On notera les deux résultats importants suivants:

- A très faible indice de modulation ($m < 0,2$) , le spectre ne comporte que deux raies latérales comme en modulation d'amplitude. La bande passante nécessaire est égale à $2f_m$.
- A très fort indice de modulation ($m > 100$) , la bande passante nécessaire est égale à $2\Delta f_{\text{crête}}$.

Pour des communications téléphoniques, on peut tolérer un taux de distorsion supérieur. La bande passante nécessaire peut alors être calculée à l'aide de la relation de Carson:

$$B \approx 2(\Delta f_{\text{crête}} + f_m) \quad \text{soit} \quad B \approx 2f_m(1+m)$$

Dans la réalité, le signal de modulation n'est pas sinusoïdal et les calculs deviennent vite complexes. Le cas du signal sinusoïdal donne néanmoins des indications précieuses.



Amplitude relative des raies en fonction de l'indice de modulation

4 MODULATION DE PHASE

On a dans ce cas $\Phi(t) = \omega_c t + \varphi_0 + k_1 f(t)$

k_1 est une constante du système
 $f(t)$ est le signal de modulation

La phase du signal porteur varie linéairement en fonction du signal de modulation.

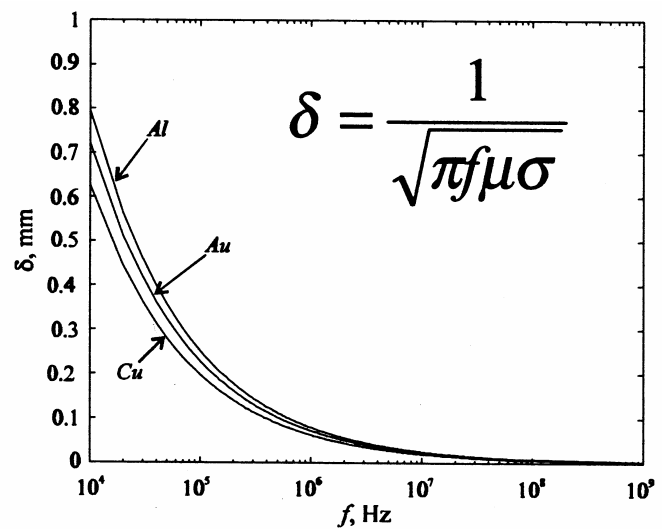
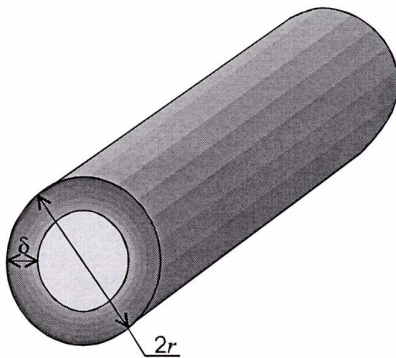
La modulation de phase, très voisine de la modulation de fréquence, est rarement utilisée en analogique.

CHAPITRE II

COMPOSANTS ACTIFS ET PASSIFS EN HAUTES FREQUENCES

1 EFFET DE PEAU

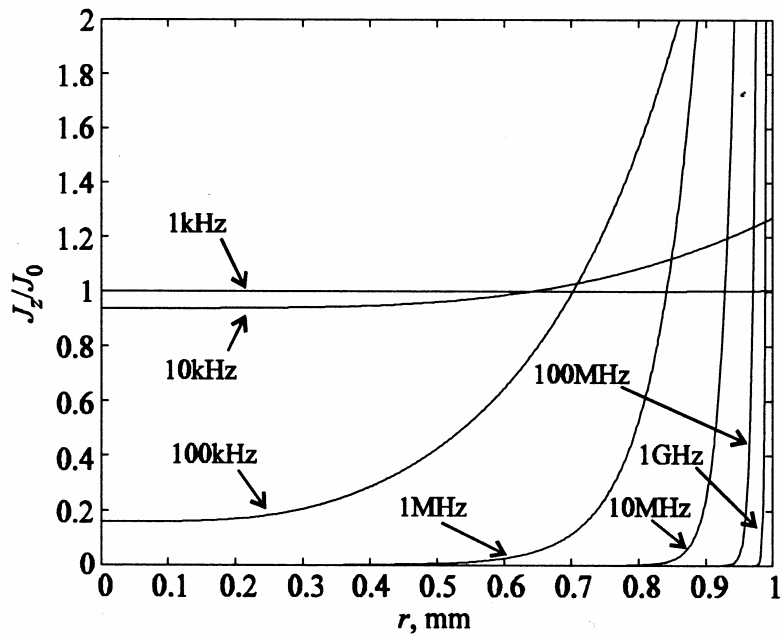
En courant continu, la densité de courant est uniforme dans toute la section d' un conducteur.
En courant alternatif et plus particulièrement en haute fréquence, ceci n'est plus vrai et le courant circule dans une couronne d'épaisseur δ à la surface du conducteur.



δ , appelée épaisseur de peau, représente la profondeur à laquelle la densité de courant chute à un facteur $1/e = 37\%$ de sa valeur en surface.

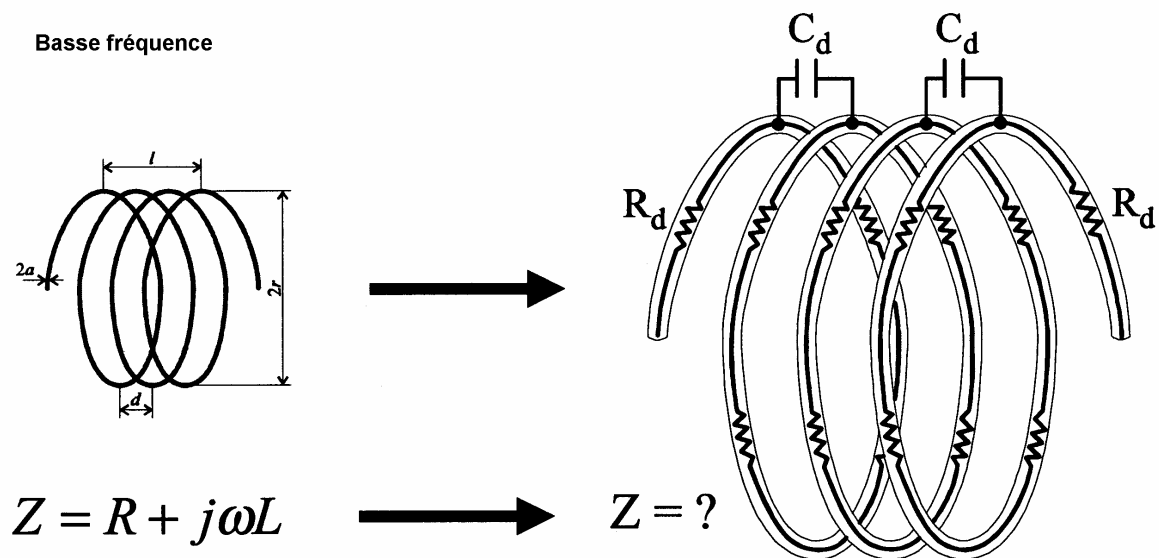
A haute fréquence, le courant circule quasiment à la surface du conducteur.

La courbe suivante montre l'évolution de la densité de courant à l'intérieur du conducteur pour diverses fréquences.



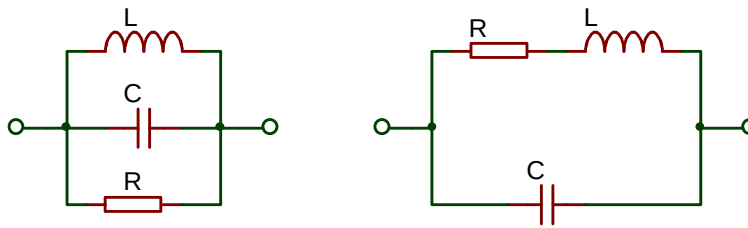
2 INDUCTANCE EN HAUTE FREQUENCE

En basse fréquence, une inductance (à air) est parfaitement définie par ses dimensions géométriques et son impédance peut être facilement calculée.



En haute fréquence, il faut tenir compte de l'effet de peau qui rend la résistance dépendante de la fréquence, des capacités parasites entre spires.

Une inductance réelle pourra donc être mise sous la forme d'un des circuits équivalents suivants:



Schémas équivalents d'une inductance en haute fréquence

3 CONDENSATEUR EN HAUTE FREQUENCE.

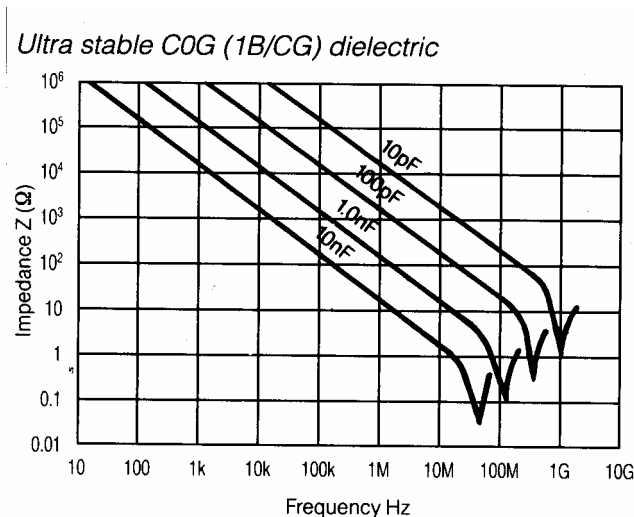
En haute fréquence, les éléments parasites (inductance des connexions et des armatures, pertes dans le diélectrique) deviennent d'une importance critique.

Le schéma équivalent d'un condensateur est le suivant.



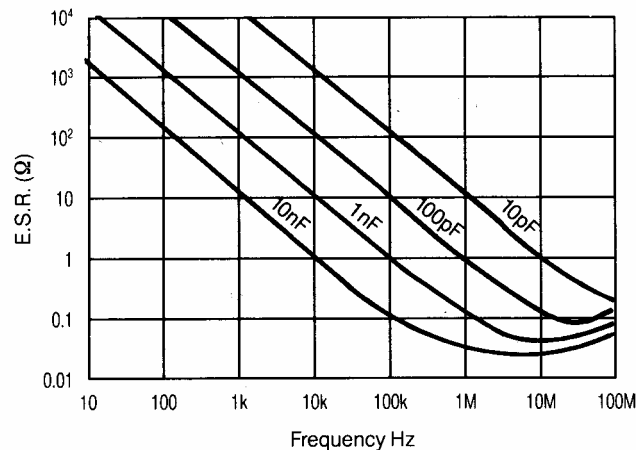
L représente l'inductance des connexions. Elle constitue avec la capacité du condensateur un circuit résonant série dont la fréquence de résonance est donnée par $LC\omega^2 = 1$.

La courbe suivante montre l'impédance d'un condensateur céramique en fonction de la fréquence au voisinage de la résonance.



Les pertes diélectriques sont d'une extrême importance à haute fréquence. Elles peuvent être caractérisée par le facteur de pertes $\tan \delta$ ou de préférence par la résistance série équivalente ESR dont l'évolution en fonction de la fréquence est présentée ci après.

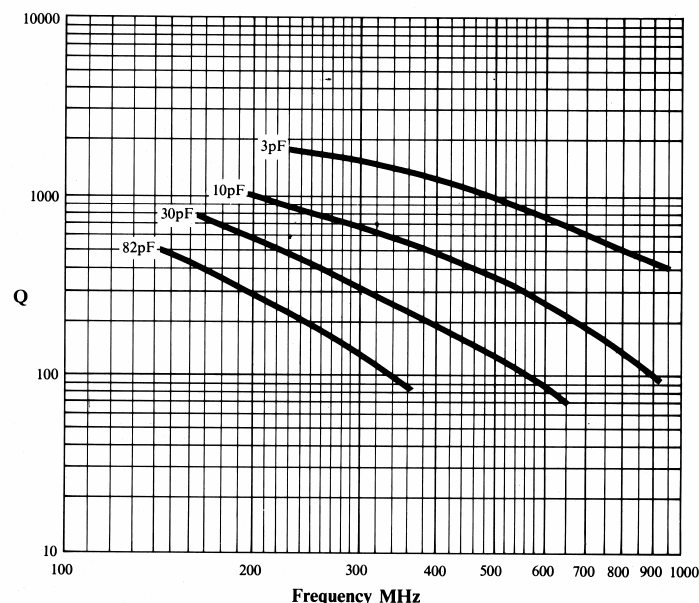
Ultra stable C0G (1B/CG) dielectric



Pour les condensateurs utilisés dans les circuits accordés, le facteur de qualité Q et le coefficient de température sont des paramètres très importants.

En haute fréquence on utilise essentiellement des condensateurs à diélectrique céramique dont les principaux types sont:

- COG : très stables, précis et à coefficient de température défini (NPO : coefficient de température = 0). Gamme de capacité de 1pF à 10nF environ, classe 1.
- X7R : Stable, capacité variable avec la température ($\pm 15\%$ dans la gamme -55°C à $+125^\circ\text{C}$). Gamme de capacité de 100 pF à $1\mu\text{F}$, classe 2.
- Z5U : usage général, instable dans le temps et en température ($+22$ à -56% dans la gamme $+10$ à $+85^\circ\text{C}$). Gamme de capacité de 1nF à $4,7\mu\text{F}$. Utilisation en filtrage, découplage.



4 DIODES

La conduction dans une diode à jonction PN résulte du mécanisme d'injection des porteurs minoritaires. De ce fait, lorsqu'une diode à jonction PN conduit, elle stocke une charge $Q_S = \tau I_F$ qu'il faut évacuer avant de pouvoir la bloquer (IF: courant dans la diode, τ : durée de vie).

Ce type de diode se prête mal à une utilisation à fréquences élevées. On lui préfère donc la diode à barrière Schottky (ou diode Schottky).

4.1 DIODES SCHOTTKY

Une diode Schottky est formée d'une barrière de potentiel créée par un contact métal-semiconducteur.

Elle est essentiellement caractérisée par:

- un seuil de conduction très inférieur à celui d'une diode à jonction PN classique.
- un temps de commutation extrêmement court (quelques dizaines de picosecondes) dû à l'absence de porteurs minoritaires.
- une capacité inverse réduite.

La tenue en tension d'une diode Schottky est normalement assez faible, de l'ordre de la dizaine de volts.

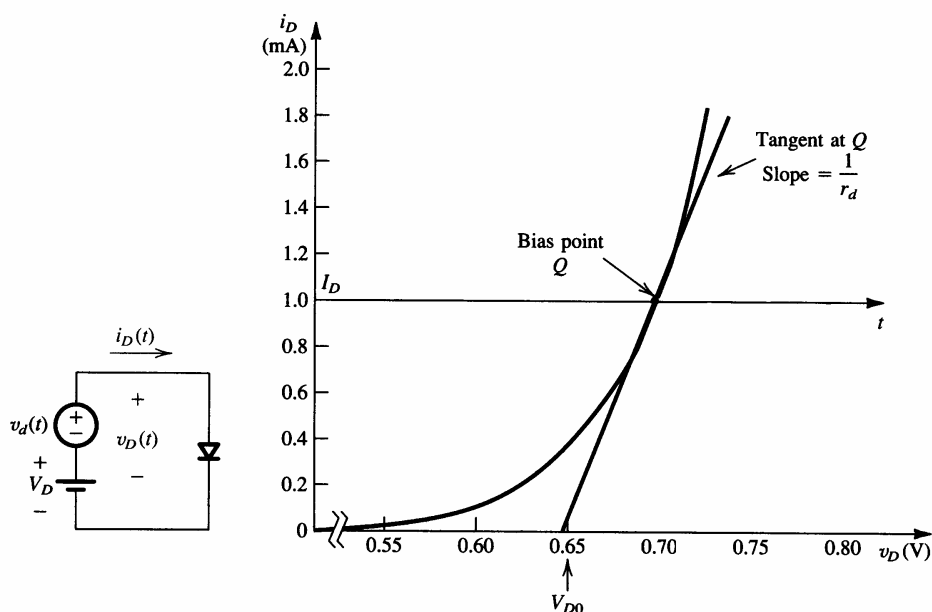
Cependant, différents procédés technologiques tels que l'adjonction d'une jonction PN à la périphérie avec optimisation du profil de gravure de la silice, la protection par oxyde passivé et contact métallique débordant permettent sur certains modèles d'atteindre la centaine de volts et d'améliorer en même temps la fiabilité.

L'équation du courant dans une diode Schottky est de la forme:

$$I_F = I_0 (\exp(V_F/nV_T) - 1)$$

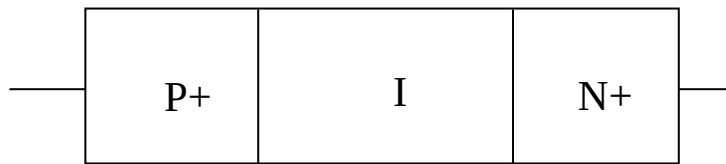
Où I_0 est le courant de saturation et n un coefficient dépendant de la technologie.

En haute fréquence, la diode lorsqu'elle conduit se comporte comme une résistance, que l'on appelle résistance dynamique R_D , qui dépend du point de polarisation et de la température.



4.2 DIODES PIN

Une diode PIN est une diode à jonction PN dans laquelle une couche de silicium intrinsèque (haute résistivité) a été placée entre les zones P et N fortement dopées.



Lorsque la diode PIN est en conduction, il y a injection de charges dans la zone intrinsèque.

Cette charge, constituée de trous et d'électrons, possède une durée de vie τ finie (temps moyen pour qu'un trou et un électron se recombinent).

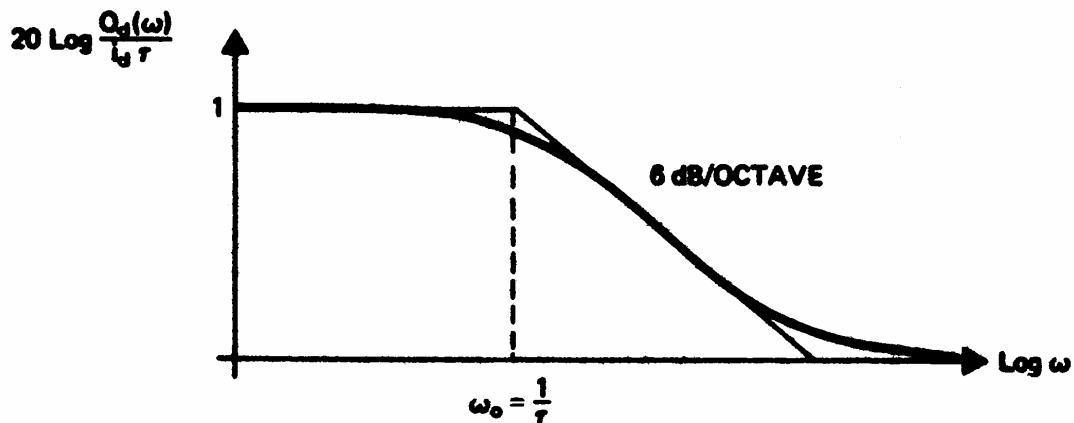
Si la diode est traversée par un courant continu (polarisation) et un courant haute fréquence superposé, l'équation du courant s'écrit:

$$I_F = Q_S/\tau + dQ_S/dt$$

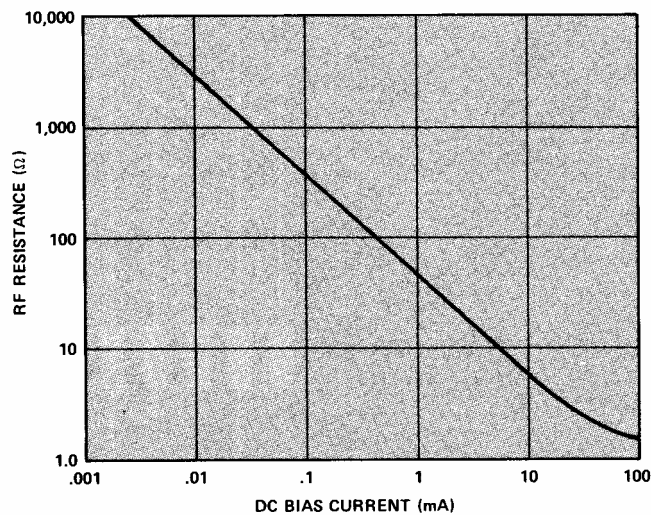
Par transformation de Laplace, on obtient l'expression de Q_S en fonction de la fréquence:

$$Q_S(j\omega) = j\omega I_F \tau / (1 + j\omega\tau)$$

Le tracé du diagramme de Bode montre qu'au dessous de $f_0 = 1/2\pi\tau$ la composante haute fréquence du courant a le même effet que la composante continue et la diode PIN se comporte comme une diode PN. Au delà de f_0 l'effet de la composante haute fréquence décroît rapidement et la diode PIN se comporte comme une résistance pure dont la valeur peut être contrôlée par la composante continue du courant.



Les diodes PIN trouvent de nombreuses applications en haute fréquence : interrupteurs, atténuateurs, limiteurs, ...



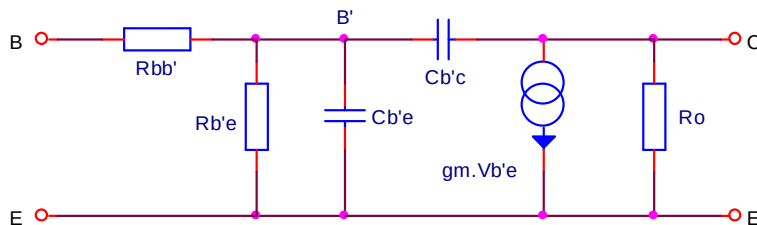
Résistance HF d'une diode PIN en fonction du courant de polarisation.

5 TRANSISTOR BIPOLAIRE

Le modèle petits signaux du transistor n'est plus suffisant à fréquence élevée car il ne tient pas compte de la capacité des jonctions émetteur-base et collecteur-base.

5.1 MODELE HAUTE FREQUENCE DU TRANSISTOR BIPOLAIRE

On utilise généralement le modèle **hybrid- π** qui est le plus commode.



$R_{bb'}$ est la résistance d'accès à la base réelle B'

$C_{b'e}$ est une capacité qui se compose de deux termes:

Une capacité de diffusion proportionnelle au courant de polarisation

Une capacité de transition qui dépend de la valeur de V_{be} .

$C_{b'c}$ est une capacité de transition (capacité de la zone de déplétion) qui dépend de la tension V_{cb} .

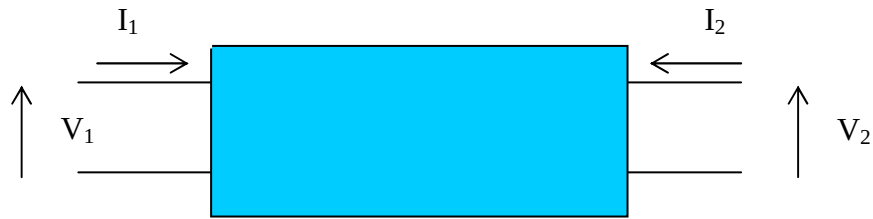
Ordres de grandeurs:

$C_{b'e} \approx 1$ à 50 pF

$C_{b'c} \approx 0,1$ à 5 pF

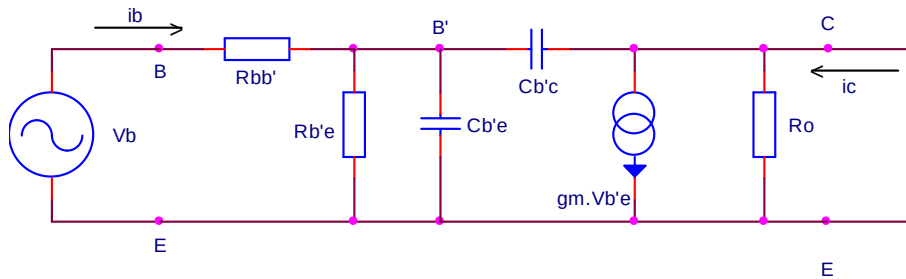
5.2 EVOLUTION DU GAIN EN COURANT β (OU h_{21}) EN FONCTION DE LA FREQUENCE.

Plaçons nous dans l'hypothèse des petits signaux. Le transistor est représenté comme un quadripôle.

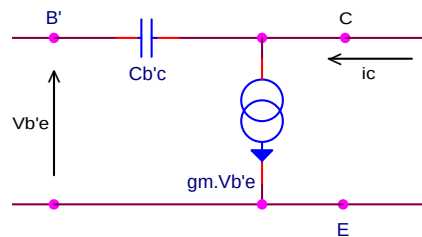


$$\begin{aligned} v_1 &= h_{11}i_1 + h_{12}v_2 \\ i_2 &= h_{21}i_1 + h_{22}v_2 \end{aligned}$$

Pour mesurer h_{21} on relie (en signal) le collecteur et l'émetteur: $v_{ce} = v_2 = 0$ d'où $i_2 = h_{21} i_1$.
Le schéma équivalent du circuit est alors le suivant:



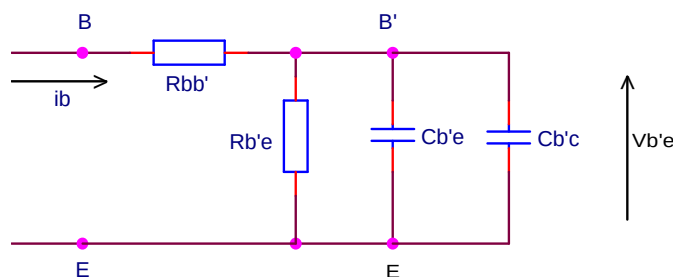
La résistance R_o étant en court circuit, on peut la supprimer.
Le circuit de sortie se simplifie et l'on a:



Le courant dans la capacité $C_{b'c}$ s'écrit: $I_{C_{b'c}} = V_{b'e} / (j\omega C_{b'c})^{-1} = j\omega C_{b'c} V_{b'e}$
La loi des nœuds au collecteur donne: $I_{C_{b'c}} + I_c = g_m V_{b'e}$
d'où l'on tire l'expression du courant collecteur:

$$I_c = (g_m - j\omega C_{b'c}) V_{b'e}$$

Calculons maintenant l'expression du courant base.



$$I_b = V_{b'e} (R_{b'e} // C_{b'e} // C_{b'c})^{-1}$$

$$I_b = (1/R_{b'e} + j\omega (C_{b'e} + C_{b'c})) V_{b'e}$$

Le gain en courant s'écrit donc:

$$h_{21} = \beta = I_c/I_b = \{(g_m - j\omega C_{b'c}) V_{b'e}\} / \{(1/R_{b'e} + j\omega (C_{b'e} + C_{b'c})) V_{b'e}\}$$

$$h_{21} = (g_m - j\omega C_{b'c}) / (1/R_{b'e} + j\omega (C_{b'e} + C_{b'c}))$$

Aux fréquences pour lesquelles le modèle est valide g_m est très grand devant $C_{b'c} \omega$.

L'expression se simplifie donc:

$$h_{21} = g_m R_{b'e} / \{1 + j\omega R_{b'e} (C_{b'e} + C_{b'c})\}$$

Posons $\beta_o = g_m R_{b'e}$ et $\omega_\beta = 1 / R_{b'e} (C_{b'e} + C_{b'c})\}$

$$h_{21} = \beta_o / (1 + j\omega/\omega_\beta)$$

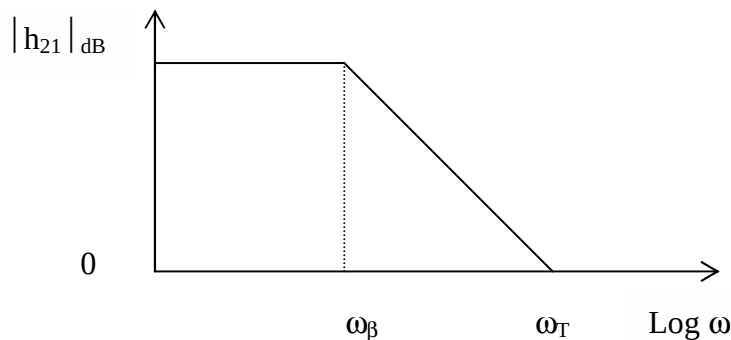
On peut représenter l'évolution de h_{21} en fonction de la fréquence dans le plan de Bode.

$$\begin{aligned} |h_{21}| &= \beta_o / \{1 + (\omega/\omega_\beta)^2\}^{1/2} \\ |h_{21}|_{dB} &= 20 \log \beta_o - 20 \log \{1 + (\omega/\omega_\beta)^2\}^{1/2} \end{aligned}$$

$$\text{Pour } \omega = 0 \Rightarrow h_{21} = \beta_o = g_m R_{b'e}$$

$$\text{Pour } \omega \gg \omega_\beta \quad |h_{21}|_{dB} = 20 \log \beta_o - 20 \log (\omega/\omega_\beta)$$

ce qui correspond à une direction asymptotique décroissant avec une pente de -6dB/octave (-20dB par décade)



5.3 FREQUENCE DE TRANSITION

Soit ω_T la pulsation correspondant à $|h_{21}|_{dB} = 0$ (c'est à dire $h_{21} = 1$)

$$|h_{21}| = \beta_o / \{1 + (\omega_T/\omega_\beta)^2\}^{1/2} = 1$$

$$(\omega_T/\omega_\beta)^2 \gg 1 \Rightarrow \beta_o \approx \omega_T / \omega_\beta$$

$$\text{d'où} \quad \omega_T = \beta_o \omega_\beta$$

ω_T est appelé produit gain-bande ou pulsation de transition.

En remplaçant β_o et ω_β par leur valeur, il vient:

$$\omega_T = g_m R_{b'e} / R_{b'e} (C_{b'e} + C_{b'c}) = g_m / (C_{b'e} + C_{b'c})$$

$f_T = \omega_T / 2\pi$ est la **fréquence de transition**.

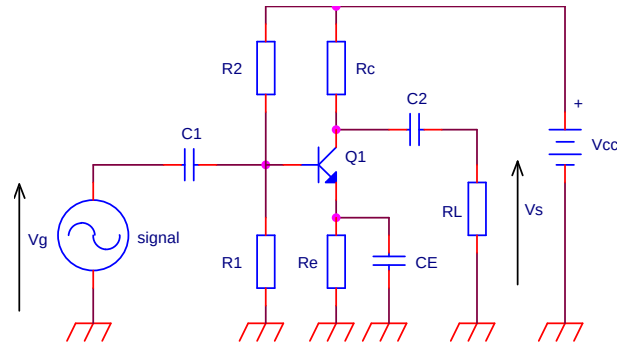
C'est un paramètre important qui est spécifié dans les notices de transistors.

Quelques exemples:

2N2222	Transistor d'usage général	$f_T = 250 \text{ MHz}$
BF199	Transistor RF	$f_T = 750 \text{ MHz}$
BFR91	Transistor UHF	$f_T = 5\,000 \text{ MHz}$

5.4 BANDE PASSANTE D'UN ETAGE AMPLIFICATEUR.

Considérons l'étage amplificateur en émetteur commun dont le schéma est le suivant:



La réponse en fréquence d'un tel amplificateur est limitée vers les basses fréquences par les capacités de liaison C_1 et C_2 ainsi que par la capacité de découplage d'émetteur C_E . Vers les hautes fréquences, c'est la chute du gain qui intervient.

Le gain en tension de l'étage s'écrit:

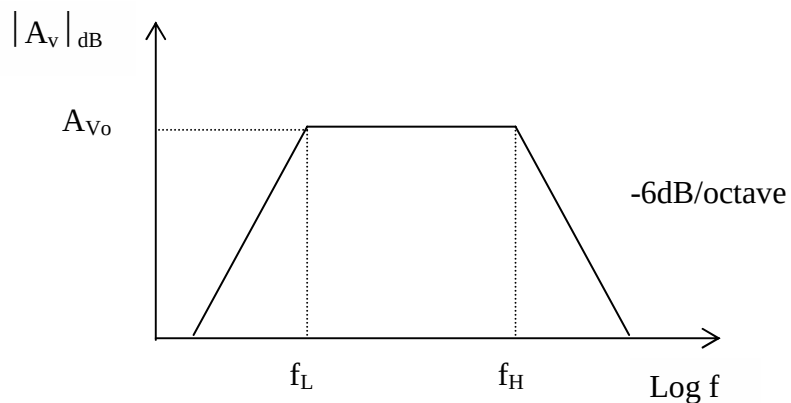
$$A_V = -\beta R_S / h_{11} \quad \text{avec } R_S = R_C // R_L$$

Soit en haute fréquence

$$A_V(\omega) = -\beta_o R_S / h_{11}(1+j(\omega/\omega_\beta))$$

On reconnaît une réponse du premier ordre avec une décroissance à -6dB/octave.

Le gain de l'amplificateur a donc l'allure suivante:



f_L est la fréquence de coupure basse (à -3dB)

f_H est la fréquence de coupure haute (à -3dB)

$f_H - f_L$ est la bande passante (à -3dB) de l'amplificateur

5.5 INFLUENCE DU BOITIER ET DES CONNEXIONS

Pour être utilisable, une puce de transistor doit être mise dans un boîtier et les différentes électrodes doivent être reliées au circuit extérieur par des fils. Ceci introduit une inductance parasite dont on doit tenir compte en haute fréquence.

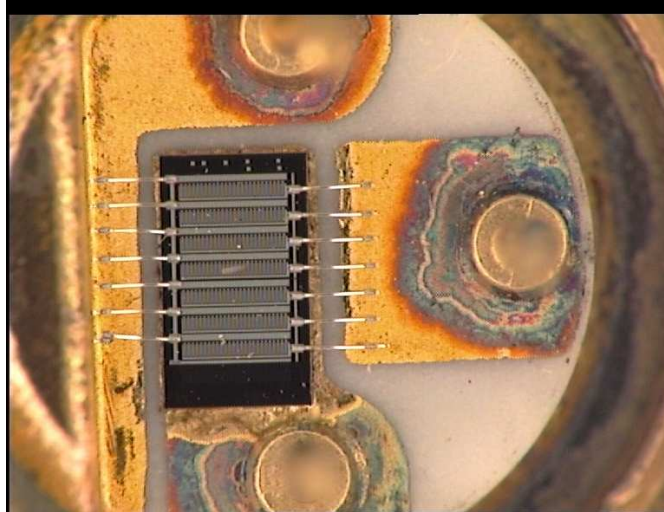
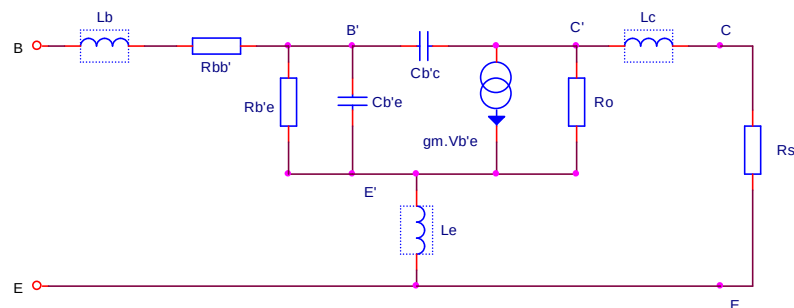


Photo interne d'un transistor de puissance RF 2N5070.

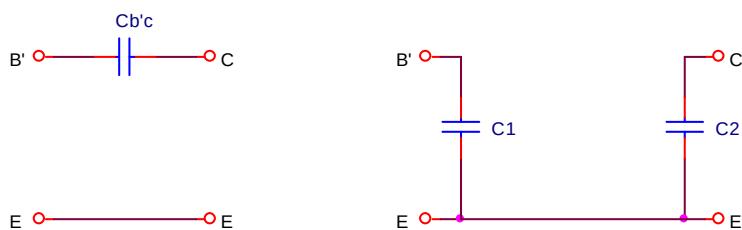
Le modèle complet du transistor dans son boîtier est alors le suivant:



Simplification du modèle

Ce modèle peut être simplifié en appliquant la transformation de Miller (voir démonstration en annexe) à la capacité $C_{b'c}$.

Cette transformation permet de remplacer $C_{b'c}$ par deux capacités connectées en dérivation à l'entrée et à la sortie du transistor (voir transformation de Miller en annexe).



Soit R_S la résistance de charge en petits signaux du transistor.

Nous avons vu que le gain de l'étage est $A_V = -\beta R_S / h_{11} = -g_m R_S$

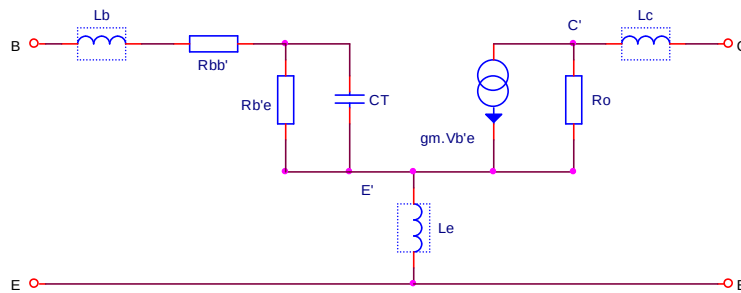
La capacité $C_{b'c}$ peut donc être remplacée par :

- A l'entrée une capacité $C_1 = C_{b'c} \cdot (1 - A_V) = C_{b'c} \cdot (1 + g_m R_S)$
- A la sortie une capacité $C_2 = C_{b'c} \cdot (1 - 1/A_V) = C_{b'c} \cdot (1 + 1/g_m R_S)$

La capacité C_2 est beaucoup plus faible que C_1 et peut être souvent négligée.

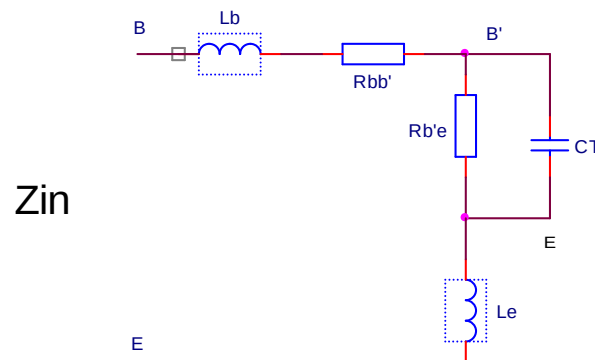
La capacité C_1 se trouve en parallèle avec la capacité $C_{b'e}$.

Nous avons donc une capacité équivalente $C_T = C_{b'e} + C_1$



5.6 IMPEDANCE D'ENTREE

Le modèle précédent permet d'établir une expression de l'impédance d'entrée du transistor



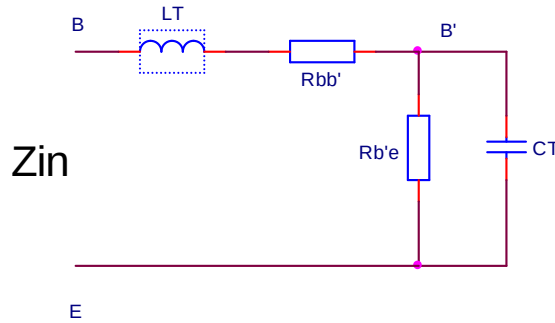
$$Z_{in} = jL_b\omega + R_{bb'} + (R_{b'e} / jC_T\omega) / (R_{b'e} + 1/jC_T\omega) + j\omega L_e$$

$$Z_{in} = j(L_b + L_e)\omega + R_{bb'} + R_{b'e} / (1 + jR_{b'e}C_T\omega)$$

En posant $L_T = L_e + L_b$ on a :

$$Z_{in} = jL_T\omega + R_{bb'} + R_{b'e} / (1 + jR_{b'e}C_T\omega)$$

Ce qui correspond au schéma équivalent suivant:



Cette expression peut se mettre sous la forme série $Z_{in} = R_S + jX_S$
ou sous la forme parallèle. $Z_{in} = R_p // jX_p$

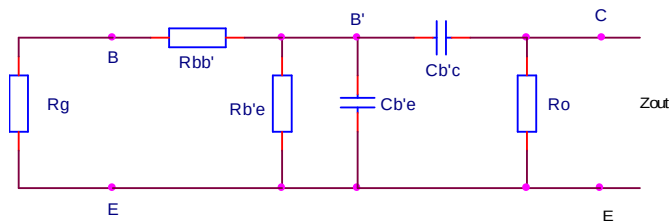


On remarque que la réactance sera soit capacitive pour des fréquences inférieures à la fréquence de résonance (donnée par $L_T C_T \omega_0^2 = 1$), soit inductive au dessus.

On notera que la valeur de la capacité C_1 ramenée à l'entrée par effet Miller dépend de la charge. Il y a donc toujours une certaine réaction de la sortie sur l'entrée par le biais de la capacité $C_{B'C}$.

5.7 IMPEDANCE DE SORTIE

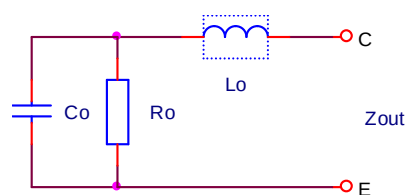
Revenons au schéma équivalent hybrid- π du transistor.



Cette représentation montre clairement que l'impédance de sortie décroît avec la fréquence et qu'elle dépend de la résistance R_g de la source (réaction de l'entrée sur la sortie).

En opérant comme précédemment, on pourra mettre l'impédance de sortie sous forme de résistance et de réactance série ou parallèle.

D'une manière générale on obtiendra le circuit équivalent suivant:



5.8 MODELE GENERAL DU TRANSISTOR

Un transistor bipolaire pourra donc être représenté comme un quadripôle présentant une impédance d'entrée complexe Z_{in} (sous forme série ou parallèle) et dont la sortie est une source de tension ou de courant dépendant de l'entrée, d'impédance complexe Z_{out} .

Le gain maximal d'un amplificateur haute fréquence à transistor ne pourra être obtenu que s'il y a adaptation des impédances entre la source et l'impédance d'entrée d'une part, et entre l'impédance de sortie et la charge d'autre part.

6 TRANSISTOR A EFFET DE CHAMP MOS

Le modèle du MOSFET utilisé pour l'étude de l'amplificateur en basse fréquence ne tient pas compte des divers effets capacitifs. Aux fréquences élevées ces effets ne sont plus négligeables et il faut modifier le schéma équivalent en conséquence.

La principale capacité du MOSFET est la capacité grille-canal.

Lorsque le MOSFET est utilisé dans sa zone de fonctionnement résistif (zone triode) c'est-à-dire à tension Drain-Source très faible, le canal est uniforme et la capacité grille-canal s'écrit :

$$C_{g \text{ canal}} = W L C_{ox}$$

Où W est la largeur du canal, L sa longueur et C_{ox} la capacité grille-substrat par unité de surface.

On modélise généralement cette capacité par deux capacités égales connectées entre grille et source et entre grille et drain.

$$C_{gs} = C_{gd} = \frac{1}{2} W L C_{ox}$$

Lorsque le MOSFET fonctionne dans sa zone linéaire, le canal n'est plus uniforme et prend une allure penchée avec pincement côté drain.

On démontre alors que la capacité grille-canal devient égale à :

$$C_{g-canal} = \frac{2}{3} W L C_{ox}$$

et on la modélise par une seule capacité entre grille et source.

En réalité, la capacité grille-source est un peu plus élevée car la métallisation de grille recouvre légèrement la source .

De la même manière, il existe une capacité grille-drain due à ce recouvrement.

La capacité drain-substrat (capacité d'une jonction bloquée) peut en général être négligée en première approximation.

Il en résulte le schéma équivalent suivant:

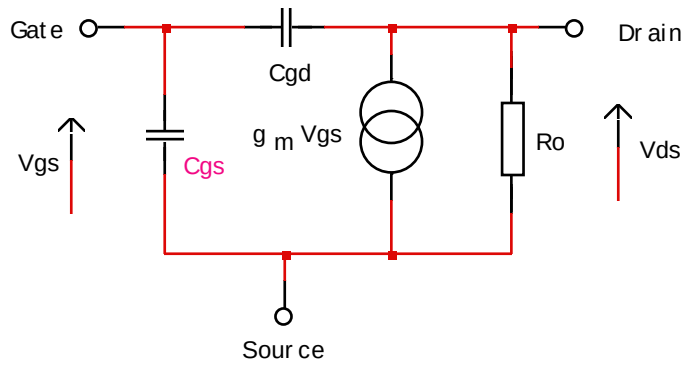


Schéma équivalent du MOSFET à haute fréquence

6.1 FREQUENCE DE TRANSITION

Pour caractériser le fonctionnement à haute fréquence d'un transistor on utilise souvent la fréquence de transition f_T qui est définie comme la fréquence à laquelle l'amplitude du gain en courant, sortie en court-circuit, devient égale à 1.

En utilisant le modèle haute fréquence on obtient le schéma suivant:

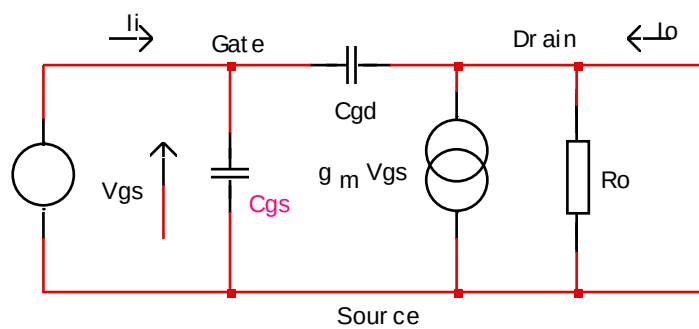


Schéma utilisé pour déterminer la fréquence de transition du MOSFET.

$$I_o = g_m V_{gs}$$

$$V_{gs} = I_i / j (\omega C_{gs} + \omega C_{gd})$$

$$I_o / I_i = g_m / j\omega (C_{gs} + C_{gd})$$

$$\text{pour } \omega = \omega_T = 2\pi f_T \text{ on a } I_o / I_i = 1$$

$$f_T = g_m / 2\pi (C_{gs} + C_{gd})$$

On note que la fréquence de transition est proportionnelle à la transconductance et inversement proportionnelle aux capacités internes du MOSFET.

f_T varie selon la technologie de quelques centaines de MHz à quelques GHz.

6.2 EFFET MILLER

Lorsqu'un MOSFET est connecté en source-commune, sa capacité grille-drain C_{gd} se trouve connectée entre l'entrée et la sortie de l'étage.

Par application du théorème de MILLER (voir annexe) , cette capacité peut-être transformée en deux capacités connectées en parallèle avec l'entrée et la sortie.

Considérons par exemple un étage amplificateur de gain -100 réalisé avec un MOSFET dont la capacité grille-drain est de 1pF.

Le théorème de Miller permet de remplacer cette capacité par une capacité C_1 connectée entre grille et source:

$$C_1 = C_{gd} (1 - K) = 1 (1 - (-100)) = 101 \text{ pF}$$

et une capacité C_2 connectée entre drain et source:

$$C_2 = C_{gd} (1 - 1/K) = 1 (1 - \frac{1}{-100})$$

$$C_2 = 1,01 \text{ pF} \quad \text{qui peut-être négligée dans la plupart des cas.}$$

Cet exemple montre l'influence extrême de la capacité "entrée-sortie".

Cet effet de multiplication de la capacité entrée-sortie porte le nom d'**effet Miller**.

6.3 REPONSE EN FREQUENCE D'UN AMPLIFICATEUR A MOSFET

Considérons l'amplificateur dont le schéma est donné ci après:

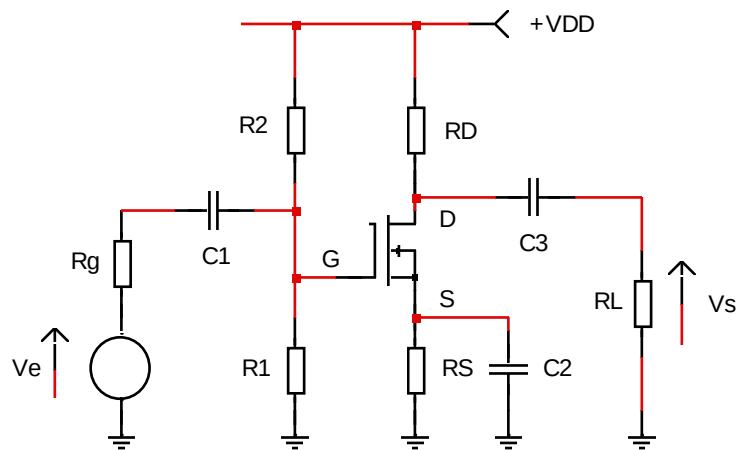


Schéma d'un étage amplificateur à MOSFET

On considère que la capacité des condensateurs C_1 , C_2 et C_3 est très grande de manière à ce que l'on puisse négliger leur impédance aux fréquences considérées.

Le schéma équivalent du circuit est le suivant:

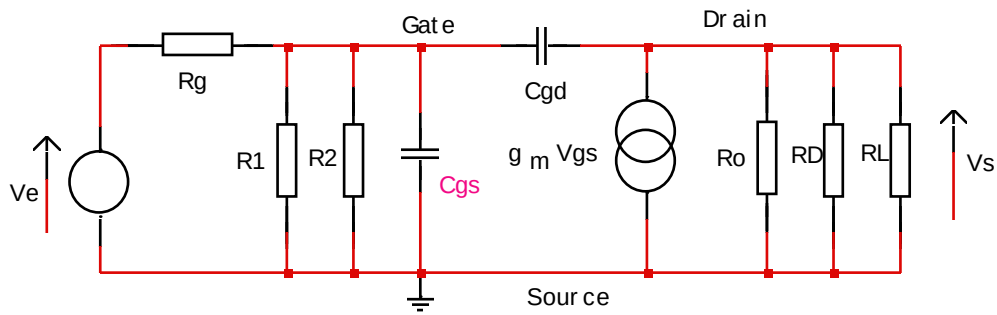


Schéma équivalent du circuit amplificateur

En combinant les résistances en parallèle :

$$R_{in} = R_1 // R_2$$

$$R'_L = R_O // R_D // R_L$$

on peut simplifier le schéma équivalent

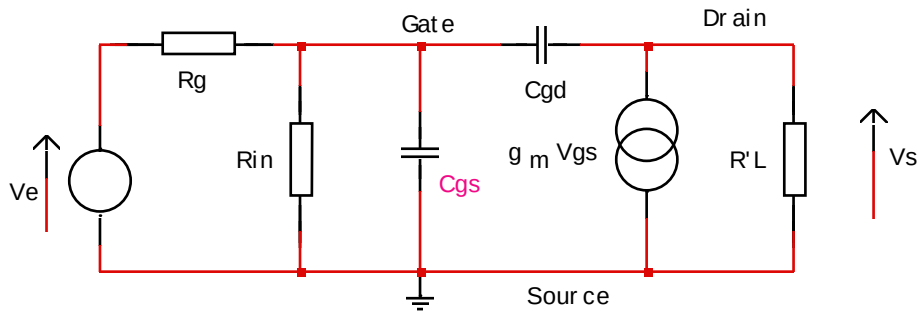


Schéma équivalent simplifié du circuit amplificateur

Par application du théorème de Thévenin à l'entrée et du théorème de Miller à la capacité C_{gd} on peut encore simplifier le schéma équivalent .

Si la capacité C_{gd} est faible, le courant qui la traverse est très faible devant la source de courant contrôlée $g_m V_{gs}$ et on peut le négliger.

La tension de sortie V_S peut alors s'écrire

$$V_S \cong - g_m V_{gs} R'_L$$

En appliquant le théorème de Miller, on peut remplacer la capacité C_{gd} par une capacité équivalente C_{eq} connectée entre grille et source :

$$C_{eq} = C_{gd} (1 + g_m R'_L)$$

Le schéma équivalent devient alors :

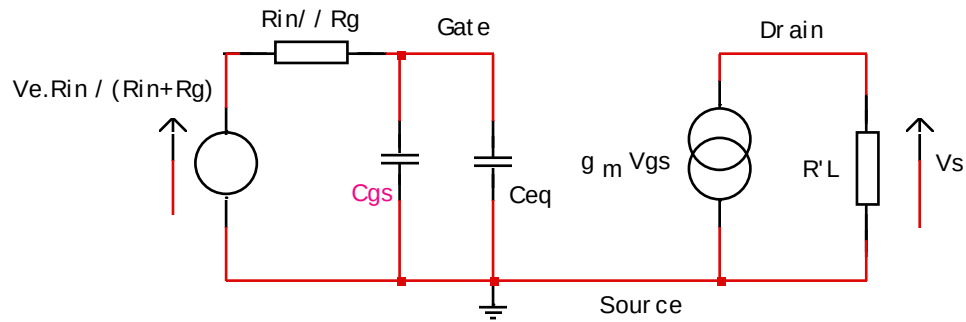


Schéma équivalent réduit de l'étage amplificateur

On voit que le circuit d'entrée constitue un filtre passe-bas du 1^{er} ordre qui détermine la réponse à haute fréquence de l'étage.

Soit $C_T = C_{eq} + C_{gs}$

$$C_T = C_{gd}(1 + g_m R'_L) + C_{gs}.$$

Nous avons $R' = R_g / R_{in}$.

La fréquence de coupure haute à - 3dB est alors :

$$f_H = \frac{1}{2\pi C_T R'}$$

Si $\omega_H = \frac{1}{C_T R'}$ est la pulsation de coupure haute, l'expression du gain à haute fréquence de l'étage s'écrira :

$$A(j\omega) = A_o \frac{1}{1 + j \frac{\omega}{\omega_H}}$$

On notera le rôle très important joué par la capacité grille-drain dans la réponse à haute fréquence de l'étage amplificateur, à cause de l'effet Miller.

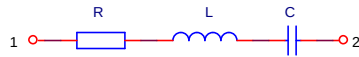
Pour des fonctionnements à très haute fréquence, on utilisera des configurations spéciales qui minimisent l'effet Miller, telles que les montages en grille commune ou les montages cascodes (MOSFET double gate).

CHAPITRE III

LES CIRCUITS RESONANTS

1 CIRCUIT RESONANT SERIE

Soit le circuit résonant série suivant :



L'impédance Z_{12} de ce circuit s'écrit :

$$Z_{12} = R + jL\omega + 1/jC\omega$$

Ou
$$Z_{12} = R + (1 - LC\omega^2)/jC\omega$$

L'impédance Z_{12} sera minimale et égale à R pour $\omega = \omega_0$ tel que $LC\omega_0^2 = 1$

Soit

$$\omega_0 = \frac{1}{\sqrt{LC}} \quad f_0 = \frac{1}{2\pi\sqrt{LC}}$$

f_0 est la **fréquence de résonance**.

1.1 COEFFICIENT DE SURTENSION Q

Nous avons :

$$Z_{12} = R + jL\omega + 1/jC\omega$$

$$LC\omega_0^2 = 1$$

On peut donc écrire :

$$Z_{12} = R + jL\omega + LC\omega_0^2/jC\omega$$

$$Z_{12} = R + jL(\omega - \omega_0^2/\omega)$$

Plaçons-nous à une fréquence $\omega = \omega_0 + \Delta\omega$ proche de la résonance.

$$Z_{12} = R + jL(\omega_0 + \Delta\omega - \omega_0^2/(\omega_0 + \Delta\omega))$$

$$Z_{12} = R + jL(\omega_0 + \Delta\omega - \omega_0 \frac{1}{1 + \frac{\Delta\omega}{\omega_0}})$$

On sait que si x est petit $1/(1+x) \approx 1 - x$, d'où pour $\Delta\omega/\omega_0$ petit

$$Z_{12} \approx R + jL[\omega_0 + \Delta\omega - \omega_0(1 - \Delta\omega/\omega_0)]$$

$$Z_{12} \approx R + 2jL\Delta\omega = R(1 + 2j\Delta\omega L/R)$$

Que l'on peut mettre sous la forme :

$$Z_{12} = R \left(1 + jL \frac{\omega_0}{R} \frac{2\Delta\omega}{\omega_0} \right)$$

Pour $\Delta\omega=0$ $Z_{12} = R$

Si

$$L \frac{\omega_0}{R} \frac{2\Delta\omega}{\omega_0} = 1$$

le module de Z_{12} vaut $|Z_{12}| = R\sqrt{2}$ et la phase de Z_{12} est égale à 45°

Posons $Q = L\omega_0/R$

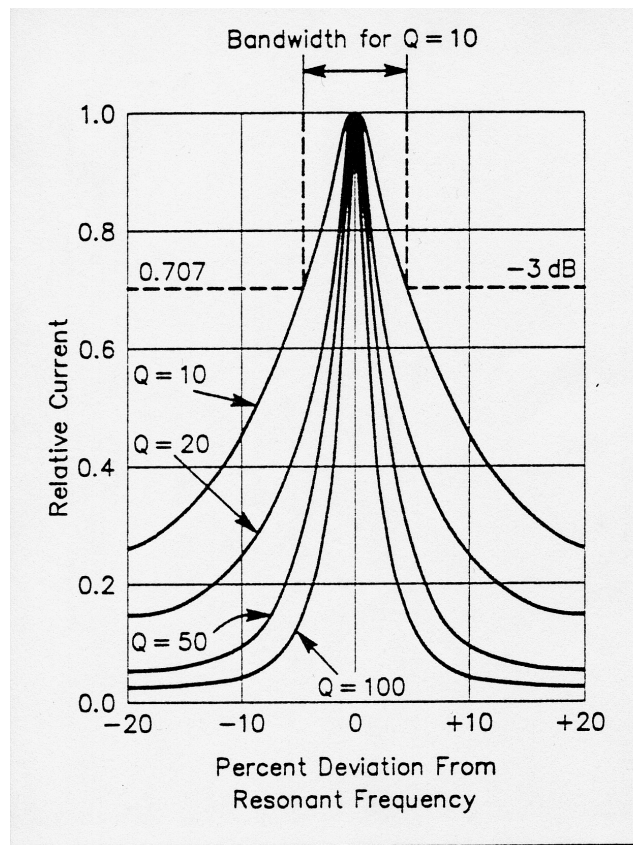
Q est appelé **facteur de surtension** (ou de qualité). Il exprime la qualité du circuit résonant.

Pour $\Delta\omega = \omega_0/2Q$ $Z_{12} = R\sqrt{2} \exp(\pm j45^\circ)$ ou $Z_{12} = R\sqrt{2} \angle \pm 45^\circ$

Si l'on applique une tension constante aux bornes du circuit, à $\omega = \omega_0 \pm \omega_0/2Q$, le courant dans le circuit résonant a chuté d'un facteur $1/\sqrt{2}$ par rapport à sa valeur à la résonance, soit une chute de -3dB. La puissance est alors divisée par 2.

On pose $B = 2\Delta\omega = \omega_0/Q$

B est la **bande passante à -3 dB** du circuit résonant.

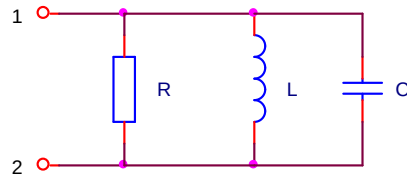


Valeur du courant dans un circuit résonant série pour diverses valeurs de Q

On notera que l'impédance Z_{12} du circuit résonant présente une réactance :

- capacitive pour $\omega < \omega_0$ ($X_s < 0$)
- inductive pour $\omega > \omega_0$ ($X_s > 0$)

2 CIRCUIT RESONANT PARALLELE



Soit $G = 1/R$ la conductance et $Y_{12} = 1/Z_{12}$ l'admittance du circuit.

$$Y_{12} = G + 1/jL\omega + jC\omega$$

$$Y_{12} = G + \frac{1 - jLC\omega^2}{jL\omega}$$

L'admittance Y_{12} sera minimale pour $\omega = \omega_0$ tel que $LC\omega_0^2 = 1$

On a alors : $Y_{12}(\omega_0) = G$ soit $Z_{12}(\omega_0) = R$

$$f_0 = \frac{1}{2\pi\sqrt{LC}}$$

f_0 est la **fréquence de résonance**

2.1 COEFFICIENT DE SURTENSION Q

L'admittance du circuit peut s'écrire :

$$Y_{12} = G + jC(\omega - \omega_0^2/\omega)$$

En se plaçant à une fréquence $\omega = \omega_0 + \Delta\omega$ proche de la résonance on a :

$$Y_{12} = G + jC \left(\omega_0 + \Delta\omega - \frac{\omega_0}{1 + \frac{\Delta\omega}{\omega_0}} \right)$$

Que l'on peut écrire : $Y_{12} \approx G + j2C\Delta\omega$

$$Y_{12} = G \left(1 + \frac{jC\omega_0}{G} \frac{2\Delta\omega}{\omega_0} \right)$$

A la résonance $\Delta\omega=0$ et $Y_{12} = G$

Si l'on se place à une fréquence telle que :

$$\frac{C\omega_0}{G} \frac{2\Delta\omega}{\omega_0} = 1$$

On a $|Y_{12}(\omega)| = G\sqrt{2}$ et donc $|Z_{12}(\omega)| = R/\sqrt{2}$

Si on alimente le circuit résonant par une source de courant, la tension aux bornes du circuit à cette fréquence chute de $1/\sqrt{2}$ par rapport à sa valeur à la résonance, soit -3dB .

La puissance est alors divisée par 2.

On définit le facteur de surtension Q par :

$$Q = C\omega_0/G = RC\omega_0$$

Sachant que

$$LC\omega_0^2 = 1$$

On en déduit que :

$$Q = RC\omega_0 = R/L\omega_0$$

En comparant ce résultat avec celui du circuit résonant série, on constate qu'il peut y avoir une ambiguïté au niveau de la définition de Q . On utilisera donc de préférence la notation suivante :

- Pour le circuit résonant série

$$Q_s = L\omega_0/R_s = 1/R_s C\omega_0$$

dans ce cas R_s est faible

- Pour le circuit résonant parallèle

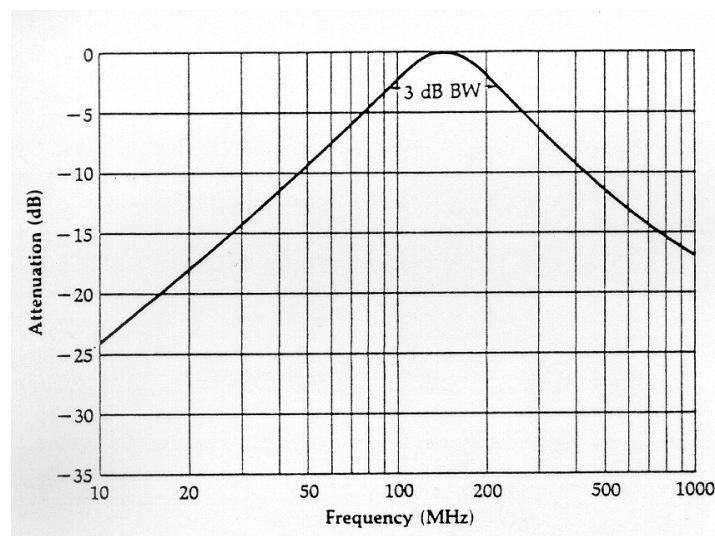
$$Q_p = R_p C\omega_0 = R_p/L\omega_0$$

dans ce cas R_p est forte

La **bande passante à -3 dB** du circuit résonant parallèle est définie par :

$$2Q_p\Delta\omega = \omega_0 \Rightarrow \Delta\omega = \omega_0/2Q_p$$

$$B = 2\Delta\omega = \omega_0/Q_p$$



Exemple de réponse en fréquence d'un circuit résonant parallèle

3 RECAPITULATION

	Résonance série	Résonance parallèle
Condition de résonance	$LC\omega_0^2 = 1$	$LC\omega_0^2 = 1$
Impédance à la résonance	$Z_s = R_s \angle 0^\circ$	$Z_p = R_p \angle 0^\circ$
Facteur de surtension	$Q_s = L\omega_0/R_s = 1/R_s C\omega_0$	$Q_p = R_p/L\omega_0 = R_p C\omega_0$
Bande passante à -3dB	$B = \omega_0/Q_s$	$B = \omega_0/Q_p$

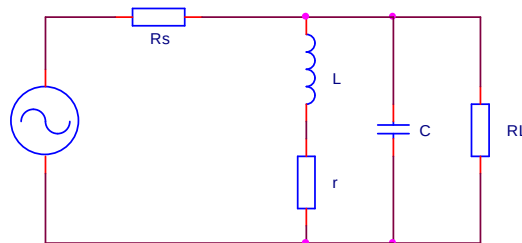
4 FACTEUR DE QUALITE EN CHARGE (LOADED Q)

Le facteur de surtension Q d'un circuit résonant série ou parallèle peut être défini comme le rapport de la fréquence de résonance sur la bande passante à -3dB .

$$Q = \omega_0/B \quad Q \text{ est appelé facteur de surtension en charge.}$$

Cette définition est intéressante car dans la pratique, un circuit résonant est toujours relié à une source et à une charge dont les impédances modifient le facteur de surtension du circuit pris tout seul.

Le Q en charge représente le facteur de surtension du circuit résonant dans son environnement.



R_s : Résistance interne de la source

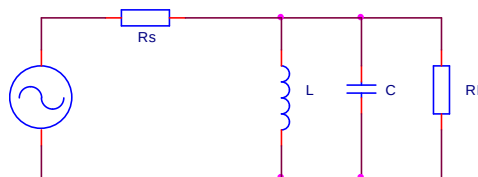
R_L : Résistance de charge

r : Résistance représentant les pertes dans l'inductance

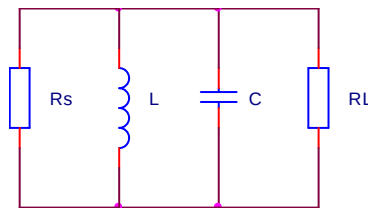
Le Q du circuit résonant tout seul est $Q_1 = L\omega_0/r$

En charge, les résistances R_s et R_L viennent se rajouter en parallèle (Thévenin) et baissent fortement la valeur de Q .

Supposons pour simplifier que le circuit soit sans perte, ce qui revient à négliger r et à se ramener au circuit suivant :



Qui est équivalent pour le calcul de Q au schéma ci dessous (application du théorème de Thévenin):



Soit $R_p = R_s // R_L$

$$R_p = R_s R_L / (R_s + R_L)$$

$$Q_p = R_p / X_p$$

$$X_p = L\omega_0 = 1/C\omega_0$$

réactance capacitive ou inductive

On voit que toute diminution de R_p diminue le Q en charge du circuit.

Exemple

On dispose d'une source sinusoïdale de fréquence $f_0 = 50\text{MHz}$ et de résistance interne $R_S = 150\Omega$, et d'une charge résistive $R_L = 1000\Omega$.

Calculer un circuit résonant parallèle présentant un $Q = 20$ à la résonance à 50 MHz .

Donner la bande passante à -3dB .

La résistance parallèle équivalente est : $R_p = 150 \cdot 1000 / (150 + 1000) = 130\Omega$

La réactance est donnée par : $Q_p = R_p / X_p \Rightarrow X_p = R_p / Q_p$

$$X_p = 130 / 20 = 6,5\Omega$$

A la résonance on a :

$$X_p = L\omega_0 = 1/C\omega_0$$

$$\omega_0 = 2\pi f_0 = 314 \cdot 10^6 \text{ rad/s}$$

On en déduit la valeur de l'inductance :

$$L = X_p / \omega_0 = 6,5 / 314 \cdot 10^6 = 20,7 \text{ nH}$$

et de la capacité :

$$C = 1 / X_p \omega_0 = 10^{-6} / 6,5 \cdot 314 = 490 \text{ pF}$$

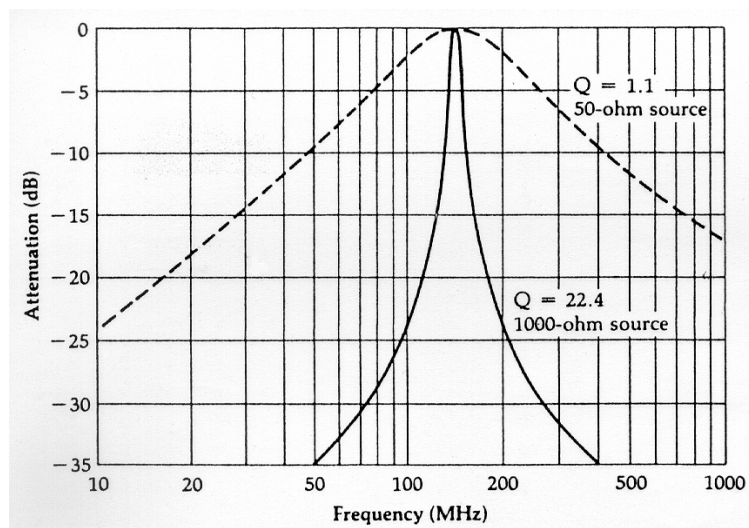
La bande passante est donnée par :

$$Q_p = \omega_0 / B_\omega = f_0 / B_f$$

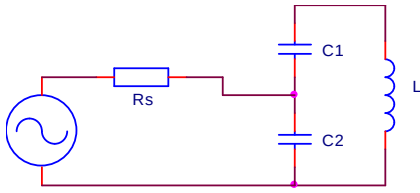
$$B = 50 \cdot 10^6 / 20 = 2,5 \text{ MHz}$$

5 TRANSFORMATION D'IMPEDANCE

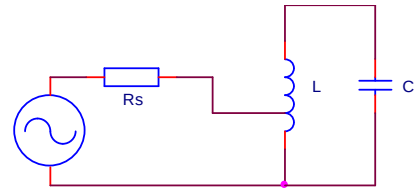
Nous avons vu que de faibles valeurs de R_S et R_L tendent à diminuer le Q du circuit, ce qui a pour effet d'élargir la bande passante.



Une solution pour remédier à ce problème consiste à créer une prise intermédiaire sur le circuit résonant.

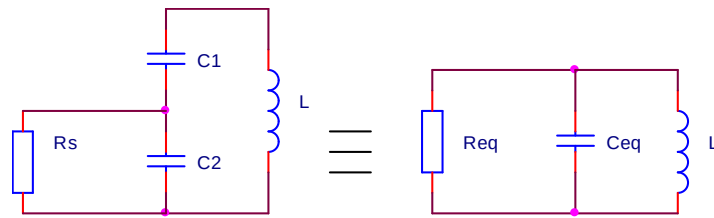


Prise intermédiaire capacitive

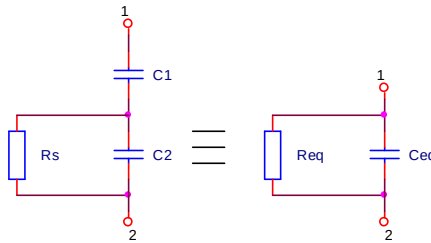


Prise intermédiaire inductive

5.1 PRISE INTERMEDIAIRE CAPACITIVE



Nous allons calculer le Q de ce circuit, qui est équivalent au circuit représenté à droite. L'inductance L et la fréquence de résonance sont identiques dans les deux circuits. Le problème revient donc au cas suivant :



Calculons l'impédance Z_{12} du circuit.

$$Z_{12} = \frac{1}{jC_1\omega} + \frac{R_s \frac{1}{jC_2\omega}}{R_s + \frac{1}{jC_2\omega}}$$

$$Z_{12} = \frac{1}{jC_1\omega} + \frac{R_s}{1 + jR_sC_2\omega}$$

$$Z_{12} = \frac{1 + jR_sC_1\omega + jR_sC_2\omega}{(1 + jR_sC_2\omega)jC_1\omega}$$

$$Z_{12} = \frac{1 + jR_sC_1\omega + jR_sC_2\omega}{jC_1\omega - R_sC_1C_2\omega^2}$$

$$Z_{12} = \frac{1 + jR_s\omega(C_1 + C_2)}{C_1\omega(j - R_sC_2\omega)}$$

L'admittance $Y_{12} = 1/Z_{12}$ s'écrit donc :

$$Y_{12} = \frac{C_1 \omega (j - R_S C_2 \omega)}{1 + j R_S \omega (C_1 + C_2)}$$

Multiplions le numérateur et le dénominateur par le conjugué du dénominateur :

$$Y_{12} = \frac{j C_1 \omega (1 + j R_S C_2 \omega) (1 - j R_S \omega (C_1 + C_2))}{1 + R_S^2 \omega^2 (C_1 + C_2)^2}$$

$$Y_{12} = \frac{R_S C_1^2 \omega^2 + j C_1 \omega (1 + R_S C_2 \omega^2 (C_1 + C_2))}{1 + R_S^2 \omega^2 (C_1 + C_2)^2}$$

Pour le circuit équivalent, nous avons :

$$Y_{12} = G_{eq} + j C_{eq} \omega \quad \text{avec} \quad G_{eq} = 1/R_{eq}$$

En identifiant les deux expressions, il vient pour la partie réelle :

$$G_{eq} = \frac{R_S C_1^2 \omega^2}{1 + R_S^2 \omega^2 (C_1 + C_2)^2}$$

Dans la pratique

$$R_S^2 \omega^2 (C_1 + C_2)^2 \gg 1$$

et donc

$$G_{eq} \approx \frac{1}{R_S} \left(\frac{C_1}{C_1 + C_2} \right)^2$$

Posons

$$n = \frac{C_1}{C_1 + C_2}$$

$$G_{eq} = n^2 / R_S \quad \text{d'où} \quad R_{eq} = R_S / n^2$$

On remarque que la résistance de source est divisée par n^2 . Le circuit se comporte comme un transformateur de rapport $1/n$.

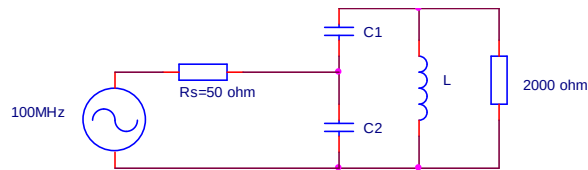
Les deux circuits ayant la même fréquence de résonance et la même inductance L , la valeur de la capacité équivalente est égale à C_1 en série avec C_2 :

$$C_{eq} = C_1 C_2 / (C_1 + C_2)$$

On obtient un résultat similaire avec une prise intermédiaire sur l'inductance L (considérer l'inductance comme un autotransformateur)

Exemple

Considérons le circuit suivant :



Calculer les éléments du circuit pour un Q de 20 à la résonance (100MHz) et un transfert maximal d'énergie (adaptation d'impédance).

Pour qu'il y ait transfert maximal d'énergie, il faut que la résistance de la source soit égale à celle de la charge, c'est à dire que la résistance équivalente ramenée en parallèle sur le circuit doit être à R_L

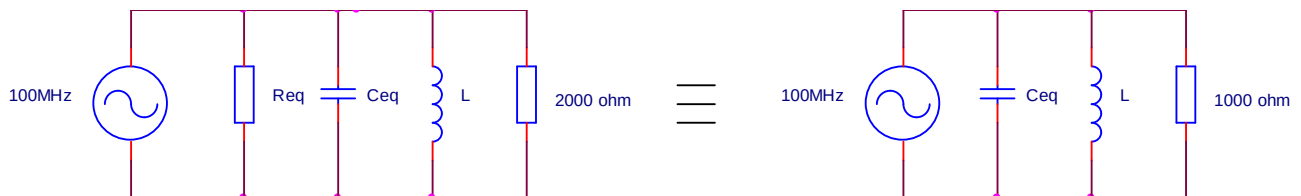
$$R_{eq} = R_L = 2\,000\Omega$$

$$R_{eq} = R_S / n^2 = R_S [(C_1 + C_2) / C_1]^2$$

$$R_{eq} = R_S (1 + C_2 / C_1)^2 \quad (1 + C_2 / C_1)^2 = R_{eq} / R_S$$

$$\frac{C_2}{C_1} = \sqrt{\frac{R_{eq}}{R_S}} - 1 \quad \frac{C_2}{C_1} = \sqrt{\frac{2000}{50}} - 1 = 5,32$$

Nous obtenons donc le schéma équivalent suivant :



$$R_{eq} = R_L = 2\,000\Omega$$

$$R_{eq} // R_L = 1\,000\Omega$$

On désire un Q en charge de 20

$$Q_P = R_P / X_P$$

$$X_P = R_P / Q_P = 1000 / 20 = 50\Omega$$

A la résonance

$$X_P = L\omega_0 = 1 / C_{eq}\omega_0$$

d'où

$$L = X_P / \omega_0 = 50 / (2\pi \cdot 100 \cdot 10^6) = 79,6 \text{ nH}$$

$$C_{eq} = 1 / X_P \omega_0 = 1 / (50 \cdot 2\pi \cdot 100 \cdot 10^6) = 31,8 \text{ pF}$$

Nous avons par ailleurs :

$$C_{eq} = C_1 C_2 / (C_1 + C_2) = 31,8 \text{ pF}$$

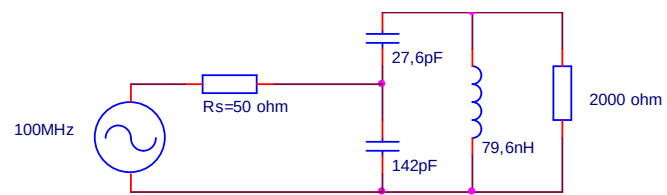
$$C_2 = 5,32 C_1$$

La résolution de ce système donne :

$$C_1 = 26,7 \text{ pF}$$

$$C_2 = 142 \text{ pF}$$

Le circuit final est donc le suivant :



CHAPITRE IV

ADAPTATION D'IMPEDANCE

1 PUISSANCE MAXIMALE FOURNIE PAR UN GENERATEUR RELIE A UNE CHARGE QUELCONQUE.

Soit un générateur de F.E.M E et de résistance interne $Z_1 = R_1 + jX_1$, relié à une charge d'impédance $Z_2 = R_2 + jX_2$.

Le courant traversant le circuit a pour valeur :

$$I = E/(Z_1 + Z_2) = E/\{(R_1 + R_2) + j(X_1 + X_2)\}$$

dont le module est :

$$|I| = \frac{E}{\sqrt{(R_1 + R_2)^2 + (X_1 + X_2)^2}}$$

La puissance utile reçue par la charge Z_2 est donc :

$$P = R_2 I^2 = \frac{R_2 E^2}{(R_1 + R_2)^2 + (X_1 + X_2)^2}$$

On voit que la puissance P sera maximale si $X_1 + X_2 = 0$, c'est à dire $X_2 = -X_1$

Les deux réactances doivent être de signe contraire, l'impédance totale du circuit est alors purement résistive.

On a alors :

$$P = \frac{R_2 E^2}{(R_1 + R_2)^2}$$

Calculons la valeur de R_2 qui rend P maximale :

$$\frac{dP}{dR_2} = E^2 \left[\frac{(R_1 + R_2)^2 - 2R_2(R_1 + R_2)}{(R_1 + R_2)^4} \right] = 0$$

$$\frac{dP}{dR_2} = E^2 \left[\frac{(R_1 + R_2)(R_1 - R_2)}{(R_1 + R_2)^4} \right] = 0$$

D'où $R_2 = R_1$

La puissance reçue par la charge Z_2 est donc maximale si : $Z_2 = R_1 - jX_1$
C'est à dire si **l'impédance de la charge est conjuguée de celle du générateur.**

On dit que les impédances sont adaptées.

La puissance maximale reçue par la charge est alors : $P_{\max} = E^2/4R_2$
Elle est égale à celle dissipée dans le générateur.

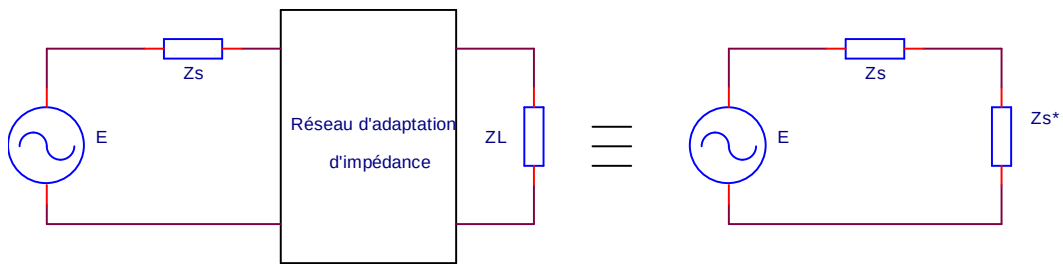
2 CIRCUITS D'ADAPTATION D'IMPEDANCE

Nous avons vu que le transfert de puissance entre une source et une charge est maximal lorsque l'impédance de la charge est complexe conjuguée de celle de la source.

$$Z_L = Z_S^*$$

Etant donné une source et une charge d'impédances quelconques, le problème est d'obtenir le transfert maximal de puissance.

On utilise pour cela un réseau d'adaptation d'impédance, dont le but est de transformer l'impédance de la charge de manière à ce que l'impédance vue par la source soit sa conjuguée.



Il existe une infinité de solutions à ce problème. La plus simple fait appel au réseau en L.

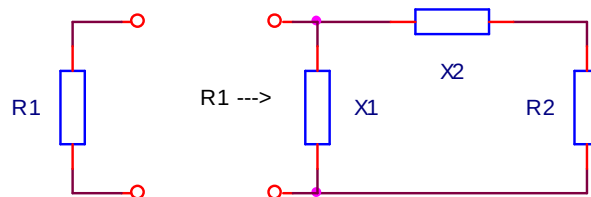
3 CIRCUIT D'ADAPTATION EN L

Ce circuit se compose de deux réactances constituant les branches d'un L.

La branche shunt doit toujours être en parallèle avec l'impédance la plus forte.

3.1 TRANSFORMATION D'UNE RESISTANCE PURE

On désire transformer une résistance R_2 en résistance R_1 en utilisant un circuit en L.



On suppose que $R_1 > R_2$

Soit n tel que $R_2 = R_1/n$ avec $n > 1$

- L'impédance du circuit vue de R_1 s'écrit : $R_1 = jX_1 // (R_2 + jX_2)$

$$R_1 = \frac{jX_1(R_2 + jX_2)}{R_2 + j(X_1 + X_2)}$$

Séparons partie réelle et partie imaginaire.

On multiplie numérateur et dénominateur par $R_2 - j(X_1 + X_2)$

$$R_1 = \frac{jX_1(R_2 + jX_2)(R_2 - j(X_1 + X_2))}{(R_2 + j(X_1 + X_2))(R_2 - j(X_1 + X_2))}$$

$$R_1 = \frac{R_2 X_1^2}{R_2^2 + (X_1 + X_2)^2} + j \frac{R_2^2 X_1 + X_1 X_2 (X_1 + X_2)}{R_2^2 + (X_1 + X_2)^2}$$

R_1 étant réelle, on a

$$R_2^2 X_1 + X_1 X_2 (X_1 + X_2) = 0$$

Soit $R_2^2 = -X_2(X_1 + X_2) \quad (1)$

- Considérons maintenant l'impédance du circuit vue de R_2 :

$$R_2 = (R_1 // X_1) + X_2$$

$$R_2 = \frac{jR_1 X_1 + jX_2(R_1 + jX_1)}{R_1 + jX_1}$$

$$R_2 = \frac{-X_1 X_2 + jR_1(X_1 + X_2)}{R_1 + jX_1}$$

$$R_2 = \frac{(-X_1 X_2 + jR_1(X_1 + X_2))(R_1 - jX_1)}{R_1^2 + X_1^2}$$

$$R_2 = \frac{R_1 X_1^2 + j(R_1^2(X_1 + X_2) + X_1^2 X_2)}{R_1^2 + X_1^2}$$

Or R_2 est réelle :

$$R_1^2(X_1 + X_2) + X_1^2 X_2 = 0$$

D'où l'équation (2) :

$$R_1^2 = -\frac{X_1^2 X_2}{X_1 + X_2}$$

L'expression de R_2 s'écrit donc :

$$R_2 = \frac{R_1 X_1^2}{R_1^2 + X_1^2}$$

On en déduit la valeur de X_1 :

$$X_1^2 = \frac{R_2 R_1^2}{R_1 - R_2}$$

$$X_1 = \mp R_1 \sqrt{\frac{R_2}{R_1 - R_2}}$$

Ou encore

$$X_1 = \mp R_1 \frac{1}{\sqrt{n-1}}$$

Avec $n = R_1/R_2$ et $n > 1$ n est le rapport de transformation du circuit.

Nous avons établi les relations suivantes (1) et (2):

$$R_2^2 = -X_2(X_1 + X_2)$$

$$R_1^2 = -\frac{X_1^2 X_2}{X_1 + X_2}$$

Effectuons le produit $R_1^2 R_2^2$, il vient : $(R_1 R_2)^2 = (X_1 X_2)^2$
 R_1 et R_2 étant réelles, le produit $R_1 R_2$ est toujours positif. Il n'y a donc qu'une seule solution
 $R_1 R_2 = -X_1 X_2$ (X_1 et X_2 sont de signes contraires), d'où :

$$R_1 R_2 = -X_1 X_2 \quad \text{et} \quad R_1/X_1 = -X_2/R_2$$

On en déduit la valeur de X_2 : $X_2 = -R_1 R_2 / X_1$

Soit en remplaçant X_1 par sa valeur :

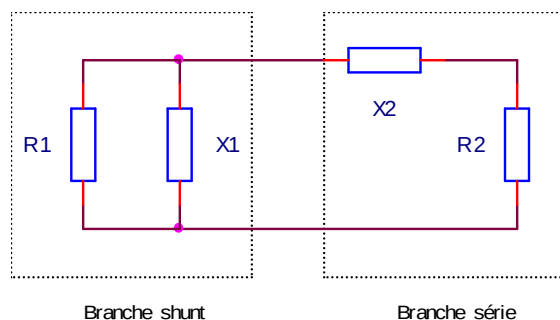
$$X_2 = \mp \frac{R_1 R_2}{R_1 \sqrt{\frac{R_2}{R_1 - R_2}}}$$

$$X_2 = \mp \sqrt{R_2 (R_1 - R_2)}$$

Ou encore :

$$X_2 = \mp R_2 \sqrt{n-1}$$

Le circuit final est donc le suivant :



3.2 FACTEUR DE QUALITE DU CIRCUIT

Le facteur de qualité de la branche série est : $Q_s = Q_2 = X_2/R_2$

Et celui de la branche shunt est $Q_p = Q_1 = R_1/X_1$

On a montré précédemment que : $R_1/X_1 = -X_2/R_2$

$$|Q_s| = |Q_p|$$

Ceci n'est valable qu'à la fréquence pour laquelle il y a adaptation.

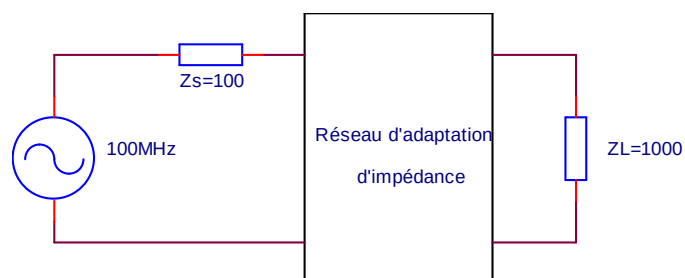
$$Q_p = Q_2 = \mp \frac{\sqrt{R_2(R_1 - R_2)}}{R_2} = \mp \sqrt{\frac{R_1}{R_2} - 1}$$

Soit d'une manière générale :

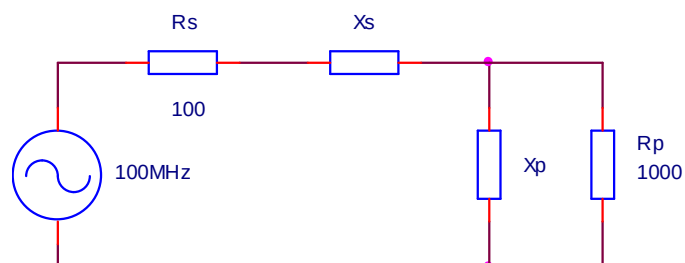
$$Q_s = Q_p = \mp \sqrt{\frac{R_p}{R_s} - 1} = \mp \sqrt{n - 1}$$

Exemple

Soit une charge de $1\,000\Omega$ à adapter à une source de 100Ω à la fréquence $f_0 = 100\text{MHz}$.



La branche shunt du circuit en L est toujours en parallèle avec l'impédance la plus forte. Le circuit est donc le suivant :



$$Q_S = Q_P = \mp \sqrt{\frac{R_P}{R_S} - 1} = \mp \sqrt{n - 1}$$

$$Q_S = Q_P = \mp \sqrt{\frac{1000}{100} - 1} = \mp \sqrt{9} = \mp 3$$

$$Q_S = X_S/R_S \quad \Rightarrow \quad X_S = Q_S R_S$$

$$X_S = \pm 3.100 = \pm 300\Omega$$

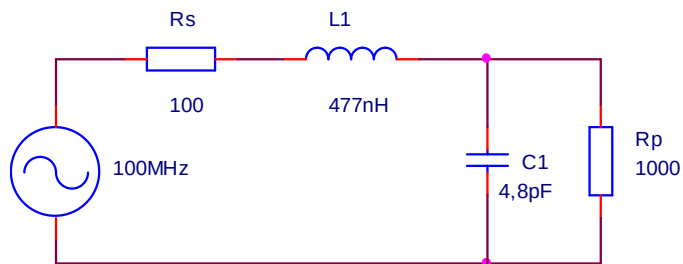
$$Q_P = R_P/X_P \quad \Rightarrow \quad X_P = R_P/Q_P$$

$$X_P = \pm 1000/3 = \pm 333\Omega$$

X_S et X_P doivent être de signe contraire, il y a donc deux solutions :

$$\begin{array}{ll} X_{S1} = 300\Omega & X_{P1} = -333\Omega \\ X_{S2} = -300\Omega & X_{P2} = 333\Omega \end{array}$$

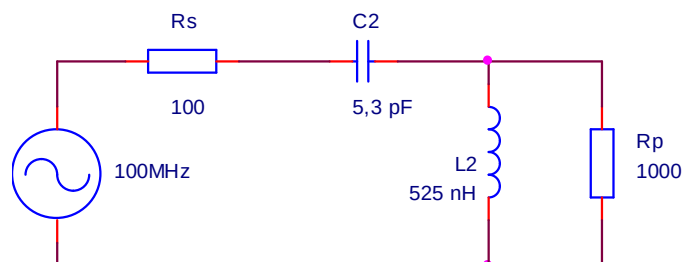
Solution 1 : circuit passe bas.



$$L_1 = X_{S1}/2\pi f_0 = 300/6,28.100.10^6 = 477.10^{-9} \text{ H} = 477\text{nH}$$

$$C_1 = 1/2\pi f_0 X_{P1} = 1/2\pi.100.10^6.333 = 4,8.10^{-12} \text{ F} = 4,8 \text{ pF}$$

Solution 2 :circuit passe haut



$$C_2 = 1/2\pi f_0 X_{S2} = 1/2\pi.100.10^6.300 = 5,3 \text{ pF}$$

$$L_2 = X_{P2}/2\pi f_0 = 333/2\pi.100.10^6 = 530.10^{-9} \text{ H} = 530 \text{ nH}$$

3.3 CAS OU LES RESISTANCES TERMINALES NE SONT PAS DES RESISTANCES PURES

On peut se ramener au cas précédent en intégrant la réactance dans le réseau d'adaptation (absorption) ou en l'annulant par une réactance de signe opposé (résonance).

Exemple

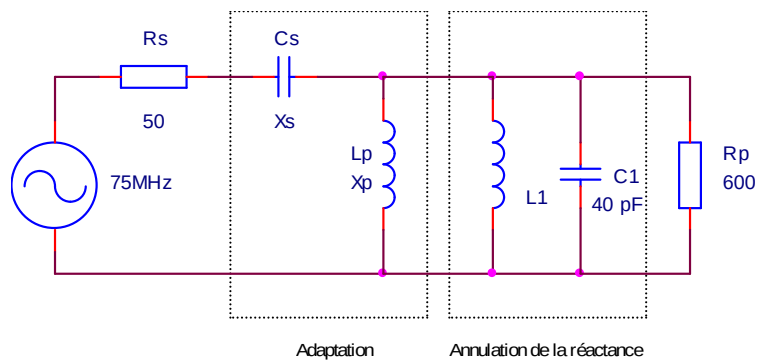
On désire adapter une charge de 600Ω en parallèle avec 40 pF à une source d'impédance 50Ω et de fréquence 75 MHz . On veut que le circuit ne laisse pas passer le courant continu.

Cette dernière condition impose d'avoir un condensateur en série, c'est à dire de choisir le circuit en L de type passe haut.

Pour annuler la composante réactive de la charge, on place en parallèle une réactance égale et de signe opposé.

Ici la charge est capacitive, on place donc une inductance L_1 en parallèle de manière à ce que l'ensemble forme un circuit résonant parallèle à 75 MHz .

Le schéma du circuit est alors le suivant :



- **Calcul de L_1**

La condition de résonance donne : $L_1 C_1 \omega^2 = 1$

$$L_1 = 1/C_1(2\pi f_0)^2 = 1/(2\pi \cdot 75 \cdot 10^6)^2 \cdot 40 \cdot 10^{-12} = 112,6\text{ nH}$$

- **Calcul du réseau d'adaptation**

Nous devons calculer un réseau en L passe haut permettant de passer de 50Ω à 600Ω .

$$Q_s = Q_p = \mp \sqrt{\frac{R_p}{R_s} - 1} = \mp \sqrt{n - 1}$$

$$Q_s = Q_p = \sqrt{\frac{600}{50} - 1} = 3,32$$

$$X_s = Q_s R_s = 3,32 \cdot 50 = 166\Omega$$

$$\text{D'où : } C_s = 1/X_s \omega = 1/X_s 2\pi f_0 = 1/2\pi \cdot 75 \cdot 10^6 \cdot 166 = 12,8\text{ pF}$$

$$X_p = R_p/Q_p = 600/3,32 = 181\Omega$$

$$L_p = X_p/2\pi f_0 = 181/2\pi \cdot 75 \cdot 10^6 = 384 \text{ nH}$$

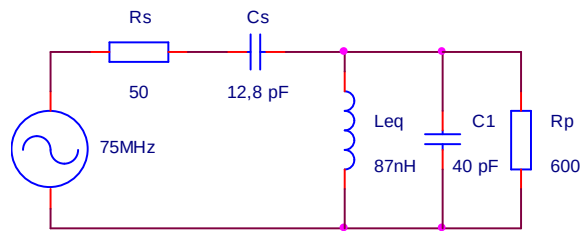
- **Calcul de l'inductance équivalente**

Nous avons deux inductances, L1 et LP, en parallèle.
L'inductance équivalente est :

$$L_{eq} = L_1 L_p / (L_1 + L_p)$$

$$L_{eq} = 112,6 \times 384 / (112,6 + 384) = 87 \text{ nH}$$

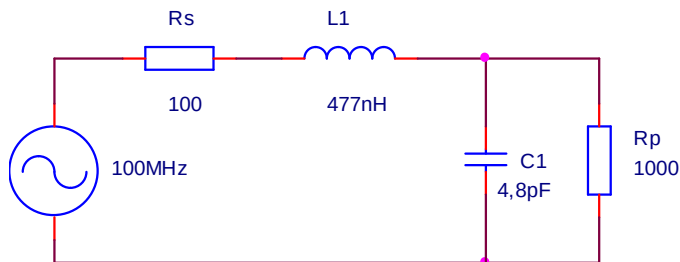
Le circuit final est donc le suivant :



3.4 BANDE PASSANTE ET FACTEUR DE QUALITE

Le circuit en L procure une adaptation d'impédance à la fréquence f_0 .
Dès que l'on s'écarte de cette fréquence, l'adaptation n'est plus parfaite, comme nous allons le voir sur un exemple.

Reprenons le circuit étudié précédemment :



$$Q_s = Q_p = 3$$

L'adaptation est calculée pour $f_0 = 100 \text{ MHz}$.

Calculons l'impédance vue par la source à $f = 95 \text{ MHz}$ et $f = 105 \text{ MHz}$.

$$Z_{in} = jL\omega + \frac{R_L \frac{1}{jC\omega}}{R_L + \frac{1}{jC\omega}}$$

$$Z_{in} = jL\omega + \frac{R_L}{1 + jCR_L\omega}$$

A $f = 95 \text{ MHz}$ on a :

$Z_{in} = 109,6 - j 27,4$ ce qui correspond à une réactance capacitive

A $f = 105 \text{ MHz}$ on a :

$Z_{in} = 91,5 + j26,6$ ce qui correspond à une réactance inductive.

Le facteur de qualité Q du circuit a une importance considérable sur la bande passante.

**Plus le facteur de qualité Q du circuit est élevé, plus la bande passante est étroite.
Pour un circuit en L, il n'existe qu'une seule valeur de Q permettant l'adaptation.**

La bande passante est donc fixée par le rapport de transformation n .

Le choix de Q ne peut s'effectuer que si l'on utilise un circuit à trois (ou plus) réactances.

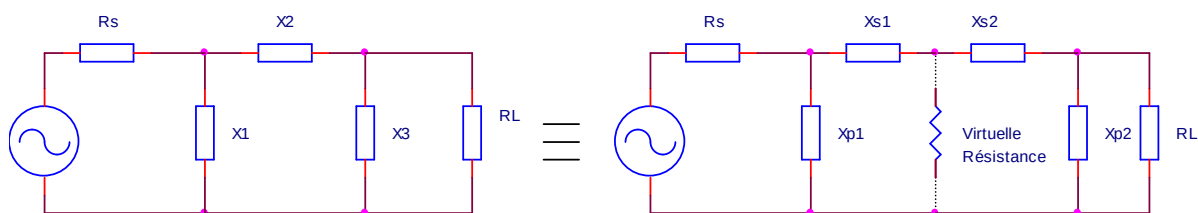
Les réseaux les plus utilisés sont les réseaux à trois réactances et T ou en Π .

4 RESEAUX D'ADAPTATION EN TE OU EN PI

Les réseaux d'adaptation en T ϵ ou en Pi permettent la réalisation de circuits d'adaptation d'impédance dont le facteur de qualité Q peut prendre n'importe quelle valeur à condition qu'elle soit supérieure à celle du réseau en L assurant la même fonction.

4.1 CIRCUIT D'ADAPTATION D'IMPEDANCE EN PI

Le réseau en Π peut être décomposé en deux réseaux en L montés en cascade et procurant une adaptation à une résistance virtuelle R_V située entre les deux.



Du fait de la configuration des circuits en L, la résistance virtuelle doit être inférieure à R_S et R_L .

Dans la pratique on peut considérer que le Q du réseau est donné par la cellule située du côté de la résistance la plus forte :

$$Q = \sqrt{\frac{R_H}{R_V} - 1}$$

R_H : La plus grande des deux résistances R_S ou R_L

R_V : Résistance virtuelle

Exemple

Calculer un réseau en Pi permettant d'adapter une source de 100Ω à une charge de 1000Ω avec un Q en charge de 15.

Cet exemple a été traité précédemment dans le cas du réseau en L. Le facteur de qualité du circuit en L étant de 3, le circuit en Pi avec Q de 15 est faisable.

$$Q = \sqrt{\frac{R_H}{R_V} - 1}$$

Ici $R_L > R_S$, donc $R_H = R_L$

$$Q = \sqrt{\frac{R_L}{R_V} - 1}$$

$$R_V = R_L/(Q^2 + 1) = 1000/(15^2 + 1) = 4,42\Omega$$

- Calculons le circuit en L côté charge.

$$X_{P2} = R_P/Q_P = R_L/Q = 1000/15 = 66,7\Omega$$

$$X_{S2} = Q_S R_S = Q R_V = 15 \cdot 4,42 = 66,3\Omega$$

- Calculons maintenant le circuit en L côté source.

Il s'agit de passer de R_S à R_V .

$$Q' = \sqrt{\frac{R_S}{R_V} - 1} = \sqrt{\frac{100}{4,42} - 1} = 4,6$$

$$X_{P1} = R_P/Q' = R_S/Q'$$

La résistance de source R_S est considérée ici comme une résistance shunt car elle se trouve en parallèle avec X_{P1} (Thévenin).

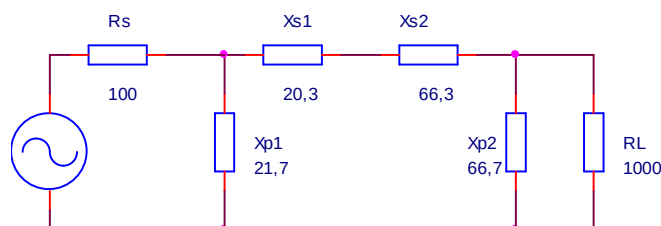
$$X_{P1} = 100/4,6 = 21,7\Omega$$

$$X_{S1} = Q' R_{Série} = Q R_V$$

La résistance R_V est considérée comme résistance série puisqu'elle se trouve connectée en série avec X_{S1}

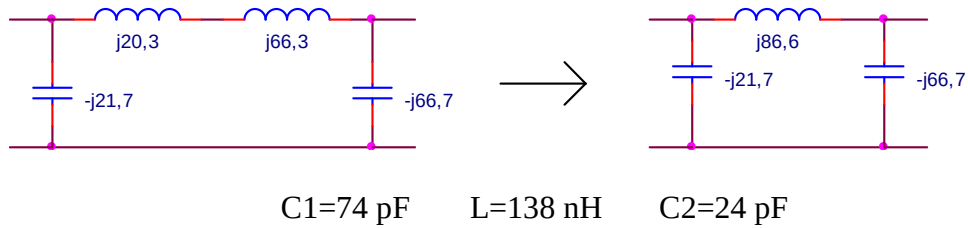
$$X_{S1} = 4,6 \cdot 4,42 = 20,3\Omega$$

Le circuit final aura donc l'allure suivante :

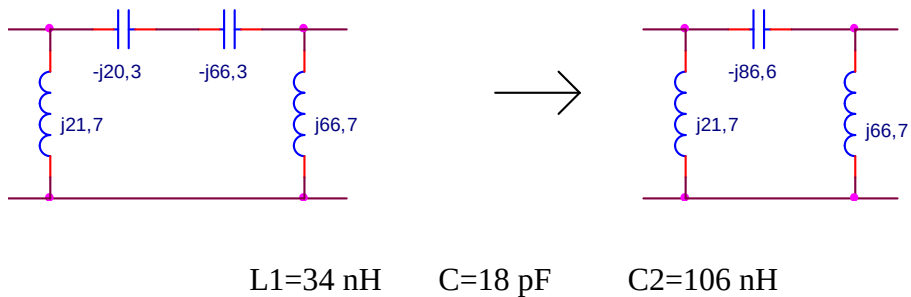


Comme il y a deux solutions pour chaque circuit en L (passe haut et passe bas), nous aurons quatre solutions en tout :

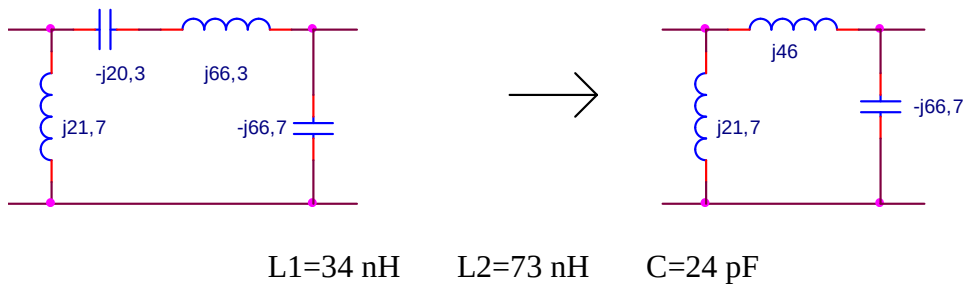
Solution 1



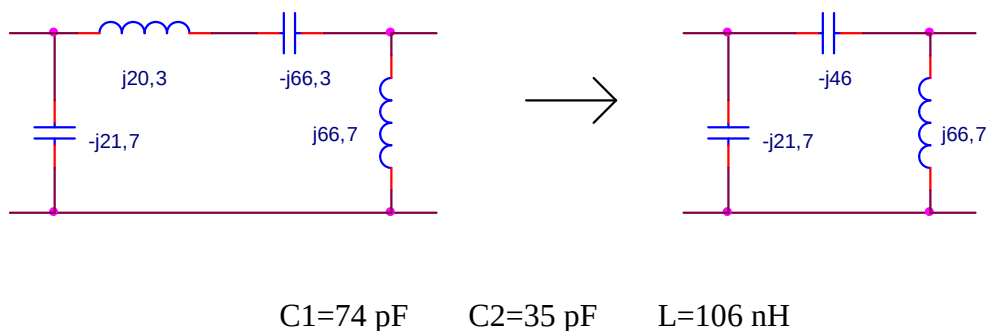
Solution 2



Solution 3



Solution 4



Calcul de la valeur des éléments (à $F = 100 \text{ MHz}$):

$-j21,7 \Omega \rightarrow 74 \text{ pF}$	$j21,7 \Omega \rightarrow 34 \text{ nH}$
$-j66,7 \Omega \rightarrow 24 \text{ pF}$	$j66,7 \Omega \rightarrow 106 \text{ nH}$
$-j86,6 \Omega \rightarrow 18 \text{ pF}$	$j86,6 \Omega \rightarrow 138 \text{ nH}$
$-j46 \Omega \rightarrow 35 \text{ pF}$	$j46 \Omega \rightarrow 73 \text{ nH}$

CHAPITRE V

CHANGEMENT DE FREQUENCE

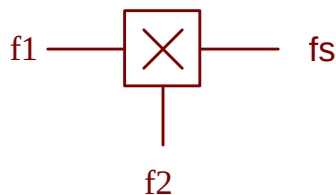
Le changement (ou transposition) de fréquence est une opération fondamentale en radiocommunications. Elle consiste à traduire le spectre d'un signal, soit vers une fréquence plus basse afin de le traiter plus facilement (down conversion), soit vers une fréquence plus élevée afin de le transmettre par voie hertzienne (up conversion).

Cette opération fait appel à un oscillateur local qui fixe la valeur de la translation en fréquence et à un dispositif mélangeur (mixer) qui fonctionne selon deux grands principes: multiplication linéaire et commutation.

1 MELANGEUR MULTIPLICATIF

Considérons deux signaux sinusoïdaux de fréquences f_1 et f_2 .

Ces deux signaux sont appliqués aux entrées d'un dispositif multiplicateur.



$$e_1(t) = A_1 \sin \omega_1 t$$

$$e_2(t) = A_2 \sin \omega_2 t$$

$$e_s(t) = e_1(t) \cdot e_2(t) = A_1 A_2 \sin \omega_1 t \cdot \sin \omega_2 t$$

d'après la relation trigonométrique $\sin a \cdot \sin b = \frac{1}{2} [\cos(a-b) - \cos(a+b)]$ on a :

$$e_s(t) = (A_1 A_2 / 2) \cos(\omega_1 - \omega_2) t - (A_1 A_2 / 2) \cos(\omega_1 + \omega_2) t$$

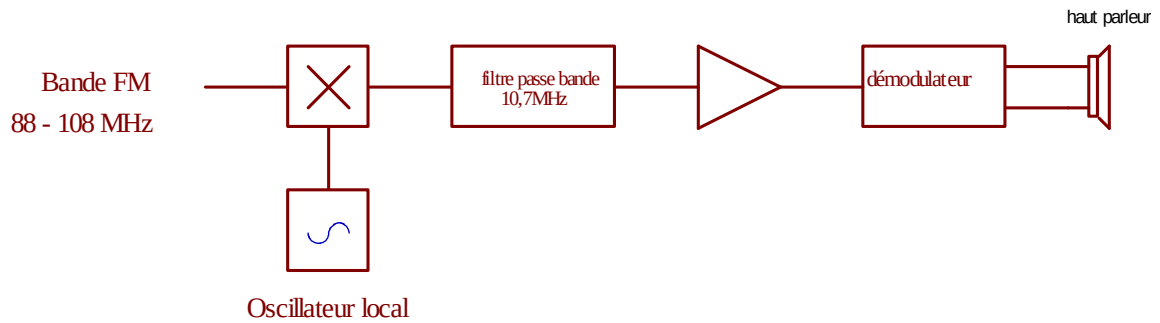
battement inférieur battement supérieur

On constate que le signal de sortie se compose de deux termes; dont les fréquences sont $f_1 - f_2$ et $f_1 + f_2$, que l'on pourra séparer par filtrage.

Si l'on sélectionne le battement inférieur $f_1 - f_2$, on obtient un changement de fréquence infradyne (downconversion). Dans le cas contraire, le changement de fréquence est dit supradyné (upconversion).

Exemple d'application : Récepteur radio FM

Le schéma synoptique d'un récepteur de radiodiffusion FM est le suivant :



Le signal capté par l'antenne est filtré, puis appliqué à un mélangeur relié à un oscillateur local dont on peut ajuster la fréquence. La sortie du mélangeur passe à travers un filtre passe bande centré sur 10,7 MHz et dont la bande passante est de quelques dizaines de kHz. Après amplification, le signal est démodulé puis appliqué à un haut parleur.

On désire par exemple recevoir une station dont la fréquence porteuse est $f_{RF} = 100\text{MHz}$.

L'oscillateur local est réglé sur $f_{OL} = 89,3\text{ MHz}$.

Le mélangeur (multiplicateur) fournit en sortie deux signaux de fréquences :

$$\begin{aligned} f_1 &= f_{RF} - f_{OL} = 100 - 89,3 = 10,7\text{ MHz} \\ f_2 &= f_{RF} + f_{OL} = 100 + 89,3 = 189,3\text{ MHz} \end{aligned}$$

Le filtre passe bande calé sur 10,7MHz élimine aisément le battement supérieur.

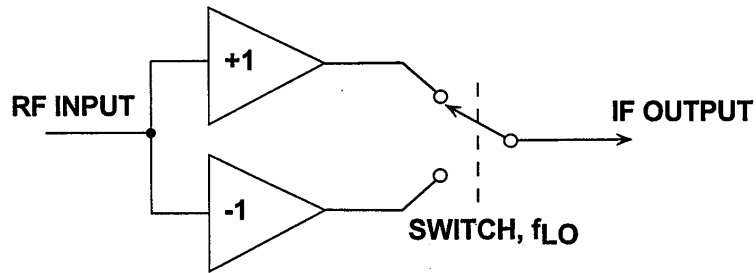
Après amplification, le battement inférieur appelé **fréquence intermédiaire** FI est démodulé et le signal audio amplifié est restitué par l'intermédiaire d'un haut parleur.

On notera qu'un signal incident de fréquence $f_{\text{image}} = f_{OL} - f_{\text{intermédiaire}}$ donnera également un battement dans la bande passante du filtre :

$$f_{\text{image}} = 89,3 - 10,7 = 78,6\text{ MHz}.$$

Cette fréquence est appelée **fréquence image** et doit être éliminée par filtrage avant le changement de fréquence.

2 MELANGEUR A COMMUTATION



Le signal RF d'entrée est séparé en deux composantes, l'une en phase et l'autre en opposition de phase. Un commutateur électronique sélectionne alternativement l'une ou l'autre voie à la fréquence de l'oscillateur local.

Nous obtenons en sortie le produit du signal d'entrée $S_{RF}(t) = A \cdot \sin \omega_{RF} t$ par un signal rectangulaire symétrique d'amplitude unitaire et de fréquence f_{LO} .

Ce dernier peut être décomposé en série de Fourier:

$$S_{LO}(t) = (4/\pi) [\sin \omega_{LO} t - (1/3) \sin 3\omega_{LO} t + (1/5) \sin 5\omega_{LO} t - \dots]$$

Le signal de sortie $S_{IF}(t) = S_{RF}(t) \cdot S_{LO}(t)$ s'écrit alors:

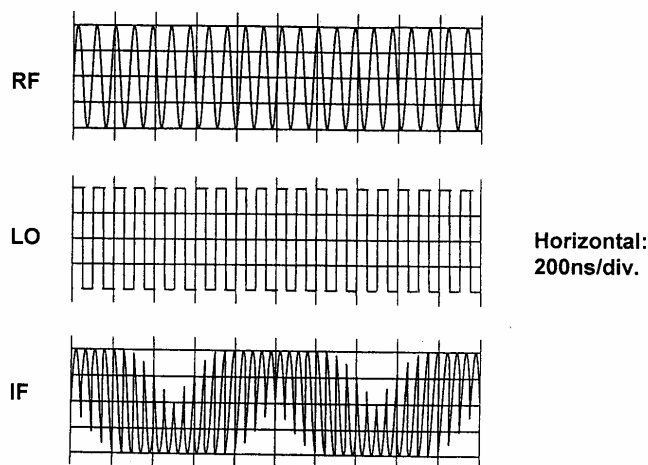
$$S_{IF}(t) = (4A/\pi) [\sin \omega_{LO} t \cdot \sin \omega_{RF} t - (1/3) \sin 3\omega_{LO} t \cdot \sin \omega_{RF} t + (1/5) \sin 5\omega_{LO} t \cdot \sin \omega_{RF} t - \dots]$$

$$S_{IF}(t) = (2A/\pi) [\cos(\omega_{RF} - \omega_{LO})t - \cos(\omega_{RF} + \omega_{LO})t - (1/3)(\cos(\omega_{RF} - 3\omega_{LO})t - \cos(\omega_{RF} + 3\omega_{LO})t) + (1/5)(\cos(\omega_{RF} - 5\omega_{LO})t - \cos(\omega_{RF} + 5\omega_{LO})t) + \dots]$$

Le coefficient $2/\pi$ montre que le mélange s'effectue avec une perte d'insertion de :

$$20 \log(2/\pi) = 3,92 \text{ dB}$$

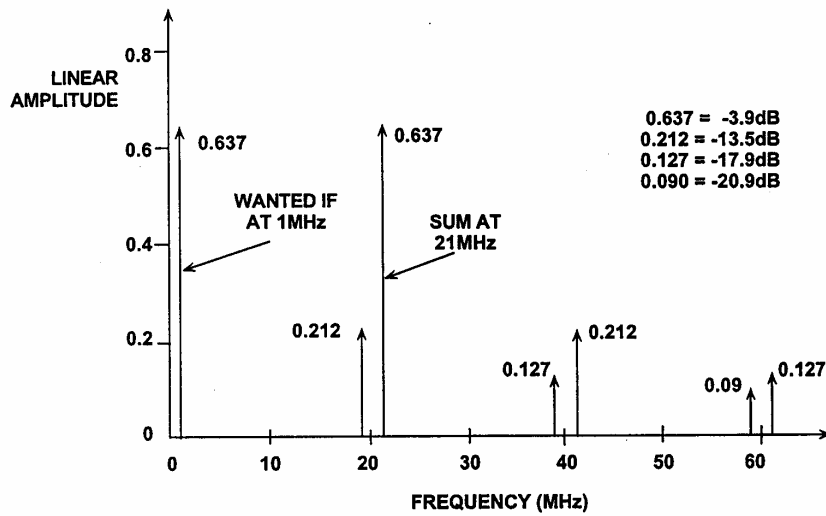
INPUTS AND OUTPUT FOR IDEAL SWITCHING MIXER FOR $f_{RF} = 11 \text{ MHz}$, $f_{LO} = 10 \text{ MHz}$



On remarque l'apparition de battements entre la fréquence d'entrée ω_{RF} et les harmoniques impairs de la fréquence de l'oscillateur local ω_{LO} .

Le spectre du signal de sortie est alors le suivant:

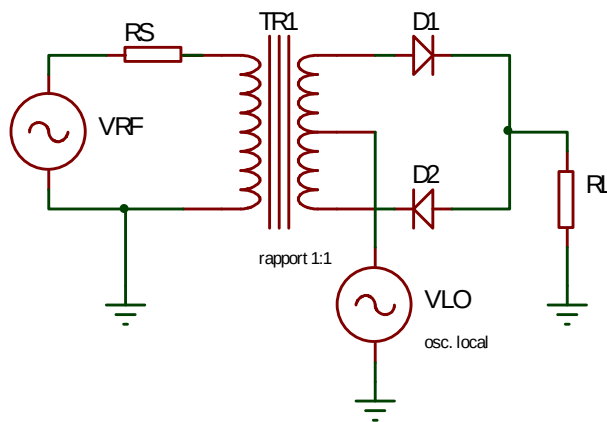
OUTPUT SPECTRUM FOR SWITCHING MIXER FOR $f_{RF} = 11\text{MHz}$ AND $f_{LO} = 10\text{MHz}$



Le circuit mélangeur doit être suivi d'un filtre permettant d'éliminer les raies indésirables.

2.1 MELANGEUR A DIODES

La réalisation d'un mélangeur à commutation peut s'effectuer simplement à l'aide du circuit suivant:

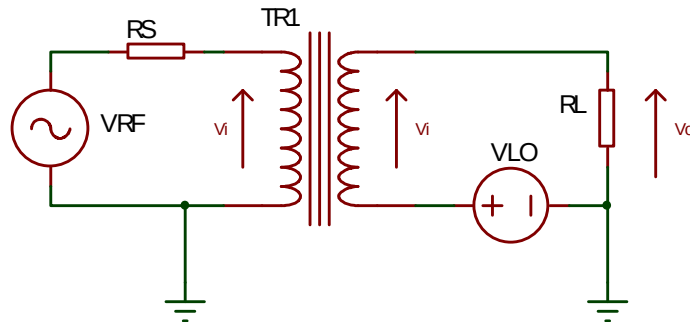


La tension de l'oscillateur local V_{LO} est supposée très grande devant celle du signal d'entrée V_{RF} et d'amplitude suffisante pour provoquer la conduction des diodes.

L'oscillateur local rend alternativement conductrices les diodes D1 et D2. Quand D1 conduit, D2 est bloquée et vice versa.

Considérons une alternance positive de V_{LO} :

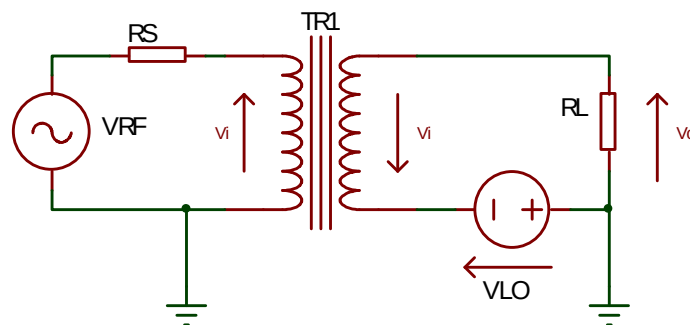
La diode D1 conduit et la diode D2 est bloquée. En supposant que les diodes sont idéales on a le schéma équivalent suivant:



La tension de sortie V_O est égale à la somme de la tension incidente V_i et de la tension V_{LO} .

$$V_O = V_{LO} + V_i$$

Sur une alternance négative de l'oscillateur local on a le schéma suivant (D2 conductrice, D1 bloquée):



d'où l'expression de la tension de sortie: $V_O = V_{LO} - V_i$

On obtient donc bien en sortie le signal d'entrée déphasé alternativement de 0° à 180° au rythme de l'oscillateur local .

En définissant une fonction $P(t) = 1$ si $V_{LO} > 0$ et $P(t) = -1$ si $V_{LO} < 0$

et en posant $V_i^* = V_i P(t)$, on peut obtenir une expression de la tension de sortie:

$$V_O = V_{LO} + V_i^*$$

Si $V_i = V \sin \omega_{RF} t$, V_O peut s'écrire:

$$V_O = V_{LO} + (2V/\pi) [\cos(\omega_{RF} - \omega_{LO})t - \cos(\omega_{RF} + \omega_{LO})t - (1/3)(\cos(\omega_{RF} - 3\omega_{LO})t - \cos(\omega_{RF} + 3\omega_{LO})t) + (1/5)(\cos(\omega_{RF} - 5\omega_{LO})t - \cos(\omega_{RF} + 5\omega_{LO})t) + \dots]$$

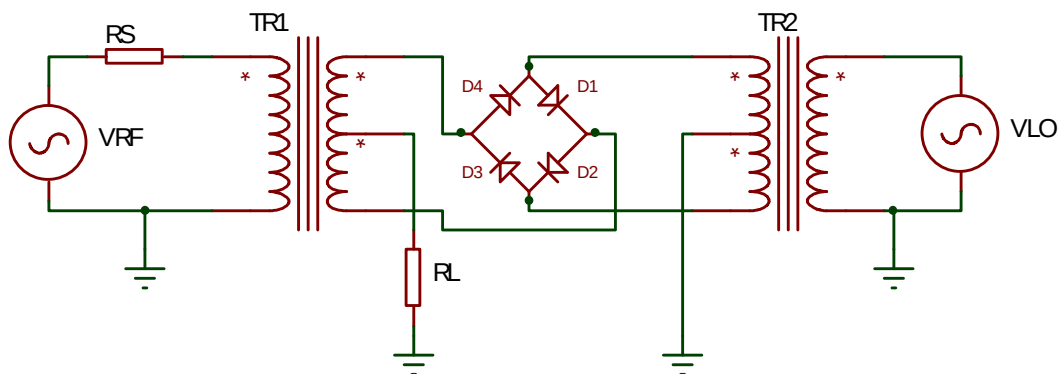
On note la présence dans le signal de sortie des composantes suivantes:

- V_{LO} : tension de l'oscillateur local ($V_{LO} \gg V_i$ par hypothèse, ce qui risque de saturer les étages suivants)
- $(2V/\pi) \cos(\omega_{RF} - \omega_{LO})t$: C'est le signal utile (cas d'un battement infradyne)
- $(2V/\pi) \cos(\omega_{RF} + \omega_{LO})t$: C'est le signal utile dans le cas d'un battement supradynne
- $(2V/\pi)[-(1/3)(\cos(\omega_{RF} - 3\omega_{LO})t - \cos(\omega_{RF} + 3\omega_{LO})t) + (1/5)(\cos(\omega_{RF} - 5\omega_{LO})t - \cos(\omega_{RF} + 5\omega_{LO})t) + \dots]$: Harmoniques indésirables qu'il faut éliminer par filtrage.

La composante V_{LO} peut être éliminée par l'utilisation d'un double circuit équilibré (Double-Balanced Mixer), dont le plus connu est le mélangeur en anneau (Ring Mixer).

2.2 MELANGEUR EN ANNEAU

On utilise un circuit à quatre diodes.

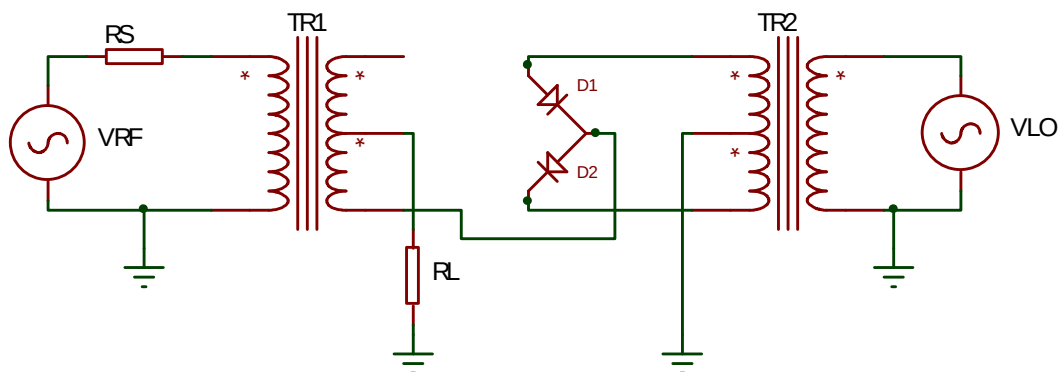


Le fonctionnement est le suivant:

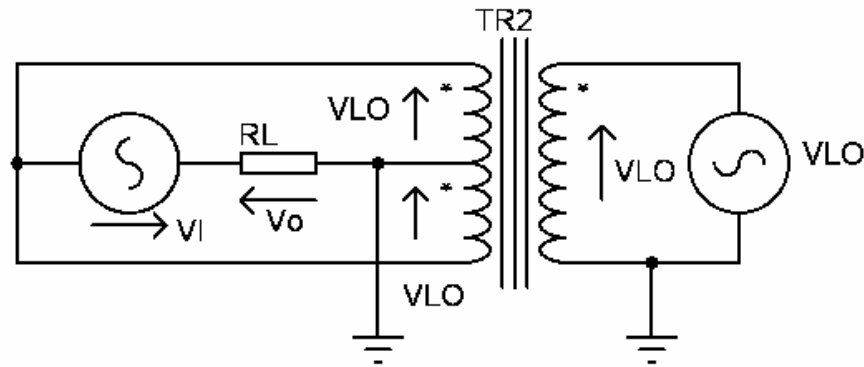
Considérons une alternance positive de l'oscillateur local.

Les diodes D1 et D2 sont rendues conductrices, alors que les diodes D3 et D4 sont bloquées.

Le circuit est alors le suivant:



Ce qui est équivalent à:



Dans la maille du haut, on peut écrire:

$$V_{LO} = -V_i + V_o$$

Et dans celle du bas:

$$-V_{LO} = -V_i + V_o$$

D'où $V_o = V_i$

Pour une alternance négative de V_{LO} , les diodes D3 et D4 sont conductrices et les diodes D1 et D2 sont bloquées. On obtient de la même manière l'expression: $V_o = -V_i$

Le signal de sortie s'écrit donc finalement:

$$V_o(t) = V_i^* = V_i P(t) \quad \text{avec } P(t) = -1 \text{ si } V_{LO} > 0 \text{ et } P(t) = 1 \text{ si } V_{LO} < 0$$

Si $V_i(t) = V \sin \omega_{RF} t$, V_o peut s'écrire:

$$V_o = - (2V/\pi) [\cos(\omega_{RF} - \omega_{LO})t - \cos(\omega_{RF} + \omega_{LO})t - (1/3)(\cos(\omega_{RF} - 3\omega_{LO})t - \cos(\omega_{RF} + 3\omega_{LO})t) + (1/5)(\cos(\omega_{RF} - 5\omega_{LO})t - \cos(\omega_{RF} + 5\omega_{LO})t) + \dots]$$

On constate que par rapport au circuit précédent, le terme V_{LO} a disparu.

Isolation entre ports

Si le circuit est parfaitement symétrique, le mélangeur en anneau présente une très grande réjection de la fréquence de l'oscillateur local.

Dans la pratique l'isolation entre les *ports* (terme générique pour désigner les entrées ou les sorties d'un multipôle) LO et IF est d'environ 40 dB.

On montre de la même manière que le mélangeur en anneau présente également une très bonne isolation entre les ports LO et RF, ce qui évite de perturber les étages précédant le mélangeur.

Perte de conversion

L'impédance d'entrée du mélangeur, vue du port RF est égale à R_L .

Le circuit étant en général adapté et l'on a : $R_S = R_L$ (50Ω le plus souvent).

Ce qui implique que $V_S = 2V_i$

A la fréquence intermédiaire $\omega_{IF} = \omega_{RF} - \omega_{LO}$ l'amplitude maximale du signal de sortie V_o est égale à V_S/π , ce qui correspond à une puissance disponible dans la résistance R_L :

$$P_o = V_S^2/\pi^2 R_L$$

La puissance incidente parvenant au port RF s'écrit:

$$P_i = V_i^2/R_L = V_S^2/4R_L$$

Le gain de conversion est alors $G = P_o/P_i = 4/\pi^2 \approx 0,4$

soit en décibel $G_{dB} = 10 \cdot \log G \approx -4dB$

Ce qui correspond à une *perte de conversion* de 4 dB.

Le signal incident à la fréquence ω_{RF} subit donc une translation en fréquence de ω_{LO} qui s'accompagne d'un affaiblissement de 4dB (valeur théorique ne tenant pas compte de la résistance dynamique des diodes. Dans la pratique la perte de conversion est de l'ordre de 6dB).

2.3 MELANGEUR ACTIF

Les mélangeurs passifs présentent quelques inconvénients: nécessité d'un fort niveau d'oscillateur local, perte de conversion, ... mais surtout, ils se prêtent mal à une intégration monolithique du fait de la présence de transformateurs.

Pour pallier ces difficultés, on peut avoir recours à un mélangeur actif dont le principe est présenté sur la figure suivante.

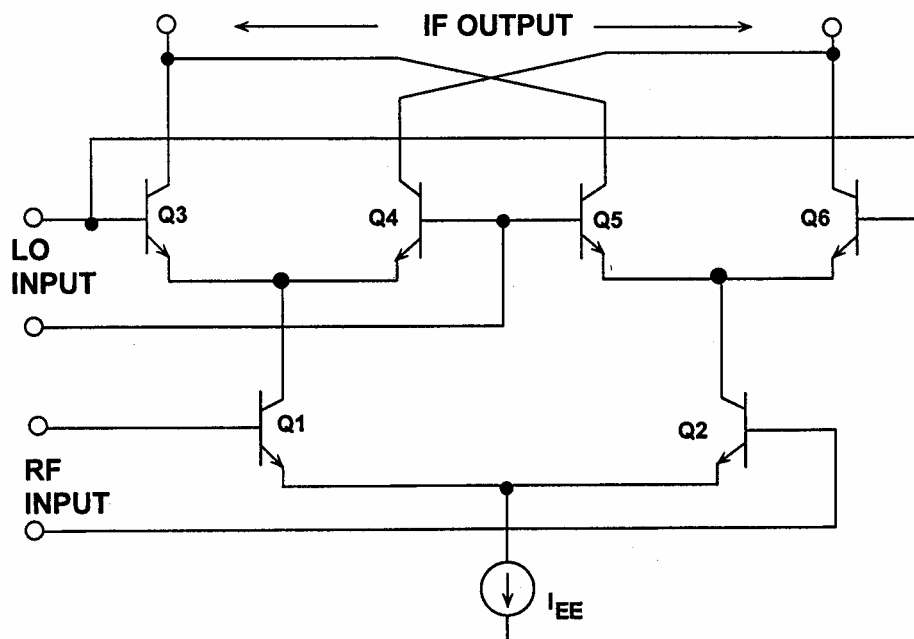


Schéma de principe d'un mélangeur actif classique.

Le circuit se compose d'un amplificateur différentiel (Q1-Q2), recevant le signal d'entrée (RF), dont les collecteurs sont commutés au rythme de l'oscillateur local.

Lorsque les transistors Q3 et Q6 conduisent, les transistors Q4 et Q5 sont bloqués et vice versa.

Les collecteurs de l'amplificateur différentiel sont donc alternativement connectés directement aux sorties, puis croisés et ainsi de suite.

Le signal amplifié, présent à la sortie est donc alternativement en phase ou en opposition de phase selon la polarité de l'oscillateur local.

Ce principe est utilisé depuis de nombreuses années pour réaliser des mélangeurs intégrés dont le plus connu est le MC1496 créé il y a plus de 30 ans par Motorola.

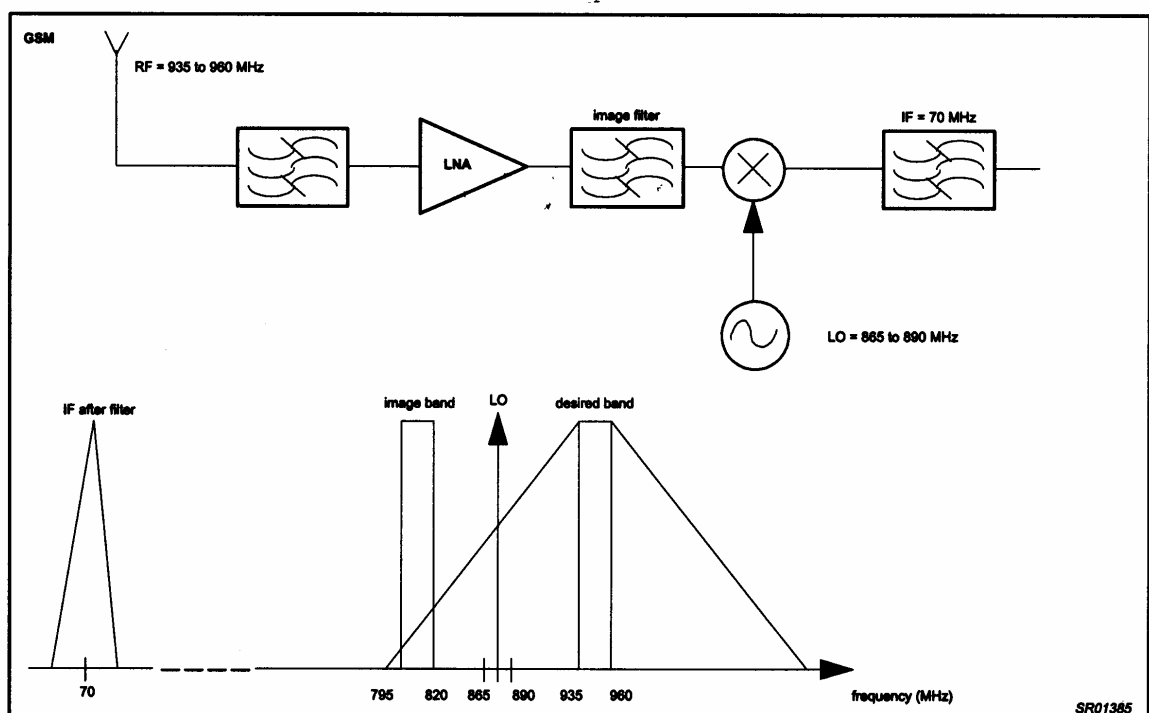
Les progrès de la technologie et la demande croissante de circuits mélangeurs hautes fréquences (notamment pour la téléphonie mobile) ont permis la mise sur le marché de circuits intégrés mélangeurs fonctionnant jusqu'à des fréquences de plusieurs GHz.

2.4 MELANGEURS A REJECTION DE LA FREQUENCE IMAGE

Dans un récepteur superhétérodyne classique, le signal incident est amplifié dans un étage HF généralement accordé, ce qui procure un premier filtrage, puis mélangé avec un oscillateur à fréquence variable afin de délivrer un signal à fréquence intermédiaire fixe.

Si la sélectivité de l'étage HF est faible et/ou si la fréquence intermédiaire est faible, la réjection de la fréquence image sera insuffisante et la réception risque d'être perturbée.

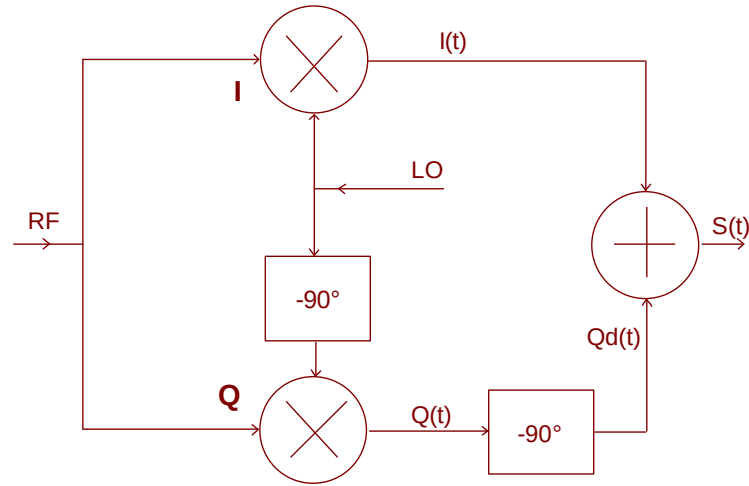
Prenons l'exemple d'un récepteur de téléphone mobile GSM fonctionnant dans la bande 935-960 MHz, dont la fréquence intermédiaire est 70MHz. L'oscillateur local couvrira donc la bande 865-890 MHz. La fréquence image se situera donc 140 MHz au dessous de la bande utile, soit entre 795 et 820 MHz.



La fréquence image doit être éliminée par filtrage, ce qui n'est pas toujours aisé à haute fréquence.

Une autre solution consiste à utiliser un mélangeur à réjection de la fréquence image.

Le principe du circuit consiste à utiliser une paire de mélangeurs attaqués en quadrature de phase par l'oscillateur local. La sortie $I(t)$ est sommée à la sortie $Q(t)$ déphasée de 90° pour donner le signal de sortie $S(t)$.



Soit $V_i = V \sin \omega_i t$ l'expression du signal RF incident et $V_{LO} = \sin(\omega_{LO} t)$ celle de l'oscillateur local. On supposera $\omega_i > \omega_{LO}$, soit $(\omega_i - \omega_{LO}) > 0$.

Le signal incident est appliqué simultanément aux deux mélangeurs I et Q.

Le signal de sortie du mélangeur I s'écrit:

$$I(t) = V \sin \omega_i t \cdot \sin \omega_{LO} t = (V/2) [\cos(\omega_i - \omega_{LO})t - \cos(\omega_i + \omega_{LO})t]$$

Le signal de sortie du mélangeur Q s'écrit:

$$Q(t) = V \sin \omega_i t \cdot \sin(\omega_{LO} t - 90^\circ) = (V/2) [\cos \{(\omega_i - \omega_{LO})t + 90^\circ\} - \cos \{(\omega_i + \omega_{LO})t - 90^\circ\}]$$

$$Q(t) = (V/2) [-\sin(\omega_i - \omega_{LO})t - \cos \{(\omega_i + \omega_{LO})t - 90^\circ\}]$$

Le signal $Q(t)$ subit ensuite un déphasage de -90° pour donner:

$$Qd(t) = (V/2) [-\sin \{(\omega_i - \omega_{LO})t - 90^\circ\} - \cos \{(\omega_i + \omega_{LO})t - 180^\circ\}]$$

$$Qd(t) = (V/2) [\cos(\omega_i - \omega_{LO})t + \cos(\omega_i + \omega_{LO})t]$$

Les signaux $I(t)$ et $Qd(t)$ sont ensuite sommés :

$$S(t) = (V/2) [\cos(\omega_i - \omega_{LO})t - \cos(\omega_i + \omega_{LO})t] + (V/2) [\cos(\omega_i - \omega_{LO})t + \cos(\omega_i + \omega_{LO})t]$$

$$S(t) = V \cos(\omega_i - \omega_{LO})t$$

On constate qu'il y a annulation du battement supradyné.

Supposons maintenant que ω_i soit la fréquence image. On a $(\omega_i - \omega_{LO}) < 0$

Le signal de sortie du mélangeur I s'écrit maintenant:

$$I(t) = V \sin \omega_i t \cdot \sin \omega_{LO} t = (V/2) [\cos(\omega_i - \omega_{LO})t - \cos(\omega_i + \omega_{LO})t]$$

$$I(t) = (V/2) [\cos(\omega_{LO} - \omega_i)t - \cos(\omega_i + \omega_{LO})t] \text{ puisque } \cos(a) = \cos(-a)$$

Et celui de Q:

$$Q(t) = (V/2) [\cos \{(\omega_i - \omega_{LO})t + 90^\circ\} - \cos \{(\omega_i + \omega_{LO})t - 90^\circ\}]$$

$$Q(t) = (V/2)[\sin (\omega_{LO} - \omega_i)t - \cos \{(\omega_i + \omega_{LO})t - 90^\circ\}]$$

Après déphasage de -90° on a:

$$Q_d(t) = (V/2)[\sin \{(\omega_{LO} - \omega_i)t - 90^\circ\} - \cos \{(\omega_i + \omega_{LO})t - 180^\circ\}]$$

$$Q_d(t) = (V/2)[- \cos (\omega_{LO} - \omega_i)t + \cos (\omega_i + \omega_{LO})t]$$

Ce qui donne en sortie:

$$S(t) = (V/2)[\cos (\omega_{LO} - \omega_i)t - \cos (\omega_i + \omega_{LO})t] + (V/2)[- \cos (\omega_{LO} - \omega_i)t + \cos (\omega_i + \omega_{LO})t]$$

$$S(t) = 0$$

Il y a annulation de la fréquence image en sortie.

Ce type de mélangeur, qui peut être totalement intégré, est très largement utilisé, notamment dans les téléphones mobiles.

CHAPITRE VI

LES OSCILLATEURS

GENERALITES

1.1 DEFINITIONS

Un oscillateur sinusoïdal est un système *autonome* qui, à partir d'une source de tension continue, délivre un signal sinusoïdal aussi pur que possible (c'est-à-dire sans harmonique), de fréquence et d'amplitude fixes (ou ajustables par l'utilisateur):

$$v_s(t) = V_s \sin \omega_0 t$$

La fréquence $f_0 = \omega_0 / 2\pi$ est la fréquence des oscillations.

1.2 PERFORMANCES D'UN OSCILLATEUR

1.2.1 Pureté spectrale

Idéalement, le signal de sortie est une sinusoïde pure. En réalité, c'est un signal périodique que l'on peut décomposer en série de Fourier:

$$v_s(t) = V_1 \sin \omega_0 t + V_2 \sin 2\omega_0 t + \dots + V_n \sin n\omega_0 t + \dots$$

La pureté spectrale de $v_s(t)$ est mesurée par le *taux de distorsion harmonique*

$$THD(\%) = \frac{100 \sqrt{V_2^2 + V_3^2 + \dots + V_n^2}}{V_1}$$

Il est parfois exprimé en décibels: $TDH_{dB} = 20 \log THD$.

1.2.2 Stabilité en fréquence

La fréquence f_0 d'un oscillateur doit rester stable. En réalité, elle varie autour de f_0 d'un écart Δf

La stabilité en fréquence d'un oscillateur s'exprime par le rapport $\Delta f / f_0$ souvent exprimé en *parties par million* (ppm).

Les causes de la dérive en fréquence peuvent être

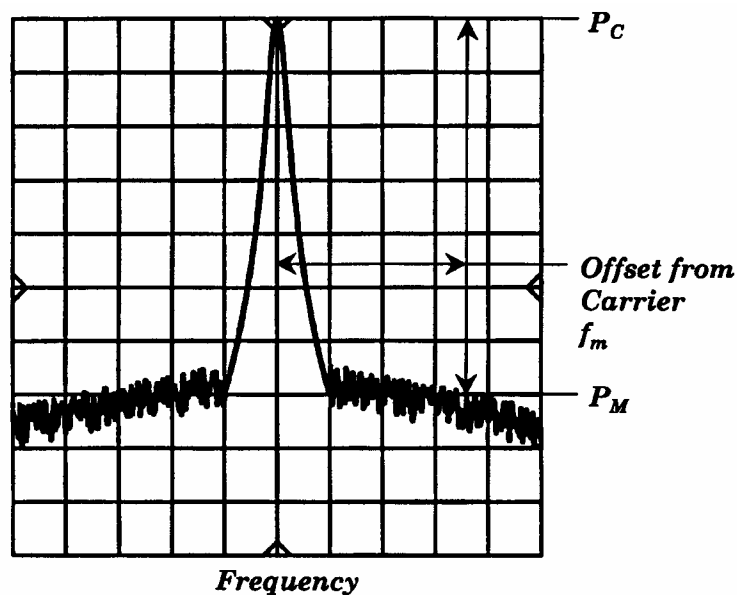
- le vieillissement des composants
- les variations de la température
- les perturbations électromagnétiques

1.2.3 Stabilité en amplitude

L'amplitude du signal de sortie doit rester aussi constante que possible. Mais des variations de la tension d'alimentation, le vieillissement des composants, ou des variations de la température sont susceptibles de modifier l'amplitude des oscillations.

1.2.4 Bruit de phase

Le bruit de phase (phase noise) décrit les fluctuations aléatoires à court terme d'un oscillateur. En effet, le signal théoriquement pur produit par l'oscillateur est modulé par du bruit ce qui donne naissance à un spectre comportant des bandes latérales de bruit (noise sidebands) que l'on caractérise en dBc/Hz (amplitude correspondant à une bande passante de 1Hz, référencée à l'amplitude de l'onde pure, appelée porteuse (carrier)).



Spectre du signal délivré par un oscillateur montrant les bandes latérales de bruit.

PRINCIPE DE FONCTIONNEMENT D'UN OSCILLATEUR

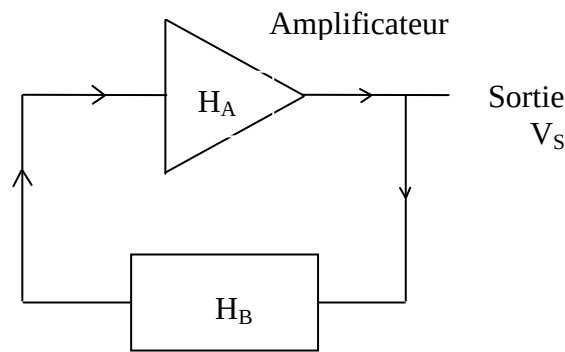
2.1 Structure d'un oscillateur

Un oscillateur comporte toujours un élément actif associé à un circuit passif qui détermine la fréquence des oscillations.

L'élément actif peut être un transistor bipolaire, un TEC ou un amplificateur opérationnel.

La structure d'un oscillateur est celle d'un circuit bouclé l'amplificateur A (élément actif) est bouclé par un circuit de réaction positive B (élément passif), le rebouclage permettant l'auto-entretien des oscillations.

Le schéma de principe d'un oscillateur est le suivant:



2.2 Condition d'oscillation

Soit $H_A(j\omega)$ la fonction de transfert de l'amplificateur A, et $H_B(j\omega)$ celle du circuit de réaction B. Le schéma ci-dessus nous conduit à écrire:

$$V_S = H_A H_B V_S$$

D'où on déduit la condition d'oscillation:

$$H_A(j\omega)H_B(j\omega) = 1$$

Cette condition portant sur des nombres complexes, elle se décompose en une condition sur les phases et une condition sur les modules:

- *condition de phase:* $\arg H_A(j\omega) + \arg H_B(j\omega) = 0$.
- La condition de phase impose généralement la fréquence d'oscillation.
- *condition de module:* $|H_A(j\omega)| \cdot |H_B(j\omega)| = 1$

La condition de module est la condition d'entretien des oscillations.

2.3 Classification des oscillateurs

On peut classer les oscillateurs en fonction de la nature des éléments du circuit de réaction. On obtient alors trois classes d'oscillateurs:

- les oscillateurs à circuit RC, utilisés pour les basses fréquences (quelques centaines de kHz au plus);
- les oscillateurs à circuit LC;
- les oscillateurs à quartz ou à résonateur céramique.

Ces derniers permettent d'atteindre des fréquences très élevées (au delà du GHz).

2.4 Condition de démarrage de l'oscillateur

Les oscillations prennent naissance à la faveur du bruit de fond qui existe dans tout montage électronique (bruit dû à l'agitation thermique).

Supposons qu'après mise sous tension du montage, la tension de bruit (apparue spontanément) soit de $10 \mu\text{V}$ à l'entrée de l'amplificateur. Appelons A le gain de l'amplificateur et B l'atténuation du circuit de réaction: $A = |H_A(j\omega)|$ $B = |H_B(j\omega)|$

Si A et B vérifient les conditions d'oscillations, on n'obtiendra, en sortie de l'oscillateur, que le signal fourni par le bruit, donc un signal d'amplitude trop faible pour être utilisable. Supposons donc que: $A = 9$ et $B = 1/3$, la tension de bruit à la sortie de l'amplificateur sera de $90 \mu\text{V}$. Le circuit de réaction l'atténue d'un facteur 3, et c'est donc une tension de $30 \mu\text{V}$ qui est appliquée à l'entrée de l'amplificateur. A sa sortie, la tension est de $270 \mu\text{V}$, et passe à $90 \mu\text{V}$ à la sortie du circuit de réaction. Et ainsi de suite...

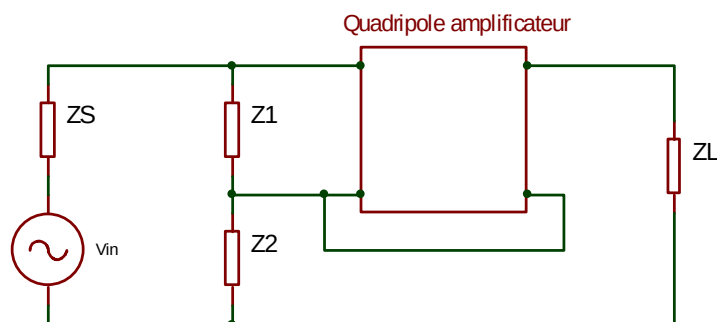
L'amplitude du signal de sortie croît donc rapidement, jusqu'à atteindre les limites de linéarité de l'amplificateur (tensions de saturation). Le signal est alors écrêté.

Si, pour une raison extérieure (variation de l'alimentation ou de la température), le gain A diminue jusqu'à une valeur inférieure à $1/B$, alors l'amplitude des oscillations va décroître jusqu'à la disparition des oscillations.

En résumé, la naissance des oscillations se produit à la faveur du bruit de fond du montage et à condition que le gain A soit supérieur à l'atténuation $1/B$.

3 OSCILLATEURS HAUTES FREQUENCES

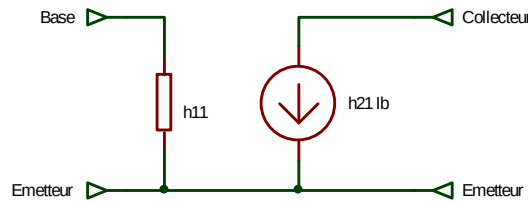
On utilise un quadripôle amplificateur (transistor bipolaire ou FET) dans la configuration suivante.



3.1 OSCILLATEUR COLPITTS A TRANSISTOR BIPOLAIRE

L'élément actif est un transistor bipolaire, Z_1 et Z_2 sont deux condensateurs C_1 et C_2 , l'impédance de charge Z_L est supposée purement résistive.

On peut établir un schéma équivalent du circuit en remplaçant le quadripôle (ici un transistor bipolaire) par son schéma équivalent (simplifié):



Le schéma du circuit est alors le suivant.

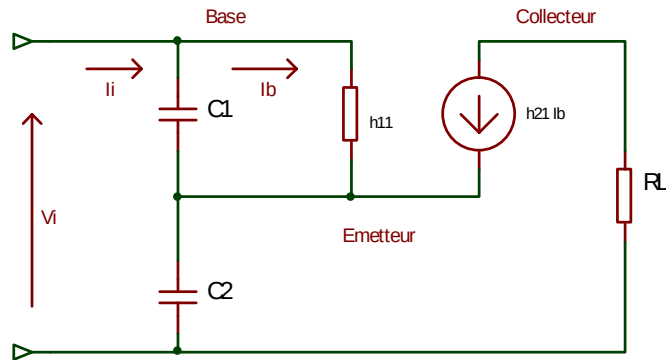


Schéma équivalent de l'oscillateur Colpitts

Nous allons montrer que l'impédance d'entrée Z_i de ce circuit est négative.

On a par définition $Z_i = V_i / I_i$

$$V_i = (Z_{C1} // h_{11}) I_i + (I_i + \beta I_b) Z_{C2}$$

$$\text{Avec : } Z_{C1} = 1/jC_1\omega \quad Z_{C2} = 1/jC_2\omega$$

$$V_i = [Z_{C1} \cdot h_{11} / (Z_{C1} + h_{11}) + Z_{C2}] I_i + \beta I_b Z_{C2} \quad (1)$$

Soit V_{C1} la tension aux bornes de C_1 .

$$V_{C1} = h_{11} I_b = [Z_{C1} \cdot h_{11} / (Z_{C1} + h_{11})] I_i$$

$$\text{D'où } I_b = [Z_{C1} / (Z_{C1} + h_{11})] I_i$$

En portant dans l'équation (1), il vient:

$$V_i = [Z_{C1} \cdot h_{11} / (Z_{C1} + h_{11}) + Z_{C2}] I_i + [\beta Z_{C2} Z_{C1} / (Z_{C1} + h_{11})] I_i$$

On en déduit la valeur de l'impédance d'entrée:

$$Z_i = Z_{C2} + (Z_{C1} \cdot h_{11} + \beta Z_{C2} Z_{C1}) / (Z_{C1} + h_{11})$$

$$Z_i = [h_{11}(Z_{C1} + Z_{C2}) + (\beta + 1)Z_{C1} \cdot Z_{C2}] / (Z_{C1} + h_{11})$$

Si $Z_{C1} \ll h_{11}$ et si $\beta \gg 1$ on a :

$$Z_i = [h_{11}(Z_{C1} + Z_{C2}) + \beta \cdot Z_{C1} \cdot Z_{C2}] / h_{11}$$

$$Z_i = Z_{C1} + Z_{C2} + Z_{C1} \cdot Z_{C2} \cdot \beta / h_{11}$$

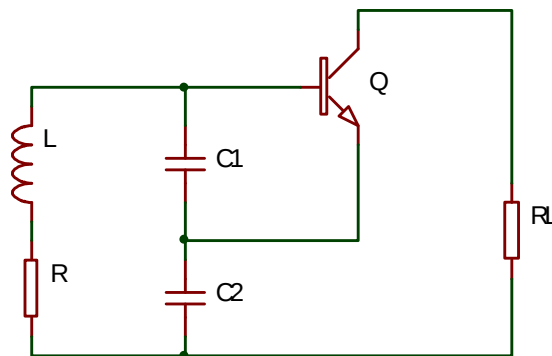
Or $\beta / h_{11} = g_m$ transconductance (ou pente) du transistor, on a donc finalement en remplaçant Z_{C1} et Z_{C2} par leur valeur:

$$Z_i = -g_m / C_1 C_2 \omega^2 - j(1/C_1 \omega + 1/C_2 \omega)$$

On constate que la partie réelle R_i de l'impédance d'entrée Z_i est négative:

$$R_i = -g_m / C_1 C_2 \omega^2 < 0$$

Si l'on connecte à ce circuit une inductance L de facteur de qualité $Q = L\omega/R$, on obtient une oscillation entretenue si la résistance négative compense les pertes du circuit oscillant.



La condition d'oscillation s'écrit donc:

$$R = g_m / C_1 C_2 \omega^2$$

Le circuit est alors sans perte et l'on a:

$$\begin{aligned} jL\omega - j(1/C_1\omega + 1/C_2\omega) &= 0 \\ j(L\omega - (C_1 + C_2) / \omega C_1 C_2) &= 0 \end{aligned}$$

Le circuit oscille à la fréquence :

$$\omega_0 = \frac{1}{\sqrt{L \frac{C_1 C_2}{C_1 + C_2}}}$$

3.2 Oscillateur Hartley

Les impédances Z_1 et Z_2 sont maintenant deux inductances L_1 et L_2 .

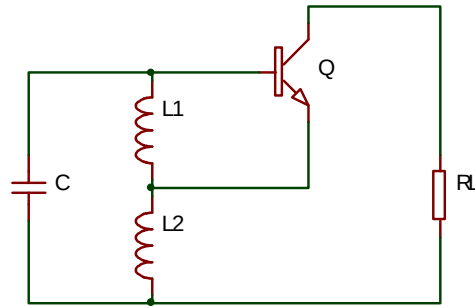


Schéma de principe de l'oscillateur Hartley

On montre de la même manière que le circuit oscille à la fréquence :

$$\omega_0 = \frac{1}{\sqrt{C(L_1 + L_2)}}$$

4 OSCILLATEURS A QUARTZ

Les oscillateurs à quartz sont particulièrement adaptés dans les applications nécessitant des signaux de haute précision, très stables en fréquence : émetteurs radiofréquences, Horloges électroniques, pilote de référence...

4.1 L 'effet piézo-électrique

- *Effet piézoélectrique direct*: lorsqu'une force mécanique est appliquée sur les faces d'un matériau piézoélectrique, il apparaît spontanément sur ces faces des charges électriques convertibles en courant ou en tension électrique.
- *Effet piézoélectrique inverse*: Inversement, lorsqu'une tension électrique est appliquée sur les faces du matériau, les dimensions géométriques du matériau sont modifiées. Le matériau piézoélectrique devient un résonateur mécanique qui vibre à la fréquence de la tension appliquée.

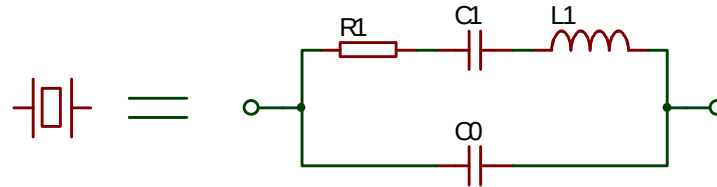
C'est cet effet qui est exploité dans les oscillateurs à quartz. Les fréquences de vibrations d'un quartz sont comprises entre 50 kHz et 150 MHz.

Le cristal de quartz a la forme d'un prisme. C'est en taillant les faces du cristal que l'on définit sa fréquence d'oscillation.

Le cristal est monté entre deux électrodes métalliques dans un boîtier fermé hermétiquement .

4.2 Schéma équivalent d'un quartz

Entre ses deux électrodes, le quartz se comporte, du point de vue électrique, comme un dipôle constitué de deux branches parallèles:



- l'une comporte la capacité statique C_0 dont la valeur est liée aux électrodes et au boîtier
- l'autre comporte le circuit R_1, L_1, C_1 , qui est lié aux caractéristiques mécaniques du quartz.
 L_1 représente la masse vibrante du quartz, C_1 représente l'élasticité du quartz et R_1 représente les pertes.

Dans un quartz, on a: $C_0 \gg C_1$

Si l'on néglige la résistance R_1 , l'impédance du quartz est une réactance

$$X = \frac{1 - L_1 C_1 \omega^2}{\omega(C_1 + C_0 - L_1 C_0 C_1 \omega^2)}$$

Cette expression s'annule pour $L_1 C_1 \omega^2 = 1$ ce qui correspond à la résonance du circuit série L_1 - C_1 .

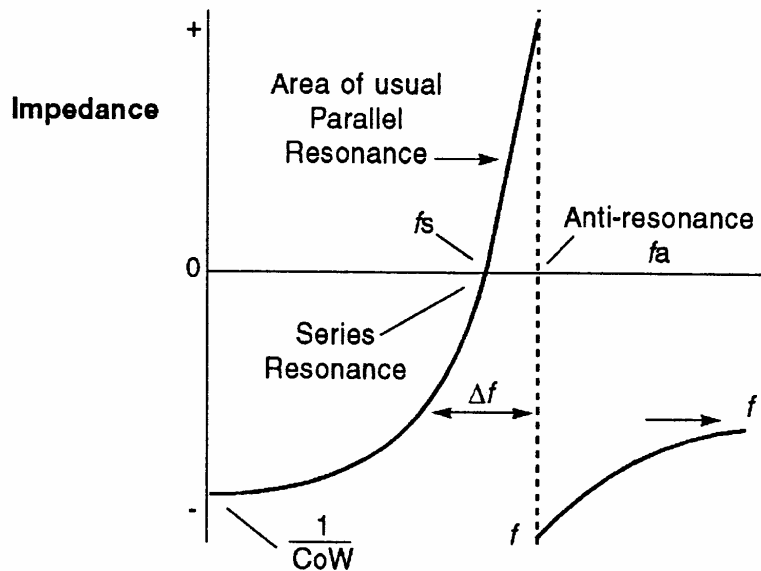
La fréquence $f_s = \frac{1}{2\pi\sqrt{L_1 C_1}}$ est appelée fréquence de résonance série.

A cette fréquence le quartz présente une impédance (théoriquement) nulle.

Le dénominateur s'annule pour une fréquence $f_p = \frac{1}{2\pi\sqrt{\frac{C_0 C_1}{C_0 + C_1} L_1}}$.

Cette fréquence est appelée fréquence de résonance parallèle du quartz.

A cette fréquence, l'impédance du quartz est (théoriquement) infinie.



Dans la bande de fréquence comprise entre f_s et f_p , le quartz est inductif, sa réactance varie très rapidement de zéro à une valeur très élevée : c'est pourquoi le quartz est utilisé pour stabiliser la fréquence des oscillateurs.

La relation liant les deux fréquences de résonance est donnée par:

$$f_p = f_s \sqrt{1 + \frac{C_1}{C_0}}$$

Comme $C_0 \gg C_1$ on peut écrire:

$$f_p \cong f_s \left(1 + \frac{C_1}{2C_0}\right)$$

4.3 Caractéristiques des cristaux de quartz

Coefficient de Qualité:

On définit le coefficient de qualité d'un quartz par le rapport

$$Q = \frac{L_1 \omega}{R_1}$$

où ω est la pulsation de résonance série du quartz.

Le quartz a un coefficient de qualité Q très élevé pouvant atteindre 100 000. Par comparaison, le facteur Q d'un circuit LC dépasse rarement 100.

Stabilité en fréquence

La stabilité en fréquence d'un quartz s'exprime en ppm (10^{-6}) dans la gamme de température spécifiée.

La fréquence nominale est généralement donnée à 25°C avec une certaine tolérance.

Dans le cas où une fréquence très précise et très stable est nécessaire on a souvent recours à un oscillateur à quartz compensé en température (TCXO) ou à un oscillateur placé dans une enceinte thermostatée (OCXO).

Quelques ordres de grandeur pour des quartz coupe AT en boîtier HC25/U:

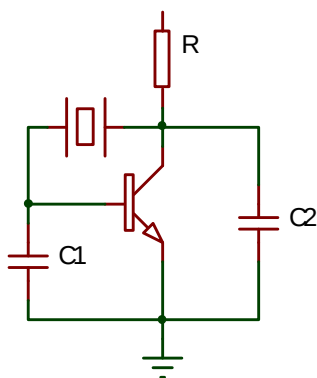
fréquence	C_0	C_1	R_1	Q
2,5 à 5 MHz	4 à 7 pF	10fF	<200 Ω	> 100 000
5 à 10 MHz	5 à 7 pF	20fF	10 à 100 Ω	> 50 000

Capacité de charge

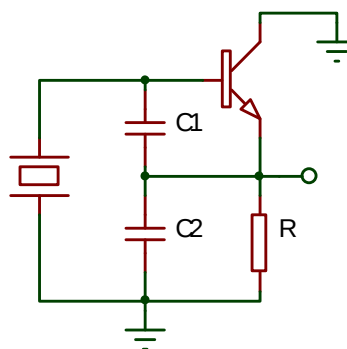
Lorsqu'un quartz est placé dans un circuit oscillateur, celui ci ramène une capacité parasite qui, dans le cas où l'on utilise la résonance parallèle, s'ajoute à la capacité C_0 et tend à diminuer la fréquence des oscillations. La fréquence de résonance parallèle est donc toujours donnée pour une capacité de charge spécifiée (de l'ordre de 30 pF).

On peut mettre à profit cette propriété pour ajuster précisément la fréquence de l'oscillateur en plaçant un condensateur ajustable en parallèle (ou en série) avec le quartz.

Quelques oscillateurs à quartz classiques



Oscillateur PIERCE



Oscillateur CLAPP

3 OSCILLATEURS COMMANDES EN TENSION

Dans de nombreuses applications, on a besoin d'un oscillateur à fréquence variable (VFO).

Le principe d'un tel oscillateur consiste à rendre variable un des paramètres du circuit oscillant qui détermine la fréquence: soit le condensateur, soit l'inductance.

On peut utiliser des condensateurs variables à commande mécanique ou des inductances à noyau plongeur, mais ces composants ne sont pas *commandables* électriquement ce qui limite leur utilisation.

5.1 Diode varicap

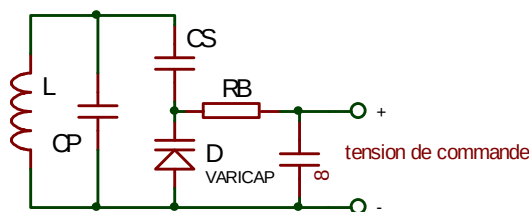
Dans les circuits modernes, on utilise la capacité d'une diode polarisée en inverse dont l'expression est de la forme:

$$C = \frac{K}{(V_d + V)^n}$$

Où V_d est lié à la barrière de potentiel, V est la tension inverse appliquée à la diode et n un coefficient dépendant de la technologie de la diode ($n = 0,33$ pour une jonction graduelle, $n=0,5$ pour une jonction abrupte, $n=0,75$ pour une jonction hyper abrupte)

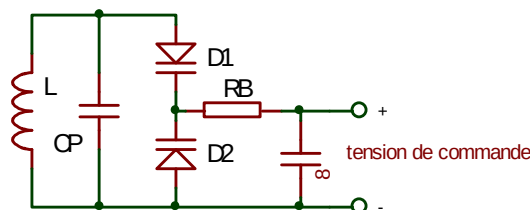
La capacité maximale que l'on peut obtenir s'étend de quelques pF à quelques centaines de pF, avec des rapports C_{max}/C_{min} allant de 5 à 15.

Le schéma suivant montre l'utilisation classique d'une diode varicap. La diode Varicap est connectée en parallèle avec le circuit oscillant principal L- C_p par l'intermédiaire d'une capacité série C_s qui permet d'ajuster la plage de variation. La tension de commande est appliquée à la diode à travers une résistance de forte valeur (pour ne pas amortir le circuit oscillant).



Lorsque le niveau d'oscillation est élevé, la diode varicap peut être mise en conduction sur les alternances positives. Il s'établit alors un courant impulsionnel générateur d'harmoniques.

Pour pallier ce défaut, on utilise souvent le circuit suivant:



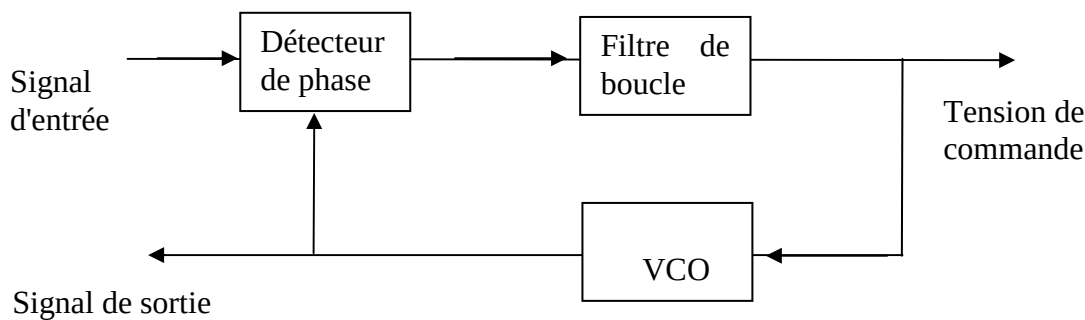
Deux diodes sont montées en série et en sens inverse, ce qui permet un niveau plus élevé de tension HF aux bornes du circuit oscillant.

4 BOUCLE A VERROUILLAGE DE PHASE (PLL)

Une boucle à verrouillage de phase (Phase Locked Loop) est un système bouclé capable de délivrer un signal dont la fréquence est asservie à un signal de référence.

Elle se compose des éléments suivants:

- Un détecteur de phase
- Un filtre passe bas
- Un oscillateur commandé en tension (VCO)



Soit $V_e = E \sin(\omega_e t + \varphi_e)$ l'expression du signal d'entrée et $V_s = E_s \sin(\omega_s t + \varphi_s)$ celle du signal de sortie.

Le comparateur de phase compare la phase du signal d'entrée avec celle du signal délivré par le VCO, et génère un signal d'erreur qui est filtré et appliqué à l'entrée de commande de l'oscillateur (VCO).

Lorsque la boucle est verrouillée (locked), la différence de phase instantanée est constante:

$$V_e = E \sin(\omega_e t + \varphi_e) = E \sin \Phi_e(t)$$

$$V_s = E_s \sin(\omega_s t + \varphi_s) = E_s \sin \Phi_s(t)$$

$$\Phi_s(t) - \Phi_e(t) = \text{constante}$$

$$d\Phi_s(t)/dt - d\Phi_e(t)/dt = 0 \quad \text{d'où} \quad \omega_s = \omega_e$$

La fréquence du VCO est égale à celle du signal d'entrée.

6.1 Principe de fonctionnement

En l'absence de signal d'entrée, le VCO délivre une tension sinusoïdale de fréquence libre f_0 (free running frequency).

Supposons que l'on utilise un multiplieur analogique comme détecteur de phase.

Si l'on injecte à l'entrée un signal sinusoïdal de fréquence f_e voisine de f_0 , le comparateur délivre un signal $e(t)$:

$$V_e(t) = A_e \sin \Phi_e(t)$$

$$V_0(t) = A_0 \sin \Phi_s(t)$$

$$e(t) = V_e(t)V_0(t) = A_e A_0 \sin \Phi_e(t) \sin \Phi_s(t)$$

d'après la relation trigonométrique $\sin a \cdot \sin b = (\frac{1}{2})[\cos(a-b) - \cos(a+b)]$ on a :

$$e(t) = (A_e A_0 / 2) \cos[\Phi_e(t) - \Phi_s(t)] - (A_e A_0 / 2) \cos[\Phi_e(t) + \Phi_s(t)]$$

Le signal d'erreur comporte des composantes à $|f_e - f_0|$ et $f_e + f_0$.

Le filtre passe bas ne laissera passer que la composante à basse fréquence $\Delta f = |f_e - f_0|$ qui est appliquée à l'entrée de commande du VCO.

Il en résulte une modulation de la fréquence du VCO par Δf qui devient ainsi une fonction du temps.

On montre que le système converge vers $\Delta f = 0$ si l'écart $f_e - f_0$ se situe à l'intérieur d'un certain domaine appelée *plage de capture* (ou d'acquisition).

Lorsque la boucle est verrouillée, $f_s = f_e$ et le comparateur de phase ne délivre qu'une tension continue V_c proportionnelle à l'écart de phase entre les signaux d'entrée et de sortie.

$$V_c = K\Delta\Phi$$

Le temps nécessaire pour obtenir le verrouillage est un paramètre important qui dépend des caractéristiques du filtre de boucle.

6.2 Etude de la boucle verrouillée

Lorsque la boucle à verrouillage de phase (PLL) est verrouillée, elle peut s'étudier comme un système asservi linéaire.

Le schéma block du système peut être représenté sous forme de fonctions de transfert comme indiqué ci après.

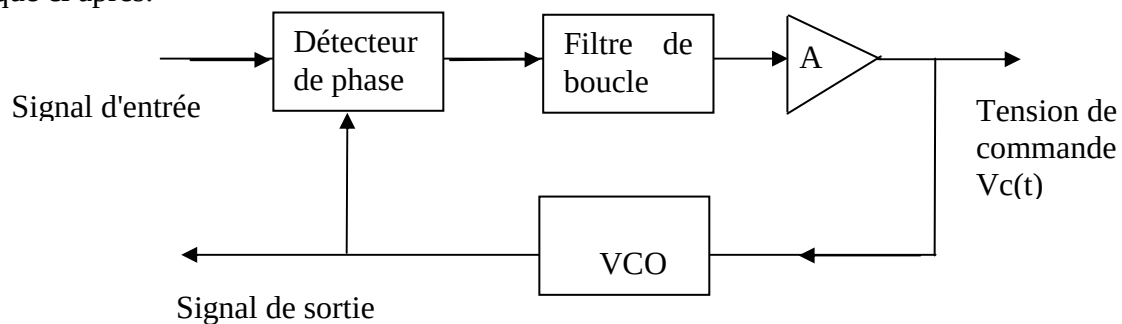
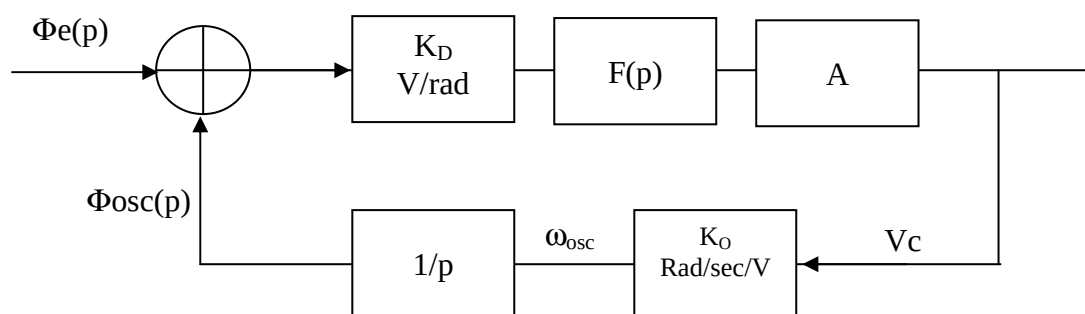


Schéma bloc de la boucle à verrouillage de phase (PLL)



Représentation en fonction de transfert

On notera que la variable d'entrée n'est pas la fréquence f_e , mais la phase $\Phi_e(t)$.

$$\omega_e = d\Phi_e/dt \quad \text{d'où:}$$

$$\Phi_e = \Phi_e(t=0) + \int_0^t \omega_e(t) dt$$

Soit en transformée de Laplace, en considérant que la phase à l'origine est nulle:

$$\Phi_e = (1/p) \omega_e(p)$$

On suppose que la fréquence du VCO est proportionnelle à la tension d'erreur V_c .

$$\omega_{osc} = \omega_0 + K_0 V_c$$

où ω_0 est la fréquence d'oscillation libre du VCO.

La sortie du VCO doit être exprimée en phase puisque le détecteur est sensible à la différence de phase entre le signal d'entrée et la sortie du VCO:

$$\Phi_{osc} = \Phi_{osc}(t=0) + \int_0^t \omega_{osc}(t) dt$$

On constate que le VCO fonctionne comme un intégrateur. Un échelon de tension d'erreur V_c correspond à un échelon de pulsation ω_{osc} et à une rampe de phase.

Ceci se traduit sur la fonction de transfert par le bloc $1/p$.

Nous avons:

$$\Phi_{osc}(p) = \omega_{osc} / p = K_0 V_c(p) / p$$

$$V_c(p) = [\Phi_e(p) - \Phi_{osc}(p)].K_D.F(p).A$$

On en déduit l'expression:

$$\frac{V_c}{\Phi_e} = \frac{pK_DFA}{p + K_0K_DFA}$$

Qui traduit la caractéristique *phase* $\rightarrow$ *tension* du système.

En utilisant la relation $\Phi_e = (1/p) \omega_e(p)$, on obtient la fonction de transfert *fréquence* $\rightarrow$ *tension* du système:

$$\frac{V_c}{\omega_e} = \frac{K_DFA}{p + K_0K_DFA}$$

En posant $K_L = K_0K_D A$, il vient:

$$\frac{V_c}{\omega_e} = \frac{K_L}{K_0} \frac{F(p)}{[p + K_L F(p)]}$$

Ou, pour la fonction de transfert *fréquence d'entrée* → *fréquence de sortie*:

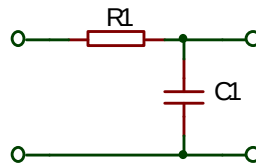
$$\frac{\omega_s}{\omega_e} = K_L \frac{F(p)}{[p + K_L F(p)]}$$

Cette fonction de transfert dépend de la fonction de transfert F(p) du filtre passe bas.

6.3 Filtre de boucle

Le filtre de boucle est de type passe bas.

Examinons le cas où l'on utilise un simple circuit R₁C₁.



La fonction de transfert de ce circuit s'écrit:

$$F(p) = 1 / (1 + R_1 C_1 p) = 1 / (1 + p/\omega_1)$$

Avec $\omega_1 = 1/R_1 C_1$ pulsation de coupure du filtre.

En remplaçant F(p) par son expression dans la fonction de transfert, on obtient:

$$\frac{\omega_s}{\omega_e} = K_L \frac{1}{p + \frac{K_L}{1 + \frac{p}{\omega_1}}} = K_L \frac{1}{\frac{1}{\omega_1} p^2 + p + K_L}$$

Nous sommes en présence d'une fonction de transfert du second ordre.

$$\frac{\omega_s}{\omega_e} = \frac{\omega_1 K_L}{p^2 + \omega_1 p + \omega_1 K_L}$$

Qui admet deux pôles:

$$p_1, p_2 = -\frac{\omega_1}{2} \left(1 \pm \sqrt{1 - \frac{4K_L}{\omega_1}} \right)$$

En identifiant à la forme canonique:

$$H(p) = \frac{K}{\frac{p^2}{\omega_n^2} + \frac{2\xi}{\omega_n} p + 1}$$

on obtient :

$$\omega_n = \sqrt{\omega_1 K_L}$$

$$\xi = \frac{1}{2} \sqrt{\frac{\omega_1}{K_L}}$$

ω_n est la *pulsation naturelle* et ξ est *l'amortissement*.

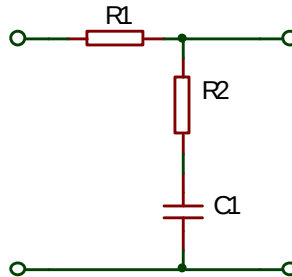
Si $\xi > 1$ la réponse est amortie, si $\xi < 1$ la réponse est oscillatoire.

Dans la pratique, on désire souvent une large plage de verrouillage de manière à pouvoir suivre de grandes variations de la fréquence d'entrée, et une bande passante faible afin de rejeter les signaux hors bande.

Ces contraintes sont relativement bien satisfaites en utilisant un filtre à avance de phase dont la fonction de transfert est:

$$F(p) = \frac{1 + \frac{p}{\omega_2}}{1 + \frac{p}{\omega_1}}$$

Avec $\omega_1 = 1/(R_1 + R_2)C_1$ et $\omega_2 = 1/R_2C_1$ $\omega_2 \gg \omega_1$



On montre alors que :

$$\omega_n = \sqrt{\omega_1 K_L}$$

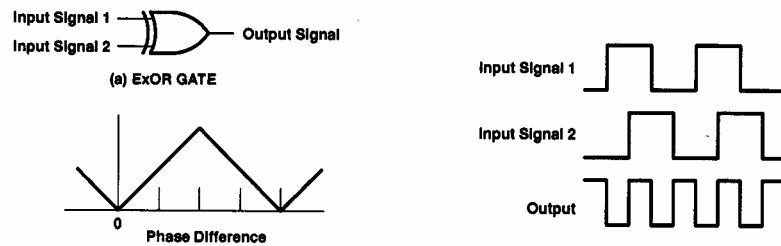
$$\xi = \frac{1}{2} \sqrt{\frac{\omega_1}{K_L}} + \frac{1}{2} \frac{\sqrt{K_L \omega_1}}{\omega_2}$$

Ce filtre permet de conserver un gain de boucle élevé et une bande passante étroite sans que le système soit sous amorti.

6.4 Comparateur de phase

Le comparateur de phase peut être analogique (multiplieur ou mélangeur à commutation), mais on utilise plus fréquemment des comparateurs numériques plus faciles à implémenter:

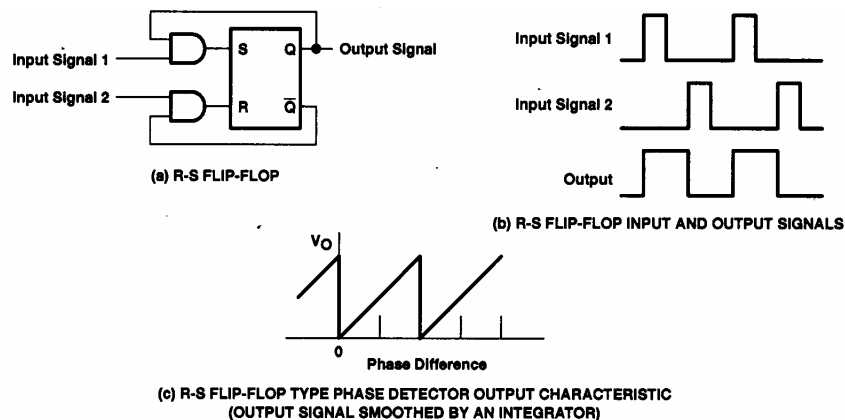
- OU exclusif



Dans le cas d'un OU exclusif, le rapport cyclique du signal de sortie dépend de la différence de phase des signaux d'entrée. Lorsque ce signal est intégré, on obtient une tension continue en relation avec la différence de phase.

On notera que ce détecteur ne fonctionne correctement que si les deux signaux d'entrée ont un rapport cyclique de 50%

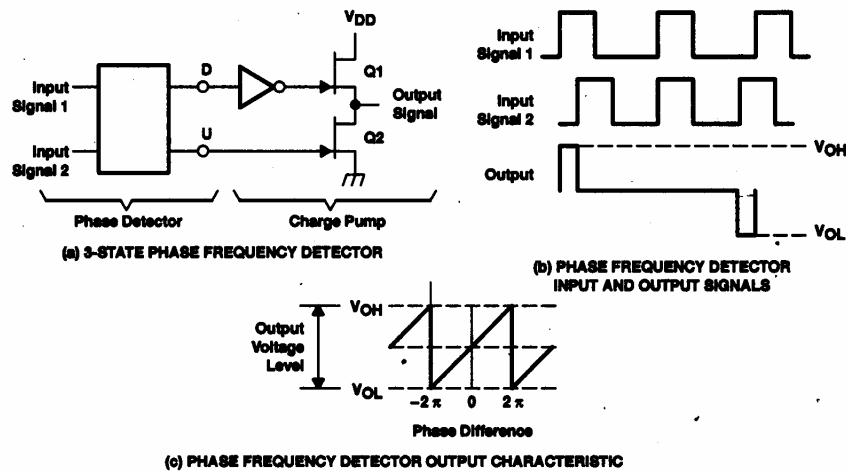
- bascule RS.



Ce circuit ne prend en compte que les fronts montants des signaux d'entrée, qui peuvent donc avoir un rapport cyclique quelconque.

- **Comparateur phase-fréquence**

Ce circuit est l'un des plus utilisés.



Lorsque le signal 2 est en retard sur le signal 1, la sortie D passe à 1 depuis le front montant du premier signal jusqu'au front montant du second signal, c'est à dire pendant la durée correspondant à la différence de phase des deux signaux.

Durant cette période, la sortie U reste à 0.

Si le signal 2 est en avance de phase par rapport au signal 1, les rôles de D et U sont inversés.

Si les deux signaux sont en phase, D et U restent à 0.

Les sorties D et U commandent une pompe de charge constituée des transistors MOS Q₁ et Q₂. Le signal de sortie peut donc prendre trois états V_{OH}, V_{OL} ou haute impédance.

Si D=1 et U=0 : Q₁ conduit, Q₂ est ouvert et la sortie vaut V_{OH}.

Si D=0 et U=1 : Q₁ est ouvert et Q₂ Conduit, la sortie vaut V_{OL}

Si D=0 et U=0 : Q₁ et Q₂ sont ouverts et la sortie devient haute impédance.

On constate que la valeur moyenne du signal de sortie est proportionnelle à la différence de phase des deux signaux d'entrée.

5 SYNTHETISEUR DE FREQUENCE

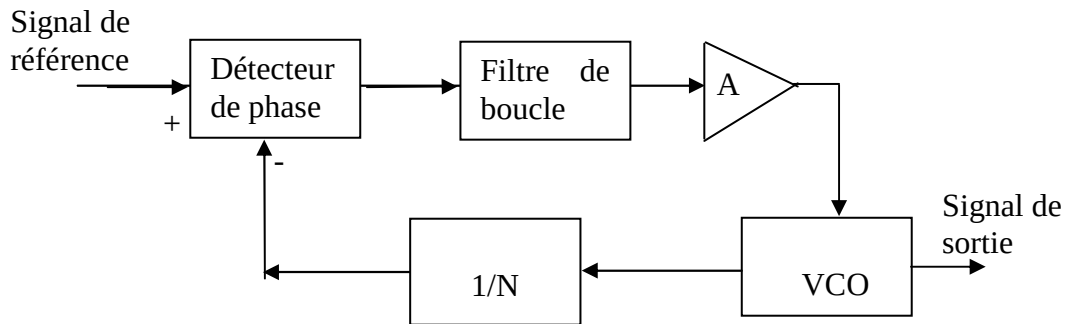


Schéma de principe d'un synthétiseur de fréquence à PLL

Un synthétiseur de fréquence se compose d'une boucle à verrouillage de phase dans laquelle:

- le signal d'entrée est une fréquence de référence stable et précise
- On a introduit un diviseur de fréquence, généralement programmable, entre le VCO et le comparateur de phase.

Le détecteur de phase compare le signal du VCO dont la fréquence est divisée par N au signal de référence. Lorsque la boucle est verrouillée on a:

$$F_{VCO} = N F_{REF}$$

Si N est programmable, la fréquence de sortie peut être rendue variable par bond de F_{REF} .

Exemple:

On désire réaliser un oscillateur pilote pour un émetteur de radiodiffusion dans la bande FM (88 - 108 MHz), sachant que les canaux doivent être espacés de 100 kHz .

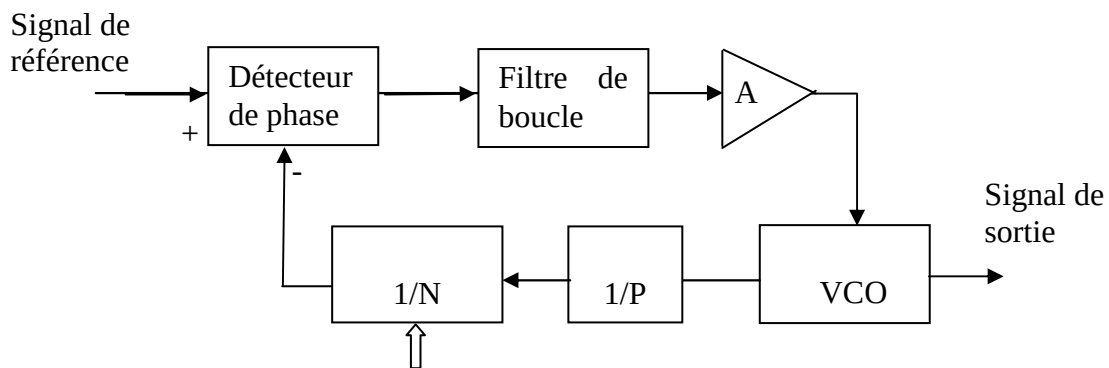
On utilise un synthétiseur de fréquence avec une référence de 100 kHz (obtenue par exemple à partir d'un oscillateur à quartz de 10 MHz suivi d'un diviseur de fréquence par 100).

Pour obtenir 88 MHz il faut diviser par $N = 880$, et pour 108 MHz par $N = 1080$.

Il suffira donc d'implanter un diviseur programmable capable de diviser par $880 < N < 1080$.

5.1 Cas des fréquences élevées.

Dans l'état actuel de la technologie, il est difficile de réaliser des diviseurs programmables fonctionnant au delà d'une centaine de MHz. On a alors recours à un pré-diviseur (prescaler) fixe par P.



Programmation

Schéma de principe d'un synthétiseur de fréquence à PLL avec prédiviseur.

La fréquence de sortie est donnée par:

$$F_{VCO} = N.P.F_{REF}$$

Si l'on incrémente N d'une unité, la fréquence de sortie devient:

$$F_{VCO} = (N+1)P.F_{REF} = N.P.F_{REF} + P.F_{REF}$$

$$F_{VCO}(N+1) = F_{VCO}(N) + P.F_{REF}$$

On constate que le pas du synthétiseur est multiplié par P, ce qui est gênant pour l'obtention de pas fins car la fréquence de comparaison peut devenir très faible (temps de réaction plus long, bruit de phase plus élevé, ...).

Exemple

Reprenons l'exemple précédent, avec un prédiviseur par P=10.

Pour obtenir un pas de 100 kHz en sortie, il faut maintenant une fréquence de référence de 10kHz. Le diviseur programmable devra toujours couvrir le même intervalle 880-1080, mais il fonctionnera à une fréquence 10 fois plus faible (8,8 - 10,8 MHz).

8.2 Synthétiseur à division fractionnaire (fractional N).

Le schéma synoptique de ce type de synthétiseur est donné ci après.

La sortie du VCO est tout d'abord envoyée dans un diviseur de fréquence capable de diviser par P ou par P+1 selon la commande qu'il reçoit.

La sortie de ce diviseur à deux modules (dual modulus) est envoyée à deux compteurs programmables M et A.

Le débordement du compteur A commande la division par P ou P+1.

La sortie du compteur M attaque le comparateur de phase.

Le débordement du compteur M commande la remise à zéro de l'ensemble.

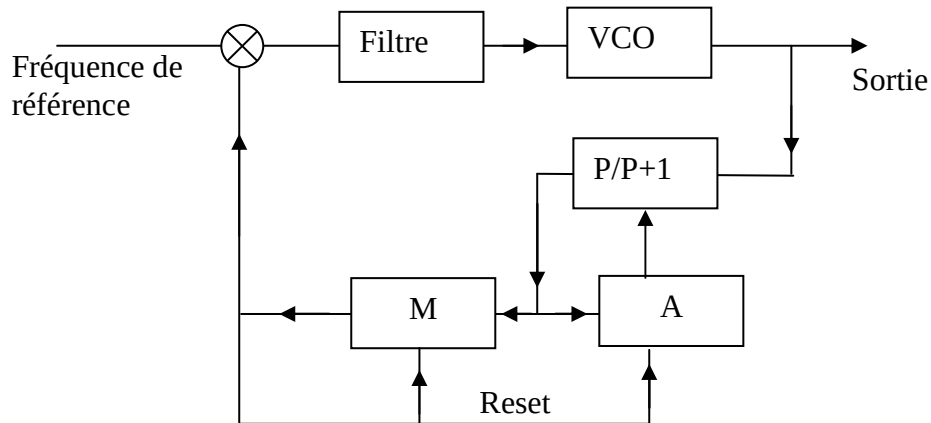


Schéma de principe d'un synthétiseur à division fractionnaire

Soit $T_0 = 1/f_{VCO}$ la période du signal de sortie.

Au début du cycle, le prédiviseur divise par P+1.

Les compteurs A et M reçoivent une horloge à la fréquence $f_{VCO} / (P+1)$

Au bout d'un temps $A(P+1)T_0$ le compteur A déborde et force le prédiviseur à diviser par P pour le reste du cycle.

Le compteur M déborde après un temps $[A(P+1) + (M-A)P]T_0 = (MP+A)T_0$ et remet le système à zéro.

Le comparateur de phase reçoit donc une fréquence $f_{comp} = 1/(MP+A)T_0$.

Lorsque la boucle est verrouillée $f_{comp} = f_{REF}$

$$f_{REF} = 1/(MP+A)T_0 = f_{VCO} / (MP+A)$$

D'où:

$$f_{VCO} = f_{REF} (MP+A)$$

Toute incrémentation d'une unité du compteur A provoque un saut de fréquence égal à la fréquence de référence.

Toute incrémentation d'une unité du compteur M provoque un saut de fréquence égal à MP fois la fréquence de référence.

On notera que A doit toujours être inférieur à M.

Le compteur A doit pouvoir compter au minimum de 0 à P-1.

$$A_{min} = 0 \quad A_{max} = P-1$$

$$M > A \rightarrow M_{min} = P$$

Il en résulte que le facteur de division minimal est:

$$D_{min} = M_{min} (P + A_{min}) = P(P+0) = P^2$$

Ce système permet l'utilisation d'un prédiviseur sans modifier le pas du synthétiseur.

Exemple:

Si l'on utilise un synthétiseur à division fractionnaire avec un prédiviseur par 10/11 dans l'exemple

précédent, on utilisera une fréquence de référence de 100kHz.

Pour 88 MHz on aura :

$$88\,000 = 100(10M + A) \Rightarrow M=88 \text{ et } A=0$$

Pour 88,1 MHz on aura: $M=88$ et $A=1$

On pourra utiliser pour A un compteur 4 bits capable de compter de 0 à 15.

Pour 89,3 MHz par exemple, on voit qu'il y aura deux solutions:

$M = 88$, $A = 13$ et $M = 89$, $A = 3$

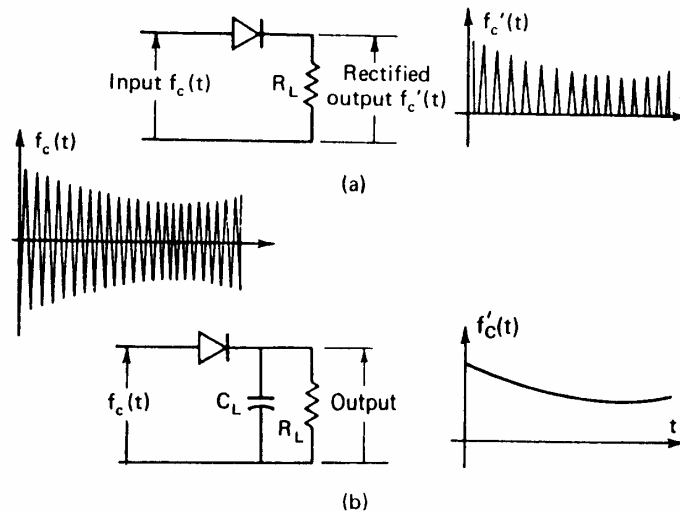
CHAPITRE VII

LA DEMODULATION

La démodulation est l'opération qui consiste à récupérer l'information utile contenue dans un signal modulé. C'est l'opération inverse de la modulation.

1 DEMODULATION D'AMPLITUDE

Un signal modulé en amplitude comporte une sinusoïde haute fréquence dont l'enveloppe varie au rythme de la modulation (dont la fréquence est toujours très faible devant celle de la porteuse). L'utilisation d'un circuit redresseur associé à un filtre passe bas permet d'extraire la modulation.



Dans la pratique, la réponse de la diode n'est pas linéaire, ce qui va introduire une certaine distorsion.

A faible niveau, la réponse de la diode peut être modélisé selon une loi quadratique:

$$i_d = k_0 + k_1 v_d + k_2 v_d^2$$

Si l'on considère que la résistance de charge est faible, la majorité de la tension d'entrée se retrouve aux bornes de la diode.

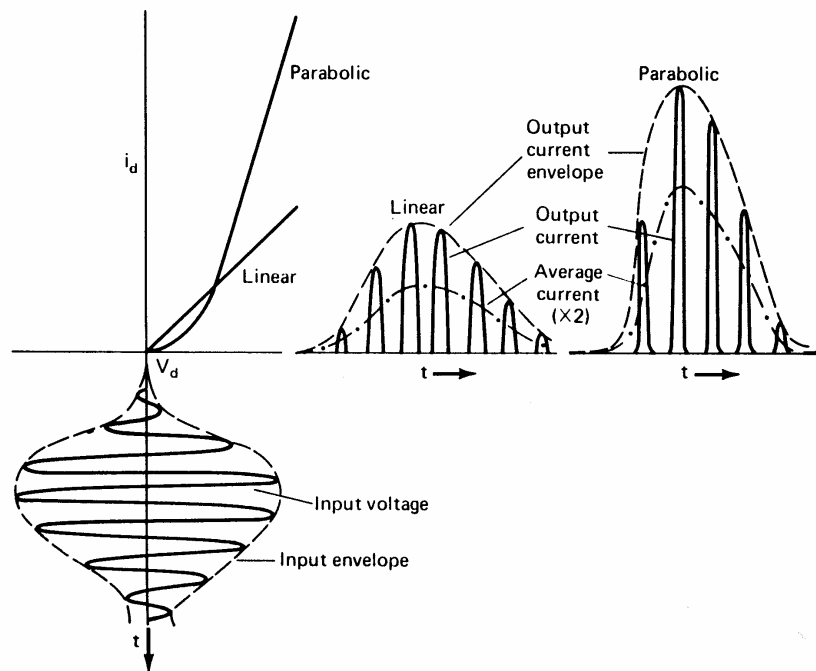
On a donc une tension v_d de la forme:

$$v_d = A[1 + m.s(t)]\cos f_c(t)$$

L'équation du courant dans la diode est alors:

$$i_d = k_0 + k_1 A[1 + m.s(t)]\cos f_c(t) + k_2 A^2[1 + m.s(t)]^2 \cos^2 f_c(t) v_d^2$$

qui fait apparaître la distorsion.



2 DEMODULATION DE FREQUENCE

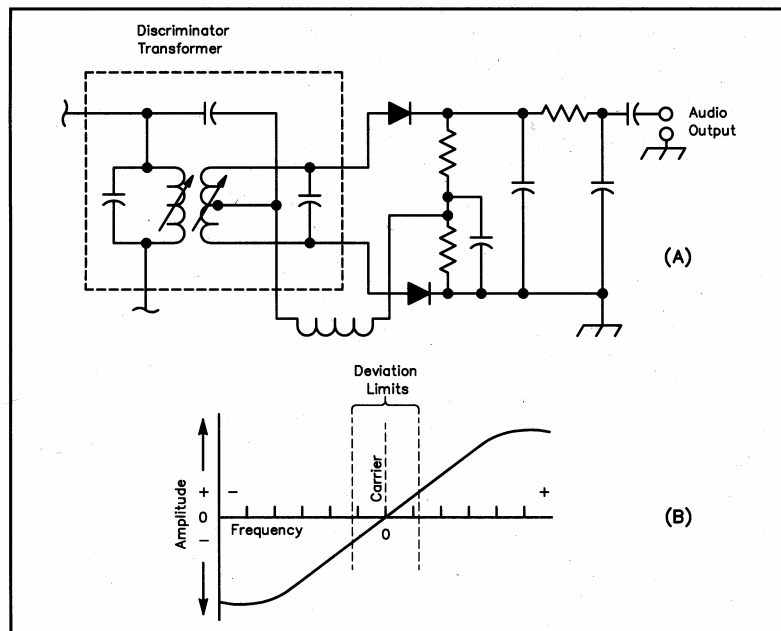
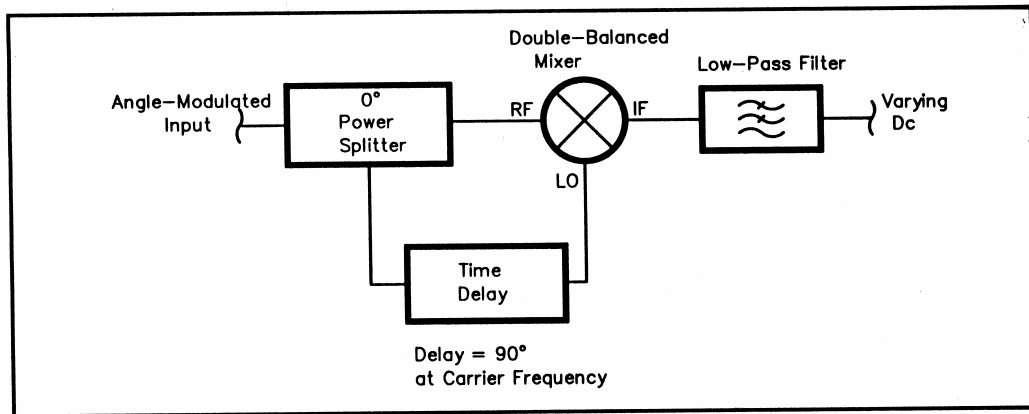


Fig 15.15—A discriminator (A) converts an angle-modulated signal's deviation into an amplitude variation (B) and envelope-detects the resulting AM signal. For undistorted demodulation, the discriminator's amplitude-versus-frequency characteristic must be linear across the input signal's deviation. A *crystal discriminator* uses a crystal as part of its frequency-selective circuitry.



Démodulateur FM en quadrature

$$A \sin(2\pi f t) \cdot A \sin 2\pi f (t + d) = \frac{A^2}{2} \cos(2\pi f d) - \frac{A^2}{2} \cos(2 \cdot 2\pi f t + 2\pi f d)$$

Le terme à deux fois la fréquence peut être facilement éliminé par un filtre passe bas. On obtient donc:

$$A \sin(2\pi f t) \cdot A \sin 2\pi f (t + d) = \frac{A^2}{2} \cos(2\pi f d)$$

Supposons que le retard d soit de 90° à la fréquence de la porteuse. Les deux signaux sont alors en quadrature. En remarquant qu'un retard de $\pi/2$ est égal à $1/4$ de la période, on peut écrire: $d = 1/4f$

Si l'on suppose que le signal incident modulé en fréquence est $f + \Delta f$ où Δf est la valeur du shift, on obtient:

$$A \sin 2\pi (f + \Delta f) t \cdot A \sin 2\pi \left((f + \Delta f) \left(t + \frac{\pi}{2} \right) \right) = \frac{A^2}{2} \cos 2\pi (f + \Delta f) \frac{1}{4f}$$

$$A \sin 2\pi (f + \Delta f) t \cdot A \sin 2\pi \left((f + \Delta f) \left(t + \frac{\pi}{2} \right) \right) = \frac{A^2}{2} \cos \left(\frac{\pi}{2} + \frac{\pi \Delta f}{2f} \right)$$

$$A \sin 2\pi (f + \Delta f) t \cdot A \sin 2\pi \left((f + \Delta f) \left(t + \frac{\pi}{2} \right) \right) = -\frac{A^2}{2} \sin \left(\frac{\pi \Delta f}{2f} \right)$$

Si $\Delta f = 0$ (porteuse non modulée), la sortie est nulle.

Si Δf est faible devant f , ce qui est le cas d'un signal modulé en fréquence, on peut assimiler le sinus à son argument. La sortie est alors proportionnelle à Δf .

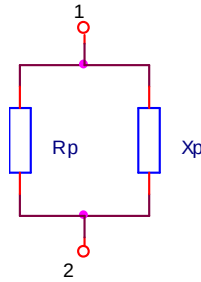
ANNEXES

- **Conversion série – parallèle d'impédance**
- **Théorème de Miller**
- **Le Bruit**
- **La Distorsion**

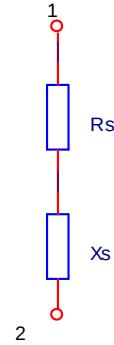
CONVERSION SERIE – PARALLELE D'IMPEDANCE

1 CONVERSION PARALLELE → SERIE

Etant donné une résistance R_p en parallèle avec une réactance X_p , on cherche l'impédance série équivalente



$$Z_p = R_p // jX_p$$



$$Z_s = R_s + jX_s$$

$$Z_p = \frac{jR_p X_p}{R_p + jX_p}$$

$$Z_p = \frac{jR_p X_p}{R_p + jX_p} = \frac{R_p X_p^2 + jR_p^2 X_p}{R_p^2 + X_p^2}$$

$$Z_p = \frac{R_p}{1 + \left(\frac{R_p}{X_p}\right)^2} + j \frac{\frac{R_p^2}{X_p}}{1 + \left(\frac{R_p}{X_p}\right)^2}$$

En identifiant Z_s et Z_p , on obtient :

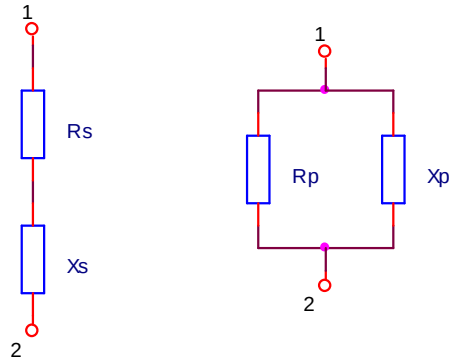
$$R_s = \frac{R_p}{1 + \left(\frac{R_p}{X_p}\right)^2} \quad X_s = \frac{\frac{R_p^2}{X_p}}{1 + \left(\frac{R_p}{X_p}\right)^2}$$

En combinant ces deux relations, il vient :

$$X_s = \frac{R_s R_p}{X_p}$$

2 CONVERSION SERIE → PARALLELE

Il s'agit du problème inverse du précédent.



Ecrivons les admittances des deux circuits.

$$Y_p = \frac{1}{R_p} + \frac{1}{jX_p} = \frac{1}{R_p} - j \frac{1}{X_p}$$

$$Y_s = \frac{1}{R_s + jX_s} = \frac{R_s - jX_s}{R_s^2 + X_s^2}$$

En divisant numérateur et dénominateur par R_s^2 :

$$Y_s = \frac{\frac{1}{R_s} - j \frac{X_s}{R_s^2}}{1 + \left(\frac{X_s}{R_s}\right)^2} = \frac{1}{R_s \left(1 + \left(\frac{X_s}{R_s}\right)^2\right)} - j \frac{X_s}{R_s^2 \left(1 + \left(\frac{X_s}{R_s}\right)^2\right)}$$

$$Y_p = 1/R_p - j/X_p$$

En identifiant, il vient :

$$\frac{1}{R_p} = \frac{1}{R_s \left(1 + \left(\frac{X_s}{R_s}\right)^2\right)} \qquad R_p = R_s \left(1 + \left(\frac{X_s}{R_s}\right)^2\right)$$

D'où en combinant ces deux relations :

$$X_p = \frac{R_s R_p}{X_s}$$

3 FACTEUR DE QUALITE

Posons

$$Q_s = \frac{X_s}{R_s} \qquad Q_p = \frac{R_p}{X_p}$$

$$Q_s = \frac{X_s}{R_s} = \frac{R_s R_p}{R_s X_p} = \frac{R_p}{X_p} = Q_p$$

On retrouve ici les définitions du facteur de qualité pour une impédance série ou parallèle.

THEOREME DE MILLER

Considérons le quadripôle de la fig. 1 a

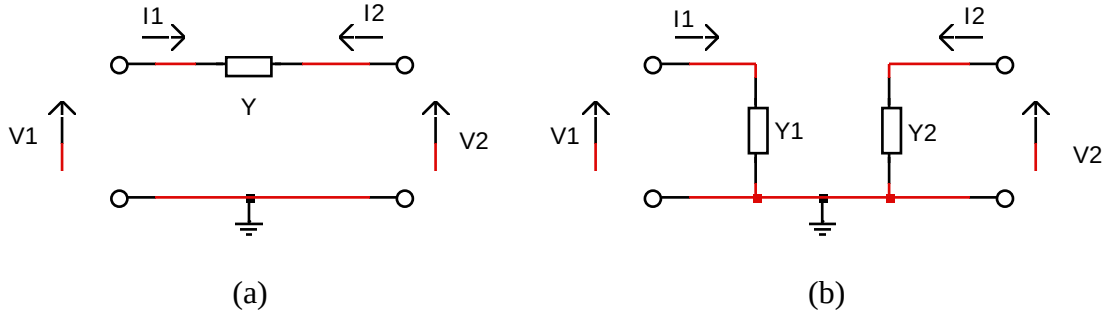


Figure 1 - Transformation de MILLER

Une admittance Y est connectée entre les noeuds d'entrée et de sortie.

Ce quadripôle est équivalent à celui de la fig. 1 b dans lequel l'admittance Y a été remplacée par deux admittances Y_1 et Y_2 reliées au point commun.

soit $K = V_2/V_1$ le gain en tension de ce quadripôle.

Le courant I_1 entrant par le noeud 1 s'écrit

$$I_1 = Y(V_1 - V_2) = YV_1(1 - \frac{V_2}{V_1}) = YV_1(1 - K)$$

dans le quadripôle équivalent nous avons $I_1 = Y_1 V_1$, en identifiant il vient :

$$Y_1 = Y (1 - K)$$

En opérant de la même manière pour le noeud de sortie, on a :

$$I_2 = Y(V_2 - V_1) = YV_2(1 - \frac{V_1}{V_2}) = YV_2(1 - \frac{1}{K})$$

en identifiant avec $I_2 = Y_2 V_2$, il vient :

$$Y_2 = Y(1 - \frac{1}{K})$$

Le théorème de Miller permet de transformer une admittance connectée entre deux points d'un circuit en deux admittances connectées en dérivation entre ces points et la masse.

LE BRUIT

Le bruit est une perturbation ayant diverses origines : agitation thermique, phénomènes naturels ou générés par l'activité humaine.

Le bruit peut altérer fortement l'intégrité d'un signal faible et le rendre inexploitable.

On distinguera deux types de bruits selon l'étendue de leur spectre et leur durée:

- le bruit blanc d'origine naturelle (thermique) dont le spectre est continu
- les bruits impulsionnels à spectre limité et/ou de courte durée tels que les parasites industriels, les phénomènes atmosphériques, les interférences électromagnétiques,...

Si les bruits impulsionnels peuvent être assez facilement éliminés, le bruit propre d'un système de réception constitue une limite au dessous de laquelle un signal ne pourra pas être correctement traité.

Bruit d'un quadripôle.

Le bruit à large spectre provient de plusieurs phénomènes tels que l'agitation thermique, le bruit de grenaille, la recombinaison des paires électron-trou dans les semiconducteurs, l'avalanche dans les jonctions PN,...

L'analyse du bruit est basée sur le concept de puissance disponible.

Considérons une source d'impédance $R_S + jX_S$ connectée à une charge d'impédance $R_L + jX_L$.

On sait que la puissance maximale que peut délivrer la source est obtenue lorsque l'impédance de la charge est complexe conjuguée de celle de la source (adaptation d'impédance).

Adaptation d'impédance : $R_S = R_L$ $X_S = -X_L$

Dans ces conditions, la moitié de la puissance est dissipée dans la source et l'autre moitié est transmise à la charge.

Si V_S est la fem de la source et V_L la tension aux bornes de la charge la puissance délivrée à la charge s'écrit:

$$P_L = V_S^2 / 4R_L = V_S^2 / 4R_S \text{ c'est la puissance disponible (available power).}$$

Bruit d'origine thermique

Le bruit thermique est généré par les vibrations aléatoires des électrons assurant la conduction électrique. Il augmente avec la température et présente un spectre continu.

On peut représenter le bruit thermique comme une infinité de générateurs élémentaires de fréquence, d'amplitude et de phase aléatoires.

Si l'on couple ce générateur de bruit à un filtre passe bande d'impédance adaptée, on obtient une certaine puissance de bruit dans la bande passante B:

$$P_W = kTB$$

Avec k : constante de Boltzmann = $1,38 \cdot 10^{-23}$ J/°K

T : température absolue

B : Bande passante équivalente de bruit

En élargissant la bande passante B du filtre on augmente la puissance de bruit disponible.

Par exemple, la puissance de bruit thermique disponible à $T = 17^\circ\text{C}$ (290°K) dans une bande passante de 1Hz est de $4 \cdot 10^{-21}$ W.

Cette puissance peut s'exprimer en dB par rapport à un niveau de référence fixé à 1mW:

$$P_{dBm} = 10 \log(P/1mW)$$

Soit -174 dBm

Gain en puissance d'un quadripôle

Le gain en puissance disponible d'un quadripôle est défini comme le rapport de la puissance disponible en sortie sur la puissance disponible de la source.

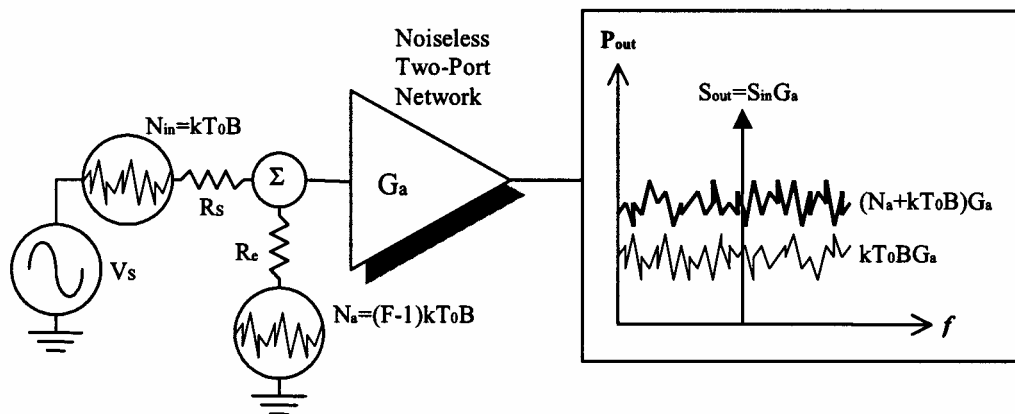
Le gain en puissance d'un quadripôle est égal à son gain en puissance disponible s'il est adapté en entrée et en sortie.

Facteur de bruit

Le facteur de bruit permet de quantifier le bruit apporté par un quadripôle.

Par définition on suppose que la température est de 290°K.

Le quadripôle est remplacé par le schéma équivalent suivant, faisant intervenir un quadripôle supposé sans bruit de gain disponible G_a et une source de bruit équivalente ramenée à l'entrée.



L'entrée est composée d'une source de signal S_{in} (sinusoïdal pour simplifier) à laquelle s'ajoute le bruit thermique $N_{in} = kT_0B$.

Si le quadripôle était sans bruit, la sortie serait la somme du signal amplifié et du bruit amplifié également:

$$S_{out} = S_{in}.G_a \quad N_{out} = N_{in}.G_a = kT_0B.G_a$$

Pour tenir compte du bruit apporté par le quadripôle lui même, on doit ajouter une source de bruit $N_a = (F-1)kT_0B$.

F est appelé figure de bruit du quadripôle (noise figure).

Dans ces conditions, le bruit en sortie devient:

$$N_{out} = (N_{in} + N_a).G_a = [kT_0B + (F-1)kT_0B].G_a = FkT_0B.G_a$$

Si la température est différente de $T_0=290^\circ K$, il suffit de remplacer T_0 par la température réelle dans l'équation.

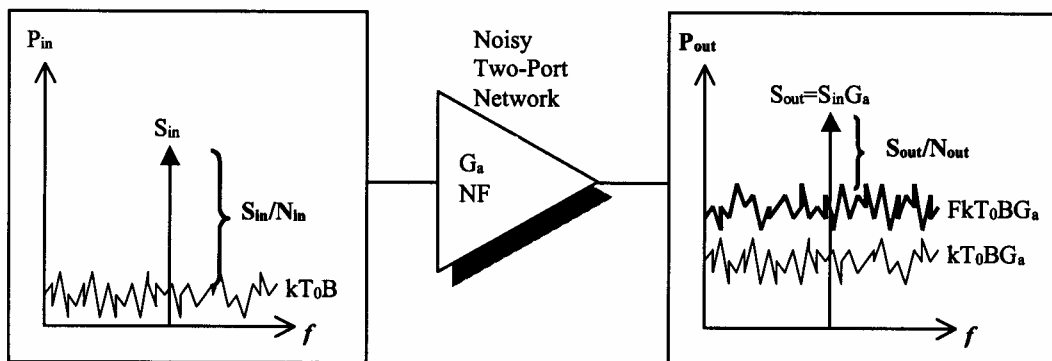
Calculons le rapport signal/bruit en entrée et en sortie à la température $T_0=290^\circ K$:

$$S_{in}/N_{in} = S_{in}/kT_0B$$

$$S_{out}/N_{out} = S_{out}/FkT_0B.G_a = G_a.S_{in}/FkT_0B.G_a$$

De ces deux équations on tire: $F = (S_{in}/N_{in}) / (S_{out}/N_{out}) = N_{out}/G_a N_{in} = (N_{in} + N_a)/N_{in}$

On constate que le rapport signal/bruit est dégradé par le passage dans le quadripôle.
La figure de bruit F permet de quantifier cette dégradation.



On définit le facteur de bruit NF exprimé en dB:

$$NF = 10\log(F) = 10\log[(S_{in}/N_{in}) / (S_{out}/N_{out})] = 10\log(S_{in}/N_{in}) - 10\log(S_{out}/N_{out})$$

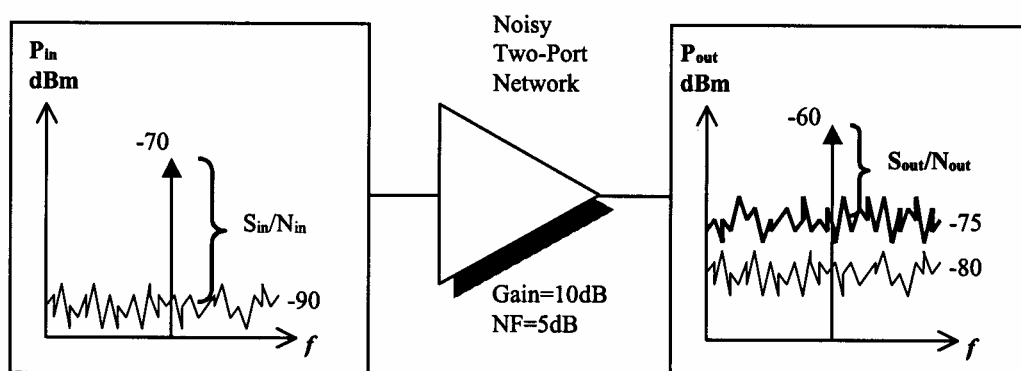
Où: S_{in} est la puissance du signal d'entrée,
 N_{in} la puissance de bruit thermique à l'entrée,
 S_{out} la puissance du signal de sortie,
 N_{out} la puissance de bruit en sortie.

On notera que le facteur de bruit est indépendant du niveau de signal et de la bande passante. En revanche, il est défini à $T_0=290^\circ K$ et dépend de la température ambiante.

Exemple:

Considérons un amplificateur dont le gain est de 10dB et le facteur de bruit 5dB, recevant un signal d'entrée d'amplitude -70dBm et un bruit thermique de -90dBm.

Le rapport signal/bruit en entrée est donc de 20dB.



En sortie le signal est amplifié jusqu'à -60dBm. Si l'amplificateur était sans bruit, le bruit en sortie serait de -80dBm et le rapport signal /bruit en sortie serait égal à celui d'entrée.

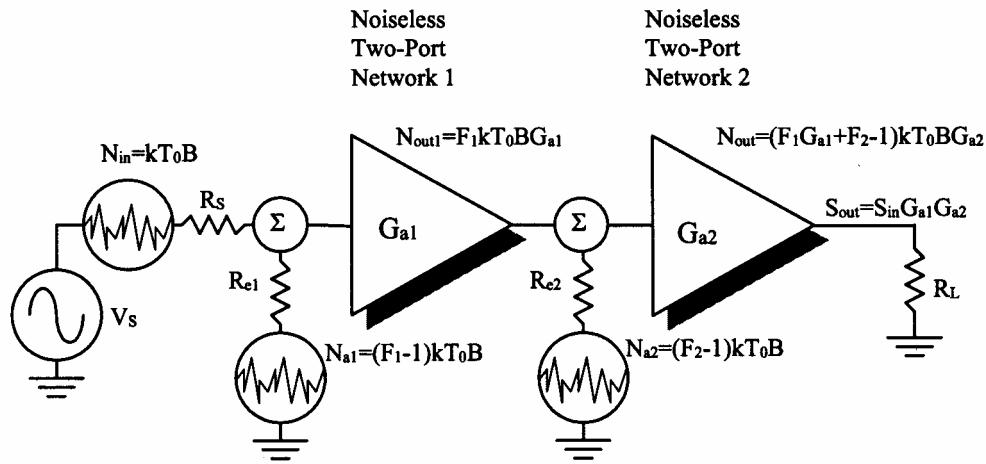
En réalité, le facteur de bruit de l'amplificateur se retranche du rapport signal /bruit d'entrée et l'on a:

rapport signal /bruit en sortie = 20dB - 5 dB = 15dB.

Cet amplificateur apporte donc un gain de 10dB mais il dégrade le rapport signal/bruit de 5dB. Cet exemple montre l'intérêt d'utiliser un amplificateur à faible bruit.

Facteur de bruit d'une chaîne d'amplification à plusieurs étages.

Un système est souvent composé de plusieurs étages montés en cascade. On se propose de calculer le facteur de bruit total du système à partir du facteur de bruit de chaque étage.
Considérons le système suivant composé de deux étages.



La puissance de bruit disponible en sortie s'écrit:

$$N_{out1} = F_1 k T_0 B G_{a1}$$

$$N_{out} = G_{a2}(N_{out1} + N_{a2}) = G_{a2}[F_1 k T_0 B G_{a1} + (F_2 - 1) k T_0 B]$$

$$N_{out} = [F_1 G_{a1} + F_2 - 1] k T_0 B G_{a2} = [F_1 + (F_2 - 1)/G_{a1}] k T_0 B G_{a1} G_{a2}$$

La figure de bruit totale est donnée par l'expression:

$$F_T = N_{out}/G_{aT}N_{in} \quad \text{où} \quad G_{aT} = G_{a1}G_{a2} \text{ est le gain total du système.}$$

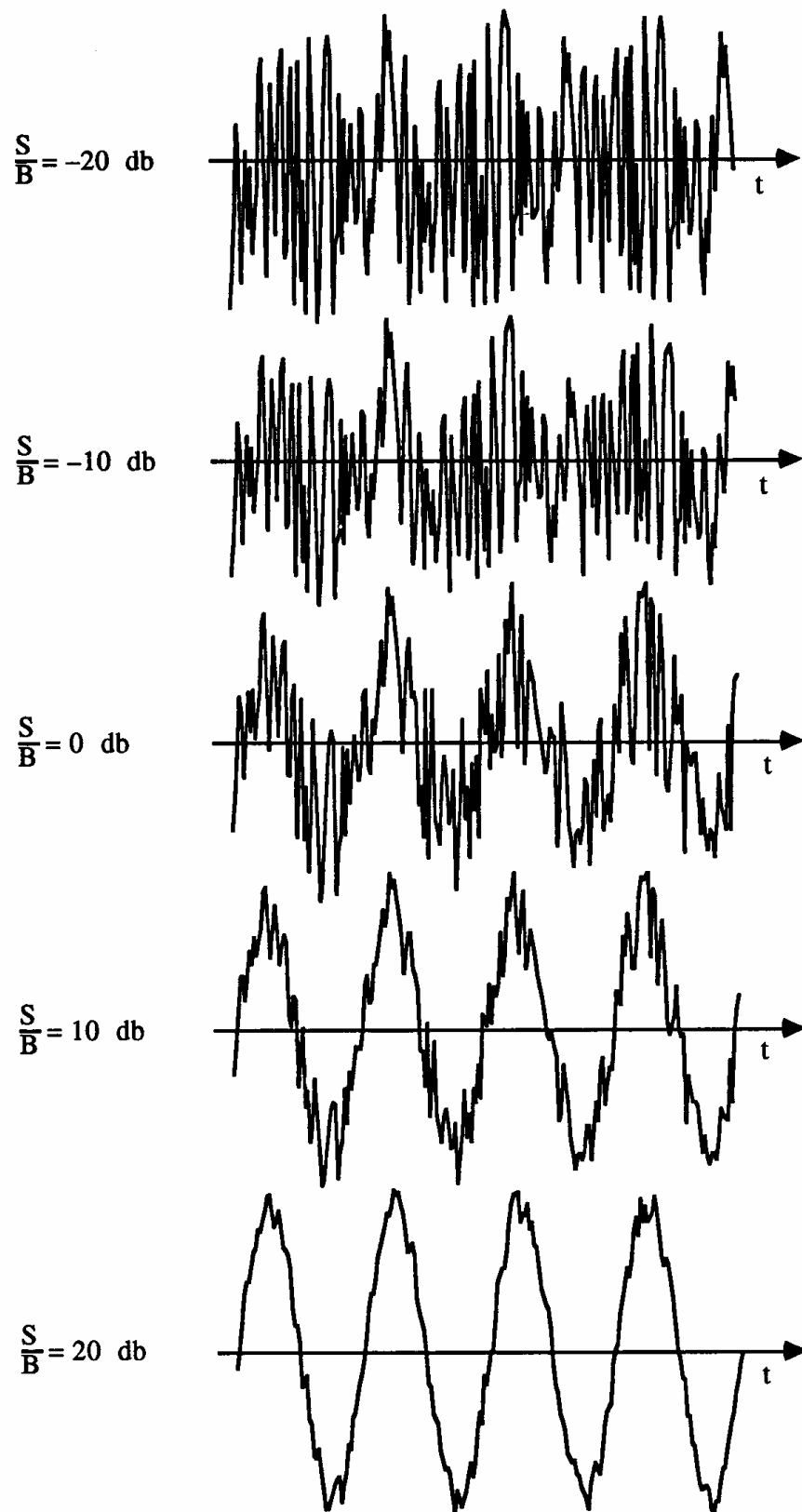
En combinant ces deux équations il vient:

$$F_T = F_1 + (F_2 - 1)/G_{a1}$$

On peut généraliser cette expression à un système à n étages:

$$F_T = F_1 + (F_2 - 1)/G_{a1} + (F_3 - 1)/G_{a1}G_{a2} + \dots + (F_n - 1)/G_{a1}G_{a2}\dots G_{a_{n-1}}$$

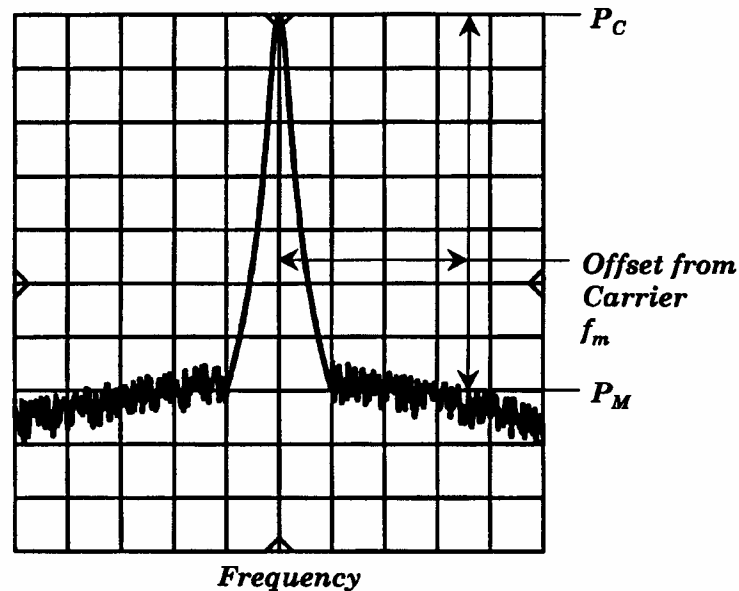
Cette équation montre que le facteur de bruit total du système dépend essentiellement des premiers étages (front end).



Représentation d'une sinusoïde bruitée (d'après M & F Biquard)

Bruit de phase

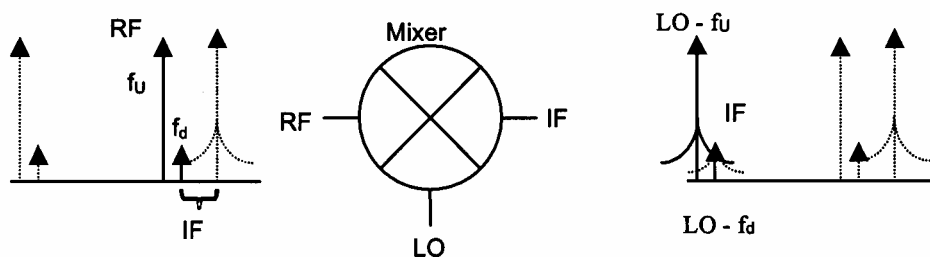
Le bruit de phase (phase noise) décrit les fluctuations aléatoires à court terme d'un oscillateur. En effet, le signal théoriquement pur produit par l'oscillateur est modulé par du bruit ce qui donne naissance à un spectre comportant des bandes latérales de bruit (noise sidebands) que l'on caractérise en dBc/Hz (amplitude correspondant à une bande passante de 1Hz, référencée à l'amplitude de l'onde pure, appelée porteuse (carrier)).



Spectre du signal délivré par un oscillateur montrant les bandes latérales de bruit.

Limitations dues au bruit de phase dans les mélangeurs.

Dans un convertisseur réel, le mélangeur reçoit le signal de l'oscillateur local (LO) avec son bruit de phase que l'on retrouvera à la sortie du mélangeur.



Supposons que l'entrée du mélangeur reçoive un signal faible de fréquence f_d , proche d'un signal f_u perturbateur beaucoup plus fort. En sortie du mélangeur on retrouve les termes $LO - f_u$ et $LO - f_d$. On voit sur la figure que le bruit de phase du signal indésirable $LO - f_u$ peut masquer le signal utile $LO - f_d$.

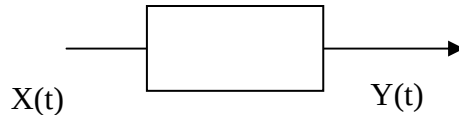
Cet exemple montre que le bruit de phase de l'oscillateur local est un paramètre crucial dans les systèmes de réception.

La distorsion

Un dispositif électronique n'est linéaire que dans une certaine plage de fonctionnement ou lorsque les signaux sont suffisamment faibles.

Par linéaire, on entend que la réponse à une somme de signaux pondérés est la somme des réponses aux différents signaux avec les mêmes coefficients de pondération.

Si cette hypothèse n'est pas vérifiée, les signaux sont dégradés, le dispositif introduit alors de la distorsion



Considérons un circuit dont la réponse n'est pas linéaire.

La fonction de transfert peut se mettre sous la forme d'un polynôme:

$$Y(t) = k_1 x(t) + k_2 x^2(t) + k_3 x^3(t) + \dots$$

Distorsion harmonique

Si l'on applique à l'entrée un signal sinusoïdal $x(t) = A \cos \omega t$, on obtient en sortie un signal :

$$Y(t) = k_1 [A \cos \omega t] + k_2 [A \cos \omega t]^2 + k_3 [A \cos \omega t]^3 + \dots$$

D'après les relations trigonométriques:

$$\cos^2 a = (1 + \cos 2a)/2 \quad \cos^3 a = (3 \cos a + \cos 3a)/4 \quad \text{il vient:}$$

$$Y(t) = k_1 A \cos \omega t + (k_2 A^2/2) [\cos 2\omega t + 1] + (k_3 A^3/4) [3 \cos \omega t + \cos 3\omega t] + \dots$$

On voit apparaître des termes de pulsation $n\omega t$. Ce sont les harmoniques du signal d'entrée.

D'une manière générale la sortie pourra s'écrire:

$$V(t) = V_0 + V_1 \cos \omega t + V_2 \cos 2\omega t + V_n \cos n\omega t + \dots$$

On définit alors le taux de distorsion harmonique par la relation:

$$THD = 100 \frac{\sqrt{\sum_{n=2}^{\infty} V_n^2}}{V_1}$$

Distorsion d'intermodulation

Supposons que le signal $x(t)$ soit composé de deux signaux sinusoïdaux x_1 et x_2 de fréquences relativement proches de manière à pouvoir considérer que la fonction de transfert est la même pour les deux signaux.

$$X(t) = A_1 \cos \omega_1 t + A_2 \cos \omega_2 t$$

$$Y(t) = k_1[A_1 \cos \omega_1 t + A_2 \cos \omega_2 t] + k_2[A_1 \cos \omega_1 t + A_2 \cos \omega_2 t]^2 + k_3[A_1 \cos \omega_1 t + A_2 \cos \omega_2 t]^3 + \dots$$

En développant et d'après les relations trigonométriques:

$$\cos^2 a = (1 + \cos 2a)/2 \quad \cos^3 a = (3 \cos a + \cos 3a)/4 \quad 2 \cos a \cos b = \cos(a+b) + \cos(a-b)$$

Il vient en se limitant au troisième ordre:

$$\begin{aligned} Y(t) = & k_1[A_1 \cos \omega_1 t + A_2 \cos \omega_2 t] \\ & + k_2 [(A_1^2/2)(1 + \cos 2\omega_1 t) + (A_2^2/2)(1 + \cos 2\omega_2 t) + A_1 A_2 (\cos(\omega_1 + \omega_2)t + \cos(\omega_1 - \omega_2)t)] \\ & + k_3 \{ A_1^3 (3 \cos \omega_1 t + \cos 3\omega_1 t)/4 + A_2^3 (3 \cos \omega_2 t + \cos 3\omega_2 t)/4 \\ & + A_1^2 A_2 [1,5 \cos(\omega_2 t) + 0,75 \cos(2\omega_1 + \omega_2)t + 0,75 \cos(2\omega_1 - \omega_2)t] \\ & + A_2^2 A_1 [1,5 \cos(\omega_1 t) + 0,75 \cos(2\omega_2 + \omega_1)t + 0,75 \cos(2\omega_2 - \omega_1)t] \} \end{aligned} \quad (1)$$

On constate que le signal de sortie comprend des termes de pulsations:

$$\omega_1, \omega_2, 2\omega_1, 2\omega_2, \omega_1 + \omega_2, \omega_1 - \omega_2, 3\omega_1, 3\omega_2, 2\omega_1 + \omega_2, 2\omega_2 + \omega_1, 2\omega_1 - \omega_2, 2\omega_2 - \omega_1$$

D'une manière générale, si l'on injecte à un quadripôle non linéaire deux signaux de pulsations ω_1 et ω_2 le signal de sortie comporte des termes de la forme $M\omega_1 \pm N\omega_2$ où M et N sont deux nombres entiers $[0, 1, 2, \dots]$.

Ces termes sont appelés **produits d'intermodulation** du $(M+N)$ ième ordre.

Exemple

Supposons qu'à l'entrée d'un récepteur FM on reçoive deux signaux de fréquences $f_1 = 100\text{MHz}$ et $f_2 = 101\text{MHz}$

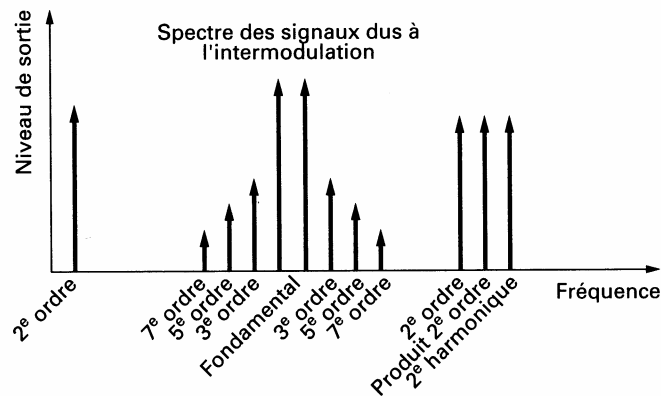
Les produits d'intermodulation du second ordre sont:

$$\begin{array}{ll} 2f_1 = 200\text{MHz} & 2f_2 = 202\text{MHz} \\ f_1 + f_2 = 201\text{MHz} & f_2 - f_1 = 1\text{MHz} \end{array}$$

Les produits d'intermodulation du troisième ordre sont:

$$\begin{array}{ll} 3f_1 = 300\text{MHz} & 3f_2 = 303\text{MHz} \\ 2f_1 - f_2 = 99\text{MHz} & 2f_2 - f_1 = 102\text{MHz} \\ 2f_1 + f_2 = 301\text{MHz} & 2f_2 + f_1 = 302\text{MHz} \end{array}$$

On remarque que les produits d'intermodulation du troisième ordre ($2f_1 - f_2$ et $2f_2 - f_1$) donnent des composantes dont la fréquence est proche des fréquences utiles f_1 et f_2 . Leur élimination par filtrage peut s'avérer impossible.



Compression du gain

D'après l'équation (1), on voit que l'expression du signal utile s'écrit:

$$[k_1 A_1 + k_3 (0,75 A_1^3 + 1,5 A_1 A_2^2)] \cos \omega_1 t$$

Dans la majorité des cas, k_3 est négatif.

Soit x_1 le signal utile de faible amplitude et x_2 un signal perturbateur d'amplitude plus élevée.

On remarque que la forte amplitude A_2 du signal x_2 peut conduire une réduction notable du signal utile x_1 .

Le gain du circuit est donc réduit. Pour éviter cet effet de **compression du gain**, il faut rendre le coefficient k_3 le plus petit possible.

On constate également que l'amplitude du signal utile dépend de l'amplitude du signal perturbateur. Ce phénomène appelé **transmodulation** est très gênant avec des signaux modulés en amplitude.

Point de compression à -1dB

Considérons le cas d'un signal unique d'amplitude A présent à l'entrée du circuit.

L'expression du signal utile en sortie est:

$$[k_1 A + k_3 (0,75 A^3)] \cos \omega_1 t$$

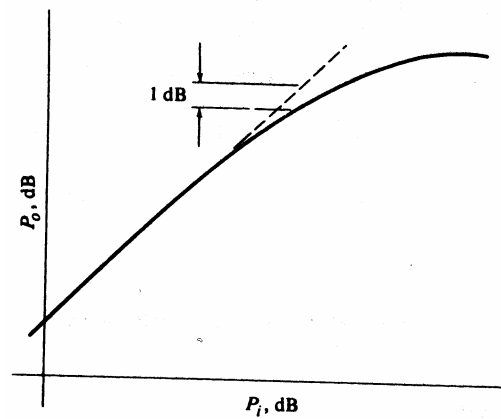
Si le circuit était parfaitement linéaire on aurait $k_1 A$.

Le rapport de ces deux termes est appelé facteur de compression du gain.

$$\text{facteur de compression du gain} = (k_1 + 0,75 k_3 A^2) / k_1$$

La figure suivante montre la déviation du gain par rapport à courbe idéale (linéaire).

Le point pour le lequel le gain en puissance chute de 1 dB par rapport à la courbe idéale est appelé point de compression à -1dB. C'est une caractéristique fondamentale de tout amplificateur.



Point de compression à -1 dB

Amplitude des produits d'intermodulation

D'après l'équation (1), on voit que les produits d'intermodulation du troisième ordre les plus gênants sont donnés par:

$$k_3 A_1^2 A_2 [0,75 \cos(2\omega_1 - \omega_2)t] + k_3 A_2^2 A_1 [0,75 \cos(2\omega_2 - \omega_1)t]$$

alors que le signal utile est:

$$[k_1 A_1 + k_3 (0,75 A_1^3 + 1,5 A_1 A_2^2)] \cos \omega_1 t$$

Supposons pour simplifier que les amplitudes A_1 et A_2 soient égales à A .

Nous observons que l'amplitude des produits d'intermodulation du troisième ordre croît comme A^3 , alors que celle du signal utile croît comme A .

De même, l'amplitude des produits d'intermodulation du second ordre croît comme A^2

Les produits d'intermodulation du troisième ordre croissent 3 fois plus vite (en dB) que le signal utile lorsqu'on augmente le niveau de puissance, et ceux du second ordre deux fois plus vite.

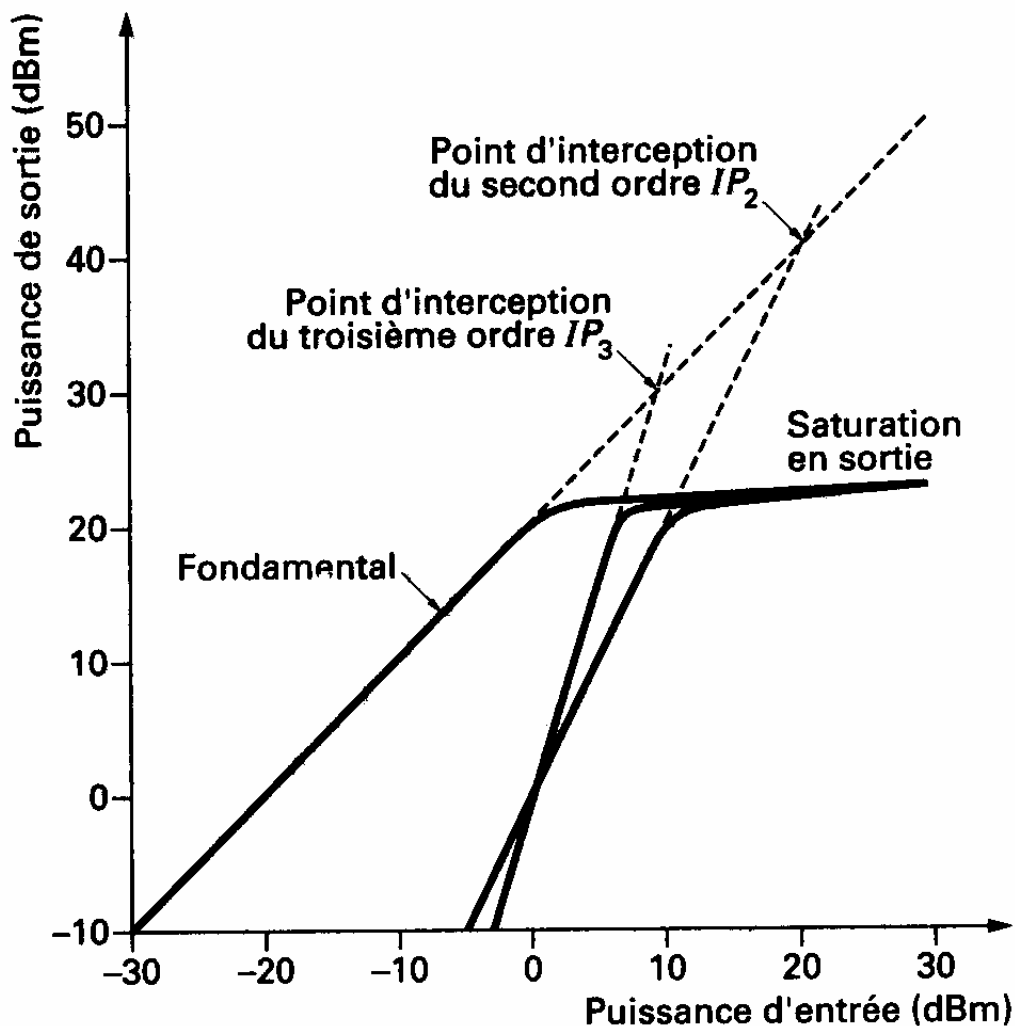
Points d'interception du second et du troisième ordre IP2 et IP3

Dans le diagramme $P_{\text{sortie}} = f(P_{\text{entrée}})$, si la pente de la courbe dans sa partie linéaire est a , les produits d'intermodulation du second ordre suivront une droite de pente $2a$ et ceux du troisième ordre une droite de pente $3a$.

Le point d'intersection des droites de pente a et $2a$ est appelé point d'interception du second ordre IP2.

Le point d'intersection des droites de pente a et $3a$ est appelé point d'interception du troisième ordre IP3.

Ces points sont des points théoriques, car ils sont situés au dessus de la saturation, mais ils caractérisent bien la linéarité du circuit. Plus les valeurs de IP2 et IP3 sont élevées, meilleure est la linéarité du circuit.



Dans une chaîne de transmission les éléments les plus critiques sont souvent les amplificateurs (notamment de puissance) et les mélangeurs.

Pour les amplificateurs, la distorsion d'intermodulation est mesurée à l'aide de deux signaux de fréquences F_1 et F_2 relativement proches et situées dans la bande passante de l'amplificateur. On relève le niveau du signal (F_1 ou F_2) et des produits d'intermodulation du troisième ordre ($2F_1-F_2$ ou $2F_2-F_1$). On en déduit le point d'interception IP_3 par simple construction graphique.

Pour les mélangeurs, on caractérise souvent l'intermodulation entre le signal RF et l'oscillateur local LO. On relève alors les niveaux du signal $|RF-LO|$ et des produits d'intermodulation du troisième ordre $2RF-LO$ ou $RF-2LO$.