

Principe de Télécommunications, Notes de cours (provisoires)

A.L. Fawe – L.Deneire

Année 1995-1996

Chapter 1

Introduction aux Télécommunications.

1.1 Définition des télécommunications.

*”Les télécommunications au sens large comprennent l’ensemble des moyens techniques nécessaires à l’acheminement aussi fidèle et fiable que possible d’informations entre deux points a priori quelconques, à une distance quelconque, avec des coûts raisonnables”*¹ et à des instants quelconques.

Il est à noter que cette définition est très large. Nous ajouterons que, dans le cadre de ce cours, les principaux moyens de transmission utilisés seront de nature électromagnétique. En outre, la théorie des télécommunications ne concerne que l’information, et non son support matériel (papier, disques, ...) ². Les notions de fidélité (conformité du message reçu et du message émis) et de fiabilité (résistance à des pannes partielles du système) seront le souci premier du concepteur. Les messages à transmettre peuvent être de nature quelconque (paroles, images, données de tous types). D’autre part, les besoins en communication ne sont pas nécessairement symétrique, on peut en effet opter pour une transmission *full-duplex*, où chaque interlocuteur émet simultanément avec l’autre; *semi-duplex*, où chaque interlocuteur émet alternativement ou *simplex*, c’est-à-dire unidirectionnelle. En outre, dans le cas de communications numériques, les débits ne sont pas nécessairement les mêmes dans les deux sens. Le cas le plus simple étant le Minitel français, où la transmission se fait à 1200 bits/s. dans le sens service ->utilisateur et à 75 bits/s. dans le sens utilisateur ->service. L’avènement des transmissions d’images (video on demand) accentue ce problème et est une des motivations du développement de l’A.T.M. (Asynchronous Transfer Mode) qui permet, sur un canal commun, la coexistence de débits très différents.

Nous avons ajouté, à la définition de Fontollet, la notion temporelle. En effet, si la simultanéité entre l’émission et la réception est souvent désirée (cas de la téléphonie), on peut désirer stocker une information pour la reprendre plus tard, à un autre endroit. Cette transmission, non plus d’un point à un autre, mais d’un temps à un autre, peut être abordée par des techniques similaires, dont les techniques de codage issues de la théorie de l’information. Un cas typique est celui des disques compacts audio ou vidéo, des CD-interactifs et de l’enseignement à distance.

1. Systèmes de télécommunications, P.-G. Fontollet, Traité d’Electricité, volume XVIII, Editions Georgi, Lausanne, 1983, p. 1

2. Le support de transmission (câble, fibre optique, ondes hertziennes) n’est pas à proprement parler un support d’information, mais bien le canal de transmission.

La chaîne de télécommunication est formée principalement de :

1. Le canal (ligne, câble coaxial, guide d'onde, fibre optique, lumière infra-rouge, canal hertzien, etc.)
2. L'émetteur, qui a comme fonction de fournir un signal (représentant le message) adapté au canal.
3. Le récepteur dont la fonction est de reconstituer le message après observation du signal présent sur le canal.

1.2 Historique.

De la naissance du télégraphe en 1838, mis au monde par Morse, aux transmissions d'images spatiales (Voyager II) en passant par la transmission en direct des premiers pas et mots de Armstrong sur la lune, les télécommunications, à l'instar de toutes les techniques de l'électronique, ont explosé dans la seconde moitié du 20ème siècle. Les quelques points de repère décrits ci-dessous vous donneront une idée de l'histoire des télécommunications [?], [?].

Année	Événement
1826	Loi d'Ohm
1838	Samuel Morse invente un système de transmission codée des lettres de l'alphabet, qui deviendra le télégraphe. Son code tient compte de la fréquence relative des lettres dans la langue anglaise pour optimiser le temps de transmission d'un message. A ce titre, c'est un précurseur intuitif de la théorie et du codage.
1858	Un câble unifilaire télégraphique est posé à travers l'Atlantique (il fonctionnera un mois).
1864	Etablissement des équations électromagnétiques de Maxwell
1870	Liaison télégraphique entre Londres et Calcutta (11.000 km)
1876	Alexandre Graham Bell dépose le brevet du téléphone. Il s'agit d'un moyen de transmettre électriquement des sons à l'aide d'une résistance variable.
1891	Premier sélecteur automatique télécommandé par le poste d'abonné (Almon Srowger)
1897	Marconi dépose le brevet d'un système de télégraphie sans fil.
1907	Lee de Forest invente l'amplificateur à triode.
1915	Première liaison téléphonique (par ondes courtes) transcontinentale par Bell System.
1920	Application de la théorie de l'échantillonnage aux communications.
1937	La modulation par impulsions codées (PCM: Pulse-code modulation), inventée par Alec Reeves, permet la représentation numérique d'informations analogiques (premier codage de la voix).
1938	Début des émissions télévisées.
1940-45	Développement du radar, application des méthodes statistiques à l'étude des signaux.
1948	Invention du transistor ; publication de la théorie mathématique de l'information par C. Shannon.
1956	Premier câble téléphonique transatlantique.
1962	Premier satellite de communications (Telstar I) appliqué à la transmission transatlantique de télévision.
1965	Premier satellite géostationnaire (Intelsat I).
1969	On a décroché la lune!
197x	Avènement des grands réseaux informatiques intercontinentaux, communications satellites commerciales, augmentation exponentielle des taux de transmission.
1980	Premières images de Jupiter et de Saturne venant d'une sonde spatiale.
1991	Réseau Numérique à Intégration de Services (RNIS) en Belgique
1994	Lancement du G.S.M. en Belgique
1995	Votre premier cours de Télécommunications ...

1.3 La chaîne de transmission.

1.3.1 Schéma de base d'une chaîne de transmission.

Le schéma de base d'une chaîne de transmission peut être représenté par la figure 1.1.

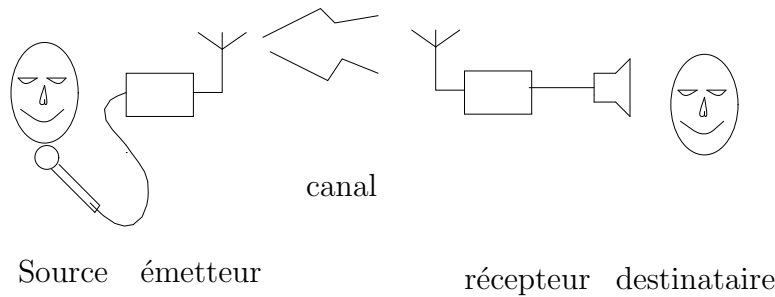


FIG. 1.1 – Schéma de base d'une chaîne de transmission.

Les cinq éléments qui y figurent sont définis comme suit :

- La source produit le message à transmettre.
- L'émetteur produit un signal adapté au canal de transmission.
- Le canal de transmission constitue le lien entre émetteur et récepteur.
- Le récepteur capte le signal et recrée le message.
- Le destinataire traite le message reçu.

La source, le destinataire, et dans une certaine mesure le canal de transmission, nous sont imposés. Tout l'enjeu consistera à choisir et concevoir adéquatement l'émetteur et le récepteur. Avant de donner un aperçu plus détaillé de ceux-ci, précisons les notions de source et de canal.

1.3.2 Le canal de transmission.

Le propre d'une transmission étant de se faire à distance, il faut utiliser un milieu physique qui assure le lien entre la source et le destinataire et un signal approprié au milieu choisi, en ce sens qu'il s'y propage bien. Citons par exemple un signal électrique dans un milieu conducteur ou un signal électromagnétique en espace libre.

La notion de canal de transmission est délicate à expliciter. Nous commencerons par des définitions assez vagues que nous préciserons au fil de ce paragraphe.

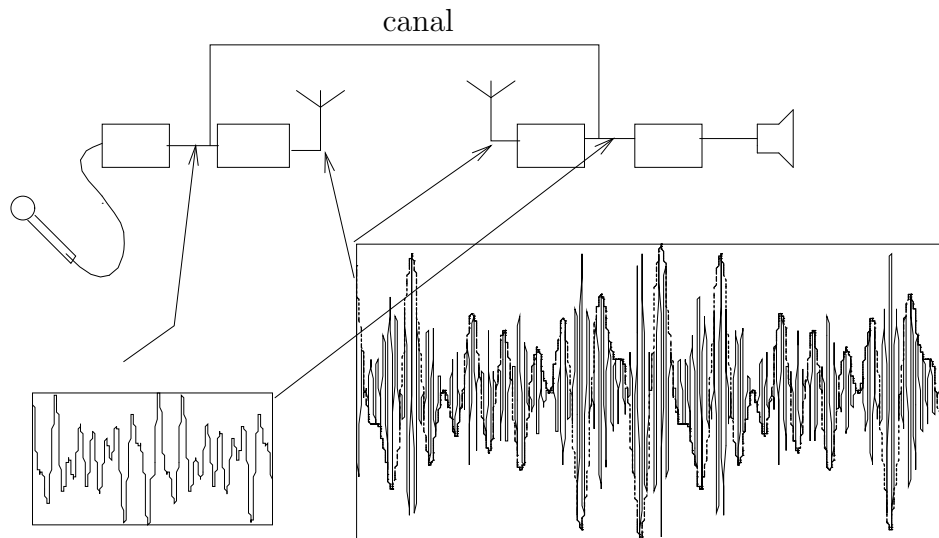
Notion de canal de transmission.

Selon le contexte, le terme de canal de transmission a des significations différentes. Le canal de transmission au sens de la propagation est la portion du milieu physique utilisée pour la transmission particulière étudiée. On parle ainsi de canal ionosphérique, de canal troposphérique, ...

Le canal de transmission au sens de la théorie des communications inclut le milieu physique de propagation et également des organes d'émission et de réception (Figure 1.2).

Les frontières qui délimitent le canal dépendent des fonctions assignées à l'émetteur et au récepteur. Sans entrer dans le détail de leurs structures, il est nécessaire d'en faire un examen préalable. Un émetteur réel peut être considéré comme ayant deux fonctions principales réalisées par deux éléments distincts :

1. l'émetteur théorique engendre un signal électrique $S(t)$ que nous définissons comme le signal émis. C'est un signal en bande de base ou sur fréquence porteuse qui véhicule

FIG. 1.2 – *Canal de transmission.*

le message numérique. Ce signal est modélisé par un processus aléatoire dont les caractéristiques exactes dépendent de celles de l'émetteur.

2. L'ensemble des organes d'émission transforme le signal $S(t)$ pour l'adapter au milieu physique dans lequel il va se propager. Parmi les traitements effectués, on peut inclure les opérations de changement de fréquence, de filtrage, d'amplification, de transformation d'un signal électrique en signal électromagnétique, d'émission proprement dite dans le milieu choisi (réalisée à l'aide de dispositifs tels que les antennes ...).

De la même façon, un récepteur réel peut être dissocié en deux blocs fonctionnels :

1. L'ensemble des organes de réception comprend les dispositifs de réception dans le milieu physique (antennes ...) et réalise les opérations de transformation de la nature du signal, d'amplification, de changement de fréquence ... Nous définissons le signal électrique $R(t)$ ainsi obtenu comme étant le signal reçu. Ce signal n'est malheureusement pas une reproduction parfaite du signal émis en raison de diverses perturbations que nous allons aborder par la suite.
2. Le récepteur théorique traite le signal afin de fournir le message au destinataire.

Le canal à bruit blanc additif gaussien (BBAG).

Les modèles utilisés pour représenter le canal de transmission défini ci-dessus sont relativement simples. Le plus simple et le plus classique est le canal à bruit blanc additif gaussien (canal BBAG). En sortie de ce canal, le signal reçu résulte de l'addition du signal émis et d'un bruit blanc. Si on excepte ce bruit, le signal émis ne subit aucune modification : nous dirons que le canal est sans distorsion. Le bruit additif est indépendant du signal. Il est modélisé par un processus aléatoire stationnaire, blanc, gaussien et centré. Sa densité spectrale bilatérale de puissance est constante.

FIG. 1.3 – Canal à bruit blanc additif gaussien (canal BBAG).

Caractéristiques du canal de transmission.

Après avoir défini le canal de transmission et son modèle, examinons les caractéristiques de ce canal afin de savoir si le modèle simple, canal BBAG, est approprié. Toute une série de facteurs interviennent pour modifier le signal électrique entre l'entrée du canal (signal émis $S(t)$) et la sortie du canal (signal reçu $R(t)$). Passons ceux-ci en revue et vérifions la validité du modèle du canal adopté (addition d'un bruit et absence de distorsion).

– **Le phénomène de propagation dans le milieu physique (atmosphère, câble, fibre, vide).**

Pour une transmission sur voie radioélectrique avec fréquence porteuse élevée, le bruit est pour l'essentiel d'origine interne. Il joue un rôle d'autant plus important que le signal reçu est faible. Dans ces conditions, pour traiter le signal, le récepteur doit opérer avant toute autre chose une amplification du signal. Etant imparfaite, celle-ci ajoute au signal amplifié un bruit thermique qui peut être modélisé par un bruit blanc additif gaussien. Ce bruit de réception est la principale cause des erreurs de transmission. Les autres défauts accentuent ce bruit et par là-même ont une influence importante sur la qualité de la transmission.

Un paramètre important qui sert à caractériser le fonctionnement du récepteur est le rapport signal à bruit, rapport des puissances de signal et de bruit reçus, évalué au niveau du récepteur. Si on se réfère au modèle canal BBAG, la puissance du signal utile est la même à l'entrée et à la sortie du canal.

Par contre, un canal réel provoque un affaiblissement de propagation qui augmente avec la distance entre émetteur et récepteur. Cet affaiblissement n'est pas pris en compte par le canal BBAG, car du point de vue théorique, il ne s'agit que d'un simple facteur d'échelle. Du point de vue pratique, il implique une amplification et éventuellement des amplificateurs et/ou répéteurs, et est donc un facteur de bruit.

De plus, le canal opère un filtrage qui est dû aux organes d'émission et de réception, et au milieu physique (un câble possède une fonction de transfert dont le module est de la forme $e^{-u\sqrt{f}}$; en propagation radioélectrique, il peut exister des trajets multiples, ...). La fonction de transfert n'apparaît pas dans le schéma du canal BBAG. Toutefois, il est possible de reprendre les raisonnements pour inclure le filtrage dû au canal, c'est ce qui sera fait dans les cours de "Questions spéciales et compléments de Télécommunications", en troisième année.

Une caractéristique du canal de transmission associée à la notion de fonction de transfert est sa bande passante. Un canal de type passe-bas permet une transmission en bande de

base, alors qu'un canal de type passe-bande nécessite l'emploi d'une fréquence porteuse. Dans la quasi-totalité des cas, il n'y a pas émission d'un signal unique à partir d'un émetteur réel. A chaque signal est alloué une certaine bande de fréquence en fonction de contraintes physiques ou réglementaires assignées au système (règles internationales, normes ...). Le canal BBAG peut-il alors être utilisé pour l'étude d'une transmission à bande limitée? Heureusement oui, puisque cette notion de canal à bande limitée n'est qu'une notion abstraite³. L'impératif de limitation de la bande de fréquence occupée explique l'importance du filtrage en transmission.

- **Les contraintes du système de transmission : le même système peut servir à la transmission conjointe d'autres signaux. L'environnement électromagnétique dans le milieu de transmission créé par d'autres systèmes utilisant le même milieu.**

La transmission conjointe de plusieurs signaux, habituellement dans des canaux fréquentiels différents, est reprise ci-dessus dans la limitation de bande du canal utilisé. Malgré les efforts de coordination entre les divers utilisateurs, sans parler de pollueurs (émissions radioélectriques parasites), les problèmes de compatibilité électromagnétique ne sont pas des plus simples à résoudre.

- **Les éventuelles transformations effectuées dans l'émetteur réel et le récepteur réel, mais que nous avons incluses dans le canal de transmission.**

Celles-ci introduisent en général un délai τ . Sa présence peut se modéliser très simplement au moyen d'une fonction de transfert de la forme $e^{-j\omega\tau}$. Le récepteur ne fonctionnant qu'à partir du signal reçu, la valeur de ce délai est sans aucune influence sur le problème théorique de réception. Par souci de simplification, nous ne le ferons donc pas apparaître dans le modèle du canal. Par contre, la valeur du délai de transmission peut revêtir une grande importance dans la définition du système pour des raisons de qualité ou de principe de fonctionnement.

Certaines imperfections de l'émetteur, inhérentes à la technologie utilisée, sont incluses dans le canal de transmission. En particulier des distorsions non linéaires apparaissent lorsqu'un ou plusieurs signaux sont amplifiés par des dispositifs non linéaires (par saturation par exemple). On parle alors de canal non linéaire, et le canal BBAG ne constitue qu'une approximation grossière de la réalité. Les non-linéarités ne sont pas abordées dans ce cours.

L'utilisation du canal BBAG reste justifiée dans le cas de milieux de transmission naturels (vide, atmosphère par temps claire), et fournit de précieux enseignements pour les autres cas. Toutefois, l'atmosphère est non stationnaire et non homogène. On constate des fluctuations d'atténuation du signal reçu et/ou des modifications très notables de la fonction de transfert du canal. Les problèmes de modification d'atténuation sont résolus par la prise de marges qui garantissent la qualité de transmission pendant un pourcentage donné du temps. Pour les fluctuations de la fonction de transfert, on opère de la même façon, ou bien on utilise des méthodes d'égalisation auto-adaptatives qui procurent des gains importants.

Le canal BBAG sera utilisé dans tout le cours de principes.

3. Au risque de faire bondir votre sens physique! Au sens des transmissions, le canal est effectivement à bande limitée. Ici, nous considérons, pour rappel, le canal après le premier élément de l'émetteur réel, et donc avant la limitation de bande.

FIG. 1.4 – *Émetteur*.

1.3.3 L'émetteur.

Comme l'illustre la figure 1.2, le canal de transmission "théorique" comprend une partie des éléments de l'émetteur réel. L'émetteur "théorique" que nous considérons dans la suite a pour fonction de transformer le message (analogique ou numérique) en un signal électrique $S(t)$, de nature analogique. Le signal émis par la source (suite de symboles, pression, image, ...) est ainsi transformé en une grandeur physique variable en fonction du temps. Le signal émis (issu de l'émetteur) doit être adapté aux contraintes imposées par le canal de transmission, contraintes au premier rang desquelles figure la nécessité de n'occuper que la bande de fréquence permise. La conception de l'émetteur doit également prendre en compte les problèmes de fonctionnement du récepteur, les interactions avec le système de codage-décodage correcteur d'erreurs s'il existe, et les impératifs de fonctionnement liés au système de transmission lui-même.

Pour répondre à ces contraintes, l'émetteur réalise deux opérations représentées schématiquement sur la figure 1.4.

Modulation.

La seconde opération est une modulation qui transforme le signal électrique précédent en un autre signal mieux adapté à un canal de transmission passe-bande. Celle modifie le signal en faisant intervenir, de façon plus ou moins explicite, une fréquence particulière dite fréquence porteuse. La modulation produit un signal de caractéristiques spectrales appropriées à la transmission sur canal de type passe-bande. On parle alors de transmission sur fréquence porteuse.

Le signal issu du modulateur, que nous appellerons signal sur fréquence porteuse, ou signal modulé, présente des différences notables avec un signal en bande de base. Nous décrirons les principaux types de modulation (analogique et numérique), avec leur densité spectrale respec-

tive.

Un élément supplémentaire, le filtrage, apparaît sur la figure 1.4, soit avec une caractéristique passe-bas pour un signal en bande de base, soit avec une caractéristique passe-bande pour un signal modulé. Le filtrage assure la mise en forme définitive du signal avant l'émission, compte tenu du codage ou de la modulation utilisée et des contraintes du canal. La partie consacrée à la transmission en bande de base fournira des informations sur la conception du filtrage passe-bas. Les chapitres consacrés à la démodulation fourniront des résultats sur les bandes utiles et le filtrage de ces signaux. La mise en forme à la fois spectrale et temporelle peut être effectuée sur le signal modulé, mais aussi en bande de base, c'est-à-dire sur le signal issu du codeur binaire à signal. A ce propos, remarquons que les frontières qui séparent les différentes fonctions de l'émetteur sont parfois difficiles à définir.

Pour en terminer avec cette partie consacrée à l'émission, examinons la question de l'adaptation du signal aux conditions de transmission. Comment savoir si le signal émis est bien adapté si ce n'est en regardant simultanément :

- la qualité du message reçu, c'est-à-dire le taux d'erreur sur les bits (BER) ou le rapport signal/bruit ($\frac{S}{N}$),
- comment sont vérifiées les autres contraintes (puissance hors bande, etc.).

Cela montre que la mise au point de l'émetteur ne peut se faire indépendamment de celle du récepteur. La définition des chaînes de transmission idéale conduira à des conditions qui portent à la fois sur l'émetteur et sur le récepteur, notamment sur leurs fonctions de filtrage. Le filtrage y apparaîtra comme une fonction essentielle qui devra répondre au double but :

- assurer le minimum de distorsion au signal qui sera traité par le récepteur,
- limiter l'occupation des fréquences à la seule bande permise.

1.3.4 Le récepteur.

Nous venons de voir que l'émetteur permet de transcrire le message en un signal afin de le transmettre sur le canal. Inversement, le récepteur doit extraire le message du signal reçu. Pour cela, il procède soit de manière séquentielle en prenant une suite de décisions sur les symboles successifs du message émis dans le cas numérique, soit par simple démodulation dans le cas analogique.

Le travail du récepteur est complexe en raison de la différence entre le signal émis et le signal reçu. Un bon fonctionnement du récepteur est lié à l'exploitation des connaissances a priori de la structure du signal émis, mais également des conditions dans lesquelles s'est déroulée la transmission.

La principale perturbation subie par le signal transmis est l'addition d'un bruit gaussien. Même en limitant la bande de fréquence du signal reçu, le récepteur doit fonctionner avec un signal utile entaché d'un signal perturbateur dont l'amplitude à un instant donné suit une loi de probabilité gaussienne. Intuitivement, on comprend que le travail du récepteur devient extrêmement délicat lors d'une pointe importante de bruit.

Cela conduit à poser la question de l'optimisation du fonctionnement du récepteur afin que celui-ci délivre au destinataire un message dont la qualité est la meilleure possible.

Nous développerons le concept de récepteur optimal dans le cas de la transmission en bande de base sur un canal BBAG, en s'appuyant sur les résultats de la théorie de la détection. La

structure des récepteurs est fondée sur l'utilisation du principe du maximum de vraisemblance, compte tenu des caractéristiques de l'émetteur utilisé et du canal de transmission.

Les résultats obtenus seront complétés par la définition de la transmission idéale qui procure au récepteur optimal l'avantage d'une structure linéaire simple et à la transmission l'avantage de la meilleure qualité possible.

La notion de transmission idéale sera étendue aux transmissions sur fréquence porteuse. Nous verrons qu'une des conditions à remplir, la connaissance de toutes les formes possibles émises, pose un problème délicat. Celui-ci ne peut être pratiquement résolu que si la modulation utilisée est linéaire et exige du récepteur une connaissance parfaite de la porteuse utilisée, c'est-à-dire de sa fréquence et de sa phase. On parlera de *réception cohérente*, et souvent même de *démodulation cohérente*.

Le récepteur cohérent doit comporter un dispositif de *synchronisation*, ou *circuit de récupération de porteuse*, qui lui permet d'acquérir la fréquence et la phase de la porteuse émise.

La conception de ce circuit de récupération pose des problèmes supplémentaires. Le désir de les éviter conduit à l'étude de structures de récepteur sans dispositif de synchronisation de porteuse. On verra alors apparaître la *démodulation non cohérente*, tels que la démodulation différentielle ou la démodulation d'enveloppe. La figure 1.5 donne les schémas de principe des récepteurs pour transmission en bande de base et pour transmission sur fréquence porteuse. Les structures représentées sont les plus simples et ont pour but de bien mettre en évidence les diverses opérations effectuées dans chaque cas par le récepteur. Il existe des structures plus complexes, notamment des structures bouclées.

La connaissance de la structure des émetteurs et récepteurs utilisés pour divers types de modulation et de démodulation permet d'effectuer des comparaisons basées sur les deux paramètres suivants :

- Le taux d'erreurs sur les bits (BER) obtenu en fonction d'un rapport de l'énergie moyenne par bit transmis E_b à la densité monolatérale de bruit N_0 .
- La largeur de bande nécessaire à la transmission du signal numérique. Ces comparaisons peuvent se faire dans le cas de la transmission idéale. Toutefois, cette transmission idéale est un objectif d'autant plus difficile à approcher que le modem (modulateur-démodulateur) est de réalisation complexe et que le canal de transmission réel n'est pas le canal BBAG. Il s'avère donc important de pouvoir faire des comparaisons, en particulier de connaître le BER dans des conditions réelles.

1.3.5 Note sur le codage de source et le codage de canal (cas numérique).

Nous avons vu apparaître implicitement les notions de codage de source et de codage en ligne. Ces deux fonctions sont bien distinctes, et il est utile de les préciser.

Codage de source.

Le codage de source, que nous situons **à l'intérieur de la source**, a pour but de rendre celle-ci aussi idéale que possible au sens de la théorie de l'information. Elle servira à réduire la redondance de la source et à transmettre un débit de symboles minimum, le plus proche possible du débit d'information réel de la source. Le code Morse est, par exemple, une tentative dans ce sens, qui optimise le temps de transmission. En transmissions d'images, on cherche à exploiter la dépendance statistique d'un point de l'image avec ses voisins ou avec le même point de l'image précédente.

FIG. 1.5 – *Récepteur.*

Codage de canal.

Le codage de canal, dont fait partie le codage en ligne et la modulation, sert à adapter le signal électrique analogique au canal de transmission. Ce codage nécessite l'introduction d'une redondance voulue, par exemple pour permettre la détection, voire la correction d'erreurs de transmission ou éliminer la composante continue d'un signal. Ce codage a lieu **dans l'émetteur**.

1.4 Les organisations internationales de Télécommunications.

L'objectif des télécommunications étant de communiquer loin (télé = $\tau\eta\lambda\eta$ = loin en grec), il est clair que c'est un domaine où la standardisation doit être la plus poussée possible. Pour atteindre ce but, des pouvoirs de normalisations nationaux (jusqu'à récemment facilités par les situations de monopoles) et internationaux ont été mis en place. Il s'agit principalement de

- L'U.I.T. (Union Internationale des Télécommunications : <http://www.itu.ch/>). Le CCITT (Comité Consultatif International pour la Téléphonie et la Télégraphie) et le CCIR (Comité Consultatif International pour les Radiocommunications) sont des organes permanents de l'IUT qui sont en charge d'établir des normes internationales).
- L'I.S.O. (International Standards Organisation : <http://www.iso.ch>)

Chapter 2

Rappels concernant Fourier et la convolution.

2.1 Définitions.

On notera les paires de transformées de Fourier $f(t) \Leftrightarrow F(\omega)$ où

$$F(\omega) = \int_{-\infty}^{\infty} f(t)e^{-j\omega t} dt \quad (2.1)$$

et

$$f(t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} F(\omega)e^{j\omega t} d\omega = \int_{-\infty}^{\infty} F(f)e^{j2\pi f} df \quad (2.2)$$

2.2 Théorème de convolution.

$$f_1(t) \otimes f_2(t) \Leftrightarrow F_1(\omega).F_2(\omega) \quad (2.3)$$

$$\int_{-\infty}^{\infty} e^{-j\omega t} \left[\int_{-\infty}^{\infty} f_1(\tau)f_2(t-\tau)d\tau \right] dt \quad (2.4)$$

En posant $t = \tau + x$ et en modifiant l'ordre d'intégration :

$$\int_{-\infty}^{\infty} f_1(\tau) \int_{-\infty}^{\infty} e^{-j\omega(\tau+x)} f_2(x) dx d\tau = \int_{-\infty}^{\infty} f_1(\tau)e^{-j\omega\tau} d\tau \int_{-\infty}^{\infty} f_2(x)e^{-j\omega x} dx \quad (2.5)$$

QED.

Exercice 2.1 *Démontrez*

$$f_1(t).f_2(t) \Leftrightarrow F_1(\omega) \otimes F_2(\omega)$$

2.2.1 Filtrage et convolution.

Le point principal à retenir est, dans le domaine temporel, l'équivalence entre filtrage et convolution. En clair, si on passe un signal $s(t)$ dans un filtre de réponse impulsionnelle $h(t)$, le signal résultant vaut

$$r(t) = \int_{-\infty}^{\infty} s(\tau)h(t-\tau)d\tau$$

. Ce qui peut être expliqué comme suit : on prend le signal temporel $s(t)$ et la réponse impulsionnelle inversée dans le temps et décalée de t : $h(t - \tau)$, la valeur du signal en sortie est la surface en dessous du produit de ces deux fonctions. Soit graphiquement :

- Un signal carré de largeur 20 ici.
- Un filtre de réponse impulsionnelle $e^{-a\tau}$ si $\tau > 0$
- La sortie du filtre.

Pour vous convaincre du graphique, faites glisser la réponse impulsionnelle inversée sur le signal carré, vous devriez “sentir” la courbe de sortie. A quel type de filtre cela vous fait-il penser, quels éléments électriques ?

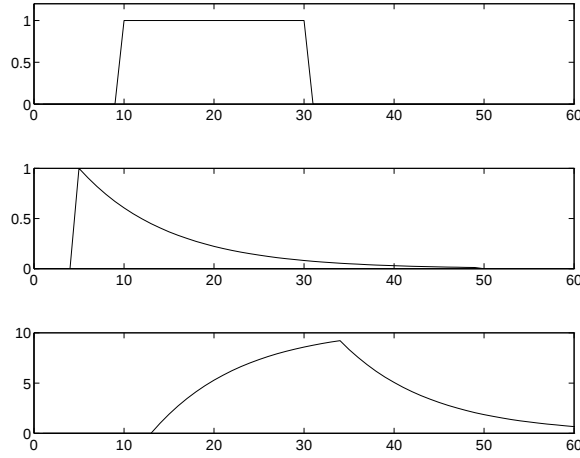


FIG. 2.1 – Convolution d’un signal carré avec un filtre de réponse exponentielle

L’équivalence entre convolution dans le domaine temporel et produit dans le domaine fréquentiel peut être illustré par le graphique suivant où la première partie représente la transformée de Fourier d’un signal, la deuxième, celle de la réponse impulsionnelle ci-dessus et la dernière le résultat du filtrage.

2.3 Théorème de Parseval.

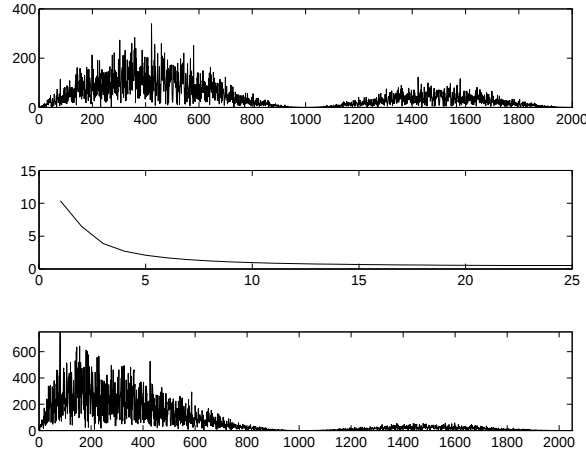
$$\int_{-\infty}^{\infty} y_1(t)y_2^*(t)dt = \frac{1}{2\pi} \int_{-\infty}^{\infty} Y_1(\omega)Y_2^*(\omega)d\omega \quad (2.6)$$

Par le théorème de convolution :

$$\int_{-\infty}^{\infty} f_1(\tau)f_2(t-\tau)d\tau \Leftrightarrow F_1(\omega)F_2(\omega) \quad (2.7)$$

en posant $y_1(\tau) = f_1(\tau)$ et $y_2^*(\tau) = f_2^*(t-\tau)$ et en remarquant :

$$TF[y_2^*(-\tau)] = \int_{-\infty}^{\infty} y_2^*(-\tau)e^{-j\omega\tau}d\tau = \int_{-\infty}^{\infty} y_2^*(\tau)e^{j\omega\tau}d\tau = Y_2^*(\omega) \quad (2.8)$$

FIG. 2.2 – *Convolution d'un signal carré avec un filtre de réponse exponentielle*

On obtient :

$$\int_{-\infty}^{\infty} y_1(\tau)y_2^*(\tau)d\tau \Leftrightarrow Y_1(\omega)Y_2^*(\omega) \text{ en } t = 0 \quad (2.9)$$

$$\int_{-\infty}^{\infty} y_1(t)y_2^*(t)dt = \frac{1}{2\pi} \int_{-\infty}^{\infty} Y_1(\omega)Y_2^*(\omega)d\omega \quad (2.10)$$

Par définition de la transformée de Fourier.

Si $y_2^*(t) = y_1(t)$, on a le résultat :

$$\int_{-\infty}^{\infty} |y_1(t)|^2 dt = \frac{1}{2\pi} \int_{-\infty}^{\infty} |Y_1(\omega)|^2 d\omega \quad (2.11)$$

L'énergie totale est égale à l'intégrale du module au carré de la transformée de Fourier du signal. $|Y_1(\omega)|^2$ peut donc être interprété comme une densité spectrale d'énergie pour les signaux certains. En effet, on peut la définir comme étant :

$$S(f)df = |Y_1(f)|^2 \quad (2.12)$$

2.4 Théorème de modulation.

Un des résultats les plus importants dans le cas des télécommunications est le théorème de modulation, qui est l'expression mathématique du mélange de fréquence, ou encore de la modulation d'amplitude à bande porteuse supprimée que nous verrons plus loin.

Soit un signal (ou fonction) $s(t)$ et une cissoïde $\cos(\omega_c t)$ où l'indice c signifie “carrier”, on forme le signal :

$$r(t) = s(t) \cdot \cos(\omega_c t) \quad (2.13)$$

et on se pose la question de connaître l'allure de la transformée de Fourier $R(\omega)$ de $r(t)$ en fonction de la transformée de Fourier $S(\omega)$ de $s(t)$.

Le calcul s'effectue aisément comme suit :

$$\begin{aligned}
\int_{-\infty}^{\infty} s(t) \cos(\omega_c) e^{-j\omega t} dt &= \frac{1}{2} \int_{-\infty}^{\infty} e^{j\omega_c t} e^{-j\omega t} dt + \frac{1}{2} \int_{-\infty}^{\infty} e^{-j\omega_c t} e^{-j\omega t} dt \\
&= \frac{1}{2} \int_{-\infty}^{\infty} e^{-j(\omega - \omega_c)t} dt + \frac{1}{2} \int_{-\infty}^{\infty} e^{-j(\omega + \omega_c)t} dt \\
&= \frac{1}{2} (S(\omega - \omega_c) + S(\omega + \omega_c))
\end{aligned} \tag{2.14}$$

La figure 2.3 illustre ce théorème. On notera que le spectre est symétrique, le centre de symétrie étant la fréquence 0. Il est clair que la multiplication par une cossoïde à une pulsation ω_c translate le spectre original autour de ω_c . **Notation** : il est d'un usage commun de désigner les “basses fréquences” en majuscules (Ω) et les fréquences transposées par ω .

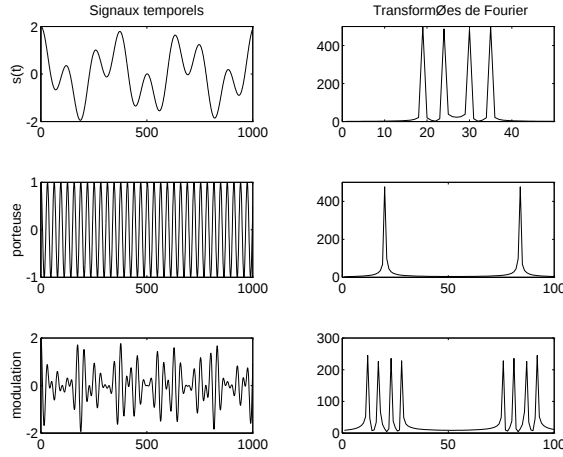


FIG. 2.3 – Illustration du théorème de modulation

Application 2.1 Mélange de fréquences.

Le mélange de fréquences consiste principalement en la translation d'un signal sur l'axe des fréquences. Si un signal en bande de base s'étend de 0 à 10 KHz, on désire le tradler à une fréquence centrale de 100 Mhz, il occupera alors la bande de 95 à 105 Mhz. Pour ce faire, la méthode la plus directe consiste à mélanger le signal en bande de base avec un signal de porteuse. Un calcul simple, et le théorème de modulation, permettent aisément de s'en convaincre, à un filtrage près.

2.5 Transformée des fonctions périodiques.

Soit $f(t)$ telle que $f(t) = f(t + T)$, montrer que $F(\omega) = 0 \forall \omega \neq \frac{2\pi n}{T}$ où n est entier.

$$\begin{aligned}
\int_{-\infty}^{\infty} f(t) e^{-j\omega t} dt &= e^{j\omega T} \int_{-\infty}^{\infty} f(t + T) e^{-j\omega(t+T)} dt \\
&\Leftrightarrow \\
F(\omega) &= e^{j\omega T} F(\omega)
\end{aligned} \tag{2.15}$$

vrai si $e^{j\omega T} = 1 \Rightarrow \omega = \frac{2\pi n}{T}$

Toute fonction périodique a une transformée de Fourier nulle en dehors des harmoniques de la fréquence fondamentale. On se ramène à la série de Fourier.

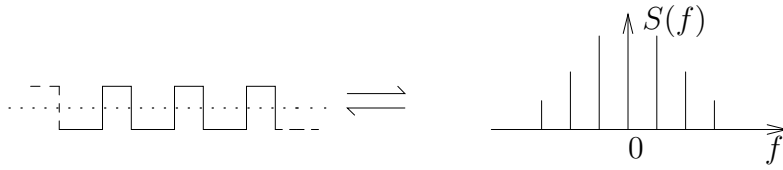


FIG. 2.4 – Transformée de Fourier d'un signal périodique.

2.6 Transformée d'une série d'impulsions.

Soit

$$\Delta(t) = \sum_{n=-\infty}^{+\infty} \delta(t - nT)$$

prouver

$$\Delta(\omega) = \omega_o \sum_{n=-\infty}^{+\infty} \delta(\omega - n\omega_o) \text{ où } \omega_o = \frac{2\pi}{T}$$

Solution. Définissons

$$\Delta_N(t) = \delta(t - nT)$$

$$\Rightarrow \Delta_N(\omega) = \sum_{n=-N}^N e^{jn\omega T} = \frac{\sin(N + \frac{1}{2}\omega T)}{\sin \frac{\omega T}{2}}$$

et

$$\lim_{N \rightarrow \infty} \frac{\sin(N + \frac{1}{2}\omega T)}{\sin \frac{\omega T}{2}} = \omega_o \sum_{n=-\infty}^{+\infty} \delta(\omega - n\omega_o)$$

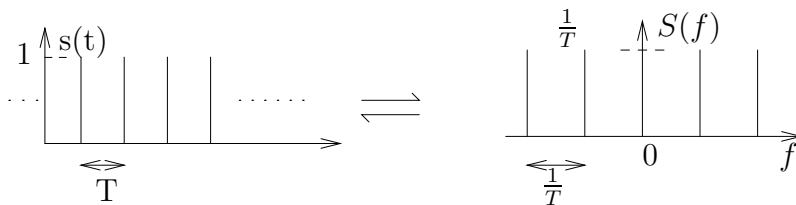


FIG. 2.5 – Transformée de Fourier d'une suite d'impulsions.

2.7 Transformée d'un signal périodisé.

Soit

$$f_1(t) = \Delta(t) \otimes \text{Rect}_a(t) \cos \Omega t$$

on demande la transformée de Fourier $F_1(\omega)$ de $f_1(t)$.

Solution.

$$\text{TF}[\text{Rect}_a(t)] =$$

$$\int_{-a}^a e^{-j\omega t} dt = \frac{\sin a\omega}{\omega/2}$$

$$\text{TF}[\text{Rect}_a(t) \cdot \cos \Omega t] = (\text{théorème de modulation})$$

$$\frac{1}{2} \left(\frac{\sin a(\omega - \Omega)}{(\omega - \Omega)/2} + \frac{\sin a(\omega + \Omega)}{(\omega + \Omega)/2} \right)$$

$$\text{TF}[\text{Rect}_a(t) \cdot \cos \Omega t \otimes \Delta(t)] =$$

$$\Delta_{\text{rect} \cdot \omega_o} \sum_{n=-\infty}^{+\infty} \delta(\omega - n\omega_o)$$

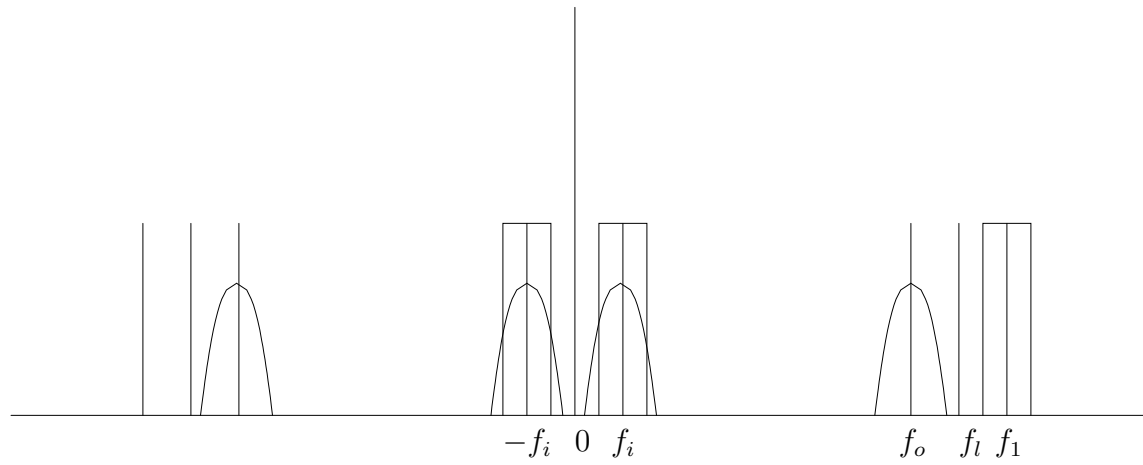
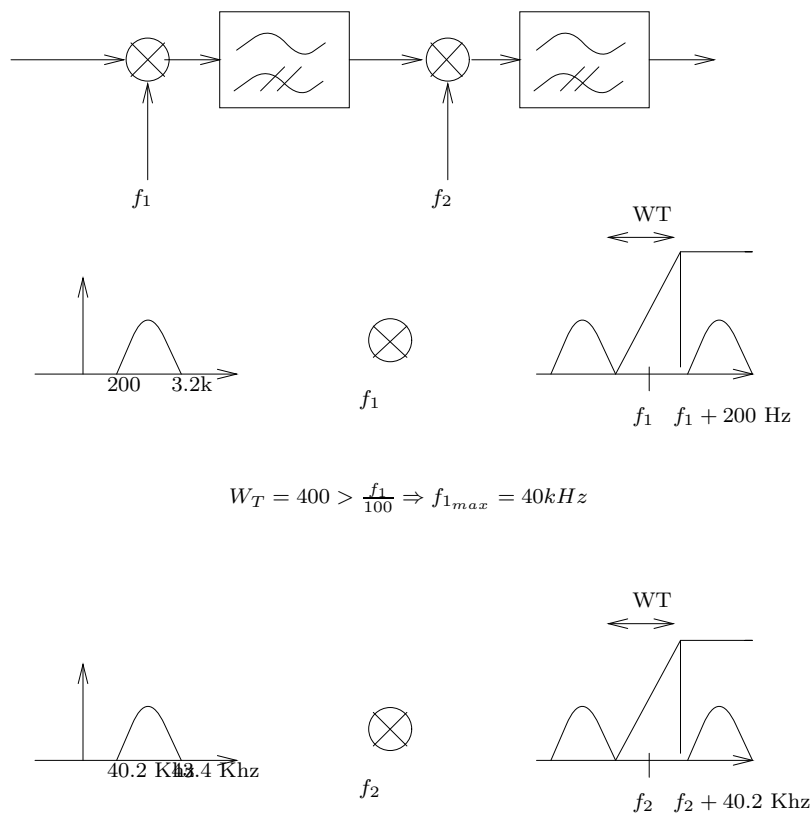
On obtient donc qu'un signal *périodisé* a comme spectre une série de Diracs espacées de l'inverse de la période de *périodisation* et de *valeurs* égales à la valeur du spectre original en cet endroit, à un facteur $2\pi/T$ près.

2.8 Problème de la fréquence image.

Plaçons nous dans le cas où l'on veut ramener un signal HF sur une fréquence IF donnée. Pour fixer les choses, on notera f_o la fréquence HF, f_i la fréquence intermédiaire et f_l la fréquence de l'oscillateur local. Par multiplication de celui-ci avec le signal HF, on désire obtenir un signal à la fréquence IF. On observe aisément qu'une fréquence $f_l = f_o + f_i$ permet de réaliser cette opération. En effet, les fréquences résultant de la multiplication des deux signaux sont $f_l \pm f_o$. Un problème subsiste cependant pour la *fréquence image* $f_1 = f_l + f_i$. En effet, le mélange de f_1 avec f_l produira une composante à $f_1 - f_l = f_i$. Donc, tout signal présent à la fréquence image sera placé à la fréquence intermédiaire et jouera le rôle de signal parasite. Il est donc impératif de filtrer ces signaux indésirables avant d'effectuer le mélange.

Application 2.2 Transposition de fréquence.

On désire transposer un signal BF s'étendant de 200 Hz à 3200 Hz. Cette opération se fera par deux opérations de mélange, chacune suivie d'un filtre passe-haut. Pour obtenir des facteurs de qualité acceptables (pas trop élevé), la largeur de décroissance des filtres W_T est limitée à $W_T \geq \frac{f_c}{100}$ où f_c est la fréquence de coupure des filtres. On demande de déterminer les fréquences maximales admissibles f_1 et f_2 des oscillateurs locaux.

FIG. 2.6 – *Problème de la fréquence image*FIG. 2.7 – *Solution*

2.9 Le récepteur superhétérodyne

Dans la plupart des cas, et en particulier dans le domaine très répandu des transmissions radio, il est nécessaire de déplacer le signal sur une portion de spectre adéquate. D'autre part, le canal radio introduit une atténuation très importante qui implique, à la réception, une amplification gigantesque. Un signal de réception typique a une amplitude de l'ordre du μV alors que les organes de production du signal (acoustique par exemple) demandent une amplitude de l'ordre du Volt. Il est clair qu'une amplification directe par un facteur 10^6 est impossible et engendrerait une oscillation de l'ampli. Il convient donc de diviser cette amplification en plusieurs étages. D'autre part, le fait d'effectuer des changements de fréquence entre ces amplifications sera également un facteur empêchant une éventuelle oscillation du système global. Ces considérations nous amènent naturellement au récepteur superhétérodyne.

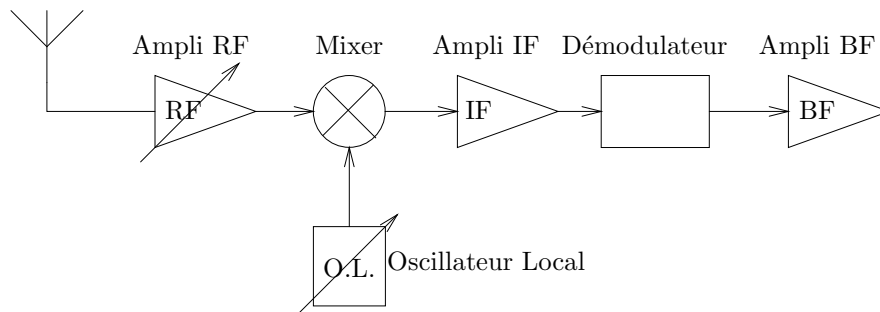


FIG. 2.8 – Récepteur superhétérodyne.

Les éléments du récepteur superhétérodyne.

1. **L'ampli RF** A la réception, une première amplification du signal RF est effectuée. Il s'agit naturellement d'une amplification sélective, c'est-à-dire d'une amplification avec filtrage, celui-ci a comme rôle, d'abord, de pallier aux inconvénients de la fréquence image (idéalement, on introduira un pôle à la fréquence image), également de diminuer la puissance de bruit captée par le récepteur et enfin d'effectuer un premier filtrage qui facilitera le filtrage IF. Il est clair également que la bande passante de ce filtre, dans le cas de la radiodiffusion par exemple, doit pouvoir être adaptée pour pouvoir capter la "station" que l'on veut.
2. **Le mélangeur et l'oscillateur local.** La démodulation pour la modulation d'amplitude (ou de fréquence) étant difficile à réaliser en haute fréquence, on ramène le signal à une *fréquence intermédiaire plus faible* pour réaliser la démodulation. Cette opération a un avantage supplémentaire : on ramène toutes les bandes HF différentes sur la même bande IF, ce qui permet avantages principaux :
 - Il est ainsi possible d'optimiser les circuits de modulation pour une seule bande de fréquence et donc d'optimiser ceux-ci.
 - Cela permet de n'avoir que deux réglages à effectuer. De plus, ceux-ci peuvent être couplés pour obtenir une concordance parfaite entre l'O.L. et la R.F.

- Cela permet d’avoir des filtres plus facilement réalisables. En effet, dans le cas de la radio F.M., les largeurs de bandes utilisées sont de 200 kHz pour une fréquence porteuse d’environ 100 MHz, ce qui impliquerait un filtre de réception ayant un facteur de qualité ($Q = \frac{\text{fréquence centrale du filtre}}{\text{largeur de bande du filtre}}$) de 500, ce qui est quasi impossible. En se ramenant à une fréquence centrale (IF) de 10.7 MHz, on a un facteur de qualité de l’ordre de 50, ce qui est élevé, mais acceptable, d’autant plus que ce filtre est à fréquence centrale fixe, et non variable comme en RF. Pour le filtre RF, il s’agit d’éviter la fréquence image. À titre d’exercice, déterminez le facteur de qualité qu’il faut dans ce cas.

L’ampli IF introduit une amplification et assure le filtrage IF, comme décrit ci-dessus.

Le démodulateur assure la transformation du signal électrique en un signal généralement proportionnel au signal de départ.

L’ampli BF ... no comment !

Chapter 3

Variables aléatoires.

3.1 Rappels de théorie des probabilités.

3.1.1 Introduction

Un signal qui peut être exprimé comme une fonction explicite du temps est un signal déterministe. Les télécommunications ayant comme objectif le transport d'informations, les signaux seront par nature non déterministes, c'est-à-dire aléatoires. De plus, soit par facilité de calcul, soit par adéquation aux signaux utilisés et/ou aux phénomènes rencontrés (idéalement par facilité **et** adéquation), nous utiliserons certaines hypothèses concernant ces signaux. Ce chapitre n'aura d'autre prétention que de faire un bref rappel de la théorie des probabilités et des processus stochastiques pour exprimer correctement les hypothèses de travail utilisées.

3.1.2 Théorie des probabilités.

3.1.3 Événements, expériences, axiomes de probabilité.

Un expérience déterminée comporte un ensemble S de tous ses résultats possibles E , appelés événements associés à cette expérience. A chaque événement associé à l'expérience considérée, on fait correspondre un nombre réel P_E appelé «probabilité de l'événement E » défini par les axiomes suivants :

1. $P_E \geq 0$
2. $P_I = 1$ un événement certain a une probabilité égale à 1
3. Si E_1 et E_2 sont mutuellement exclusifs, i.e. si $E_1 \cap E_2 = \emptyset$
 $P_{E_1 \cup E_2} = P_{E_1} + P_{E_2}$ (axiome de la somme)
4. $P_{E_1 \cap E_2} = P_{E_1} \cdot P_{E_2|E_1}$ (axiome du produit) où $P_{E_2|E_1}$ est la probabilité que E_2 survienne dans l'hypothèse où E_1 arrive.

On peut déduire de ces axiomes que P_E est toujours compris entre 0 et 1 et que la probabilité de l'événement impossible est nulle (la réciproque n'étant pas nécessairement vraie).

3.1.4 Indépendance statistique.

Deux événements E_1 et E_2 sont statistiquement indépendants si

$$P_{E_1 \cap E_2} = P_{E_1} \cdot P_{E_2} \quad (3.1)$$

$E_1 \dots E_n$ sont n événements statistiquement indépendants si

$$P_{E_i \cap E_j} = P_{E_i} \cdot P_{E_j} \quad (3.2)$$

$$P_{E_i \cap E_j \cap \dots \cap E_k} = P_{E_i} \cdot P_{E_j} \cdot \dots \cdot P_{E_k} \quad (3.3)$$

3.1.5 Loïs de composition

1. $P_{\bar{E}} = 1 - P_E$ où $\bar{E}$ est le complément de E
2. $P_{E_1 \cup E_2} = P_{E_1} + P_{E_2} - P_{E_1 \cap E_2}$
3. $P_{E_1 \cap E_2 \cap \dots \cap E_n} = P_{E_1} \cdot P_{E_2|E_1} \cdot P_{E_3|E_2 \cap E_1} \cdot \dots \cdot P_{E_n|E_{n-1} \cap \dots \cap E_2 \cap E_1}$
4. Soient $E_1, \dots, E_n$ n événements indépendants.
 $P_{E_1 \cup E_2 \cup \dots \cup E_n} = 1 - P_{\bar{E}_1} \cdot P_{\bar{E}_2} \cdot \dots \cdot P_{\bar{E}_n}$
5. Soient $E_1, \dots, E_n$ n événements indépendants.
 $P_{E_1 \cup E_2 \cup \dots \cup E_n} = 1 - P_{E_1} \cdot P_{E_2} \cdot \dots \cdot P_{E_n}$

3.1.6 Probabilités a posteriori

Soient $H_1, \dots, H_n$, un ensemble d'événements mutuellement exclusifs ($H_i \cap H_j = \emptyset$) tels que $H_1 \cup H_2 \cup \dots \cup H_n = I$ et non indépendants de E . Les H_i sont appelés hypothèses et peuvent être des causes de l'événement E .

$$P_E = \sum_{i=1}^n P_{H_i} \cdot P_{E|H_i} \quad (3.4)$$

Et nous avons la formule de Bayes :

$$P_{H_i|E} = \frac{P_{E \cap H_i}}{P_E} = \frac{P_{H_i} \cdot P_{E|H_i}}{P_E} \quad (3.5)$$

où

- $P_{H_i|E}$ est la probabilité a posteriori ;
- P_{H_i} est la probabilité a priori ;
- $P_{E|H_i}$ est la probabilité conditionnelle.

3.2 Variables aléatoires.

La classe S des événements associés à une expérience peut toujours être décrite à l'aide d'un ensemble d'événements mutuellement exclusifs appelés *événements élémentaires* de cette expérience.

Tout événement E consiste en la réunion d'un certain nombre d'événements élémentaires et sa probabilité est la somme des probabilités de ceux-ci.

Une **variable aléatoire** est définie par correspondance biunivoque avec un ensemble d'événements élémentaires et est caractérisée par la distribution de probabilité de celui-ci. Elle peut être à une ou plusieurs dimensions, discrète ou continue.

3.2.1 Fonction de répartition.

On considère une variable aléatoire X dont le domaine de définition sera pris systématiquement $[-\infty, \infty]$, bien que l'arrivée dans certains sous-domaines puisse être impossible. Cette variable aléatoire est entièrement définie par sa *fonction de répartition* (en anglais: *c.d.f.*: *cumulative distribution function*)

$$F(x) = P\{X \leq x\} \quad (3.6)$$

Cette fonction possède les propriétés suivantes :

1. $F(-\infty) = 0$; $\lim_{x \rightarrow \infty} F(x) = 1$
2. $F(b) - F(a) = P\{a < X \leq b\}$
3. $F(x)$ est une fonction monotone non décroissante.

La variable aléatoire X est dite *discrète* si la fonction $F(x)$ est une fonction en escalier, c'est-à-dire de la forme :

$$F(x) = \sum_i P_i u(x - x_i) \quad ; \quad p_i > 0 \quad ; \quad \sum_i P_i = 1 \quad (3.7)$$

où $u()$ est la fonction échelon. Une telle variable ne peut prendre que les valeurs x_i et ce, avec les probabilités p_i .

La variable X est dite *continue* si la fonction de répartition $F(x)$ est continue.

Dans le cas général, la fonction $F(x)$ peut contenir des discontinuités par saut brusque positif.

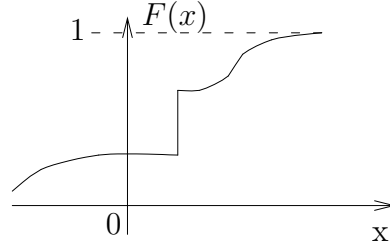
3.2.2 Densité de probabilité.

Pour une fonction de répartition continue, on définit la densité de probabilité de la variable aléatoire (*p.d.f.*: *probability density function*)

$$p(x) = \frac{dF}{dx} \quad (3.8)$$

On déduit aisément la signification de $p(x)$ de la propriété :

$$P\{a < X \leq b\} = F(b) - F(a) = \int_a^b p(x) dx \quad (3.9)$$

FIG. 3.1 – *Fonction de Répartition*

d'où $\int_{-\infty}^{\infty} p(x)dx = 1$ et, en particulier,

$$P\{x < X \leq x + dx\} = p(x)dx \quad (3.10)$$

ce qui illustre bien la dénomination de densité : plus $p(x)$ est grande, plus la probabilité que la variable tombe au voisinage de x est grande.

A condition d'admettre l'introduction de fonctions impulsions de Dirac, la notion de densité de probabilité peut être étendue aux variables aléatoires non continues, par exemple, pour une variable discrète :

$$p(x) = \sum_i P_i \delta(x - x_i) \quad (3.11)$$

3.2.3 Moments d'une variable aléatoire.

L'opérateur *espérance mathématique* $E\{f(X)\}$ fait correspondre à une fonction donnée f de la variable aléatoire X un nombre, et cela par la définition :

$$E\{f(X)\} = \int_{-\infty}^{\infty} f(x)p(x)dx \quad (3.12)$$

C'est une moyenne pondérée de la fonction f , la fonction de poids étant la densité de probabilité.

Dans le cas d'une variable aléatoire discrète, on a :

$$E\{f(x)\} = \sum_i P_i f(x_i) \quad (3.13)$$

Les *moments de la variable aléatoire* X sont les espérances mathématiques des puissances X^n

$$m_n = E\{X^n\} = \int_{-\infty}^{\infty} x^n p(x)dx \quad x \text{ continu} \quad (3.14)$$

$$m_n = E\{X^n\} = \sum_i x_i^n P_i \quad x \text{ discret} \quad (3.15)$$

En particulier, on distingue le moment du premier ordre m_1 qui est l'espérance mathématique de la variable elle-même :

$$m_1 = E\{X\} = \int_{-\infty}^{\infty} xp(x)dx \quad x \text{ continu} \quad (3.16)$$

$$m_1 = E\{X\} = \sum_i x_i P_i \quad x \text{ discret} \quad (3.17)$$

que l'on appelle *moyenne* ou *valeur la plus probable*. La moyenne donne la valeur de X autour de laquelle les résultats d'essais devraient se disperser.¹ Une variable aléatoire est dite *centrée* si sa moyenne est nulle. On travaille souvent avec les variables aléatoires centrées $(X - m_1)$. Les moments d'ordre supérieur à 1 de la nouvelle variable, notés $\mu_n = E\{(x - m_1)^n\}$ sont souvent plus utiles que les moments m_n .

$$\text{On a } \mu_n = \sum_{k=0}^n (-1)^k C_n^k m_{n-k} m_1^k.$$

En particulier le moment centré d'ordre deux est la *variance* :

$$\mu_2 = \sigma_x^2 = E\{(x - m_1)^2\} = m_2 - m_1^2$$

La racine carrée de la variance σ_x est appelée *dispersion* ou *écart-type*. Elle donne une mesure de la dispersion vraisemblable de résultats d'essais autour de la moyenne, ainsi que le montre l'inégalité de Bienaymé-Tchebychef :

$$p\{|X - m_1| > \alpha\} < \left(\frac{\sigma_x}{\alpha}\right)^2 \quad (3.18)$$

Il peut être utile de considérer la variable centrée et réduite $\frac{X-m_1}{\sigma}$ dont la variance est l'unité.

3.2.4 Variables réelles à plusieurs dimensions.

On peut étendre les notions vues précédemment à une variable aléatoire à n dimensions $X = (X_1, X_2, \dots, X_n)$. On définit la fonction de répartition (fonction scalaire) :

$$F(x_1, x_2, \dots, x_n) = P\{X_1 \leq x_1, X_2 \leq x_2, \dots, X_n \leq x_n\} \quad (3.19)$$

La densité de probabilité est alors définie par :

$$p(x_1, x_2, \dots, x_n) = \frac{\partial^n F(x_1, x_2, \dots, x_n)}{\partial x_1 \partial x_2 \dots \partial x_n} \quad (3.20)$$

L'extension des notions précitées est immédiate.

On exprime les lois marginales par

$$F_{X_1}(x_1) = F_{X_1 \dots X_n}(x_1, \infty, \dots, \infty) \quad (3.21)$$

et par conséquent :

$$p_{X_1}(x_1) = \frac{\partial F_{X_1}(x_1)}{\partial x_1} = \int_{-\infty}^{\infty} \dots \int_{-\infty}^{\infty} p_{X_1 \dots X_n}(x_1, x_2, \dots, x_n) dx_2 \dots dx_n \quad (3.22)$$

Notons également l'introduction possible de lois conditionnelles du type :

1. Il ne faut pas la confondre avec la *médiane* qui est la valeur $x_{1/2}$ telle que $F(x_{1/2}) = \frac{1}{2}$ et à laquelle elle n'est pas toujours égale.

$$F_{X_1, \dots, X_k | X_{k+1}, \dots, X_n}(x_1, \dots, x_k | x_{k+1}, \dots, x_n) \\ = P\{X_1 \leq x_1, \dots, X_k \leq x_k | X_{k+1} = x_{k+1}, \dots, X_n = x_n\} \quad (3.23)$$

$$p_{X_1, \dots, X_k | X_{k+1}, \dots, X_n}(x_1, \dots, x_k | x_{k+1}, \dots, x_n) \\ = \frac{\partial^k F_{X_1, \dots, X_k | X_{k+1}, \dots, X_n}(x_1, x_2, \dots, x_n)}{\partial x_1 \partial x_2 \dots \partial x_n} \quad (3.24)$$

Et on déduit la généralisation de la formule de Bayes :

$$p_{X_1, \dots, X_k | X_{k+1}, \dots, X_n} = \frac{p_{X_1, \dots, X_n}}{p_{X_1, \dots, X_k}} \quad (3.25)$$

Les moments et moyennes sont alors définis grâce à l'opérateur espérance mathématique :

$$E\{f(X_1, \dots, X_n)\} = \int_{-\infty}^{\infty} \dots \int_{-\infty}^{\infty} f(x_1, \dots, x_n) p(x_1, \dots, x_n) dx_1 \dots dx_n \quad (3.26)$$

Les moments des deux premiers ordres sont alors :

$$m_{10} = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} x_1 p(x_1, x_2) dx_1 dx_2 = \int_{-\infty}^{\infty} x_1 p_{X_1}(x_1) dx_1 \quad (3.27)$$

$$m_{01} = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} x_2 p(x_1, x_2) dx_1 dx_2 = \int_{-\infty}^{\infty} x_2 p_{X_2}(x_2) dx_2 \quad (3.28)$$

$$m_{20} = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} x_1^2 p(x_1, x_2) dx_1 dx_2 = \int_{-\infty}^{\infty} x_1^2 p_{X_1}(x_1) dx_1 \quad (3.29)$$

$$m_{02} = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} x_2^2 p(x_1, x_2) dx_1 dx_2 = \int_{-\infty}^{\infty} x_2^2 p_{X_2}(x_2) dx_2 \quad (3.30)$$

$$m_{11} = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} x_1 x_2 p(x_1, x_2) dx_1 dx_2 \quad (3.31)$$

On considère souvent les moments centrés d'ordre 2, dont les variances :

$$\sigma_1^2 = E\{(X_1 - m_{10})^2\} = m_{20} - m_{10}^2 \quad (3.32)$$

$$\sigma_2^2 = E\{(X_2 - m_{01})^2\} = m_{02} - m_{01}^2 \quad (3.33)$$

la *corrélation mutuelle* (*cross-correlation*)

$$r_{12} = E\{(X_1)(X_2)\} = m_{11} \quad (3.34)$$

l'*autocorrélation*

$$r_{11} = E\{(X_1)(X_1)\} = m_{20} = \sigma_1^2 + m_{10}^2 \quad (3.35)$$

et la *covariance mutuelle*

$$\mu_{12} = E\{(X_1 - m_{10})(X_2 - m_{01})\} = m_{11} - m_{10}.m_{01} \quad (3.36)$$

On peut aisément étendre cette théorie des moments aux variables à plus de deux dimensions. En particulier, on utilise couramment les moments du premier ordre ou moyennes, permettant de centrer les variables. On forme avec les moments centrés d'ordre deux une matrice de covariance ($n \times n$) qui est symétrique et dont les éléments de la diagonale principale sont les variances :

$$C = \begin{bmatrix} \sigma_1^2 & \mu_{12} & \dots & \mu_{1n} \\ \mu_{21} & \sigma_2^2 & \dots & \mu_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ \mu_{n1} & \mu_{n2} & \dots & \sigma_n^2 \end{bmatrix} \quad (3.37)$$

3.2.5 Fonction caractéristique

La fonction caractéristique de la variable aléatoire à n dimension X est définie par :

$$\psi_x(q_1, \dots, q_n) = E\{e^{j \sum_{k=1}^n q_k X_k}\} = \int_{-\infty}^{\infty} e^{jq_1 x_1} dx_1 \int_{-\infty}^{\infty} e^{jq_2 x_2} dx_2 \dots \int_{-\infty}^{\infty} e^{jq_n x_n} p(x_1, \dots, x_n) dx_n \quad (3.38)$$

On peut aussi définir les fonctions caractéristiques marginales et conditionnelles.

Si tous les moments sont finis, on dispose du développement en série :

$$\psi_x(q_1, \dots, q_n) = 1 + \frac{(j)^k}{k!} \sum_{i_1 + \dots + i_n = k} E\{X_1^{i_1} \dots X_n^{i_n}\} q_1^{i_1} \dots q_n^{i_n} \quad (3.39)$$

En particulier, pour $n = 2$:

$$\psi_{x_1, x_2}(q_1, q_2) = 1 + j(m_{10}q_1 + m_{01}q_2) + \frac{j^2}{2!}(m_{20}q_1^2 + 2m_{11}q_1q_2 + m_{02}q_2^2) + \dots \quad (3.40)$$

et l'on a, sous réserve de différentiabilité :

$$E\{X_1^i X_2^k\} = \frac{1}{(j)^{i+k}} \left\{ \frac{\partial^{i+k} \psi_{x_1 x_2}}{\partial q_1^i \partial q_2^k} \right\}_{q_1=q_2=0} \quad (3.41)$$

3.3 Extension aux variables complexes

3.3.1 Fonction de répartition et densité de probabilité

Une variable aléatoire complexe X étant l'ensemble des nombres réels ordonnés (A, B) , sa fonction de répartition et sa densité de probabilité sont définies comme étant celles de la variable bidimensionnelle (A, B) . De même, pour une variable complexe à plusieurs dimensions, on devra considérer la variable réelle de dimension double constituée des parties réelles et imaginaires.

3.3.2 Moments

Espérance mathématique de la variable ou moyenne.

Soit $x = A + jB$ une variable aléatoire complexe à une ou à plusieurs dimensions. La moyenne de X est définie par

$$E\{X\} = m_x = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (a + jb)p_{AB}(a, b)da db \quad (3.42)$$

Comme on a, par exemple :

$$m_A = \int_{-\infty}^{\infty} a p_A(a)da \quad ; \quad \int_{-\infty}^{\infty} p(a, b)db = p(a) \quad (3.43)$$

on a aussi :

$$m_x = m_A + jm_B \quad (3.44)$$

Comme précédemment, la variable $(X - m_X)$ est dite *centrée*.

Moments du second ordre.

On considère une variable aléatoire complexe à une dimension $X = A + jB$ de moyenne m_x . Par définition, sa variance est :

$$\sigma_x^2 = E\{|X - m_x|^2\} \quad (3.45)$$

On voit que l'on a veillé à ce que la variance soit un nombre *réel et positif*.

Comme $|X - m_x|^2 = (A - m_A)^2 + (B - m_B)^2$, on a

$$\sigma_x^2 = \sigma_A^2 + \sigma_B^2 = E\{|X|^2\} - |m_x|^2 \quad (3.46)$$

Considérons maintenant une *variable complexe* à n dimensions $\mathbf{X} = \mathbf{A} + j\mathbf{B}$. Par définition, sa matrice ($n \times n$) de covariance est

$$\mathbf{C}_{\mathbf{X}} \equiv E\{(\mathbf{X} - m_{\mathbf{X}})(\mathbf{X} - m_{\mathbf{X}})^H\} \quad (3.47)$$

où le H signifie la transposée hermitienne (conjugué complexe du transposé). Comme $(\mathbf{X} - m_{\mathbf{X}})(\mathbf{X} - m_{\mathbf{X}})^H = [\mathbf{A} - m_{\mathbf{A}} + j(\mathbf{B} - m_{\mathbf{B}})][\mathbf{A}^T - m_{\mathbf{A}}^T - j(\mathbf{B}^T - m_{\mathbf{B}}^T)]$, on a

$$\mathbf{C}_{\mathbf{X}} = \mathbf{C}_{\mathbf{A}} + \mathbf{C}_{\mathbf{B}} + j(\mathbf{C}_{\mathbf{BA}} - \mathbf{C}_{\mathbf{AB}}) \quad (3.48)$$

De la propriété $\mathbf{C}_{\mathbf{BA}} = \mathbf{C}_{\mathbf{AB}}^T$ vue pour les variables réelles résulte la propriété

$$\mathbf{C}_{\mathbf{X}} = \mathbf{C}_{\mathbf{X}}^H \quad (3.49)$$

C'est-à-dire que la *matrice de covariance est hermitienne*. De même, elle est *définie non négative*.

On adapte d'une manière analogue la définition de la matrice de covariance mutuelle de deux variables à plusieurs dimensions

$$\begin{aligned} \mathbf{C}_{\mathbf{XY}} &= E\{(\mathbf{X} - m_{\mathbf{X}})(\mathbf{Y} - m_{\mathbf{Y}})^H\} & \mathbf{X}(n \times 1) \\ & & \mathbf{Y}(m \times 1) \\ & & \mathbf{C}_{\mathbf{XY}}(n \times m) \end{aligned} \quad (3.50)$$

La démonstration de la propriété :

$$\mathbf{C}_{\mathbf{YX}} = \mathbf{C}_{\mathbf{XY}}^H \quad (3.51)$$

est analogue à celle de l'hermiticité de $\mathbf{C}_{\mathbf{X}}$.

3.4 Quelques lois de probabilité importantes

3.4.1 Loi à deux valeurs

Densité de probabilité.

La variable réelle à une ou plusieurs dimensions X peut prendre les valeurs a et b avec des probabilités respectives p_a et p_b telles que $p_a + p_b = 1$:

$$p(x) = p_a \delta(x - a) + p_b \delta(x - b) \quad (3.52)$$

Rappelons qu'en coordonnées cartésiennes à n dimensions

$$\sigma(\mathbf{x} - \mathbf{a}) = \delta(x_1 - a_1) \delta(x_2 - a_2) \dots \delta(x_n - a_n) \quad (3.53)$$

Moments.

$$m_1 = E\{X\} = a.p_a + b.p_b \quad (3.54)$$

Plus généralement : $m_n = E\{X^n\} = a^n p_a + b^n p_b$; (pour une dimension)

$$\sigma_x^2 = p_a p_b |b - a|^2 ; \text{ (pour une dimension)} \quad (3.55)$$

$$C_{\mathbf{x}} = p_a p_b (\mathbf{b} - \mathbf{a})(\mathbf{b} - \mathbf{a})^H ; \text{ (pour plusieurs dimensions)} \quad (3.56)$$

Fonction caractéristique

$$\psi_x(q) = p_a e^{jaq} + p_b e^{jbq} ; \text{ (pour une dimension)} \quad (3.57)$$

3.4.2 Loi binomiale

Densité de probabilité.

On fait n essais indépendants relatifs à un événement E ayant une probabilité p . La variable aléatoire X "nombre d'arrivées de l'événement au cours des n essais" peut prendre les valeurs $0, 1, \dots, n$. La densité de probabilité de la variable discrète X est :

$$p(x) = \sum_{k=0}^n p_k \delta(x - k) \quad (3.58)$$

avec

$$p_k = C_n^k p^k (1 - p)^{n-k} \quad (3.59)$$

Cette loi est unimodale, le mode étant le plus grand entier inférieur ou égal à $[(n + 1)p]$.

Moments.

$$\begin{aligned}
m_1 = E\{X\} &= np & \mu_2 = \sigma_x^2 &= np(1-p) \\
m_2 = E\{X^2\} &= np(1-p) + n^2p^2 & \mu_3 &= np(1-p)(1-2p) \\
& & \mu_4 &= 3n^2p^2(1-p)^2 + np(1-p)(1-6p+6p^2)
\end{aligned} \tag{3.60}$$

Fonction caractéristique

$$\psi_x(q) = (1 - p + pe^{jq})^n \tag{3.61}$$

Propriétés.

1. Stabilité² : si X_1 et X_2 suivent des lois binomiales de paramètres respectifs (n_1, p) et (n_2, p) – le même p – alors $X_1 + X_2$ suit une loi binomiale de paramètres $(n_1 + n_2, p)$.
2. Si $n \rightarrow \infty$, X est asymptotiquement normale de moyenne np et de variance $\sigma^2 = np(1-p)$.
3. Si $n \rightarrow \infty$ et simultanément $p \rightarrow 0$ de telle manière que $np \rightarrow \lambda$, on obtient la loi de Poisson. On peut approcher la loi de Poisson dès que $p < 0,1$ et $np > 1$.

3.4.3 Loi de Poisson**Densité de Probabilité.**

La loi de Poisson :

$$p(x) = \sum_{k=0}^{\infty} p_k \delta(x - k) \quad \text{avec} \quad p_k = \frac{\lambda^k e^{-\lambda}}{k!} \tag{3.62}$$

est relative à une variable réelle à une dimension ne pouvant prendre que les valeurs entières positives (Fig. 3.2).

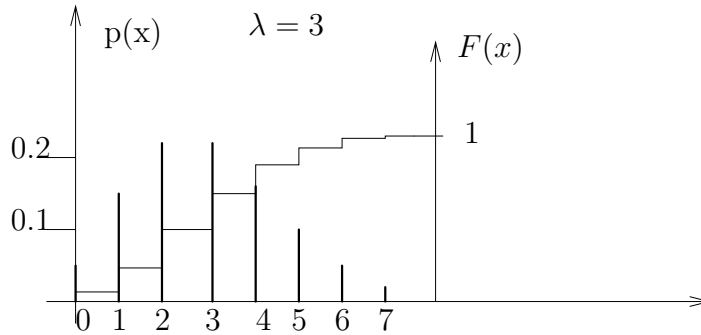


FIG. 3.2 – *Loi de Poisson*

2. Une loi de probabilité est dite stable si la somme de variables indépendantes suivant une loi de ce type obéit elle aussi à une loi de ce type.

Interprétation :

Exemple : un nombre $n \gg 1$ de téléviseurs de même marque et de même âge sont en service dans une ville. La probabilité que l'un quelconque de ces appareils tombe en panne demain est $p \ll 1$. La variable X = nombre d'appareils qui tomberont en panne demain est une loi de Poisson de paramètre $\lambda = np$.

b) Moments

$$m_1 = E\{X\} = \lambda$$

$$\sigma_x^2 = \lambda$$

Fonction caractéristique

$$\psi_x(q) = e^{-\lambda(e^{jq}-1)} \quad (3.63)$$

Propriété de stabilité.

La somme de deux variables aléatoires obéissant à des lois de Poisson de paramètres λ_1 et λ_2 suit une loi de Poisson de paramètre $\lambda_1 + \lambda_2$.

Exemple : le nombre total de téléviseurs de tous âges et marques tombant en panne demain dans telle ville suit une loi de Poisson.

Loi de Poisson à plusieurs dimensions.

$$p(x_1, \dots, x_n) = \sum e^{-(\lambda_1 + \dots + \lambda_n)} \frac{\lambda_1^{k_1} \dots \lambda_n^{k_n}}{k_1! \dots k_n!} \delta(x_1 - k_1) \dots \delta(x_n - k_n) \quad (3.64)$$

c'est-à-dire que les composantes $x_1, \dots, x_n$ sont indépendantes, mais suivent chacune une loi de Poisson.

3.4.4 Loi uniforme

Densité de probabilité.

X étant une variable réelle à une dimension

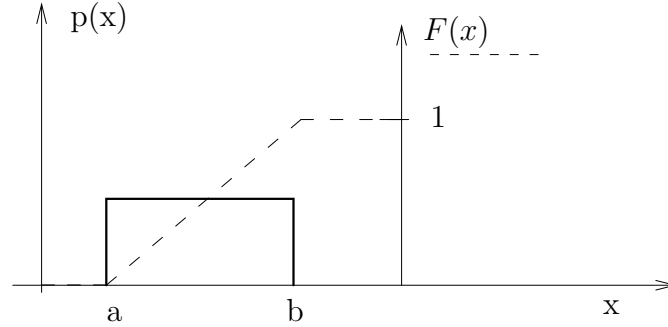
$$p(x) = \begin{cases} \frac{1}{b-a} & \text{si } a \leq x \leq b \\ 0 & \text{ailleurs} \end{cases} \quad (3.65)$$

c'est-à-dire que X arrive au hasard entre a et b . On dit que X est une *variable de chance*. La fonction de répartition est :

$$F(X) = \begin{cases} 0 & \text{si } x \leq a \\ \frac{x-a}{b-a} & \text{si } a \leq x \leq b \\ 1 & \text{si } x > b \end{cases} \quad (3.66)$$

Moments.

$$\begin{aligned}
m_1 = E\{X\} &= \frac{a+b}{2} & \sigma_x^2 &= \frac{(b-a)^2}{12} \\
m_n &= \frac{1}{n+1} \sum_{k=0}^n b^k a^{n-k}
\end{aligned} \tag{3.67}$$

FIG. 3.3 – *Loi uniforme***Fonction caractéristique**

$$\phi_x(q) = \frac{1}{j(b-a)q} (e^{jqb} - e^{jq a}) \tag{3.68}$$

3.4.5 Loi de Gauss ou loi normale**Variable à une dimension**

Densité de probabilité. La variable aléatoire scalaire réelle X est dite gaussienne ou normale si sa densité de probabilité est de la forme

$$p(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x-m)^2}{2\sigma^2}} \tag{3.69}$$

Cette loi est unimodale et symétrique autour de m . (Fig. 3.4)

En utilisant la fonction

$$\operatorname{erf}(z) = -\operatorname{erf}(-z) = \sqrt{\frac{2}{\pi}} \int_0^z e^{-\frac{\xi^2}{2}} d\xi; z \geq 0 \tag{3.70}$$

on peut écrire la fonction de répartition sous la forme :

$$F(x) = \frac{1}{2} \left[1 + \operatorname{erf} \frac{x-m}{\sigma\sqrt{2}} \right] \tag{3.71}$$

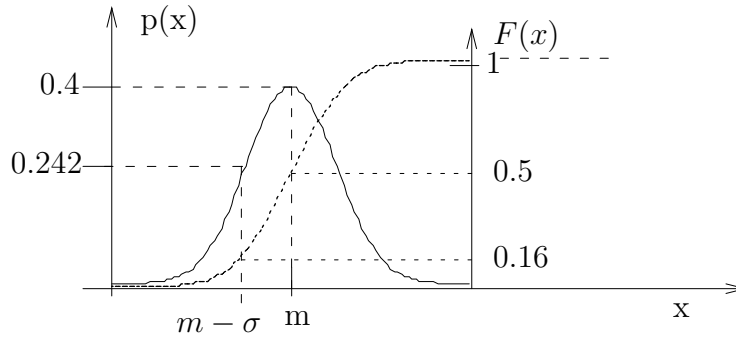


FIG. 3.4 – Loi Gaussienne

Moments.

$$\begin{aligned}
 m_1 &= E\{X\} = m \\
 m_2 &= E\{X^2\} = \sigma^2 + m^2; & \mu_2 &= \sigma^2 \\
 & & \mu_3 &= 0 \\
 & & \mu_4 &= 3\sigma^4 \\
 & & \dots & \\
 & & \mu_{2n-1} &= 0 \\
 & & \mu_{2n} &= (2n-1)!! \sigma^{2n}
 \end{aligned} \tag{3.72}$$

Fonction caractéristique

$$\psi_x(q) = e^{jmq} e^{-\frac{\sigma^2 q^2}{2}} \tag{3.73}$$

Propriétés

- Stabilité : si X_1 et X_2 sont des variables gaussiennes de paramètres (m_1, σ_1) et (m_2, σ_2) , respectivement, $X_1 + X_2$ est gaussienne de paramètres $(m_1 + m_2, \sqrt{\sigma_1^2 + \sigma_2^2})$.
- On a vu qu'en général, la connaissance des moments de tous les ordres est nécessaire pour définir une variable aléatoire. *Dans le cas de la loi de Gauss, les deux premiers moments suffisent.*
- Si X est normale, $Y = aX + b$ l'est aussi, a et b étant certains.

Variable à plusieurs dimensions.

Densité de probabilité. La variable aléatoire réelle vectorielle $\mathbf{X}(n \times 1)$ est dite normale si

$$p(x_1, \dots, x_n) = \sqrt{\frac{\text{Det } \mathbf{C}^{-1}}{(2\pi)^n}} e^{-\frac{1}{2}(\mathbf{x}-\mathbf{m})^H \mathbf{C}^{-1}(\mathbf{x}-\mathbf{m})} \tag{3.74}$$

où $\mathbf{m}$ (vecteur $n \times 1$) est l'espérance mathématique de $\mathbf{X}$

Propriétés.

1. Une variable vectorielle normale est entièrement définie par sa moyenne et sa matrice de covariance.
2. Les densités de probabilité conditionnelles et marginales sont toutes gaussiennes.

NB : L'inverse n'est pas toujours vrai. Par exemple si X_1 est normale ($m = 0, \sigma = 1$) et si

$$\begin{aligned} X_2 &= X_1 \text{ avec une probabilité } 1/2 \\ &-X_1 \text{ avec une probabilité } 1/2 \end{aligned} \quad (3.75)$$

X_1 et X_2 sont normales, mais la variable bidimensionnelle (X_1, X_2) ne l'est pas.

3. La condition nécessaire et suffisante pour que les composantes $X_1, \dots, X_n$ d'une distribution normale à n dimensions soient indépendantes est qu'elles ne soient pas corrélées (C diagonale). *L'indépendance statistique et l'indépendance en moyenne sont donc équivalentes pour la loi normale.*
4. Toute transformation linéaire sur des variables gaussiennes conserve le caractère normal. En particulier, il existe une transformation linéaire qui rend les nouvelles variables indépendantes.

Exemple : $n = 2$, variables centrées

$$\begin{aligned} m &= \begin{bmatrix} m_1 \\ m_2 \end{bmatrix} = \begin{bmatrix} 0 \\ 0 \end{bmatrix}; \quad \mathbf{C} = \begin{bmatrix} \sigma_1^2 & \mu \\ \mu & \sigma_2^2 \end{bmatrix}; \quad \mathbf{C}^{-1} = \frac{1}{\sigma_1^2 \sigma_2^2 - \mu^2} \begin{bmatrix} \sigma_1^2 & -\mu \\ -\mu & \sigma_2^2 \end{bmatrix}; \\ \text{Det } \mathbf{C}^{-1} &= (\sigma_1^2 \sigma_2^2 - \mu^2)^{-2}; \\ p(x_1, x_2) &= \frac{1}{2\pi \sqrt{\sigma_1^2 \sigma_2^2 - \mu^2}} e^{\frac{1}{2} \cdot \frac{\sigma_2^2 x_1^2 - 2\mu x_1 x_2 + \sigma_1^2 x_2^2}{\sigma_1^2 \sigma_2^2 - \mu^2}} \end{aligned} \quad (3.76)$$

3.4.6 Loi de Rayleigh

Densité de probabilité.

$$p(x) = \begin{cases} 0 & \text{si } x < 0 \\ \frac{x}{\sigma^2} e^{-\frac{x^2}{2\sigma^2}} & \text{si } x \geq 0 \end{cases} \quad (3.77)$$

Fonction de répartition :

$$F(x) = 1 - e^{-\frac{x^2}{2\sigma^2}} \text{ pour } x \geq 0 \quad (3.78)$$

Moments

$$\begin{aligned} m_1 &= \sigma \sqrt{\frac{\pi}{2}} \\ &\dots \quad \mu_2 = \sigma_x^2 = 2\sigma^2 \\ m_n &= (2\sigma^2)^{\frac{n}{2}} \Gamma(1 + \frac{n}{2}) \end{aligned} \quad (3.79)$$

Fonction Caractéristique

$$\psi_x(q) = 1 + \sqrt{\frac{\pi}{2}} p e^{\frac{p^2}{2}} \left[1 - \operatorname{erf} \frac{p}{\sqrt{2}} \right] \quad (3.80)$$

avec $p = \frac{iq}{\sigma}$

Interprétation

Si X_1 et X_2 sont deux variables normales centrées de variance σ^2 , la variable $X = \sqrt{X_1^2 + X_2^2}$ suit la loi de Rayleigh.

Exemple : tir sur une cible avec des dispersions égales suivant les axes vertical et horizontal ; le rayon du point d'impact suit la loi de Rayleigh. **Autre Exemple : Canal caractérisé par des chemins multiples.** Dans le cas d'un canal physique réel, les réflexions, de quelque type qu'elles soient, font emprunter plusieurs chemins au signal. L'expression du signal reçu devient alors :

$$r(t) = \sum_{i=0}^L \alpha_i(t) e^{-i2\pi(f_c \tau_i(t) + \phi_i(t))} \delta(\tau - \tau_i(t)) \quad (3.81)$$

où $\alpha_i(t)$ sont les amplitudes des chemins i et $\tau_i(t)$ les délais. Dans ce cas, on peut exprimer le signal équivalent passe-bas par :³

$$r(t) = \sum_n \alpha_n(t) e^{-j\omega_c \tau_n(t)} u[t - \tau_n(t)] \quad (3.82)$$

La réponse impulsionnelle du canal équivalent passe-bas a donc l'expression :

$$C(\tau; t) = \sum_n \alpha_n(t) e^{-j\omega_c \tau_n(t)} \delta[\tau - \tau_n(t)] \quad (3.83)$$

Dans le cas où $u(t) = 1$, le signal reçu devient :

$$r(t) = \sum_n \alpha_n(t) e^{-j\theta_n(t)} \quad (3.84)$$

On voit clairement que cette somme, en fonction des angles $\theta_n(t)$, peut engendrer des affaiblissements importants par interférences destructives.⁴ On parle dans ce cas de canal à évanouissement (*fading channel*). En général, on peut approximer $C(\tau; t)$ par un processus gaussien complexe de moyenne nulle. De par ce qui a été dit ci-dessus, on déduit immédiatement que l'enveloppe $|C(\tau; t)|$ suit une loi de Rayleigh. On parle dans ce cas de **canal de Rayleigh**. De plus, si des réflecteurs ou obstacles fixes existent, $C(\tau; t)$ n'est plus à moyenne nulle et on a un **canal de Rice**.

3. En bref, le signal équivalent passe-bas correspond au signal passe-bande en ce sens que son spectre a la même allure, mais qu'il est translaté pour avoir une «fréquence centrale» nulle. On a alors, si $s(t)$ est le signal émis passe-bande, $u(t)$, le signal émis équivalent passe-bas sera relié à $s(t)$ par : $s(t) = \operatorname{Re}[u(t)e^{j\omega_c t}]$.

4. Un cas simple est celui où on a deux chemins de même amplitude et de phase opposée

3.4.7 Loi de Rice

Densité de probabilité.

Soit $Y = X_1^2 + X_2^2$ où X_1 et X_2 sont des variables gaussiennes centrées et indépendantes de moyenne m_1 et m_2 et de variance σ^2 . On note $s^2 = m_1^2 + m_2^2$. La densité de probabilité de Y vaut :

$$p(y) = \frac{1}{2\sigma^2} e^{-\frac{s^2+y}{2\sigma^2}} . I_0 \left(\sqrt{y} \frac{s}{\sigma^2} \right); \quad y \geq 0 \quad (3.85)$$

En définissant $R = \sqrt{Y}$, on obtient :

$$p(r) = \frac{r}{\sigma^2} e^{-\frac{s^2+r^2}{2\sigma^2}} . I_0 \left(\frac{rs}{\sigma^2} \right); \quad r \geq 0 \quad (3.86)$$

Qui est la densité de probabilité d'une variable aléatoire Ricienne.

Fonction de répartition :

$$F(r) = 1 - e^{-(\frac{s^2}{\sigma^2} + \frac{r^2}{\sigma^2})/2} \sum_{k=0}^{\infty} \left(\frac{s}{r} \right)^k I_k \left(\frac{rs}{\sigma^2} \right) \quad (3.87)$$

où $I_k()$ sont les fonctions de Bessel d'ordre k .

Chapter 4

Fonctions aléatoires

4.1 Notion de fonction aléatoire.

Le calcul des probabilités traite de variables aléatoires qui ne dépendent pas, du moins explicitement, du temps ou d'autres paramètres tels que les coordonnées spatiales ; on traite de variables et non de fonctions.

On peut introduire la notion de fonction aléatoire comme étant une fonction du temps et d'un certain nombre de variables aléatoires, e.g.

$$X(t) = A \sin(\omega t + \phi) \quad (4.1)$$

où A, ω et ϕ sont des variables aléatoires. En multipliant le nombre de paramètres aléatoires, on peut arriver à définir de cette manière des fonctions aléatoires très générales, par exemple :

$$X(t) = \sum_{k=1}^n A_k \sin(\omega_k t + \phi_k) \quad (4.2)$$

Si l'on considère l'ensemble des fonctions définies de cette manière, on appelle une **réalisation** de la fonction aléatoire $X(t)$, une fonction $X_r(t)$ où une épreuve a été faite sur les paramètres.

Une autre manière de définir la notion de fonction aléatoire est de dire qu'elle représente un processus où le hasard peut intervenir à tout moment. On arrive alors naturellement à la définition suivante :

Une fonction aléatoire (réelle ou complexe, à une ou plusieurs dimensions) de l'argument t est, pour toute valeur de t , une variable aléatoire (réelle ou complexe, à une ou plusieurs dimensions)

L'argument t sera considéré comme une variable réelle à une dimension, généralement le temps. On peut étendre l'étude à des arguments à plusieurs dimensions, par exemple les coordonnées spatiales.

Il est évident que la théorie à établir ne devra pas seulement décrire la variable aléatoire $X(t_1)$, mais aussi les interdépendances qui existent entre les variables aléatoires $X(t_1), X(t_2), \dots$ pour divers instants.

4.2 Fonctions de répartition et densités de probabilité

Soit une fonction aléatoire scalaire $X(t)$, on peut caractériser la variable aléatoire $X_{t_1} = X(t_1)$ par sa fonction de répartition ou sa densité de probabilité $p_{X_{t_1}}(x_1, t_1)$, fonction des deux variables x_1 et t_1 . Mais ce n'est pas suffisant pour déterminer la fonction aléatoire, car on ne sait rien de l'interdépendance des valeurs prises à des instants différents.

Si l'on considère n valeurs de l'argument $t_1, t_2, \dots, t_n$, on obtient la variable aléatoire à n dimensions $[X_{t_1}, \dots, X_{t_n}]$ que l'on doit caractériser par la densité de probabilité

$$p_{X_{t_1} \dots X_{t_n}}(x_1, t_1; x_2, t_2; \dots; x_n, t_n) \quad (4.3)$$

appelée *densité de probabilité du n^{eme} ordre*.

La connaissance de la densité de probabilité du n^{eme} ordre fournit automatiquement celle des densités d'ordre inférieur à n , car ce sont des densités marginales pouvant se calculer par la formule :

$$\begin{aligned} & p_{X_{t_1} \dots X_{t_k}}(x_1, t_1; x_2, t_2; \dots; x_n, t_k) \\ = & \int_{-\infty}^{+\infty} \dots \int_{-\infty}^{+\infty} p_{X_{t_1} \dots X_{t_n}}(x_1, t_1; x_2, t_2; \dots; x_n, t_n) dx_{k+1} \dots dx_n \end{aligned} \quad (4.4)$$

On peut aussi faire intervenir les densités de probabilité conditionnelle et écrire :

$$\begin{aligned} & p_{X_{t_1} \dots X_{t_n}}(x_1, t_1; x_2, t_2; \dots; x_n, t_n) \\ = & p_{X_{t_1}}(x_1, t_1) p_{X_{t_2}}(x_2, t_2 | x_1, t_1) \dots p_{X_{t_n}}(x_n, t_n | x_1, t_1; x_2, t_2; \dots; x_{n-1}, t_{n-1}) \end{aligned} \quad (4.5)$$

On admettra qu'une fonction aléatoire est complètement définie par sa densité de probabilité d'ordre $n \rightarrow \infty$. Il suffit donc, en principe de procéder aux extensions de la théorie des variables aléatoires pour un nombre de variables infini.

Ces notions peuvent être extrapolées sans difficulté aux fonctions aléatoires multidimensionnelles ou complexes.

Il est certain que l'on ne connaîtra que très rarement les densités de probabilité de tous ordres, mais on verra que de nombreux problèmes relatifs à la transmission de l'énergie peuvent être résolus si l'on connaît la densité de probabilité d'ordre 2. Cette connaissance est d'ailleurs suffisante pour caractériser complètement les fonctions aléatoires gaussiennes.

4.3 Classification des fonctions aléatoires selon leurs propriétés statistiques.

4.3.1 Fonction aléatoire à valeurs indépendantes.

La fonction aléatoire $X(t)$ est dite à valeurs indépendantes si, pour tout ensemble de valeurs $(t_1, \dots, t_n)$, n quelconque, différentes de l'argument, les variables aléatoires $X_{t_1}, \dots, X_{t_n}$ sont indépendantes. On a alors :

$$p_{X_{t_1} \dots X_{t_n}}(x_1, t_1; x_2, t_2; \dots; x_n, t_n) = p_{X_{t_1}}(x_1, t_1) p_{X_{t_2}}(x_2, t_2) \dots p_{X_{t_n}}(x_n, t_n) \quad (4.6)$$

Une telle fonction aléatoire est entièrement définie par sa première densité de probabilité. A tout instant, le futur est indépendant du présent et du passé, car on peut écrire :

$$p_{X_{t_n}}(x_n, t_n | x_1, t_1; x_2, t_2; \dots; x_{n-1}, t_{n-1}) = p_{X_{t_n}}(x_n, t_n) \quad (4.7)$$

4.3.2 Fonction aléatoire à valeurs non corrélées ou orthogonales.

La fonction aléatoire $X(t)$ est dite à valeurs non corrélées ou orthogonales si, pour tout couple de valeurs différentes de l'argument (t_1, t_2) , les variables aléatoires X_{t_1} et X_{t_2} sont non corrélées (orthogonales), c'est-à-dire :

$$\text{non corrélation} \quad \text{Cov}[X_{t_1}, X_{t_2}] = 0 \quad (4.8)$$

$$\text{orthogonalité} \quad E\{X_{t_1}\} = E\{X_{t_2}\} = 0 \text{ et } \text{Cov}[X_{t_1}, X_{t_2}] = 0 \quad (4.9)$$

La notion de non-corrélation est un affaiblissement de celle d'indépendance et est surtout utile dans la théorie du second ordre¹. Les deux notions se confondent pour des fonctions aléatoires normales.

4.3.3 Fonction aléatoire additive.

C'est une fonction aléatoire à accroissements indépendants : pour tout ensemble de n couples de valeurs différentes $(t'_k, t''_k), k = 1, \dots, n$ (n quelconque), les accroissements $\Delta_k = X_{t''_k} - X_{t'_k}$ sont des variables aléatoires indépendantes.

4.3.4 Fonction aléatoire gaussienne

La fonction aléatoire $X(t)$ est normale ou gaussienne si, pour tout ensemble de valeurs $(t_1, \dots, t_n)$ de l'argument (n quelconque), la variable aléatoire vectorielle à n dimensions $\mathbf{X} = [X_{t_0} \dots X_{t_n}]^T$ est normale, c'est-à-dire :

$$p_{X_{t_1} \dots X_{t_n}}(x_1, t_1; x_2, t_2; \dots; x_n, t_n) = \sqrt{\frac{\text{Det} \mathbf{C}^{-1}}{(2\pi)^n}} e^{-\frac{1}{2}(\mathbf{x} - \mathbf{m})^H \mathbf{C}^{-1}(\mathbf{x} - \mathbf{m})} \quad (4.10)$$

où

$$\mathbf{m} = E\{\mathbf{X}\} \quad (4.11)$$

4.4 Moments d'une fonction aléatoire

4.4.1 Moyenne ou espérance mathématique

Soit une fonction aléatoire réelle ou complexe, scalaire ou vectorielle $X(t)$. La moyenne ou espérance mathématique de $X(t)$ est une fonction certaine réelle ou complexe, scalaire ou vectorielle de même dimension que $X(t)$, égale à tout instant t à la moyenne de la variable aléatoire $X(t)$.

On note :

$$\boxed{m_x(t) \text{ ou } E\{X(t)\}} \quad (4.12)$$

1. i.e. si on étudie uniquement les densités de probabilité d'ordre un et deux

Il s'agit d'une **moyenne d'ensemble** définie par

$$m_x(t) = \int_{-\infty}^{\infty} xp_X(x, t)dx \quad (X \text{ scalaire réelle}) \quad (4.13)$$

où le temps n'intervient que comme variable muette. Par exemple, si l'on étudie la tension de bruit à la sortie d'un amplificateur, on peut s'imaginer que l'on dispose de n amplificateurs macroscopiquement identiques et que l'on fasse à l'instant t la moyenne des tensions de sortie des amplificateurs. Pour $n \rightarrow \infty$, cette moyenne tend vers $m_x(t)$, en probabilité.

4.4.2 Variance. Covariance.

Covariance d'une fonction aléatoire.

Soit une fonction aléatoire réelle ou complexe, scalaire ou vectorielle $X(t)$ de dimension $n \times 1$. La matrice de covariance de $X(t)$ est une fonction certaine réelle ou complexe, scalaire ou matricielle (de dimension $n \times n$) fonction de t et t' égale pour tout couple (t, t') à la covariance des variables aléatoires $X(t)$ et $X(t')$.

$$C_x(t, t') = E\{[X(t) - m_x(t)][X(t') - m_x(t')]^H\} \quad (4.14)$$

Variance d'une fonction aléatoire.

Il s'agit de la covariance pour $t = t'$

$$\sigma_x^2(t) = C_x(t, t) \quad (4.15)$$

Covariance mutuelle de deux fonctions aléatoires.

Soient deux fonctions aléatoires $X(t)$ et $Y(t)$ de dimensions $(n \times 1)$ et $(m \times 1)$, respectivement. Leur matrice de covariance mutuelle est, par définition, la matrice $(n \times m)$ fonction de t et t'

$$C_{XY}(t, t') = E\{[X(t) - m_x(t)][Y(t') - m_y(t')]^H\} \quad (4.16)$$

4.4.3 Propriétés des variances et covariances

Fonctions aléatoires scalaires, réelles ou complexes.

$$\sigma_x^2(t) = C_X(t, t) \geq 0 \quad (4.17)$$

$$C_X(t, t') = C_X^*(t', t) \quad (4.18)$$

$$|C_X(t, t')| \leq \sqrt{\sigma_x^2(t)\sigma_x^2(t')} \quad (4.19)$$

Ceci est une conséquence immédiate des propriétés d'hermiticité et de non négative définition de la matrice de covariance d'une variable aléatoire vectorielle appliquées à la variable à deux dimensions $\begin{bmatrix} X(t) \\ X(t') \end{bmatrix}$.

La fonction de covariance est définie non négative, c'est-à-dire que, pour toute fonction continue $\xi(t)$, on a

$$\int_{\tau} \int_{\tau} \xi^*(t) C_x(t, t') \xi(t') dt dt' \geq 0 \quad (\text{réel !}) \quad (4.20)$$

Fonctions aléatoires vectorielles, réelles ou complexes.

On peut assez aisément étendre aux fonctions aléatoires vectorielle les propriétés qui viennent d'être démontrées pour les fonctions scalaires : les variances, covariances et covariances mutuelles sont des matrices telles que :

$$\mathbf{C}_{\mathbf{X}}(t, t') = \mathbf{C}_{\mathbf{X}}(t', t)^H \quad ; \quad \sigma_{\mathbf{x}}^2(t) = \sigma_{\mathbf{x}}^{2H}(t) \quad (4.21)$$

$$|[\mathbf{C}_{\mathbf{X}}(t, t')]_{ij}| \leq \sqrt{[\mathbf{C}_{\mathbf{X}}(t, t)]_{ii} [\mathbf{C}_{\mathbf{X}}(t', t')]_{jj}} \quad (4.22)$$

$$\mathbf{C}_{\mathbf{XY}}(t, t') = \mathbf{C}_{\mathbf{YX}}^H(t', t) \quad (4.23)$$

Si $\xi(t)$ est un vecteur à n dimensions :

$$\int_{\tau} \int_{\tau} \xi^H(t) \mathbf{C}_{\mathbf{X}}(t, t') \xi(t') dt dt' \geq 0 \quad (4.24)$$

4.4.4 Fonctions aléatoires non corrélées ou orthogonales.

Pour une fonction vectorielle, la non corrélation revient à la nullité de la matrice de covariance ($\mathbf{C}_{\mathbf{X}}(t, t')$) pour $t \neq t'$. On aura non corrélation de deux fonctions aléatoires $\mathbf{X}(t)$ et $\mathbf{Y}(t)$ si la matrice de covariance mutuelle $\mathbf{C}_{\mathbf{XY}}(t, t')$ est identiquement nulle, même pour $t = t'$ et orthogonalité si, en outre, $\mathbf{m}_{\mathbf{X}} = \mathbf{m}_{\mathbf{Y}} = \mathbf{0}$.

4.5 Stationnarité et Ergodisme

4.5.1 Stationnarité.

Il est courant, dans la pratique, que les propriétés statistiques d'une fonction aléatoire ne semblent pas évoluer au cours du temps ; c'est souvent le cas pour une tension de bruit observée à l'oscilloscope. Mathématiquement, on exprime cette permanence en introduisant les concepts de stationnarité. Cette notion est cependant réservée aux fonctions aléatoires s'étendant du domaine $-\infty < t < \infty$, de sorte que, rigoureusement, de telles fonctions n'existent pas. Cependant, si les caractéristiques statistiques d'une fonction aléatoire ne se modifient pas sur la durée d'observation, on peut toujours imaginer qu'il en est de même sur tout le passé et le futur.

Stationnarité stricte.

Une fonction aléatoire est dite strictement stationnaire si toutes ses densités de probabilité ne dépendent pas du choix de l'origine des temps, c'est-à-dire ne changent pas lorsqu'on remplace t par $t + t_0$, t_0 quelconque.

Ceci implique que la première densité de probabilité est indépendante de t et que la $n^{\text{ème}}$ densité de probabilité (n quelconque) ne dépend des t_i que par les seules différences $t_2 - t_1, t_3 - t_1, \dots, t_n - t_1$ (en prenant t_1 comme référence).

Cette notion n'étant guère utilisable en pratique, on introduit la stationnarité faible.

Stationnarité faible (au sens large).**Définition.**

- La fonction aléatoire réelle ou complexe, scalaire ou vectorielle $X(t)$ est dite faiblement stationnaire si son espérance mathématique est indépendante de t et si sa covariance ne dépend de t et t' que par la différence $t - t'$.

$$m_x(t) = m \quad \text{constante} \quad (4.25)$$

$$\text{Cov}[X(t)] = C_X(\tau) \quad \text{avec } \tau = t - t' \quad (4.26)$$

- Deux fonctions aléatoires $X(t)$ et $Y(t)$ sont dites mutuellement stationnaires au sens faible si elles sont chacune faiblement stationnaires et si en outre leur covariance mutuelle ne dépend que de $\tau = t - t'$.

Propriétés des variances et covariances. Les propriétés générales des variances et covariances se transposent de la manière suivante pour les fonctions aléatoires faiblement stationnaires :

1. $\sigma_x^2 = C_X(0)$ est une constante (scalaire ou matrice).
2. Hermiticité : $C_X(\tau) = C_X(-\tau)^H$; $C_{XY}(\tau) = C_{YX}(-\tau)^H$
3. Positive-définition

- Pour une fonction scalaire :

$$\begin{aligned} \sigma_x^2 &= C_X(0) \geq 0 \quad (\text{nécessairement réel}) \\ |C_X(\tau)| &\leq C_X(0) \end{aligned}$$

Pour une fonction vectorielle, outre cette propriété valable pour chaque composante, i.e. pour les éléments diagonaux de la matrice de covariance comparés à ceux de la matrice de variance, on a

$$|[C_X(\tau)]_{ij}| \leq \sqrt{[\sigma_x^2]_{ii} [\sigma_x^2]_{jj}} \quad (4.27)$$

- La matrice $\S_X(\omega) = \mathcal{F}\{C_X(\tau)\}$, dont les éléments sont les transformées de Fourier de ceux de la matrice de covariance, est définie non négative. En particulier, pour une fonction scalaire réelle ou complexe : $\S_X(\omega) \geq 0$ et pour deux fonctions scalaires réelles ou complexes

$$|S_{XY}(\omega)| = |S_{YX}^*(\omega)| \leq \sqrt{S_X(\omega) S_Y(\omega)}$$

Cas des fonctions aléatoires périodiques. Une fonction aléatoire $X(t)$ est dite périodique, de période T , si son espérance mathématique est périodique et si sa fonction de covariance $C_x(t, t')$ est périodique en t et en t' .

Il en est ainsi si toutes les réalisations possibles de la fonction le sont. Inversement, si $m_x(t)$ et $C_X(t, t')$ sont périodiques, $X(t)$ et $X(t - T)$ sont égales avec une probabilité unité.

Bien sur, une fonction aléatoire périodique n'est pas nécessairement stationnaire au sens faible, mais pour la rendre telle, il suffit de la décaler sur l'axe t d'une quantité t_0 qui est une variable de chance sur le domaine $[0, T]$, c'est-à-dire de rendre l'origine de la période purement aléatoire.

4.5.2 Ergodisme.

Dans la pratique, on dispose souvent d'une seule réalisation de la fonction aléatoire à l'étude. On imagine difficilement devoir construire un millier d'amplificateurs pour déterminer les caractéristiques statistiques de ceux-ci. Dès lors, on est confronté à la question suivante : dans quelle mesure peut-on déduire de la réalisation dont on dispose (la seule et unique) certaines caractéristiques statistiques de la fonction aléatoire, par exemple des moments qui sont des moyennes faites sur l'ensemble des réalisations ? La théorie de l'ergodisme tente de répondre à cette question, principalement pour des fonctions aléatoires.

En bref, cette théorie essaiera de rapprocher les moyennes d'ensemble (moyenne des variables aléatoires X_{t_i} des moyennes temporelles (notées $\langle X(t_0) \rangle$)

Typiquement, on considère une variable aléatoire définie à partir d'une fonction aléatoire, généralement une fonction des valeurs prises par la fonction aléatoire à divers instants :

$$\eta = f[X_{t_1}, \dots, X_{t_n}] \quad (4.28)$$

Par exemple X_{t_1} ou le produit $X_{t_1} X_{t_2}$. En faisant l'hypothèse que $X(t)$ est stationnaire, on peut espérer qu'en faisant "glisser" sur l'axe des temps l'échantillon $X_r(t)$ disponible, i.e. en considérant $X_r(t + t_0)$, et ce pour tout t_0 , on obtienne un ensemble de fonctions présentant les mêmes propriétés que l'ensemble des réalisations de la fonction aléatoire $X(t)$.

Nous obtenons alors, pour η :

$$\eta(t_0) = f[X_{t_0+t_1}, \dots, X_{t_0+t_n}] \quad (4.29)$$

et on peut calculer des moyennes de $\eta(t_0)$ sur l'ensemble des fonctions glissées :

$$\langle \eta \rangle \equiv \lim_{T \rightarrow \infty} \frac{1}{T} \int_0^T \eta(t_0) dt_0 \quad (4.30)$$

dont on peut espérer qu'elle tende vers la moyenne d'ensemble $E\{f\} = E\{\eta\}$.

Evidemment, η étant une variable aléatoire, $\eta(t_0)$ est une fonction aléatoire ; pour obtenir la moyenne temporelle $\langle \eta \rangle$, il faut intégrer et prendre un passage à la limite.

Comme nous étudions surtout la théorie du second ordre, nous nous intéresserons uniquement à la possibilité de déterminer $m_x(t)$ et $C_X(t, t')$ dans le cadre cohérent d'hypothèses :

– **convergence en moyenne quadratique :**

- On dit que la suite de fonctions aléatoires $X_n(t)$ définies sur $t \in \tau$ converge en moyenne quadratique pour $n \rightarrow \infty$ vers la fonction aléatoire $X(t)$ définie sur τ , et l'on écrit :

$$l.i.m. \quad X_n(t) = X(t) \quad (4.31)$$

si, pour tout $t \in \tau$, la variable aléatoire $X_n(t)$ converge en moyenne vers $X(t)$, i.e. si :

$$\lim_{n \rightarrow \infty} E\{|X_n(t) - X(t)|^2\} = 0 \quad (4.32)$$

- On dit que la fonction aléatoire $X(t)$ converge en moyenne quadratique pour $t \rightarrow t_0$ vers la variable aléatoire X_0 , et l'on écrit :

$$\lim_{n \rightarrow \infty} l.i.m. \quad X(t) = X_0 \quad (4.33)$$

si

$$\lim_{t \rightarrow t_0} E\{|X(t) - X_0|^2\} = 0 \quad (4.34)$$

– **intégrale en moyenne quadratique**

- **Définition :** L'intégrale en moyenne quadratique de la fonction aléatoire $X(t)$ est, quand elle existe, la variable aléatoire :

$$\int_a^b X(t) dt = \lim_{n \rightarrow \infty} l.i.m. \quad \sum_{k=1}^n X(t'_{k,n})(t_{k,n} - t_{k-1,n}) \quad (4.35)$$

avec

$$a = t_{0,n} < t_{1,n} < \dots < t_{n,n} = b \text{ partition de } (a,b)$$

$$(t_{k,n} - t_{k-1,n}) \rightarrow 0 \text{ quand } n \rightarrow \infty$$

$$t'_{k,n} \text{ point quelconque de } (t_{k-1,n}, t_{k,n})$$

- **Théorème :** La condition nécessaire et suffisante d'existence de $\int_a^b X(t) dt$ est l'existence des intégrales certaines :

$$\int_a^b m_x(t) dt \quad ; \quad \int_a^b \int_a^b C_X(t, t') dt dt' \quad (4.36)$$

- stationnarité faible.

Moyennes temporelles.

Moyenne temporelle d'une fonction aléatoire $X(t)$ C'est, quand elle existe,

$$\langle X(t_0) \rangle = \lim_{T \rightarrow \infty} l.i.m. \quad \left[Y(T, t_0) \equiv \frac{1}{T} \int_{t_0}^{t_0+T} X(t) dt \right] \quad (4.37)$$

On notera qu'en effectuant un changement de variable :

$$Y(T, t_0) = \frac{1}{T} \int_0^T X(t + t_0) dt \quad (4.38)$$

$Y(T, t_0)$ est une variable aléatoire dont la limite $\langle X(t_0) \rangle$ est en général aléatoire et dépendante de t_0 , ayant les mêmes dimensions que $X(t)$.

Valeur quadratique moyenne (temporelle) d'une fonction aléatoire $X(t)$. C'est, quand elle existe :

$$\langle |X(t_0)|^2 \rangle = \lim_{T \rightarrow \infty} \left[Y(T, t_0) \equiv \frac{1}{T} \int_{t_0}^{t_0+T} |X(t)|^2 dt \right] \quad (4.39)$$

C'est une variable aléatoire scalaire positive, dépendant de t_0 . Cette notion est surtout utilisée pour une fonction aléatoire scalaire, auquel cas elle prend souvent la signification d'énergie moyenne de la réalisation $X(t)$, à un facteur constant près.

Fonction de corrélation d'une fonction aléatoire $X(t)$ $\langle X(t_0) \rangle$ étant supposé existant, c'est, quand elle existe :

$$R_X(\tau, t_0) = \lim_{T \rightarrow \infty} \left[Y(T, t_0) \equiv \frac{1}{T} \int_{t_0}^{t_0+T} [X(t+\tau) - \langle X(t_0+\tau) \rangle][X(t) - \langle X(t_0) \rangle]^H dt \right] \quad (4.40)$$

C'est en général une matrice aléatoire ($n \times n$) si $X(t)$ est un vecteur ($n \times 1$), fonction aléatoire de τ et de t_0 .

Fonction de corrélation mutuelle de deux fonctions aléatoires $X(t)$ et $Y(t)$ On suppose $\langle X(t_0) \rangle$ et $\langle Y(t_0) \rangle$ existants. C'est, quand elle existe :

$$R_{XY}(\tau, t_0) = \lim_{T \rightarrow \infty} \left[Y(T, t_0) \equiv \frac{1}{T} \int_{t_0}^{t_0+T} [X(t+\tau) - \langle X(t_0+\tau) \rangle][Y(t) - \langle Y(t_0) \rangle]^H dt \right] \quad (4.41)$$

C'est en général une matrice aléatoire ($n \times m$).

Ergodisme au sens large d'une fonction aléatoire.

Définition : une fonction aléatoire (non nécessairement stationnaire) est ergodique au sens large si $\langle X(t_0) \rangle$ existe et est une variable aléatoire indépendante de t_0 .

Théorème : si $X(t)$ est faiblement stationnaire et intégrable en moyenne quadratique sur tout intervalle fini, elle est ergodique au sens large.

Ergodisme au sens strict d'une fonction aléatoire.

Définition : une fonction aléatoire (non nécessairement stationnaire) est ergodique au sens strict si

1. $\langle X(t_0) \rangle$ existe et est un nombre certain indépendant de t_0 . (on le note $\langle X \rangle$);
2. $\langle X \rangle = \lim_{t \rightarrow \infty} m_x(t)$

Cette limite étant supposée existante. Ce n'est pas toujours le cas, même lorsque $m_x(t)$ est fini. Par exemple $X(t) = A \sin(\omega t)$ où A est aléatoire, a une espérance mathématique $m_x(t) = m_A \sin(\omega t)$.

Ergodisme relativement à la fonction de corrélation.

Définition : une fonction aléatoire scalaire faiblement stationnaire $X(t)$ est ergodique relativement à sa fonction de corrélation si :

1. elle est strictement ergodique ;
2. sa fonction de corrélation $R_X(\tau, t_0)$ existe, est certaine et indépendante de t_0 . On la note $R_X(\tau)$;
3. $R_X(\tau) = C_X(\tau)$. La fonction de corrélation est égale à la fonction de covariance

Conclusions.

Dans les applications, on suppose fréquemment qu'il y a ergodisme relativement à la fonction de corrélation, bien que souvent on ne puisse le démontrer du fait d'une connaissance insuffisante du modèle mathématique de la fonction aléatoire stationnaire. Cette *hypothèse ergodique* revient à admettre l'égalité des moyennes d'ensemble et des moyennes temporelles jusqu'au deuxième ordre, ces dernières étant calculées sur la réalisation dont on dispose.

$$\begin{aligned} m_x &= \langle X \rangle \\ C_X(\tau) &= R_X(\tau) \\ \sigma_x^2 &= C_X(0) = R_X(0) = \langle |X - m_x|^2 \rangle \end{aligned}$$

En particulier, l'égalité des fonctions de corrélation (facile à estimer) et de covariance (plus difficile) est d'une importance capitale pour la détermination de la bande passante (transformée de Fourier de la covariance) et des propriétés d'indépendance et d'orthogonalité des signaux.

Chapter 5

Modulations analogiques : les principes.

5.1 Rôle de la modulation

Nous avons déjà indiqué que la plupart des signaux de source devaient être adaptés de manière à pouvoir être envoyé sur le canal, c'est le rôle principal de la modulation. Dans la première partie de ce cours, nous nous attachons aux modulations analogiques, ce qui veut dire que le message (ou signal de source) est de nature analogique et qu'elle n'est pas transformée au préalable en signal numérique. On peut détailler les différents rôles de la modulation de la manière suivante.

Facilité d'émission La longueur d'onde d'une onde électromagnétique, et donc la taille de l'antenne nécessaire à l'émission, est inversement proportionnelle à sa fréquence. Il est donc évident que pour avoir des tailles d'antennes acceptables, il est impératif de moduler les signaux BF (parole de 0 à 4 kHz). D'autre part, des considérations de distances de rayonnement (plus faibles à 100 MHz qu'à 10 MHz par exemple) guideront le choix des fréquences utilisées et éventuellement des types de modulation

Réduction du bruit et des interférences Dans la suite de ce cours, nous verrons comment la modulation FM par exemple permet de réduire l'influence du bruit additif, au dépens d'une largeur de bande supérieure.

Choix d'un canal fréquentiel Voir le récepteur hétérodyne !

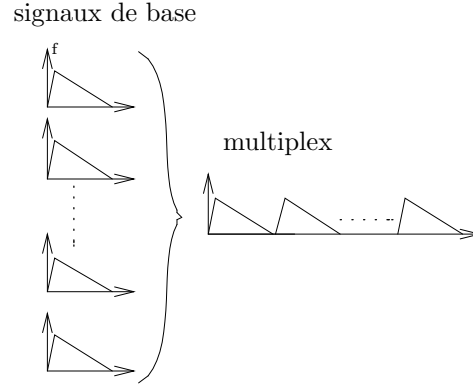
Multiplexage fréquentiel Par exemple dans le cas de la téléphonie, on groupera un nombre important de communications sur un seul canal (un câble coaxial par exemple) en les décalant en fréquence sur la bande passante disponible (figure 5.1)

Si on considère que l'on désire un signal "autour d'une fréquence centrale", nous obtenons une expression générique de ce signal de la manière suivante :

$$s(t) = A(t)[\sin(\omega_c(t)t + \phi(t)] \quad (5.1)$$

On peut alors caractériser les différents types de modulation (qu'elles soient numériques ou analogiques d'ailleurs):

La modulation d'amplitude où on fait varier $A(t)$ en fonction du temps, les autres paramètres (ω_c et ϕ) étant constants.

FIG. 5.1 – *Exemple de multiplex fréquentiel.*

La modulation de fréquence où on fait varier $\omega_c(t)$

La modulation de phase où on fait varier $\phi(t)$.

Rien ne nous empêche évidemment de croiser les modulations. Nous verrons que dans le cas numérique, on utilise des modulations d'amplitude et de phase en même temps.

5.2 La modulation d'amplitude

5.2.1 Modulation DSB (Double Side Band).

Soit un message $m(t)$, la signal modulé en amplitude est donné par l'expression :

$$s_{dsb}(t) = A[1 + Km(t)] \sin(\omega_c t) \quad (5.2)$$

où l'on a fait l'hypothèse que $|m(t)| < 1$. L'amplitude instantannée est donc bien proportionnelle au signal modulant dans ces conditions.

K est appelé l'indice de modulation et déterminera la profondeur de modulation. Idéalement $K < 1$, sans quoi la détection d'enveloppe que nous verrons plus loin est impossible. Le graphique ci-dessous illustre l'allure temporelle des signaux.

Spectre du signal DSB

En utilisant le théorème de modulation, on démontre aisément que le spectre du signal DSB vaut :

$$S_{dsb}(\omega) = \pi A[\delta(\omega + \omega_c) + \delta(\omega - \omega_c)] + \frac{AK}{2}[M(\omega + \omega_c) + M(\omega - \omega_c)] \quad (5.3)$$

On y reconnaît les deux bandes latérales dans $M(\omega + \omega_c)$ aux fréquences positives et $M(\omega - \omega_c)$ aux fréquences négatives ainsi que la porteuse ($\delta(\omega + \omega_c)$ et $\delta(\omega - \omega_c)$). La figure 5.3 illustre cela.

En fait, "l'apparition" des deux bandes latérales est probablement la première bonne justification à la "réalité mathématique" des fréquences négatives. Celles-ci, quoique intuitivement

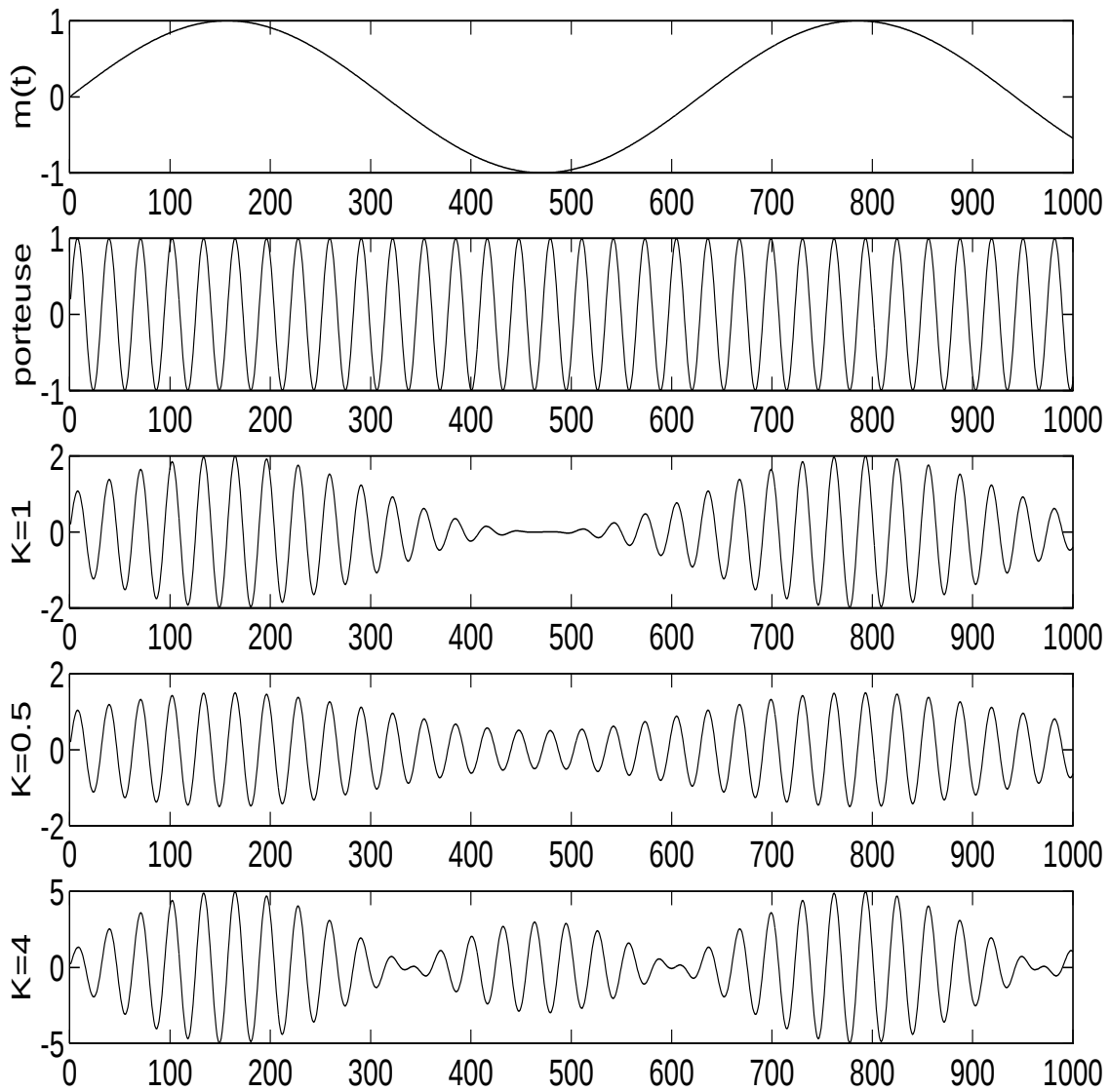


FIG. 5.2 – Modulation DSB, indice de modulation

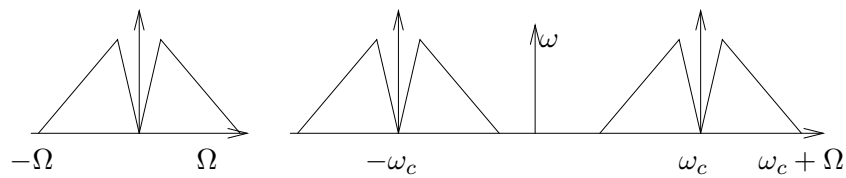


FIG. 5.3 – Modulation DSB, spectre.

assez peu explicables, se retrouvent au moment de la modulation. Une autre manière (grossière) de les “expliquer” est que, au partir de signaux réels, nous faisons une transformée qui nous mène dans le monde complexe, il paraît donc logique que cette transformation amène des choses à priori inattendues.

Emetteur DSB - Principe.

Le principe de base de l'émetteur DSB consiste à faire passer le message et la porteuse dans une non linéarité, de manière à faire apparaître le produit de la porteuse et du message en sortie.

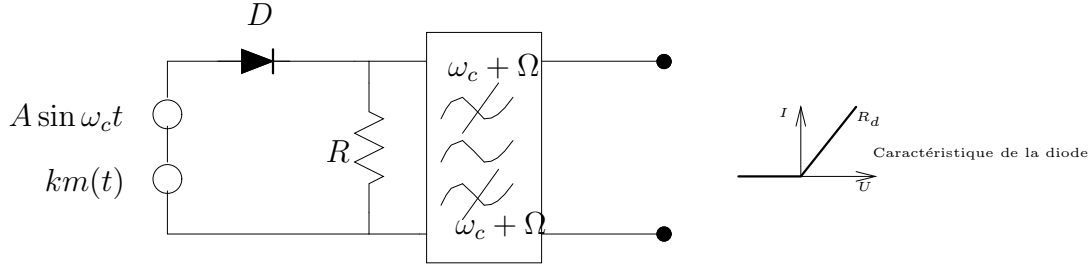


FIG. 5.4 – *Emetteur AM (principe)*

Dans cet émetteur, nous faisons l'hypothèse que le signal $m(t)$ est d'amplitude faible par rapport à la porteuse $A \sin \omega_c t$. La caractéristique idéale de la diode nous permet d'écrire :

$$\begin{aligned} \text{si } A \sin \omega_c t + km(t) > 0 &\Rightarrow V_R = \frac{R}{R+R_d} [A \sin \omega_c t + km(t)] \\ \text{si } A \sin \omega_c t + km(t) < 0 &\Rightarrow V_R = 0 \end{aligned} \quad (5.4)$$

La sortie peut donc être écrite sous la forme

$$r(t) = [A \sin \omega_c t + km(t)]p(t) \quad (5.5)$$

où $p(t)$ est un signal carré à la fréquence ω_c , illustré à la figure 5.5.

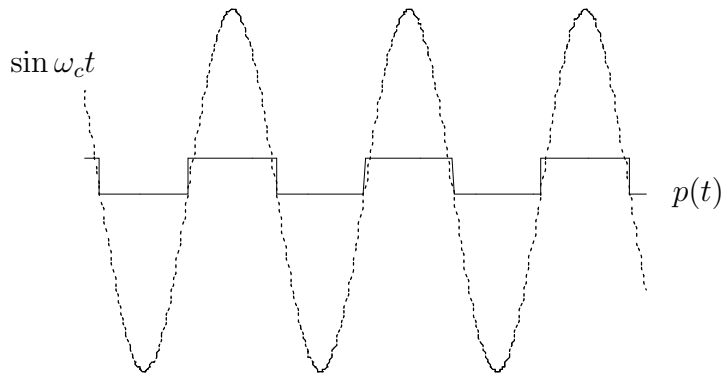


FIG. 5.5 – *Porteuse et signal p(t) en génération DSB*

On peut écrire $p(t)$ sous la forme :

$$p(t) = \frac{1}{2} + \frac{2}{\pi} \left(\cos \omega_c t - \frac{\cos 3\omega_c t}{3} + \frac{\cos 5\omega_c t}{5} - \dots \right) \quad (5.6)$$

Après filtrage passe-bande autour de ω_c on obtient donc bien un signal du type $[cste + m(t)] \sin \omega_c t$.

Démodulateur DSB : le détecteur d'enveloppe

On pose l'hypothèse $|Km(t)| \leq 1 \Rightarrow 1 + Km(t) \geq 0$, où on suppose généralement que $m(t)$ est à moyenne nulle. Dans ce cas, on a le détecteur d'enveloppe de la figure 5.6.

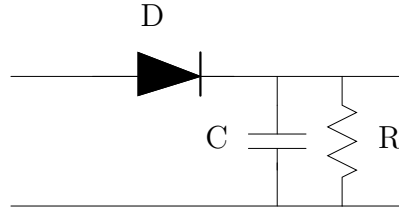


FIG. 5.6 – Détecteur d'enveloppe : montage série

Les hypothèses sur le circuit sont :

1. La source du signal de réception a une impédance nulle.
2. La diode est idéale ($R = R_d$ dans le sens passant)
3. La constante de temps RC est grande par rapport à $\frac{1}{\omega_c}$

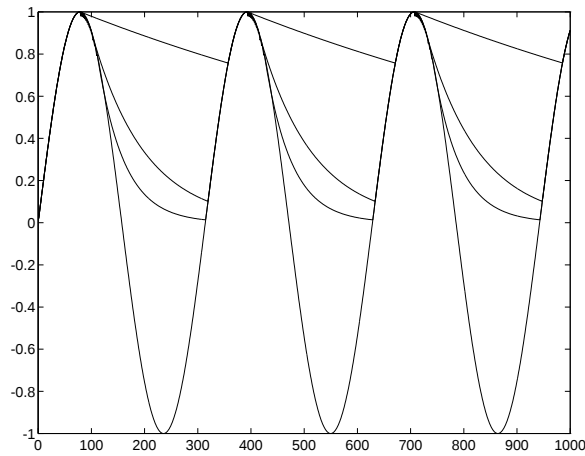


FIG. 5.7 – Principe de fonctionnement du détecteur d'enveloppe

Le principe de fonctionnement est illustré à la figure 5.7. Quand la tension d'entrée est plus grande que la tension de sortie, la diode est passante et le signal de sortie suit le signal d'entrée fidèlement, pour peu que la résistance de diode R_d ne soit pas trop grande. Quand la tension d'entrée est plus faible, la diode devient bloquante et on assiste simplement à la décharge du réseau RC . Les conditions sur celui-ci sont simplement:

$$\frac{1}{f_c} \ll RC \ll \frac{1}{F_0} \text{ où } F_0 \text{ est la plus grande fréquence contenue dans } m(t) \quad (5.7)$$

Si la constante RC est trop faible, on a une décharge trop rapide comme indiqué en figure ??, dans le cas contraire, on observe un décrochage de l'enveloppe

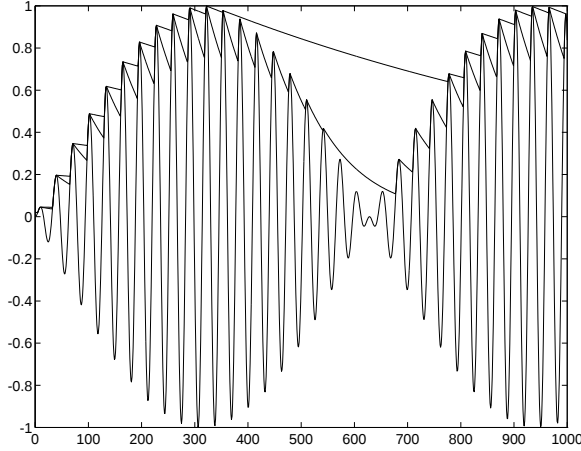


FIG. 5.8 – Décrochage du détecteur d'enveloppe

Répartition de la puissance dans le spectre

Dans la mesure où la porteuse est à une pulsation élevée par rapport à la largeur de bande (i.e. $\omega_c \gg \Omega_O$ où Ω_0 est la pulsation la plus élevée de $m(t)$), on peut considérer que la puissance instantanée du signal vaut ($T_c = 1/\omega_c$):

$$\begin{aligned} P(t) &= \frac{1}{T_c} \int_t^{t+T_c} r^2(t) dt \\ &\approx \frac{1}{2} A^2(t) \end{aligned} \quad (5.8)$$

$$ds_{approx} P_0 [1 + Km(t)]^2$$

Cette puissance fluctue entre $P_0(1 - K)^2$ et $P_0(1 + K)^2$. On peut calculer la puissance moyenne :

$$P = P_0 [1 + 2K \overline{m(t)} + K^2 \overline{m^2(t)}] \quad (5.9)$$

où le surlignage signifie la valeur moyenne temporelle. Dans le cas classique d'un message à moyenne nulle, on a donc :

$$P = P_0[1 + K^2 m_{eff}^2] \quad (5.10)$$

Ce qui signifie que dans le meilleur des cas ($K = 1$ si on a effectué la normalisation $|m(t)| \leq 1 \Rightarrow m_{eff}^2 = \frac{1}{2}$ pour un signal sinusoïdal) la porteuse utilise 66 % de la puissance totale. Dans des cas plus réaliste, la puissance de la porteuse est de plus de 90 % de la puissance totale. Ceci amène naturellement à la modulation DSB-SC, où on a simplement éliminé la porteuse. Il est cependant à noter qu'en radio-diffusion, c'est toujours la modulation DSB qui est utilisée, dans le cas des émissions "AM"

5.2.2 Modulation DSB-SC (Double Side Band - Suppressed Carrier)

Le principe de cette modulation est extrêmement simple : on multiplie le message par la porteuse :

$$s_{dsb-sc}(t) = Am(t) \cdot \cos \omega_c t \quad (5.11)$$

où l'on a fait l'hypothèse que $|m(t)| < 1$.

Spectre du signal DSB-SC

En utilisant le théorème de modulation, on démontre aisément que le spectre du signal DSB-SC vaut :

$$S_{dsb-sc}(\omega) = \frac{A}{2} [M(\omega + \omega_c) + M(\omega - \omega_c)] \quad (5.12)$$

On y reconnaît les deux bandes latérales dans $M(\omega + \omega_c)$ et $M(\omega - \omega_c)$ tandis que la porteuse a disparu. La figure 5.9 illustre cela.

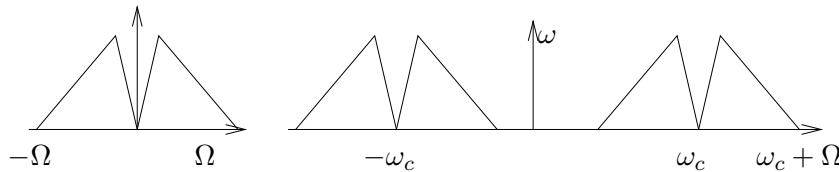
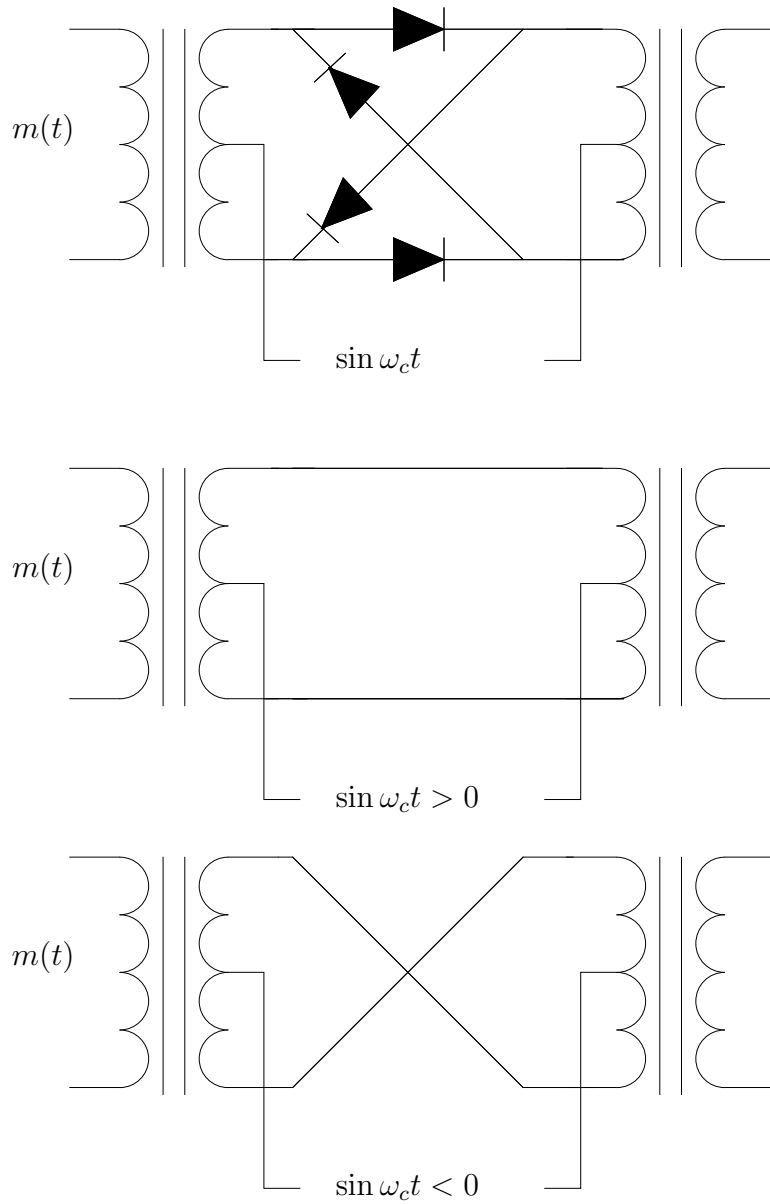


FIG. 5.9 – Modulation DSB-SC, spectre.

Emetteur DSB-SC - Le modulateur balancé.

Le principe de base de l'émetteur DSB-SC consiste également à faire passer le message et la porteuse dans une non linéarité, de manière à faire apparaître le produit de la porteuse et du message en sortie, mais cette fois-ci sans la porteuse. D'un point de vue "schéma bloc", on parlera simplement d'un mélangeur. De tels circuits existent dans le commerce, nous allons décrire le modulateur balancé, qui est utilisé pour les mélangeurs "basse fréquence" (jusque quelque 10 MHz) sous une forme un peu plus évoluée que ci-dessous.

FIG. 5.10 – *Emetteur DSB-SC (principe)*

Dans cet émetteur, nous faisons l'hypothèse que le signal $m(t)$ est d'amplitude faible par rapport à la porteuse $A \sin \omega_c t$. Quand la porteuse est positive, les diodes croisées sont bloquantes et les autres sont passantes : on a donc $+m(t)$ en sortie du modulateur ; dans le cas inverse, ce sont les diodes croisées qui sont passantes et on a $-m(t)$ en sortie.

Les caractéristiques idéales de la diode nous permettent d'écrire :

$$r(t) = m(t)p(t) \quad (5.13)$$

où $p(t)$ est un signal carré à la fréquence ω_c , illustré à la figure 5.11.

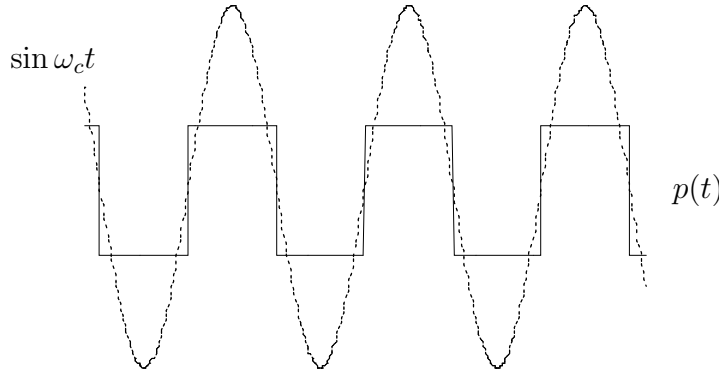


FIG. 5.11 – Porteuse et signal $p(t)$ en génération DSB-SC

On peut écrire $p(t)$ sous la forme :

$$p(t) = \frac{4}{\pi} \left(\cos \omega_c t - \frac{\cos 3\omega_c t}{3} + \frac{\cos 5\omega_c t}{5} - \dots \right) \quad (5.14)$$

Après filtrage passe-bande autour de ω_c on obtient donc bien un signal du type $m(t) \cos \omega_c t$.

Démodulation DSB-SC

La démodulation se fera simplement par le même principe, puisque si on prend :

$$d(t) = r(t) \cdot \cos \omega_c t = m(t) \cos^2 \omega_c t = \frac{m(t)}{2} [1 + \cos 2\omega_c t] \quad (5.15)$$

on obtient bien le signal $m(t)$ après filtrage passe-bas, pour éliminer $\cos 2\omega_c t$. D'autre part, comme un exercice a pu l'illustrer, il est important d'effectuer la démodulation avec une porteuse en phase, sous peine d'observer une atténuation importante. C'est ce que l'on appelle la *démodulation cohérente*. La difficulté de l'opération consistera donc principalement en la *récupération de porteuse*. Un moyen simple d'effectuer cette récupération est d'élever le signal au carré :

$$d^2(t) = m^2(t) \cos^2(\omega_c t + \phi) = \frac{m^2(t)}{2} [1 + \cos 2(\omega_c t + \phi)] \quad (5.16)$$

Comme $m^2(t)$ comprend nécessairement une composante continue, on a obligatoirement une raie à la fréquence $2f_c$ (voir le théorème de modulation). On peut acquérir cette raie à l'aide

d'un filtre passe-bande étroit (ou d'une boucle à verrouillage de phase) et effectuer une division par deux. L'équation précédente montre bien que la phase sera ainsi récupérée correctement.

5.2.3 Modulation à Bande Latérale Unique (BLU, SSB: Single Side-Band)

Il paraît assez naturel, pour économiser de la puissance et de la bande passante, de passer à une seule bande latérale. Une méthode simple de réalisation consiste à filtrer la bande non désirée. Cependant, dans la plupart des cas, cela mènera à des filtres de facteur de qualité prohibitif. Une des solutions à ce problème a été exposée dans le cadre du rappel sur les transformées de Fourier. L'autre solution consiste à chercher l'expression analytique du signal BLU et d'en déduire un schéma de génération d'un signal BLU. C'est ce que nous faisons ci-dessous.

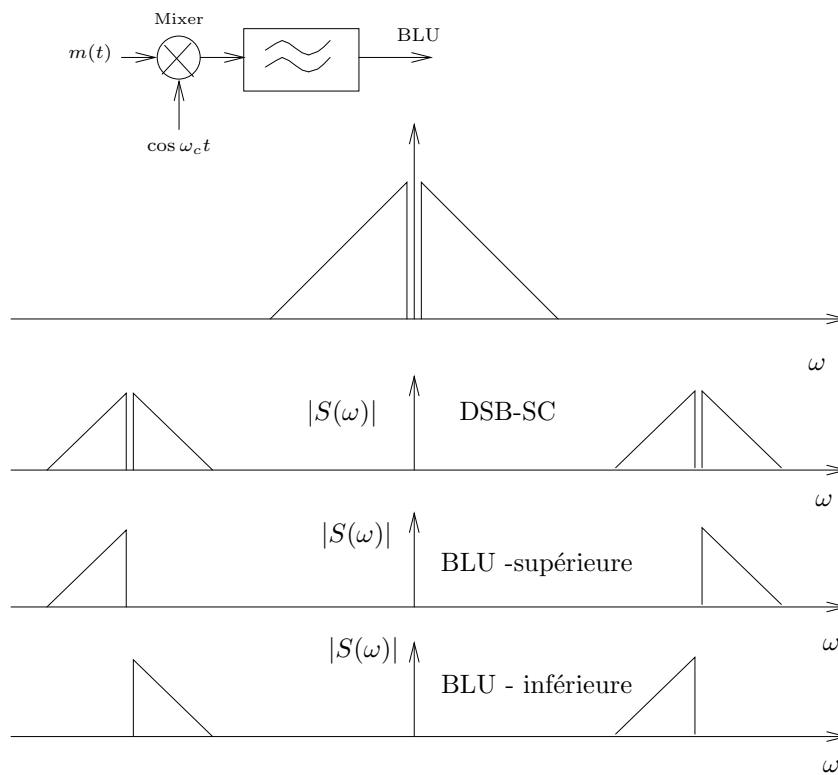


FIG. 5.12 – Génération d'un signal BLU par filtrage

Expression analytique d'un signal BLU

La Transformée de Hilbert Nous introduisons ici la transformée de Hilbert, qui est un outil classique en télécommunications et est utilisé pour la BLU.

Soit un filtre de fonction de transfert $H(f) \Leftrightarrow h(t)$:

$$H(f) = -j \operatorname{sgn} f \text{ où } \operatorname{sgn} f = \begin{cases} 1, & f > 0 \\ 0, & f = 0 \\ -1 & f < 0 \end{cases} \quad (5.17)$$

D'autre part, on peut démontrer :

$$h(t) = \frac{1}{\pi t} \quad (5.18)$$

On appelle *transformée de Hilbert* de $x(t)$:

$$\hat{x}(t) = h(t) \otimes x(t) = \text{TF}^{-1}[-j \operatorname{sgn} f X(f)] \quad (5.19)$$

Ce qui, dans le domaine temporel, peut s'écrire :

$$\hat{x}(t) = \int_{-\infty}^{\infty} \frac{x(\tau)}{\pi \cdot (t - \tau)} d\tau = \int_{-\infty}^{\infty} \frac{x(t - \tau)}{\pi \cdot \tau} d\tau \quad (5.20)$$

Exemple 5.1 *Transformée de Hilbert d'une sinusoïde*

soit $x(t) = \cos \omega_c t$. En fréquentiel, on peut écrire :

$$X(f) = \frac{1}{2} [\delta(f - f_c) e^{-j\pi/2} + \delta(f + f_c) e^{j\pi/2}] \quad (5.21)$$

En prenant la transformée de Fourier inverse, on obtient :

$$\hat{x}(t) = \cos(\omega_c t - \pi/2) = \sin \omega_c t \quad (5.22)$$

Cet exemple confirme bien l'idée selon laquelle **la transformée de Hilbert correspond à un déphasage de 90 degrés**

◁

Supposons que l'on désire démoduler un signal BLU-inférieure, cela correspond à un filtrage par :

$$H(\omega) = \begin{cases} 1, & |\omega| \leq \omega_c \\ 0, & \text{ailleurs} \end{cases} \quad (5.23)$$

La réponse impulsionnelle de ce filtre vaut $h(t) = \frac{\sin \omega_c t}{\pi t}$. Le signal de sortie vaut donc, par convolution :

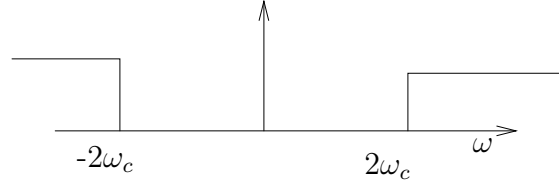
$$\begin{aligned} s(t) &= \int_{-\infty}^{\infty} m(\tau) \cos \omega_c \tau \frac{\sin \omega_c (t - \tau)}{\pi (t - \tau)} d\tau \\ &= \frac{1}{2} \sin \omega_c t \int_{-\infty}^{\infty} \frac{m(\tau)}{\pi (t - \tau)} d\tau + \int_{-\infty}^{\infty} m(\tau) \frac{\sin \omega_c (t - 2\tau)}{\pi (t - \tau)} d\tau \\ &= \frac{1}{2} \hat{m}(t) \sin \omega_c t + \int_{-\infty}^{\infty} m(\tau) \frac{\sin \omega_c (2t - 2\tau - t)}{\pi (t - \tau)} d\tau \\ &= \frac{1}{2} \hat{m}(t) \sin \omega_c t + \cos \omega_c t \int_{-\infty}^{\infty} \frac{\sin 2\omega_c (t - \tau)}{\pi (t - \tau)} d\tau - \sin \omega_c t \int_{-\infty}^{\infty} \frac{\cos 2\omega_c (t - \tau)}{\pi (t - \tau)} d\tau \end{aligned} \quad (5.24)$$

On peut aisément reconnaître en $\int_{-\infty}^{\infty} \frac{\sin 2\omega_c(t-\tau)}{\pi(t-\tau)} dt$ un filtrage passe-bas du signal $m(t)$ par un filtre de fréquence de coupure $2\omega_c$ dont le résultat est forcément $m(t)$, le message ne pouvant pas avoir de contenu spectral au-delà de ω_c .

La deuxième intégrale $\int_{-\infty}^{\infty} \frac{\cos 2\omega_c(t-\tau)}{\pi(t-\tau)} dt$ consiste en un filtrage par un filtre de réponse

$$\frac{\cos 2\omega_c t}{\pi t} = \frac{e^{2j\omega_c t} + e^{-2j\omega_c t}}{2\pi t} \Leftrightarrow -\frac{j}{2} [\text{sgn}(\omega - 2\omega_c) + \text{sgn}(\omega + 2\omega_c)] \quad (5.25)$$

Cette fonction de transfert a l'allure :



Nous obtenons donc le résultat (où LSB signifie Lower SideBand et USB Upper SideBand):

$$\begin{aligned} s_{LSB}(t) &= \frac{A}{2} [m(t) \cos \omega_c t + \hat{m}(t) \sin \omega_c t] \\ s_{USB}(t) &= \frac{A}{2} [m(t) \cos \omega_c t - \hat{m}(t) \sin \omega_c t] \end{aligned} \quad (5.26)$$

Et qui nous permet de concevoir un modulateur SSB comme indiqué en figure 5.13. Il est clair que le déphasage de $\pi/2$ sur le message est une approximation de la transformée de Hilbert et n'est valable que si la bande passante est faible par rapport à la fréquence porteuse.

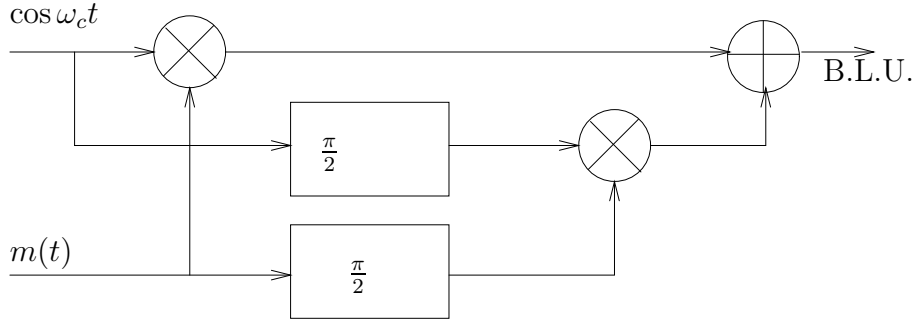


FIG. 5.13 – Modulateur BLU

5.2.4 Démodulation d'un signal BLU

Démodulation synchrone

En multipliant le signal BLU par $4 \cos(\omega_c t + \phi)$, où ϕ est l'erreur de phase, on obtient:

$$\begin{aligned} r(t) &= \frac{A}{2} [m(t) \cos \omega_c t \pm \hat{m}(t) \sin \omega_c t] \cdot 4 \cos(\omega_c t + \phi) \\ &= A m(t) \cos \phi + A m(t) \cos(2\omega_c t + \phi) \mp A \hat{m}(t) \sin \phi \pm A \hat{m}(t) \sin(2\omega_c t + \phi) \end{aligned} \quad (5.27)$$

Soit, après filtrage :

$$r(t) = Am(t) \cos \phi \mp A\hat{m}(t) \sin \phi \quad (5.28)$$

Soit une démodulation synchrone correcte, mais sujette au défaut de cohérence, pour éviter ce problème, on peut éventuellement insérer une porteuse qui permet de faire de la démodulation par détection d'enveloppe.

5.3 Modulation de Fréquence (FM)

Soit le signal

$$s_{FM}(t) = A \cos(\omega_c + \Delta\omega \int_{-\infty}^t m(\tau) d\tau) \quad (5.29)$$

Définissons la *fréquence instantannée*:

$$f_i = \frac{\omega_i}{2\pi} = \frac{1}{2\pi} \frac{d\phi(t)}{dt} \quad (5.30)$$

Dans ce cas, on obtient

$$\omega_i = \frac{d\phi(t)}{dt} = \omega_c + \Delta\omega m(t) \quad (5.31)$$

qui justifie l'appellation "fréquence modulée" puisque la fréquence instantannée est proportionnelle, à une constante près, au message. L'étude de la FM dans le cas d'un signal $m(t)$ quelconque est très complexe, on se limitera donc au cas du signal sinusoïdal.

5.3.1 Déviation de fréquence, indice de modulation

Dans le cas sinusoïdal, on peut écrire le signal FM sous la forme :

$$s_{FM}(t) = A \cos(\omega_c t + \beta \sin \Omega_0 t) \quad (5.32)$$

où β est l'*indice de modulation* et vous pouvez vérifier

$$\beta = \frac{\Delta\omega}{\Omega_0} = \frac{\Delta f}{F_0} \quad (5.33)$$

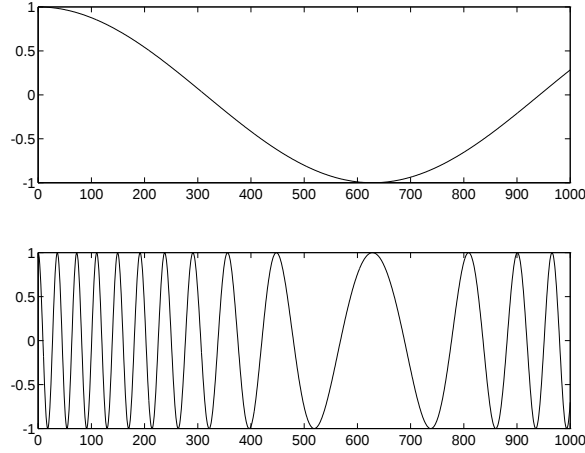
Soit :

$$s_{FM}(t) = A \cos(2\pi f_c t + \frac{\Delta f}{F_0} \sin \Omega_0 t) \quad (5.34)$$

et la fréquence instantannée vaut :

$$f_i = \frac{1}{2\pi} \frac{d\phi(t)}{dt} = f_c + \Delta f \cos \Omega_0 t \quad (5.35)$$

Ce qui justifie l'appellation de "déviation de fréquence" pour Δf . Notons que ce n'est pas parce que la fréquence instantannée varie de $f_c - \Delta f$ à $f_c + \Delta f$ que le contenu spectral est limité à ces fréquences.

FIG. 5.14 – *Signal à fréquence modulée*

5.3.2 Spectre d'un signal FM

Toujours dans le cas d'un signal sinusoïdal, nous obtenons :

$$\begin{aligned} s_{FM}(t) &= A \cos(\omega_c t + \beta \sin \Omega_0 t) \\ &= A [\cos \omega_c t \cos(\beta \sin \Omega_0 t) - \sin \omega_c t \sin(\beta \sin \Omega_0 t)] \end{aligned} \quad (5.36)$$

On peut alors développer $\cos(\beta \sin \Omega_0 t)$ en série de Fourier :

$$\cos(\beta \sin \Omega_0 t) = J_0(\beta) + 2J_2(\beta) \cos 2\Omega_0 t + 2J_4(\beta) \cos 4\Omega_0 t + \dots + 2J_{2n}(\beta) \cos(2n\Omega_0 t) + \dots \quad (5.37)$$

$$\begin{aligned} \sin(\beta \sin \Omega_0 t) &= 2J_1(\beta) \sin \Omega_0 t + 2J_3(\beta) \sin 3\Omega_0 t + 2J_5(\beta) \sin 5\Omega_0 t \\ &\quad + \dots + 2J_{2n+1}(\beta) \sin((2n+1)\Omega_0 t) + \dots \end{aligned} \quad (5.38)$$

où $J_n(\beta)$ sont les fonctions de Bessel de premier ordre. On en retiendra principalement que le spectre d'un signal FM est un spectre de raies équiespacées, théoriquement de largeur infinie. Cependant, on peut définir la largeur de bande “pratique” d'un signal FM en la définissant comme étant la largeur de bande dans laquelle 99 % de la puissance est incluse.

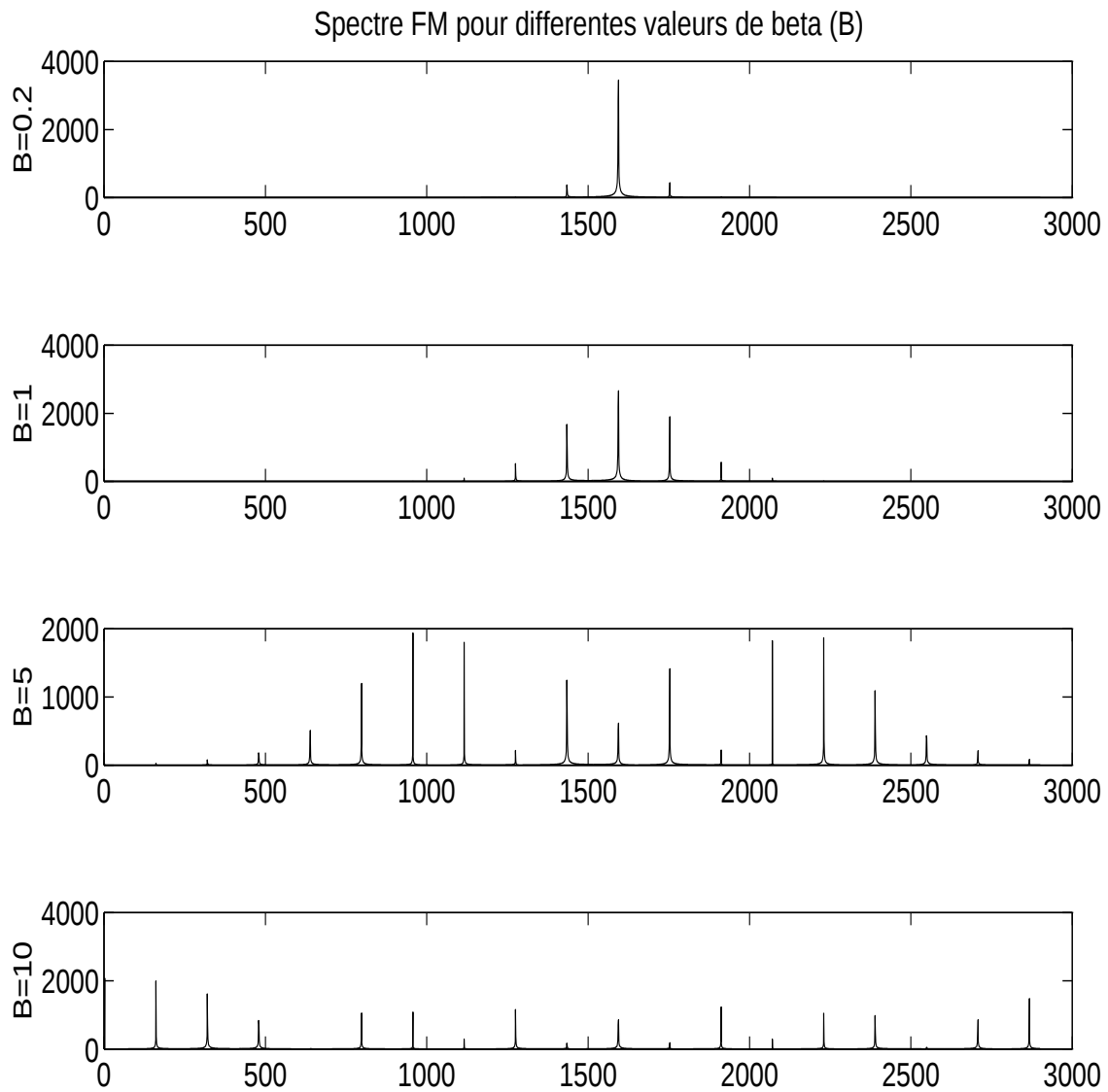
La **règle de Carson** nous donne cette valeur :

$$B = 2(\Delta f + F_0) \quad (5.39)$$

Dans le cas d'un signal $m(t)$ quelconque, F_0 représente la plus haute fréquence contenue dans ce signal.

5.3.3 Modulation de fréquence à faible indice

Si l'indice de modulation est faible ($\beta \leq 0.2$), on a une largeur de bande faible (on parle de *NBFM* : *Narrow Band Frequency Modulation* et, de plus, on obtient, de par les propriétés des fonctions de Bessel, un spectre du type (figure 5.15). Cette allure de spectre ressemble furieusement à un spectre de modulation DSB. Cependant, le diagramme phasorial permet de

FIG. 5.15 – Spectre d'un signal FM, influence de β

bien voir la différence entre les deux modulations (avec entre autres le fait que l'amplitude du signal FM est constante).

En effet, si β est petit, nous pouvons écrire :

$$\begin{aligned} s_{NBFM}(t) &= \cos(\omega_c t + \beta \sin \Omega_0 t) \\ &\approx \cos \omega_c t - \frac{\beta}{2} \cos(\omega_c - \Omega_0) + \frac{\beta}{2} \cos(\omega_c + \Omega_0) \end{aligned} \quad (5.40)$$

En adoptant un système de coordonnées tournant dans le sens trigonométrique à une vitesse angulaire ω_c , le phaseur de la porteuse est fixe et nous le plaçons en position horizontale. Dans le même système, le terme $\frac{\beta}{2} \cos(\omega_c + \Omega_0)$ tourne dans le sens trigonométrique à une vitesse Ω_0 tandis que l'autre terme tourne dans l'autre sens, avec la même vitesse. En faisant le même type de diagramme pour la DSB-SC, on peut se rendre compte que le spectre doit avoir la même allure, puisqu'on a à chaque fois un phaseur important représentant la porteuse et deux phaseurs tournant aux mêmes vitesses, au signe près. Par contre, on voit également que l'amplitude du signal NBFM ne varie pas tandis que l'amplitude du signal AM varie (heureusement!).

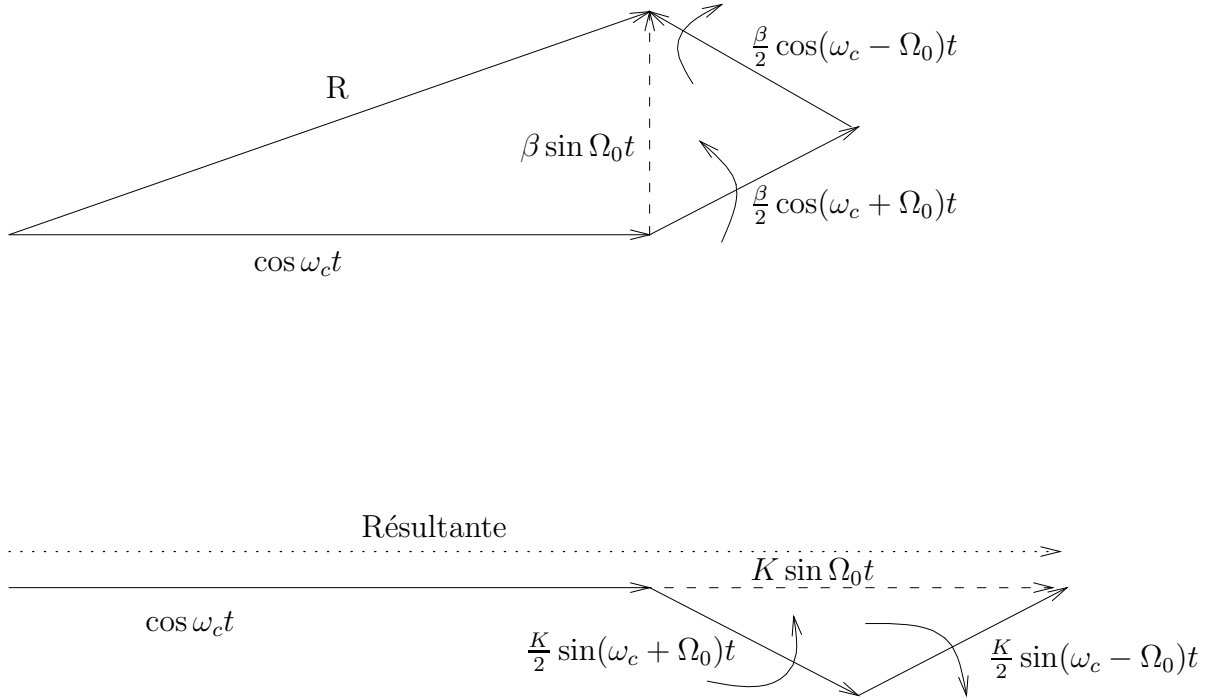
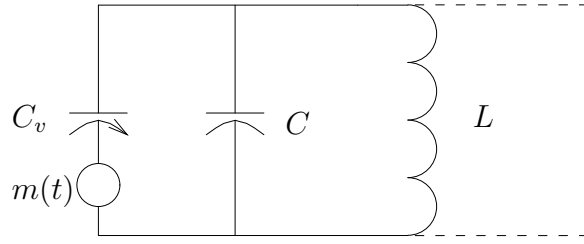


FIG. 5.16 – Comparaison des phaseurs pour la NBFM et la DSB

5.3.4 Modulation FM par variation de paramètres

Une manière simple de générer un signal FM, quoique peu praticable, est de faire varier la fréquence d'accord d'un oscillateur LC par variation de la capacité. On obtient le générateur de la figure 5.17.

Ce circuit permet d'obtenir une pulsation de sortie $\omega = (\sqrt{L(C + C_v)})^{-1}$. La *varicap* C_v est une capacité variable commandée par la tension et est souvent une diode polarisée en inverse dont la capacité varie en fonction de la tension inverse appliquée.

FIG. 5.17 – *Modulateur FM simple*

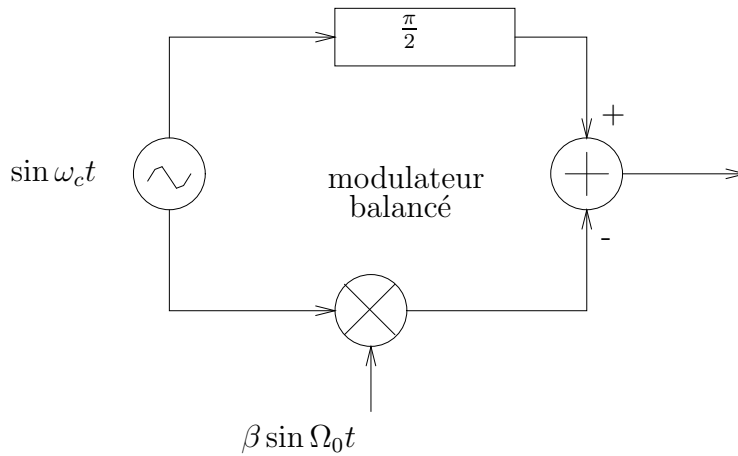
5.3.5 Modulateur d'Armstrong

Ayant vu que la NBFM a une expression relativement simple, il est naturel de vouloir adopter un circuit qui reproduit cette expression.

En effet, on a

$$\begin{aligned} s_{NBFM}(t) &= \cos(\omega_c t + \beta \sin \Omega_O t) \\ &\approx \cos \omega_c t + \sin \omega_c t \cdot \beta \sin \Omega_O t \end{aligned} \quad (5.41)$$

Ce qui se traduit directement par la figure 5.18.

FIG. 5.18 – *Modulateur d'Armstrong : principe*

Cependant, l'application directe de cette technique est inacceptable. Considérons un signal de parole, que l'on veut transmettre avec un indice de modulation maximal de $\beta \leq 0.5$ de manière à ce que l'approximation de l'équation 5.41 soit valable. L'oscillateur local (ω_c) est fixé à 200 kHz et on désire transmettre un signal de fréquence allant de 50 Hz à 15 kHz. L'indice de modulation valant $\beta = \frac{\Delta f}{F_0}$ pour une sinusoïde, il est maximal pour la fréquence minimale (ici 50 Hz) et nous amène à choisir une déviation de fréquence $\Delta f = 25 \text{ Hz}$.

D'autre part, pour des raisons de qualité de transmission (rapport Signal/Bruit à la réception), et d'occupation spectrale acceptable, on désire une largeur de bande de 180 kHz centré

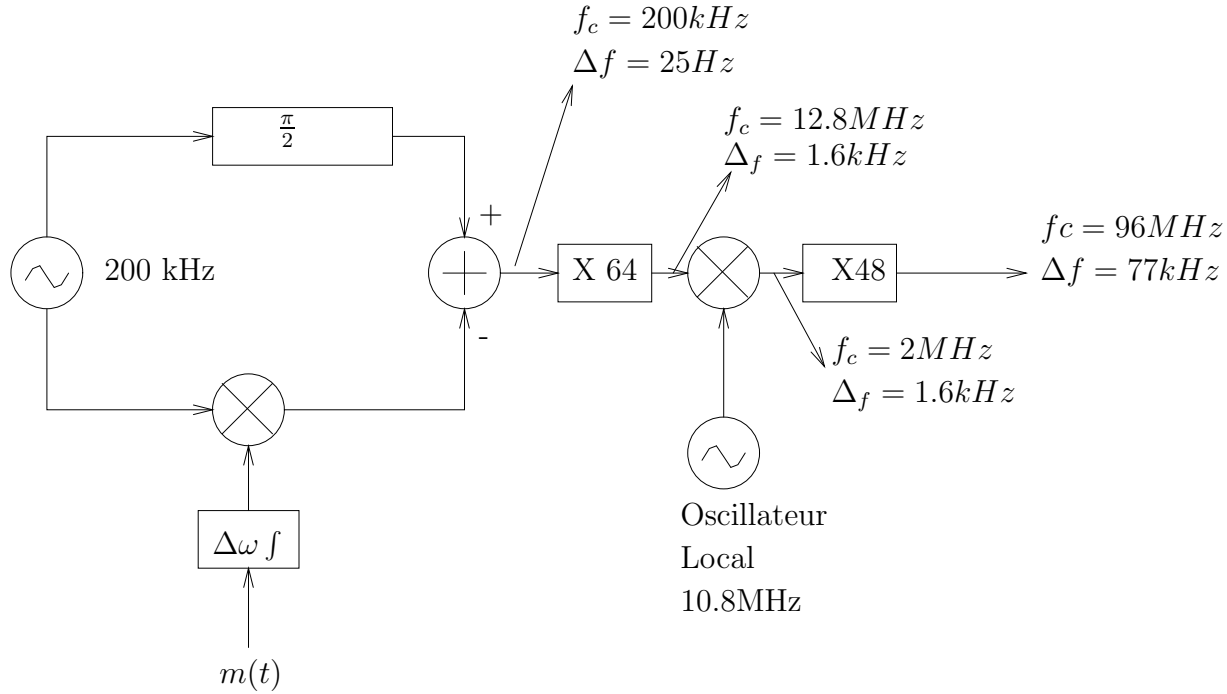


FIG. 5.19 – Modulateur FM à partir d'un modulateur d'Armstrong

aux alentours de 100 MHz. Par multiplication de fréquence, cela mènerait à une multiplication par 500 et donc un Δf de 12.5 kHz, ce qui est insuffisant. D'autre part, pour obtenir une largeur de bande de 180 kHz, avec une fréquence maximale de 15 kHz, la règle de Carson nous amène à un Δf de 75 kHz, ce qui, si on adoptait un multiplicateur de fréquence, nous amènerait à une porteuse à 1.2 GHz.

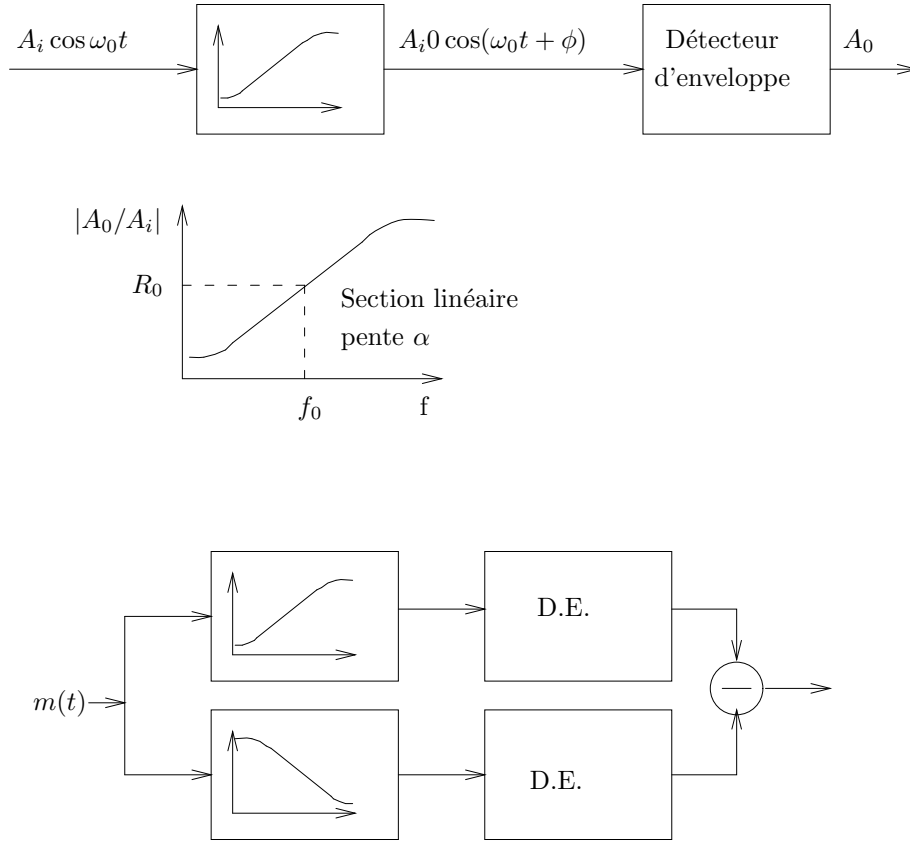
La solution à ce problème consiste à faire deux multiplications séparées par un décalage de fréquence. La figure 5.19 explicite les opérations suivies.

5.3.6 Démodulateur FM

Le principe d'un démodulateur simple peut être représenté par le haut de la figure 5.20. Un signal de fréquence f_0 est appliqué à une boîte noire qui a comme sortie une amplitude proportionnelle à la fréquence d'entrée. Cette sortie est appliquée à un détecteur d'enveloppe qui, si f_0 est constant, sort une amplitude A_0 proportionnelle à la fréquence d'entrée.

Si le signal d'entrée est un signal à fréquence variable, la sortie A_0 sera proportionnelle à la fréquence d'entrée, si celle-ci est contenue dans la bande passante de la "boîte noire".

Ce que l'on demande à un démodulateur FM, c'est de fournir une sortie A_0 qui soit proportionnelle à la déviation fréquence instantanée du signal, ce qui nous amène à limiter le plus possible l'amplitude d'entrée du système de manière à éviter une influence des modulations d'amplitudes parasites dues à des imperfections du système ou au canal. D'autre part, on désire une zone de linéarité dans la fonction de transfert "fréquence-tension" qui soit la plus grande possible. Pour ce faire, on utilisera deux filtres de fonction de transfert opposées (voir la figure) de manière à améliorer la linéarité de l'ensemble.

FIG. 5.20 – *Démodulateur FM*

En effet, si la fonction de transfert est du type

$$A_0 = R_0 A_i + \alpha A_i (f - f_0) + \beta A_i (f - f_0)^2 + \dots \quad (5.42)$$

où les termes en $(f - f_0)^i, i > 1$ rendent compte de la non-idéalité du filtre, en utilisant des filtres opposés, on a, pour le second filtre :

$$A_0 = R_0 A_i - \alpha A_i (f - f_0) + \beta A_i (f - f_0)^2 + \dots \quad (5.43)$$

Soit, en faisant la différence des deux :

$$A_0 = 2\alpha A_i (f - f_0) + \text{Termes en } (f - f_0)^i, i \text{ impair} \quad (5.44)$$

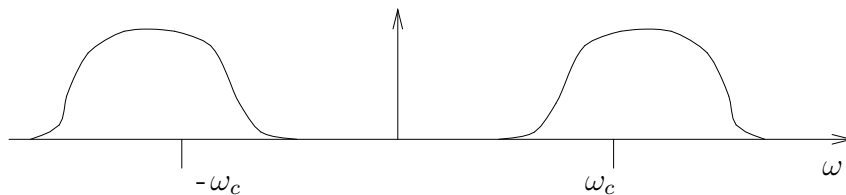
Ce qui nous donne une meilleure caractéristique de linéarité, le terme d'ordre 2 étant éliminé et les termes ultérieurs étant à priori de plus faible importance.

Chapter 6

Performances des systèmes de modulation

6.1 Notion de signal équivalent passe-bas

Soit un signal ayant l'allure spectrale suivante :



On peut imaginer que ce signal est la résultante d'une modulation d'amplitude du type :

$$s(t) = m_1(t) \cos \omega_c t \text{ ou } s(t) = m_2(t) \sin \omega_c t \quad (6.1)$$

soit, en combinant les deux :

$$s(t) = m_1(t) \cos \omega_c t + m_2(t) \sin \omega_c t \quad (6.2)$$

où $m_1(t)$ et $m_2(t)$ sont des signaux passe-bas.

Cette vision des choses nous donne l'intuition qu'une telle décomposition pour tout signal passe-bas (c'est-à-dire si son spectre ne s'étend ni jusqu'à la fréquence nulle, ni jusqu'à la fréquence infinie).

6.1.1 Signal analytique

Le *signal analytique* de $x(t)$, noté $x_a(t)$ est défini par :

$$X_a(\omega) = \begin{cases} 2X(\omega) & \text{si } \omega > 0 \\ 0 & \text{si } \omega < 0 \end{cases} \quad (6.3)$$

Le signal analytique est donc la sortie d'un filtre ne laissant passer que les fréquences positives. De ce fait, sa transformée de Fourier n'est pas paire et $x_a(t)$ est donc un signal complexe.

La fonction de transfert du filtre qui génère le signal analytique peut s'écrire :

$$H(\omega) = 1 + j[-j\text{sgn}(\omega)] = H_1(\omega) + jH_2(\omega) \quad (6.4)$$

avec $H_1(\omega) = 1$ et $H_2(\omega) = -j\text{sgn}(\omega)$, ce qui nous permet d'exprimer le signal analytique :

$$x_a(t) = x(t) + j\hat{x}(t) \quad (6.5)$$

où $\hat{x}(t)$ est la transformée de Hilbert de $x(t)$.

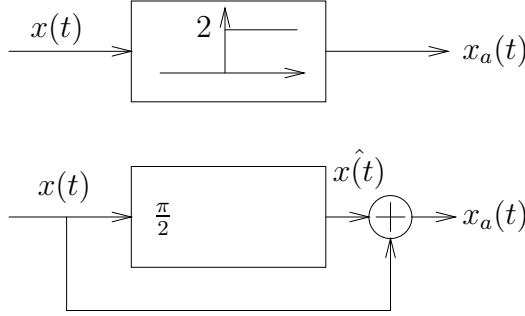


FIG. 6.1 – *Signal analytique*

6.1.2 Signal équivalent passe-bas

Soit $x(t)$ un signal passe-bande, on appelle signal équivalent passe-bas¹ :

$$\begin{aligned} u(t) &= x_a(t)e^{-j\omega_c t} \\ \text{soit } U(\omega) &= X_a(\omega - \omega_c) \end{aligned} \quad (6.6)$$

Le signal $x(t)$ étant la partie réelle de $x_a(t)$, on peut écrire :

$$x(t) = \Re[x_a(t)] = \Re[u(t)e^{j\omega_c t}] \quad (6.7)$$

Ce qui permet de confirmer l'intuition selon laquelle on peut exprimer :

$$\begin{aligned} x(t) &= \Re[u(t)] \cos \omega_c t - \Im[u(t)] \sin \omega_c t \\ &= m_1(t) \cos \omega_c t - m_2(t) \sin \omega_c t \\ &= a(t)e^{-j\omega_c t + \phi(t)} \end{aligned} \quad (6.8)$$

Tout signal passe-bande peut donc être considéré soit comme deux modulations d'amplitudes en quadrature, soit comme une modulation combinée d'amplitude et de fréquence.

6.1.3 Filtre équivalent passe-bas

Soit un filtre passe-bande $H(\omega)$ centré sur ω_c , on définit :

$$F(\omega - \omega_c) = \begin{cases} H(\omega) & \omega > 0 \\ 0 & \omega < 0 \end{cases} \quad (6.9)$$

Sachant que le filtre est un système de réponse impulsionnelle $h(t)$ réelle, on a que sa transformée de Fourier a la propriété $H(\omega) = H^*(-\omega)$ et $H(\omega)$ a l'expression :

1. encore appelé enveloppe complexe

$$H(\omega) = H^*(-\omega) \quad (6.10)$$

ce qui se traduit par la figure :

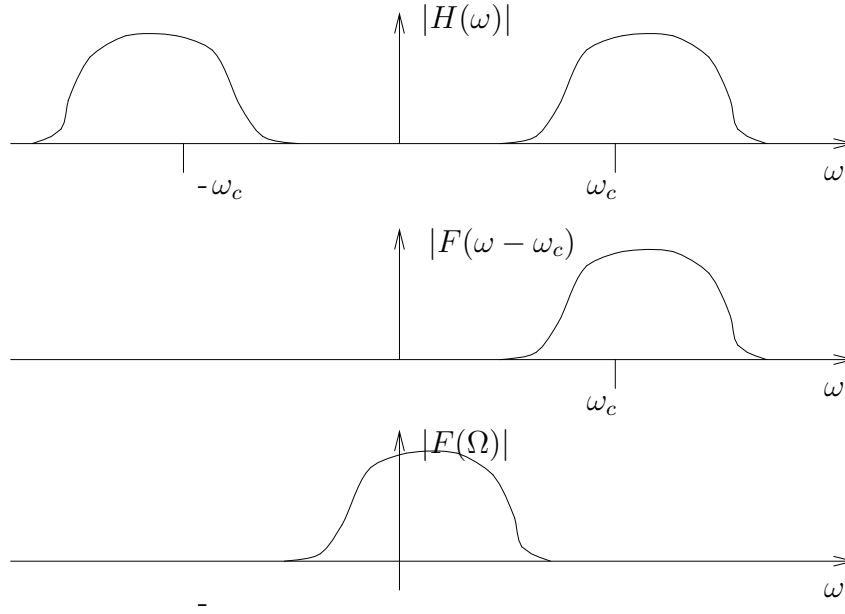


FIG. 6.2 – *Filtre équivalent passe-bas*

6.1.4 Réponse d'un filtre passe-bande

Il est clair que les développements précédents n'auraient aucun sens si l'opération de filtrage ne pouvait se faire de façon équivalente en basses fréquences (en équivalent passe-bas) et en bande transposée.

Le petit développement ci-dessous, en fréquentiel, démontre cette équivalence. On notera les "basses fréquences" $\Omega = \omega - \omega_c$

Ce qui donne :

$$\begin{aligned}
 F(\Omega)U(\omega) &= F(\Omega).X_a(\omega - \omega_c) \\
 &= F(\Omega).X(\omega - \omega_c)[1 + \text{sgn}(\omega)] \\
 &= \frac{F_a(\omega - \omega_c)}{2}.X(\omega - \omega_c)[1 + \text{sgn}(\omega)] \\
 &= \frac{H(\omega - \omega_c)}{2}[1 + \text{sgn}(\omega)].X(\omega - \omega_c)[1 + \text{sgn}(\omega)] \\
 &= H(\omega - \omega_c).X(\omega - \omega_c)[1 + \text{sgn}(\omega)]
 \end{aligned} \quad (6.11)$$

La dernière expression nous indiquant que $F(\Omega)U(\omega)$ est le signal équivalent passe-bas à $H(\omega).X(\omega)$.

Outre l'élégance théorique de ces développements, ils permettent d'une part d'avoir un outil permettant de ramener toutes les études en bande de base et d'autre part, d'un point de vue simulation, de simuler en bande de base également, ce qui permet de faire un échantillonnage (pour les ordinateurs, il faut bien échantillonner) à une fréquence acceptable.

6.2 Densité spectrale de puissance

6.2.1 Cas des signaux certains

Dans le cas des signaux certains, du théorème de Parseval exprimant l'énergie :

$$E = \int_{-\infty}^{\infty} |x(t)|^2 dt = \int_{-\infty}^{\infty} |X(f)|^2 df \quad (6.12)$$

on peut interpréter $|X(f)|^2$ comme étant une densité spectrale d'énergie.

La puissance étant définie par (E_T représentant l'énergie d'un signal de durée T) :

$$\begin{aligned} P &= \lim_{T \rightarrow \infty} \frac{E_T}{T} \\ &= \lim_{T \rightarrow \infty} \frac{1}{T} \int_{-\frac{T}{2}}^{\frac{T}{2}} |x(t)|^2 dt \\ &= \lim_{T \rightarrow \infty} \frac{1}{T} \int_{-\infty}^{\infty} |X(f)|^2 df \end{aligned} \quad (6.13)$$

et $\lim_{T \rightarrow \infty} \frac{1}{T} |X(f)|^2$ peut être interprété comme étant la densité spectrale de puissance.

Ainsi, en définissant l'autocorrélation du signal certain $x(t)$ par

$$R(\tau) = \lim_{T \rightarrow \infty} \frac{1}{T} \int_{-\frac{T}{2}}^{\frac{T}{2}} x(t)x(t+\tau) dt \quad (6.14)$$

et sa transformée de Fourier $S(f)$, en utilisant la transformée de Fourier inverse :

$$\text{TF}^{-1} \left[\lim_{T \rightarrow \infty} \frac{1}{T} |X(f)|^2 \right]$$

$$\int_{-\infty}^{\infty} \left[\lim_{T \rightarrow \infty} \frac{1}{T} |X(f)|^2 \right] df$$

avec (Fourier inverse) $X_T(f) =$

$$\begin{aligned} &= \lim_{T \rightarrow \infty} \frac{1}{T} \int_{-\infty}^{\infty} df \int_{-\frac{T}{2}}^{\frac{T}{2}} x(t) e^{-j\omega t} dt \int_{-\frac{T}{2}}^{\frac{T}{2}} x(t') e^{-j\omega t'} e^{j\omega \tau} dt' \\ &= \lim_{T \rightarrow \infty} \frac{1}{T} \int_{-\frac{T}{2}}^{\frac{T}{2}} x(t) dt \underbrace{\int_{-\frac{T}{2}}^{\frac{T}{2}} x(t') dt' \int_{-\infty}^{\infty} e^{j\omega(t'-t+\tau)} df}_{\substack{= \delta(t'-t+\tau) \\ = x(t-\tau)}} \\ &= \lim_{T \rightarrow \infty} \frac{1}{T} \int_{-\frac{T}{2}}^{\frac{T}{2}} x(t)x(t-\tau) dt \\ &= R(-\tau) = R(\tau) \end{aligned} \quad (6.15)$$

On a donc le résultat très important que la densité spectrale de puissance est la transformée de Fourier de la fonction d'autocorrélation.

$$S(f) = \text{TF}[R(\tau)] \quad (6.16)$$

6.2.2 Cas des signaux aléatoires

Dans le cas d'un signal aléatoire faiblement stationnaire et ergodique relativement à la fonction de corrélation (voir le chapitre concernant les fonctions aléatoires) il y a égalité entre la fonction de corrélation temporelle (du signal certain qu'est une réalisation d'un signal aléatoire) et la fonction de corrélation au sens stochastique (espérance mathématique ...).

On peut donc écrire :

$$R_{xx}(\tau) = \int_{-\infty}^{\infty} S_{xx}(f) e^{j\omega\tau} df \quad (6.17)$$

et donc la puissance vaut :

$$m_x^2 + \sigma_x^2 = \int_{-\infty}^{\infty} S(f) df \quad (6.18)$$

et, de même que pour les signaux certains,

$$S_{xx}(f) = \text{TF}[R_{xx}(\tau)]$$

On peut également définir, si on considère deux signaux $x(t)$ et $y(t)$, la densité spectrale de puissance mutuelle.

$$S_{yx}(\omega) = \text{TF}[R_{yx}(\tau)] = \text{TF}[E\{y(t)x(t+\tau)\}] \quad (6.20)$$

6.2.3 Théorème de Wiener-Kintchine

Si $y(t)$ est la réponse d'un système linéaire, permanent et stable de fonction de transfert $H(\omega)$, à une excitation $x(t)$, signal aléatoire faiblement stationnaire, on a

$$S_{yy}(\omega) = |H(\omega)|^2 S_{xx}(\omega) \quad (6.21)$$

$$S_{yx}(\omega) = H(\omega) S_{xx}(\omega) \quad (6.22)$$

6.3 Représentation du bruit blanc

6.3.1 Bruit blanc passe-bas

Un bruit blanc $n(t)$ est défini par :

$$\begin{aligned} R_{nn}(\tau) &= \frac{N_O}{2} \delta(t - \tau) \\ \text{et donc } S_{nn}(f) &= \frac{N_O}{2} \end{aligned} \quad (6.23)$$

où N_0 est appelée densité spectrale unilatérale de bruit. En effet, si on utilise une représentation unilatérale du bruit, on aura $2 \cdot \frac{N_O}{2}$ comme densité.

On s'intéresse au bruit blanc filtré par un filtre passe-bas. On a

$$S_{nn}(f) = \begin{cases} \frac{N_O}{2} & |f| \leq B \\ 0 & \text{ailleurs} \end{cases} \quad (6.24)$$

Et on trouve aisément, par la TF inverse :

$$R_{nn}(\tau) = BN_0 \frac{\sin 2B\tau}{2B\tau} = BN_0 \text{sinc} 2B\tau \quad (6.25)$$

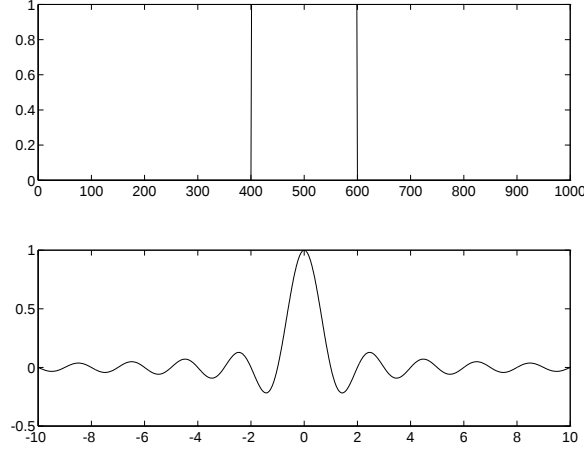


FIG. 6.3 – Spectre et autocorrélation d'un bruit blanc filtré passe-bas

6.3.2 Bruit blanc filtré passe-bande

Dans le cas des modulations, nous utiliserons un filtre en entrée qui limitera le bruit blanc à la bande passante utile. Celui-ci aura donc l'allure spectrale :

$$S_{nn}(f) = \begin{cases} ds \frac{N_0}{2} & |f| \omega_c - \Omega < \omega < \omega_c + \Omega \\ ds \frac{N_0}{2} & |f| - \omega_c - \Omega < \omega < -\omega_c + \Omega \\ 0 & \text{ailleurs} \end{cases} \quad (6.26)$$

On peut aisément déduire :

$$R_{nn}(\tau) = 2FN_0 \frac{\sin \Omega\tau}{\omega\tau} \cos \omega_c\tau \quad (6.27)$$

Représentation en quadrature

On peut également représenter le bruit $n(t)$ passe-bande par ses deux composantes en quadrature, comme indiqué en début de chapitre.

$$n(t) = n_c(t) \cos \omega_c t - n_s(t) \sin \omega_c t \quad (6.28)$$

L'avantage principal de cette représentation est le suivant : prenons un signal AM bruité :

$$A[1 + Km(t)] \cos \omega_c t + n(t) \quad (6.29)$$

A la réception, le signal est filtré, et le bruit $n(t)$ est donc bien un bruit filtré passe-bande. On peut donc exprimer le signal AM bruité par :

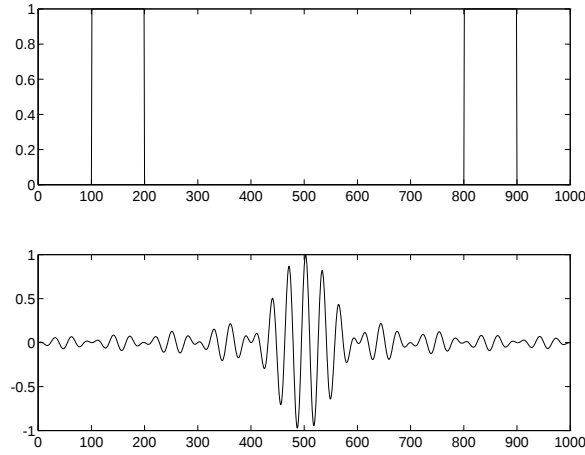


FIG. 6.4 – Spectre et autocorrélation d'un bruit blanc filtré passe-bande

$$\underbrace{A[1 + Km(t)] + n_c(t)}_{A_1} \cos \omega_c t + \underbrace{n_s(t)}_{A_2} \sin \omega_c t \quad (6.30)$$

Le détecteur d'enveloppe nous fournira une tension $\sqrt{A_1^2 + A_2^2}$ où le second terme est très faible si on a un rapport signal/bruit ($0.5A_1^2/\sigma_n^2$) en entrée élevé. La sortie du détecteur d'enveloppe sera alors approximativement

$$A[1 + Km(t)] + n_c(t) \quad (6.31)$$

et donc le signal est principalement perturbé par la composante $n_c(t)$ en phase avec le signal. On peut démontrer les caractéristiques des composantes de bruit suivantes :

$$E\{n_c^2(t)\} = E\{n_s^2(t)\} = E\{n^2(t)\} \quad (6.32)$$

Le bruit blanc étant stationnaire (N_0 est constant) on a

$$E\{n_c^2(t)\} = \sigma_{nc}^2 = \dots \quad (6.33)$$

6.3.3 Représentation Equivalent passe-bas de signaux stochastiques stationnaires

Soit un signal

$$n(t) = x(t) \cos \omega_c t - y(t) \sin \omega_c t \quad (6.34)$$

$$= \Re[z(t)e^{j\omega_c t}] \quad (6.35)$$

où $n(t)$ est à moyenne nulle et stationnaire.

– $n(t)$ étant à moyenne nulle, on déduit immédiatement que $n_c(t)$ et $n_s(t)$ sont à moyenne nulle.

– De la stationnarité de $n(t)$

$$E \{ (n(t)n(t+\tau)) \} = E \{ [x(t) \cos \omega_c t - y(t) \sin \omega_c t] \quad (6.36)$$

$$[x(t+\tau) \cos \omega_c(t+\tau) - y(t+\tau) \sin \omega_c(t+\tau)] \} \quad (6.37)$$

$$= R_{xx}(\tau) \cos \omega_c t \cos \omega_c(t+\tau) + R_{yy}(\tau) \sin \omega_c t \sin \omega_c(t+\tau) \quad (6.38)$$

$$- R_{xy}(\tau) \sin \omega_c t \cos \omega_c(t+\tau) - R_{yx}(\tau) \cos \omega_c t \sin \omega_c(t+\tau) \quad (6.39)$$

$\Rightarrow$

$$E \{ (n(t)n(t \pm \tau)) \} = \frac{1}{2} [R_{xy}(\tau) + R_{yy}(\tau)] \cos \omega_c \tau \quad (6.40)$$

$$+ \frac{1}{2} [R_{xx}(\tau) - R_{yy}(\tau)] \cos \omega_c(2t + \tau) \quad (6.41)$$

$$- \frac{1}{2} [R_{yx}(\tau) - R_{xy}(\tau)] \sin \omega_c \tau \quad (6.42)$$

$$- \frac{1}{2} [R_{yx}(\tau) + R_{xy}(\tau)] \sin \omega_c(2t + \tau) \quad (6.43)$$

$n(t)$ étant stationnaire, $E \{ (n(t)n(t+\tau)) \} = R_{nn}(\tau)$ ne peut dépendre de t et les termes en $\cos \omega_c(2t + \tau)$ et $\sin \omega_c(2t + \tau)$ doivent être nuls

$$\Rightarrow R_{xx}(\tau) = R_{yy}(\tau) \quad (6.44)$$

$$R_{yx}(\tau) = -R_{xy}(\tau) \quad (6.45)$$

et donc

$$R_{nn}(\tau) = R_{xx}(\tau) \cos \omega_c \tau - R_{yx}(\tau) \sin \omega_c \tau \quad (6.46)$$

En définissant la fonction d'autocorrélation du signal équivalent passe-bas $z(t) = x(t) + jy(t)$ par

$$R_{zz}(\tau) = \frac{1}{2} E \{ z(t) z^*(t+\tau) \} \quad (6.47)$$

$$= \frac{1}{2} E \{ (x(t) + jy(t))(x(t+\tau) - jy(t+\tau)) \} \quad (6.48)$$

$$= \frac{1}{2} [R_{xx}(\tau) + R_{yy}(\tau) - jR_{xy}(\tau) + jR_{yx}(\tau)] \quad (6.49)$$

avec les relations (6.44), on obtient :

$$R_{zz}(\tau) = R_{xx}(\tau) + jR_{yx}(\tau) \quad (6.50)$$

et

$$R_{nn}(\tau) = \text{Re}[R_{zz}(\tau)e^{j\omega_c \tau}] \quad (6.51)$$

que l'on peut exprimer en disant que la fonction d'autocorrélation du signal équivalent passe-bas est l'équivalent passe-bas de la fonction d'autocorrélation du signal passe-bande.

En outre, la densité spectrale de puissance étant la transformée de Fourier de la fonction d'autocorrélation, on a :

$$S_{nn}(f) = \int_{-\infty}^{+\infty} \Re[R_{zz}(\tau)e^{j\omega_c\tau}]e^{-j\omega\tau}d\tau \quad (6.52)$$

$$= \frac{1}{2}[R_{zz}(f - f_c) + R_{zz}(-f - f_c)] \quad (6.53)$$

D'autre part, $E\{z(t)z^*(t + \tau)\} = E\{z(t - \tau)z^*(t)\}^*$, c'est-à-dire $R_{zz}(\tau) = R_{zz}^*(-\tau) \rightarrow R_{zz}(f)$ est une fonction à valeurs réelles.

6.3.4 Propriétés des composantes en quadrature

Toute fonction de corrélation satisfait la condition :

$$R_{yx}(\tau) = R_{xy}(-\tau) \quad (6.54)$$

des relations (6.44), on déduit

$$R_{yx}(\tau) = -R_{xy}(-\tau) \quad (6.55)$$

et donc $R_{yx}(0) = 0$

c'est-à-dire que $x(t)$ et $y(t)$, étant à moyenne nulle, sont décorrélés ($R_{xy}(0) = 0$), ce qui implique $S_{zz}(f) = S_{zz}(-f)$.

Dans le cas du bruit blanc passe-bande, le signal $z(t)$ équivalent passe-bas a les caractéristiques :

$$S_{zz}(f) = \begin{cases} N_o & |\omega| \leq \Omega \\ 0 & |\omega| > \Omega \end{cases} \quad (6.56)$$

et donc $R_{zz}(\tau) = N_o \frac{\sin \frac{\Omega\tau}{2}}{\frac{\Omega\tau}{2}}$

on en déduit $R_{yz}(\tau) = 0$

d'où

$$\sigma_z^2 = R_{zz}(0) = R_{xx}(0) = R_{yy}(0) = R_{nn}(0) \quad (6.57)$$

$$= \sigma_{nc}^2 = \sigma_{ns}^2 = \sigma_n^2 \quad (6.58)$$

6.3.5 Densités spectrales

de $R_{nn}(\tau) = R_{nc}(\tau) \cos \omega_o \tau$

et du théorème de modulation, on en déduit

$$S_{nn}(f) = \frac{1}{2}S_{nc}(f - f_c) + \frac{1}{2}S_{ns}(-f - f_c) \quad (6.59)$$

et donc, si

$$S_{nn}(f) = \begin{cases} \eta & |f_c - \frac{F}{2}| |\omega| < (\frac{F}{2} + F_c) \\ 0 & \text{ailleurs} \end{cases} \quad (6.60)$$

on a

$$S_{nc}(f) = S_{ns}(f) = \begin{cases} 2\eta & |f| < F \\ 0 & \text{ailleurs} \end{cases} \quad (6.61)$$

6.4 Rapport signal/bruit en AM

En utilisant la représentation des composantes en quadrature du bruit, on a le signal AM à la réception :

$$S_{AM}(t) = A[1 + Km(t)] \cos \omega_c t + n_c(t) \cos \omega_c t + n_s(t) \sin \omega_c t \quad (6.62)$$

$$= \underbrace{A[1 + Km(t) + n_c(t)]}_{A_1} \cos \omega_c t + \underbrace{n_s(t)}_{A_2} \sin \omega_c t \quad (6.63)$$

Le détecteur d'enveloppe nous fournit $\sqrt{A_1^2 + A_2^2}$ où $A_2^2 \ll A_1^2$ dans l'hypothèse d'un rapport signal/bruit en entrée élevé. Nous pouvons donc approximer la sortie du détecteur par

$$s_{DE}(t) \simeq A[1 + Km(t)] + n_c(t) \quad (6.64)$$

La puissance de sortie peut être aisément calculée, d'où on aura enlevé la composante continue (par un condensateur en série par exemple)

$$E\{m(t)\} = 0 \quad (6.65)$$

$$E\{A^2 K^2 n^2(t)\} = A^2 K^2 E\{m^2(t)\} = A^2 K^2 \sigma_m^2 \quad (6.66)$$

$$\begin{aligned} \rightarrow S_0 &= A^2 K^2 \sigma_n^2 \\ N_0 &= \sigma_{nc}^2 = 2\eta F \end{aligned} \quad (6.67)$$

$$\rightarrow \left(\frac{S}{N}\right)_o = \frac{A^2 K^2 \sigma_m^2}{2\eta F} \quad (6.68)$$

En entrée

$$S_i = \frac{A^2}{2} E\{[1 + Km(t)]^2\} = \frac{A^2(1 + K^2 \sigma_m^2)}{2} \quad (6.69)$$

$$N = En^2(t) = En_c^2(t) = 2\eta F \quad (6.70)$$

$$\rightarrow \left(\frac{S}{N}\right)_i = \frac{A^2}{2} \frac{1 + K^2 \sigma_m^2}{2\eta F} \quad (6.71)$$

$$\rightarrow \left(\frac{S}{N}\right)_o = \frac{2K^2 \sigma_m^2}{1 + K^2 \sigma_m^2} \left(\frac{S}{N}\right)_i = \frac{K^2 \sigma_m^2}{1 + K^2 \sigma_m^2} \frac{S_i}{2\eta F} \quad (6.72)$$

Comme on a, en AM, $|Km(t)| < 1$, la nature de la parole nous donne $K^2 \sigma_m^2 \simeq 0.1$ et donc $\left(\frac{S}{N}\right)_o \simeq 0.2 \left(\frac{S}{N}\right)_i$, soit une perte de qualité de 7 dB en terme de rapport signal/bruit et donc, si on désire un signal de sortie de bonne qualité, une puissance d'entrée élevée.

Dans le cas de la DSB-SC,

$$s_{DSB-SC}(t) = Am(t) \cos \omega_c t + n_c(t) \cos \omega_c t + n_s(t) \sin \omega_c t \quad (6.73)$$

soit, après démodulation synchrone et filtrage

$$s_o(t) = Am(t) + n_c(t) \quad (6.74)$$

on obtient immédiatement

$$\left(\frac{S}{N} \right)_o = \frac{S_i}{2\eta F} \quad (6.75)$$

Pour la SSB, les puissance du signal ainsi que du bruit sont simplement divisés par deux, en entrée comme en sortie, nous avons donc la même relation que pour la DSB-SC.

6.5 Rapport signal/bruit en FM

6.5.1 Signal d'entrée

Nous avons un signal

$$s_{FM}(t) = A \sin[\omega_c t + \Delta\omega \int m(\tau) d\tau] \quad (6.76)$$

soit $S_i = \frac{A^2}{2}$ où nous faisons l'hypothèse que l'amplitude A est constante en fonction du temps.

6.5.2 Signal de sortie

A la réception

$$s_o(t) = A + \sin[\omega_c t + \Delta\omega \int m(t) dt] + n'_c(t) \cos \omega_c t - n_s(t) \sin \omega_c t \quad (6.77)$$

Soit, après dérivation

$$s_1(t) = A[\omega_c + \Delta\omega m(t)] \cos[\omega_c t + \Delta\omega \int m(t) dt] - n'_c(t) \omega_c \sin \omega_c t - n'_s(t) \omega_c \cos \omega_c t \quad (6.78)$$

En faisant l'hypothèse que ce bruit n'intervient pas dans la puissance du signal, on a, après détection d'enveloppe $s_2(t) = A[\omega_c + \Delta\omega m(t)]$ et donc, en éliminant la composante continue

$$S_o = E \left\{ (A\Delta\omega m(t))^2 \right\} = A^2 \Delta\omega^2 \sigma_m^2 \quad (6.79)$$

6.5.3 Bruit de sortie

Après détection d'enveloppe de $s_1(t)$, en tenant compte du bruit, on obtient, en faisant l'hypothèse $m(t) = 0$

$$s_3(t) = \omega_c \sqrt{[A - n'_s(t)]^2 + n'^2_c(t)} \quad (6.80)$$

$$\simeq \omega_c \sqrt{[A - n'_s(t)]^2} = \omega_c [A - n'_s(t)] \quad (6.81)$$

Dans cette expression, nous devons comparer $n'_s(t)$ à A , le facteur multiplicatif ω_c affectant les deux parties (signal et bruit), il ne doit pas être pris en compte. La composante continue étant éliminée, on a, au signe près

$$s_3(t) \simeq n'_s(t) \quad (6.82)$$

$$\Rightarrow N_{out} = R_{n'_s}(\tau = 0) = \frac{1}{2\pi} \int S_{n'_s}(\omega) \delta\omega \quad (6.83)$$

La dérivée temporelle correspond à un filtre de transmittance $j\omega$.

La dérivée de $n'_s(t)$ a donc comme densité spectrale de puissance

$$dsp(n'_s) = dsp(n(s))|j\omega|^2 = 2\eta|j\omega|^2 = 2\eta\omega^2 \quad (6.84)$$

$$\begin{aligned} \Rightarrow \left(\frac{S}{N}\right)_0 &= \frac{A^2 \Delta\omega^2 \sigma_m^2}{2\eta} \left(\frac{1}{2\pi} \int_{-\Omega}^{+\Omega} \omega^2 \Delta\omega \right)^{-1} \\ &= 3\pi \frac{A^2 \Delta\omega^2 \sigma_m^2}{2\eta \Omega^3} \\ &= 3\pi \frac{A^2 \sigma_m^2}{2\eta} \left(\frac{\Delta\omega}{\Omega} \right)^2 \cdot \frac{1}{\Omega} \end{aligned} \quad (6.85)$$

en tenant compte de $S_i = \frac{A^2}{2}$

$$\left(\frac{S}{N}\right)_0 = 3\sigma_m^2 \left(\frac{\Delta f}{F}\right)^2 \frac{S_i}{2\eta F} \quad (6.86)$$

6.5.4 Le facteur de mérite $\frac{S_i}{2\eta F}$

Cette quantité est ainsi appelée parce qu'elle contient trois paramètres importants qui caractérisent la transmission, à savoir la puissance du signal en entrée, la largeur de bande du message et la densité spectrale de puissance du bruit.

Il est donc logique de comparer les rapports signal/bruit par rapport à cette quantité, le paramètre restant étant la bande passante utilisée.

6.5.5 Comparaison des modulations

Notons, pour la facilité, $\gamma = \frac{(S/N)_0}{S_i/2\eta F}$, on a alors, pour un $m(t)$ quelconque ($\sigma_m^2 = 0.1$) et $\Delta f = 75kHz$ pour la FM (largeur de Bande = $2(75kHz + 15kHz) = 180kHz$)

$$\begin{aligned} \gamma_{DSB-SC} &= 1 = 0\text{dB} \\ \gamma_{DSB} &\simeq 0.1 = -10\text{dB} \\ \gamma_{FM} &\simeq 7.5 = +9\text{dB} \end{aligned} \quad (6.87)$$

Le passage à la FM permet donc de gagner près de 20 dB par rapport à l'AM classique.

FIG. 6.5 – *diagramme phasoriel*FIG. 6.6 – *trajectoire*

6.5.6 Effet de seuil en FM

On peut représenter le bruit dans le diagramme phasoriel, en faisant l'hypothèse $m(t) = 0$

Dans le cas où $n_s(t)$ devient plus grand que A et d'orientation opposée, et que $n_c(t)$ change de signe, le vecteur résultant effectue la trajectoire indiquée ci-dessous.

Ce qui, si on observe la tension de bruit résultante, donne une impulsion du type

Ce type de bruit, appelé "anormal" (car non gaussien) n'a pas été pris en compte dans notre analyse précédente et a une probabilité d'apparition d'autant plus grande que le rapport signal/bruit d'entrée est faible. On observera alors un seuil en-deçà duquel la loi $\frac{(\frac{S}{N})_0}{S_i/2\eta F}$ linéaire n'est plus valable. La figure(*) représente cet "effet de seuil", ainsi que l'influence de $(\beta = \frac{\Delta f}{F})$ sur la caractéristique de bruit.

6.6 Préaccentuation et désaccentuation en FM

Nous avons, après détection, une densité spectrale de puissance du bruit parabolique:

D'autre part, la nature du signal (parole + musique) fait que le contenu spectral est faible aux fréquences élevées. La combinaison de ces deux faits provoque une qualité de transmission médiocre en "hautes" fréquences. Il semble donc logique de préaccentuer les fréquences élevées

FIG. 6.7 – *Impulsion résultante*

FIG. 6.8 – *bruit*FIG. 6.9 – *Schéma général*

avant émission, à charge du récepteur de désaccentuer de manière à récupérer correctement le message.

En considérant le schéma général

où $N(t)$ est le bruit de sortie du discriminateur et $N_F(t)$ le bruit à la sortie du désaccentuateur, on a

$$S_{NF}(\omega) = \frac{S_N(\omega)}{|H^2(\omega)|} \quad (6.88)$$

La situation idéale pour nous sera celle d'un bruit qui affecterait toutes les composantes de bruit de manière uniforme, c'est-à-dire un bruit blanc. Il est donc logique d'opter pour $H_p(\omega) = j\omega \Rightarrow H_d(\omega) = \frac{1}{j\omega} \Rightarrow S_{NF}(\omega) = \frac{\eta\omega^2}{\omega^2} = \eta$

En effet, il s'agit du bruit après détection FM, filtré par le filtre de désaccentuation ($H_d(\omega)$)

Préaccentuation en FM commerciale

On adopte les filtres RC suivants:

soit

FIG. 6.10 – *Filtre RC*

FIG. 6.11 – *transmittance*FIG. 6.12 – *puissance*

$$H_d(f) = \frac{1}{1 + j\frac{f}{f_1}} \quad (6.89)$$

où

$$f_1 = \frac{1}{2\pi RC} \quad (6.90)$$

et

$$H_p(f) = \frac{r}{R}(1 + j\omega RC) = \frac{r}{R}(1 + j\frac{f}{f_1}) \quad (6.91)$$

Ce qui donne les allures de transmittance

où $f_2 = \frac{1}{2\pi rC}$ est choisi en-dehors de la bande utile, de manière à n'avoir pas d'influence sur le signal.

Pour travailler dans des conditions similaires il faut que la puissance du signal ne soit pas modifiée par les différents étages de filtrage.

On doit donc avoir la situation suivante du point de vue puissance du signal.

c'est-à-dire

$$P_m = \sigma^2 m = \int_{-F}^{+F} |H_p(f)|^2 S_m(f) df = \int_{-F}^{+F} S_m(f) df \quad (6.92)$$

On doit donc introduire un facteur de proportionnalité pour tenir compte de cette contrainte. Dans le cas qui nous occupe, on peut modéliser le signal par

$$S_m(f) = \begin{cases} S \frac{1}{(1+\frac{f}{f_1})^2} & |f| \leq F \\ 0 & \text{ailleurs} \end{cases} \quad (6.93)$$

et

$$H_p(f) = k(1 + j \frac{f}{f_1}) \quad (6.94)$$

Pour la simplicité, on adopte $f' = f_1$
on obtient alors aisément l'équation

$$P_m = \sigma_m^2 = \int_{-F}^{+F} \frac{S df}{1 + (f/f_1)^2} = \int_{-F}^{+F} K^2 S df \quad (6.95)$$

qui nous donne comme solution

$$K^2 = \frac{f_1}{F} \arctan \frac{F}{f_1} \quad (6.96)$$

On peut alors calculer la puissance de bruit après désaccentuation:

$$N_{od} = \int_{-F}^{+F} |H_\delta(f)|^2 \cdot 2\eta\omega^2 d\omega \quad (6.97)$$

$$= \frac{8\pi^2\eta}{K^2} \int_{-F}^{+F} \frac{f^2}{1 + \left(\frac{f}{f_1}\right)^2} df \quad (6.98)$$

de

$$\int \frac{x^2}{a^2 + x^2} dx = x - a \arctan \frac{x}{a} \quad (6.99)$$

on obtient

$$N_{od} = 8\pi^2\eta \frac{F - f_1 \arctan \frac{F}{f_1}}{\frac{f_1}{F} - \arctan \frac{F}{f_1}} \quad (6.100)$$

Comme la puissance du signal est gardée intacte, on a :

$$\frac{\left(\frac{S}{N}\right)_{od}}{\left(\frac{S}{N}\right)_o} = \frac{N_o}{N_{od}} \text{ où } N_o = 4\pi^2\eta \int_{-F}^F df = 8\pi^2\eta \quad (6.101)$$

d'où

$$\frac{N_o}{N_{od}} = \frac{\arctan \frac{F}{f_1}}{3\frac{f_1}{F}[1 - \arctan \frac{F}{f_1}]} \quad (6.102)$$

Soit, dans les conditions de la FM commerciale ($F = 15kHz$ et $f_1 = 2.1kHz$), un rapport d'environ 9.7 dB.

Chapter 7

Vers les communications numériques

7.1 Introduction

Nous venons de parcourir rapidement les principales modulations analogiques et leurs performances. Une question qui s'est rapidement posée est de savoir si la transmission numérique peut avoir des avantages importants pour la transmission de messages de nature analogique du type parole ou image par exemple.

Un des premiers critères est la qualité du message reconstitué. Plusieurs questions permettent de bien cerner ce problème :

1. **Dans le cas d'un message analogique, la double transformation analogique-numérique et numérique-analogique n'altère-t-elle pas trop le message analogique?**

Cette question pose le problème de la numérisation. Des études ont été conduites :

- pour caractériser la qualité du message délivré, ce qui amène à définir un seuil de qualité pour chaque type de message.
- pour analyser dans quelles conditions la double transformation évoquée précédemment permet d'atteindre cet objectif de qualité.

Parmi les résultats théoriques les plus classiques relevant de ce dernier point, on peut citer le théorème d'échantillonnage ainsi que de nombreux travaux sur la quantification (entre autres la quantification de la parole et des images).

2. **La suite de symboles délivrée au destinataire ne diffère-t-elle pas trop de celle fournie par la source?**

En raison des perturbations présentes sur le canal, les symboles délivrés au destinataire ne sont pas tous identiques aux symboles fournis par la source : il y a apparition d'erreurs de transmission. On définit le taux d'erreur sur les symboles comme le rapport du nombre de symboles erronés au nombre de symboles transmis. On définit de façon semblable le taux d'erreur sur les bits (TEB, BER : Bit Error Rate) lorsque les symboles sont des éléments binaires ou des mots binaires. Les mesures faites sur des durées limitées permettent d'en obtenir des estimations plus ou moins précises. La qualité de la transmission est d'autant

meilleure que le taux d'erreur est faible. L'idéal est d'obtenir un taux d'erreur nul, mais cet idéal ne peut être qu'approché au prix de dépenses diverses (complexité, puissance, bande, délai ...). Selon la nature du message, on fixe un taux d'erreur à ne pas dépasser de façon à assurer le destinataire d'une qualité minimale du message utile. La détermination de ce seuil de taux d'erreur doit prendre en compte l'influence des erreurs de transmission sur les résultats des opérations effectuées par le destinataire.

Remarquons que de ce point de vue, dans certains cas, le taux d'erreur n'est peut-être pas un paramètre suffisant pour caractériser valablement la qualité de la transmission. On définit également le taux d'erreur par blocs, etc.

3. Le destinataire a-t-il reçu le message dans un délai convenable?

Ce délai est à mettre au compte de la propagation dans le milieu physique utilisé, et également au compte de tous les traitements subis par le message au cours de la transmission. Selon la nature du message, le délai est un paramètre plus ou moins important de la qualité de la transmission. Par exemple, une conversation téléphonique est désagréable si ce délai dépasse quelques dixièmes de seconde.

L'existence d'un délai de transmission a une influence importante dans la définition de certains systèmes complexes qui doivent s'en accommoder sous peine de mauvais fonctionnement. Cependant, il ne sera pas traité dans ce cours.

7.1.1 Les transmissions en plusieurs bonds.

En général, une liaison entre un point A, où est situé l'émetteur, à un point B, où est situé le récepteur, se fait par l'intermédiaire de points C,D,E, ..., où sont situés des répéteurs.

En transmission numérique, on rencontre deux types de répéteurs :

1. Le répéteur-régénérateur propre aux transmissions numériques.

Le répéteur-régénérateur n'est en fait qu'un récepteur qui reconstitue le message numérique émis, suivi d'un émetteur. La liaison globale, en N bonds avec répéteurs-régénérateurs, peut être considérée comme formée de N liaisons numériques ayant chacune une certaine qualité. L'analyse de la liaison globale repose sur l'hypothèse selon laquelle les perturbations sur les N liaisons élémentaires sont indépendantes. Ceci conduit à effectuer N analyses séparées de liaisons numériques et à calculer, par exemple, le taux d'erreur de la liaison globale par addition des taux d'erreur élémentaires (approximation valable pour des taux d'erreur faibles). Compte tenu des variations du taux d'erreur en fonction du rapport signal à bruit, le phénomène d'addition des erreurs est moins dommageable que le cumul des distorsions et des bruits qui se produit lors d'une liaison avec répéteurs sans régénération.

2. Le répéteur sans régénération, utilisé également en transmission analogique.

L'utilisation de répéteurs sans régénération, inévitable en transmission analogique, semble donc sans intérêt en transmission numérique. Cependant, la réalisation d'un répéteur-régénérateur peut être trop onéreuse ou se heurter à de grandes difficultés technologiques. C'est (encore) le cas en télécommunications spatiales et dans les systèmes hybrides de transmission sur câbles. Ceux-ci sont cependant appelés à disparaître.

Exemple 7.1 A

Imaginons que l'on désire transmettre une information numérique sur une distance de 10.000 km par tronçons de 1.000 km. Le canal est tel que sur un tronçon de 1.000 km, on a un taux d'erreurs $\epsilon \ll 1$. Dans le premier cas, on a une probabilité de transmission sans erreurs sur un tronçon qui vaut $1 - \epsilon$. Le bruit responsable des erreurs étant supposé indépendant entre une section et les autres, la probabilité de transmission sans erreurs sur les 10.000 km est simplement le produit des probabilités individuelles soit $(1 - \epsilon)^{10} \simeq 1 - 10\epsilon$, soit un taux d'erreur total valant 10ϵ .

Dans le cas d'un répéteur sans régénération (amplificateur), on peut considérer, en adoptant les hypothèses adéquates, que la puissance du signal à la réception est égale à la puissance du signal émis. Par contre, la puissance de bruit σ_n^2 s'ajoute à chaque tronçon. En consultant alors un graphique typique de l'évolution du taux d'erreur en fonction du rapport signal/bruit, on se rend compte que dans ce cas-ci, où l'on a une dégradation du rapport signal/bruit de 10 dB, que l'évolution du taux d'erreurs est nettement plus catastrophique. $\triangleleft$

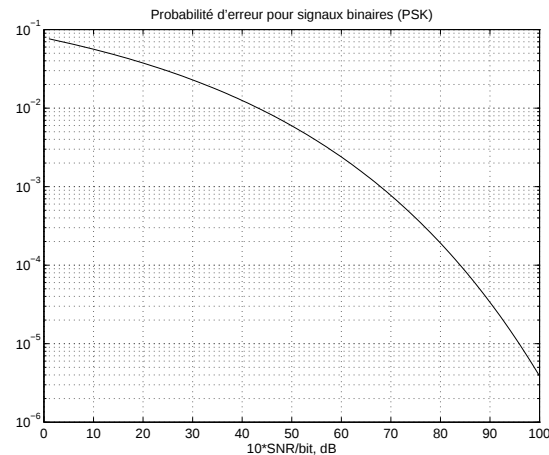


FIG. 7.1 – Évolution typique du taux d'erreur en fonction du rapport signal/bruit.

7.1.2 Largeur de bande

Enfin, la bande passante étant, avec la puissance du signal, la quantité à minimiser dans tout système de communication, il conviendra de déterminer les largeurs de bande utilisées et de minimiser celles-ci en utilisant des modulations numériques adéquates.

7.2 Conversion analogique/numérique

La première étape à la numérisation d'un signal analogique est l'échantillonnage. Il consiste à prélever, à un instant précis, un échantillon du signal analogique, on obtient alors une suite de nombres réels et de temps qui séparent les instants d'échantillonnage. Nous nous limiterons ici aux signaux *unidimensionnels* (une tension en fonction du temps par exemple) et à l'échantillonnage *uniforme*, l'intervalle de temps T_e entre les instants d'échantillonnage étant constant.

La deuxième étape est la quantification. Cette opération consiste à discrétiser l'ensemble des valeurs prises par les échantillons, ce qui revient à approximer la valeur réelle s de l'échantillon

par une valeur s_q qui ne prend qu'un ensemble discret de valeurs appelées *niveaux de quantification*.

La troisième et dernière étape est le codage, qui consiste à établir une correspondance biunivoque entre les N niveaux de quantification et une suite de N représentations binaires.

7.2.1 L'échantillonnage

Nous nous posons la question de savoir dans quelles conditions l'opération d'échantillonnage est une opération réversible. En d'autres mots, si on a un signal $x(t)$ et ses échantillons $x(nT_e)$, quand peut-on dire que les deux signaux portent la même information. La réponse est donnée par le théorème de Shannon.

Théorème de SHANNON

Enoncé: Une fonction $x(t)$ à spectre limité, c'est-à-dire dont la transformée de Fourier $X(\omega)$ est nulle pour $|\omega| > \omega_0 \triangleq 2\pi f_0$, est parfaitement déterminée par ses échantillons $x(nT_e)$ pour autant que la fréquence d'échantillonnage $f_e \triangleq 1/T_e$ soit supérieure au double de la fréquence maximale f_0 . Donc

$$f_e \triangleq \frac{1}{T_e} \geq 2f_0 \quad (7.1)$$

Démonstration

Considérons ce que nous appellerons la **fonction échantillonnée**

$$x_e(t) = \sum_{n=-\infty}^{\infty} x(nT_e)\delta(t - nT_e) \quad (7.2)$$

C'est donc un train d'impulsions de Dirac de moments $x(nT_e)$. Il est clair que, tout en étant une fonction à support t continu, $x_e(t)$ ne contient d'autre information que celle de la séquence $x(nT_e)$. Il nous suffit donc de montrer que l'on peut reconstruire $x(t)$ à partir de $x_e(t)$. En vertu des propriétés de l'impulsion de Dirac, on peut encore écrire (7.2) sous la forme

$$x_e(t) = x(t)p(t) \quad (7.3)$$

où

$$p(t) = \sum_{n=-\infty}^{\infty} \delta(t - nT) \quad (7.4)$$

Comme le montre la figure (7.2), on peut considérer que $x_e(t)$ résulte de la modulation de l'amplitude du train d'impulsions périodiques $p(t)$ par la fonction $x(t)$.

De (7.3), il résulte que la transformée de Fourier de $x_e(t)$ est donnée par la convolution

$$X_e(\omega) = \frac{1}{2\pi} X(\omega) \otimes P(\omega) \quad (7.5)$$

Or, on sait que

$$P(\omega) = 2\pi \sum_{k=-\infty}^{\infty} \delta(\omega - k\omega_e) \quad (7.6)$$

où l'on a posé

FIG. 7.2 – *Echantillonnage de $x(t)$* FIG. 7.3 – *Spectre de la fonction*

$$\omega_e = \frac{2\pi}{T_e} \quad (7.7)$$

Le calcul de convolution est donc aisé et l'on obtient

$$X_e(\omega) = \frac{1}{T_e} \sum_{k=-\infty}^{\infty} X(\omega - k\omega_e) \quad (7.8)$$

Ce résultat est remarquablement simple: *la transformée de la fonction échantillonnée est obtenue par la répétition de celle de $x(t)$ à tous les multiples de la fréquence d'échantillonnage*, comme l'illustre la figure (7.3)

Il est clair que la reconstruction de $x(t)$ à partir de $x_e(t)$ équivaut à celle de $X(\omega)$ à partir de $X_e(\omega)$. Une condition *suffisante* pour que cette reconstruction soit possible est que les spectres partiels ne se recouvrent pas, car on peut alors recouvrer $X(\omega)$ par filtrage. Cela sera toujours possible si $\omega_e \geq 2\omega_0$: le théorème de Shannon est ainsi démontré.

7.2.2 Formule d'interpolation de Whittaker

Il résulte de la figure 7.3 que la reconstruction de $x(t)$ à partir de sa version échantillonnée $x_e(t)$ peut se faire en faisant passer cette dernière dans un filtre passe-bas idéal de gain unitaire et de pulsation de coupure ω_c .

En effet, l'ensemble des fonctions $\text{sinc}[\pi(t - nT_e)/T_e]$ est orthogonal et toute la théorie qui précède montre qu'il est complet sur le domaine des fonctions à spectre seul pour $|\omega| > \omega_e/2$. Comme ces fonctions sont nulles pour $t = kT_e$, sauf pour $k = n$, auquel cas elles prennent la valeur unité, les coefficients du développement sont les échantillons $x(nT_e)$.

FIG. 7.4 – *Reconstruction de $x(t)$ par la formule de Whittaker*FIG. 7.5 – *Repli du spectre du au sous-échantillonnage*

7.2.3 Commentaires

Le théorème de Shannon énonce une condition suffisante pour que l'on puisse reconstruire une fonction $x(t)$ à partir de ses échantillons $x(nT_e)$: spectre limité à une fréquence maximale f_0 et utilisation d'une fréquence d'échantillonnage $f_e \geq 2f_0$. La fréquence d'échantillonnage minimale $2f_0$ est appelée **fréquence de Nyquist**.

Lorsque la fréquence d'échantillonnage est supérieure à ce minimum, on dit qu'il y a *suréchantillonnage*. Suréchantillonner un signal analogique destiné à subir un traitement numérique présente un inconvénient: comme le processeur numérique doit effectuer un certain nombre d'opérations mathématiques sur la durée T_e , il devra être plus rapide. Mais cela présente aussi un certain avantage. Le signal de sortie du processeur consiste souvent dans la séquence des échantillons $y(nT_e)$ d'une fonction $y(t)$ dont le spectre est limité à la même fréquence maximale f_0 que le signal à traiter. Techniquement, la sortie du processeur pourra être une suite de brèves impulsions électriques d'amplitude proportionnelle à $y(nT_e)$, et dont le spectre est du type représenté sur la figure (1.2(b)). Pour reconstruire $y(t)$, il faut appliquer ces impulsions à l'entrée d'un filtre passe-bas. Un filtre passe-bas idéal est évidemment irréalisable. Pratiquement, il faudra réaliser un filtre analogique dont l'affaiblissement passe d'une valeur très faible en $\omega = \omega_0$ à une valeur très grande en $\omega = \omega_e - \omega_0$. La complexité d'un filtre étant directement dépendante de la raideur du flanc de la courbe d'affaiblissement, on tirera donc un certain avantage à utiliser une fréquence d'échantillonnage quelque peu supérieure à la fréquence de Nyquist.

Lorsque la fréquence d'échantillonnage du signal est inférieure à la fréquence de Nyquist, il y a sous-échantillonnage. Dans ce cas, les spectres répétés de $X(\omega)$ se superposent et un filtrage passe-bas du signal échantillonné $x_e(t)$ ne restitue pas exactement $x(t)$, quelle que soit la fréquence de coupure de ce filtre (par exemple ω_0 ou $\omega_e/2$). L'erreur sur le signal reconstitué correspond à un *repli du spectre* (aliasing error), comme l'illustre la figure (7.5).

Il est parfois difficile de prédire quels seront les effets du repli du spectre lorsque l'on sous-

échantillonne un signal analogique en vue de lui faire subir un traitement numérique. Aussi, si les conditions imposent une fréquence d'échantillonnage inférieure à la fréquence de Nyquist, préfère-t-on en général limiter le spectre du signal entrant par un filtre passe-bas de fréquence de coupure inférieure à la moitié de la fréquence d'échantillonnage imposée. Un tel filtre sera de nature analogique et devra bien entendu précéder l'échantillonneur. Il s'impose dans tous les cas où le signal utile est accompagné d'un signal parasite (bruit) occupant des fréquences supérieures à ω_0 .

Enfin, il convient de rappeler que le théorème de Shannon donne une condition suffisante pour qu'on puisse reconstruire une fonction à partir de ses échantillons. Cette condition devient nécessaire si l'on ne possède d'autre information à propos de la fonction que la borne supérieure de son spectre. Dans certains cas, on dispose d'informations complémentaires et il est permis de sous-échantillonner.

Exemple:

- on montre qu'un signal dont le spectre est du type passe-bande, c'est-à-dire est nul en dehors de l'intervalle de fréquence $(f_0 - B/2, f_0 + B/2)$ possède une représentation du type

$$x(t) = x_1(t) \cos \omega_0 t + x_2(t) \sin \omega_0 t \quad (7.9)$$

où $x_1(t)$ et $x_2(t)$ ne contiennent pas de fréquences supérieures à $B/2$. Cette représentation, dite de Rice, est beaucoup utilisée en télécommunications. Selon le théorème de Shannon, on devrait échantillonner $x(t)$ à une fréquence supérieure à $2f_0 + B$. Cependant $x(t)$ est entièrement défini par $x_1(t)$ et $x_2(t)$, et donc par les échantillons de ceux-ci pris à une fréquence supérieure à B . Ceux-ci peuvent être obtenus directement en échantillonnant $x(t)$ à des paires d'instants $(t'_n, t''_n = t_n + \frac{\pi}{2\omega_0})$ tels que $\cos \omega_0 t'_n = 1$, $\sin \omega_0 t''_n = 1$, la saisir de ces paires d'échantillons devant se faire à une fréquence supérieure ou égale à B . C'est là bien un sous-échantillonnage de $x(t)$, d'autant plus intéressant que l'on peut avoir $B \ll 2f_0 + B$.

7.2.4 La quantification

Quantification uniforme

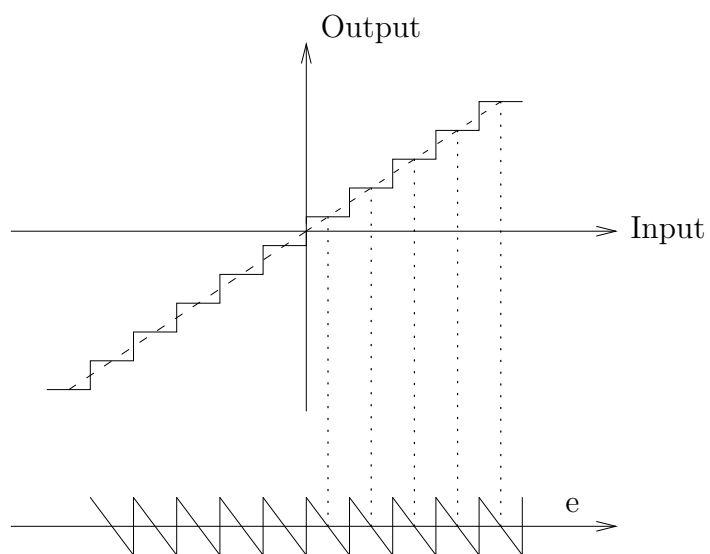
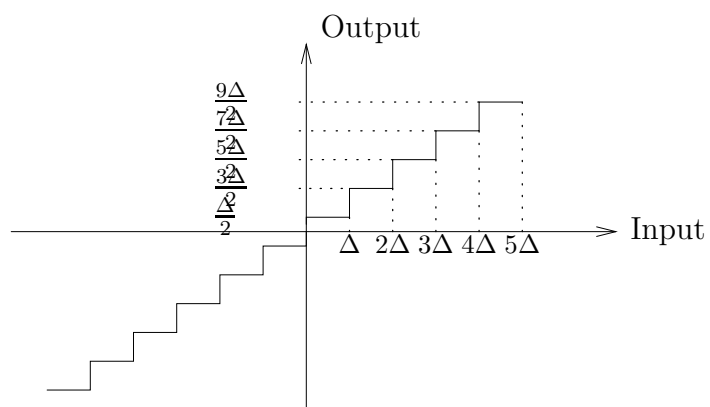
A la suite de l'échantillonnage se trouve la quantification. Si on considère un signal $m(t)$ et sa suite d'échantillons $\{m_n\} = \{m(nT)\}$ où T est la période d'échantillonnage, on va remplacer m_n par sa version quantifiée m_{q_n} . La "fonction de transfert" du quantificateur a l'allure suivante :

Dans ce cas, on a choisi de diviser le range du signal d'entrée en M intervalles égaux (on parle de *quantification uniforme*) de valeur Δ . La question importante devient alors : **quelle est l'erreur de quantification**, c'est-à-dire, connaissant la nature du signal analogique utilisé, quelles sont la moyenne et surtout la variance de l'erreur commise en considérant le signal quantifié plutôt que le signal non quantifié. Dans le cas simple de la quantification uniforme, nous pouvons illustrer l'erreur de quantification par la figure suivante :

où on a représenté l'erreur $e = m_q - m$, en considérant toutes la valeurs possible du signal m . La ligne pointillée représente bien la fonction de transfert idéale ($m_{out} = m_{in}$).

En considérant alors la densité de probabilité des amplitudes de m : $f(m)$, la variance de l'erreur de quantification (ou encore l'erreur quadratique moyenne de quantification) vaut :

$$\overline{e^2} = \sigma_e^2 = \int_{m_1 - \Delta/2}^{m_1 + \Delta/2} f(m)(m - m_1)^2 dm + \int_{m_2 - \Delta/2}^{m_2 + \Delta/2} f(m)(m - m_1)^2 dm + \dots \quad (7.10)$$



où $m_1, m_2, \dots$ sont les niveaux de quantification. En faisant l'hypothèse (vraisemblable si le pas de quantification Δ est suffisamment faible) que $f(m) = f(m_i)$ est constant sur un intervalle, on obtient :

$$\begin{aligned}\sigma_e^2 &= (f(m_1) + f(m_2) + \dots) \int_{-\Delta/2}^{\Delta/2} x^2 dx = (f(m_1) + f(m_2) + \dots) \frac{\Delta^3}{12} \\ &= (f(m_1)\Delta + f(m_2)\Delta + \dots) \frac{\Delta^2}{12}\end{aligned}\quad (7.11)$$

où $f(m_i)\Delta$ est la probabilité que le signal m aie une valeur comprise dans le $i^{\text{ème}}$ intervalle. La somme entre parenthèse vaut donc l'unité et l'erreur quadratique moyenne de quantification vaut :

$$\sigma_e^2 = \frac{\Delta^2}{12} \quad (7.12)$$

Quantification non-uniforme

Dans le cas du signal de parole, on remarque que la distribution d'amplitude est loin d'être linéaire, et est plus particulièrement concentrée vers les amplitudes faibles. Il semble donc naturel d'adopter une quantification non-uniforme, avec plus de niveaux de quantification aux faibles amplitudes. Il convient donc de caractériser ce type de quantification et de calculer l'erreur quadratique moyenne.

Vous trouverez ci-dessous une loi de compression typique (entrée en abscisse et sortie en ordonnée) illustrée par $l(m)$.

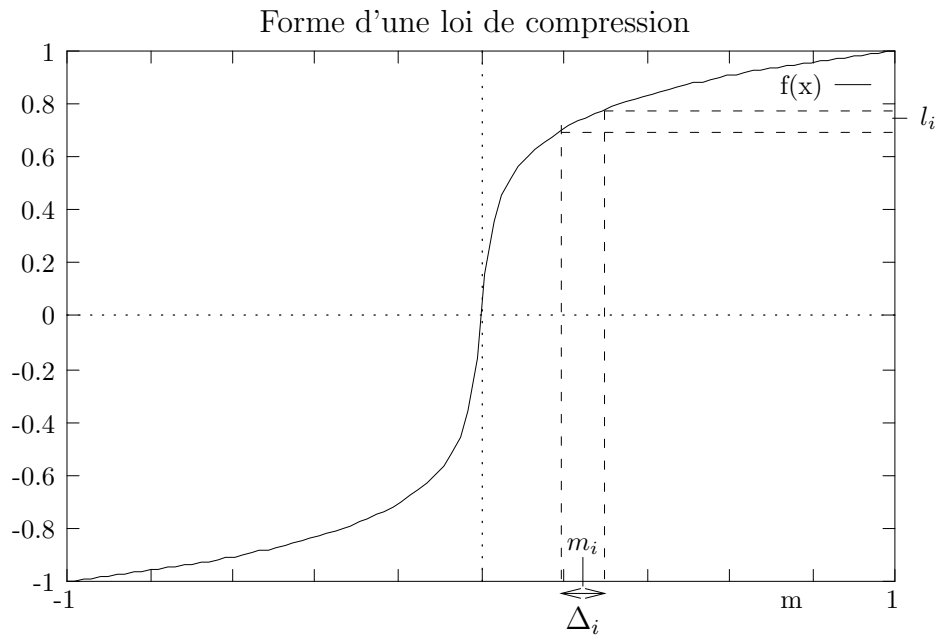


FIG. 7.6 – 7.2.4 Loi de compression typique

Comme dans le cas uniforme, on suppose que les pas de quantification (variables) sont suffisamment petits pour considérer que $f(m) = f(m_i)$ est constant sur l'intervalle considéré.

L'erreur quadratique moyenne peut encore s'écrire :

$$\overline{e^2} = \sigma_e^2 = \int_{m_1 - \Delta_1/2}^{m_1 + \Delta_1/2} f(m)(m - m_1)^2 dm + \int_{m_2 - \Delta_2/2}^{m_2 + \Delta_2/2} f(m)(m - m_1)^2 dm + \dots \simeq \sum_{i=1}^M \frac{p_i \Delta_i^2}{12} \quad (7.13)$$

où p_i est la probabilité que le signal m aie une amplitude comprise dans l'intervalle $[m_i - \frac{\Delta_i}{2}, m_i + \frac{\Delta_i}{2}]$.

De la figure , on déduit aisément :

$$\left. \frac{dl}{dm} \right|_{m=m_i} \simeq \frac{2}{M \Delta_i} \quad (7.14)$$

et donc

$$\overline{e^2} \simeq \sum_{i=1}^M \frac{p_i \Delta_i^2}{12} = \sum_{i=1}^M \frac{p_i}{3M^2} \left(\left. \frac{dl}{dm} \right|_{m=m_i} \right)^{-2} \quad (7.15)$$

En supposant encore les pas de quantification suffisamment petits, on peut passer de la somme à l'intégrale :

$$\overline{e^2} \simeq \frac{1}{3M^2} \int_{-1}^1 \left(\frac{dl}{dm} \right)^{-2} f(m) dm \quad (7.16)$$

On peut alors aisément déduire le rapport Signal/Bruit de quantification :

$$3M^2 \frac{\sigma_m^2}{\int_{-1}^1 \left(\frac{dl}{dm} \right)^{-2} f(m) dm} \quad (7.17)$$

A partir de cette expression, on peut optimiser la loi de compression $l(m)$ de manière à minimiser le rapport signal/bruit de quantification. Cependant, cela sous-entendrait que l'on connaisse la distribution d'amplitude du signal. D'autre part, ce signal étant souvent non stationnaire (par exemple dans le cas de la parole), les performances seraient variables en fonction du type de son (voyelle, consonne), de l'interlocuteur, etc. Il est donc plus utile de choisir une loi de compression qui soit indépendante de la distribution d'amplitude, et donc obtenir un rapport signal/bruit indépendant du message.

On peut réécrire la rapport signal/bruit de quantification sous la forme :

$$3M^2 \frac{\int_{-1}^1 f(m) m^2 dm}{\int_{-1}^1 \left(\frac{dl}{dm} \right)^{-2} f(m) dm} \quad (7.18)$$

Clairement, l'objectif sera atteint si on adopte :

$$\frac{dm}{dl} = \frac{1}{k} |m| \text{ avec } l(1) = 1, l(-1) = -1 \quad (7.19)$$

ou encore :

$$l = (1 + k \ln |m|) \text{sgn}(m) \quad (7.20)$$

En pratique, l'adoption de cette loi induirait des valeurs infinies en $m = 0$, ce qui est inacceptable.

Le CCITT recommande la loi-A :

$$l(m) = \begin{cases} k.A.m & 0 \leq |m| \leq \frac{1}{A} \\ (1 + k \cdot \ln(m)) \text{sgn}(m) & \frac{1}{A} \leq |m| \leq 1 \end{cases} \quad (7.21)$$

avec $k = \frac{1}{1 + \ln A}$

On peut alors dériver l'expression du rapport signal/bruit qui devient :

$$\left(\frac{S}{N}\right)_o = \frac{3M^2 k^2}{1 + \frac{1}{\sigma_m^2} \int_{-\frac{1}{A}}^{\frac{1}{A}} \frac{1}{A} f(m) \left(\frac{1}{A^2} - m^2\right) dm} \quad (7.22)$$

On voit “immédiatement” que le choix de A est un compromis entre :

- Diminuer l'intervalle linéaire (augmenter A) et obtenir ainsi un SNR le plus indépendant possible de $f(m)$.
- Diminuer A , ce qui augmente k et le SNR

En pratique, on a adopté $A = 87.6$ et $M = 256$, ce qui donne un SNR de 38 dB et la courbe vue précédemment.

Notons que dans ce cas-là, pour transmettre un message téléphonique (de fréquence maximale égale à 3400 Hz), on échantillonne à 8 kHz, ce qui, avec $M = 256$ conduit à $8(\text{bits}) * 8 \text{ kHz} = 64 \text{ kbits/sec}$.

Le PCM (Pulse Coded Modulation) ou MIC (Modulation par Impulsions Codées), qui est le codage décrit ci-dessus est entre autres le standard adopté en RNIS (Réseau Numérique à Intégration de Services). Dans son service de base, celui-ci propose 2 canaux à 64 kbits/sec et un canal (données et/ou signalisation) à 16 kbits/sec ; on parle de $2B + D$.

7.3 Delta Modulation

7.3.1 Foreword

The purpose of this chapter is to investigate the performances of delta modulation (DM) as compared to pulse code modulation (PCM).

We first consider the basic circuit. They show that, when some slope overload is tolerated (as in PCM some clipping must be accepted), the signal to quantization noise ratio compares favourably. Further a little more complexity, such as filtering in the feedback loop, improves the performances. A second integrator (a special case of lowpass filtering) is considered to simplify the analysis. Finally delta modulation with slope matching is shown to lead to both adequate signal-to-noise ratio and dynamic range.

Further both coders are considered with imperfect digital transmission. Here the superiority of delta modulation appears definitely since higher error rates are tolerated.

Currently available PCM and DM integrated circuits and the implementation of a speech communication link based on delta modulation are considered in the subsequent chapters.

7.3.2 A reminder of PCM performances

It can be shown¹ that the general formula for the PCM signal to noise ratio is

$$\left(\frac{S}{N}\right)_o = 3N^2 \frac{\int_{-1}^{+1} f(m)m^2 dm}{\int_{-1}^{+1} f(m)\left(\frac{dl}{dm}\right)^2 dm} \quad (7.23)$$

where

$f(m)$ is the probability density function of the message samples

$l(m)$ is the companding law

N is the number of quantized levels (256) and the peak value of the message has been normalized to unity.

It follows that a pure logarithmic companding law leads to $\left(\frac{S}{N}\right)_o$ independant of the characteristics of the message (i.e. $f(m)$), a very desirable property where speech is considered. In practice an approximately logarithmic law can be implemented :

$$l = kAm \quad \text{for } 0 \leq |m| \leq A^1 \rightarrow (1 + klm|m|)sgn m \quad \text{for } A^1 \leq |m| \leq 1 \quad (7.24)$$

with

$$k = (1 + \ln A)^{-1}$$

and A defined by the CCITT ($A = 87.6$). The above is known as the A-law (a similar law known as the μ -law is also used).

With the A-law (within an error smaller than 3 dB)

$$\left(\frac{S}{N}\right)_o \simeq \frac{3N^2}{(1 + \ln A)^2} \quad \text{for } \sigma_n \geq A^1 \quad (7.25)$$

where σ_m is the standard deviation of the message sample.

Finally, over a message dynamic range of more or less 30 dB², the signal to noise ratio is remains close to 38 dB.

7.3.3 Basic DM circuit

The simplest and well-known bloc-diagram of an ideal delta modulator is given in figure 7.7.

The output of the integrator at each sampling instant varies by an amount $\pm s$ and provide the change of the message with time is smaller than s/T (no slope overload), follows $m(t)$ with an error in the range $[-s, +s]$.

The obvious advantage of such an encoder is its simplicity : a few operational amplifiers and a digital circuit (as we shall see later). The decoder is even more simple than the encoder : the

1. see for instance A.L. FAWC, "Telecommunications"; lectures notes.

2. The actual dynamic range depends on the amount of clipping which is not taken into account here. However, increasing the dynamic range would decrease the signal to noise ratio.

FIG. 7.7 – *Basic DM circuit.*

feedback loop and to reduce the quantization noise, a lowpass filter with cut-off frequency equal to the highest message frequency.

Another point of interest in the DM is the equal weight associated to the bits, in other words, the absence of hierarchy in the digital signal. This is a definite advantage from two points of view :

- synchronization (there is no word, but bit synchronization);
- behaviour when a transmission error occurs.

The derivation of the signal to quantization noise ratio is based on four assumptions :

- H_1 : there is no slope overload
- H_2 : the quantization noise is uniformly distributed in $[-s, +s]$
- H_3 : the quantization noise power density is uniform distributed in $[-T_1, +T_1]$
- H_4 : the message is gaussian or sinusoidal.

(Note that T^{-1} is the bandwidth or rather the first zero-crossing of the power spectrum, of the actually transmitted (baseband) digital waveform).

Then

$$N_q = \frac{s^2}{3} \quad (7.26)$$

After ideal lowpass filtering in the decoder

$$N_q = \frac{s^2}{3} FT \quad (7.27)$$

where F is the highest message frequency, and

$$\frac{S}{N_q} = 3\left(\frac{\sigma_m}{s}\right)^2 (FT)^{-1} \quad (7.28)$$

where σ_m^2 is the message variance or power.

7.3.4 Signal to noise ratio for Gaussian messages

The greatest $\frac{S}{N_q}$ is obtained for the highest value of σ_m . However the higher the variance, the higher the probability of slope overload. Here we shall assume that $m(t)$ is gaussian and that an acceptable probability is, nat, $p = 1\%$.

$$p = P\{|m'(t)| > \frac{s}{T}\} = \frac{2}{\sqrt{2\pi}\sigma_{m'}} \int_{\frac{s}{T}}^{\infty} e^{-\frac{x^2}{2\sigma_{m'}^2}} dx \quad (7.29)$$

or, with a classical way to denote the integral in digital communications,

$$p = 2Q\left(\frac{s}{\sigma_{m'}T}\right)$$

with

$$\sigma_{m'}^2 = \frac{1}{2\pi} \int_{-\Omega}^{+\Omega} S_m(w) w^2 dw \quad (7.30)$$

This step therefore implies a knowledge of the message power spectrum. For speech signals the spectrum is usually approximated by

$$S_m(w) = \frac{S_o}{1 + \left(\frac{w}{w_c}\right)^2} \quad (7.31)$$

i.e. the spectrum is almost plot up to a corner and decreases by 6 dB per octave above. To simplify³ we could assume that f_c is the lowest transmitted frequency (300 Hz for telephony). These in the band of interest

$$S_m(w) \simeq S_l\left(\frac{w_c}{w}\right)^2 \quad (7.32)$$

and

$$\sigma_{m'}^2 = \frac{1}{2\pi} \int_{-\Omega}^{\Omega} S_o w_c^2 dw = \frac{S_o w_c^2 \Omega}{\pi} \quad (7.33)$$

On the other hand

$$\sigma_m^2 = \frac{1}{2\pi} \int_{-\Omega}^{+\Omega} S_m(w) dw = \frac{S_o w_c^2}{\pi} \left(\frac{1}{w_c} - \frac{1}{\Omega}\right) \quad (7.34)$$

since $S_m(w) = o$ below w_c .

Therefore

$$\sigma_{m'}^2 = \sigma_m^2 \Omega^2 \left(\frac{\Omega}{w_o} - 1\right)^{-1} \quad (7.35)$$

and

$$p(\text{given}) = 2Q\left[\frac{s}{\sigma_m} \frac{\left(\frac{\Omega}{w_c} - 1\right)^{\frac{1}{2}}}{\Omega T}\right] \quad (7.36)$$

$$Q(x) \simeq b e^{-axl} \quad (7.37)$$

3. The exact calculation is given in appendix A.

with an error smaller than 10% if $b = 0.1550$ and $a = 0.52$

$$p \simeq 2b_e - a \frac{s^2}{\sigma_m^2} \left(\frac{\frac{\Omega}{w_c} - 1}{\Omega^2 T^2} \right) \quad (7.38)$$

$$\left(\frac{\sigma_m}{s} \right)^2 = \frac{Q}{\ln \frac{2b}{p}} \frac{\left(\frac{\Omega}{w_c} - 1 \right)}{\Omega^2 T^2} \quad (7.39)$$

and finally

$$\boxed{\frac{S}{N_q} = \frac{\sigma \pi a \left(\frac{\Omega}{w_c} - 1 \right)}{\ln \frac{2b}{p} \Omega^3 T^3} \text{ for DM with simple integration}} \quad (7.40)$$

We shall now examine two cases :

A) F = 3,400 Hz

$f_c (= f, \text{ the smallest message frequency}) = 300 \text{ Hz}$

$p = 1\%$

$$\frac{S}{N_q} = 0.120 (FT)^3 \quad (7.41)$$

$$\boxed{\left(\frac{S}{N_q} \right)_{dB} = -9.2 - 30 \log_{10} FT} \quad (7.42)$$

For $T^{-1} = 64 \text{ kbit/s}$, $\frac{S}{N_q} : 2\text{q dB}$

Also the signal to noise ratio decreases by 9 dB when we halve the data rate.

Note that the result is not very sensitive to the choice of p . For instance for $p = 1\%$ the signal to noise ratio is 2 dB smaller than the previous value.

B)

$F = 7,000 \text{ Hz}$ (the norm proposed for high

$f_c = 50\text{Hz}$ quality telephony)

$p = 1\%$

$$\left(\frac{S}{N_q} \right) = 1.62 (FT)^{-3}$$

$$\left(\frac{S}{N_q} \right)_{dB} = 2.1 - 30 \log FT$$

For $T^{-1} = 64 \text{ kbit/s}$, $\frac{S}{N_q} = 30.9 \text{ dB}$

7.3.5 Signal-to-noise ratio for sinewaves

Another approach to the problem is based on the experimental result that for speech signals no noticeable slope distortion occurs if the power is limited to the same value found for an 800 Hz sine wave.

For a pure sine wave, $A \sin \omega t$, there is no slope overload if

$$A \leq \frac{s}{\omega T} \quad (7.43)$$

Therefore, based on the experimental result,

$$\frac{S}{N_q} = \frac{3\pi}{\omega^2 \Omega T^3} \quad (7.44)$$

For $F = 3,400$ Hz and $f = 850$ Hz (to simplify the calculation)

$$\frac{S}{N_q} = \frac{6}{\pi^2} (FT)^{-3} \quad (7.45)$$

$$= 0.608 (FT)^{-3}$$

(instead of $0.120 (FT)^{-3}$ by the previous method.)

$$\left(\frac{S}{N_q}\right)_{dB} = -2.24 - 30 \log_{10} FT \quad (7.46)$$

i.e. a difference of 7 dB. This gives for instance a signal to noise ratio of 36 dB at 64 kbit/sec, and the previous result seems rather conservative.

From another point of view, the condition (21) for no slope overload, it is tempting to say that the delta modulator is well matched to speech signals since their spectrum decreases (approximately) by 6 dB per octave. However this argument might be somewhat questionable (... all the frequencies are in the speech signals at the same time).

7.3.6 Driving the DM into slope overload

We may increase the message power beyond the maximum derived from the slope overload condition. Then the signal-to-noise ratio becomes

$$\left(\frac{S}{N}\right)_o = \frac{S_o}{N_q + N_s}$$

where N_s denotes the power of the overload error (figure 7.8).

N_s finally counteracts this effect of the message level increases.

Calculation of N_s would be tedious and not worth the exact knowledge of the signal-to-noise ratio improvement (a few dB). However we should remember that driving the delta modulator slightly into overload is interesting to obtain the best performances, particularly because this type of error is a rare event which does not have as much effect on intelligibility as the quantization noise itself (a problem similar to the anomalous errors due to the imperfect channel; cf the study of the threshold)... and the result quoted above should therefore be considered as the minimum of the achievable signal-to-noise ratio.

FIG. 7.8 – *Overload error*

7.4 Delta modulation with double integration in the feedback loop

The block diagram is the same as fig. 7.7 but with two integrators in the feedback loop as well as in the receiver.

The output of the first integrator is again a staircase waveform with level quantized to multiples??, while the output of the second integrator is made of ramps (a closer approximation of $m(t)$ with slope quantized to multiples of s). Now the rate of change of the slope (instead of the rate of change of the amplitude) is bounded by $\frac{s}{T}$ i.e.

$$|m''(t)| \leq \frac{s}{T} \quad (7.47)$$

The signal to quantization noise ratio can be found in a similar way if no overload occurs. In the most case (fig.7.8) we are just below (above) $m(kT)$ and the slope of $m(t)$ is not changing significantly. The error reaches $-(+)sT$. With the same hypothesis (H , to H_4) as before, with “slope” replaced by “slope change” and s by sT , the quantization noise after lowpass filtering is

$$N_q = \frac{s^2 T^2}{3} FT \quad (7.48)$$

and

$$\frac{S}{N_q} = 3 \left(\frac{\sigma_m}{sT} \right)^2 (fT)^{-1} \quad (7.49)$$

where σ_m is again obtained from the overload condition.

7.4.1 Signal to noise ratio for gaussian messages

With the same approximation as in (10)

$$S_m^{(w)} \simeq S_o \left(\frac{w_c}{w} \right)^2 \quad (7.50)$$

$$S_m^{(w)} \simeq S_o (w_c w)^2 \quad (7.51)$$

and remembering that the spectrum?? has outside ($w_1 = w_c, \Omega$)

$$\sigma_{m''}^2 = \frac{Q_o w_e^2}{2\pi} \int -\Omega^\Omega w^2 dw = \frac{S_o w_c^2}{3\pi} (\Omega^3) w_c$$

$$\sigma_m^2 = \frac{S_o w^2}{\pi} \left(\frac{1}{w_c} - \frac{1}{\Omega} \right) \quad (cf \ (12)) \quad (7.52)$$

and

$$\sigma_{2m''}^2 = \frac{\sigma_{2m}^2 w_c \Omega}{3(\Omega - w_c)} (\omega^3 - w_c^3) \quad (7.53)$$

$$p = P\{|m''_n| > \frac{s}{T}\} = 2Q\left(\frac{s}{\sigma_{m''} T}\right) \simeq 2b e^{-\frac{as^2}{\sigma_{m''}^2 T^2}} \quad (7.54)$$

$$\left(\frac{\sigma - m''}{s} \right)^2 = \frac{a}{T^2 \ln \frac{2b}{n}} \quad (7.55)$$

$$\frac{S}{N_q} = \frac{9a}{16\pi 4 \ln \frac{2b}{p}} \frac{(\frac{\Omega}{w_c} - 1)}{[1 - (\frac{w_c}{\Omega})^3]} (FT)^5 \text{ for DM with integration} \quad (7.56)$$

For $F = 3,400, Hz$, $f_c = 300 Hz$ and $p = 1\%$

$$\frac{S}{N_q} = 0.00914 (FT)^{-5} \quad (7.57)$$

$$\left(\frac{S}{N_q} \right)_{dB} = -20.4 - 50 \log_{10} FT \quad (7.58)$$

At $T^{-1} = 64 kbit/s$, $\frac{S}{S_q} = 43.3 dB$ i.e. a better S/N_q as expected, but also a greater sensitivity to the bit rate (-15 dB for halving the bit rate).

On the other hand for $F = 7,000 Hz$, $f_c = 50 Hz$ and $p = 1\%$.

$$\frac{S}{N_q} = .123 (FT)^{-5} \quad (7.59)$$

$$\left(\frac{S}{N_q} \right)_{dB} = -9.1 - 50 \log_{10} FT \quad (7.60)$$

At $T^{-1} = 64 kbit/s$, $\frac{S}{N_q} = 39 dB$.

FIG. 7.9 – *Delta sigma modulator*

7.4.2 Signal to noise ratio for sinewaves

The condition of no overload becomes

$$A w^2 \leq \frac{s}{T} \quad (7.61)$$

$$\begin{aligned} \frac{S}{N_q} &= \frac{3}{2} \frac{1}{w^4 T^4} (FT)^{-1} \\ &= \frac{3}{32\pi^4} \left(\frac{F}{f}\right)^4 (FT)^{-5} \end{aligned} \quad (7.62)$$

With $F = 3,400 \text{ Hz}$ and $f = 850 \text{ Hz}$ (see 1.4.2)

$$\frac{S}{N_q} = 0.246 (FT)^{-5} \quad (7.63)$$

$$\left(\frac{S}{N_q}\right)_{dB} = -6.08 - 50 \log_{10} FT \quad (7.64)$$

At 64 kbit/sec , $\frac{S}{N_q} = 57.7 \text{ dB}$.

From such a result it is tempting to say that the sine wave approach is quite questionable.

To conclude this point 1.5 we should add that Delta modulation with double integration has some oscillatory behaviour. Thus, it cannot?? (see Steele, Ch. 2, for the details). The main interest is to point out the improvement by low pass filtering in the feedback loop (of which integration is a special case).

7.5 Delta-sigma modulation (DSM)

Let us assume that we put an integrator of the delta modulator (and a differentiator in point of the delta demodulator). Now the two integrators can be replaced by single one in the direct path. At the other end of the link the differentiator and integrator cancel and the receiver reduces to a low pass (or bandpass) filter (fig. 7.9 et 7.10)

We can see that the **delta signal modulator is made of the same components as the delta modulator but rearranged, while the demodulator is even simple.**

Now the overload occurs when

FIG. 7.10 – *Delta sigma modulator*

$$\frac{d}{dt} \left[\int_{-\infty}^t m(t) dt \right] > \frac{s}{T} \quad (7.65)$$

or

$$m(t) > \frac{s}{T} \quad (7.66)$$

7.5.1 Signal to main rate for Gaussian messages

$$p = P\{|m(t)| \geq \frac{s}{T}\} = 2Q\left(\frac{s}{\sigma - mT}\right) \quad (7.67)$$

or

$$\frac{s^2}{\sigma_m^2} \cong \frac{I^2 \ln \frac{2b}{p}}{a} \quad (7.68)$$

Note that here the result (and therefore $\frac{S}{N_q}$) is **independent of the message spectrum**.

To understand the calculation of the signal to noise ratio we must refer to the delta sigma modulator in its initial form (fig. 3a). In the receiver, the output of the integrator is $\int_{-\infty}^t m(t) dt$ with an error uniformly distributed in $[-s, +s]$ where spectrum is uniform in $(-T^{-1}, +T^{-1})$. Therefore at this point the noise power density is

$$\frac{s^2}{3} \frac{1}{2T^{-1}} = \frac{s^2 T}{6} \quad (7.69)$$

which becomes

$$\frac{s^2 T w^2}{6} \quad (7.70)$$

at the output of the differentiator. Then

$$N_q = \frac{1}{T} \int_{w_1}^{w_2} \frac{s^2 T w^2}{6} dw$$

$$N_q = \frac{s^2 T}{18\pi} (w_2^3 - w_1^3) \quad (7.71)$$

Finally

$$\frac{S}{N_q} = \frac{\sigma_m^2}{s^2} \frac{18\pi}{T} (w_2^3 - w_1^3)$$

$$= \frac{18\pi a}{T^3 \ln \frac{2b}{p}} (w_2^3 - w_1^3)^{-1} \quad (7.72)$$

$$\frac{S}{N_q} = \frac{9a}{4M^2 \ln \frac{2b}{p}} [1 - (\frac{w_1}{w_2})^3]^{-1} (FT) \quad (7.73)$$

$$\frac{S}{N_q} = \frac{9a}{4\pi^2 \ln \frac{2b}{p}} (FT)^{-3} \text{ for DSM (ind. of the power spectrum)} \quad (7.74)$$

$$= .0349 (FT)^{-3}$$

$$(\frac{S}{N_q})_{dB} = -14.6 - 30 \log_{10} FT \quad (7.75)$$

for

$$f_2 = F = 3,400 \text{ Hz}, f_1 = 300 \text{ Hz}, p = 1\%$$

At 54 kbit/s:

$$\frac{S}{N_q} = 23.6 \text{ dB for } f_2 = F = 3,400$$

$$\frac{S}{N_q} = 14.2 \text{ dB for } f_2 = 7,000 \text{ Hz}, f_1 = 50 \text{ Hz}, p = 1\%$$

Again let us underline that here the performances are independent of the shape of the message spectrum.

7.5.2 Signal to noise ratio for sinewaves

Here the overload condition becomes

$$A > \frac{s}{T} \quad (7.76)$$

and is independent of frequency. For this reason, it is generally said that DSM is well matched to messages with a uniform spectrum. Actually it is the power of the message, independently of the shape of the message spectrum which drives the modulator into overload.

$$\frac{S}{N_q} = \frac{a^2}{2} \frac{18\pi}{s^2 T} w_2^3 \quad (7.77)$$

where we already take $w_1^3 \ll w_2^3$ into account.

$$\frac{S}{N_q} = \frac{9\pi}{w_2^3 t^3} = \frac{9}{8\pi^2} (FT)^{-3} \quad (7.78)$$

FIG. 7.11 – *DSM with an integrator in the feedback loop*

$$\frac{S}{N_q} = .114(FT)^{-3}$$

$$\left(\frac{S}{N_q}\right)_{dB} = -9.4 - 30 \log FT$$

i.e. 5.2 dB above the estimate with Gaussian hypothesis.

For $F = 3,400 \text{ Hz}$ and $T = 64 \text{ kbit/s}$, $\frac{S}{N_q} = 28.8 \text{ dB}$.

For $F = 7,000 \text{ Hz}$ and $F = 64 \text{ kbit/s}$, $\frac{S}{N_q} = 19.4 \text{ dB}$

7.6 DSM with an integrator in the feedback loop

As before, we may expect an improvement of the performances with an additional integrator in the DSM feedback loop. This case (not considered in the literature) is equivalent to an integrator in front of a Delta modulator with double integration in the loop.

The overload condition is now :

$$p = P\{|m'(t)| \geq \frac{s}{T}\}$$

7.6.1 Signal to noise ratio for Gaussian messages

(From appendix A.1.)

$$\left(\frac{\sigma_m}{s}\right)^2 = \frac{a}{w_c^2 T^2 \ln \frac{2b}{t}} \left(\frac{w_2 - w_1}{\tan^{-1} w_2 - \tan^{-1} w_1} - 1 \right)^{-1}$$

normalized with respect to f_c

With the error after double integration uniformly distributed in $[-sT, +sT]$ and its spectrum uniformly distributed in $(-T^{-1}, +T^{-1})$, the noise power density at the differentiation output is now

$$\frac{s^2 T}{6} w^2$$

Therefore, with the normalized frequencies

$$N_q = \frac{s^2 T^3}{6\pi} w_c^3 \int_{w_1}^{w_2} w^2 dw = \frac{s^2 T^3}{18\pi} w_c^3 (w_2^3 - w_1^3)$$

$$\frac{S}{N_9} = \frac{18\pi a}{w_c^5 T^5 \ln \frac{2b}{p}} [(w_2^3 - w_1^3) (\frac{w_2 - w_1}{\tan^{-1} w_2 - \tan^{-1} w_1} - 1)]^{-1}$$

and since $w_2^3 \gg w_1^3$

$$\boxed{\frac{S}{N_q} = \frac{9aw_2}{16\pi^4 \ln \frac{2b}{p}} [(\frac{w_2 - w_1}{\tan^{-1} w_2 - \tan^{-1} w_1} - 1)]^{-1} (FT)} \quad (7.79)$$

for DSM with integration in the feedback loop

$$f_2 = 3,400Hz, f_1 = 300Hz, F_c = 850Hz, p = 1\%$$

$$\frac{S}{N_4} = 0.00524(FT)^{-5}$$

$$(\frac{S}{N_q})_{dB} = -22.8 - 50 \log FT$$

i.e. 40.9 dB at 64 k bit/s.

$$f_2 = 7,000Hz, f_1 = 50Hz, f_c = 850Hz, t = 1\%$$

$$\frac{S}{N_4} = 0.0123(FT)^{-5}$$

$$(\frac{S}{N})_q = -19.1 - 50 \log_{10} FT$$

i.e. 29.0 dB at 64 kbit/s.

7.6.2 Signal to noise ratio for sine waves

$$aw = \frac{s}{T} \text{ or } Aw = \frac{s}{w_c T} \text{ with normalized frequencies}$$

$$N_q = \frac{s^2 T^3}{18\pi} w_c^3 (w_c^3 - w_1^3) \simeq \frac{s^2 T^3 w_c^3}{18\pi} w_c^3$$

$$\frac{S}{N_q} = \frac{A^2}{2N_q} = \frac{9\pi}{fw^2 w_c^5 T^5 w_2^3} = \frac{9w_2^2}{32\pi^4 l^2} (FT)^{-5}$$

$$= \frac{9}{32\pi^4} (\frac{w_2}{w})^2 (FT)^{-5}$$

For $f_2 = 3,400Hz, f_1 = 300Hz, f = 850Hz$

$$(\frac{S}{N_q})_{dB} = -13.4 - 50 G_{10} FT$$

i.e. $f_2 = 7,000Hz, f_1 = 50Hz, f = 850Hz$.

$$(\frac{S}{N_q})_{dB} = -7.1 - 50 GFT$$

i.e. 41.9 dB at 64 kbit/s.

7.7 Summary of the previous results and conclusions

All the results are for 64 kbits. It should be remembered that the change of the signal to noise ratio with the data rate is 9 dB (15 dB) by factor 2 for single (double) integration. D(S)M denotes the delta (sigma) modulator with single (double) integration in the feedback loop for the initial form (i.e. with an integrator in front of a delta modulator). G denotes the estimation of $\frac{S}{N_q}$ for a Gaussian message, S for a sinewave. The value in brackets denotes the results with the simplified formula.

	$F(Hz)$	DM_1	DM_2
G	3,400	<u>26.6</u> (29.0) 36.0	<u>39.7</u> (43.3) 57.7
G	7,000	20.9 (30.9)	28.3 (39.0)
S		32.9	54.6
		DSM_1	DSM_2
G	3,400	<u>23.6</u>	<u>40.9</u>
S		28.8	50.3
G	7,000	14.2	29.0
S		19.4	41.0

Table 1

Signal to quantization noise ratio (in dB) for various configurations of delta modulation.

- It is worth to use the exact formula for 3,400 Hz and indispensable for 7,000 Hz.
- Replacing the message by an equivalent sinewave is to be rejected (the overestimation is at least 9dB and increases with the system complexity and message bandwidth).
- The signal to noise ratio is about the same for both type of speech (3,400 and 7,000 Hz) if T^{-1}/F is the same for instance 64 Kbit/s for 3,400 Hz and 128 kbit/s for 7,000 Hz).
- Lowpass filtering in the feedback loop (as illustrated by a second integrator) is most important : gains of 13 dB and 17 dB are obtained for DM and DSM respectively for 3,400 Gaussian messages and integration.
- The best results obtained so far (39.7 and 40.9) for 3,400 message are quite comparable to the PCM performances with companding (35 dB at 64 kbit/sec), if we take into account (to be considered next) leading to a similar dynamic range. The main, but drastic, difference resides in the behaviour of the signal-to-quantization noise ratio with bit rate : (in decibels) linear in PCM and logarithmic DM.

APPENDIX A

A.1. **Exact calculation of $\frac{S}{N_q}$ for basic DM.** Here we shall denote by the normalized frequency : $\frac{f}{f_c} \rightarrow f$.

$$S_m(w) = \frac{S_o}{1+w^2}$$

$$S_{m'}(w) = w_c^2 \frac{S_o w^2}{1+w^2}$$

$$\sigma_m^2 = \frac{S_o w_c}{\pi} \int \frac{dw}{1+w^2}$$

$$= \frac{S_o w_c}{\pi} (\tan^{-1} w_c - \tan^{-1} w_1)$$

$$\sigma_{m'}^2 = \frac{S_o w_c^3}{\pi} \int \frac{w^2 dx}{1+w^2} = \frac{S_o w_c^3}{\pi} \left(\int dw - \int \frac{dw}{1+w^2} \right)$$

$$= \frac{S_o w_c^3}{\pi} [-w_1 - (\tan^{-1} w_2 - \tan^{-1} w_1)]$$

$$p = 2Q\left(\frac{s}{\sigma_{m'} T}\right) \simeq 2be^{-\frac{as^2}{\sigma_{m'}^2 T^2}}$$

$$\left(\frac{\sigma_{m'}}{s}\right)^2 = \frac{a}{T^2 \ln \frac{2b}{p}}$$

$$\sigma_{m'}^2 = \sigma_m^2 w_c^2 \left(\frac{w_2 - w_1}{\tan^{-1} w_2 - \tan^{-1} w_1} - 1 \right)$$

$$\left(\frac{\sigma_m}{s}\right)^2 = \frac{a}{w_c^2 T^2 \ln \frac{2b}{p}} \left(\frac{w_2 - w_1}{\tan^{-1} w_2 - \tan^{-1} w_1} - 1 \right)^{-1}$$

$$\frac{S}{N_q} = 3 \left(\frac{\sigma_m}{s}\right)^2 (FT)^{-1}$$

$$\frac{S}{N_q} = \frac{3aw_2^2}{4\pi^2 \ln \frac{2b}{p}} \left(\frac{w_2 - w_1}{\tan^{-1} w - \tan^{-1} w_1} - 1 \right)^{-1} (FT)$$

for DM with simple integration

(note that F is not normalized in the above formula.

For the (true) frequencies $f_2 = 3,400Hz$, $f_1 = 300Hz$, $f_c = 850$ the literature proposes 800 Hz, but 850 Hz will make the calculation somewhat easier) and $p = 1\%$.

$$\frac{S}{N_q} = 0.0690(FT)^{-3}$$

$$\left(\frac{S}{N_q}\right)_{dB} = -11.6 - 30\log_{10}FT$$

or 2.4 dB less than the approximate result given by (20) i.e. 26.6 dB at 64 kbit/s.
For $f_2 = 7,000Hz$, $f_1 = 50Hz$, $f_c = 850Hz$ and $p = 1\%$

$$\frac{S}{N_q} = 0.162(FT)^{-3}$$

$$\left(\frac{S}{N_q}\right)_{dB} = -7.9 - 30\log_{10}FT$$

i.e. 20.9 dB at 64 kbit/s. (certainly more reliable result than (20 bis) (30.9 dB) because of strong approximation in that case.

A.2. Exact calculation of $\frac{S}{N_q}$ for DM double integration

In the previous section we already found:

$$\sigma_m^2 = \frac{S_o w_c}{\pi} (\tan^{-1} w_2 - \tan^{-1} w_1)$$

On the other hand

$$S_{m''}(w) = w_c^4 \frac{S_o w^4}{1 + w^2}$$

$$\begin{aligned} \sigma_{m''}^2 &= \frac{S_o w_c^5}{2\pi} \int \frac{w^4 dw}{1 + w^2} \\ &= \frac{S_o w_c^5}{2\pi} \left(\int \frac{w^4 - 1 + 1}{w^2 + 1} dw \right) \\ &= \frac{S_o w_c^5}{2\pi} \int (w^2 - 1 + \frac{1}{1 + w^2}) dw' \\ &= \frac{S_o w_c^5}{\pi} \left[\frac{w_2^3 - w_1^3}{3} - (w_2 - w_1) + (\tan^{-1} w_2 - \tan^{-1} w_1) \right] \\ &= \sigma_m^2 w_c^4 \left(\frac{\frac{w_2^3 - w_1^3}{3} - (w_2 - w_1)}{\tan^{-1} w_2 - \tan^{-1} w_1} + 1 \right) \end{aligned}$$

Therefore with

$$\frac{S}{N_q} = 3 \left(\frac{\sigma_m}{s\pi} \right)^2 (FT)^1 \quad (7.80)$$

and

$$\left(\frac{\sigma_{m^v}}{s}\right)^2 = \frac{a}{T^2 \ln \frac{2b}{p}} \quad (7.81)$$

$$\frac{S}{N_q} = \frac{3a}{w_c^4 T^4 \ln \frac{2b}{p}} \left(\frac{\frac{w_2^3 - w_1^3}{3} - (w_2 - w_1)}{\tan^{-1} w_2 - \tan^{-1} w_1} + 1 \right)^{-1} (FT)^{-1}$$

$$\frac{S}{N_q} = \frac{3aw_2^4}{16\pi^4 \ln \frac{2b}{p}} \left(\frac{\frac{w_2^3 - w_1^3}{3} - (w_2 - w_1)}{\tan^{-1} w_2 - \tan^{-1} w_1} + 1 \right)^{-1} (FT)^{-1}$$

For $f_2 = 3,400$, $f_1 = 300Hz$, $f_c = 850Hz$ and $p = 1\%$

$$\frac{S}{N_q} = 0.00399(FT)^{-5}$$

$$\left(\frac{S}{N_q}\right) = -24.0 - 50 \log_{10} FT$$

$$\frac{S}{N_q} = 39.7dB \text{ at } 64 \text{ kbit/s}$$

Or 3.6 dB less than the approximate result.

For $f_2 = 7000Hz$, $f_2 = 50Hz$, $f_c = 850Hz$ and $p = 1\%$

$$\frac{S}{N_q} = 0.0105(FT)^{-5}$$

$$\left(\frac{S}{N_q}\right)_{dB} = -19.8 - 50 \log_{10} FT$$

28.3 dB at 64 kbit/s.

Again we should retain this result and note the approximate one (39 dB).

Chapter 8

Modulations numériques : les signaux

8.1 Représentation des signaux numériques modulés

Un signal numérique étant représenté par une suite de nombre, il convient de construire un interface entre ces nombres et le canal, c'est le rôle de la modulation numérique, qui fournira un signal forcément analogique au canal. La procédure habituelle est, étant donné un alphabet M de symboles différents, que l'on peut représenter par $k = \log_2 M$ bits, on fait correspondre (biunivoquement) à chacun des M symboles de la séquence d'information $\{a_n\}$ un signal $s_m(t)$ pris dans un ensemble $\{s_m(t)\}, m = 1, 2, \dots, M$. Les signaux $s_m(t)$ sont supposés être à énergie finie et, forcément, déterministes.

De la même manière que dans le cas analogique, nous avons principalement le choix entre les modulations d'amplitude, de phase ou de fréquence, ou croisées (principalement amplitude/phase) :

$$s(t) = \underbrace{A(t)}_{PAM: Pulse Amplitude Modulation} \cos\left[\underbrace{\omega(t)}_{FSK: Frequency Shift Keying} + \underbrace{\phi(t)}_{PSK: Phase Shift Keying} \right] \quad (8.1)$$

On distingue les modulations :

sans mémoire si, la *rapidité de modulation* (nombre de symboles transmis par seconde) étant $\frac{1}{T}$, $s_m(t) = 0$ si $t < 0$ et $t > T$. En d'autres termes, le signal présent à la sortie du modulateur ne dépend que d'un seul symbole a_n à la fois.

avec mémoire dans le cas contraire.

linéaire si le principe de superposition est applicable à la sortie du modulateur (la sortie peut s'écrire sous la forme d'une somme d'impulsions $s_m(t)$).

non-linéaire , dans le cas contraire.

8.1.1 Modulations linéaires sans mémoire

Dans un premier temps, rappelons-nous que nous travaillons toujours en "signaux équivalents passe-bas", comme décrit au début du chapitre 6. On représente alors le signal par :

$$s(t) = \Re[v(t)e^{j\omega_c t}] \quad (8.2)$$

où $f_c = \frac{\omega_c}{2\pi}$ est la fréquence porteuse et $v(t)$ est le signal équivalent passe-bas.

Modulation d'amplitude

Encore appelée en français MDA (modulation par déplacement d'amplitude) ou, en anglais PAM ou ASK (Amplitude Shift Keying), il s'agit simplement d'associer aux symboles la série de signaux :

$$s_m(t) = \Re[A_m u(t)e^{j\omega_c t}] \quad (8.3)$$

où $\{A_m, m = 1, 2, \dots, M\}$ représentent les M amplitudes possibles.

Exemple 8.1 Modulation d'amplitude à deux états

Simplifions cet exemple à l'extrême, en adoptant :

1. $M = 2$, i.e. deux symboles représentables par un seul bit
2. $A_1 = 1, A_2 = -1$
3. $u(t) = \begin{cases} 1 & 0 < t \leq T \\ 0 & \text{ailleurs} \end{cases}$

Ce qui donne les allures de signaux de la figure 8.1. ◀

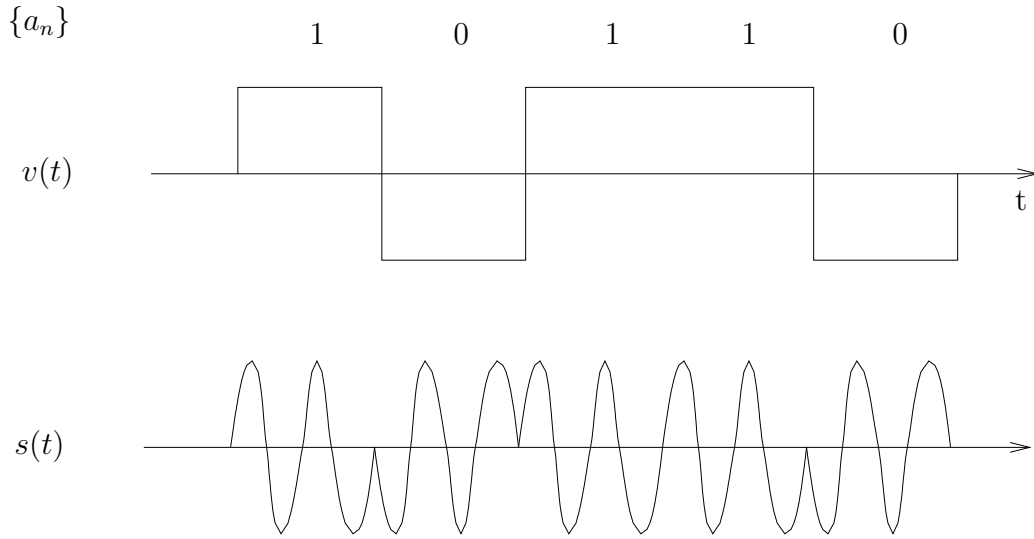


FIG. 8.1 – Modulation d'amplitude binaire à deux états

Le rôle de l'**impulsion de base** $u(t)$ est de transformer le signal discret (présent seulement en des endroits discrets du temps) en un signal analogique. Selon la forme de celui-ci, on a une modulation à mémoire ou non, d'autre part, cette impulsion permet de déterminer, comme nous allons le voir plus loin, la forme de la densité spectrale à la sortie du modulateur.

Le *diagramme d'état* (figure 8.2) d'une modulation représente les états possibles de la sortie dans des axes représentant les fonctions de bases des signaux. Dans ce cas-ci, la fonction de base est $\cos \omega_c t = \Re[e^{j\omega_c t}]$, et donc on a un seul axe.

En termes de signal équivalent passe-bas, on obtient :

$$v(t) = A_m(t) = A_m u(t), \quad m = 1, 2, \dots, M \quad (8.4)$$

Dans ce formalisme, $u(t)$ est la fonction de base et A_m les états possibles du système.

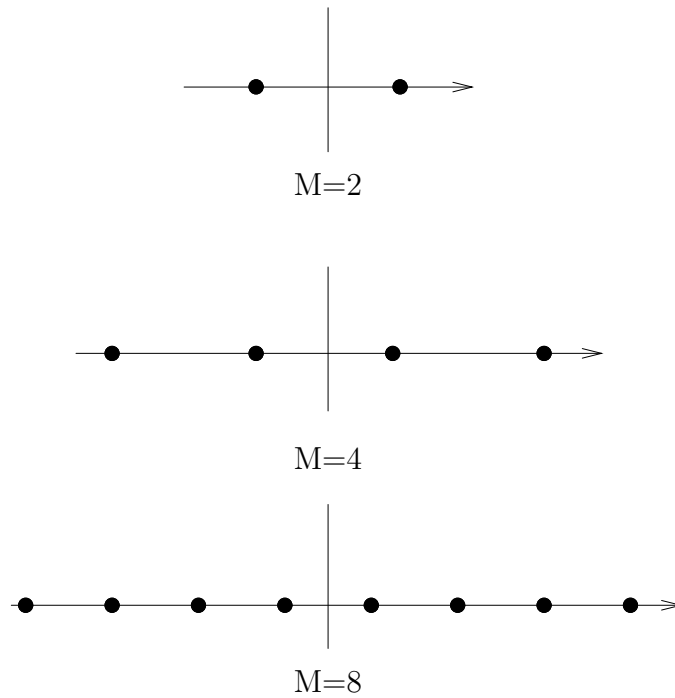


FIG. 8.2 – *Diagramme d'état du signal PAM-M*

Modulation d'amplitude en quadrature

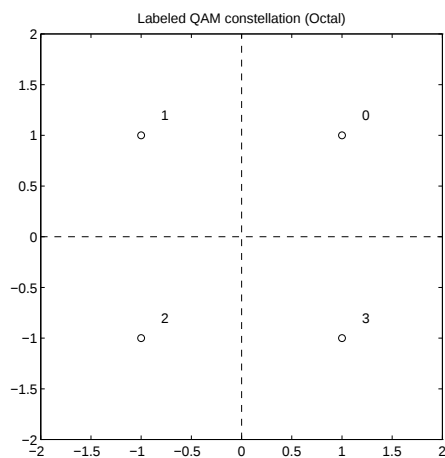
La modulation d'amplitude simple consiste à multiplier la porteuse par une amplitude variable au gré des symboles. Il semble tout-à-fait naturel de vouloir utiliser une seconde porteuse en quadrature. On obtient alors les signaux de base :

$$s_m(t) = A_{mc}u(t) \cos \omega_c t - A_{ms} \sin \omega_c t, \quad m = 1, 2, \dots, M \quad (8.5)$$

Cela correspond simplement à

$$s_m(t) = \Re[(A_{mc} + jA_{ms})u(t)e^{j\omega_c t}] \quad (8.6)$$

Clairement, nous aurons cette fois-ci un diagramme d'état bidimensionnel. Les figures suivantes illustrent le cas du QAM-4 et du QAM-16 (qui portent respectivement 2 et 4 bits par symbole).

FIG. 8.3 – *Constellation du QAM-4*FIG. 8.4 – *Constellation du QAM-16*

Modulation de phase (PSK : Phase Shift Keying)

La modulation de phase consiste à affecter la porteuse d'une phase variable au gré des symboles. On obtient alors les signaux de base :

$$\begin{aligned} s_m(t) &= \Re[u(t)e^{j\theta_m}e^{j\omega_c t}] \\ &= u(t) \cos[\omega_c t + \frac{2\pi}{M}(m-1)] \quad m = 1, 2, \dots, M \end{aligned} \quad (8.7)$$

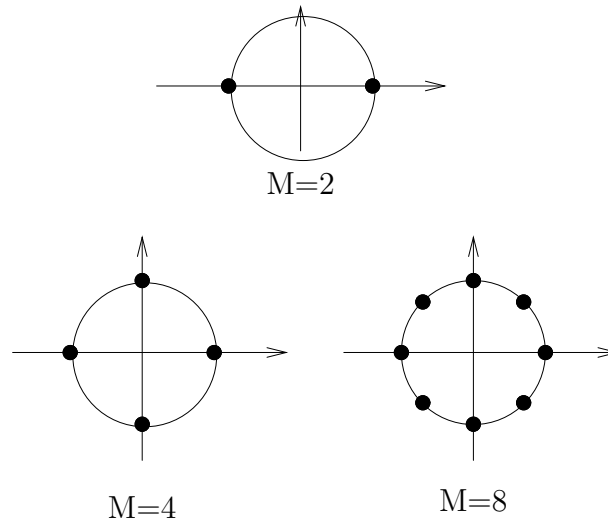


FIG. 8.5 – Diagramme d'état du signal PSK-M

Comme d'habitude, $u(t)$ est une impulsion de base qui sert à déterminer la forme du spectre. Quand celle-ci est constante, le signal PSK est un signal d'amplitude constante. On peut également combiner la modulation d'amplitude avec le PSK.

8.1.2 Modulations non-linéaires à mémoire

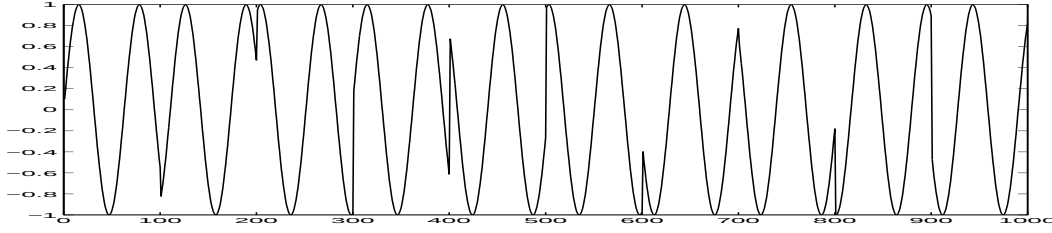
CPFSK : Continuous-Phase Frequency Shift Keying

Un signal FSK est généré en modifiant la fréquence en fonction des données d'une grandeur : $f_n = (\Delta f/2)I_n$, $I_n = \pm 1, \pm 3, \dots, \pm(M-1)$. Le paramètre important est évidemment Δf . On pourrait générer ces différentes fréquences par l'utilisation de $M = 2^k$ oscillateurs, ce qui pourrait provoquer des discontinuités importantes au niveau du signal, comme indiqué ci-dessous.

Il est clair que ces discontinuités de phase généreront un contenu spectral important en dehors de la bande désirée. Nous devons donc nous limiter à un passage d'une fréquence à l'autre à phase continue : **CPFSK** : *Continuous-Phase* FSK.

La représentation d'un signal FSK passe par la définition d'un signal intermédiaire (dit de données) PAM :

$$d(t) = \sum_n I_n u(t - nT) \quad (8.8)$$



où les amplitudes $\{I_n\}$ valent $\pm 1, \pm 3, \dots, \pm(M-1)$ en fonction de la séquence d'information $\{a_n\}$ et $u(t)$ est ici une impulsion rectangulaire d'amplitude $\frac{1}{2T}$ (de manière à avoir $\int_T u(t) = 1/2$) et de durée T . On exprime alors le signal CPFSK équivalent passe-bas par :

$$v(t) = A \exp \left\{ j \left(4\pi T f_d \int_{-\infty}^t d(\tau) d\tau + \phi_0 \right) \right\} \quad (8.9)$$

où f_d est la déviation de fréquence maximale et ϕ_0 est une constante.

Le signal passe-bande peut alors être exprimé sous la forme :

$$s(t) = A \cos[2\pi f_c t + \phi(t; \mathbf{I}) + \phi_0] \quad (8.10)$$

où $\phi(t; \mathbf{I})$ est la phase variable, définie par

$$\begin{aligned} \phi(t; \mathbf{I}) &= 4\pi T f_d \int_{-\infty}^t d(\tau) d\tau \\ &= 4\pi T f_d \int_{-\infty}^t \left(\sum_n I_n u(\tau - nT) \right) d\tau \end{aligned} \quad (8.11)$$

L'intégrale de $d(t)$ est continue et, partant, le signal est bien à phase continue. La phase peut d'ailleurs, en développant l'intégrale, s'exprimer sur l'intervalle $nT \leq t \leq (n+1)T$ par :

$$\begin{aligned} \phi(t; \mathbf{I}) &= 2\pi T f_d \sum_{k=-\infty}^{n-1} I_k + 2\pi f_d (t - nT) I_n \\ &= \theta_n + 2\pi h I_n q(t - nT) \end{aligned} \quad (8.12)$$

où $h, \theta_n, q(t)$ sont définis par :

$$h = 2f_d T \quad (8.13)$$

$$\theta_n = \pi h \sum_{k=-\infty}^{n-1} I_k \quad (8.14)$$

$$q(t) = \begin{cases} 0 & t < 0 \\ t/2T & 0 \leq t \leq T \\ 1/2 & t > T \end{cases} \quad (8.15)$$

On appelle h l'indice de modulation. On observe que θ_n contient une constante qui représente l'accumulation de tous les symboles émis jusque $(n-1)T$.

CPM : Modulation de Phase Continue

On peut généraliser la CPFSK en utilisant simplement un signal à phase variable avec :

$$\phi(t; \mathbf{I}) = 2\pi \sum_{k=-\infty}^n I_k h_k q(t - kT) \quad nT \leq t \leq (n+1)T \quad (8.16)$$

où l'on peut exprimer $q(t)$ en fonction d'une impulsion $u(t)$ par :

$$q(t) = \int_0^t u(\tau) d\tau \quad (8.17)$$

Quand $h_k = h \forall k$, l'indice de modulation est constant pour tous les symboles, sinon, on parle de *modulation CPM multi- h* , dans ce cas, les h_k varient de façon cyclique dans un ensemble fini d'indices de modulation.

Dans le cas du CPFSK, il est intéressant de dessiner les trajectoires de phase possible, c'est ce que fait la figure 8.6 pour le cas binaire. On appelle ce diagramme *l'arbre de phase*.

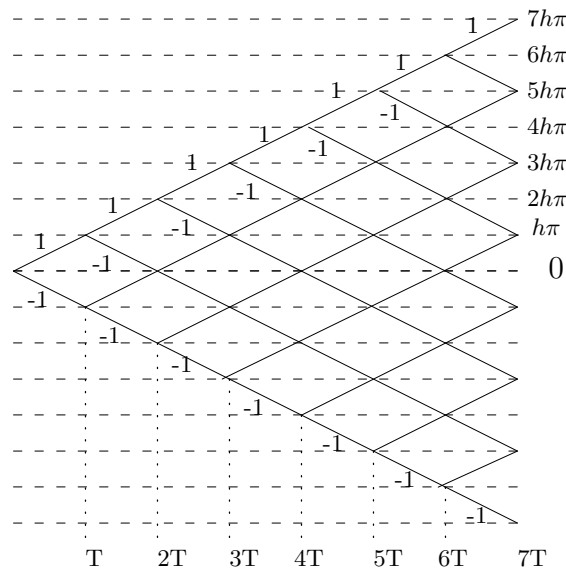


FIG. 8.6 – Arbre de phase pour le CPFSK-2

Cet arbre de phase croît indéfiniment avec le temps, pour pouvoir revenir à des proportions plus acceptables, il convient de ramener les phases entre $-\pi$ et π d'une part, et de choisir h de manière à ce que la phase passe par un multiple de 2π à certains instant kT , k étant un entier. L'arbre de phase devient alors un *treillis de phase* ou *treillis d'état*, les différentes phases aux instants kT se confondant avec des états du système.

La figure 8.7 représente le treillis pour le CPFSK2 avec $h = \frac{1}{2}$.

MSK : Minimum Shift Keying

Dans le cas où l'indice de modulation $h = \frac{1}{2}$, on obtient une modulation particulière appelée *Minimum Shift Keying*. La phase du signal, dans l'intervalle $nT \leq t \leq (n+1)T$ vaut :

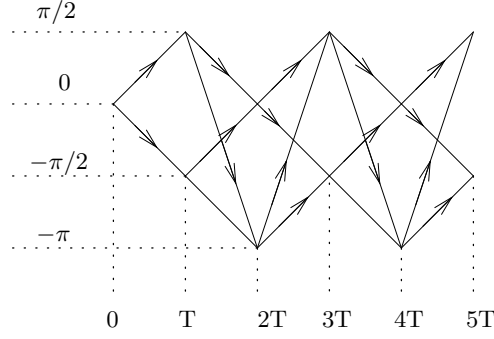


FIG. 8.7 – Treillis de phase pour le CPFSK-2

$$\begin{aligned}\phi(t; \mathbf{I}) &= \frac{\pi}{2} \sum_{k=-\infty}^{n-1} I_k + \pi I_n q(t - nT) \\ &= \theta_n + \frac{\pi}{2} I_n \left(\frac{t - nT}{T} \right)\end{aligned}\quad (8.18)$$

Le signal modulé vaut donc :

$$\begin{aligned}s(t) &= A \cos \left[\omega_c t + \theta_n + \frac{\pi}{2} I_n \left(\frac{t - nT}{T} \right) \right] \\ &= A \cos \left[2\pi \left(f_c + \frac{1}{4T} I_n \right) t - \frac{n\pi}{2} I_n + \theta_n \quad nT \leq t \leq (n+1)T \right]\end{aligned}\quad (8.19)$$

Cette dernière expression montre clairement que le CPFSK-2 est une sinusoïde ayant deux fréquences possibles dans un intervalle de temps donné :

$$\begin{aligned}f_1 &= f_c - \frac{1}{4T} \\ f_2 &= f_c + \frac{1}{4T}\end{aligned}\quad (8.20)$$

La différence de fréquence vaut $\Delta f = f_2 - f_1 = 1/2T$. On peut montrer que c'est la différence minimale pour assurer l'orthogonalité entre les signaux $s_1(t)$ et $s_2(t)$ sur un intervalle de temps de symbole, ce qui justifie l'appellation *Minimum Shift Keying*.

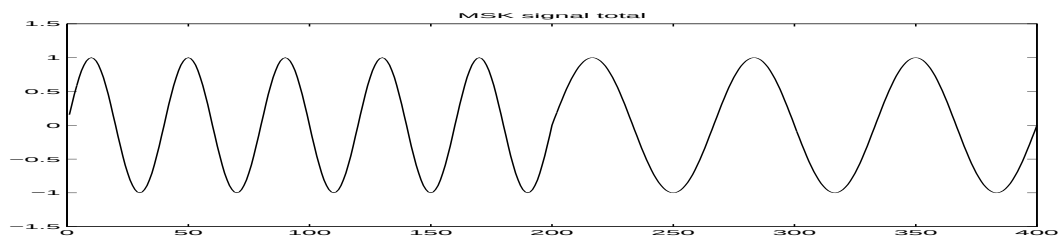
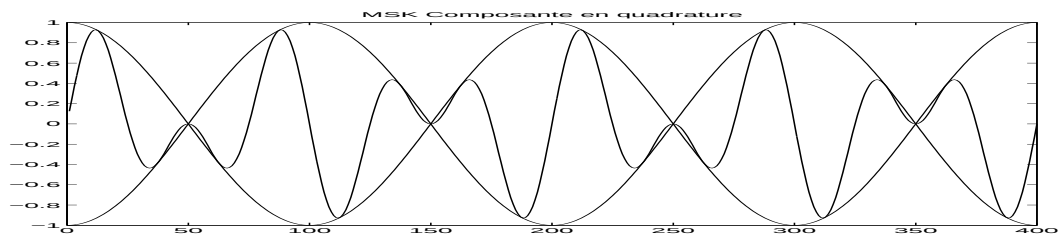
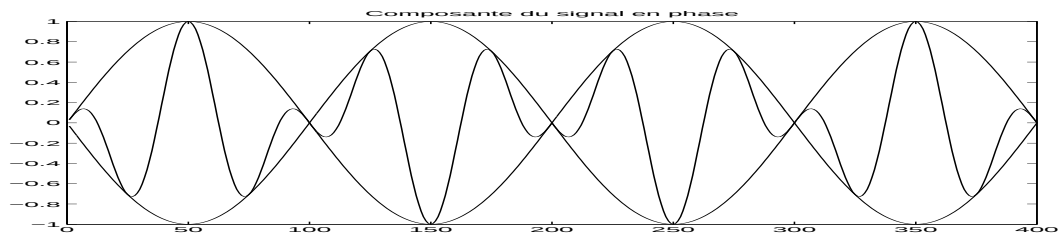
Le MSK peut également être représenté par un PSK-4 à impulsion de base sinusoïdale, soit, en équivalent passe-bas :

$$v(t) = \sum_{n=-\infty}^{\infty} [I_{2n} u(t - 2nT) - j I_{2n+1} u(t - 2nT - T)] \quad (8.21)$$

avec l'impulsion de base

$$u(t) = \begin{cases} \sin \frac{\pi t}{2T} & 0 \leq t \leq 2T \\ 0 & \text{ailleurs} \end{cases} \quad (8.22)$$

Le signal modulé peut encore s'écrire :



$$s(t) = A \left\{ \left[\sum_{n=-\infty}^{\infty} I_{2n} u(t - 2nT) \right] \cos \omega_c t + \left[\sum_{n=-\infty}^{\infty} I_{2n+1} u(t - 2nT - T) \right] \sin \omega_c t \right\} \quad (8.23)$$

et est également appelé *OQPSK* : *Offset Quadrature PSK*, ici, avec impulsion de base sinusoïdale. La figure suivante montre bien les différences entre MSK, OQPSK à impulsion de base rectangulaire et PSK-4 classique. Dans le premier cas, les sauts de phase de 90 degrés se confondent avec le changement de fréquence. Dans le deuxième cas, ces sauts correspondent à une discontinuité dans le signal tandis que pour le QPSK, les sauts de phase sont deux fois moins fréquents et peuvent être de 180 degrés.

8.2 Densités spectrales des signaux modulés

Dans ce chapitre, les deux notions fondamentales sont, d'une part, qu'en fonction du type de modulation et de sa complexité, on peut avoir un débit d'informations (en bits/s.) plus élevé que le débit de symboles (en symboles/s.), d'autre part, que la densité spectrale, et donc la largeur de bande, peut être déterminée par la forme de l'impulsion de base.

En effet, dans le cas des modulations linéaires, on peut voir le système d'émission comme suit :

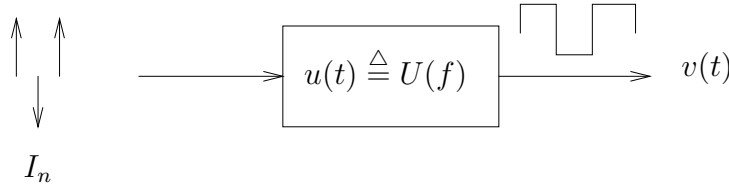


FIG. 8.8 – *Système d'émission*

Les données sont présentes sous formes d'une série d'impulsions de dirac à l'entrée d'une boîte noire qui met ces données en forme. Cette boîte est donc un filtre de réponse impulsionnelle $u(t)$ et fréquentielle $U(f)$. La forme de la densité spectrale dépend donc des caractéristiques des données d'une part et de l'impulsion de base $u(t)$ d'autre part.

Pour calculer les spectres, partons du signal modulé :

$$s(t) = \Re[v(t)e^{j\omega_c t}] \quad (8.24)$$

Sa fonction d'autocorrélation vaut :

$$R_{ss}(\tau) = \Re(R_{vv}(\tau)e^{j\omega_c \tau}) \quad (8.25)$$

Et, par transformation de Fourier, on obtient la densité spectrale :

$$S_{ss}(f) = \frac{1}{2}[S_{vv}(f - f_c) + S_{vv}(-f - f_c)] \quad (8.26)$$

Le signal équivalent passe-bas, dans le cas des modulations linéaires, peut s'écrire :

$$v(t) = \sum_{n=-\infty}^{\infty} I_n u(t - nT) \quad (8.27)$$

dont il suffit de calculer l'autocorrélation.

$$\begin{aligned} R_{vv}(t + \tau; t) &= \frac{1}{2} \mathbb{E} \{v(t + \tau)v^*(t)\} \\ &= \frac{1}{2} \sum_{n=-\infty}^{\infty} \sum_{m=-\infty}^{\infty} \mathbb{E} \{I_n^* I_m\} u^*(t - nT) u(t + \tau - mT) \end{aligned} \quad (8.28)$$

La séquence d'informations $\{I_n\}$ est supposée stationnaire au sens large et de séquence d'autocorrélation $r_{ii}(m) = \frac{1}{2} \mathbb{E} \{I_n^* I_{n+m}\}$. L'autocorrélation de $v(t)$ devient alors :

$$\begin{aligned} R_{vv}(t + \tau; t) &= \sum_{n=-\infty}^{\infty} \sum_{m=-\infty}^{\infty} r_{ii}(m - n) u^*(t - nT) u(t + \tau - mT) \\ &= \sum_{m=-\infty}^{\infty} r_{ii}(m) \sum_{n=-\infty}^{\infty} u^*(t - nT) u(t + \tau - mT - nT) \end{aligned} \quad (8.29)$$

où le terme $\sum_{n=-\infty}^{\infty} u^*(t - nT) u(t + \tau - mT - nT)$ est périodique de période T . $R_{vv}(t + \tau; t)$ l'est donc également, i.e.

$$R_{vv}(t + \tau; t) = R_{vv}(t + \tau + T; t + T) \quad (8.30)$$

et

$$\mathbb{E} \{v(t)\} = \mu_i \sum_{n=-\infty}^{\infty} u(t - nT) \quad (8.31)$$

En clair, cela signifie que $v(t)$ est un processus stochastique ayant sa moyenne et sa fonction d'autocorrélation périodiques, c'est ce qu'on appelle *un processus cyclostationnaire au sens large*.

Pour pouvoir calculer correctement la densité spectrale d'un processus cyclostationnaire, il faut rendre sa fonction d'autocorrélation indépendante de t (ce qui revient à dire qu'il faut le "stationnariser"). Pour ce faire, on va simplement considérer la moyenne de sa fonction d'autocorrélation sur une période T :

$$\begin{aligned} \overline{R_{vv}}(\tau) &= \frac{1}{T} \int_{-T/2}^{T/2} R_{vv}(t + \tau; t) dt \\ &= \sum_{m=-\infty}^{\infty} r_{ii}(m) \sum_{n=-\infty}^{\infty} \frac{1}{T} \int_{-T/2}^{T/2} u^*(t - nT) u(t + \tau - mT - nT) dt \\ &= \sum_{m=-\infty}^{\infty} r_{ii}(m) \sum_{n=-\infty}^{\infty} \frac{1}{T} \int_{-T/2-nT}^{T/2-nT} u^*(t) u(t + \tau - mT) dt \end{aligned} \quad (8.32)$$

En se rappelant que la fonction d'autocorrélation (déterministe) de $u(t)$ vaut :

$$R_{uu}(\tau) = \int_{-\infty}^{\infty} u^*(t) u(t + \tau) dt \quad (8.33)$$

On obtient la relation sympathique :

$$\overline{R_{vv}}(\tau) = \frac{1}{T} \sum_{m=-\infty}^{\infty} r_{ii}(m) R_{uu}(\tau - mT) \quad (8.34)$$

Et donc, la densité spectrale de puissance étant la transformée de Fourier de la fonction d'autocorrélation :

$$S_{vv}(f) = \frac{1}{T} |U(f)|^2 S_{ii}(f) \quad (8.35)$$

En fait, si on travaille en fréquences normalisées ($|\omega| \leq \frac{1}{2} \Rightarrow T = 1$, en vertu du théorème d'échantillonnage), on obtient $S_{vv}(f) = |U(f)|^2 S_{ii}(f)$ par le théorème de Wiener-Kintchine, et, en dénormalisant, on retrouve la relation initiale. Cependant, cette démarche n'est strictement correcte que pour des signaux stationnaires, et le petit calcul qui précède est donc nécessaire pour établir le résultat en toute rigueur.

La densité spectrale de puissance des données vaut :

$$S_{ii} = \sum_{n=-\infty}^{\infty} r_{ii}(m) e^{-j\omega mT} = \sigma_i^2 + \mu_i^2 \sum_{m=-\infty}^{\infty} e^{-j\omega mT} \quad (8.36)$$

En se rappelant des notations :

$$r_{ii}(m) = \begin{cases} \sigma_i^2 + \mu_i^2 & m = 0 \\ \mu_i^2 & m \neq 0 \end{cases} \quad (8.37)$$

L'expression $\sum_{m=-\infty}^{\infty} e^{-j\omega mT}$ est périodique de période $1/T$ et vaut $1/T \sum_{m=-\infty}^{\infty} \delta(f - m/T)$ et donc :

$$S_{ii} = \sigma_i^2 + \mu_i^2 \frac{1}{T} \sum_{m=-\infty}^{\infty} \delta\left(f - \frac{m}{T}\right) \quad (8.38)$$

On obtient finalement

$$R_{vv}(f) = \frac{\sigma_i^2}{T} |U(f)|^2 + \frac{\mu_i^2}{T} \sum_{m=-\infty}^{\infty} \left| U\left(\frac{m}{T}\right) \right|^2 \delta\left(f - \frac{m}{T}\right)$$

(8.39)

Chapter 9

Modulations numériques : les performances.

Nous nous plaçons dans le modèle simple du canal à bruit blanc additif gaussien.

9.1 Démodulateur optimal.

Sans entrer dans le détail du développement de Karhunen-Loève¹, il nous suffira ici de savoir qu'un signal temporel peut être exprimé dans un espace de fonctions de base orthonormales en nombre infini : en appelant $r(t)$ le signal reçu et $s_i(t)$ les signaux émis (dans le cas du PSK-2 par exemple, on aurait $s_0(t) = \cos \omega_c t$ et $s_1(t) = -\cos \omega_c t$).

$$r(t) = s_i(t) + n(t) \quad (9.1)$$

Le développement de Karhunen-Loève nous dit que

$$r(t) = \lim_{N \rightarrow \infty} \sum_{k=1}^N r_k f_k(t) \quad (9.2)$$

où $\{f_k(t)\}$ forment un ensemble de fonctions orthonormales (ou encore base orthonormale) définies sur $[0, T]$:

$$\int_0^T f_m(t) f_n^*(t) dt = \delta_{mn} \quad (9.3)$$

et $\{r_k\}$ représentent les projections de $r(t)$ sur la base précitée. $r(t)$ étant un processus stochastique, les r_k sont des variables aléatoires.

On peut alors montrer :

$$r_k = s_{ik} + n_k, \quad k = 1, 2, \dots \quad (9.4)$$

où s_{ik} et n_k sont respectivement les projections de $s_i(t)$ et de $z(t)$ sur l'espace des $f_k(t)$.

La séquence des $\{n_k\}$ est une séquence gaussienne de moyenne nulle et de covariance $E\{n_m n_k^*\} = \sigma^2 \delta_{km}$.

Il est clair que les v.a. r_k sont gaussiennes, centrées sur s_{ik} et le vecteur $\mathbf{r}_N = [r_1 r_2 \dots r_N]$ a une densité de probabilité conditionnée sur $\mathbf{s}_i = [s_{i1} s_{i2} \dots s_{iN}]$ donnée par :

1. Voir l'appendice 4A de [?]

$$p(\mathbf{r}_N | \mathbf{s}_i) = \frac{1}{\prod_{k=1}^N 2\pi\sigma^2} \exp\left(-\frac{1}{2} \sum_{k=1}^N \frac{|r_k - s_{ik}|^2}{\sigma^2}\right) \quad (9.5)$$

L'objectif fixé est de trouver une règle de décision (en observant $r(t)$, décider quel $s_i(t)$ a été le plus probablement envoyé). En admettant que les $s_i(t)$ et $r(t)$ puissent être représentés dans un plan, (ou un espace de dimension plus grande) cela revient, étant donné un signal $r(t)$ dans ce plan, de lui associer un $s_i(t)$, ce qui revient à diviser ce plan en M régions R_M , M étant le nombre de signaux $s_i(t)$.

La probabilité d'erreur de symboles peut alors s'écrire :

$$P_S = \sum_{i=1}^M p_i \int p(\mathbf{r}_N | \mathbf{s}_i) F_i(\mathbf{r}_N) d(\mathbf{r}_N) \quad (9.6)$$

où $F_i(\mathbf{r}_N)$ est une fonction *indicatrice* valant l'unité partout sauf sur la région R_i correspondant au symbole $\mathbf{s}_i$ qui a été émis et p_i représente la probabilité a priori que $\mathbf{s}_i$ ait été émis. L'intégrale se fait sur l'espace du signal.

En effet, $\int p(\mathbf{r}_N | \mathbf{s}_i) (1 - F_i(\mathbf{r}_N)) d(\mathbf{r}_N)$ représente la probabilité que, $\mathbf{s}_i$ ayant été émis, $\mathbf{r}_N$ soit compris dans la portion d'espace R_i , $1 - \int p(\mathbf{r}_N | \mathbf{s}_i) (1 - F_i(\mathbf{r}_N)) d(\mathbf{r}_N) = \int p(\mathbf{r}_N | \mathbf{s}_i) F_i(\mathbf{r}_N) d(\mathbf{r}_N)$ est donc la probabilité que $\mathbf{s}_i$ ayant été envoyé, $\mathbf{r}_N$ n'appartienne pas à R_i , c'est-à-dire la probabilité d'erreur. La probabilité d'erreur de symbole globale est la somme des probabilités d'erreur individuelle, pondérée par les probabilités d'émission a priori.

Les termes de la somme étant positifs, il suffit de décider que l'on ait émis $\mathbf{s}_i$ (i.e. $F_i(\mathbf{r}_N) = 0$ sur R_i) si

$$p_i p(\mathbf{r}_N | \mathbf{s}_i) \geq p_j p(\mathbf{r}_N | \mathbf{s}_j) \quad \forall j \neq i \quad (9.7)$$

Soit, **en adoptant des symboles équiprobables**

$$p(\mathbf{r}_N | \mathbf{s}_i) \geq p(\mathbf{r}_N | \mathbf{s}_j) \quad \forall j \neq i \quad (9.8)$$

C'est-à-dire, en introduisant l'équation 9.5

$$\frac{1}{\prod_{k=1}^N 2\pi\sigma^2} \exp\left(-\frac{1}{2} \sum_{k=1}^N \frac{|r_k - s_{ik}|^2}{\sigma^2}\right) \geq \frac{1}{\prod_{k=1}^N 2\pi\sigma^2} \exp\left(-\frac{1}{2} \sum_{k=1}^N \frac{|r_k - s_{jk}|^2}{\sigma^2}\right) \quad (9.9)$$

Les quantités étant positives et le logarithme étant monotonément croissant, c'est équivalent à :

$$\sum_{k=1}^N |r_k - s_{ik}|^2 \leq \sum_{k=1}^N |r_k - s_{jk}|^2 \quad \forall i \neq j \quad (9.10)$$

Au passage à la limite, les sommes deviennent des intégrales :

$$\int_{-\infty}^{\infty} |r(t) - s_i(t)|^2 dt \leq \int_{-\infty}^{\infty} |r(t) - s_j(t)|^2 dt \quad \forall i \neq j \quad (9.11)$$

Interprétation géométrique : en représentant les signaux dans un espace de signal, cela revient simplement à choisir le signal $s_i(t)$ le plus proche de $r(t)$, au sens de la distance euclidienne.

En continuant le développement :

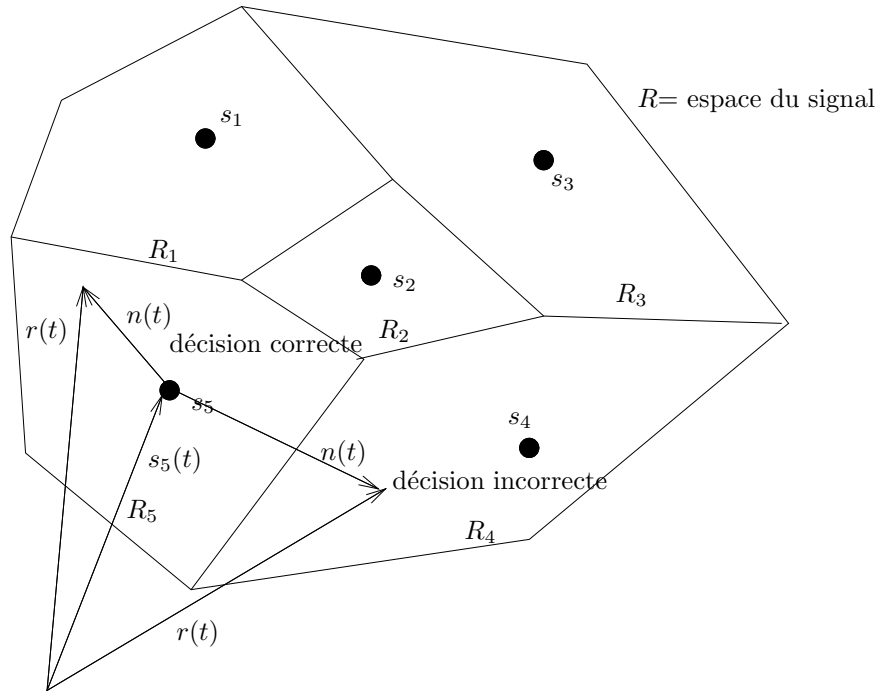


FIG. 9.1 – Représentation géométrique du processus de décision

$$2\Re \int_{-\infty}^{\infty} r(t)s_i^*(t)dt - \int_{-\infty}^{\infty} |s_i(t)|^2 dt \leq 2\Re \int_{-\infty}^{\infty} r(t)s_j^*(t)dt - \int_{-\infty}^{\infty} |s_j(t)|^2 dt \quad \forall i \neq j \quad (9.12)$$

où $\int_{-\infty}^{\infty} |s_i(t)|^2 dt$ représente l'énergie du signal. En adoptant les mêmes énergies, on obtient.

$$\int_{-\infty}^{\infty} r(t)s_i^*(t)dt \leq \int_{-\infty}^{\infty} r(t)s_j^*(t)dt \quad \forall i \neq j \quad (9.13)$$

En optant pour des signaux $s_i(t)$ nuls en dehors de l'intervalle $[0, T]$:

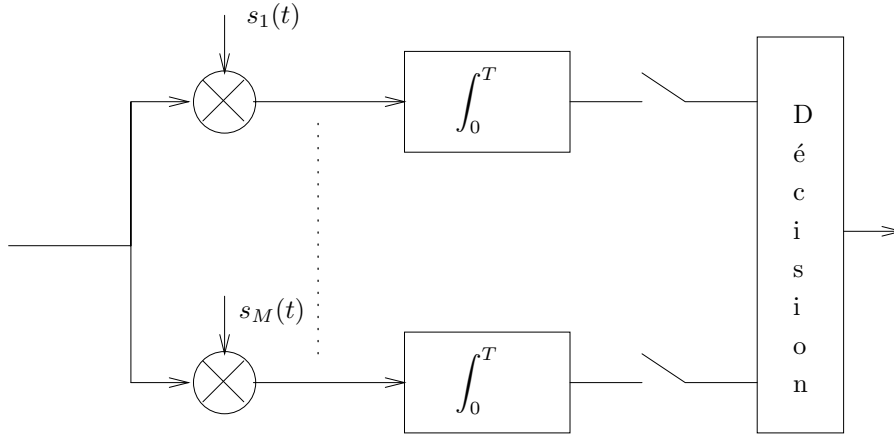
$$\int_0^T r(t)s_i^*(t)dt \leq \int_0^T r(t)s_j^*(t)dt \quad \forall i \neq j \quad (9.14)$$

9.1.1 Récepteur à corrélation

L'application directe de la formule (9.14) nous mène naturellement au récepteur à corrélation. Celui-ci consiste à effectuer une corrélation entre le signal reçu et les différents signaux possibles $s_i(t)$, d'intégrer entre 0 et T , d'échantillonner en T , ce qui donne les variables de décisions :

$$S_i = \int_0^T r(t)s_i^*(t)dt \quad (9.15)$$

Il convient alors de choisir le signal qui correspond à la valeur la plus élevée, soit le schéma de la figure 9.2.

FIG. 9.2 – *Démodulateur à corrélation*

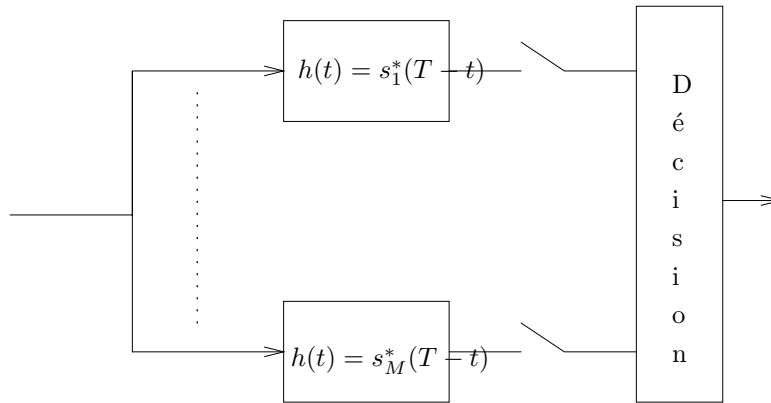
9.1.2 Démodulateur à filtre adapté

L'intégrale (9.14) nous est familière et correspond à la sortie du **filtre adapté** excité par le signal reçu, et échantillonné au taux $1/T$.

$$S_i = \int_{-\infty}^{\infty} r(t) s_i^*(t) dt \quad (9.16)$$

En effet, si on considère le filtre de réponse impulsionnelle $s_i^*(T-t)$, la convolution du signal $r(t)$ avec ce filtre donne bien l'expression (9.14).

L'autre solution est donc :

FIG. 9.3 – *Récepteur à filtre adapté*

Application 9.1 Cas du PAM-2 (= PSK-2)

Dans le cas du PSK-2, on peut écrire le signal sous la forme :

$$s_0(t) = \pm \sqrt{\frac{2E_b}{T}} \sin \omega_c t \quad 0 \leq t \leq T \quad (9.17)$$

où E_b est l'énergie d'un bit (= un symbole ici)

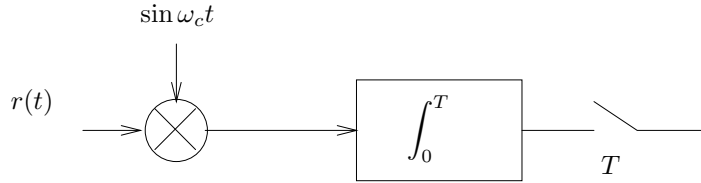


FIG. 9.4 – Récepteur PSK-2 à corrélation

Dans ce cas, l'espace du signal est de dimension 1, où la fonction de base de cet espace vaut $\sin \omega_c t$. De plus, les deux signaux pouvant s'exprimer simplement $s_0(t) = -s_1(t)$, il suffit d'utiliser un seul corrélateur et de décider pour 0 si la sortie est négative et pour 1 si la sortie est positive. Notons au passage que cela demande une réception synchrone (même fréquence pour $s_i(t)$ à l'émission et à la réception) et cohérente (même phase).

9.2 Détermination du taux d'erreurs pour transmissions binaires

On considère les deux signaux $s_0(t)$ et $s_1(t)$ donnés par :

$$s_m(t) = \Re[u_m(t)e^{j\omega_c t}] \quad m = 0, 1; 0 \leq t \leq T \quad (9.18)$$

Ces signaux ont la même énergie

$$E_b = E_s = \int_0^T s_m^2(t) dt = \frac{1}{2} \int_0^T |u_m(t)|^2 dt \quad (9.19)$$

De plus, on peut les caractériser par le coefficient de corrélation mutuelle (cross-correlation) :

$$\rho = \frac{1}{E_b} \int_0^T s_0(t)s_1^*(t) dt \quad (9.20)$$

Supposons que le signal $s_0(t)$ ait été transmis entre $0 \leq t \leq T$, le signal reçu est :

$$r(t) = s_0(t) + n(t) \quad 0 \leq t \leq T \quad (9.21)$$

les variables de décisions valent :

$$\begin{aligned} S_0 &= E_b + N_0 \\ S_1 &= E_b \rho + N_1 \end{aligned} \quad (9.22)$$

avec $N_m = \int_0^T n(t)s_m^*(t)dt$, des composantes de bruit gaussiennes à moyenne nulle.
La probabilité d'erreur est la probabilité que $S_1 > S_0$:

$$P_e = P(S_1 > S_0) = P(S_1 - S_0 > 0) = P(S_0 - S_1 < 0) \quad (9.23)$$

Notons $V = S_0 - S_1 = E_b(1 - \rho) + N_0 - N_1$

Il faut d'abord définir les grandeurs en jeu, c'est-à-dire la moyenne et la variance de cette variable gaussienne :

$$m_v = E\{V\} = E_b(1 - \rho) \quad (9.24)$$

puisque la seule partie aléatoire est $N_0 - N_1$ et est de moyenne nulle.

$$\begin{aligned} \sigma_v^2 &= E\{(N_0 - N_1)^2\} \\ &= E\left\{\left(\int_0^T n(t)s_0^*(t)dt - \int_0^T n(t)s_1^*(t)dt\right)^2\right\} \\ &= E\left\{\int_0^T \int_0^T n(t)n^*(t')(s_0(t) - s_1(t))^*(s_0(t') - s_1(t'))dt dt'\right\} \\ &= \int_0^T \int_0^T \underbrace{E\{n(t)n^*(t')\}}_{=\frac{N_0}{2}\delta(t-t')=\sigma_n^2\delta(t-t')} (s_0(t) - s_1(t))^*(s_0(t') - s_1(t'))dt dt' \\ &= \frac{N_0}{2} \int_0^T |s_0(t) - s_1(t)|^2 dt \\ &= N_0 E_b(1 - \rho) \end{aligned} \quad (9.25)$$

La probabilité d'erreur peut alors être calculée par :

$$\begin{aligned} P(V > 0) &= \int_{-\infty}^0 p(v)dv \\ &= \frac{1}{\sqrt{2\pi\sigma_v^2}} \int_{-\infty}^0 e^{-\frac{(v-m_v)^2}{2\sigma_v^2}} dv \\ &= Q\left(\sqrt{\frac{E_b}{N_0}}(1 - \rho)\right) \\ &= \frac{1}{2} \operatorname{erfc}\left(\sqrt{\frac{E_b}{2N_0}}(1 - \rho)\right) \end{aligned} \quad (9.26)$$

Par un procédé semblable, on obtient aisément que la probabilité d'erreur si $s_1(t)$ avait été envoyé est la même et donc, la probabilité d'erreur totale vaut :

$$P_e = p_0 P_{1|0} + p_1 P_{0|1} \quad (9.27)$$

où p_0 et p_1 sont les probabilités a priori d'émission de $s_0(t)$ et $s_1(t)$ et $P_{1|0}$ est la probabilité d'erreur si $s_0(t)$ a été envoyé. (lire probabilité de décider 1 si 0 a été envoyé). Si les symboles sont équiprobables, on obtient bien :

$$P_e = \frac{1}{2} \operatorname{erfc}\left(\sqrt{\frac{E_b}{2N_0}}(1 - \rho)\right) \quad (9.28)$$

Exemple 9.1 comparaison des taux d'erreur en PSK-2 et en FSK

Dans le cas du PSK-2, on a simplement $s_0(t) = -s_1(t)$ et donc $\rho = -1$, d'où :

$$P_{e_{PSK2}} = \frac{1}{2} \operatorname{erfc}\left(\sqrt{\frac{E_b}{N_0}}\right) \quad (9.29)$$

Dans le cas de signaux orthogonaux ($\rho = 0$), on peut écrire

$$\begin{aligned} s_0(t) &= \sqrt{E_b} \psi_0(t) \\ s_1(t) &= \sqrt{E_b} \psi_1(t) \end{aligned} \quad (9.30)$$

avec $\int_0^T \psi_0(t) \psi_1^*(t) dt = 0$

Dans le cas du FSK-2 on a :

$$\begin{aligned} s_0(t) &= \sqrt{E_b} \sin \omega_0 t \\ s_1(t) &= \sqrt{E_b} \sin \omega_1 t \end{aligned} \quad (9.31)$$

où les pulsations ω_i ont été choisie de manière appropriée. Dans ce cas, on obtient :

$$P_{e_{FSK2}} = \frac{1}{2} \operatorname{erfc}\left(\sqrt{\frac{E_b}{2N_0}}\right) \quad (9.32)$$

Soit une perte de 3 dB par rapport au PSK-2. Ce qui veut dire que pour un même taux d'erreurs, il faut dans le second cas un rapport $\frac{E_b}{N_0}$ deux fois plus grand (soit un signal 3 dB plus puissant à densité spectrale de bruit identique).

remarque De manière à obtenir une approximation rapide des taux d'erreurs, on peut utiliser l'approximation :

$$Q(x) \simeq be^{-ax^2} \quad (9.33)$$

avec $b = 0.155$ et $a = 0.5256$, ce qui donne une précision meilleure que 10 dans l'intervalle $[10^{-2}, 10^{-10}]$, précision % tout-à-fait acceptable. ◀

Notons que pour le cas du FSK-2, nous avons supposé l'utilisation d'une réception cohérente. Etant donné la perte de 3 dB en performances par rapport au PSK-2, il est clair que le FSK ne sera pas utilisé dans ces conditions. Par contre, on peut aisément imaginer un schéma de réception du FSK-2 par filtrage des signaux et comparaison de la puissance de sortie des filtres comme indiqué en figure 9.5.

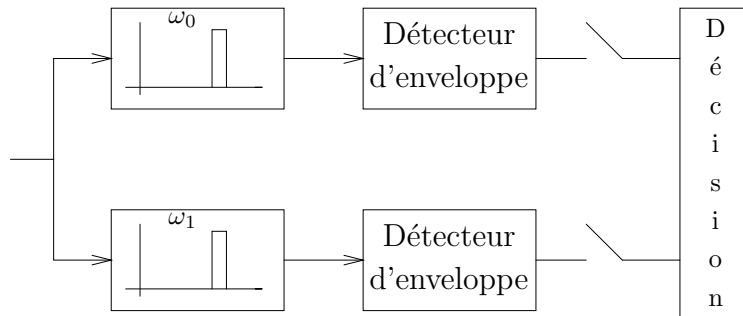


FIG. 9.5 – Récepteur FSK-2 non cohérent

Dans ce cas, on trouve :

$$P_{e_{FSK2 \text{ non cohérent}}} = \frac{1}{2} \exp\left(-\frac{2E_b}{N_0}\right) \quad (9.34)$$

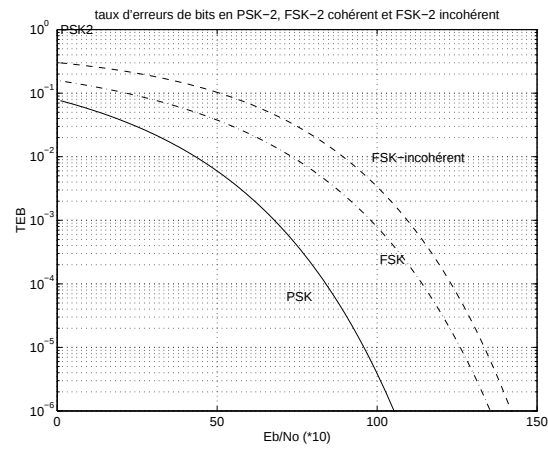


FIG. 9.6 – Comparaison des taux d'erreurs

Chapter 10

Communications dans un canal à bande limitée

10.1 Les canaux.

10.1.1 Le canal linéaire.

Un des modèles de canal les plus utilisés, et valable dans un grand nombre de cas, consiste simplement à approximer le canal par un **filtre linéaire** ayant une réponse fréquentielle équivalent passe-bas $C(f) \Leftrightarrow c(t)$. Dans ce cas, la réponse impulsionnelle équivalent passe-bas d'un signal $s(t) = \Re[v(t)e^{j2\pi f_c t}]$ sera donnée par $r(t) = \int_{-\infty}^{\infty} v(\tau)c(t - \tau)d\tau + z(t)$ où $z(t)$ représente le bruit additif.

Usuellement, pour des raisons évidentes, le canal est limité à une largeur de bande B Hz, c'est-à-dire $C(f) = 0$ pour $|f| > B$.

Si on exprime la réponse fréquentielle $C(f)$ par :

$$C(f) = |C(f)|e^{j\theta(f)} \quad (10.1)$$

où $|C(f)|$ et $\theta(f)$ sont respectivement l'amplitude et la phase de la réponse fréquentielle. On définit le *délai de groupe* par :

$$\tau(f) = -\frac{1}{2\pi} \frac{d\theta(f)}{df} \quad (10.2)$$

Un canal est dit *idéal* si sa réponse fréquentielle est constante en amplitude et linéaire en phase (dans la bande passante), c'est-à-dire que le délai de groupe est constant. En effet, si on définit un canal idéal, en temporel, comme un gain (ou une atténuation) pur et un délai, la transformée de Fourier de $A\delta(t - \tau)$ étant $Ae^{-j2\pi f\tau}$, la réponse fréquentielle doit bien avoir la caractéristique susdite. Dans le cas contraire, la forme du signal est distordue, soit en amplitude, soit en phase. La conséquence de ceci est une déformation des impulsions qui peuvent avoir une longueur plus importante qu'à l'émission et donc introduire une *interférence entre symboles*. (*ISI: InterSymbol Interference*)

La manière la plus évidente de visualiser celle-ci est de dessiner la réponse impulsionnelle du canal. Dans un cas comme celui indiqué ci-dessous, on observe directement, non seulement la distorsion introduite, mais également la superposition des symboles successifs qui, aux instants d'échantillonnage T , introduira cette *interférence entre symboles*. Le rôle de l'*égaliseur/organe de décision*, placé derrière le canal, sera de compenser cette ISI, soit, comme dans le cas des

égaliseurs en essayant de filtrer la sortie de manière à retrouver la forme d'impulsion d'entrée ou, en tous cas, de diminuer l'ISI présente aux instants kT , soit en identifiant le canal et en tenant compte de cette information dans le processus de *décision*.

Application 10.1 Dans le cas d'un délai de groupe $\tau(f)$ linéaire en fréquence, étant donné une impulsion de départ ne donnant pas lieu à interférence entre symboles, on constate que la sortie présente bien de l'ISI. La sortie de l'égaliseur nous donnera, idéalement, l'impulsion de source.

application

Une autre manière d'aborder l'égalisation est d'observer la réponse fréquentielle du canal, par son amplitude et son délai de groupe. Ensuite, on cherche à compenser l'"inidéalité" du canal par un filtre. La figure ci-dessous vous donne les caractéristiques temporelles et fréquentielles typiques d'un canal téléphonique.

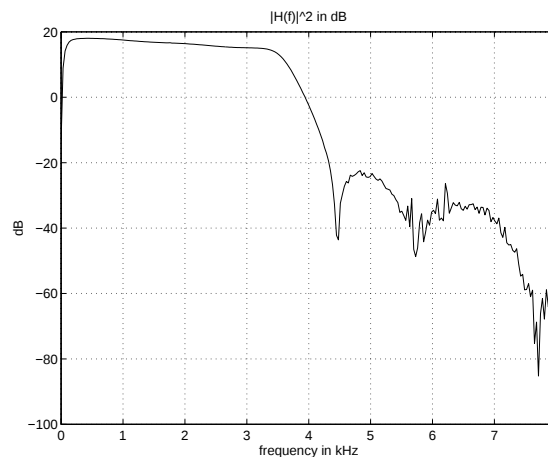


FIG. 10.1 – Réponse fréquentielle typique d'un canal téléphonique

D'autres types de distorsion peuvent apparaître sur les canaux :

Les distorsions non-linéaires, principalement dues aux non-linéarités des amplificateurs et des mélangeurs, elles sont très difficilement prises en compte au niveau traitement numérique du signal et doivent être réduites le plus possible à la source (d'où l'importance de la conception des amplificateurs et autres éléments électroniques du système d'émission et de réception).

Les décalages fréquentiels, présents dans les systèmes présentant des modulations de type SSB et dus à une démodulation imparfaite (porteuse résiduelle et/ou décalée). Ceux-ci peuvent être compensés soit par une boucle à verrouillage de phase, soit par un traitement numérique adéquat.

Le bruit de phase (phase jitter), qui peut être considéré, dans certains cas, comme une modulation de fréquence de faible indice de modulation. C'est souvent une spécification que les systèmes de réception doivent tenir, de manière à ne pas perturber outre mesure l'égalisation/décision.

Le bruit impulsionnel, additif, qui s'ajoute éventuellement au **bruit blanc** ou coloré.

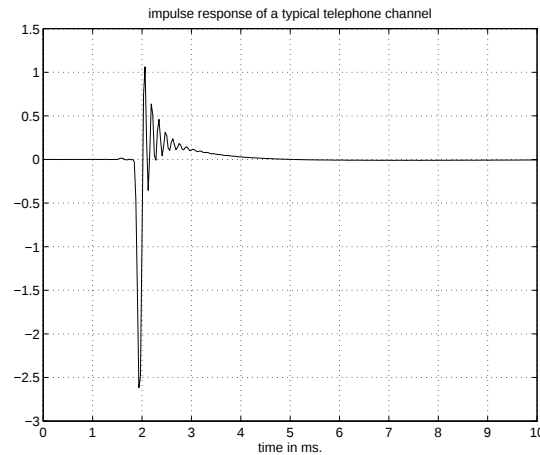


FIG. 10.2 – Réponse temporelle typique d'un canal téléphonique.

10.2 Le critère de Nyquist.

10.2.1 Notion d'Interférence Entre Symboles (ISI: Inter Symbol Interference).

Dans un premier temps, nous considérons le cas d'un canal linéaire sans bruit additif. Dans le cas classique des modulations linéaires, on peut représenter les signaux de source équivalent passe-bas par :

$$\sum_{n=0}^{\infty} I_n g(t - nT) \quad (10.3)$$

où $\{I_n\}$ est la séquence de symboles d'information et $g(t)$ l'impulsion de mise en forme du signal, à bande passante limitée à B ($g(t) \Leftrightarrow G(f)$ avec $|G(f)| = 0$ pour $|f| > B$). Ce signal est transmis sur un canal ($C(f)$) ayant la même limitation de bande passante. En conséquence, on représente le signal reçu par

$$\sum_{n=0}^{\infty} I_n h(t - nT) + z(t) \quad (10.4)$$

où

$$h(t) = \int_{-\infty}^{\infty} g(\tau) c(t - \tau) d\tau \quad (10.5)$$

et $z(t)$ représente le bruit additif Gaussien.

La théorie de la décision nous apprend qu'il suffit d'échantillonner la sortie du canal, préalablement filtrée par un filtre adapté (de réponse fréquentielle $H^*(f)$) à un taux de $1/T$ éch./sec. On peut donc écrire

$$y(t) = \sum_{n=0}^{\infty} I_n x(t - nT) + \nu(t) \quad (10.6)$$

où $x(t)$ est la réponse impulsionnelle du filtre $h(t) \otimes h(-t)$ et $\nu(t)$ le bruit blanc filtré par $H^*(f)$. Après échantillonnage, les variables aléatoires de sortie sont :

$$y(kT + \tau_o) = y_k = \sum_{n=0}^{\infty} I_n x(kT - nT + \tau_o) + \nu(kT + \tau_o) \quad (10.7)$$

où τ_o est un délai introduit par le canal.

$$y_k = \sum_{n=0}^{\infty} I_n x_{k-n} + \nu_k \quad k = 0, 1, \dots \quad (10.8)$$

En normalisant x_o à 0, on obtient :

$$y_k = I_k + \sum_{\substack{n=0 \\ n \neq k}}^{\infty} I_n x_{k-n} + \nu_k \quad (10.9)$$

où le second terme représente l'interférence entre symboles (influence sur la sortie à l'instant k des symboles d'indice différent de k) et ν_k représente le bruit additif coloré par le filtre de réception.

10.2.2 Le critère de Nyquist.

L'objectif étant d'avoir $y_k = I_k$, c'est-à-dire d'éliminer l'interférence entre symboles, nous allons aborder le problème d'un point de vue temporel dans un premier temps.

La condition pour éliminer l'ISI dans l'équation (10.9) est :

$$x(t = kT) = x_k = \begin{cases} 1 & k = 0 \\ 0 & k \neq 0 \end{cases} \quad (10.10)$$

Le théorème d'échantillonnage des signaux à bande passante limitée nous apprend que $x(t)$ peut être exprimé par :

$$x(t) = \sum_{n=-\infty}^{\infty} x\left(\frac{n}{2B}\right) \frac{\sin 2\pi B(t - n/2B)}{2\pi B(t - n/2B)} \quad (10.11)$$

et

$$x\left(\frac{n}{2B}\right) = \int_{-B}^B X(f) e^{j2\pi n/2B} df \quad (10.12)$$

ce qui correspond à une période d'échantillonnage $T = \frac{1}{2B}$. Si on fait le choix particulier de ce débit de symboles (i.e. on envoie un symbole I_k toutes les T secondes), on obtient :

$$x(t) = \sum_{n=-\infty}^{\infty} x(nT) \frac{\sin \pi(t - nT)/T}{\pi(t - nT)/T} \quad (10.13)$$

Si l'on veut réduire l'ISI à néant, c'est-à-dire obtenir $x(nT) = 0 \forall n \neq 0$, nous devons adopter :

$$x(t) = \frac{\sin \pi t/T}{\pi t/T} \quad (10.14)$$

soit

$$X(f) = \begin{cases} T & |f| \leq \frac{1}{2T} \\ 0 & |f| > \frac{1}{2T} \end{cases} \quad (10.15)$$

Ce type de signal est de toute évidence impossible à réaliser, puisque l'impulsion $x(t)$ ainsi définie est totalement anti-causale et de durée infinie. D'autre part, la décroissance de $x(t)$ est en $1/t$ et, un échantillonnage non idéal (en un temps $kT + \tau$) auront comme effet pour l'échantillon de se voir affecter d'un terme provenant d'une somme de tous les symboles à un facteur près :

$$\sum_{\substack{n=0 \\ n \neq k}}^{\infty} I_n \frac{\sin \pi(nT + \tau)/T}{\pi(nT + \tau)/T} \quad (10.16)$$

Cette somme peut diverger (ce qui est non physique, comme l'est l'impulsion considérée!). Il est donc naturel d'adopter un débit de symboles $1/T$ inférieur à $2B$.

Dans ce cas, en se souvenant de la propriété de périodisation du spectre des signaux échantillonnés :

$$X_e(f) = \frac{1}{T} \sum_k X\left(f + \frac{k}{T}\right) \quad (10.17)$$

On obtient la relation spectrale entre la sortie et l'entrée du canal de la forme :

$$Y(f) = \frac{1}{T} \sum_k X\left(f + \frac{k}{T}\right) \quad (10.18)$$

Et la condition pour ne pas avoir d'ISI s'exprime par :

$$X_e(f) = T \quad \forall f \quad (10.19)$$

Les impulsions ci-dessous satisfont le critère de Nyquist et sont appelées *impulsions de Nyquist*.

Habituellement, on choisit le débit de symboles $B < 1/T < 2B$. Une impulsion couramment utilisée l'impulsion appelée *en cosinus surélevé*.

$$X(f) = \begin{cases} T & 0 \leq |f| \leq (1 - \beta)/2T \\ \frac{T}{2} [1 - \sin(f - 1/2T)/\beta] & (1 - \beta)/2T \leq |f| \leq (1 + \beta)/2T \end{cases} \quad (10.20)$$

où β est appelé le *paramètre de décroissance*. L'impulsion correspondante à l'allure :

$$x(t) = \frac{\sin \pi t/T}{\pi t/T} \frac{\cos \beta \pi t/T}{1 - 4\beta^2 t^2/T^2} \quad (10.21)$$

La décroissance de l'impulsion est cette fois-ci en $1/t^3$, ce qui conduit à une somme convergente pour l'ISI.

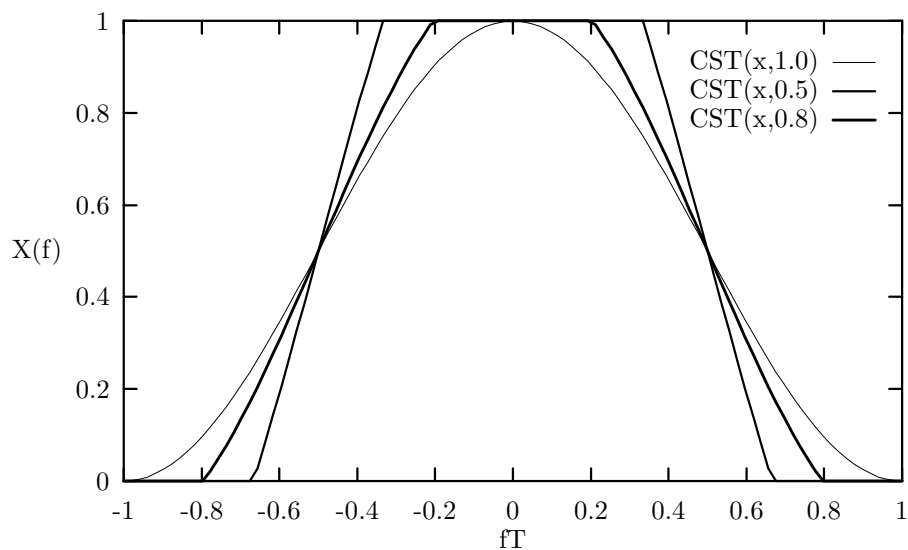


FIG. 10.3 – Réponse fréquentielle en cosinus surélevé

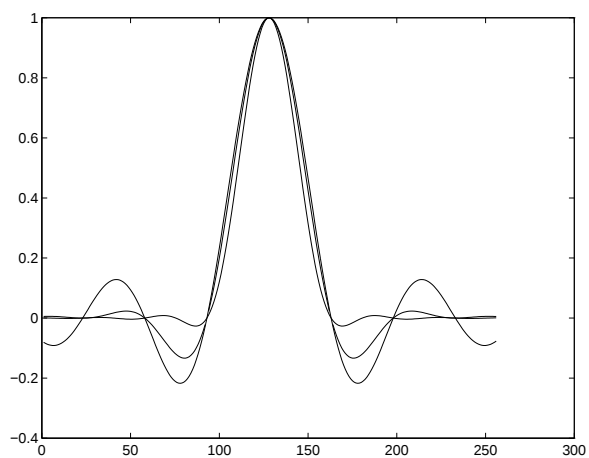
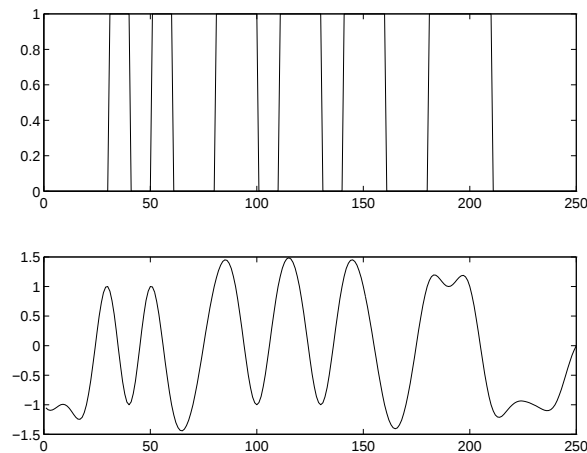


FIG. 10.4 – Réponse temporelle de filtres en cosinus surélevé

FIG. 10.5 – *Train d'impulsions filtré par un filtre de Nyquist*

10.2.3 Lien avec le théorème d'échantillonnage.

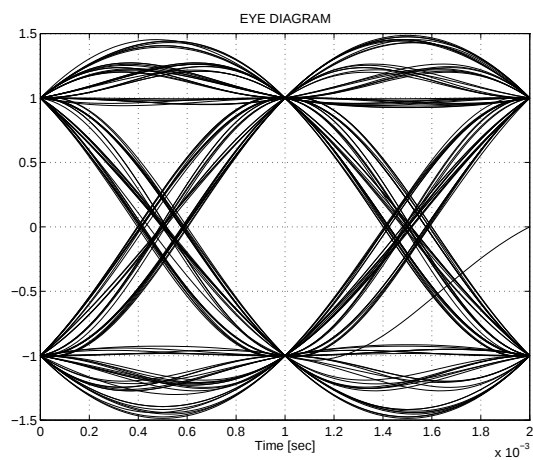
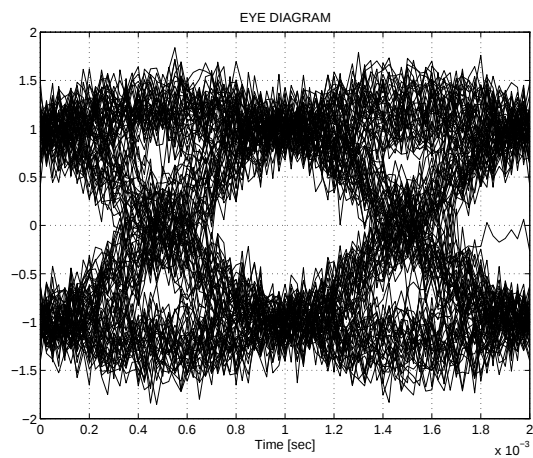
Dans le cas du théorème d'échantillonnage, on nous apprend que si un signal (passe-bande) a une bande passante B , il faut échantillonner à un taux supérieur à $T = 1/2B$. En d'autres termes, la quantité d'information contenue dans un signal continu est équivalente à la quantité d'information contenue dans un signal échantillonné correctement.

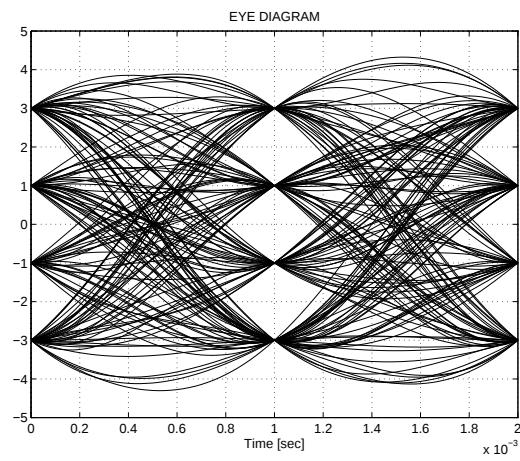
Dans le cas du critère de Nyquist, on cherche un signal continu qui contienne au moins l'information du signal discret, il est donc logique que l'on doive adopter une bande passante supérieure $B = 1/2T$. Il s'agit en quelque sorte de deux théorèmes duaux.

10.2.4 Le diagramme en oeil.

Une manière classique de caractériser l'ISI est le diagramme en oeil. Il s'agit simplement d'afficher le signal continu $y(t)$ sur un oscilloscope en synchronisant la base de temps sur un multiple de T (habituellement $2T$). La figure obtenue a l'allure d'un oeil dont :

1. L'ouverture verticale reflète la résistance au bruit sous échantillonnage idéal.
2. L'ouverture horizontale reflète la sensibilité au désalignement de l'échantillonnage.
3. L'épaisseur des traits, à l'instant d'échantillonnage idéal, reflète la quantité d'ISI présente. Dans le cas où cette épaisseur est très grande, on a un oeil "fermé". Le rôle d'un égaliseur sera d'"ouvrir" le plus possible cet oeil.
4. La pente de l'oeil indique la sensibilité de l'ouverture verticale à de faibles désalignement d'échantillonnage.

FIG. 10.6 – *Diagrammes en oeil (binaire)*FIG. 10.7 – *Diagrammes en oeil (binaire bruité)*

FIG. 10.8 – *Diagrammes en oeil (PAM-4)*

Contents

1	Introduction aux Télécommunications.	1
1.1	Définition des télécommunications.	1
1.2	Historique.	2
1.3	La chaîne de transmission.	3
1.3.1	Schéma de base d'une chaîne de transmission.	3
1.3.2	Le canal de transmission.	4
1.3.3	L'émetteur.	8
1.3.4	Le récepteur.	9
1.3.5	Note sur le codage de source et le codage de canal (cas numérique).	10
1.4	Les organisations internationales de Télécommunications.	12
2	Rappels concernant Fourier et la convolution.	13
2.1	Définitions.	13
2.2	Théorème de convolution.	13
2.2.1	Filtrage et convolution.	13
2.3	Théorème de Parseval.	14
2.4	Théorème de modulation.	15
2.5	Transformée des fonctions périodiques.	16
2.6	Transformée d'une série d'impulsions.	17
2.7	Transformée d'un signal périodisé.	18
2.8	Problème de la fréquence image.	18
2.9	Le récepteur superhétérodyne	20
3	Variables aléatoires.	23
3.1	Rappels de théorie des probabilités.	23
3.1.1	Introduction	23
3.1.2	Théorie des probabilités.	23
3.1.3	Événements, expériences, axiomes de probabilité.	23
3.1.4	Indépendance statistique.	24
3.1.5	Lois de composition	24
3.1.6	Probabilités a posteriori	24
3.2	Variables aléatoires.	25
3.2.1	Fonction de répartition.	25
3.2.2	Densité de probabilité.	25
3.2.3	Moments d'une variable aléatoire.	26
3.2.4	Variables réelles à plusieurs dimensions.	27
3.2.5	Fonction caractéristique	29

3.3	Extension aux variables complexes	29
3.3.1	Fonction de répartition et densité de probabilité	29
3.3.2	Moments	29
3.4	Quelques lois de probabilité importantes	31
3.4.1	Loi à deux valeurs	31
3.4.2	Loi binomiale	31
3.4.3	Loi de Poisson	32
3.4.4	Loi uniforme	33
3.4.5	Loi de Gauss ou loi normale	34
3.4.6	Loi de Rayleigh	36
3.4.7	Loi de Rice	38
4	Fonctions aléatoires	39
4.1	Notion de fonction aléatoire.	39
4.2	Fonctions de répartition et densités de probabilité	40
4.3	Classification des fonctions aléatoires selon leurs propriétés statistiques.	40
4.3.1	Fonction aléatoire à valeurs indépendantes.	40
4.3.2	Fonction aléatoire à valeurs non corrélées ou orthogonales.	41
4.3.3	Fonction aléatoire additive.	41
4.3.4	Fonction aléatoire gaussienne	41
4.4	Moments d'une fonction aléatoire	41
4.4.1	Moyenne ou espérance mathématique	41
4.4.2	Variance. Covariance.	42
4.4.3	Propriétés des variances et covariances	42
4.4.4	Fonctions aléatoires non corrélées ou orthogonales.	43
4.5	Stationnarité et Ergodisme	43
4.5.1	Stationnarité.	43
4.5.2	Ergodisme.	45
5	Modulations analogiques : les principes.	49
5.1	Rôle de la modulation	49
5.2	La modulation d'amplitude	50
5.2.1	Modulation DSB (Double Side Band).	50
5.2.2	Modulation DSB-SC (Double Side Band - Suppressed Carrier)	55
5.2.3	Modulation à Bande Latérale Unique (BLU, SSB: Single Side-Band)	58
5.2.4	Démodulation d'un signal BLU	60
5.3	Modulation de Fréquence (FM)	61
5.3.1	Déviation de fréquence, indice de modulation	61
5.3.2	Spectre d'un signal FM	62
5.3.3	Modulation de fréquence à faible indice	62
5.3.4	Modulation FM par variation de paramètres	64
5.3.5	Modulateur d'Armstrong	65
5.3.6	Démodulateur FM	66
6	Performances des systèmes de modulation	69
6.1	Notion de signal équivalent passe-bas	69
6.1.1	Signal analytique	69
6.1.2	Signal équivalent passe-bas	70

6.1.3	Filtre équivalent passe-bas	70
6.1.4	Réponse d'un filtre passe-bande	71
6.2	Densité spectrale de puissance	72
6.2.1	Cas des signaux certains	72
6.2.2	Cas des signaux aléatoires	73
6.2.3	Théorème de Wiener-Kintchine	73
6.3	Représentation du bruit blanc	73
6.3.1	Bruit blanc passe-bas	73
6.3.2	Bruit blanc filtré passe-bande	74
6.3.3	Représentation Equivalent passe-bas de signaux stochastiques stationnaires	75
6.3.4	Propriétés des composantes en quadrature	77
6.3.5	Densités spectrales	77
6.4	Rapport signal/bruit en AM	78
6.5	Rapport signal/bruit en FM	79
6.5.1	Signal d'entrée	79
6.5.2	Signal de sortie	79
6.5.3	Bruit de sortie	79
6.5.4	Le facteur de mérite $\frac{S_i}{2\eta F}$	80
6.5.5	Comparaison des modulations	80
6.5.6	Effet de seuil en FM	81
6.6	Préaccentuation et désaccentuation en FM	81
7	Vers les communications numériques	85
7.1	Introduction	85
7.1.1	Les transmissions en plusieurs bonds.	86
7.1.2	Largeur de bande	87
7.2	Conversion analogique/numérique	87
7.2.1	L'échantillonnage	88
7.2.2	Formule d'interpolation de Whittaker	89
7.2.3	Commentaires	90
7.2.4	La quantification	91
7.3	Delta Modulation	95
7.3.1	Foreword	95
7.3.2	A reminder of PCM performances	96
7.3.3	Basic DM circuit	96
7.3.4	Signal to noise ratio for Gaussian messages	98
7.3.5	Signal-to-noise ratio for sinewaves	100
7.3.6	Driving the DM into slope overload	100
7.4	Delta modulation with double integration in the feedback loop	101
7.4.1	Signal to noise ratio for gaussian messages	101
7.4.2	Signal to noise ratio for sinewaves	103
7.5	Delta-sigma modulation (DSM)	103
7.5.1	Signal to main rate for Gaussian messages	104
7.5.2	Signal to noise ratio for sinewaves	105
7.6	DSM with an integrator in the feedback loop	106
7.6.1	Signal to noise ratio for Gaussian messages	106
7.6.2	Signal to noise ratio for sine waves	107
7.7	Summary of the previous results and conclusions	108

8	Modulations numériques : les signaux	113
8.1	Représentation des signaux numériques modulés	113
8.1.1	Modulations linéaires sans mémoire	113
8.1.2	Modulations non-linéaires à mémoire	117
8.2	Densités spectrales des signaux modulés	122
9	Modulations numériques : les performances.	125
9.1	Démodulateur optimal.	125
9.1.1	Récepteur à corrélation	127
9.1.2	Démodulateur à filtre adapté	128
9.2	Détermination du taux d'erreurs pour transmissions binaires	129
10	Communications dans un canal à bande limitée	133
10.1	Les canaux.	133
10.1.1	Le canal linéaire.	133
10.2	Le critère de Nyquist.	135
10.2.1	Notion d'Interférence Entre Symboles (ISI: Inter Symbol Interference). . .	135
10.2.2	Le critère de Nyquist.	136
10.2.3	Lien avec le théorème d'échantillonnage.	139
10.2.4	Le diagramme en oeil.	139