

REPUBLIQUE ALGERIENNE DEMOCRATIQUE ET POPULAIRE
Ministère de l'Enseignement Supérieur et de la Recherche Scientifique
Université Mouloud MAMMERI de Tizi-Ouzou



Faculté de Génie Electrique et Informatique
Département d'Electronique

Mémoire de fin d'études

Présenté en vue de l'obtention Du diplôme D'Ingénieur d'Etat en Electronique

Option : Communication

Thème

*Réduction du Facteur de Crête d'un
signal OFDM par la méthode PTS*

Proposé et dirigé par :

M^r AIT SAADI.H

M^r AIT BACHIR.Y

Réalisé par :

M^{lle} : BEN FERHAT Ouerdia

M^{lle} : BELABDI Samia

Promotion : 2009/2010

Remerciement

*Nous tenons à remercier notre promoteur
M^r Ait Bachir pour avoir accepté de nous
encadrer, et leur suivi tout le long de ce travail.*

*Nos vifs remerciements vont à M^r Ait Saadi et
M^r Zair pour l'aide précieuse qu'ils nous ont
apportée afin de mener à bien ce travail.*

*Nos remerciements iront également au président
et aux membres de jury, qui ont accepté de juger
notre travail.*

*Enfin, nous remercions tout le personnel du
département d'électronique en particulier nos
enseignants qui nous ont guidés tout au long de
notre cursus.*

Dédicaces

Je dédie ce modeste travail à :

Mes très chers parents qui se sont sacrifiés pour moi.

Mes frères et sœurs.

Tous ceux que j'aime et ceux qui m'aiment.

Querdia

Dédicaces

Je dédie ce modeste travail à :

A la lumière de ma vie, mes très chers parents.

*A mes très chères sœurs : Hayet, Fadila, Nabila et leurs
familles.*

*A mes très chères frères : Arezki, Toufik, Yacine,
Nadjim et sa famille.*

A toute la famille de BELABDI.

A mes meilleurs (es) amis (es).

A tous ceux que j'aime et qui m'aiment.

Samia

Sommaire

<i>Liste des figures</i>	1
--------------------------------	---

Liste des tableaux

<i>Introduction générale</i>	4
------------------------------------	---

Chapitre I : Généralités :

I.1. Introduction.....	6
I.2. Chaîne de transmission numérique	6
I.2.1. Codage source et codage canal.....	7
I.2.2. Codage binaire à symbole.....	7
I.2.3. Codage symboles à signal (CSS).....	9
I.2.4. Transposition de fréquence et amplification.....	10
I.2.5. Canal, réception et démodulation.....	10
I.2.6. Décodage symbole à signal.....	11
I.2.7. Décodage binaire à symboles, décodage de canal et décodage de source.....	11
I.3. Canaux de communication.....	12
I.3.1. Canal à bruit blanc gaussien.....	12
I.3.2. Canal à évanouissement.....	14
I.4. Modulations Numériques.....	16
I.4.1. La modulation à déplacement d'amplitude (MDA).....	18
I.4.2. modulation par déplacement de phase (MDP).....	19
I.4.3. Modulation d'amplitude en quadrature (MAQ).....	20
I.5. Conclusion.....	22

Chapitre II : La modulation multiporteuse OFDM

II.1. Introduction.....	23
II.2. Pourquoi l'OFDM ?.....	23
II.3. Principe des modulations multiporteuses.....	24
II.3.1. Notions d'orthogonalités.....	25
II.3.2. Principe de la modulation OFDM.....	26
II.3.3. Principe de démodulation.....	28
II.3.4. Réalisation numérique des opérations de modulation et de démodulation.....	29
II.3.4.1. Implantation numérique du modulateur.....	31
II.3.4.2. Implantation numérique de démodulateur.....	32
II.3.5. Interférence entre symboles OFDM.....	32
II.3.6. Interférences entre sous-canaux.....	34
II.4. Avantages et inconvénients.....	36
II.4.1. Avantages.....	36
II.4.2. Inconvénients.....	37
II.5. conclusion	37

Chapitre III : Méthodes de réduction du facteur de crête :

III.1. Introduction.....	38
III.2. Le facteur de crête.....	38
III.3. Caractéristiques d'un dispositif non linéaire.....	41
III.3.1. Les harmoniques	43
III.3.2. Les produits d'intermodulation.....	44
III.3.3. Le point de compression à 1db.....	46
III.4. Sources de non linéarités (Amplificateur de puissance).....	47

III.4.1. Caractérisation du la non-linéarité d'amplitude et de la phase d'amplificateur de puissance	47
III.4.2. Les types d'amplificateurs de puissance.....	51
III.5. Principe des méthodes de réduction du PAPR	52
III.5.1. Écrêtage plus filtrage.....	52
III.5.2. Tone Reservation.....	54
III.5.3. Le Selective Mapping (SLM).....	57
IV.5.4. La méthode “Partial Transmit Séquences ” (PTS).....	59
III.6. Conclusion	62
 <i>Chapitre IV: Application de la méthode PTS pour réduire le PAPR :</i>	
IV .1. Introduction.....	63
IV.2. Effets non linéaires de l'amplification.....	63
IV.2.1. Effet sur le spectre OFDM.....	64
IV.2.2. Effet sur le taux d'erreur binaire (BER).....	67
IV.3. Implémentation de la méthode « Partial Transmit Séquence » (PTS).....	70
IV.3.1. Application de la méthode PTS sous MATLAB.....	71
IV.3.2. Résultats de simulation.....	73
IV.4. Conclusion.....	74
<i>Conclusion générale .</i>	75

Liste des figures

Fig. I.1 : Schéma bloc d'une chaîne de transmission numérique.

Fig. I.2 : Exemple de constellations.

Fig. I.3 : Spectre avant et après transposition de fréquence signaux.

Fig. I.4 : Frontières des zones de décision sur les constellations MDP8 et MAQ16.

Fig. I.5 : Représentation d'un canal à BBAG.

Fig. I.6 : comparaison montrant un canal sélectif en fréquence en (a) et non sélectif en (b).

Fig. I.7 : Illustration du phénomène de trajets multiples sur le canal radio-mobil.

Fig. I.8. Forme générale du modulateur numérique.

Fig. I.9 : constellation des symboles en modulation de phases MDP-8.

Fig. I.9 : Constellation MAQ-16 et MAQ-64.

Fig.II.1: principe du non sélectivité en fréquence.

Fig II.2 : Base orthogonale temporelle.

Fig.II.3 : Base orthogonale fréquentielle.

Fig. II.4 : Schéma de principe d'un modulateur OFDM.

Fig.II.5 : allure de l'ensemble des spectres des porteuses d'un symbole OFDM.

Fig II.6 : Schéma de principe d'un démodulateur OFDM.

Fig.II.7 : Modulateur OFDM numérique.

Fig. II. 8 : Démodulateur OFDM numérique.

Fig. II.9 : Intervalle de garde.

Fig.II .10 : Interférence entre porteuses dans le domaine temporel (a) et fréquentiel (b).

Fig.II.12 : Schéma synoptique d'un émetteur et d'un récepteur OFDM.

Fig.III.1 : signaux multiporteuses (haut) et monoporteuse (bas).

Fig.III.2 : Effet de sur-échantillonnage.

Fig.III.3 : Caractéristique non-linéaire de transfert, f [.]

Fig.III.4 : Produit d'intermodulation .

Fig.III.5 : point de compression à 1 Db.

Fig.III.7 : caractéristique de la relation entrée-sortie d'un amplificateur de puissance.

Fig.III.8 : Calcul d'ACPR.

Fig.III.1 : Insertion de préfixe cyclique.

Fig. III .9 : Insertion du filtrage.

Fig.IV.10 : Schéma général pour le « Tone Reservation ».

Fig.III.11 : Impact de la Méthode « Tone Reservation » sur le « PAPR » ($N=512$).

Fig.III.12 : Schéma d'un modulateur « SelectiveMapping ».

Fig.III.13 : Schéma d'un modulateur « Partial Transmit Sequences ».

Fig.III.14 : Exemple de partitionnement d'un symbole en sous blocs pour application de la technique PTS.

Fig.IV.1 : Organigramme de simulation d'une chaîne OFDM (partie émission).

Fig.IV.2 : Effets non linéaires d'amplification sur le spectre pour différentes valeurs d'IBO.

Fig.IV.3 : Organigramme de simulation d'une chaîne OFDM (partie réception).

Fig. IV.4 : « BER » avant et après amplification en fonction de (E_b/N_0) .

Fig. IV.5 : Impact du non linéarité du « PA » modèle TWT illustré dans les constellations.

Fig. IV.6 : CCDF des symboles OFDM pour N variable.

Fig.IV.6 : Organigramme de simulation de la chaîne OFDM avec PTS.

Fig. IV.8 : Réduction du « PAPR » à l'aide de la méthode PTS.

Fig. IV.11 : Génération de nouveaux pics après le module de PTS.

Liste des tableaux

Tab. IV.1 : les paramètres de la modulation OFDM.

Tab. IV.2 : Mesure de l' « ACPR » pour différentes valeurs du recul « IBO ».

Tab. IV.3 : Les paramètres de la modulation OFDM.

Introduction

Les systèmes de télécommunications ont connus une évolution spectaculaire au cours de ces dernières années vu de développement de l'électronique numérique.

A l'heure actuelle des technologies permettant de transmettre et/ou d'accéder a un volume d'information très important, le plus vite possible avec une souplesse et mobilité sans limite, comme pour l'internet haut débit ADSL, les réseaux locaux sans fils : WIFI et WIMAX , ainsi que l'UMTS , un standard de la troisième génération (3G) qu'aujourd'hui on retrouve sur les « téléphones portables » permettant d'utiliser les services vidéo, internet, etc....

La transmission haute débit est possible grâce a l'utilisation de la modulation multiporteuse dite OFDM, présentant des avantages concernant la robustesse de signal vis-à-vis de canal multi-trajet avec évanouissement et encombrement spectrale optimal. Par contre, elle présente un principale inconvénient qui est représenté par les fortes fluctuations en amplitude de l'enveloppe de signal modulé est donc par des variations importantes de la puissance instantanée. Le PAPR qui tient compte ces variations, est un paramètre indispensable dans la caractérisation des modulations à enveloppe non constante.

La grande dynamique en amplitude et en puissance du signal OFDM a des effets sur les performances de l'amplificateur de puissance en émission (effet non linéaire). Par conséquent l'AP doit avoir un recul en entrée « IBO » suffisant pour ne pas saturer le signal OFDM à amplifier, mais dans ce cas l'amplificateur aura un rendement plus optimal on doit diminuer l'IBO ce qui veut dire, ce reprocher de la zone de saturation.

En ce qui concerne les méthodes de réduction du « PAPR » dans le cas de l'OFDM, il existe aujourd'hui différentes techniques. Certaines techniques agissent sur l'amplificateur afin d'éviter la saturation du signal d'entrée et d'autres techniques se basent plutôt sur un traitement réalisé directement au niveau du signal pour abaisser son « PAPR ». IL existe évidemment des conséquences pour chaque méthode comme la complexité, l'augmentation de puissance moyenne, la dégradation du taux d'erreur binaire, la remonte ou encore la diminution du débit utile. Parmi ces techniques, nous en avons retenue une : partial transmit séquences, cette technique réduisant le PAPR mais l'inconvénient principale de la technique PTS réside dans la complexité de la recherche des vecteurs de phases $\Phi^{(v)}$ pour minimiser le PAPR.

Ce rapport se représente sur quatre chapitres comme suite :

Le chapitre I : présente les éléments de communication numériques nécessaire à la bonne compréhension de ce travail, les différents canaux de communication et les modulations numériques.

Le chapitre II : nous allons présenter la modulation OFDM qui est la modulation multiporteuse sur la quelle porte ce travail. Le chapitre commence par le principe de l'OFDM puis le passage du modulateur OFDM analogique au numérique..

Le chapitre III : consiste à étudier le facteur de crête, la non linéarité commençons par des principaux éléments caractérisant un dispositif-non linéaire, l'amplificateur de puissance ainsi que les techniques de réduction du « PAPR » qui font l'objet de ce chapitre, parmi les nombreuses techniques existantes, nous avons alors choisi quatre techniques qui agissent sur le signal d'entrée pour obtenir une réduction des fluctuations d'enveloppe. qui sont : « *l'écrêtage classique plus filtrage* », « *tone réservation* », « *sélective mapping* » et « *partial transmit sequences* » qui est l'objet de ce mémoire.

Le chapitre IV : c'est le dernier chapitre où nous présentons un ensemble des simulations et interprétations.

Et enfin, nous clôturons ce travail par une conclusion générale.

I.1. Introduction :

L'objectif de ce chapitre est d'introduire quelques généralités sur les communications numériques qui serviront à la bonne compréhension. Nous allons dans un premier temps décrire le fonctionnement d'une chaîne de transmission numérique de la source d'information binaire au destinataire, par les étapes successives de codage de modulation, d'amplification, de transmission dans un canal physique et démodulation.

Après avoir indiqué brièvement le fonctionnement d'une chaîne de transmission et le rôle des principaux éléments nous introduirons les canaux de transmission et les types de modulations numériques.

I.2. Chaîne de transmission numérique :

Les systèmes de transmission numérique véhiculent de l'information entre une source et un destinataire en utilisant un support physique comme le câble, la fibre optique ou encore la propagation sur un canal radioélectrique.

Les signaux transportés peuvent être soit directement d'origine numérique, comme dans les réseaux de données, soit d'origine analogique (parole, image) mais convertis sous une forme numérique. Le principe du système de transmission est alors d'acheminer l'information de la source vers le destinataire avec le plus de fiabilité possible.

Les différentes étapes seront explicitées successivement dans ce chapitre. La figure I.1 résume l'ensemble de ces étapes dans le schéma bloc d'une chaîne de transmission numérique.

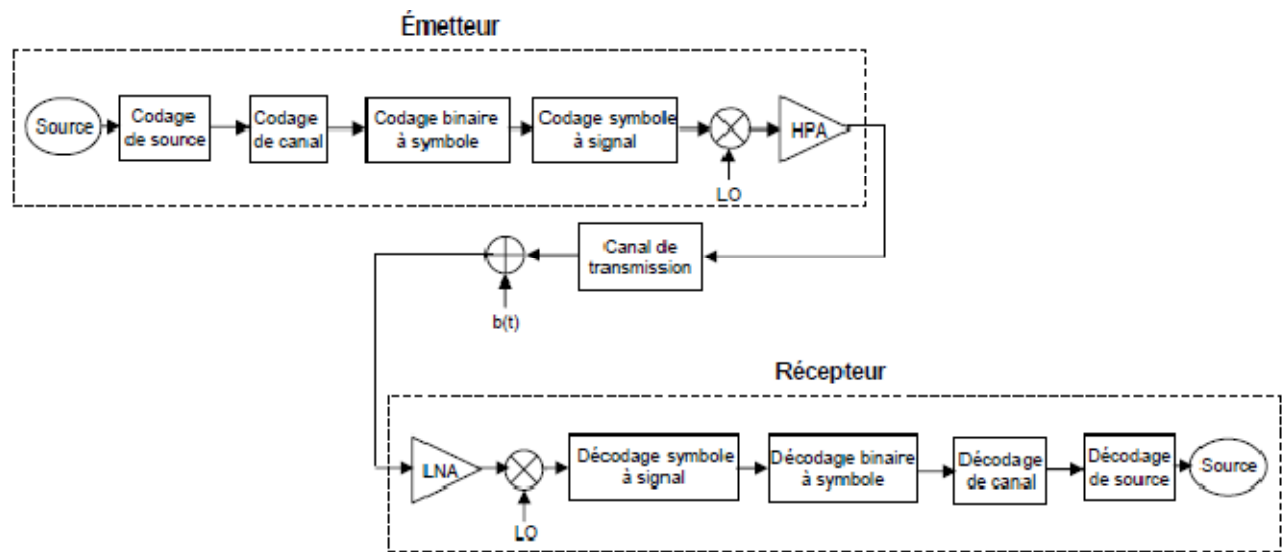


Fig. I.1 : Schéma bloc d'une chaîne de transmission numérique

I.2.1. Codage source et codage canal :

Le but du codage de source est de représenter la source, qu'elle soit analogique ou numérique avec un minimum de bits sans diminuer la quantité d'information. Cette séquence binaire en sortie de codeur de source est appelée séquence d'information qui est caractérisé par ce qu'on appelle le débit numérique $D_b = \frac{1}{T_b} \text{ bit/s}$, [où T_b est le temps bit en s].

Une transmission est dite synchrone, si l'émission d'éléments binaires délivrés par la source s'effectue à cadence constante. Si la cadence est variable elle est dite asynchrone.

En pratique des erreurs peuvent se produire durant la communication, et elles sont principalement dues au bruit et aux interférences produites par le canal de transmission lui-même. Pour y remédier, on utilise un codage correcteur d'erreurs : des bits de redondance sont ajoutés aux informations numériques à transmettre, et ceux-ci permettent au récepteur de détecter et/ou corriger des erreurs.

I.2.2. Codage binaire à symbole :

Le codage binaire à symbole est l'étape qui associe les éléments binaires à des symboles C_k appelés symboles numériques, il est constitué de n éléments binaires.

Un codage binaire peut se représenter de manière graphique, appelée constellation, dont chaque point correspond à un symbole c_k , à côté duquel on indique éventuellement la donnée numérique que le symbole code. Par exemple les constellations des codages MDP8 et MAQ16 peuvent être représentées de la forme illustrée sur la figure I.2

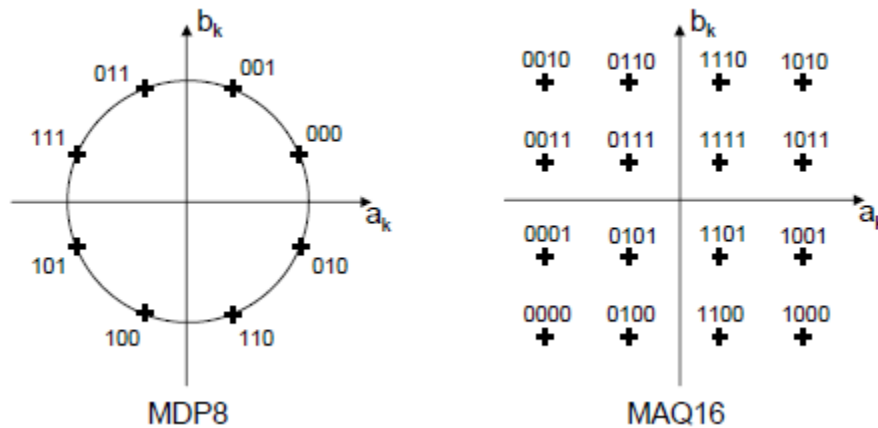


Fig. I.2 : Exemple de constellations

Le choix de la répartition des points dépend des critères suivants :

- Pour pouvoir distinguer deux symboles, il faut respecter une distance minimale d_{min} , entre les points représentatifs de ces symboles. Plus cette distance est grande et plus la probabilité d'erreur sera faible. La distance minimale entre tous les symboles est :

$$d_{min} = \min_{i \neq j} d_{ij} \quad (\text{I.1})$$

$$\text{avec } d_{ij} = |c_i - c_j|^2$$

- A chaque symbole émis correspond un signal élémentaires $m_k(t)$ et par là même une énergie nécessaire à la transmission de ce symbole. Dans la constellation, la distance entre un point et l'origine est proportionnelle à la racine carrée de l'énergie qu'il faut fournir pendant l'intervalle de temps $[kT, (k+1)T[$ pour émettre ce symbole.

Supposons que la source délivre des éléments binaires toutes les T_b secondes, la période symbole est définie par $T_s = nT_b$. La rapidité de modulation est $R = \frac{1}{T_s}$ s'exprime en bauds et correspond au nombre de changements d'états par seconde d'un ou de plusieurs paramètres modifiés simultanément.

I.2.3. Codage symboles à signal (CSS) :

Le canal de transmission étant un milieu continu avant de pouvoir y transmettre les symboles C_k il faut obtenir un signal continu par interpolation. Les symboles sont cadencés par une horloge à la fréquence $F_s = \frac{1}{T_s}$ où T_s est la durée d'un symbole, et une forme d'onde $h(t)$ permet d'interpoler le signal discret. $h(t)$ est une fonction non nulle sur $[0, T_s[$ et comme son nom l'indique donne la forme au signal continu :

$$C(t) = \sum_{k=0}^{\infty} C_k h(t - kT_s) \quad (\text{I.2})$$

Une forme d'onde classique est tout simplement le rectangle de durée T_s :

$$h(t) = \begin{cases} 1 & 0 \leq t < T_s \\ 0 & \text{ailleurs} \end{cases} \quad (\text{I.3})$$

le signal ainsi généré à un spectre infini .En pratique il est impossible d'utiliser tout le spectre dans un canal réel, et il est alors nécessaire d'ajouter après cette forme d'onde un filtre basse fréquence qui servira à limiter la bande passante de signal émis .Ce filtre est appelé filtre d'émission, il est placé après la forme d'onde.

I.2.4. Transposition de fréquence et amplification :

La transposition de fréquence est nécessaire dans le cas d'une transmission radio. En effet un canal radio est caractérisé par une bande de fréquence précise, et afin de ne pas perturber les communications sur les autres canaux radio, il faut s'assurer que la transmission n'utilise que cette bande de fréquence. La largeur de cette bande Δf est souvent faible devant sa fréquence centrale f_0 , et ainsi le signal qui est propagé est dit à bande étroite. Le signal provenant du filtre d'émission est quand à lui un signal basse fréquence, dit signal en bande de base. La modulation où transposition de fréquence consiste donc à décaler la fréquence centrale du signal pour respecter les caractéristiques imposées par le canal. La figure I.3 montre la forme des densités spectrales de puissance (DSP) du signal avant et après transposition de fréquence.

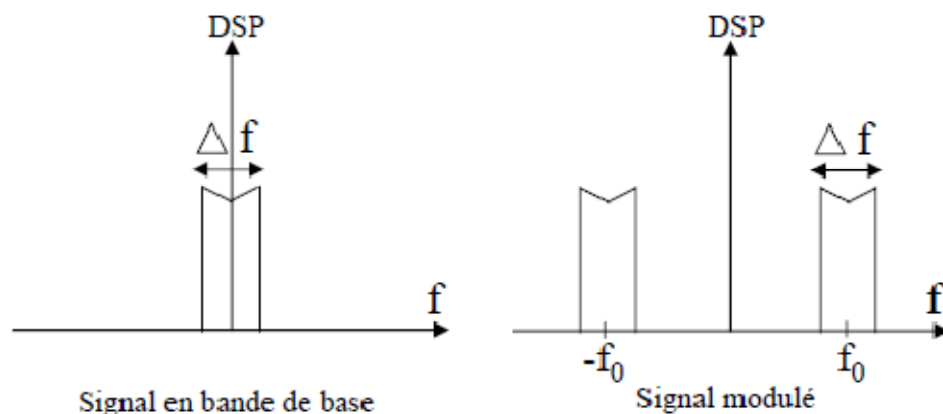


Fig. I.3 : Spectre avant et après transposition de fréquence signaux

Un amplificateur est enfin nécessaire pour augmenter la puissance du signal afin que son niveau soit suffisant à la réception, compensant ainsi les pertes en espace libre. Dans le cas d'un canal radio, l'amplificateur est relié à une antenne qui rayonne et crée ainsi le signal radio. Dans le cas d'un canal filaire l'amplificateur est relié au câble.

I.2.5. Canal, réception et démodulation :

Le canal de propagation perturbe le signal, en le déformant et en y ajoutant du bruit nous supposons dans un premier temps que le canal est sans bruit, le récepteur recueille le signal transmis par l'intermédiaire d'une antenne pour un canal radio ou directement depuis

le câble pour une transmission filaire. Une fois le signal est ré-amplifié, il est nécessaire de le démoduler, c'est-à-dire de faire une nouvelle transposition de fréquence afin d'obtenir un signal en bande de base.

I.2.6. Décodage symbole à signal :

Le signal démodulé est un signal continu, mais le récepteur va devoir réaliser un échantillonnage afin de déterminer les éléments binaires transmis. Cependant avant l'échantillonnage il faut réaliser un filtre adaptée à l'émetteur pour une réception optimale des symboles transmis. Dans le cas où aucun filtre d'émission n'est employé, le filtre de réception adapté à la forme d'onde $h(t)$ à pour réponse impulsionnelle :

$$g_r(t) = h(t_0 - t) \quad (\text{I.4})$$

Où t_0 est l'instant d'échantillonnage. Le système de réception est très simple dans ce cas, car le signal reçue à un instant donné après ce filtrage correspond directement à un unique symbole, et celui-ci peut être décodé. Par contre lorsqu'un filtre d'émission est utilisé la réponse de filtre est généralement plus longue que T_s et le signal reçu à un instant t ne dépend plus d'un seul symbole émis, mais également des autres symboles. Ce phénomène est appelé interférence entre symboles (ISI). Pour annuler ISI il faut qu'à l'instant d'échantillonnage on ne prélève que le symbole émis, et annuler l'influence due aux autres symboles.

I.2.7. Décodage binaire à symboles, décodage de canal et décodage de source :

L'étape suivante consiste à déterminer les bits correspondant au symbole reçu d_k après le décodage à symbole. Ce symbole peut être différent du symbole qui avait été envoyé C_k à cause des perturbations introduites par le canal. La détection par maximum de vraisemblance est le critère optimal permettant de déterminer le symbole qui a été envoyé avec la plus grande probabilité. Pour cela on sélectionne le point de la constellation le plus proche du symbole reçu, et les bits qui sont associés à ce point de la constellation sont les bits qui ont été émis avec la plus grande vraisemblance. Le plan complexe est ainsi partitionné en zones

de décision, chacune correspondant à un symbole de la constellation et donc un ensemble de bits particulier.

On peut représenter les limites de ces zones par des traits pointillés (on suppose que tous les symboles sont équiprobables) comme le montre la figure I.4.

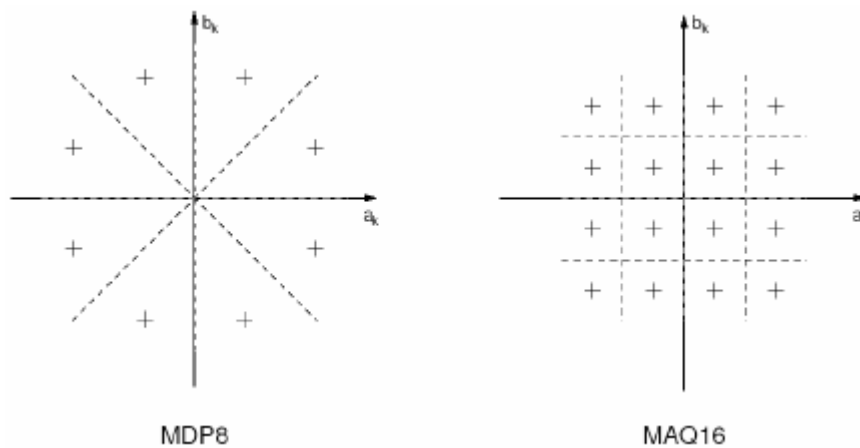


Fig. I.4 : Frontières des zones de décision sur les constellations

MDP8 et MAQ16

Le signal décidé, sous forme binaire, sera décodé grâce au décodeur canal. Ce décodeur correspond au codeur canal qui a été utilisé dans l'émetteur pour ajouter de la redondance aux informations transmises. Cette redondance est utilisée par le décodeur canal pour détecter des erreurs dans le flux binaire et éventuellement les corriger. Dans le cas d'un système FEC (Forward Error Correction) les erreurs sont corrigées directement par le décodeur, et dans le cas d'un système ARQ (Automatique Repeat ReQuest). Le message numérique résultant est finalement passé à travers le décodeur de source rendant ainsi un signal compatible avec son traitement par le destinataire.

I.3. Canaux de communication :

I.3.1. Canal à bruit blanc gaussien :

Le modèle de canal le plus fréquemment utilisé pour la simulation de transmission numérique, qui est aussi un des plus faciles à générer et analyser, est le canal à bruit blanc additif gaussien (BBAG) (la figure I.5). Ce bruit est généré par des signaux parasites

transitant sur le même canal et par le bruit thermique des composants électroniques. Le signal reçu s'écrit sous la forme :

$$y(t) = x(t) + b(t) \quad (\text{I.5})$$

Où b représente le BBAG, caractérisé par un processus aléatoire gaussien de moyenne nulle, de variance σ_b^2 , et de densité spectrale de puissance bilatérale $\frac{N_0}{2}$.

La densité de probabilité conditionnelle de y est donnée par l'expression :

$$p(y/x) = \frac{1}{\sqrt{2\pi\sigma_b^2}} e^{-\frac{(y-x)^2}{2\sigma_b^2}} \quad (\text{I.6})$$

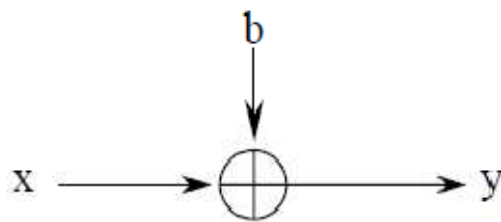


Fig. I.5 : Représentation d'un canal à BBAG

- **Le rapport signal à bruit :**

Lorsqu'il s'agit d'un bruit qui s'ajoute au signal étudié $x(t)$, le rapport signal à bruit SNR (Signal to Noise Ratio) est le rapport de la puissance moyenne du signal $y(t)$ et de la puissance moyenne de bruit $b(t)$, il est donné par l'équation I.7 :

$$SNR = \frac{E_b}{N_0} + 10 \log_{10}(m) \quad (\text{I.7})$$

$\frac{E_b}{N_0}$: Energie par bit à la densité de bruit et m est le nombre de bit par symbole

Ce rapport caractérise la performance du récepteur. Plus il est grand, moins le bruit perturbe le signal d'où la diminution de la probabilité d'erreur P_e par bit, qui est une valeur théorique indiquant une estimation de taux d'erreur binaire « BER » (Bit Error Rate)

$$BER = \frac{\text{Nombre de bits faux}}{\text{Nombre de bits transmis}} \quad (\text{I.8})$$

I.3.2. Canal à évanouissement :

Les communications radio ont souvent besoin d'un modèle plus élaboré prenant en compte les différences de propagation du milieu, appelées encore atténuations ou évanouissements, qui affectent la puissance du signal. Cette atténuation du signal est principalement due à un environnement de propagation riche en échos et donc caractérisé par de nombreux multi-trajets, mais aussi au mouvement relatif de l'émetteur et de récepteur entraînant des variations temporelles du canal. En ce qui concerne les variations temporelles du canal on peut distinguer deux classes, l'étalement temporel et l'effet Doppler.

- **Définition de l'étalement temporel :**

Le récepteur radio-mobile dispose de plusieurs répliques du signal émis, issues de trajets différents et retardés les uns par rapport aux autres. Le temps séparant l'arrivée de premier trajet de l'arrivée du dernier noté T_m est appelé étalement temporel maximal. Pour caractériser la dispersion temporelle du canal, on utilise la notion de bande de cohérence du canal, notée B_c définit comme la gamme de fréquence dans laquelle les amplitudes des différentes composantes fréquentielles subissent des atténuations semblables. La bande de cohérence du canal est de même ordre de grandeur que l'inverse de l'étalement temporel maximal: $B_c \sim \frac{1}{T_m}$.

Notons T_s La durée symbole de signal et w la bande de fréquence occupée par le signal. Si w est très inférieur à la bande cohérence du canal B_c (ou $T_s \gg T_m$) alors toutes les composantes fréquentielles du signal subissent des atténuations semblables : le canal est dit non sélectif en fréquence (la figure I.6). Dans le cas contraire, ceci se traduit par la présence d'interférence entre symboles (ISI). En pratique, on cherche à rendre $w \ll B_c$ afin d'éviter ce phénomène.

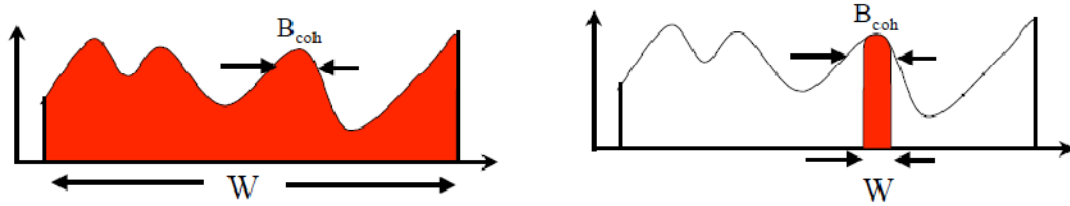


Fig. 1.6 : comparaison montrant un canal sélectif en fréquence en (a) et non sélectif en (b)

- **Effet Doppler :**

Lorsque un récepteur se déplace à une vitesse v , la fréquence reçue est modifiée d'une quantité f_D , dite fréquence Doppler, fonction à la fois de longueur d'onde émise et la vitesse de déplacement du mobile :

$$f_D = \frac{v}{\lambda} \cos \theta \quad (\text{I.9})$$

L'angle θ est pris entre la direction de v et celle suivie par l'onde électromagnétique émise.

Le signe de f_D est déterminé par le $\cos \theta$ et il est positif si le récepteur s'éloigne de l'émetteur et négatif s'il s'en approche.

Une fois définie la fréquence Doppler f_D , nous pouvons caractériser le canal en termes de ses variations temporelles en introduisant le temps de cohérence T_c :

$$T_c \approx \frac{1}{f_D} \quad (\text{I.10})$$

Ce paramètre représente l'intervalle temporel durant lequel les distorsions sont négligeables.

Ainsi, un canal à évanouissement temporel rapide possède un T_c faible. Pour éviter toute sélectivité temporelle sur symbole émis, il faudrait donc avoir un temps de cohérence nettement supérieur au temps symbole T_s ($T_c \gg T_s$).

- **Canal à multi trajets :**

Ce canal est illustré dans la figure 1.7. Nous considérons que le canal subit des évanouissements lents, c'est-à-dire que la durée d'un symbole est très inférieure au temps de cohérence du canal, et que le signal reçu ne varie donc pas ou très peu sur la durée d'un symbole. En tenant compte du bruit blanc additif gaussien, le signal équivalent en bande de

base reçu à la sortie de ce canal à évanouissements lents comportant N_p trajets multiples s'exprime alors :

$$y(t) = \sum_{n=0}^{N_p-1} \alpha_n s(t - \tau_n) + b(t) \quad (\text{I.11})$$

Ou le bruit BBAG complexe est représenté par $b(t)$, et α_n et τ_n caractérisent respectivement l'atténuation complexe et le retard affectant chaque trajet.

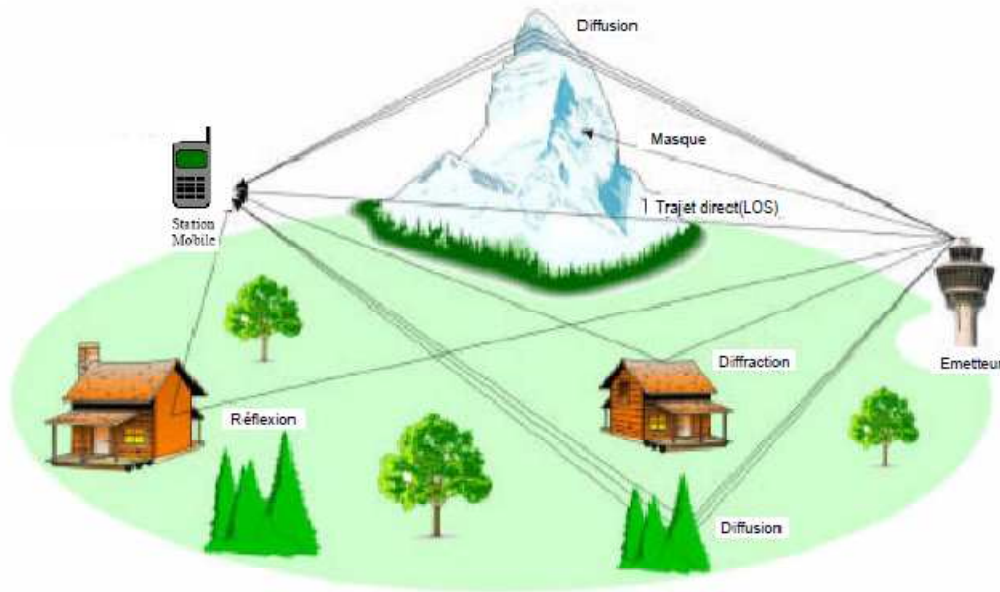


Fig. I.7 : Illustration du phénomène de trajets multiples sur le canal radio-mobil

I.4. Modulations Numériques :

La modulation a pour objectif d'adapter le signal à émettre au canal de transmission. Dans le cas de transmission sur porteuse, l'opération consiste à modifier un ou plusieurs paramètres d'une onde porteuse de forme sinusoïdale d'expression général $s(t) = A \cos(w_0 t + \varphi_0)$ centrée sur la bande de fréquence de canal, les paramètres modifiables dans cette expression sont :

- l'amplitude de l'onde : A .
- la fréquence porteuse : $f_o = \frac{w_o}{2\pi}$
- la phase : φ_o

Dans les procédés de modulation binaire, l'information est transmise à l'aide d'un paramètre qui ne prend que deux valeurs possibles.

Dans les procédés de modulation M-aire, l'information est transmise à l'aide d'un paramètre qui prend M valeurs. Ceci permet d'associer à un état de modulation un mot de n digits binaires. Le nombre d'états est donc $M = 2^n$. Ces n digits proviennent du découpage en paquets de n digits du train binaire issu du codeur.

❖ Principes de modulations numériques :

Le signal modulant, obtenu après codage, est un signal en bande de base, éventuellement complexe, qui s'écrit sous la forme :

$$c(t) = \sum_k c_k \cdot h(t - kT) = \sum_k c_k(t) \quad \text{avec} \quad c_k(t) = a_k(t) + jb_k(t) \quad (\text{I.12})$$

La modulation transforme ce signal $c(t)$ en un signal modulé $m(t)$ tel que :

$$m(t) = \text{Re}[\sum_k c_k(t) \cdot e^{j(\omega_0 t + \varphi_0)}] \quad (\text{I.13})$$

Si les $c_k(t) = a_k(t) + jb_k(t)$ sont réels ($b_k(t)=0$), la modulation est dite unidimensionnelle, et s'ils sont complexes la modulation est dite bidimensionnelle.

Le signal modulé s'écrit aussi plus simplement :

$$m(t) = a(t) \cdot \cos(\omega_0 t + \varphi_0) - b(t) \cdot \sin(\omega_0 t + \varphi_0) \quad (\text{I.14})$$

Le schéma théorique du modulateur est représenté sur la figure I.8.

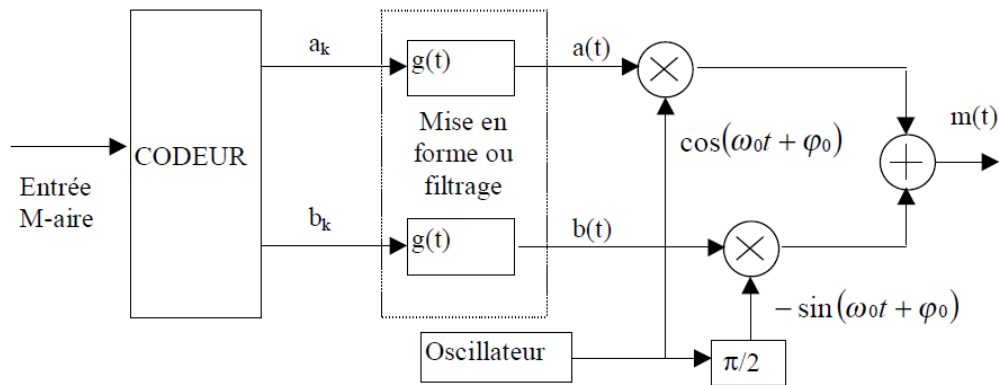


Fig. I.8. Forme générale du modulateur numérique

Les types de modulations les plus fréquents rencontrés sont :

I.4 .1. La modulation à déplacement d'amplitude (MDA) :

Les Modulations par Déplacement d'Amplitude (MDA) sont aussi souvent appelées par leur abréviation anglaise : ASK pour "Amplitude Shift Keying".

Dans ce cas, la modulation ne s'effectue que sur la porteuse en phase $\cos(\omega_0 + \varphi_0)$. Il n'y a pas de porteuse en quadrature. Cette modulation est parfois dite mono dimensionnelle. Le signal modulé s'écrit alors :

$$m(t) = \sum_k a_k g(t - kT) \cdot \cos(\omega_0 t + \varphi_0) \quad (\text{I.15})$$

Ce type de modulation est simple à réaliser mais assez peu employé pour $M > 2$ car ces performances sont moins bonnes que celles d'autres modulations, notamment pour sa résistance au bruit.

I.4.2. modulation par déplacement de phase (MDP) :

Les Modulations par déplacement de phase (MDP) sont aussi souvent appelées par leur abréviation anglaise : PSK pour "Phase Shift Keying". Les signaux élémentaires $a_k(t)$ et $b_k(t)$ utilisent la même forme d'onde $h(t)$ qui est ici une impulsion rectangulaire, de durée T et d'amplitude égale à A si t appartient à l'intervalle $[0, T[$ et égale à 0 ailleurs. Dans le cas présent, les symboles c_k sont répartis sur un cercle, et par conséquent :

$$c_k = a_k + jb_k = e^{j\varphi_k} \quad (\text{I.16})$$

$$\text{d'où : } a_k = \cos(\varphi_k) \quad b_k = \sin(\varphi_k)$$

$$\text{et : } a_k(t) = \cos(\varphi_k) \cdot g(t - kT) \quad b_k(t) = \sin(\varphi_k) \cdot g(t - kT)$$

On pourrait imaginer plusieurs MDP-M pour la même valeur de M où les symboles seraient disposés de façon quelconque sur le cercle ! Pour améliorer les performances par rapport au bruit, on impose aux symboles d'être répartis régulièrement sur le cercle. L'ensemble des phases possibles se traduit alors par les expressions suivantes :

$$\varphi_k = \frac{\pi}{M} + k \frac{2\pi}{M} \quad \text{lorsque } M > 2 \quad \text{avec } k = 0, 1, \dots, M-1$$

$$\text{et : } \varphi_k = 0 \text{ ou } \pi \quad \text{lorsque } M = 2$$

On peut considérer que a_k et b_k prennent simultanément leurs valeurs dans l'alphabet $\{\cos(\varphi_0)\}$ et $\{\sin(\varphi_0)\}$.

Le signal modulé devient :

$$m(t) = \text{Re}[\sum_k e^{j\varphi_k} \cdot g(t - kT) \cdot e^{j(\omega_0 t - \varphi_0)}] = \text{Re}[\sum_k g(t - kT) \cdot e^{j(\omega_0 t - \varphi_0 + \varphi_k)}] \quad (\text{I.17})$$

Soit, plus simplement en ne considérons que l'intervalle de temps $[kT, (k+1)T[$:

$$m(t) = \text{Re}[A e^{j(\omega_0 t - \varphi_0 + \varphi_k)}]$$

$$m(t) = A \cdot \cos(\omega_0 t + \varphi_0) \cos(\varphi_k) - A \cdot \sin(\omega_0 t + \varphi_0) \sin(\varphi_k) \quad (\text{I.18})$$

Cette dernière expression montre que la phase de la porteuse est modulée par l'argument φ_k de chaque symbole ce qui explique le nom donné à la MDP.

On appelle MDP-M une modulation par déplacement de phase (MDP) correspondant à des symboles M-aires. La figure I.9 montre la constellation de MDP pour M= 8.

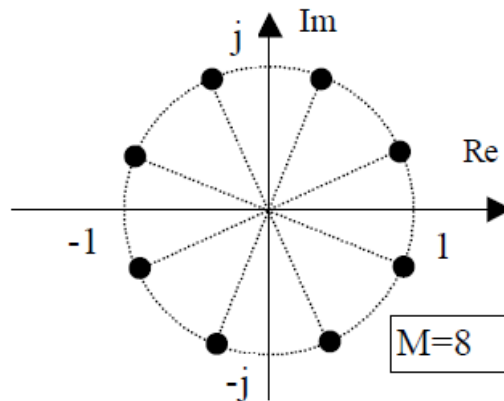


Fig. I.9 : constellation des symboles en modulation de phases MDP-8

La complexité de l'ensemble émission réception de la MDP augmente avec M, mais reste raisonnable ce qui on fait une modulation fréquemment utilisée pour M allant de 2 à 16 avec de bonnes performances dans les inconvénients de la MDP citons l'existence de sauts de phase importants qui font apparaître des discontinuités d'amplitude.

I.4.3. Modulation d'amplitude en quadrature (MAQ) :

Les modulations précédentes ne constituent pas une solution satisfaisantes pour utiliser efficacement l'énergie émise lorsque le nombre de point M est grand.

En effet dans la MDA les points de la constellation sont sur une droite et dans la MDP les points sont sur un cercle. Or la probabilité d'erreur est fonction de la distance minimale entre les points de la constellation, et la meilleure modulation est celle qui maximise cette distance pour une puissance moyenne donnée.

Pour ce faire on écrit le signal modulé $m(t)$ sous la forme suivante :

$$m(t) = a(t) \cdot \cos(\omega_0 t + \varphi_0) - b(t) \cdot \sin(\omega_0 t - \varphi_0) \quad (\text{I.19})$$

Les deux signaux $a(t)$ et $b(t)$ ont pour expression :

$$a(t) = \sum_k a_k \cdot g(t - kT) \quad \text{et} \quad b(t) = \sum_k b_k \cdot g(t - kT) \quad (\text{I.20})$$

Le signal modulé $m(t)$ est donc la somme de deux porteuses en quadrature, modulées en amplitude par les deux signaux $a(t)$ et $b(t)$.

Les symboles a_k et b_k prennent respectivement leur valeurs dans l'alphabet $\{\pm 3d, \pm 5d, \dots, \pm (M-1)d\}$ ou d est une constante, donnant ainsi naissance à une modulation possédant $E = M^2$ états. Chaque état est donc représenté par un couple (a_k, b_k) ou ce qui revient au même par un symbole complexe $c_k = a_k + jb_k$

Dans le cas particulier mais très fréquent où M peut s'écrire $M = 2^n$. Alors les a_k représentent un mot de n bit et les b_k représentent un mot de n bit. Le symbole complexe $c_k = a_k + jb_k$ peut par conséquent représenter un mot de $2n$ bit. Intérêt de cette configuration est que le signal $m(t)$ est alors obtenu par une combinaison de deux porteuses en quadrature modulées en amplitude. Cette modulation prend naturellement le nom de la modulation d'amplitude en quadrature (MAQ) et si sa constellation comporte E états, on la note MAQ- E .

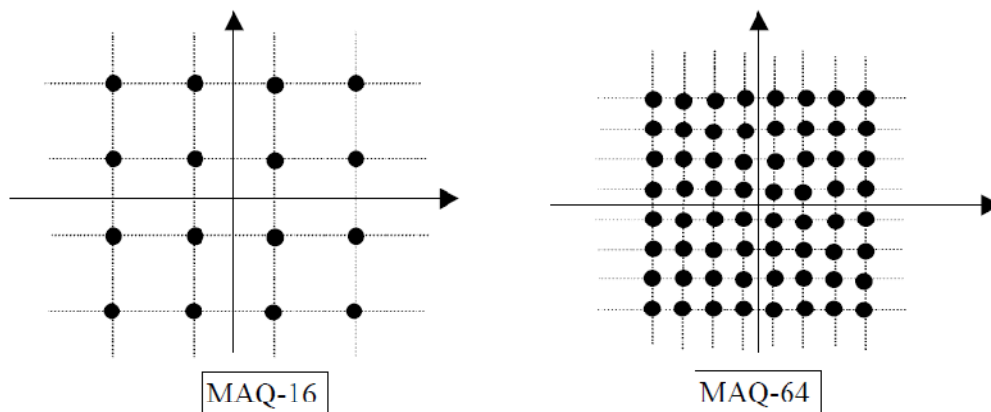


Fig. I.9 : Constellation MAQ-16 et MAQ-64

En effet, pour le modulateur le train binaire entrant $\{i_k\}$ est facilement divisé en deux trains $\{a_k\}$ et $\{b_k\}$.

La réception d'un signal MAQ fait appel à une démodulation cohérente et par conséquent nécessite l'extraction d'une porteuse synchronisée en phase et en fréquence avec la porteuse à l'émission. Le signal reçu est démodulé dans deux branches parallèles, sur l'une avec la porteuse en phase sur l'autre avec la porteuse en quadrature. Les signaux démodulés sont convertis par deux convertisseurs analogiques numériques (CAN) puis une logique de décodage détermine les symboles et régénère le train de bits reçus.

I.5. Conclusion :

Dans ce chapitre, nous avons introduit le concept des communications numérique ; des généralités ont été données, fournissant une bonne compréhension de l'ensemble de ce mémoire.

A l'heure actuelle, les nouvelles technologies cherchent toujours à améliorer les performances des systèmes de communication, surtout en termes de qualité, vitesse de transmission et quantité de données à transmettre dans la plus petite bande de fréquence possible. Ces nouvelles technologies comme à titre d'exemple WIMAX, ADSL, DVB, DVB-T et WIFI utilisent la modulation multiporteuse OFDM.

Cette modulation OFDM présente des avantages admirables permettant une meilleure transmission des données. Afin de comprendre son principe, ses avantages et ses inconvénients, un chapitre complet lui en consacre.

II.1. Introduction :

Ce chapitre a pour but d'introduire les modulations multiporteuses, et en particulier la modulation par répartition orthogonale des fréquences (OFDM) qui est une technique de transmission très performante pour les réseaux sans fil à hauts débits numériques.

Les techniques multiporteuses consistent à transmettre des données numériques en les modulant sur un grand nombre de porteuses en même temps. Le gain d'intérêt actuel réside dans l'amélioration apportée pour augmenter l'efficacité spectrale en orthogonalisant les porteuses ce qui permet d'implémenter la modulation et la démodulation à l'aide des circuits performants de transformé de Fourier rapide (FFT).

II.2. Pourquoi l'OFDM ?

L'un des plus grands problèmes qu'on rencontre lors d'une transmission haut débit est le cas d'un canal sélectif en fréquence, l'amplification du signal dans certaines bandes de fréquences et l'atténuation dans d'autres bandes. Comme un très grand débit impose une grande bande passante et si cette bande passante couvre une partie du spectre comportant des creux (due à la sélectivité fréquentielle du canal à trajet multiples), il y a perte totale de l'information pour la fréquence correspondante.

L'OFDM apparaît comme une solution admirable afin de lutter contre évanouissement. L'idée de l'OFDM est de répartir l'information sur un grand nombre de porteuses, créant ainsi des sous canaux très étroits pour lesquels la réponse fréquentielle du canal peut être considérée comme constante; sachant que chaque sous bande est plus petite que la bande de cohérence du canal, ce qui implique la variation de chaque sous porteuse dans le canal est linéaire.

Ainsi, pour ces sous canaux, le canal est non-sélectif en fréquence et s'il y a un creux, il n'affectera que certaines sous porteuses qui pourront être récupérées grâce à une égalisation simple à réaliser en plus de codage et de l'entrelacement astucieux.

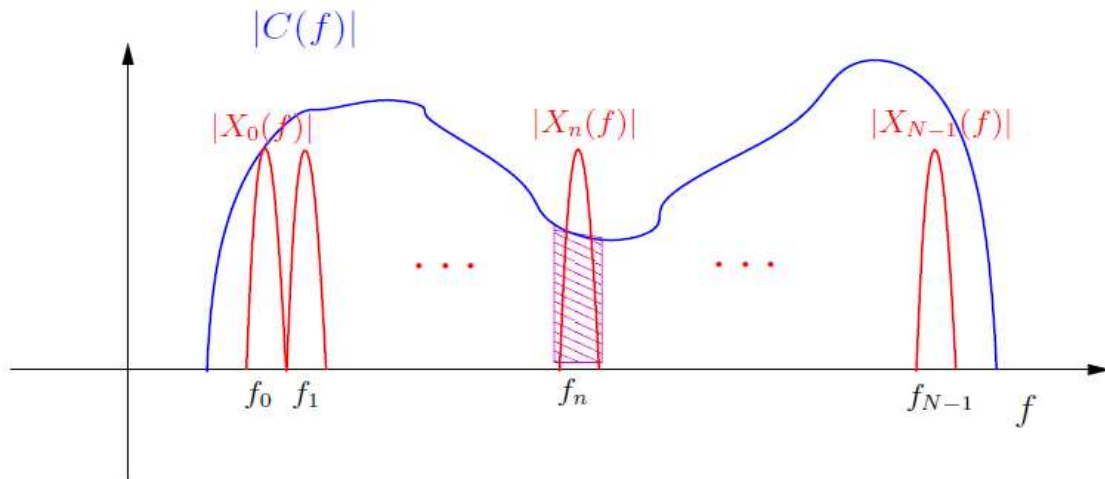


Fig. II.1: principe du non sélectivité en fréquence

Pour obtenir une efficacité spectrale optimale, il faut que les sous-bandes se chevauchent tout en restant orthogonales pour faciliter la séparation. On donne à présent les notions de base sur l'orthogonalité afin d'entamer le procédé OFDM.

II.3. Principe des modulations multiporteuses :

La modulation multiporteuse permet de simplifier le problème de l'égalisation dans le cas d'un canal sélectif en fréquence, le principe est de transmettre simultanément plusieurs symboles en parallèle sur différentes porteuses. En modulant sur N porteuses, il est possible d'utiliser des symboles N fois plus longs tout en conservant le même débit qu'avec une modulation monoporteuse ; En choisissant une valeur assez grande pour N , la durée des symboles deviennent grande devant l'étalement des retards, et les perturbations liées aux échos deviennent négligeables.

On définit l'efficacité spectrale comme étant le débit binaire par unité de fréquence. Plus l'efficacité spectrale est importante, plus il sera possible de transmettre un débit important sur un canal donné. Le choix des porteuses et leur écartement va influencer sur cette

efficacité spectrale. Pour garder la même efficacité qu'avec la modulation monoporteuse équivalente, il faut choisir soigneusement les fréquences des porteuses utilisées. La méthode la plus répandue est l'utilisation de porteuses orthogonales

II.3.1. Notions d'orthogonalités

Deux fonctions $f(t)$ et $g(t)$ sont orthogonales sur $[a, b]$ si :

$$\int_a^b f(t)g(t)dt = 0 \quad (\text{II.1})$$

Où l'intégrale définit le produit scalaire de deux fonctions et l'intervalle $[a, b]$ représente le domaine sur lequel porte l'étude. Dans ses conditions, ces deux fonctions sont disjointes sur le segment $[a, b]$ et l'interfèrent donc pas l'une avec l'autre.

Nous considérons des fenêtres rectangulaire espacées avec un intervalle de garde sur un intervalle de temps t entre a et b . Pour assurer l'orthogonalité dans le domaine temporelle, on doit rajouter un intervalle de garde afin d'éviter l'interférence entre les portes comme le montre la figure II.2.

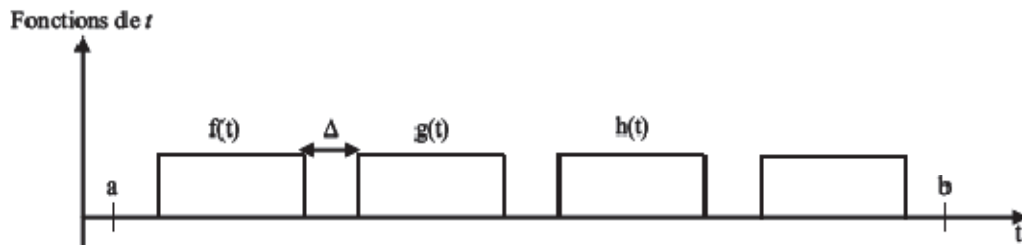


Fig. II.2 : Base orthogonale temporelle

Il est évident que $\int_a^b f(t)g(t)dt = 0$ et que $\int_a^b g(t)h(t)dt = 0$ ect...

La transformée de Fourier de la fonction porte d'amplitude A et de largeur T_u est un sinus cardinal comme en équation (II.2) :

$$TF[A \Pi_{T_u}(t)] = AT_u \text{sinc}(fT_u) \quad (\text{II.2})$$

Il est donc possible d'associer à une base orthogonale temporelle de fonction porte une base

orthogonale fréquentielle de sinus cardinaux par transformation de Fourier de chaque porte.

La figure II.3 représente un exemple de base orthogonale en fréquence dérivée de la base orthogonale en temps. L'espacement entre les N sinus cardinaux (sous porteuses) est défini par $\Delta f = \frac{1}{T_u}$

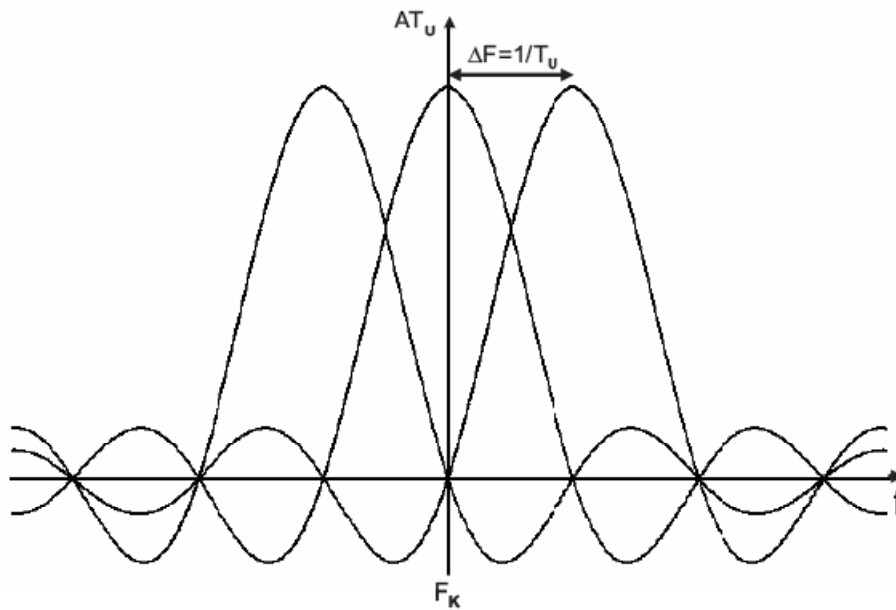


Fig. II.3 : Base orthogonale fréquentielle

II.3.2. Principe de la modulation OFDM :

Le principe de multiplexage en fréquence est de grouper des données numériques C_k par paquets de N , qu'on appellera symbole OFDM et de moduler par chaque donnée une porteuse différente en même temps ; Noté que les C_k sont des nombres complexes définis à partir des éléments binaires par une constellation souvent de modulation MAQ à 4, 16, 64, à 2^q états, ces données sont des symboles q -aires formés par groupement de q bits.

Considérons une séquence de N données : $C_0, C_1, \dots, C_{N-1}$.

Appelons T_s la durée symbole.

Chaque donnée C_k module un signal à la fréquence f_k .

Le signal individuel s'écrit sous forme complexe : $C_k e^{2j\pi f_k t}$.

Le signal $s(t)$ total correspondant à toutes les données d'un symbole OFDM est la somme des signaux individuels :

$$s(t) = \sum_{k=0}^{N-1} C_k e^{2j\pi f_k t} \quad (\text{II.3})$$

le multiplexage est orthogonal si l'espace entre les fréquences est $\frac{1}{T_s}$.

$$\text{Alors } f_k = f_0 + \frac{k}{T_s} \quad k = 0, 1, \dots, N-1$$

Voici le schéma de principe de la modulation :

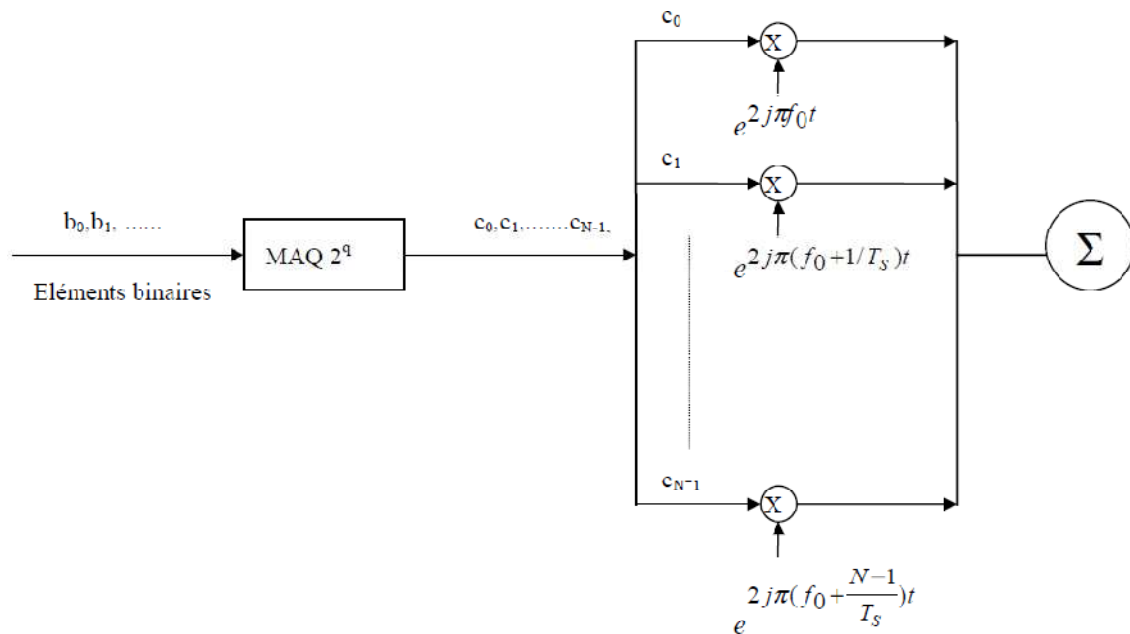


Fig. II.4 : Schéma de principe d'un modulateur OFDM

De l'équation II.3 on peut déduire l'expression réelle du signal OFDM :

$$\text{Si } C_k = a_k + jb_k$$

$$\begin{aligned}
 S(t) = \text{Re}(s(t)) &= \sum_{k=0}^{N-1} (a_k + jb_k) e^{2j\pi(f_0 + \frac{k}{T_s})t} \\
 &= \sum_{k=0}^{N-1} \left[a_k \cos\left(2\pi\left(f_0 + \frac{k}{T_s}\right)t\right) - b_k \sin\left(2\pi\left(f_0 + \frac{k}{T_s}\right)t\right) \right] \quad (\text{II.4})
 \end{aligned}$$

Chaque porteuse modulant une donnée numérique pendant une fenêtre rectangulaire de durée T_s , son spectre en fréquence est un sinus cardinal, fonction qui s'annule tous les multiples de $\frac{1}{T_s}$.

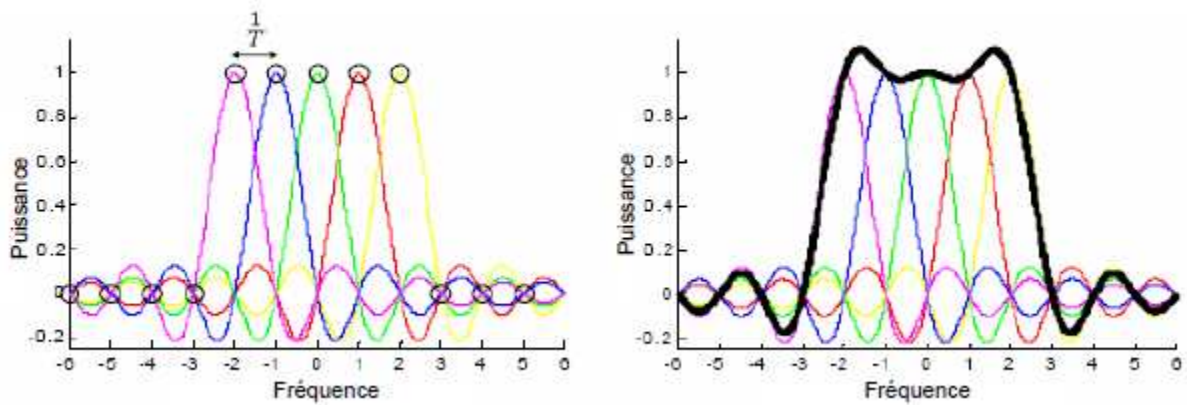


Fig.II.5 : allure de l'ensemble des spectres des porteuses d'un symbole OFDM

II.3.3. Principe de démodulation :

Le signal parvenant au récepteur s'écrit sur une durée symbole T_s :

$$y(t) = \sum_{k=0}^{N-1} C_k H_k(t) e^{j2\pi(f_0 + \frac{k}{T_s})t} \quad (\text{II.5})$$

$H_k(t)$ est la fonction de transfert du canal autour de la fréquence f_k et au temps t . Cette fonction varie lentement et on peut la supposer constante sur la période T_s ($T_s \ll \frac{1}{B_c}$)

La démodulation classique consisterait à démoduler le signal suivant les N sous-porteuses suivant le schéma classique dans la figure II.6:

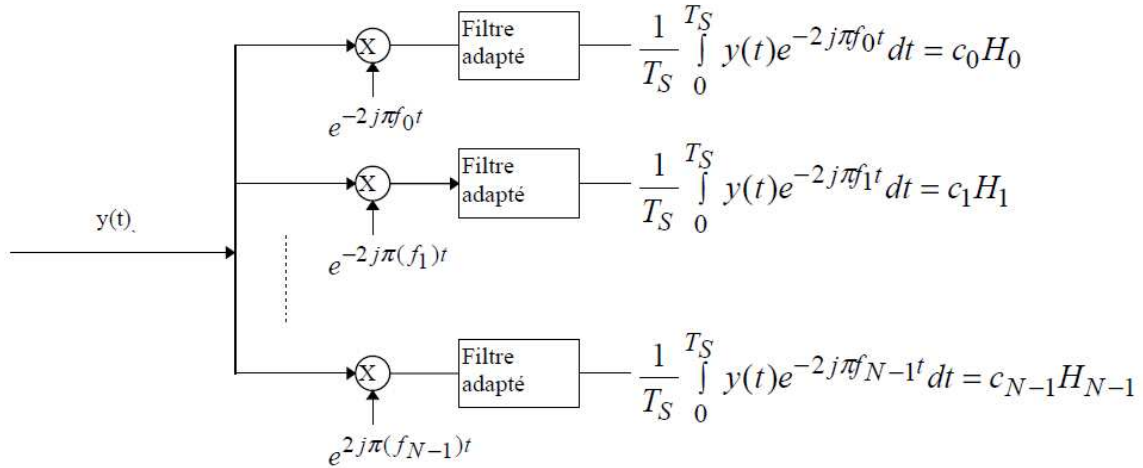


Fig. II.6 : Schéma de principe d'un démodulateur OFDM

La condition d'orthogonalité nous montre que :

$$\frac{1}{T_s} \int_0^{T_s} y(t) e^{-2j\pi f_i t} dt = \frac{1}{T_s} \sum_{k=0}^{N-1} C_k H_k \int_0^{T_s} e^{2j\pi(k-i)\frac{t}{T_s}} dt = C_i H_i \quad (\text{II.6})$$

Car :

$$\frac{1}{T_s} \int_0^{T_s} e^{2j\pi(k-i)\frac{t}{T_s}} dt = 0 \quad \text{si } k \neq i, 1 \text{ si } k = i$$

II.3.4. Réalisation numérique des opérations de modulation et de

démodulation :

La réalisation pratique de la modulation OFDM de façon directe (avec des oscillateurs et des mélangeurs) implique un circuit d'une complexité prohibitive.

Heureusement, il est possible de réaliser respectivement le modulateur et le démodulateur par des Transformées de Fourier Discrète Inverse et Directe (IDFT et DFT, via l'algorithme de

l'IFFT et FFT, si N est une puissance de 2). La complexité de ces opérations est de l'ordre de $N \log_2 N$ par symbole OFDM.

➤ Transformée de Fourier discrète :

La transformée de Fourier discrète (TFD) est un outil mathématique de traitement du signal numérique qui est l'équivalent discret de la transformée de Fourier continue qui est utilisée pour le traitement du signal analogique.

Sa définition mathématique pour un signal s de N échantillons est la suivante :

$$S(k) = \sum_{n=0}^{N-1} s(n) \cdot e^{-2i\pi k \frac{n}{N}} \quad \text{pour} \quad 0 \leq k < N \quad (\text{II. 7})$$

La transformée inverse TFDI (IDFT en anglais Inverse Discrete Fourier Transform) est donnée par :

$$s(n) = \frac{1}{N} \sum_{k=0}^{N-1} S(k) \cdot e^{-2i\pi k \frac{n}{N}} \quad (\text{II. 8})$$

La TFD ne calcule pas le spectre continu d'un signal continu. Elle permet seulement d'évaluer une représentation spectrale discrète d'un signal discret sur une fenêtre de temps finie.

La transformée de Fourier rapide (sigle anglais : *FFT* ou *Fast Fourier Transform*) est un algorithme de calcul de la transformée de Fourier discrète (TFD).

Sa complexité varie en $n \ln n$ avec le nombre de points n , alors que la complexité du calcul de base s'exprime en n^2 .

Cet algorithme est couramment utilisé en traitement numérique du signal pour transformer des données discrètes du domaine temporel dans le domaine fréquentiel, en particulier dans les analyseurs de spectre.

II .3.4.1. Implantation numérique du modulateur :

$s(t)$ est sous la forme :

$$s(t) = e^{2j\pi f_0 t} \sum_{k=0}^{N-1} C_k e^{2j\pi f_k t} \quad (\text{II.9})$$

En discrétisant ce signal et en le ramenant en bande de base pour l'étude numérique on obtient une sortie $s(n)$ sous la forme :

$$s(n) = \sum_{k=0}^{N-1} c_k e^{j2\pi \frac{nk}{N}} \quad (\text{II.10})$$

Les $s(n)$ sont donc obtenus par une Transformée de Fourier Inverse Discret des c_k . En choisissant le nombre de porteuses N tel que $N = 2^n$ (ou n est un nombre entier), le calcul de la Transformée de Fourier Inverse se simplifie et peut se réaliser avec une simple IFFT nous conduisant au schéma numérique présenté sur la figure II.7.

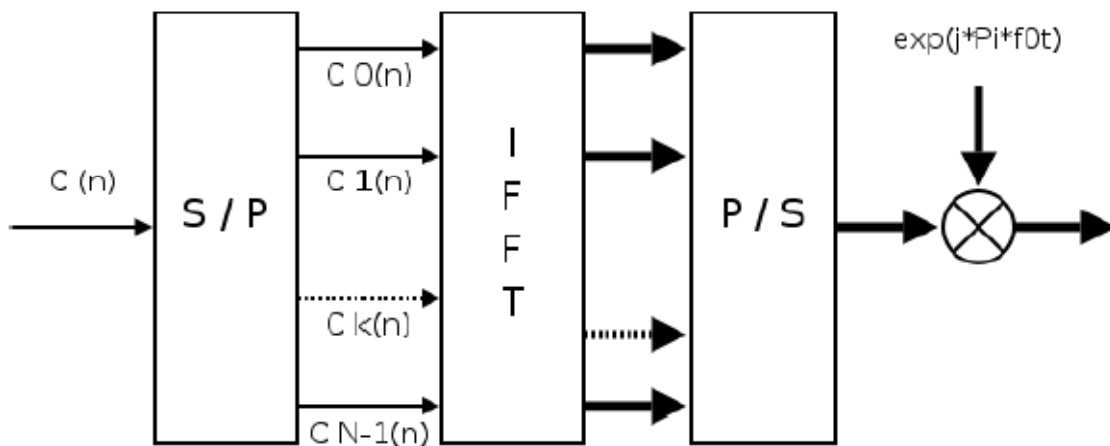


Fig.II.7 : Modulateur OFDM numérique

II.3.4.2. Implantation numérique de démodulateur :

L'analyse théorique définit le signal discrétisé reçu au niveau du démodulateur sous la forme :

$$Z(t_n) = Z_n \left(\frac{nT_s}{N} \right) = \sum_{k=0}^{N-1} C_k H_k e^{2j\pi \frac{K_n}{N}} \quad (\text{II.11})$$

Où Z_n est la transformée de Fourier Discrète Inverse de $C_k H_k$ et t_n : le temps d'échantillonnage consiste donc à effectuer une Transformée de Fourier Directe Discrète de $[Z_0, \dots, Z_{N-1}]$. Le nombre de porteuses ayant été choisi tel que $N = 2^n$, on peut réaliser ce calcul à l'aide d'une FFT. On obtient alors le schéma de principe dans la figure. II.8.

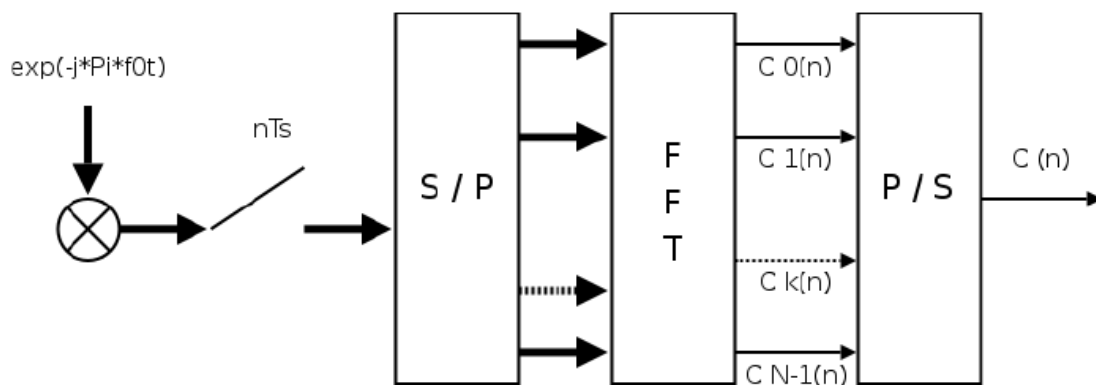


Fig. II. 8 : Démodulateur OFDM numérique

II.3.5. Interférence entre symboles OFDM :

Les interférences entre symboles sont dues au comportement multi trajets du canal, le signal reçu provenant de la contribution du trajet direct et des trajets multiples introduisant des retards, ces derniers pouvant être du même ordre de grandeur que la durée d'un symbole.

Pour remédier à ce problème on ajoute entre deux trames OFDM, un préfixe cyclique ou intervalle de garde $\Delta = T_g$.

➤ Intervalle de garde :

Le signal OFDM émis à travers un canal multi trajets, subit des échos et donc un symbole émis lors d'une période iT_s parvient au récepteur sous forme de plusieurs symboles atténués et retardés, pouvant se superposer à un écho provenant du symbole émis à la période

$(i - 1)T_s$, il se produit alors des interférences donc on ajoute un intervalle de garde T_g entre symboles OFDM d'une durée supérieure à l'étalement des retards T_m ainsi les derniers échos du symbole OFDM ne sera plus perturbé par le précédent d'où élimination des interférences entre symboles (ISI) qui subsistent malgré l'orthogonalité des porteuses (figure II.9).

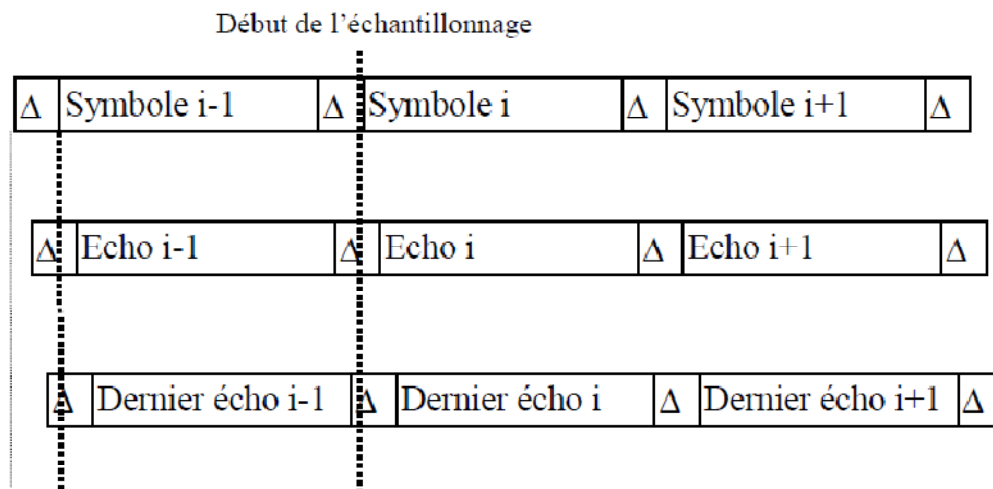


Fig. II.9 : Intervalle de garde

Le préfix est ajouté à l'émission après la IFFT, puis enlevé à la réception avant le module de la FFT.

En pratique $T_g = \frac{1}{4}T_s$ ce qui représente un bon compromis entre diminutions des erreurs et perte de débit utile.

Entre la durée de symbole, la durée utile et l'intervalle de garde s'instaurent donc la relation suivante :

$$T_u = T_g + T_s \quad (\text{II.11})$$

II.3.6. Interférences entre sous-canaux :

si le préfixe inséré au début du symbole OFDM est muet, des interférences inter porteuses, ou ICI (inter carrier interférence) vont se produire. Soit une transmission OFDM à N sous porteuses à travers un canal à deux trajets, dont le retard du trajet indirect δ , inférieur à la longueur du préfixe Δ .

Observons sur la figure II.10 les chronogrammes de deux voix particulières correspondants aux sous porteuses de fréquences respectives f_i et f_{i+1} .

En réception, après suppression du préfixe cyclique on réalise la FFT sur la durée T_s du symbole OFDM, la transformée de Fourier d'une sinusoïde de fréquence f_i , convoluée par une fonction porte de largeur T_s correspondra au zéro d'un autre et inversement.

Par contre, pour les signaux ayant subi une ou plusieurs réflexions donc décalés dans le temps, la sinusoïde de fréquence n'est présente que sur une durée $T_r < T_s$ ceci entrera une modification de la fonction caractérisant le contenu spectral de puissance de signal, dont les passages par zéros se produiront donc pour les valeurs différentes de celles associées au trajet direct. Lors de l'échantillonnage : il n'y aura plus d'orthogonalité entre les sous porteuses une autre.

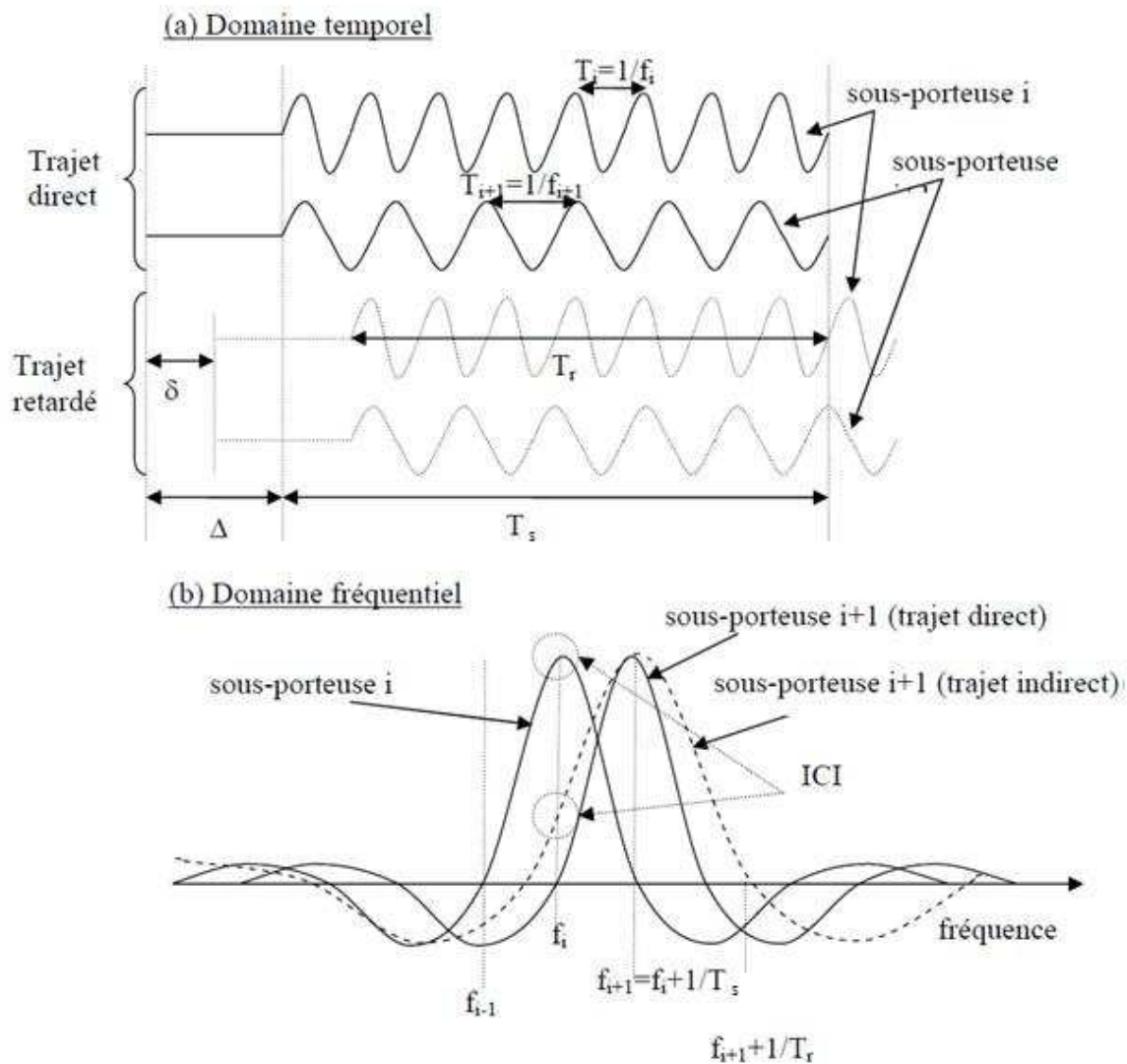


Fig.II .10 : Interférence entre porteuses dans le domaine temporel (a) et fréquentiel (b)

Afin d'éviter ces interférences, le préfixe ne doit pas être muet, mais être la copie des « L » derniers symboles numériques des symboles OFDM.

L'avantage de cette copie est que chaque signal, issu d'un trajet multiple possédera toujours un nombre entier de sinusoides sur la durée T_s . Dans le domaine fréquentiel et grâce au préfixe cyclique. La sommation des signaux de la sous porteuse f_i issus des divers trajets ne détruit donc pas l'orthogonalité des sous porteuse, mais introduit seulement un déphasage. Le synoptique de la figure II11. Illustre les différents traitements qui ont été exposé

précédemment et présente donc les modules de la chaîne de transmission et de réception OFDM.

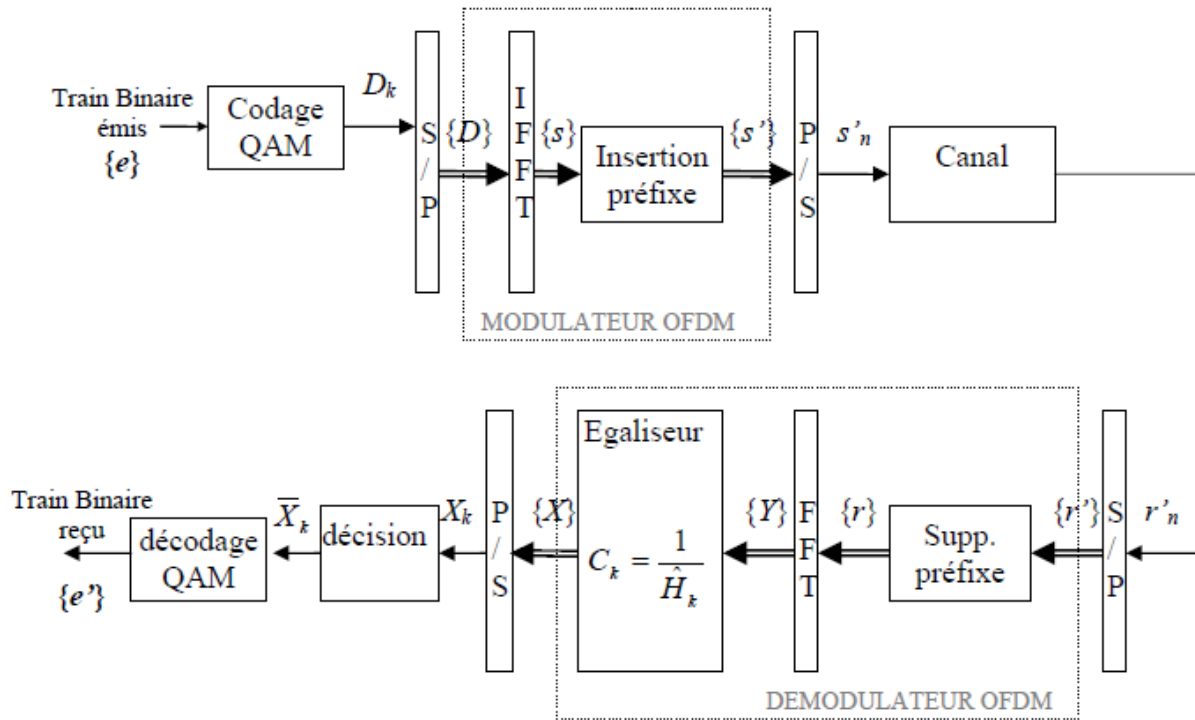


Fig. II.12 : Schéma synoptique d'un émetteur et d'un récepteur OFDM

II.4. Avantages et inconvénients :

II.4.1. Avantages :

Faible IES :

Le fait d'ajouter un intervalle de garde augmente la robustesse du signal OFDM au trajet multiples. Cela permet d'avoir en réception une IES acceptable.

Encombrement spectral optimal :

L'orthogonalité entre les N sous porteuses permet de faire chevaucher leurs respectives bandes fréquentielles et donc d'optimiser l'occupation spectrale du signal modulé

Canal invariant localement :

La bande passante de chaque sous-porteuses est petite devant la totalité de la bande passante du signal OFDM. Nous pouvons considérer que la réponse fréquentielle du canal de transmission est plate au niveau de chaque sous-porteuse.

Egalisation fréquentielle simple :

L'égalisation se fait par simple multiplication

II.4.2. Inconvénients :**Synchronisation émetteur /récepteur :**

Le décalage en fréquence (frequency offset) due à l'effet Doppler engendre l'interférence entre sous porteuses qui peut détruire l'orthogonalité des sous porteuse.

Dans le second cas, les erreurs de synchronisation induisant un déphasage sur les symboles reçus.

Fluctuations d'enveloppe :

Un signal de type OFDM présente des fortes fluctuations d'enveloppe et donc un PAPR élevé. Cela exige une grande linéarité au niveau de l'amplificateur de puissance qui présentera, alors un rendement divergent.

La caractéristique de transfert non-linéaire de l'amplificateur génère une distorsion dans la bande du signal OFDM qui aura un impact sur les N sous-porteuses qui interféreront entre elles avec une dégradation des performances en « BER » du système de transmission OFDM.

II.5. conclusion :

Dans ce chapitre nous avons décrit et caractérisé le signal OFDM. Le système multiporteuses permet de surmonter efficacement les dégradations introduites par le canal comme la sélectivité en fréquence et le bruit. Grâce aux progrès dans la fabrication des circuits numériques, la réalisation du système OFDM devient possible.

En revanche elle souffre du problème du facteur de crête élevé ou en encore « Peak-to-Average Power Ratio » PAPR qui fait l'objet de chapitre suivant.

III.1. Introduction :

La dynamique en amplitude d'un signal quelconque $s(t)$ représente un paramètre intéressant à prendre en compte surtout si ce signal est traité par des systèmes non- linéaires.

Les dispositifs non-linéaires qui présentent une caractéristique ce transfert avec saturation, font partie de la chaine de transmission, le problème se pose en émission et il prend de l'ampleur au niveau de l'amplificateur de puissance, car c'est l'élément qui consomme plus d'énergie sur une chaine de transmission, on comprend alors qu'il est nécessaire d'optimiser sa consommation surtout dans des terminaux mobiles où la consommation est un facteur important.

Dans ce chapitre notre attention se focalisera sur le facteur de crête de signal OFDM , la non-linéarité précisément dans l'amplificateur de puissance, et ensuite on verra les différentes méthodes proposées pour réduire ce problème.

III.2. Le facteur de crête :

Le signal OFDM temporel est la somme de N porteuses modulées. D'après le théorème de la limite centrale, si N tend vers l'infini, la distribution des valeurs prises par le signal OFDM tend vers une variable aléatoire normale.

Dans le cas des modulations monoporteuses, une telle sommation de porteuses n'intervient pas. Ainsi, si l'on compare les amplitudes instantanées des signaux mono et multiporteuse, on constate des différences importantes. Le signal multiporteuse a une dynamique plus grande et on peut remarquer la présence de « pics » d'amplitude importante, contrairement au signal monoporteuse. Sur la figure III.1, le signal du haut est un signal de type OFDM, et le signal de bas est un signal de type monoporteuse.

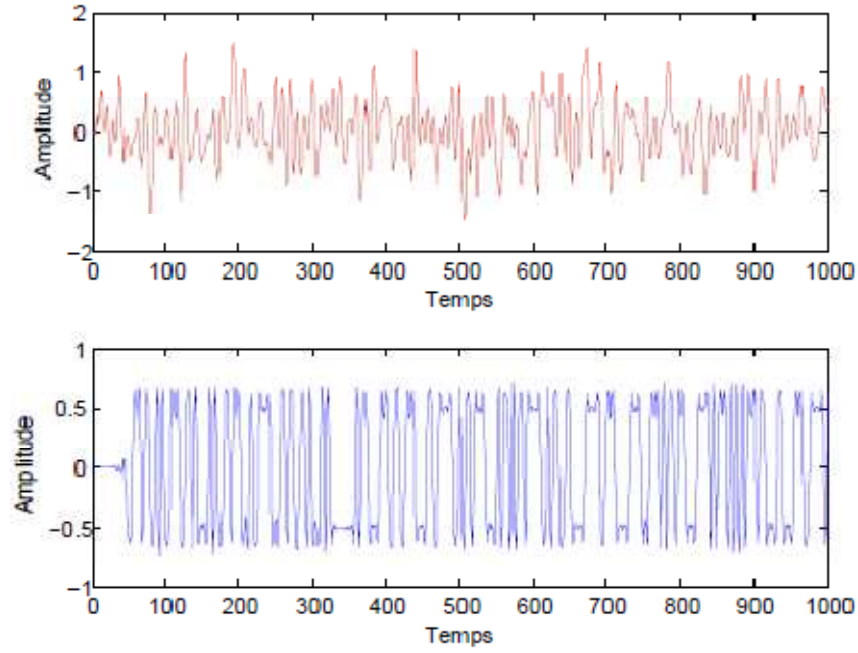


Fig.III.1 : signaux multiporteuses (haut) et monoporteuse (bas)

On peut quantifier cette caractéristique avec une grandeur appelée PAPR (Peak to Average Power Ratio) qui est le rapport entre la puissance maximale et la puissance moyenne d'un signal temporel $s(t)$ dans un intervalle T .

$$PAPR = \frac{\max_{t \in [0, T]} |s(t)|^2}{\frac{1}{T} \int_0^T |s(t)|^2 dt} \quad (\text{III. 1})$$

Ainsi, on définit le facteur de CR (Crête Factor) par la formule :

$$CF = \sqrt{PAPR} \quad (\text{III. 2})$$

Dans le cas de l'OFDM, le PAPR est donnée par

$$PAPR(\tilde{x}(t)) = \frac{\max_{t \in [0, T]} |\tilde{x}(t)|^2}{\frac{1}{T} \int_0^T |\tilde{x}(t)|^2 dt} \quad (\text{III.3})$$

Où $\tilde{x}(t)$ représente le signal en bande de base d'un symbole OFDM à N sous-porteuse exprimé par :

$$\tilde{x}(t) = \tilde{x}_I(t) + j\tilde{x}_Q(t) = \sum_{k=0}^{N-1} C_k e^{2j\pi \frac{k}{T} t}$$

C_k est le symbole complexe de la k^{ime} sous-porteuse et T et le temps d'un symbole OFDM.

En outre un signal numérique sur-échantillonné représente au mieux le signal analogique. Le sur-échantillonnage peut donc démasquer d'éventuels pics d'amplitude autrement perdus [figure (III.2)].

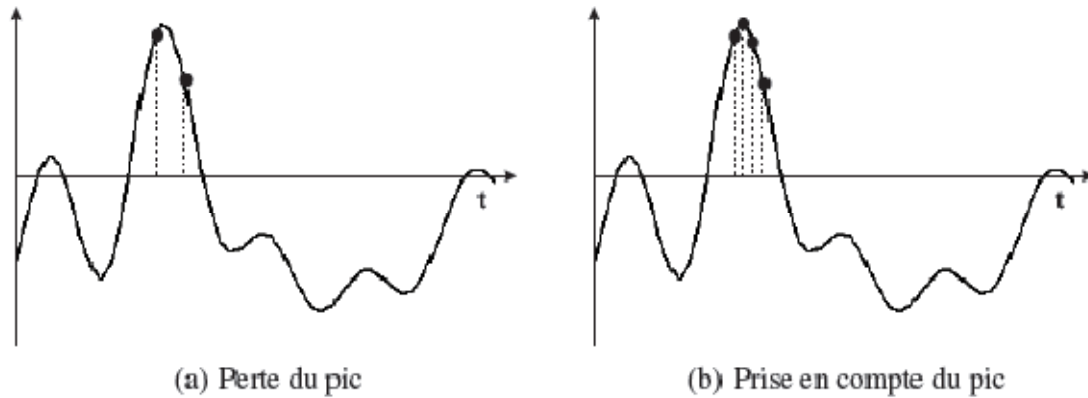


Fig.III.2 : Effet de sur-échantillonnage

Lorsque le signal OFDM échantillonné, il est donné par :

$$\tilde{x}(n/L) = \tilde{x}_I(n/L) + j\tilde{x}_Q(n/L) = \sum_{k=0}^{NL-1} C_k e^{2j\pi \frac{k}{NL} n}$$

Où $0 \leq n \leq NL - 1$, où L est le facteur de sur-échantillonnage.

Dans ce cas le PAPR est donné par la relation :

$$PAPR\{\tilde{x}\} = \frac{\max_{n \in [0, NL-1]} |\tilde{x}(n)|^2}{\frac{1}{NL} \sum_{n=0}^{NL-1} |\tilde{x}(n)|^2} \quad (\text{III. 4})$$

Le signal OFDM est une variable aléatoire ayant de grandes variations d'amplitude, le seul calcul de la quantité définie par le PAPR n'est pas suffisant. Une étude statistique s'impose et

cela passe par la détermination de la fonction de probabilité que le PAPR dépasse un seuil donné γ_0 . Cette fonction est connue sous le nom de CCDF (Complementary Cumulative Distribution Function), elle est donnée par l'équation (III.5).

$$CCDF(PAPR) = Prob(PAPR > \gamma_0) \quad (\text{III.5})$$

De nombreux travaux ont permis de déterminer soit de façon empirique ou soit de façon analytique la fonction de répartition du PAPR. En supposant les échantillons mutuellement indépendants et décorellés. Van Nee montre dans [15] que la CCDF peut être donnée par la relation :

$$CCDF_{théorique}(PAPR) = Prob(PAPR > \gamma_0) = 1 - (1 - e^{\gamma_0})^N \quad (\text{III.6})$$

Remarque :

Au cours des dernières années de nombreux travaux ont porté sur l'étude du PAPR surtout dans un contexte de signaux OFDM. Ainsi donc, la définition du PAPR n'était toujours pas la même d'un auteur à l'autre.

Dans la thèse le PAPR sera défini pour des signaux en bande de base parce que la plupart des techniques de réduction du PAPR que nous traitons interviennent en bande de base.

Le signal OFDM temporel ayant un facteur de crête fluctuant celui-ci peut subir des distorsions dus aux performances des dispositifs non linéaires, ce problème se pose surtout en émission et il prend de l'ampleur au niveau de l'amplificateur de puissance.

III.3. Caractéristiques d'un dispositif non linéaire :

Un système quelconque est caractérisé par sa fonction de transfert qui définie dans le domaine fréquentiel, comme le rapport entre la grandeur de sortie $y(jw)$ et celle d'entrée $X(jw)$:

$$H(jw) = \frac{Y(jw)}{X(jw)} = |H(jw)| \exp [j\phi(w)] \quad (\text{III.7})$$

Cette fonction de transfert prend en compte les distorsions linéaires d'amplitudes et de phase en fonction de la fréquence. D'autre distorsions liées à la présence d'éléments non-linéaires peuvent apparait dans le système.

Dans ce cas une simple fonction de transfert $H(jw)$ ne suffit plus pour décrire le comportement de système. Il est alors nécessaire d'exprimer le signal de sortie $y(t)$ comme une fonction $f[.]$ du signal d'entrée $x(t)$. Si la sortie à l'instant t ne dépend que de l'entrée au même instant, le système non linéaire est alors défini sans mémoire et on peut écrire :

$$y(t) = f[x(t)] \quad (\text{III.8})$$

Si les grandeurs d'entrée et de sortie sont des tensions, alors l'équation (III.8) peut se réécrire comme suit :

$$V_s(t) = f[V_e(t)] \quad (\text{III.9})$$

La fonction $f[.]$ représentée dans la figure III.3 est exemple caractéristique non linéaire du système.

Focalisons alors notre attention dans le modèle polynomial représentant la non linéarité du système. Dans ce cas la fonction $f[.]$ peut donc s'écrire sous la forme d'un polynôme d'ordre n tel que :

$$V_s(t) = a_1 V_e(t) + a_2 V_e^2(t) + a_3 V_e^3(t) + \dots + a_n V_e^n(t) \quad (\text{III.10})$$

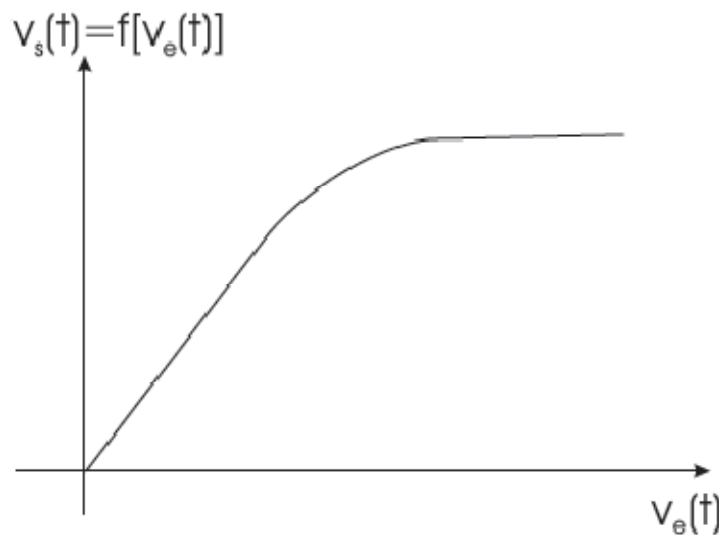


Fig.III.3 : Caractéristique non-linéaire de transfert, $f[.]$

De manière générale les coefficients a_i peuvent s'écrire comme $(\alpha_i + j\beta_i)$. On dit alors que la non-linéarité de la fonction $f[.]$ introduit juste une distortion d'amplitude si les coefficients a_i sont réels, tandis que la distortion est d'amplitude et de phase s'ils sont complexes. En outre, les principales grandeurs caractérisant un dispositif non linéaire sont les suivantes :

III.3.1. Les harmoniques :

Supposons que la non-linéarité est toujours représentable par une fonction $f[.]$ de type polynomiale comme dans l'équation (III.10).

Considérons que les non-linéarités d'ordre supérieure à 3 sont négligeables et que le signal d'entrée $V_e(t)$ est à un seul ton, c'est-à-dire une sinusoïde pure [équation (III.11)].

$$V_e(t) = A \cos(w_1 t) \quad (\text{III.11})$$

Alors nous pouvons en déduire que le signal à la sortie du dispositif non-linéaire s'écrit en remplaçant l'équation (III.11) dans (III.10), sous la forme suivante :

$$\begin{aligned} V_s(t) &= a_1 A \cos(w_1 t) + a_2 A^2 \cos^2(w_1 t) + a_3 A^3 \cos^3(w_1 t) \\ &= a_1 A \cos(w_1 t) + a_2 A^2 \left[\frac{1}{2} + \frac{1}{2} \cos(2w_1 t) \right] \\ &\quad + a_3 A^3 \left[\frac{3}{4} \cos(w_1 t) + \frac{1}{4} \cos(3w_1 t) \right] \end{aligned}$$

et encore :

$$\begin{aligned} V_s(t) &= \frac{1}{2} a_2 A^2 && \text{composante continue} \\ &+ \left[a_1 A + \frac{3}{4} a_3 A^3 \right] \cos(w_1 t) && \text{fondamental} \\ &+ \frac{1}{2} a_2 A^2 \cos(2w_1 t) && 1^{\text{ère}} \text{ harmonique} \\ &+ \frac{1}{4} a_3 A^3 \cos(3w_1 t) && 2^{\text{ème}} \text{ harmonique} \end{aligned} \quad (\text{III.11})$$

Dans l'équation (III.11) nous retrouvons l'expression du fondamentale ainsi que l'expression de la composante continue. Les termes dont la fréquence est un multiple de la fréquence du fondamental représente les harmoniques du signal, générées par la non linéarité du dispositif.

III.3.2. Les produits d'intermodulation :

L'effet de l'intermodulation apparait lorsque plusieurs signaux à des fréquences différentes traversent un composant non linéaire.

Lorsque le signal d'entrée d'un système non-linéaire est un signal à deux tons, c. à d. la somme de deux sinusoides, de nouveaux termes apparaissent en sortie. Ce ne sont ni le fondamental, ni les harmoniques du signal. Considérons maintenant un signal d'entrée à deux tons :

$$V_e(t) = A_1 \cos(w_1 t) + A_2 \cos(w_2 t) \quad (\text{III.12})$$

Si l'on insère le signal $V_e(t)$ dans l'équation (III.10), on obtient en sortie du système non-linéaire le signal suivant :

$$\begin{aligned}
 V_s(t) = & a_1 A_1 \cos(w_1 t) + a_2 A_2 \cos(w_2 t) && \text{fondamental} \\
 & + \frac{a_1 A_1^2}{2} + \frac{a_2 A_2^2}{2} && \text{composante continue} \\
 & + \frac{a_2 A_1^2}{2} \cos(2w_1 t) + \frac{a_2 A_2^2}{2} \cos(2w_2 t) && 2^{\text{ème}} \text{ harmonique} \\
 & + a_2 A_1 A_2 [\cos(w_1 + w_2)t + \cos(w_1 - w_2)t] && IM\ 2 \\
 & + \left(\frac{3a_3 A_1 A_2^2}{2} + \frac{3a_3 A_1^3}{4} \right) \cos(w_1 t) + \dots && (\text{III.13}) \\
 & \dots + \left(\frac{3a_3 A_1^2 A_2}{2} + \frac{3a_3 A_2^3}{4} \right) \cos(w_2 t) && \text{fondamental} \\
 & + \frac{a_3 A_1^3}{4} \cos(3w_1 t) + \frac{a_3 A_2^3}{4} \cos(3w_2 t) && 3^{\text{ème}} \text{ harmonique} \\
 & + \frac{3a_3 A_1 A_2^2}{4} [\cos(2w_1 + w_2)t + \cos(2w_1 - w_2)t] + \dots
 \end{aligned}$$

$$\dots + \frac{3a_3 A_1^2 A_2}{4} [\cos(2w_2 + w_1)t + \cos(2w_2 - w_1)t] \quad IM\ 3$$

+... termes d'ordre supérieur à 3

Dans l'équation (III.13) nous voyons apparaître les termes du fondamental ainsi que d'autres termes appelés harmoniques et produits d'intermodulation (*IM*) dont la fréquence est multiple ou combinaison linéaire des fréquences fondamentales. Une non-linéarité d'ordre 2 provoque le produit d'intermodulation d'ordre 2 (*IM2*) tandis qu'une non-linéarité d'ordre 3 génère un produit d'intermodulation d'ordre 3 (*IM3*).

Les amplitudes des produits d'intermodulation décroissent avec l'ordre de l'intermodulation et ceux qui se situent en fréquence à côté du fondamental seront les plus gênants.

La figure (III.4) représente les termes d'intermodulation dans le domaine fréquentiel

lorsque l'amplitude des deux composantes du signal d'entrée est la même ($A_1 = A_2 = A$).

Donc, nous pouvons en conclure que les produits d'intermodulation d'ordre impair sont les plus gênants, en particulier l'ordre 3 (*IM3*), car ils sont les plus proches des fréquences fondamentales (f_1 et f_2). En revanche, les produits d'intermodulation d'ordre pair et les produits d'intermodulation d'ordre impair, somme des fréquences harmoniques, sont rejetés loin des signaux aux fréquences fondamentales. Ces derniers peuvent donc être éliminés par filtrage.

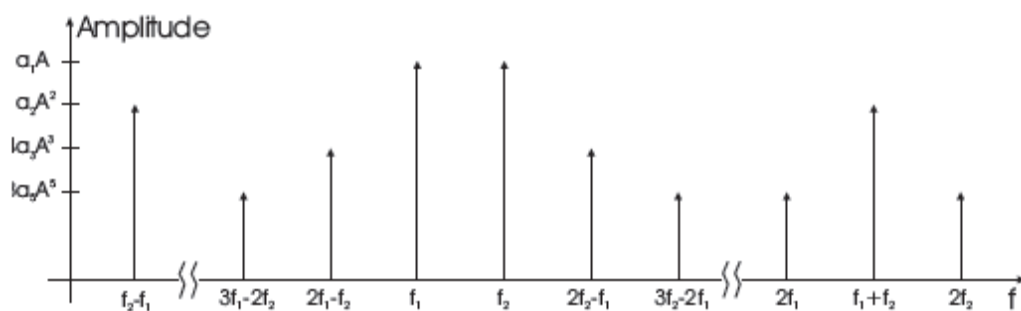


Fig.III.4 : Produit d'intermodulation

III.3.3. Le point de compression à 1db :

D'après l'équation (III.13) l'amplitude du fondamental du signal de sortie est égale à :

$$A_{fond} = a_1 A + \frac{3}{4} a_3 A^3 \quad (\text{III.14})$$

Cette grandeur est inférieure à l'amplitude du signal amplifié linéairement ($a_1 A$) si $a_3 < 0$ (compression de gain), et elle est supérieure à $a_1 A$ si $a_3 > 0$ (expansion de gain).

La plupart des dispositifs travaillent en compression, c'est-à-dire avec $a_3 < 0$.

On définit alors le gain du dispositif à la fréquence fondamentale qui est donné par l'équation (III.15).

$$G_{fon} = 20 \log \left(\frac{a_1 A + \frac{3}{4} a_3 A^3}{A} \right) = 20 \log \left(a_1 + \frac{3}{4} a_3 A^2 \right) \quad (\text{III.15})$$

Le gain linéaire G_{lin} vaut :

$$G_{lin} = 20 \log \left(\frac{a_1 A}{A} \right) = 20 \log(a_1) \quad (\text{III.16})$$

D'où nous en déduisant le gain à 1 dB de compression définit comme la compression d'1 dB sur le fondamental par rapport au gain linéaire [équation (III.17)].

$$G_{1dB} = G_{lin} - 1 \text{ dB} \quad (\text{III.18})$$

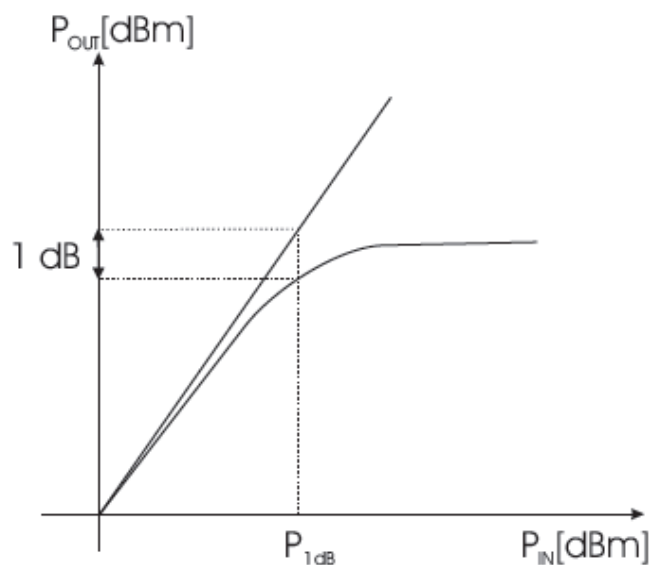


Fig.III.5 : point de compression à 1 dB

Au point de compression d'1 dB, l'équation (III.18) peut alors se réécrire comme suit :

$$20 \log \left(a_1 + \frac{3}{4} a_3 A_{1dB}^2 \right) = 20 \log(a_1) - 1 \text{ dB} \quad (\text{III.19})$$

D'où l'on déduit la valeur du point de compression en tension :

$$A_{1dB} = \sqrt{0.145 \left| \frac{a_1}{a_3} \right|} \quad (\text{III.20})$$

Cette grandeur est une mesure du niveau maximal de signal d'entrée que le dispositif peut traiter. Au delà de cette valeur, le signal est de plus en plus compressé jusqu'à arriver à la saturation.

III.4. Sources de non linéarités (Amplificateur de puissance) :

Pour assurer correctement l'acheminement des informations, les émetteurs ont besoins d'amplificateurs de puissance pour fournir une certaine puissance aux signaux (radiofréquence) et éviter qu'ils ne s'affaiblissent fortement lors de leur propagation dans l'espace libre.

III.4.1. Caractérisation du la non-linéarité d'amplitude et de la phase

d'amplificateur de puissance

La relation entrée-sortie, appelée aussi la caractéristique de transfert, a une allure typique pour tous amplificateurs de puissance. Les partie (a) et (b) de la figure III.6 relatent la variation de la puissance de sortie en fonction de la puissance d'entrée, appelée aussi la caractéristique Amplitude /Amplitude ou encore la compression AM/AM. La partie (c) de cette figure présente le gain en puissance de l'amplificateur en fonction de la puissance d'entrée. Les caractéristiques se divisent en trois zones :

1) Zone linéaire : dans cette zone, l'amplificateur à un comportement proche d'un système linéaire. La puissance de sortie est proportionnelle à la puissance d'entrée selon un rapport appelé gain de l'amplificateur. Les puissances d'entrées sont faibles.

2) Zone de compression : dans cette zone, la puissance de sortie n'est déjà plus proportionnelle à la puissance d'entrée. La courbe commence à s'incurver (par rapport à la droite linéaire). Les distorsions du signal apparaissent et sont de plus en plus importantes. Le gain de l'amplificateur diminue pour de fortes puissances d'entrée. On parle de zone de compression du gain. Un point important est situé dans cette zone. Il s'agit du point où l'écart entre la courbe de gain et le gain linéaire vaut 1dB ($P_{e,1dB}$), c'est un point caractéristique de l'amplificateur de puissance.

3) Zone de saturation : à partir d'une certaine puissance d'entrée, la puissance de sortie devient quasiment constante et la courbe de gain décroît linéairement. La saturation se manifeste par un écrêtage du signal de sortie. La puissance de saturation en sortie est elle aussi une caractéristique de l'amplificateur désignée par $P_{s,sat}$.

La courbe exprimant le déphasage entre la sortie et l'entrée est appelée caractéristique Amplitude/Phase, ou aussi conversion AM/PM. La conversion varie suivant la technique de conception et les conditions de fonctionnement de l'amplificateur.

La partie (d) de la figure III.7 montre que la courbe AM/PM reste constante dans la zone linéaire. Des variations commencent à apparaître autour du point à 1db de compression de gain P_{1db} et peuvent être très importantes à proximité de la saturation.

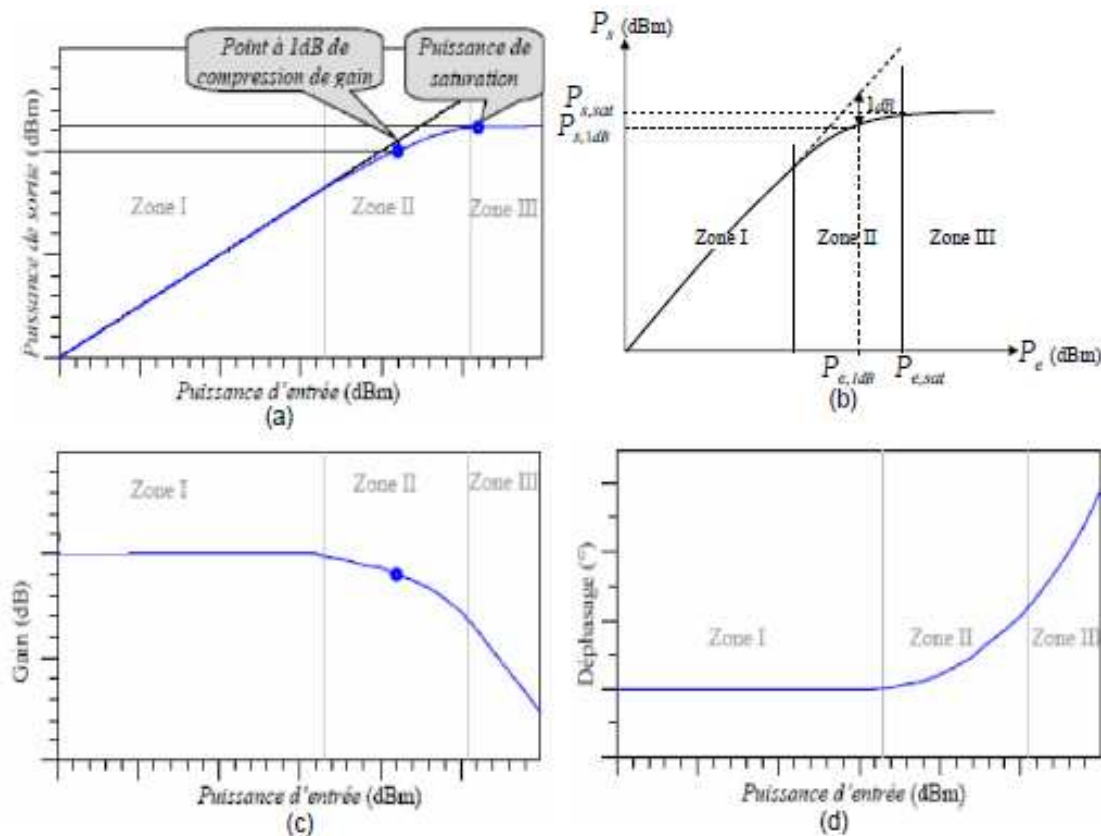


Fig.III.7 : caractéristique de la relation entrée-sortie d'un amplificateur
de puissance

Pour éviter au du moins diminuer les effets néfastes dus à la non-linéarité des amplificateurs, on est souvent amené à surdimensionner l'amplificateur ou en d'autres termes à prendre un certain recul et d'être dans la zone linéaire ou proche de cette zone.

On définit alors les grandeurs qui sont le recul d'entrée IBO (Input Back Off) et le recul de sortie OBO (Output Back Off).

La grandeur IBO est le rapport entre la puissance de saturation ramenée à l'entrée sur la puissance moyenne d'entrée P_e du signal. Elle est exprimée en db.

$$IBO = \frac{P_{e,sat}}{P_e} \quad \text{ou} \quad IBO(dB) = P_{e,sat}(dB_m) - P_e(dB_m) \quad (\text{III. 21})$$

De la même façon on définit le paramètre OBO qui est le rapport entre la puissance de saturation et la puissance moyenne P_s de sortie du signal.

$$OBO = \frac{P_{s,sat}}{P_s} \quad \text{ou} \quad OBO(db) = P_{s,sat}(db) - P_s(db) \quad (\text{III.22})$$

Plus le recul d'entrée (ou de sortie) est élevé, plus l'amplificateur est surdimensionné, et par conséquent moins il y a des distorsions. Cette solution n'est cependant pas idéale dans la mesure où le rendement dans ce cas est faible. D'où l'intérêt de chercher des solutions peuvent concilier aux mieux la linéarité et le rendement.

- **ACPR :**

Les déformations des lobes adjacents par les non-linéarités du dispositif, sont caractérisées par « l'Adjacent Channel Power Ratio » « ACPR ». L'ACPR est défini comme le rapport entre la puissance dans le canal adjacent (P_{BA}) et celle dans le canal principal (P_{BV}). On peut alors parler de « ACPR » supérieur, calculé en considérant le canal adjacent supérieur, et de « ACPR » inférieur « low » calculé en considérant le canal adjacent inférieur il est calcul par l'équation (III.23):

$$ACPR_{(up,low)}(db) = 10 \log \frac{P_{BA}}{P_{BV}} \quad (\text{III.23})$$

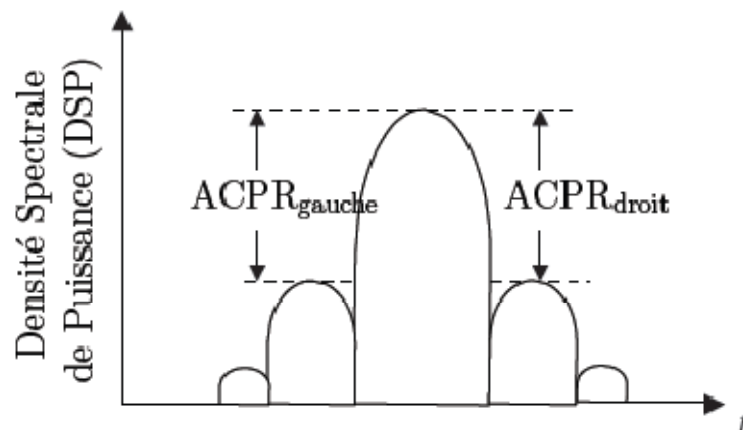


Fig.III.8 : Calcul d'ACPR

III.4.2. Les types d'amplificateurs de puissance :

D'une manière générale, deux types d'amplificateurs de puissance sont utilisés dans les systèmes de communications. Tout d'abord nous exprimons l'entrée de l'amplificateur sous forme de l'équation III.24 et la sortie de l'amplificateur est donnée par l'équation III.25

$$x(t) = |x(t)|e^{j \arg(x(t))} = \rho(t)e^{j \phi(t)} \quad (\text{III.24})$$

$$\tilde{x}(t) = A(\rho(t))e^{j(\phi(t)+\phi(\rho(t)))} \quad (\text{III.25})$$

Avec $A(\rho(t))$ et $\phi(\rho(t))$ sont respectivement les conversions *AM/AM* et *AM/PM* de l'amplificateur non linéaire.

➤ l'amplificateur à tubes à ondes progressives « *TWTA* » (Travelling Wave Tube Amplificateur) plus utilisé par exemple dans les transmissions satellites pour transmettre de fortes puissances.

Donc pour ce modèle les fonctions $A(\rho(t))$ et $\phi(\rho(t))$ sont données par les expressions suivantes :

$$A(\rho(t)) = \nu \rho(t)/(1 + \beta_\alpha \rho^2(t)) \quad (\text{III.26})$$

$$\phi(\rho(t)) = \frac{\pi}{3} * \rho^2(t)/(A_{sat}^2 + \rho^2(t)) \quad (\text{III.27})$$

$$A_0 = \max(A(\rho(t))) = \nu A_{sat}/2 \quad (\text{III.28})$$

Avec : $\begin{cases} \nu \text{ est le gain d'amplification} \\ A_{sat} = \frac{1}{\sqrt{\beta_\alpha}} \text{ est l'entrée de saturation de l'AP} \\ A_0 \text{ est l'amplitude max de la sortie} \end{cases}$

➤ l'amplificateur à l'état solide « *SSPA* » (Solide State Power Amplificateur) utilisé dans les transmissions radio terrestres pour transmettre de faibles puissances. C'est le cas par exemple du téléphone mobile. Donc pour ce modèle les fonctions $A(\rho(t))$ et $\phi(\rho(t))$ sont données par les expressions suivantes

$$A(\rho(t)) = \nu \rho(t) / [(1 + (\nu \rho(t)/A_0)^{2p})]^{1/2p} \quad (\text{III.29})$$

$$\phi(\rho(t)) \cong 0 \quad (\text{III.30})$$

Avec p permet d'ajuster la transition entre la zone linéaire et la zone de saturation en tendant p vers l'infini, on obtient un limiteur.

III.5. Principe des méthodes de réduction du PAPR :

Comme déjà mentionné auparavant, la grande dynamique en amplitude et en puissance peut saturer le signal OFDM traversant un amplificateur de puissance. Différentes techniques ont été découvertes par les chercheurs afin de réduire ces fluctuations, c'est-à-dire réduire le PAPR. Nous nous sommes principalement intéressés aux méthodes agissant sur le signal OFDM émis. Parmi ces techniques nous retrouvons celle comme l'écrêtage plus filtrage et « Tone Reservation » qui garantissent une compatibilité descendante et celle comme le « selective –mapping » SLM et « Partial Transmit Sequence » PTS qui ne l'a garantissent pas.

III.5.1. Écrêtage plus filtrage :

Cette méthode a été proposée dès la mise en œuvre de l'OFDM terrestre DVB-T (Digital Vidéo Broadcasting-Terrestrial), elle consiste à éliminer les forts pics du signal OFDM en faisant un écrêtage au signal par seuil E prédéterminé. On peut dire aisément que la variation de l'amplitude du signal écrêté sera moins important que celle de signal sans écrêtage ; par conséquent, la sensibilité du signal aux non linéarité d'amplification diminue mais cette limite du signal dégrade ses performances s'il reste compatible, on a donc l'augmentation du taux d'erreur binaires. Soient $x(n)$ et $y(n)$ les signaux OFDM avant et après l'écrêtage respectivement, on a l'expression suivante :

$$y(n) = \begin{cases} +E & \text{si } x(n) > +E \\ x(n) & \text{si } |x(n)| \leq E \\ -E & \text{si } x(n) < -E \end{cases} \quad (\text{III.31})$$

On définit le facteur CR (Clipping Ratio), le rapport entre le niveau d'écrêtage E [volts] et la moyenne quadratique du symbole OFDM, qui n'est pas que la valeur efficace du signal (mesuré en volts), on donne ainsi l'expression :

$$CR = \frac{E}{\sqrt{\frac{1}{N} \sum_{i=0}^{N-1} c_i^2}} \quad (\text{III. 32})$$

Sachant que les c_i sont les symboles numériques d'un symbole OFDM.

L'implémentation de la fonction d'écrêtage dans l'émetteur OFDM est mise entre l'IFFT et le préfixe cyclique, ensuite on amplifie le signal résultant, comme le montre la figure III.8.

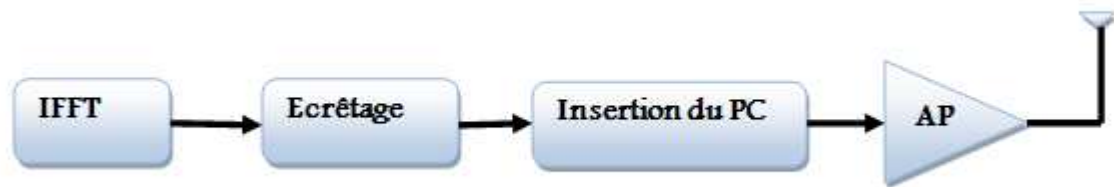


Fig.III.1 : Insertion de préfixe cyclique

L'écrêtage génère les trois types de problèmes du non linéarité qui sont :

- La remontée des lobes secondaires due aux produits d'intermodulations et les harmoniques hors la bande utile.
- Un bruit dans la bande utile du spectre OFDM, due aux produits d'intermodulations d'ordre impair.
- Déformations significatives du signal OFDM qui implique une dégradation en terme de « BER ».

Le premier problème est réglé par un filtre en fréquence juste après l'écrêtage (voir la figure IV.2)

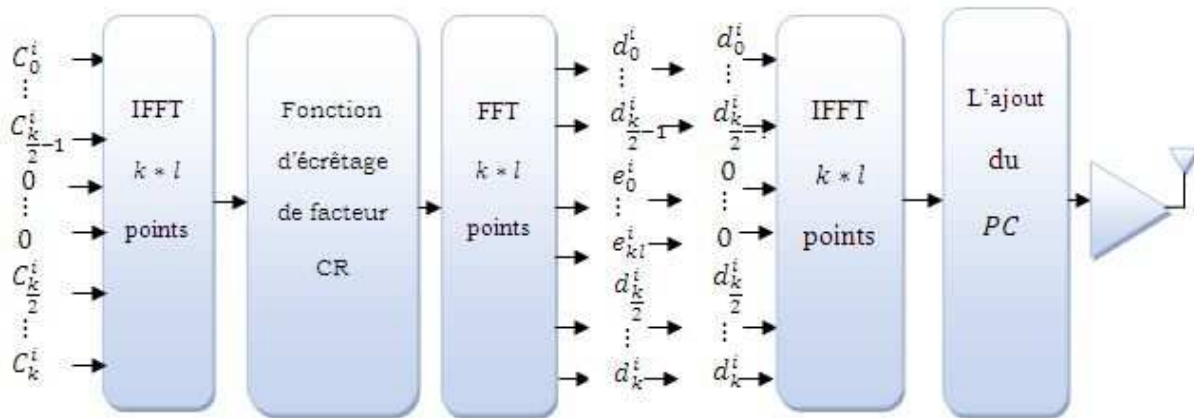


Fig. III .9 : Insertion du filtrage

On rappelle que C_i^m signifie le $m^{ème}$ symbole numérique du $i^{ème}$ symbole OFDM, les zéros ajoutés au milieu du paquet font le sur-échantillonnage de facteur l . La fonction d'écrêtage déforme le symbole OFDM, pour cela on fait une FFT pour visualiser l'effet non linéaire sur le spectre, afin de remettre les symboles e_n^i à zéros ; par conséquent on a éliminé les lobes secondaires. Par contre c'est très compliqué de supprimer le bruit d'intermodulations d'ordre impair.

La fonction d'écrêtage permet de réduire les fortes fluctuations d'amplitude et de puissance : par conséquent le PAPR de signal OFDM. Cependant l'écrêtage est lui aussi une opération non linéaire qui engendre des distorsions en fait plus le niveau d'écrêtage E est petit, est plus on a une dégradation importante en terme de « BER » ; puis de l'autre coté une diminution considérable du PAPR. L'avantage par contre par l'utilisation de l'écrêtage est que les distorsions en dehors de la bande utile (mesuré par l'ACPR) peuvent être facilement éliminées grâce à un simple filtrage, chose que l'on ne peut pas faire directement après un amplificateur de puissance. L'amplificateur peut jouer son rôle après d'une manière transparente, c'est-à-dire : sans affecter négativement le signal.

III.5.2 . Tone Reservation :

La technique « Tone Reservation » est une méthode de réduction du « PAPR » dont le principe repose sur l'ajout d'un signal au signal original afin de diminuer son « PAPR ».

L'idée de base de cette technique consiste à réserver des sous porteuses pour générer le signal servant à réduire le « PAPR » des signaux OFDM. Jose-Tellado-Mourelo[12] en est le précurseur.

Contrairement à certaines techniques dites « ajout de signal », la technique « Tone Reservation » n'a pas pour objectif de transformer le signal original en un signal à enveloppe constante mais plutôt à réduire au maximum son « PAPR » .

En fait, l'idée de base de cette technique est de réserver un certain nombre de sous –porteuses pour générer les signaux devant servir à réduire le « PAPR » de sorte à garantir en même temps la compatibilité descendante.

Pour ce faire, le signal ajouté et le signal utile sont à supports fréquentiels disjoints. Notons qu'initialement, la technique « Tone Reservation » fut proposée pour les signaux OFDM d'où le schéma illustratif de la figure III-10.

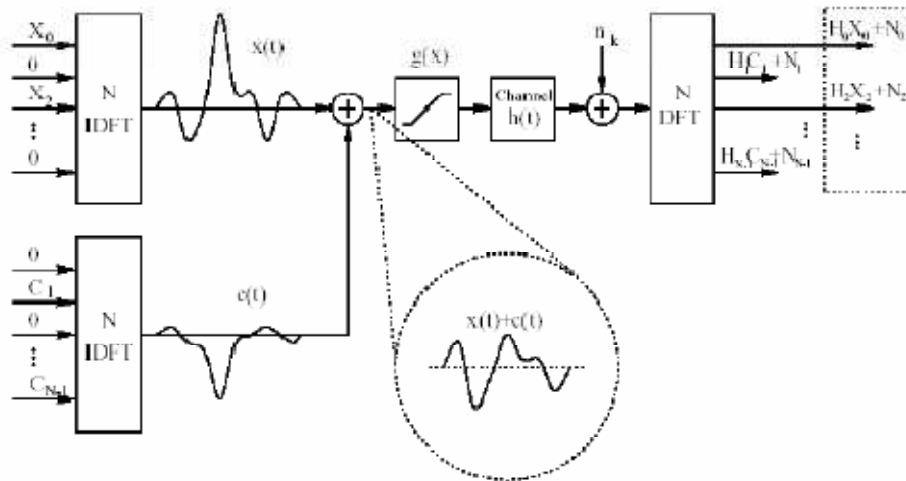


Fig.IV.10 : Schéma général pour le « Tone Reservation »

D'un point de vue analytique on a : $\mathbf{X}_k^m \mathbf{C}_k^m = \mathbf{0}$ où $\mathbf{X}_k^m$ désigne le $k^{\text{ème}}$ symbole numérique de $m^{\text{ème}}$ symbole OFDM et $\mathbf{C}_k^m$ le $k^{\text{ème}}$ symbole numérique de $m^{\text{ème}}$ signal « artificiel ».

Posons $P = \{0, 1, 2, \dots, N-1\}$ l'ensemble des positions de toutes les N sous-porteuses du signal multiporteuse, $R = \{i_0, i_1, \dots, i_{r-1}\}$ l'ensemble ordonné des positions des sous porteuses réservées pour le signal « artificiel » et R^c le complémentaire de R dans P , c'est-à-dire $P = R \cup R^c$. Par suite :

$$X_k^m + C_k^m = \begin{cases} C_k^m & \text{si } k \in R \\ X_k^m & \text{si } k \in R^C \end{cases} \quad (\text{III. 32})$$

Le signal multiportuese s'écrit alors :

$$\begin{aligned} \bar{x} &= x^m[n/L] + c^m[n/L] \\ &= \sum_{k \in R^C} X_k^m e^{\frac{2j\pi kn}{NL}} + \sum_{k \in R} C_k^m e^{\frac{2j\pi kn}{NL}} \\ &= \text{IFFT}(X_k^m + C_k^m) \end{aligned} \quad (\text{III. 33})$$

Où L est le facteur de sur_échantillonnage. Le « PAPR » ainsi obtenu est donnée par la relation :

$$\text{PAPR}[x^m + c^m] = \frac{\|x^m + c^m\|_\infty^2}{E(|x^m[N/L]|^2)} \quad (\text{III. 34})$$

Où $\|v\|_\infty$ représente la norme $-\infty$ du vecteur v . Puisque le dénominateur n'est pas fonction du signal ajoutée (si c'était le cas le « PAPR » définit ne serait pas une mesure fiable car elle aurait été plus faible que le « PAPR » de signal original). Le problème d'optimisation du PAPR se résume au calcul de la valeur $C^{m,opt}$ qui minimise le numérateur, c'est-à-dire :

$$\min_c \|x^m + c\|_\infty = \min_c \|x^m + Q_L C\|_\infty \quad (\text{III. 35})$$

Un seuil de « PAPR » est nécessaire pour le calcul de $C^{m,opt}$. Le choix du seuil se fait en général en fonction de la distribution de la puissance du signal. Dans la figure III.11, on peut voir par exemple que pour un seuil pour lequel la probabilité est inférieure à 10^{-5} et pour rapport $R/N = 5\%$ le « PAPR » peut être réduit de 15db à 9db (soit de 6db).

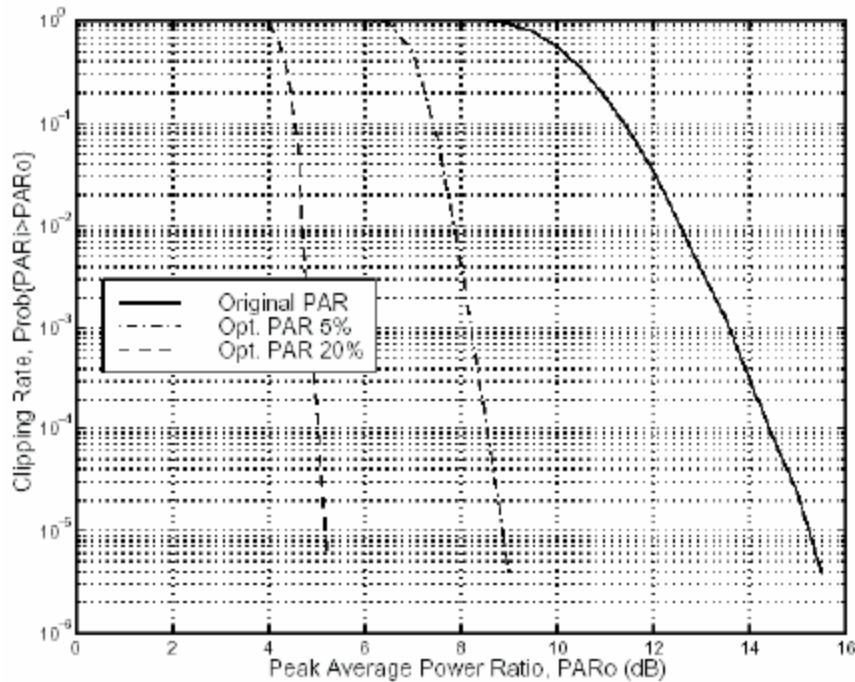


Fig.III.11 : Impact de la Méthode « Tone Reservation » sur le « PAPR » ($N=512$)

- **Avantage :**

L'avantage de la technique « Tone Reservation » est qu'elle permet d'obtenir une bonne réduction du « PAPR » tout en assurant une compatibilité descendante.

- **Inconvénients :**

Les difficultés de la technique « Tone Reservation » résident essentiellement dans le choix de l'algorithme d'optimisation et dans le choix de seuil. Un autre inconvénient de cette technique est qu'elle augmente sensiblement la puissance moyenne relative.

III.5.3. Le Selective Mapping (SLM):

L'idée est de multiplier la séquence de symboles complexe issus de la modulation numérique par une série de L différents vecteurs de façon à ce que ne soit retenue que le produit aux « PAPR » le plus faible (après la IFFT).

Cette méthode nécessite néanmoins la transmission d'une information de redondance pour que le récepteur identifie le vecteur optimal. On doit cette technique à R-Bauml, R-Fisher et

J.Huber[13]. Elle a été ensuite détaillée et agrémentée de précision par S.H Muller et J.B Huber.

Cette méthode s'applique à l'OFDM pour un nombre quelconque de sous-porteuses et pour une modulation numérique quelconque.

Soit $X = \{X_k\}$, $k = 0, \dots, N-1$, le vecteur symbole OFDM dans le domaine fréquentiel. L'idée de la technique « Selective Mapping » est de multiplier le vecteur X par un vecteur $\Phi^{(u)} = \{\phi_k^{(u)}\}$, $k = 0, 1, \dots, N-1$ les $\phi_k^{(u)}$ sont de la forme :

$$\phi_k^{(u)} = e^{j\varphi_k^{(u)}} \quad \text{avec} \quad \varphi_k^{(u)} \in [0, 2\pi], u = 0, 1, \dots, U-1$$

Le nouveau signal OFDM dans le domaine fréquentiel après pondération s'écrit :

$$X^{(u)} = X \cdot \Phi^{(u)} \quad (\text{III. 36})$$

On obtient ainsi U signaux différents de N composantes. Finalement, le signal OFDM temporel transmit s'écrit :

$$x^{(u^*)} = \text{IDFT}(X^{(u^*)}) \quad (\text{III. 37})$$

où u^* est l'indice correspond au signal OFDM dont le PAPR est le plus faible. La valeur de l'indice u^* sera transmise au receveur pour la reconstruction via un code correcteur d'erreurs. Les auteurs proposent que le nombre de bit sur lequel doit être codé cet indice soit de l'ordre de $\log_2 U$. Le principe de la technique est illustré par la figure III.12.

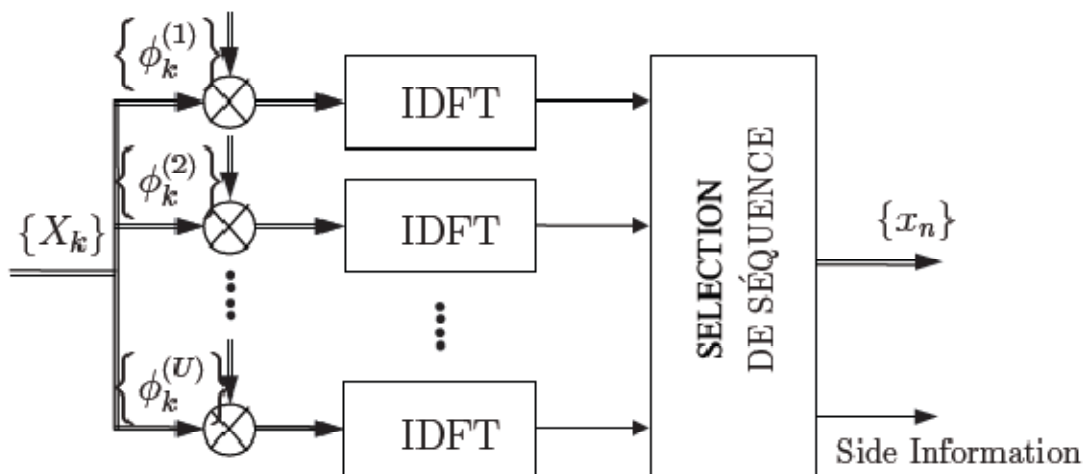


Fig.III.12 : Schéma d'un modulateur « Selective Mapping »

- **Inconvénients :**

L'inconvénient majeur de cette technique est sa complexité, du fait de l'utilisation de plusieurs (U) opérations d'IDFT. De plus, cette méthode nécessite la transmission de séquences d'information (« Side Information ») pour que le récepteur identifie la séquence qui a permis de générer le PAPR le plus faible. L'inconvénient de transmettre une information entre l'émetteur et le récepteur est double : d'une part à cause du risque que cette séquence soit entachée d'erreurs via le canal de transmission et d'autre part à cause d'une diminution de débit utile.

IV.5.4. La méthode “Partial Transmit Séquences ” (PTS) :

La technique “Partial Transmit Séquences ” s'inscrit dans la continuité de celle de “ Selective Mapping”. Elle à été proposé par S.H Muller et J.B Huber [14].

L'idée de cette méthode est de tronquer le train des N porteuses en V blocs de $\frac{N}{V}$ porteuses. Une porteuse utilisée dans un bloc particulier sera mise à zéros dans tous les autres. Une fois ces $\frac{N}{V}$ blocs formés, l'idée de “Selective Mapping” est appliquée.

Un vecteur $\Phi^{(v)} = \{\phi_k^{(v)}\}$, $v = 1, \dots, V$ effectuera une pondération de chacun des V blocs après IDFT pour former le signal final au PAPR le plus faible.

Comme illustré sur la figure III.13 et la figure III.14, l'algorithme de PTS est comme suit :

- (i) Le symbole OFDM fréquentiel X de N porteuses est tranqué en V

sous-blocs disjoints $X^{(v)}$ de $\frac{N}{V}$ porteuses tel que : $X = \sum_{v=1}^V X^{(v)}$.

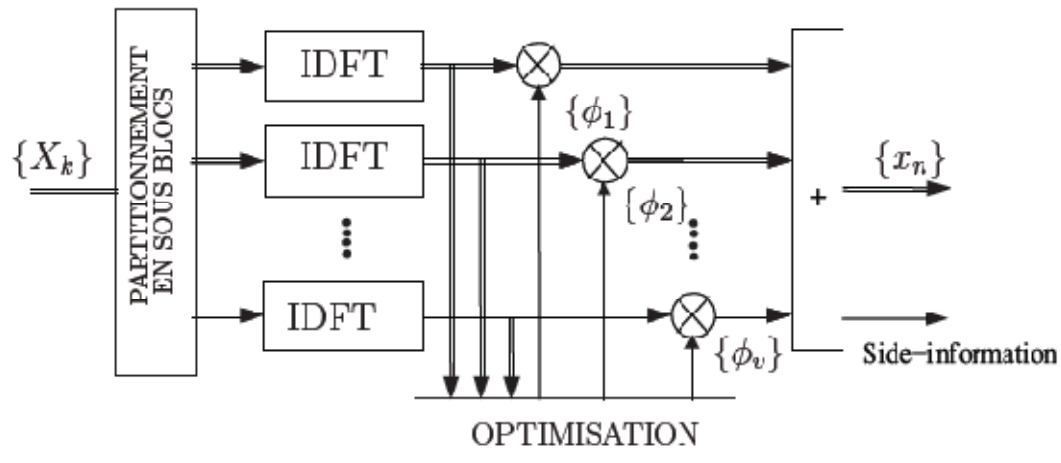


Fig.III.13 : Schéma d'un modulateur « Partial Transmit Sequences »

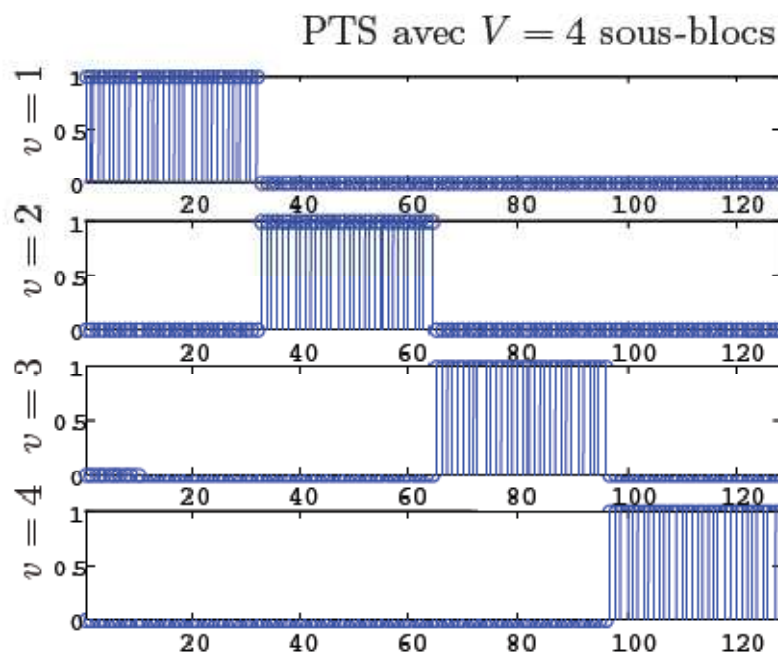


Fig.III.14 : Exemple de partitionnement d'un symbole en sous blocs pour application de la technique PTS

- (i) A chaque sous-blocs disjoints $X^{(v)}$, on applique un décalage de phase et le nouvel symbole OFDM fréquentiel s'écrit :

$$X = \sum_{v=1}^V X^{(v)} \cdot \phi^{(v)}, \quad \phi^{(v)} = e^{j\varphi^{(v)}}, \quad v = 1, \dots, V. \quad (\text{III. 38})$$

Avec $\varphi^{(v)} \in [0, 2\pi]$.

Remarque :

Afin de simplifier l'implémentation de cette technique on utilise le vecteur pseudo-aléatoire $\phi^{(v)} \in \{\pm 1, \pm i\}$

(ii) Le symbole OFDM temporel x s'écrit alors :

$$x = \text{IDFT} \left(\sum_{v=1}^V X^{(v)} \cdot \phi^{(v)} \right) = \sum_{v=1}^V \phi^{(v)} \cdot \text{IDFT}(X^{(v)}) = \sum_{v=1}^V \phi^{(v)} \cdot x^{(v)}. \quad (\text{III. 39})$$

La taille d'IDFT obtenant par (III.39) est de NL , où L est coefficient de

Sur-échantillonnage ; ainsi l'équation (III.39) peut être réécrite sous la forme suivante :

$$x = \begin{bmatrix} x_{1,1} & \cdots & x_{1,V} \\ \vdots & \ddots & \vdots \\ x_{NL,1} & \cdots & x_{NL,V} \end{bmatrix} \begin{bmatrix} e^{j\phi^{(1)}} \\ \vdots \\ e^{j\phi^{(V)}} \end{bmatrix} \quad (\text{IV. 14})$$

Le receveur réalisera la même récomposition, et multipliera par l'inverse des valeurs $\phi^{(v)}$.

- **inconvénients**

La façon dans les symboles est partitionnée en sous-blocs à une influence sur la performance et la complexité de la technique. L'inconvénient majeur de la technique PTS réside dans la complexité de la recherche des vecteurs de pondération $\Phi^{(v)}$ pour minimiser le PAPR. En effet, en considérons V sous blocs et des vecteurs $\Phi^{(v)}$ ($\Phi^{(v)}$, $v = 1, 2$ sont uniquement composés de 1 ou -1), le nombre de combinaison possible est de 2^V qui doivent être toutes passées en revue pour déterminer le jeu de vecteur qui minimise le PAPR. L'idée proposée dans [90] par A.D.S.Jayalath et C.Tellambura est alors de stopper le processus de recherche de vecteur $\Phi^{(v)}$ dès lors que le PAPR voulu est atteint. Un autre inconvénient est qu'elle nécessite la transmission de « Side Information » (SI) pour que le récepteur identifie la séquence qui a permis de générer le PAPR le plus faible.

III.6. Conclusion :

Dans ce chapitre, quatre grandes classes de techniques de réduction du PAPR ont été montrées différents schémas et équations ont été présentés, permettant la bonne compréhension du chapitre suivant et de tirer une conclusion sur le problème des forts pics pour les systèmes multi-porteuses OFDM.

IV .1. Introduction :

Après avoir étudié et présenté la technique OFDM, les méthodes de réduction du PAPR, on va essayer d'évaluer les performances d'un système multiporteur OFDM surtout en termes de « BER », « ACPR » et « PAPR ».

Pour cela on essayera de concevoir une plate forme de simulation informatique permettant de mieux comprendre ces phénomènes.

L'objet de ce chapitre est donc de simuler et présenter dans un premier temps, les effets de non linéaires introduits par les amplificateurs de puissance (PA). On va tout d'abord exposer l'influence de la caractéristique non linéaire de l'amplificateur sur le spectre de signal OFDM, l'effet sur le taux d'erreur binaire (BER) et les diagrammes de constellation. Ensuite dans un deuxième temps, on implémente la méthode Partial Transmit Séquences (PTS) pour diminuer le « PAPR » du signal.

Les simulations ont été réalisées à l'aide du logiciel « MATLAB ».

IV.2. Effets non linéaires de l'amplification :

Comme nous l'avant déjà évoqué dans le chapitre III, l'inconvénient majeur du signal OFDM est qu'il a une enveloppe non constante avec une fluctuation importante d'amplitude et de puissance, cette particularité rend le signal OFDM très sensible au non linéarités des composantes analogiques (en particulier la non linéarité de l'amplificateur de puissance) qui présentent des caractéristiques de transfert avec saturation de la chaine de transmission. Donc dans cette partie de simulation on va étudier avant et après amplification du signal : le spectre du signal OFDM, le «BER » et « CCDF » pour différentes valeur d' « IBO ».

Pour ce faire, nous avons fait un programme qui simule une chaine OFDM en émission et en réception. Les paramètres de la modulation OFDM utilisés dans notre application sont illustrés dans le tableau IV.1.

Nombre de porteuses N	128
Nombre de porteuses utiles	64
Nombre de porteuses nulles	64
Nombre de symboles OFDM	480
T_s : la période d'un symbole OFDM	$2.44 \mu s$
Préfixe Cyclique (PC)	16 échantillons
Type de modulation numérique	MAQ-16

Tab. IV.1 : les paramètres de la modulation OFDM

IV.2.1. Effet sur le spectre OFDM :

Pour visualiser l'effet de l'amplificateur sur le spectre OFDM nous avons mesuré la densité spectrale de puissance (DSP) du signal OFDM avant et après l'amplification. Le programme qui réalise la partie émission de la chaîne OFDM, est donné par l'organigramme de la figure IV.1

L'amplificateur de puissance utilisé est de type SSPA, ces caractéristiques de transferts AM/AM et AM/PM sont données par les équations (III.29 et III.30). Ses principaux paramètres sont le « Input Back Off » équivalent, $IBO = [0, 1, 2]$ dB, le gain d'amplificateur $= 1$.

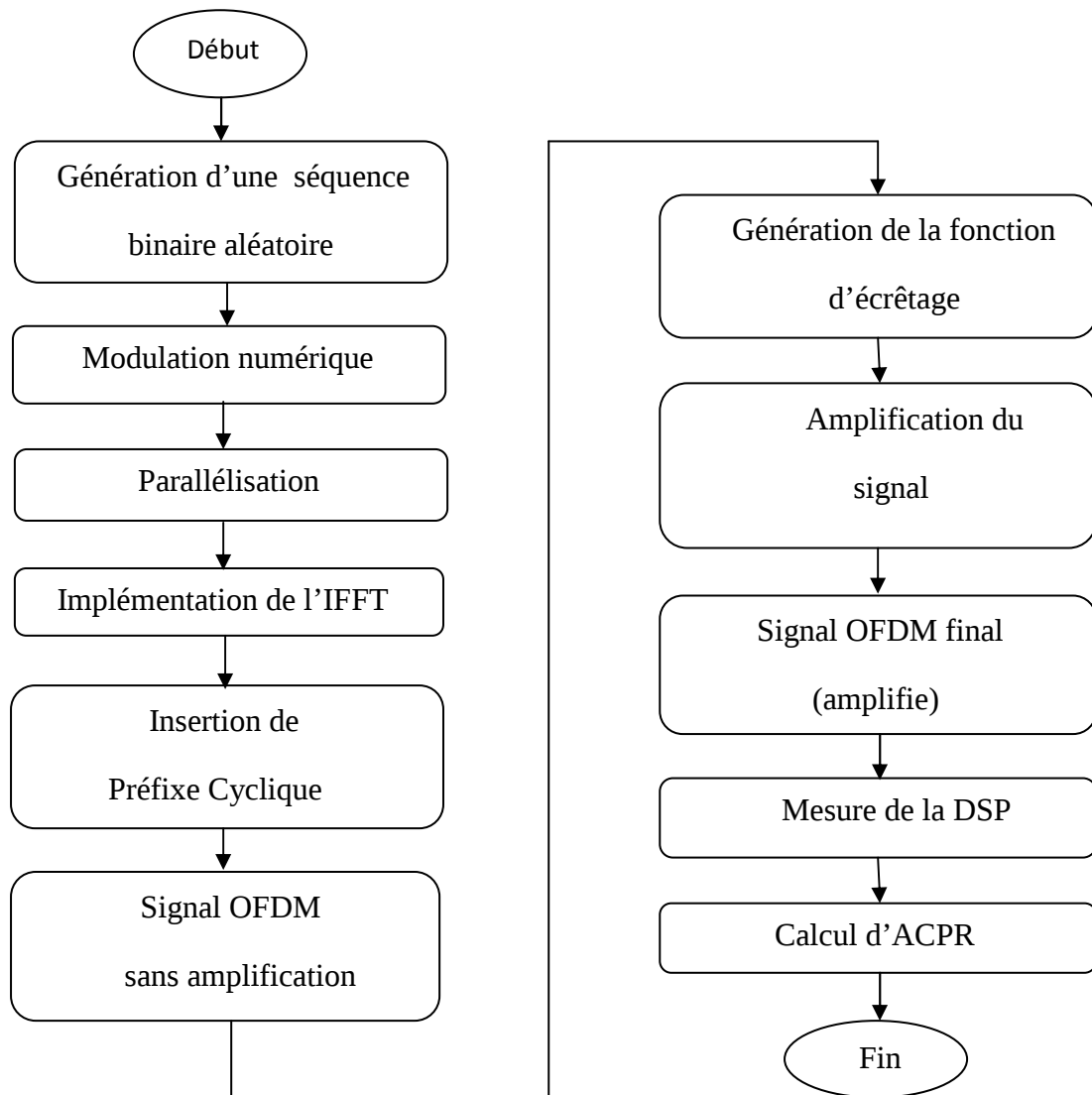


Fig.IV.1 : Organigramme de simulation d'une chaîne OFDM (partie émission)

Le résultat de simulation est représenté dans la figure IV.2. Nous illustrons le phénomène de remontée spectrale après amplification du signal ce qui peut entraîner des interférences avec les canaux adjacents, ainsi nous remarquons la remontée spectrale augmente quand le recul diminue ; En effet, des valeurs faible de l' « IBO » signifient que l'amplificateur de puissance fonctionne en limite de sa zone de saturation. C'est dans cette zone que les signaux subissent le plus de distorsions ce qui explique la remontée spectrale de plus en plus importante lorsque l' « IBO » devient de plus en plus faible. Le facteur qui permet de mesurer ces distorsions est l'ACPR qui est défini dans l'équation III.26. Le tableau IV. 2 montre les résultats de calcul

d'ACPR pour différentes valeurs de l' « IBO », nous remarquons plus le recul diminue plus l'ACPR très important.

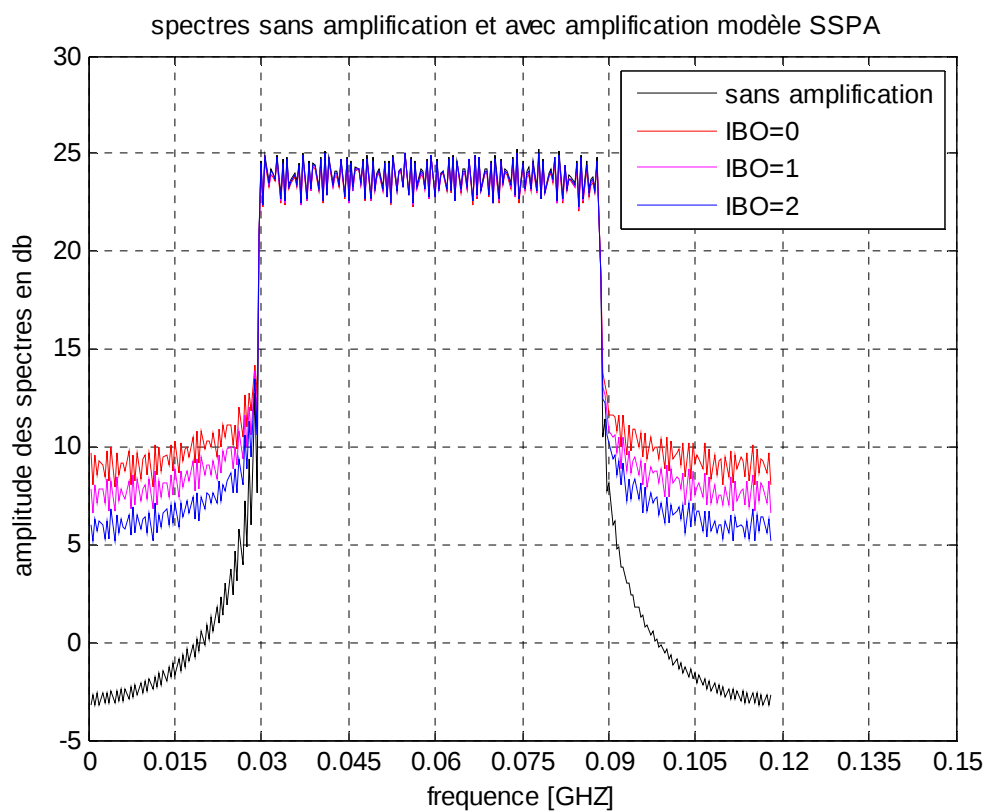


Fig.IV.2 : Effets non linéaires d'amplification sur le spectre pour différentes valeurs d'IBO

IBO [dB]	0	1	2
ACPR [dB]	-22.53	-23.16	-24.02

Tab. IV.2 : Mesure de l' « ACPR » pour différentes valeurs du recul « IBO »

IV.2.2. Effet sur le taux d'erreur binaire (BER) :

La saturation du signal OFDM n'a pas que des effets sur le spectre du signal mais aussi sur le taux d'erreur binaire « BER ». Pour ce faire nous avons calculé en réception le « BER » en fonction de (E_b/N_0) sans et avec amplification du signal OFDM généré en émission.

Dans cette partie l'amplificateur utilisé est de type TWT dont les conversions AM/AM et AM/PM sont données par les expressions (III.26 et III.27) respectivement. Ses principaux paramètres sont : le $IBO = [0, 3, 6, 9]$ dB, le gain d'amplificateur $\nu = 1$. Le signal OFDM généré est le même que le précédent et le canal de transmission est un canal bruit blanc additif gaussien (BBAG).

Les étapes de simulation sont résumées dans l'organigramme de la figure IV.3.

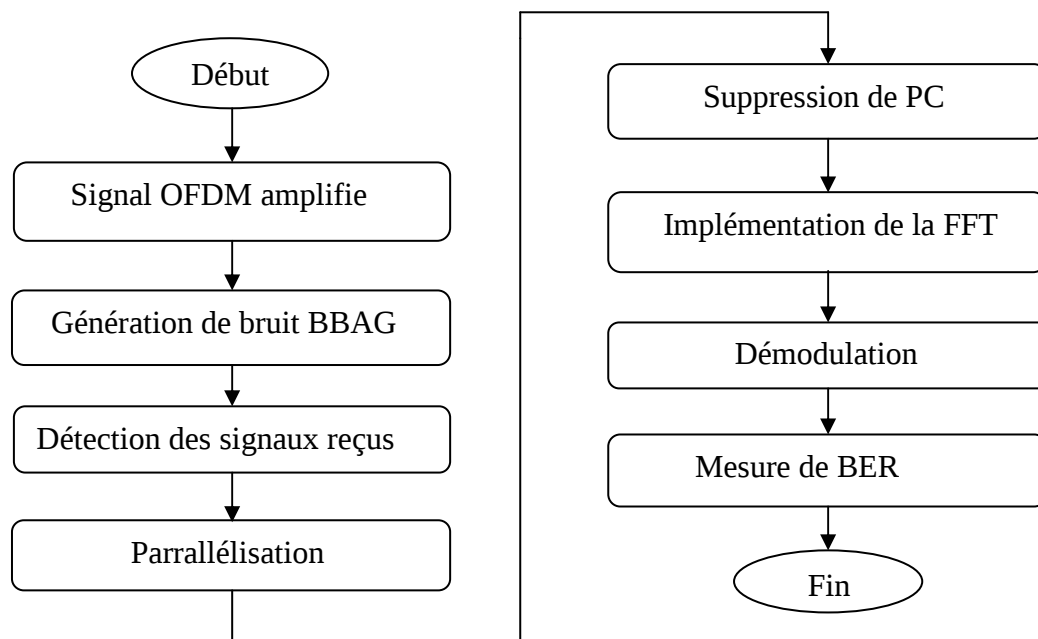


Fig.IV.3 : Organigramme de simulation d'une chaîne OFDM (partie réception)

Le résultat de simulation est représenté dans la figure IV.4 qui présente le « BER » en fonction de (E_b/N_0) pour différentes valeurs d'« IBO ». Nous pouvons constater une augmentation du taux d'erreurs binaire pour de faibles valeurs de recul. L'amplificateur travaille dans sa zone non-linéaire. Cependant lorsque l'« IBO » devient important, le BER

tend à se confondre avec la courbe du signal sans amplification. C'est la preuve qu'il y a moins en moins des perturbations liées aux non linéarité dans sa région non linéaire.

Nous donnons également à travers la figure IV.5 la constellation de signal de sortie pour différentes valeurs d'« IBO », nous remarquons un chevauchement dans les quatre cas de reculs. Nous observons aussi que plus l'« IBO » est petit plus le chevauchement est important, donc la probabilité d'erreurs augmente.

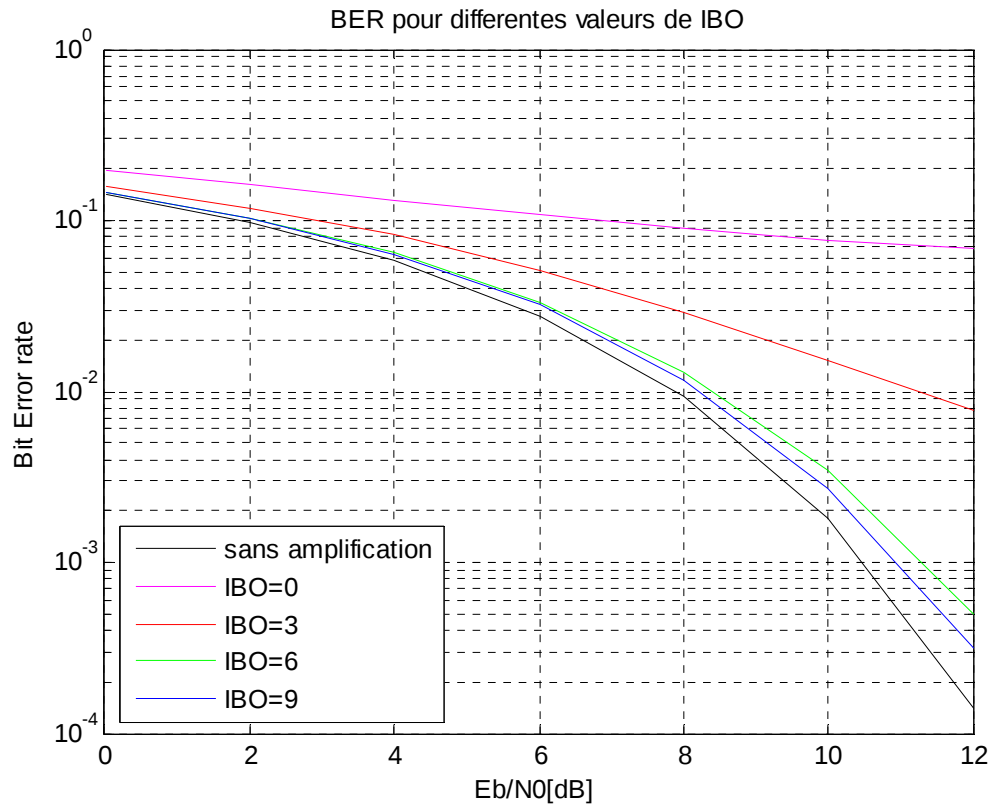


Fig. IV.4 : « BER » avant et après amplification en fonction de (E_b/N_0)

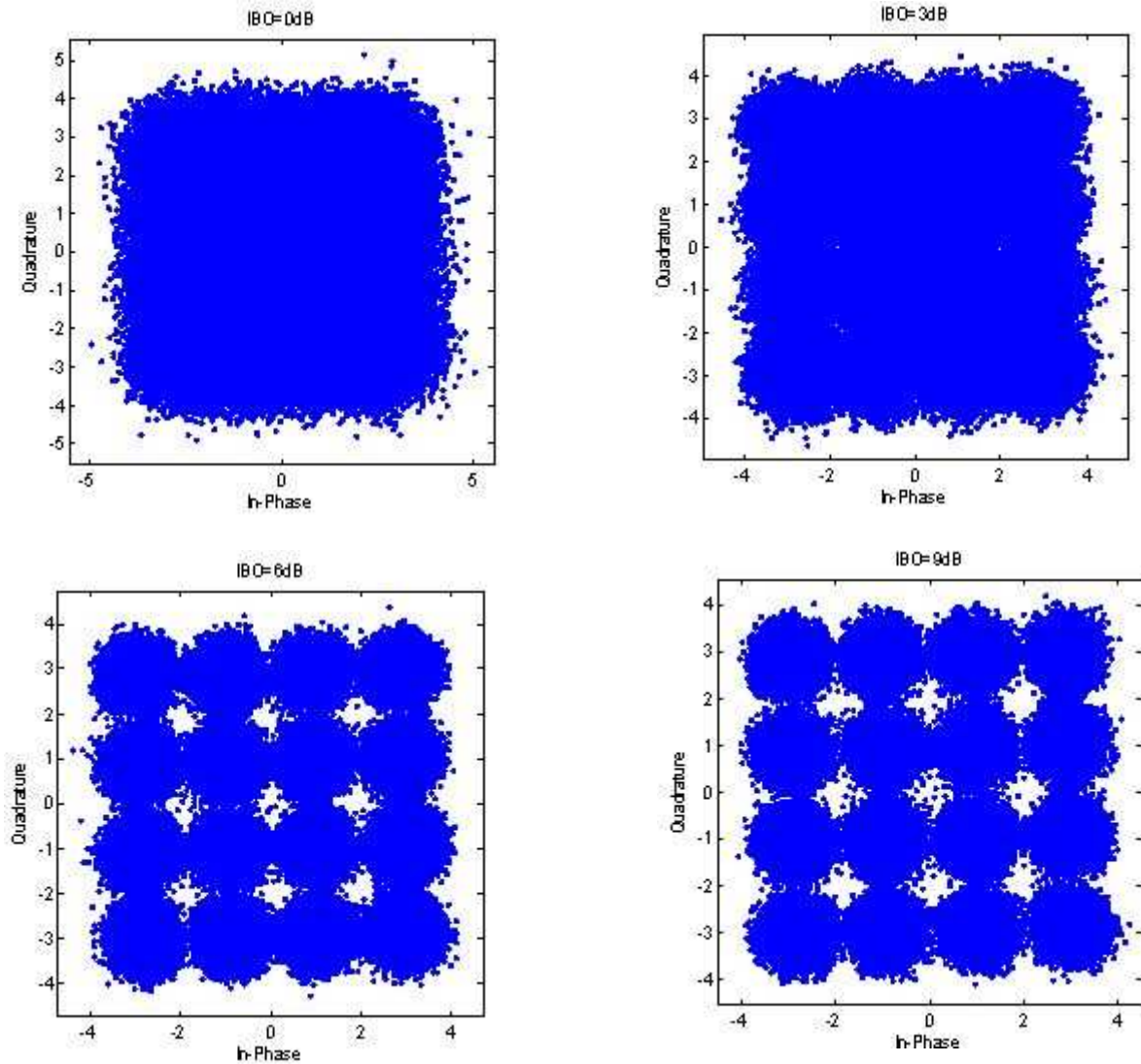


Fig. IV.5 : Impact du non linéarité du « PA » modèle TWT illustré dans les constellations

Remarque :

Le nombre de sous porteuses dans une trame OFDM influe sur la fonction de répartition du « PAPR » (CCDF) la relation entre ces deux est donnée par l'équation suivante :

$$Prob(PAPR > \gamma_0) = 1 - (1 - e^{-\gamma_0})^N$$

La figure IV.6 est obtenue après le calcul du CCDF pour des trames OFDM de sous porteuses N différentes. Elle montre que plus le nombre de sous porteuses est grands plus la probabilité d'avoir un $(PAPR > \text{papr seuil})$ augmente, car une trame OFDM de N sous

porteuses c'est une somme de N exponentielles, donc si N augmente la probabilité d'avoir des pics augmente aussi

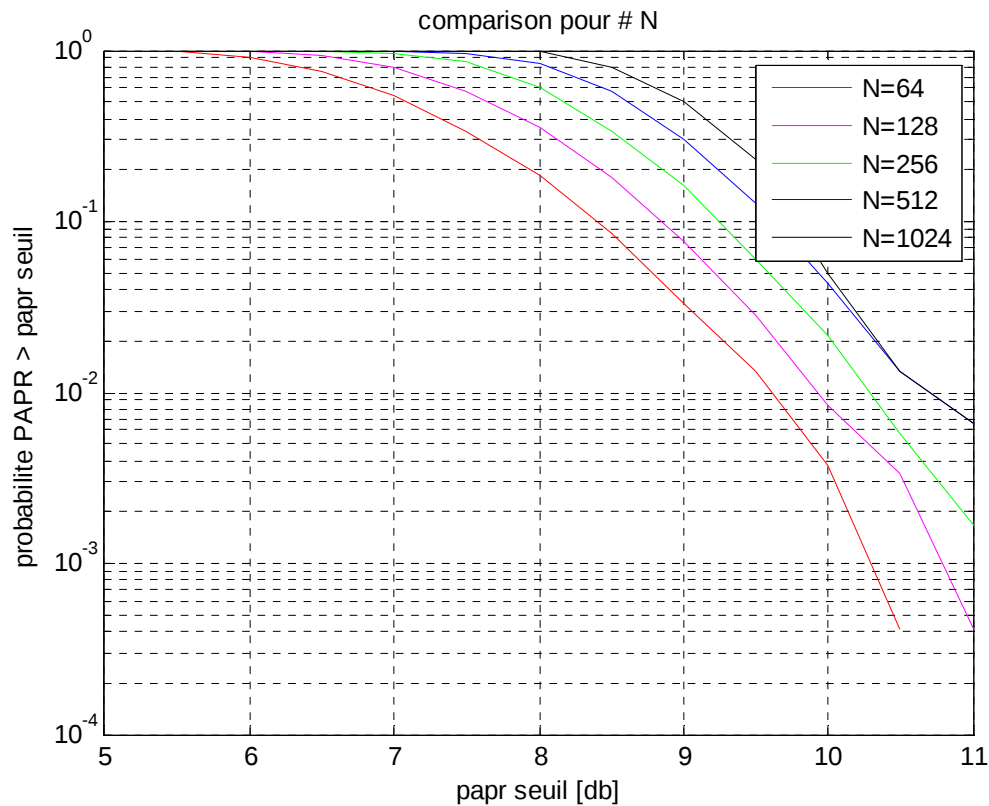


Fig. IV.6 : CCDF des symboles OFDM pour N variables

IV.3. Implémentation de la méthode « Partial Transmit Séquence » (PTS) :

Dans l'étude précédente, nous avons vu qu'en présence de l'amplificateur de puissance, les signaux à enveloppes non constante ou à fort PAPR subissent des distorsions. Malheureusement, plus le PAPR est élevé plus les distorsions sont importantes. Donc pour éviter ces problèmes, on doit travailler dans la zone fortement linéaire du « PA ».

On peut penser donc à augmenter l'« IBO », mais cela va accroître la consommation énergétique. Afin de se reprocher de la zone sans trop saturer le signal d'entrée, il faut alors réduire les fluctuations d'enveloppe de signal OFDM, donc son PAPR, qui est l'objectif de notre travail.

Dans notre application nous avons choisi la méthode « Partial Transmite Séquence » (PTS).

Cette méthode utilise les vecteurs pseudo-aléatoires pour minimiser le « PAPR ».

IV.3.1. Application de la méthode PTS sous MATLAB :

Nous avons fixé le nombre de sous blocs $V = 2$ et pour simplifier l'implémentation on utilise le vecteur aléatoire $\phi^{(v)} \in \{-1, +1, -j, +j\}$ avec $v = 1 \dots V$ alors le nombre de combinaisons possibles est $M = 4^V$ qui doivent être toutes passées en revue pour déterminer le vecteur qui minimise le PAPR mais nous en arrêtons lorsque le nouveau PAPR (PAPR_{new}) soit inférieur au PAPR original (PAPR_0).

L'Organigramme de la figure IV.6 résume les étapes de simulation de la chaîne OFDM en émission en utilisant la méthode « PTS ».

Les paramètres de la modulation de signal OFDM original sont résumés dans le tableau IV.3.

Paramètres	valeurs
Modulation	MAQ-16
T_u (durée symbole)	$6.4 \mu s$
Nombre de porteuse	128
Nombre de symboles OFDM	9600

Tab. IV.3 : Les paramètres de la modulation OFDM

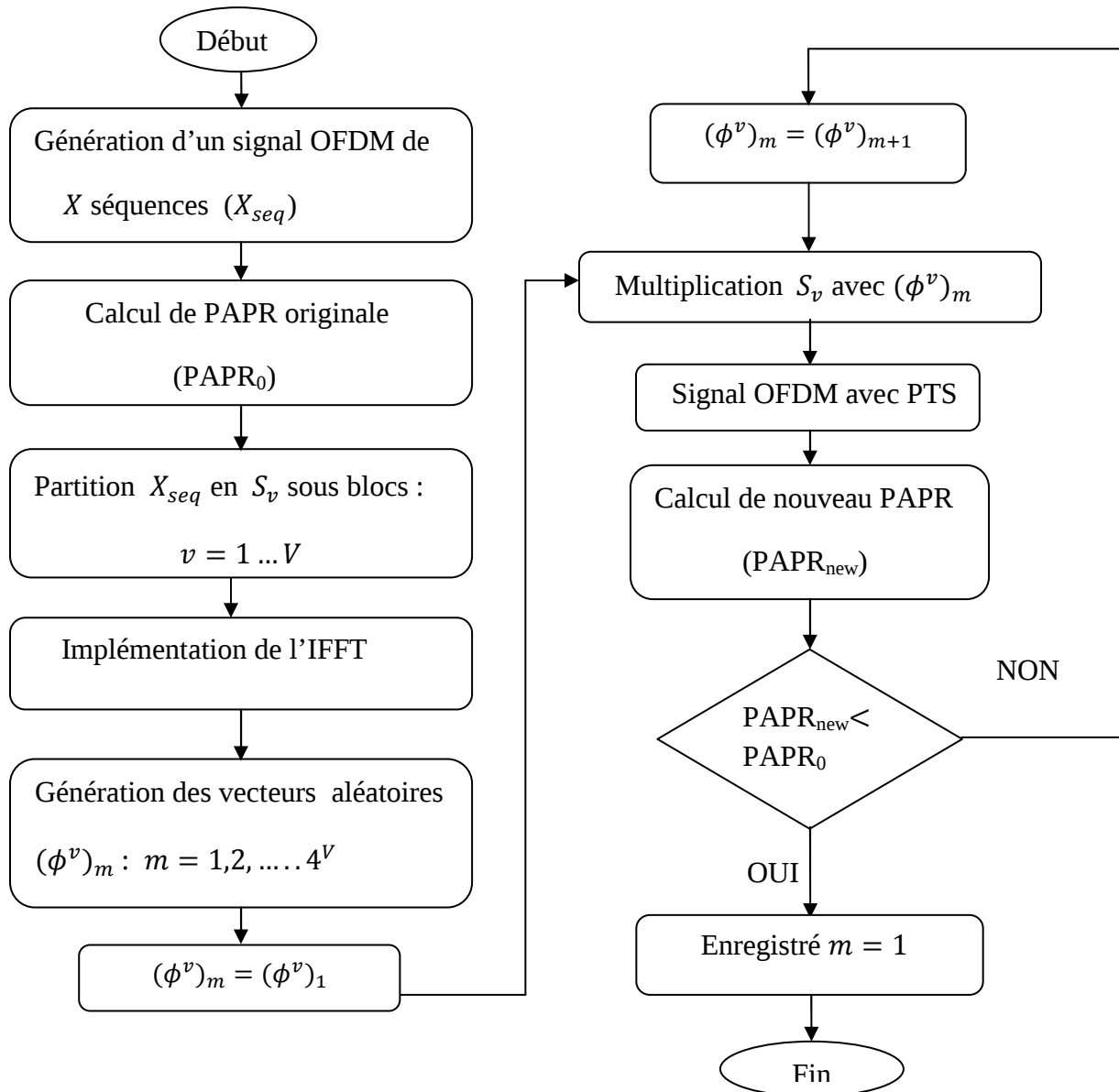


Fig.IV.7 : Organigramme de simulation de la chaîne OFDM avec PTS

IV.3.2. Résultats de simulation :

- **Effet de la méthode PTS sur le PAPR :**

La méthode PTS opère une réduction du « PAPR » du signal OFDM à émettre. Cette impact sur le « PAPR » est présenté en figure IV.8 et pour un nombre de vecteurs $V = 2$ par exemple on trouve que la réduction est 2 dB pour une « CCDF » entre 10^{-3} et 10^{-4} . Cela confirme bien l'efficacité de la méthode PTS comme technique de réduction du PAPR de signal OFDM.

Ensuite nous avons augmenté le nombre de vecteur à $V = 4$, on remarque que plus on augmente le nombre de vecteur plus la probabilité d'avoir ($\text{PAPR}_{\text{new}} > \text{papr}_{\text{seuil}}$) diminue car la probabilité de trouver un vecteur qui minimise le « PAPR » augmente.

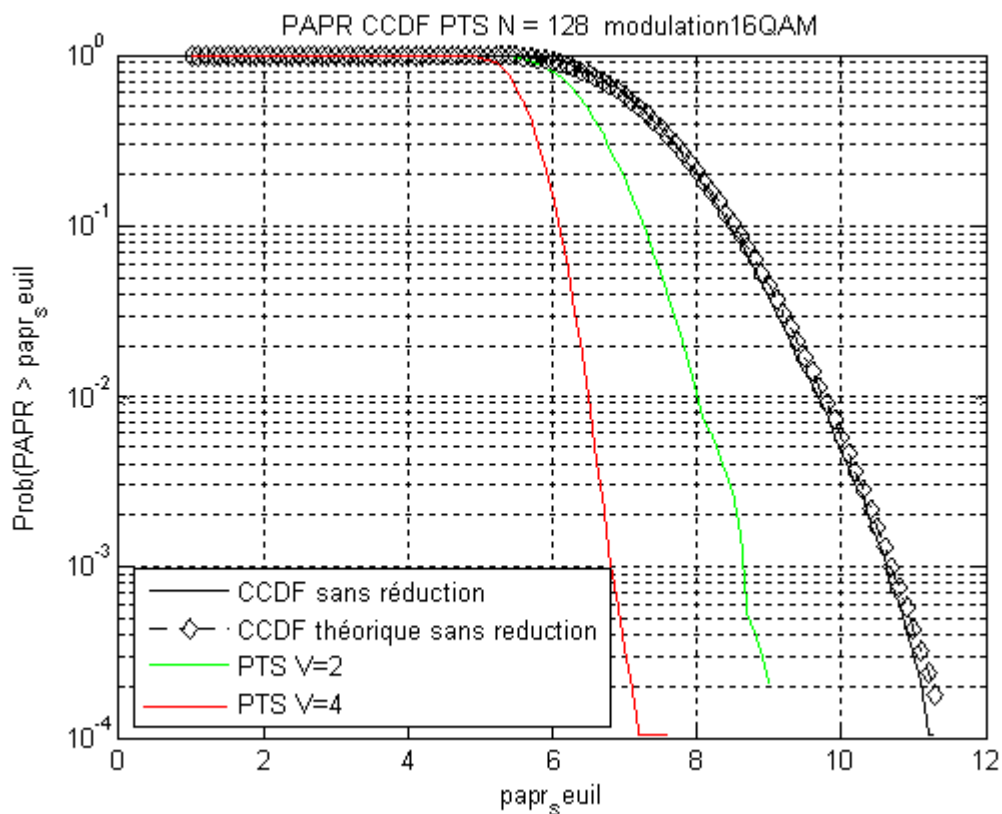


Fig. IV.8 : Réduction du « PAPR » à l'aide de la méthode PTS.

On trace un échantillon de signal sans puis avec la technique PTS : on se rend compte que la technique PTS génère des forts pics, dans le PAPR reste inférieur au PAPR initial, ce qui est constaté dans la figure IV.8. La figure IV.9 montre l'apparition de nouveaux pics dont l'amplitude est inférieure a celle du fort pic de signal initial.

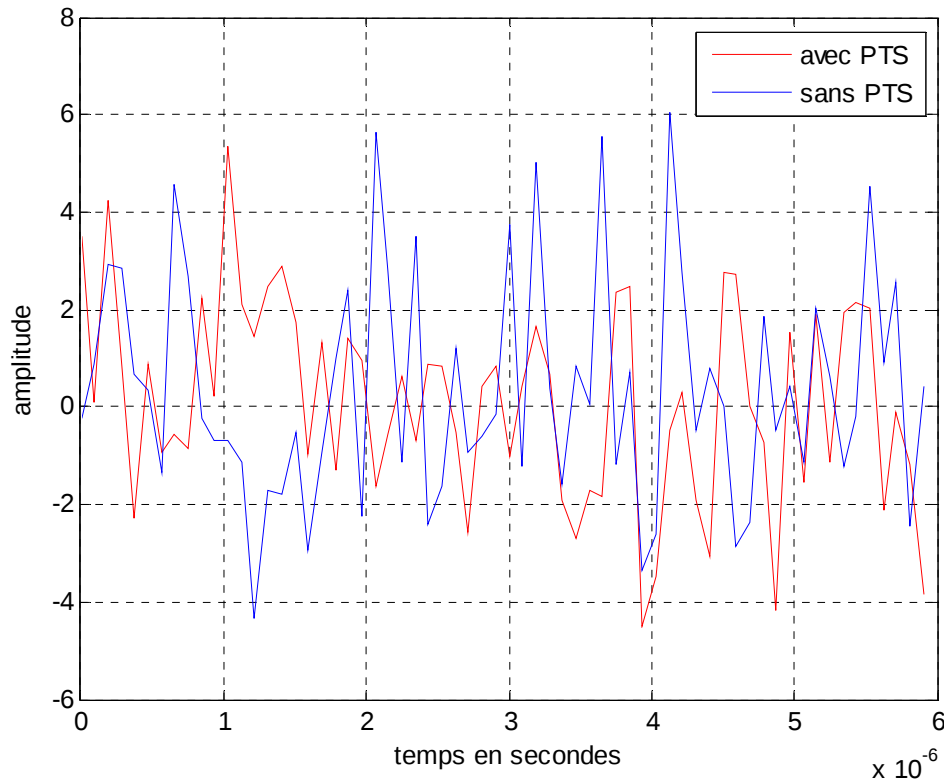


Fig. IV.9 : Génération de nouveaux pics après le module de PTS

IV.4. Conclusion :

Dans ce chapitre nous avons implémenté la nouvelle technique de réduction de PAPR Partial Transmit Séquence pour des signaux OFDM qui a des fortes fluctuations en amplitude dans un contexte d'amplification non-linéaire de ces signaux.

Cette méthode nous permet de choisir le vecteur qui minimise le PAPR. Les résultats de simulation ont montrée cela.

Conclusion

Ce mémoire a été consacré à l'étude de la réduction du « PAPR » par la méthode Partial Transmit Séquences.

Aujourd'hui plusieurs standards reposent sur la modulation OFDM, car elle permet de traiter simplement le phénomène d'interférences dû à la sélectivité fréquentielle, en transformant le canal de propagation en canaux à évanouissements plats, cependant le signal OFDM temporel possède un fort « PAPR ». La solution sous optimale du recul était alors privilégiée jusqu'à l'application des nouvelles méthodes de réduction du « PAPR » plus efficaces. Nous nous sommes intéressées particulièrement à la méthode PTS, qui consiste à décomposer la séquence de symboles OFDM en sous-blocs et les multiplier par des vecteurs de phases ensuite sélectionner le vecteur qui minimise le PAPR.

Finalement, les résultats obtenus dans l'étude pratique justifient ce qui est annoncé dans l'étude théorique, donc le problème sur lequel repose ce travail est traité et le but de ce dernier est atteint.

Rappels sur l'enveloppe complexe

L'expression de l'enveloppe complexe $\tilde{x}(t)$ associée au signal modulé $s(t)$ est définie par l'équation suivante :

$$\tilde{x}(t) = \dot{x}(t) e^{-2j\pi f_c t} = (s(t) + j\tilde{x}(t)) e^{-2j\pi f_c t} \quad (\text{A.1})$$

Ici $\dot{x}(t)$ représente le signal analytique associé à $s(t)$. Ce signal complexe a sa partie réelle qui coïncide avec le signal $s(t)$,

$$s(t) = \mathcal{R}\{\dot{x}(t)\} \quad (\text{A.2})$$

et sa partie imaginaire est égale à :

$$\hat{s}(t) = \frac{1}{\pi} \int_{-\infty}^{+\infty} \frac{s(\theta)}{t-\theta} d\theta \quad (\text{A.3})$$

Le signal $\hat{s}(t)$ de l'équation (A.3) est la transformée de Hilbert de $s(t)$. Cette transformée est une opération linéaire, invariante dans le temps, définie par une fonction de transfert en fréquence $H(f)$ et une réponse impulsionnelle $h(t)$.

$$H(f) = -j \operatorname{sign}(f), \quad h(t) = \frac{1}{\pi} vp\left(\frac{1}{t}\right) \quad (\text{A.4})$$

La notation $vp()$ représente la distribution valeur principale. D'après ces définitions, l'expression de l'équation (A.3) peut être vue comme le produit de convolution entre le signal $s(t)$ et la réponse impulsionnelle $h(t)$

Relation entre le « PAPR » en bande de base et en radiofréquences

Fonction d'autocorrélation et puissance

On réécrit d'abord l'enveloppe complexe $\tilde{x}(t)$ du signal modulé en fonction de ces composantes en phase et en quadrature [équation (B.1)] et sous forme de $a(t)e^{j\theta(t)}$ [équation (B.2)].

$$\tilde{x}(t) = I(t) + jQ(t) \quad (\text{B.1})$$

$$= a(t)e^{j\theta(t)} \quad (\text{B.2})$$

où

$$a(t) = \sqrt{I^2(t) + Q^2(t)} \quad (\text{B.3})$$

$$\theta(t) = \arctan\left(\frac{Q(t)}{I(t)}\right) \quad (\text{B.4})$$

Ensuite le signal modulé $s(t)$ peut s'écrire comme suit :

$$\begin{aligned} s(t) &= \mathcal{R}[\tilde{x}(t)e^{jw_c t}] = \mathcal{R}[a(t)e^{j\theta(t)}e^{jw_c t}] \\ &= a(t)\cos(w_c t + \theta(t)) \end{aligned} \quad (\text{B.5})$$

Nous allons alors calculer les fonctions d'autocorrélation $\phi_{ss}(\tau)$ et $\phi_{II}(\tau)$ et $\phi_{QQ}(\tau)$ pour les signaux $s(t)$, $I(t)$ et $Q(t)$ respectivement. Cela nous permet d'en déduire les puissances moyennes des signaux respectifs si $\tau = 0$.

On sait que la fonction d'autocorrélation se calcule de la manière suivante :

$$\phi_{ss}(\tau) = E[s(t + \tau)S^*(t)] \quad (\text{B.6})$$

Après quelques calculs, on peut montrer que les fonctions d'autocorrélation de $I(t)$ et $Q(t)$ se relient à la fonction d'autocorrélation de $s(t)$ ($\phi_{ss}(\tau)$) et à sa transformée de

Hilbert ($\hat{\phi}_{ss}(\tau)$) :

$$\phi_{II}(\tau) = \phi_{QQ}(\tau) = \phi_{ss}(\tau) \cos(w_c \tau) + \hat{\phi}_{ss}(\tau) \sin(w_c \tau) \quad (\text{B.7})$$

si $\tau = 0$

$$\phi_{II}(0) = \phi_{QQ}(0) = \phi_{ss}(0) \quad (\text{B.8})$$

On sait qu'à $\tau = 0$ ces fonction d'autocorrélation correspondent aux puissances moyennes des différents signaux. Donc la puissance de signal modulé $s(t)$ est égale à celle des composantes en phase ($I(t)$) et en quadrature ($Q(t)$) de l'enveloppe complexe.

$$P_{moy}(I) = P_{moy}(Q) = P_{moy}(S) \quad (\text{B.9})$$

En plus depuis l'équation (B.2),

$$\begin{aligned} P_{moy}(\hat{x}) &= E[|\hat{x}(t)|^2] = E[|a(t)|^2] = P_{moy}(a) \\ &= E[I^2(t) + Q^2(t)] = E[I^2(t)] + E[Q^2(t)] \\ &= P_{moy}(I) + P_{moy}(Q) \end{aligned} \quad (\text{B.10})$$

Et depuis l'équation (B.9),

$$P_{moy}(\hat{x}) = P_{moy}(a) = 2P_{moy}(s) \quad (\text{B.11})$$

Donc la puissance moyenne du signal complexe en bande de base $\tilde{x}(t)$ est égale à 2 fois la puissance moyenne du signal modulé RF $s(t)$.

« PAPR »

Le « PAPR » classique (en radiofréquences) est défini comme le rapport entre la puissance maximale et la puissance moyenne du signal $s(t)$:

$$PAPR_{RF} = \frac{P_{max}(s)}{P_{moy}(s)} \quad (\text{B.12})$$

$$= \frac{\max_{t \in [0, T]} |s(t)|^2}{P_{moy}(s)} \quad (\text{B.13})$$

Et depuis l'équation (C.1),

$$PAPR_{RF} = \frac{\max_{t \in [0, T]} |\mathcal{R}[x(t)e^{jw_c t}]|^2}{P_{moy}(s)} \quad (\text{B.14})$$

D'après les auteurs de [pal 2005b] :

$$|\mathcal{R}[\hat{x}(t)e^{jw_c t}]|^2 \leq |\hat{x}(t)|^2 \quad (\text{B.15})$$

Et donc

$$\max_{t \in [0, T]} |\mathcal{R}[\hat{x}(t)e^{jw_c t}]|^2 \leq \max_{t \in [0, T]} |\hat{x}(t)|^2 \quad (\text{B.16})$$

Et si on considère aussi que $P_{moy}(s) = 1/2 \cdot P_{moy}(\hat{x})$ alors en déduit depuis l'équation (B.14) que le $PAPR_{RF}$ vaut :

$$PAPR_{RF} \leq 2 \cdot \frac{\max_{t \in [0, T]} |\hat{x}(t)|^2}{P_{moy}(\hat{x})} \quad (\text{B.17})$$

Le deuxième terme de l'équation (B.17) représente bien 2 fois le « PAPR » en bande de base ($PAPR_{BdB}$), d'où

$$PAPR_{RF} \leq 2 \cdot PAPR_{BdB} \quad (\text{B.18})$$

Et enfin en dB,

$$PAPR_{RF}[dB] \leq 2.PAPR_{dB}[dB] + 3dB \quad (\text{B.19})$$

Enfin, nous tenons à préciser que l'identité est atteinte lorsque les puissances instantanées en radiofréquence et en bande de base ont le même maximum au même instant.

Ceci est facilement réalisable lorsque $f_c \geq 1 / T_s$ où f_c représente la fréquence de la porteuse RF et T_s représente le temps symbole. Cette condition est toujours vérifiée pour les télécommunications. C'est pourquoi l'on peut affirmer qu'en pratique le $PAPR_{RF}[dB]$ est égal au $PAPR_{dB}[dB] + 3dB$.

LISTE D'ABREVIATIONS

ACPR	Adjacent Channel P ower R atio.
AM/AM	Amplitude M odulation/Amplitude M odulation
AM/PM	Amplitude M odulation/ P hase M odulation
ADSL	Assymetrique D igital Subscriber L ine.
AWGN	Additive W hite G aussian N oise.
BER	B it E rror R ate
BBAG	B ruit B lanc A dditif G aussien
CDMA	C ode D ivision M ultiple A ccess.
CNA	Convertisseur N umérique A nalogique
CAN	Convertisseur A nalogique N umérique
CF	C rest F actor.
CCDF	C omplementary C umulative D istribution F unction
CP	C yclic P refix.
CR	C lipping R atio.
DAB	D igital A udio B roadcasting.
DFT	D iscrete F ourier T ransform.
DSP	D ensity S pectral of P ower.
DVB	D igital V ideo B roadcasting.
DVB-T	D igital V ideo B roadcasting – T errestrial.
FFT	F ast F ourier T ransform.
IBO	I nter B ack O f.
IDFT	I nverse D iscrete F ourier T ransform.
IFFT	I nverse F ast F ourier T ransform.
IM	I nter M odulation.

ICI	Inter-Carrier-Interference
ISI	Inter Symbol Interferen
OBO	Output Back Of.
OFDM	Orthogonal Frequency Division Multiplexing.
PA	Puissance Amplifier.
PAPR	Peak to Avereage Power Ratio.
PTS.	Partial Transmit Sequences
QAM	Quadrature Amplitude Modulation
Q-PSK	Quaternary Phase Shift Keying
SLM	Selected Mapping
SNR	Signal to Noise Ratio
SSPA	Solid-State Power Amplifier
TR	Tone Reservation
TWT	Travelling-Wave Tube.
WIFI	Wireless Fidelity
WIMAX	Worldwide Interoperability for Microwave Access

NOTATION

A_{sat}	L'entrée de saturation du PA.
A_{fond}	L'amplitude fondamentale du signal de sortie
B_c	La bande de cohérence.
C_k	Un symbole numérique.
D_b	Le débit numérique
d_{min}	La distance minimale entre les symboles
d_k	Le symbole numérique reçu
f_D	Fréquence doppler.
f_i	La fréquence de la porteuse
f_k	La fréquence de signal modulé
f_0	La fréquence centrale de canal
G_{fon}	Le gain fondamental
G_{lin}	Le gain linéaire
N	Le nombres des porteuses dans une trame OFDM.
N_p	Les nombre des trajets multiples.
N_0	Puissance du bruit
L	Le facteur de sur-échantillonnage.
P_i	La probabilité sélection d'un individu.
T_b	La durée d'un élément binaire.
t_0	L'instant d'échantillonnage
T_g	La durée de l'intervalle de garde
T	La durée de forme d'onde
T_c	Le temps de cohérence.
T_m	L'étalement temporel maximal.

T_s	La durée d'un symbole OFDM.
T_u	La durée utile d'un symbole OFDM.
τ_n	Le retard effectuant chaque trajet.
p_0	Un réel
P_e	La puissance d'entrée
P_s	Puissance de sortie
$P_{s,sat}$	La puissance de saturation en sortie
$P_{e,sat}$	La puissance de saturation en entrée
Q_L	La matrice de Fourier
V	Nombre de sous-blocs disjoints
W	Une bande de fréquence.
Z_n	La transformée de Fourier Discret Inverse
Δ	La durée d'un CP
Δf	L'espacement entre les N sinus cardinaux (sous porteuses)
α_n	L'atténuation complexe
τ_n	Le retard affectant chaque trajet.
γ_0	Un seuil donné
$\phi^{(v)}$	Vecteur de phase
$\ \cdot\ _\infty$	Norme infinie

BIBLIOGRAPHIE

- [1] : Olivier Berder « Optimisation et stratégie d'allocation de puissance de système de transmission multi-antennes » thèse doctorat, université de Bretagne occidentale ,2002 .
- [2] : Basel Rihawi « Analyse et réduction du power radio des systèmes de radio communication multi-antenne », thèse doctorat, université de Renne I, Mars 2008.
- [3] : Salvator Ragusa « Ecrêtage inversible pour l'amplification non-linéaire des signaux OFDM dans les terminaux mobiles », thèse doctorat, université Joseph Fourier, Juin 2006.
- [4] : Sidkièta Zabre « Amplification non linéaire d'un multiplex de porteuse modulée a fort facteur de crête », thèse doctorat, université de Rennes I, Avril 2007.
- [5] : Annick Le Glaunec « Modulations multiporteuses ».
- [6] : Sylvain TERTOIS « Réduction des effets de non linéarités dans une modulation multiporteuse à l'aide de réseaux de neurones », thèse de doctorat décembre 2003, Rennes.
- [7] : Patrice KADIONIK « Les modulations numériques »ENSEIRB 2000, Bordeaux
- [8]: Peng Liu,W.P.Zhu and M.O.Ahmad «A Phase Adjustment Based Partial Transmit Sequence Scheme For PAPR Reduction » , Birkhäuser Boston 2004.
- [9]: Pierre GRUYER et Simon PAILLARD « Modélisation d'un modulateur et démodulateur OFDM », thèse de doctorat décembre 2005, Bretagne.
- [10] : Désiré GUEL « Etude de nouvelles techniques de réduction de « facteur de crête » à compatibilité descendante pour les systèmes multiporteuses » , thèse de doctorat 2009.
- [11] : DAMERDJI Tidjani « Etude et Simulations des méthodes de réduction du PAPR pour les systèmes multiporteuses OFDM », thèse d'ingénieur en électronique USTHB 2008 .
- [12] : J.Téllado-Mourelo « Peak to average power reduction for multicarrier modulation » Thèse de doctorat , Université de Stanford, Sept. 1999.

[13]: R. Bäuml, R.Fisher,J.Huber «Reducing the peak-to-Average power Ratio of Multicarrier Modulation by Selecting Mapping» Electronics Letters,vol. 32. No.22.pp.2056-2057, Oct.1996.

[14]: S.Muller and J. Huber «OFDM with reduced peak-to-average power ratio by optimum combination of partial transmit sequences » Electronics Letters, vol. 33,pp.368-369, February 1997.

[15]: R. Van Nee and R. Prasad, « OFDM for wireless multimedia communications », universal personal communications, Artech House publishers, Chapter 6, janvier 2000