



Communications au-delà de la cadence de Nyquist sur canal radiomobile.

Alexandre Marquet

► To cite this version:

Alexandre Marquet. Communications au-delà de la cadence de Nyquist sur canal radiomobile.. Réseaux et télécommunications [cs.NI]. Université Grenoble Alpes, 2017. Français. NNT: . tel-01800392v1

HAL Id: tel-01800392

<https://theses.hal.science/tel-01800392v1>

Submitted on 25 May 2018 (v1), last revised 6 Jun 2018 (v2)

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

THÈSE

Pour obtenir le grade de

DOCTEUR DE L'UNIVERSITÉ DE GRENOBLE

Spécialité : **signal, image, parole, télécommunications (SIPT)**

Arrêté ministériel : 25 mai 2016

Présentée par

M. Alexandre Marquet

Thèse dirigée par **M. Laurent Ros**

supervisée par **M. Cyrille Siclet** et **M. Damien Roque**

préparée au sein du **Laboratoire Grenoble Images, Parole, Signal, Automatique (Gipsa-lab)**

dans l'**École Doctorale Électronique, Électrotechnique, Automatique, Traitement du Signal (EEATS)**

Communications au-delà de la cadence de Nyquist sur canal radiomobile

Thèse soutenue publiquement le **21 décembre 2017**,
devant le jury composé de :

Mme. Geneviève Baudoin

Professeur, ESIEE Paris, Présidente

Mme. Catherine Douillard

Professeur, IMT Atlantique Bretagne-Pays de la Loire, Rapporteur

M. Michel Terré

Professeur, CNAM Paris, Rapporteur

M. Jean-Marc Brossier

Professeur, Grenoble INP, Examineur

M. Romain Tajan

Maître de conférences, Bordeaux INP, Examineur

M. Laurent Ros

Maître de conférences, Grenoble INP, Directeur de thèse (pour la 3^{ème} année)

M. Cyrille Siclet

Maître de conférences, Université Grenoble Alpes, Encadrant de thèse

M. Damien Roque

Maître de conférences, ISAE-Supaéro - Toulouse, Encadrant de thèse

M. Pierre Siohan

Chercheur retraité, Orange labs Cesson-Sevigné, Invité



Remerciements

En prenant du recul sur ce travail de thèse, je ne peux m'empêcher de noter une contradiction intéressante : si le travail en lui-même relève d'un investissement majoritairement individuel, il me paraît toutefois inenvisageable de le réussir sans soutiens de première qualité. Le lecteur trouvera donc ci-dessous une liste non-exhaustive des gens sans lesquels ce travail n'existerait tout simplement pas.

Tout d'abord, je remercie tous les membres de mon jury de thèse. Mme. Geneviève Baudoin pour avoir présidé le jury. Mme. Catherine Douillard et M. Michel Terré, en tant que rapporteurs, pour leurs retours pertinents sur ce manuscrit, mais aussi M. Pierre Siohan pour les mêmes raisons. MM. Jean-Marc Brossier et Romain Tajan pour leur participation en tant qu'examinateurs. Je leur suis tous particulièrement reconnaissant pour l'échange cordial et extrêmement intéressant que nous avons eu pendant ma soutenance.

J'adresse également mes sincères remerciements à mes encadrants. Cyrille pour ses grandes qualités humaines, toujours à l'écoute, toujours de bon conseil, et pour son amour des hypothèses clairement énoncées et des équations bien posées. Damien pour son investissement passionné dans le domaine des communications numériques (dans lequel s'inscrit bien évidemment ce travail), pour le plaisir de monter des antennes sur les toits des bâtiments de l'ISAE, et pour sa bonne humeur. Laurent pour sa gentillesse et ses conseils toujours pertinents. J'adresse enfin une pensée particulière à Pierre Comon, sans qui ma thèse n'aurait pas pu commencer en premier lieu.

Merci à mes collègues du D1110. Marc, avec qui j'ai partagé les joies et les galères du doctorat du début à la fin. Guanghan, la vétérante du bureau. Pierre Maho, « Junior » (à ce qu'il paraît), le petit nouveau à qui on prend plaisir à donner de bons tuyaux. Et, bien sûr, la chercheuse de talent avec qui j'ai l'honneur de partager aujourd'hui ma vie : Quyên.

Merci aussi à tous les autres collègues, galériens de doctorat ou de post-doctorat¹. Pierre N. (ne devrait-on pas, par cohérence, l'appeler « Major » ?), le bricoleur de la bande, et Aziliz qui nous aura fait découvrir le breton (mais pas la Bretagne). Anthony, qui nous aura, lui, fait découvrir l'Alsace, et de la meilleure des manières : par ses vignobles ! Raph, l'astronome passionné, mais jamais dans les étoiles. Marion et ses excellents gâteaux ! Et, bien-sûr, Fabrice, qui n'aura jamais refusé un verre de vin, aussi infect soit-il (désolé...). Emmanuelle, à qui je ne tiendrai vraisemblablement jamais la promesse d'installer un *vrai* klaxon sur son fauteuil... Désolé aussi ! Jeanne, qui aura sans doute fini de convertir l'ensemble des ≈ 300 personnes du labo au féminisme à l'heure où j'écris ces lignes (c'est en tous cas ce que je lui souhaite !), et Louis, le toulousain acquis à la cause grenobloise.

Une mentions spéciale pour les collègues avec qui j'ai partagé l'aventure Gipsa-doc. L'honneur revenant aux « anciens » : Pascal, Céline, Cindy, Maëlle, Mael, qui m'ont convaincu de m'investir dans l'aventure. Raph (encore !), pour sa motivation sans faille, et sans qui le Gipsadoc n'aurait pas le même visage aujourd'hui. Marc (encore !), pour sa gestion exemplaire de

1. Ou les deux, mais pas en même temps !

la trésorerie. Alex. H., pour sa bonne humeur constante. Sophie, aux idées toujours fraîches. Pierre N. (encore !), Miguel et Nicolas (avec qui je partage en sus le goût pour la guitare) pour les bonnes idées et l'investissement sans faille. Merci enfin à ceux du Gipsadoc que j'ai moins eu l'occasion de côtoyer pour ma troisième année de thèse, mais dont la contribution à l'association n'en est pas moins importante : Rémi, Bharrat, Thubaut, Li Liu, Diego et Nadia.

Merci à mes amis de toujours : Alex. H. (encore !), Nicolas, Alex R., Aurore, Anna, Doubblas, Petit, Chapio, Maxime, Olivier, Alwin, Coline, MoMo, Romain, Jboon et Julien, pour avoir toujours été là. J'en profite par ailleurs pour souhaiter tous mes vœux de réussite à Alex H., Alex R. et Anna pour leur thèses en cours. Merci également aux amis d'école et d'IUT : Rémi, Geoffrey, Philippe (Ryou), François, Clément, Vincent (et Anne). Grâce à vous, je n'aurais pas complètement perdu de vue la technique. Merci aussi à « l'équipe des couples de docteurs Grenoblois » (on devrait pouvoir l'écrire comme ça, maintenant que j'ai soutenu !) : Georgia + Mircea et Guanghan (encore !) + Zhong Yang, pour le soutien et les bons plans sortie !

J'adresse également une pensée à mes anciens collègues du LAN : Matthieu, Guillaume, Nicolas, Vincent, Gilles, Cécile, Thierry et Hervé. Je remercie plus particulièrement Fabien et Anouar, à l'époque doctorants en CIFRE, qui m'ont donné envie de commencer cette thèse.

Enfin merci à toute ma famille, mon frère, mes parents et mes grands-parents pour m'avoir porté jusque-là. Merci à mes oncles et tantes, avec une mention particulière pour « Tonton Fred » qui m'aura donné le goût pour les nouvelles technologies. Enfin, je porte une reconnaissance éternelle à mes parents qui ont toujours soutenu et accompagné mes choix, et m'ont permis de faire les études qui me plaisaient le mieux.

Je terminerai par ces quelques lignes en hommage à mon grand-père, Maurice Marquet, qui m'a transmis sa curiosité, son sens de l'inventivité et du faire par soi-même, mais qui nous a malheureusement quitté cette année. Je lui dédie ce travail de thèse.

Table des matières

Notations	v
Table des sigles et acronymes	ix
Introduction	1
1 Communications numériques	7
1.1 Introduction	7
1.2 Modélisation de la chaîne de communication	10
1.3 Modélisation des canaux de transmission	13
1.4 Transmissions multiporteuses	22
1.5 Communications à haute efficacité spectrale	36
1.6 Conclusion	43
2 Transmissions au-delà de la cadence de Nyquist	45
2.1 Introduction	45
2.2 FTN monoporteuse et FTN multiporteuse	46
2.3 Limite de Mazo	48
2.4 Techniques de détection	51
2.5 Conclusion	62
3 Système multiporteuse linéaire au-delà de la cadence de Nyquist optimal en présence de BABG	65
3.1 Introduction	65
3.2 Prototypes maximisant le RSIB	66
3.3 Analyse statistique du terme d'auto-interférence	70
3.4 Performances théoriques et simulations	77
3.5 Conclusion	85
4 Turbo-égalisation à faible complexité des systèmes multiporteuses au-delà de la cadence de Nyquist	89
4.1 Introduction	89
4.2 Système MEQM à base de suppression successive d'interférence et de frames étroites en présence de BABG	90
4.3 Suppression successive d'interférence en présence de canal radiomobile	102
4.4 Conclusion	110

5	Impact de la signalisation au-delà de la cadence de Nyquist sur le PAPR des modulations multiporteuses	113
5.1	Introduction	113
5.2	Définition du PAPR	114
5.3	Conditions d'optimalité	114
5.4	Prototypes d'émission optimaux	115
5.5	Simulations	116
5.6	Conclusion	118
	Conclusion	121
A	Estimation au sens de la minimisation de l'erreur quadratique moyenne	125
	Bibliographie	134

Notations

Mathématiques, probabilité

$\mathbb{Z}$	Ensemble des entiers relatifs.
$\mathbb{R}$	Ensemble des réels.
$\mathbb{C}$	Ensemble des complexes.
j	Le nombre complexe tel que $j^2 = -1$.
z^*	Conjugué du nombre complexe $z \in \mathbb{C}$.
$ z $	Module du nombre complexe $z \in \mathbb{C}$.
$\mathcal{R}\{z\}$	Partie réelle du nombre complexe $z \in \mathbb{C}$.
$\mathcal{I}\{z\}$	Partie imaginaire du nombre complexe $z \in \mathbb{C}$.
$\mathbf{M}^T$	Matrice transposée de $\mathbf{M}$.
$\mathbf{M}^H$	Matrice hermitienne (transposée conjuguée) de $\mathbf{M}$.
$\mathcal{L}_2(\mathbb{R})$	Fonctions de carré intégrable : $f \in \mathcal{L}_2(\mathbb{R}) \Leftrightarrow \int_{-\infty}^{+\infty} f(t) ^2 dt < \infty$.
$\ell_2(\mathbb{Z})$	Fonctions de carré sommable sur $\mathbb{Z}$: $f \in \ell_2(\mathbb{Z}) \Leftrightarrow \sum_{n \in \mathbb{Z}} f[n] ^2 < \infty$.
$\ell_2(\mathbb{Z}^2)$	Fonctions de carré sommable sur $\mathbb{Z}^2$: $f \in \ell_2(\mathbb{Z}^2) \Leftrightarrow \sum_{m,n \in \mathbb{Z}} f[m,n] ^2 < \infty$.
$\ell_2(\Lambda)$	Fonctions de carré sommable sur $\Lambda \subset \mathbb{Z}^2$: $f \in \ell_2(\Lambda) \Leftrightarrow \sum_{m,n \in \Lambda} f[m,n] ^2 < \infty$.
$\langle x; y \rangle$	Produit scalaire entre x et y . Si $x, y \in \mathcal{L}_2(\mathbb{R})$: $\langle x; y \rangle = \int_{-\infty}^{+\infty} x^*(t)y(t)dt$. Si $x, y \in \ell_2(\mathbb{Z})$: $\langle x; y \rangle = \sum_{n \in \mathbb{Z}} x^*[n]y[n]$.
$(x * y)$	Produit de convolution entre x et y . Si $x, y \in \mathcal{L}_2(\mathbb{R})$: $(x * y)(t) = \int_{-\infty}^{+\infty} x(\tau)y(t - \tau)d\tau$. Si $x, y \in \ell_2(\mathbb{Z})$: $(x * y)[k] = \sum_{n \in \mathbb{Z}} x[n]y[k - n]$.
$\mathbb{P}\{X = x\}$	Probabilité pour la variable aléatoire discrète X de prendre la valeur x .
$f_X(x)$	Densité de probabilité de la variable aléatoire X évaluée en x .
$F_X(x)$	Fonction de répartition de la variable aléatoire X évaluée en x .
$Q(x)$	Fonction de répartition complémentaire d'une variable aléatoire X suivant une loi normale centrée réduite, évaluée en x .
$\mathbb{E}\{X\}$	Espérance mathématique de la variable aléatoire X .
$\gamma_s(f)$	Densité spectrale de puissance du signal s évaluée en f .
$\mathcal{F}\{x\}$	Transformée de Fourier de la fonction $x \in \mathcal{L}_2(\mathbb{R})$.
$\mathcal{F}^{-1}\{x\}$	Transformée de Fourier inverse de la fonction $x \in \mathcal{L}_2(\mathbb{R})$.
$\mathbf{I}$	Opérateur identité.

$\delta(t)$	Distribution de Dirac.
δ_n	Delta de Kronecker, tel que $\delta_n = 1$ si $n = 0$ et $\delta_n = 0$ sinon.
$ \mathcal{A} $	Nombre d'éléments de l'ensemble fini $\mathcal{A}$.

Théorie des frames

$T_{\check{\mathbf{g}}}$	Opérateur d'analyse utilisant la frame $\check{\mathbf{g}}$: soit $\check{\mathbf{g}} = \{\check{g}_{m,n}(t)\}_{(m,n) \in \mathbb{Z}^2}$ une frame de $\mathcal{L}_2(\mathbb{R})$ et $f \in \mathcal{L}_2(\mathbb{R})$, alors $T_{\check{\mathbf{g}}}f = \{\langle \check{g}_{m,n}; f \rangle\}_{(m,n) \in \mathbb{Z}^2}$.
$T_{\mathbf{g}}^*$	Opérateur de synthèse utilisant la frame $\mathbf{g}$: soit $\mathbf{g} = \{g_{m,n}(t)\}_{(m,n) \in \mathbb{Z}^2}$ une frame de $\mathcal{L}_2(\mathbb{R})$ et $\mathbf{c} = \{c_{m,n}\}_{(m,n) \in \mathbb{Z}^2} \in \ell_2(\mathbb{Z}^2)$, alors $T_{\mathbf{g}}^*\mathbf{c} = \sum_{(m,n) \in \mathbb{Z}^2} c_{m,n}g_{m,n}(t)$.
$S_{\mathbf{g}}$	Opérateur de frame utilisant la frame $\mathbf{g}$: soit $f \in \mathcal{L}_2(\mathbb{R})$, alors $S_{\mathbf{g}}f = T_{\mathbf{g}}^*T_{\mathbf{g}}f$.
$A_{g,\check{g}}$	Fonction d'interambiguïté de $g \in \mathcal{L}_2(\mathbb{R})$ et $\check{g} \in \mathcal{L}_2(\mathbb{R})$: $A_{g,\check{g}}(\nu, \tau) = \int_{-\infty}^{+\infty} g^*(t)\check{g}(t - \tau)e^{j2\pi\nu t}dt \quad \forall \nu, \tau \in \mathbb{R}$.

Communications numériques

C	Capacité de Shannon (en bits/s).
η	Efficacité spectrale (en bits/s/Hz).
E_b	Énergie d'un bit (en Joules).
E_s	Énergie d'un symbole (en Joules).
N_0	Densité spectrale monolatérale du bruit (en Joules).
D_b	Débit binaire.
R_c	Rendement de codage du codeur canal.
L_b	Longueur du message binaire à transmettre.
L_c	Longueur de la séquence binaire en sortie du codeur canal.
b_k	k -ième bit à transmettre.
$\hat{b}_k$	Estimation du k -ième bit par le récepteur.
b_k^c	k -ième bit en sortie du codeur canal.
b_k^i	k -ième bit en sortie de l'entrelaceur.
$\mathbf{b}$	Séquence binaire à transmettre (message) : $\mathbf{b} = \{b_k\}_{k \in [0 \dots L_b - 1]}$.
$\hat{\mathbf{b}}$	Estimation du message binaire par le récepteur : $\hat{\mathbf{b}} = \{\hat{b}_k\}_{k \in [0 \dots L_b - 1]}$.
$\mathbf{b}^c$	Séquence binaire en sortie du codeur canal : $\mathbf{b}^c = \{b_k^c\}_{k \in [0 \dots L_c - 1]}$.
$\mathbf{b}^i$	Séquence binaire en sortie de l'entrelaceur : $\mathbf{b}^i = \{b_k^i\}_{k \in [0 \dots L_c - 1]}$.

Notations

R	Débit symbole.
$\mathcal{A}$	Alphabet de modulation (constellation).
N_b	Nombre de bits par symbole.
M	Nombre de sous-porteuses.
K	Nombre de symboles multiporteuses.
Λ	Ensemble des indices temporels et fréquentiels : $\Lambda = [0 \dots M - 1] \times [0 \dots K - 1]$.
T_0	Espacement inter-symbole (en s).
F_0	Espacement inter-porteuse (en Hz).
g	Prototype d'émission.
$\check{g}$	Prototype de réception.
$g_{m,n}$	Fonction d'émission de la porteuse m à l'indice temporel n .
$\check{g}_{m,n}$	Fonction de réception de la porteuse m à l'indice temporel n .
$\mathbf{g}$	Famille d'émission : $\mathbf{g} = \{g_{m,n}\}_{(m,n) \in \Lambda}$.
$\check{\mathbf{g}}$	Famille de réception : $\check{\mathbf{g}} = \{\check{g}_{m,n}\}_{(m,n) \in \Lambda}$.
$c_{m,n}$	Symbole émis sur la porteuse m à l'indice temporel n .
$\tilde{c}_{m,n}$	Symbole en sortie du démodulateur pour la porteuse m à l'indice temporel n .
$\bar{c}_{m,n}$	Symbole de la porteuse m à l'indice temporel n égalisé par le récepteur.
$\hat{c}_{m,n}$	Estimation du symbole émis sur la porteuse m à l'indice temporel n par le récepteur.
$\mathbf{c}$	Séquence de symboles émis (sortie du convertisseur bit à symbole) : $\mathbf{c} = \{c_{m,n}\}_{(m,n) \in \Lambda}$.
$\tilde{\mathbf{c}}$	Séquence de symboles en sortie du démodulateur. $\tilde{\mathbf{c}} = \{\tilde{c}_{m,n}\}_{(m,n) \in \Lambda}$.
$\bar{\mathbf{c}}$	Séquence de symboles égalisés par le récepteur. $\bar{\mathbf{c}} = \{\bar{c}_{m,n}\}_{(m,n) \in \Lambda}$.
$\hat{\mathbf{c}}$	Séquence de symboles estimés par le récepteur. $\hat{\mathbf{c}} = \{\hat{c}_{m,n}\}_{(m,n) \in \Lambda}$.
$b_l(c)$	l -ième élément de la séquence binaire associée au symbole $c \in \mathcal{A}$.
$i_{m,n}$	Terme d'interférence sur la porteuse m à l'indice temporel n .
$\hat{i}_{m,n}$	Estimation du terme d'interférence sur la porteuse m à l'indice temporel n par le récepteur.
$\nu_{m,n}$	Somme du terme d'interférence et du terme de bruit filtré sur la porteuse m à l'indice temporel n : $\nu_{m,n} = i_{m,n} + \langle \check{g}_{m,n}; b \rangle$.
$\mathbf{i}$	Séquence des termes d'interférence.

$\mathbf{i}$	$\mathbf{i} = \{i_{m,n}\}_{(m,n) \in \Lambda}$.
$\hat{\mathbf{i}}$	Séquence des termes d'interférence estimés par le récepteur. $\hat{\mathbf{i}} = \{\hat{i}_{m,n}\}_{(m,n) \in \Lambda}$.
s	Signal modulé.
r	Signal à l'entrée du démodulateur.
b	Bruit ajouté par le canal.
$\mathcal{H}$	Opérateur de canal.
h	Réponse impulsionnelle évolutive du canal.
L_h	Fonction de transfert évolutive du canal.
S_h	Fonction d'étalement du canal.
H_h	Fonction caractéristique du canal.
$L(\mathbf{b})$	Séquence des logarithmes de rapport de vraisemblance des bits constituant la séquence $\mathbf{b}$.
$L_a(\mathbf{b})$	Séquence des logarithmes de rapport de vraisemblance <i>a priori</i> des bits constituant la séquence $\mathbf{b}$.
$L_{ext}(\mathbf{b})$	Séquence des logarithmes de rapport de vraisemblance extrinsèques des bits constituant la séquence $\mathbf{b}$.

Table des sigles et acronymes

ADN	Acide DésoxyriboNucléique
ADSL	<i>Asymmetric Digital Subscriber Line</i> (Ligne d'abonnée numérique asymétrique)
BABG	Bruit Additif Blanc Gaussien
BPSK	<i>Binary Phase Shift Keying</i> (Modulation par saut de phase binaire)
CP-OFDM	<i>Cyclic Prefixed OFDM</i> (OFDM avec préfixe cyclique)
DSP	<i>Digital Signal Processor</i> (Processeur de traitement du signal)
EDGE	<i>Enhanced Data rates for GSM Evolution</i> Évolution du GSM pour un meilleur débits de données.
EQM	Erreur Quadratique Moyenne
EHB	Prototype conçu pour minimiser l'Énergie Hors Bande
EUTRAN	<i>Enhanced Universal Terrestrial Radio Access Network</i> (Réseau d'accès radio terrestre universel amélioré)
FBMC	<i>FilterBank MultiCarrier</i> (Modulations multiporteuses à base de bancs de filtres)
FFT	<i>Fast Fourier Transform</i> (Transformée de Fourier rapide)
FSK	<i>Frequency Shift Keying</i> (Modulation par déplacement de fréquence)
FTN	<i>Faster-Than-Nyquist</i> (Au-delà de la cadence de Nyquist)
GSM	<i>Global System for Mobile Communications</i> (Système global pour les communications mobiles)
HEVC	<i>High Efficiency Video Coding</i> (Codage vidéo à haute efficacité)
HSPA	<i>High Speed Packet Access</i> (Accès paquet à haute vitesse)
IEP	Interférence Entre Porteuses
IES	Interférence Entre Symboles
IID	Indépendants et Identiquements Distribués
JPEG	<i>Joint Photographic Experts Group</i> (Groupe de concertation d'experts en photographies)
LDPC	<i>Low Density Parity Check codes</i> (Codes de contrôle de parité à faible densité)
LFI	<i>Linear Frequency Invariant</i> (Linéaire invariant en fréquence)
LMS	<i>Least Mean Square</i> (Moindre carrés)
LTE	<i>Long Term Evolution</i> (Évolution à long terme)
LTF	Prototype conçu pour maximiser la Localisation Temps-Fréquence

LTI	<i>Linear Time Invariant</i> (Linéaire invariant en temps)
LTV	<i>Linear Time Variant</i> (Linéaire variant en temps)
LRV	Logarithme de Rapport de Vraisemblances
MAP	Maximum <i>A Posteriori</i>
MEQM	Minimisation de l'Erreur Quadratique Moyenne
MP3	<i>MPEG-1 or 2, layer III audio</i> (MPEG-1 ou 2, couche III audio)
MPEG	<i>Moving Picture Experts Group</i> (Groupe d'experts en images mouvantes)
MV	Maximum de Vraisemblance
OFDM	<i>Orthogonal frequency-division multiplexing</i> (Multiplexage fréquentiel orthogonal)
OQAM	<i>Offset QAM</i> (QAM décalé)
PAPR	<i>Peak-to-Average Power Ratio</i> (Rapport puissance moyenne à puissance pic)
PPM	<i>Pulse-Position Modulation</i> (Modulation en position d'impulsions)
PSK	<i>Phase Shift Keying</i> (Modulation par saut de phase)
PRS	<i>Partial Response Signaling</i> (Signalisation à réponse partielle)
QAM	<i>Quadrature Amplitude Modulations</i> (Modulation d'amplitude en quadrature)
QPSK	<i>Quaternary Phase Shift Keying</i> (Modulation par saut de phase quaternaire)
RSB	Rapport signal à bruit
RSI	Rapport signal à interférence
RSIB	Rapport signal à interférence plus bruit
RIF	Réponse impulsionnelle finie
STN	<i>Slower-Than-Nyquist</i> (En deçà de la cadence de Nyquist)
SISO	<i>Soft Input-Soft Output</i> (Entrée/sortie pondérée)
SOVA	<i>Soft Output Viterbi Algorithm</i> (Algorithme de Viterbi à sortie pondérée)
TEB	Taux d'Erreur Binaire
VDSL	<i>Very-high-bit-rate Digital Subscriber Line</i> (Ligne d'abonné numérique à très haut débit)
WCP-OFDM	<i>Weighted Cyclic Prefixed OFDM</i> (OFDM avec préfixe cyclique pondéré)
WiFi	<i>Wireless Fidelity</i> (Fidélité sans fil)

Introduction

La tendance générale veut que chaque nouvelle révision d'un standard de communication amène une augmentation du débit maximal proposé. Par exemple :

- en téléphonie mobile : 1 Mbit/s (2G EDGE²), 42 Mbit/s (3G HSPA³⁺), puis 300 Mbit/s (4G EUTRAN⁴) ;
- pour les lignes d'abonné numérique : 12 Mbit/s (ADSL⁵), 24 Mbit/s (ADSL2+), 55 Mbit/s (VDSL⁶), puis 100 Mbit/s (VDSL2) ;
- pour les courants porteurs en ligne : 14 Mbit/s (HomePlug), 200 Mbit/s (HomePlug AV), puis 1 Gbit/s (HomePlug AV2).

Dans un monde de plus en plus communiquant, la ressource physique supportant ces communications toujours plus rapides tend à saturer, mettant en exergue le besoin de techniques de communication présentant une meilleure efficacité spectrale. L'efficacité spectrale est donc le nombre de bits maximal qu'un système de communication peut transmettre dans une zone de 1 s.Hz (elle est donc mesurée en bit/s/Hz). Idéalement, cette efficacité spectrale maximale doit pouvoir être accessible à des rapports signal-à-bruit aussi faibles que possible, de manière à permettre une meilleure efficacité énergétique.

Pour répondre au double problème de l'efficacité spectrale et de l'efficacité énergétique, on utilise classiquement deux techniques complémentaires. La première consiste à regrouper plusieurs bits dans un seul symbole. De cette manière, à débit symbole fixe et sans changer la bande occupée, plus on augmente le nombre de bits par symbole, plus on augmente l'efficacité spectrale⁷. Cette technique a cependant le désavantage d'augmenter la sensibilité du système au bruit à mesure que le nombre de bits par symbole augmente, et nécessite des rapports signal-à-bruit généralement éloignées des limites théoriques établies par la théorie de l'information. La deuxième technique consiste à ajouter de la redondance au message. De cette manière, il est possible en utilisant les meilleures techniques actuelles (turbo-codes, codes LDPC⁸) de communiquer à des rapports signal-à-bruit très faibles, et proches des limites théoriques. Cependant, cet ajout de redondance implique une moins bonne efficacité spectrale de la part de ces systèmes (dits « codés ») par rapport à leur équivalent non codé. En complément de ces deux techniques, les communications au-delà de la cadence de Nyquist permettent, dans certaines conditions, d'augmenter l'efficacité spectrale de la transmission sans pour autant augmenter la sensibilité au bruit.

L'idée fondamentale des transmissions au-delà de la cadence de Nyquist est de rapprocher

-
2. *Enhanced Data rates for GSM Evolution* Évolution du GSM pour un meilleur débits de données.
 3. *High Speed Packet Access* (Accès paquet à haute vitesse)
 4. *Enhanced Universal Terrestrial Radio Access Network* (Réseau d'accès radio terrestre universel amélioré)
 5. *Asymmetric Digital Subscriber Line* (Ligne d'abonnée numérique asymétrique)
 6. *Very-high-bit-rate Digital Subscriber Line* (Ligne d'abonné numérique à très haut débit)
 7. Du moins lorsque les symboles discrets sont directement utilisés comme poids des impulsions de mises en formes analogiques à transmettre, soit dans le cas des modulations dites « linéaires ».
 8. *Low Density Parity Check codes* (Codes de contrôle de parité à faible densité)

les impulsions de mise en forme, support physique des symboles, en temps (et, éventuellement, dans le cas des communications multiporteuses, en fréquence) et au-delà de la limite établie par Nyquist en 1928 [Nyq28a], limite indispensable pour garantir une communication sans interférence entre impulsions de mise en forme. La première publication ayant révélé l'intérêt des communications au-delà de la cadence de Nyquist est intitulée « Faster-Than-Nyquist signaling » et fut publiée par J. E. Mazo, en 1975 [Maz75]. En effet, ses travaux suggèrent la possibilité de rapprocher les impulsions de mise en forme au-delà de la limite de Nyquist, sans nécessairement dégrader la qualité de la transmission, à condition de ne pas dépasser une certaine limite depuis dénommée « limite de Mazo » et d'utiliser un détecteur optimal. Cependant, la nécessité de disposer d'une détection optimale a laissé cette technique peu attrayante durant de nombreuses années. Ce n'est qu'au début des années 2000 que l'on a pu assister à un regain d'activité de recherche relative aux communications au-delà de la cadence de Nyquist, sous l'impulsion d'auteurs tels qu'Anderson, Rusek et Liveris [ARÖ13] ; [RA08] ; [LG03].

« As Jack Sparrow might say, it turns out that Nyquist's limit is more of a guideline than a rule. »

Alan Gatherero.⁹

Comme nous l'avons évoqué plus haut, l'outrepassement du critère de Nyquist se paye par l'apparition d'interférence entre impulsions de mise en forme, interférence devant être traitée par le récepteur afin de garantir la probabilité d'erreur requise. Fort heureusement, en monoporteuse, ce problème est équivalent à celui de l'égalisation d'un signal perturbé par un canal sélectif en fréquence, lequel est déjà bien couvert par la littérature (une différence notable apparaît néanmoins concernant les propriétés statistiques du bruit vu par le récepteur). De nombreuses adaptations de techniques d'égalisation ont donc été proposées dans le contexte des communications monoporteuses au-delà de la cadence de Nyquist, permettant l'exploitation de l'efficacité spectrale supplémentaire permise par cette technique. L'état de l'art est en revanche bien différent du côté des modulations multiporteuses, lesquelles ont justement été popularisées grâce au fait qu'elles ne nécessitent pas de recourir à des techniques d'égalisation avancées pour compenser l'effet d'un canal sélectif en fréquence. De plus, la transposition des techniques d'égalisation utilisées en monoporteuse pour les modulations multiporteuses se révèle généralement très inefficace, algorithmiquement parlant.

Dans ce manuscrit, nous choisissons de traiter le problème de l'interférence entre impulsions de mise en forme non pas seulement du point de vue du récepteur, mais aussi de celui de l'émetteur, en cherchant des impulsions de mise en forme facilitant l'égalisation. Nous montrons que choisir de telles impulsions permet de profiter à la fois des apports des transmissions au-delà de la cadence de Nyquist sur l'efficacité spectrale, et de la capacité propre aux modulations

9. Extrait de *Running Faster than Nyquist : An Idea Whose Time May Have Come*, publié en février 2017 dans IEEE Comsoc Technology News : <http://www.comsoc.org/ctn/running-faster-nyquist-idea-whose-time-may-have-come>. La traduction est volontairement laissée à la discrétion du lecteur.

Introduction

multiporteuses à proposer un très bon compromis performance/complexité sur canal sélectif en fréquence et radiomobile.

Synthèse des contributions

Les travaux de thèse décrits dans ce manuscrit ont engendré quatre contributions principales.

1. On propose une famille de paires d'impulsions de mise en forme (en émission) et de filtres de réception permettant la maximisation du rapport signal-à-interférence-plus-bruit sur canal à bruit additif blanc gaussien.
2. On démontre l'optimalité de cette famille en terme de minimisation de l'erreur quadratique moyenne, dans un contexte de turbo-égalisation à suppression souple de l'interférence, toujours sur canal à bruit additif blanc gaussien.
3. On démontre l'optimalité de cette famille sur le facteur de crête du signal modulé et, dans ce cas de figure, observons que ce dernier diminue à mesure que la densité du système augmente.
4. On adapte le système de turbo-égalisation à suppression souple de l'interférence pour fonctionner sur canal sélectif en fréquence et radiomobile, et l'on montre que les méthodes pour choisir les paramètres du système multiporteuse en fonction de la sélectivité temps-fréquence du canal de transmission est inchangée par rapport aux systèmes respectant le critère de Nyquist.

Publications

Journaux

1. Alexandre Marquet, Cyrille Siclet, Damien Roque, Pierre Siohan, « Analysis of a FTN Multicarrier System : Interference Mitigation Based on Tight Gabor Frames ». In : *EAI Endorsed Transactions on Cognitive Communications*, vol. 17, No 10 (2017), DOI : 10.4108/eai.23-2-2017.152191.
2. Alexandre Marquet, Cyrille Siclet, Damien Roque « Analysis of the faster-than-Nyquist optimal linear multicarrier system ». In : *Comptes Rendus Physique*, vol. 18, issue 2 (février 2017), pp. 167-177, DOI : 10.1016/j.crhy.2016.11.006.
3. Alexandre Marquet, Damien Roque, Cyrille Siclet, Pierre Siohan, « FTN multicarrier transmission based on tight Gabor frames ». In : *EURASIP Journal on Wireless Communications and Networking*, (2017), p. 97, DOI : 10.1186/s13638-017-0878-3.

Conférences

1. Cyrille Siclet, Damien Roque, Alexandre Marquet, Laurent Ros « Improving spectral efficiency while reducing PAPR using faster-than-Nyquist multicarrier signaling ». In :

Communication Technologies for Vehicles : 12th International Workshop, Nets4Cars / Nets4Trains / Nets4Aircraft 2017 proceedings, pp. 21-32 (Mai 2017), Toulouse, France. DOI : https://doi.org/10.1007/978-3-319-56880-5_3.

2. Albert Abelló, Damien Roque, Cyrille Siclet, Alexandre Marquet « On zero-forcing equalization for short-filtered multicarrier faster-than-Nyquist signaling ». In : *2016 50th Asilomar Conference on Signals, Systems and Computers*, pp. 904-908, (novembre 2016), Pacific Grove, États-unis, DOI : 10.1109/ACSSC.2016.7869180.
3. Alexandre Marquet, Cyrille Siclet, Damien Roque, Pierre Siohan, « Analysis of a multicarrier communication system based on overcomplete Gabor frames ». In : *Cognitive Radio Oriented Wireless Networks*, Springer International Publishing, p. 387 (2016), Lecture Notes of the Institute for Computer Sciences, Social Informatics and Telecommunications Engineering, DOI : 10.1007/978-3-319-40352-6_32.
4. Alexandre Marquet, Cyrille Siclet, Damien Roque « Analysis of the optimal linear system for multicarrier FTN communications ». In : *Actes des journées scientifiques 2016 de l'URSI-France Energie et radiosciences*, (mars 2016), Cesson-Sévigné, France.

Structure du document

Le document est constitué de 5 chapitres, dont les 3 derniers présentent nos contributions. Le chapitre 1 est tout d'abord dédié à l'introduction aux éléments fondamentaux de la discipline du traitement du signal pour les communications numériques. Nous précisons ensuite le périmètre de l'étude, puis détaillons les différents modèles et outils mathématiques qui seront utilisés dans la suite du manuscrit : modélisation déterministe des canaux de transmission, principes et modèles pour les transmissions multiporteuses, éléments d'analyse temps-fréquence. Le chapitre se termine sur la présentation de plusieurs méthodes classiques de transmission à haute efficacité spectrale, et l'introduction des techniques au-delà de la cadence de Nyquist permettant, sous certaines conditions, d'augmenter l'efficacité spectrale sans perte en terme de probabilité d'erreur binaire.

L'état de l'art des techniques de transmission au-delà de la cadence de Nyquist est l'objet du chapitre 2. Après un bref rappel historique, on y détaille un modèle des relations d'entrée/sortie pour les systèmes monoporteuses et multiporteuses. Ces modèles montrent que ce type de système induit nécessairement la présence d'interférence entre impulsions de mise en forme. Vient ensuite la présentation de la limite de Mazo et de son extension pour les modulations multiporteuses, quantifiant la densité de signalisation maximale atteignable avant l'effondrement des performances en terme de probabilité d'erreur binaire. À partir de ces bases théoriques, on détaille tout d'abord les techniques de détection optimales, puis sous-optimales, proposées à ce jour dans la littérature, avec une emphase sur la complexité, trop élevée pour permettre une quelconque application pratique, des techniques de détection optimales pour les modulations multiporteuses.

Dans le chapitre 3, nous donnons les conditions permettant d'obtenir des prototypes maximisant le rapport-signal-à-bruit ou le rapport signal-à-interférence. En corollaire, nous proposons

Introduction

des prototypes d'émission (ou impulsions de mise en forme) et de réception maximisant le rapport signal-à-bruit-plus-interférence. Nous vérifions ensuite de manière expérimentale les propriétés statistiques de l'interférence, et montrons dans quelles conditions une approximation de cette dernière par une loi normale est réaliste. Les performances du système multiporteuse non codé au-delà de la cadence de Nyquist utilisant ces prototypes sont enfin évaluées par le biais de simulations. On y observe des taux d'erreurs binaires trop élevés pour susciter des applications pratiques, appelant à des techniques de traitement de l'interférence plus sophistiquées.

Afin de permettre de meilleurs taux d'erreur binaire, une compensation plus efficace de l'interférence induite par les modulations multiporteuses, et une complexité algorithmique raisonnable, nous proposons dans le chapitre 4 un système multiporteuse au-delà de la cadence de Nyquist turbo-égalisé par suppression souple de l'interférence. On montre que ce système, lorsqu'utilisé avec les prototypes optimaux en termes de rapport signal-à-interférence-plus-bruit proposés dans le chapitre 3, permet de minimiser l'erreur quadratique moyenne en sortie de l'anneau d'interférence. Les simulations sur canal à bruit additif blanc gaussien montrent que ce système permet effectivement un gain significatif en efficacité spectrale par rapport au système codé original. Dans un second temps, nous proposons une extension de ce système pour permettre son usage sur canal sélectif en fréquence et/ou en temps. Les simulations montrent là encore un gain en efficacité spectrale, mais aussi que les choix des paramètres de l'émetteur/récepteur multiporteuse (tels que le nombre de porteuses et le choix de l'impulsion de mise en forme) se font selon les mêmes critères, et ont le même impact que pour les systèmes en deçà de la cadence de Nyquist.

Enfin, le chapitre 5 montre que le facteur de crête des signaux issus des modulateurs multiporteuses au-delà de la cadence de Nyquist est minimal lorsque ces derniers font usage des impulsions de mise en forme proposées dans le chapitre 3. Dans ce cas, ce facteur de crête est d'autant plus faible que la densité du système augmente.

Communications numériques

Sommaire

1.1	Introduction	7
1.2	Modélisation de la chaîne de communication	10
1.3	Modélisation des canaux de transmission	13
1.3.1	Canaux sélectifs en fréquence	16
1.3.2	Canaux sélectifs en temps	17
1.3.3	Canaux doublement sélectifs	19
1.4	Transmissions multiporteuses	22
1.4.1	Frames et familles de Gabor	24
1.4.2	Chaîne de communication multiporteuse	32
1.5	Communications à haute efficacité spectrale	36
1.5.1	Limites théoriques	37
1.5.2	Mises en œuvres pratiques	38
1.6	Conclusion	43

1.1 Introduction

Avant d'aborder l'aspect numérique, il convient de définir la notion de communication. Étymologiquement parlant, le verbe communiquer nous vient du latin *communicare* [Aca17] formé de *cum* (« avec » que l'on retrouve par exemple dans « commun ») et *munio* (« fonction » que l'on retrouve par exemple dans « munir ») [Gaf34] et revêt deux principaux sens dans son usage moderne :

- « transmettre », sous-entendu, de manière unilatérale (une communication dans un colloque, par exemple) ;
- « être, se mettre en relation » ou « être relié par un passage », ce qui qualifie une transmission bilatérale (une communication entre deux pièces, par exemple).

Dans les deux cas, la communication se rapporte à un objet, que nous appellerons par la suite « message », et implique plusieurs entités chargées d'émettre et/ou de recevoir ledit message. Notons enfin que le message ne fait que porter l'« information », notion plus abstraite, que nous n'essaierons pas de définir ici.

L'adjectif « numérique », que nous apposerons tout au long de ce manuscrit à celui de « communication », qualifie justement ledit message. Il s'agit d'un dérivé du mot latin *numerus*,

qui signifie « nombre » [Aca17]. À partir de là, il serait tentant de définir un message numérique comme étant constitué d'une suite de valeurs numériques. Une telle définition posséderait cependant au moins deux désavantages :

- elle impose que chaque élément du message soit un chiffre ;
- il n'y a pas forcément de limite au nombre de valeurs que peut prendre une valeur numérique (typiquement, si elle est prise dans l'ensemble des entiers naturels), or en disposant d'un nombre infini de valeurs, il devient possible de représenter n'importe quelle information à l'aide d'un seul élément, éliminant dans de nombreux cas le besoin d'une suite d'éléments pour représenter une information particulière.

Ainsi, de la même manière que [Ber+07] ; [JG07] ; [PS08], nous proposons de définir un message numérique comme une suite d'éléments prenant valeurs dans un ensemble dénombrable appelé « alphabet », supposé de taille finie. Ainsi munis de cette définition, nous pouvons montrer que les messages numériques ne sont pas l'exclusivité de l'informatique et de ses célèbres « 0 » et « 1 » à travers quelques exemples :

- l'acide désoxyribonucléique (ADN) utilise un alphabet de quatre éléments (ici quatre molécules : l'adénine, la cytosine, la guanine et la thymine), dont la suite le long des brins de l'ADN porte un message se rapportant à de l'information génétique (voir figure 1.1a) ;
- l'écriture utilise un alphabet (plus ou moins vaste selon les langues et la typographie) permettant de former des messages se rapportant à l'information pensée par le rédacteur (voir figure 1.1b) ;
- le code Morse utilise un alphabet composé de trois éléments usuellement représentés par un tiret, un point et un espace, permettant lui aussi de former des messages se rapportant à l'information pensée par le rédacteur (voir figure 1.1c).

En revanche, puisque l'on considère des alphabets de taille finie, il est possible d'associer de manière unique un élément à un nombre. Cela signifie qu'il est possible d'associer n'importe quel alphabet fini à un alphabet numérique, ce qui nous permettra de le modéliser plus aisément à l'aide du formalisme mathématique.



FIGURE 1.1 – Exemples de messages numériques basés sur différents alphabets.

1.1. Introduction

On retrouve donc de nombreuses manifestations, naturelles ou non, de messages numériques. Pourtant, le monde tel que nous le percevons est fondamentalement analogique. Dit autrement, les grandeurs physiques mesurables évoluent généralement de manière continue, et peuvent prendre une infinité de valeurs elles aussi continues. L'explication de la multiplication des outils de traitements numériques est à trouver ailleurs. En particulier, les représentations numériques permettent des traitements algorithmiques, qui comportent plusieurs avantages :

- flexibilité : il est possible d'effectuer plusieurs traitements différents avec le même matériel (microprocesseur, microcontrôleur, etc.) ;
- stockage uniformisé : en exprimant tout message sous forme numérique, on peut stocker indifféremment tout type de message (image, son, texte, etc.) sur un même support (disque dur, mémoire « flash », bande magnétique, etc.) ;
- miniaturisation : à quelques exceptions près, pour une fonction donnée, la puissance de calcul ramenée au volume des unités de calculs numériques permet une implémentation moins volumineuse en numérique qu'en analogique.

D'autre part, le célèbre théorème d'échantillonnage de Shanon–Nyquist [Sha49] ; [Nyq28a] ainsi que d'autres travaux plus récents dans le domaine de l'acquisition comprimée [CRT06] ; [Don06] nous indiquent qu'il est possible, sous certaines conditions, de passer d'une représentation temporelle continue à une représentation discrète sans perte d'information. Par ailleurs, si le passage au discret des grandeurs mesurées entraîne pour sa part un ajout de bruit (bruit de quantification), on pourra en revanche le considérer négligeable à partir du moment où l'alphabet utilisé est de taille suffisamment importante. Cela explique pourquoi à l'heure actuelle, on préférera convertir, stocker, traiter, mais aussi communiquer la plupart des signaux physiques mesurables de manière numérique.

Ce chapitre a pour objectif d'introduire les éléments généraux de communication numérique qui nous seront utiles pour développer les contributions de cette thèse. Nous introduirons tout d'abord le modèle de chaîne de communication qui sera utilisé dans la suite du manuscrit. Viendra ensuite une partie plus détaillée dédiée à la modélisation d'un bloc particulier de cette chaîne : le canal de transmission. Cette partie nous permettra de mettre en évidence la pertinence des émetteurs récepteurs « multiporteuses » dans le cadre des communications radiomobiles. Ainsi, la quatrième partie sera dédiée aux outils d'analyse et à la modélisation de ce type de systèmes. Enfin, nous introduirons la technique de transmission au-delà de la cadence de Nyquist, et la positionnerons par rapport à d'autres techniques classiques d'augmentation de l'efficacité spectrale. Ce dernier point servira de transition pour le second chapitre « Transmissions au-delà de la cadence de Nyquist », dont l'objectif sera de faire un tour d'horizon de l'état de l'art sur ce sujet.

1. Travail dérivé de : MesserWoland – DNA_structure_and_bases_PL.svg, CC BY-SA 3.0, <https://commons.wikimedia.org/w/index.php?curid=5323405>

1.2 Modélisation de la chaîne de communication

On se restreint désormais au cas où l'on souhaite transmettre un message numérique entre un seul émetteur et un seul récepteur au travers d'un « canal », qui représente un milieu plus ou moins hostile à la transmission du message (voir figure 1.2). Le rôle de l'émetteur est donc de « mettre en forme » le message, modélisé sous la forme d'un vecteur d'éléments binaires $\mathbf{m}$, par un signal $s(t)$ nécessairement analogique, puisqu'adapté au canal de transmission, analogique par nature. Le rôle du récepteur est quant à lui de produire une estimation $\hat{\mathbf{m}}$ du message émis, en se basant sur le signal disponible $r(t)$ en sortie du canal. Entrons désormais plus en détails sur ces trois éléments (émetteur, récepteur, canal) fondamentaux d'une chaîne de communication numérique.

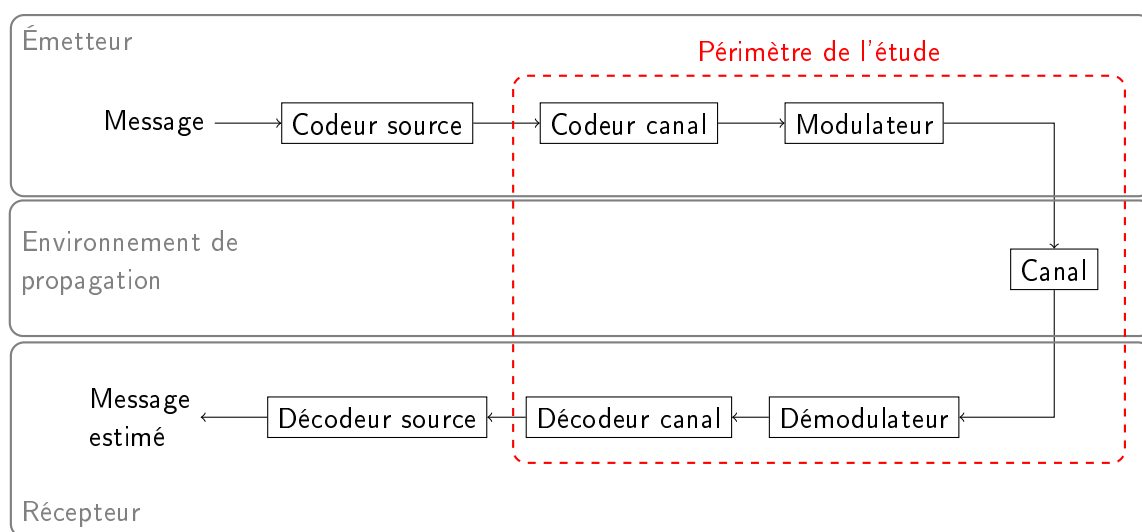


FIGURE 1.2 – Synoptique du modèle de chaîne de communication.

On distingue usuellement trois principaux composants dans l'émetteur, ainsi que leurs deux dans le récepteur : le codeur source, le codeur de canal, et le modulateur (respectivement décodeur source, décodeur canal, et démodulateur en réception).

Premièrement, le codage source a pour objectif de supprimer :

- la redondance présente dans le message à communiquer ;
- éventuellement, une partie du message (considérée inutile pour l'application visée).

Certains algorithmes de codage source, par exemple basés sur du codage entropique ou de la compression par dictionnaire, sont généralement indépendants du type de données (texte, image, etc.). On citera le codage de Huffman [Huf52] et l'algorithme LZ77 [ZL77], utilisés par exemple dans l'algorithme de compression par défaut du format d'archivage « zip » (algorithme Deflate [Deu96]). D'autres en revanche sont propres à un certain type de message et en suppriment la redondance intrinsèque. Par exemple, en compression vidéo, la technique de compensation du mouvement permet d'exprimer les pixels d'une image comme une translation de ceux présents

1.2. Modélisation de la chaîne de communication

dans l'image suivante ou précédente. Il est aussi possible d'effacer une partie du message qui sera soit imperceptible par l'utilisateur final, soit considérée comme non-indispensable à l'application visée. Ainsi, en compression d'image par exemple, la technique de sous-échantillonnage de la chrominance est basée sur l'observation que l'observateur humain est moins sensible à la précision de la chrominance qu'à celle de la luminance. Également, en téléphonie, on ne conserve généralement que la bande de fréquence entre 300 et 3400 Hz du signal capté par le microphone du téléphone [ITU09]. On notera qu'en pratique, les algorithmes de codages sources spécialisés à certains types de données (vidéo, son, etc.) combinent suppression de la redondance intrinsèque, suppression de données, et compression par codage entropique ou par dictionnaire (c'est par exemple le cas des formats MP3² [ISO93], JPEG³ [ITU92] et HEVC⁴ [ITU16]). Cet élément de la chaîne de transmission ne sera pas abordé plus en profondeur dans le cadre de ce manuscrit.

Vient ensuite le codage canal, dont le rôle est d'ajouter une redondance maîtrisée (connue de l'émetteur et du récepteur), et indépendante du type de message à transmettre. Cette redondance a pour but de contrecarrer l'effet du canal, de manière à produire, en réception, une estimée plus fiable du message effectivement émis. En particulier, les travaux de Shannon sur la théorie de l'information [Sha48] nous indiquent qu'à conditions de canal fixées, il est possible de transmettre sans aucune perte d'information, à condition :

- que le débit d'information binaire D_b de la transmission ne dépasse pas la capacité C du canal de transmission pour ces conditions ;
- d'utiliser de « bons codes ».

L'article de Shannon ne dit pas comment trouver de bons codes dans le cas général. En revanche, sur canal à bruit additif blanc gaussien, les turbo-codes et les codes LDPC sont connus pour permettre d'approcher à moins d'1 dB la capacité de Shannon [Ber+07, chap. 7 et 9]. Nous explorerons plus en détail cet élément de la chaîne de transmission dans la section 1.5.2.

Enfin, le modulateur est l'interface avec le monde physique (le canal). Son rôle est principalement d'adapter le contenu binaire à son entrée au canal de transmission. Par exemple, pour une communication sur fibre optique, il faudra traduire les bits en impulsions lumineuses. Avant toute chose, précisons que nous nous intéresserons ici aux signaux sur onde porteuse, c'est-à-dire centrés en fréquence autour d'une fréquence porteuse f_0 . On peut donc écrire $s_{\text{RF}}(t) \in \mathcal{L}_2(\mathbb{R})$, le signal réel sur onde porteuse, sous la forme :

$$s_{\text{RF}}(t) = A(t) \cos(2\pi f_0 t + \theta(t)), \quad (1.1)$$

avec $A(t) \in \mathbb{R}_+$ son amplitude et $\theta(t) \in [-\pi; \pi]$ sa phase. On remarque que l'on peut agir sur deux paramètres indépendants et à valeurs dans des intervalles continus : l'amplitude et la phase du signal, tandis que l'on peut seulement faire varier l'amplitude et le signe (qui ne peut prendre que deux valeurs) d'un signal en bande de base ($f_0 = 0$). Ces deux paramètres peuvent être regroupés au sein d'une même fonction complexe $s(t)$, permettant de réécrire (1.1) par :

$$s_{\text{RF}}(t) = |s(t)| \cos(2\pi f_0 t + \arg\{s(t)\}) = \mathcal{R} \left\{ s(t) e^{j2\pi f_0 t} \right\}. \quad (1.2)$$

2. *MPEG-1 or 2, layer III audio* (MPEG-1 ou 2, couche III audio)

3. *Joint Photographic Experts Group* (Groupe de concertation d'experts en photographies)

4. *High Efficiency Video Coding* (Codage vidéo à haute efficacité)

Enfin, si la condition de bande étroite, stipulant que la fréquence porteuse f_0 doit être grande devant la bande B occupée par $s_{\text{RF}}(t)$ ($f_0 \gg B$), est respectée, alors on peut travailler directement, et sans perte de généralité, sur $s(t)$ qui est alors le signal équivalent complexe en bande de base de $s_{\text{RF}}(t)$. Dans la suite de ce manuscrit, afin de simplifier le modèle, et sauf mention contraire, nous travaillerons exclusivement sur des représentations équivalentes complexes en bande de base (pour le signal émis, reçus, et pour le canal) ce qui signifie que la condition de bande étroite sera toujours supposée respectée. Plus de détails sur cette représentation peuvent être trouvés dans [JG07, Chapitre 4][Roq12, Chapitre 2].

Il existe de nombreuses familles de modulation selon qu'elles soient avec ou sans mémoire, linéaires ou non, et dont on pourra trouver une bonne description dans [PS08, Chapitre 3]. Dans ce manuscrit, nous étudierons des modulations *linéaires*, pour leur capacité à atteindre de hautes efficacités spectrales⁵, et *sans mémoire*, car ces dernières possèdent une complexité de démodulation plus faible. Pour les modulations sans mémoire, le modulateur associera un signal pris parmi M_s à un ensemble de N_b bits ($M_s = 2^{N_b}$), et transmettra ces signaux en série sur le canal de transmission, espacés d'un intervalle T_0 . Cela signifie que l'on peut exprimer l'équivalent complexe en bande de base du signal en sortie du démodulateur par

$$s(t) = \sum_{k \in \Lambda} u_k(t - kT_0), \quad t \in \mathbb{R}, \quad (1.3)$$

où :

- $\Lambda = [0; K - 1]$, avec K le nombre de signaux à transmettre (directement lié à la taille du message binaire à transmettre) ;
- $u_k(t) \in \mathcal{S}$ et $\mathcal{S} = \{s_0(t) \dots s_{M_s-1}(t)\} \subset \mathcal{L}_2(\mathbb{R})$ est l'espace des signaux.

L'espace des signaux engendré par les modulations linéaires est formé par des variations de l'amplitude complexe d'une impulsion de mise en forme notée $g(t)$. Les éléments de $\mathcal{S}$ sont alors de la forme :

$$s_m(t) = A_m g(t), \quad m \in [0; M_s - 1], \quad (1.4)$$

où $\mathcal{A} = \{A_0 \dots A_{M_s-1}\} \subset \mathbb{C}$. $\mathcal{A}$ est alors nommé alphabet de modulation ou constellation, et ses éléments sont appelés « symboles ». Les différents types de modulation linéaires sont généralement nommés en fonction de l'alphabet de modulation utilisé. Ainsi, comme l'illustre la figure 1.3 dans le plan complexe, les modulations d'amplitude en quadrature (QAM⁶ en utilisant l'acronyme anglo-saxon) utilisent une constellation dont les valeurs sont représentées dans le plan complexe par des points répartis sur une grille, tandis que pour les modulations par saut de phase (PSK⁷ en anglais) ces points sont répartis sur un cercle.

En réception, on retrouve le dual des trois fonctions de l'émetteur : le démodulateur permet le passage du monde physique au monde numérique ; le décodeur canal est chargé de retrouver les bits reçus par le codeur canal, sachant la sortie du démodulateur et le code canal ; enfin le décodeur source a pour rôle de retrouver le message émis sachant la sortie du décodeur canal et le code source.

5. L'efficacité spectrale représente le débit de la transmission, divisé par la bande occupée par le signal modulé.

6. *Quadrature Amplitude Modulations* (Modulation d'amplitude en quadrature)

7. *Phase Shift Keying* (Modulation par saut de phase)

1.3. Modélisation des canaux de transmission

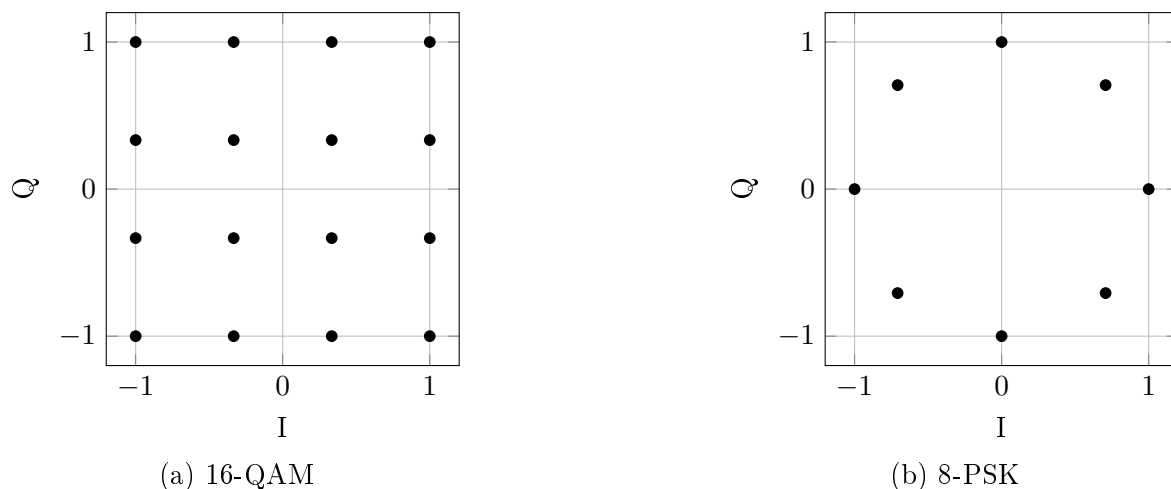


FIGURE 1.3 – Représentations de constellations dans le plan complexe. « I » (pour *in-phase*) désigne la voie en phase, (axe des réels) et « Q » désigne la voie en quadrature (axe des imaginaires).

1.3 Modélisation des canaux de transmission

Le signal engendré par le modulateur subit différents types de perturbations liées à la nature physique du canal de transmission. En communications numériques, on retrouve principalement trois types de perturbations :

- les déformations linéaires ;
- les déformations non-linéaires ;
- le bruit additif.

La déformation linéaire la plus simple est le gain du canal. Ce dernier est généralement lié à la distance entre l'émetteur et le récepteur, et à la présence d'amplificateurs entre l'émetteur et le récepteur. Viennent ensuite les déformations qui sont fonctions de la fréquence. Par exemple, pour une propagation en espace libre, le gain du canal G est lié à la fois à la distance d entre l'émetteur et le récepteur, et à la longueur d'onde $\lambda = c/f$ (avec c la célérité de l'onde dans le milieu considéré et f la fréquence) par :

$$G = \left(\frac{\lambda}{4\pi d} \right)^2 = \left(\frac{c}{2\pi} \right)^2 \times \underbrace{\frac{1}{d^2}}_{\text{Gain lié à la distance}} \times \underbrace{\frac{1}{f^2}}_{\text{Gain lié à la fréquence}}. \quad (1.5)$$

La situation est plus complexe en présence de propagation à trajets multiples, et/ou en présence d'effet Doppler (voir figure 1.4). Néanmoins, ces phénomènes constituent eux-aussi une transformation linéaire du signal, comme nous le verrons dans les parties 1.3.1 et 1.3.2.

En communication radiofréquence, les non-linéarités sont la plupart du temps engendrées par les amplificateurs. En effet, quelle que soit la technologie utilisée, on imagine volontiers qu'un

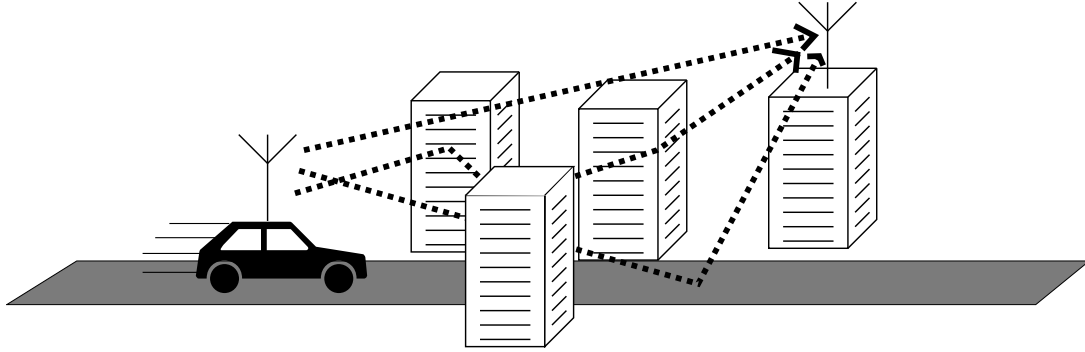


FIGURE 1.4 – Illustration d’une communication radiomobile : propagation multitrajet (plusieurs chemins de propagation entre l’émetteur et le récepteur) et effet Doppler (l’émetteur et/ou le récepteur sont mobiles).

amplificateur alimenté par une puissance P_{alim} , et avec lequel on souhaite amplifier un signal de puissance P_{in} ne pourra pas délivrer plus que $P_{\text{out}} = P_{\text{alim}} + P_{\text{in}}$. Ainsi, comme présenté en figure 1.5, un amplificateur ne présentera un comportement linéaire (ou plutôt, approché comme tel) que pour un certain intervalle de puissances d’entrées. Pour les applications radiomobiles qui nous intéressent, les non-linéarités sont principalement le fait des amplificateurs. Dans ce manuscrit, nous considérerons que ces derniers sont utilisés dans leurs zones linéaires, et nous considérerons donc des canaux exempts de non-linéarités. Cela posera alors le problème du PAPR⁸, qui doit rester faible pour atteindre de bons rendements énergétiques et que nous traiterons dans le chapitre 5. Notons que dans le cas des communications sur fibre optique, il existe de nombreuses non-linéarités liées directement à l’interaction entre la lumière et le medium, et qu’il est par conséquent délicat de ne pas les prendre en compte dans l’idée d’en tirer le meilleur parti [Tou05].

Vient enfin le bruit additif, qui peut être vu comme le terme regroupant tous les signaux reçus indésirables, tels que les signaux parasites émis par d’autres dispositifs de communication et les bruits inhérents au récepteur. En communications radiomobiles, de tous ces bruits, ce sont le bruit thermique (ou bruit de Johnson–Nyquist) et le bruit de grenaille qui sont prépondérants [Nyq28b] ; [Yea11]. Le premier est lié au mouvement de porteurs de charge (typiquement, des électrons) dans un milieu conducteur (typiquement, une résistance électrique). On peut montrer qu’il est blanc, et que sa puissance est $P_b = kTB$ avec :

- k la constante de Boltzmann (en J.K^{-1}) ;
- T la température du milieu conducteur (en K) ;
- B la bande de fréquence considérée (en Hz).

De plus, étant la somme de plusieurs processus aléatoires indépendants (le mouvement de chaque électron dans la résistance électrique), et en vertu du théorème central limite, le bruit thermique est modélisé par une variable aléatoire gaussienne. Le bruit de grenaille est quant à lui lié à la nature granulaire des porteurs de charges, et est fonction de l’intensité du courant. Il est

8. *Peak-to-Average Power Ratio* (Rapport puissance moyenne à puissance pic)

1.3. Modélisation des canaux de transmission

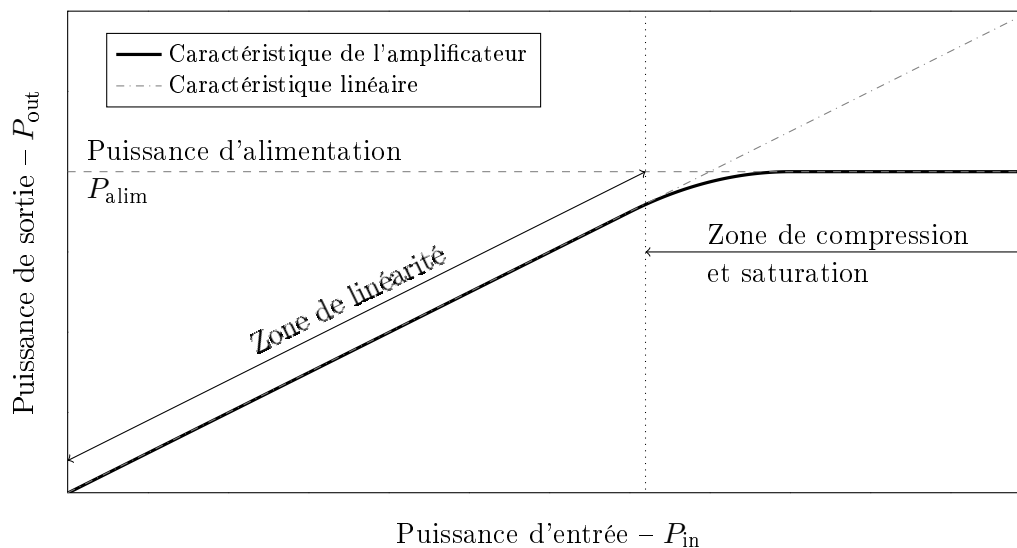


FIGURE 1.5 – Modèle simplifié d'un amplificateur.

également blanc mais suit en revanche une loi de Poisson. Cependant, pour un grand nombre de charges, cette loi poissonnienne peut-être approximée par une loi gaussienne. En pratique, les composants électroniques voient leur bruit caractérisé par un « facteur de bruit ». Connaissant le facteur de bruit des composants utilisés dans un récepteur, on peut calculer sa « température équivalente de bruit » T_{eq} . Cette dernière permet alors de calculer la puissance de bruit générée par le récepteur via la formule $P_b = kT_{eq}B$, comme s'il n'était qu'en présence de bruit thermique à une température T_{eq} . Enfin, et bien que cela ne concerne pas les communications radiomobiles (utilisant généralement des fréquences comprises entre le MHz et le GHz), notons que ces modélisations des bruits thermiques et de grenaille ne valent que pour des fréquences ne dépassant pas le terahertz (au-delà, les modèles de la physique dite « classique » ne suffisent plus à décrire correctement le phénomène).

Ainsi, dans la suite de ce document, nous nous intéresserons uniquement aux canaux linéaires à bruit additif blanc gaussien. Le modèle équivalent en bande de base du signal reçu s'écrit donc sous la forme suivante :

$$r(t) = \mathcal{H}s(t) + b(t), \quad (1.6)$$

où :

- $\mathcal{H}$ est l'opérateur de canal ($\mathcal{H} : \mathcal{L}_2(\mathbb{R}) \rightarrow \mathcal{L}_2(\mathbb{R})$) ;
- $b(t)$ est un bruit additif blanc gaussien de densité spectrale de puissance bilatérale $\gamma_b(f) = 2N_0 \forall f \in \mathbb{R}$, avec $N_0 = kT_{eq}$.

Enfin, précisons que l'opérateur de canal $\mathcal{H}$ sera ici supposé déterministe. Des modélisations stochastiques de ce dernier existent, mais ne seront par présentées dans ce manuscrit [Bel63].

1.3.1 Canaux sélectifs en fréquence

Un canal est dit sélectif en fréquence lorsqu'il présente un gain ne dépendant que de la fréquence du signal présent à son entrée. Par dualité temps-fréquence, on dira qu'il est dispersif en temps. Cela signifie que s'il est excité par une impulsion, il produira à sa sortie une ou plusieurs répliques atténuées, déphasées et retardées de cette impulsion. Avec ce type de canal, la sélectivité fréquentielle (ou de manière équivalente, la dispersivité temporelle) est toujours la même, quel que soit l'instant. Il est ainsi dit « linéaire invariant en temps » (LTI⁹ pour « linear time invariant »).

Physiquement, on trouve deux principales occurrences de ce type de canal. La première est la propagation à trajets multiples. Dans ce scénario, le récepteur reçoit autant de copies que de trajets entre l'émetteur et le récepteur (chaque trajet ayant un gain qui lui est propre). La deuxième est la présence de dispositifs de filtrage linéaire, intentionnels ou non, sur le trajet de propagation. C'est le cas par exemple lors d'une transmission sur une ligne téléphonique, ou à travers une grille métallique.

Mathématiquement, l'opérateur de canal prend la forme d'une convolution temporelle :

$$\begin{aligned} \mathcal{H} : \mathcal{L}_2(\mathbb{R}) &\rightarrow \mathcal{L}_2(\mathbb{R}) \\ s(t) &\mapsto r(t) = (s * h)(t). \end{aligned} \quad (1.7)$$

Par conséquent, dans le domaine fréquentiel, l'opérateur de canal s'écrit comme une simple multiplication :

$$\begin{aligned} \mathcal{H} : \mathcal{L}_2(\mathbb{R}) &\rightarrow \mathcal{L}_2(\mathbb{R}) \\ S(f) &\mapsto R(f) = S(f) \cdot H(f), \end{aligned} \quad (1.8)$$

où $H(f)$, $S(f)$ et $R(f)$, sont les transformées de Fourier de $h(t)$, $s(t)$ et $r(t)$, respectivement. $h(t)$ est alors la « réponse impulsionnelle » du canal, tandis que $H(f)$ est la « fonction de transfert » ou « réponse fréquentielle » du canal (voir figure 1.6).

Nous nous intéressons désormais à la structure propre de $\mathcal{H}$. Comme nous l'avons vu dans l'équation (1.8), l'opérateur de canal s'exprime comme une simple multiplication dans le domaine de Fourier. Cela signifie que les exponentielles complexes sont fonctions propres de $\mathcal{H}$:

$$\mathcal{H}e^{j2\pi f_0 t} = \int_{-\infty}^{+\infty} h(\tau) e^{j2\pi f_0(t-\tau)} d\tau = e^{j2\pi f_0 t} \mathcal{F}\{h\}(f_0) = H(f_0) e^{j2\pi f_0 t}, \quad t, f_0 \in \mathbb{R}, \quad (1.9)$$

et $H(f_0)$ sont les valeurs propres associées. En termes de systèmes de communication, cela signifie qu'une manière de s'adapter au canal est de répartir l'information sur des exponentielles complexes, par exemple en construisant un signal d'émission de la forme suivante :

$$s(t) = \sum_{n \in \mathbb{Z}} c_n e^{j2\pi n F_0 t}, \quad t \in \mathbb{R}, \quad (1.10)$$

avec $c_n \in \mathcal{A} \forall n \in \mathbb{Z}$, et $F_0 \in \mathbb{R}^+$. L'inconvénient majeur d'une telle modulation est que chaque symbole doit être envoyé sur une fréquence pure, ce qui est infaisable en pratique. Pour contourner cet inconvénient, on peut utiliser des fonctions propres approchées. En particulier, d'après [KM97], si :

9. *Linear Time Invariant* (Linéaire invariant en temps)

1.3. Modélisation des canaux de transmission

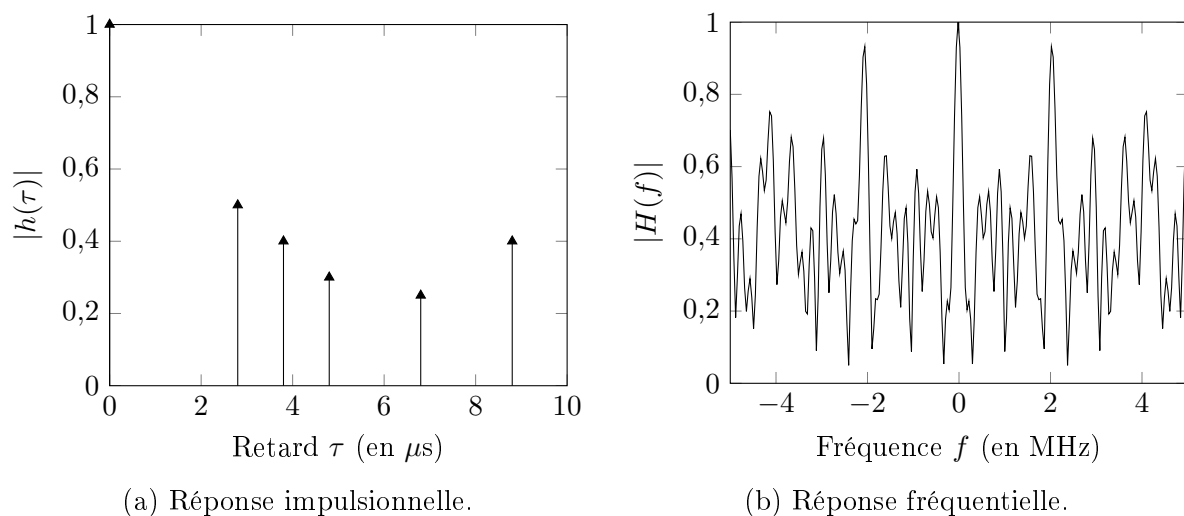


FIGURE 1.6 – Exemple de canal LTI avec 6 chemins de propagation.

- la dispersion temporelle du canal est bornée : $\forall \tau \notin [-\tau_0, \tau_0], h(\tau) = 0, t_0 \in \mathbb{R}^*$;
- le gain maximal du canal est unitaire ;
- $g(t) \in \mathcal{L}_2(\mathbb{R})$ est un filtre passe-bas de norme unitaire et à support fréquentiel borné en $[-\frac{\epsilon}{2\pi\tau_0}, \frac{\epsilon}{2\pi\tau_0}]$, $\epsilon \in \mathbb{R}^+$;

alors toute version de $g(t)$ translatée fréquentiellement en f_0 , est une fonction propre ϵ -approchée associée à la valeur propre $H(f_0)$, c'est-à-dire vérifiant :

$$\left\| \mathcal{H}g(t)e^{j2\pi f_0 t} - H(f_0)g(t)e^{j2\pi f_0 t} \right\|^2 \leq \epsilon. \quad (1.11)$$

1.3.2 Canaux sélectifs en temps

Un canal est sélectif en temps lorsqu'il présente un gain n'évoluant qu'avec le temps. Par dualité temps-fréquence, on dira qu'il est dispersif en fréquence. Autrement dit, en plaçant une fréquence pure à son entrée, on obtiendra à sa sortie une ou plusieurs (voire un continuum de) répliques atténuées, déphasées et fréquentiellement translatées de cette fréquence. Pour ce type de canal, la sélectivité temporelle (ou de manière équivalente, la dispersivité fréquentielle) est identique quelle que soit la fréquence du signal à son entrée. Il est ainsi dit « linéaire invariant en fréquence » (LFI¹⁰ pour « linear frequency invariant »).

Physiquement, l'occurrence la plus connue de ce genre de canal correspond à la présence d'effet Doppler. Ce dernier se caractérise par un décalage de la fréquence du signal vu par le récepteur lorsque la distance entre l'émetteur et le récepteur varie : elle est translatée vers les hautes fréquences lorsque la distance diminue, et vers les basses lorsque la distance augmente.

10. *Linear Frequency Invariant* (Linéaire invariant en fréquence)

On les rencontre également en prenant en compte les imperfections des oscillateurs locaux des émetteurs/récepteurs, induisant des décalages fréquentiels du signal émis et/ou reçu.

Mathématiquement, l'opérateur de canal prend alors la forme d'une multiplication dans le domaine temporel :

$$\begin{aligned} \mathcal{H} : \mathcal{L}_2(\mathbb{R}) &\rightarrow \mathcal{L}_2(\mathbb{R}) \\ s(t) &\mapsto r(t) = s(t).h(t). \end{aligned} \quad (1.12)$$

Il prend donc la forme d'une convolution dans le domaine fréquentiel :

$$\begin{aligned} \mathcal{H} : \mathcal{L}_2(\mathbb{R}) &\rightarrow \mathcal{L}_2(\mathbb{R}) \\ S(f) &\mapsto R(f) = (S * H)(f). \end{aligned} \quad (1.13)$$

$h(t)$ est simplement la « réponse temporelle » du canal, tandis que $H(f)$ est son « spectre Doppler » (voir figure 1.7).

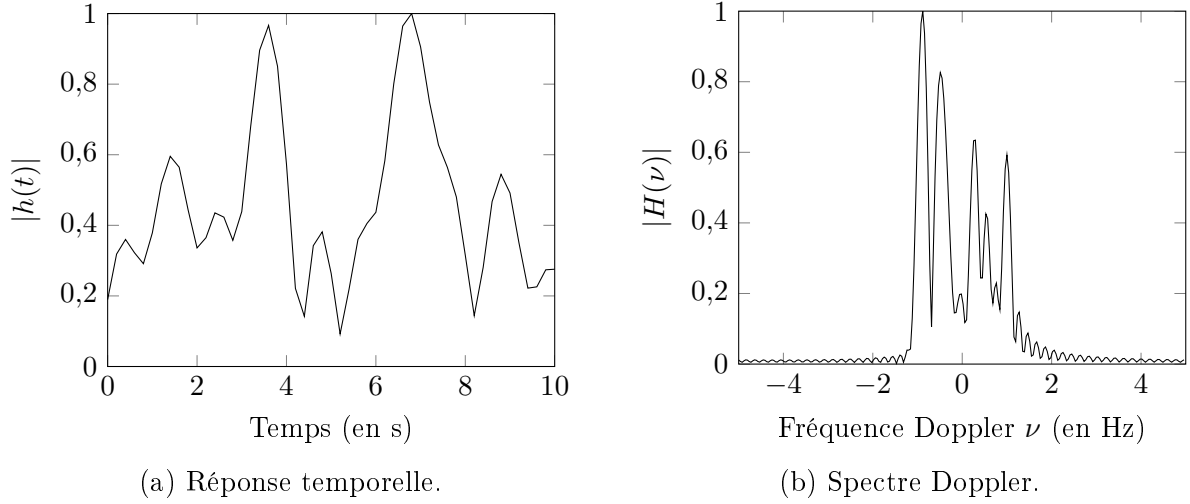


FIGURE 1.7 – Exemple de canal LFI correspondant à un modèle de Rayleigh avec un trajet de propagation et une fréquence Doppler maximale de 1 Hz.

Nous nous intéressons désormais à la structure propre de $\mathcal{H}$. Comme nous l'avons vu dans l'équation (1.12), l'opérateur de canal s'exprime comme une simple multiplication dans le domaine temporel. Cela signifie que les distributions de Dirac sont fonctions (ou distributions, en l'occurrence) propres de $\mathcal{H}$:

$$\mathcal{H}\delta(t - t_0) = h(t).\delta(t - t_0) = h(t_0).\delta(t - t_0), \quad t, t_0 \in \mathbb{R}, \quad (1.14)$$

et $h(t_0)$ sont les valeurs propres associées. En termes de systèmes de communication, cela signifie qu'une manière de s'adapter à ce type de canal est de répartir l'information sur des impulsions temporellement translatées, par exemple en construisant un signal d'émission de la forme suivante :

$$s(t) = \sum_{n \in \mathbb{Z}} c_n \delta(t - nT_0), \quad t \in \mathbb{R}, \quad (1.15)$$

1.3. Modélisation des canaux de transmission

cependant, il n'est pas physiquement envisageable de reproduire une impulsion de Dirac parfaite. On peut en revanche utiliser des fonctions propres approchées. En effet, d'après [KM97], si :

- la dispersion fréquentielle du canal est bornée : $\forall \nu \notin [-\nu_0, \nu_0], H(\nu) = 0$;
- le gain maximal du canal est unitaire ;
- $g(t)$ est un filtre passe-bas de norme unitaire à support temporel borné en $[-\frac{\epsilon}{2\pi\nu_0}, \frac{\epsilon}{2\pi\nu_0}]$, $\epsilon \in \mathbb{R}$;

alors toute version de $g(t)$ translaté temporellement en t_0 est une fonction propre ϵ -approchée associée à la valeur propre $h(t_0)$, c'est-à-dire vérifiant :

$$\|\mathcal{H}g(t - t_0) - h(t_0)g(t - t_0)\|^2 \leq \epsilon. \quad (1.16)$$

1.3.3 Canaux doublement sélectifs

Un canal doublement sélectif, ou doublement dispersif, est un canal sélectif (ou dispersif, ce qui revient au même) à la fois en temps et en fréquence. On peut donc le voir comme un canal LTI dont la sélectivité fréquentielle varie au cours du temps. Il n'est donc plus invariant en temps et, pour cette raison, nous l'appellerons par la suite canal « linéaire variant en temps » (LTV¹¹ pour « linear time variant » en anglais).

Physiquement, on pourrait dire qu'on les retrouve partout, puisqu'aucun système physique n'est parfaitement invariant en temps. Cependant, on parlera plutôt dans ce cas de canaux LTI ou LFI à « variations lentes », selon si la composante multitrajet ou Doppler est dominante. Par exemple, un canal sélectif en fréquence dont la fonction de transfert n'évolue que de manière très marginale pendant la durée de transmission d'un symbole ne sera pas qualifié ici de LTV, mais de LTI à variation lente. Ainsi, la dénomination de LTV ne dépend finalement pas que du canal, mais aussi du système de transmission, et en particulier de la bande de fréquence et de l'intervalle temporel occupés par la transmission. En environnement radiomobile (soit, en présence de trajets multiples et de mobilité entre l'émetteur et le récepteur), on considèrera que les échelles de variations temporelles et fréquentielles du canal seront telles qu'elles pourront varier de manière significative sur la bande et le temps occupés par chaque symbole.

Mathématiquement, il y a plusieurs manières de modéliser un tel canal :

- par une réponse impulsionnelle évolutive, notée $h(t, \tau)$ avec t le temps et τ le retard (plan temps-retard) ;
- par une fonction de transfert évolutive, notée $L_h(f, t)$ avec f la fréquence et t le temps (plan temps-fréquence) ;
- par une fonction d'étalement, notée $S_h(\nu, \tau)$ avec ν le décalage Doppler et τ le retard (plan retard-Doppler) ;
- par une fonction caractéristique, notée $H_h(\nu, f)$ avec ν le décalage Doppler et f la fréquence.

La figure 1.8 illustre la manière dont ces quatre fonctions représentent le canal LTV résultant de l'association des canaux LTI et LFI des figures 1.6 et 1.7.

11. *Linear Time Variant* (Linéaire variant en temps)

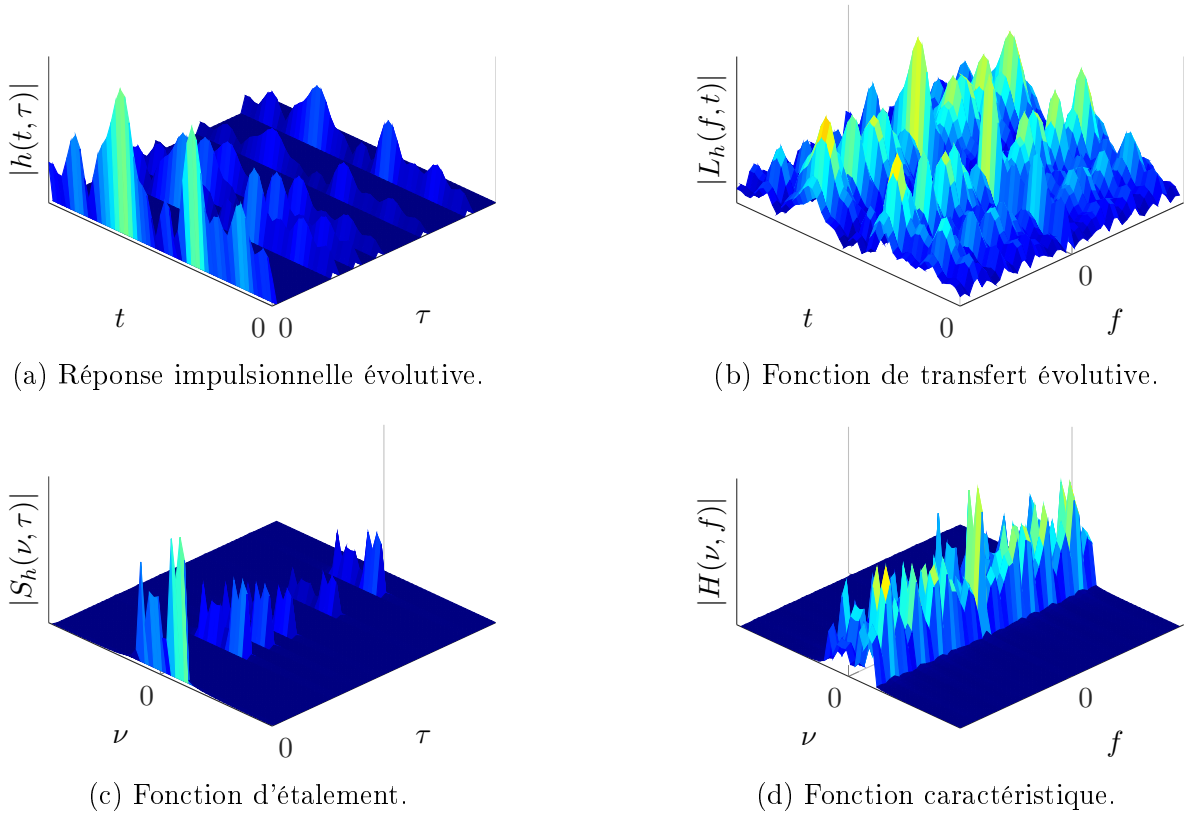


FIGURE 1.8 – Exemple de canal LTV correspondant à l'association des canaux des figures 1.6 et 1.7.

1.3. Modélisation des canaux de transmission

En utilisant la fonction de transfert évolutive, l'opérateur de canal revient à une convolution temporelle, dont le noyau de convolution évolue dans le temps :

$$\begin{aligned} \mathcal{H} : \mathcal{L}_2(\mathbb{R}) &\rightarrow \mathcal{L}_2(\mathbb{R}) \\ s(t) &\mapsto r(t) = \int_{-\infty}^{+\infty} h(t, \tau) s(t - \tau) d\tau. \end{aligned} \quad (1.17)$$

La fonction de transfert évolutive s'obtient en prenant la transformée de Fourier sur la variable représentant le retard de la réponse impulsionnelle évolutive :

$$L_h(f, t) = \int_{-\infty}^{+\infty} h(t, \tau) e^{-j2\pi f\tau} d\tau, \quad (1.18)$$

et on peut montrer que l'opérateur de canal revient à effectuer la transformée de Fourier inverse de la représentation fréquentielle du signal, pondérée de manière différente selon l'instant considéré, du signal émis :

$$\begin{aligned} \mathcal{H} : \mathcal{L}_2(\mathbb{R}) &\rightarrow \mathcal{L}_2(\mathbb{R}) \\ s(t) &\mapsto r(t) = \int_{-\infty}^{+\infty} L_h(f, t) S(f) e^{j2\pi ft} df. \end{aligned} \quad (1.19)$$

La fonction d'étalement est la transformée de Fourier sur la variable représentant le temps dans la réponse impulsionnelle évolutive :

$$S_h(\nu, \tau) = \int_{-\infty}^{+\infty} h(t, \tau) e^{-j2\pi \nu t} dt, \quad (1.20)$$

et on peut montrer que l'opérateur de canal revient à sommer des répliques du signal émis translatées en chaque point du plan retard-Doppler et pondérées par la valeur de la fonction d'étalement en ce point :

$$\begin{aligned} \mathcal{H} : \mathcal{L}_2(\mathbb{R}) &\rightarrow \mathcal{L}_2(\mathbb{R}) \\ s(t) &\mapsto r(t) = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} S_h(\nu, \tau) s(t - \tau) e^{j2\pi \nu \tau} d\nu d\tau. \end{aligned} \quad (1.21)$$

Enfin, la fonction caractéristique s'obtient par double transformation de Fourier (plus précisément, par transformation de Fourier symplectique) de la réponse impulsionnelle évolutive :

$$H_h(\nu, f) = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} h(t, \tau) e^{-j2\pi \nu t} e^{-j2\pi f\tau} dt d\tau, \quad (1.22)$$

son effet sur le signal émis est plus facile à voir dans le domaine fréquentiel. En effet, s'agissant du dual fréquentiel de la réponse impulsionnelle évolutive, on peut montrer que l'opérateur de canal consiste en une convolution de la représentation fréquentielle du signal émis avec un noyau de convolution évoluant en fonction de la fréquence :

$$\begin{aligned} \mathcal{H} : \mathcal{L}_2(\mathbb{R}) &\rightarrow \mathcal{L}_2(\mathbb{R}) \\ S(f) &\mapsto R(f) = \int_{-\infty}^{+\infty} H_h(f - \nu, \nu) S(\nu) d\nu. \end{aligned} \quad (1.23)$$

En tant que cas particuliers de canaux LTV, les canaux LTI et LFI donnent des formes particulières de ces quatre fonctions, comme le montre le tableau 1.1. On vérifie alors aisément

qu'un canal LTI opère une simple pondération des fréquences du signal émis (équation 1.19 avec $L_h(f, t) = H(f)$), et qu'un canal LFI opère une simple pondération du signal émis dans le domaine temporel (équation 1.17 avec $h(t, \tau) = h(t)\delta(\tau)$). On ne peut cependant pas exprimer le canal LTV de manière aussi triviale avec ces outils, et une généralisation de l'expression de la structure propre d'un canal LTV, n'existe pas. On peut en revanche construire des fonctions propres approchées selon une méthode très semblable à celles adoptées dans les parties précédentes (1.3.1 et 1.3.2). En effet, toujours d'après [KM97], si :

- la dispersivité du canal, passé un certain retard et un certain décalage Doppler, est négligeable (on la considèrera alors nulle) : $\forall(\nu, \tau) \notin [-\nu_0, \nu_0] \times [-\tau_0, \tau_0], S_h(\nu, \tau) = 0$;
- le gain maximal du canal est unitaire ;
- $g(t) \in \mathcal{L}_2(\mathbb{R})$ est de norme unitaire et tel que la plupart de son énergie soit concentrée dans la surface temps-fréquence $[-\nu_0, \nu_0] \times [-\tau_0, \tau_0]$: $|A_g(\nu, \tau) - 1| < \epsilon$, avec $A_g(\nu, \tau) = \int_{-\infty}^{+\infty} g^*(t)g(t - \tau)e^{j2\pi\nu t}dt$ la fonction d'ambiguïté de g ;

alors toute version de $g(t)$ translatée en (t_0, f_0) dans le plan temps-fréquence est une fonction propre ϵ -approchée associée à la valeur propre $L_h(f_0, t_0)$, c'est-à-dire vérifiant :

$$\left\| \mathcal{H}g(t - t_0)e^{j2\pi f_0 t} - L_h(t_0, f_0)g(t - t_0)e^{j2\pi f_0 t} \right\|^2 \leq \epsilon. \quad (1.24)$$

TABLE 1.1 – Valeurs particulières de fonctions de modélisation des canaux idéal, LTI et LFI.

	$h(t, \tau)$	$L_h(f, t)$	$S_h(\nu, \tau)$	$H_h(\nu, f)$
Idéal	$\delta(\tau)$	1	$\delta(\tau)\delta(\nu)$	$\delta(\nu)$
LTI	$h(\tau)$	$H(f)$	$h(\tau)\delta(\nu)$	$H(f)\delta(\nu)$
LFI	$h(t)\delta(\tau)$	$h(t)$	$H(\nu)\delta(\tau)$	$H(\nu)$

En termes de systèmes de communication, cela signifie qu'une manière de s'adapter au canal de manière approchée est de répartir l'information dans le plan temps-fréquence, en la mettant en forme à l'aide d'une fonction g dont la localisation temps-fréquence est adaptée à la dispersivité du canal :

$$s(t) = \sum_{(m,n) \in \mathbb{Z}^2} c_{m,n} g(t - nT_0) e^{j2\pi n F_0 t}, \quad t \in \mathbb{R}, \quad (1.25)$$

avec $c_{m,n} \in \mathcal{A}$, $\forall(m, n) \in \mathbb{Z}^2$, et $F_0, T_0 \in \mathbb{R}_+^*$. Or un signal construit de la sorte correspond à un message construit par l'intermédiaire d'un modulateur dit « multiporteuse », avec un espacement entre porteuses de F_0 Hz et un espacement entre symboles multiporteuses de T_0 s. Dans ce travail de thèse, dédié aux communications (au delà de la cadence de Nyquist) sur canal radiomobile, nous nous intéresserons donc à ce type de modulation.

1.4 Transmissions multiporteuses

Comme nous l'avons vu dans la partie précédente, les modulations multiporteuses se prêtent volontiers aux transmissions sur canal radiomobile. Cependant, elles ont été originellement développées pour s'accomoder efficacement des canaux sélectifs en fréquence. En effet, mettre en

1.4. Transmissions multiporteuses

forme les symboles par des impulsions de mise en forme très étroites en fréquence (et donc, de grande durée temporelle) permet d'approcher les fonctions propres d'un canal LTI (voir partie 1.3.1). D'autre part, on peut montrer qu'en utilisant une mise en forme rectangulaire et un préfixe cyclique, il est possible :

- de tolérer un certain montant de chevauchement fréquentiel entre impulsions de mise en forme, tout en conservant l'orthogonalité entre ces dernières, garantissant une bonne efficacité spectrale tout en conservant une complexité de réception raisonnable (récepteur linéaire) ;
- d'assimiler, de manière exacte, l'effet du canal LTI par un gain sur chaque porteuse, par l'addition d'un préfixe cyclique (qui doit alors avoir une longueur supérieure ou égale au retard maximal du canal), simplifiant alors grandement le travail du récepteur sur canal sélectif en fréquence ;
- d'implémenter de manière algorithmiquement efficace le modulateur et le démodulateur multiporteuses à l'aide de transformées de Fourier rapides (FFT ¹²).

Pour ces raisons, les modulations multiporteuses orthogonales à base de mise en forme rectangulaires (OFDM ¹³) et avec préfixe cyclique (CP-OFDM ¹⁴) sont aujourd'hui au cœur de nombreux standards de communication (ADSL [ITU99], LTE ¹⁵ [ETSI10], WiFi ¹⁶ 802.11n [IEEE09], etc.).

Pour certaines applications cependant, la mauvaise localisation spectrale des impulsions rectangulaires et/ou la perte d'efficacité spectrale due à la présence d'un préfixe cyclique pose problème. Une première approche consiste à filtrer les impulsions rectangulaire de manière à lisser ses transitions. On parle alors d'OFDM fenêtré (utilisé par exemple dans les standard VDSL [ITU04] et HomePlug AV/AV2 [Hom07] ; [Hom14]). Une autre technique, d'implémentation encore plus légère, consiste à pondérer le préfixe cyclique (WCP-OFDM ¹⁷). Pour cette dernière, plusieurs travaux de recherche ont permis de déterminer des pondérations optimisées selon le critère de la minimisation de l'énergie hors-bande et de la localisation temps-fréquence [PS11] ; [PS13]. D'autre part, le comportement de ce type de modulations sur canal radiomobile a été étudié dans [Roq12]. Ce type de stratégie permet de conserver la plupart des avantages du CP-OFDM (égalisation simplifiée, faible complexité de l'émetteur et du récepteur), tout en réduisant les lobes secondaires du spectre du signal.

Il arrive cependant que les contraintes imposées au système ne soient pas satisfaites par une « simple » fenêtre rectangulaire lissée, ni par la présence d'un préfixe cyclique. C'est par exemple le cas de l'occupation de l'intervalle de garde entre canaux télévisés, où la bande de fréquence disponible est faible (dans ce cas, retirer le préfixe cyclique permet de gagner en débit utile), et où les contraintes du régulateur sont très strictes quant à la perturbation des canaux adjacents, nécessitant un spectre bien localisé en fréquence. L'utilisation de modulations multi-

12. *Fast Fourier Transform* (Transformée de Fourier rapide)

13. *Orthogonal frequency-division multiplexing* (Multiplexage fréquentiel orthogonal)

14. *Cyclic Prefixed OFDM* (OFDM avec préfixe cyclique)

15. *Long Term Evolution* (Évolution à long terme)

16. *Wireless Fidelity* (Fidélité sans fil)

17. *Weighted Cyclic Prefixed OFDM* (OFDM avec préfixe cyclique pondéré)

porteuses filtrées (FBMC¹⁸), utilisant des impulsions de mise en forme présentant une meilleure localisation fréquentielle (et donc, à localisation temporelle constante, une meilleure localisation temps-fréquence), apparaît alors comme une solution intéressante. Néanmoins, par application du théorème de Balian-Low, il n'est pas possible de jouir simultanément des propriétés (i) d'orthogonalité entre impulsions en mise en forme et (ii) d'une bonne localisation temps-fréquence de ces impulsions. Une solution populaire pour se soustraire à cette limitation est de changer la condition d'orthogonalité, pour ne l'imposer que sur la partie réelle, et menant aux modulations OFDM-OQAM¹⁹ [FAB95].

Les travaux que nous allons détailler à partir du chapitre 2 se passent de la condition d'orthogonalité à des fins d'augmentation de l'efficacité spectrale. Cela permet également de se soustraire au théorème de Balian-Low, et donc de profiter d'impulsions de mise en forme bien localisées en temps et en fréquence sans recourir aux modulations multiporteuses avec OQAM. Ces dernières ne seront par conséquent pas traitées dans ce manuscrit. D'autre part, les modulations multiporteuses étudiées ici utiliseront un écart fréquentiel F_0 et temporel T_0 constants entre impulsions de mise en forme, ce qui revient à paver le plan temps-fréquence de manière régulière. Il est néanmoins intéressant de noter que d'autres types de pavages temps-fréquence peuvent être envisagés, tels que des pavages hexagonaux [HZ07], ou issus de la théorie des ondelettes [LIA07] (voir figure 1.9).

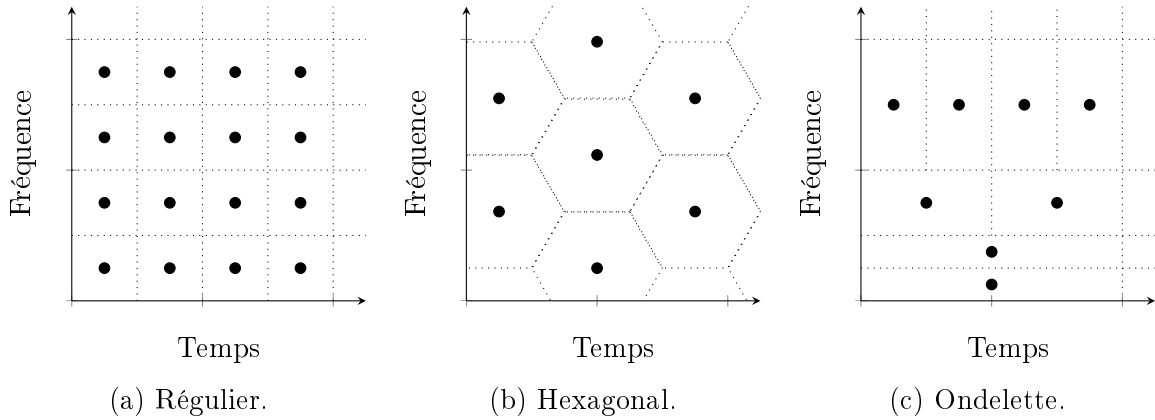


FIGURE 1.9 – Exemples de répartition des symboles dans le plan temps-fréquence.

1.4.1 Frames et familles de Gabor

Comme évoqué plus haut, les modulations multiporteuses considérées dans ce manuscrit répartissent l'information de manière régulière dans le plan temps-fréquence. Elles nécessitent donc d'être étudiées conjointement en temps et en fréquence. Pour ce faire, on s'appuiera sur

18. *FilterBank MultiCarrier* (Modulations multiporteuses à base de bancs de filtres)

19. *Offset QAM* (QAM décalé)

1.4. Transmissions multiporteuses

deux outils mathématiques issus du domaine de l'analyse temps-fréquence du signal et de la mécanique quantique : les frames et les familles de Gabor.

Commençons par les familles de Gabor (ou de Weyl–Heisenberg). Il s'agit de familles contenant toutes les T_0 -translatées en temps et F_0 -translatées en fréquence d'une fonction dite « prototype » (ou encore « atome de Gabor ») $g(t) \in \mathcal{L}_2(\mathbb{R})$:

$$\mathbf{g} = \{g_{m,n}(t) = g(t - nT_0)e^{j2\pi mF_0t}\}_{(m,n) \in \mathbb{Z}^2}. \quad (1.26)$$

Sous cette forme, on dit que $\mathbf{g}$ est une famille de Gabor de paramètre (T_0, F_0) et de prototype $g(t)$. On définit également sa densité :

$$\rho = 1/(F_0T_0). \quad (1.27)$$

Ces familles ont été introduites par Denis Gabor (avec g une gaussienne isotrope) à des fins d'analyse temps-fréquence de signaux [Gab45]. On retrouve ainsi sans surprise ce type de famille dans la définition de la transformée de Gabor (qui correspond à une transformée de Fourier à court terme discrète) :

$$c_{m,n}(f) = \int_{-\infty}^{+\infty} f(t)\check{g}^*(t - nT_0)e^{-j2\pi mF_0t}dt = \langle \check{g}_{m,n}; f \rangle, \quad \forall (m,n) \in \mathbb{Z}^2, \quad (1.28)$$

avec $f(t) \in \mathcal{L}_2(\mathbb{R})$ le signal à analyser, $\check{g}(t) \in \mathcal{L}_2(\mathbb{R})$ la fonction de fenêtrage, T_0 la résolution temporelle et F_0 la résolution fréquentielle. Les coefficients $\{c_{m,n}\}_{(m,n) \in \mathbb{Z}^2}$ peuvent être interprétés comme les coordonnées temps-fréquence du signal $f(t)$. Avec une telle interprétation, il est tentant de vouloir reconstruire le signal à partir de ses coefficients temps-fréquence de la manière suivante :

$$f(t) = \sum_{(m,n) \in \mathbb{Z}^2} c_{m,n}(f)g(t - nT_0)e^{j2\pi mF_0t}, \quad (1.29)$$

avec $g(t) \in \mathcal{L}_2(\mathbb{R})$, ce qui mène à l'interrogation suivante : comment choisir $\check{g}(t)$, $g(t)$, T_0 et F_0 pour garantir une reconstruction stable à partir de la décomposition en (1.28) et en utilisant (1.29) ? En particulier, si deux fonctions s et s' sont « proches » ($\|s - s'\| \rightarrow 0$), alors il serait souhaitable que les coefficients associés, issus de (1.28), le soient aussi ($\|\{c_{m,n}(s)\}_{(m,n) \in \mathbb{Z}^2} - \{c_{m,n}(s')\}_{(m,n) \in \mathbb{Z}^2}\| \rightarrow 0$). Une réponse à cette question a été apportée par Daubechies dans [Dau90]. Sa solution consiste à choisir $\check{\mathbf{g}}$ formant une frame, c'est-à-dire respectant l'inégalité :

$$A\|f\|^2 \leq \sum_{(m,n) \in \mathbb{Z}^2} |\langle \check{g}_{m,n}; f \rangle|^2 \leq B\|f\|^2, \quad \forall f \in \mathcal{L}_2(\mathbb{R}), \quad (1.30)$$

où $A > 0$ et $B < +\infty$ sont les bornes de la frame, et $\mathbf{g}$ une frame duale de $\check{\mathbf{g}}$. En effet on rappelle, d'après (1.28), que $c_{m,n}(f) = \langle \check{g}_{m,n}; f \rangle$. L'inégalité de frame (1.30) permet donc de garantir que la norme de la décomposition temps-fréquence d'un signal donné est relativement proportionnelle à la norme du signal lui-même. Cela est d'autant plus vrai que le rapport A/B tend vers 1. Le cas limite, c'est-à-dire $A = B$ définit les frames dites « étroites », lesquelles

possèdent différentes propriétés, décrites ci-après, qui nous seront utiles pour nos contributions (voir chapitres 3, 4 et 5).

Le paragraphe précédent demande à définir plusieurs notions liées à la théorie des frames, et plus particulièrement des frames de Gabor. Soit $\check{\mathbf{g}}$ une famille de Gabor de paramètres (T_0, F_0) et de prototype $\check{g}(t) \in \mathcal{L}_2(\mathbb{R})$. Définissons tout d'abord l'opérateur d'analyse, dont le rôle est de calculer les coefficients de la décomposition temps-fréquence :

$$\begin{aligned} T_{\check{\mathbf{g}}} : \mathcal{L}_2(\mathbb{R}) &\rightarrow \ell_2(\mathbb{Z}^2) \\ f(t) &\mapsto \{\langle \check{g}_{m,n}; f \rangle\}_{(m,n) \in \mathbb{Z}^2}. \end{aligned} \quad (1.31)$$

Vient ensuite l'opérateur de synthèse, permettant de reconstruire le signal d'origine à partir de sa décomposition temps-fréquence, et qui est aussi l'adjoint de l'opérateur d'analyse :

$$\begin{aligned} T_{\mathbf{g}}^* : \ell_2(\mathbb{Z}^2) &\rightarrow \mathcal{L}_2(\mathbb{R}) \\ \{c_{m,n}\}_{(m,n) \in \mathbb{Z}^2} &\mapsto \sum_{(m,n) \in \mathbb{Z}^2} c_{m,n} g_{m,n}(t). \end{aligned} \quad (1.32)$$

À partir de ces deux opérateurs, on définit la frame duale $\mathbf{g}$ de $\check{\mathbf{g}}$ telle que la relation :

$$T_{\mathbf{g}}^* T_{\check{\mathbf{g}}} = I, \quad (1.33)$$

soit respectée, avec I l'opérateur identité. Parmi toutes les frames duales existantes, il en existe une ayant la propriété de minimiser la norme des coefficients engendrés par l'opérateur d'analyse : la frame duale canonique [Chr08, Lemme 5.3.6]. En définissant $S_{\mathbf{g}} = T_{\mathbf{g}}^* T_{\mathbf{g}}$ l'opérateur de frame (auto-adjoint, défini positif), et son inverse $S_{\mathbf{g}}^{-1}$ (auto-adjoint, défini positif), alors la frame duale canonique duale de $\mathbf{g}$ s'écrit simplement $\check{\mathbf{g}} = S_{\mathbf{g}}^{-1} \mathbf{g}$. On vérifie rapidement que la frame duale canonique est bien une frame duale :

$$T_{\mathbf{g}}^* T_{\check{\mathbf{g}}} f = T_{\mathbf{g}}^* \langle \check{g}_{m,n}; f \rangle = T_{\mathbf{g}}^* \langle S_{\mathbf{g}}^{-1} g_{m,n}; f \rangle = T_{\mathbf{g}}^* \langle g_{m,n}; S_{\mathbf{g}}^{-1} f \rangle = T_{\mathbf{g}}^* T_{\mathbf{g}} S_{\mathbf{g}}^{-1} f = S_{\mathbf{g}} S_{\mathbf{g}}^{-1} f = f. \quad (1.34)$$

Notons que l'opérateur de frame nous permet également de réécrire l'inégalité de frame sous la forme suivante :

$$A \|f\|^2 \leq \langle S_{\check{\mathbf{g}}} f; f \rangle \leq B \|f\|^2, \quad \forall f \in \mathcal{L}_2(\mathbb{R}). \quad (1.35)$$

Comme nous l'avons évoqué plus haut, un cas particulier de frame survient lorsque la borne supérieure est égale à la borne inférieure. On parle alors de frame étroite :

$$\langle S_{\check{\mathbf{g}}} f; f \rangle = A \|f\|^2, \quad \forall f \in \mathcal{L}_2(\mathbb{R}), \quad (1.36)$$

avec $0 < A < +\infty$. Ce qui signifie que l'opérateur de frame et son inverse prennent des formes simples : $S_{\check{\mathbf{g}}} = AI$ et $S_{\check{\mathbf{g}}}^{-1} = A^{-1}I$. Notons alors que la frame duale canonique d'une frame étroite est alors obtenue par simple mise à l'échelle $\mathbf{g} = A^{-1} \check{\mathbf{g}}$.

Enfin, notons que pour une frame de Gabor, les bornes sont liées au prototype et à la densité [Dau90, II-B-2-a]. Soit $\mathbf{g}$ une frame de Gabor de paramètre (F_0, T_0) et de prototype $g(t)$, notons $\rho = 1/(F_0 T_0)$ sa densité, A sa borne inférieure et B sa borne supérieure, alors

$$A \leq \rho \|g\|^2 \leq B. \quad (1.37)$$

1.4. Transmissions multiporteuses

Maintenant que l'on en sait un peu plus sur la théorie des frames et les frames de Gabor, il convient de rappeler que toute famille de Gabor n'est pas forcément une frame. En particulier, pour une famille de Gabor $\mathbf{g}$ de densité ρ , si :

- $\rho < 1$, $\mathbf{g}$ ne peut être ni une frame, ni une base de $\mathcal{L}_2(\mathbb{R})$ ²⁰ ;
- $\rho = 1$, si $\mathbf{g}$ est une frame, alors c'est aussi une base (dite de Riesz, car elle respecte l'inégalité de frame (1.30)) ;
- $\rho > 1$, $\mathbf{g}$ peut être une frame, mais pas une base de $\mathcal{L}_2(\mathbb{R})$.

Il peut être tentant de systématiquement choisir $\rho = 1$, permettant l'unicité de la décomposition temps-fréquence. Cependant un tel choix, d'après le théorème de Balian-Low, implique un prototype mal localisé dans le plan temps-fréquence. Notons qu'une famille de Gabor avec $\rho > 1$ peut aussi très bien ne pas être une frame. De plus, il n'existe pas de procédure qui, sachant un prototype, permettent de trouver tous les paramètres F_0 et T_0 menant à des frames de Gabor.

Nous allons maintenant montrer que dans le cas de prototypes à support temporel ou fréquentiel borné, il suffit de trouver les paramètres permettant de former des familles biorthogonales, par conséquence du théorème de Wexler-Raz. On rappelle que deux familles $\mathbf{g} = \{g_{m,n}(t)\}_{(m,n) \in \mathbb{Z}^2}$ et $\check{\mathbf{g}} = \{\check{g}_{m,n}(t)\}_{(m,n) \in \mathbb{Z}^2}$ sont biorthogonales si et seulement si :

$$\langle \check{g}_{m,n}; g_{p,q} \rangle = \delta_{m-p} \delta_{n-q} \quad \forall m, n, p, q \in \mathbb{Z}. \quad (1.38)$$

Notons que, par abus de langage, nous parlerons aussi de biorthogonalité en présence d'une relation du type $\langle \check{g}_{m,n}; g_{p,q} \rangle = \alpha \delta_{m-p} \delta_{n-q}$, avec $\alpha \in \mathbb{C}^*$. Établissons tout d'abord la définition d'une séquence de Bessel. Soit $\check{\mathbf{g}} = \{\check{g}_{m,n}\}_{(m,n) \in \mathbb{Z}^2}$ une séquence dans un espace de Hilbert $\mathcal{H}$. Il s'agit alors d'une séquence de Bessel de borne B si et seulement si :

$$\sum_{(m,n) \in \mathbb{Z}^2} |\langle \check{g}_{m,n}; f \rangle|^2 \leq B \|f\|^2, \quad \forall f \in \mathcal{H}, \quad (1.39)$$

où $B < \infty$ est la borne de la séquence de Bessel. Le théorème de Wexler-Raz s'énonce alors comme ci-dessous.

Théorème 1.1 (Wexler-Raz)

Soit $g(t) \in \mathcal{L}_2(\mathbb{R})$, $\check{g}(t) \in \mathcal{L}_2(\mathbb{R})$ et $T_0 > 0$, $F_0 > 0$.

Si $\{g_{mF_0, nT_0}(t) = g(t - nT_0)e^{j2\pi mF_0 t}\}_{(m,n) \in \mathbb{Z}^2}$ et $\{\check{g}_{mF_0, nT_0}(t) = \check{g}(t - nT_0)e^{j2\pi mF_0 t}\}_{(m,n) \in \mathbb{Z}^2}$ sont deux séquences de Bessel, alors elles constituent aussi deux frames duales si et seulement si :

$$\langle g; \check{g}_{m/T_0, n/F_0} \rangle = F_0 T_0 \delta_m \delta_n = \frac{1}{\rho} \delta_m \delta_n \quad \forall (m, n) \in \mathbb{Z}^2. \quad (1.40)$$

Notons qu'en utilisant ce théorème avec $g(t) = \check{g}(t)$, on obtient une frame étroite de borne $A = 1$ (borne qui peut-être ajustée, en utilisant (1.36) et (1.37), par simple multiplication de $g(t)$ par un scalaire). Plusieurs prototypes respectant (1.40) ont été développés dans le cadre des modulations multiporteuses biorthogonales, en posant $F'_0 = 1/T_0$ l'écart entre porteuses, et $T'_0 = 1/F_0$ la durée d'un symbole multiporteuses (on rappelle que pour garantir la biorthogonalité

20. $\mathbf{g}$ peut en revanche être une base de Riesz d'un sous-espace de $\mathcal{L}_2(\mathbb{R})$, c'est-à-dire une séquence de Riesz.

d'un tel système, il faut imposer $\rho' = 1/(F_0' T_0') \leq 1$. Ces prototypes peuvent donc être utilisés pour construire des frames, à condition de pouvoir former des séquences de Bessel. Or, d'après [Chr08, Theorem 9.1.6], si $g(t) \in \mathcal{L}_2(\mathbb{R})$ est bornée et à support borné, alors les systèmes de Gabor qu'il engendre sont des séquences de Bessel.

Théorème 1.2 ([Chr08, Theorem 9.1.6])

Soit $g(t) \in \mathcal{L}_2(\mathbb{R})$ bornée et à support borné.

Alors $\{g_{mF_0, nT_0}(t) = g(t - nT_0)e^{j2\pi mF_0 t}\}_{(m,n) \in \mathbb{Z}^2}$ est une séquence de Bessel quel que soit le choix de $F_0 > 0$, $T_0 > 0$.

Ce résultat peut facilement être étendu aux fonctions bornées dans le domaine fréquentiel, et à support fréquentiel borné. En effet, si $G(f) \in \mathcal{L}_2(\mathbb{R})$ est bornée et à support borné, les systèmes de Gabor $\{G_{nT_0, mF_0}(f) = G(f - mF_0)e^{j2\pi nT_0 f}\}_{(m,n) \in \mathbb{Z}^2}$ qu'il engendre sont des séquences de Bessel quels que soient $F_0 > 0$ et $T_0 > 0$:

$$\sum_{(m,n) \in \mathbb{Z}^2} |\langle G_{nT_0, mF_0}; H \rangle|^2 \leq B \|H\|^2, \quad \forall H \in \mathcal{L}_2(\mathbb{R}). \quad (1.41)$$

Notons $G(f)$ la transformée de Fourier de $g(t)$, alors

$$\mathcal{F}^{-1}\{G_{nT_0, mF_0}(f)\}(t) = g(t + nT_0)e^{j2\pi mF_0(t+nT_0)}, \quad (1.42)$$

permettant d'écrire :

$$\sum_{(m,n) \in \mathbb{Z}^2} |\langle \mathcal{F}^{-1}\{G_{nT_0, mF_0}\}; h \rangle|^2 = \sum_{(m,n) \in \mathbb{Z}^2} \left| \int_{-\infty}^{+\infty} g^*(t + nT_0)e^{-j2\pi mF_0(t+nT_0)} h(t) dt \right|^2 \quad (1.43)$$

$$= \sum_{(m,n) \in \mathbb{Z}^2} \left| \int_{-\infty}^{+\infty} g^*(t - nT_0)e^{-j2\pi mF_0 t} h(t) dt \right|^2 \quad (1.44)$$

$$= \sum_{(m,n) \in \mathbb{Z}^2} |\langle g_{mF_0, nT_0}; h \rangle|^2, \quad (1.45)$$

avec $g_{mF_0, nT_0}(t) = g(t - nT_0)e^{j2\pi mF_0 t}$. Ainsi, par conservation du produit scalaire et de la norme par la transformée de Fourier, on a :

$$\sum_{(m,n) \in \mathbb{Z}^2} |\langle G_{nT_0, mF_0}; H \rangle|^2 = \sum_{(m,n) \in \mathbb{Z}^2} |\langle g_{mF_0, nT_0}; h \rangle|^2 \leq B \|H\|^2 = B \|h\|^2, \quad \forall h \in \mathcal{L}_2(\mathbb{R}), \quad (1.46)$$

en notant H la transformée de Fourier de h , ce qui signifie que les systèmes de Gabor $\mathbf{g} = \{g_{mF_0, nT_0}(t)\}_{(m,n) \in \mathbb{Z}^2}$ sont aussi des séquences de Bessel quels que soient $F_0 > 0$ et $T_0 > 0$.

Nous présentons maintenant les prototypes qui seront utilisés dans la suite de ce manuscrit, ainsi que les conditions pour lesquelles ils engendrent des frames duales et, éventuellement, étroites. Pour ce faire, soient $g(t) \in \mathcal{L}_2(\mathbb{R})$ et $\check{g} \in \mathcal{L}_2(\mathbb{R})$, on considérera les deux systèmes de Gabor suivants :

$$\mathbf{g} = \{g_{m,n}(t) = g(t - nT_0)e^{j2\pi mF_0 t}\}_{(m,n) \in \mathbb{Z}^2} \quad (1.47)$$

1.4. Transmissions multiporteuses

et

$$\check{\mathbf{g}} = \{\check{g}_{m,n}(t) = \check{g}(t - nT_0)e^{j2\pi m F_0 t}\}_{(m,n) \in \mathbb{Z}^2} \quad (1.48)$$

Notons que la densité $\rho = 1/(F_0 T_0)$ de ces deux familles est strictement supérieure à l'unité, puisque c'est une condition nécessaire pour obtenir une frame. À des fins de clarté, nous définissons également les deux familles

$$\mathbf{g}' = \left\{ g'_{m,n}(t) = g(t - nT'_0)e^{j2\pi m F'_0 t} \right\}_{(m,n) \in \mathbb{Z}^2} = \left\{ g'_{m,n}(t) = g(t - n/F_0)e^{j2\pi \frac{m}{T'_0} t} \right\}_{(m,n) \in \mathbb{Z}^2} \quad (1.49)$$

et

$$\check{\mathbf{g}}' = \left\{ \check{g}'_{m,n}(t) = \check{g}(t - nT'_0)e^{j2\pi m F'_0 t} \right\}_{(m,n) \in \mathbb{Z}^2} = \left\{ \check{g}'_{m,n}(t) = \check{g}(t - n/F_0)e^{j2\pi \frac{m}{T'_0} t} \right\}_{(m,n) \in \mathbb{Z}^2} \quad (1.50)$$

La densité de ces deux familles est donc $\rho' = 1/F'_0 T'_0 = F_0 T_0 < 1$. Ainsi, si $\mathbf{g}$ et $\check{\mathbf{g}}$ sont des familles de Bessel, alors pour montrer qu'il s'agit de frames duales en vertu du théorème de Wexler-Raz, il suffira de montrer que $\mathbf{g}'$ et $\check{\mathbf{g}}'$ sont biorthogonales :

$$\langle g'_{p,q}; \check{g}'_{m,n} \rangle = \rho' \delta_{m-p} \delta_{n-q} \quad \forall m, n, p, q \in \mathbb{Z}. \quad (1.51)$$

Prototypes rectangulaires

Notons $\Pi_T(t)$ la fonction porte de largeur $T > 0$, définie par :

$$\Pi_T(t) = \begin{cases} 1 & \text{si } |t| < \frac{T}{2}, \\ 0 & \text{sinon.} \end{cases} \quad (1.52)$$

Ces prototypes sont utilisés dans le cadre de modulations multiporteuses orthogonales, en particulier en OFDM avec ou sans préfixe cyclique. Dans ce cas, la condition de biorthogonalité entre $\mathbf{g}'$ et $\check{\mathbf{g}}'$ est respectée en choisissant $g(t) = \sqrt{1/T'_0} \Pi_{1/F'_0}(t)$ et $\check{g}(t) = \sqrt{1/T'_0} \Pi_{T'_0}(t)$, avec $\rho' \leq 1$. On remarque que la condition est aussi respectée pour $\check{g}(t) = g(t) = \sqrt{1/T'_0} \Pi_{1/F'_0}(t)$.

Ainsi, en choisissant $\check{g}(t) = g(t) = \sqrt{F_0} \Pi_{T_0}(t)$ (borné par $\sqrt{F_0}$ et à support borné en $[-T_0/2; T_0/2]$, voir figure 1.10), on obtient, pour $\rho > 1$, une famille $\mathbf{g} = \check{\mathbf{g}}$ constituant une frame étroite de borne $A = 1$.

De même, en choisissant $g(t) = \sqrt{F_0} \Pi_{T_0}(t)$ et $\check{g}(t) = \sqrt{F_0} \Pi_{1/F_0}(t)$, $\mathbf{g}$ et $\check{\mathbf{g}}$ sont deux frames duales, mais puisque $g(t) \neq \check{g}(t)$, $\check{\mathbf{g}}$ n'est pas étroite.

Enfin, si l'on choisit $g(t) = \sqrt{F_0} \Pi_{1/F_0}(t)$, alors $\mathbf{g}$ ne forme pas une frame étroite. On peut donc construire sa duale canonique à partir du prototype $\check{g}(t) = S_{\mathbf{g}}^{-1} g(t)$ (voir figure 1.11). Dans ce cas, là-encore, l'opérateur de frame inverse est évalué numériquement via la fonction `gabdual` de la toolbox Matlab LTFAT (<http://ltfat.sourceforge.net/doc/gabor/gabdual.php>).

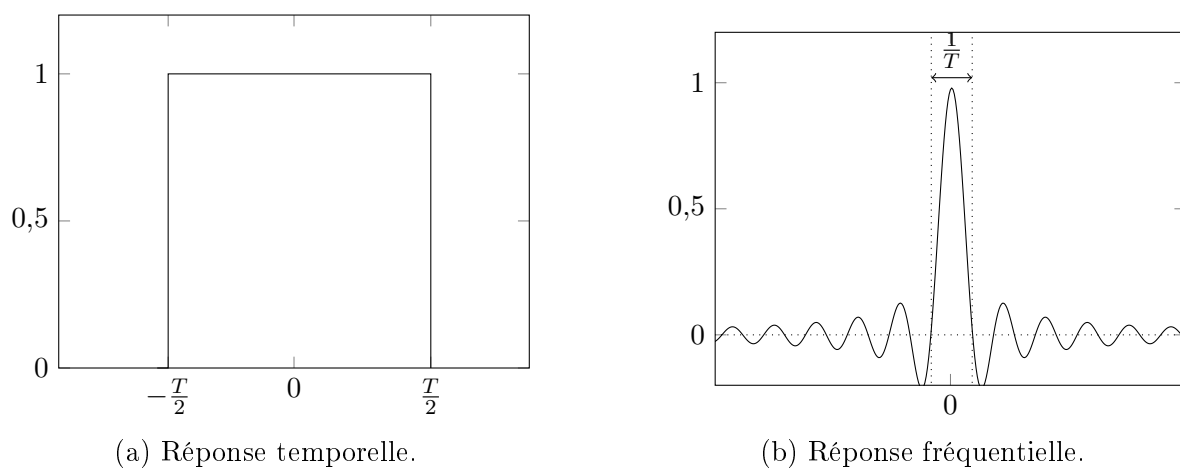


FIGURE 1.10 – Réponses temporelle et fréquentielle normalisées du prototype rectangulaire de largeur T : $\Pi_T(t)$.

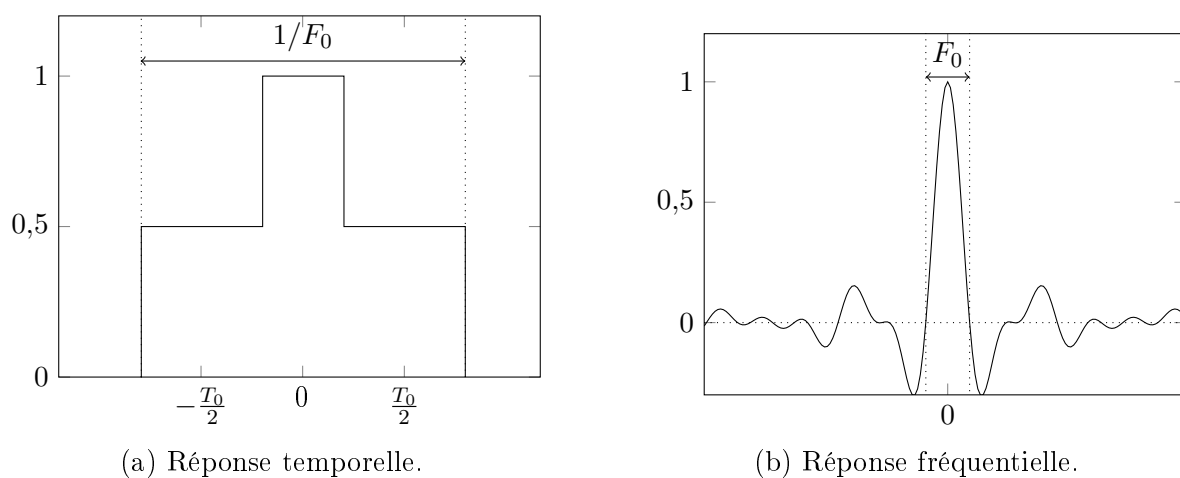


FIGURE 1.11 – Réponses temporelle et fréquentielle normalisées du prototype de la frame duale canonique pour la frame de Gabor $\mathbf{g}$, de prototype $g(t) = \Pi_{1/F_0}(t)$ et de densité $\rho = 1.6$: $S_{\mathbf{g}}^{-1}g(t)$.

1.4. Transmissions multiporteuses

Prototypes en racine de cosinus surélevé

Notons $\text{RCS}_{T,\alpha}(t)$ le prototype en racine de cosinus surélevé de coefficient d'amortissement $0 \leq \alpha \leq 1$, et $T > 0$, défini par :

$$\text{RCS}_{T,\alpha}(t) = \begin{cases} (1 + \alpha (\frac{4}{\pi} - 1)) & \text{si } t = 0, \\ \frac{\alpha}{\sqrt{2}} \left((1 + \frac{2}{\pi}) \sin(\frac{\pi}{4\alpha}) + (1 - \frac{2}{\pi}) \cos(\frac{\pi}{4\alpha}) \right) & \text{si } t \pm \frac{T}{4\alpha}, \\ \frac{T \sin(\pi t(1-\alpha)/T) + 4\alpha t \cos(\pi t(1+\alpha)/T)}{\pi t(1-(4\alpha t/T)^2)} & \text{sinon.} \end{cases} \quad (1.53)$$

Dans le cadre de modulations multiporteuses biorthogonales, la condition de biorthogonalité entre $\mathbf{g}'$ et $\check{\mathbf{g}}'$ est respectée pour $\rho' \leq 1$ en choisissant $g(t) = \check{g}(t) = F'_0 \text{RCS}_{T'_0,\alpha}(t)$ et $\alpha \leq 1/\rho' - 1$.

Ainsi, si $\check{g}(t) = g(t) = \text{RCS}_{1/F_0,\alpha}(t)/\sqrt{T_0}$ (de spectre borné par $1/\sqrt{\rho}$ et à support fréquentiel borné en $[-(1+\alpha)F_0/2; (1+\alpha)F_0/2]$, voir figure 1.12), alors $\mathbf{g} = \check{\mathbf{g}}$ est une frame étroite de borne $A = 1$.

On peut également choisir $g(t) = \text{RCS}_{T_0,\alpha}(t)$, dans ce cas $\mathbf{g}$ ne forme pas une frame étroite. On peut donc construire sa duale canonique à partir du prototype $\check{g}(t) = \mathbf{S}_{\mathbf{g}}^{-1}g(t)$ (voir figure 1.11). Dans ce cas, l'opérateur de frame inverse est évalué numériquement via la fonction `gabddual` de la toolbox Matlab LTFAT (<http://lftat.sourceforge.net/doc/gabor/gabddual.php>).

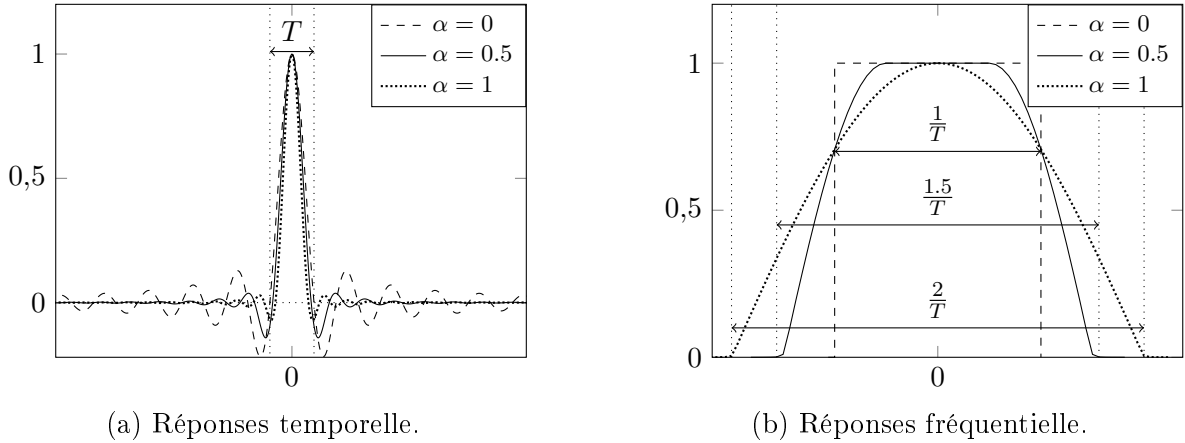


FIGURE 1.12 – Exemple de réponses temporelles et fréquentielles normalisées du prototype en racine de cosinus surélevé de support fréquentiel $[-(1+\alpha)/(2T); (1+\alpha)/(2T)]$: $\text{RCS}_{T,\alpha}(t)$.

Prototypes courts minimisant l'énergie hors-bande (EHB) et maximisant la localisation temps-fréquence (LTF)

Les prototypes EHB et LTF sont deux prototypes optimisés de manière à minimiser l'énergie hors-bande et maximiser la localisation temps-fréquence, respectivement, tout en garantissant l'orthogonalité entre $\mathbf{g}'$ et $\check{\mathbf{g}}'$ pour $\rho' \leq 1$. Leurs expressions sont données dans [PS11] et [PS13].

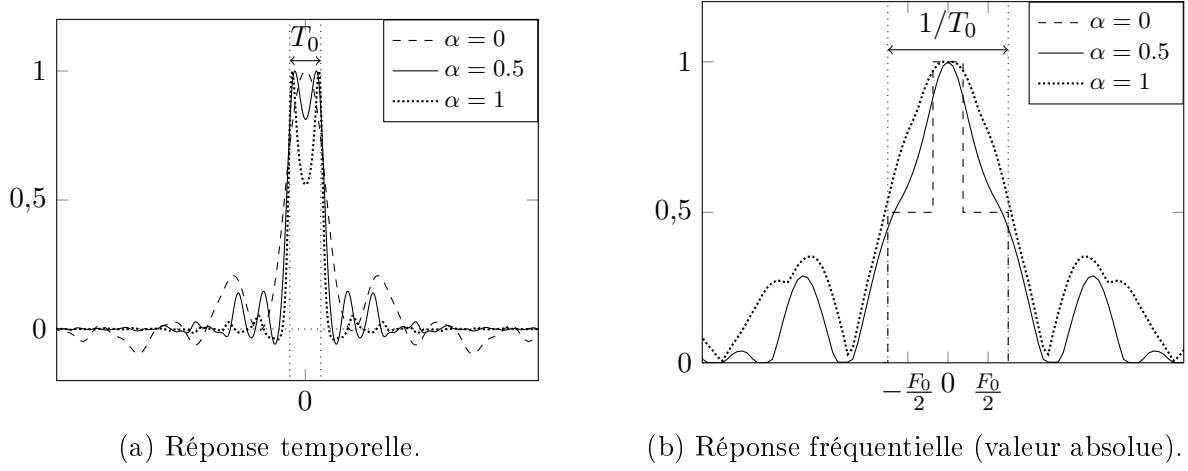


FIGURE 1.13 – Exemples de réponses temporelle et fréquentielle normalisées du prototype de la frame duale canaonique pour la frame de Gabor $\mathbf{g}$, de prototype $g(t) = \text{RCS}_{T_0, \alpha}(t)$ et de densité $\rho = 1.6 : S_{\mathbf{g}}^{-1}g(t)$.

Ils sont utilisés en particulier dans les modulations multiporteuses à préfixe cyclique pondéré [Roq12]. Ils présentent également un support temporel borné (voir figure 1.14) et sont dits « courts » puisque dans le cadre de modulations multiporteuses biorthogonales, leur support correspond à la durée d'un symbole multiporteuse (T'_0).

Notons $\text{EHB}_T(t)$ et $\text{LTF}_T(t)$ les prototypes EHB et LTF, respectivement, de support $[-T/2; T/2]$ et de norme unitaire. Alors, par construction, si $g(t) = \check{g}(t) = \sqrt{\rho} \text{EHB}_{1/F_0}(t)$ ou si $g(t) = \check{g}(t) = \sqrt{\rho} \text{LTF}_{1/F_0}(t)$, alors $\mathbf{g} = \check{\mathbf{g}}$ est une frame étroite de borne $A = 1$. En revanche, pour tout $\rho > 1$, ils s'écartent des solutions optimales pour les critères de minimisation de l'énergie hors bande et localisation temps-fréquence. Les valeurs possibles de ρ sont :

- $21/20 \leq \rho \leq 5/3$ pour $g(t) = \check{g}(t) = \sqrt{\rho} \text{EHB}_{1/F_0}(t)$ [PS11] ;
- $21/20 \leq \rho \leq 2$ pour $g(t) = \check{g}(t) = \sqrt{\rho} \text{LTF}_{1/F_0}(t)$ [PS13].

1.4.2 Chaîne de communication multiporteuse

Puisque les modulations multiporteuses ont pour but de répartir l'information de manière régulière dans le plan temps-fréquence, l'opération de modulation multiporteuse peut être vue comme le dual de l'analyse temps-fréquence d'un signal : on dispose de symboles à différentes coordonnées temps-fréquence, et on cherche à les convertir en signal de telle sorte qu'il soit possible de les reconstruire à partir de ce signal.

Dans la suite de cette partie, on considérera la modulation d'une séquence à énergie finie $\mathbf{c}$ de symboles $c_{m,n}$ indexés par n en temps et m en fréquence, appartenant à une constellation $\mathcal{A}$, centrée, indépendants et identiquement distribués et de variance σ_c^2 :

- $\mathbf{c} = \{c_{m,n}\}_{(m,n) \in \Lambda} \in \ell_2(\Lambda)$;

1.4. Transmissions multiporteuses

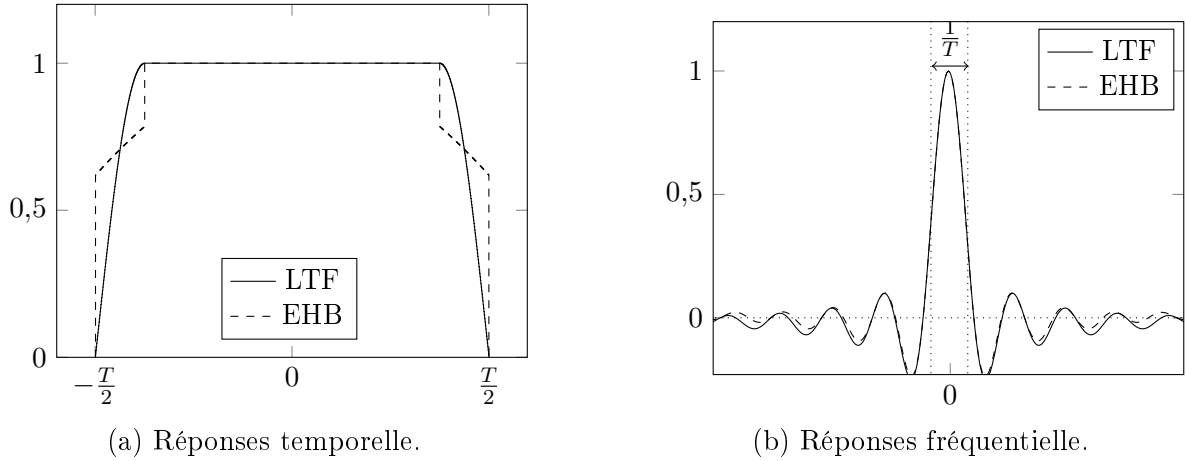


FIGURE 1.14 – Exemples de réponses temporelles et fréquentielles normalisées des prototypes EHB et LTF de largeur T : $\text{EHB}_T(t)$ et $\text{LTF}_T(t)$ construits ici pour former une frame étroite à $\rho = 8/7$ ($T = 1/F_0$).

$$\mathbb{E} \{c_{m,n} c_{p,q}^*\} = \sigma_c^2 \delta_{m-p} \delta_{n-q}, \forall (m,n), (p,q) \in \Lambda.$$

En pratique, une transmission se limitera à l'utilisation de M porteuses et de K symboles multiporteuses, tel que $\Lambda = [0; M-1] \times [0; K-1]$. Cependant, lorsque ces deux nombres prennent des valeurs élevées, et pour simplifier les calculs, on substituera parfois Λ par $\mathbb{Z}^2$.

Le modulateur consiste en un banc de synthèse utilisant un prototype d'émission $g(t) \in \mathcal{L}_2(\mathbb{R})$ à temps continu et $g[k] \in \ell_2(\mathbb{Z})$ à temps discret. On parlera parfois d'impulsion de mise en forme en lieu et place de prototype d'émission. Le signal en sortie du modulateur est noté $s(t) \in \mathcal{L}_2(\mathbb{R})$ à temps continu et $s[k] \in \ell_2(\mathbb{Z})$ à temps discret. Le démodulateur consiste en un banc d'analyse utilisant un prototype de réception $\check{g}(t)$ à temps continu et $\check{g}[k]$ à temps discret. Ce dernier reçoit un signal $r(t) \in \mathcal{L}_2(\mathbb{R})$ à temps continu et $r[k] \in \ell_2(\mathbb{Z})$ à temps discret qui est le résultat de la corruption par le canal du signal émis. Dans le cas général, le signal reçu s'écrit $r(t) = \mathcal{H}s(t) + b(t)$, (respectivement $r[k] = \mathcal{H}s[k] + b[k]$ à temps discret) avec $\mathcal{H}$ l'opérateur de canal (cf. partie 1.3), et $b(t)$ (respectivement $b[k]$) un bruit additif blanc gaussien de variance σ_b^2 . La sortie du démodulateur consiste en des estimées de $\mathbf{c}$ que nous noterons $\tilde{\mathbf{c}}$. Ces notations sont illustrées en figure 1.15.

Enfin, pour permettre une comparaison juste de différentes modulations, nous comparerons leurs performances en fonction du rapport signal à bruit par bit E_b/N_0 . Cela permet notamment de ne pas être biaisé par l'amplification liée à la norme du prototype d'émission. Ainsi, l'énergie par bit est définie par :

$$E_b = \frac{\sigma_c^2 \|g\|^2}{2N_b}, \quad (1.54)$$

où $N_b = \log_2(|\mathcal{A}|)$ est le nombre de bits par symbole. La notation N_0 signifie que l'on considère un bruit de densité spectrale de puissance bilatérale du bruit sur onde porteuse égale à $N_0/2$.

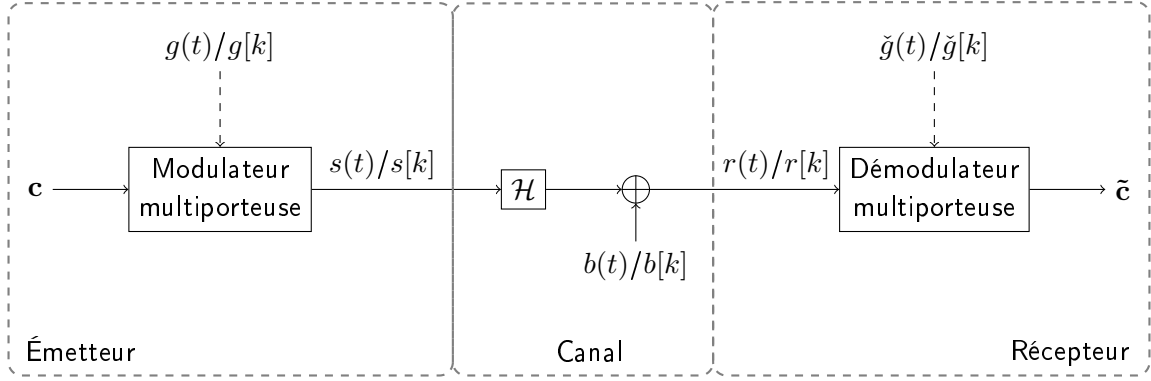


FIGURE 1.15 – Synoptique d’une chaîne de communication multiporteuse.

On la lie à la variance du bruit en bande de base par :

$$N_0 = \frac{\sigma_b^2}{2B}, \quad (1.55)$$

Avec B la bande occupée par le signal.

1.4.2.1 À temps continu

Notons F_0 l’espacement entre porteuses, et T_0 l’espacement entre symboles multiporteuses. La densité du système est donc $\rho = 1/(F_0 T_0)$. On définit la famille de Gabor d’émission, de prototype $g(t) \in \mathcal{L}_2(\mathbb{R})$ par $\mathbf{g} = \{g_{m,n}(t) = g(t - nT_0)e^{j2\pi m F_0 t}\}_{(m,n) \in \Lambda}$, et la famille de réception, de prototype $\check{g}(t) \in \mathcal{L}_2(\mathbb{R})$ par $\check{\mathbf{g}} = \{\check{g}_{m,n}(t) = \check{g}(t - nT_0)e^{j2\pi m F_0 t}\}_{(m,n) \in \Lambda}$.

L’opération de modulation s’écrit :

$$s(t) = \sum_{(m,n) \in \Lambda} c_{m,n} g_{m,n}(t) = \sum_{(m,n) \in \Lambda} c_{m,n} g(t - nT_0) e^{j2\pi m F_0 t}. \quad (1.56)$$

On remarque que lorsque $\Lambda = \mathbb{Z}^2$, il s’agit d’une opération de synthèse : $s(t) = \mathbf{T}_{\mathbf{g}}^* \mathbf{c}$. Dans ce cas, il est intéressant que $\mathbf{g}$ soit une séquence de Bessel, puisque cela garantit que $\mathbf{T}_{\mathbf{g}}^*$ est un opérateur borné [Chr08, Théorème 3.1.3].

En réception, la démodulation s’exprime par :

$$\tilde{c}_{p,q} = \langle \check{g}_{p,q}; r \rangle = \int_{-\infty}^{+\infty} \check{g}_{p,q}^*(t) r(t) dt = \int_{-\infty}^{+\infty} \check{g}^*(t - qT_0) r(t) e^{-j2\pi p F_0 t} dt. \quad (1.57)$$

On remarque que pour $\Lambda = \mathbb{Z}^2$ il s’agit de l’analyse de $r(t)$: $\tilde{\mathbf{c}} = \mathbf{T}_{\check{\mathbf{g}}} r(t)$. Dans ce cas, il est intéressant que $\check{\mathbf{g}}$ soit une frame puisque, comme nous l’avons vu au début de la partie précédente, cela garantit que plus le signal reçu $r(t)$ est proche du signal émis $s(t)$ (ou, de manière équivalente, moins le signal émis est modifié par le canal), plus les symboles estimés

1.4. Transmissions multiporteuses

par le démodulateur seront proches de ceux estimés en présence d'un canal parfait (sans bruit et non-dispersif).

En développant $r(t) = \mathcal{H}s(t) + b(t)$ dans (1.57), et en exploitant la linéarité du produit scalaire et de $\mathcal{H}$, on obtient :

$$\begin{aligned} \tilde{c}_{p,q} &= \langle \check{g}_{p,q}; \mathcal{H}s \rangle + \langle \check{g}_{p,q}; b \rangle \\ &= \underbrace{c_{p,q} \langle \check{g}_{p,q}; \mathcal{H}g_{p,q} \rangle}_{\text{Signal utile}} + \underbrace{\sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} c_{m,n} \langle \check{g}_{p,q}; \mathcal{H}g_{m,n} \rangle}_{\text{Interférence : } i_{p,q}} + \underbrace{\langle \check{g}_{p,q}; b \rangle}_{\text{Bruit filtré}}. \end{aligned} \quad (1.58)$$

Or, comme nous l'avons vu dans la partie 1.3.3, si la dispersivité du canal est bornée et si $g(t)$ possède une bonne localisation temps-fréquence, on peut écrire $\mathcal{H}g_{m,n}(t) \approx L_h(mF_0, nT_0)g_{m,n}(t)$. Dans ce cas, (1.58) devient :

$$\begin{aligned} \tilde{c}_{p,q} &\approx c_{p,q} L_h(pF_0, qT_0) \langle \check{g}; g \rangle \\ &\quad + \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} c_{m,n} L_h(mF_0, nT_0) \langle \check{g}_{p,q}; g_{m,n} \rangle + \langle \check{g}_{p,q}; b \rangle. \end{aligned} \quad (1.59)$$

Enfin, si $\mathbf{g}$ et $\check{\mathbf{g}}$ sont biorthogonales, alors le terme d'interférence s'annule de telle sorte que les symboles en sortie du démodulateur s'écrivent :

$$\tilde{c}_{p,q} \approx c_{p,q} L_h(pF_0, qT_0) \langle \check{g}; g \rangle + \langle \check{g}_{p,q}; b \rangle. \quad (1.60)$$

Ainsi, l'égalisation peut très simplement s'effectuer par division de $\tilde{c}_{p,q}$ par $L_h(pF_0, qT_0) \langle \check{g}; g \rangle$, où $\langle \check{g}; g \rangle$ est connu, et $L_h(pF_0, qT_0)$ est estimé par l'insertion de symboles déterministes et connus dans $\mathbf{c}$ (insertion de symboles et porteuses « pilotes »). Cette simplicité d'égalisation, conjuguée à des implémentations efficaces de l'émetteur/récepteur à base de FFT et/ou de filtrage polyphase [Sic02], ont majoritairement participé à la démocratisation des systèmes de transmission à base de modulations multiporteuses.

Cependant, pour respecter la condition de biorthogonalité, il est nécessaire d'avoir une correspondance unique entre chaque séquence de symboles $\mathbf{c}$ possible et chaque signal possible $s(t)$ en sortie du démodulateur. De manière équivalente, on peut dire que chaque vecteur $s(t)$ généré doit être associé à une unique décomposition $\mathbf{c} = \mathbf{T}_{\mathbf{g}} s(t)$ dans $\ell_2(\mathbb{Z})$. Or, pour ce faire, il faut que $\mathbf{g}$ soit une séquence de Riesz. Cependant, comme évoqué dans la partie précédente, si $\rho > 1$, alors $\mathbf{g}$ ne peut pas être une séquence de Riesz, ce qui induit systématiquement un terme d'interférence entre impulsions de mise en forme. Ainsi, pour les modulations multiporteuses, on généralisera la notion de cadence de Nyquist de la manière suivante :

- si $\rho < 1$, la modulation opère en deçà de la cadence de Nyquist (STN ²¹) ;
- si $\rho = 1$, la modulation opère à la cadence de Nyquist ;
- si $\rho > 1$, la modulation opère au-delà de la cadence de Nyquist (FTN ²²).

21. *Slower-Than-Nyquist* (En deçà de la cadence de Nyquist)

22. *Faster-Than-Nyquist* (Au-delà de la cadence de Nyquist)

1.4.2.2 À temps discret

À des fins de simulation et autres implémentations numériques de l'émetteur/récepteur multiporteuse, nous détaillons ici leurs relations d'entrée/sortie à temps discret.

On se place dans le cas où $\Lambda = [0; M-1] \times [0; K-1]$. Considérons que le prototype $g(t)$ (supposé passe-bas et centré en 0) occupe une bande bilatérale B_g , on obtient alors que le signal équivalent complexe en bande de base $s(t)$ en sortie du démodulateur occupe une bande bilatérale $B_s = (M-1)F_0 + B_g$. En supposant un grand nombre de porteuses (plus précisément : $|B_g - F_0|/(MF_0) \ll 1$), la bande occupée par le signal modulé est bien approchée par $B_s = MF_0$.

Ainsi, on peut choisir la fréquence d'échantillonnage de manière à effectuer un échantillonnage critique de $s(t)$: $F_e = 1/T_e = B_s$. On a donc $F_0 = F_e/M$. D'autre part, en notant N le nombre d'échantillons par symbole multiporteuse, on a $T_0 = NT_e$. On constate alors que la densité s'écrit alors $\rho = 1/(F_0T_0) = M/N$.

Exprimons maintenant le signal modulé à temps discret :

$$s[k] = \sqrt{T_e} s(kT_e) = \sqrt{T_e} \sum_{(m,n) \in \Lambda} c_{m,n} g(kT_e - nT_0) e^{j2\pi m F_0 k T_e} \quad (1.61)$$

$$= \sum_{(m,n) \in \Lambda} c_{m,n} \sqrt{T_e} g((k - nN)T_e) e^{j2\pi \frac{m}{M} k} \quad (1.62)$$

$$= \sum_{(m,n) \in \Lambda} c_{m,n} g_{m,n}[k], \quad (1.63)$$

où le facteur $\sqrt{T_e}$ permet la conservation de l'énergie entre le domaine continu et le domaine discret, et $g_{m,n}[k] = \sqrt{T_e} g((k - nN)T_e) e^{j2\pi \frac{m}{M} k}$. En définissant $g[k] = \sqrt{T_e} g(kT_e)$, on obtient :

$$g_{m,n}[k] = g[k - nN] e^{j2\pi \frac{m}{M} k}. \quad (1.64)$$

De la même manière qu'à temps continu, on obtient les symboles en sortie du démodulateur par projection du signal reçu sur la famille de réception :

$$\tilde{c}_{p,q} = \langle \check{g}_{p,q}; r \rangle = \sum_{k=-\infty}^{\infty} \check{g}_{p,q}^*[k] r[k] = \sum_{k=-\infty}^{\infty} \check{g}^*[k - qN] e^{j2\pi \frac{p}{M} k} r[k], \quad (1.65)$$

avec $\check{g}_{p,q}[k] = \check{g}[k - qN] e^{j2\pi \frac{p}{M} k}$ et $\check{g}[k] = \sqrt{T_e} \check{g}(kT_e)$.

Des études portant sur l'implémentation efficace de tels systèmes selon les prototypes utilisés sont par exemple à trouver dans [Sic02].

1.5 Communications à haute efficacité spectrale

Afin d'être en mesure de choisir le schéma de communication le mieux adapté à une application ou un scénario donné, il faut disposer de critères de comparaison pertinents. Ceci n'est cependant pas chose aisée. En effet, ces dernières peuvent être caractérisées par différents critères tels que, par exemple, la consommation énergétique, la complexité algorithmique, le débit ou

1.5. Communications à haute efficacité spectrale

encore la probabilité d'erreur. On comprend alors aisément qu'une bonne modulation pour une application donnée devra respecter plusieurs de ces critères.

Une modulation ayant pour but fondamental de transmettre de l'information, un critère de comparaison intéressant consiste à quantifier la quantité d'information transmise par unité de ressource physique, à conditions de canal fixées. Dans le cadre de ce manuscrit, la ressource physique que nous souhaitons économiser est la surface temps-fréquence occupée par la modulation. Cela nous mène à les comparer en terme d'*efficacité spectrale*, en $\text{bits.s}^{-1}.\text{Hz}^{-1}$. Cette mesure peut aussi s'exprimer comme le débit d'information binaire divisé par l'occupation spectrale, ce qui nous mène à définir l'efficacité spectrale par :

$$\eta = \frac{D_b}{B}, \quad (1.66)$$

avec D_b le débit d'information binaire de la transmission et B la bande occupée par le signal émis.

1.5.1 Limites théoriques

Pour une bande fréquentielle et un canal donné, il existe une limite fondamentale au débit maximal atteignable tout en garantissant la possibilité théorique d'une communication « fiable » : la capacité. Évoquée en premier par Shannon dans [Sha48] ; [Sha49], la capacité C est propre à un canal, et se définit de telle manière que si D_b est le débit d'information binaire de la transmission, alors il ne sera possible d'obtenir, en réception, une probabilité d'erreur arbitrairement faible que si $D_b \leq C$. Le fait d'être en mesure d'obtenir une probabilité d'erreur arbitrairement faible définit la notion de fiabilité, au sens de Shannon. Cette possibilité implique d'utiliser un système de communication adéquat.

La formule de la capacité d'un canal à bruit additif blanc gaussien est sans doute la plus connue. En effet, même s'il est possible de calculer la capacité de canaux linéaires tels que présentés dans la partie 1.3, la capacité sur canal BABG présente l'avantage d'être un point de comparaison universel, puisqu'il s'agit de la borne supérieure des capacités sur canaux sélectifs [PS08, Chapitre 14]. Elle est donnée par :

$$C_{\text{BABG}} = B \log_2 (1 + \text{RSB}), \quad (1.67)$$

avec B la bande fréquentielle occupée par le signal réel modulé et RSB le rapport signal à bruit défini par l'équation (1.68). Pour obtenir une communication fiable sur canal BABG²³, on cherchera donc à respecter $D_b \leq C_{\text{BABG}}$. La densité spectrale bilatérale du bruit réel autour de la fréquence porteuse étant $N_0/2$, on en déduit que la puissance de bruit dans la bande B occupée par le signal réel modulé est $P_b = N_0 B$. D'autre part, par définition de l'énergie d'un bit, on a $E_b = P_s T_b$, avec P_s la puissance du signal et $T_b = 1/D_b$ le temps d'un bit. On peut donc écrire :

$$\text{RSB} = \frac{P_s}{P_b} = \frac{E_b D_b}{N_0 B} = \eta \frac{E_b}{N_0}, \quad (1.68)$$

23. Bruit Additif Blanc Gaussien

avec η l'efficacité spectrale de la modulation. En injectant (1.68) dans (1.67), on obtient la condition sur le rapport signal à bruit par bit garantissant une transmission fiable à une efficacité spectrale donnée :

$$D_b \leq C_{\text{BAGB}} \Leftrightarrow \eta \leq \log_2 \left(1 + \eta \frac{E_b}{N_0} \right) \Leftrightarrow \frac{E_b}{N_0} \geq \frac{2^\eta - 1}{\eta}. \quad (1.69)$$

En analysant (1.67), on peut distinguer deux régimes distincts. Premièrement, quand $\text{RSB} \rightarrow +\infty$, on a $C_{\text{BAGB}} \approx B \log_2(\text{RSB})$: la capacité évolue de manière logarithmique avec le RSB et de manière linéaire avec l'occupation spectrale B . Le système est alors dit en régime de bande contrainte, ou encore spectralement efficace. À l'inverse, si $\text{RSB} \rightarrow 0$, alors en rappelant que $\log_2(x) = \ln(x)/\ln(2)$ et $\ln(1+x) \rightarrow x$ quand $x \rightarrow 0$, alors $C_{\text{BAGB}} \approx B \cdot \text{RSB} / \ln(2) = P_s / N_0 / \ln(2)$. La capacité varie alors linéairement avec la puissance du signal, et est indépendante de la bande occupée. Le système est alors dit en régime d'énergie contrainte, ou encore énergétiquement efficace. Par extension, on considère généralement que les systèmes présentant une efficacité spectrale supérieure à l'unité sont spectralement efficaces, tandis qu'à l'inverse, on considère un système énergétiquement efficace à partir du moment où son efficacité spectrale est inférieure à l'unité [PS08, Chapitre 4]. Ainsi, dans la suite de ce manuscrit, on cherchera à construire, par l'intermédiaire de la technique FTN, des systèmes spectralement efficaces.

1.5.2 Mises en œuvres pratiques

Rappelons tout d'abord que le lien établi entre efficacité spectrale et rapport E_b/N_0 en (1.69) est une limite théorique qui suppose un système optimal dont on ne connaît pas de réalisation pratique. Il s'agit de la meilleure efficacité spectrale atteignable, tracée en figure 1.16. Afin de comparer des systèmes de communications réels à cette limite, on calculera tout d'abord son efficacité spectrale limite η_0 pour un rapport signal à bruit infini. On déterminera ensuite le plus petit rapport $(E_b/N_0)_0$ permettant d'atteindre une probabilité d'erreur P_0 (très faible, typiquement 10^{-5}). Munis de ces deux informations, nous placerons enfin un point aux coordonnées $((E_b/N_0)_0, \eta_0)$ dont on pourra alors évaluer la distance à la limite théorique sur canal BAGB.

Nous allons maintenant évoquer quelques techniques permettant de s'approcher de la limite de Shannon. Notons enfin que la possibilité, pour un système de communication, de prendre différentes positions dans le plan efficacité spectrale– E_b/N_0 présente également un intérêt substantiel, puisque cela implique une meilleure capacité d'adaptation aux changements de conditions de canal.

Augmentation de la taille de la constellation

Pour les modulations linéaires, une méthode très simple pour augmenter l'efficacité spectrale consiste à augmenter le nombre de bits par symbole (donc d'augmenter la taille de la constellation, voir figure 1.17). Cependant, à énergie par symbole E_s constante, cela diminue la distance entre les points de la constellation, et augmente par conséquent la sensibilité au bruit, comme nous le verrons dans la partie 2.3. Par conséquent, pour obtenir une probabilité d'erreur P_0

1.5. Communications à haute efficacité spectrale

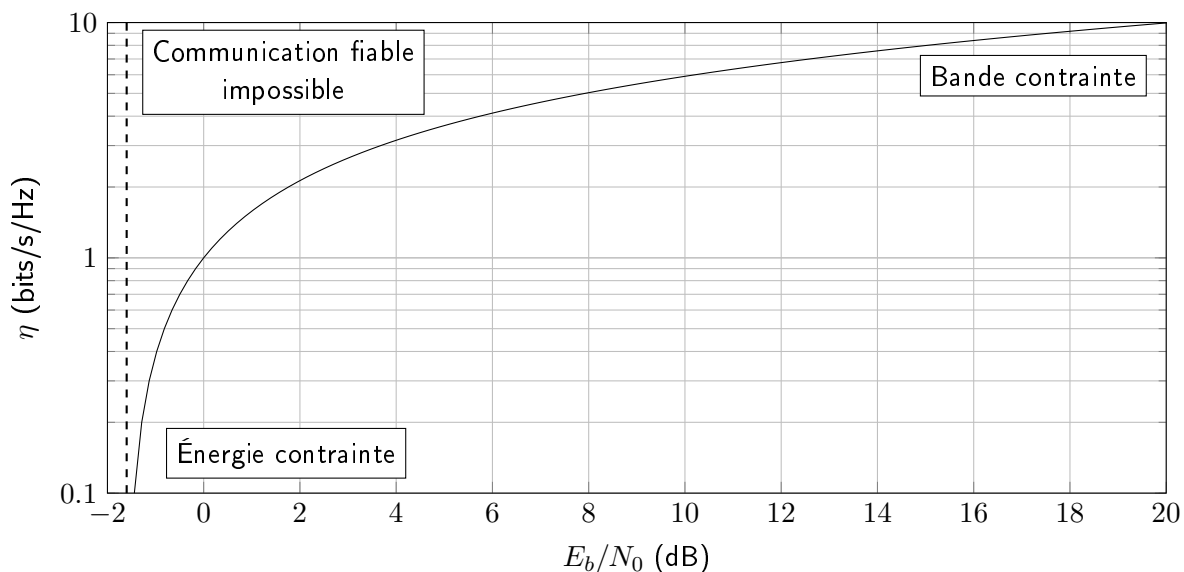


FIGURE 1.16 – Efficacité spectrale maximale d’une transmission sur canal BABG en fonction du rapport E_b/N_0 en dB. On remarque que pour $E_b/N_0 < \ln(2)$ (environ -1.6 dB, trait en pointillés) il n’est en aucun cas possible d’établir une communication fiable.

identique en augmentant la taille de la constellation, il faut également augmenter le rapport E_b/N_0 .

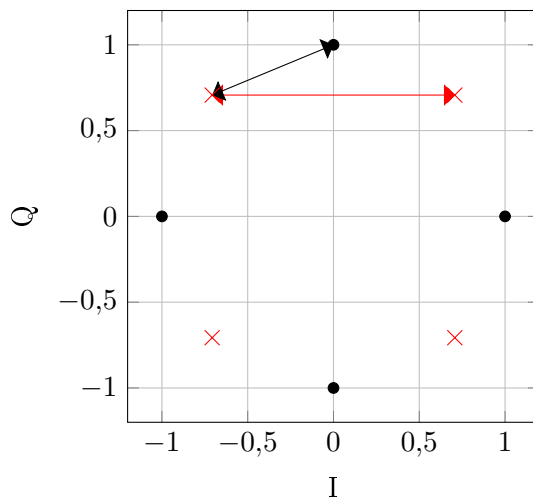


FIGURE 1.17 – Passage d’une constellation QPSK²⁴ (croix) à une 8-PSK (croix et points). On remarque que la distance entre les symboles diminue en augmentant le nombre de bits par symbole, ce qui entraîne une diminution de la résilience au bruit.

Cette méthode n’augmente que très peu la complexité de l’émetteur et du récepteur. Ce-

pendant, en ne permettant qu'une variation d'efficacité spectrale par pas de $1 \text{ bits.s}^{-1}.\text{Hz}^{-1}$ au mieux, et vu les écarts de rapport signal à bruit que cela nécessite (voir figure 1.18) cette technique manque de flexibilité.

Codage canal

Une méthode plus sophistiquée et plus efficace en terme de distance à la limite de Shannon consiste à utiliser du codage canal, et en particulier de « bons codes ». Le lien entre codage et capacité est en effet extrêmement étroit. Cela est bien mis en évidence par le fait que Shannon utilise des codes aléatoires²⁵ dans la démonstration de l'existence de la capacité [Sha48].

Le fait d'utiliser un codage correcteur d'erreur dans le cadre de communications à haute efficacité spectrale peut paraître contre-intuitif au premier abord. En effet, en ajoutant de la redondance dans un schéma de transmission donné, on diminue le nombre de bits utiles (par opposition aux bits de redondance) transmis par unité de temps. Néanmoins, l'ajout de redondance permet généralement de diminuer la probabilité d'erreur à un E_b/N_0 donné. Par conséquent, en fiabilisant la transmission à un rapport signal à bruit plus faible, ils rendent possible l'utilisation de constellations d'ordre élevé à des rapports signal à bruit atteignables en pratiques.

Il existe deux principales familles de codes utilisées en pratique : les codes en blocs et les codes convolutifs. Les premiers sont plutôt adaptés aux transmissions par blocs, tandis que les seconds se prêtent plus naturellement à la protection d'un flot continu de données. De manière plus rigoureuse, les codes en blocs agissent sur un vecteur ligne $\mathbf{m}$ de K bits en utilisant une matrice d'éléments binaires dite « génératrice » $\mathbf{G}$ de taille $K \times N$, de manière à former un mot de code $\mathbf{c}$ de longueur $N > K$ bits défini par :

$$\mathbf{c} = \mathbf{mG}. \quad (1.70)$$

En plus de leur capacité de correction d'erreur, ils sont aussi beaucoup utilisés dans le cadre de contrôle d'intégrité (détection d'erreur). Il existe de nombreuses sous-familles de codes en blocs (codes cycliques, codes produits, codes de Reed-Solomon, etc.), possédant différentes propriétés et pour lesquels différents algorithmes de réception sont à mettre en œuvre. Notons que les codes en blocs utilisés en pratique possèdent une taille des mots de codes N généralement assez élevée, de telle manière qu'il n'est généralement pas possible d'utiliser un algorithme de décodage optimal sans se heurter à des problèmes de complexité d'implémentation. Ce type de code ne sera pas utilisé dans la suite de ce manuscrit. Un codeur convolutif quant à lui associera, à chaque bit m_k mis à son entrée à l'instant k , L_g bits en sortie notés $c_{k,0}$ à c_{k,L_g-1} . On définit L_g réponses impulsionnelles binaires possédant L coefficients : $\mathbf{g}^{(i)} = (g_0^{(i)} \dots g_{L-1}^{(i)})$ ($i \in [0; L_g - 1]$),

25. Admettons que la taille des séquences binaires en entrée d'un codeur soit de N_{in} éléments, et que celle des séquences binaires à sa sortie (mots de code) soit de $N_{\text{out}} > N_{\text{in}}$ éléments, alors un codeur aléatoire sera construit en associant chaque séquence binaire d'entrée possible à une séquence binaire de taille N_{out} choisie aléatoirement parmi les $2^{N_{\text{out}}}$ possibles.

1.5. Communications à haute efficacité spectrale

de manière à ce que la relation d'entrée-sortie de l'encodeur s'écrive :

$$c_{k,i} = \sum_{l=0}^{L-1} g_l^{(i)} m_{k-l}, \quad i \in [0; L_g - 1], \quad (1.71)$$

où la sommation s'effectue modulo 2. La ressemblance avec l'opération de convolution discrète vaut à cette famille de codes son nom. Notons enfin que les réponses impulsionnelles sont souvent données en représentation octale. Par exemple, le code convolutif (5, 7) possède deux réponses impulsionnelles : $\mathbf{g}^{(0)} = (101)$ et $\mathbf{g}^{(1)} = (111)$. De nos jours, le décodage des codes convolutifs est quasi-systématiquement effectué par l'algorithme de Viterbi (éventuellement à entrée et/ou sortie souple) [Vit67] ; [Bat87]. Pour les deux familles, on définit le rendement de codage par le rapport du nombre de bits dans le message source et dans le mot de code : $R = K/N$ pour un code en bloc, et $R = 1/L_g$ pour un code convolutif.

Il est à noter que tous les codes correcteurs d'erreurs ne se valent pas. En particulier, ils ne présentent pas le même pouvoir de correction pour un rendement de codage donné. De plus, même un très bon code, pour être utilisable, doit pouvoir être décodé de manière efficace par le récepteur. C'est-à-dire que l'on doit pouvoir disposer de décodeurs de complexité raisonnable. Par exemple, on ne connaît pas d'algorithme de décodage efficace pour les codes aléatoires évoqués plus haut. Ainsi, ce n'est que relativement récemment (1991) que des codes correcteurs d'erreurs d'implémentation pratiques et approchant de près la capacité du canal BABG pour une modulation BPSK (à 0.5 dB contre au moins 3 dB pour les meilleures solutions d'alors) ont été développés [Ber+07, Chapitre 4] [BGT93]. Ces codes, plutôt que de former une nouvelle famille à proprement parler, consistent en une concaténation de codes constituants classiques (concaténation parallèle ou série pour les codes convolutifs, et codes produits pour les codes en blocs), et leur décodage s'effectue en associant les décodeurs (qui doivent obligatoirement être à entrée/sortie pondérée – SISO ²⁶) des codes constituants en un système bouclé [Ber+07, Chapitre 6]. Peu de temps après cette découverte, les codes LDPC et le décodage par propagation de croyance, qui permettent eux aussi d'approcher de près la capacité du canal BABG et développés par Gallager en 1963, ont été redécouverts [Gal63] ; [MN95]. On observe bien sur la figure 1.18 la proximité à la limite de Shanon permise par ces deux types de codes, comparés ici à un code convolutif classique.

Les techniques de décodage des turbo-codes et des LDPC ont depuis été étendues pour être utilisables avec des constellations de plus de deux symboles. Pour cela, on utilisera généralement soit le principe de « modulations codées pragmatiques » établi par Viterbi pour les codes convolutifs et dont l'approche est de séparer la démodulation du décodage, soit le principe de modulations codées en treillis, introduit par Ungerboeck et dont l'approche est d'optimiser conjointement le codage canal et la conversion bit à symbole (qui peut être vue comme une forme de codage) [Vit+89] ; [GGB94] ; [Ung82] ; [RW95] [Ber+07, Chapitre 10].

Enfin, signalons que la technique du « poinçonnage » consistant simplement à ne pas envoyer certains bits de redondance du mot de code, permet de faire varier relativement finement l'effi-

26. *Soft Input-Soft Output* (Entrée/sortie pondérée)

cacité spectrale du système en fonction des conditions du canal, mais implique une diminution de la résilience au bruit.

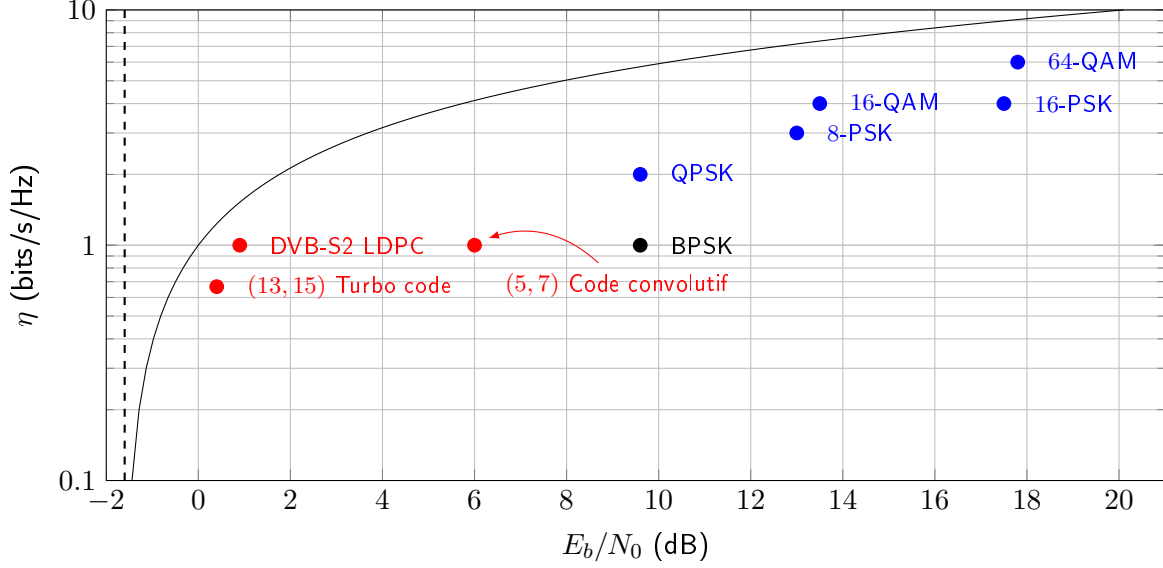


FIGURE 1.18 – Comparaison de l'effet sur efficacité spectrale de différentes techniques de transmission, pour une probabilité d'erreur cible $P_e = 10^{-5}$. Les techniques de codage canal ont ici été appliquées à une modulation QPSK, avec des blocs de 16384 bits pour le code convolutif et le turbo-code, et de 32400 bits pour le code LDPC. Le code LDPC est issu de [ETS14].

Densification de la signalisation

Classiquement, pour éviter d'avoir à compenser le terme d'interférence dans (1.58), on choisit de respecter le critère de Nyquist. Cela limite néanmoins la densité du système, et par conséquent, pour une constellation donnée, son efficacité spectrale. Illustrons cela en monoporteuse, puis en multiporteuse. En monoporteuse, le respect du critère de Nyquist implique que le rythme symbole R soit inférieur ou égal à la bande B occupée par le signal. En notant N_b le nombre de bits par symboles et D_b le débit d'information binaire, on a :

$$R \leq B \quad \Leftrightarrow \quad \frac{D_b}{N_b} \leq B \quad \Leftrightarrow \quad \eta \leq N_b. \quad (1.72)$$

Le principe du FTM, en monoporteuse, est de permettre $R > B$, et donc $\eta > N_b$. En multiporteuse, en notant M le nombre de sous-porteuses et T_0 le temps d'un symbole multiporteuse, le débit d'information binaire s'écrit $D_b = N_b \cdot M / T_0$. D'autre part, comme nous l'avons vu dans 1.4.2.2, en supposant un grand nombre de porteuses, et en notant F_0 l'écart inter-porteuses, alors la bande occupée par le signal multiporteuse est bien approchée par $B = M F_0$. On peut

1.6. Conclusion

donc exprimer l'efficacité spectrale en multiporteuse par :

$$\eta = \frac{D_b}{B} = \frac{N_b}{T_0 F_0} = N_b \rho, \quad (1.73)$$

avec $\rho = 1/(T_0 F_0)$ la densité du système multiporteuse. Comme nous l'avons vu dans la partie 1.4.2.1, et à l'instar du comportement en monoporteuse, choisir $\rho > 1$ permet d'augmenter l'efficacité spectrale au prix d'une interférence entre impulsions de mise en forme qui devra être compensée par le récepteur. On remarque que ρ étant réel, l'on dispose de toute latitude pour ajuster finement l'efficacité spectrale à la hausse en fonction des conditions du canal, à condition de pouvoir correctement compenser l'interférence.

Évoquons ici le fait que le FTN est un cas particulier de signalisation à réponse partielle PRS²⁷. La signalisation à réponse partielle regroupe toutes les techniques consistant à introduire de manière délibérée de l'interférence entre impulsions de mise en forme. Cependant, le terme PRS est plutôt utilisé lorsque cette introduction d'interférence a pour but d'obtenir des propriétés spectrales particulières. On pourra trouver plus de détail sur cette technique dans la référence suivante [KP75].

Enfin, de la même manière qu'en codage canal, la capacité à augmenter l'efficacité spectrale via la technique FTN dépendra aussi de la complexité de démodulation admissible. En particulier, nous verrons dans la partie 2.4.1 que le récepteur FTN optimal possède une complexité rédhibitoire en multiporteuse. Un des principaux objectifs de ce travail de thèse sera donc de proposer des structures sous-optimales proposant un bon compromis entre probabilité d'erreur et complexité algorithmique.

1.6 Conclusion

Dans ce chapitre, nous avons défini le concept de communication comme le moyen de transmettre un message au travers d'un canal. Nous avons ensuite expliqué que les avancées technologiques dans le domaine du traitement numérique des données font des communications numériques l'un des moyens de communication les plus incontournables, faisant l'objet de travaux de recherche encore très actifs.

Nous avons ensuite présenté un modèle classique de chaîne de communication numérique, sur lequel nous travaillerons en ne considérant que les fonctions de codage et de décodage canal, chargées de protéger le message par l'ajout de redondance, de modulation et de démodulation, chargées d'effectuer la conversion entre signal numérique et signal physique, ainsi que le canal.

Ce dernier, imposé par l'application et son contexte de mise en œuvre, a pour effet d'introduire de l'incertitude au niveau du récepteur sur le message envoyé. Nous avons introduit une modélisation déterministe pour les canaux sélectifs en fréquence (LTI), présentant de l'effet Doppler (LFI), et variants dans le temps (LTV) c'est-à-dire présentant à la fois une sélectivité fréquentielle et un décalage Doppler. L'introduction des notions de fonctions propres et de

27. *Partial Response Signaling* (Signalisation à réponse partielle)

fonctions propres approchées pour ces trois types de canaux nous a permis de conclure que les modulations multiporteuses sont bien adaptées à la transmission sur canaux LTI et LTV.

Les modulations multiporteuses ont pour effet de répartir les symboles élémentaires de la transmission dans le plan temps-fréquence. Afin de permettre leur analyse, nous avons donc introduit la théorie des frames, empruntée à la communauté de l'analyse temps-fréquence, et en particulier les frames de Gabor. Par la suite, un modèle à temps continu ainsi qu'un modèle à temps discret de ce type de modulation ont été donnés. Ce modèle met en avant l'existence d'une densité, paramètre permettant de choisir à quel point les symboles élémentaires de la transmission sont comprimés dans le plan temps-fréquence. Ce paramètre permet donc d'agir sur l'efficacité spectrale de la transmission. En particulier, choisir une densité sur-unitaire permet d'augmenter l'efficacité spectrale, au prix d'une complexité de réception plus élevée. Cela correspond à une transmission au-delà de la cadence de Nyquist (FTN).

Enfin, en introduisant la capacité de Shannon, nous avons mis en avant que l'efficacité spectrale d'une transmission, à conditions de canal données, ne peut pas être choisie arbitrairement élevée. Nous avons ensuite expliqué que l'augmentation de la taille de la constellation permet d'augmenter l'efficacité spectrale, au prix d'une sensibilité accrue au bruit. À l'inverse, l'ajout de redondance permet de diminuer la sensibilité au bruit, au prix d'une baisse de l'efficacité spectrale. Une autre manière d'agir sur l'efficacité spectrale est de modifier la densité de la transmission. En particulier, nous verrons dans le chapitre suivant que la technique FTN permet même, sous certaines conditions, d'augmenter l'efficacité spectrale jusqu'à un certain point sans pour autant nécessiter de meilleurs rapports signal à bruit.

Transmissions au-delà de la cadence de Nyquist

Sommaire

2.1	Introduction	45
2.2	FTN monoporteuse et FTN multiporteuse	46
2.2.1	FTN monoporteuse	46
2.2.2	FTN multiporteuses	47
2.3	Limite de Mazo	48
2.4	Techniques de détection	51
2.4.1	Égalisation	51
2.4.2	Turbo-égalisation	56
2.5	Conclusion	62

2.1 Introduction

La limite de Nyquist a originellement été introduite pour permettre des transmissions à bande limitée sans interférence entre impulsions de mise en forme après échantillonnage au temps symbole¹, sur canal à BABG [Nyg28a]. En monoporteuse, elle stipule que le débit symbole $R = 1/T_0$ (avec T_0 l'espacement temporel entre symbole) doit impérativement être inférieur ou égal à la bande bilatérale occupée par l'équivalent complexe en bande de base du signal à transmettre : $R \leq B$. Comme nous l'avons vu dans la partie 1.4.2, une telle limite existe également pour les modulations multiporteuses, et porte à la fois sur l'espacement en temps T_0 et sur l'espacement en fréquence F_0 entre impulsions de mise en forme : $\rho = 1/(F_0 T_0) \leq 1$.

Le critère de Nyquist est systématiquement respecté dans les standards de communication actuels, et est rarement remis en question dans la mesure où il simplifie sensiblement les relations d'entrée/sortie. Pourtant, dès 1975, Mazo montre que dans certaines conditions, il est possible de transmettre à un débit symbole R_{Mazo} 25% plus élevé que la limite de Nyquist R sans pénalité en termes de TEB², mais à condition d'utiliser un détecteur optimal au sens de la minimisation de la probabilité d'erreur [Maz75] ; [For73]. À cette époque, bien que l'algorithme de Viterbi ait

1. Lorsque le critère de Nyquist est respecté, les impulsions de mise en forme peuvent interférer dans le domaine analogique, mais pas après échantillonnage au temps symbole.

2. Taux d'Erreur Binaire

déjà été découvert (en 1967 [Vit67]) et étudié [For73], une détection au maximum de vraisemblance d'un signal présentant de l'interférence entre symbole restait infaisable en pratique. Plus tard, Mazo dira avoir effectué cette étude par simple curiosité [Das+14, Préface]. Cependant, avec l'augmentation des capacités de calcul, et devant des ressources spectrales toujours plus contraintes, la technique FTN a été redécouverte au début des années 2000. Les travaux de Mazo ont alors été étendus à différentes impulsions de mise en forme et différents alphabets de modulation et, plus important en ce qui nous concerne, aux modulations multiporteuses [RA05] ; [RA08]. De même, de nombreux efforts ont été déployés pour proposer des techniques de détection à complexité réduite [TPB16] ; [Lah17] ; [DRO10]. Nous verrons cependant que, dans le cas du FTN multiporteuse, la plupart des algorithmes de détection optimaux ou quasi-optimaux sont algorithmiquement trop demandeurs pour permettre une implémentation pratique.

Ce chapitre est organisée de la manière suivante. Premièrement nous aborderons les différences fondamentales entre FTN monoporteuse et FTN multiporteuse en termes de modèles. Nous verrons également que l'hypothèse de bruit blanc en sortie du récepteur linéaire est impossible à satisfaire dans le cas général dans ce contexte. Nous aborderons également deux modélisations distinctes pour un système FTN multiporteuses : la première ayant été introduite par Rusek et Anderson [RA05], et la seconde permettant de s'appuyer plus naturellement sur la théorie des frames. Ensuite, nous définirons le concept de limite de Mazo, en monoporteuse et en multiporteuse, et présenterons les principaux travaux de recherche sur la théorie des systèmes FTN. Les deux parties précédentes permettent d'expliquer les besoins exigeants en termes d'optimalité de la détection des systèmes FTN, et posent les bases des modèles mathématiques permettant de développer différentes techniques de détection, optimales et sous-optimales, dans la troisième partie. La question de la détection et du décodage conjoint des transmissions FTN codées, via des techniques de turbo-égalisation, y sera également traitée. Nous y développerons en détail les raisons qui font que la plupart des schémas de détection en FTN monoporteuse ne sont pas envisageables en FTN multiporteuse, ce qui permettra la transition vers le chapitre suivant, dont l'objectif est de proposer et d'étudier des structures de réception présentant un bon compromis performance-complexité.

2.2 FTN monoporteuse et FTN multiporteuse

Cette section a pour objectif de montrer les différences fondamentales de modélisation du FTN monoporteuse et multiporteuse. Elle servira de base à la section suivante, où nous discuterons des techniques de détection envisageables.

2.2.1 FTN monoporteuse

En monoporteuse, une façon simple et très utilisée d'obtenir un système FTN consiste à partir d'un système respectant le critère de Nyquist au rythme symbole T_0 , et de simplement accélérer ce rythme symbole par un facteur $0 < \tau < 1$. La nouvelle durée symbole obtenue s'écrit alors $T'_0 = \tau T_0 < T_0$.

2.2. FTN monoporteuse et FTN multiporteuse

Considérons $g(t) \in \mathcal{L}_2(\mathbb{R})$ le filtre d'émission, $\check{g}(t) \in \mathcal{L}_2(\mathbb{R})$ le filtre de réception et $h(t) = (g * \check{g})(t)$. Typiquement, $g(t)$ est un filtre en « racine de Nyquist » tel qu'un racine de cosinus surélevé, et $\check{g}(t)$ le filtre adapté à $g(t)$ (permettant de maximiser le RSB³), de sorte que $h(t)$ permette de respecter le critère de Nyquist à la cadence T_0 .

Avec ces notations, en considérant l'émission d'une séquence de symboles $\mathbf{c} = \{c_l\}_{l \in \mathbb{Z}}$ sur canal BABG, le signal reçu après filtrage de réception s'écrit :

$$r(t) = \sum_l c_l h(t - l\tau T_0) + (b * \check{g})(t), \quad (2.1)$$

avec $b(t)$ un bruit blanc gaussien complexe circulaire. Un modèle à temps discret s'obtient en échantillonnant tous les τT_0 :

$$r[n] = \sum_l c_l h[n - l] + (b * \check{g})[n] = s[n] + (b * \check{g})[n] = \tilde{c}_n. \quad (2.2)$$

À des fins de simplification, on considèrera que $h[n]$ est à RIF⁴. D'autre part, $h[n]$ possède forcément plus d'un coefficient non-nul (dans le cas contraire, il n'y aurait pas d'interférence entre impulsions de mise en forme et le critère de Nyquist ne serait donc pas violé).

L'expression du signal reçu (2.2) prend la même forme que lors d'une transmission monoporteuse à la cadence de Nyquist sur canal LTI, à la différence que ce « canal » $h[k]$ est ici parfaitement connu. Une autre différence concerne le traitement du terme de bruit $(\check{g} * b)[n]$. En effet, lors d'une transmission sur canal LTI, on filtre généralement le signal $r[n]$ par un filtre blanchisseur $f[n]$, de sorte que le bruit filtré après échantillonnage $(f * \check{g} * b)[n]$ soit blanc. Cependant, dans le cadre d'une transmission FTN, le filtre blanchisseur est généralement mal défini, ce qui demandera d'adapter certaines des techniques de détection issues de l'égalisation de canaux LTI [ARÖ13].

2.2.2 FTN multiporteuses

Sur le même modèle que (2.1), plusieurs auteurs se basent sur un système multiporteuse orthogonal à $\rho = 1$, et introduisent deux facteurs d'accélération : un en temps $0 < \tau < 1$, et un en fréquence $0 < \nu < 1$ [DRO10] ; [ARÖ13] ; [Lin+15]. En reprenant les notations de la partie 1.4, le signal vu par le démodulateur s'écrit :

$$r(t) = \sum_{(m,n) \in \Lambda} c_{m,n} g(t - n\tau T_0) e^{j2\pi m\nu F_0 t} = \sum_{(m,n) \in \Lambda} c_{m,n} g_{m\nu, n\tau}(t). \quad (2.3)$$

Les prototypes d'émission $g(t)$ et de réception $\check{g}(t)$ sont choisis tels que, pour $\tau = \nu = 1$, on ait :

$$\langle \check{g}_{p,q}; g_{m,n} \rangle = \delta_{m-p} \delta_{n-q}. \quad (2.4)$$

En revanche, à partir du moment où $\rho' = 1/(\tau\nu) > 1$, alors $\langle \check{g}_{p\nu, q\tau}; g_{m\nu, n\tau} \rangle \neq \delta_{m-p} \delta_{n-q}$ par construction (voir sous-section 1.4.1). Cette approche a cependant pour désavantage que les

3. Rapport signal à bruit

4. Réponse impulsionnelle finie

prototypes d'émission et de réception ne peuvent pas jouir d'une bonne localisation temps-fréquence. Cela est une conséquence directe du théorème de Balian-Low, les prototypes devant respecter (2.4) quand $\rho' = 1$. Pour contourner ce problème, il est possible d'appliquer le même principe en partant d'un système OFDM/OQAM orthogonal.

Cependant, dans le cas FTN, puisque la contrainte d'orthogonalité entre impulsions de mise en forme est relâchée, il est possible de se soustraire au théorème de Balian-Low. Il est donc possible, par un choix judicieux de $g(t)$, d'obtenir une bonne localisation temps-fréquence avec $\rho > 1$. Ainsi, notre approche dans ce manuscrit sera de considérer un signal multiporteuse de la forme suivante :

$$s(t) = \sum_{(m,n) \in \Lambda} c_{m,n} g(t - nT_0) e^{j2\pi m F_0 t}. \quad (2.5)$$

avec $\rho = 1/(F_0 T_0) > 1$, pour lequel on se laissera toute latitude dans le choix du prototype.

Sur canal BABG, l'expression des symboles estimés se déduit aisément de (1.58) :

$$\tilde{c}_{p,q} = \sum_{(m,n) \in \Lambda} c_{m,n} \langle \check{g}_{p,q}; g_{m,n} \rangle + \langle \check{g}_{p,q}; b \rangle, \quad (2.6)$$

que l'on réécrit comme suit, de manière à simplifier certains développements de la partie 2.4

$$\tilde{c}_{p,q} = \sum_{(m,n) \in \Lambda} c_{m,n} \sum_l \check{g}_{p,q}^*[l] g_{m,n}[l] + \langle \check{g}_{p,q}; b \rangle \quad (2.7)$$

$$= \sum_{(m,n) \in \Lambda} c_{m,n} \sum_l \check{g}^*[l - qN] g[l - nN] e^{j2\pi \frac{m-p}{M} l} + \langle \check{g}_{p,q}; b \rangle \quad (2.8)$$

$$= \sum_{(m,n) \in \Lambda} c_{m,n} \sum_l \check{g}^*[l] g[l - (n - q)N] e^{j2\pi \frac{m-p}{M} l} e^{j2\pi \frac{m-p}{M} qN} + \langle \check{g}_{p,q}; b \rangle \quad (2.9)$$

$$= \sum_{(m,n) \in \Lambda} c_{m,n} A_{\check{g},g}[m - p, n - q] e^{j2\pi \frac{m-p}{M} qN} + \langle \check{g}_{p,q}; b \rangle, \quad (2.10)$$

avec $A_{\check{g},g}[m, n] = \sum_l \check{g}^*[l] g[l - nN] e^{j2\pi \frac{m}{M} l}$. On remarque ici une certaine ressemblance avec le cas d'une transmission multiporteuse sur canal radiomobile (1.59). De la même manière qu'en monoporteuse, on utilisera donc des techniques d'égalisation afin de traiter l'interférence. Notons enfin qu'ici aussi, le terme de bruit filtré $\langle \check{g}_{p,q}; b \rangle$ n'est pas blanc, ce qui appellera à quelques modifications des égaliseurs.

2.3 Limite de Mazo

Dans cette section, nous détaillons une partie des justifications théoriques qui font du FTN une technique de transmission intéressante.

Définissons tout d'abord le concept de distance euclidienne minimale entre signaux $d_{\min}$, qui nous permettra d'évaluer les performances asymptotiques en termes de probabilité d'erreur d'un format de modulation donné, dans le cas où un détecteur optimal au sens de la minimisation

2.3. Limite de Mazo

de la probabilité d'erreur est utilisé. Soient deux séquences de symboles $\mathbf{c} = \{c_{m,n}\}_{(m,n) \in \Lambda}$ et $\mathbf{c}' = \{c'_{m,n}\}_{(m,n) \in \Lambda}$ telles que $\mathbf{c} \neq \mathbf{c}'$, et leurs signaux modulés associés $s(t)$ et $s'(t)$, alors la distance euclidienne minimale entre signaux est définie par :

$$d_{\min}^2 = \min_{\mathbf{c}, \mathbf{c}' \in \mathcal{A}, \mathbf{c} \neq \mathbf{c}'} \|s(t) - s'(t)\|^2 = \min_{\mathbf{c}, \mathbf{c}' \in \mathcal{A}, \mathbf{c} \neq \mathbf{c}'} \left\| \sum_{(m,n) \in \Lambda} (c_{m,n} - c'_{m,n}) g_{m,n}(t) \right\|^2 \quad (2.11)$$

$$= \min_{\epsilon \in \mathcal{E} \setminus \{0\}} \left\| \sum_{(m,n) \in \Lambda} \epsilon_{m,n} g_{m,n}(t) \right\|^2, \quad (2.12)$$

où $\mathcal{E} = \left\{ \epsilon = \{c_{m,n} - c'_{m,n}\}_{(m,n) \in \Lambda} \mid c_{m,n}, c'_{m,n} \in \mathcal{A} \quad \forall (m,n) \in \Lambda \right\}$. Remarquons que dans le cas où la famille d'émission $\mathbf{g}$ est orthonormale (ce qui n'est possible que lorsque le critère de Nyquist est respecté), alors la distance minimale ne dépend que de l'alphabet de modulation $\mathcal{A}$ utilisé :

$$d_{\min}^2 = \min_{\epsilon \in \mathcal{E} \setminus \{0\}} \left\| \sum_{(m,n) \in \Lambda} \epsilon_{m,n} g_{m,n}(t) \right\|^2 \quad (2.13)$$

$$= \min_{\epsilon \in \mathcal{E} \setminus \{0\}} \int_{-\infty}^{+\infty} \left| \sum_{(m,n) \in \Lambda} \epsilon_{m,n} g_{m,n}(t) \right|^2 dt \quad (2.14)$$

$$= \min_{\epsilon \in \mathcal{E} \setminus \{0\}} \sum_{(m,n) \in \Lambda} \sum_{(m',n') \in \Lambda} \epsilon_{m,n} \epsilon_{m',n'}^* \int_{-\infty}^{+\infty} g_{m,n}(t) g_{m',n'}^*(t) dt \quad (2.15)$$

$$= \min_{\epsilon \in \mathcal{E} \setminus \{0\}} \sum_{(m,n) \in \Lambda} \sum_{(m',n') \in \Lambda} \epsilon_{m,n} \epsilon_{m',n'}^* \langle g_{m',n'}; g_{m,n} \rangle \quad (2.16)$$

$$= \min_{\epsilon \in \mathcal{E} \setminus \{0\}} \sum_{(m,n) \in \Lambda} |\epsilon_{m,n}|^2 \quad (2.17)$$

$$= \min_{\epsilon \in \mathcal{E} \setminus \{0\}} \|\epsilon_{m,n}\|^2. \quad (2.18)$$

Ainsi, une spécificité des modulations FTM est que la distance euclidienne minimale entre signaux est lié à l'alphabet de modulation d'une part, mais aussi au prototype d'émission utilisé. Le concept de distance euclidienne minimale est fondamental dans l'étude du comportement d'une modulation en termes de probabilité d'erreur atteignable. En effet, en présence de bruit blanc gaussien, il a été montré par Forney que la probabilité d'erreur symbole P_e atteinte à l'aide du détecteur au maximum de vraisemblance⁵ est bornée par :

$$K_L Q \left(\frac{d_{\min}}{\sqrt{2N_0}} \right) \leq P_e \leq K_U Q \left(\frac{d_{\min}}{\sqrt{2N_0}} \right), \quad (2.19)$$

5. Notons que le détecteur au maximum de vraisemblance minimise la probabilité d'erreur en présence de symboles IID.

où $Q(\cdot)$ est la fonction de répartition complémentaire d'une variable aléatoire normale centrée réduite, tandis que K_L et K_U sont deux constantes réelles [For73]. Notons que le rapport K_L/K_U tend vers 1 lorsque le rapport E_b/N_0 tend vers l'infini, de telle manière que (2.19) est très utilisé pour évaluer le comportement asymptotique du détecteur optimal.

La distance euclidienne minimale est donc un paramètre déterminant pour les performances d'un système de communication. Or, dans son article de 1975, Mazo a montré que dans le cas monoporteuse, en utilisant une impulsion de mise en forme en sinus cardinal de manière à respecter le critère de Nyquist à $\tau = 1$ et un alphabet de modulation BPSK⁶, alors il est possible de choisir $\tau > 0.8$ sans modifier la distance euclidienne entre les signaux émis [Maz75]. Ce résultat a par ailleurs été démontré dans [ML88]. Soit un modulateur monoporteuse fonctionnant à la cadence de Nyquist, et notons d_0 la distance euclidienne minimale entre les signaux engendrés par ce dernier. Alors le facteur d'accélération minimal tel que la distance euclidienne minimale entre symboles reste d_0 généralement nommé « limite de Mazo », est définie par :

$$\tau_{\text{Mazo}} = \min_{0 < \tau \leq 1} \{\tau | d_{\min} = d_0\}. \quad (2.20)$$

Concrètement, cela signifie que tant que la limite de Mazo n'est pas atteinte, il est possible d'augmenter le débit symbole du système orthogonal sans augmenter la probabilité d'erreur. Cependant, comme nous le verrons en section 2.4, cela implique une complexification significative du récepteur, ce qui explique que la technique FTN ait été peu considérée jusqu'au début des années 2000. En 2003, Liveris observe par simulation l'existence d'une limite de Mazo en utilisant des impulsions de mise en forme en racine de cosinus surélevé [LG03]. En 2008, Rusek propose un algorithme permettant la recherche de $d_{\min}$ de manière efficace, et montre l'existence d'une limite de Mazo pour des alphabets de modulation d'ordre supérieur (en utilisant une mise en forme en racine de cosinus surélevé) [RA08].

Enfin, en 2005, Rusek montre l'existence d'une limite de Mazo pour les modulations multiporteuses filtrées à base de prototypes en racine de cosinus surélevés. Il indique en outre que le FTN multiporteuse semble permettre de meilleurs gains en efficacité spectrale que le FTN monoporteuse [RA05]. Considérons le formalisme FTN multiporteuse de (2.3), notons d_0 la distance minimale du système orthogonal de départ, alors la limite de Mazo dite « à deux dimensions » est définie par l'ensemble des couples de facteurs d'accélération temporels et fréquentiels $(\tau_{\text{Mazo}}, \nu_{\text{Mazo}})$ tels que $d_{\min} = d_0$, et tels que $d_{\min} < d_0$ si $(\tau, \nu) = (\tau_{\text{Mazo}}, \nu_{\text{Mazo}} - \zeta)$ ou $(\tau, \nu) = (\tau_{\text{Mazo}} - \Psi, \nu_{\text{Mazo}}) \forall \zeta, \Psi > 0$.

Comme nous l'avons évoqué dans la chapitre 1 (partie 1.5.2), la densification de la signalisation, tant qu'elle ne dépasse pas la limite de Mazo, a le potentiel d'augmenter l'efficacité spectrale des modulations basées sur le respect du critère de Nyquist, sans pour autant nécessiter des rapports signaux-à-bruit plus importants. Cela signifie qu'en théorie, tous les points de la figure 1.18 pourraient être décalés vers de meilleures valeurs d'efficacité spectrale. Cependant, pour tirer le meilleur parti des modulations FTN, il faut utiliser un détecteur optimal au sens de la minimisation de la probabilité d'erreur. Dans la partie suivante, nous donnerons

6. *Binary Phase Shift Keying* (Modulation par saut de phase binaire)

2.4. Techniques de détection

des algorithmes implémentant de tels détecteurs. Des solutions sous-optimales seront également présentées, et nous finirons par une discussion sur les compromis performances/complexité de ces différentes solutions.

2.4 Techniques de détection

Comme nous l'avons vu dans la partie précédente, l'effet d'une transmission FTM se modélise de manière similaire à l'effet d'un canal sur une transmission à la cadence de Nyquist. Il est ainsi naturel de se tourner vers les techniques d'égalisation, développées dans ce contexte. De plus, considérant les possibilités ouvertes par les techniques de décodages probabilistes et itératives des codes concaténés (voir partie 1.5), et vu l'omniprésence de codes correcteurs et/ou détecteurs d'erreurs dans les systèmes de communication numérique, nous présenterons également les techniques de turbo-égalisation, dont l'approche initiale est de considérer une transmission codée sur canal, de modéliser ce dernier comme un code, et d'appliquer des techniques de turbo-décodage sur la concaténation du code correcteur d'erreur et du canal [Dou+95].

2.4.1 Égalisation

2.4.1.1 Égalisation au maximum a posteriori

Intéressons-nous tout d'abord à l'égalisation d'une séquence au maximum *a posteriori* (MAP⁷). Il s'agit d'une approche probabiliste dans laquelle la séquence de symboles émis $\mathbf{c} \in \mathcal{A}^{K \times M}$ (avec K le nombre de symboles multiporteuses et M le nombre de sous-porteuses, ce qui implique $\Lambda = [0; M-1] \times [0; K-1]$) est modélisée comme une variable aléatoire. Le but est alors de maximiser la probabilité que la séquence de symboles égalisés $\tilde{\mathbf{c}} \in \mathcal{A}^{K \times M}$ soit égale à la séquence émise :

$$\tilde{\mathbf{c}} = \underset{\tilde{\mathbf{c}} \in \mathcal{A}^{K \times M}}{\operatorname{argmax}} \mathbb{P} \{ \mathbf{c} = \tilde{\mathbf{c}} | \mathbf{Y} = \tilde{\mathbf{c}} \}, \quad (2.21)$$

où $\tilde{\mathbf{c}} \in \mathbb{C}^{K \times M}$, l'observation, est la séquence de symboles estimée après filtrage, donnée par (2.2) en monoporteuse et (2.10) en multiporteuse, et $\mathbf{Y}$ est la variable aléatoire associée. Notons que lorsque les symboles sont IID⁸, (2.21) se réduit à trouver $\tilde{\mathbf{c}}$ selon le critère du maximum de vraisemblance (MV⁹).

Dans le cas général, la résolution du problème en (2.21) requiert une recherche exhaustive, dont la complexité est en $\mathcal{O}(|\mathcal{A}|^{K \times M})$. Une telle complexité est inenvisageable en pratique, sauf pour une transmission de peu de symboles (K et M faibles), avec une constellation à faible nombre d'états (telle qu'une BPSK).

Cependant, dans le cas monoporteuse ($M = 1$), puisque $h[k]$ est à RIF, alors le système peut être représenté sous la forme de registres à décalages et d'une fonction combinant les sorties de ces registres (voir figure 2.1c). À chaque instant n , le système est donc dans un état $r_n \in \mathcal{R}$

7. Maximum *A Posteriori*

8. Indépendants et Identiquement Distribués

9. Maximum de Vraisemblance

représenté par la valeur stockée dans les registres ($\mathcal{R}$ étant l'ensemble des états que peut prendre le système), et l'état suivant, sachant l'état courant, ne dépend que de l'entrée du système à l'instant $n + 1$: $\mathbf{c}_{n+1} = \{c_{m,n}\}_{m \in [0;M-1]}$. Le système se modélise donc comme une chaîne de Markov, laquelle peut aussi être représentée par son diagramme d'état (voir figure 2.1b). Une autre représentation, faisant intervenir le temps et introduite par Forney [For73], est le « treillis » (voir figure 2.1d).

À partir de cette représentation, si on arrive à définir une métrique (ou distance) pour chaque branche de ce treillis, alors le problème (2.21) se résume à une recherche du plus court chemin dans un graphe dirigé et acyclique : un treillis. Ceci est très exactement le fonctionnement de l'algorithme de Viterbi [Vit67] ; [For73]. Notons qu'il existe deux métriques selon le type de bruit :

- lorsque le bruit est blanc (ou s'il existe un filtre blanchisseur), on utilise le modèle de Forney, pour lequel les métriques de branches sont données dans [For73, Partie III]¹⁰ ;
- lorsque le bruit est coloré, on utilise le modèle d'Ungerboeck, pour lequel les métriques de branches sont données dans [Ung74, Équation 23] [RCS15, Équation 9].

Notons que le modèle d'Ungerboeck a jusque-là suscité peu d'intérêt dans la communauté des communications numériques. Il est néanmoins incontournable pour les communications FTN, où le filtre blanchisseur est généralement mal défini [ARÖ13].

L'algorithme de Viterbi analyse chacune des branches du treillis, sa complexité est donc linéaire avec la longueur de la séquence et le nombre de transitions possible entre deux instants. En pratique, le nombre de branches maximal entre deux instants étant le carré du nombre d'états, la valeur de complexité généralement retenue est $\mathcal{O}(K|\mathcal{R}|^2)$. En monoporteuse, le nombre d'états est lié à la taille de la réponse impulsionnelle de $h[k]$, notée L_h , et à la taille de l'alphabet de modulation par $|\mathcal{R}| = |\mathcal{A}|^{L_h-1}$, de telle sorte que la complexité de l'algorithme de Viterbi, appliqué au problème d'égalisation, est en $\mathcal{O}(K|\mathcal{A}|^{2(L_h-1)})$. Une telle égalisation peut alors être envisagée pour des alphabets à faible nombre d'états, et une réponse impulsionnelle modérément longue (selon la capacité de calcul disponible), ce qui implique un facteur d'accélération τ modérément faible.

Pour analyser la complexité de l'algorithme de Viterbi en multiporteuses, il convient tout d'abord d'exprimer la relation d'entrée/sortie du système multiporteuse de manière à être compatible avec une représentation en treillis :

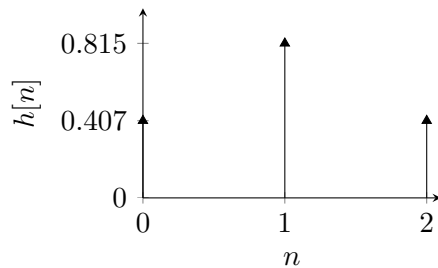
$$\tilde{c}_{p,q} = \sum_{(m,n) \in \Lambda} c_{m,n} A_{\tilde{g},g}[m-p, n-q] e^{j2\pi \frac{m-p}{M} qN} + \langle \tilde{g}_{p,q}; b \rangle \quad (2.22)$$

$$\tilde{\mathbf{c}}_{\mathbf{q}} = \sum_{n=0}^{K-1} \mathbf{A}_{\tilde{\mathbf{g}},\mathbf{g}}[n-q, q] \mathbf{c}_{\mathbf{n}} + \mathbf{b}_{\mathbf{q}} = \sum_{n=0}^{K-1} \overleftarrow{\mathbf{A}}_{\tilde{\mathbf{g}},\mathbf{g}}[q-n, q] \mathbf{c}_{\mathbf{n}} + \mathbf{b}_{\mathbf{q}}, \quad (2.23)$$

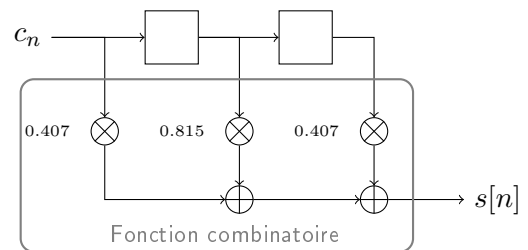
avec $\tilde{\mathbf{c}}_{\mathbf{n}} = (\tilde{c}_{0,n}, \dots, \tilde{c}_{M-1,n})^T$, $\mathbf{c}_{\mathbf{n}} = (c_{0,n}, \dots, c_{M-1,n})^T$, $\mathbf{b}_{\mathbf{q}} = (\langle \tilde{g}_{0,q}; b \rangle), \dots, (\langle \tilde{g}_{M-1,q}; b \rangle)^T$,

10. Lorsqu'en plus, la distribution du bruit suit une loi gaussienne est que les symboles sont IID, on retrouve l'algorithme de Viterbi sous sa forme « classique », optimale au sens du maximum de vraisemblance, et où les métriques de branche correspondent à la distance euclidienne entre l'observation $\tilde{c}_n$ et le symbole d'entrée c_n associé à la branche considérée.

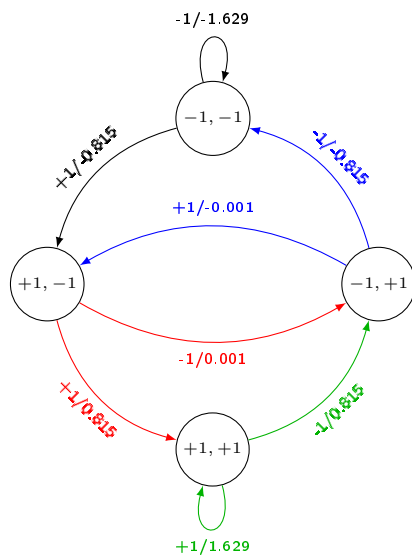
2.4. Techniques de détection



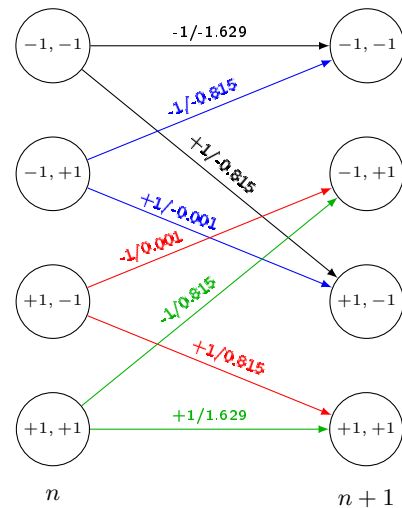
(a) Réponse impulsionnelle.



(c) Registres à décalage et fonction combinatoire.



(b) Chaîne de Markov. Les étiquettes de branche sont sous la forme entrée/sortie : $c_n/s[n]$.



(d) Treillis. Les étiquettes de branche sont sous la forme entrée/sortie : $c_n/s[n]$.

FIGURE 2.1 – Plusieurs manières de modéliser une transmission sur canal discret équivalent. On considère ici une transmission BPSK ($\mathcal{A} = \{-1; +1\}$) monoporteuse sur canal Proakis B [PS08, Chapitre 9].

$\overleftarrow{\mathbf{A}}_{\tilde{\mathbf{g}},\mathbf{g}}[n, q] = \mathbf{A}_{\tilde{\mathbf{g}},\mathbf{g}}[-n, q]$ et $\mathbf{A}_{\tilde{\mathbf{g}},\mathbf{g}}[n, q]$ une matrice Toeplitz définie par :

$$\mathbf{A}_{\tilde{\mathbf{g}},\mathbf{g}}[n, q] = \begin{pmatrix} A_{\tilde{g},g}[0, n] & A_{\tilde{g},g}[1, n]e^{j2\pi\frac{1}{M}qN} & \cdots & A_{\tilde{g},g}[M-1, n]e^{j2\pi\frac{M-1}{M}qN} \\ A_{\tilde{g},g}[-1, n]e^{j2\pi\frac{M-1}{M}qN} & A_{\tilde{g},g}[0, n] & \cdots & A_{\tilde{g},g}[M-2, n]e^{j2\pi\frac{M-2}{M}qN} \\ \vdots & \vdots & \ddots & \vdots \\ A_{\tilde{g},g}[1-M, n]e^{j2\pi\frac{1-M}{M}qN} & A_{\tilde{g},g}[2-M, n]e^{j2\pi\frac{2-M}{M}qN} & \cdots & A_{\tilde{g},g}[0, n] \end{pmatrix}. \quad (2.24)$$

Notons, sans que cela ne modifie le fonctionnement de l'algorithme, que la dépendance en q dans le terme $\overleftarrow{\mathbf{A}}_{\tilde{\mathbf{g}},\mathbf{g}}[q-n, q]$ de (2.23) rend le treillis variant dans le temps. Ici, les registres à décalage ont vocation à recevoir les vecteurs $\mathbf{c}_{\mathbf{n}}$, à valeurs dans $|\mathcal{A}|^M$. Cela signifie qu'en définissant L_a tel que $\mathbf{A}_{\tilde{\mathbf{g}},\mathbf{g}}[n, q] = \mathbf{0}$ si $n \notin [0; L_a - 1]$, alors le nombre d'états dans le treillis est $|\mathcal{R}| = |\mathcal{A}|^{M(L_a-1)}$, menant à une complexité de l'algorithme de Viterbi appliquée au problème d'égalisation du FTN multiporteuse exponentielle avec le nombre de sous-porteuses $\mathcal{O}(K|\mathcal{A}|^{2M(L_a-1)})$. Or une telle complexité est inacceptable pour un format de modulation dont le principe est justement d'utiliser de nombreuses sous-porteuses.

Enfin, notons qu'il est aussi possible de poser le problème de l'égalisation de chacun des symboles d'une séquence au sens du maximum a posteriori :

$$\bar{c}_{p,q} = \underset{\bar{c}_{p,q} \in \mathcal{A}}{\operatorname{argmax}} \mathbb{P}\{c_{p,q} = \bar{c}_{p,q} | \mathbf{Y} = \tilde{\mathbf{c}}\}, \quad (2.25)$$

ce qui permet de légèrement meilleurs TEB à faible RSB (voir figure 2.2) [LC04, Chapitre 12]. Là aussi, il existe un algorithme permettant de résoudre ce problème en s'appuyant sur une représentation du système sous forme de treillis : il s'agit de l'algorithme BCJR (du nom de ses créateurs Bahl, Cocke, Jelinek et Raviv) [Bah+74]. Il nécessite cependant de parcourir deux fois le treillis¹¹. De ce fait, on peut dire en première approximation que sa complexité est doublée par rapport à l'algorithme de Viterbi. D'autre part, il fait appel à des multiplications et des calculs d'exponentielles, qui rendent son implémentation matérielle plus délicate comparé à l'algorithme de Viterbi [RW95]. Ainsi, le faible gain de performance vis-à-vis de l'augmentation en complexité induite par l'algorithme BCJR font qu'il a été longtemps ignoré. Sa capacité à calculer les probabilités a posteriori $\mathbb{P}\{c_{p,q} = \bar{c}_{p,q} | \tilde{\mathbf{c}}\}$ en font cependant un candidat de choix lorsqu'une sortie souple (autrement dit, probabiliste) est nécessaire, comme c'est le cas notamment dans les systèmes turbo (voir partie 2.4.2).

2.4.1.2 Approches à plus faible complexité

Comme nous l'avons vu précédemment, il est rarement envisageable d'utiliser l'algorithme Viterbi ou BCJR, tant le nombre d'états dans le treillis grandit vite en fonction de l'alphabet

11. Étapes dites « forward » et « backward », de telle sorte que l'algorithme BCJR est aussi dénommé « forward-backward »

2.4. Techniques de détection

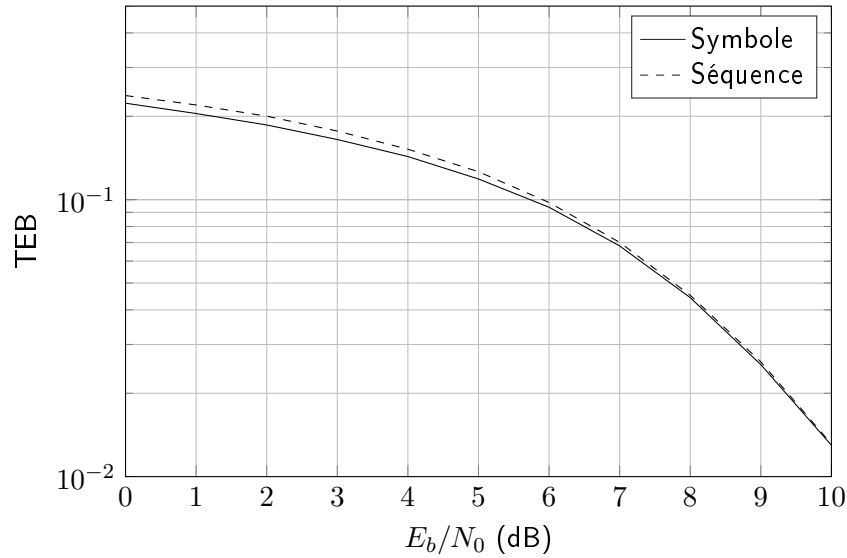


FIGURE 2.2 – Égalisation MAP d’une séquence (via l’algorithme de Viterbi) comparé à l’égalisation MAP de chaque symbole (via l’algorithme BCJR) pour une transmission monoporteuse BPSK sur canal Proakis B [PS08, Chapitre 9].

de modulation, de la mémoire introduite par le système FTN, et du nombre de sous-porteuses. Cependant, pour être sûr de pouvoir profiter de l’invariance de la distance minimale entre les différents signaux possibles lorsque la densité du système est inférieure à la limite de Mazo, il faut pouvoir égaliser de manière optimale (c’est-à-dire, au sens du MAP ou du MV), ou du moins quasi-optimale.

Dans le contexte monoporteuse, différentes stratégies ont été proposées [ARÖ13] :

- la réduction de treillis, qui consiste à ne considérer qu’une partie de la réponse impulsionnelle de $h[k]$ dans le treillis ;
- ajouter un filtre $f[k]$ avant l’algorithme de Viterbi, dont le rôle est de réduire la taille de la réponse impulsionnelle $(h * f)[k]$ vue par ce dernier (technique de « raccourcissement de canal » [RE13] ;
- ne retenir que les transitions pour lesquelles la métrique de branche est au dessus d’un seuil (algorithme M-BCJR) ;
- égalisation à retour de décision (égalisation linéaire, et suppression d’interférence à partir des symboles précédemment décidés).

En multiporteuse, la complexité inhérente aux approches par treillis, et la capacité initiale des systèmes multiporteuses orthogonaux à faciliter l’égalisation font que peu de stratégies de compensation basées sur des treillis de l’interférence entre impulsions de mises en forme ont été étudiées. Citons tout de même la solution de [Mat+98], qui consiste à effectuer une première étape d’égalisation ne prenant en compte que l’IEP¹², et dont le résultat est utilisé à l’entrée

12. Interférence Entre Porteuses

du second égaliseur dont le rôle est de traiter l'IES¹³. Enfin, notons que de la même manière qu'un lien très fort existe entre codage convolutif et égalisation en monoporteuse [For73], il est à supposer que les travaux sur le décodage des codes convolutifs à 2 dimensions et l'égalisation des systèmes FTN multiporteuses sont appelés à s'enrichir mutuellement [FV94] ; [CT03]. Enfin, d'autres approches adoptent une approche non-statistique et cherchent à minimiser la norme de l'erreur de reconstruction via des méthodes d'optimisation [HJZ14] ; [SHN16].

2.4.2 Turbo-égalisation

La plupart des systèmes de communication actuels utilisent des codes correcteurs d'erreurs afin d'améliorer la robustesse du système au bruit présent sur le canal, au prix d'une efficacité spectrale moindre (voir figure 1.18). Utiliser de tels systèmes en FTN pourrait permettre de compenser la perte d'efficacité spectrale induite par le code correcteur d'erreur, à condition qu'il soit possible d'atteindre des performances comparables en termes de probabilité d'erreur.

Dans le contexte des communications codées sur canal sélectif en fréquence, la technique de turbo-égalisation a montré sa capacité à compenser parfaitement l'IES sur plusieurs types de canaux [Dou+95]. La turbo-égalisation consiste à appliquer le concept de turbo décodage des codes concaténés série aux communications codées et entrelacées sur canal sélectif (voir figure 2.3). Notons que l'entrelacement a pour effet de disperser l'information, ce qui évite les paquets d'erreurs. Du point de vue de la théorie des codes, il permet d'augmenter la distance minimale entre les signaux produits par le code formé par la concaténation du code convolutif et du canal [Ber+07, Chapitre 6].

Par la suite, nous noterons $\mathbf{b}$ la séquence binaire constituant le message à transmettre, $\mathbf{b}^c$ est la séquence binaire codée de manière à ajouter de la redondance à $\mathbf{b}$ et $\mathbf{b}^i$ est simplement constitué des éléments de $\mathbf{b}^c$, dont la position dans le vecteur a été modifiée par l'entrelaceur.

Le turbo-égaliseur est principalement constitué de deux décodeurs à entrées/sorties souples (SISO) arrangés de manière à former un système bouclé (voir figure 2.4). Dans la boucle, l'égaliseur et le décodeur échangent leurs estimations probabilistes des mots de codes émis de manière itérative, permettant le traitement conjoint de l'IES et du décodage canal.

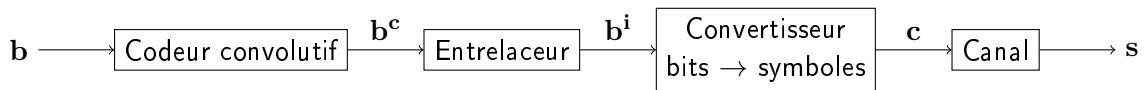


FIGURE 2.3 – Émetteur pour une modulation codée et entrelacée. Dans le cadre de la turbo égalisation, le canal est vu comme un code.

2.4.2.1 Turbo-égalisation optimale

Dans sa version la plus optimale, l'égaliseur et le décodeur du turbo-égaliseur sont implémentés à l'aide de l'algorithme BCJR. Le décodeur doit être capable de calculer les probabilités

13. Interférence Entre Symboles

2.4. Techniques de détection

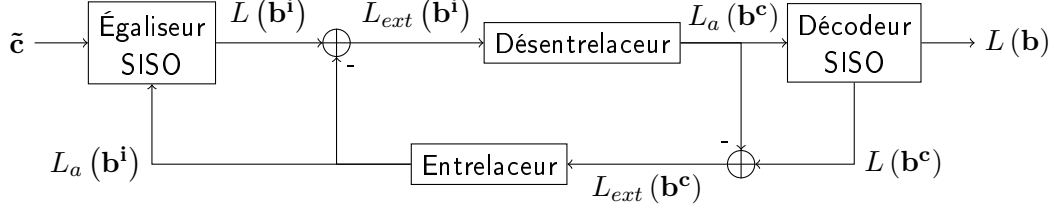


FIGURE 2.4 – Structure générale d'un turbo-égaliseur.

a posteriori des bits du message $\mathbf{b}$, mais aussi des bits codés $\mathbf{b}^c$ (donc de l'entrée et de la sortie du codeur). Cette fonctionnalité permet d'utiliser les probabilités sur les bits codés comme *a priori* pour l'égaliseur. En pratique l'échange d'informations s'effectue sur la forme de LRV¹⁴ extrinsèques, dont nous expliquons l'intérêt ci-après. Admettons que l'on utilise l'algorithme BCJR sur un codeur dont la séquence binaire d'entrée est notée $\mathbf{e} = \{e_k\}_{k \in [0; L_e]}$, la séquence binaire de sortie $\mathbf{s} = \{s_n\}_{n \in [0; L_s]}$, l'observation $\mathbf{y}$, et la variable aléatoire associée à l'observation $\mathbf{Y}$, alors les LRV *a posteriori* sont définis par :

$$L(e_k|\mathbf{y}) = \ln \frac{\mathbb{P}\{e_k = 0|\mathbf{Y} = \mathbf{y}\}}{\mathbb{P}\{e_k = 1|\mathbf{Y} = \mathbf{y}\}} \quad \text{et} \quad L(s_n|\mathbf{y}) = \ln \frac{\mathbb{P}\{s_n = 0|\mathbf{Y} = \mathbf{y}\}}{\mathbb{P}\{s_n = 1|\mathbf{Y} = \mathbf{y}\}}. \quad (2.26)$$

Ces LRV *a posteriori* peuvent se décomposer en une partie dite « extrinsèque » et une partie *a priori*. Montrons cela sur les LRV des bits d'entrée :

$$\begin{aligned} L(e_k|\mathbf{y}) &= \ln \frac{\mathbb{P}\{e_k = 0|\mathbf{Y} = \mathbf{y}\}}{\mathbb{P}\{e_k = 1|\mathbf{Y} = \mathbf{y}\}} = \ln \frac{p_{e_k, \mathbf{Y}}(0, \mathbf{y}) / p_{\mathbf{Y}}(\mathbf{y})}{p_{e_k, \mathbf{Y}}(1, \mathbf{y}) / p_{\mathbf{Y}}(\mathbf{y})} = \ln \frac{p_{\mathbf{Y}|e_k}(\mathbf{y}, 0) \mathbb{P}\{e_k = 0\}}{p_{\mathbf{Y}|e_k}(\mathbf{y}, 1) \mathbb{P}\{e_k = 1\}} \\ &= \underbrace{\ln \frac{p_{\mathbf{Y}|e_k}(\mathbf{y}, 0)}{p_{\mathbf{Y}|e_k}(\mathbf{y}, 1)}}_{L_{ext}(e_k)} + \underbrace{\ln \frac{\mathbb{P}\{e_k = 1\}}{\mathbb{P}\{e_k = 0\}}}_{L_a(e_k)}, \end{aligned} \quad (2.27)$$

où $L_{ext}(\cdot)$ donne l'information extrinsèque et $L_a(\cdot)$ l'information *a priori*. Or, dans la structure bouclée, l'information *a priori* disponible à l'entrée de l'égaliseur provient du décodeur, et celle du décodeur provient de l'égaliseur. Dans ce cas, d'après le « principe turbo » il est préférable de ne transmettre que l'information extrinsèque pour permettre une convergence optimale du turbo-égaliseur [Ber+07, Chapitre 7][Dou+95]. En particulier, il a été prouvé que cette structure de turbo-égalisation (n'échangeant que l'information extrinsèque) converge vers le détecteur au maximum de vraisemblance quand le RSB est suffisant [SRF08]. Ainsi, à condition que le canal ne modifie pas la distance minimale entre signaux, le turbo-égaliseur permet de retrouver les performances du système codé (sans canal sélectif). En pratique, on remarque que cela intervient à des RSB relativement faibles, voir [Ber+07, Chapitre 11].

Dans la suite de ce document, nous utiliserons le code convolutif $(7, 5)_8$, sauf mention contraire, dont le treillis est donné en figure 2.5. En effet, avec seulement quatre états, la complexité de l'algorithme BCJR appliqué à ce code est suffisamment faible pour nous permettre

14. Logarithme de Rapport de Vraisemblances

d'effectuer des simulations dans un temps raisonnable. D'autre part, il règne une sorte de consensus autour du choix de ce code pour la comparaison de différents turbo-égaliseurs [ARÖ13].

Notons enfin qu'à l'instar du cas où la transmission s'effectue sans code, la complexité de l'algorithme BCJR appliquée à l'égalisation d'un système multiporteuse FTN est généralement prohibitive, ce qui rend la turbo-égalisation optimale hors de portée.

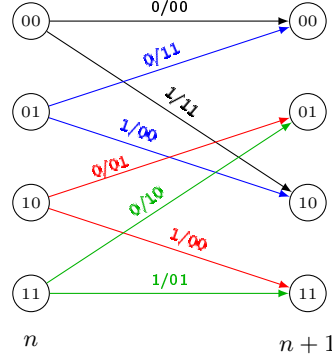


FIGURE 2.5 – Treillis du code convolutif $(7,5)_8$. Les étiquettes de branche sont sous la forme entrée/sortie : $b_n/b_{2n}^c b_{2n+1}^c$.

2.4.2.2 Turbo-égalisation à plus faible complexité

Un aspect intéressant des techniques de turbo-égalisation, est que l'égaliseur BCJR peut être remplacé par une version sous-optimale, du moment qu'elle est SISO. Dans ce cas, le seuil de convergence, c'est à dire le RSB à partir duquel le turbo-égaliseur retrouve les performances du système sans canal, se voit plus ou moins augmenté, selon la pertinence de l'égaliseur utilisé.

Ainsi, toutes les techniques de réduction de la complexité de l'égalisation évoquées en 2.4.1 peuvent être utilisées. Nous allons les séparer en catégories :

- les algorithmes qui sont SISO par construction, tels que les dérivés de l'algorithme BCJR et de l'algorithme de Viterbi (réduction de treillis, M-BCJR, SOVA¹⁵ [Bat87]) ;
- les algorithmes basés sur des approches linéaires à retour de décision, lesquels nécessitent des blocs d'adaptation (ou de conversion SISO), comme illustré en figure 2.6.

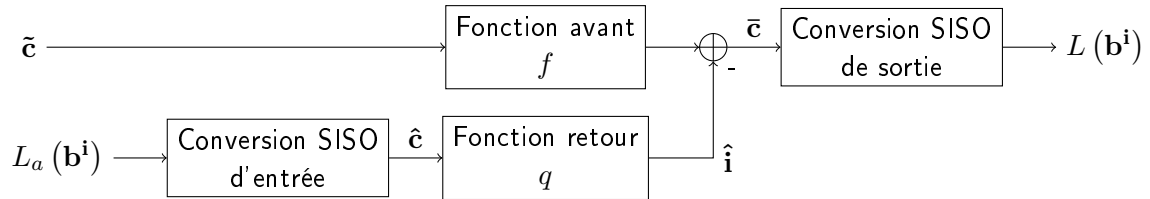


FIGURE 2.6 – Structure générale d'un égaliseur « linéaire » avec information *a priori*.

15. *Soft Output Viterbi Algorithm* (Algorithme de Viterbi à sortie pondérée)

2.4. Techniques de détection

L'idée de réutiliser la riche littérature de l'égalisation linéaire à retour de décision dans un contexte turbo a été introduite par Laot *et al.* [LGL01]. Cette approche a l'avantage de présenter une complexité linéaire avec la mémoire introduite par le canal, et constante en fonction de la taille de la constellation $|\mathcal{A}|$ (on omet ici volontairement de prendre en compte les blocs d'adaptation SISO, qui seront détaillés plus tard). Sur le synoptique de la figure 2.6, on observe la structure générale d'un égaliseur « linéaire¹⁶ » SISO.

Afin d'obtenir de plus faibles complexités, on impose une contrainte de linéarité sur les fonctions avant f et retour q . La fonction avant f a pour but de compenser l'interférence (par exemple, selon le critère MEQM¹⁷), tandis que la fonction retour q est chargée de calculer une estimée du terme d'interférence $\mathbf{i}$ résiduelle en sortie de f . Cette estimée est retranchée à la sortie de la fonction avant pour produire les symboles égalisés $\bar{\mathbf{c}}$.

Le convertisseur SISO d'entrée permet de calculer des estimations $\hat{\mathbf{c}}$ des symboles envoyés $\mathbf{c}$, qui permettront le calcul de l'interférence résiduelle via q . Il est généralement construit comme l'estimateur minimisant l'EQM¹⁸ entre les symboles estimés $\hat{\mathbf{c}}$ et les symboles émis $\mathbf{c}$ à partir des LRV *a priori* $L_a(\mathbf{b}^i)$ sur les bits entrelacés :

$$\hat{c}_{m,n} = \underset{\hat{c}_{m,n} \in \mathbb{C}}{\operatorname{argmin}} \mathbb{E} \left\{ |c_{m,n} - \hat{c}_{m,n}|^2 \middle| L_a(\mathbf{b}^i) \right\} \quad (2.28)$$

$$= \mathbb{E} \left\{ c_{m,n} \middle| L_a(\mathbf{b}^i) \right\}, \quad (2.29)$$

où la démonstration du passage de la première à la deuxième ligne est donnée en annexe A. Ce convertisseur SISO sera utilisé dans tout le reste du manuscrit, nous détaillons donc l'expression des symboles estimés $\hat{c}_{m,n}$ ici :

$$\hat{c}_{m,n} = \mathbb{E} \left\{ c_{m,n} \middle| L_a(\mathbf{b}^i) \right\} = \sum_{c \in \mathcal{A}} c \mathbb{P} \left\{ c_{m,n} = c \middle| L_a(\mathbf{b}^i) \right\}, \quad (2.30)$$

où on rappelle que les symboles sont formés à partir de bits supposés IID, grâce à l'action de l'entrelaceur. Notons $b_c(l)$ ($l \in [0; N_b - 1]$) les bits composant un symbole $c \in \mathcal{A}$, et rappelons que les symboles émis $c_{m,n}$ sont construits à partir de bits entrelacés : $b_l(c_{m,n}) \in \mathbf{b}^i$. On obtient alors :

$$\mathbb{P} \left\{ c_{m,n} = c \middle| L_a(\mathbf{b}^i) \right\} = \prod_{l=0}^{N_b-1} \mathbb{P} \left\{ b_l(c_{m,n}) = b_l(c) \middle| L_a(\mathbf{b}^i) \right\} \quad (2.31)$$

$$= \prod_{l=0}^{N_b-1} \mathbb{P} \left\{ b_l(c_{m,n}) = b_l(c) \middle| L_a(b_l(c_{m,n})) \right\}, \quad (2.32)$$

or d'après la définition d'un LRV, on a :

$$L(b) = \ln \frac{\mathbb{P}\{b=1\}}{\mathbb{P}\{b=0\}} \Leftrightarrow \mathbb{P}\{b=0\} = \frac{1}{1 + e^{-L(b)}} \text{ et } \mathbb{P}\{b=1\} = \frac{e^{-L(b)}}{1 + e^{-L(b)}}, \quad (2.33)$$

16. L'appellation « linéaire » constitue ici un abus de langage, puisque l'égaliseur dans son ensemble, en prenant en compte les blocs d'adaptation SISO, est clairement non-linéaire.

17. Minimisation de l'Erreur Quadratique Moyenne

18. Erreur Quadratique Moyenne

en introduisant la fonction tangente hyperbolique :

$$\tanh(z) = \frac{1 - e^{-2z}}{1 + e^{-2z}}, \quad \forall z \in \mathbb{C} \setminus \{j\pi(\mathbb{Z} + 1/2)\}, \quad (2.34)$$

on remarque que :

$$\mathbb{P}\{b = 0\} = \frac{1}{2} (1 + \tanh(L(b)/2)) \quad \text{et} \quad \mathbb{P}\{b = 1\} = \frac{1}{2} (1 - \tanh(L(b)/2)), \quad (2.35)$$

ce qui peut être généralisé par :

$$\mathbb{P}\{b = x\} = \frac{1}{2} (1 + (1 - 2x) \tanh(L(b)/2)). \quad (2.36)$$

En combinant (2.36) et (2.32) on obtient l'expression des probabilités *a priori* sur les symboles :

$$\mathbb{P}\{c_{m,n} = c | L_a(\mathbf{b}^c)\} = \frac{1}{2^{N_b}} \prod_{l=0}^{N_b-1} \frac{1 + (1 - 2b_l(c_{m,n})) \tanh\left(\frac{L(b_l(c_{m,n}))}{2}\right)}{2}, \quad (2.37)$$

ce qui, injecté dans (2.30), donne l'expression des bits estimés :

$$\hat{c}_{m,n} = \frac{1}{2^{N_b}} \sum_{c \in \mathcal{A}} c \prod_{l=0}^{N_b-1} \frac{1 + (1 - 2b_l(c_{m,n})) \tanh\left(\frac{L(b_l(c_{m,n}))}{2}\right)}{2}. \quad (2.38)$$

Cette expression se simplifie plus ou moins selon la constellation utilisée, et des tables permettant d'effectuer plus facilement la conversion SISO d'entrée pour différentes constellations classiques (BPSK, QPSK, etc.) ont été calculées par Tüchler dans [TSK02]. Notons que cette opération peut être vue comme une conversion bit à symbole souple.

Le convertisseur SISO de sortie permet de convertir les symboles égalisés $\bar{\mathbf{c}}$ en LRV. Hormis dans le cas du raccourcissement de canal, où l'approche hybride filtrage linéaire/égalisation BCJR permet de calculer exactement les LRV [RE13], une approche symbole par symbole à plus faible complexité est généralement utilisée.

$$L(b_l(c_{p,q}) | \bar{c}_{p,q}) = \ln \frac{\mathbb{P}\{b_l(c_{p,q}) = 0 | \bar{c}_{p,q}\}}{\mathbb{P}\{b_l(c_{p,q}) = 1 | \bar{c}_{p,q}\}} = \ln \frac{\sum_{c^0 \in \mathcal{A}, b_l(c^0)=0} \mathbb{P}\{c_{p,q} = c^0 | \bar{c}_{p,q}\}}{\sum_{c^1 \in \mathcal{A}, b_l(c^1)=1} \mathbb{P}\{c_{p,q} = c^1 | \bar{c}_{p,q}\}} \quad (2.39)$$

$$= \ln \frac{\sum_{c^0 \in \mathcal{A}, b_l(c^0)=0} \mathbb{P}\{\bar{c}_{p,q} | c_{p,q} = c^0\} \mathbb{P}\{c_{p,q} = c^0\}}{\sum_{c^1 \in \mathcal{A}, b_l(c^1)=1} \mathbb{P}\{\bar{c}_{p,q} | c_{p,q} = c^1\} \mathbb{P}\{c_{p,q} = c^1\}}. \quad (2.40)$$

Ici les termes $\mathbb{P}\{c_{p,q} = c^0\}$ et $\mathbb{P}\{c_{p,q} = c^1\}$ représentent l'information a priori sur les symboles. On les remplacera donc par $\mathbb{P}\{c_{p,q} = c^0 | L_a(b^c)\}$ et $\mathbb{P}\{c_{p,q} = c^1 | L_a(b^c)\}$, respectivement et déjà calculés en (2.37). Exprimons désormais les symboles égalisés $\bar{c}_{p,q}$ comme étant la somme d'un terme utile, et d'un terme prenant en compte l'interférence résiduelle et le bruit :

$$\bar{c}_{p,q} = \alpha_{p,q} c_{p,q} + \nu_{p,q}, \quad (2.41)$$

2.4. Techniques de détection

avec $\alpha_{p,q} \in \mathbb{C}$. On fait ensuite l'hypothèse que le terme $\nu_{p,q}$ est gaussien complexe circulaire et centré, de variance $\sigma_{\nu_{p,q}}^2$ de telle sorte que sa densité de probabilité est donnée par :

$$p_{\nu_{p,q}}(\nu) = \frac{1}{\pi \sigma_{\nu_{p,q}}^2} \exp\left(-\frac{|\nu|^2}{\sigma_{\nu_{p,q}}^2}\right). \quad (2.42)$$

Ainsi, d'après (2.41) la variable aléatoire $\bar{c}_{p,q}$ sachant la valeur de $c_{p,q}$ est simplement la somme de la variable aléatoire $\nu_{p,q}$ et d'une constante $\alpha_{p,q}c_{p,q}$. Elle suit donc la même loi, à l'exception près de son espérance qui vaut alors $\alpha_{p,q}c_{p,q}$:

$$\mathbb{P}\{\bar{c}_{p,q}|c_{p,q} = c\} = \frac{1}{\pi \sigma_{\nu_{p,q}}^2} \exp\left(-\frac{|\bar{c}_{p,q} - \alpha_{p,q}c|^2}{\sigma_{\nu_{p,q}}^2}\right). \quad (2.43)$$

On obtient alors l'expression finale des LRV *a posteriori* en injectant (2.43) dans (2.40) :

$$L(b_l(c_{p,q})|\bar{c}_{p,q}) = \ln \frac{\sum_{c^0 \in \mathcal{A}, b_l(c^0)=0} \exp\left(-|\bar{c}_{p,q} - \alpha_{p,q}c^0|^2/\sigma_{\nu_{p,q}}^2\right) \mathbb{P}\{c_{p,q} = c^0|L_a(b^c)\}}{\sum_{c^1 \in \mathcal{A}, b_l(c^1)=1} \exp\left(-|\bar{c}_{p,q} - \alpha_{p,q}c^1|^2/\sigma_{\nu_{p,q}}^2\right) \mathbb{P}\{c_{p,q} = c^1|L_a(b^c)\}}. \quad (2.44)$$

Il est là aussi possible d'interpréter cette opération comme une conversion symbole à bit souple.

Bien que cette approche mettant à profit des structures linéaires d'égalisation permettent un gain effectif en complexité, il est à noter que celle des convertisseur SISO n'est pas anodine. En effet, concernant la conversion SISO d'entrée, si des simplifications existent pour plusieurs types de constellation [TSK02], l'expression générale fait tout de même intervenir la fonction tangente hyperbolique, ainsi que de nombreuses multiplications dont le nombre augmente avec la taille de la constellation. De même, la conversion SISO de sortie fait intervenir le logarithme naturel de sommes d'exponentielles, où le nombre de termes à sommer augmente avec la taille de la constellation. Pour ce dernier problème, il serait possible d'utiliser des approches similaires à celles adoptées dans les algorithmes log-BCJR et max-log-BCJR [RVH95]. On en conclut qu'à l'inverse des égaliseurs (classiques) linéaires, les turbo-égaliseurs « linéaires » voient leur complexité augmenter avec la taille de la constellation, à travers les convertisseurs SISO.

Notons que les expressions de $\alpha_{p,q}$ et $\sigma_{\nu_{p,q}}^2$ diffèrent selon les fonctions avant/arrière utilisées et le canal. À partir de cette présentation générale, on peut retrouver la plupart des égaliseurs linéaires SISO de la littérature.

En monoporteuse citons le cas où le canal est supposé parfaitement estimé, l'égaliseur MEQM où les coefficients de f et q sont calculés de manière à minimiser l'EQM entre les symboles émis et les symboles égalisés. Le Bidan en a proposé plusieurs versions dans sa thèse, pour un canal statique : longueur de la réponse impulsionnelle de f et q finie, infinie, ou approximation de la version à réponse impulsionnelle infinie via transformée de Fourier. Il a également démontré la faisabilité de l'implémentation d'un tel système sur DSP¹⁹ [Le 03]. Tüchler a également proposé une version pour les canaux variant dans le temps [TSK02]. Dans les deux cas, y compris

19. *Digital Signal Processor* (Processeur de traitement du signal)

lorsque le canal est statique, la réponse impulsionnelle de f et q doit être recalculée à chaque itération. Dans sa version originale, proposée par Glavieux, les coefficients de f et q étaient calculés de manière adaptative à l'aide de l'algorithme LMS²⁰ [GLL97]. Cette approche reste intéressante dans le contexte du suivi d'un canal variant lentement dans le temps, elle peut également permettre de réduire la taille de la séquence d'apprentissage [Hél+03]. Enfin, l'égalisation à espacement fractionnaire, notamment connue pour sa capacité à être plus robuste aux erreurs de synchronisation [GW81], a aussi été développée en version turbo [Otn04].

En multiporteuses, le fait d'avoir deux dimensions à égaliser rendent l'application de critères classiques, tels que la minimisation de l'EQM ou le forçage à zéro²¹ plus difficile à mener, et éventuellement à des complexités algorithmiques élevées. Citons l'approche développée par Lahbabi dans sa thèse, consistant à égaliser l'interférence dans une des dimensions (soit l'IES, soit l'IEP) suivant le critère MEQM, de manière similaire à l'approche monoporteuse de Le Bidan et Tüchler, tandis que dans l'autre dimension, l'interférence estimée est simplement soustraite [Lah17, Chapitre 3]. Une illustration du système traitant l'IES selon le critère MEQM est à trouver en figure 2.7. Dans le même esprit, un égaliseur traitant l'IEP par forçage à zéro, et à annulation de l'IES a été développé par Abello [Abe+16]. Enfin, en multiporteuse, une solution efficace et peu complexe consiste à ne faire que de l'annulation d'interférence souple. Cette technique, en se passant de filtre avant f [DRO10], permet également à la fonction q de rester invariante d'une itération à l'autre. Dans le chapitre 4, nous verrons comment un choix judicieux des prototypes d'émission et de réception g et $\check{g}$ du système multiporteuse permet à un système à suppression souple de l'interférence de satisfaire le critère de la minimisation de l'EQM en présence de BABG, sans nécessiter de mise-à-jour de la fonction q . Notons que la plupart des stratégies de turbo-égalisation présentées dans ce paragraphe ont été développées pour des systèmes FTN basés sur des modulations FBMC/OQAM, là où nos travaux portent sur des systèmes FTN FBMC/QAM tels que présentés dans la partie 1.4.2.

2.5 Conclusion

Dans ce chapitre, nous avons vu les spécificités des communications FTN, et plus particulièrement FTN multiporteuses, en termes de modélisation, de limites théoriques, et de techniques de réception.

En particulier, nous avons évoqué le fait que le bruit en sortie du démodulateur linéaire n'est pas blanc dans le cas général, impliquant une adaptation des algorithmes classiques d'égalisation. D'autre part, nous avons vu que les algorithmes d'égalisation basés sur des représentations en treillis sont inenvisageables dans le cadre de modulations multiporteuses, pour des raisons de complexité algorithmique.

En revanche, la grande flexibilité de configuration offerte par la turbo-égalisation « linéaire »

20. *Least Mean Square* (Moindre carrés)

21. L'égalisation suivant le critère du forçage à zéro consiste à trouver le filtre f compensant totalement le canal, de manière à ce qu'il n'y ait plus d'interférence. Une telle approche nécessite un canal inversible et a tendance à rehausser le niveau de bruits sur les canaux dont la réponse fréquentiel présente de forts évanouissements.

2.5. Conclusion

en font un candidat de choix pour le développement de systèmes de réception arborant un bon compromis performance-complexité, y compris en FTN multiporteuse. La recherche de tels systèmes est justement l'objet du chapitre 4, dans lequel nous verrons notamment qu'un choix judicieux des paramètres du système linéaire permet de construire un turbo-égaliseur à la fois simple et présentant de bonnes performances en termes de TEB. Mais avant cela, nous nous intéressons aux systèmes linéaires non-codés dans le chapitre 3. Nous y proposons en particulier des impulsions de mise en forme permettant, pour une densité donnée, de limiter la quantité d'auto-interférence sur canal à BABG.

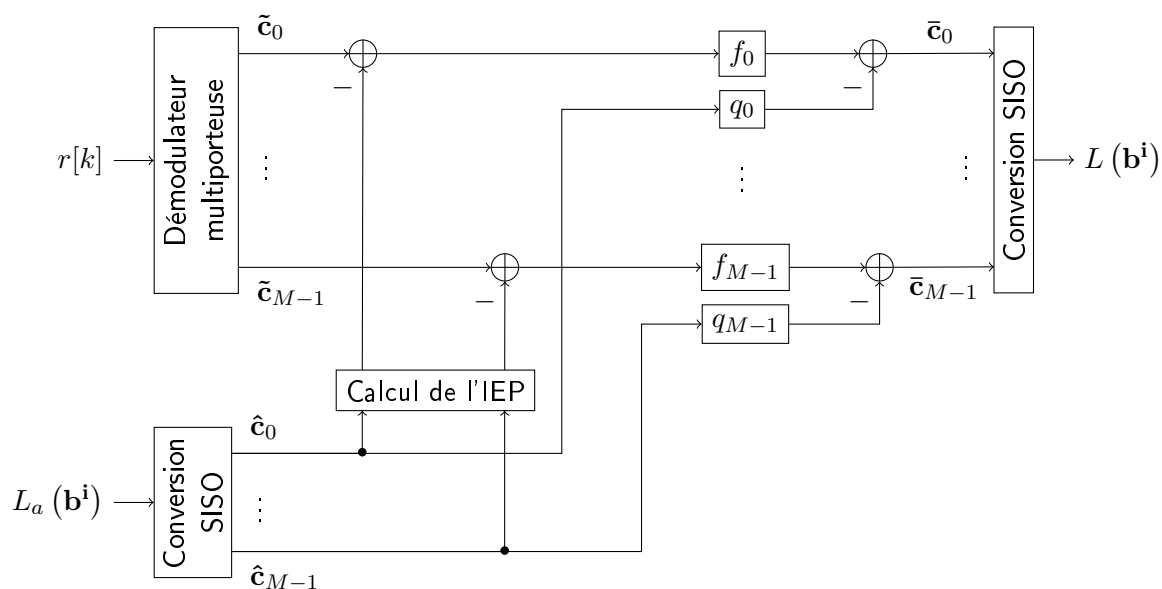


FIGURE 2.7 – Synoptique d'un égaliseur traitant l'IES au sens MEQM, et à annulation de l'IEP avec information *a priori*.

Système multiporteuse linéaire au-delà de la cadence de Nyquist optimal en présence de BABG

Sommaire

3.1	Introduction	65
3.2	Prototypes maximisant le RSIB	66
3.2.1	Prototypes maximisant le RSB	68
3.2.2	Prototypes maximisant le RSI	68
3.2.3	Prototypes maximisant conjointement le RSB et le RSI	69
3.3	Analyse statistique du terme d'auto-interférence	70
3.3.1	Espérance	71
3.3.2	Égalité de la variance des parties réelles et imaginaires	71
3.3.3	Décorrélation de la partie réelle et imaginaire	73
3.3.4	Distribution de la partie réelle et imaginaire	73
3.3.5	Résumé	77
3.4	Performances théoriques et simulations	77
3.4.1	Comportement en RSB, RSI et RSIB	78
3.4.2	Comportement en TEB	81
3.5	Conclusion	85

3.1 Introduction

Une des conséquences des modulations FTM est l'apparition d'un terme d'interférence en sortie du démodulateur. Comme nous l'avons évoqué dans le chapitre 2, il existe de nombreuses manières, plus ou moins performantes et algorithmiquement exigeantes, permettant de diminuer l'impact de cette interférence côté récepteur. Cependant il est également intéressant de chercher à optimiser les émetteurs, de manière à minimiser l'importance du terme d'interférence pour une densité ρ donnée. Ce chapitre a vocation à étudier cet aspect, qui semble relativement négligé dans la littérature.

Dans ce chapitre, en se plaçant dans le contexte d'un émetteur/récepteur linéaire et d'un canal à BABG, nous déterminerons premièrement les familles d'émission et de réception permettant de maximiser conjointement le rapport entre la puissance du signal utile et la somme

de la puissance de l'interférence et du bruit filtré : le RSIB¹. On donnera également une expression théorique du RSIB maximal. Dans un deuxième temps, nous étudierons les propriétés statistiques de l'interférence, et en particulier, nous verrons les cas dans lesquels il est possible de l'approcher par une loi normale complexe circulaire. Enfin, nous commenterons le comportement théorique, en termes de RSIB, de RSI² et du RSB du système utilisant les prototypes optimaux. Nous dériverons une formule d'approximation de la probabilité d'erreur binaire, et vérifierons la validité de nos résultats théoriques par le biais de simulations.

3.2 Prototypes maximisant le RSIB

Dans cette partie, nous cherchons à déterminer les conditions sur les prototypes d'émission et de réception permettant de maximiser le RSIB. Nous verrons que les prototypes que nous proposons permettent en particulier de maximiser simultanément le RSI et le RSB, ce qui signifie qu'il n'est pas possible dans ce cas d'obtenir un meilleur RSI en relâchant la contrainte sur l'optimalité du RSB, et vice-versa.

Avant de maximiser le RSIB, il nous faut tout d'abord le définir. Pour cela, on réécrit l'expression des symboles en sortie du démodulateur multiporteuse (1.58) pour l'adapter au cas où le canal est simplement à BABG (pas de sélectivité en temps et/ou en fréquence) :

$$\tilde{c}_{p,q} = \underbrace{c_{p,q} \langle \check{g}_{p,q}; g_{p,q} \rangle}_{\text{Signal utile}} + \underbrace{\sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} c_{m,n} \langle \check{g}_{p,q}; g_{m,n} \rangle}_{\text{Interférence : } i_{p,q}} + \underbrace{\langle \check{g}_{p,q}; b \rangle}_{\text{Bruit filtré}}. \quad (3.1)$$

On rappelle également que les symboles sont supposés IID, centrés et de variance σ_c^2 , et que le bruit $b(t)$ est supposé blanc, gaussien, complexe circulaire, centré et de variance σ_b^2 . En s'appuyant sur l'expression (3.1), on peut déterminer que l'espérance du signal utile est nulle :

$$\mathbb{E} \{ c_{p,q} \langle \check{g}_{p,q}; g_{p,q} \rangle \} = \mathbb{E} \{ c_{p,q} \} \langle \check{g}_{p,q}; g_{p,q} \rangle = 0, \quad (3.2)$$

puis calculer sa variance :

$$\mathbb{E} \left\{ |c_{p,q} \langle \check{g}_{p,q}; g_{p,q} \rangle|^2 \right\} = \mathbb{E} \{ |c_{p,q}|^2 \} |\langle \check{g}_{p,q}; g_{p,q} \rangle|^2 = \sigma_c^2 |\langle \check{g}_{p,q}; g_{p,q} \rangle|^2. \quad (3.3)$$

Le terme d'interférence a lui aussi une espérance nulle :

$$\mathbb{E} \{ i_{p,q} \} = \mathbb{E} \left\{ \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} c_{m,n} \langle \check{g}_{p,q}; g_{m,n} \rangle \right\} = \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} \mathbb{E} \{ c_{m,n} \} \langle \check{g}_{p,q}; g_{m,n} \rangle = 0, \quad (3.4)$$

1. Rapport signal à interférence plus bruit
2. Rapport signal à interférence

3.2. Prototypes maximisant le RSIB

et sa variance est donnée par :

$$\mathbb{E} \{ |i_{p,q}|^2 \} = \mathbb{E} \left\{ \left| \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} c_{m,n} \langle \check{g}_{p,q}; g_{m,n} \rangle \right|^2 \right\} \quad (3.5)$$

$$= \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} \sum_{(m',n') \in \Lambda \setminus \{(p,q)\}} \mathbb{E} \{ c_{m,n} c_{m',n'}^* \} \langle \check{g}_{p,q}; g_{m,n} \rangle \langle \check{g}_{p,q}; g_{m',n'} \rangle^* \quad (3.6)$$

$$= \sigma_c^2 \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} |\langle \check{g}_{p,q}; g_{m,n} \rangle|^2 \quad (3.7)$$

$$= \sigma_c^2 \left(\sum_{(m,n) \in \Lambda} |\langle \check{g}_{p,q}; g_{m,n} \rangle|^2 - |\langle \check{g}_{p,q}; g_{p,q} \rangle|^2 \right). \quad (3.8)$$

Enfin, pour le terme de bruit filtré, on a là encore une espérance nulle :

$$\mathbb{E} \{ \langle \check{g}_{p,q}; b \rangle \} = \mathbb{E} \left\{ \int_{-\infty}^{+\infty} \check{g}_{p,q}^*(t) b(t) dt \right\} = \int_{-\infty}^{+\infty} \check{g}_{p,q}^*(t) \mathbb{E} \{ b(t) \} dt = 0, \quad (3.9)$$

et une variance donnée par :

$$\mathbb{E} \{ |\langle \check{g}_{p,q}; b \rangle|^2 \} = \mathbb{E} \left\{ \left| \int_{-\infty}^{+\infty} \check{g}_{p,q}^*(t) b(t) dt \right|^2 \right\} \quad (3.10)$$

$$= \mathbb{E} \left\{ \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \check{g}_{p,q}^*(t) \check{g}_{p,q}(t') b(t) b^*(t') dt dt' \right\} \quad (3.11)$$

$$= \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \check{g}_{p,q}^*(t) \check{g}_{p,q}(t') \mathbb{E} \{ b(t) b^*(t') \} dt dt' \quad (3.12)$$

$$= \sigma_b^2 \int_{-\infty}^{+\infty} |\check{g}_{p,q}(t)|^2 dt \quad (3.13)$$

$$= \sigma_b^2 \|\check{g}_{p,q}\|^2. \quad (3.14)$$

Munis de ces définitions, on peut donner les expressions du RSB, du RSI ainsi que du RSIB, toutes trois dépendantes de (p, q) :

$$\text{RSB}_{p,q} = \frac{\sigma_c^2 |\langle \check{g}_{p,q}; g_{p,q} \rangle|^2}{\sigma_b^2 \|\check{g}_{p,q}\|^2}, \quad (3.15)$$

$$\text{RSI}_{p,q} = \frac{|\langle \check{g}_{p,q}; g_{p,q} \rangle|^2}{\sum_{(m,n) \in \mathbb{Z}^2} |\langle \check{g}_{p,q}; g_{m,n} \rangle|^2 - |\langle \check{g}_{p,q}; g_{p,q} \rangle|^2}, \quad (3.16)$$

$$\text{RSIB}_{p,q} = \frac{1}{\text{RSI}_{p,q}^{-1} + \text{RSB}_{p,q}^{-1}} \quad (3.17)$$

$$= \frac{|\langle \check{g}_{p,q}; g_{p,q} \rangle|^2}{\sum_{(m,n) \in \Lambda} |\langle \check{g}_{p,q}; g_{m,n} \rangle|^2 - |\langle \check{g}_{p,q}; g_{p,q} \rangle|^2 + \|\check{g}_{p,q}\|^2 \sigma_b^2 / \sigma_c^2}. \quad (3.18)$$

Nous présenterons par la suite les conditions pour maximiser de manière séparée le RSB et le RSI, et nous montrerons que l'association de ces conditions est possible, menant à la maximisation du RSIB.

3.2.1 Prototypes maximisant le RSB

Pour maximiser le RSB, il suffit d'appliquer la relation de Cauchy–Schwarz au numérateur :

$$|\langle \check{g}_{p,q}; g_{p,q} \rangle|^2 \leq \|\check{g}_{p,q}\|^2 \|g_{p,q}\|^2 \Leftrightarrow \text{RSB}_{p,q} \leq \sigma_c^2 \|g_{p,q}\|^2 / \sigma_b^2, \quad (3.19)$$

où la borne $\text{RSB}_{\max} = \sigma_c^2 \|g_{p,q}\|^2 / \sigma_b^2 = \sigma_c^2 \|g\|^2 / \sigma_b^2$ est atteinte lorsque $\check{g}_{p,q}(t) = \mu_{p,q} g_{p,q}(t)$, $\mu_{p,q} \in \mathbb{C}^* \forall (p, q) \in \mathbb{Z}^2$.

Théorème 3.1 (Maximisation du RSB)

Soit un système de transmission multiporteuse FTN linéaire tel que défini dans la partie 1.4.2. On considère la transmission d'une séquence de symboles centrés IID $\mathbf{c} = \{c_{m,n}\}_{(m,n) \in \mathbb{Z}^2}$ de variance σ_c^2 , sur un canal BABG dont on notera σ_b^2 la variance du bruit (centré).

Notons $\mathbf{g} = \{g_{m,n}\}_{(m,n) \in \mathbb{Z}^2}$ la famille de Gabor d'émission, de prototype $g(t) \in \mathcal{L}_2(\mathbb{R})$, $\check{\mathbf{g}} = \{\check{g}_{m,n}\}_{(m,n) \in \mathbb{Z}^2}$ la famille de réception, et $\rho > 1$ la densité du système.

Alors, le RSB est maximisé lorsque $\check{g}_{m,n}(t) = \mu_{m,n} g_{m,n}(t)$, avec $\mu_{m,n} \in \mathbb{C}^* \forall (m, n) \in \mathbb{Z}^2$.

Le RSB maximal est alors donné par $\text{RSB}_{\max} = \sigma_c^2 \|g\|^2 / \sigma_b^2$.

3.2.2 Prototypes maximisant le RSI

Pour maximiser le RSI, il convient tout d'abord d'observer que l'ensemble des termes $\langle g_{m,n}; \check{g}_{p,q} \rangle \forall (m, n) \in \mathbb{Z}^2$, peuvent être obtenus à l'aide de l'opérateur d'analyse défini dans la partie 1.4.1 :

$$\{\langle g_{m,n}; \check{g}_{p,q} \rangle\}_{(m,n) \in \mathbb{Z}^2} = \mathbf{T}_{\mathbf{g}} \check{g}_{p,q}(t), \quad (p, q) \in \mathbb{Z}^2. \quad (3.20)$$

Dans ce contexte, le terme $\sum_{(m,n) \in \mathbb{Z}^2} |\langle \check{g}_{p,q}; g_{m,n} \rangle|^2 = \sum_{(m,n) \in \mathbb{Z}^2} |\langle g_{m,n}; \check{g}_{p,q} \rangle|^2$ correspond à la norme de l'opérateur d'analyse appliqué à $\check{g}_{p,q}(t)$. On peut donc réécrire le RSI sous la forme suivante :

$$\text{RSI}_{p,q} = \frac{|\langle \check{g}_{p,q}; g_{p,q} \rangle|^2}{\|\mathbf{T}_{\mathbf{g}} \check{g}_{p,q}\|^2 - |\langle \check{g}_{p,q}; g_{p,q} \rangle|^2}. \quad (3.21)$$

Décomposons maintenant $\langle \check{g}_{p,q}; g_{p,q} \rangle$ de manière à faire apparaître un terme en $\mathbf{T}_{\mathbf{g}} \check{g}_{p,q}$:

$$\langle \check{g}_{p,q}; g_{p,q} \rangle = \langle \check{g}_{p,q}; \mathbf{S}_{\mathbf{g}} \mathbf{S}_{\mathbf{g}}^{-1} g_{p,q} \rangle \quad (3.22)$$

$$= \langle \check{g}_{p,q}; \mathbf{T}_{\mathbf{g}}^* \mathbf{T}_{\mathbf{g}} \mathbf{S}_{\mathbf{g}}^{-1} g_{p,q} \rangle \quad (3.23)$$

$$= \langle \mathbf{T}_{\mathbf{g}} \check{g}_{p,q}; \mathbf{T}_{\mathbf{g}} \mathbf{S}_{\mathbf{g}}^{-1} g_{p,q} \rangle, \quad (3.24)$$

et appliquons-lui la relation de Cauchy–Schwarz :

$$|\langle \check{g}_{p,q}; g_{p,q} \rangle|^2 \leq \|\mathbf{T}_{\mathbf{g}} \check{g}_{p,q}\|^2 \|\mathbf{T}_{\mathbf{g}} \mathbf{S}_{\mathbf{g}}^{-1} g_{p,q}\|^2 \Leftrightarrow \|\mathbf{T}_{\mathbf{g}} \check{g}_{p,q}\|^2 \geq \frac{|\langle \check{g}_{p,q}; g_{p,q} \rangle|^2}{\|\mathbf{T}_{\mathbf{g}} \mathbf{S}_{\mathbf{g}}^{-1} g_{p,q}\|^2}, \quad (3.25)$$

3.2. Prototypes maximisant le RSIB

on obtient donc une borne inférieure pour $\|\mathbf{T}_{\mathbf{g}}\check{g}_{p,q}\|^2$, qui est atteinte lorsque $\mathbf{T}_{\mathbf{g}}\check{g}_{p,q}$ et $\mathbf{T}_{\mathbf{g}}\mathbf{S}_{\mathbf{g}}^{-1}g_{p,q}$ sont proportionnels. Remarquons que maximiser le RSI est équivalent à minimiser son inverse. Ce problème d'optimisation s'écrit :

$$(\mathbf{g}, \check{\mathbf{g}}) \in \underset{\mathbf{g}, \check{\mathbf{g}}}{\operatorname{argmin}} \frac{\|\mathbf{T}_{\mathbf{g}}\check{g}_{p,q}\|^2 - |\langle \check{g}_{p,q}; g_{p,q} \rangle|^2}{|\langle \check{g}_{p,q}; g_{p,q} \rangle|^2} \quad (3.26)$$

$$\Leftrightarrow (\mathbf{g}, \check{\mathbf{g}}) \in \underset{\mathbf{g}, \check{\mathbf{g}}}{\operatorname{argmin}} \frac{\|\mathbf{T}_{\mathbf{g}}\check{g}_{p,q}\|^2}{|\langle \check{g}_{p,q}; g_{p,q} \rangle|^2}, \quad (3.27)$$

or, d'après (3.25), choisir $\mathbf{T}_{\mathbf{g}}\check{g}_{p,q} = \beta_{p,q}\mathbf{T}_{\mathbf{g}}\mathbf{S}_{\mathbf{g}}^{-1}g_{p,q}$, avec $\beta_{p,q} \in \mathbb{C}^* \forall (p,q) \in \mathbb{Z}^2$, est une solution à ce problème. Réécrivons cette condition :

$$\mathbf{T}_{\mathbf{g}}\check{g}_{p,q}(t) = \beta_{p,q}\mathbf{T}_{\mathbf{g}}\mathbf{S}_{\mathbf{g}}^{-1}g_{p,q}(t) \quad (3.28)$$

$$\Leftrightarrow \mathbf{S}_{\mathbf{g}}\check{g}_{p,q}(t) = \beta_{p,q}\mathbf{S}_{\mathbf{g}}\mathbf{S}_{\mathbf{g}}^{-1}g_{p,q}(t) \quad (3.29)$$

$$\Leftrightarrow \mathbf{S}_{\mathbf{g}}\check{g}_{p,q}(t) = \beta_{p,q}g_{p,q}(t) \quad (3.30)$$

$$\Leftrightarrow \check{g}_{p,q}(t) = \beta_{p,q}\mathbf{S}_{\mathbf{g}}^{-1}g_{p,q}(t). \quad (3.31)$$

Le RSI est donc maximisé lorsque $\check{\mathbf{g}}$ est proportionnelle à la frame duale canonique de $\mathbf{g}$. Remarquons maintenant que :

- $|\langle \check{g}_{p,q}; g_{p,q} \rangle|^2 = |\langle g_{p,q}; \check{g}_{p,q} \rangle|^2 = |\beta_{p,q}|^2 |\langle g_{p,q}; \mathbf{S}_{\mathbf{g}}^{-1}g_{p,q} \rangle|^2$;
- $\|\mathbf{T}_{\mathbf{g}}\check{g}_{p,q}\|^2 = |\beta_{p,q}|^2 \|\mathbf{T}_{\mathbf{g}}\mathbf{S}_{\mathbf{g}}^{-1}g_{p,q}\|^2$;
- $\|\mathbf{T}_{\mathbf{g}}\mathbf{S}_{\mathbf{g}}^{-1}g_{p,q}\|^2 = \langle \mathbf{T}_{\mathbf{g}}\mathbf{S}_{\mathbf{g}}^{-1}g_{p,q}; \mathbf{T}_{\mathbf{g}}\mathbf{S}_{\mathbf{g}}^{-1}g_{p,q} \rangle = \langle \mathbf{S}_{\mathbf{g}}\mathbf{S}_{\mathbf{g}}^{-1}g_{p,q}; \mathbf{S}_{\mathbf{g}}^{-1}g_{p,q} \rangle = \langle g_{p,q}; \mathbf{S}_{\mathbf{g}}^{-1}g_{p,q} \rangle$.

De plus, $\mathbf{g}$ et $\check{\mathbf{g}}$ étant des frames de Gabor duales canoniques, elles partagent le même réseau temps fréquence, de telle manière que $\langle g_{p,q}; \mathbf{S}_{\mathbf{g}}^{-1}g_{p,q} \rangle = \langle g; \mathbf{S}_{\mathbf{g}}^{-1}g \rangle$. Enfin, $\mathbf{g}$ et $\mathbf{S}_{\mathbf{g}}^{-1}\mathbf{g}$ étant duales alors, en vertu du théorème de Wexler–Raz (théorème 1.1), on a $\langle g; \mathbf{S}_{\mathbf{g}}^{-1}g \rangle = 1/\rho$. Ainsi, le RSI maximal est indépendant de l'indice temporel q et fréquentiel p , et s'écrit :

$$\text{RSI}_{\max} = \frac{1}{\rho - 1}. \quad (3.32)$$

Théorème 3.2 (Maximisation du RSI)

Soit un système de transmission multiporteuse FTN linéaire tel que défini dans la partie 1.4.2. Alors, en se plaçant dans les mêmes conditions que dans le théorème 3.1, le RSI est maximisé lorsque $\check{\mathbf{g}}$ est proportionnelle à la frame duale canonique de $\mathbf{g}$: $\check{g}_{m,n}(t) = \beta_{m,n}\mathbf{S}_{\mathbf{g}}^{-1}g_{m,n}(t)$, avec $\beta_{m,n} \in \mathbb{C}^ \forall (m,n) \in \mathbb{Z}^2$.*

Le RSI maximal est alors donné par $\text{RSI}_{\max} = 1/(\rho - 1)$.

3.2.3 Prototypes maximisant conjointement le RSB et le RSI

On remarque qu'il est possible d'associer les conditions des théorèmes 3.1 et 3.2 :

$$\begin{cases} \check{g}_{m,n}(t) = \beta_{m,n}\mathbf{S}_{\mathbf{g}}^{-1}g_{m,n}(t) \\ \check{g}_{m,n}(t) = \mu_{m,n}g_{m,n}(t) \end{cases} \Leftrightarrow \begin{cases} \mathbf{S}_{\mathbf{g}}g_{m,n}(t) = \frac{\beta_{m,n}}{\mu_{m,n}}g_{m,n}(t) \\ \check{g}_{m,n}(t) = \mu_{m,n}g_{m,n}(t) \end{cases}, \quad (3.33)$$

posons $\beta_{m,n}/\mu_{m,n} = \beta, \forall (m,n) \in \mathbb{Z}^2$:

$$\begin{cases} \mathbf{S}_{\mathbf{g}} g_{m,n}(t) &= \beta g_{m,n}(t) \\ \check{g}_{m,n}(t) &= \mu_{m,n} g_{m,n}(t) \end{cases} . \quad (3.34)$$

On obtient ainsi que le RSIB est maximisé lorsque la famille d'émission est une frame étroite, et lorsque la famille de réception est proportionnelle à la famille d'émission. Dans ce cas, le RSIB maximal est donné par :

$$\text{RSIB}_{\max} = \frac{1}{\text{RSI}_{\max}^{-1} + \text{RSB}_{\max}^{-1}} = \frac{1}{\rho - 1 + \sigma_b^2/(\sigma_c^2 \|g\|^2)} . \quad (3.35)$$

Théorème 3.3 (Maximisation du RSIB)

Soit un système de transmission multiporteuse FTN linéaire tel que défini dans la partie 1.4.2. Alors, en se plaçant dans les mêmes conditions que dans le théorème 3.1, le RSIB est maximisé lorsque :

1. $\mathbf{g}$ est une frame étroite ;
2. $\check{g}_{m,n}(t) = \mu_{m,n} g_{m,n}(t)$, avec $\mu_{m,n} \in \mathbb{C}^* \forall (m,n) \in \mathbb{Z}^2$.

Le RSIB maximal est alors donné par (3.35).

3.3 Analyse statistique du terme d'auto-interférence

Dans cette partie nous nous attardons sur les propriétés statistiques de la variable aléatoire correspondant à l'interférence, lorsque les conditions permettant de maximiser le RSI sont respectées (voir théorème 3.2). La connaissance de cette distribution permettra alors de donner une formule théorique de la probabilité d'erreur binaire en sortie du récepteur multiporteuse linéaire en fonction du RSIB. De plus, dans le chapitre 4, on se basera sur la connaissance de la densité de probabilité de l'interférence pour calculer les probabilités *a posteriori* des symboles, intervenant dans le calcul des LRV.

Rappelons l'expression de la variable aléatoire $i_{p,q}, (p,q) \in \Lambda$:

$$i_{p,q} = \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} c_{m,n} \langle \check{g}_{p,q}; g_{m,n} \rangle , \quad (3.36)$$

où les termes $c_{m,n}$ sont IID et suivent une loi discrète complexe circulaire uniforme à valeurs dans l'alphabet de modulation $\mathcal{A}$, d'espérance nulle et de variance σ_c^2 . En observant cette expression, on observe qu'il s'agit de la somme de variables aléatoires indépendantes et suivant le même type de loi (loi uniforme discrète), mais avec des paramètres différents, à cause du facteur $\langle \check{g}_{p,q}; g_{m,n} \rangle$. Par conséquent, il n'est pas possible d'appliquer le théorème central limite, et la loi suivie par l'interférence ne converge pas vers une loi normale, comme le confirment les tests de normalité du tableau 3.1. Cependant, afin d'évaluer la probabilité d'erreur binaire à l'aide de formules classiques (voir partie 3.4.2) ou encore, dans un contexte de turbo-égalisation linéaire, de pouvoir utiliser le convertisseur SISO de sortie décrit en partie 2.4.2.2, il est important d'évaluer

3.3. Analyse statistique du terme d'auto-interférence

la pertinence d'une approximation de la loi suivie par l'interférence via une loi normale complexe circulaire centrée.

Dans la suite, nous commencerons par démontrer que la loi suivie par l'interférence est centrée, que ses parties réelle et imaginaire sont de variances égales, et de corrélation nulle. Enfin, en déterminant les conditions pour lesquelles la densité de la partie réelle et imaginaire de l'interférence sont bien approchées par une gaussienne, on saura sous quelles conditions la loi suivie par l'interférence est bien approchée par une loi normale complexe circulaire centrée.

3.3.1 Espérance

On montre facilement la nullité de l'espérance de l'interférence, à partir du fait que les symboles sont centrés :

$$\mathbb{E} \{i_{p,q}\} = \mathbb{E} \left\{ \sum_{(m,n) \in \mathbb{Z}^2 \setminus \{(p,q)\}} c_{m,n} \langle S_{\mathbf{g}}^{-1} g; g_{m-p,n-q} \rangle \right\} \quad (3.37)$$

$$= \sum_{(m,n) \in \mathbb{Z}^2 \setminus \{(p,q)\}} \mathbb{E} \{c_{m,n}\} \langle S_{\mathbf{g}}^{-1} g; g_{m-p,n-q} \rangle \quad (3.38)$$

$$= 0. \quad (3.39)$$

3.3.2 Égalité de la variance des parties réelles et imaginaires

Dans cette partie, afin d'alléger le développement des calculs, on adopte les notations suivantes :

- $\mathcal{R} \{c_{m,n}\} = c_{m,n}^R, \forall (m,n) \in \mathbb{Z}^2$;
- $\mathcal{I} \{c_{m,n}\} = c_{m,n}^I, \forall (m,n) \in \mathbb{Z}^2$;
- $\mathcal{R} \{\langle \check{g}_{p,q}; g_{m,n} \rangle\} = \alpha_{m,n,p,q}^R, \forall (m,n), (p,q) \in \mathbb{Z}^2$;
- $\mathcal{I} \{\langle \check{g}_{p,q}; g_{m,n} \rangle\} = \alpha_{m,n,p,q}^I, \forall (m,n), (p,q) \in \mathbb{Z}^2$.

Calculons tout d'abord la variance de la partie réelle de l'interférence :

$$\begin{aligned} & \mathbb{E} \left\{ \mathcal{R} \{i_{p,q}\}^2 \right\} \\ &= \mathbb{E} \left\{ \left(\sum_{(m,n) \in \mathbb{Z}^2 \setminus \{(p,q)\}} \mathcal{R} \{c_{m,n} \langle \check{g}_{p,q}; g_{m,n} \rangle\} \right)^2 \right\} \end{aligned} \quad (3.40)$$

$$= \mathbb{E} \left\{ \left(\sum_{(m,n) \in \mathbb{Z}^2 \setminus \{(p,q)\}} (c_{m,n}^R \alpha_{m,n,p,q}^R - c_{m,n}^I \alpha_{m,n,p,q}^I) \right)^2 \right\} \quad (3.41)$$

$$= \sum_{(m,n), (m',n') \in \mathbb{Z}^2 \setminus \{(p,q)\}} \mathbb{E} \left\{ (c_{m,n}^R \alpha_{m,n,p,q}^R - c_{m,n}^I \alpha_{m,n,p,q}^I) (c_{m',n'}^R \alpha_{m',n',p,q}^R - c_{m',n'}^I \alpha_{m',n',p,q}^I) \right\}, \quad (3.42)$$

or, puisque les symboles sont supposées IID, avec partie réelle et imaginaire également IID, on a $\mathbb{E}\{c_{m,n}^R c_{m,n}^R\} = \mathbb{E}\{c_{m,n}^I c_{m,n}^I\} = \sigma_c^2/2\delta_{m-m'}\delta_{n-n'}$, et $\mathbb{E}\{c_{m,n}^R c_{m',n'}^I\} = 0 \ \forall (m,n), (m',n') \in \mathbb{Z}^2$, donc

$$\mathbb{E}\{\mathcal{R}\{i_{p,q}\}^2\} = \sum_{(m,n) \in \mathbb{Z}^2 \setminus \{(p,q)\}} \left(\frac{\sigma_c^2}{2} \alpha_{m,n,p,q}^R{}^2 + \frac{\sigma_c^2}{2} \alpha_{m,n,p,q}^I{}^2 \right) \quad (3.43)$$

$$= \frac{\sigma_c^2}{2} \sum_{(m,n) \in \mathbb{Z}^2 \setminus \{(p,q)\}} |\langle \check{g}_{p,q}; g_{m,n} \rangle|^2. \quad (3.44)$$

Concernant la partie imaginaire, on a :

$$\begin{aligned} & \mathbb{E}\{\mathcal{I}\{i_{p,q}\}^2\} \\ &= \mathbb{E}\left\{\left(\sum_{(m,n) \in \mathbb{Z}^2 \setminus \{(p,q)\}} \mathcal{I}\{c_{m,n} \langle \check{g}_{p,q}; g_{m,n} \rangle\}\right)^2\right\} \end{aligned} \quad (3.45)$$

$$= \mathbb{E}\left\{\left(\sum_{(m,n) \in \mathbb{Z}^2 \setminus \{(p,q)\}} (c_{m,n}^R \alpha_{m,n,p,q}^I - c_{m,n}^I \alpha_{m,n,p,q}^R)\right)^2\right\} \quad (3.46)$$

$$= \sum_{(m,n), (m',n') \in \mathbb{Z}^2 \setminus \{(p,q)\}} \mathbb{E}\{(c_{m,n}^R \alpha_{m,n,p,q}^I - c_{m,n}^I \alpha_{m,n,p,q}^R)(c_{m',n'}^R \alpha_{m',n',p,q}^I - c_{m',n'}^I \alpha_{m',n',p,q}^R)\} \quad (3.47)$$

$$= \sum_{(m,n) \in \mathbb{Z}^2 \setminus \{(p,q)\}} \left(\frac{\sigma_c^2}{2} \alpha_{m,n,p,q}^I{}^2 + \frac{\sigma_c^2}{2} \alpha_{m,n,p,q}^R{}^2 \right) \quad (3.48)$$

$$= \frac{\sigma_c^2}{2} \sum_{(m,n) \in \mathbb{Z}^2 \setminus \{(p,q)\}} |\langle \check{g}_{p,q}; g_{m,n} \rangle|^2 \quad (3.49)$$

$$= \mathbb{E}\{\mathcal{R}\{i_{p,q}\}^2\}. \quad (3.50)$$

Donc la variance de la partie réelle de l'interférence est effectivement égale à la celle de sa partie imaginaire.

3.3. Analyse statistique du terme d'auto-interférence

3.3.3 Décorrélation de la partie réelle et imaginaire

En réutilisant les notations de la partie précédente, on montre qu'il n'y a aucune corrélation entre la partie réelle et imaginaire de l'interférence :

$$\begin{aligned} & \mathbb{E} \{ \mathcal{R} \{ i_{p,q} \} \mathcal{I} \{ i_{p,q} \} \} \\ &= \sum_{(m,n), (m',n') \in \mathbb{Z}^2 \setminus \{(p,q)\}} \mathbb{E} \{ (c_{m,n}^R \alpha_{m,n,p,q}^R - c_{m,n}^I \alpha_{m,n,p,q}^I) (c_{m',n'}^R \alpha_{m',n',p,q}^I + c_{m',n'}^I \alpha_{m',n',p,q}^R) \} \end{aligned} \quad (3.51)$$

$$= \sum_{(m,n) \in \mathbb{Z}^2 \setminus \{(p,q)\}} \left(\frac{\sigma_c^2}{2} \alpha_{m,n,p,q}^R \alpha_{m,n,p,q}^I - \frac{\sigma_c^2}{2} \alpha_{m,n,p,q}^I \alpha_{m,n,p,q}^R \right) \quad (3.52)$$

$$= 0. \quad (3.53)$$

3.3.4 Distribution de la partie réelle et imaginaire

Comme souligné au début de cette partie, l'expression du terme d'interférence (3.36) ne permet pas d'invoquer le théorème central limite, puisque la variance des termes $c_{m,n} \langle \check{g}_{p,q}; g_{m,n} \rangle = c_{m,n} \langle \check{g}_{p,q}; g_{m,n} \rangle$ dans la somme dépend de la fonction d'interambiguïté du prototype d'émission et de réception. Cependant, une approximation gaussienne permettrait de grandement simplifier l'expression du TEB, ainsi que le développement du convertisseur SISO de sortie évoqué dans le chapitre 2.4.2.2 et qui sera utilisé dans le chapitre 4. Dans cette partie, nous évaluerons la pertinence d'une telle approximation.

Nous avons également choisi de restreindre cette étude aux familles respectant le théorème 3.2. Plus particulièrement, on choisit $\check{\mathbf{g}}$ et $\mathbf{g}$ générés à partir des prototypes $\check{g}(t)$ et $g(t)$ tels que $\check{g}(t) = \mathbf{S}_{\mathbf{g}}^{-1} g(t)$. Sachant cela, il sera suffisant d'indiquer le prototype d'émission pour caractériser les familles d'émission et de réception utilisées. D'autre part, pour chaque prototype $g(t)$ étudié et chaque valeur de ρ , toutes les statistiques présentées ici (densités et tests statistiques) ont été effectués à partir de 640×10^3 échantillons du terme d'interférence $i_{p,q}$ obtenus après simulations de transmissions multiporteuses FTN sur canal parfait avec $M = 128$ porteuses, $K = 5000$ symboles multiporteuses et une constellation QPSK.

Observons tout d'abord le tableau 3.1. On remarque que les tests de Kolmogorov–Smirnov et du χ^2 rejettent l'hypothèse que les échantillons de $\mathcal{R} \{ i_{p,q} \}$ et $\mathcal{I} \{ i_{p,q} \}$ proviennent d'une variable aléatoire normale. Cependant, l'étude des figures 3.1 et 3.2 indiquent que, pour $\rho < 2$, les densités de $\mathcal{R} \{ i_{p,q} \}$ et $\mathcal{I} \{ i_{p,q} \}$ (notées respectivement $f_{\mathcal{R} \{ i_{p,q} \}}(x)$ et $f_{\mathcal{I} \{ i_{p,q} \}}(x)$) s'éloignent peu d'une gaussienne. À l'inverse, pour une densité $\rho \geq 2$ on remarque, d'une part, un écart significatif entre les densités de probabilité empiriques et la gaussienne et, d'autre part, des comportements différents en fonction des prototypes utilisés (en particulier à $\rho = 2$), appuyant la dépendance de la loi de $i_{p,q}$ à la fonction d'interambiguïté entre le prototype d'émission et de réception.

3. L'hypothèse nulle que les échantillons de la partie réelle ($\mathcal{R}$) et de la partie imaginaire ($\mathcal{I}$) de l'interférence proviennent d'une loi normale est ici rejetée pour un seuil de signification de 5%.

4. La p -valeur est la probabilité d'obtenir la même valeur (métrique) du test statistique si l'hypothèse nulle était vraie.

TABLE 3.1 – Tests statistiques du χ^2 et de Kolmogorov–Smirnov (KS) pour l’hypothèse nulle que les échantillons de la partie réelle ($\mathcal{R}$) et de la partie imaginaire ($\mathcal{I}$) de l’interférence proviennent d’une loi normale.

$g(t)$	ρ	KS rejeté ³ ?		KS p -valeur ⁴		χ^2 rejeté ?		χ^2 p -valeur	
		$\mathcal{R}$	$\mathcal{I}$	$\mathcal{R}$	$\mathcal{I}$	$\mathcal{R}$	$\mathcal{I}$	$\mathcal{R}$	$\mathcal{I}$
$\text{LTF}_{1/F_0}(t)$	16/15	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$
	4/3	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$
	8/5	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$
$\text{EHB}_{1/F_0}(t)$	16/15	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$
	4/3	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$
	8/5	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$
$\text{RCS}_{1/F_0, \rho-1}(t)$	16/15	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$
	4/3	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$
	8/5	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$
$\text{RCS}_{T_0, \rho-1}(t)$	16/15	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$
	4/3	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$
	8/5	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$
$\Pi_{T_0}(t)$	16/15	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$
	4/3	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$
	8/5	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$
$\Pi_{1/F_0}(t)$	16/15	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$
	4/3	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$
	8/5	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$	Oui	Oui	$< 10^{-6}$	$< 10^{-6}$

3.3. Analyse statistique du terme d'auto-interférence

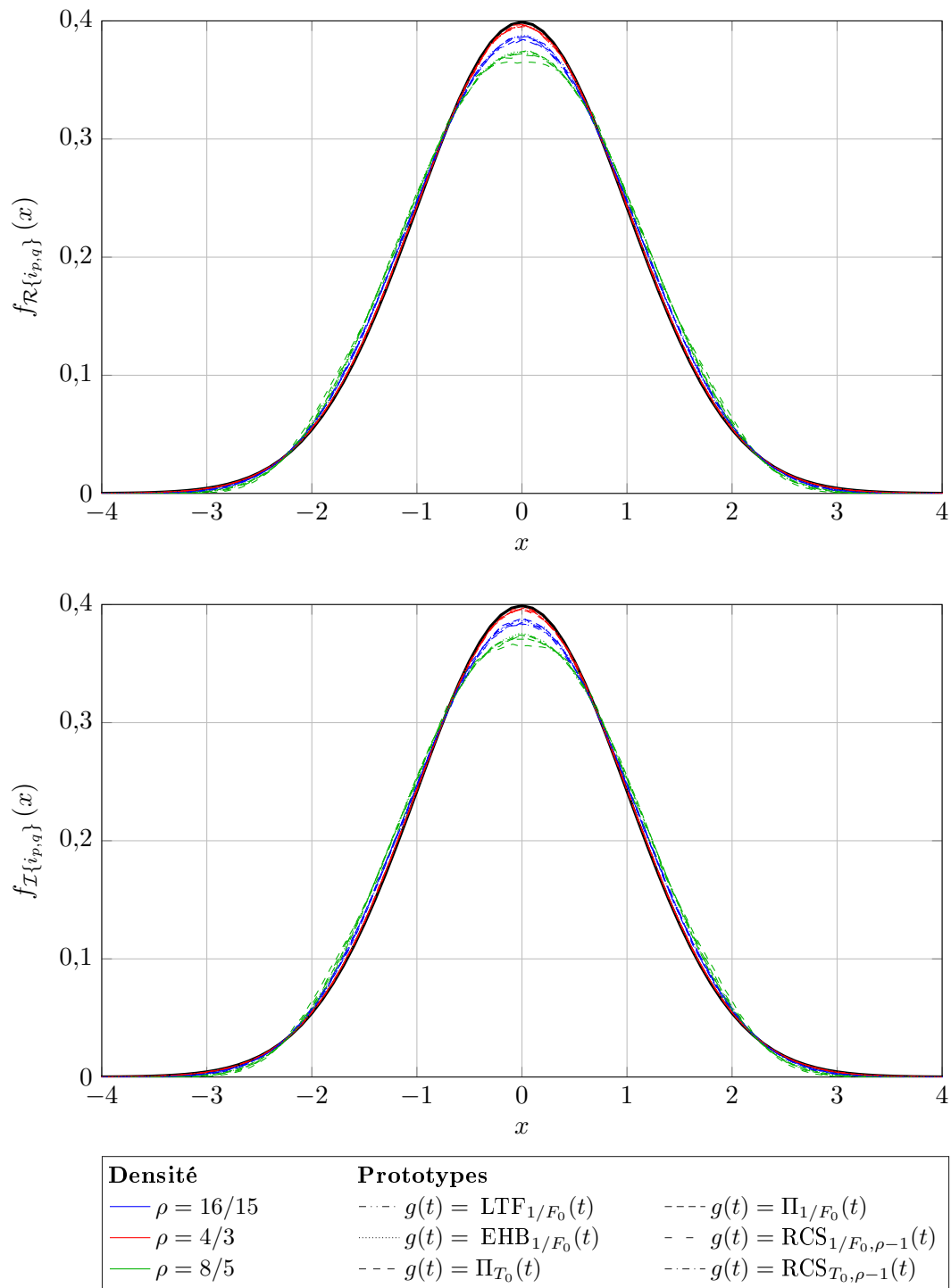


FIGURE 3.1 – Densités de probabilités des parties réelle et imaginaire de l'interférence (centrées et réduites) pour différents prototypes et valeurs de densité $\rho < 2$, comparés à la densité d'une variable aléatoire suivant une loi normale centrée et réduite (en trait plein noir).

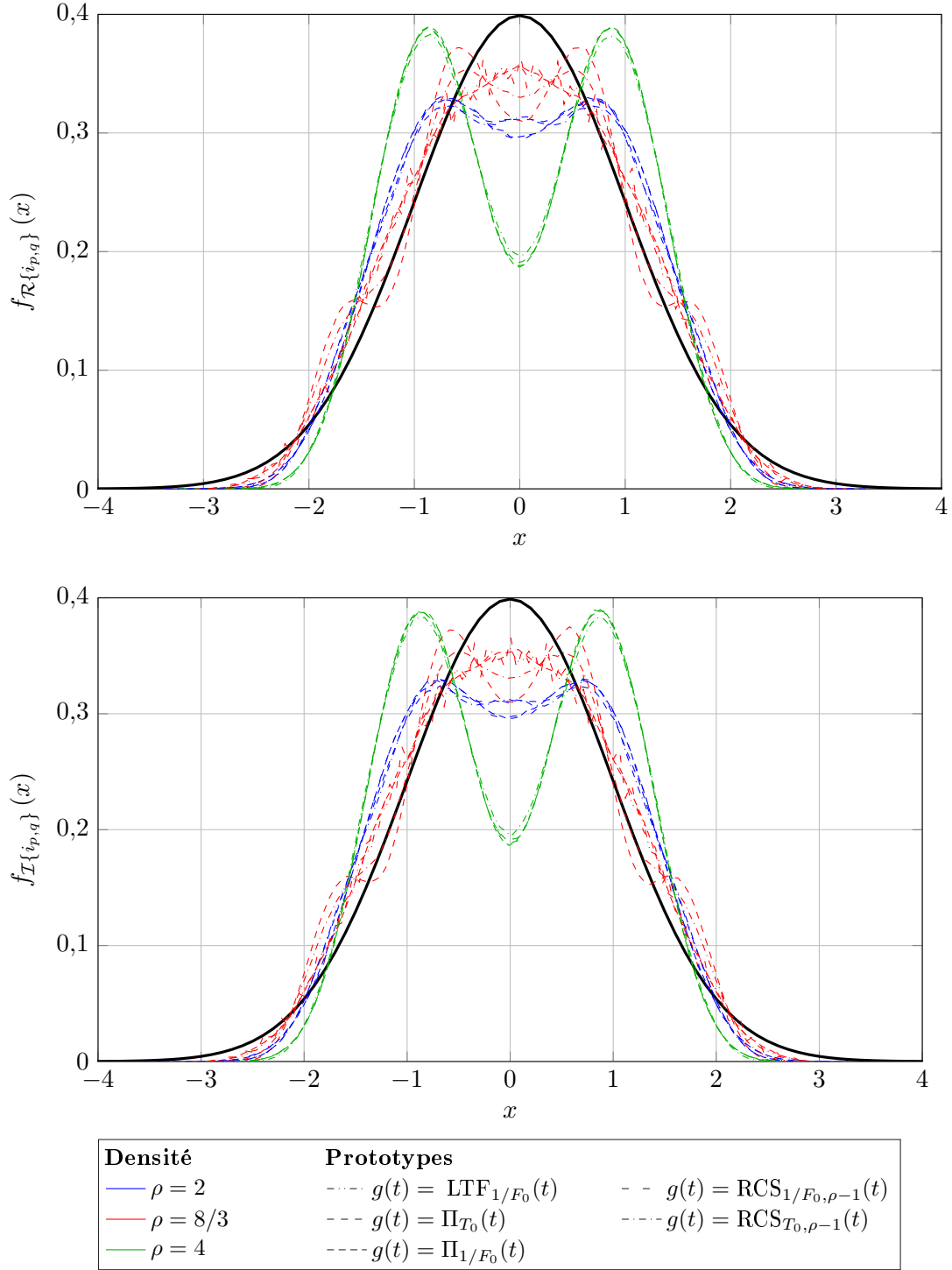


FIGURE 3.2 – Densités de probabilités des parties réelle et imaginaire de l'interférence (centrées et réduites) pour différents prototypes et valeurs de densité $\rho \geq 2$, comparés à la densité d'une variable aléatoire suivant une loi normale centrée et réduite (en trait plein noir).

3.4. Performances théoriques et simulations

3.3.5 Résumé

On conclut ainsi que l'approximation de la loi suivie par l'interférence peut être décentement approchée par une loi normale, centrée, complexe circulaire pour $\rho < 2$ et quel que soit le prototype utilisé. Pour $\rho \geq 2$ par contre, seule une analyse au cas par cas permettrait de déterminer une approximation fiable.

3.4 Performances théoriques et simulations

Dans les deux parties précédentes, nous avons donné les conditions de maximisation du RSB, du RSI et du RSIB, ainsi que des expressions pour la valeur maximale de ces quantités. Dans cette partie, nous nous attarderons tout d'abord à commenter le comportement de ces quantités en fonction de la densité ρ et du rapport $E_s/N_0 = \sigma_c^2 \|g\|^2 / \sigma_b^2$ (avec E_s l'énergie par symbole). Ensuite, nous vérifierons par simulation que seuls les systèmes respectant les conditions énoncées dans les théorèmes 3.1, 3.2 et 3.3 atteignent la borne maximale du RSB, du RSI et du RSIB, respectivement. Enfin, à l'aide de l'expression du RSIB maximal, et grâce aux propriétés statistiques du terme d'interférence exprimées dans la partie 3.3, nous proposerons une méthode pour dériver des formules de calcul de la probabilité d'erreur binaire en fonction du RSIB, à partir des formules classiques des systèmes STN sur canal à BABG. La pertinence de ces formules sera vérifiée par le biais de simulations.

Les simulations présentées dans cette partie se basent sur la transmission de $K = 4096$ symboles multiporteuses sur $M = 512$ porteuses, en utilisant une constellation QPSK, pour différents prototypes et valeurs de densité. Quatre types de relation entre les prototypes d'émission et de réception sont comparés :

- les prototypes d'émission formant des frames étroites, auquel cas on choisit $\check{g}(t) = g(t)$ pour maximiser le RSIB ;
- les prototypes d'émission et de réception formant des paires de frames duales canoniques⁵ (non-étroites, on prendra $\check{g}(t) = \mathbf{S}_{\mathbf{g}}^{-1}g(t)$), et qui maximisent le RSI ;
- les prototypes d'émission proportionnels aux prototypes de réception (on prendra $\check{g}(t) = g(t)$), permettant de maximiser le RSB ;
- les prototypes d'émission et de réception formant des paires de frames duales (ni canoniques ni étroites).

Ainsi, par la suite, plutôt que de préciser la paire de prototypes d'émission et de réception utilisés, on précisera le prototype d'émission $g(t)$ et le type de relation entre ce dernier et le prototype de réception $\check{g}(t)$. Notons enfin deux cas particuliers :

- le cas du prototype en racine de cosinus surélevé de coefficient d'amortissement $\alpha = 1$, et permettant l'absence d'IES en monoporteuse à la cadence $R = F_0$, noté $g(t) = \check{g}(t) = \text{RCS}_{1/F_0,1}(t)$, fait partie des paires de prototypes proportionnels pour $\rho < 2$, mais forme une paire de frames étroites pour $\rho \geq 2$ (on rappelle que la famille de Gabor de paramètres

5. Le prototype de la frame duale canonique $\check{g}(t) = \mathbf{S}_{\mathbf{g}}^{-1}g(t)$ sera alors calculé via la fonction `gabdual` de la toolbox Matlab LTFAT (<http://ltfat.sourceforge.net/doc/gabor/gabdual.php>).

T_0, F_0 et de prototype $\text{RCS}_{1/F_0, \alpha}(t)$ est une frame étroite pour tout $\alpha \leq \rho - 1$, voir partie 1.4.1) ;

- le cas du prototype rectangulaire de largeur $1/F_0$, noté $g(t) = \Pi_{1/F_0}(t)$, associé en réception au prototype permettant de former sa frame duale canonique, noté $\check{g}(t) = S_g^{-1}g(t)$, forme une paire de frames duales canoniques dans le cas général, mais forme une paire de frames étroites lorsque $\rho = 2$.

3.4.1 Comportement en RSB, RSI et RSIB

Observons tout d'abord le comportement théorique du système optimal en termes de RSIB. La figure 3.3 met en valeur la décroissance du RSIB en $1/\rho$. On remarque en particulier que ce dernier passe rapidement de $\text{RSIB}_{\max} = \text{RSB}_{\max}$ quand $\rho = 1$ à $\text{RSIB}_{\max} \leq 0$ dB quand $\rho = 2$. En fonction du rapport E_s/N_0 , on remarque sur la figure 3.4 que même les densités très peu supérieures à 1 (par exemple, $\rho = 16/15 \approx 1,07$), et offrant donc un très faible gain en efficacité spectrale, ont un fort impact sur le RSIB. Enfin, on voit que lorsque $\rho \gg 1$, disposer d'un meilleur RSB ne permet pas de compenser la perte en RSIB due à l'interférence.

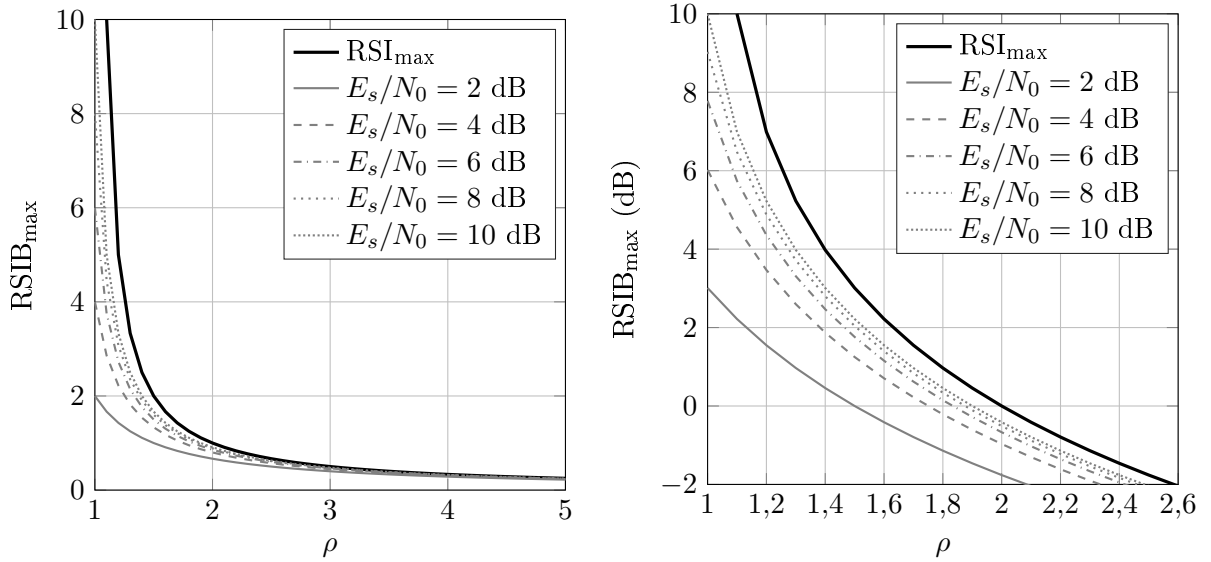


FIGURE 3.3 – Évolution du RSIB maximal avec la densité de la transmission. On observe bien, en échelle linéaire, la décroissance en $1/\rho$.

Vérifions maintenant les prédictions théoriques des théorèmes 3.1, 3.2 et 3.3 par simulation. Observons tout d'abord le comportement du système en fonction du rapport E_s/N_0 , et précisons qu'il est inutile d'observer celui du RSI dans ce cas, puisque par définition (3.16), il ne dépend pas du rapport E_s/N_0 . En termes de RSIB, on remarque sur la figure 3.5 que les frames étroites permettent d'atteindre $\text{RSIB}_{\max}$ pour tout E_s/N_0 , et qu'utiliser une paire de frames duales canoniques ne permet de l'atteindre qu'à fort E_s/N_0 . Les prototypes proportionnels s'éloignent

3.4. Performances théoriques et simulations

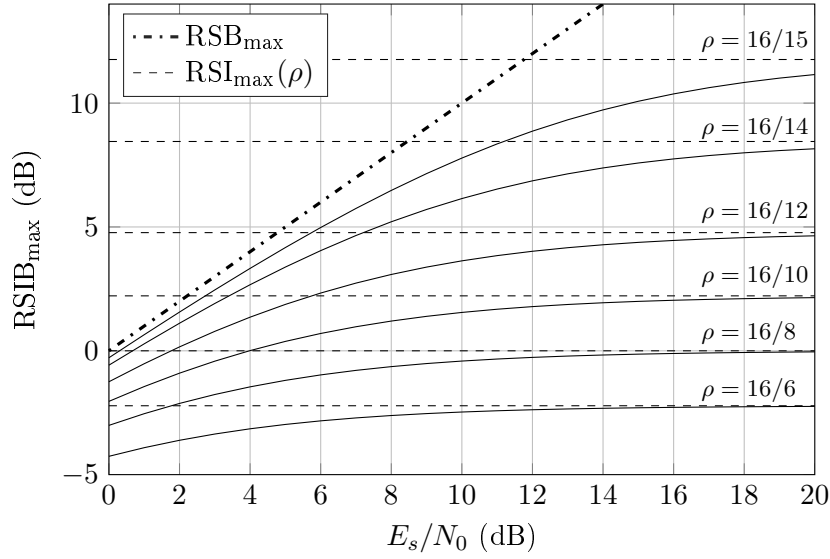


FIGURE 3.4 – Évolution du RSIB maximal avec le rapport E_s/N_0 . On remarque que lorsque ce dernier tend vers 0, $\text{RSIB}_{\max}$ tend vers $\text{RSB}_{\max}$. À l'inverse, $\text{RSIB}_{\max}$ tend rapidement vers l'asymptote horizontale $\text{RSI}_{\max}(\rho)$ lorsque E_s/N_0 augmente.

quant à eux significativement du RSIB optimal à mesure que E_s/N_0 augmente. Sur la figure 3.6, on remarque un comportement linéaire du RSB pour tous les prototypes d'émission et de réception utilisés. Cependant seuls les systèmes où le prototype de réception est proportionnel à celui d'émission (incluant ceux générant des frames étroites) atteignent $\text{RSB}_{\max}$. Notons que sur cette courbe en particulier, les performances de $g(t) = \Pi_{1/F_0}(t)$ et $g(t) = \text{RCS}_{T_0, \rho-1}(t)$ sont identiques, néanmoins, comme le montre la figure 3.8, cela n'est pas forcément le cas selon la densité ρ choisie. Enfin, dans les deux cas, on remarque qu'utiliser en réception une frame duale autre que la duale canonique ne présente pas d'intérêt particulier.

Analysons maintenant le comportement du système en fonction de la densité ρ . La figure 3.7 décrivant l'évolution du RSIB montre que tous les prototypes formant des frames étroites ont un comportement identique et prédictible par la formule du RSIB maximal (3.35). On remarque également que le prototype $g(t) = \text{RCS}_{1/F_0, 1}(t)$, qui ne forme une frame étroite que pour $\rho \geq 2$ présente dans les faits un comportement quasi optimal dès $\rho = 1,6$. De même, le comportement du prototype $g(t) = \Pi_{1/F_0}$ semble atteindre un RSIB très proche de $\text{RSB}_{\max}$, non seulement pour $\rho = 2$ (cas particulier où il forme une frame étroite), mais aussi pour $\rho \geq 2$. Cela s'explique à l'étude de la figure 3.8, où l'on voit que ce prototype (qui maximise le RSI par construction) présente un RSB proche de $\text{RSB}_{\max}$ pour $\rho \geq 2$. Cette figure montre également que seuls les prototypes d'émission et de réception proportionnels (incluant les frames étroites) atteignent le RSB maximal, et que le RSB n'est pas forcément constant vis-à-vis de la densité lorsque les prototypes d'émission et de réception ne sont pas proportionnels. Enfin, la figure 3.9 confirme que seuls les prototypes d'émission et de réception formant une paire de frames duales canoniques

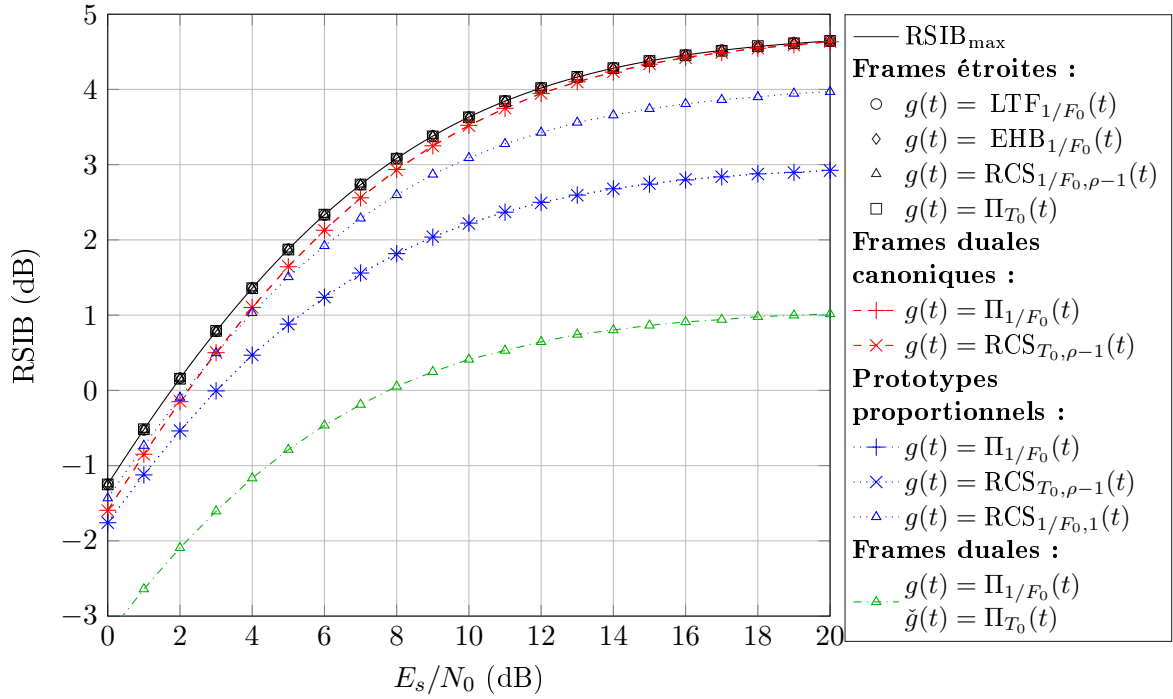


FIGURE 3.5 – Évolution du RSIB en fonction du rapport E_s/N_0 pour différents prototypes, $M = 512$ et $\rho = 4/3$.

3.4. Performances théoriques et simulations

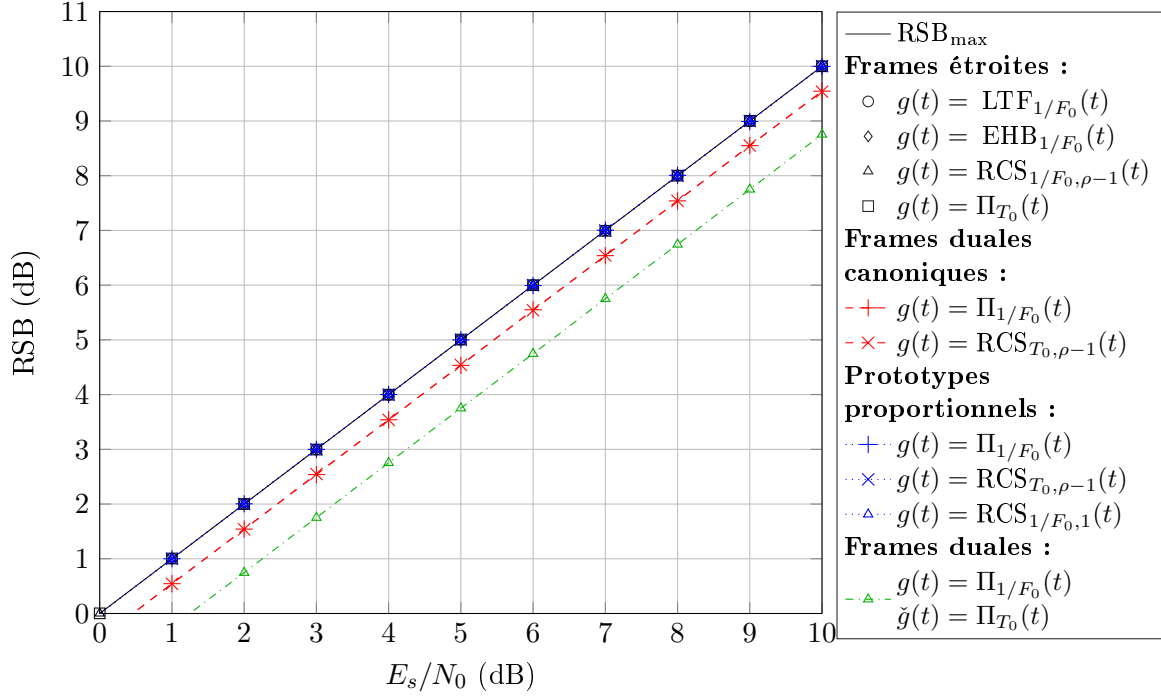


FIGURE 3.6 – Évolution du RSB en fonction du rapport E_s/N_0 pour $M = 512$ et $\rho = 4/3$.

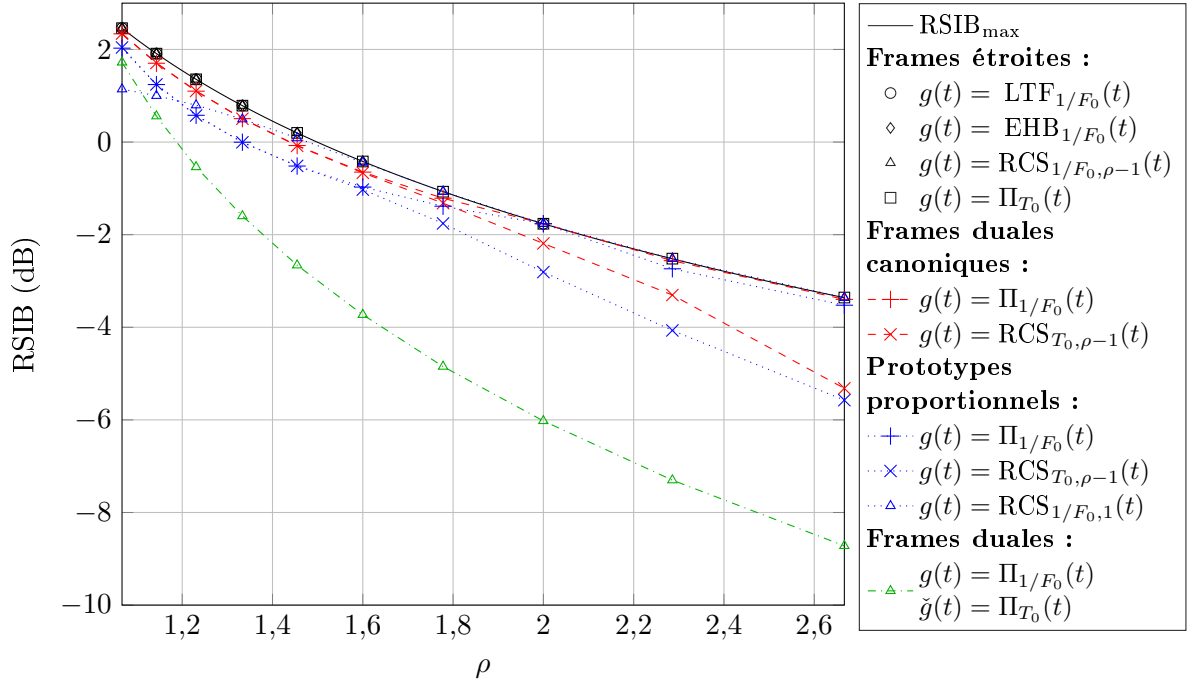
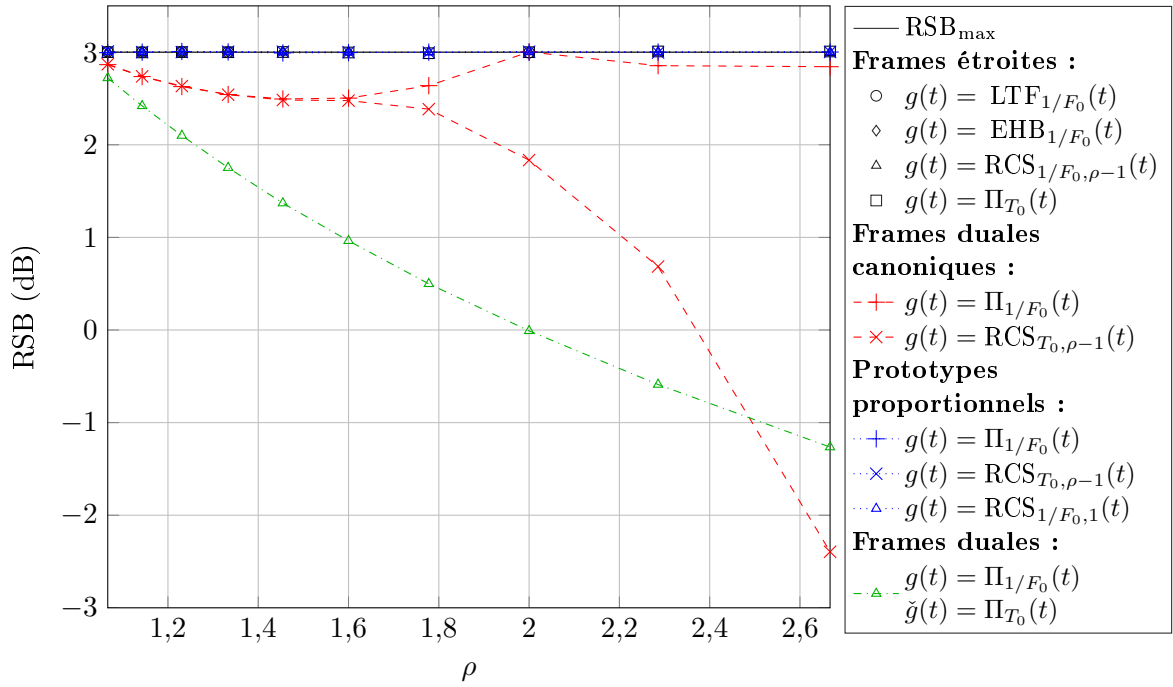
(incluant les frames étroites) atteignent le RSI maximal. On y retrouve les comportements particuliers du prototype $g(t) = \Pi_{1/F_0}(t)$ qui forme une frame étroite à $\rho = 2$, et présente un RSI très proche de $\text{RSI}_{\max}$ pour $\rho \geq 2$, et du prototype $g(t) = \text{RCS}_{1/F_0, 1}$ qui rejoint la courbe du RSI maximal dès $\rho = 1, 6$.

En conclusion, même si certains prototypes ne formant pas de frames étroites peuvent présenter de très bonnes performances pour certaines valeurs de la densité et certains rapports E_s/N_0 , seuls ceux permettant de générer des frames étroites garantissent un comportement prédictible et optimal du RSIB (et, séparément, du RSI et du RSB) pour tout E_s/N_0 et toute densité ρ .

3.4.2 Comportement en TEB

En associant l'hypothèse de normalité du terme d'interférence, telle qu'établie dans la partie 3.3, à celle de normalité du bruit, on peut considérer que la variable aléatoire $\nu_{p,q} = i_{p,q} + \langle \tilde{g}_{p,q}; b \rangle$ suit une loi normale complexe circulaire centrée. Cette hypothèse permet d'établir des formules de probabilité d'erreur pour le système multiporteuse FTN utilisant des frames étroites, en réutilisant les formules classiques établies pour les systèmes STN sur canal à BABG (voir par exemple [CY02]).

Ainsi, dans le cas où on utilise une conversion bit à symbole respectant le codage de Gray et une constellation P -aire, on passe de la formule STN sur canal à BABG (généralement donnée en fonction de E_b/N_0) à la formule s'appliquant au système FTN multiporteuse utilisant des frames

FIGURE 3.7 – Évolution du RSIB en fonction de la densité, pour $M = 512$ et $E_s/N_0 = 3$ dB.FIGURE 3.8 – Évolution du RSB en fonction de la densité, pour $M = 512$ et $E_s/N_0 = 3$ dB.

3.4. Performances théoriques et simulations

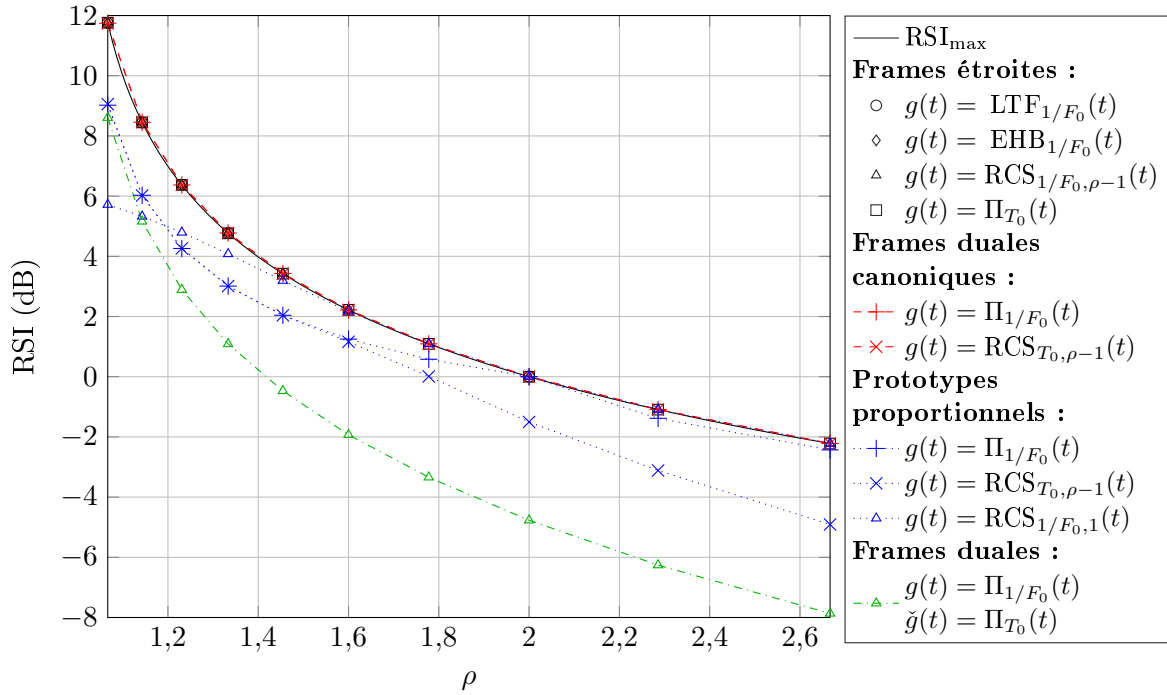


FIGURE 3.9 – Évolution du RSI en fonction de la densité, pour $M = 512$ et $E_s/N_0 = 3$ dB.

étroites, en remplaçant simplement E_b/N_0 par $\text{RSI}_{\max}/\log_2(P)$ ($P = |\mathcal{A}|$ étant le nombre de symboles constituant la constellation). Par exemple, dans le cas d'une constellation QPSK, on obtient :

$$P_{e,\text{QPSK}} = Q(\sqrt{\text{RSI}_{\max}}) = Q\left(\sqrt{\frac{1}{\rho - 1 + \frac{N_0}{E_s}}}\right), \quad (3.54)$$

avec $Q(\cdot)$ la fonction de répartition complémentaire d'une variable aléatoire normale centrée et réduite.

Observons les probabilités d'erreur théoriques en fonction de ρ et du rapport E_b/N_0 sur les figures 3.10 et 3.11. On observe une dégradation très nette de la probabilité d'erreur à mesure que la densité augmente. De plus, sur la figure 3.11, on remarque en particulier que même à $\rho = 16/15 \approx 1,07$, la probabilité d'erreur est au-dessus de 10^{-5} quel que soit le rapport E_s/N_0 . Cela est très éloigné des TEB obtenus avec leurs équivalents STN (on rappelle que 10^{-5} est la probabilité cible utilisée dans ce manuscrit pour comparer les différents formats de modulation, tel que sur la figure 1.18, sur canal à BABG). On y remarque aussi que la probabilité d'erreur tend vers une asymptote horizontale de valeur $P_{e,\text{QPSK}}^{\min} = Q(\sqrt{\text{SIR}_{\max}})$ lorsque le rapport E_b/N_0 tend vers l'infini.

En termes de simulations, on confirme sur la figure 3.12 que l'approximation gaussienne de l'interférence permet une bonne estimation du TEB, en particulier pour $\rho < 2$, et reste pertinente au-delà. Pour les prototypes ne formant pas des frames étroites on observe une dégradation

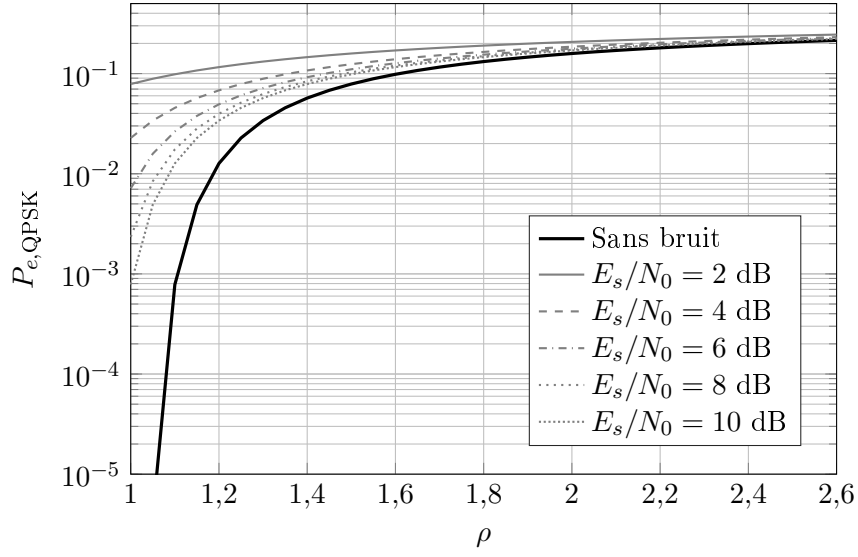


FIGURE 3.10 – Évolution de la probabilité d'erreur en fonction de la densité du système multiporteuse linéaire FTN optimal avec une constellation QPSK, et en supposant l'interférence gaussienne.

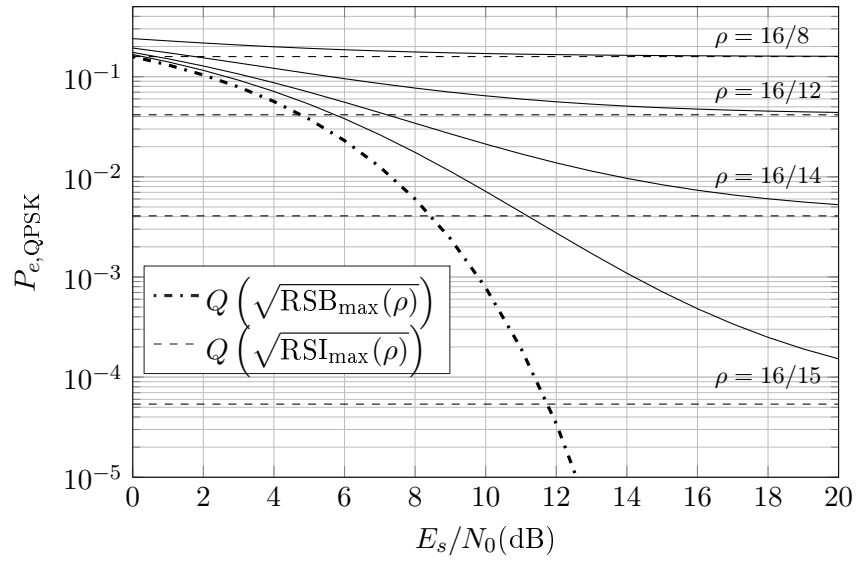


FIGURE 3.11 – Évolution de la probabilité d'erreur en fonction du rapport E_s/N_0 du système multiporteuse linéaire FTN optimal avec une constellation QPSK, et en supposant l'interférence gaussienne.

3.5. Conclusion

plus ou moins importante, mais non prédictible, du TEB. À l'inverse, les prototypes formant des frames étroites engendrent tous les mêmes performances. Sur la figure 3.13 on observe une bonne adéquation entre le TEB des prototypes formant des frames étroites et $P_{e,\text{QPSK}}$ sur toute la gamme de rapports E_s/N_0 . De la même manière que dans la partie précédente, on remarque que les prototypes d'émission et de réception formant une paire de frames duales canoniques rattrapent les performances optimales lorsque le rapport E_s/N_0 devient grand. Enfin, utiliser un prototype d'émission et de réception formant un paire de frames duales (ni étroites, ni canoniques) semble ne présenter, là encore, aucun intérêt.

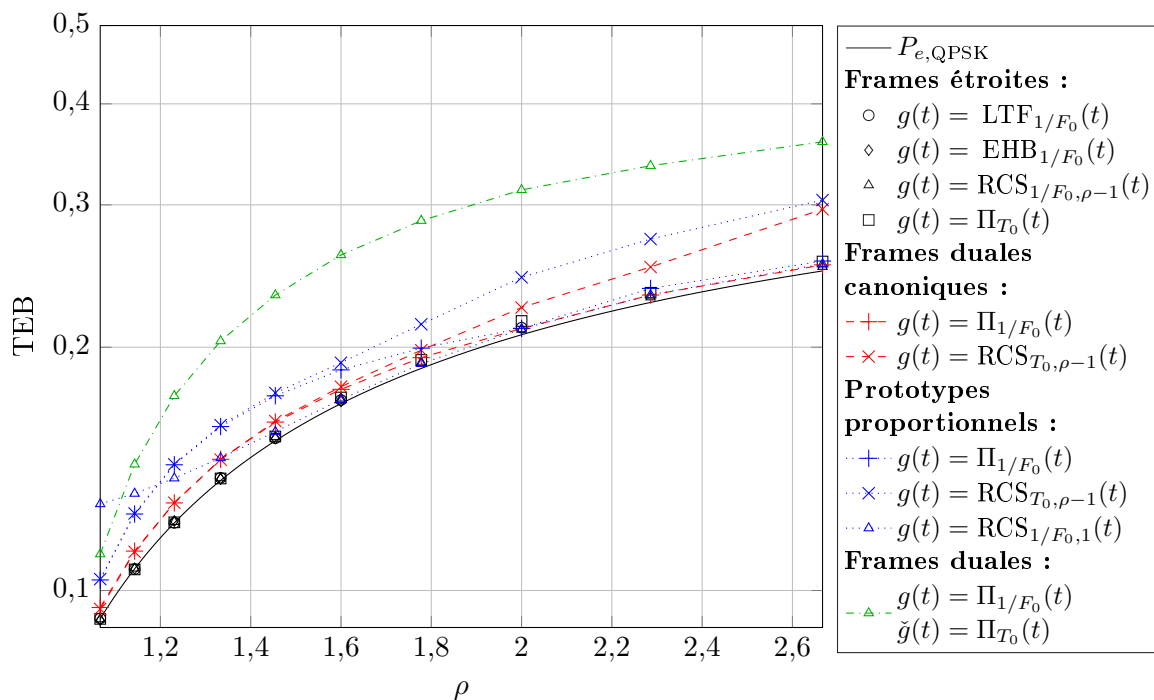


FIGURE 3.12 – Évolution de la probabilité d'erreur en fonction de la densité pour une constellation QPSK, $M = 512$ et $E_s/N_0 = 3$ dB.

3.5 Conclusion

Dans ce chapitre, nous avons vu qu'utiliser des prototypes formant des frames étroites en émission et en réception permet de maximiser le RSIB. Or les prototypes formant des frames étroites pour un espacement entre symboles T_0 et un espacement entre sous-porteuses F_0 , sont les mêmes que ceux permettant de respecter le critère de Nyquist pour un espacement entre symboles $1/F_0$ et un espacement entre sous-porteuses $1/T_0$. Ainsi, parmi les prototypes satisfaisant cette contrainte, on retrouve notamment deux classiques du domaine des communications numériques : le filtre rectangulaire dont la longueur correspond au temps symbole, et le filtre en racine de

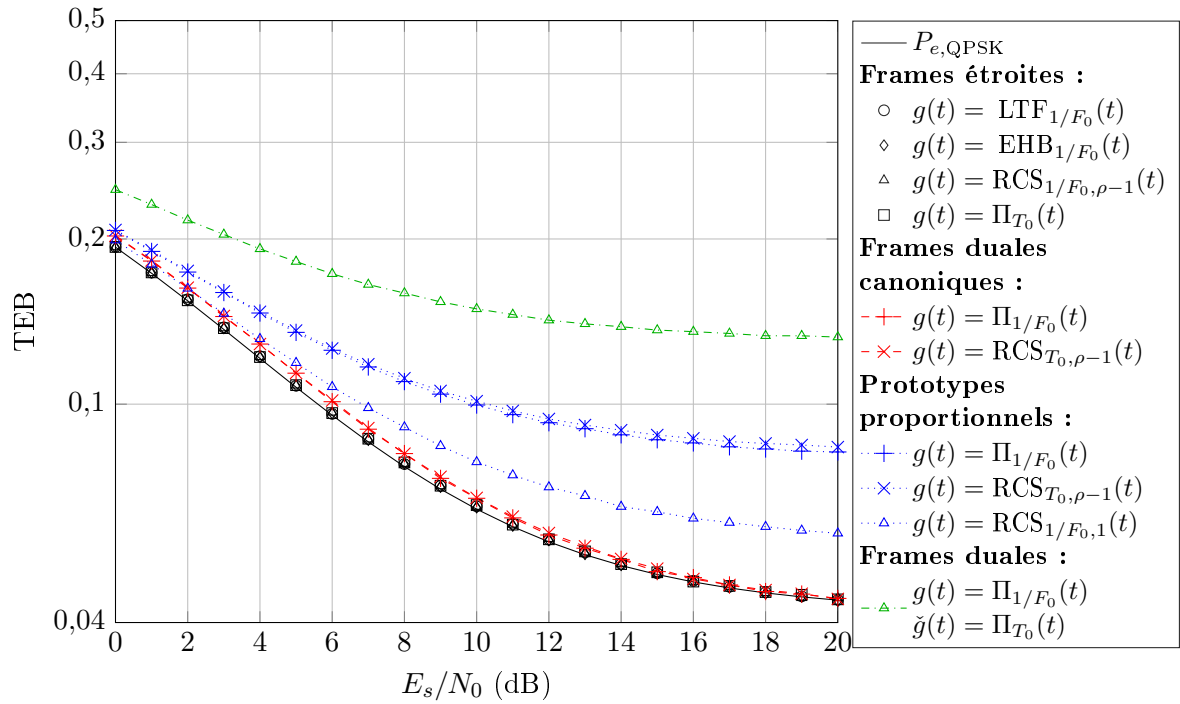


FIGURE 3.13 – Évolution de la probabilité d'erreur en fonction du rapport E_s/N_0 pour une constellation QPSK, $M = 512$ et $\rho = 4/3$.

3.5. Conclusion

cosinus surélevé (qui doit être conçu de la même manière que lorsqu'on souhaite obtenir un système monoporteuse sans IES à la cadence $R = F_0$), dont le coefficient d'amortissement α doit être choisi tel que $\alpha \leq \rho - 1$.

Dans la mesure où l'on a également montré, dans ce cas, que la loi de probabilité suivie par l'interférence est approximativement normale pour des valeurs de densité $\rho < 2$ (rappelons que, strictement parlant, l'interférence ne suit pas une loi normale, il ne s'agit que d'une approximation), cette optimalité en RSIB se traduit en optimalité du TEB. De plus, l'utilisation de ce type de famille d'émission et de réception permet de dériver des expressions théoriques pour le RSIB et la probabilité d'erreur binaire.

Cependant, nous avons aussi observé que les performances en TEB, bien qu'optimales parmi les systèmes linéaires, restent mauvaises comparées aux modulations classiques STN, et ce même pour des densités proches de l'unité. Ce constat constitue la motivation pour la proposition et l'étude des systèmes du chapitre 4, utilisant le principe de la turbo-égalisation pour tendre à supprimer l'interférence.

Publications associées à ce chapitre

Journaux

1. Alexandre Marquet, Cyrille Siclet, Damien Roque, Pierre Siohan, « Analysis of a FTM Multicarrier System : Interference Mitigation Based on Tight Gabor Frames ». In : *EAI Endorsed Transactions on Cognitive Communications*, vol. 17, No 10 (2017), DOI : 10.4108/eai.23-2-2017.152191.
2. Alexandre Marquet, Cyrille Siclet, Damien Roque « Analysis of the faster-than-Nyquist optimal linear multicarrier system ». In : *Comptes Rendus Physique*, vol. 18, issue 2 (février 2017), pp. 167-177, DOI : 10.1016/j.crhy.2016.11.006.
3. Alexandre Marquet, Damien Roque, Cyrille Siclet, Pierre Siohan, « FTM multicarrier transmission based on tight Gabor frames ». In : *EURASIP Journal on Wireless Communications and Networking*, (2017), p. 97, DOI : 10.1186/s13638-017-0878-3.

Conférences

1. Alexandre Marquet, Cyrille Siclet, Damien Roque, Pierre Siohan, « Analysis of a multicarrier communication system based on overcomplete Gabor frames ». In : *Cognitive Radio Oriented Wireless Networks*, Springer International Publishing, p. 387 (2016), Lecture Notes of the Institute for Computer Sciences, Social Informatics and Telecommunications Engineering, DOI : 10.1007/978-3-319-40352-6_32.
2. Alexandre Marquet, Cyrille Siclet, Damien Roque « Analysis of the optimal linear system for multicarrier FTM communications ». In : *Actes des journées scientifiques 2016 de l'URSI-France Energie et radiosciences*, (mars 2016), Cesson-Sévigné, France.

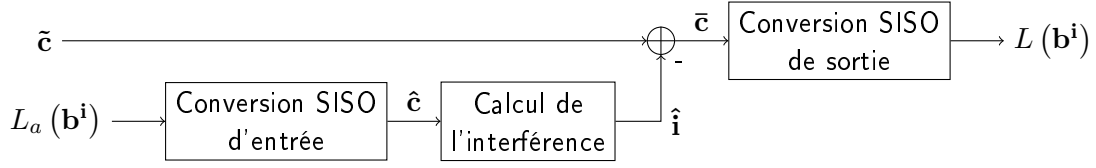
Turbo-égalisation à faible complexité des systèmes multiporteuses au-delà de la cadence de Nyquist

Sommaire

4.1	Introduction	89
4.2	Système MEQM à base de suppression successive d'interférence et de frames étroites en présence de BABG	90
4.2.1	Calcul de l'interférence	90
4.2.2	Prototypes minimisant le RSIB	92
4.2.3	Conditions pour la minimisation de l'erreur quadratique moyenne	94
4.2.4	Conversion SISO de sortie	96
4.2.5	Simulations	97
4.3	Suppression successive d'interférence en présence de canal radio-mobile	102
4.3.1	Calcul de l'interférence	102
4.3.2	Conversion SISO de sortie	103
4.3.3	Simulations	104
4.4	Conclusion	110

4.1 Introduction

Comme nous l'avons vu dans le chapitre 2 (partie 2.4), l'aspect bidimensionnel des modulations multiporteuses (temps et fréquence) rend l'égalisation et, par extension, la turbo-égalisation de l'auto-interférence induit par les modulations FTN difficiles à implémenter sous leurs formes classiques (MAP, MEQM, forçage à zéro). De plus, dans le chapitre précédent, nous avons vu qu'une approche linéaire de la maximisation du RSIB est insuffisante pour obtenir des performances acceptables en termes de TEB dans un contexte FTN. Ainsi, afin d'obtenir des performances et une complexité algorithmique acceptables, nous proposons dans ce chapitre d'étudier un égaliseur à suppression d'interférence avec information *a priori*. Le principe de ce dernier est illustré en figure 4.1 et consiste simplement à estimer le terme d'interférence de manière souple, à l'aide de l'information extrinsèque fournie par un décodeur canal SISO, ce qui implique donc la mise en œuvre d'un code correcteur d'erreur côté émetteur.

FIGURE 4.1 – Structure d'un égaliseur à suppression d'interférence avec information *a priori*.

Ce chapitre est organisé comme suit. Nous verrons dans la première partie que combiner un système à suppression successive d'interférence avec les prototypes d'émission/réception proposés dans le chapitre 3 permet de minimiser l'EQM sur canal à BABG. Nous montrerons ensuite qu'il est possible de profiter des facilités d'adaptation au canal propres aux modulations multiporteuses (évoqués dans le chapitre 1 partie 1.4) pour obtenir de bonnes performances pour un faible surcoût algorithmique, dans le cas de transmissions FTN sur canaux sélectifs en fréquence et, éventuellement, en temps.

4.2 Système MEQM à base de suppression successive d'interférence et de frames étroites en présence de BABG

L'objectif de cette partie est de développer l'égaliseur SISO à suppression successive de l'interférence dans le cadre d'une transmission sur canal à BABG. Tout d'abord, on détaille la manière dont le terme d'interférence $\hat{i}_{p,q}$ ($p, q \in \Lambda$) est estimé à partir de la séquence de symboles souples $\hat{\mathbf{c}} = \{\hat{c}_{m,n}\}_{(m,n) \in \Lambda}$ fournie par le convertisseur d'entrée SISO (présenté en 2.4.2.2, équation (2.38)) à partir d'information extrinsèque (généralement fournie par un décodeur canal). Ensuite, on donnera l'expression de la variance du terme d'interférence résiduelle et de bruit $\nu_{p,q}$, ainsi que le biais $\alpha_{p,q}$ du terme de signal utile, permettant d'utiliser le convertisseur SISO de sortie détaillé en 2.4.2.2, équations (2.43) et (2.44). Enfin, on montrera qu'utiliser les familles d'émission et de réception préconisées dans le chapitre précédent pour maximiser le RSIB permet dans ce contexte de minimiser l'EQM en sortie de l'égaliseur, ainsi que de simplifier l'implémentation du convertisseur SISO de sortie.

4.2.1 Calcul de l'interférence

Rappelons l'expression des symboles en sortie du démodulateur multiporteuse dans le cas d'un canal à BABG :

$$\tilde{c}_{p,q} = \underbrace{c_{p,q} \langle \check{g}_{p,q}; g_{p,q} \rangle}_{\text{Signal utile}} + \underbrace{\sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} c_{m,n} \langle \check{g}_{p,q}; g_{m,n} \rangle}_{\text{Interférence : } i_{p,q}} + \underbrace{\langle \check{g}_{p,q}; b \rangle}_{\text{Bruit filtré}}. \quad (4.1)$$

De cette expression, on extrait le terme d'interférence :

$$i_{p,q} = \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} c_{m,n} \langle \check{g}_{p,q}; g_{m,n} \rangle. \quad (4.2)$$

Ainsi, on estime ce terme d'interférence à partir de la séquence de symboles estimés par :

$$\hat{i}_{p,q} = \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} \hat{c}_{m,n} \langle \check{g}_{p,q}; g_{m,n} \rangle, \quad (4.3)$$

où $\hat{c}_{m,n}$ est défini par (2.38). Observons que cette estimation du terme d'interférence minimise l'erreur quadratique moyenne entre $\hat{i}_{p,q}$ et $i_{p,q}$ sachant les LRV *a priori* $L_a(\mathbf{b}^i)$ sur les bits entrelacés. En effet, en accord avec l'annexe A, l'estimateur minimisant l'EQM, sachant l'observation $L_a(\mathbf{b}^i)$ est donné par :

$$\mathbb{E} \left\{ i_{p,q} | L_a(\mathbf{b}^i) \right\} = \mathbb{E} \left\{ \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} c_{m,n} \langle \check{g}_{p,q}; g_{m,n} \rangle | L_a(\mathbf{b}^i) \right\} \quad (4.4)$$

$$= \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} \mathbb{E} \left\{ c_{m,n} | L_a(\mathbf{b}^i) \right\} \langle \check{g}_{p,q}; g_{m,n} \rangle \quad (4.5)$$

$$= \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} \hat{c}_{m,n} \langle \check{g}_{p,q}; g_{m,n} \rangle. \quad (4.6)$$

Les symboles en sortie de l'anneur d'interférence sont donc donnés par

$$\bar{c}_{p,q} = \tilde{c}_{p,q} - \hat{i}_{p,q} \quad (4.7)$$

$$= \underbrace{c_{p,q} \langle \check{g}_{p,q}; g_{p,q} \rangle}_{\text{Signal utile}} + \underbrace{\sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} (c_{m,n} - \hat{c}_{m,n}) \langle \check{g}_{p,q}; g_{m,n} \rangle}_{\text{Interférence résiduelle : } i_{p,q} - \hat{i}_{p,q}} + \underbrace{\langle \check{g}_{p,q}; b \rangle}_{\text{Bruit filtré}}. \quad (4.8)$$

Remarquons que lorsqu'aucune estimation des symboles émis n'est disponible, alors le terme d'interférence estimé est nul. Dans le contexte d'un turbo-égaliseur, cela signifie qu'à la première itération, la seule fonction de l'égaliseur à suppression successive d'interférence est de calculer les LRV à destination du décodeur canal. À l'inverse, en présence d'une estimation parfaitement fiable des symboles émis, le terme d'interférence est complètement annulé, et les performances du turbo-égaliseur correspondent alors aux performances d'un système STN.

Enfin, il est possible de calculer la séquence $\hat{\mathbf{i}} = \left\{ \hat{i}_{m,n} \right\}_{(m,n) \in \Lambda}$ de termes d'interférences estimés de manière efficace en fournissant la séquence de symboles $\hat{\mathbf{c}} = \{ \hat{c}_{m,n} \}_{(m,n) \in \Lambda}$ à un modulateur multiporteuse suivi d'un démodulateur multiporteuse et en retranchant les termes $\hat{c}_{m,n} \langle \check{g}_{m,n}; g_{m,n} \rangle$:

$$\hat{\mathbf{i}} = \text{DEMOD} \{ \text{MOD} \{ \hat{\mathbf{c}} \} (t) \} - \{ \hat{c}_{m,n} \langle \check{g}_{m,n}; g_{m,n} \rangle \}_{(m,n) \in \Lambda}, \quad (4.9)$$

avec $\text{MOD}\{\cdot\}$ et $\text{DEMOD}\{\cdot\}$ les opérations de modulation et de démodulation multiporteuse telles que définies par (1.56) et (1.57) à temps continu, ou par (1.63) et (1.65) à temps discret. De par la linéarité du récepteur multiporteuse, on peut finalement mutualiser l'opération de démodulation dans le calcul de l'interférence, et celle du signal reçu $r(t)$, de manière à simplifier encore le récepteur :

$$\bar{\mathbf{c}} = \text{DEMOD} \{ r(t) - \text{MOD} \{ \hat{\mathbf{c}} \} (t) \} - \{ \hat{c}_{m,n} \langle \check{g}_{m,n}; g_{m,n} \rangle \}_{(m,n) \in \Lambda}, \quad (4.10)$$

avec $\bar{\mathbf{c}} = \{ \bar{c}_{m,n} \}_{(m,n) \in \Lambda}$.

4.2.2 Prototypes minimisant le RSIB

Après avoir détaillé le calcul de l'interférence souple, on peut montrer qu'utiliser des prototypes d'émission et de réception engendrant des frames étroites permet ici aussi de maximiser le RSIB après annulation d'interférence. Nous en déduisons alors les conditions permettant de minimiser l'EQM.

Développons tout d'abord l'expression du RSIB après annulation d'interférence. Pour ce faire, de la même manière qu'en 2.4.2.2, équation (2.41), on note :

$$\bar{c}_{p,q} = \alpha_{p,q} c_{p,q} + \nu_{p,q}, \quad (4.11)$$

où $\alpha_{p,q} = \langle \check{g}_{p,q}; g_{p,q} \rangle$ et $\nu_{p,q}$ prend en compte l'interférence résiduelle et le bruit :

$$\nu_{p,q} = i_{p,q} - \hat{i}_{p,q} + \langle \check{g}_{p,q}; b \rangle \quad (4.12)$$

$$= \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} (c_{p,q} - \hat{c}_{p,q}) \langle \check{g}_{p,q}; g_{m,n} \rangle + \langle \check{g}_{p,q}; b \rangle. \quad (4.13)$$

Dans ce contexte, et sachant que la variance du terme utile est donnée en 3.2, équation (3.3), le RSIB prend la forme suivante :

$$\text{RSIB}_{p,q} = \frac{\sigma_c^2 |\langle \check{g}_{p,q}; g_{p,q} \rangle|^2}{\sigma_{\nu_{p,q}}^2}, \quad (4.14)$$

avec $\sigma_{\nu_{p,q}}^2 = \mathbb{E} \{ |\nu_{p,q}|^2 \}$. Calculons cette expression :

$$\sigma_{\nu_{p,q}}^2 = \mathbb{E} \{ |\nu_{p,q}|^2 \} = \mathbb{E} \{ |(i_{p,q} - \hat{i}_{p,q}) + \langle \check{g}_{p,q}; b \rangle|^2 \} \quad (4.15)$$

$$= \mathbb{E} \{ |i_{p,q} - \hat{i}_{p,q}|^2 \} + \mathbb{E} \{ |\langle \check{g}_{p,q}; b \rangle|^2 \} + 2\mathcal{R} \left\{ \mathbb{E} \left\{ (i_{p,q} - \hat{i}_{p,q}) \langle \check{g}_{p,q}; b \rangle^* \right\} \right\}, \quad (4.16)$$

où $\mathbb{E} \{ |\langle \check{g}_{p,q}; b \rangle|^2 \}$ est donné en 3.2, équation (3.10). Développons le terme $\mathbb{E} \{ |i_{p,q} - \hat{i}_{p,q}|^2 \}$:

$$\mathbb{E} \{ |i_{p,q} - \hat{i}_{p,q}|^2 \} \quad (4.17)$$

$$\begin{aligned} &= \mathbb{E} \left\{ \left| \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} (c_{m,n} - \hat{c}_{m,n}) \langle \check{g}_{p,q}; g_{m,n} \rangle \right|^2 \right\} \\ &= \sum_{(m,n), (m',n') \in \Lambda \setminus \{(p,q)\}} \mathbb{E} \{ (c_{m,n} - \hat{c}_{m,n})(c_{m',n'}^* - \hat{c}_{m',n'}^*) \} \langle \check{g}_{p,q}; g_{m,n} \rangle \langle \check{g}_{p,q}; g_{m',n'} \rangle^*, \end{aligned} \quad (4.18)$$

or on peut montrer que $\mathbb{E} \{ c_{m,n} \hat{c}_{m',n'}^* \} = \mathbb{E} \{ \hat{c}_{m,n} c_{m',n'}^* \} = \mathbb{E} \{ \hat{c}_{m,n} \hat{c}_{m',n'}^* \}$ [LLL05]. Ainsi, on a :

$$\mathbb{E} \{ (c_{m,n} - \hat{c}_{m,n})(c_{m',n'}^* - \hat{c}_{m',n'}^*) \} = \mathbb{E} \{ c_{m,n} c_{m',n'}^* \} - \mathbb{E} \{ \hat{c}_{m,n} \hat{c}_{m',n'}^* \}. \quad (4.19)$$

4.2. Système MEQM à base de suppression successive d'interférence et de frames étroites en présence de BABG 93

En supposant que les symboles émis tout comme les symboles estimés sont IID (ce qui nécessite en pratique un entrelaceur en émission, comme l'illustre la figure 2.3 [LLL05]), de variances respectives σ_c^2 et $\sigma_{\hat{c}}^2$, on obtient :

$$\mathbb{E} \left\{ (c_{m,n} - \hat{c}_{m,n})(c_{m',n'}^* - \hat{c}_{m',n'}^*) \right\} = \sigma_c^2 \delta_{m-m'} \delta_{n-n'} - \sigma_{\hat{c}}^2 \delta_{m-m'} \delta_{n-n'}. \quad (4.20)$$

Enfin, en injectant (4.24) dans (4.20), on arrive à :

$$\mathbb{E} \left\{ |i_{p,q} - \hat{i}_{p,q}|^2 \right\} = (\sigma_c^2 - \sigma_{\hat{c}}^2) \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} |\langle \check{g}_{p,q}; g_{m,n} \rangle|^2. \quad (4.21)$$

Notons qu'il est extrêmement difficile d'obtenir une valeur analytique pour $\sigma_{\hat{c}}^2$. En effet, cette variance dépend des propriétés statistiques des LRV *a priori*, lesquels dépendent eux-mêmes, dans un turbo-égaliseur, de la relation d'entrée/sortie du décodeur utilisé. En pratique, on estimera donc cette variance de la manière suivante :

$$\sigma_{\hat{c}}^2 = \frac{1}{|\Lambda|} \sum_{(m,n) \in \Lambda} |\hat{c}_{m,n}|^2, \quad (4.22)$$

avec $|\Lambda|$ le cardinal de Λ . Passons enfin au dernier terme :

$$\mathbb{E} \left\{ (i_{p,q} - \hat{i}_{p,q}) \langle \check{g}_{p,q}; b \rangle^* \right\} = \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} (\mathbb{E} \{ c_{m,n} \langle \check{g}_{p,q}; b \rangle^* \} - \mathbb{E} \{ \hat{c}_{m,n} \langle \check{g}_{p,q}; b \rangle^* \}) \langle \check{g}_{p,q}; g_{m,n} \rangle. \quad (4.23)$$

Ici, les symboles émis étant indépendants du bruit, on peut écrire $\mathbb{E} \{ c_{m,n} \langle \check{g}_{p,q}; b \rangle^* \} = 0$. À l'opposé, on ne peut pas considérer les symboles estimés à partir des LRV comme indépendants du bruit dans le cas général. En effet, si dans les deux cas suivants :

- aucune estimation des symboles n'est disponible ($\hat{c}_{m,n} = 0 \ \forall (m,n) \in \Lambda$) ;
- une estimation parfaite des symboles est disponible ($\hat{c}_{m,n} = c_{m,n} \ \forall (m,n) \in \Lambda$) ;

l'expression (4.23) s'annule, il est difficile d'en conclure de même dans les situations intermédiaires, puisque le module des LRV en sortie du décodeur représente la fiabilité du décodage, et est donc lié aux conditions du canal. En pratique, cependant, on observe que négliger ce terme n'a pas d'impact significatif sur la valeur de $\sigma_{\nu_{p,q}}^2$. On le considèrera donc nul par la suite :

$$\mathbb{E} \left\{ (i_{p,q} - \hat{i}_{p,q}) \langle \check{g}_{p,q}; b \rangle^* \right\} = - \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} (\mathbb{E} \{ \hat{c}_{m,n} \langle \check{g}_{p,q}; b \rangle^* \}) \langle \check{g}_{p,q}; g_{m,n} \rangle \approx 0. \quad (4.24)$$

Ainsi, en combinant (3.10) (4.21) (4.24) dans (4.16), on obtient l'expression de la variance de l'interférence et du bruit résiduel :

$$\sigma_{\nu_{p,q}}^2 = (\sigma_c^2 - \sigma_{\hat{c}}^2) \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} |\langle \check{g}_{p,q}; g_{m,n} \rangle|^2 + \sigma_b^2 \|\check{g}_{p,q}\|^2, \quad (4.25)$$

et le RSIB prend alors une forme très semblable à (3.17) :

$$\text{RSIB}_{p,q} = \frac{\sigma_c^2 |\langle \check{g}_{p,q}; g_{m,n} \rangle|^2}{(\sigma_c^2 - \sigma_{\hat{c}}^2) \left(\sum_{(m,n) \in \Lambda} |\langle \check{g}_{p,q}; g_{m,n} \rangle|^2 - |\langle \check{g}_{p,q}; g_{p,q} \rangle|^2 \right) + \sigma_b^2 \|\check{g}_{p,q}\|^2}, \quad (4.26)$$

et sa maximisation est également garantie par les conditions du théorème 3.3, c'est-à-dire en utilisant un prototype formant une frame étroite en émission, ainsi qu'un prototype de réception proportionnel à celui d'émission. Dans ce cas, le RSIB maximal prends, là encore, une forme très semblable à (3.35) :

$$\text{RSIB}_{p,q} = \text{RSIB}_{\max} = \frac{\sigma_c^2}{(\sigma_c^2 - \sigma_{\tilde{c}}^2)(\rho - 1) + \sigma_b^2/\|g\|^2}. \quad (4.27)$$

4.2.3 Conditions pour la minimisation de l'erreur quadratique moyenne

Donnons tout d'abord l'expression de l'erreur quadratique moyenne :

$$\mathbb{E} \{|c_{p,q} - \bar{c}_{p,q}|^2\} = \mathbb{E} \{|c_{p,q} - \alpha_{p,q}c_{p,q} - \nu_{p,q}|^2\} \quad (4.28)$$

$$= \mathbb{E} \{|c_{p,q}(1 - \alpha_{p,q}) - \nu_{p,q}|^2\} \quad (4.29)$$

$$= \sigma_c^2 |1 - \alpha_{p,q}|^2 + \sigma_{\nu_{p,q}}^2 - \mathbb{E} \{c_{p,q} \nu_{p,q}^*\}, \quad (4.30)$$

où le terme $\mathbb{E} \{c_{p,q} \nu_{p,q}^*\}$ est nul : rappelons en effet que $\mathbb{E} \{c_{p,q} \hat{c}_{m,n}^*\} = \mathbb{E} \{\hat{c}_{p,q} \hat{c}_{m,n}^*\}$ [LLL05], et que le bruit est indépendant des symboles émis. On obtient donc :

$$\mathbb{E} \{c_{p,q} \nu_{p,q}^*\} = \mathbb{E} \{c_{p,q} (i_{p,q}^* - \hat{i}_{p,q}^* + \langle \check{g}_{p,q}; b \rangle)\} \quad (4.31)$$

$$= \mathbb{E} \{c_{p,q} (i_{p,q}^* - \hat{i}_{p,q}^*)\} \quad (4.32)$$

$$= \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} (\mathbb{E} \{c_{p,q} c_{m,n}^*\} - \mathbb{E} \{c_{p,q} \hat{c}_{m,n}^*\}) \langle \check{g}_{p,q}; g_{m,n} \rangle \quad (4.33)$$

$$= 0. \quad (4.34)$$

À partir de ce résultat, et en rappelant que $\alpha_{p,q} = \langle \check{g}_{p,q}; g_{p,q} \rangle$, on peut faire apparaître le RSIB dans l'expression de l'erreur quadratique moyenne :

$$\mathbb{E} \{|c_{p,q} - \bar{c}_{p,q}|^2\} = \sigma_c^2 |\alpha_{p,q}|^2 \left(\left| \frac{1 - \alpha_{p,q}}{\alpha_{p,q}} \right|^2 + \text{RSIB}^{-1} \right). \quad (4.35)$$

Comme indiqué précédemment, RSIB^{-1} est minimisé en respectant les conditions énoncées dans 3.3. De plus, dans ce cas, on peut choisir arbitrairement $\mu_{p,q} \in \mathbb{C}^*$ tel que $\check{g}_{p,q}(t) = \mu_{p,q} g_{p,q}(t)$, ce qui signifie que l'on peut construire n'importe quelle valeur $\alpha_{p,q} = \langle \check{g}_{p,q}; g_{p,q} \rangle = \mu_{p,q}^* \|g\|^2$ dans $\mathbb{C}^*$ à partir de $\mu_{p,q} \forall (p,q) \in \Lambda$. Dans ces conditions, il ne reste qu'à trouver $\alpha_{p,q}$ tel que :

$$\alpha_{p,q} \in \underset{\alpha_{p,q} \in \mathbb{C}^*}{\text{argmin}} |\alpha_{p,q}|^2 \left(\left| \frac{1 - \alpha_{p,q}}{\alpha_{p,q}} \right|^2 + \text{RSIB}_{\max}^{-1} \right), \quad (4.36)$$

et en déduire la valeur de $\mu_{p,q}$ à imposer. Pour cela, il suffit de chercher une valeur de $\alpha_{p,q}$ annulant le gradient suivant :

$$\nabla_{\alpha_{p,q}} |\alpha_{p,q}|^2 \left(\left| \frac{1 - \alpha_{p,q}}{\alpha_{p,q}} \right|^2 + \text{RSIB}_{\max}^{-1} \right) = \nabla_{\alpha_{p,q}} |1 - \alpha_{p,q}|^2 + \text{RSIB}_{\max}^{-1} \nabla_{\alpha_{p,q}} |\alpha_{p,q}|^2 \quad (4.37)$$

$$= (1 + \text{RSIB}_{\max}^{-1}) \nabla_{\alpha_{p,q}} |\alpha_{p,q}|^2 - 2 \nabla_{\alpha_{p,q}} \mathcal{R} \{ \alpha_{p,q} \}, \quad (4.38)$$

avec

$$\nabla_{\alpha_{p,q}} |\alpha_{p,q}|^2 = \frac{\partial |\alpha_{p,q}|^2}{\partial \mathcal{R} \{\alpha_{p,q}\}} + j \frac{\partial |\alpha_{p,q}|^2}{\partial \mathcal{I} \{\alpha_{p,q}\}} = 2\alpha_{p,q}, \quad (4.39)$$

et

$$\nabla_{\alpha_{p,q}} \mathcal{R} \{\alpha_{p,q}\} = \frac{\partial \mathcal{R} \{\alpha_{p,q}\}}{\partial \mathcal{R} \{\alpha_{p,q}\}} = 1, \quad (4.40)$$

ce qui revient à résoudre l'équation ci-dessous :

$$2(1 + \text{RSIB}_{\max}^{-1}) \alpha_{p,q} - 2 = 0 \quad (4.41)$$

$$\alpha_{p,q} = \frac{1}{1 + \text{RSIB}_{\max}^{-1}}. \quad (4.42)$$

Pour minimiser l'EQM de ce système multiporteuse FTN avec annulation souple de l'interférence, il faut donc garantir $\check{g}_{p,q}(t) = \mu_{p,q} g_{p,q}(t)$, avec :

$$\mu_{p,q}^* = \frac{\alpha_{p,q}}{\|g\|^2} = \frac{1}{\|g\|^2 (1 + \text{RSIB}_{\max}^{-1})} = \mu, \quad (4.43)$$

et $\mathbf{g}$ une frame étroite. Ces conditions sont résumées dans le théorème 4.1.

Théorème 4.1 (Minimisation de l'EQM)

Soit un émetteur/récepteur multiporteuse FTN tel que défini par (1.56) et (1.57), adossé à un système de suppression souple d'interférence tel que définie par (4.7).

On considère la transmission d'une séquence de symboles IID $\mathbf{c} = \{c_{m,n}\}_{(m,n) \in \mathbb{Z}^2} \in \ell_2(\mathbb{Z}^2)$, de variance σ_c^2 , sur un canal à BABG dont on notera σ_b^2 la variance du bruit (supposé centré).

Notons σ_c^2 la variance des symboles estimés à partir de l'information a priori par (2.37), $\mathbf{g} = \{g_{m,n}\}_{(m,n) \in \mathbb{Z}^2}$ la famille de Gabor d'émission, de prototype $g(t) \in \mathcal{L}_2(\mathbb{R})$, $\check{\mathbf{g}} = \{\check{g}_{m,n}\}_{(m,n) \in \mathbb{Z}^2}$ la famille de réception, et $\rho > 1$ la densité du système.

Alors, pour maximiser l'EQM après suppression d'interférence, il est suffisant de réunir les conditions suivantes :

1. $\mathbf{g}$ est une frame étroite ;
2. $\check{g}_{m,n}(t) = \mu g_{m,n}(t) \quad \forall (m,n) \in \mathbb{Z}^2$;
3. $\mu = 1 / [\|g\|^2 (1 + \text{RSIB}_{\max}^{-1})]$;

où $\text{RSIB}_{\max}$ est donné par (4.27).

Quelques remarques sur ce résultat :

- $\Lambda = \mathbb{Z}^2$ est fondamentalement incompatible avec le fonctionnement par bloc de données d'un système turbo, nous verrons cependant par simulation que le choix du nombre de symboles multiporteuses K et le nombre de sous-porteuses M est finalement peu contraint ;
- lorsqu'une information *a priori* parfaite est disponible ($\sigma_c^2 = \sigma_c^2$), le RSIB maximal devient égal au RSB maximal que l'on obtiendrait par filtrage adapté dans un système STN, cela signifie que dans ce contexte, lorsque le turbo-égaliseur converge, il converge vers les performances du système STN codé sur canal à BABG.

4.2.4 Conversion SISO de sortie

Dérivons l'expression de $\alpha_{p,q}$ et $\sigma_{\nu_{p,q}}^2$ afin d'utiliser le convertisseur SISO classique détaillé dans la partie 2.4.2.2.

Commençons par $\alpha_{p,q}$. Lorsque les conditions du théorème 4.1 sont respectées, on a :

$$\alpha_{p,q} = \langle \check{g}_{p,q}; g_{p,q} \rangle = \mu \langle g_{p,q}; g_{p,q} \rangle, \quad (4.44)$$

avec $\mu = 1 / [\|g\|^2 (1 + \text{RSIB}_{\max}^{-1})]$. $\mathbf{g}$ étant une famille de Gabor, on a $\langle g_{p,q}; g_{p,q} \rangle = \|g\|^2$. Ainsi, le terme $\alpha_{p,q}$ ne dépend pas dans ce cas ni l'indice temporel q , ni de l'indice de sous-porteuse p :

$$\alpha_{p,q} = \alpha = \mu \|g\|^2 \quad (4.45)$$

Quant à $\sigma_{\nu_{p,q}}^2$, dont on rappelle l'expression générale donnée en (4.25) :

$$\sigma_{\nu_{p,q}}^2 = (\sigma_c^2 - \sigma_{\check{c}}^2) \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} |\langle \check{g}_{p,q}; g_{m,n} \rangle|^2 + \sigma_b^2 \|\check{g}_{p,q}\|^2, \quad (4.46)$$

lorsque les conditions du théorème 4.1 sont respectées on obtient, d'une part :

$$\|\check{g}_{p,q}\|^2 = \mu^2 \|g_{p,q}\|^2 = \mu^2 \|g\|^2, \quad (4.47)$$

et

$$\begin{aligned} \sum_{(m,n) \in \mathbb{Z}^2 \setminus \{(p,q)\}} |\langle \check{g}_{p,q}; g_{m,n} \rangle|^2 &= \mu^2 \left(\sum_{(m,n) \in \mathbb{Z}^2} |\langle g_{p,q}; g_{m,n} \rangle|^2 - |\langle g_{p,q}; g_{p,q} \rangle|^2 \right) \\ &= \mu^2 \left(\sum_{(m,n) \in \mathbb{Z}^2} |\langle g; g_{m,n} \rangle|^2 - \|g\|^4 \right), \end{aligned} \quad (4.48)$$

ainsi que, d'autre part,

$$\sum_{(m,n) \in \mathbb{Z}^2} |\langle g; g_{m,n} \rangle|^2 = \langle \mathbf{S}_{\mathbf{g}} g; g \rangle = \rho \|g\|^4. \quad (4.49)$$

Ainsi, en combinant (4.47), (4.48) et (4.49) dans (4.46), on obtient finalement une expression ne dépendant ni de l'indice temporel q , ni de l'indice de sous-porteuse p :

$$\sigma_{\nu_{p,q}}^2 = \sigma_{\nu}^2 = \mu^2 \|g\|^4 [(\sigma_c^2 - \sigma_{\check{c}}^2)(\rho - 1) + \sigma_b^2 / \|g\|^2], \quad (4.50)$$

où on identifie le dénominateur de $\text{RSIB}_{\max}$, de manière à simplifier l'expression de σ_{ν}^2 comme suit :

$$\sigma_{\nu}^2 = \mu^2 \|g\|^4 \sigma_c^2 \text{RSIB}_{\max}^{-1}. \quad (4.51)$$

Enfin, on peut exprimer $\text{RSIB}_{\max}^{-1}$ en fonction de μ :

$$\mu = \frac{1}{\|g\|^2 (1 + \text{RSIB}_{\max}^{-1})} \Leftrightarrow \text{RSIB}_{\max}^{-1} = (\mu \|g\|^2)^{-1} - 1, \quad (4.52)$$

ce qui, en réintroduisant $\alpha = \mu \|g\|^2$, donne finalement :

$$\sigma_\nu^2 = \sigma_c^2 \alpha (1 - \alpha). \quad (4.53)$$

En résumé, la conversion SISO s'effectue via l'association des équations (2.43) et (2.44) avec, $\forall(p, q) \in \mathbb{Z}^2$:

- $\alpha_{p,q} = \alpha = 1 / (1 + \text{RSIB}^{-1})$;
- $\sigma_{\nu_{p,q}}^2 = \sigma_\nu^2 = \sigma_c^2 \alpha (1 - \alpha)$.

Ainsi, en plus de permettre la minimisation de l'EQM, l'utilisation de frames étroites dans ce contexte de suppression souple d'interférence dans un système multiporteuses FTN permet également de ne calculer qu'une fois par itération, et pour tous les symboles, les paramètres du convertisseur SISO de sortie.

4.2.5 Simulations

Dans cette partie, nous vérifierons les résultats théoriques développés auparavant, et évaluons les performances du système à suppression d'interférence souple MEQM sur canal à BABG. Pour ce faire, nous mettons en œuvre une transmission multiporteuse FTN codée et entrelacée. On utilise pour ce faire le code convolutif (7,5) et un entrelaceur aléatoire agissant sur une séquence de 32768 bits codés. Les bits entrelacés seront ensuite convertis en symboles QPSK. En réception, le système d'annulation souple de l'interférence est implémenté de concert avec un décodeur BCJR dans une boucle turbo.

Notons M le nombre de sous-porteuses et K , le nombre de symboles par porteuse, la figure 4.2 évalue les différences de performances pour $M = 2$ et $K = 8192$, $M = 512$ et $K = 32$, $M = 8192$ et $K = 2$, de telle sorte que le produit $K.M$ est constant. On observe un comportement identique, même dans les cas limites où $M = 2$ et $K = 2$, permettant un choix complètement libre de K et M pour une valeur fixée de $K.M$.

Afin d'effectuer la conversion SISO de sortie, nous considérons que la distribution suivie par la somme de l'interférence et du bruit $\nu_{p,q}$ suit une distribution gaussienne. Or, nous l'avons vu dans le chapitre 3 (partie 3.3.4), cette approximation est relativement mauvaise pour $\rho \geq 2$. Dans les simulations présentées en figure 4.3, on compare les performances du système pour $\rho = 2$ lorsque le convertisseur SISO de sortie utilise une approximation gaussienne, et lorsque la densité de probabilité empirique de $\nu_{p,q}$ est utilisée. De manière intéressante, on n'observe pas de différence notable entre les deux cas. En particulier, pour $\rho = 8/3$, on remarque qu'aucun des deux systèmes ne permet la convergence du turbo-égaliseur.

Sur la figure 4.4, on remarque que tous les systèmes utilisant des frames étroites présentent des performances semblables en termes de TEB. Lorsque des frames duales canoniques sont utilisées et que le système turbo converge, on remarque un décalage avec les performances du système orthogonal. Ce décalage correspond au fait que ces prototypes ne permettent pas de maximiser le RSB (voir chapitre 3). À l'inverse, lorsque le système turbo converge avec des prototypes d'émission et de réception proportionnels, on devrait effectivement retrouver les performances du système orthogonal. Cependant, dans nos simulations, cette configuration ne permet pas au système turbo de converger.

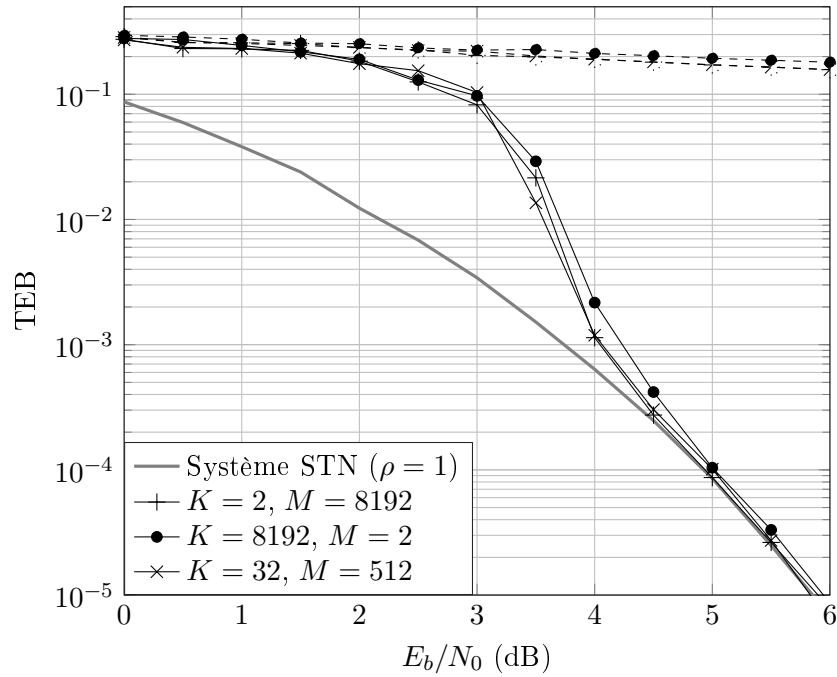


FIGURE 4.2 – Performances du système FTN turbo à suppression souple de l'interférence pour différentes valeurs de K et M , avec $K.M = 16384$, $\rho = 2$ et le code convolutif $(7, 5)$. Les courbes en tirets représentent la première itération du turbo-égaliseur, celles en trait plein la dixième.

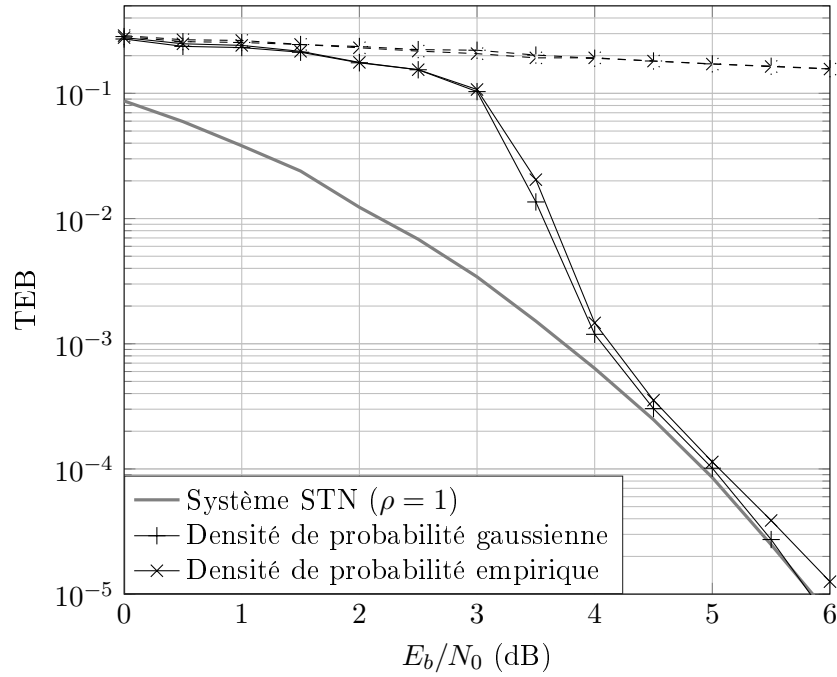


FIGURE 4.3 – Performances comparées du système FTN turbo à suppression souple de l'interférence lorsque la conversion SISO de sortie s'effectue en prenant en compte la densité de probabilité empirique de l'interférence (voir figure 3.2), et lorsque cette dernière est approchée par une gaussienne. Ces simulations sont effectuées avec $K = 32$ symboles multiporteuses, $M = 512$ porteuses pour une densité $\rho = 2$, et en utilisant le code convolutif (7, 5). Les courbes en tirets représentent la première itération du turbo-égaliseur, celles en trait plein la dixième.

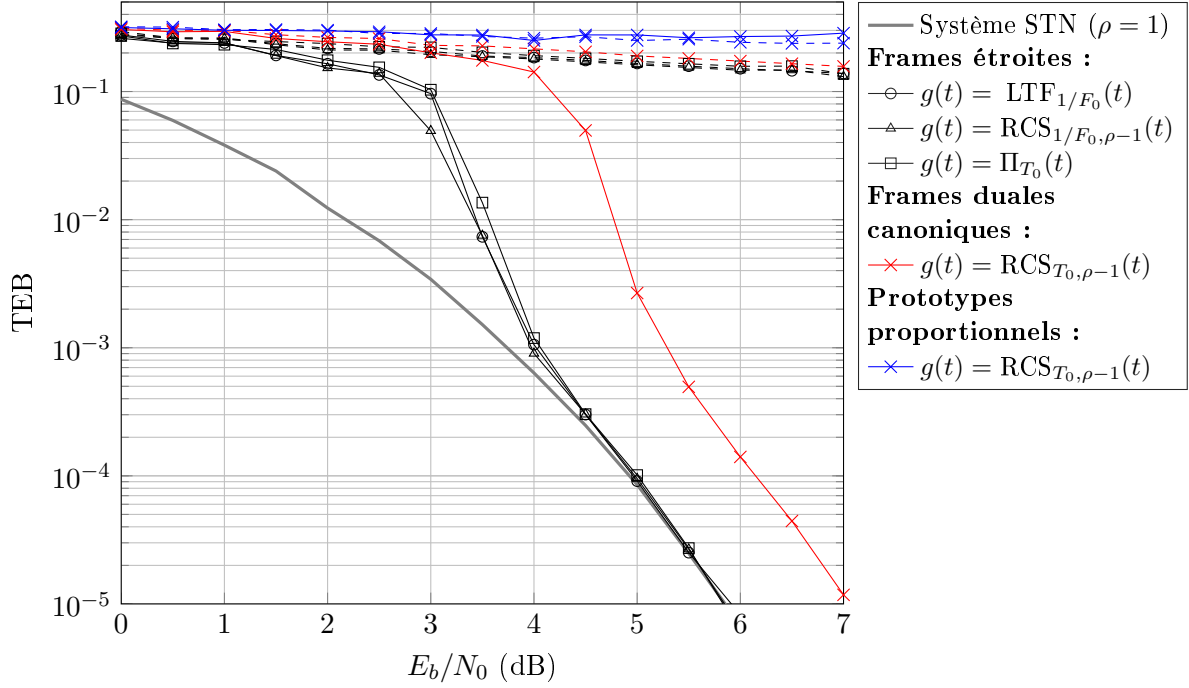


FIGURE 4.4 – Comparaison des performances du système FTN turbo à suppression souple de l'interférence pour différents prototypes. Ces simulations sont effectuées avec $K = 32$ symboles multiporteuses, $M = 512$ porteuses pour une densité $\rho = 2$, et en utilisant le code convolutif $(7, 5)$. Les courbes en tirets représentent la première itération du turbo-égaliseur, celles en trait plein la dixième.

Enfin, on observe en figure 4.5 que ce système présente de très bonnes performances en QPSK, puisqu'il permet ici d'augmenter l'efficacité spectrale de 14% par rapport à une QPSK non-codée, tout en diminuant le rapport E_b/N_0 nécessaire pour atteindre le TEB cible de 3,6 dB. Les performances sont plus mitigées pour des constellations d'ordre supérieures. Remarquons cependant que lorsque la structure turbo converge, le gain en efficacité spectrale se fait toujours à E_b/N_0 constant par rapport au système codé orthogonal. À l'inverse, les techniques de poinçonnage de code par exemple, permettent d'augmenter l'efficacité spectrale du système codé mais au prix d'une augmentation du rapport E_b/N_0 nécessaire pour obtenir le TEB cible. Évidemment, de meilleures performances sont possibles en utilisant de meilleurs codes. Cependant, pour conserver une complexité acceptable, il serait sans doute préférable de ne pas utiliser de codes qui nécessiteraient d'imbriquer deux systèmes itératifs (un pour l'égalisation et l'autre pour le décodage) dans la mesure où cela induirait un nombre total d'itération élevé pour le décodeur. Dans ce contexte, les turbo-codes et les codes LDPC ne sont donc pas de bons candidats. Notons qu'en monoporteuse, des travaux ont été menés par Anderson *et al.* pour trouver des codes convolutifs adaptés aux transmissions FTM [AZ14].

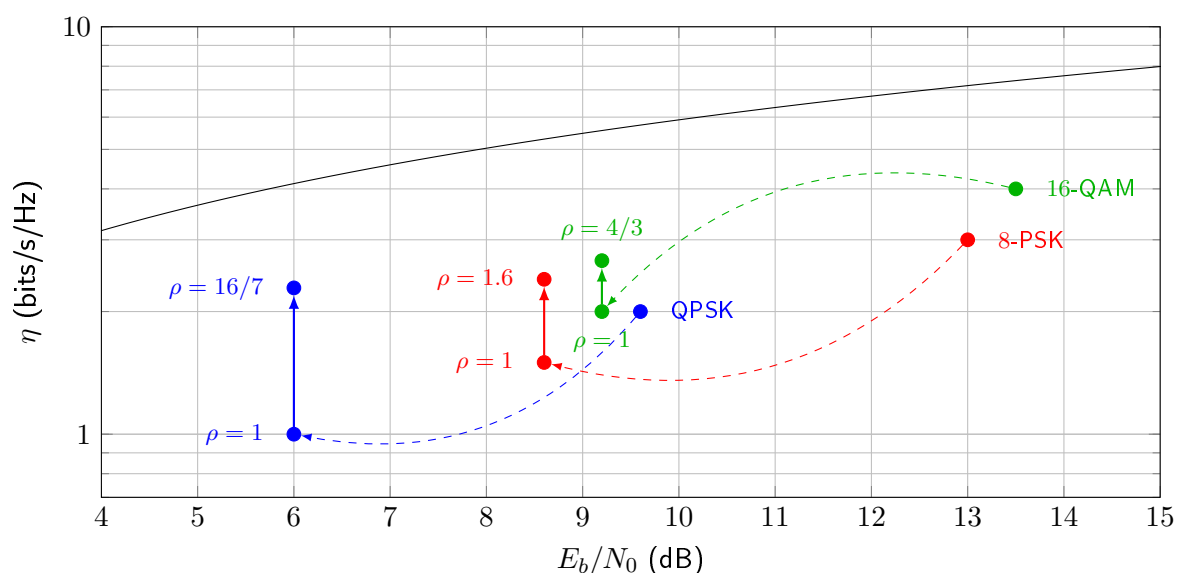


FIGURE 4.5 – Impact sur l'efficacité spectrale du système multiporteuse FTM turbo à annulation souple de l'interférence sur canal à BABG pour un TEB cible de 10^{-5} . Pour chaque constellation, trois points sont donnés. Le premier (identifié par le nom de la constellation) représente la performance du système orthogonal non-codé, le deuxième (identifié par $\rho = 1$) représente la performance du système orthogonal codé utilisant un code convolutif (7, 5), et le dernier (identifié par la valeur maximale de densité permettant la convergence du système turbo) représente la performance du système FTM turbo à annulation souple de l'interférence utilisant le code (7, 5).

4.3 Suppression successive d'interférence en présence de canal radiomobile

Dans cette partie, étendons le système à suppression souple d'interférence présenté précédemment dans le cadre d'un canal à BABG au cas des canaux LTI et LTV. En particulier, nous verrons que les propriétés d'adaptation au canal de transmission propres aux modulations multiporteuses (telle que l'approximation des canaux LTI par un coefficient par sous-porteuse) sont conservées, et que les règles de choix des prototypes, en termes de localisation temps-fréquence notamment, restent les mêmes qu'en STN.

4.3.1 Calcul de l'interférence

Rappelons l'expression des symboles en sortie du démodulateur multiporteuses dans le cas d'une transmission sur canal LTV, noté $\mathcal{H}$:

$$\tilde{c}_{p,q} = c_{p,q} \langle \check{g}_{p,q}; \mathcal{H}g_{p,q} \rangle + \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} c_{m,n} \langle \check{g}_{p,q}; \mathcal{H}g_{m,n} \rangle + \langle \check{g}_{p,q}; b \rangle. \quad (4.54)$$

Dans le cas où les conditions données en partie 1.3.3 sont respectées (c'est-à-dire que sa fonction d'étalement est négligeable passé un certain retard et une certain décalage Doppler, et que le prototype est bien localisé dans le plan temps-fréquence), on peut simplifier cette expression en approchant le canal avec un coefficient par symbole :

$$\tilde{c}_{p,q} \approx c_{p,q} L_h(pF_0, qT_0) \langle \check{g}_{p,q}; g_{p,q} \rangle + \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} c_{m,n} L_h(mF_0, nT_0) \langle \check{g}_{p,q}; g_{m,n} \rangle + \langle \check{g}_{p,q}; b \rangle. \quad (4.55)$$

Pour simplifier les notations, nous écrivons $L_{m,n} = L_h(mF_0, nT_0) \forall (m,n) \in \Lambda$, ce qui donne :

$$\tilde{c}_{p,q} \approx \underbrace{c_{p,q} L_{p,q} \langle \check{g}_{p,q}; g_{p,q} \rangle}_{\text{Signal utile}} + \underbrace{\sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} c_{m,n} L_{m,n} \langle \check{g}_{p,q}; g_{m,n} \rangle}_{\text{Interférence : } i_{p,q}} + \underbrace{\langle \check{g}_{p,q}; b \rangle}_{\text{Bruit filtré}}, \quad (4.56)$$

Ainsi, en supposant une connaissance parfaite du canal par le récepteur (ce qui nécessite en pratique une étape préliminaire d'estimation du canal, qui ne sera pas traitée dans le cadre de nos travaux), on estime le terme d'interférence à partir de la séquence de symboles estimés $\hat{\mathbf{c}}$ via :

$$\hat{i}_{p,q} = \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} \hat{c}_{m,n} L_{m,n} \langle \check{g}_{p,q}; g_{m,n} \rangle. \quad (4.57)$$

Les symboles en sortie de l'anneur d'interférence sont donc donnés par :

$$\bar{c}_{p,q} \approx \underbrace{c_{p,q} L_{p,q} \langle \check{g}_{p,q}; g_{p,q} \rangle}_{\text{Signal utile}} + \underbrace{\sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} (c_{m,n} - \hat{c}_{m,n}) L_{m,n} \langle \check{g}_{p,q}; g_{m,n} \rangle}_{\text{Interférence résiduelle : } i_{p,q} - \hat{i}_{p,q}} + \underbrace{\langle \check{g}_{p,q}; b \rangle}_{\text{Bruit filtré}}. \quad (4.58)$$

4.3. Suppression successive d'interférence en présence de canal radiomobile

Tout comme précédemment, cette expression se calcule de manière efficace en fournissant la séquence de symboles $\hat{\mathbf{c}}' = \{\hat{c}'_{m,n}\}_{(m,n) \in \Lambda}$ avec $\hat{c}'_{m,n} = L_{m,n} \hat{c}_{m,n} \forall (m,n) \in \Lambda$ à un modulateur multiporteuse suivi d'un démodulateur multiporteuse et en retranchant les termes $\hat{c}'_{m,n} \langle \check{g}_{m,n}; g_{m,n} \rangle$:

$$\hat{\mathbf{i}} = \text{DEMOD} \{ \text{MOD} \{ \hat{\mathbf{c}}' \} \} - \{ \hat{c}'_{m,n} \langle \check{g}_{m,n}; g_{m,n} \rangle \}_{(m,n) \in \Lambda}. \quad (4.59)$$

En mutualisant l'opération de démodulation dans le calcul de l'interférence, et celle du signal reçu $r(t)$, on obtient :

$$\bar{\mathbf{c}} = \text{DEMOD} \{ r(t) - \text{MOD} \{ \hat{\mathbf{c}}' \}(t) \} - \{ \hat{c}'_{m,n} \langle \check{g}_{m,n}; g_{m,n} \rangle \}_{(m,n) \in \Lambda}. \quad (4.60)$$

4.3.2 Conversion SISO de sortie

Dérivons l'expression de $\alpha_{p,q}$ et $\sigma_{\nu_{p,q}}^2$ afin d'utiliser le convertisseur SISO classique détaillé dans la partie 2.4.2.2.

Commençons par $\alpha_{p,q}$. Sa valeur se déduit de (4.44) en remplaçant $g_{p,q}(t)$ par $L_{p,q} g_{p,q}(t) \forall (p,q) \in \Lambda$:

$$\alpha_{p,q} = L_{p,q} \langle \check{g}_{p,q}; g_{p,q} \rangle = \mu_{p,q} L_{p,q} \|g\|^2, \quad (4.61)$$

on remarque que cette expression dépend de la valeur de la fonction de transfert évolutive évaluée en (p,q) , sauf si l'on choisit $\mu_{p,q} = 1/L_{p,q}$. Cependant, dans ce cas, il est nécessaire de s'assurer que $L_{p,q}$ est non-nul, quel que soit (p,q) .

L'expression de $\sigma_{\nu_{p,q}}^2$ se déduit de (4.46) de la même manière :

$$\sigma_{\nu_{p,q}}^2 = (\sigma_c^2 - \sigma_{\check{c}}^2) \sum_{(m,n) \in \Lambda \setminus \{(p,q)\}} |L_{m,n}|^2 |\langle \check{g}_{p,q}; g_{m,n} \rangle|^2 + \sigma_b^2 \|\check{g}_{p,q}\|^2 \quad (4.62)$$

$$= |\mu_{p,q}|^2 (\sigma_c^2 - \sigma_{\check{c}}^2) \left(\sum_{(m,n) \in \Lambda} |L_{m,n}|^2 |\langle g_{p,q}; g_{m,n} \rangle|^2 - |L_{p,q}|^2 |\langle g_{p,q}; g_{p,q} \rangle|^2 \right) + |\mu_{p,q}|^2 \sigma_b^2 \|g\|^2 \quad (4.63)$$

$$= |\mu_{p,q}|^2 (\sigma_c^2 - \sigma_{\check{c}}^2) \left(\sum_{(m,n) \in \Lambda} |L_{m,n}|^2 |\langle g; g_{m-p,n-q} \rangle|^2 - |L_{p,q}|^2 \|g\|^4 \right) + |\mu_{p,q}|^2 \sigma_b^2 \|g\|^2. \quad (4.64)$$

Ici, il n'est pas possible de simplifier davantage en utilisant, par exemple, les propriétés des frames. Nous proposons néanmoins plusieurs pistes pour effectuer efficacement ce calcul dans un contexte de turbo-égalisation. Premièrement, on remarque qu'il n'est nécessaire de calculer $\mu_{p,q} \sigma_b^2 \|g\|^2$ et $\mu_{p,q} \left(\sum_{(m,n) \in \Lambda} |L_{m,n}|^2 |\langle g; g_{m-p,n-q} \rangle|^2 - |L_{p,q}|^2 \|g\|^4 \right)$ qu'une fois pour toutes les itérations. Deuxièmement, $\sum_{(m,n) \in \Lambda} |L_{m,n}|^2 |\langle g; g_{m-p,n-q} \rangle|^2$ correspond à la convolution bidimensionnelle de $f(m,n) = |\langle g; g_{-m,-n} \rangle|^2$ et $h(m,n) = |L_{m,n}|^2 \forall (m,n) \in \Lambda$. Des implémentations logicielles et matérielles efficaces d'une telle opération, faisant notamment intervenir la

transformée de Fourier rapide à deux dimensions, sont à trouver dans le domaine du traitement d'images [Wos+98]; [MMN92]. Enfin, les termes de la famille $\{|\langle g; g_{m,n} \rangle|^2\}_{(m,n) \in \Lambda}$ devraient être pré-calculés pour alléger encore la charge de calcul.

4.3.3 Simulations

4.3.3.1 Canal sélectif en fréquence (LTI)

Observons le comportement du système multiporteuse FTN à suppression souple de l'interférence sur canal sélectif en fréquence. Pour cela, on évalue les performances de ce système sur canal Proakis B [PS08], dont les réponses impulsionnelle et fréquentielle sont données en figure 4.6. Il s'agit d'un canal relativement agressif, comme en témoignent les fréquences pour lesquelles le gain du canal est proche voire égal à zéro.

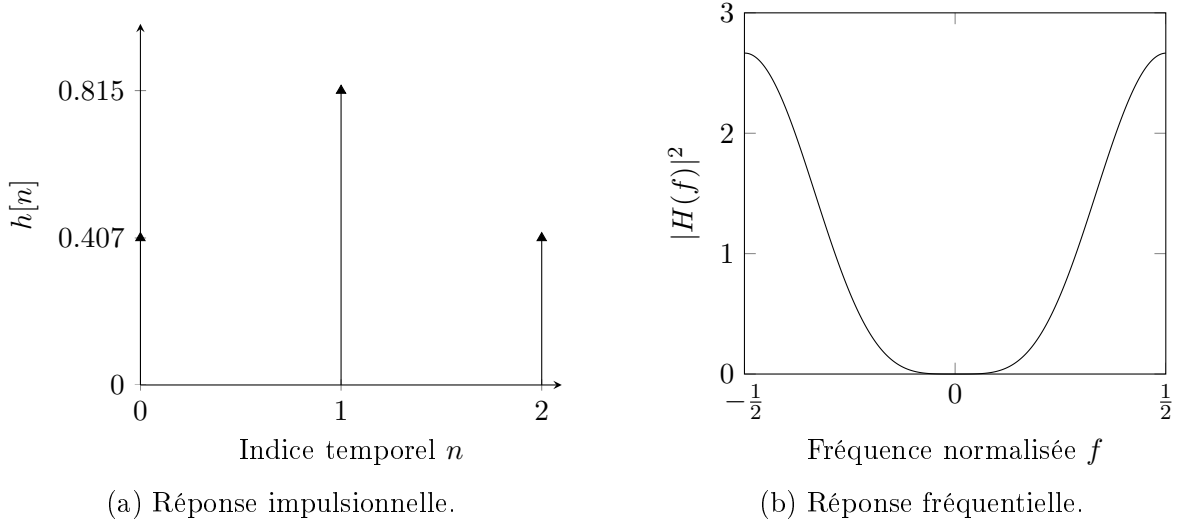


FIGURE 4.6 – Canal Proakis B.

Similairement aux simulations présentées précédemment, on simule ici la transmission de $K = 16384$ bits, codés par le code convolutif (7, 5), puis convertis en symboles QPSK. La densité du système FTN est ici fixée à $\rho = 1,6$ ce qui correspond, étant donné le code canal et la constellation utilisée, à une efficacité spectrale de $\eta = 1,6$ bits/s/Hz. On comparera ces performances en TEB à celle du système multiporteuse à la cadence de Nyquist ($\rho = 1$), et dont l'efficacité spectrale est de $\eta = 1$ bits/s/Hz.

On compare tout d'abord l'impact du nombre de sous-porteuses sur les performances du système sur la figure 4.7. On rappelle que pour un système multiporteuse STN, le fait d'augmenter le nombre de porteuses permet une meilleure adaptation aux canaux sélectifs en fréquence. On retrouve cette tendance avec le système FTN à suppression souple de l'interférence. Il est intéressant de noter qu'un nombre de sous-porteuses plus faible affecte plus durement les perfor-

4.3. Suppression successive d'interférence en présence de canal radiomobile

mances du système FTN que celles du système STN. Cela s'explique par le fait que le calcul de l'auto-interférence est forcément imparfait si l'approximation à 1 coefficient par sous-porteuse est mauvaise (même en présence d'information *a priori* parfaite). Dans ces conditions, où le nombre de sous-porteuses est insuffisant, le système FTN ne pourra pas converger vers les performances du système STN.

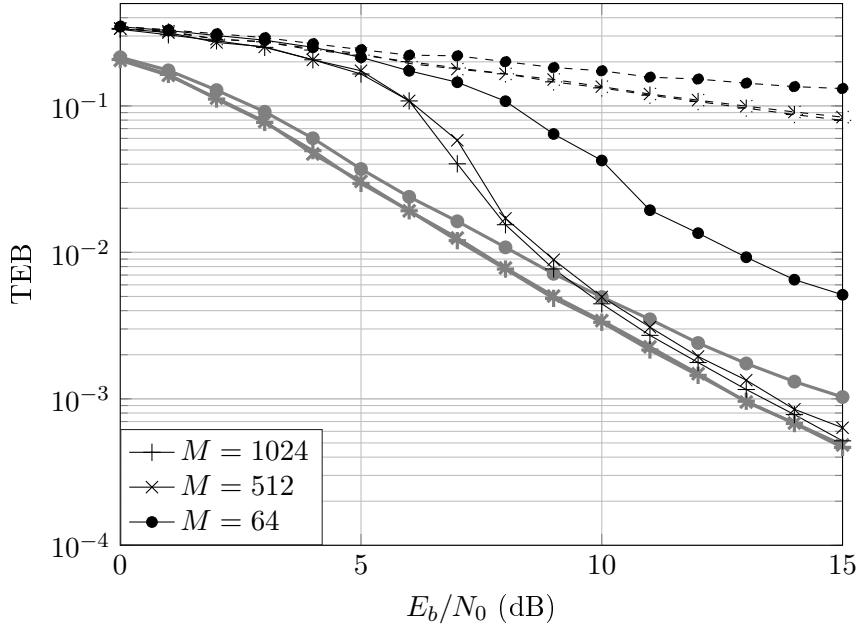


FIGURE 4.7 – Impact du nombre de sous-porteuses sur les performances du système FTN turbo à suppression souple de l'interférence pour $\rho = 1.6$ et $g(t) = \Pi_{T_0}(t)$. On remarque qu'un nombre insuffisant de porteuses impacte plus durement le système FTN (traits noirs) que le système STN (traits gris). Les courbes en tirets représentent la première itération du turbo-égaliseur, celles en trait plein la dixième.

L'impact des prototypes d'émission et de réception est évalué en figure 4.8. On y observe les mêmes tendances que dans le chapitre 3 concernant les prototypes menant à des paires de frames duales canoniques et les prototypes proportionnels. C'est-à-dire que ces derniers présentent un point de convergence pour un E_b/N_0 plus élevé que les autres prototypes, mais ils convergent vers les performances du système orthogonal. À l'inverse, les prototypes formant des paires de frames duales canoniques ont leur point de convergence à plus faible rapport E_b/N_0 , mais présentent un décalage avec la courbe du système STN en zone de convergence (en raison d'un RSB qui n'est pas maximal). Enfin, concernant les frames étroites, si on observe qu'elles combinent bien point de convergence à faible rapport E_b/N_0 et convergence vers les performances du système STN, on remarque en revanche qu'elles n'ont pas exactement le même comportement. En effet, entre le point de convergence et le début de la zone de convergence, on observe que les prototypes bien localisés en fréquence présentent de légèrement meilleures performances.

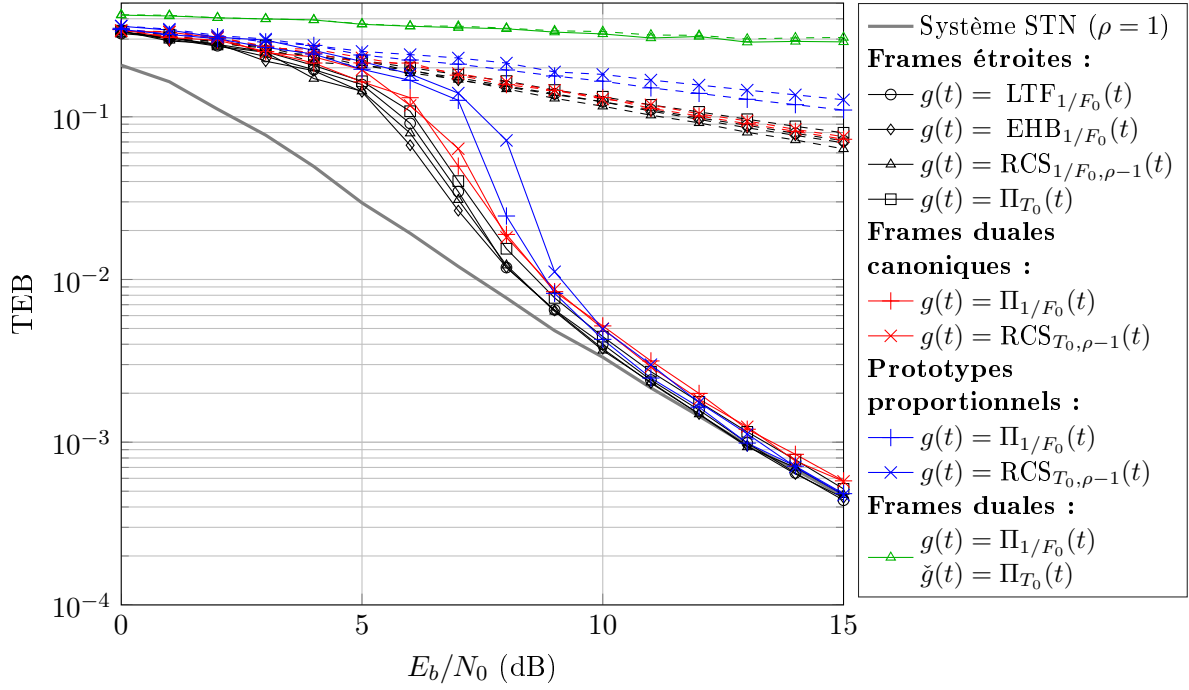


FIGURE 4.8 – Impact des prototypes d’émission et de réception sur les performances du système FTN turbo à suppression souple de l’interférence pour $\rho = 1.6$, $M = 1024$ sous-porteuses, et en présence d’un canal Proakis B. Les prototypes formant des frames étroites présentent les meilleures performances. On remarque que les performances de ces prototypes, lorsqu’ils présentent une bonne localisation fréquentielle, sont légèrement meilleures. Les courbes en tirets représentent la première itération du turbo-égaliseur, celles en trait plein la dixième.

4.3. Suppression successive d'interférence en présence de canal radiomobile

4.3.3.2 Canal radiomobile

Nous évaluons pour finir les performances du système multiporteuse FTN à suppression souple de l'interférence sur canal radiomobile. Pour ce faire, on utilise le modèle COST 207 TUx6, qui représente un canal urbain « typique » (TU signifiant Typical Urban) [Fai88]. Nous étudions un scénario de faible mobilité (type piéton), où la vitesse maximale de déplacement du récepteur est de $v_{\max} = 3$ km/h, et un scénario de forte mobilité avec $v_{\max} = 350$ km/h, sachant que le système occupe une bande $B = 8$ MHz, centrée autour d'une fréquence porteuse $f_0 = 600$ MHz.

Similairement aux simulations présentées précédemment, on simule la transmission de $K = 16384$ bits, codés par le code convolutif $(7, 5)$, puis convertis en symboles QPSK. La densité du système FTN est fixée à $\rho = 1,6$ ce qui correspond ici à une efficacité spectrale de $\eta = 1,6$ bits/s/Hz. À la différence des simulations présentées jusque là, il n'est pas possible ici de comparer les performances du système à une unique référence (où l'interférence aurait été parfaitement annulée). En effet, ici, la pertinence de l'approximation à un coefficient par symbole dépend du profil temps-fréquence du prototype utilisé, ainsi que de l'écart entre porteuses F_0 et l'écart entre symboles multiporteuses T_0 . Ainsi, les différents prototypes ne tendent pas vers les mêmes performances lorsque l'interférence est complètement éliminée, et plus encore, ne tendent pas vers celles des systèmes STN utilisant les mêmes prototypes (puisque les paramètres F_0 et T_0 ne sont pas les mêmes). C'est la raison pour laquelle, dans les simulations présentées ici, on se contentera de comparer les systèmes FTN entre-eux, sans analyser leurs comportements vis-à-vis de leurs équivalents idéaux (annulant parfaitement l'interférence).

Les figures 4.9 et 4.10 présentent les performances dans le scénario de faible mobilité, avec $M = 1024$ et $M = 4096$ sous-porteuses, respectivement. On remarque que tous les prototypes permettent la convergence du système turbo et, en raison d'une faible sélectivité temporelle, l'augmentation du nombre de sous-porteuses a ici un effet modéré sur les performances. Les meilleures performances en termes de TEB sont obtenues en utilisant des frames étroites, mais on remarque que toutes les frames étroites ne donnent pas les meilleures performances. Ainsi, le prototype rectangulaire $\Pi_{T_0}(t)$ présente ici de moins bonnes performances que le prototype $\Pi_{1/F_0}(t)$, lequel ne permet pas de former de frame étroite mais présente une meilleure sélectivité fréquentielle. Suivant la même tendance, les prototypes en racine de cosinus surélevé $\text{RCS}_{1/F_0, \rho-1}(t)$ et minimisant l'énergie hors bande $\text{EHB}_{1/F_0}(t)$, en raison de leur capacité à former des frames étroites et de leur bonne localisation fréquentielle, présentent les meilleures performances.

Les figures 4.11 et 4.12 présentent les performances dans le scénario de forte mobilité, avec $M = 1024$ et $M = 4096$ sous-porteuses, respectivement. On observe tout d'abord que tous les prototypes en racine de cosinus surélevé présentent des performances médiocres, en raison de leur très mauvaise localisation temporelle. Les prototypes rectangulaires, ceux minimisant l'énergie hors bande (EHB), et ceux maximisant la localisation temps-fréquence (LTF) permettent quant à eux la convergence du système turbo. De plus, on remarque ici encore que les meilleures performances sont détenues par les prototypes permettant de former de frames étroites. Enfin, l'impact négatif de l'augmentation du nombre de sous-porteuses (augmenter le nombre de sous-

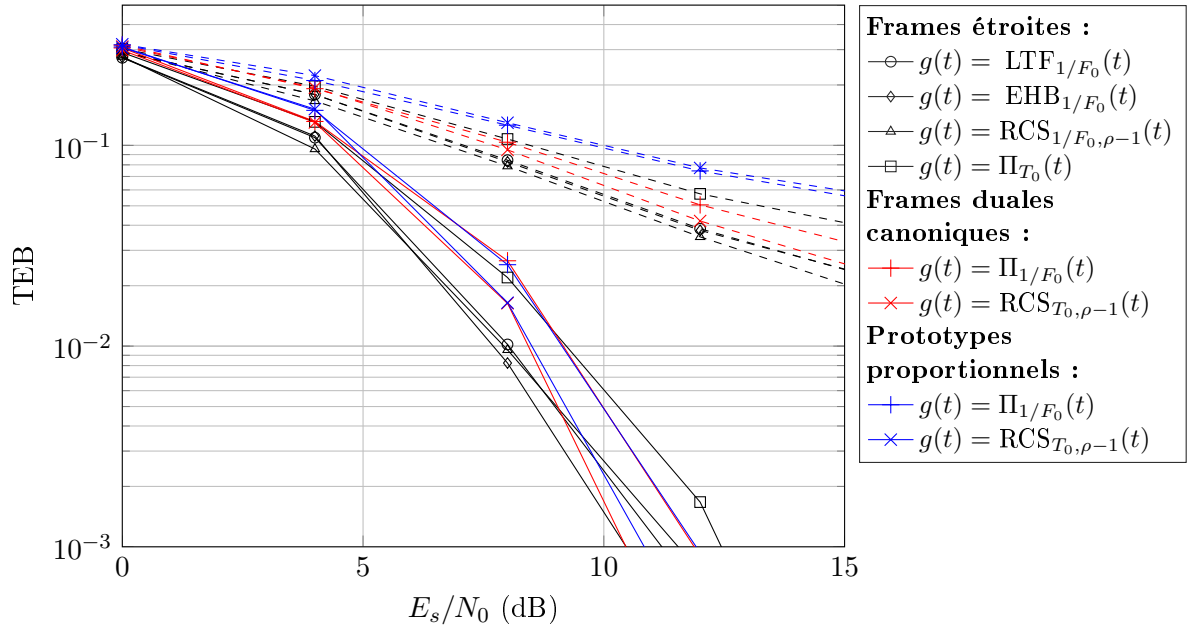


FIGURE 4.9 – Impact des prototypes d’émission et de réception sur les performances du système FTN turbo à suppression souple de l’interférence pour $\rho = 1.6$, $M = 1024$ sous-porteuses, et en présence d’un canal COST 207 TUX6 en situation de faible mobilité ($v_{\max} = 3$ km/h). Les courbes en tirets représentent la première itération du turbo-égaliseur, celles en trait plein la dixième.

4.3. Suppression successive d'interférence en présence de canal radiomobile

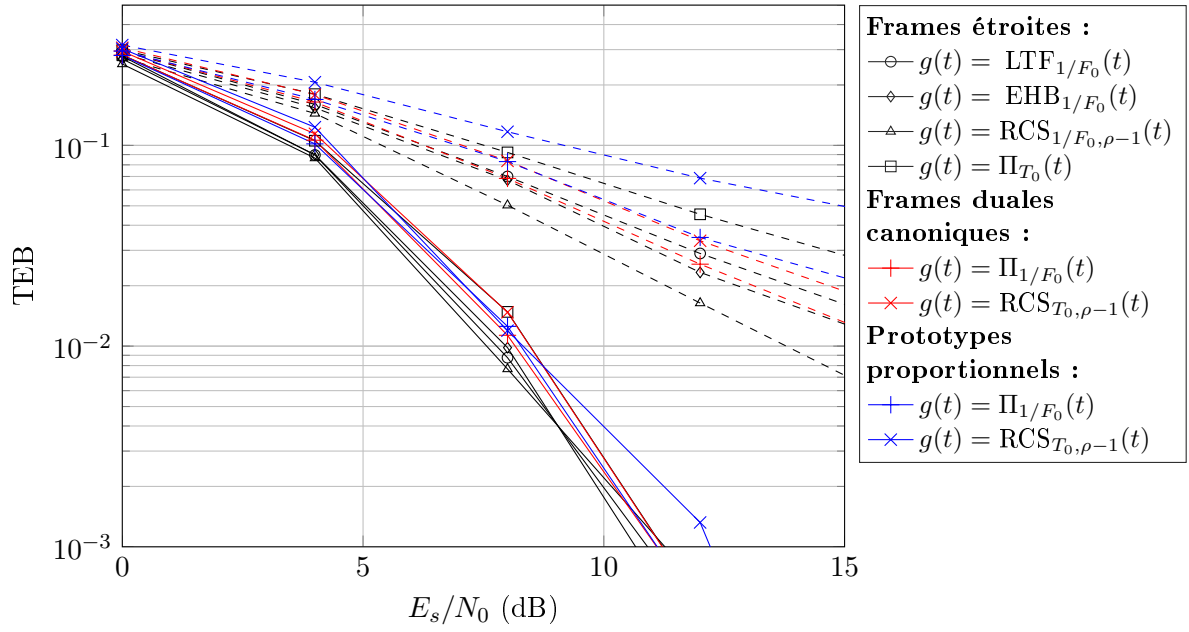


FIGURE 4.10 – Impact des prototypes d’émission et de réception sur les performances du système FTN turbo à suppression souple de l’interférence pour $\rho = 1.6$, $M = 4096$ sous-porteuses, et en présence d’un canal COST 207 TUX6 en situation de faible mobilité ($v_{\max} = 3$ km/h). Les courbes en tirets représentent la première itération du turbo-égaliseur, celles en trait plein la dixième.

porteuses ayant pour effet d'allonger la durée des symboles multiporteuses, entraînant une plus mauvaise localisation temporelle) est ici bien visible à la comparaison des deux figures.

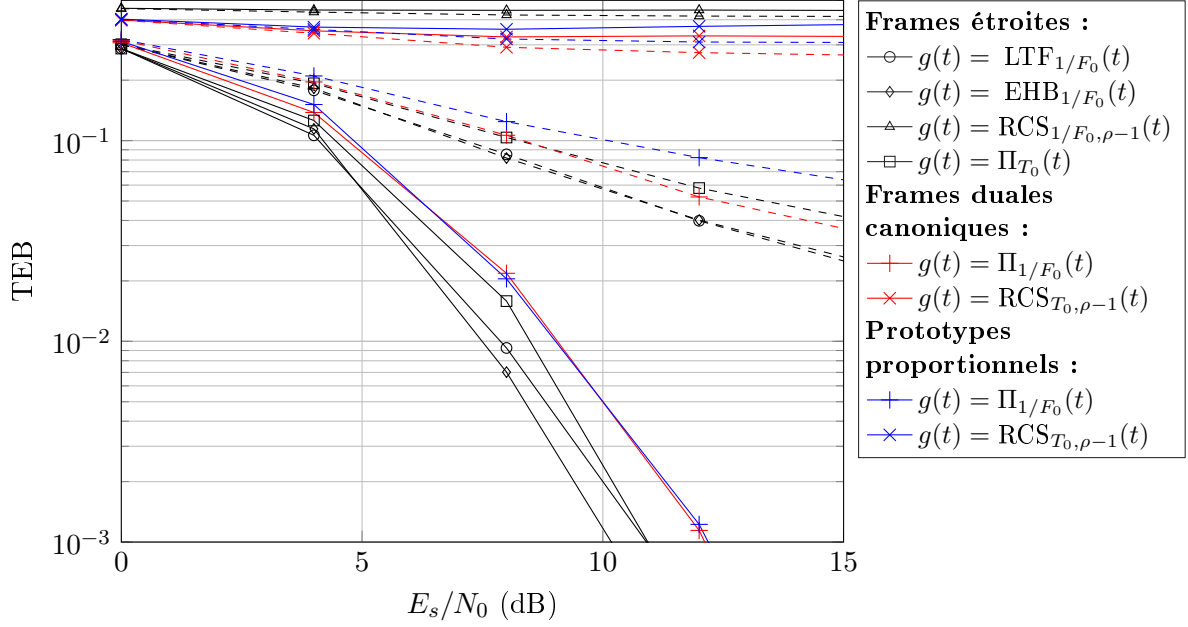


FIGURE 4.11 – Impact des prototypes d'émission et de réception sur les performances du système FTN turbo à suppression souple de l'interférence pour $\rho = 1.6$, $M = 1024$ sous-porteuses, et en présence d'un canal COST 207 TUx6 en situation de forte mobilité ($v_{\max} = 350$ km/h). Les courbes en tirets représentent la première itération du turbo-égaliseur, celles en trait plein la dixième.

4.4 Conclusion

Dans ce chapitre, nous avons montré que choisir des frames étroites en émission et en réception permet au système multiporteuse FTN avec suppression souple de l'interférence de minimiser l'erreur quadratique moyenne sur canal à bruit additif blanc gaussien. Les simulations associées montrent que ce système, une fois associé à un décodeur canal dans une boucle turbo, permet d'augmenter significativement l'efficacité spectrale sans pénalité en terme de taux d'erreur binaire, et pour une complexité algorithmique modeste.

Sur canal sélectif en fréquence et doublement sélectif (en temps et en fréquence), on montre par simulation que le système à suppression souple de l'interférence présente là aussi de meilleures performances lorsque des frames étroites sont utilisées en émission et en réception (bien que cela ne permette pas de minimiser l'erreur quadratique moyenne). On remarque surtout que les stratégies d'adaptation et de compensation du canal restent inchangées par rapport aux systèmes STN. En effet, sur canal sélectif en fréquence, l'approximation du canal à un coefficient par

4.4. Conclusion

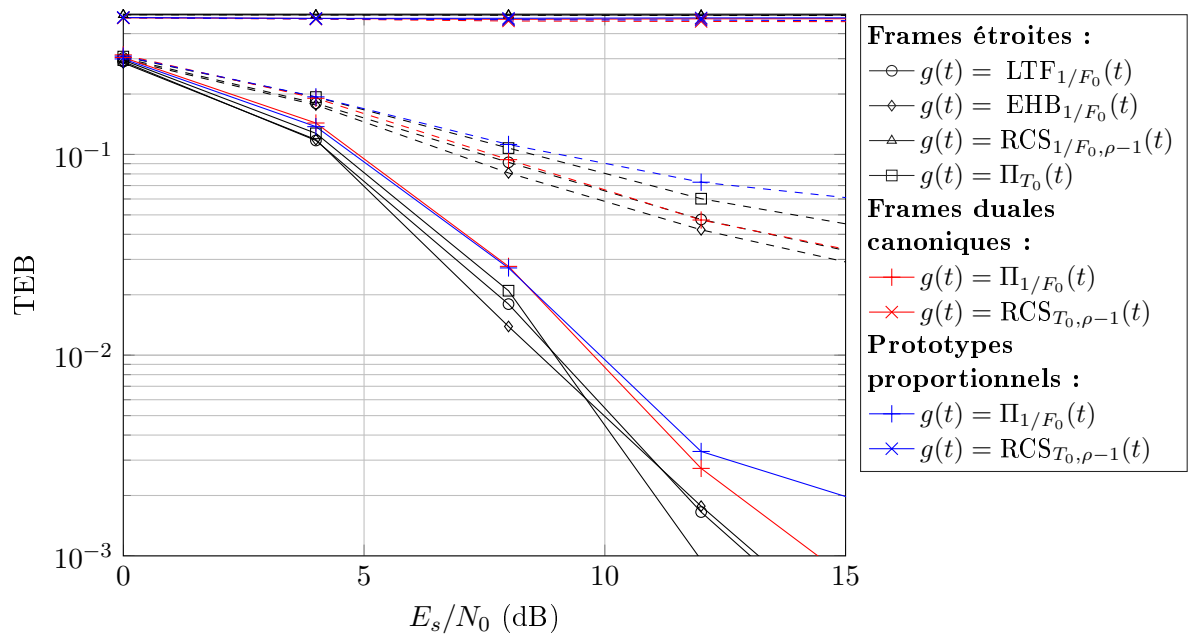


FIGURE 4.12 – Impact des prototypes d’émission et de réception sur les performances du système FTN turbo à suppression souple de l’interférence pour $\rho = 1.6$, $M = 4096$ sous-porteuses, et en présence d’un canal COST 207 TUx6 en situation de forte mobilité ($v_{\max} = 350$ km/h). Les courbes en tirets représentent la première itération du turbo-égaliseur, celles en trait plein la dixième.

porteuse est pertinente à condition d'utiliser suffisamment de sous-porteuses (à choisir en fonction du canal considéré), et les prototypes bien localisés en fréquence présentent de meilleures performances à nombre de porteuses constant. De même, sur canal radiomobile, il est nécessaire d'adapter le choix du nombre de sous-porteuses au profil temps-fréquence du canal afin que l'approximation du canal à un coefficient par symbole soit pertinente, et les prototypes présentant la meilleure localisation temps-fréquence présentent les meilleures performances. Dans ce contexte, nous préconisons l'utilisation des formes d'ondes EHB¹ et LTF² qui, en plus d'avoir une réponse impulsionnelle courte, ont de bonnes propriétés en termes de localisation temps-fréquence, et présentent dans nos simulations les meilleures performances comparées à celles des prototypes rectangulaires et en racine de cosinus surélevé.

Dans ce chapitre et celui qui le précède, nous avons étudié et proposé des techniques permettant d'augmenter l'efficacité spectrale des systèmes multiporteuses via la méthode FTN. L'utilisation de cette méthode a un impact sur le signal en sortie du modulateur et, compte tenu du mauvais comportement des systèmes multiporteuses classiques en termes de facteur de crête, nous étudions l'effet de la signalisation FTN sur ce dernier dans le chapitre 5.

1. Prototype conçu pour minimiser l'Énergie Hors Bande

2. Prototype conçu pour maximiser la Localisation Temps-Fréquence

Impact de la signalisation au-delà de la cadence de Nyquist sur le PAPR des modulations multiporteuses

Sommaire

5.1	Introduction	113
5.2	Définition du PAPR	114
5.3	Conditions d'optimalité	114
5.4	Prototypes d'émission optimaux	115
5.5	Simulations	116
5.6	Conclusion	118

5.1 Introduction

Comme nous l'avons évoqué dans l'introduction du chapitre 1, le PAPR est un élément fondamental du système de transmission puisqu'il influe sur le choix des amplificateurs de puissance, et sur l'efficacité énergétique du système de transmission. Comparé par exemple aux modulations à amplitude constante, dont le PAPR est minimal, ou même aux autres modulations linéaires classiques, les modulations multiporteuses, à l'instar d'une somme de sinusoïdes de fréquences différentes, présentent par nature un PAPR élevé. Dans la mesure où la puissance maximale est généralement limitée par des contraintes réglementaires et/ou d'intégration, la puissance moyenne des signaux multiporteuses doit être ajustée en conséquence, ce qui pénalise leur bilan de liaison. Ce fort PAPR est encore aujourd'hui un frein à l'adoption des modulations multiporteuses dans des environnements énergétiquement contraints, tels que les systèmes embarqués ou satellitaires.

Dans ce chapitre, on montre qu'un choix judicieux du prototype d'émission permet non seulement d'obtenir un PAPR optimal pour une modulation multiporteuse, mais aussi de garantir que ce dernier diminue à mesure que la densité ρ du système multiporteuse augmente. En particulier, dans un contexte FTN ($\rho > 1$), ce choix judicieux engendre un PAPR systématiquement meilleur que celui des systèmes STN, à nombre de sous-porteuses égal.

Nous présentons ces résultats en définissant tout d'abord le modèle de PAPR utilisé, puis les conditions à respecter pour garantir l'optimalité du PAPR d'un système multiporteuse. Nous

montrerons ensuite comment choisir des formes d'ondes d'émission respectant ces conditions dans un contexte FTN. Enfin, ces résultats seront vérifiés et analysés par le biais de simulations numériques.

5.2 Définition du PAPR

En reprenant l'expression du signal modulé à temps continu $s(t)$, et en supposant que son support temporel est borné à $[0; T]$ (on prendra $T \geq T_0$ et, la plupart du temps $T = T_0$ [SSJ06]), on définit simplement le PAPR par :

$$\text{PAPR} = \frac{\max_{t \in [0; T]} |s(t)|^2}{\mathbb{E} \left\{ \frac{1}{T} \int_0^T |s(t)|^2 dt \right\}}. \quad (5.1)$$

À temps discret, et à échantillonnage critique, on adapte la formule précédente en utilisant l'expression du signal modulé $s[k]$ à temps discret, dont on considère le support borné à $[0; N_T[$, avec $N_T = T/T_e$. L'expression du PAPR devient alors :

$$\text{PAPR}_d = \frac{\max_{k \in [0; N_T[} |s[k]|^2}{\mathbb{E} \left\{ \frac{1}{N_T} \sum_{k=0}^{N_T-1} |s[k]|^2 \right\}}. \quad (5.2)$$

Notons que PAPR_d est généralement significativement inférieur à PAPR. Pour pallier à ce problème, il suffit d'interpoler $s[k]$ par un facteur N_i , de telle sorte que PAPR_d tende vers PAPR à mesure que N_i devient grand ($N_i = 4$ est généralement suffisant).

5.3 Conditions d'optimalité

Il peut être intéressant de chercher comment choisir les paramètres d'un système multiporteuse (prototype, densité), afin d'optimiser le PAPR. Cependant, il est à noter que le PAPR n'est pas une mesure déterministe, du moins pas pour un système multiporteuse. Ainsi, afin de comparer et d'optimiser les performances d'un système en PAPR, on utilisera sa fonction de répartition complémentaire : $\mathbb{P} \{ \text{PAPR} > \gamma \}$.

Les travaux de Skrzypczak [SSJ06] ont permis de déterminer une expression analytique de la fonction de répartition complémentaire du PAPR dans le cas discret. Ainsi, si

1. les symboles $c_{m,n}$ sont indépendants et identiquement distribués ;
2. le prototype d'émission $g[k]$ est à support temporel borné ;
3. $\sum_{n \in \mathbb{Z}} |g[k - nN]|^2 > 0$ pour tout $k \in \mathbb{Z}$;
4. les échantillons du signal modulé $s[k]$ sont indépendants ;

alors, pour M suffisamment grand (typiquement, $M > 8$), une bonne approximation de la fonction de répartition complémentaire du PAPR est donnée par :

$$\mathbb{P} \{ \text{PAPR}_d > \gamma \} \approx 1 - \prod_{k=0}^{N-1} \left(1 - e^{-\gamma x[k]} \right), \quad (5.3)$$

5.4. Prototypes d'émission optimaux

avec :

$$x[k] = \frac{\|g\|^2}{N \sum_{n \in \mathbb{Z}} |g[k - nN]|^2}. \quad (5.4)$$

Notons que les hypothèses 1 et 2 sont presque toujours respectées en pratique. Dans le cadre des simulations présentées dans la partie 5.5 elles seront imposées : les symboles seront générés de telle manière à être IID, les prototypes utilisés seront tronqués si nécessaire, de manière à ce que leur support soit inférieur ou égal à $L_g = 32N$ (on rappelle que N est le nombre d'échantillons par symbole multiporteuse). Quant au nombre de sous-porteuses, il sera porté à $M = 1024$.

L'expression (5.3) est minimisée lorsque $x[k] = 1 \ \forall k \in \mathbb{Z}$. Lorsque cette condition supplémentaire est vérifiée, on obtient

$$\mathbb{P}\{\text{PAPR}_d > \gamma\} \approx 1 - (1 - e^{-\gamma})^N. \quad (5.5)$$

Il est à noter que cette limite est atteinte par l'OFDM sans préfixe cyclique (prototype rectangulaire) [Cha16, Chapitre 6]. On remarque également qu'à M fixé, si les conditions 1 à 4 sont respectées avec $x[k] = 1$, alors plus N est petit, plus le PAPR est faible. Cela signifie que l'augmentation de la densité permet, dans certaines conditions, de diminuer le PAPR.

5.4 Prototypes d'émission optimaux

Intéressons-nous tout d'abord à la condition 4, spécifiant que les échantillons $s[k]$ du signal modulé doivent être indépendants. Sous l'hypothèse que les $s[k]$ proviennent d'une distribution gaussienne, cela équivaut à dire que la densité spectrale de puissance $\gamma_s(f)$ du signal modulé doit être constante dans la bande $B_s = MF_0$. Or, d'après [Cho05], dans le cas où les symboles sont IID de variance σ_c^2 , cette densité spectrale s'écrit :

$$\gamma_s(f) = \sigma_c^2 \sum_{m \in \mathbb{Z}} |G_{m,0}(f)|^2 = \sigma_c^2 \sum_{m \in \mathbb{Z}} |G(f - mF_0)|^2, \quad f \in \mathbb{R}, \quad (5.6)$$

avec $G_{m,0}(f)$ et $G(f)$ la transformée de Fourier de $g_{m,0}(t)$ et $g(t)$, respectivement. La propriété d'indépendance des échantillons du signal émis ne dépend donc plus que du prototype d'émission. Remarquons que la condition $\sum_{m \in \mathbb{Z}} |G(f - mF_0)|^2 = A$ (avec $A \in \mathbb{R}^{+*}$ une constante) revient à imposer le respect du critère de Nyquist par $g(t)$ à la cadence $1/F_0$. Or si $\mathbf{g} = \{g(t - nT_0)e^{j2\pi m F_0 t}\}_{(m,n) \in \mathbb{Z}^2}$ est une frame étroite alors, d'après le théorème de Wexler–Raz (théorème 1.1), la famille $\{g(t - n/F_0)e^{j2\pi m t/T_0}\}_{(m,n) \in \mathbb{Z}^2}$ est une famille orthonormale, ce qui implique que $g(t)$ respecte le critère de Nyquist à la cadence $1/F_0$. On en conclut que sous l'hypothèse de normalité de l'enveloppe complexe du signal modulé échantillonné $s[k]$, la condition 4 est validée si $\mathbf{g}$ est une frame étroite. De plus, dans le contexte de l'OFDM classique, il a été montré que la distribution de l'enveloppe complexe du signal à temps continu converge vers une gaussienne lorsque le nombre de sous-porteuses tend vers l'infini [WKG10]. Pour les systèmes multiporteuses FTN tels que défini dans la partie 1.4.2 et étudiés dans ce manuscrit, il serait possible de déterminer la pertinence de l'hypothèse de normalité de $s[k]$ par une analyse empirique semblable à celle menée en partie 3.3.

Intéressons-nous maintenant à la condition 3. Si $\mathbf{g}$ est une frame étroite de borne A , alors [Chr08, Théorème 9.2.1] indique que :

$$\sum_{n \in \mathbb{Z}} |g(t - nT_0)|^2 = AF_0 \quad \forall t \in \mathbb{R}. \quad (5.7)$$

$\mathbf{g}$ étant une frame étroite de Gabor, on a $A = \rho \|g\|^2 = \|g\|^2 / (F_0 T_0)$, donc :

$$\sum_{n \in \mathbb{Z}} |g(t - nT_0)|^2 = \frac{\|g\|^2}{T_0} \quad \forall t \in \mathbb{R}. \quad (5.8)$$

En discrétisant cette expression de la même manière que dans le chapitre 1, partie 1.4.2.2, on obtient :

$$\sum_{n \in \mathbb{Z}} |g[k - nN]|^2 = \frac{\|g\|^2}{N} > 0 \quad \forall k \in \mathbb{Z}. \quad (5.9)$$

La condition 3 est donc respectée en utilisant une frame étroite.

Enfin, on vérifie que la condition $x[k] = 1 \quad \forall k \in \mathbb{Z}$ est respectée par l'usage d'un prototype d'émission engendrant une frame étroite en injectant (5.9) dans l'expression de $x[k]$:

$$x[k] = \frac{\|g\|^2}{N \sum_{n \in \mathbb{Z}} |g[k - nN]|^2} = \frac{\|g\|^2}{N \frac{\|g\|^2}{N}} = 1, \quad \forall k \in \mathbb{Z}. \quad (5.10)$$

Ainsi, quand les conditions 1, 2 sont respectées par l'émetteur multiporteuse, que le nombre de sous-porteuses est strictement supérieur à 8, et sous l'hypothèse que les échantillons du signal émis suivent une distribution gaussienne, utiliser une frame de Gabor étroite en émission permet de minimiser l'approximation de la probabilité d'obtenir des valeurs élevées de PAPR_d . On peut alors écrire (voir 5.1) :

$$\mathbb{P}\{\text{PAPR}_d > \gamma\} \approx 1 - (1 - e^{-\gamma})^N. \quad (5.11)$$

Ainsi, à un nombre de sous-porteuses M donné, et à condition d'utiliser un prototype d'émission engendrant une frame étroite, la probabilité d'obtenir un PAPR_d élevé diminue à mesure que la densité ρ augmente.

5.5 Simulations

L'affirmation de l'optimalité des prototypes d'émission engendrant des frames étroites en terme de PAPR, démontrée dans la partie précédente, est ici confirmée par le biais de simulations.

Les figures 5.1 et 5.2 comparent les fonctions de répartition complémentaires de PAPR_d . On remarque tout d'abord une très bonne correspondance entre les fonctions de répartition complémentaires obtenues par simulation et leurs approximations théoriques, données par (5.11). On observe également la tendance selon laquelle le PAPR s'améliore à mesure que la densité ρ augmente pour les systèmes utilisant des frames étroites. On note enfin que le prototype ne permettant pas de former une frame étroite utilisé ici présente un PAPR significativement dégradé.

5.5. Simulations

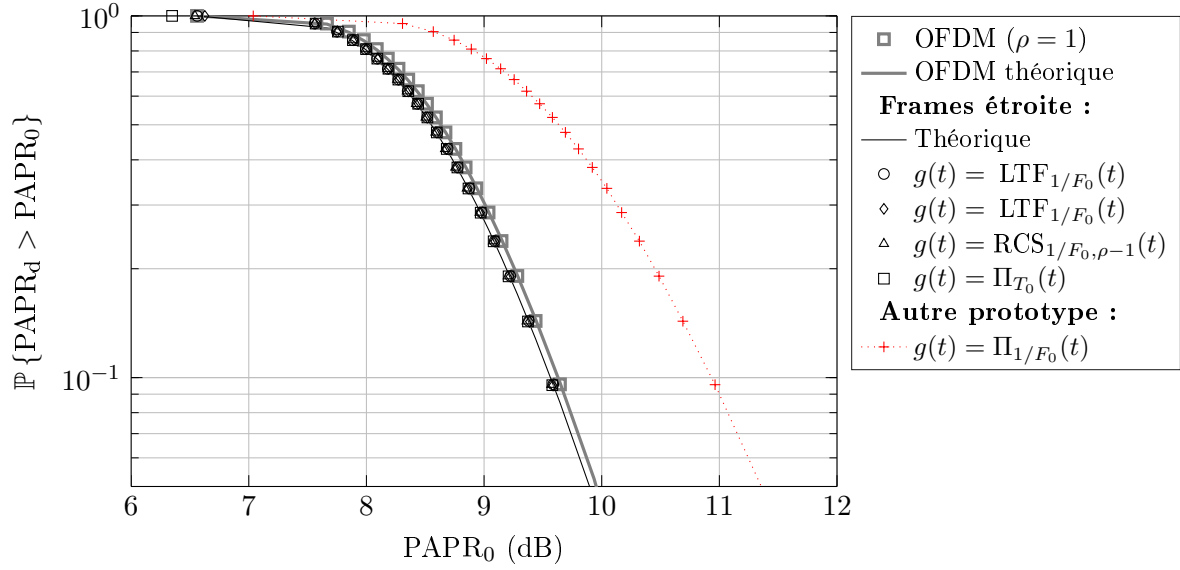


FIGURE 5.1 – Comparaison des fonctions de répartition complémentaires du PAPR discret pour différents prototypes, $M = 1024$ sous-porteuses et une densité $\rho = 16/14$. Les courbes théoriques sont données par (5.11). On remarque la similitude du comportement en PAPR des prototypes utilisant des frames étroites, lequel est ici légèrement meilleur qu'en OFDM.

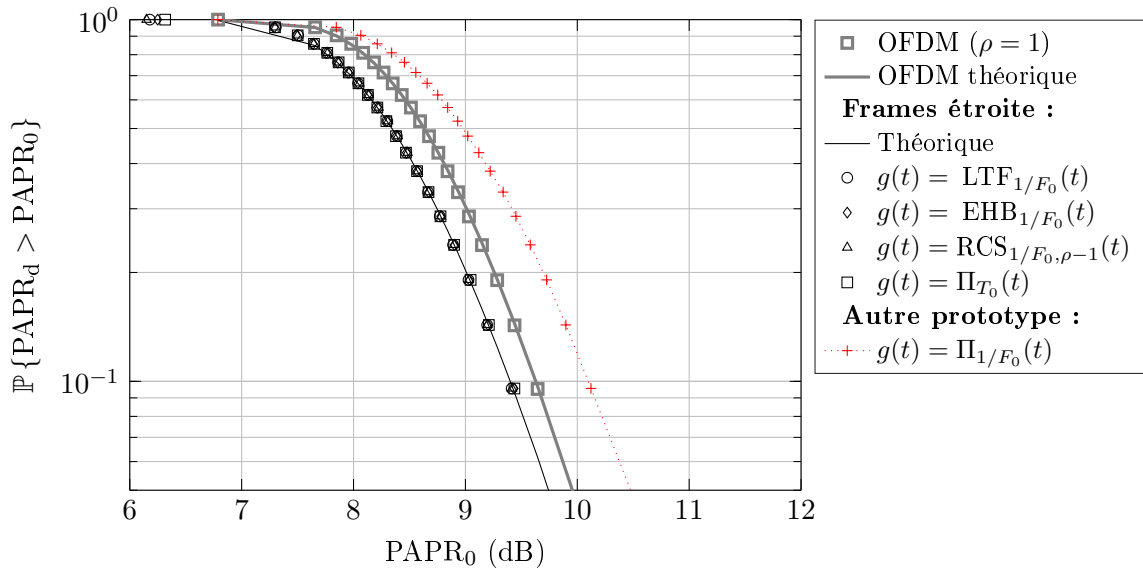


FIGURE 5.2 – Comparaison des fonctions de répartition complémentaires du PAPR discret pour différents prototypes, $M = 1024$ sous-porteuses et une densité $\rho = 16/10$. Les courbes théoriques sont données par (5.11). Comme prévu, comparé à la figure 5.1, l'écart de performances entre le PAPR des frames étroites et celui de l'OFDM est plus marqué, grâce à une densité plus élevée.

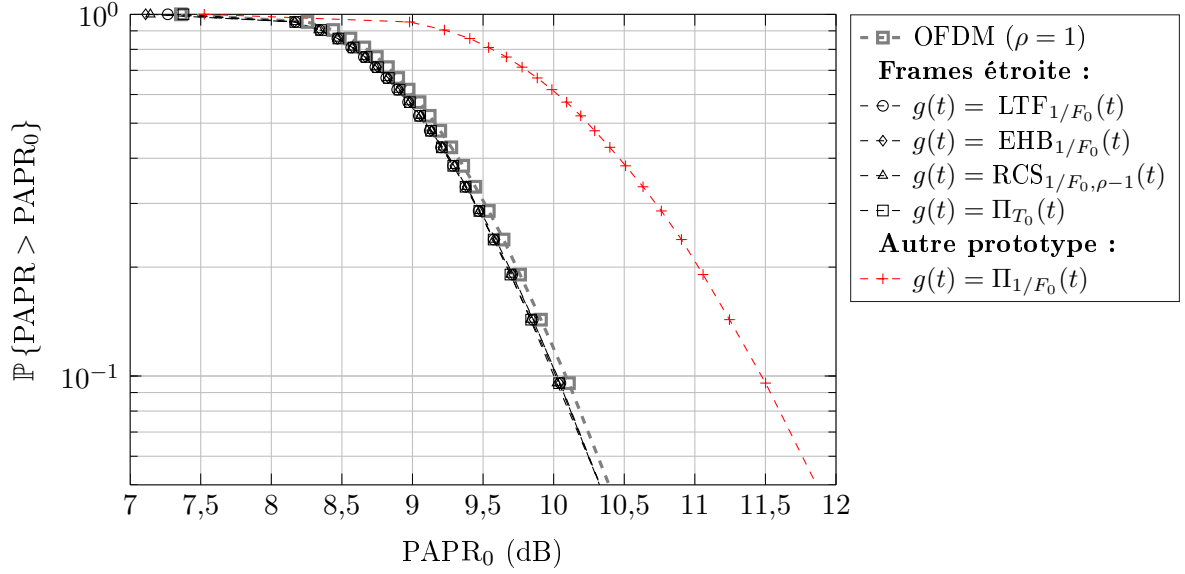


FIGURE 5.3 – Comparaison des fonctions de répartition complémentaires du PAPR du signal $s[k]$ interpolé d'un facteur $N_i = 4$ pour différents prototypes, $M = 1024$ sous-porteuses et une densité $\rho = 16/14$. On observe les mêmes tendances que sur la figure 5.1.

Pour approcher plus finement le comportement en PAPR du système dans des conditions réalistes, on compare en figures 5.3 et 5.4 les fonctions de répartition complémentaires du PAPR obtenues après interpolation du signal $s[k]$ par un facteur $N_i = 4$. On y remarque que les tendances sont exactement les mêmes que pour PAPR_d , c'est-à-dire que les prototypes d'émission engendrant des frames étroites permettent de surpasser le comportement en PAPR de l'OFDM classique, et ce d'autant plus que ρ est grand devant 1. Le prototype rectangulaire de longueur $1/F_0$, ne formant pas une frame étroite, possède quant à lui un comportement en PAPR dégradé, même comparé à l'OFDM classique.

5.6 Conclusion

Dans ce chapitre, nous avons montré que choisir des prototypes d'émission permettant de former des frames étroites permet d'améliorer le comportement en PAPR des modulations multiporteuses en augmentant la densité de transmission ρ (pour un nombre de sous-porteuses M donné). Ces prototypes sont les mêmes que ceux utilisés dans les chapitres 3 et 4 afin de traiter le terme d'auto-interférence propre aux modulations FTN. On y retrouve en particulier les prototypes rectangulaires (de largeur T_0), et en racine de cosinus surélevé (utilisant un facteur d'amortissement $\alpha \leq \rho - 1$, et à construire de la même manière que si l'on souhaitait obtenir une transmission monoporteuse sans IES à la cadence $R = F_0$). Dans ce contexte, la distribution du PAPR pour un système multiporteuse FTN avec une densité ρ et un nombre de sous-porteuses

5.6. Conclusion

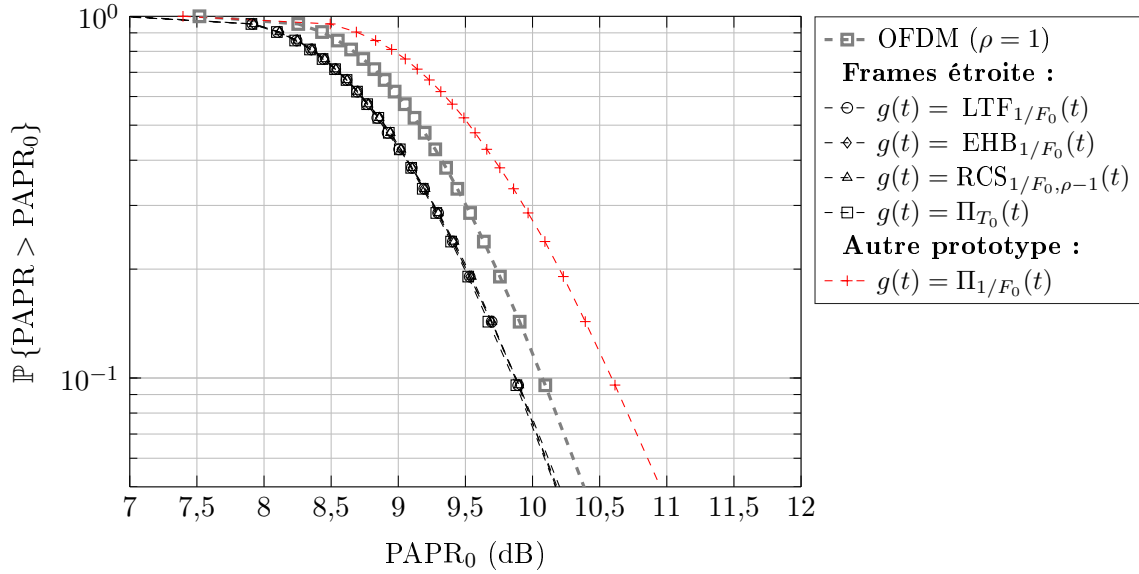


FIGURE 5.4 – Comparaison des fonctions de répartition complémentaires du PAPR du signal $s[k]$ interpolé d'un facteur $N_i = 4$ pour différents prototypes, $M = 1024$ sous-porteuses et une densité $\rho = 16/10$. On observe les mêmes tendances que sur la figure 5.2.

M est la même que celle d'un système OFDM utilisant $M/\rho < M$ porteuses.

Bien que le comportement en PAPR d'un système multiporteuse FTN reste objectivement mauvais comparé, par exemple, aux systèmes monoporteuses à enveloppe constante, son amélioration au travers de l'augmentation de ρ constitue un effet secondaire appréciable.

Publications associées à ce chapitre

Conférences

1. Cyrille Siclet, Damien Roque, Alexandre Marquet, Laurent Ros « Improving spectral efficiency while reducing PAPR using faster-than-Nyquist multicarrier signaling ». In : *Communication Technologies for Vehicles : 12th International Workshop, Nets4Cars / Nets4Trains / Nets4Aircraft 2017 proceedings*, pp. 21-32 (Mai 2017), Toulouse, France. DOI : https://doi.org/10.1007/978-3-319-56880-5_3.

Conclusion et perspectives

L'objet principal de ce travail de thèse était de proposer des systèmes multiporteuses au-delà de la cadence de Nyquist à faible complexité algorithmique, et d'évaluer leur comportement sur canal radiomobile.

Les systèmes travaillant au-delà de la cadence de Nyquist permettent, en théorie et sous certaines conditions, d'augmenter l'efficacité spectrale de la modulation tout en conservant les performances en terme de probabilité d'erreur binaire du même système travaillant à la cadence de Nyquist. Dans un contexte où les ressources spectrales et énergétiques se font de plus en plus rares, disposer de systèmes de communication présentant une haute efficacité spectrale à des rapports signal-à-bruit aussi faibles que possible est indispensable. Dans un scénario radio-mobile, il est également nécessaire que ces systèmes de communication permettent une bonne adaptation au canal de transmission, afin de diminuer la complexité algorithmique des récepteurs et de permettre miniaturisation et meilleure autonomie. Pour répondre à ce besoin, on utilise généralement des modulations multiporteuses, connues pour leur capacité d'adaptation aux canaux de transmission sélectifs en fréquence et/ou en temps, ce qui leur permet en effet d'utiliser des stratégies d'égalisation du canal à très faible complexité en réception (typiquement, à un coefficient par sous-porteuse sur canal sélectif en fréquence, et à un coefficient par symbole sur canal radiomobile). Au cours de cette thèse, nous avons montré qu'il est possible de concevoir des systèmes multiporteuses au-delà de la cadence de Nyquist permettant un gain appréciable en efficacité spectrale, tout en conservant une bonne capacité d'adaptation au canal de transmission, ainsi qu'une complexité algorithmique raisonnable.

Synthèse du travail et principales contributions

Dans le chapitre 1, nous avons évoqués quelques éléments fondamentaux du domaine du traitement du signal appliqué aux communications numériques. Nous avons développé les modèles et théories utilisés dans le reste du manuscrit. Ainsi, nous avons premièrement décrit la modélisation de la chaîne de communication, issue des travaux de Shannon, puis une modélisation empirique des canaux de transmission sélectifs en fréquence et/ou en temps. Nous avons ensuite présenté le principe des modulations multiporteuses, et mis en exergue le lien étroit entre cette technique de transmission et les techniques d'analyse temps-fréquence. Par conséquent, nous nous sommes ensuite attardés sur une brève présentation de la théorie des frames (omniprésente en analyse temps-fréquence), et de quelques outils qui en sont issus et qui sont utilisés dans la suite du manuscrit. Ensuite, nous avons décrit le type de modulation/démodulation multiporteuse adopté dans cette thèse, sa modélisation associée, et les conditions sous lesquelles il est possible d'effectuer une égalisation à faible complexité. Enfin, nous avons abordé la question des communications à haute efficacité spectrale en introduisant tout d'abord la limite de Shannon sur canal à bruit additif blanc gaussien, puis en présentant plusieurs techniques permettant d'augmenter l'efficacité spectrale et/ou de se rapprocher de la limite de Shannon : l'augmentation

du nombre de bits par symbole, le codage canal, et l'augmentation de la densité de signalisation (menant aux techniques au-delà de la cadence de Nyquist lorsque cette densité dépasse l'unité).

Après avoir présenté le principe général sur lequel repose les communications au-delà de la cadence de Nyquist, nous avons proposé un tour d'horizon de la littérature existante sur le sujet dans le chapitre 2. Pour ce faire, et étant donné que la plupart des travaux existant ont été effectués dans le cadre de modulations monoporteuses, nous avons donné les relations d'entrée/sortie discrètes équivalentes en bande de base pour les systèmes au-delà de la cadence de Nyquist, à la fois en monoporteuse et en multiporteuse. Nous avons ensuite défini la notion de distance euclidienne minimale entre signaux et son lien direct avec la meilleure probabilité d'erreur atteignable (minimisant la probabilité d'erreur). Nous avons pu remarquer que, contrairement aux communications en deçà de la cadence de Nyquist où cette distance ne dépend que de la constellation choisie, elle dépend également de l'impulsion de mise en forme choisie dans le cas FTN. Une autre particularité des communications FTN, découverte par Mazo en 1975, réside dans le fait que la distance euclidienne minimale entre signaux reste constante (ou presque) lorsque la densité dépasse l'unité mais reste inférieure à une valeur critique, nommée depuis « limite de Mazo ». Nous avons enfin évoqué l'extension de cette limite aux modulations multiporteuses, introduite par Rusek en 2005. Par la suite, nous avons décrit quelques techniques de détection des modulations FTN, et avons mis l'accent sur le fait que parmi ces techniques, seules celles qui sont optimales ou quasi-optimales vis-à-vis de la probabilité d'erreur binaire permettent d'exploiter les propriétés des modulations FTN en terme de distance minimale. Nous avons notamment montré que les approches par treillis, permettant justement d'obtenir des performances optimales ou quasi-optimales, sont accessibles aux modulations monoporteuses à condition de rester à des valeurs de densité modérées et d'utiliser une constellation à faible nombre d'états. Elles sont en revanche inaccessibles à la plupart des systèmes multiporteuses, dans la mesure où le nombre d'états du treillis croît de manière exponentielle avec le nombre de sous-porteuses. Ainsi, dans l'optique de permettre une détection quasi-optimale sans pour autant recourir à des approches à forte complexité algorithmique, nous avons présenté la technique de turbo-égalisation, qui traite conjointement décodage source et égalisation selon un processus itératif. Nous nous sommes en particulier intéressés aux techniques de turbo égalisation faisant intervenir des égaliseurs linéaires, et avons donné les expressions générales des éléments constitutifs de tels égaliseurs.

Détermination de prototypes d'émission/réception optimaux au sens de la maximisation du rapport signal-à-interférence-plus-bruit sur canal à bruit additif blanc gaussien

Le chapitre 3 est dédié à la détermination d'impulsions de mise en forme (ou prototype d'émission) et de prototypes de réception optimaux au sens de la maximisation du rapport signal-à-interférence-plus-bruit sur canal à bruit additif blanc gaussien. Nous avons montré que choisir un prototype d'émission formant une frame étroite en conjonction avec un prototype de réception proportionnel à celui d'émission permet cette optimalité. Notons que des prototypes de cette sorte peuvent être obtenus à partir d'impulsions de mise en forme de systèmes en deçà

Conclusion

de la cadence de Nyquist, en vertu du théorème de Wexler–Raz énoncé dans le chapitre 1. En sus de cette démonstration, nous avons mis en avant le fait que les prototypes de réception proportionnels aux prototypes d’émission (ce qui revient à effectuer un filtre adapté) maximisent le rapport signal-à-bruit (mais pas nécessairement le rapport signal-à-interférence), et que les prototypes de réception tels que la famille de réception est la frame duale canonique de la famille (nécessairement une frame) d’émission permettent de maximiser le rapport signal-à-interférence (mais pas nécessairement le rapport signal-à-bruit).

Optimalité des frames étroites en terme de minimisation de l’erreur quadratique moyenne dans un contexte de turbo-égalisation à suppression souple de l’interférence sur canal à bruit additif blanc gaussien

La première partie du chapitre 4 prouve que l’utilisation de frames étroites, dont on a montré l’optimalité vis-à-vis de la maximisation du rapport signal-à-interférence-plus-bruit dans le chapitre 3, en émission et en réception permet de minimiser l’erreur quadratique moyenne en sortie d’une structure de suppression souple de l’interférence en présence d’un canal (non sélectif) à bruit additif blanc gaussien. Nous avons ensuite évalué les performances de ce système en termes de taux d’erreur binaire en le simulant en collaboration avec un décodeur canal dans une boucle turbo. Ces simulations ont montré la capacité de ce système à compenser, au moins en partie, la perte d’efficacité spectrale liée à l’introduction du codage canal. Dans le meilleur cas, en QPSK, il permet un gain en efficacité spectrale de 1,14 par rapport au système non-codé à la cadence de Nyquist, et ce pour un rapport signal-à-bruit plus faible de 3,6 dB. Par rapport au système codé à la cadence de Nyquist, il permet un gain en efficacité spectrale de 2,29 par rapport au système codé, pour la même valeur de rapport signal-à-bruit.

Invariance des méthodes pour choisir les paramètres du système multiporteuse en fonction de la sélectivité temps-fréquence du canal de transmission

Cette contribution est traitée dans la deuxième partie du chapitre 4. Elle débute en étendant le système à annulation souple de l’interférence au cas général des canaux radiomobiles (ce qui couvre par conséquent le cas des canaux uniquement sélectifs en fréquence ou uniquement sélectifs en temps), sous l’hypothèse que ce dernier est correctement approché par un coefficient par symbole. Nous avons ensuite simulé ce système associé à un décodeur canal dans une configuration turbo, et avons montré que l’influence des paramètres du système multiporteuse (nombre de sous-porteuses, localisation temps-fréquence du prototype d’émission) a le même impact que dans le cas d’un système en deçà de la cadence de Nyquist. Ainsi, sur canal sélectif en fréquence, nous avons observé qu’il est souhaitable d’avoir un grand nombre de porteuses ainsi qu’un prototype d’émission bien localisé en fréquence. Sur canal radiomobile, nous avons remarqué qu’il est avant tout intéressant de disposer de prototypes d’émission présentant une bonne localisation temps-fréquence, et que le nombre de sous-porteuses doit être adapté au profil temps-fréquence du canal (on pourra, par exemple, se permettre un plus grand nombre de porteuses si le canal est plus sélectif en fréquence qu’en temps, tandis que dans la situation inverse, il sera nécessaire de

diminuer le nombre de porteuses). Enfin, nous avons observé de meilleures performances lorsque les systèmes utilisent des frames étroites en émission et en réception.

Optimalité des frames étroites au sens de la minimisation du facteur de crête

Le facteur de crête, généralement évalué sous la forme du rapport entre la puissance moyenne et de la puissance maximale du signal et désigné par l'acronyme anglo-saxon PAPR, est l'un des principaux points faibles des systèmes multiporteuses, limitant notamment son utilisation dans des terminaux embarqués. Dans le chapitre 5, nous avons tout d'abord montré que les approximations de la fonction de répartition complémentaire du PAPR dérivées dans le cadre des communications multiporteuses en deçà de la cadence de Nyquist restent valable au delà. Ensuite, nous avons démontré que l'utilisation de prototypes d'émission formant des frames étroites permet d'atteindre le minimum de cette fonction de répartition complémentaire. Dans ce cas, le PAPR est systématiquement plus faible que pour les systèmes en deçà de la cadence de Nyquist, et l'est d'autant plus que la densité augmente.

Perspectives

Dans tous les systèmes turbo-égalisés proposés dans cette thèse, nous avons utilisé le même code convolutif $(7, 5)$, qui a l'avantage d'être utilisé par de nombreux auteurs, mais qui n'est pas forcément le plus adapté à notre application. Il serait intéressant d'évaluer l'impact de différents codes sur les performances du système, voire de chercher des codes adaptés aux transmissions multiporteuses au-delà de la cadence de Nyquist, dans la lignée de [AZ14].

Concernant le système à suppression souple de l'interférence sur canal sélectif en fréquence et/ou en temps, il conviendrait tout d'abord d'évaluer la sensibilité du système aux erreurs d'estimation du canal. Il serait également intéressant de développer des techniques de suivi du canal, voir d'estimation aveugle de ce dernier dans le contexte particulier d'une transmission au-delà de la cadence de Nyquist.

Il serait également intéressant de mener une optimisation des prototypes d'émission similaire à celle menée dans ce manuscrit, mais pour les modulations OFDM/OQAM, en complément des travaux de thèse menés par Lahbabi [Lah17].

Enfin, nous avons vu dans le chapitre 4 qu'utiliser des frames étroites ne permet pas de minimiser l'erreur quadratique moyenne en sortie du suppresseur souple de l'interférence si le canal est sélectif en fréquence et/ou en temps. Il serait intéressant d'évaluer le gain en performances permis, par exemple, par un filtrage de Wiener à deux dimensions. Cette solution apparaît de plus en plus abordable au regard des progrès en termes de performances et de miniaturisation des processeurs graphiques d'une part, et du développement de solutions à base de radio logicielle d'autre part.

Estimation au sens de la minimisation de l'erreur quadratique moyenne

Dans cette annexe, nous développons l'expression générale de l'estimateur minimisant l'erreur quadratique moyenne.

Notons $\theta \in \mathbf{E}$ un paramètre à estimer, et $\hat{\theta} \in \mathbf{V}$ son estimateur à partir de l'observation $y \in \mathbf{Y}$. $\hat{\theta}$ est donc une fonction de l'observation (on notera $\hat{\theta} = g(y)$), et on cherche l'estimateur optimal $\hat{\theta}_{\text{opt}} = g_{\text{opt}}(y)$ au sens de l'EQM :

$$\hat{\theta}_{\text{opt}} \in \underset{\hat{\theta} \in \mathbf{V}}{\operatorname{argmin}} \mathbb{E} \left\{ \left| \theta - \hat{\theta} \right|^2 \right\}, \quad (\text{A.1})$$

ce qui suppose que les moments d'ordre 1 et 2 de θ et $\hat{\theta}$ sont définis. On peut réécrire la fonction de coût de la manière suivante :

$$\mathbb{E} \left\{ \left| \theta - \hat{\theta} \right|^2 \right\} = \mathbb{E}_y \left\{ \mathbb{E} \left\{ \left| \theta - \hat{\theta} \right|^2 \middle| y \right\} \right\}, \quad (\text{A.2})$$

de manière à ce que le problème d'optimisation puisse s'écrire :

$$\hat{\theta}_{\text{opt}} \in \underset{\hat{\theta} \in \mathbf{V}}{\operatorname{argmin}} \mathbb{E} \left\{ \left| \theta - \hat{\theta} \right|^2 \middle| y \right\}, \quad \forall y \in \mathbf{Y}. \quad (\text{A.3})$$

Développons maintenant $\mathbb{E} \left\{ \left| \theta - \hat{\theta} \right|^2 \middle| y \right\}$ en remplaçant $\hat{\theta}$ par $g(y)$ pour mieux identifier les termes qui ne sont pas aléatoires vis-à-vis de l'espérance conditionnelle sur y :

$$\mathbb{E} \left\{ \left| \theta - \hat{\theta} \right|^2 \middle| y \right\} = \mathbb{E} \left\{ |\theta|^2 \middle| y \right\} + \mathbb{E} \left\{ |g(y)|^2 \middle| y \right\} - \mathbb{E} \left\{ \theta^* g(y) \middle| y \right\} - \mathbb{E} \left\{ \theta g^*(y) \middle| y \right\} \quad (\text{A.4})$$

$$= \mathbb{E} \left\{ |\theta|^2 \middle| y \right\} + |g(y)|^2 - g(y) \mathbb{E} \left\{ \theta^* \middle| y \right\} - g^*(y) \mathbb{E} \left\{ \theta \middle| y \right\} \quad (\text{A.5})$$

$$= \mathbb{E} \left\{ |\theta|^2 \middle| y \right\} + \left| \hat{\theta} \right|^2 - \hat{\theta} \mathbb{E} \left\{ \theta^* \middle| y \right\} - \hat{\theta}^* \mathbb{E} \left\{ \theta \middle| y \right\}. \quad (\text{A.6})$$

Pour minimiser cette quantité, on cherche les valeurs $\hat{\theta}$ pour lesquelles le gradient s'annule. Calculons tout d'abord ce gradient :

$$\nabla_{\hat{\theta}} \mathbb{E} \left\{ \left| \theta - \hat{\theta} \right|^2 \middle| y \right\} = \nabla_{\hat{\theta}} \left| \hat{\theta} \right|^2 - \mathbb{E} \left\{ \theta^* \middle| y \right\} \nabla_{\hat{\theta}} \hat{\theta} - \mathbb{E} \left\{ \theta \middle| y \right\} \nabla_{\hat{\theta}} \hat{\theta}^*, \quad (\text{A.7})$$

avec :

$$\nabla_{\hat{\theta}} |\hat{\theta}|^2 = \nabla_{\hat{\theta}} \left(\mathcal{R} \{ \hat{\theta} \}^2 + \mathcal{I} \{ \hat{\theta} \}^2 \right) = \frac{\partial \left(\mathcal{R} \{ \hat{\theta} \}^2 + \mathcal{I} \{ \hat{\theta} \}^2 \right)}{\partial \mathcal{R} \{ \hat{\theta} \}} + j \frac{\partial \left(\mathcal{R} \{ \hat{\theta} \}^2 + \mathcal{I} \{ \hat{\theta} \}^2 \right)}{\partial \mathcal{I} \{ \hat{\theta} \}} \quad (\text{A.8})$$

$$= 2\mathcal{R} \{ \hat{\theta} \} + j2\mathcal{I} \{ \hat{\theta} \} = 2\hat{\theta}, \quad (\text{A.9})$$

$$\nabla_{\hat{\theta}} \hat{\theta} = \nabla_{\hat{\theta}} \left(\mathcal{R} \{ \hat{\theta} \} + j\mathcal{I} \{ \hat{\theta} \} \right) = \frac{\partial \left(\mathcal{R} \{ \hat{\theta} \} + j\mathcal{I} \{ \hat{\theta} \} \right)}{\partial \mathcal{R} \{ \hat{\theta} \}} + j \frac{\partial \left(\mathcal{R} \{ \hat{\theta} \} + j\mathcal{I} \{ \hat{\theta} \} \right)}{\partial \mathcal{I} \{ \hat{\theta} \}} \quad (\text{A.10})$$

$$= 1 + j^2 = 0, \quad (\text{A.11})$$

$$\nabla_{\hat{\theta}} \hat{\theta}^* = \nabla_{\hat{\theta}} \left(\mathcal{R} \{ \hat{\theta} \} - j\mathcal{I} \{ \hat{\theta} \} \right) = \frac{\partial \left(\mathcal{R} \{ \hat{\theta} \} - j\mathcal{I} \{ \hat{\theta} \} \right)}{\partial \mathcal{R} \{ \hat{\theta} \}} + j \frac{\partial \left(\mathcal{R} \{ \hat{\theta} \} - j\mathcal{I} \{ \hat{\theta} \} \right)}{\partial \mathcal{I} \{ \hat{\theta} \}} \quad (\text{A.12})$$

$$= 1 - j^2 = 2, \quad (\text{A.13})$$

ce qui, une fois réinjecté dans (A.7), donne :

$$\nabla_{\hat{\theta}} \mathbb{E} \left\{ \left| \theta - \hat{\theta} \right|^2 \middle| y \right\} = 2\hat{\theta} - 2\mathbb{E} \{ \theta | y \}. \quad (\text{A.14})$$

Cette expression étant nulle quand $\hat{\theta} = \mathbb{E} \{ \theta | y \}$, on en déduit que l'estimateur $\hat{\theta}$ de θ minimisant l'erreur quadratique moyenne est :

$$\hat{\theta} = \mathbb{E} \{ \theta | y \} \quad \forall y \in \mathbf{Y}. \quad (\text{A.15})$$

Bibliographie

- [Abe+16] A. ABELLO et al. « On zero-forcing equalization for short-filtered multicarrier faster-than-Nyquist signaling ». In : *2016 50th Asilomar Conference on Signals, Systems and Computers*. Nov. 2016, p. 904–908. DOI : 10.1109/ACSSC.2016.7869180 (cf. p. 62).
- [Aca17] ACADÉMIE FRANÇAISE. *Le Dictionnaire de l'Académie française*. 9^e éd. [En ligne : <http://atilf.atilf.fr/academie9.htm>]. 1992-2017 (cf. p. 7, 8).
- [ARÖ13] J. B. ANDERSON, F. RUSEK et V. ÖWALL. « Faster-Than-Nyquist Signaling ». In : *Proceedings of the IEEE* 101.8 (août 2013), p. 1817–1830. ISSN : 0018-9219. DOI : 10.1109/JPROC.2012.2233451 (cf. p. 2, 47, 52, 55, 58).
- [AZ14] J. B. ANDERSON et M. ZEINALI. « Analysis of best high rate convolutional codes for faster than Nyquist turbo equalization ». In : *2014 IEEE International Symposium on Information Theory*. Juin 2014, p. 1987–1991. DOI : 10.1109/ISIT.2014.6875181 (cf. p. 101, 124).
- [Bah+74] L. BAHL et al. « Optimal decoding of linear codes for minimizing symbol error rate (Corresp.) ». In : *IEEE Transactions on Information Theory* 20.2 (mar. 1974), p. 284–287. ISSN : 0018-9448. DOI : 10.1109/TIT.1974.1055186 (cf. p. 54).
- [Bat87] Gérard BATTAIL. « Pondération des symboles décodés par l'algorithme de Viterbi ». In : *Annales des Télécommunications* 42.1 (1987), p. 31–38. ISSN : 1958-9395. DOI : 10.1007/BF02996166 (cf. p. 41, 58).
- [Bel63] P. BELLO. « Characterization of randomly time-variant linear channels ». In : *IEEE Transactions on Communications Systems* 11.4 (1963), p. 360–393. DOI : 10.1109/TCOM.1963.1088793 (cf. p. 15).
- [Ber+07] Claude BERROU et al. *Codes et turbocodes (sous la direction de Claude Berrou)*. Iris. Paris : Springer, 2007. ISBN : 978-2-287-32739-1 (cf. p. 8, 11, 41, 56, 57).
- [BGT93] C. BERROU, A. GLAVIEUX et P. THITIMAJSHIMA. « Near Shannon limit error-correcting coding and decoding : Turbo-codes. 1 ». In : *Communications, 1993. ICC '93 Geneva. Technical Program, Conference Record, IEEE International Conference on*. T. 2. Mai 1993, 1064–1070 vol.2. DOI : 10.1109/ICC.1993.397441 (cf. p. 41).
- [Cha16] Marwa CHAFII. « Study of a new multicarrier waveform with low PAPR ». Thesis. CentraleSupélec, oct. 2016 (cf. p. 115).
- [Cho05] W. CHONGBUREE. « Analysis of power spectral density of digitally-modulated combined pulse trains ». In : *Annual Conf. on Elect. Eng./Electron., Comput., Telecommun. and Information Technology, ECTI*. 2005 (cf. p. 115).
- [Chr08] O. CHRISTENSEN. *Frames and bases : An introductory course*. Birkhauser, 2008. ISBN : 0817646779 (cf. p. 26, 28, 34, 116).

-
- [CRT06] E. J. CANDÉS, J. ROMBERG et T. TAO. « Robust uncertainty principles : exact signal reconstruction from highly incomplete frequency information ». In : *IEEE Transactions on Information Theory* 52.2 (fév. 2006), p. 489–509. ISSN : 0018-9448. DOI : 10.1109/TIT.2005.862083 (cf. p. 9).
 - [CT03] C. CHAROENLARNOPPARUT et S. TANTARATANA. « Multidimensional convolutional code : progresses and bottlenecks ». In : *Circuits and Systems, 2003. ISCAS '03. Proceedings of the 2003 International Symposium on*. T. 3. Mai 2003, DOI : 10.1109/ISCAS.2003.1205112 (cf. p. 56).
 - [CY02] K. CHO et D. YOON. « On the general BER expression of one- and two-dimensional amplitude modulations ». In : *IEEE Transactions on Communications* 50.7 (juil. 2002), p. 1074–1080. ISSN : 0090-6778. DOI : 10.1109/TCOMM.2002.800818 (cf. p. 81).
 - [Das+14] Deepak DASALUKUNTE et al. *Faster than Nyquist Signaling*. Springer International Publishing, 2014. DOI : 10.1007/978-3-319-07031-5 (cf. p. 46).
 - [Dau90] I. DAUBECHIES. « The wavelet transform, time-frequency localization and signal analysis ». In : *IEEE Transactions on Information Theory* 36.5 (sept. 1990), p. 961–1005. ISSN : 0018-9448. DOI : 10.1109/18.57199 (cf. p. 25, 26).
 - [Deu96] P. DEUTSCH. *DEFLATE Compressed Data Format Specification version 1.3*. RFC 1951 (Informational). RFC. Fremont, CA, USA : RFC Editor, mai 1996. DOI : 10.17487/RFC1951 (cf. p. 10).
 - [Don06] D. L. DONOHO. « Compressed sensing ». In : *IEEE Transactions on Information Theory* 52.4 (avr. 2006), p. 1289–1306. ISSN : 0018-9448. DOI : 10.1109/TIT.2006.871582 (cf. p. 9).
 - [Dou+95] Catherine DOUILLARD et al. « Iterative correction of intersymbol interference : turbo-equalization ». In : *European transactions on telecommunications* 6.5 (sept. 1995), p. 507–512 (cf. p. 51, 56, 57).
 - [DRO10] D. DASALUKUNTE, F. RUSEK et V. OWALL. « An Iterative Decoder for Multicarrier Faster-Than-Nyquist Signaling Systems ». In : *Communications (ICC), 2010 IEEE International Conference on*. Mai 2010, p. 1–5. DOI : 10.1109/ICC.2010.5502554 (cf. p. 46, 47, 62).
 - [ETS10] ETSI. *Evolved Universal Terrestrial Radio Access (E-UTRA); Physical Channels and Modulation*. 3GPP TS 36.211. Technical specification. Mar. 2010 (cf. p. 23).
 - [ETS14] ETSI. *Digital Video Broadcasting (DVB); Second generation framing structure, channel coding and modulation systems for Broadcasting, Interactive Services, News Gathering and other broadband satellite applications (DVB-S2)*. EN 302 307. Nov. 2014 (cf. p. 42).

Bibliographie

-
- [FAB95] B. Le FLOCH, M. ALARD et C. BERROU. « Coded orthogonal frequency division multiplex [TV broadcasting] ». In : *Proceedings of the IEEE* 83.6 (juin 1995), p. 982–996. ISSN : 0018-9219. DOI : 10.1109/5.387096 (cf. p. 24).
- [Fai88] M. FAILLI. « COST 207 digital land mobile radio communications ». In : *Commission of the European Communities* (1988), p. 137 (cf. p. 107).
- [For73] G. David FORNEY. « The viterbi algorithm ». In : *Proceedings of the IEEE* 61.3 (mar. 1973), p. 268–278. ISSN : 0018-9219. DOI : 10.1109/PROC.1973.9030 (cf. p. 45, 46, 50, 52, 56).
- [FV94] E. FORNASINI et M. E. VALCHER. « Algebraic aspects of two-dimensional convolutional codes ». In : *IEEE Transactions on Information Theory* 40.4 (juil. 1994), p. 1068–1082. ISSN : 0018-9448. DOI : 10.1109/18.335967 (cf. p. 56).
- [Gab45] D. GABOR. « Theory of communication. Part 1 : The analysis of information ». In : *Electrical Engineers - Part III : Radio and Communication Engineering, Journal of the Institution of* 93.26 (sept. 1945), p. 429–441. DOI : 10.1049/ji-3-2.1946.0074 (cf. p. 25).
- [Gaf34] Félix GAFFIOT. *Dictionnaire illustré latin-français*. [En ligne : <http://www.lexilogos.com/latin/gaffiot.php>]. Hachette, 1934 (cf. p. 7).
- [Gal63] R. G. GALLAGER. « Low-Density Parity-Check Codes ». Thèse de doct. Cambridge, 1963 (cf. p. 41).
- [GGB94] S. Le GOFF, A. GLAVIEUX et C. BERROU. « Turbo-codes and high spectral efficiency modulation ». In : *Communications, 1994. ICC '94, SUPERCOMM/ICC '94, Conference Record, 'Serving Humanity Through Communications.'* *IEEE International Conference on*. Mai 1994, 645–649 vol.2. DOI : 10.1109/ICC.1994.368804 (cf. p. 41).
- [GLL97] Alain GLAVIEUX, Christophe LAOT et Joël LABAT. « Turbo Equalization over a frequency selective channel ». In : *Proc. 1st Symp. Turbo Codes*. 1997, p. 96–102 (cf. p. 62).
- [GW81] R. D. GITLIN et S. B. WEINSTEIN. « Fractionally-spaced equalization : An improved digital transversal equalizer ». In : *The Bell System Technical Journal* 60.2 (fév. 1981), p. 275–296. ISSN : 0005-8580. DOI : 10.1002/j.1538-7305.1981.tb00240.x (cf. p. 62).
- [Hél+03] M. HÉLARD et al. « On the Performance of a Turbo-Equalizer including Blind Equalizer over Time and Frequency-selective Channel. Comparison with an OFDM System. » In : *Proc. Int. Symposium on Turbo Codes and Related Topics*. Sept. 2003, p. 419–422 (cf. p. 62).

-
- [HJZ14] F. M. HAN, M. JIN et H. X. ZOU. « Binary Symbol Recovery Via ℓ_∞ Minimization in Faster-Than-Nyquist Signaling Systems ». In : *IEEE Transactions on Signal Processing* 62.20 (oct. 2014), p. 5282–5293. ISSN : 1053-587X. DOI : 10.1109/TSP.2014.2347920 (cf. p. 56).
 - [Hom07] HOMEPLUG. *HomePlug AV Specification*. HomePlug AV Specification Version 1.1. Standard. Mai 2007 (cf. p. 23).
 - [Hom14] HOMEPLUG. *HomePlug AV Specification*. HomePlug AV Specification Version 2.1. Standard. Fév. 2014 (cf. p. 23).
 - [Huf52] D. A. HUFFMAN. « A Method for the Construction of Minimum-Redundancy Codes ». In : *Proceedings of the IRE* 40.9 (sept. 1952), p. 1098–1101. ISSN : 0096-8390. DOI : 10.1109/JRPROC.1952.273898 (cf. p. 10).
 - [HZ07] F. M. HAN et X. D. ZHANG. « Hexagonal Multicarrier Modulation : A Robust Transmission Scheme for Time-Frequency Dispersive Channels ». In : *IEEE Transactions on Signal Processing* 55.5 (mai 2007), p. 1955–1961. ISSN : 1053-587X. DOI : 10.1109/TSP.2006.890884 (cf. p. 24).
 - [IEE09] IEEE-SA. *IEEE 802.11n-2009 – Amendment 5 : Enhancements for Higher Throughput*. IEEE Std 802.11n-2009. Standard. New York, nov. 2009 (cf. p. 23).
 - [ISO93] ISO/IEC. *Information technology – Coding of moving pictures and associated audio for digital storage media at up to about 1,5 Mbit/s – Part 3 : Audio*. ISO/IEC 11172-3 :1993. Août 1993 (cf. p. 11).
 - [ITU04] ITU-T. *Émetteurs-récepteurs pour lignes d’abonné numérique à très grande vitesse*. ITU-T G.993.1. Recommendation. Genève, juin 2004 (cf. p. 23).
 - [ITU09] ITU-T. *Transmission characteristics for narrow-band digital loudspeaking and hands-free telephony terminals*. ITU-T P.342. Recommendation. Genève, juin 2009 (cf. p. 11).
 - [ITU16] ITU-T. *High efficiency video coding*. ITU-T H.265. Recommendation. Genève, déc. 2016 (cf. p. 11).
 - [ITU92] ITU-T. *Information technology - Digital compression and coding of continuous-tone still images - Requirements and guidelines*. ITU-T T.81. Recommendation. Genève, sept. 1992 (cf. p. 11).
 - [ITU99] ITU-T. *Émetteurs-récepteurs de ligne d’abonné numérique asymétrique*. ITU-T G.992.1. Recommendation. Genève, juin 1999 (cf. p. 23).
 - [JG07] M. JOINDOT et A. GLAVIEUX. *Introduction aux communications numériques : Cours et exercices corrigés*. Sciences sup. Dunod, 2007. ISBN : 9782100513345 (cf. p. 8, 12).
 - [KM97] W. KOZEK et A. F. MOLISCH. « On the eigenstructure of underspread WSSUS channels ». In : *Proc. First IEEE Signal Processing Workshop Signal Processing Advances in Wireless Communications*. 1997, p. 325–328. DOI : 10.1109/SPAWC.1997.630380 (cf. p. 16, 19, 22).

Bibliographie

-
- [KP75] P. KABAL et S. PASUPATHY. « Partial-Response Signaling ». In : *IEEE Transactions on Communications* 23.9 (sept. 1975), p. 921–934. ISSN : 0090-6778. DOI : 10.1109/TCOM.1975.1092918 (cf. p. 43).
- [Lah17] Naila LAHBABI. « Faster than Nyquist transceiver design : algorithms for transmission and reception and global enhancements ». Thèse de doct. IMT Atlantique Bretagne, 2017 (cf. p. 46, 62, 124).
- [LC04] Shu LIN et Daniel J. COSTELLO. *Error Control Coding, Second Edition*. Upper Saddle River, NJ, USA : Prentice-Hall, Inc., 2004. ISBN : 0130426725 (cf. p. 54).
- [Le 03] Raphaël LE BIDAN. « Turbo-Equalization for bandwidth-efficient digital communications over frequency-selective channels ». Thesis. INSA de Rennes, nov. 2003 (cf. p. 61).
- [LG03] A. D. LIVERIS et C. N. GEORGHIADES. « Exploiting faster-than-Nyquist signaling ». In : *IEEE Transactions on Communications* 51.9 (sept. 2003), p. 1502–1511. ISSN : 0090-6778. DOI : 10.1109/TCOMM.2003.816943 (cf. p. 2, 50).
- [LGL01] C. LAOT, A. GLAVIEUX et J. LABAT. « Turbo equalization : adaptive equalization and channel decoding jointly optimized ». In : *IEEE Journal on Selected Areas in Communications* 19.9 (sept. 2001), p. 1744–1752. ISSN : 0733-8716. DOI : 10.1109/49.947038 (cf. p. 59).
- [LIA07] S. L. LINFOOT, M. K. IBRAHIM et M. M. AL-AKAIDI. « Orthogonal Wavelet Division Multiplex : An Alternative to OFDM ». In : *IEEE Transactions on Consumer Electronics* 53.2 (mai 2007), p. 278–284. ISSN : 0098-3063. DOI : 10.1109/TCE.2007.381686 (cf. p. 24).
- [Lin+15] H. LIN et al. « An efficient FTN implementation of the OFDM/OQAM system ». In : *2015 IEEE International Conference on Communications (ICC)*. Juin 2015, p. 4787–4792. DOI : 10.1109/ICC.2015.7249080 (cf. p. 47).
- [LLL05] Christophe LAOT, R. LE BIDAN et D. LEROUX. « Low-complexity MMSE turbo equalization : a possible solution for EDGE ». In : *Wireless Communications, IEEE Transactions on* 4.3 (mai 2005), p. 965–974. ISSN : 1536-1276. DOI : 10.1109/TWC.2005.847095 (cf. p. 92–94).
- [Mat+98] K. MATHEUS et al. « Two-dimensional (recursive) channel equalization for multicarrier systems with soft impulse shaping (MCSIS) ». In : *Global Telecommunications Conference, 1998. GLOBECOM 1998. The Bridge to Global Integration. IEEE. T. 2. 1998*, 956–961 vol.2. DOI : 10.1109/GLOCOM.1998.776871 (cf. p. 55).
- [Maz75] James E. MAZO. « Faster Than Nyquist Signaling ». In : *Bell System Technical Journal* 54 (oct. 1975), p. 1451–1462 (cf. p. 2, 45, 50).
- [ML88] J. E. MAZO et H. J. LANDAU. « On the minimum distance problem for faster-than-Nyquist signaling ». In : *IEEE Transactions on Information Theory* 34.6 (nov. 1988), p. 1420–1427. ISSN : 0018-9448. DOI : 10.1109/18.21281 (cf. p. 50).

-
- [MMN92] M. S. MOHAMMED, S. L. MATILIA et L. NOZAL. « Fast 2D convolution filter based on look up table FFT ». In : *[1992] Proceedings of the IEEE International Symposium on Industrial Electronics*. Mai 1992, 446–449 vol.1. DOI : 10.1109/ISIE.1992.279637 (cf. p. 104).
- [MN95] David J. C. MACKAY et Radford M. NEAL. « Good codes based on very sparse matrices ». In : sous la dir. de Colin BOYD. Berlin, Heidelberg : Springer Berlin Heidelberg, 1995, p. 100–111. ISBN : 978-3-540-49280-1. DOI : 10.1007/3-540-60693-9_13 (cf. p. 41).
- [Nyq28a] H. NYQUIST. « Certain Topics in Telegraph Transmission Theory ». In : *American Institute of Electrical Engineers, Transactions of the* 47.2 (1928), p. 617–644. ISSN : 0096-3860. DOI : 10.1109/T-AIEE.1928.5055024 (cf. p. 2, 9, 45).
- [Nyq28b] H. NYQUIST. « Thermal Agitation of Electric Charge in Conductors ». In : *Phys. Rev.* 32 (1 juil. 1928), p. 110–113. DOI : 10.1103/PhysRev.32.110 (cf. p. 14).
- [Otn04] R. OTNES. « Fractionally spaced linear MMSE turbo equalization ». In : *2004 12th European Signal Processing Conference*. Sept. 2004, p. 465–468 (cf. p. 62).
- [PS08] J.G. PROAKIS et M. SALEHI. *Digital Communications*. McGraw-Hill International Edition. McGraw-Hill, 2008. ISBN : 9780071263788 (cf. p. 8, 12, 37, 38, 53, 55, 104).
- [PS11] D. PINCHON et P. SIOHAN. « Closed-Form Expressions of Optimal Short PR FMT Prototype Filters ». In : *2011 IEEE Global Telecommunications Conference - GLOBECOM 2011*. Déc. 2011, p. 1–5. DOI : 10.1109/GLOCOM.2011.6133592 (cf. p. 23, 31, 32).
- [PS13] D. PINCHON et P. SIOHAN. « Derivation of analytical expressions for flexible PR low complexity FBMC systems ». In : *21st European Signal Processing Conference (EUSIPCO 2013)*. Sept. 2013, p. 1–5 (cf. p. 23, 31, 32).
- [RA05] F. RUSEK et J.B. ANDERSON. « The two dimensional Mazo limit ». In : *Information Theory, 2005. ISIT 2005. Proceedings. International Symposium on*. Sept. 2005, p. 970–974. DOI : 10.1109/ISIT.2005.1523482 (cf. p. 46, 50).
- [RA08] F. RUSEK et J. B. ANDERSON. « Non Binary and Precoded Faster Than Nyquist Signaling ». In : *IEEE Transactions on Communications* 56.5 (mai 2008), p. 808–817. ISSN : 0090-6778. DOI : 10.1109/TCOMM.2008.060075 (cf. p. 2, 46, 50).
- [RCS15] F. RUSEK, G. COLAVOLPE et C. E. W. SUNDBERG. « 40 Years with the Ungerboeck Model : A Look at its Potentialities [Lecture Notes] ». In : *IEEE Signal Processing Magazine* 32.3 (mai 2015), p. 156–161. ISSN : 1053-5888. DOI : 10.1109/MSP.2014.2374221 (cf. p. 52).
- [RE13] F. RUSEK et O. EDFORS. « An information theoretic characterization of channel shortening receivers ». In : *2013 Asilomar Conference on Signals, Systems and Computers*. Nov. 2013, p. 2108–2112. DOI : 10.1109/ACSSC.2013.6810679 (cf. p. 55, 60).

Bibliographie

- [Roq12] Damien ROQUE. « Modulations multiporteuses WCP-OFDM : évaluation des performances en environnement radiomobile ». Thèse de doct. Université de Grenoble, 2012 (cf. p. 12, 23, 32).
- [RVH95] P. ROBERTSON, E. VILLEBRUN et P. HOEHER. « A comparison of optimal and sub-optimal MAP decoding algorithms operating in the log domain ». In : *Communications, 1995. ICC '95 Seattle, 'Gateway to Globalization', 1995 IEEE International Conference on*. T. 2. Juin 1995, 1009–1013 vol.2. DOI : 10.1109/ICC.1995.524253 (cf. p. 61).
- [RW95] P. ROBERTSON et T. WORZ. « Coded modulation scheme employing turbo codes ». In : *Electronics Letters* 31.18 (août 1995), p. 1546–1547. ISSN : 0013-5194. DOI : 10.1049/el:19951064 (cf. p. 41, 54).
- [Sha48] C. E. SHANNON. « A mathematical theory of communication ». In : *The Bell System Technical Journal* 27.3 (juil. 1948), p. 379–423. ISSN : 0005-8580. DOI : 10.1002/j.1538-7305.1948.tb01338.x (cf. p. 11, 37, 40).
- [Sha49] C. E. SHANNON. « Communication in the Presence of Noise ». In : *Proceedings of the IRE* 37.1 (jan. 1949), p. 10–21. ISSN : 0096-8390. DOI : 10.1109/JRPROC.1949.232969 (cf. p. 9, 37).
- [SHN16] H. SASAHARA, K. HAYASHI et M. NAGAHARA. « Symbol Detection for Faster-Than-Nyquist Signaling by Sum-of-Absolute-Values Optimization ». In : *IEEE Signal Processing Letters* 23.12 (déc. 2016), p. 1853–1857. ISSN : 1070-9908. DOI : 10.1109/LSP.2016.2625839 (cf. p. 56).
- [Sic02] C. SICLET. « Application de la théorie des bancs de filtres à l'analyse et à la conception de modulations multiporteuses orthogonales et biorthogonales ». Thèse de doct. Université de Rennes 1, nov. 2002 (cf. p. 35, 36).
- [SRF08] N. SELLAMI, A. ROUMY et I. FIJALKOW. « A Proof of Convergence of the MAP Turbo-Detector to the AWGN Case ». In : *IEEE Transactions on Signal Processing* 56.4 (avr. 2008), p. 1548–1561. ISSN : 1053-587X. DOI : 10.1109/TSP.2007.910467 (cf. p. 57).
- [SSJ06] A. SKRZYPCZAK, P. SIOHAN et J. P. JAVAUDIN. « Analysis of the Peak-To-Average Power Ratio of the Oversampled OFDM ». In : *Int. Conf. on Acoustics Speech and Signal Process. (ICASSP)*. T. 4. IEEE, mai 2006, p. IV–IV. DOI : 10.1109/ICASSP.2006.1660967 (cf. p. 114).
- [Tou05] J. TOULOUSE. « Optical nonlinearities in fibers : review, recent examples, and systems applications ». In : *Journal of Lightwave Technology* 23.11 (nov. 2005), p. 3625–3641. ISSN : 0733-8724. DOI : 10.1109/JLT.2005.855877 (cf. p. 14).
- [TPB16] R. TAJAN, C. POUILLIAT et M. L. BOUCHERET. « Circular Faster Than Nyquist : Transmitter and iterative receiver design ». In : *2016 9th International Symposium on Turbo Codes and Iterative Information Processing (ISTC)*. Sept. 2016, p. 241–245. DOI : 10.1109/ISTC.2016.7593113 (cf. p. 46).

-
- [TSK02] M. TÜCHLER, A.C. SINGER et R. KOETTER. « Minimum mean squared error equalization using a priori information ». In : *Signal Processing, IEEE Transactions on* 50.3 (mar. 2002), p. 673–683. ISSN : 1053-587X. DOI : 10.1109/78.984761 (cf. p. 60, 61).
 - [Ung74] G. UNGERBOECK. « Adaptive Maximum-Likelihood Receiver for Carrier-Modulated Data-Transmission Systems ». In : *IEEE Transactions on Communications* 22.5 (mai 1974), p. 624–636. ISSN : 0090-6778. DOI : 10.1109/TCOM.1974.1092267 (cf. p. 52).
 - [Ung82] G. UNGERBOECK. « Channel coding with multilevel/phase signals ». In : *IEEE Transactions on Information Theory* 28.1 (jan. 1982), p. 55–67. ISSN : 0018-9448. DOI : 10.1109/TIT.1982.1056454 (cf. p. 41).
 - [Vit+89] A. J. VITERBI et al. « A pragmatic approach to trellis-coded modulation ». In : *IEEE Communications Magazine* 27.7 (juil. 1989), p. 11–19. ISSN : 0163-6804. DOI : 10.1109/35.31452 (cf. p. 41).
 - [Vit67] A. VITERBI. « Error bounds for convolutional codes and an asymptotically optimum decoding algorithm ». In : *IEEE Transactions on Information Theory* 13.2 (avr. 1967), p. 260–269. ISSN : 0018-9448. DOI : 10.1109/TIT.1967.1054010 (cf. p. 41, 46, 52).
 - [WGK10] S. WEI, D. L. GOECKEL et P. A. KELLY. « Convergence of the Complex Envelope of Bandlimited OFDM Signals ». In : *IEEE Transactions on Information Theory* 56.10 (oct. 2010), p. 4893–4904. ISSN : 0018-9448. DOI : 10.1109/TIT.2010.2059550 (cf. p. 115).
 - [Wos+98] M. WOSNITZA et al. « A high precision 1024-point FFT processor for 2D convolution ». In : *1998 IEEE International Solid-State Circuits Conference. Digest of Technical Papers, ISSCC. First Edition (Cat. No.98CH36156)*. Fév. 1998, p. 118–119. DOI : 10.1109/ISSCC.1998.672398 (cf. p. 104).
 - [Yea11] Chen-Pang YEANG. « Tubes, randomness, and Brownian motions : or, how engineers learned to start worrying about electronic noise ». In : *Archive for History of Exact Sciences* 65.4 (2011), p. 437–470. DOI : 10.1007/s00407-011-0082-5 (cf. p. 14).
 - [ZL77] J. ZIV et A. LEMPEL. « A universal algorithm for sequential data compression ». In : *IEEE Transactions on Information Theory* 23.3 (mai 1977), p. 337–343. ISSN : 0018-9448. DOI : 10.1109/TIT.1977.1055714 (cf. p. 10).

Résumé — Avec la multiplication des terminaux mobiles et le foisonnement des objets dits « connectés », on assiste à la montée d'un besoin de moyens de communication à tout endroit et en toute situation, accompagné d'un encombrement spectral toujours plus important. Dans ce contexte, si la capacité d'adaptation au canal des modulations multiporteuses permet de bien s'accommoder du besoin de communication en tout endroit, les techniques actuelles, en particulier l'OFDM, souffrent d'une mauvaise localisation fréquentielle et d'un mauvais facteur de crête, ce qui limite leur utilisation dans un contexte embarqué et/ou en présence de fortes contraintes spectrales.

Dans cette thèse, nous étudions les modulations multiporteuses au-delà de la cadence de Nyquist. En augmentant la densité de signalisation, ces dernières permettent d'augmenter l'efficacité spectrale. Cela est cependant contrebalancé par l'apparition d'auto-interférence, ce qui rend la réception plus délicate. Sur canal à bruit additif blanc gaussien, on montre comment choisir des impulsions de mise en forme maximisant le rapport signal à interférence plus bruit. On montre que ces dernières permettent d'obtenir une turbo-égalisation linéaire de l'auto-interférence minimisant l'erreur quadratique moyenne. Nos travaux mettent en évidence que ces mêmes impulsions permettent également de réduire le facteur de crête à mesure que la densité augmente. Enfin, sur canal sélectif en fréquence, on vérifie que l'approximation du canal par un coefficient par sous-porteuse est toujours possible. Ces résultats montrent que ce nouveau type de modulation permet d'augmenter l'efficacité spectrale tout en conservant la capacité d'adaptation au canal intrinsèque aux modulations multiporteuses.

Mots clés : Modulations au-delà de la cadence de Nyquist, modulations multiporteuses, canal radiomobile, efficacité spectrale, faible complexité, facteur de crête, turbo-égalisation, systèmes de Gabor, théorie des frames.

Abstract — With an increasing number of mobile terminals coupled with a large spreading of so-called “smart devices”, we can see a growing demand for effective communication means in any place and in any situation. This goes with a more and more overcrowded spectrum. In this context, multicarrier modulations are good candidates to allow effective communication in any place. However current techniques, OFDM in particular, suffer from a bad time-frequency localization and peak-to-average power ratio, limiting their relevancy in an embedded context, or in scenarios with severe spectral constraints.

In this thesis, we study faster-than-Nyquist multicarrier modulations. This kind of modulation allow for an increase in spectral efficiency by means of an increase in signaling density. This, in compensation, comes at the price of unavoidable self-interference, which makes demodulation harder. On an additive white Gaussian noise channel, we show how to carefully chose pulse-shapes that maximize signal-to-interference-plus-noise ratio. We show that these particular pulse-shapes yields a linear turbo-equalization of self-interference minimizing the mean squared error. Next, our work highlights the capability of these optimal pulse-shapes to reduce peak-to-average power ratio as density rises. Lastly, on frequency selective channels, we confirm that low complexity equalization using one tap by subcarrier is still possible. These results show how this new modulation technique helps increasing spectral efficiency while keeping what made multicarrier modulations popular : good adaptation to transmission channels.

Keywords : Faster-than-Nyquist modulations, multicarrier modulations, radiomobile channel, spectral efficiency, low-complexity, peak-to-average power ratio, turbo-equalization, Gabor systems, frame theory.
