



Stratégies d'allocation de ressources dans des contextes mono et multiutilisateurs pour des communications à très haut débit sur lignes d'énergie

Ali Maiga

► To cite this version:

Ali Maiga. Stratégies d'allocation de ressources dans des contextes mono et multiutilisateurs pour des communications à très haut débit sur lignes d'énergie. Traitement du signal et de l'image [eess.SP]. INSA de Rennes, 2010. Français. NNT : . tel-00605144

HAL Id: tel-00605144

<https://theses.hal.science/tel-00605144v1>

Submitted on 30 Jun 2011

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.



THESE INSA Rennes
sous le sceau de l'Université européenne de Bretagne
pour obtenir le titre de
DOCTEUR DE L'INSA DE RENNES
Spécialité : Electronique

présentée par
Ali MAIGA
ECOLE DOCTORALE
MATISSE

**Stratégies d'allocation de
ressources dans des
contextes mono et multi-
utilisateurs pour des
communications à très haut
débit sur lignes d'énergie**

Thèse soutenue le 13.12.2010
devant le jury composé de :

Luc Vandendorpe
Professeur à l'Université Catholique de Louvain / *président*
Gilles Burel
Professeur à l'Université de Bretagne Occidentale / *rapporteur*
Pierre Siohan
Docteur-Ingénieur à Orange Labs / *rapporteur*
Ahmed Zeddami
Docteur-Ingénieur à Orange Labs / *examinateur*
Inbar Fijalkow
Professeur à l'ENSEA / *examinateur*
Jean-Yves Baudais
Chargé de Recherches CNRS-IETR / *Co-encadrant de thèse*
Jean-François Héliard
Professeur à l'INSA de Rennes / *Directeur de thèse*

Stratégies d'allocation de ressources dans des contextes mono et multi-utilisateurs pour des communications à très haut débit sur lignes d'énergie

Ali MAIGA



En partenariat avec

		The logo for OMEGA features a stylized orange Greek letter Ω (Omega) with a small square inside it. Below the symbol, the word 'OMEGA' is written in a grey, sans-serif font.		
--	--	--	--	--

à Filsan

*Si tu veux réussir dans la vie d'ici-bas, cherches le savoir,
si tu veux réussir dans la vie de l'au-delà, cherches le savoir,
si tu veux réussir dans les deux vies, cherches le savoir.*

Le prophète Mouhammad

Remerciements

En premier lieu, je tiens à exprimer ma profonde et très sincère reconnaissance à Jean-François Héland, professeur à l'INSA de Rennes, pour m'avoir proposé cette thèse et en avoir dirigé les travaux. J'exprime également toute ma gratitude à Jean-Yves Baudais, chargé de recherche CNRS à l'IETR, pour avoir co-encadré l'ensemble de ces travaux de recherche. Je les remercie tous les deux pour la confiance qu'ils m'ont témoignée pendant ces trois années, et pour m'avoir fait bénéficier de leurs compétences et de leurs conseils. Mais au delà des aspects techniques, je tiens à souligner leurs qualités humaines qui ont fait de cet encadrement en complémentarité une expérience très positive.

J'adresse tout naturellement mes remerciements à l'ensemble des membres du jury, sans qui mes travaux de recherche n'auraient pu donné lieu à cette thèse. C'est ainsi que je remercie Gilles Burel Professeur à l'Université de Bretagne Occidentale, et Pierre Siohan, Docteur-ingénieur à Orange Labs, pour avoir accepté de participer au jury en tant que rapporteurs et pour l'attention qu'ils ont accordée à la lecture de ce mémoire. Je remercie également Luc Vandendorpe, Professeur à l'Université Catholique de Louvain, président du jury, Ahmed Zeddami, Docteur-Ingénieur à Orange Labs, et Inbar Fijalkow, Professeur à l'ENSEA, pour avoir pris de leur temps et avoir participé au jury en tant qu'examinateurs.

J'exprime ma très grande gratitude à ma famille et en particulier à mes parents pour m'avoir soutenu, depuis l'Afrique, tout au long de mes études. Je sais les sacrifices que ces longues années ont représentés et les remercient d'avoir appuyé mes choix et d'avoir toujours su m'encourager. Je remercie du fond du cœur mon épouse pour son soutien et sa patience pendant ces trois années. Cette thèse est certes une récompense personnelle, mais aussi bel et bien le fruit d'un effort vécu à deux, dans la complexité de l'adéquation entre réussite personnelle et professionnelle. Je ne la remercierai jamais assez pour avoir su me comprendre et me soutenir au quotidien et particulièrement pendant le marathon final de la rédaction. Merci également à ma belle mère et à ma petite Maryam pour n'avoir jamais tenu rigueur de mes absences répétées et avoir su m'apporter la joie et la gaieté dont elle a le secret.

Par ailleurs, je tiens à remercier Najmeddine Kout pour le sérieux de son travail effectué durant son stage de fin d'étude de Master et d'ingénieur. Ses résultats ont été très utiles à mon travail et son encadrement s'est révélé une expérience particulièrement enrichissante. Pour leur bonne humeur au quotidien, j'adresse un grand merci à l'ensemble des permanents, doctorants et stagiaires que j'ai côtoyés durant ses trois années, et plus particulièrement, à Fahad Muhammad, Ayman Khalil, Hassan, Abdallah Hamini, anciens comme nouveaux, qui ont contribué à la bonne ambiance des journées au labo. Toutes mes excuses à ceux que j'aurais oubliés ...

Table des matières

Remerciements	iii
Table des figures	ix
Liste des tableaux	xi
Introduction	1
1 Le réseau électrique pour la transmission de données	7
1.1 Un peu d'histoire	7
1.2 La technologie CPL	8
1.3 Organismes de réglementations et consortiums	10
1.3.1 Compatibilité électromagnétique	11
1.3.2 Organismes de normalisations	11
1.3.2.1 CEI / CISPR	11
1.3.2.2 CENELEC	11
1.3.2.3 FCC partie 15	12
1.3.3 Les consortiums	13
1.3.3.1 Alliance Homeplug	14
1.3.3.2 UPA	14
1.3.3.3 CEPCA	14
1.4 Projet OMEGA	14
1.4.1 Objectifs	15
1.4.2 CPL OMEGA	16
1.5 Caractéristiques du canal de propagation	17
1.5.1 Réponse du canal	18
1.5.1.1 Modèle de canaux CPL	19
1.5.2 Sources de bruit	22
1.5.2.1 Bruit stationnaire	24
1.5.2.2 Bruit de fond	25
1.5.2.3 Bruits à bande étroite	26
1.5.2.4 Bruits impulsifs	26
1.6 Conclusion	27

2	Spécifications du système et l'allocation des ressources	29
2.1	Rappels sur les techniques de modulations multiporteuses et d'étalement de spectre	29
2.2	Techniques d'accès multiple	31
2.3	Système LP-OFDM	32
2.3.1	Expression des signaux LP-OFDM	33
2.3.2	Choix de la technique d'égalisation	34
2.3.2.1	Distorsion crête	35
2.3.2.2	Erreur quadratique moyenne	35
2.4	Allocation des ressources	36
2.4.1	Modulations adaptatives	37
2.4.1.1	Capacité et débit associés à un canal non dispersif	37
2.4.1.2	Marge de RSB des modulations MAQ	37
2.4.2	Politiques d'optimisation : problème général	39
2.5	Conclusion	41
3	Maximisation du débit dans un contexte mono-utilisateur	43
3.1	Introduction	43
3.2	Modulations adaptatives et OFDM : état de l'art	44
3.3	Modulations adaptatives en LP-OFDM	47
3.3.1	Système et degrés de liberté	48
3.3.2	Information mutuelle	50
3.3.2.1	Choix de la technique d'égalisation	51
3.3.3	Optimisation du débit	52
3.3.3.1	Maximisation du débit suivant le critère ZF	53
3.3.3.2	Maximisation du débit suivant le critère EQM	55
3.3.3.3	Extension des résultats au cas multibloc	57
3.3.4	Simulation et performances	61
3.3.4.1	Prise en compte des fonctions de codage de canal	61
3.3.4.2	Débits atteignables sur les canaux OMEGA	62
3.3.4.3	Exploitation de la DSP	64
3.3.4.4	Influence de la longueur des séquences de précodage	64
3.3.5	Conclusion	66
3.4	Maximisation du débit sous contrainte de TEB	66
3.4.1	Maximisation du débit sous la contrainte de TEB crête	67
3.4.2	Maximisation du débit sous la contrainte de TEB moyen	69
3.4.3	Simulation et performances	72
3.5	Conclusion	74
4	Allocation des ressources dans un contexte multicast	75
4.1	Introduction	75
4.1.1	Description du système	77
4.2	Allocation des ressources en multicast monodébit	78
4.2.1	Notion de canal équivalent	79
4.2.2	Application de la méthode LCG	79
4.2.2.1	Limitation du débit avec la méthode LCG	79

4.2.2.2	Amélioration du débit multicast offert par la méthode LCG	81
4.2.3	Allocation des ressources en LP-OFDM multicast	82
4.2.3.1	Solution combinatoire	83
4.2.3.2	Solution LBCG (<i>low block channel gain</i>)	84
4.2.4	Comportement asymptotique du débit multicast total	85
4.2.5	Application de l'algorithme du <i>bit-loading</i> avec le critère EQM	87
4.3	Allocation des ressources en multicast multidébit	88
4.3.1	Description du système multicast multidébit	89
4.3.1.1	Réalisations possibles de l'encodage à débits variables	90
4.3.1.2	Mesure de l'équité	91
4.3.2	Séparation des destinataires dans le domaine fréquentiel	93
4.3.2.1	Formulation du problème	94
4.3.2.2	Solution au problème d'optimisation	94
4.3.2.3	Prise en compte de l'équité entre les utilisateurs	95
4.3.3	Séparation des destinataires dans le domaine temporel	97
4.4	Simulations et performances	100
4.4.1	Étude d'un système multicast à 9 utilisateurs	101
4.4.1.1	Influence de la longueur des séquences de précodage	103
4.4.1.2	Étude de l'équité entre les utilisateurs	105
4.4.2	Évolution du débit avec le nombre d'utilisateurs	105
4.4.3	Discussion sur l' <i>overhead</i> des solutions proposées	107
4.5	Conclusion	108
5	Allocation des ressources dans un contexte d'accès	109
5.1	Allocation centralisée des ressources	110
5.1.1	Problèmes d'allocation des ressources en OFDMA : état de l'art	111
5.1.1.1	Critères d'optimisation	112
5.1.1.2	Maximisation du débit sous contrainte en OFDMA	112
5.1.1.3	Maximisation du débit total sous contraintes de débits minimums	114
5.1.1.4	Détermination des sous-ensembles de sous-canaux	114
5.1.2	Allocation des ressources en LP-OFDMA	117
5.1.3	Simulation et performances	118
5.1.4	Conclusion	123
5.2	Allocation décentralisée des ressources	123
5.2.1	Quelques éléments de la théorie des jeux	124
5.2.1.1	Définition d'un jeu en forme stratégique	125
5.2.1.2	Équilibre de Nash	126
5.2.1.3	Existence et unicité de l'équilibre de Nash	127
5.2.2	Minimisation de la puissance sous contraintes de débit et de DSP	128
5.2.2.1	Existence et unicité de l'équilibre de Nash	131
5.2.2.2	Mise en œuvre de l'allocation décentralisée des ressources	131
5.2.2.3	Simulation	132
5.2.3	Conclusion	134
5.3	Conclusion du chapitre	135

Conclusion générale	137
A Preuve de l'existence et de l'unicité de l'équilibre de Nash	141
Bibliographie	143

Table des figures

1.1	Principe de couplage du modem CPL sur le réseau électrique.	9
1.2	Structure de la boucle locale électrique.	9
1.3	Réseau domestique à très haut débit – OMEGA	15
1.4	Proposition de masque DSP pour le système CPL OMEGA.	18
1.5	Fonctions de répartition des capacités des canaux issus des différentes classes pour $\Delta_f = 24,414$ kHz et $N_0 = -140$ dBm/Hz [1].	21
1.6	Réponses du canal électrique pour les trois classes sélectionnées : classe 2, classe 5 et classe 9. À gauche, les fonctions de transfert, à droite les réponses impulsionnelles.	23
1.7	Ensemble des types de bruits additifs rencontrés sur les lignes électriques	24
1.8	Modèle de bruit stationnaire sur ligne d'énergie	25
2.1	Représentation schématique du système LP-OFDM dans un contexte multi-utilisateur.	33
2.2	Points de fonctionnement des modulations MAQ non codés pour $P_{eb} = P_{es} = 10^{-3}$	39
3.1	Comparaison des résultats d'allocation des bits en granularité infinie et finie, dans le cadre de la maximisation du débit en OFDM.	46
3.2	Fonction de transfert du modèle de référence du canal CPL proposé par Zimmermann.	47
3.3	Représentation schématique du paramétrage du signal LP-OFDM par les algorithmes d'allocation dynamique des ressources.	48
3.4	Principe d'allocation adaptative des ressources dans un système LP-OFDM mono-utilisateur.	49
3.5	Algorithme de maximisation du débit pour le système monobloc.	56
3.6	Principe de l'algorithme d'allocation avec le critère EQM.	58
3.7	Comparaison des résultats d'allocation des bits en granularité infinie et finie, dans le cadre de la maximisation du débit pour le système LP-OFDM multibloc suivant le critère ZF (a) et le critère EQM (b). La longueur de précodage choisie est $L = 32$	60
3.8	Exemple de RSB par sous-porteuse pour un canal de la classe 5.	61
3.9	Fonctions de répartition des débits atteignables sur des canaux issus des 9 classes de canaux pour les systèmes OFDM et LP-OFDM avec $L = 32$	63
3.10	Évolution de la puissance utilisée en fonction de la longueur des séquences de précodage sur un canal de classe 5.	65

3.11	Évolution du débit atteignable en fonction de la longueur des séquences de précodage pour les systèmes LP-OFDM-ZF et LP-OFDM-EQM sur un canal de classe 5.	65
3.12	Fonctions de répartition des débits réalisables avec les différents algorithmes de <i>bit-loading</i> sur des canaux de classe 2 (a), de classe 5 (b) et de classe 9 (c).	73
4.1	Types de communications	75
4.2	Exemple de communication multicast dans un réseau CPL.	78
4.3	Comparison des résultats analytiques et de simulation pour le débit multicast total dans $\mathbb{R}$, en bit par sous-porteuse, avec la méthode LCG ($E = 1$, $N_0 = 1$, and $\sigma^2 = 1$).	81
4.4	Moyenne des débits multicast totaux en bit par sous-porteuse en fonction du nombre d'utilisateurs avec $\frac{E}{\Gamma N_0} = 20$ dB et $\sigma = 1$	87
4.5	Module inter-couche PHY-MAC du côté de l'émetteur.	89
4.6	Types d'encodage des données sources	91
4.7	(a) Fonction de répartition de la capacité du système MDHT, (b) indice de la répartition optimale des utilisateurs sous l'hypothèse (4.52).	99
4.8	Fonction de répartition du débit total en multicast multidébit, (a) dans le domaine temporel et (b) dans le domaine fréquentiel pour $L = 32$	101
4.9	Fonction de répartition du débit minimum en multicast multidébit pour $L = 32$	103
4.10	Débit total en kbit par symbole OFDM versus la longueur des séquences de précodage.	104
4.11	Fonction de répartition de l'indice d'équité entre les utilisateurs, (a) dans le domaine temporel et (b) dans le domaine fréquentiel pour $L = 32$	106
4.12	Évolution des débits (débit (a) total et (b) minimum) en fonction du nombre d'utilisateurs pour $L = 32$	107
5.1	Un exemple de communications centralisées dans un réseau CPL.	111
5.2	Schéma global des procédures d'allocation des ressources.	117
5.3	Fonctions de répartition des débits totaux réalisables avec les différents algorithmes pour 9 utilisateurs.	120
5.4	Évolution du débit total (a) et du nombre d'utilisateurs insatisfaits en fonction du nombre d'utilisateurs du système.	122
5.5	Système distribué avec 2 utilisateurs.	125
5.6	Taux d'existence d'un équilibre de Nash versus le rapport entre la capacité du canal des utilisateurs et le débit requis.	133
5.7	Évolution de la somme des puissances allouées aux différentes sous-porteuses avec $R_{\min}^1 = 8,1$ kb/symbole OFDM et $R_{\min}^2 = 6,3$ kb/symbole OFDM.	134
5.8	Évolution de la répartition de la puissance par sous-porteuse avec $R_{\min}^1 = 8,1$ kb/symbole OFDM et $R_{\min}^2 = 6,9$ kb/symbole OFDM.	134

Liste des tableaux

1.1	Limites pour les perturbations conduites dans les ports CPL (fonction de télécommunication inactive) [2].	12
1.2	Limites des perturbations rayonnées à une distance de 10 mètres [2].	12
1.3	comparaison des principaux paramètres des spécifications de produits CPL existants.	13
1.4	Principaux paramètres des systèmes HPAV et OMEGA.	17
1.5	Répartition des fonctions de transfert par site.	19
1.6	Pourcentages d'apparition des classes, sites correspondants de mesure et atténuations moyennes.	22
3.1	Débits transmis, et gains absolus et relatifs à la valeur 0,5 de la fonction de répartition.	63
3.2	Temps de calculs moyen (maximal) en seconde pour différentes valeurs de RSB moyen sous MATLAB avec Intel Core2@2.66GHz.	71
4.1	Nombre total de définitions possibles de D pour $N = 128$	84
4.2	Liste exhaustive des répartitions possibles des différents utilisateurs en 2 sous-groupes.	98
4.3	Liste exhaustive des répartitions possibles des différents utilisateurs en 3 sous-groupes.	99
4.4	Allocation des intervalles de temps (IT) aux différents sous-groupes sur 12 intervalles de temps pour les mode 2 et 3.	100
4.5	Principaux paramètres de simulation.	101
5.1	Stratégies de maximisation du débit.	113
5.2	Pourcentage de cas testés où respectivement 0, 1, 2, 3 et 4 utilisateurs sur 9 sont insatisfaits.	121
5.3	La matrice des gains du dilemme du prisonnier	127

Introduction

Développées à l'origine dans le cadre d'applications bas débit de télémétrie, de contrôle d'infrastructure ou encore de mesure de consommation, les communications par courant porteur en ligne (CPL) connaissent depuis peu un regain d'intérêt manifeste au sein de la communauté scientifique. Les avancées importantes réalisées ces dernières années sur les techniques de modulation et de traitement du signal permettent en effet d'envisager aujourd'hui l'utilisation du réseau des lignes d'énergie pour le développement de réseaux domestiques et l'acheminement de données multimédia à haut débit. La demande de services sur le réseau CPL est en forte croissance du fait, entre autres, de sa potentialité améliorée, des hauts débits supportés, de la facilité de connexion entre les équipements et de leur faible coût de déploiement en comparaison avec les technologies sans fils comme le WIFI. Ce type de réseau supporte plusieurs trafics non homogènes et satisfait plusieurs contraintes de qualité de service. Dans ce contexte, une gestion efficace des ressources radios, par l'intermédiaire de stratégies d'allocations, s'avère indispensable pour mieux exploiter l'ensemble des ressources disponibles sur l'interface radio des réseaux CPL. Ces stratégies consistent à définir des règles de partage des ressources dans le but d'optimiser les débits d'utilisation et de satisfaire les multiples contraintes de qualité de service.

Le canal de propagation CPL est cependant peu favorable à la transmission de données à haut débit puisqu'il n'a, à l'origine, pas été conçu dans ce but. Des campagnes de mesures ont permis de caractériser les principales composantes du canal de transmission CPL, à savoir la réponse du canal qui rend compte des phénomènes qui viennent modifier la forme des ondes émises, ainsi que les brouilleurs qui viennent s'ajouter au signal reçu et dont l'origine peut être multiple [1]. Afin d'exploiter ce canal difficile, des techniques de transmission efficaces et robustes doivent alors être utilisées, parmi lesquelles on trouve naturellement les modulations à porteuses multiples. La modulation à porteuses multiples réduit la sensibilité du système à la sélectivité du canal de transmission, en divisant le canal en une collection de sous-canaux non sélectifs en fréquences et à occupation spectrale minimale. Pour améliorer les performances du système, nous proposons pour la transmission des données d'utiliser une forme d'onde combinant la

technique de précodage linéaire aux modulations à porteuses multiples de type OFDM et conduisant à la solution LP-OFDM (*linear precoded OFDM*). Sous l'hypothèse d'une connaissance parfaite de la réponse du canal, cette combinaison permet une exploitation plus efficace de la puissance disponible [3]. Les débits de transmission sont alors augmentés en adaptant les ordres de modulation, les niveaux de puissance et la répartition des ressources temps-fréquences aux conditions des liens de transmission.

Les travaux de cette thèse visent à étudier et à optimiser les stratégies d'allocation des ressources temps-fréquences, pour un ou plusieurs utilisateurs, en vue de la maximisation des débits. Les ressources mises en jeu sont les intervalles de temps de transmission, les sous-canaux du système à porteuses multiples ou les séquences de précodage du système LP-OFDM, ainsi que les bits et puissances attribués à ces sous-canaux, ou séquences. L'étude de méthodes d'allocation des ressources, principalement dans un contexte mono-utilisateur en CPL, a fait l'objet de précédents travaux de thèses [3,4]. Ce présent travail propose de nouveaux algorithmes d'allocation des ressources à la fois pour le contexte mono et multi-utilisateur. L'étude a été menée au sein du groupe « Communications-Propagation-Radar » de l'Institut d'électronique et de télécommunications de Rennes (IETR) et s'inscrit dans le cadre du projet européen OMEGA (*home gigabit access*). Ce projet, regroupant 24 partenaires européens industriels, opérateurs de télécommunications et universitaires, vise à développer un réseau domestique à très haut débit. Ce réseau sera capable de fournir des services très hauts débits à la vitesse du gigabit par seconde à travers des technologies de communications hétérogènes, y compris les technologies filaires et sans fils [5].

Les études menées dans cette thèse peuvent être regroupées en trois grandes parties. La première partie est consacrée au problème de maximisation du débit dans un contexte mono-utilisateur et sert de base pour le contexte multi-utilisateur. L'objectif est de trouver des techniques permettant d'accroître les débits transmissibles sur le support en vue d'un éventuel partage entre plusieurs utilisateurs. Ainsi, un nouvel algorithme d'allocation des ressources pour le système LP-OFDM avec une mise en œuvre de l'égalisation suivant le critère de l'erreur quadratique moyenne est proposé, ce qui constitue une première contribution originale. En effet, les travaux antérieurs sur l'allocation des ressources pour le système LP-OFDM utilisent le critère de distorsion crête à l'égalisation. Il est connu que le critère de l'erreur quadratique moyenne offre de meilleures performances que le critère de distorsion crête [6]. Nous mettrons en évidence le gain en débit apporté par l'utilisation d'un critère d'erreur quadratique moyenne. De plus, deux autres nouveaux algorithmes d'adaptation des ordres de modulation et des niveaux de puissance, de faible complexité, sont proposés pour maximiser le débit total tout en respectant une contrainte de taux d'erreur binaire (TEB). Les procédures d'optimisation des débits utilisent généralement la contrainte classique du taux d'erreur symbole (TES).

Le premier algorithme résout le problème d'allocation des ressources sous la contrainte d'un TEB crête pour le système LP-OFDM, à savoir que le TEB cible est fixé pour chaque sous-porteuse et toutes les sous-porteuses doivent respecter la valeur de TEB donnée. Le second algorithme vise à maximiser le débit des systèmes OFDM sous la contrainte d'un TEB moyen sur l'ensemble du symbole OFDM. Nous verrons que le passage d'une contrainte de TES crête à une contrainte de TEB crête ou TEB moyen permet d'augmenter le débit des systèmes à porteuses multiples, qu'ils soient OFDM ou LP-OFDM.

La seconde partie concerne l'optimisation de la couche physique des communications CPL dans un contexte multicast. Nous proposons un système LP-OFDM qui permet de mieux exploiter la diversité des liens de transmission pour augmenter les débits des utilisateurs, ce qui constitue la seconde contribution originale de cette thèse. Ce système est analysé dans les contextes multicast monodébit (dans le cas où tous les utilisateurs reçoivent le même débit) et multidébit (dans le cas où les utilisateurs reçoivent des débits différents). Comparées aux méthodes d'allocation des ressources pour le système OFDM multicast, nous verrons que les solutions retenues permettent d'améliorer grandement les débits du système, et nous mettrons en évidence l'apport de la composante de précodage linéaire.

Enfin, dans la dernière partie, la possibilité pour plusieurs utilisateurs de transmettre simultanément des données dans un réseau CPL est analysée. Les systèmes actuels CPL sont caractérisés par des procédés d'accès multiple où les différents utilisateurs transmettent leurs signaux dans des intervalles de temps distincts [7]. De nouveaux algorithmes d'allocation des ressources sont alors proposés et analysés pour des communications simultanées sur le même support physique dans le réseau, dans un contexte centralisé ou décentralisé. Dans un premier temps, nous proposons une nouvelle technique d'accès multiple LP-OFDMA combinant la technique de précodage linéaire et l'OFDMA (*orthogonal frequency division multiple access*). Les résultats obtenus mettent à nouveau en évidence l'intérêt de la solution LP-OFDM. Dans un second temps, nous étudions une modélisation, sous forme de jeu non coopératif, du problème de minimisation des puissances de transmission sous contrainte de débits minimaux. Les développements analytiques, dans le cas simple de deux utilisateurs, ont permis d'obtenir des conditions suffisantes qui garantissent l'existence et l'unicité d'un équilibre de Nash. Sous ces conditions d'existence d'un équilibre, les différents utilisateurs peuvent communiquer simultanément sur le support sans perte de données.

Publications, communications et rapports

Publication internationale

- A. MAIGA, J.-Y. BAUDAIS et J.-F. HÉLARD, « Bit rate optimization with MMSE detector for multicast LP-OFDM systems », soumis à *IEEE transactions on power delivery*.

Communications internationales

- A. MAIGA, J.-Y. BAUDAIS et J.-F. HÉLARD, « Very high bit rate power line communications for home networks », in *IEEE International Symposium on Power Line Communications and Its Applications*, (Dresden, Germany), p. 313–318, mars 2009.
- A. MAIGA, J.-Y. BAUDAIS et J.-F. HÉLARD, « An efficient channel condition aware proportional fairness resource allocation for powerline communications », in *International Conference on Telecommunications*, (Marrakech, Morocco), p. 286–291, mai 2009.
- A. MAIGA, J.-Y. BAUDAIS et J.-F. HÉLARD, « An efficient bit-loading algorithm with peak BER constraint for the band-extended PLC », in *International Conference on Telecommunications*, (Marrakech, Morocco), p. 292–297, mai 2009.
- A. MAIGA, J.-Y. BAUDAIS et J.-F. HÉLARD, « Increase in multicast OFDM data rate in PLC network using adaptive LP-OFDM », in *International Conference on Adaptive Science Technology*, (Accra, Ghana), p. 384–389, déc. 2009.
- A. MAIGA, J.-Y. BAUDAIS et J.-F. HÉLARD, « Subcarrier, bit and time slot allocation for multicast precoded OFDM systems », in *IEEE International Conference on Communications*, (Cap Town, South Africa), p. 1–6, mai 2010.

Communications nationales

- A. MAIGA, J.-Y. BAUDAIS et J.-F. HÉLARD, « Allocation des ressources basée sur le précodage linéaire pour les systèmes OFDM multicast », in *Colloque GRETSI*, (Dijon, France), sept. 2009.
- A. MAIGA, J.-Y. BAUDAIS et J.-F. HÉLARD, « Maximisation du débit des systèmes OFDM multicast dans un contexte de courant porteur en ligne », in *Colloque MajecSTIC*, (Avignon, France), nov. 2009.

Rapports techniques – projet OMEGA

- P. SIOHAN, A. ZEDDAM, G. AVRIL, P. PAGANI, S. PERSON, M. LE BOT, E. CHEVREAU, O. ISSON, F. ONADO, X. MONGABOURE, F. PECILE, A. TONELLO, S. D'ALES-

- SANDRO, S. DRAKUL, M. VUKSIC, J.-Y. BAUDAIS, A. MAIGA et J.-F. HÉLARD, « State of the art, application scenario and specific requirements for PLC », rap. tech., projet OMEGA, avril 2008.
- M. TLICH, P. PAGANI, G. AVRIL, F. GAUTHIER, A. ZEDDAM, A. KARTIT, O. ISSON, A. TONELLO, F. PECILE, S. D’ALESSANDRO, T. ZHENG, M. BIONDI, G. MIJIC, K. KRIZNAR, J.-Y. BAUDAIS et A. MAIGA, « PLC channel characterization and modelling », rap. tech., projet OMEGA, déc. 2008.
 - R. RAZAFFERSON, P. PAGANI, A. ZEDDAM, J.-Y. BAUDAIS, A. MAIGA et O. ISSON, « Report on electromagnetic compatibility of power line communications », rap. tech., projet OMEGA, déc. 2009.
 - A. TONELLO, S. D’ALESSANDRO, M. ANTONIALI, M. BIONDI, F. VERSOLATTO, A. MAIGA, J.-Y. BAUDAIS, F. S. MUHAMMAD, P. SIOHAN, M. L. BOT, H. LIN, P. ACHAICHIA, G. NDO, G. MIJIC, B. CERATO, S. DRAKUL, E. VITERBO et O. ISSON, « Performance report of optimized PHY algorithms », rap. tech., projet OMEGA, juin 2010.
 - A. MAIGA, J.-Y. BAUDAIS, O. ISSON, A. TONELLO et S. D’ALESSANDRO, « Optimized MAC algorithms and performance report », rap. tech., projet OMEGA, mai 2010.

Chapitre 1

Le réseau électrique pour la transmission de données

Ce premier chapitre a pour but de présenter le contexte de l'étude et se décompose pour cela en deux parties complémentaires. La première décrit la technologie courant porteur en ligne (CPL), son origine et son cadre de déploiement. Ainsi, après quelques informations historiques, nous nous attardons sur la description de la technologie CPL. Les aspects normatifs concernant notamment les problèmes de compatibilité électromagnétique sont ensuite abordés : il s'agit là d'un point crucial du développement de la technologie CPL. Un survol des systèmes définis par différents consortiums est ensuite proposé. La seconde partie traite de la caractérisation du canal de propagation. Nous mettons en évidence les caractéristiques des canaux CPL qui sont par la suite utiles au dimensionnement du système. Les réponses impulsionnelles utilisées en simulation sont également introduites. Pour clore ce chapitre, une description du contexte de bruit des lignes électriques est enfin effectuée.

1.1 Un peu d'histoire

À l'origine, le réseau électrique a été conçu pour le transport et la distribution de l'énergie électrique. Toutefois, l'idée d'utiliser les lignes électriques à des fins de communication n'est pas si nouveau. En effet, au début du XX^e siècle déjà, les premiers systèmes CPL virent le jour aux États-Unis dans le cadre d'applications de télémétrie et de télécontrôle. Ainsi, les fournisseurs d'énergie pouvaient effectuer le contrôle et la maintenance des infrastructures, et la mesure de consommation. Dans les années 1980, la technologie CPL s'est peu à peu ouverte au grand public par la voie de la domotique. Différents industriels commercialisaient alors des modules CPL permettant de piloter tout type d'appareil électrique à l'intérieur d'un bâtiment ou d'une maison individuelle.

Ces systèmes permettaient de faire communiquer différents appareils en réseau sans avoir à rajouter de liens physiques. Les applications domestiques les plus courantes sont l'allumage de lampes, le réglage d'un système de chauffage ou encore la surveillance de locaux. La technologie CPL a aujourd'hui dépassé le cadre d'applications bas débit et sert à l'acheminement de données (Internet, vidéos, données, audio) à haut débit. En effet, devant l'explosion de la demande en connexions Internet privées, combinée au fort potentiel du marché des télécommunications, les CPL connaissent depuis peu un regain d'intérêt manifeste au sein de la communauté scientifique et auprès des industriels du secteur des télécommunications. Le principal atout de la technologie CPL, mis en avant par le plus grand nombre, réside dans la densité et l'omniprésence de l'infrastructure électrique. Le réseau de distribution électrique est non seulement présent à l'extérieur et à l'intérieur des bâtiments suivant un maillage extrêmement riche, mais il est en vérité bien plus répandu sur l'ensemble du globe que le réseau des lignes téléphoniques. Lorsque l'on sait que la réduction des coûts de déploiement est un facteur clef dans la réalisation de nouveaux réseaux de communications, il n'est alors pas étonnant que l'on s'intéresse à la technologie CPL aujourd'hui. Pourtant, en raison de l'hostilité du milieu de propagation *a priori* non adapté à la transmission de données, les industriels ont longtemps boudé la technologie CPL. Il a fallu plusieurs années avant que les recherches dans les domaines des communications numériques et du traitement du signal permettent d'obtenir des résultats satisfaisants en terme de débit et de robustesse des systèmes de communication.

1.2 La technologie CPL

La technologie CPL consiste à exploiter le réseau de distribution de l'énergie électrique pour véhiculer des signaux de communications. Lors de la mise en place d'une transmission par courant porteur, on cherche donc à faire cohabiter sur la grille de distribution d'énergie des ondes courtes à hautes fréquences (HF) avec les signaux électriques de fréquence égale à 50 ou 60 Hz selon les pays. Nous verrons que les signaux de communication CPL empruntent des bandes de fréquences pouvant s'étendre jusqu'à 100 MHz. La superposition est obtenue par une opération de couplage inductif ou capacitif qui permet le transfert de l'information sur les lignes d'énergie [8]. Le coupleur doit assurer une séparation galvanique optimale entre les lignes électriques et les appareils de communications, et agit en réception comme un filtre passe-haut afin d'extraire les signaux d'information aux signaux de puissance. Le principe du couplage et de la superposition de ces signaux est représenté sur la figure 1.1. On y retrouve aussi les principaux éléments présents chez l'abonné d'un réseau CPL : le coupleur, le modem, et les appareils connectés au réseau tels que l'ordinateur, la télévision, le téléphone, etc.

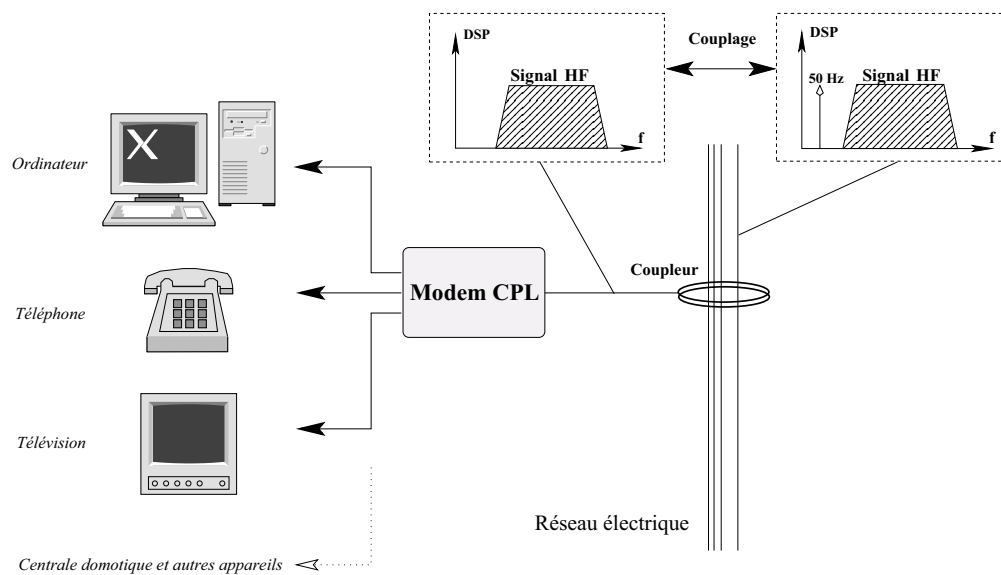


FIGURE 1.1 – Principe de couplage du modem CPL sur le réseau électrique.

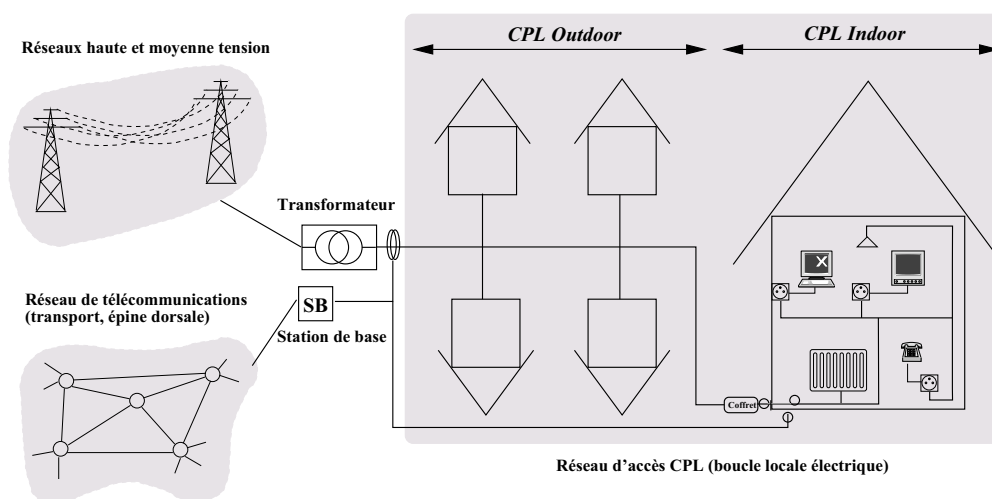


FIGURE 1.2 – Structure de la boucle locale électrique.

Le réseau électrique global est structuré en trois sous-réseaux que l'on identifie classiquement à partir du voltage correspondant : haute tension (HT), moyenne tension (MT) et basse tension (BT). En considérant le réseau électrique comme un système de communication, deux aspects de la technologie CPL sont à distinguer. La partie du réseau composée des lignes extérieures est appelée réseau d'accès (*outdoor*), et la partie correspondant aux installations privées est appelée réseau résidentiel (*indoor*¹). Une représentation schématique de cette nouvelle boucle locale est donnée figure 1.2. Le réseau *outdoor* est connecté à l'épine dorsale du réseau de télécommunications par l'intermédiaire d'un coupleur et d'une station de base placés au pied du transformateur MT/BT. Il s'agit pour la station de base de coupler une arrivée Internet haut débit (obtenue par une arrivée satellite ou fibre optique au niveau du transformateur MT/BT) au réseau électrique local. Ainsi, tous les utilisateurs desservis par ce transformateur peuvent bénéficier de cet accès haut débit via le réseau électrique. Le réseau *indoor* part du principe qu'une connexion haut débit existe déjà dans l'habitation, elle peut avoir été amenée par le câble, la fibre optique, les ondes ou même le CPL *outdoor*. Un modem CPL permet alors de convertir les données reçues de la connexion haut débit sous une forme adaptée à leur transmission sur les lignes d'énergie. Ces données seront alors accessibles par les autres équipements (ordinateur, TV, imprimante, etc) connectés aussi au réseau électrique via un modem CPL. Il faut noter qu'avec l'essor des services d'automatisation, non seulement pour leur application dans le secteur industriel et commercial, et dans les grands bâtiments, mais aussi pour leur application dans les ménages privés, le réseau CPL *indoor* est une solution raisonnable pour la réalisation de réseaux avec un grand nombre de terminaux. En effet, les systèmes fournissant les services d'automatisation tels que la surveillance, la sécurité, la commande de chauffage, le contrôle automatique de la lumière, doivent connecter un grand nombre de terminaux tels que les capteurs, les caméras, les moteurs électriques, l'éclairage, etc.

1.3 Organismes de réglementations et consortiums

L'un des rôles des consortiums et des organismes de régulations est de définir des règles qui peuvent être acceptées par les différents acteurs (opérateurs, fournisseurs de services Internet, les équipementiers et les intégrateurs réseaux) concernés par l'utilisation possible des réseaux et des équipements CPL. Comme la capacité dépend de la largeur de bande et du niveau de puissance du signal, ces 2 paramètres sont d'une importance primordiale.

1. Les termes anglophones sont ici d'usage plus que courant, c'est pourquoi nous les utiliserons dans la suite du document.

1.3.1 Compatibilité électromagnétique

Rappelons tout d'abord que les lignes électriques n'ont pas été conçues pour autre chose que de transporter l'énergie électrique avec le moins de pertes possible aux fréquences de 50 ou 60 Hz selon les pays. Utiliser ces lignes pour mettre en place des communications CPL signifie qu'elles devront transmettre des signaux à des fréquences allant de quelques kilohertz à plusieurs dizaines de mégahertz. De plus, ces lignes se comportent comme des antennes donnant naissance à un champ électromagnétique qui vient perturber l'environnement. Le champ induit par les lignes agit comme un perturbateur du point de vue des autres systèmes de communication et son niveau doit être limité à un certain seuil pour des raisons de compatibilité électromagnétique (CEM), afin de ne pas entraver leur bon fonctionnement. Inversement, les contraintes imposées par la CEM doivent aussi permettre aux équipements CPL de fonctionner correctement sous l'influence de la pollution électromagnétique environnante. Tous ces aspects sont pris en compte par les organismes de régulation pour définir les normes.

1.3.2 Organismes de normalisations

1.3.2.1 CEI / CISPR

Le CISPR (Comité international spécial des perturbations radioélectriques) est responsable de l'élaboration et de la mise à jour de plusieurs normes internationales sur la CEM (émission et immunité) pour les grandes familles de produits électriques ou électroniques dont le but principal est la protection des services radio dans la gamme de fréquences de 9 kHz à 400 GHz. Le CISPR est divisé en plusieurs sous-comités dont le sous-comité « CISPR/I » qui est en charge du développement des méthodes de mesure et de l'établissement des limites d'émission pour les équipements des technologies de l'information. Les équipements CPL sont soumis aux prescriptions de la norme CEM CISPR 22. Un projet de normalisation sur les limites d'émission et la méthode de mesure pour les équipements de télécommunication à large bande sur lignes électriques dans la bande de fréquences de 150 kHz à 30 MHz propose les limites données dans le tableau 1.1. Les limites de rayonnement et la méthode de mesure dans la gamme de fréquences supérieures à 30 MHz restent inchangées et ne sont par conséquent pas prises en compte dans ce projet de normalisation. Ces limites sont rappelées dans le tableau 1.2.

1.3.2.2 CENELEC

Le CENELEC (Comité européen de normalisation électrique) est l'organisation européenne chargée de la normalisation électrique et électronique. Pour améliorer l'efficacité de nombreuses normes qui sont maintenant définies au niveau international, le CENE-

TABLE 1.1 – Limites pour les perturbations conduites dans les ports CPL (fonction de télécommunication inactive) [2].

Bande de fréquences (MHz)	Limites (dB (μV))	
	Quasi-crête	Moyenne
1,605 – 5	56	46
5 – 30	60 – 56	50

TABLE 1.2 – Limites des perturbations rayonnées à une distance de 10 mètres [2].

Bande de fréquences (MHz)	Limites quasi-crête (dB (μV))
30 – 230	30
230 – 1000	37

LEC travaille en étroite collaboration avec la Commission électrotechnique internationale (CEI) et le CISPR. Il existe plusieurs commissions dont la commission SC 205 A qui est en charge d'établir des normes harmonisées pour les systèmes de communications utilisant les lignes d'énergie basse tension ou le câblage des bâtiments comme moyen de transmission dans la bande de 3 kHz à 30 MHz. Cette tâche comprend aussi l'attribution des bandes de fréquences pour la transmission des signaux. Plusieurs groupes de travail sont rattachés à cette commission pour définir les exigences et les normes pour l'utilisation du CPL en respect des normes déjà existantes. Ces travaux portent sur, entre autres, les exigences d'immunité et les limites d'émission.

1.3.2.3 FCC partie 15

La FCC (*Federal communications commission*) est la principale agence gouvernementale américaine chargée de la planification des fréquences. La partie 15 de la FCC est un essai de norme commune pour la plupart des équipements électroniques. Cette partie couvre les caractéristiques techniques, les exigences administratives et d'autres conditions relatives à la commercialisation d'appareils relevant de ladite partie. Selon le type d'équipements, la vérification, la déclaration de conformité, ou la certification est le processus mis en place dans la partie 15 de la FCC [9].

Cette partie couvre aussi le CPL indoor. Les règles actuelles de la FCC exigent que les systèmes à courant porteur (y compris le CPL) doivent respecter les limites générales de rayonnement d'une émission intentionnelle. Un émetteur intentionnel est celui qui transmet un signal radio dans le cadre de son fonctionnement normal. En dessous de 30 MHz de fréquence, les systèmes CPL sont limités à une intensité de champ rayonné de 30 μV par mètre mesurée à 30 mètres de la source du signal. Par contre, au delà de 30 MHz, les systèmes CPL sont limités à une intensité de champ rayonné de 100 μV par mètre mesurée à 3 mètres de la source du signal. Dans la plupart des cas, la source sera

TABLE 1.3 – comparaison des principaux paramètres des spécifications de produits CPL existants.

	CEPCA	Homeplug	UPA
Modulation	wavelet OFDM	windowed OFDM	windowed OFDM
Channel coding	RS-CC, LDPC	Parallel-concatenated TCC	RS + 4D-TCM concatenation
Mapping	PAM 2-32	QAM 2, 4, 8, 16, 64, 256, 1024	ADPSK 2-1024
FFT/FB size	512 (extendable to 2048)	3072	NC
Max number of carriers	NC	1536	1536
Sample frequency	62,5 MHz	75 MHz	NC
Frequency band	4-28 MHz	2-28 MHz (2-28 MHz optional)	0-30 MHz (0-20 MHz optional)
PHY Rate	190 Mbps	200 Mbps	200 Mbps
Power Spectral Density	NC	-56 dBm/Hz	-56 dBm/Hz
Media Access Method	TDMA-CSMA/CA	TDMA-CSMA/CA	ADTDM
MAX number of nodes	64	NC	64

le câblage électrique à l'intérieur d'un bâtiment ou les lignes électriques qui passent à proximité des résidences et des entreprises aux États-Unis [10].

1.3.3 Les consortiums

Les industriels du CPL se sont regroupés dans des consortiums pour accompagner les travaux de normalisation. Ces consortiums leur permettent également de mettre en commun leurs points de vues, leurs intérêts et de proposer leurs propres normes. Parmi ces consortiums, ceux qui proposent des produits sur le marché sont *Homeplug*, *UPA* et *CEPCA*. Le tableau 1.3 regroupe les principales différences entre les produits de CEPCA (Panasonic), de Homeplug et de l'UPA. On peut voir que ces consortiums utilisent deux types de modulations multiporteuses : Homeplug et UPA utilisent l'OFDM (*orthogonal frequency division multiplex*) fenêtré (*windowed OFDM*) et CEPCA utilise l'OFDM basée sur une transformation en ondelettes. On note une différence majeure dans les techniques de codage de canal. Les trois solutions s'appuient sur des techniques d'accès multiple de type TDM (*time division multiplex*), offrant des débits maximums comparables de l'ordre de 200 Mb/s.

1.3.3.1 Alliance Homeplug

L'alliance internationale Homeplug, créée en mars 2000 et qui compte aujourd'hui plus de 75 membres, travaille à créer des programmes de spécifications et de certification pour une utilisation fiable du réseau CPL [11]. L'alliance accélère la demande d'autorisation de mise sur le marché de ses produits et services Homeplug dans le monde entier grâce à des programmes de parrainage et de formation sur le marché. L'alliance a créé plusieurs spécifications pour les normes CPL comme HomePlug 1.0, HomePlug AV (HPAV) et HomePlug BPL.

1.3.3.2 UPA

L'UPA (*Universal powerline association*), fondée en 2004 pour intégrer le CPL dans le paysage des télécommunications, est une organisation internationale travaillant en collaboration avec le projet de recherche européen OPERA. L'UPA favorise le développement de modules DS2 et a créé la norme DHS (*digital home standard*), dont le but est de fournir une spécification complète pour les fournisseurs de composants électroniques en vue de la conception de circuits intégrés pour la voix, la vidéo et la distribution de données sur les lignes électriques [12].

1.3.3.3 CEPCA

La CEPCA (*Consumers electronics powerline communication alliance*) est une autre alliance qui veut promouvoir le CPL. Ses quatorze membres sont principalement des fabricants japonais (Sony, Mitsubishi, Panasonic, etc.). L'objectif de la CEPCA est de permettre aux différents systèmes CPL de coexister. Panasonic fait la promotion du module HD-PLC (*high definition powerline communication*), qui utilise une modulation OFDM combinée avec les ondelettes, et un filtre coupe-bande programmable qui empêche les interférences avec d'autres communications radios telles que les radios amateurs.

1.4 Projet OMEGA

Le développement des réseaux locaux domestiques capables de fournir des services très hauts débits à la vitesse du gigabit est un élément clé dans la vision qu'a l'Union européenne de l'Internet du futur. Les consommateurs exigeront des réseaux domestiques devant être simples à installer, sans aucun ajout de nouveaux câbles, et assez faciles à utiliser afin que les services de communication en cours d'exécution sur le réseau soit « juste un autre utilitaire », comme, par exemple, l'électricité, l'eau et le gaz le sont aujourd'hui [5]. Le projet OMEGA vise à combler le fossé entre le réseau domestique et le réseau d'accès, pourvoyant ainsi une connectivité au gigabit par seconde (Gb/s).

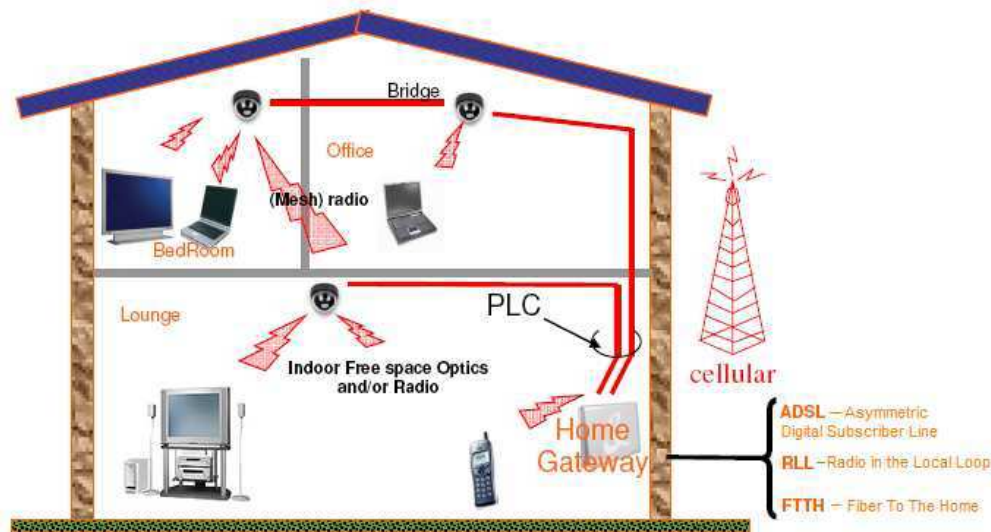


FIGURE 1.3 – Réseau domestique à très haut débit – OMEGA

OMEGA est un projet intégré co-financé par la Commission européenne dans le cadre du programme de ICT FP7, prévu pour une durée de 3 ans de janvier 2008 à décembre 2010. Le consortium pluridisciplinaire est constitué de 20 partenaires européens industriels, opérateurs de télécommunications et universitaires.

1.4.1 Objectifs

L'objectif principal du projet OMEGA est d'assurer une capacité de 1 Gb/s et une faible latence du réseau domestique vers le réseau d'accès et vice versa permettant ainsi l'accès et le développement de services nouveaux et innovants. Ce réseau local devra permettre des communications à très haut débit via différents systèmes de communication, incluant les technologies filaires (CPL) et sans fil (ultra large bande (ULB), WIFI, systèmes optiques, technologie 60 GHz). La figure 1.3 présente une vision de OMEGA du réseau domestique à très haut débit. Les données en provenance des réseaux d'accès (ADSL, fibre optique, réseau mobile) entrent dans la maison et sont acheminées par la passerelle domestique. La passerelle à son tour est reliée aux équipements OMEGA, qui peuvent supporter des communications à la vitesse de 1 Gb/s par l'utilisation des technologies 60 GHz ou optiques sans fil en ligne de visée (LOS, *ligne of sight*). La passerelle peut alors utiliser les communications par CPL à très haut débit pour se connecter aux équipements OMEGA. Les communications dans la maison sont assurées par la technologie ULB et la diffusion par l'utilisation des communications en lumière visible (VLC, *visible-light communications*) [13].

Une nouvelle couche MAC (appelée couche inter-MAC) assure l'interopérabilité de

ces différentes technologies afin de fournir les services et la connectivité à un certain nombre d'équipements auxquels l'utilisateur souhaite se connecter et ce de n'importe quelle pièce du domicile. Cette couche gère tous les aspects du réseau domestique OMEGA, y compris la qualité de service, l'équilibrage des charges, l'intra et l'inter-transfert de technologies et assure également aux différents terminaux la meilleure connexion possible. OMEGA ouvre ainsi la voie à une approche entièrement nouvelle puisque ce sont les premiers travaux de recherche sur la convergence de ces diverses technologies sans fil et filaires dans des applications exigeant du très haut débit dans le réseau domestique.

Le projet OMEGA est divisé en 8 groupes de travail afin d'atteindre les objectifs attendus. Les travaux de cette thèse ont été menés dans le cadre du groupe de travail 3 concernant la technologie CPL.

1.4.2 CPL OMEGA

La technologie CPL bénéficie de l'avantage de la disponibilité du réseau électrique pour les communications robustes à haut débit et remplit facilement un des principaux critères du projet OMEGA qui est désigné par la périphrase « sans aucun ajout de nouveaux câbles ». Dans un contexte où les services nécessitent des communications à très haut débit, un seul point d'accès dans le réseau domestique s'avère être insuffisant et le CPL a cette capacité de connecter différents segments du réseau envisagé. Par contre, cette technologie doit être considérablement améliorée pour fournir de tels services à très haut débit et être intégrée dans un tel réseau domestique. Les équipements CPL actuels offrent des débits théoriques de l'ordre de 200 Mb/s. L'objectif principal du groupe de travail sur le CPL est d'étudier et d'implémenter une solution CPL capable de fournir des services à la vitesse du gigabit par seconde. Pour atteindre cet objectif, la possibilité d'augmenter la bande passante jusqu'à 100 MHz est étudiée. Il faut rappeler que les équipements CPL actuellement disponibles sur le marché utilisent la bande de 2 à 30 MHz. De plus, des techniques avancées de modulation, de codage de canal et de multiplexage dans le but d'optimiser les couches PHY et MAC sont étudiées. Pour la couche PHY, l'accent est mis sur le développement d'une interface de transmission compatible avec les spécifications actuelles de la norme HPAV. De nouvelles méthodes de modulations multiporteuses telles que l'OFDM combinée avec la technique OQAM (*offset quadrature amplitude modulation*) ou la technique de précodage linéaire, et la modulation FMT (*filtered multi tone*) ont été évaluées pour obtenir le gigabit par seconde [14]. Pour la couche MAC, une approche cognitive est envisagée afin d'assurer la compatibilité et l'interopérabilité avec la norme HPAV dans la sous bande 0–30 MHz. Dans l'autre sous-bande (au-dessus des 30 MHz), l'approche étudiée devra assurer la coexistence avec des solutions concurrentes non HPAV ou d'autres technologies. Dans les deux cas, des stratégies spécifiques permettant le partage des ressources tout en garantissant une qualité de ser-

TABLE 1.4 – Principaux paramètres des systèmes HPAV et OMEGA.

Paramètres	Bande passante	Δ_f	N_{utile}	N_{FFT}	Bitcap
HPAV	[2; 28] MHz	24,414 kHz	917	3072	10
OMEGA	[0; 100] MHz	24,414 kHz	3345	8192	14

vice donnée sont étudiées [15]. Les techniques d'accès multiple telles que le FDMA (*frequency-division multiple access*), le TDMA (*time-division multiple access*), le CDMA (*code-division multiple access*) ou encore l'OFDMA (*orthogonal frequency-division multiple access*) sont envisagées pour allouer les ressources (les bandes de fréquence, les intervalles de temps et les séquences de précodage) tout en exploitant la diversité des canaux de transmission [15].

Des campagnes de mesures du canal de propagation jusqu'à 100 MHz ont permis d'avoir des modèles de canaux de propagation pour le CPL à bande étendue [1]. Les impacts d'un tel système sur les contraintes CEM et les niveaux de puissance des signaux pouvant être transmis ont été étudiés [10]. Les premières observations de ces études montrent que le niveau de DSP du signal dans la bande 30–100 MHz doit être de -80 dBm/Hz afin de respecter les contraintes CEM. La figure 1.4 présente le masque de densité spectrale de puissance (DSP) allouée sur toute la bande passante. Le masque dans la bande 0–30 MHz est conforme au masque défini par la norme HPAV et a été construit en mettant à zéro toutes les sous-porteuses comprises dans les bandes de fréquences allouées aux radio amateurs. Ce masque est utilisé dans la suite pour les études de performances. Le tableau 1.4 résume les principaux paramètres du système OMEGA et du système HPAV. Le bitcap désigne le nombre maximum de bits qui peuvent être alloués à une sous-porteuse, Δ_f est l'espace inter-porteuse, N_{utile} le nombre total de sous-porteuses utilisées dans le système de transmission et N_{FFT} la taille de la FFT. Dans le but d'avoir un système OMEGA interopérable avec le système HPAV, l'espace inter-porteuse a été choisi identique à celui de HPAV. La norme HPAV utilise une taille de FFT surdimensionnée à 3072 en considérant la bande passante de 0–37,5 MHz, les sous-porteuses dans la bande au delà des 28 MHz n'étant pas utilisées.

1.5 Caractéristiques du canal de propagation

D'après le paradigme proposé par Shannon, toute chaîne de communications peut être décomposée en trois blocs, à savoir l'émetteur, le milieu de transmission appelé canal de transmission et le récepteur. Du point de vue de la théorie des communications, le canal vu par le système comporte non seulement le médium à travers lequel se propage le message, les filtres d'émission et de réception présents dans toute chaîne de communica-

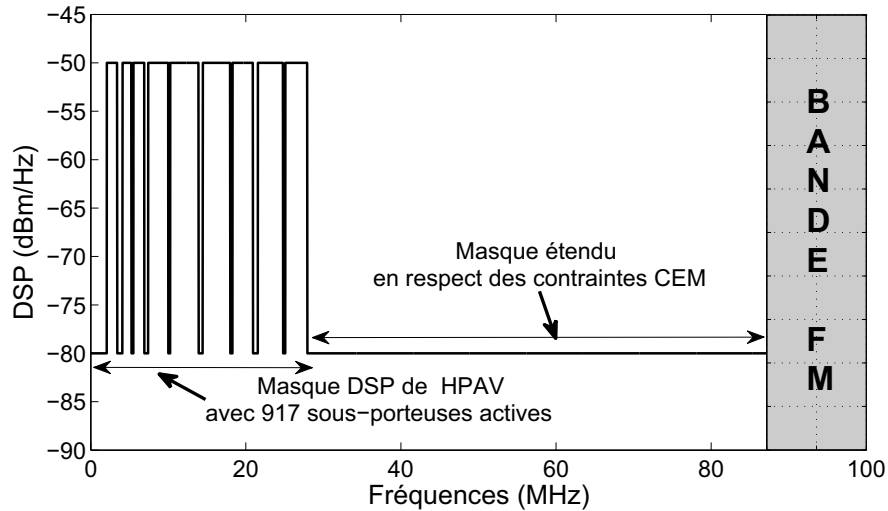


FIGURE 1.4 – Proposition de masque DSP pour le système CPL OMEGA.

tions, mais aussi les sous-ensembles qui permettent au message d'accéder à ce médium. Afin de développer des systèmes CPL efficaces et de proposer des améliorations aux technologies existantes, il est donc nécessaire de caractériser l'infrastructure électrique. Les principales composantes du canal de transmission CPL, à savoir la réponse du canal qui rend compte des phénomènes qui viennent modifier la forme des ondes émises, ainsi que les brouilleurs qui viennent s'ajouter au signal reçu et dont l'origine peut être multiple, sont caractérisées dans le projet OMEGA [1].

1.5.1 Réponse du canal

Lors de leur propagation à travers le canal de transmission, les ondes émises sont sujettes à différents phénomènes qui viennent modifier leur forme, c'est-à-dire leur amplitude et leur phase. Dans le cas le plus général, il peut s'agir de phénomènes d'atténuation, de déphasage, de réflexion, de diffraction ou encore de diffusion, selon les interactions entre les ondes et le support physique. Le canal CPL est de plus caractérisé par d'autres différences par rapport aux autres médias car ses interférences et ses niveaux de bruit sont plus élevés. Plusieurs approches, pour caractériser le canal CPL, ont été proposées dans la littérature. Une approche intéressante décrit le canal CPL en modélisant les multitrajets [16, 17]. D'autres travaux tentant de modéliser le canal CPL comme une ligne de transmission bifilaire (à deux conducteurs) [18] ou trifilaire (à trois conducteurs) [19] ont également été publiés. Ces approches ne décrivent que partiellement la physique sous-jacente de la propagation du signal CPL et par conséquent ne permettent pas de prendre en compte les propriétés générales du canal CPL [1]. De plus, l'approche reposant sur la modélisation des multitrajets est basée sur un modèle para-

TABLE 1.5 – Répartition des fonctions de transfert par site.

Numéro du site	Type du site	Nombre de fonctions de transfert
1	Maison, ville	19
2	Nouvelle maison, ville	13
3	Appartement récemment restauré, ville	12
4	Maison récente, ville A	28
5	Maison récente, ville B	34
6	Maison récente, campagne	22
7	Ancienne maison, campagne	16

métrique où la plupart des paramètres ne peuvent être estimés qu’après avoir mesuré la réponse impulsionnelle (RI) du canal, limitant ainsi la capacité *a priori* à modéliser le canal CPL. Dans [20], un modèle déterministe du canal CPL est proposé à partir de la représentation complexe des multitrajets avec la théorie des matrices. Le modèle développé reste incomplet dans la mesure où ni les valeurs des impédances à l’extrémité des branches, ni leur variation dans le temps ne sont prises en compte.

L’approche choisie dans le projet OMEGA est de modéliser le canal de propagation CPL à partir d’études statistiques sur un grand nombre de mesures de la réponse du canal dans la bande étendue à 100 MHz [1]. Dans ce qui suit, nous présentons les modèles de canaux qui ont découlé des campagnes de mesure.

1.5.1.1 Modèle de canaux CPL

Les modèles de canaux CPL proposés découlent de mesures réalisées dans la bande de 30 kHz à 100 MHz dans divers environnements domestiques (campagnes et villes, bâtiments neufs et anciens, appartements et maisons), comme indiqué dans le tableau 1.5. Ces campagnes de mesure ont été menées par notre partenaire ORANGE LABS et servent de base d’étude pour le projet OMEGA. Les fonctions de transfert étudiées ont été relevées sur 7 différents sites et un total de 144 fonctions de transfert ont été mesurées. Pour chaque site, les fonctions de transfert sont mesurées entre une prise de courant principal et toutes les autres prises de la maison. Les fonctions de transfert mesurées ont été obtenues dans le domaine fréquentiel à l’aide d’un analyseur de réseau vectoriel.

L’observation des résultats de mesure des fonctions de transfert des canaux CPL permet de distinguer deux catégories de canaux :

- Les fonctions de transfert où les prises de l’émetteur et du récepteur se rapportent au même circuit électrique, c’est-à-dire se trouvent en série sur la même branche correspondant à un fusible dans le boîtier électrique ;
- Les fonctions de transfert où les prises de l’émetteur et du récepteur se rapportent à deux circuits électriques différents.

Pour chaque couple composé d'une catégorie de canal et d'un site de mesure, les fonctions de transfert sont presque identiques et indépendantes des positions des prises [1]. En effet, les observations des mesures montrent que les pics et les évanouissements du canal sont quasiment aux mêmes fréquences. Ces observations conduisent à l'idée de classer les canaux CPL mesurés en fonction de leur potentiel de transmission. En raison de l'impossibilité de calculer les distances séparant les émetteurs des récepteurs, les canaux CPL ont été regroupés en plusieurs classes par ordre croissant de leur capacité. Les capacités sont calculées à partir de la formule de Shannon avec les mêmes niveaux de bruit et le même masque de puissance d'émission défini à la figure 1.4. Cette formule est

$$C = \Delta_f \sum_{n=1}^N \log_2 \left(1 + \frac{P(f_n) |H(f_n)|^2}{N_0} \right), \quad (1.1)$$

avec f_n les fréquences dans la bande passante, $P(f_n)$ le niveau de puissance admis pour la fréquence f_n , N_0 le niveau de bruit supposé constant pour toutes les fréquences, Δ_f et N sont donnés dans le tableau 1.4. Les fonctions de transfert de la classe 1 sont celles qui conduisent au plus faible débit de transmission, les fonctions de transfert de la classe 2 sont celles qui conduisent à des débits supérieurs à ceux de la classe 1 mais inférieurs à celles de la classe 3, etc. Le processus de classification des canaux CPL a conduit à 9 classes de canaux [1]. La figure 1.5 présente les fonctions de répartitions des capacités des canaux pour les différentes classes. Les canaux ont été générés par le logiciel WITS (*Wideband Indoor Transmission channel Simulator*), basé sur des modèles de canaux mesurés [1]. Les pourcentages d'apparition des classes, les sites correspondants de mesure et les modèles d'atténuations moyennes sont détaillés dans le tableau 1.6. Nous pouvons observer que, pour le niveau de bruit blanc choisi, les capacités varient de 276 à 542 Mb/s pour la classe 1 et de 1516 à 1796 Mb/s pour la classe 9, avec la DSP définie figure 1.4. Pour la valeur 0,5 de la fonction de répartition, l'écart minimal entre les capacités est supérieur à 100 Mb/s. Les fonctions de transfert sont distribuées sur l'ensemble des 9 classes avec un nombre à peu près constant d'apparitions pour les classes 4, 5, 6 et 7 (entre 9,79% et 12,58%), un nombre plus conséquent pour les classes 2 et 3 (autour de 17%) et un nombre plus faible pour les autres classes (environ 3% pour la classe 1 et 7% pour les classes 8 et 9). Mise à part la classe 1, les fonctions de transfert pour les autres classes ont été mesurées sur au moins 3 sites différents. Ces sites (voir tableau 1.5) sont variables en termes de taille (appartements, maisons) et de date de construction (récentes et anciennes installations électriques). Au regard des objectifs du projet OMEGA en terme de débit pour le CPL, les canaux des classes 1 à 4 ne permettent pas d'atteindre le gigabit par seconde. Ce gigabit est toujours atteint pour les classes 7 à 9 et susceptible d'être atteint à 98% pour la classe 6 et 20% pour la classe 5.

Sur la figure 1.6, les réponses impulsionnelles et les fonctions de transfert de canaux

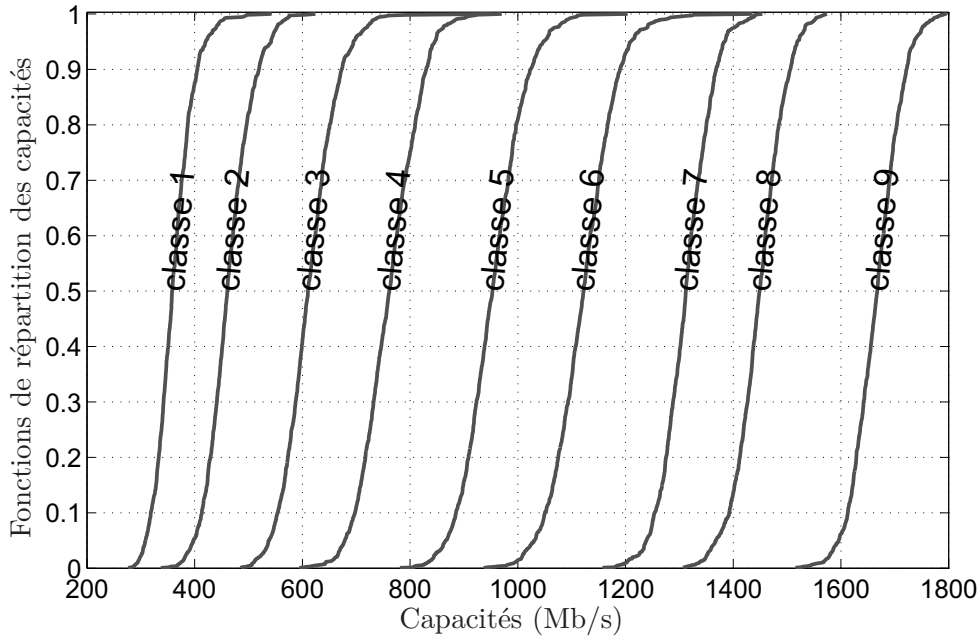


FIGURE 1.5 – Fonctions de répartition des capacités des canaux issus des différentes classes pour $\Delta_f = 24,414$ kHz et $N_0 = -140$ dBm/Hz [1].

de transmission d'une classe de catégorie « mauvaise » (classe 2), d'une classe de catégorie « moyenne » (classe 5) et d'une classe de catégorie « bonne » (classe 9) sont représentées. Comme attendu, la réponse présentant le moins de distorsion est la réponse issue de la classe 9, avec très peu d'évanouissements fréquentiels profonds. Ce canal est également celui dont l'étalement des retards est le plus faible, limité à environ $0,6 \mu\text{s}$. Le canal de la classe 2 présente le niveau de distorsion le plus élevé des trois canaux. Sa réponse impulsionnelle est très riche et s'étale sur environ $4 \mu\text{s}$. L'atténuation fréquentielle de ce canal est très forte avec des évanouissements de l'ordre de -80 dB. Le canal de la classe 5 affiche des évanouissements moins marqués que le canal de la classe 2 et un étalement de sa réponse impulsionnelle d'environ $2 \mu\text{s}$.

Les réponses présentées sont utilisées dans les simulations menées par la suite, en tant qu'échantillons représentatifs du comportement du canal CPL. En général, les canaux de transmission CPL sont supposés comme variant lentement au fil du temps. Dans de nombreux travaux, les canaux CPL semblent être quasi-statiques. En effet, leur réponse fréquentielle varie moins que la fonction de transfert d'un canal radiomobile [21]. Cependant, dans la littérature (par exemple, [22, 23]), des cas sont répertoriés où le canal de propagation CPL est variant dans le temps, en particulier, lorsque des équipements à impédances variables au cours du temps sont proches du réseau CPL où est prélevé la réponse du canal. Dans la suite de ce document, nous considérerons lors de nos simulations

TABLE 1.6 – Pourcentages d'apparition des classes, sites correspondants de mesure et atténuations moyennes.

Classes	Pourcentages de canaux	Sites	Atténuations moyennes
1	3,49%	6	$-80 + 30 \times \cos\left(\frac{f}{5,5 \cdot 10^7} - 0,5\right)$
2	16,78%	1, 5 et 6	$-43 + 25 \times \exp\left(-\frac{f}{3 \cdot 10^6}\right) - \frac{15}{10^8}f$
3	18,18%	1, 3, 4, 5, 6 et 7	$-38 + 25 \times \exp\left(-\frac{f}{3 \cdot 10^6}\right) - \frac{14}{10^8}f$
4	11,88%	1, 3, 4 et 7	$-32 + 20 \times \exp\left(-\frac{f}{3 \cdot 10^6}\right) - \frac{15}{10^8}f$
5	11,88%	1, 3, 4, 5 et 7	$-27 + 17 \times \exp\left(-\frac{f}{3 \cdot 10^6}\right) - \frac{15}{10^8}f$
6	12,58%	2, 4, 5 et 7	$-38 + 17 \times \cos\left(\frac{f}{7 \cdot 10^7}\right)$
7	9,79%	2, 4, 5 et 6	$-32 + 17 \times \cos\left(\frac{f}{7 \cdot 10^7}\right)$
8	7,69%	2, 3, 4 et 6	$-20 + 9 \times \cos\left(\frac{f}{7 \cdot 10^7}\right)$
9	7,69%	1, 2, 3, 5, 6 et 7	$-13 + 17 \times \cos\left(\frac{f}{4,5 \cdot 10^7} - 0,5\right)$

que les canaux CPL sont quasi-statiques et statiques à l'échelle de la communication.

1.5.2 Sources de bruit

Outre les distorsions apportées par les réponses des canaux sur la forme des signaux propagés sur les lignes électriques, il faut aussi considérer comme élément perturbateur la part de bruit, pris au sens large du terme, qui vient s'ajouter à l'énergie utile transmise. À la différence de la plupart des canaux de communications, le bruit présent à

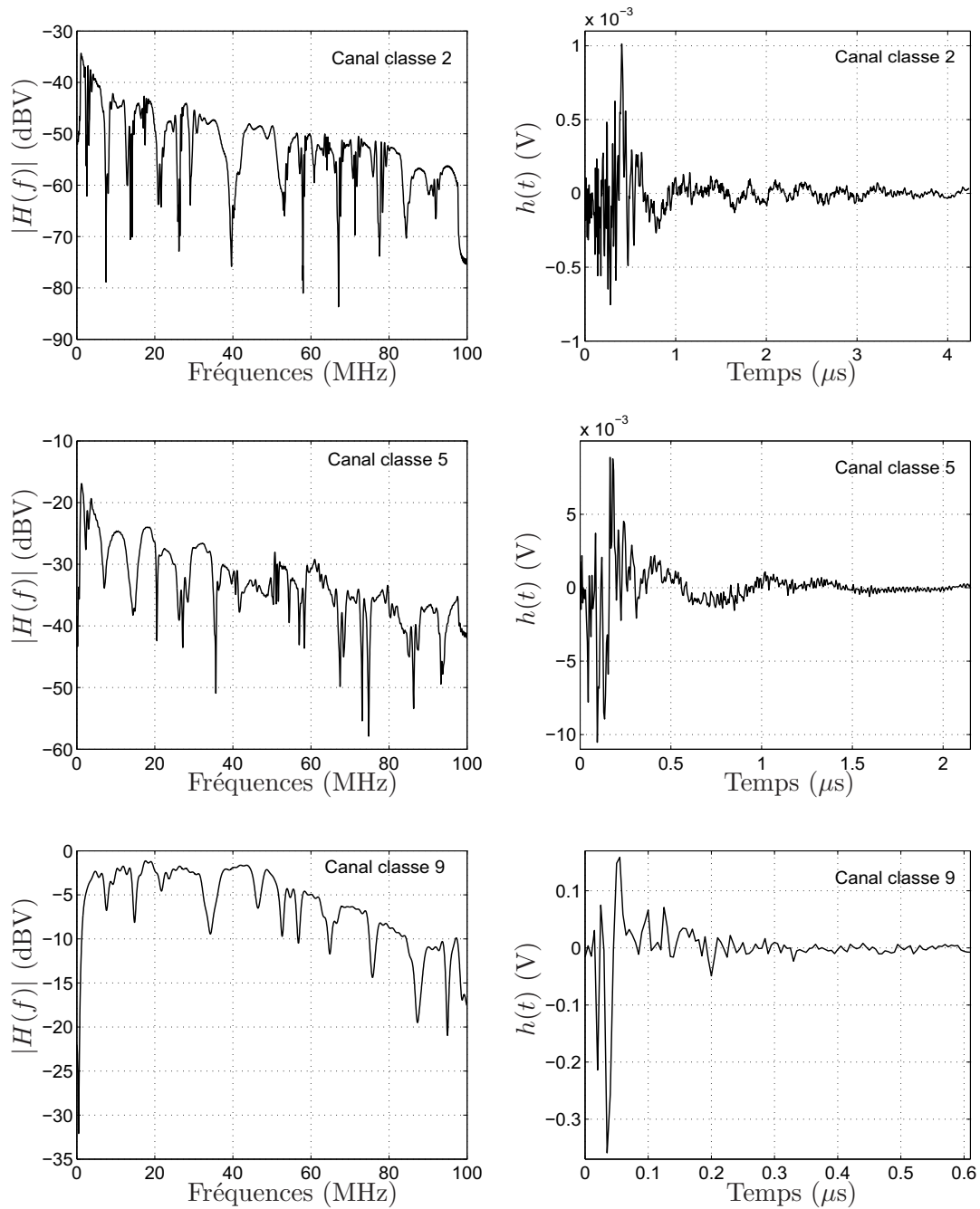


FIGURE 1.6 – Réponses du canal électrique pour les trois classes sélectionnées : classe 2, classe 5 et classe 9. À gauche, les fonctions de transfert, à droite les réponses impulsionnelles.

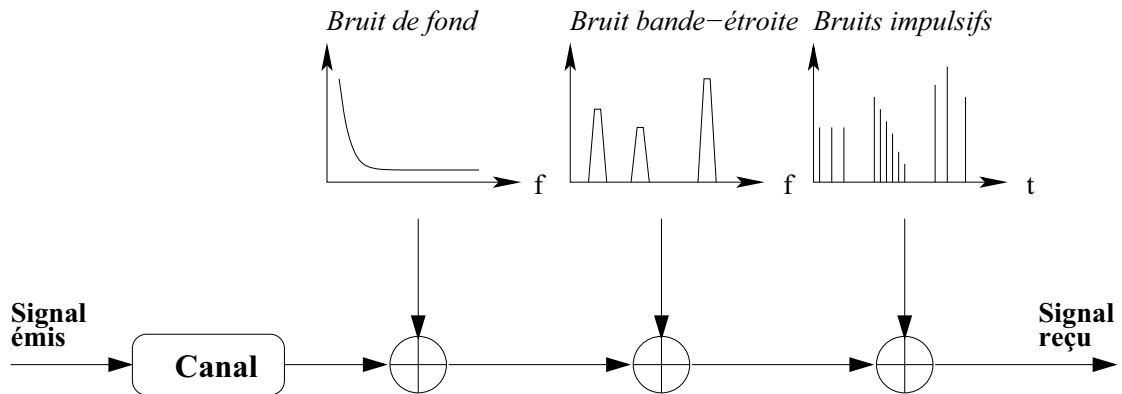


FIGURE 1.7 – Ensemble des types de bruits additifs rencontrés sur les lignes électriques

l'entrée d'un récepteur CPL ne se réduit pas à la seule contribution du bruit thermique, encore appelé bruit blanc additif gaussien (AWGN, *additive white gaussian noise*) [3]. On doit cette spécificité à la grande variété d'appareils connectés au réseau électrique, ainsi qu'à la multiplicité des perturbations captées par rayonnement. Des travaux précédents comme par exemple ceux publiés dans [1, 24], permettent de classer les bruits rencontrés en trois grandes catégories. La représentation schématique de cette classification est donnée figure 1.7. On distingue le bruit de fond, les bruits à bande étroite et les bruits impulsifs. Les mesures disponibles sur ces différentes sources de bruit ont généralement montré que les deux premières sources de bruit et une partie des bruits impulsifs demeuraient stationnaires sur des périodes temporelles pouvant s'étendre à plusieurs minutes, voire plusieurs heures. Au contraire, les autres sources de bruit impulsif possèdent des caractéristiques variables en quelques millisecondes [25]. Dans le projet OMEGA, plusieurs campagnes de mesure ont été réalisées afin de caractériser toutes les sources de bruit.

1.5.2.1 Bruit stationnaire

Le bruit stationnaire est composé du bruit de fond, des bruits à bande étroite et du bruit impulsif périodique asynchrone. Les impulsions qui composent ce bruit impulsif ont généralement une fréquence de répétition comprise entre 100 et 200 kHz. Ce type de bruit est le plus souvent engendré par les blocs d'alimentation à découpage rencontrés dans beaucoup d'équipements domestiques d'aujourd'hui. À cause de la forte occurrence des impulsions, les fréquences occupées sont proches et constituent des groupements de raies qui peuvent être assimilées à une forme de bruit à bande étroite. Leur puissance est cependant bien plus faible que celle des bruits engendrés par les activités de radiodiffusion, voire parfois à peine supérieure au niveau du bruit de fond.

L'analyse fréquentielle des bruits stationnaires mesurés révèle une forme décroissante

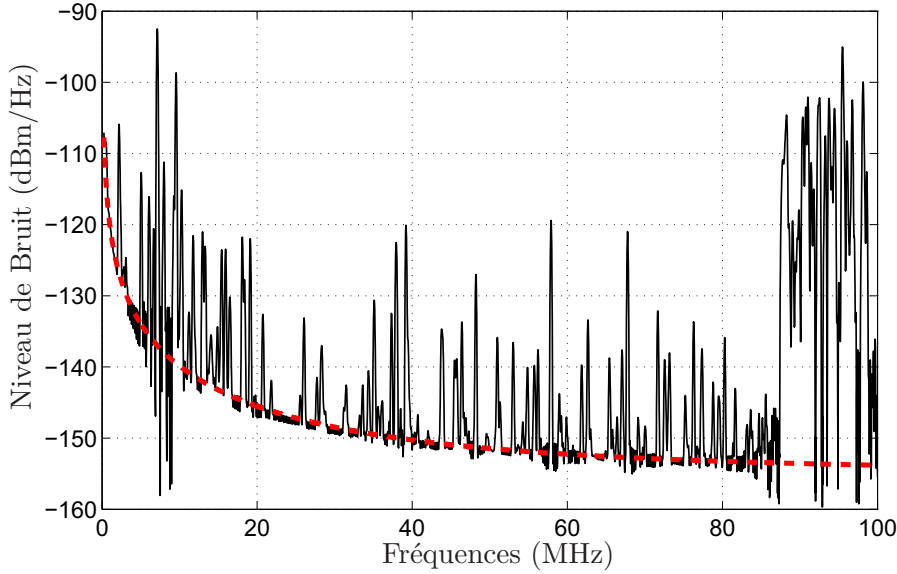


FIGURE 1.8 – Modèle de bruit stationnaire sur ligne d'énergie

en $1/f^2$ [1]. Ce résultat a permis de mettre en place un modèle de bruit stationnaire. La figure 1.8 présente un exemple de modèle de bruit stationnaire. La courbe en trait pointillé donne la forme en $1/f^2$ que suit le modèle et représente le bruit de fond. Cette courbe est définie par la fonction

$$f \mapsto \frac{1}{f^2} + 10^{\frac{-155}{10}}. \quad (1.2)$$

Il convient de noter que le modèle décrit est une première proposition conçue pour répondre aux observations expérimentales menées dans le projet OMEGA. Ce premier modèle sera utilisé pour évaluer les performances des algorithmes proposés.

1.5.2.2 Bruit de fond

Il est présent sur les lignes électriques et possède une densité spectrale de puissance relativement basse et décroissante avec la fréquence. Ce type de bruit résulte de la superposition d'une grande variété de sources de bruit de faible intensité présentes dans l'environnement des lignes. Son niveau de puissance varie à l'échelle des minutes voire des heures. Par opposition au bruit blanc qui possède une densité spectrale de puissance uniforme, le bruit de fond rencontré ici est un bruit coloré qui affiche une nette dépendance en fréquence principalement dans la partie basse du spectre. Au delà de 40 MHz, cette dépendance s'avère négligeable, et l'on peut considérer que la DSP devient plate. Le niveau moyen relevé par mesure est établi entre -155 et -145 dBm/Hz, ce qui est

de 15 à 25 dB au dessus du bruit provenant de l'agitation thermique des électrons qui est égal à -174 dBm/Hz [1, 26].

1.5.2.3 Bruits à bande étroite

Ils sont le résultat de la captation par les lignes électriques des émissions de radio-diffusion. Il s'agit donc de brouilleurs persistants qui apparaissent souvent sous la forme d'un signal sinusoïdal modulé en amplitude et occupent les sous-bandes correspondant aux diffusions grandes et moyennes ondes. En fonction de leur distance, les bruits à bande étroite peuvent être de 30 à 40 dB au-dessus du bruit de fond. Leur amplitude varie lentement au cours de la journée et devient plus importante la nuit lorsque les propriétés de l'atmosphère sont les plus propices à la réflexion des ondes. Des niveaux particulièrement élevés de bruit sont remarqués dans les basses fréquences (jusqu'à 10 MHz) ainsi que dans les hautes fréquences (bande FM à partir de 87,5 MHz) (cf. figure 1.8).

1.5.2.4 Bruits impulsifs

Parmi l'ensemble des sources de bruits, c'est de loin les bruits de type impulsif qui sont les plus défavorables aux communications sur les lignes électriques. En effet, le bruit impulsif est la principale source d'interférences qui provoquent des déformations du signal, conduisant à des erreurs lors de la transmission de données. Les origines des bruits impulsifs sont multiples : interrupteurs, alimentations et plus généralement les appareils électroménagers. Plusieurs approches ont été suivies pour la caractérisation du bruit impulsif en CPL. Dans [27, 28], les modèles proposés sont basés sur la classification du bruit en fonction de différents critères : la durée, la bande passante et l'intervalle de temps entre les impulsions. Dans [24], on distingue trois types de bruits impulsifs selon qu'ils sont périodiques ou apériodiques, synchrones ou asynchrones à la fréquence principale, à savoir 50 ou 60 Hz. Les modèles proposés sont développés en considérant que le bruit impulsif est un bruit imprévisible mesuré au niveau du récepteur. Les auteurs se retrouvent alors confrontés à des milliers de modèles de bruits impulsifs dont la pluralité proviendrait très probablement de la diversité des parcours suivis par le bruit d'origine [1].

Dans le projet OMEGA, une approche innovante de modélisation est appliquée aux bruits impulsifs, qui sont désormais étudiés directement à leurs sources. Le bruit au niveau du récepteur est considéré comme le modèle de bruit à la source filtré par le canal CPL [1]. Dans la classification qui est proposée, on distingue six classes de bruits impulsifs. Les illustrations correspondantes à ces différentes classes de bruits impulsifs sont détaillées dans [1].

- Classe 1 : allumage d'interrupteur électrique et de thermostat. Les bruits impulsifs

de cette classe ont de grandes amplitudes et sont composés d'une unique impulsion et caractérisés par une succession de petites impulsions décousues. L'amplitude maximale et la durée maximale sont respectivement de 5,65 V et 14 ms. L'amplitude des impulsions montantes est écrêtée à 5,65 V ;

- Classe 2 : extinction d'interrupteur électrique et de thermostat. Les bruits impulsifs de cette classe ont de fortes amplitudes, et se caractérisent par deux courtes impulsions successives séparées par un bruit dense. L'amplitude maximale et la durée maximale sont respectivement de 5,65 V et 9 ms ;
- Classe 3 : branchement d'une prise électrique. Les bruits impulsifs de cette classe ont de fortes amplitudes, et se caractérisent par deux impulsions successives très proches. L'amplitude maximale et la durée maximale sont respectivement de 5,65 V et 11 ms ;
- Classe 4 : débranchement d'une prise électrique. Les bruits impulsifs de cette classe sont plus faibles que ceux des premières classes. Comme dans le cas du branchement d'une prise, ces bruits se caractérisent par deux impulsions successives mais encore plus proches. L'amplitude maximale et la durée maximale sont respectivement de 1,41 V et 0,8 ms ;
- Classe 5 : démarrage d'un moteur électrique. Les bruits impulsifs de cette classe ont de fortes amplitudes, sont dispersés et très longs. L'amplitude maximale et la durée maximale sont respectivement de 5,65 V et 47 ms ;
- Classe 6 : divers bruits faibles. Cette classe comprend une diversité de bruits impulsifs faibles. Ces bruits sont caractérisés par leur très faible amplitude et apparaissent généralement sous deux principales formes : des courtes impulsions isolées ou deux larges impulsions. L'amplitude maximale et la durée maximale sont respectivement de 0,25 V et 50 ms.

1.6 Conclusion

Ce premier chapitre nous a donné l'occasion de nous familiariser avec le sujet de l'étude, et de présenter l'état de l'art sur les communications par courant porteur. Nous avons présenté le projet européen OMEGA, le cadre dans lequel s'est inscrit cette thèse. L'objectif principal du groupe de travail sur le CPL est d'étudier et d'implémenter une solution CPL capable de fournir des services à la vitesse du gigabit par seconde. Pour cela, le système CPL OMEGA exploite la bande $[0; 100]$ MHz tout en respectant la limitation d'un certain masque de puissance.

Dans un second temps, nous nous sommes attardés sur la caractérisation du canal de propagation et la présentation des modèles de canaux issus de campagnes de mesure et utilisés dans la suite de l'étude. On retient alors principalement qu'il existe neuf

classes de canaux CPL avec des atténuations fréquentielles différentes dans la bande $[0; 100]$ MHz. Par ailleurs, il apparaît que la réponse du canal peut être considérée comme quasi-statique à l'échelle des communications, caractéristique qui sera largement mise à profit dans la suite du document. Enfin, l'étude du contexte de bruit a montré que les lignes électriques sont altérées par un bruit de fond de niveau faible et qu'elles sont perturbées par des brouilleurs à bande étroite et par différents types de bruits impulsifs imprévisibles.

Chapitre 2

Spécifications du système et l'allocation des ressources

Ce deuxième chapitre est dédié à la présentation des spécifications du système mis en œuvre. Une vue globale de l'allocation des ressources y est aussi apportée. L'étude qui est menée dans cette thèse utilise comme trame de fond la mise en place de communications à haut débit sur lignes d'énergie en combinant les techniques de modulations multiporteuses et d'étalement de spectre. Nous allons commencer par rappeler les techniques de modulations multiporteuses et d'étalement de spectre, qui constituent les bases des techniques qui nous intéressent. Les techniques de modulations multiporteuses comme l'OFDM ont été retenues pour assurer des débits de transmission élevés dans les milieux très sélectifs en fréquence. Afin d'améliorer les performances du système, nous utiliserons une technique combinant l'OFDM et l'étalement de spectre. Comme nous allons le voir, cette combinaison permet d'accroître l'exploitation de la ressource disponible en puissance pour augmenter les débits de transmission du système. Nous parlerons alors d'OFDM précodé linéairement ou LP-OFDM (*linear precoded OFDM*).

2.1 Rappels sur les techniques de modulations multiporteuses et d'étalement de spectre

Le principe de modulation multiporteuse repose sur la parallélisation en fréquence de l'information à transmettre. Les données, de débit initial $1/T_d$ élevé, sont réparties sur plusieurs sous-canaux fréquentiels élémentaires modulés à bas débit, les sous-porteuses. Si N est le nombre de sous-porteuses utilisées, les symboles transmis par chacune d'elles ont une durée $T_s = NT_d$, si bien que le débit global du signal obtenu reste identique à celui d'une modulation monoporteuse avec une même occupation spectrale [29]. Le signal est composé de N sous-porteuses formant une base orthogonale en fréquence et

à occupation spectrale minimale. Dans le domaine fréquentiel, les distorsions du signal introduites par le canal sont de cette manière limitées puisque chaque sous-bande devient suffisamment étroite pour considérer la réponse du canal comme plate localement. Dans le domaine temporel, le signal obtenu se décompose en symboles de durée T_s résultant de la superposition de N signaux sinusoïdaux de fréquences différentes. En augmentant suffisamment le nombre de sous-porteuses, la durée des symboles peut être rendue bien supérieure à l'étalement des retards de la réponse impulsionnelle du canal, ce qui tend à minimiser les effets d'interférence entre symboles. En pratique, une augmentation trop importante ne peut cependant pas être envisagée en raison des limites imposées par le temps de cohérence du canal. Par conséquent, on a le plus souvent recours à l'utilisation d'un intervalle de garde. L'intervalle de garde constitue un laps de temps pendant lequel aucune donnée utile n'est émise. Inséré en tant que préfixe de chaque symbole OFDM, son rôle est d'absorber l'interférence inter-symbole résiduelle, pour peu que sa durée soit choisie supérieure ou égale à l'étalement maximal des retards de la réponse impulsionnelle du canal. Les modulations multiporteuses ont été largement étudiées dans la littérature et apparaissent dans plusieurs travaux de thèse de notre laboratoire [3, 4, 30]. Pour notre part, on suppose que le signal est adapté au canal de transmission : l'intervalle de garde, le nombre de sous-porteuses et l'espace inter-porteuse sont sélectionnés pour parfaitement absorber les trajets multiples causés par le canal et pour limiter la perte d'efficacité spectrale engendrée par l'intervalle de garde. Par ailleurs, l'utilisation de la modulation adaptative OFDM permet de profiter pleinement de la connaissance du canal à l'émission et d'appliquer le principe du *water-filling* via les algorithmes d'allocation de l'information ou *bit loading*. L'information est alors distribuée sur les sous-porteuses du signal OFDM en fonction de la valeur du rapport entre la puissance du signal et celle du bruit propre à chaque sous-porteuse. Cependant, avec une densité spectrale de puissance (DSP) du signal émis limitée comme dans le contexte CPL, une perte de débit liée à la quantification des ordres de modulation est inévitable. Cette perte peut être limitée en regroupant simplement les sous-porteuses à l'aide de séquences de précodage [31]. Ce regroupement revient à combiner les techniques de modulations multiporteuses et d'étalement de spectre.

Le principe de l'étalement de spectre se justifie par la relation de Shannon qui décrit la dépendance entre la quantité maximale d'information C qu'il est possible de transmettre sans erreur sur un canal donné, la densité spectrale du signal E_s , la largeur de bande W de ce canal perturbé par un bruit blanc additif gaussien de densité spectrale de puissance N_0 . Cette relation bien connue s'écrit

$$C = W \log_2 \left(1 + \frac{E_s}{N_0} \right) . \quad (2.1)$$

D'après cette relation, on déduit que la capacité d'un canal C en bit par seconde nécessite une densité spectrale de puissance à l'émission d'autant plus faible que la bande utilisée est large. C'est l'idée maîtresse des systèmes à étalement de spectre, pour lesquels le signal est émis sur une bande fréquentielle largement supérieure à celle du signal utile et avec une densité spectrale de puissance réduite, souvent inférieure à celle du bruit de fond. Les différents procédés permettant de réaliser l'opération d'étalement ont recours pour la plupart à l'utilisation de séquences pseudo-aléatoires [32]. Parmi tous ces procédés, nous allons nous intéresser plus particulièrement à l'étalement par séquence directe qui est employé, entre autres, dans les systèmes multiporteuses à spectre étalé. Cette technique consiste à multiplier le message d'information numérique par une séquence pseudo-aléatoire dont le débit numérique est supérieur à celui du message. De manière générale, chaque élément du code, appelé bribe ou *chip*, prend sa valeur dans un alphabet fini de valeurs complexes, même si la plupart des codes utilisés sont des codes binaires de valeurs $\{-1, +1\}$.

2.2 Techniques d'accès multiple

Dans l'optique d'une meilleure répartition des ressources entre plusieurs utilisateurs d'un même système de communication, plusieurs techniques d'accès multiple peuvent être utilisées. Les principales techniques couramment mises en œuvre sont le FDMA (*frequency-division multiple access*), le TDMA (*time-division multiple access*), le CDMA (*code-division multiple access*) et l'OFDMA (*orthogonal frequency-division multiple access*) [15]. Pour les systèmes où la réponse du canal est inconnue à l'émission, chaque jeu d'éléments (slots temporels, sous-porteuses, codes) est attribué de façon figée et arbitraire à chacun des utilisateurs. Au contraire, lorsque la réponse du canal est une donnée connue de l'émetteur, celui-ci peut envisager d'adapter le partage des ressources entre les utilisateurs en fonction du canal de chacun. On aboutit de cette manière à une utilisation bien plus efficace des ressources temps-fréquence-code [33]. Pour plus de détails sur ces techniques, le lecteur pourra donc se reporter aux références [8, 34].

La technique FDMA consiste à diviser le spectre de fréquence en sous-canaux individuels qui seront attribués aux différents utilisateurs. Cette technique peut être facilement mise en œuvre puisque les utilisateurs peuvent être séparés au niveau du récepteur à l'aide d'un simple filtre.

La technique TDMA permet à plusieurs utilisateurs de partager la même fréquence en divisant le signal en différents intervalles de temps très courts (slots temporels). Les utilisateurs transmettent successivement, l'un après l'autre, chacun utilisant son propre intervalle de temps. Cette technique est généralement plus difficile à appliquer que le FDMA car elle nécessite la synchronisation du réseau.

Avec la technique CDMA, plusieurs utilisateurs sont en mesure de transmettre des données simultanément sur la même bande de fréquences. Les signaux des utilisateurs sont distingués par des codes pseudo-aléatoires différents qui doivent être connus au niveau du récepteur. Une sélection judicieuse des codes pseudo-aléatoires avec de bonnes propriétés d'inter et d'auto-corrélation est nécessaire pour obtenir des performances optimales pour le système.

La technique OFDMA quant à elle est la combinaison des modulations adaptatives OFDM avec la technique FDMA. Elle consiste tout simplement à exploiter la parallélisation fréquentielle de l'OFDM pour transmettre des messages appartenant à des utilisateurs différents. L'accès multiple est réalisé en attribuant des sous-ensembles de sous-porteuses aux différents utilisateurs.

2.3 Système LP-OFDM

Depuis leur développement au début des années 1990, les techniques combinant l'OFDM et l'étalement de spectre ont fait l'objet d'un grand nombre de travaux qui ont permis d'en avoir aujourd'hui une connaissance détaillée. L'utilisation conjointe de l'OFDM et de l'étalement de spectre peut donner lieu à un grand nombre de variantes, regroupées sous l'appellation générique MC-SS (multicarrier spread spectrum). L'analyse de ces variantes a été réalisée dans des précédentes thèses (voir [3, 4]) et a conduit au choix du système LP-OFDM appelé aussi SS-MC-MA (*spread spectrum multi carrier multiple access*) dans un contexte radiomobile. Son intérêt est de réduire la complexité de l'estimation de canal sur les communications en voie montante d'un réseau point-multipoint [35]. En effet, une même sous-porteuse n'est affectée que par un seul canal relatif à un seul utilisateur, alors que dans un système MC-CDMA (*multi carrier CDMA*), une sous-porteuse est affectée par les différents canaux des différents utilisateurs.

Comme nous l'avons vu plus haut, l'utilisation des modulations adaptatives OFDM dans un contexte où la DSP du signal émis est limitée entraîne une perte de débit liée à la quantification des ordres de modulation. Cette perte de débit est minimisée en transmettant l'information non plus sur des sous-porteuses isolées mais sur des groupes de sous-porteuses reliées par des séquences de précodage linéaire. Par la suite, ces groupes de sous-porteuses seront appelés *blocs* de sous-porteuses. Dans le cas du système LP-OFDM, l'étalement des données est effectué avant l'opération de transformation de Fourier. Ainsi, le signal généré est un signal à porteuses multiples et il hérite en cela des propriétés de l'OFDM. Ce système réalise un changement de base du système OFDM, les symboles sont cette fois portés par les séquences de précodage et non par les sous-porteuses. La technique de précodage linéaire permet ainsi de combiner les énergies des sous-porteuses sans changer la forme d'onde initiale du signal. Contrairement au cas

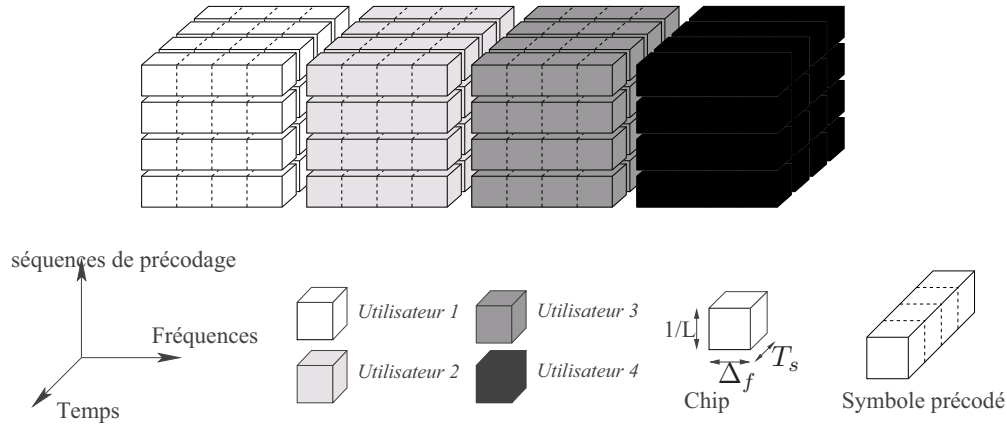


FIGURE 2.1 – Représentation schématique du système LP-OFDM dans un contexte multi-utilisateur.

CDMA, la composante de précodage n'est pas mise en œuvre pour réaliser un accès multiple à proprement parler car toutes les séquences de précodage sont utilisées pour la même liaison point-à-point dans un contexte mono-utilisateur [36].

Dans un contexte multi-utilisateur, le système LP-OFDM peut être vu comme une simple extension des techniques OFDMA avec ajout de la composante de précodage linéaire selon l'axe des fréquences, comme l'indique la figure 2.1. Cette figure donne une représentation schématique du système LP-OFDM pour quatre utilisateurs. L'accès multiple est alors fourni par la dimension fréquentielle. Les séquences de précodage sont utilisées pour multiplexer les différents symboles de données d'un même utilisateur. Il est à noter que les sous-porteuses dans un même bloc ne sont pas forcément adjacentes, comme sur la figure 2.1.

2.3.1 Expression des signaux LP-OFDM

En considérant l'émission d'une trame de N_s symboles LP-OFDM successifs, l'équation générique permettant d'exprimer le signal LP-OFDM en sortie du modulateur OFDM peut s'écrire

$$\underbrace{\mathbf{S}}_{N \times N_s} = \underbrace{\mathbf{F}^H}_{N \times N} \underbrace{\mathbf{D}}_{N \times N} \underbrace{\mathbf{C}}_{N \times N_L} \underbrace{\mathbf{X}}_{N_L \times N_s}, \quad (2.2)$$

où $\mathbf{S}$ est la matrice des N_s symboles LP-OFDM émis, composés chacun de N échantillons temporels. $\mathbf{X}$ est la matrice des symboles de modulations à transmettre et $\mathbf{F}$ est la matrice de Fourier qui réalise l'opération de transformé de Fourier. $\mathbf{D}$ est une matrice de distribution utilisée pour répartir les données sur la grille fréquentielle. C'est donc la matrice qui définit le *chip mapping*. Il s'agit d'une matrice de permutation, c'est-à-dire

qu'un élément $d_{i,j}$ de cette matrice vaut 1 si l'élément de la j^{e} colonne de la matrice issue du produit $\mathbf{C}\mathbf{X}$ doit être émis sur la i^{e} sous-porteuse [3]. Enfin, $\mathbf{C}$ est la matrice de précodage linéaire appliquée à $\mathbf{X}$ qui précode les N_L symboles complexes des N_s symboles LP-OFDM à transmettre sur les N sous-porteuses. Il a été prouvé que l'utilisation de codes orthogonaux issus des matrices de Hadamard permet de maximiser la capacité du système de communication [36]. Le produit $\mathbf{C}\mathbf{X}$ s'écrit

$$\mathbf{C}\mathbf{X} = \begin{pmatrix} \mathbf{C}_1 & & & \\ & \ddots & & \mathbf{0} \\ & & \mathbf{C}_b & \\ & \mathbf{0} & & \ddots \\ & & & & \mathbf{C}_B \end{pmatrix} \begin{pmatrix} \mathbf{X}_1 \\ \vdots \\ \mathbf{X}_b \\ \vdots \\ \mathbf{X}_B \end{pmatrix}, \quad (2.3)$$

où B est le nombre de blocs de sous-porteuses dans la sous-bande avec $B \times L = N$, $\mathbf{C}_b$ est la matrice contenant les N_L séquences de précodage du bloc b ($b \in [1, \dots, B]$), et enfin $\mathbf{X}_b$ est la matrice des N_s vecteurs de N_L symboles complexes transmis dans le bloc b .

Au niveau de la réception et après les opérations de démodulation OFDM, l'opération de déprécodage linéaire s'effectue tout simplement en appliquant la même matrice de précodage $\mathbf{C}$ qu'à l'émission. Elle permet de récupérer les symboles de données d'un symbole LP-OFDM à partir des L chips de ce même symbole qui ont été transmis sur les emplacements de la trame OFDM définis par l'opération de *chip mapping*. Le signal reçu après l'opération de déprécodage linéaire s'écrit

$$\mathbf{Y} = \mathbf{C}^H \mathbf{G} \mathbf{H} \mathbf{C} \mathbf{X} + \mathbf{C}^H \mathbf{G} \mathbf{Z}, \quad (2.4)$$

où $\mathbf{H}$ et $\mathbf{G}$ sont respectivement les matrices des coefficients du canal et de correction du canal, constituées d'éléments diagonaux respectifs h_n et g_n avec $n \in [1, \dots, N]$. $\mathbf{Z}$ le vecteur des échantillons du bruit blanc additif gaussien complexe.

2.3.2 Choix de la technique d'égalisation

L'ensemble des chips transmis simultanément sur une même sous-porteuse subit les mêmes distorsions h_n introduites par le canal. L'opération d'égalisation consiste à corriger les effets du canal de transmission afin de restaurer les données transmises et de réduire les interférences entre les symboles. Différentes techniques d'égalisation peuvent être envisagées et sont notamment décrites et étudiées dans [6]. Afin d'obtenir des formules analytiques pour la répartition des bits et des énergies, nous nous limitons à des

structures d'égalisation simples nécessitant une simple multiplication complexe par un coefficient g_n pour chaque sous-porteuse. Ces méthodes de détection sont qualifiées de mono-utilisateur et peuvent être envisagées dans un système LP-OFDM. On parle alors de détection mono-utilisateur étant donné que le détecteur ne nécessite la connaissance des séquences de précodage que d'un seul utilisateur dont on veut restaurer les données.

Dans notre étude, nous nous limiterons à l'utilisation des critères de distorsion crête (*zero forcing* (ZF) en anglais) et d'erreur quadratique moyenne (EQM). Il est connu que le critère EQM offre de meilleures performances que le critère ZF. Cependant, du fait de la complexité des formules de capacité, l'optimisation des paramètres ne permet pas de dégager une structure algorithmique simple contrairement à la détection ZF [6]. Dans des précédentes thèses [3, 4, 30], il a été choisi de travailler uniquement avec le critère ZF. Comme nous allons le voir dans le chapitre suivant, la répartition dans les blocs de sous-porteuses ayant des amplitudes équivalentes permet de réduire les distorsions entre ces sous-porteuses. Avec les canaux utilisés dans ces travaux, la différence de performance entre la détection ZF et la détection EQM est très faible. Par contre dans un contexte multi-utilisateur (avec accès en fréquence), il est difficile de trouver des politiques de répartition des sous-porteuses qui minimisent les distorsions entre les sous-porteuses d'un même bloc. Dans cette thèse, nous étudierons le gain en débit que peut apporter l'utilisation du critère EQM dans un contexte multi-utilisateur.

2.3.2.1 Distorsion crête

Le critère ZF, ou critère de distorsion-crête, permet de supprimer les interférences en réception et consiste à inverser les coefficients du canal de manière à compenser totalement l'atténuation qu'ils introduisent et ainsi annuler totalement l'interférence entre les séquences de précodage mais au prix d'une augmentation du bruit [37]. Ainsi,

$$g_n = \frac{1}{h_n} = \frac{h_n^*}{|h_n|^2}. \quad (2.5)$$

2.3.2.2 Erreur quadratique moyenne

Le critère EQM réalise un compromis entre la minimisation des interférences et l'augmentation du facteur de bruit. Le coefficient d'égalisation g_n est calculé dans le but de minimiser l'erreur quadratique moyenne entre le signal émis et le signal reçu obtenu après l'égalisation [38], et s'écrit

$$g_n = \frac{h_n^*}{|h_n|^2 + \frac{N_0}{E_s}}. \quad (2.6)$$

Cette technique nécessite l'estimation du rapport signal à bruit (RSB) pour chaque sous-porteuse, induisant une complexité supplémentaire. Afin de s'affranchir de cette estimation, il est possible d'appliquer un coefficient ϵ fixé en fonction du point limite de fonctionnement du système. Une technique hybride entre la technique ZF et EQM existe, elle est appelée égalisation TORC (threshold orthogonality restoring combining) et propose un compromis entre les deux techniques d'égalisation [39, 40].

2.4 Allocation des ressources

La gestion des ressources au sein d'un système de communications s'impose comme une question de premier plan dès lors que l'on cherche à optimiser ses performances. Tout système doit en effet composer avec une certaine somme de limitations physiques et technologiques additionnées à des contraintes supplémentaires imposées par la qualité de service [3]. A l'inverse, on dénombre un certain nombre de degrés de liberté qui peuvent servir à configurer le système dans le respect de ces limitations et contraintes. L'adaptation de ces paramètres libres doit prendre en compte les caractéristiques du canal, et les caractéristiques de la liaison. Dans l'hypothèse d'une connaissance parfaite de la réponse du canal, on peut alors optimiser ces paramètres libres du système en fonction du comportement du canal et des contraintes afin d'assurer un certain niveau de qualité de service. Les ressources qui doivent être gérées dépendent du système mis en œuvre et peuvent être la fréquence, le temps, la puissance d'émission, etc. En effet, tout système de communication a une bande de fréquences limitée et travaille sous des contraintes de puissance fixées par les normes et les contraintes technologiques. L'utilisation des modulations adaptatives dans un système à porteuses multiples, rend possible l'attribution d'un nombre différent de bits et d'une puissance différente à chaque sous-porteuse ou séquence de précodage du système. Dans un contexte multi-utilisateur, nous avons évoqué au paragraphe 2.2, le fait que l'accès au *medium* pour chaque utilisateur peut être mis en œuvre dans le domaine temporel, le domaine fréquentiel ou le domaine des codes en utilisant les techniques TDMA, FDMA, CDMA et OFDMA. L'allocation des bits et des puissances peut ensuite être effectuée pour chaque utilisateur du système sur son propre jeu d'éléments. Nous verrons que ces principes peuvent être appliqués au système OFDM ainsi qu'au système LP-OFDM proposé. En définitive, dans un contexte comme le nôtre, les algorithmes d'allocation doivent à la fois combiner le partage dynamique des ressources entre utilisateurs et l'allocation dynamique des ressources individuelles de chaque utilisateur.

2.4.1 Modulations adaptatives

2.4.1.1 Capacité et débit associés à un canal non dispersif

Pour un canal non dispersif d'atténuation α donnée, encore qualifié de canal uniforme, perturbé par la seule présence du bruit blanc additif gaussien de densité de puissance unilatérale N_0 , les deux grandeurs, bande de fréquence W et densité spectrale de puissance d'émission E_s , fixent à elles seules la capacité envisageable, en vertu de la relation de Shannon. On rappelle que cette relation s'écrit,

$$C = W \log_2 \left(1 + \frac{\alpha E_s}{N_0} \right) = W \log_2 (1 + \text{RSB}) , \quad (2.7)$$

où C est donnée en bit par seconde (b/s). La capacité du canal croît donc avec le RSB de façon logarithmique et le débit binaire que peut transmettre un canal ne peut pas dépasser la limite donnée par la capacité C . Pour atteindre cette limite, un code théorique d'une complexité infinie, dont le délai de codage/décodage est lui aussi infini, doit être utilisé. Bien évidemment, les codes mis en œuvre dans les systèmes réels sont sous-optimaux par rapport à ce code théorique infini et ne permettent pas d'atteindre la capacité C du canal. On introduit alors un paramètre Γ , connu sous le nom de marge de RSB (*SNR gap*), et qui constitue une mesure de la performance relative d'un système utilisant un schéma de codage donné, par rapport à la capacité du canal. Cette marge de RSB dépend du schéma de codage choisi, de l'ordre de modulation et de la probabilité d'erreur cible. Le nombre r de bits par symbole pouvant être véhiculés est alors donné par [41]

$$r = \log_2 \left(1 + \frac{1}{\Gamma} \frac{\alpha E_s}{N_0} \right) . \quad (2.8)$$

Cette équation montre concrètement que le débit maximal pouvant être atteint sur le canal est inférieur à sa capacité. À noter que Γ est parfois appelé le RSB normalisé, car on peut en effet écrire $\Gamma = \text{RSB}/(2^r - 1)$.

Finalement, avec la connaissance de la réponse du canal de transmission à l'émission et de la marge de RSB du schéma de codage utilisé, on peut déterminer précisément la quantité d'information, en bits par symbole, qu'il est possible de transmettre sur le canal considéré pour une énergie par symbole E_s fixée.

2.4.1.2 Marge de RSB des modulations MAQ

Dans notre étude, les modulations numériques de type MAQ vont être utilisées pour transmettre les données sur les différentes sous-porteuses. Il est donc nécessaire de quantifier la valeur de Γ examinant l'efficacité du schéma de transmission relatif aux mo-

dulations MAQ. La marge de RSB peut être déterminée à partir des courbes de taux d'erreur des modulations, pour un point de fonctionnement donné, à savoir un taux d'erreur symbole (TES) ou un taux d'erreur binaire (TEB). Suivant l'analyse classique de la marge de RSB [41], Γ a une valeur constante pour tous les ordres de modulations MAQ non codées et pour un TES cible fixé. Cette marge Γ est donc approchée par

$$\Gamma \approx \frac{1}{3} \left(Q^{-1} \left(\frac{P_{\text{es}}}{4} \right) \right)^2, \quad (2.9)$$

où Q^{-1} est l'inverse de la fonction Q définie par

$$Q(a) = \frac{1}{\sqrt{2\pi}} \int_a^{+\infty} e^{-\frac{x^2}{2}} dx \quad (2.10)$$

On note que cette relation ne fait pas intervenir le nombre r de bits des constellations MAQ. En revanche, pour une probabilité d'erreur binaire P_{eb} fixée plutôt que pour une probabilité d'erreur symbole P_{es} , Γ devient variable. Classiquement, l'approximation de Γ pour un TEB cible donné se fait en utilisant le fait que $P_{\text{eb}} \cong P_{\text{es}}/r$, expression qui fait réapparaître le nombre r de bits relatif à la taille des constellations MAQ. La marge peut alors être réécrite dans ce cas, en fonction de r

$$\Gamma(r) \approx \frac{1}{3} \left(Q^{-1} \left(\frac{r P_{\text{eb}}}{4} \right) \right)^2, \quad (2.11)$$

Par contre, cette approximation devient imprécise pour les faibles valeurs de RSB. Des expressions exactes et générales du TEB et du TES ont été développées dans [42] pour des schémas de modulations MAQ à une ou deux dimensions. Ces expressions, assez complexes, permettent de tracer les courbes de taux d'erreur d'où peuvent être déduites les différentes valeurs de Γ pour chaque ordre de modulation. La figure 2.2 donne la capacité effective des modulations MAQ non codées de type carré pour une probabilité d'erreur symbole P_{es} et binaire P_{eb} cible de 10^{-3} . La marge de RSB Γ de chaque modulation est obtenue en évaluant la distance entre les points de fonctionnement et la courbe de capacité théorique de Shannon. On constate que les approximations sont bonnes pour les ordres pairs de modulation. Pour les ordres impairs de modulations, la différence entre la courbe analytique et la courbe approchée est d'environ 1 dB. Dans la suite, nous nous limiterons à l'utilisation des approximations de Γ pour un TES ou un TEB fixé. On note bien une valeur constante de la marge pour un TES cible fixé, soit $\Gamma = 6$ dB. Par contre, pour un TEB cible fixé, $\Gamma = 6$ dB pour une modulation MAQ-2 et $\Gamma = 4.2$ dB pour une modulation MAQ-1024, et cette différence (environ 2 dB) apporte une différence significative dans les performances des algorithmes d'allocation de ressources.

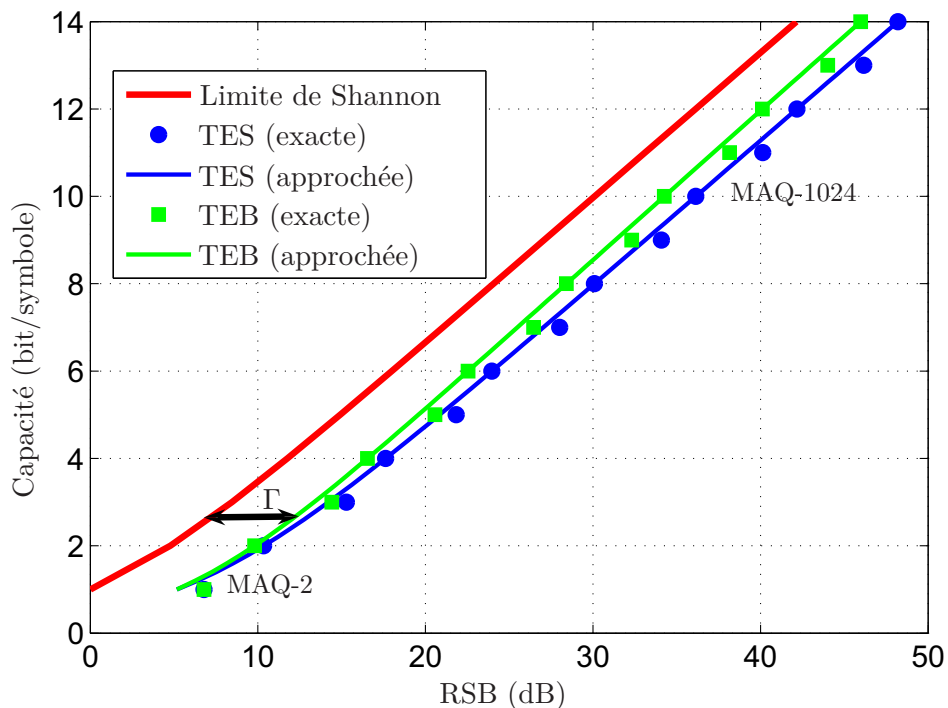


FIGURE 2.2 – Points de fonctionnement des modulations MAQ non codées pour $P_{\text{eb}} = P_{\text{es}} = 10^{-3}$.

Les avantages de l'utilisation de l'approche de TEB cible seront discutés dans le chapitre suivant. Grâce à l'introduction de la marge de RSB, il est donc possible d'adapter la modulation MAQ choisie à la réponse du canal. Il s'agit là du principe de modulation adaptative que nous mettrons en œuvre par la suite dans le cas de signaux à porteuses multiples, et en particulier avec le système LP-OFDM étudié.

2.4.2 Politiques d'optimisation : problème général

Nous venons de voir que la connaissance du canal à l'émission permet à l'émetteur de déterminer la quantité maximale d'information pouvant être transmise par le canal de propagation, à puissance d'émission fixée. On parle ici de puissance d'émission au sens large. Il peut en fait s'agir d'une limitation en puissance totale ou d'une limitation en DSP. Dans cette thèse, nous nous limiterons seulement aux formulations sous la contrainte de DSP, contrainte correspondant au masque de puissance imposé dans les systèmes de communication CPL. Selon les exigences de la qualité de service, la stratégie à adopter peut se décliner sous deux formes : la maximisation du débit de transmission sous la contrainte de puissance d'émission ou la maximisation de la marge de bruit du système sous la contrainte de débit cible fixé. Les algorithmes développés dans la suite

de ce document cherchent à maximiser le débit.

Dans les systèmes pratiques, le problème d'allocation des ressources est traité par les algorithmes d'allocation des bits et des énergies, qui sont utilisés pour répartir, entre les sous-porteuses ou les utilisateurs, le nombre total de bits et l'énergie totale disponible, dans le but de maximiser les performances du système et de satisfaire les contraintes de qualité de service. De manière générale pour un système à N sous-canaux (qui peuvent être les sous-porteuses, les séquences de précodage ou de façon générale les modes propres), nous notons $\mathcal{R}$, $\mathcal{P}$ et $\mathcal{S}$ les politiques d'allocation des ressources mises en œuvre au sein du système. $\mathcal{R}$ donne la politique de répartition des bits, $\mathcal{P}$ la politique de répartition des puissances et $\mathcal{S}$ celle des sous-porteuses, en fonction des degrés de liberté du système. On suppose dans ce contexte général que le système comporte plusieurs utilisateurs, le cas mono-utilisateur étant un cas particulier. Pour une probabilité d'erreur cible donnée, on désigne par $f(\mathcal{R}, \mathcal{P}, \mathcal{S})$ la fonction donnant le débit obtenu après paramétrage du système conformément aux politiques d'allocation $\mathcal{R}$, $\mathcal{P}$ et $\mathcal{S}$. La fonction f représente l'information mutuelle entre les variables aléatoires relatives aux données de la source (l'émetteur) et aux données du destinataire (le récepteur) [3].

Sous la contrainte d'une puissance d'émission fixée et d'une probabilité d'erreur fixée, on peut alors exprimer la stratégie de maximisation du débit comme suit,

$$\max_{\mathcal{R}, \mathcal{P}, \mathcal{S}} f(\mathcal{R}, \mathcal{P}, \mathcal{S}) \quad (2.12)$$

Suivant les politiques mises en œuvre pour $\mathcal{R}$, $\mathcal{P}$ et $\mathcal{S}$, le résultat de la maximisation du débit donne pour l'utilisateur k

$$R_k = \sum_{n=1}^N \rho_{k,n} r_{k,n}, \quad (2.13)$$

où R_k est le débit alloué à l'utilisateur k . La variable binaire $\rho_{k,n}$, résultant de la politique $\mathcal{S}$ de répartition des sous-porteuses, indique si la sous-porteuse n est allouée à l'utilisateur k ou non, et peut être définie par

$$\rho_{k,n} = \begin{cases} 1 & \text{si } n \text{ est allouée à } k \\ 0 & \text{sinon} \end{cases} \quad (2.14)$$

Le nombre $r_{k,n}$ résulte des politiques $\mathcal{S}$ et $\mathcal{P}$ de répartition des bits et des puissances, et s'écrit

$$r_{k,n} = \log_2 \left(1 + \frac{\text{rsb}_{k,n}}{\Gamma(r_{k,n})} \right), \quad (2.15)$$

où $\text{rsb}_{k,n}$ est le rapport signal à bruit sur le sous-canal n de l'utilisateur k , $\Gamma(r_{k,n})$ la marge de RSB pouvant être variable selon la probabilité d'erreur cible choisie. Cette marge de bruit ne prend pas en compte le gain de codage. La prise en compte du schéma de codage de canal dans le processus d'allocation des ressources, en LP-OFDM, a fait l'objet d'études dans une précédente thèse [4]. Dans notre étude, nous travaillerons avec des modulations non codées. Les politiques d'allocation menant à ces résultats seront détaillées dans les chapitres suivants à la fois dans un contexte mono et multi-utilisateur.

2.5 Conclusion

La présentation des spécifications du système dans ce chapitre a permis de décrire et d'expliquer le principe des modulations multiporteuses et du système LP-OFDM qui sera utilisé par la suite. Le système LP-OFDM exploite l'étalement de spectre à des fins de multiplexage des données et le FDMA pour le multiplexage des utilisateurs. De part sa structure, le LP-OFDM possède toute la souplesse nécessaire à l'adaptation dynamique des ressources entre les utilisateurs, tout en étant une solution adaptée à un environnement fortement bruité.

Ce chapitre nous a aussi permis d'introduire les différents principes exploités dans les problèmes d'optimisation des ressources, avec en particulier les notions de modulations adaptatives et de marge de rapport signal à bruit. Sous l'hypothèse de la connaissance du canal de transmission par l'émetteur, l'allocation des ressources peut être efficacement réalisée pour optimiser les performances du système et garantir un certain niveau de qualité de service.

Chapitre 3

Maximisation du débit dans un contexte mono-utilisateur

3.1 Introduction

Le chapitre précédent a permis de décrire le système LP-OFDM et le principe général de l'allocation des ressources pour les modulations multiporteuses adaptatives. Dans ce chapitre, nous allons étudier différentes stratégies de maximisation du débit pour les systèmes OFDM et LP-OFDM dans un contexte mono-utilisateur. Ces études seront étendues par la suite au contexte multi-utilisateur dans les chapitres suivants. L'objectif ici est de proposer des algorithmes d'allocation des ressources capables de fournir des hauts débits sous les contraintes de densité spectrale de puissance et de qualité de service pour le réseau CPL. Dans une première phase, nous reviendrons sur l'application des politiques de répartition des ressources pour les systèmes OFDM. Dans la seconde phase, nous nous attarderons sur l'allocation des ressources pour les systèmes LP-OFDM avec mise en œuvre de l'égalisation suivant les critères de la distorsion crête (*zero forcing* (ZF)) et de l'erreur quadratique moyenne (EQM), dans le cadre d'une optimisation du débit. Précisons que l'optimisation du système LP-OFDM avec l'utilisation d'un détecteur ZF en réception a fait l'objet d'études antérieures au sein du laboratoire dans le contexte des transmissions CPL [3]. Nous proposons d'en faire le rappel ici. Les algorithmes alors développés constituaient des résultats majeurs permettant de donner la répartition optimale des bits et des puissances dans les différents blocs de sous-porteuses. Ces résultats nous serviront de base pour le développement d'autres algorithmes, en particulier ceux reposant sur le critère EQM qui constituent une première contribution originale de cette thèse.

Dans la troisième phase, deux nouvelles contraintes de taux d'erreur binaire (TEB) sont utilisées afin d'obtenir de meilleurs débits pour les systèmes à porteuses multiples

en comparaison des solutions existantes. Les procédures d'optimisation dans les phases précédentes utilisent la contrainte classique du taux d'erreur symbole (TES). En général, la limite du taux d'erreur est imposée par les couches supérieures du réseau (les couches transport et application) et cette limite est exprimée en terme de TEB et non de TES. Les symboles de données au niveau des couches basses (MAC, PHY) étant différents des symboles au niveau des couches supérieures, il est donc plus judicieux de travailler avec une contrainte de TEB que de TES. Les études théoriques sur l'allocation des ressources utilisent classiquement une approximation de la probabilité d'erreur [43], qui entraîne des violations de la contrainte de taux d'erreur dans certains cas [44]. Tout d'abord, le problème d'allocation des ressources sous la contrainte d'un TEB crête est analysé, à savoir que le TEB cible est fixé pour chaque sous-porteuse et toutes les sous-porteuses doivent respecter la valeur de TEB donnée. Cette approche est peu différente de l'approche classique du TES, en ce sens qu'au lieu de fixer le TES, c'est plutôt le TEB cible qui est fixé pour chaque sous-porteuse. Les algorithmes sont développés pour le système OFDM et le système LP-OFDM. Dans la seconde approche, le débit du système OFDM est maximisé sous la contrainte d'un TEB moyen sur l'ensemble du symbole OFDM. De cette manière, les différentes sous-porteuses dans le symbole OFDM sont autorisées à ne pas respecter la limite de TEB fixée, la limite étant désormais imposée sur l'ensemble du symbole OFDM. Cette dernière approche est développée uniquement pour le système OFDM.

3.2 Modulations adaptatives et OFDM : état de l'art

Avant de se consacrer à l'allocation des ressources pour les systèmes LP-OFDM, nous allons rappeler les mécanismes d'application des modulations adaptatives aux systèmes OFDM. Le système adaptatif obtenu, plus connu sous le nom de DMT (*discrete multitone*) pour les communications filaires, servira de référence à notre étude, en tant que système mis en œuvre au sein des modems CPL actuels. Nous allons tout d'abord considérer que le système adaptatif OFDM est exploité en contexte mono-utilisateur point-à-point, avant d'étendre les concepts présentés au contexte multi-utilisateur point-à-multipoint dans les chapitres suivants.

Pour une liaison point-à-point, le problème d'optimisation sous contrainte de l'allocation de l'information sur les sous-porteuses est bien connu aujourd'hui. Sous l'hypothèse de connaissance parfaite du canal de transmission à l'émission, la maximisation de la capacité du réseau conduit alors à répartir judicieusement l'information à l'émission. En considérant le cas de vecteurs spéciaux gaussiens pour les symboles émis, et avec une connaissance du canal à l'émission et à la réception, l'exploitation de l'information

mutuelle permet d'écrire la capacité C_{OFDM} de la forme d'onde OFDM [3]

$$C_{\text{OFDM}} = \sum_{n=1}^N \log_2 \left(1 + \underbrace{|H_n|^2 \frac{E_n}{N_0}}_{\text{rsb}_n} \right). \quad (3.1)$$

Il apparaît que cette capacité s'exprime comme une somme de capacités d'une collection de canaux gaussiens chacun étant caractérisé par un RSB propre. Ainsi, en introduisant la marge de bruit des modulations MAQ Γ (cf. (2.9)), le débit total atteignable est la somme des débits sur l'ensemble des sous-canaux et s'écrit

$$R_{\text{OFDM}} = \sum_{n=1}^N r_n = \sum_{n=1}^N \log_2 \left(1 + \frac{1}{\Gamma} |H_n|^2 \frac{E_n}{N_0} \right). \quad (3.2)$$

Cette écriture définit la fonction f introduite à l'équation (2.12), et sur laquelle vont porter les optimisations. Dans ce paragraphe, les formules sont données pour une marge de RSB constante pour toutes les constellations, c'est-à-dire pour une probabilité d'erreur symbole cible fixée (voir paragraphe 2.4.1.2). Concernant la contrainte de puissance, celle-ci peut se traduire par une limitation de la puissance totale E_T , soit $\sum_n E_n \leq E_T$, ou par une contrainte sur la DSP, soit $E_n \leq E_{\text{DSP}}$, où E_n est le niveau maximal d'énergie autorisée sur la sous-porteuse n . La contrainte de puissance totale a été le sujet d'un grand nombre d'études dans la littérature. C'est sous cette contrainte que la solution bien connue du *water-filling* a notamment vu le jour, par optimisation du débit suivant la méthode des multiplicateurs de Lagrange [45, 46]. Dans le cas d'une granularité infinie c'est-à-dire où les nombres r_n de bits transmis par symbole prennent leurs valeurs dans $\mathbb{R}_+$, la technique *water-filling* fournit la solution optimale. Par contre dans le cas des systèmes réels pour lesquels ces nombres r_n sont des valeurs entières (dans $\mathbb{N}$), imposées par la granularité finie des modulations MAQ, d'autres solutions ont été proposées pour allouer les bits et les puissances sur les différentes sous-porteuses [47–49]. Dans cette étude, l'approche qui nous intéresse concerne la limitation en DSP, imposée par les normes d'émission. C'est donc cette contrainte qui sera appliquée dans le cadre des diverses optimisations menées. Dans le cas où les débits peuvent prendre des valeurs réelles positives, la solution au problème de maximisation du débit consiste à allouer la puissance maximale autorisée sur chaque sous-porteuse. La répartition des bits r_n découle de l'allocation des énergies, puisque les débits individuels transmis sur chaque sous-porteuse sont reliés de manière univoque à la répartition des puissances. Les politiques

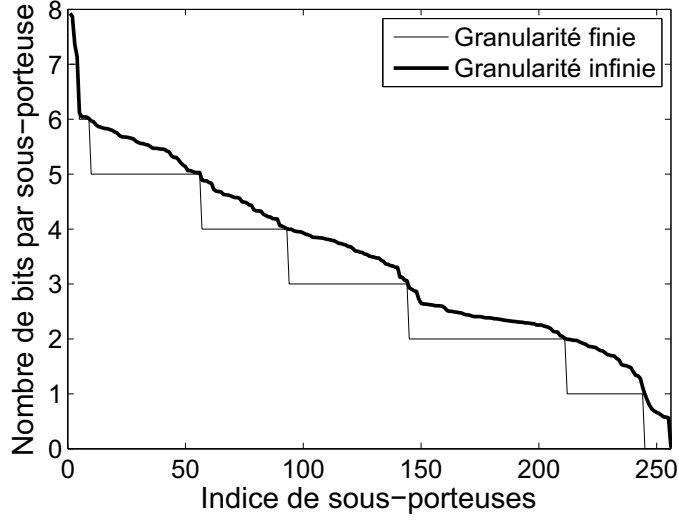


FIGURE 3.1 – Comparaison des résultats d'allocation des bits en granularité infinie et finie, dans le cadre de la maximisation du débit en OFDM.

d'allocation des énergies et des bits s'écrivent donc [36]

$$\begin{cases} \mathcal{P} : \forall n \ E_n^* = E_{\text{DSP}} , \\ \mathcal{R} : \forall n \ r_n^* = \log_2 \left(1 + \frac{1}{\Gamma} |H_n|^2 \frac{E_{\text{DSP}}}{N_0} \right) . \end{cases} \quad (3.3)$$

Dans la suite, la notation x^* sera utilisée pour marquer le fait que x est optimal. L'extension au cas de débits à valeurs entières est immédiate. La fonction valeur entière inférieure est en effet simplement appliquée sur l'allocation des bits afin d'obtenir des valeurs compatibles avec la mise en œuvre des modulations MAQ. Par conséquent, l'allocation des énergies peut être réévaluée de manière à fournir à chaque sous-porteuse la quantité d'énergie exactement nécessaire à la transmission de nouveaux débits calculés, au taux d'erreur cible envisagé au sein du système. La relation (3.3) est exploitée pour cela, si bien que les politiques d'allocation des ressources peuvent être exprimées comme suit

$$\begin{cases} \mathcal{R} : \forall n \ r_n^* = \lfloor r_n^* \rfloor , \\ \mathcal{P} : \forall n \ E_n^* = \Gamma \frac{N_0}{|H_n|^2} \left(2^{r_n^*} - 1 \right) . \end{cases} \quad (3.4)$$

où $\lfloor \cdot \rfloor$ indique que x est un paramètre relatif à un contexte de granularité finie et $\lfloor \cdot \rfloor$ est la fonction partie entière. Notons que l'énergie dépensée sur chaque sous-porteuse est telle que $E_n \leq E_{\text{DSP}}$.

La figure 3.1 permet d'illustrer la différence entre les débits résultant des politiques d'allocation avec une granularité finie et infinie. Le modèle de référence du canal CPL

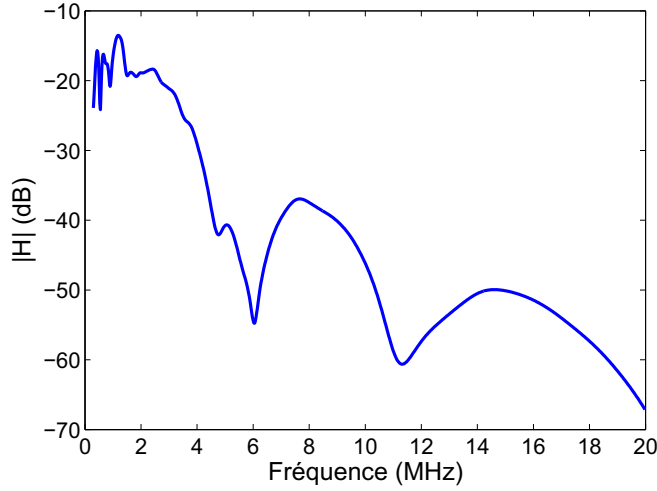


FIGURE 3.2 – Fonction de transfert du modèle de référence du canal CPL proposé par Zimmermann.

à multiples trajets proposé par Philipps [16] et Zimmermann [17] est utilisé dans cette simulation. La figure 3.2 présente ce modèle de référence qui est la réponse fréquentielle d'un câble de 100 mètres comportant 15 trajets, et est donné par

$$H(f) = \sum_{i=1}^{15} g_i \cdot e^{(a_0 + a_1 f^k)^{d_i}} \cdot e^{-2j\pi f(\tau_i)} . \quad (3.5)$$

où τ_i est le retard de propagation du trajet i , g_i est le facteur de pondération du trajet i , d_i est la distance en mètre du trajet i et $\{a_0; a_1; k\}$ sont les paramètres ajustables du modèle. Les paramètres du modèle de référence sont donnés dans [17].

Pour faciliter la lecture graphique, les sous-porteuses ont été classées par ordre décroissant de leur amplitude et seulement 256 sous-porteuses sont représentées. Naturellement, lorsque la granularité est finie, la courbe obtenue est en escalier, rendant ainsi compte de l'utilisation de la fonction valeur entière. En présence d'une limitation en DSP, le système OFDM n'est donc pas capable d'exploiter la totalité de l'énergie mise à sa disposition. Dans l'exemple choisi ici, on remarque d'ailleurs que les 10 dernières sous-porteuses ne sont pas du tout exploitables par le système, leur débit respectif dans $\mathbb{R}$ étant inférieure à 1 bit par symbole.

3.3 Modulations adaptatives en LP-OFDM

Dans ce paragraphe, nous allons chercher à appliquer les principes d'allocation dynamique des ressources au système LP-OFDM. L'objectif est de mettre au point des politiques d'allocation des ressources similaires à celles développées dans le paragraphe

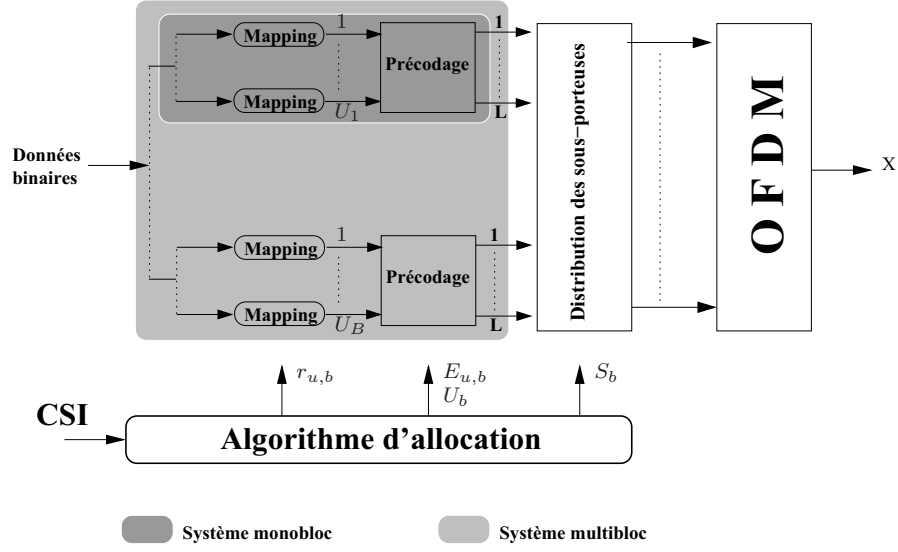


FIGURE 3.3 – Représentation schématique du paramétrage du signal LP-OFDM par les algorithmes d'allocation dynamique des ressources.

précédent. D'abord, nous allons revenir sur la description du système LP-OFDM en s'attachant à définir ses degrés de liberté et à déterminer sa capacité. Ensuite, les algorithmes d'allocation seront développés pour le système LP-OFDM avec une mise en œuvre de l'égalisation suivant les critères ZF et EQM, dans le cadre d'une optimisation du débit. Enfin, nous nous attarderons à évaluer les performances des algorithmes mis en œuvre, et en particulier à comparer les résultats à ceux obtenus avec un système OFDM. Les expressions analytiques obtenues pour le critère ZF sont tirées de [50] et servent de base pour le développement des procédures d'optimisation pour le critère EQM.

3.3.1 Système et degrés de liberté

Pour mener à bien l'étude sur l'allocation des ressources, commençons par nous concentrer sur la structure de l'émetteur LP-OFDM. La figure 3.3 propose une représentation de cet émetteur, mettant en évidence les différentes opérations intervenant lors du paramétrage d'un signal LP-OFDM. Il s'agit là d'un système multibloc pour lequel les données d'un même utilisateur sont précodées linéairement à l'aide de séquences de précodage. Comme cela est mis en évidence sur la figure, l'organe d'allocation dynamique des ressources vient configurer le système en fonction de la connaissance de l'état du canal (CSI, *channel state information*). Cette connaissance sera supposée parfaite dans la suite de l'étude. Elle comprend non seulement la connaissance de la réponse des différents canaux mais aussi la connaissance du niveau des bruits et brouilleurs.

Une représentation schématique d'un exemple de répartition des séquences de précodage en fonction des sous-porteuses est donnée figure 3.4. Le principe d'allocation en

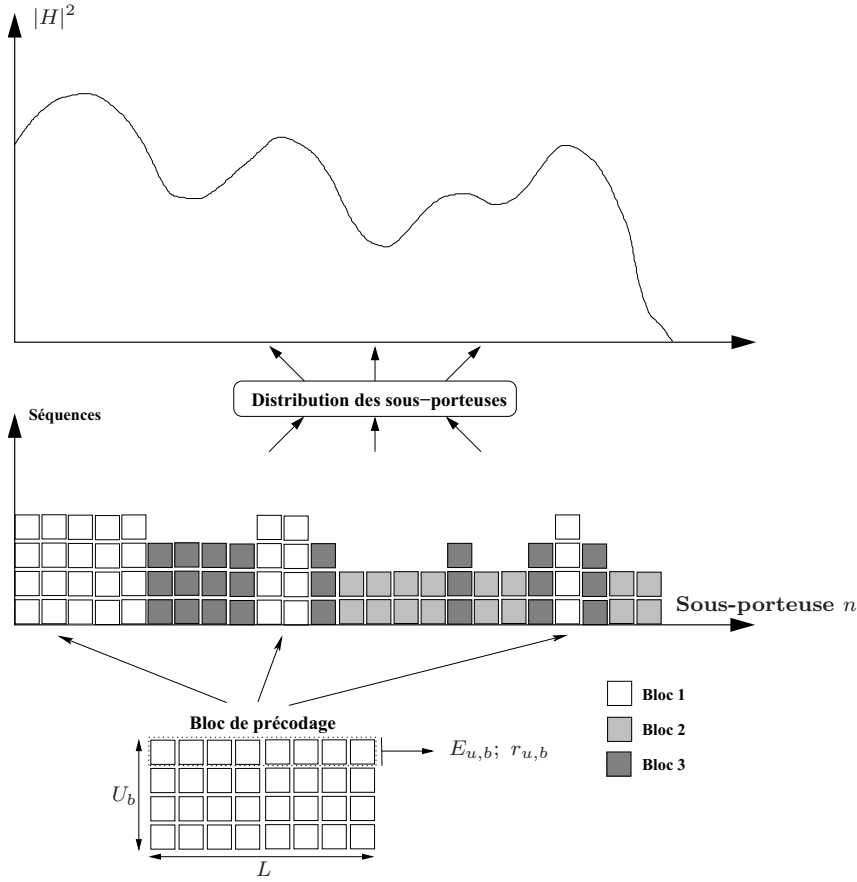


FIGURE 3.4 – Principe d'allocation adaptative des ressources dans un système LP-OFDM mono-utilisateur.

mode LP-OFDM est alors clairement illustré. Sur ce schéma on retrouve les différents paramètres que l'organe d'allocation a pour mission d'adapter en fonction du canal et des contraintes du système. En considérant que la longueur L des séquences de précodage est constante d'un bloc de sous-porteuses à l'autre, on peut lister les différentes données paramétrables, à savoir

- le nombre de séquences U_b utilisées au sein de chaque bloc b de sous-porteuses ;
- la puissance $E_{u,b}$ attribuée à la u^e séquence du b^e bloc ;
- l'ordre de modulation $r_{u,b}$ associé à cette même séquence ;
- les jeux S_b d'indices de sous-porteuses relatifs à chaque bloc b .

Les séquences de précodage utilisées sont des séquences orthogonales si bien que le nombre de séquences exploitables par bloc de sous-porteuses est tel que $U_b \leq L$, $\forall b$. Si on ne se restreint pas aux seules constructions de Sylvester, la longueur L de ces séquences peut valoir n'importe quel multiple de 4 [51, 52]. Par ailleurs, on rappelle que les modulations employées sont des MAQ dont l'ordre peut varier de 1 à 14, et donc

$r_{u,b} \in [1, \dots, 14]$. La limitation de puissance en DSP s'écrit quant à elle,

$$\sum_{u=1}^{U_b} E_{u,b} \leq E_{\text{DSP}}, \quad (3.6)$$

où E_{DSP} est le niveau de la DSP limite. On rappelle enfin que le nombre B de blocs de sous-porteuses qu'il est possible de former dans le cas d'une longueur de précodage L est tel que $B = \left\lfloor \frac{N}{L} \right\rfloor$.

Comme indiqué sur la figure 3.3, le système LP-OFDM considéré peut être divisé en deux sous systèmes de complexité croissante. Le premier sous système, dit système monobloc, constitue l'élément de base du système global. Ce système est équivalent à un système MC-CDMA dans lequel un seul utilisateur se verrait attribuer tous les codes. Le deuxième sous-système, dit système multibloc, est l'extension du système monobloc au cas d'utilisation de plusieurs blocs de sous-porteuses. Nous allons tout d'abord développer les différentes équations pour le système élémentaire monobloc avant de les étendre par la suite au système multibloc. Puisque ce système n'exploite qu'un seul bloc de sous-porteuses, l'indice de bloc b sera omis dans ces développements et on considèrera les indices de sous-porteuses l tels que $l \in [1, \dots, L]$.

3.3.2 Information mutuelle

Comme dans l'étude effectuée dans le paragraphe 2.4.2 sur l'allocation des ressources en général, nous avons besoin ici d'exprimer la fonction f permettant de donner le débit du système considéré en fonction de ses paramètres. Les procédures d'optimisation pourront alors être entreprises à partir de cette expression. C'est donc l'information mutuelle qu'il nous faut déterminer. On rappelle qu'on considère le cas de récepteurs dits mono-utilisateurs, c'est-à-dire que chaque séquence est traitée de manière indépendante. Le signal reçu après l'opération de déprécodage linéaire, pour un système monobloc, s'écrit

$$\mathbf{Y} = \mathbf{C}^H \mathbf{G} \mathbf{H} \mathbf{C} \mathbf{X} + \mathbf{C}^H \mathbf{G} \mathbf{Z}. \quad (3.7)$$

Les matrices $\mathbf{H}$ de la réponse du canal et $\mathbf{G}$ de l'égalisation étant diagonales, le symbole reçu après le déprécodage de la séquence v est [6]

$$y_v = \underbrace{\sum_{l=1}^L c_{v,l} g_l h_l c_{l,v} x_v}_{A_1} + \underbrace{\sum_{l=1}^L \sum_{\substack{u=1 \\ u \neq v}}^U c_{v,l} g_l h_l c_{l,u} x_u}_{A_2} + \underbrace{\sum_{l=1}^L c_{v,l} g_l z_l}_{A_3}. \quad (3.8)$$

Dans cette expression on distingue, de gauche à droite, le terme A_1 relatif au signal utile, un terme d'interférence A_2 et un terme de bruit A_3 . Sous l'hypothèse de séquences

orthogonales, la capacité du système s'exprime comme la somme des capacités apportées par chacune des séquences. Il suffit alors de calculer l'information mutuelle $\mathcal{I}_v$ entre les processus y_v et x_v . Il faut noter que la capacité du canal est donnée par le maximum de l'information mutuelle entre les signaux d'entrée et de sortie du canal, où la maximisation est faite par rapport à la distribution de l'entrée [53]. En supposant que le terme d'interférence suit une loi gaussienne, ce qui est vrai pour une longueur de séquences suffisamment grande, on peut écrire [3]

$$\mathcal{I}_v = \log_2 \left(1 + \frac{\mathbb{E} [A_1 A_1^H]}{\mathbb{E} [A_2 A_2^H] + \mathbb{E} [A_3 A_3^H]} \right). \quad (3.9)$$

3.3.2.1 Choix de la technique d'égalisation

Comme nous l'avons vu dans le chapitre précédent, différentes techniques d'égalisation peuvent être envisagées pour le système LP-OFDM. Dans notre étude, nous nous limitons à l'utilisation des critères ZF et EQM pour l'égalisation.

Critère ZF

Rappelons que le critère du forçage à zéro permet de supprimer les interférences en réception. Le coefficient d'égalisation est $g_l = 1/h_l$. Comme les séquences de précodage sont orthogonales, $A_2 = 0$ et l'équation (3.8) du signal égalisé par la technique ZF pour le système LP-OFDM s'écrit alors

$$y_v = x_v + \sum_{l=1}^L c_{v,l} \frac{1}{h_l} z_l. \quad (3.10)$$

L'information mutuelle $\mathcal{I}$, correspondant à la somme des informations mutuelles $\mathcal{I}_v$ entre les processus y_v et x_v , s'écrit [3]

$$\mathcal{I} = \sum_{u=1}^U \log_2 \left(1 + \frac{L^2}{\sum_{l=1}^L \frac{1}{|h_l|^2}} \frac{E_u}{N_0} \right) \quad (3.11)$$

Critère EQM

Selon le critère EQM, le coefficient d'égalisation g_n est calculé de sorte à minimiser l'erreur quadratique moyenne entre le signal émis et le signal reçu obtenu après l'égalisation.

sation. L'erreur quadratique moyenne est alors définie par

$$\begin{aligned}\mathbb{E} \left[|\varepsilon_l|^2 \right] &= \mathbb{E} \left[\left| \sum_{u=1}^U g_l h_l c_{l,u} x_u + g_l z_l - \sum_{u=1}^U c_{l,u} x_u \right|^2 \right] \\ &= \sum_{u=1}^U |g_l|^2 |h_l|^2 |c_{l,u}|^2 \mathbb{E} \left[|x_u|^2 \right] + |g_l|^2 \mathbb{E} \left[|z_l|^2 \right] \\ &\quad + \sum_{u=1}^U |c_{l,u}|^2 \mathbb{E} \left[|x_u|^2 \right] - 2\Re \left(\sum_{u=1}^U g_l h_l |c_{l,u}|^2 \mathbb{E} \left[|x_u|^2 \right] \right),\end{aligned}\tag{3.12}$$

avec les hypothèses habituelles $\mathbb{E} [x_u x_v] = \delta_{u,v} E_u$ et $\mathbb{E} [z_l z_j] = \delta_{l,j} N_0$, δ étant ici le symbole de Kronecker, et $|c_{l,u}|^2 = 1$. Pour un système à pleine charge avec $U = L$, la minimisation de cette erreur quadratique moyenne conduit à [6]

$$g_l = \frac{\bar{h}_l}{|h_l|^2 + \frac{N_0}{\sum_{u=1}^U E_u}} = \frac{h_l^*}{|h_l|^2 + \frac{N_0}{E_{\text{DSP}}}}.\tag{3.13}$$

Le développement des termes d'espérance de l'information mutuelle donne,

$$\mathbb{E} \left[A_1 A_1^H \right] = \mathbb{E} \left[\left| \sum_{l=1}^L c_{v,l} g_l h_l c_{l,v} x_v \right|^2 \right] = \left| \sum_{l=1}^L \frac{|h_l|^2}{|h_l|^2 + \frac{N_0}{E_{\text{DSP}}}} \right|^2 E_v;\tag{3.14}$$

$$\mathbb{E} \left[A_2 A_2^H \right] = \mathbb{E} \left[\left| \sum_{l=1}^L \sum_{\substack{u=1 \\ u \neq v}}^U c_{v,l} g_l h_l c_{l,u} x_u \right|^2 \right] = \sum_{\substack{u=1 \\ u \neq v}}^U \left| \sum_{l=1}^L \frac{|h_l|^2}{|h_l|^2 + \frac{N_0}{E_{\text{DSP}}}} c_{v,l} c_{u,l} \right|^2 E_u;\tag{3.15}$$

$$\mathbb{E} \left[A_3 A_3^H \right] = \mathbb{E} \left[\left| \sum_{l=1}^L c_{v,l} g_l z_l \right|^2 \right] = \sum_{l=1}^L \frac{|h_l|^2}{\left(|h_l|^2 + \frac{N_0}{E_{\text{DSP}}} \right)^2} N_0.\tag{3.16}$$

On arrive finalement à l'expression de l'information mutuelle $\mathcal{I}$

$$\mathcal{I} = \sum_{v=1}^U \log_2 \left(1 + \frac{\left| \sum_{l=1}^L \frac{|h_l|^2}{|h_l|^2 + \frac{N_0}{E_{\text{DSP}}}} \right|^2 E_v}{\sum_{\substack{u=1 \\ u \neq v}}^U \left| \sum_{l=1}^L \frac{|h_l|^2}{|h_l|^2 + \frac{N_0}{E_{\text{DSP}}}} c_{v,l} c_{u,l} \right|^2 E_u + \sum_{l=1}^L \frac{|h_l|^2}{\left(|h_l|^2 + \frac{N_0}{E_{\text{DSP}}} \right)^2} N_0} \right).\tag{3.17}$$

3.3.3 Optimisation du débit

Le paragraphe précédent a permis de calculer les informations mutuelles utiles aux développements du problème d'optimisation de la répartition des ressources en vue de la

maximisation du débit. Les contraintes considérées sont la limitation de DSP et l'ordre entier des modulations. Nous considérerons que le masque de puissance définit une DSP plate de niveau E_{DSP} pour la bande « autorisée ». Cela signifie que les algorithmes composeront uniquement avec les sous-porteuses des bandes « autorisées », les sous-porteuses des bandes dites « interdites » étant tout simplement écartées. Enfin, on supposera que le taux d'erreur cible du système est donné par les contraintes de qualité de service. Dans ce paragraphe, la marge de RSB Γ relative aux modulations MAQ sera considérée fixe, ce qui signifie que nous travaillerons sous la contrainte d'un TES crête. Nous allons traiter le problème d'optimisation pour un système monobloc, d'où nous déduirons par la suite le cas multibloc. Les algorithmes seront développés pour un système LP-OFDM avec mise en œuvre de l'égalisation suivant les critères ZF et EQM.

3.3.3.1 Maximisation du débit suivant le critère ZF

A partir de (3.11) et en tenant compte de la marge de RSB Γ , le problème de l'optimisation au sens du débit réalisable s'écrit

$$\left\{ \begin{array}{l} \max_{S, U, E_u \forall u} \sum_{u=1}^U \log_2 \left(1 + \frac{1}{\Gamma} \frac{L^2}{\sum_{n \in S} \frac{1}{|h_n|^2}} \frac{E_u}{N_0} \right), \\ \text{sous la contrainte,} \quad \sum_{u=1}^U E_u \leq E_{\text{DSP}}. \end{array} \right. \quad (3.18)$$

Le débit réalisable prend ses valeurs dans $\mathbb{R}$ dans le cas de la granularité infinie et dans $\mathbb{N}$ dans le cas de la granularité finie. La résolution de ce problème consiste à déterminer la politique $\mathcal{S}$ de sélection des sous-porteuses dans un bloc, le nombre de séquences de précodage U à utiliser ainsi que la politique $\mathcal{P}$ de répartition de la puissance entre les séquences. Cette dernière conduira de manière univoque à la politique $\mathcal{R}$ de répartition des bits r_u sur l'ensemble des séquences.

Choix des sous-porteuses

La politique de sélection des sous-porteuses est définie de sorte à maximiser le débit du système. On remarque que la fonction à maximiser, décrite en (3.18) est une fonction décroissante de la somme des puissances inverses relatives aux sous-porteuses. Il suffit pour cela de chercher une politique qui permet de minimiser $\sum_{n \in S} 1/|h_n|^2$. Ainsi, la stratégie à adopter consiste tout simplement à sélectionner les L sous-porteuses de plus fortes amplitudes $|h_n|^2$, parmi les N disponibles. En notant S^* le jeu de sous-porteuses

ainsi constitué, la politique de sélection des sous-porteuses est donc la suivante

$$\{\mathcal{S} : S^* = \{l\}, \text{ tq. } \forall n \notin S^* \quad |h_l| \geq |h_n| \}. \quad (3.19)$$

Granularité infinie

Connaissant la politique de choix des sous-porteuses, la résolution du problème d'optimisation sous la contrainte de l'inégalité (3.18) à l'aide de la méthode des multiplicateurs de Lagrange a conduit à la détermination de la répartition optimale de la puissance sur les U séquences de précodage utilisées [3]. Le débit maximal dans $\mathbb{R}$ du système monobloc est obtenu pour $E_u = E_{\text{DSP}}/U$ et s'écrit

$$R^* = L \log_2 \left(1 + \frac{1}{\Gamma} \frac{L}{\sum_{l \in S^*} \frac{1}{|h_l|^2}} \frac{E_{\text{DSP}}}{N_0} \right). \quad (3.20)$$

Il faut noter que pour $E_u = E_{\text{DSP}}/U$, le nombre optimal de séquences est $U = L$. Les politiques d'allocation des ressources en granularité infinie peuvent alors être énoncées comme suit

$$\begin{cases} \mathcal{P} : \forall l \in [1, \dots, L] & E_u^* = E_{\text{DSP}}/L, \\ \mathcal{R} : \forall l \in [1, \dots, L] & r_u^* = R^*/L. \end{cases} \quad (3.21)$$

La puissance E_u associée au débit r_u s'écrit aussi

$$E_u^* = (2^{r_u^*} - 1) \frac{\Gamma N_0}{L^2} \sum_{l \in S^*} \frac{1}{|h_l|^2}. \quad (3.22)$$

Notons qu'on atteint de cette manière la capacité du système monobloc. Ces résultats nous permettent de maximiser le débit quelle que soit la longueur des séquences de précodage. En revanche, il n'existe pas de formule analytique pour déterminer la longueur optimale de ces séquences de précodage. En utilisant toutes les constructions possibles des matrices de Hadamard [52], les longueurs possibles des séquences sont $\{1, 2, 4n | n \in \mathbb{N}\}$. Il y a donc un nombre limité de cas à comparer. Dans nos simulations, nous testerons un certain nombre de ces cas afin de déterminer la longueur optimale.

Granularité finie

Si on restreint à présent les ordres de modulations au corps des entiers, les politiques précédentes ne peuvent plus être appliquées et les procédures d'optimisation doivent prendre en compte la contrainte $r_u \in \mathbb{N}$. Le débit maximal réalisable dans $\mathbb{N}$ est alors obtenu en allouant $\lfloor R^*/L \rfloor + 1$ bits à n_b séquences et $\lfloor R^*/L \rfloor$ bits aux $L - n_b$ autres

séquences, avec R^* le débit optimal dans $\mathbb{R}$ donné par l'équation (3.21) [3]. A partir de cette répartition, l'allocation des puissances se déduit simplement en utilisant la relation (3.22). On peut alors énoncer les politiques d'allocation des ressources comme suit

$$\left\{ \begin{array}{l} \forall l \in [1, \dots, n_b] \quad \dot{r}_u^* = \lfloor R^*/L \rfloor + 1, \\ \mathcal{R} : \forall l \in [n_b + 1, \dots, L] \quad \dot{r}_u^* = \lfloor R^*/L \rfloor, \\ \text{avec, } n_b = \left\lfloor L \left(2^{R^*/L - \lfloor R^*/L \rfloor} - 1 \right) \right\rfloor \\ \mathcal{P} : \forall l \in [1, \dots, L] \quad \dot{E}_u^* = (2^{\dot{r}_u^*} - 1) \frac{\Gamma N_0}{L^2} \sum_{l \in S^*} \frac{1}{|h_l|^2}. \end{array} \right. \quad (3.23)$$

Notons que la procédure d'allocation nécessite le calcul de simplement deux valeurs de débit et de deux valeurs de puissance pour réaliser l'allocation de l'ensemble des séquences. Le débit total, dans $\mathbb{N}$, atteint par le système monobloc vaut donc finalement,

$$\dot{R}^* = L \lfloor R^*/L \rfloor + \left\lfloor L \left(2^{R^*/L - \lfloor R^*/L \rfloor} - 1 \right) \right\rfloor. \quad (3.24)$$

Remarquons qu'on a $\dot{R}^* \leq R^*/L$, où l'égalité est obtenue pour $R^*/L \in \mathbb{N}$. Cela signifie que la capacité du système monobloc n'est pas atteinte en granularité finie, tout comme c'était le cas en OFDM. Nous verrons cependant que la différence entre la capacité et le débit réalisé n'atteint pas les mêmes proportions dans les deux systèmes.

L'algorithme d'allocation dynamique des ressources qui regroupe l'ensemble des politiques mises en œuvre pour le système monobloc est présenté figure 3.5. Notons qu'il suffit de remplacer le calcul de $\dot{r}_c^*$ et $\dot{E}_c^*$ par celui de r_c^* et E_c^* pour obtenir les résultats en granularité infinie.

3.3.3.2 Maximisation du débit suivant le critère EQM

A partir de (3.17) et en tenant compte de la marge de RSB Γ , le problème de l'optimisation au sens du débit réalisable pour le système LP-OFDM monobloc avec mise en œuvre de l'égalisation suivant le critère EQM, s'écrit

$$\left\{ \begin{array}{l} \max_{S, U, E_u \forall u} \sum_{u=1}^U \log_2 \left(1 + \frac{1}{\Gamma} \frac{\phi E_u}{\sum_{\substack{v=1 \\ v \neq u}}^U \varphi_{u,v} E_v + \lambda N_0} \right), \\ \text{sous les contraintes,} \quad \sum_{u=1}^U E_u \leq E_{\text{DSP}} \text{ et } E_u \geq 0, \end{array} \right. \quad (3.25)$$

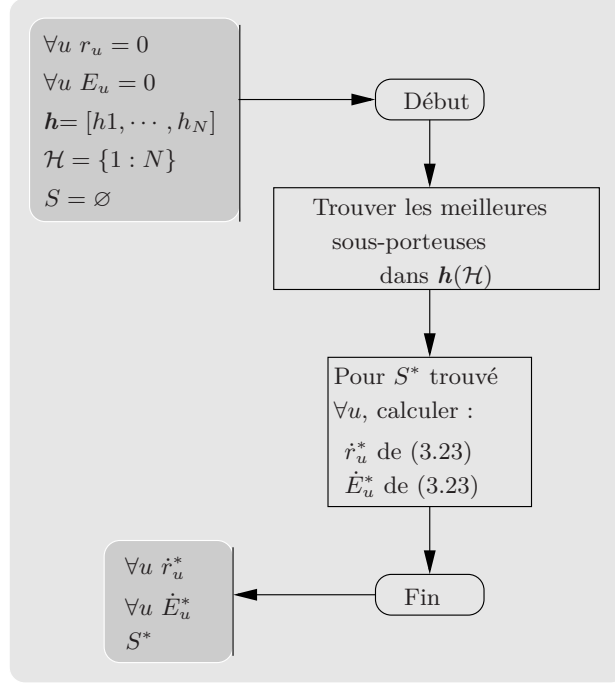


FIGURE 3.5 – Algorithme de maximisation du débit pour le système monobloc.

où

$$\phi = \left| \sum_{l=1}^L \frac{|h_l|^2}{|h_l|^2 + \frac{N_0}{E_{\text{DSP}}}} \right|^2 = \text{Tr} \left(\mathbf{H}\bar{\mathbf{H}}(\mathbf{H}\bar{\mathbf{H}} + \frac{N_0}{E_{\text{DSP}}})^{-1} \right)^2 = \varphi_{u,u}, \quad (3.26)$$

$$\varphi_{u,v} = \left| \sum_{l=1}^L \frac{|h_l|^2}{|h_l|^2 + \frac{N_0}{E_{\text{DSP}}}} c_{u,l} c_{v,l} \right|^2 = \left({}^t \mathbf{C}_u \mathbf{H}\bar{\mathbf{H}}(\mathbf{H}\bar{\mathbf{H}} + \frac{N_0}{E_{\text{DSP}}})^{-1} \mathbf{C}_v \right)^2, \quad (3.27)$$

$$\lambda = \sum_{l=1}^L \frac{|h_l|^2}{\left(|h_l|^2 + \frac{N_0}{E_{\text{DSP}}} \right)^2} = \text{Tr} \left(\mathbf{H}\bar{\mathbf{H}}(\mathbf{H}\bar{\mathbf{H}} + \frac{N_0}{E_{\text{DSP}}})^{-2} \right). \quad (3.28)$$

Contrairement au cas ZF, il n'existe pas de solution analytique au problème d'optimisation du débit. Par contre, pour une répartition donnée des bits il est possible de déterminer la répartition de l'énergie sur les différentes séquences pour un égaliseur EQM. Ainsi, nous utilisons l'allocation des bits réalisée avec l'égaliseur ZF comme point de départ pour l'allocation des bits pour un égaliseur EQM. Nous considérons par la même occasion que la politique $\mathcal{S}$ de sélection des sous-porteuses pour le bloc et le nombre de séquences de précodage U à utiliser sont ceux trouvés avec le critère ZF. Pour une répartition $\{r_u\}_{u \in [1,U]}$ de bits donnée, les énergies $\{E_u\}_{u \in [1,U]}$ allouées aux

séquences vérifient

$$\frac{1}{\Gamma} \phi \frac{E_u}{E_{\text{DSP}}} - (2^{r_u} - 1) \sum_{\substack{v=1 \\ v \neq u}}^U \varphi_{u,v} \frac{E_v}{E_{\text{DSP}}} = (2^{r_u} - 1) \lambda \frac{N_0}{E_{\text{DSP}}}. \quad (3.29)$$

Soit Φ , Λ et Ψ les matrices définies par

$$\left\{ \begin{array}{l} \Phi_{u,u} = \frac{\phi}{\Gamma}; \\ \forall u \neq v, \Phi_{u,v} = -(2^{r_u} - 1) \varphi_{u,v}; \\ \Lambda_u = (2^{r_u} - 1) \lambda \frac{N_0}{E_{\text{DSP}}} \\ \Psi_u = \frac{E_u}{E_{\text{DSP}}}. \end{array} \right. \quad (3.30)$$

La matrice Φ est à diagonale dominante et donc est inversible [54]. Ainsi, pour une répartition $\{r_u\}_{u \in [1,U]}$ de bits donnée, la répartition des énergies relatives $\{\Psi_u\}_{u \in [1,U]}$ vérifie

$$\Psi = \Phi^{-1} \Lambda. \quad (3.31)$$

L'allocation en granularité finie des bits est réalisée par un algorithme itératif. Partant de la répartition initiale de bits et tant que la somme des énergies relatives Ψ_u est inférieure à 1, c'est-à-dire que la contrainte de DSP est vérifiée, la répartition des bits est mise à jour. Cette mise à jour s'effectue en ajoutant un bit supplémentaire à la répartition $\{r_u\}$ tout en minimisant la dispersion des valeurs. Le principe de l'algorithme d'allocation que nous proposons de mettre en œuvre pour déterminer les répartitions des bits et des énergies est présenté figure 3.6.

3.3.3.3 Extension des résultats au cas multibloc

Comme nous le verrons dans les résultats de simulations, l'utilisation de séquences de précodage trop longues fait perdre le gain en débit apporté par le regroupement des sous-porteuses. Dans un tel contexte, la distorsion subie par les séquences en raison de la fonction de transfert sélective du canal de transmission devient trop importante. Il est donc intéressant de remplacer le système monobloc par un système multibloc. Le système en devient plus souple, avec un nombre de sous-porteuses utilisées qui n'est plus strictement égal à la longueur des séquences. Dans notre approche, nous nous limitons à une longueur L des séquences identique pour tous les blocs. Une généralisation est possible en prenant une longueur des séquences propre à chaque bloc, mais il n'existe pas *a priori* de formule analytique donnant la solution optimale.

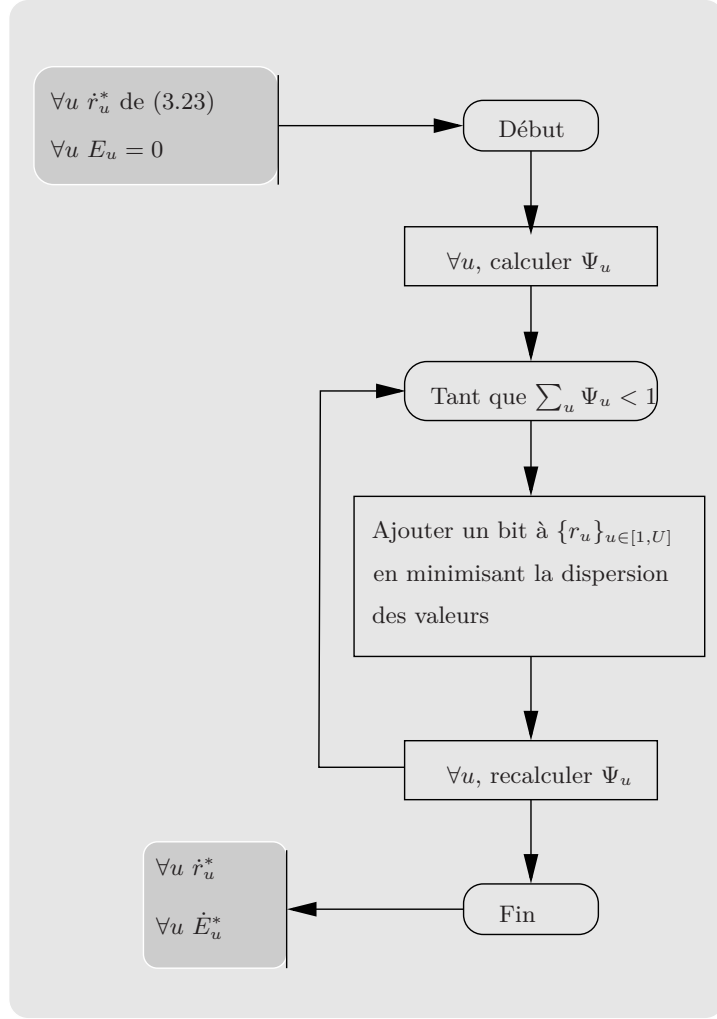


FIGURE 3.6 – Principe de l'algorithme d'allocation avec le critère EQM.

Pour le système multibloc, les algorithmes d'allocation dynamique des ressources doivent gérer le partage du spectre en B blocs de sous-porteuses deux-à-deux disjoints, déterminer le nombre de séquences U_b à mettre en jeu sur chacun de ces blocs, et donner les politiques de répartition de la puissance et des débits sur ces blocs. Les formules analytiques obtenues dans le cas monobloc peuvent encore s'appliquer dans le cas multibloc. Ainsi, La solution optimale consiste à choisir les sous-porteuses de façon à maximiser les débits bloc par bloc. Il suffit alors d'ordonner les sous-porteuses par ordre décroissant de leur amplitude, puis d'attribuer les L premières sous-porteuses au premier bloc, les L suivantes au deuxième bloc, etc [3]. Le débit total R_t^* réalisable s'écrit

$$R_t^* = \sum_{b=1}^B R_b^* = \sum_{b=1}^B \sum_{u=1}^{U_b} r_{u,b}^*. \quad (3.32)$$

Ce débit est donné dans $\mathbb{R}$ et pourrait s'écrire de la même façon dans $\mathbb{N}$. Cette solution obtenue n'est plus optimale dans $\mathbb{N}$, mais elle offre un très bon compromis. La solution optimale nécessiterait d'échanger des sous-porteuses entre les blocs de façon à réduire l'énergie résiduelle, et d'effectuer un grand nombre de calculs de débit.

La figure 3.7 permet d'illustrer l'allocation des débits résultant des politiques développées pour le système LP-OFDM avec mise en œuvre de l'égalisation suivant les critères ZF et EQM. Les mêmes conditions de simulation que pour le système OFDM qui est donné en référence, sont également utilisées ici (cf. fonction de transfert figure 3.2). La différence de comportement des deux systèmes peut alors être mise en évidence. La longueur des séquences de précodage est arbitrairement fixée à $L = 32$. Contrairement au cas de l'OFDM, les courbes de débits sont en escalier à la fois pour la granularité infinie et finie. Comme les sous-porteuses sont classées par ordre croissant d'amplitude pour faciliter la représentation, chaque palier correspond à un bloc S de sous-porteuses. Les indices de sous-porteuses au sein d'un bloc doivent donc être vus comme des indices de séquences de précodage. La largeur des paliers est alors égale à la longueur L de ces séquences. Précisons que la valeur prise par chaque palier, en granularité finie, correspond au débit moyen attribué à chaque séquence, si bien que le débit total véhiculé par un bloc donné est égal à L fois le niveau du palier. Ceci explique que les valeurs obtenues ne soient pas entières, contrairement au cas de l'OFDM.

On remarque que le débit en granularité infinie du système OFDM est supérieur au débit en granularité infinie du système LP-OFDM quel que soit le critère d'égalisation choisi. Par contre, lorsqu'on travaille avec des ordres entiers de modulation, le système LP-OFDM est nettement plus performant que le système OFDM. On remarque que le système LP-OFDM tente d'approcher au mieux la limite de capacité en ajustant le niveau des paliers au plus proche de la courbe de granularité infinie. Pour les dernières sous-porteuses, on remarque de plus que le système est capable de transmettre un débit non nul contrairement à ce que faisait le système OFDM dans la même configuration. Ce comportement est rendu possible grâce à la composante de précodage linéaire qui permet de façon naturelle d'exploiter collectivement la puissance disponible sur chaque sous-porteuse d'un même bloc. Au contraire, le système OFDM n'est capable d'exploiter la puissance que de façon individuelle sur chacune des sous-porteuses. Comme nous l'avons mis en évidence au paragraphe 3.2, la limitation à des ordres entiers de modulations ne lui permet pas d'exploiter la totalité de la puissance disponible sur chaque sous-porteuse. Le système LP-OFDM peut quant à lui mutualiser les bribes de puissance non exploitées par l'OFDM de manière à transmettre des bits supplémentaires. En particulier, cette mutualisation de la puissance permet d'atteindre un débit non nul sur les sous-porteuses où l'OFDM ne pouvait rien transmettre. Notons que ce principe rejoint l'idée proposée dans [31] sous le nom de « fusion de sous-porteuses » (*carrier merging*) et qui

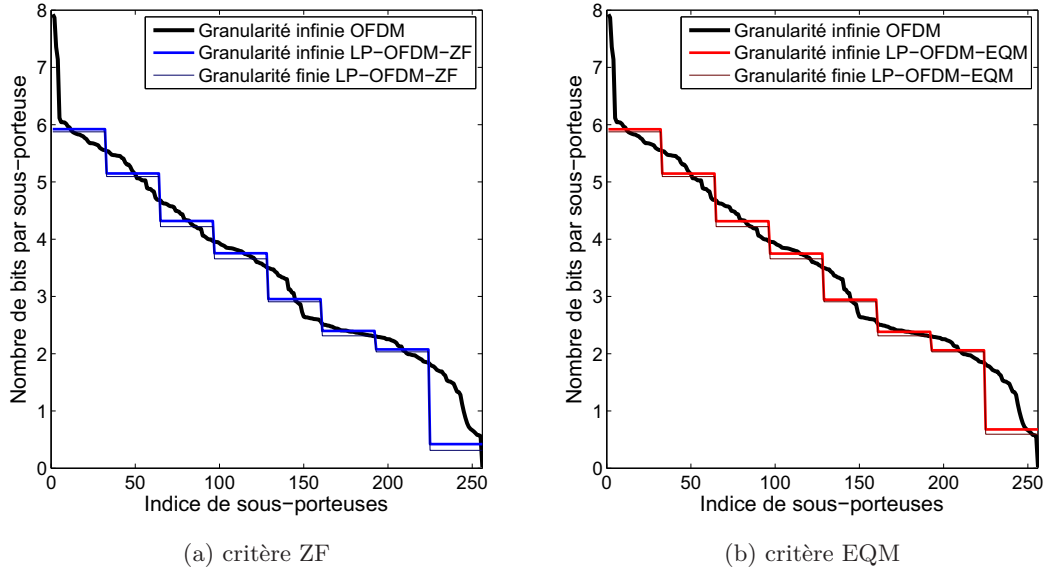


FIGURE 3.7 – Comparaison des résultats d'allocation des bits en granularité infinie et finie, dans le cadre de la maximisation du débit pour le système LP-OFDM multibloc suivant le critère ZF (a) et le critère EQM (b). La longueur de précodage choisie est $L = 32$.

consiste à appliquer une séquence de précodage sur un jeu de sous-porteuses dont les RSB individuels sont trop faibles pour permettre la transmission d'un quelconque ordre de modulation MAQ que ce soit. En vertu du principe de mutualisation de la puissance décrit à l'instant, la fusion des sous-porteuses rend alors possible la transmission de quelques bits supplémentaires dès lors que la ressource cumulée au sein du bloc de sous-porteuses le permet [3].

En outre, l'utilisation du critère EQM apporte une légère amélioration du débit du système LP-OFDM. Pour des raisons de lisibilité des courbes, les systèmes LP-OFDM-ZF (a) et LP-OFDM-EQM (b) ont été séparées. La différence de débits entre ces deux systèmes peut se lire principalement sur les 32 dernières sous-porteuses. Cet exemple confirme les résultats de travaux antérieurs sur le faible gain apporté par l'utilisation du critère EQM dans un contexte mono-utilisateur.

Sur cet exemple simple, nous venons donc de mettre en évidence le fait que le système LP-OFDM adaptatif proposé est capable de dépasser les performances du système OFDM. Cet aspect sera confirmé par de plus amples simulations sur des canaux CPL dans le paragraphe suivant.

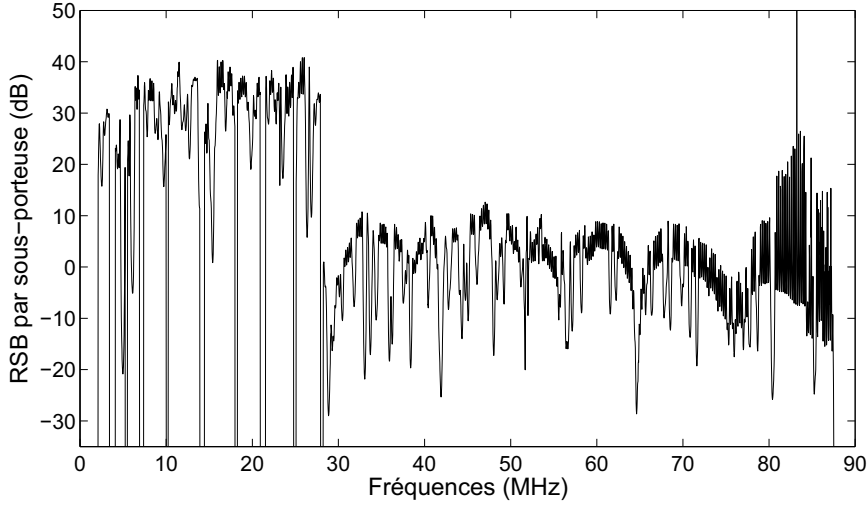


FIGURE 3.8 – Exemple de RSB par sous-porteuse pour un canal de la classe 5.

3.3.4 Simulation et performances

Dans ce paragraphe, nous présentons les performances des algorithmes d'allocation de ressources pour la maximisation du débit dans un contexte mono-utilisateur. Les performances du système LP-OFDM avec mise en œuvre de l'égalisation suivant les critères ZF et EQM seront comparées avec celles du système OFDM.

L'ensemble des simulations va être mené sur des modèles de canaux du projet OMEGA. Rappelons que ces modèles de canaux, issus de campagnes de mesure, ont été répartis en neuf classes de canaux avec des atténuations fréquentielles différentes dans la bande $[0; 100]$ MHz. On se place ici sous l'hypothèse de la connaissance parfaite du RSB sur chaque sous-porteuse à l'émission. Le modèle de bruit stationnaire décrit au paragraphe 1.5.2 sera utilisé pour définir la densité spectrale de puissance de bruit par sous-porteuse. Concernant les puissances d'émission, nous utiliserons le masque de DSP défini au paragraphe 1.4.2. Les autres paramètres de simulation sont définis dans le tableau 1.4 pour le système OMEGA. La figure 3.8 donne un exemple de rapport signal à bruit par sous-porteuse pour un canal de la classe 5 après application du masque de DSP et de la prise en compte du modèle de bruit stationnaire. Les valeurs de RSB en dessous des -30 dB sont celles correspondant aux fréquences non utilisées. Par ailleurs, la probabilité d'erreur symbole cible est fixée à 10^{-3} en sortie de l'égaliseur.

3.3.4.1 Prise en compte des fonctions de codage de canal

La définition de la marge de bruit Γ permet également de prendre en compte le schéma de codage de canal dans les processus d'allocation des ressources. Cette marge

constitue une mesure de la performance relative d'un système utilisant un schéma de codage donné, par rapport à la capacité du canal. Comme nous l'avons vu pour un système non codé, une approximation consiste à considérer Γ indépendant de l'ordre de modulation. La même approche peut être utilisée pour un système codé. Dans ce cas, l'approximation consiste à considérer Γ indépendant du schéma de codage choisi et avec

$$\Gamma = \Gamma_{\text{modulation}} + \Gamma_{\text{codage}}. \quad (3.33)$$

À une probabilité d'erreur symbole cible en sortie de l'égaliseur correspond une probabilité d'erreur en sortie du décodeur, le gain de codage est donné par Γ_{codage} .

Une autre approche consiste à prendre la valeur exacte de la marge de bruit définie pour chaque couple de modulation et de schéma de codage canal. Cette approche a été étudiée dans [4]. La marge de bruit est variable, prend en compte le codage de canal, et l'algorithme correspondant est alors similaire à celui développé au paragraphe 3.4.

Quelle que soit l'approche utilisée, approximation ou valeur exacte de Γ , la prise en compte du schéma de codage ne modifie pas la stratégie et la politique d'allocation des ressources. Cette prise en compte conduit simplement à une modification des valeurs de la ou des marges de bruits. Les valeurs de débits obtenus pourront alors être différentes.

3.3.4.2 Débits atteignables sur les canaux OMEGA

Dans ce paragraphe, nous nous intéressons aux débits transmissibles dans le cas mono-utilisateur sur les différents canaux OMEGA. Les performances des systèmes LP-OFDM avec égalisation ZF et EQM sont comparées à celles d'un système OFDM. La longueur des séquences de précodage est fixée à 32 dans ces premières simulations. La figure 3.9 décrit les fonctions de répartition obtenues pour les débits maximaux sur des canaux issus des 9 classes de canaux. On note en premier lieu que le système LP-OFDM offre les meilleurs débits transmissibles, ce qui rejoint les commentaires du paragraphe 3.3.3.3 sur les débits en granularité finie pour les deux systèmes. Les différences de débits entre les systèmes LP-OFDM-ZF et LP-OFDM-EQM sont faibles et ne permettent pas de dissocier leurs fonctions de répartition. On remarque que plus on monte en numéro de classe de canaux (donc en RSB moyen), plus le gain absolu (la différence de débits entre le LP-OFDM et l'OFDM) augmente. Le tableau 3.1 récapitule les débits transmis, les gains absolus et relatifs à la valeur 0,5 de la fonction de répartition. Les gains absolus et relatifs sont calculés de la façon suivante

$$\begin{aligned} \text{gain absolu} &= \text{débit LP-OFDM} - \text{débit OFDM}, \\ \text{gain relatif} &= \frac{\text{débit LP-OFDM} - \text{débit OFDM}}{\text{débit OFDM}}. \end{aligned}$$

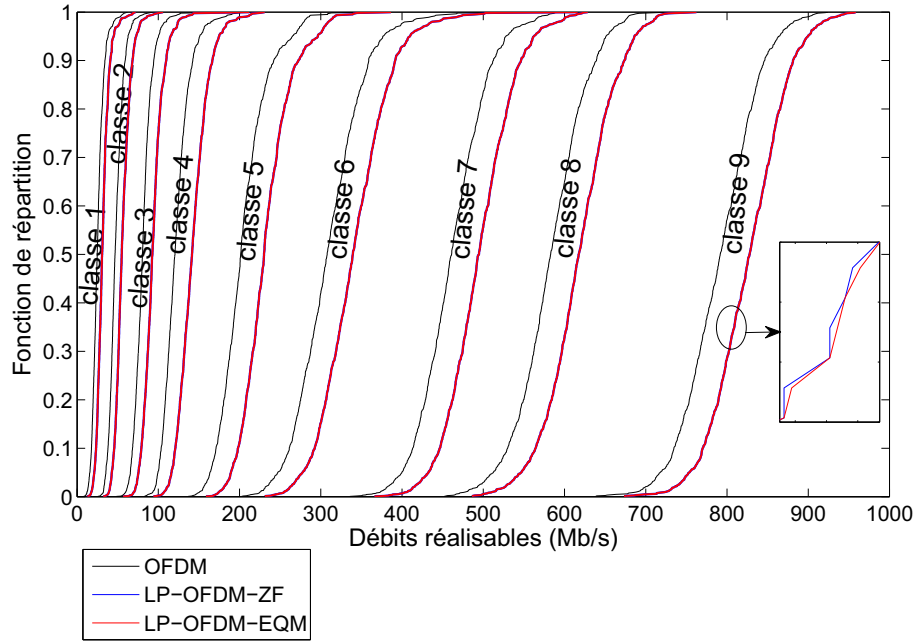


FIGURE 3.9 – Fonctions de répartition des débits atteignables sur des canaux issus des 9 classes de canaux pour les systèmes OFDM et LP-OFDM avec $L = 32$.

On remarque bien que le gain absolu augmente en fonction du numéro de la classe de canaux considéré alors que les valeurs du gain relatif montrent une tendance générale à la baisse.

TABLE 3.1 – Débits transmis, et gains absolus et relatifs à la valeur 0,5 de la fonction de répartition.

	Classes de canaux								
	1	2	3	4	5	6	7	8	9
Débit OFDM (Mb/s)	24,76	46,7	80,52	120,5	201,4	307,9	461	582,9	791,8
Débit LP-OFDM (Mb/s)	30,81	54,96	92,68	141,5	230,1	339,5	494,8	617,7	827,9
Gain absolu (Mb/s)	6,05	8,26	12,16	21	28,7	31,6	33,8	34,8	36,1
Gain relatif (%)	24,43	17,68	15,10	17,42	14,25	10,26	7,33	5,97	4,56

Au regard de ces résultats sur des cas pratiques, on peut dire que la mise en œuvre de la technique de précodage linéaire permet d'augmenter le débit du système en granularité finie. La raison de la meilleure performance du LP-OFDM par rapport à l'OFDM est expliquée dans le paragraphe suivant, où les puissances utilisées sont analysées. On note aussi que l'utilisation d'un égaliseur basé sur le critère de l'erreur quadratique moyenne offre les meilleures performances comparée à l'utilisation du critère du *zero forcing*. Néanmoins, la différence de performance entre ces deux critères est faible sur les canaux qui nous intéressent et dans le contexte mono-utilisateur.

3.3.4.3 Exploitation de la DSP

Nous avons vu précédemment que le système LP-OFDM permettait de mieux exploiter la ressource en puissance que le système OFDM. Nous proposons ici de mettre cette propriété en évidence en représentant la quantité de puissance dépensée par le LP-OFDM en fonction de la longueur L des séquences de précodage. Comme les différences de performance entre les systèmes LP-OFDM-ZF et LP-OFDM-EQM sont faibles, les allocations de puissances donnent quasiment les mêmes résultats. Pour mettre en évidence la puissance utilisée en fonction de la longueur des séquences, nous donnons les résultats de comparaison pour le système LP-OFDM (qui correspond au système LP-OFDM-ZF). Ces résultats sont donnés figure 3.10 pour un unique utilisateur exploitant un canal de classe 5 et pour seulement 1024 sous-porteuses. Pour mettre en évidence la répartition de la puissance, ces sous-porteuses sont triées par ordre décroissant de leur amplitude. La puissance utilisée est la puissance minimale requise permettant la transmission du débit maximal. Les sauts dans les courbes correspondent au changement d'ordre de modulation. Il est évident que le système OFDM n'exploite pas toute la puissance disponible. On constate effectivement que la dépense de puissance est très différente d'une sous-porteuse à l'autre. Dans certains cas, la totalité de la DSP est en effet exploitée tandis que seule la moitié de la puissance exploitable est utilisée dans d'autres cas. Il réside finalement une forte quantité de puissance non allouée comme on le voit sur la figure. C'est cette puissance perdue que le système LP-OFDM met à profit. La composante de précodage linéaire accumule les puissances résiduelles d'un bloc donné de sous-porteuses pour transmettre des bits supplémentaires. On comprend alors que plus la longueur des séquences est élevée, et plus il est possible d'augmenter le débit à partir de quantités de puissances infimes. Sur la figure 3.10, on voit clairement que la puissance disponible est d'autant mieux exploitée que L augmente, la DSP d'émission tendant alors vers la DSP limite dans ce cas.

L'analyse de cette illustration permet donc de mettre en évidence la capacité du système proposé à exploiter plus efficacement la ressource en puissance disponible.

3.3.4.4 Influence de la longueur des séquences de précodage

Nous allons à présent mettre en évidence le gain en débit apporté par la composante de précodage linéaire. Pour cela nous nous intéressons aux résultats obtenus lorsque L varie. La figure 3.11 donne l'évolution du débit atteignable pour les systèmes LP-OFDM-ZF et LP-OFDM-EQM sur un canal de classe 5. Le débit pour le système OFDM est donné pour $L = 1$. On note déjà sur cette figure la confirmation de la faible différence entre les débits atteignables dans le cas de systèmes LP-OFDM avec une égalisation basée sur le critère ZF et EQM. En outre, deux principales tendances sont à relever sur

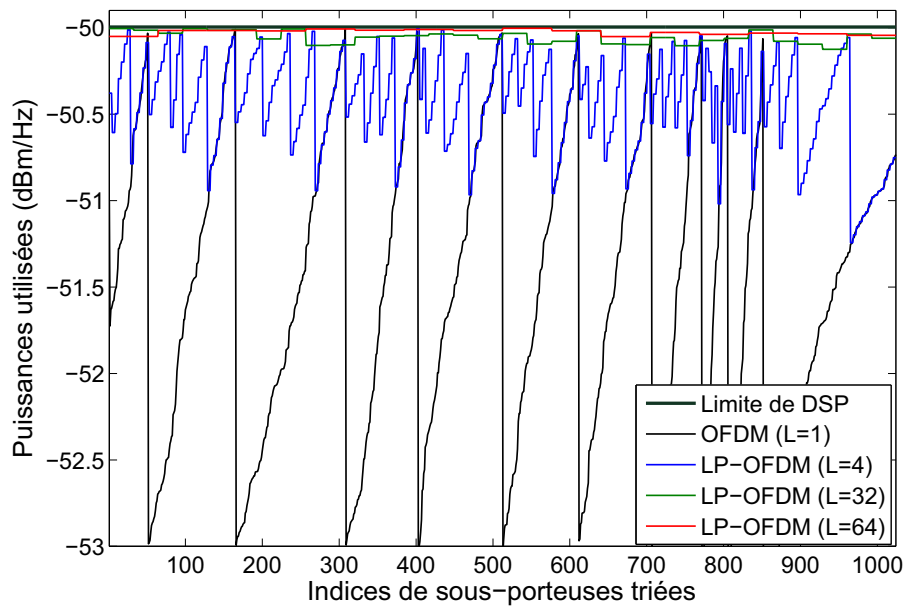


FIGURE 3.10 – Évolution de la puissance utilisée en fonction de la longueur des séquences de précodage sur un canal de classe 5.

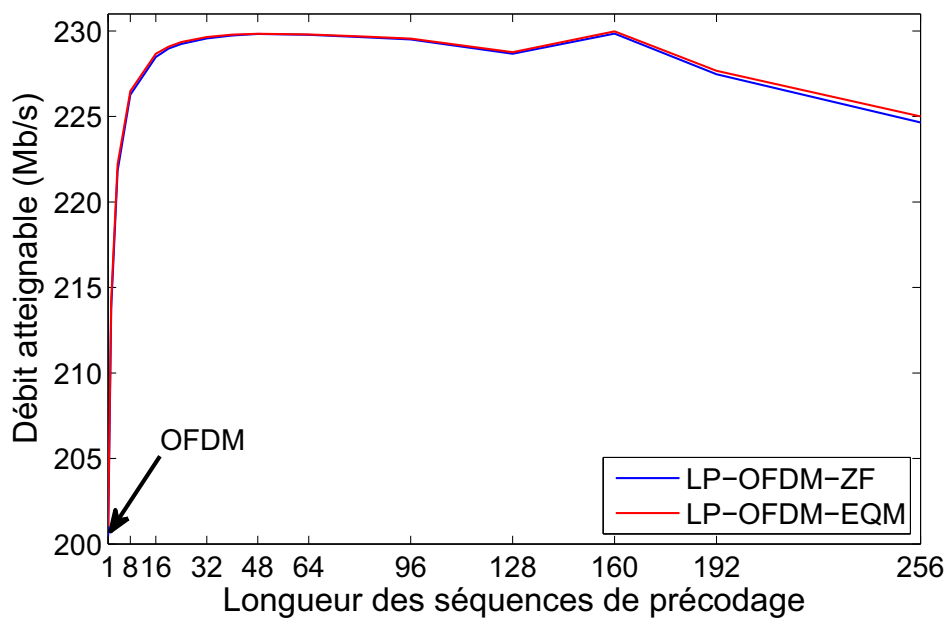


FIGURE 3.11 – Évolution du débit atteignable en fonction de la longueur des séquences de précodage pour les systèmes LP-OFDM-ZF et LP-OFDM-EQM sur un canal de classe 5.

ce tracé. Tout d'abord, il apparait clairement que le débit obtenu avec le système LP-OFDM, quel que soit le critère d'égalisation, augmente lorsque la longueur des séquences est supérieure à 1, atteignant un palier à partir de $L = 32$. On note en particulier que le débit obtenu passe d'environ 200 Mb/s pour l'OFDM à près de 229,6 Mb/s avec $L = 32$. Cela représente un gain absolu de 29,9 Mb/s, soit plus de 14,8% de gain relatif. Ces résultats permettent alors de valider l'analyse faite précédemment, à savoir que l'exploitation de la puissance disponible est plus efficace avec le LP-OFDM qu'avec l'OFDM.

La seconde observation concerne l'évolution des débits lorsque la longueur des séquences de précodage devient grande. On remarque en effet que le débit atteignable commence à décroître pour une longueur supérieure à 160. Une longueur trop importante des séquences de précodage entraîne trop de distorsions entre les sous-porteuses d'un bloc donné. Ce qui a pour effet d'accroître l'interférence entre les séquences de précodage portées par ce bloc. Cette interférence conduit à la réduction de débit pour de très longues séquences de précodage. Finalement, en combinant cette conclusion à celle émise précédemment, il vient qu'on doit faire un compromis sur la longueur des séquences de précodage.

3.3.5 Conclusion

Dans cette phase, nous nous sommes basés sur les résultats connus de l'optimisation des systèmes LP-OFDM avec un détecteur de type ZF pour proposer un nouvel algorithme avec un détecteur de type EQM afin de maximiser le débit total. Les simulations réalisées à partir des modèles de canaux et de bruits stationnaires OMEGA ont permis de comparer les systèmes LP-OFDM avec les deux techniques de détection et le système OFDM. Nous avons pu montrer que les systèmes LP-OFDM permettaient de mieux exploiter la DSP disponible que le système OFDM. Par conséquent, nous avons pu évaluer des gains en débit entre 4 et 25% pour les différentes classes de canaux OMEGA à la valeur 0,5 de la fonction de répartition des débits. En outre, on a pu noter également la faible différence de performances entre les deux critères d'égalisation sur ces mêmes canaux. Le système LP-OFDM avec mise en œuvre de l'égalisation suivant le critère EQM donne les meilleures performances, mais reste relativement plus complexe que le système LP-OFDM-ZF au vu des formules de débits obtenus.

3.4 Maximisation du débit sous contrainte de TEB

Jusqu'ici les problèmes de maximisation du débit utilisent l'approche classique du taux d'erreur symbole cible. Comme discuté dans le chapitre précédent, l'approximation de la marge de RSB Γ (cf. (2.9)), dérivée du TES cible, donne une valeur constante pour

tous les ordres de modulation. En prenant en compte une valeur de TEB, Γ (cf. (2.11)) n'est plus constant pour tous les ordres de modulation. Les politiques d'allocation des ressources doivent alors être modifiées pour tenir compte de la variation de Γ . Les problèmes d'optimisation alors considérés consistent à optimiser les performances du système sous une contrainte, soit de TEB crête soit de TEB moyen. Plusieurs algorithmes d'allocation des ressources sous la contrainte de TEB en vue de la maximisation du débit ont été proposés dans la littérature [55–57]. La plupart des algorithmes de *bit-loading* proposés sont des algorithmes itératifs, qui peuvent augmenter le temps des calculs. En outre, la minimisation du TEB moyen du système pour un débit cible donné afin d'améliorer sa robustesse a fait l'objet d'études dans des précédentes thèses [4, 30].

Dans cette partie, nous proposons deux différentes stratégies d'allocation de ressource basées sur les contraintes de taux d'erreur binaire. L'idée est d'optimiser les débits des systèmes multiporteuses en prenant en compte soit une contrainte de TEB crête soit une contrainte de TEB moyen. Ces nouveaux algorithmes d'allocation binaire tentent d'atteindre l'allocation optimale des bits avec une faible complexité de calcul, tout en respectant une contrainte sur le taux d'erreur binaire crête ou moyen. Dans ce chapitre, les études sont faites dans un contexte mono-utilisateur. Le contexte multi-utilisateur sera analysé dans le chapitre 5. Il faut noter que l'algorithme proposé pour une contrainte de TEB crête est développé pour le système LP-OFDM donc pour le système OFDM. Par contre, l'algorithme proposé pour une contrainte de TEB moyen est développé uniquement pour le système OFDM. Comme nous l'avons vu plus haut, les nombres de bits alloués aux différentes séquences de précodage dans un bloc de sous-porteuses peuvent être différents, compliquant ainsi la recherche d'algorithmes de *bit-loading* avec la contrainte de TEB moyen sur l'ensemble du symbole LP-OFDM. Dans [4], un algorithme itératif de *bit-loading* basé sur la contrainte de TEB moyen a été proposé pour les systèmes LP-OFDM. Cet algorithme est basé sur les hypothèses suivantes : toutes les séquences dans un bloc ont le même nombre de bits et la probabilité d'erreur sur un bloc est la somme des probabilités d'erreur des différentes séquences.

3.4.1 Maximisation du débit sous la contrainte de TEB crête

Dans une première approche, nous étudions le problème de maximisation du débit sous la contrainte de taux d'erreur binaire crête. Le développement des procédures d'optimisation sera fait pour un système LP-OFDM avec une égalisation de type ZF. Le système OFDM s'obtient en fixant L à 1. Le problème de maximisation du débit sous la

contrainte d'un TEB crête P_{eb} , pour un bloc S de L sous-porteuses, s'écrit

$$\max_{r_u, E_u} \sum_{u=1}^U r_u = \max_{r_u, E_u} \sum_{u=1}^U \log_2 \left(1 + \frac{1}{\Gamma(r_u)} \frac{L^2}{\sum_{n \in S} |h_n|^2} \frac{E_u}{N_0} \right), \quad (3.34)$$

sous la contrainte de DSP

$$\sum_{u=1}^U E_u \leq E_{\text{DSP}} \quad (3.35)$$

et de TEB crête

$$\frac{4}{r_u} Q \left(\sqrt{3\Gamma(r_u)} \right) \leq P_{eb}. \quad (3.36)$$

Les paramètres ajustables sont alors la puissance E_u et le nombre r_u de bits alloués à chaque séquence de précodage. L'interdépendance entre le nombre r_u de bits alloués et la marge de RSB $\Gamma(r_u)$ rend difficile l'obtention d'expressions analytiques comme dans le cas d'une contrainte classique de TES crête. Dans [56], un algorithme itératif de *bit-loading* a été proposé dans le but de maximiser le débit total du système sous la contrainte d'un TEB crête. Le principe de cet algorithme est d'allouer itérativement les bits aux séquences de précodage tant que les contraintes de DSP et de TEB crête sont respectées. Dans ce paragraphe, nous proposons un nouvel algorithme efficace de *bit-loading* pour le système LP-OFDM. Cet algorithme est relativement simple à mettre en œuvre car il ne pose aucun problème de convergence au fil des itérations. Une table de correspondance est utilisée afin de stocker, pour chaque ordre de modulation admissible, les amplitudes du canal nécessaires pour garantir le TEB cible. Dans cette optique, nous utilisons les solutions au problème de maximisation du débit en granularité infinie avec une contrainte de TES (cf. paragraphe 3.3.3.2). À partir de l'équation (3.21), la somme des gains inverses du canal relatifs aux sous-porteuses dans le bloc S , nécessaire pour transmettre r_u bits, à pleine charge, s'écrit

$$\sum_{n \in S} \frac{1}{|h_n|^2} = \frac{E_{\text{DSP}}}{\Gamma(r_u) N_0} \frac{L}{(2^{r_u} - 1)}, \quad (3.37)$$

Soit q la fonction définie par

$$q : r_u \mapsto \frac{E_{\text{DSP}}}{\Gamma(r_u) N_0} \frac{L}{(2^{r_u} - 1)}. \quad (3.38)$$

Pour une valeur constante du TEB crête, l'approximation utilisée pour Γ permet de montrer que la fonction q décroît avec r_u . En effet, la fonction $Q(\cdot)$ donnée en (2.10) est strictement décroissante, d'où son inverse $Q^{-1}(\cdot)$ est croissante. De plus, la fonction $r_u \mapsto (2^{r_u} - 1)$ est croissante.

Puisque $\dot{r}_u \leq r_u \leq \dot{r}_u + 1$, avec $\dot{r}_u = \lfloor r_u \rfloor$, il vient que

$$q(\dot{r}_u) \leq \sum_{n \in S} \frac{1}{|h_n|^2} \leq q(\dot{r}_u + 1). \quad (3.39)$$

Ainsi, le débit atteignable R^* en granularité infinie sur le bloc S est calculé à partir de la formule (3.20) en posant $\Gamma = \Gamma(\dot{r}_u)$. La répartition des bits en granularité finie ne peut plus être calculée avec (3.23) puisque la marge de RSB est variable. En s'appuyant sur la répartition des bits obtenue pour un TES crête, la nouvelle répartition optimale des bits pour un TEB crête est

$$\dot{r}_u^* = \begin{cases} \lfloor R^*/L \rfloor + 1 & (1 \leq u \leq n_u), \\ \lfloor R^*/L \rfloor & (n_u < u \leq L), \end{cases} \quad (3.40)$$

où n_u est obtenue en réécrivant la contrainte de DSP

$$n_u = \left\lfloor \frac{L \left(2^{R^*/L} - 2^{\lfloor R^*/L \rfloor} \right)}{(2^{\lfloor R^*/L \rfloor + 1} - 1) \frac{\Gamma(\dot{r}_u^* + 1)}{\Gamma(\dot{r}_u^*)} - (2^{\lfloor R^*/L \rfloor} - 1)} \right\rfloor. \quad (3.41)$$

La contribution en puissance de chaque séquence de précodage à la puissance totale du bloc est donnée par

$$E_u = (2^{\dot{r}_u^*} - 1) \frac{\Gamma(\dot{r}_u^*) N_0}{L^2} \sum_{n \in S} \frac{1}{|h_n|^2}, \quad (3.42)$$

qui satisfait la contrainte de DSP (3.35). L'algorithme du *bit-loading* avec une contrainte de TEB crête pour le bloc S de L sous-porteuses est donné par l'algorithme 1. Comme nous l'avons dit plus haut, l'algorithme que nous proposons ne nécessite aucune itération. L'utilisation d'une table de correspondance permet de réduire la complexité d'implémentation de cet algorithme car il permet de déterminer pour un nombre fini de valeur de r_u les amplitudes du canal nécessaires pour satisfaire la limite de TEB crête. En outre, le passage au système multi-bloc est immédiat. Après la définition des différents blocs de sous-porteuses, l'algorithme est alors appliqué sur chaque bloc.

3.4.2 Maximisation du débit sous la contrainte de TEB moyen

Dans ce paragraphe, nous proposons un algorithme de *bit-loading* pour les systèmes OFDM sous les contraintes de DSP et de taux d'erreur binaire moyen. Comme indiqué précédemment, afin de parvenir à un taux d'erreur cible, les différentes sous-porteuses sont autorisées à être affectées par des valeurs différentes de TEB et une limite sur le TEB moyen est alors fixée pour l'ensemble du symbole OFDM [55, 58]. Le problème de

Algorithme 1 : Algorithme du *bit-loading* avec contrainte de TEB crête.

```

1: pour chaque ordre de modulation  $\dot{r}_u$ , ( $\dot{r}_u = 1, \dots, \dot{r}_{max}$ ) faire
2:   calculer  $\Gamma(\dot{r}_u)$  à partir de (2.11)
3:   calculer  $q(\dot{r}_u)$  à partir de (3.38) et sauvegarder dans une table
4: fin pour
5: pour le bloc  $S$  faire
6:   calculer  $s = \sum_{n \in S} 1/|H_n|^2$ 
7:   si  $s < q(\dot{r}_{max})$  alors
8:      $\forall u \dot{r}_u^* = \dot{r}_{max}$ 
9:   sinon
10:    trouver  $\dot{r}_u$  tel que  $q(\dot{r}_u) \leq s < q(\dot{r}_u + 1)$ 
11:    calculer  $R^*$  à partir de (3.20) en posant  $\Gamma = \Gamma(\dot{r}_u)$ 
12:    calculer  $n_u, \dot{r}_u^*$ 
13:   fin si
14: fin pour

```

maximisation du débit sous la contrainte d'un TEB moyen P_{eb} s'écrit

$$\max_{r_n} \sum_{n=1}^N r_n = \max_{r_n} \sum_{n=1}^N \log_2 \left(1 + \frac{1}{\Gamma(r_n)} \frac{E_{\text{DSP}} |h_n|^2}{N_0} \right), \quad (3.43)$$

sous les contraintes de DSP $E_n \leq E_{\text{DSP}}$ et de TEB moyen

$$\frac{\sum_{n=1}^N r_n P_{eb}(r_n)}{\sum_{n=1}^N r_n} \leq P_{eb}, \quad (3.44)$$

où $P_{eb}(r_n)$ est la valeur du TEB de la sous-porteuse n et est donnée par

$$P_{eb}(r_n) = \frac{4}{r_n} Q \left(\sqrt{3\Gamma(r_n)} \right) = \frac{4}{r_n} Q \left(\sqrt{3 \frac{E_{\text{DSP}} |h_n|^2}{N_0(2^{r_n} - 1)}} \right). \quad (3.45)$$

Ce type de problème peut simplement être résolu en mettant en œuvre un algorithme glouton (*greedy* en anglais), qui constitue une procédure pour laquelle on atteint une solution du problème par une maximisation pas à pas de chaque fonction locale. Cet algorithme peut être appliqué de deux façons différentes. La première façon est de commencer par allouer l'ordre de modulation le plus élevé à toutes les sous-porteuses. La taille de constellation de la sous-porteuse la moins performante est réduite à chaque itération jusqu'à ce que la contrainte de TEB moyen soit respectée [55]. La deuxième façon consiste à commencer par allouer l'ordre de modulation le plus faible à toutes les

sous-porteuses. La taille de constellation de la sous-porteuse la plus performante est augmentée à chaque itération, tant que la contrainte de TEB moyen est respectée [47]. Dans cette dernière méthode, la taille de constellation de la sous-porteuse n^* est augmentée, où

$$n^* = \arg \min_n \{(r_n + 1)P_{eb}(r_n + 1) - r_n P_{eb}(r_n)\}. \quad (3.46)$$

Bien que cet algorithme donne le débit optimal, la complexité de calcul est encore relativement élevée. Plus précisément, le nombre d'itérations varie entre N et $r_{\max} \times N$. Il est donc nécessaire de trouver des algorithmes d'allocation binaire, qui tentent d'atteindre l'allocation optimale avec une faible complexité de calcul, tout en respectant une contrainte sur le TEB moyen. Un algorithme de faible complexité, qui exploite la relation entre le TEB crête autorisé par sous-porteuse et le TEB moyen, a été proposé par Wyglinski dans [55]. En faisant varier la limite de TEB crête, le TEB moyen se trouve alors modifié. Ainsi, en modifiant de façon intelligente la limite maximale du TEB crête, on aboutit à une allocation des bits plus rapide que les algorithmes incrémentaux avec une solution plus proche de l'optimal. Après avoir déterminé la limite maximale du TEB crête, l'algorithme demeure incrémental.

L'algorithme que nous proposons est une version modifiée de l'algorithme de Wyglinski. En dehors des itérations pour déterminer la limite maximale du TEB crête, l'algorithme ne nécessite aucune autre itération supplémentaire. Les étapes permettant de déterminer la limite maximale et l'allocation finale des bits sont décrites dans l'algorithme 2. Jusqu'à l'étape 6, cet algorithme est le même que l'algorithme de détermination de la limite initiale de TEB crête qui sera utilisé dans l'algorithme itératif d'allocation des bits [55, 57]. La détermination de cette limite initiale de TEB crête permet d'initialiser l'algorithme itératif pour réduire le nombre d'itérations. Dans notre algorithme, la détermination du paramètre I (dans l'algorithme 2) suffit pour trouver l'allocation finale des bits en vertu de l'équation (3.48). Ainsi, cet algorithme est plus rapide que celui de Wyglinski, comme le montre le tableau 3.2. Dans ce tableau, on note que les temps de calculs augmentent en fonction du RSB moyen. L'algorithme que nous proposons engendre un temps de calcul inférieur d'au moins une seconde par rapport à l'algorithme de Wyglinski.

TABLE 3.2 – Temps de calculs moyen (maximal) en seconde pour différentes valeurs de RSB moyen sous MATLAB avec Intel Core2@2.66GHz.

Algorithme	20 dB	35 dB	50 dB
Proposé dans [55, 57]	1.9813 (2.2188)	2.8375 (3.1875)	2.7750 (3.1719)
Notre proposition	1.1453 (1.2031)	1.4641 (1.5313)	1.4719 (1.5313)

Algorithme 2 : Algorithme du *bit-loading* avec contrainte de TEB moyen.

- 1: Pour chaque sous-porteuse n , calculer $P_{eb}(r_n)$ correspondant aux différents ordres de modulation r_n ($r_u = 1, \dots, r_{max}$)
- 2: Trouver β_n , la plus grande valeur des $P_{eb}(r_n)$ qui ne dépasse pas P_{eb} et l'ordre de modulation correspondant r_n^β
- 3: Trouver α_n , la plus faible valeur des $P_{eb}(r_n)$ qui ne dépasse pas P_{eb} et l'ordre de modulation correspondant r_n^α
- 4: Déterminer la valeur ΔP_{eb} telle que

$$P_{eb} = \frac{\sum_{n=1}^N r_n^\beta \beta_n + \Delta P_{eb}}{\sum_{n=1}^N r_n^\beta} \Leftrightarrow \Delta P_{eb} = \sum_{n=1}^N r_n^\beta (P_{eb} - \beta_n) \quad (3.47)$$

- 5: Remplacer β_n par α_n tant que la contrainte de TEB moyen est respectée. En supposant que les α_n sont triées dans l'ordre croissant, sans perdre de généralité, on peut écrire

$$P_{eb} \geq \frac{\sum_{n=I+1}^N r_n^\beta \beta_n + \sum_{n=1}^I r_n^\alpha \alpha_n}{\sum_{n=I+1}^N r_n^\beta + \sum_{n=1}^I r_n^\alpha} \quad (3.48)$$

et par la suite,

$$\Delta P_{eb} + \left(\sum_{n=1}^I (r_n^\alpha - r_n^\beta) \right) P_{eb} \geq \sum_{n=1}^I (r_n^\alpha \alpha_n - r_n^\beta \beta_n) \quad (3.49)$$

- 6: Déterminer la plus grande valeur de I pour laquelle la relation (3.48) est vraie.
- 7: L'allocation finale des bits est donc la suivante :

$$r_n = \begin{cases} r_n^\alpha & (1 \leq n \leq I) , \\ r_n^\beta & (I < n \leq N) . \end{cases} \quad (3.50)$$

3.4.3 Simulation et performances

Dans ce paragraphe, les performances des algorithmes de *bit-loading* proposés avec prise en compte du taux d'erreur binaire sont évaluées. Ces algorithmes sont également comparés avec les stratégies d'allocation des ressources suivant l'approche classique considérant une contrainte de taux d'erreur symbole crête. Les conditions de simulation sont identiques à celles décrites au paragraphe 3.3.4. En outre, les contraintes de TEB et de TES sont respectivement fixées à 10^{-3} et 2.10^{-3} .

La figure 3.12 donne les fonctions de répartition des débits réalisables avec les différents algorithmes de *bit-loading* sur des canaux de classe 2 (figure 3.12a), de classe 5 (figure 3.12b) et de classe 9 (figure 3.12c). *Primo*, on note que le passage d'une contrainte de TES crête à une contrainte de TEB crête ou de TEB moyen permet d'augmenter le débit des systèmes OFDM. Le gain en débit offert par les algorithmes avec contrainte

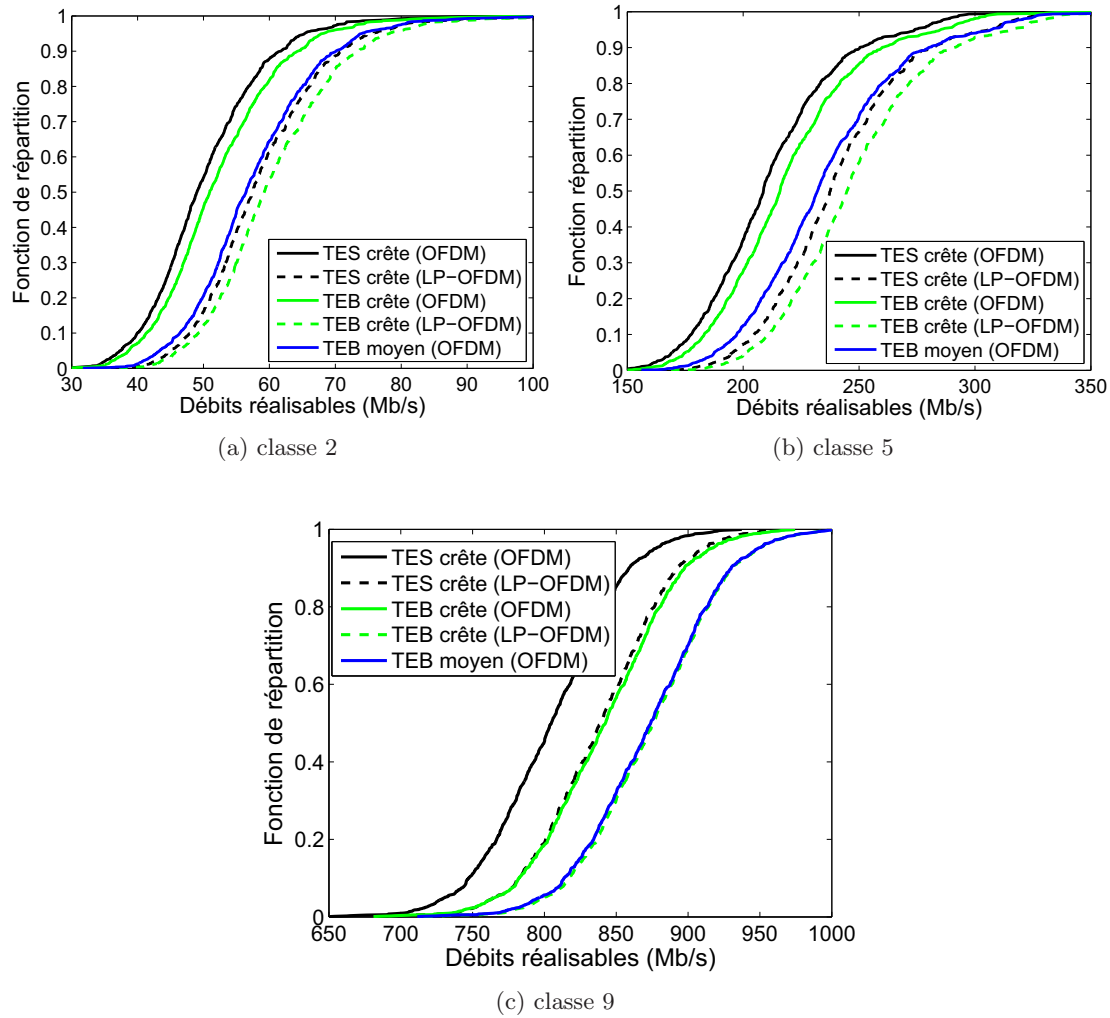


FIGURE 3.12 – Fonctions de répartition des débits réalisables avec les différents algorithmes de *bit-loading* sur des canaux de classe 2 (a), de classe 5 (b) et de classe 9 (c).

de TEB croît en absolu avec le numéro de la classe des canaux, donc du RSB moyen. De plus, l'utilisation d'une contrainte de TEB moyen donne de meilleurs résultats qu'avec le TEB crête pour les systèmes OFDM. *Secundo*, le gain en débit apporté par la technique de précodage linéaire dans une approche classique de TES crête, pour les classes 2 et 5, assure un meilleur débit total par rapport à l'application des approches de TEB sur les systèmes OFDM. L'utilisation conjointe de la technique de précodage et l'approche de TEB crête offre le meilleur débit total quelle que soit la classe de canaux utilisée.

3.5 Conclusion

Dans ce chapitre, nous avons présenté plusieurs stratégies d'allocation des ressources visant à maximiser le débit des systèmes OFDM et LP-OFDM dans un contexte mono-utilisateur. Dans une première partie, nous avons rappelé les mécanismes d'application des modulations adaptatives aux systèmes OFDM et les procédures d'optimisation des ressources correspondantes. Dans la seconde partie, nous avons aussi rappelé les principaux résultats de l'optimisation du système LP-OFDM avec l'utilisation d'un détecteur *zero forcing* en réception dans le contexte des transmissions CPL. Puis, nous avons développé un nouvel algorithme d'allocation pour le système LP-OFDM avec mise en œuvre de l'égalisation suivant le critère de l'erreur quadratique moyenne, dans le cadre d'une optimisation du débit, ce qui constitue une première contribution originale.

Lors des simulations sur les canaux OMEGA, nous avons mis en évidence les résultats déjà connus, à savoir que le système OFDM n'avait pas la capacité d'exploiter efficacement la ressource en puissance disponible, contrairement au système LP-OFDM. Les résultats présentés montrent que la composante de précodage permet de mutualiser les puissances disponibles sur les différentes sous-porteuses d'un même bloc, et de les exploiter collectivement pour augmenter le débit du système. On a également pu noter le faible gain en débit apporté par l'utilisation d'un égaliseur basé sur le critère EQM comparé à l'utilisation du critère ZF.

Dans la dernière partie, nous avons proposé deux nouveaux algorithmes de *bit-loading* de faible complexité, qui cherchent à maximiser le débit du système, tout en respectant une contrainte de taux d'erreur binaire crête ou moyen. Les résultats de simulations montrent que le passage d'une contrainte de TES crête à une contrainte de taux d'erreur binaire crête ou moyen permet d'augmenter le débit des systèmes OFDM. De plus, la mise en œuvre des algorithmes de *bit-loading* avec la contrainte de TEB crête pour les systèmes LP-OFDM conduit à de meilleures performances.

Finalement, on peut généraliser et affirmer que la capacité de la fonction d'étalement à collecter, mutualiser et exploiter toutes les puissances disponibles permet au système LP-OFDM d'offrir, sous contrainte de masque de DSP, de meilleurs résultats en débit que le système OFDM dans un contexte mono-utilisateur. Nous allons dans les chapitres qui suivent, tenter de tirer profit des bonnes performances du système LP-OFDM dans un contexte multi-utilisateur.

Chapitre 4

Allocation des ressources dans un contexte multicast

4.1 Introduction

En informatique, le terme *multicast* (multidiffusion) est utilisé pour désigner une méthode de diffusion de l'information d'un émetteur (source unique) vers un groupe de récepteurs (plusieurs supports/médias). On dit aussi diffusion multipoint ou diffusion de groupe. L'avantage de ce système par rapport au classique unicast (connexion réseau point à point) devient évident quand on veut diffuser de la vidéo. En transmission de flux continu et en unicast, on envoie une image autant de fois qu'on a de connexions simultanées ce qui conduit à une perte de temps, de ressources du serveur et surtout de bande passante. *A contrario* de l'unicast, en multicast le paquet n'est émis qu'une seule fois, et sera routé vers tous les destinataires du groupe de diffusion. Lorsque le paquet est destiné à tous les éléments du réseau, on parle alors de broadcast. La figure 4.1 présente les différents types de communications dans un réseau.

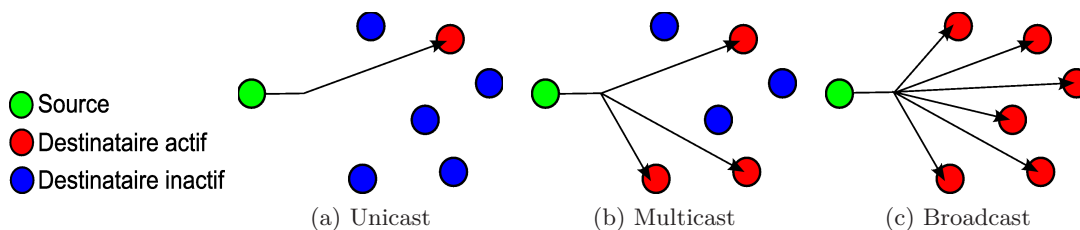


FIGURE 4.1 – Types de communications

Dans ce chapitre, nous étudierons la problématique de l'allocation des ressources dans un contexte multicast. Le routage en multicast a fait l'objet de nombreuses études dans la littérature pour les réseaux sans fils et câblés [59–61]. Dans ce chapitre, nous

étudions une approche couche physique (PHY) des communications en multicast. Nous nous limiterons donc à l'adaptation de la couche PHY pour les systèmes à porteuses multiples en multicast sur lignes d'énergie. Le réseau CPL pouvant être utilisé pour fournir des services *triple play* (l'Internet à haut débit, la téléphonie et la télévision sur IP), le multicast s'avère être une méthode intéressante dans ce contexte. Les paramètres de la couche PHY doivent être adaptés de sorte à satisfaire tous les destinataires. La méthode classique en OFDM multicast consiste à ajuster ces paramètres pour servir l'utilisateur qui présente le plus mauvais canal de transmission. Comme nous allons le voir, cette méthode offre un même débit, à tous les utilisateurs, limité par l'état du canal du plus « mauvais » utilisateur [62].

Afin de mieux correspondre aux conditions des liens et d'augmenter le débit total en multicast, il a été proposé d'utiliser le multicast à débits hétérogènes [63]. Un tel système sera nommé système multicast multidébit. Le système où tous les utilisateurs reçoivent le même débit sera nommé système multicast monodébit. Dans un système multicast multidébit, les utilisateurs sont regroupés en sous-groupes et les destinataires dans chaque sous-groupe reçoivent un service en fonction de l'état de leurs canaux de transmission. Par conséquent, les systèmes multicast à débits hétérogènes ont un grand avantage sur les systèmes multicast à débit uniforme en s'adaptant aux exigences des différents destinataires et à l'hétérogénéité du réseau [64]. Une méthode pour construire des systèmes multicast à débits hétérogènes est la compression des données multicast en plusieurs couches suivant une hiérarchie de qualités [65]. Les couches allouées aux utilisateurs dépendent de la qualité des canaux de transmission. L'utilisation de l'approche multicast multidébit a permis d'augmenter significativement le débit multicast total par rapport à la méthode classique mais au prix d'une dégradation de l'équité entre les utilisateurs. Pour remédier à ce problème, une autre méthode prenant en compte l'équité a été proposée [62, 65]. Cette méthode adapte les débits des utilisateurs en fonction des débits alloués précédemment. Par conséquent, l'indice d'équité entre les utilisateurs se trouve amélioré mais le débit total est réduit.

Dans ce chapitre, nous proposons un système LP-OFDM pour les communications dans un contexte multicast afin d'augmenter les débits des utilisateurs et d'améliorer l'indice d'équité. Pour exploiter la diversité des liens de transmission des utilisateurs, les méthodes proposées utilisent la technique de précodage linéaire, ce qui constitue la deuxième contribution originale de cette thèse. Dans un premier temps, nous étudierons le problème d'allocation des ressources pour les systèmes LP-OFDM dans un contexte multicast monodébit. Dans un second temps, les méthodes développées seront analysées dans un contexte multicast multidébit.

4.1.1 Description du système

Comme indiqué précédemment, le multicast est la méthode qui permet de délivrer une même donnée à plusieurs utilisateurs à travers une seule transmission. Cette technique est particulièrement intéressante pour la transmission de données multimédias à haut débit car elle permet de réduire l'utilisation des ressources du réseau [62]. La figure 4.2 illustre un simple cas de communication multicast dans un contexte CPL où la source S envoie des données multimédias à trois destinataires $D1$, $D2$ et $D3$. Un coordinateur central, CCo (*central coordinator*), contrôle les activités du réseau et envoie périodiquement des trames balises contenant des informations de planification allouant ainsi du temps à chaque communication dans un contexte CPL [7]. Puisque la couche MAC du CPL est orientée connexion, toutes les communications de données se font à travers des connexions « logiques ». Dans les télécommunications, une communication orientée connexion décrit un moyen de transmission des données dans lequel

- les appareils à chaque extrémité, utilisent un protocole préliminaire pour établir une connexion de bout en bout avant que les données soient envoyées ;
- les données sont envoyées le long d'un même chemin pendant la communication.

Ces connexions « logiques » sont identifiées individuellement par un identifiant global de connexion et aussi des spécifications de qualité de service. Les exigences de qualité de service comprennent la garantie de bande passante, la quasi absence d'erreurs de transmission, la latence fixe et le contrôle de la gigue [7]. La source S communique ses besoins et ses exigences au CCo et demande au CCo une allocation appropriée de temps de connexion pour les communications avec les destinataires $D1$, $D2$ et $D3$. Si le CCo peut prendre en compte les besoins de S , il demande aux stations de sonder leurs canaux de transmission [7]. Cela permet aux stations d'effectuer l'estimation initiale du canal. On suppose qu'un trajet de retour fait état des amplitudes $|H_{k,n}|$ des sous-porteuses (d'indice n) des canaux de propagation entre chaque destinataire (d'indice k) et la source. Dans notre étude, on suppose également que la source connaît parfaitement les amplitudes du canal de chaque destinataire, ce qui n'est pas forcément le cas dans les réseaux CPL actuels. En effet, seules les liaisons point à point, nécessitant la connaissance d'un seul canal, sont possibles dans ces réseaux. Sur la base de la connaissance des différents canaux et des exigences de qualité de service, la source peut alors effectuer l'allocation des ressources du système multicast.

Les caractéristiques fréquentielles du signal LP-OFDM sont les mêmes que celles d'un signal OFDM. La composante OFDM du signal est supposée adaptée au canal de transmission. L'intervalle de garde, le nombre de sous-porteuses, et l'espace entre ces sous-porteuses sont correctement choisis afin d'absorber parfaitement les échos de tous les canaux de propagation, et de limiter la perte d'efficacité spectrale liée à cet inter-

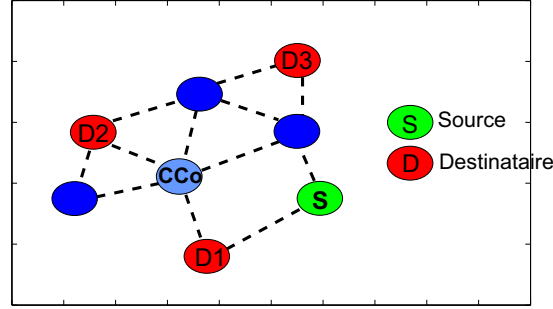


FIGURE 4.2 – Exemple de communication multicast dans un réseau CPL.

valle de garde. Due à la contrainte de DSP dans les systèmes CPL, tous les utilisateurs ont la même contrainte de puissance crête sur chaque sous-porteuse. Les séquences de précodage utilisées sont les séquences orthogonales de Walsh-Hadamard de longueur L . Pour simplifier les calculs, on suppose ici que chaque bloc de précodage utilise la même longueur de séquences.

4.2 Allocation des ressources en multicast monodébit

Dans ce paragraphe, nous analysons l'allocation des ressources dans le cas multicast monodébit c'est-à-dire où tous les utilisateurs reçoivent le même débit de transmission. Pour assurer un tel débit identique à tous les utilisateurs, la méthode d'allocation des ressources doit adapter les paramètres de la couche PHY à l'état du canal du plus « mauvais » utilisateur sur chaque sous-canal. Les sous-canaux peuvent être les sous-porteuses dans le cadre de l'OFDM ou les blocs de sous-porteuses dans le cadre du LP-OFDM. L'hétérogénéité des canaux de transmission entre la source et les destinataires complique l'adaptation de la couche PHY (ordre de modulation, codage de canal, etc.) à l'état du canal de chaque utilisateur. Nous présenterons tout d'abord la méthode classique d'allocation des ressources en OFDM multicast et ses limites. Ensuite, nous proposerons d'utiliser conjointement la technique de précodage linéaire et une adaptation de la méthode classique d'allocation des ressources pour augmenter le débit des systèmes multicast à débits homogènes sans hiérarchiser les données. Il est à noter que les équations proposées dans ce paragraphe pour le système LP-OFDM sont basées sur l'expression du débit en granularité infinie avec un détecteur de type ZF (cf. paragraphe 3.3.3). L'expression du débit obtenue pour une égalisation de type EQM s'avère relativement complexe, donc difficilement compatible avec le développement d'expressions analytiques lors des procédures d'optimisation. L'algorithme d'allocation des ressources basé sur le critère EQM sera néanmoins appliqué à la fin de ce paragraphe sur la solution qui donne

le meilleur résultat dans le cas ZF.

4.2.1 Notion de canal équivalent

Du fait qu'en multicast monodébit tous les utilisateurs reçoivent les mêmes ressources, le débit multicast peut être considéré comme le débit calculé sur un canal équivalent de transmission. Ce canal est créé à partir de la combinaison des différents canaux des utilisateurs et est défini par

$$H_b^{\text{eq}} = f(\{H_{k,i}\}_{i \in S_b, k \in [1, K]}), \quad (4.1)$$

où b représente l'indice des sous-porteuses ou des blocs de sous-porteuses, I le nombre total de sous-porteuses ou de blocs de sous-porteuses et K le nombre total de destinataires. Après construction de ce canal équivalent, l'allocation des ressources en multicast devient équivalente à l'allocation des ressources en mono-utilisateur (point-à-point). Ainsi, les algorithmes d'allocation des ressources en mono-utilisateur peuvent être directement appliqués sur ce canal équivalent.

4.2.2 Application de la méthode LCG

En multicast, la méthode classique LCG (*low channel gain*, [65]) est une méthode simple qui permet d'allouer l'information multicast tout en satisfaisant les qualités de service des usagers. Avec cette méthode, l'ordre de modulation sur chaque sous-porteuse est calculé à partir de la plus faible amplitude des utilisateurs sur la sous-porteuse. Ainsi, dans le cadre de l'OFDM, le canal équivalent pour chaque sous-porteuse est défini par l'amplitude minimale des utilisateurs sur cette sous-porteuse. Dans la suite, ce canal sera nommé le canal équivalent OFDM (CeO),

$$|H_n^{\text{eq}}| = \min_k |H_{k,n}|. \quad (4.2)$$

Le débit total réalisé par la méthode LCG sur la sous-porteuse n dans un contexte CPL s'écrit

$$\mathcal{R}_n^{\text{LCG}} = K \log_2 \left(1 + \frac{E_{\text{DSP}}}{\Gamma N_0} |H_n^{\text{eq}}|^2 \right). \quad (4.3)$$

4.2.2.1 Limitation du débit avec la méthode LCG

Supposons que les canaux de propagation des différents utilisateurs sont des canaux de Rayleigh indépendants et identiquement distribués (iid) de moment d'ordre deux σ^2 . Le gain $|H_{k,n}|^2$ suit une loi χ_2 à deux degrés de liberté [66]. L'analyse suivante est faite pour une seule sous-porteuse et l'indice de la sous-porteuse sera supprimé. La densité de

probabilité P_H et la fonction de répartition F_H pour l'utilisateur k s'écrivent

$$\begin{aligned} P_{H_k}(h) &= \frac{1}{2\sigma^2} \exp\left(-\frac{h}{2\sigma^2}\right), \quad h > 0; \\ F_{H_k}(h) &= 1 - \exp\left(-\frac{h}{2\sigma^2}\right), \quad h > 0. \end{aligned} \quad (4.4)$$

Le débit multicast total dans $\mathbb{R}$ sur une sous-porteuse avec la méthode LCG s'écrit

$$\begin{aligned} \mathcal{R}_n^{\text{LCG}} &= K \left(\min_k \log_2 \left(1 + \frac{E}{\Gamma N_0} |H_k|^2 \right) \right) \\ &= K \log_2 \left(1 + \frac{E}{\Gamma N_0} \left(\min_k |H_k|^2 \right) \right) \end{aligned} \quad (4.5)$$

Soit $H_m = \min_k |H_k|^2$, la statistique d'ordre 1. La densité de probabilité de la variable H_m s'écrit [67]

$$P_{H_m}(h) = \frac{K}{2\sigma^2} \exp\left(-\frac{Uh}{2\sigma^2}\right), \quad h > 0. \quad (4.6)$$

Ainsi,

$$\begin{aligned} \mathbb{E}[\mathcal{R}_n^{\text{LCG}}] &= \mathbb{E} \left[K \log_2 \left(1 + \frac{E}{\Gamma N_0} H_m \right) \right] \\ &= \int_{\mathbb{R}_+} K \log_2 \left(1 + \frac{E}{\Gamma N_0} h \right) \frac{K}{2\sigma^2} \exp\left(-\frac{Uh}{2\sigma^2}\right) dh. \end{aligned} \quad (4.7)$$

Comme [68]

$$\begin{aligned} \int_{\mathbb{R}_+} \exp(-\alpha h) \log(1 + \beta h) &= -\frac{1}{\alpha} \exp\left(\frac{\alpha}{\beta}\right) \mathbf{E}_i\left(-\frac{\alpha}{\beta}\right), \\ \text{pour } \arg(\beta) < \pi, \text{ et } \alpha > 0, \text{ où } \mathbf{E}_i(-h) &= -\int_h^{+\infty} \frac{\exp(-t)}{t} dt \end{aligned} \quad (4.8)$$

il s'en suit que

$$\begin{aligned} \mathbb{E}[\mathcal{R}_n^{\text{LCG}}] &= -\frac{K}{\log(2)} \exp\left(\frac{K\Gamma N_0}{2E\sigma^2}\right) \mathbf{E}_i\left(-\frac{K\Gamma N_0}{2E\sigma^2}\right) \\ &= \frac{K}{\log(2)} \exp\left(\frac{K\Gamma N_0}{2E\sigma^2}\right) \int_{\frac{K\Gamma N_0}{2E\sigma^2}}^{+\infty} \frac{\exp(-t)}{t} dt, \end{aligned} \quad (4.9)$$

et donc

$$\lim_{K \rightarrow \infty} \mathbb{E}[\mathcal{R}_n^{\text{LCG}}] = \frac{2}{\log(2)} \frac{E\sigma^2}{\Gamma N_0}. \quad (4.10)$$

La figure 4.3 présente les résultats analytiques, calculés avec (4.9), et de simulation pour l'espérance mathématique du débit multicast total dans $\mathbb{R}$ avec la méthode LCG. Pour une convergence rapide mais sans pour autant perte de généralité, les DSP du signal

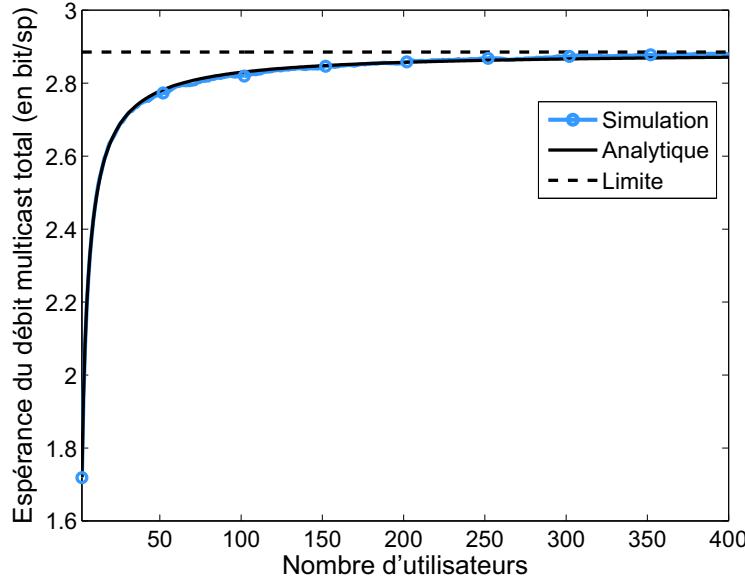


FIGURE 4.3 – Comparison des résultats analytiques et de simulation pour le débit multicast total dans $\mathbb{R}$, en bit par sous-porteuse, avec la méthode LCG ($E = 1$, $N_0 = 1$, and $\sigma^2 = 1$).

et du bruit sont fixées à 1. Ces résultats montrent que le débit total dans $\mathbb{R}$ du système OFDM multicast est borné lorsque le nombre d'utilisateurs augmente. Les résultats obtenus dans $\mathbb{N}$ sont différents de ceux obtenus ici et seront détaillés au paragraphe 4.2.4 pour l'ensemble des solutions proposées.

4.2.2.2 Amélioration du débit multicast offert par la méthode LCG

On a vu qu'avec la méthode LCG, le débit total est le débit en simple lien calculé sur le canal CeO. Pour augmenter le débit multicast LCG, l'algorithme du *bit-loading* en simple lien pour le LP-OFDM est directement appliqué sur le canal CeO. Cette méthode est considérée comme la méthode LP-LCG (*linear precoding based LCG*). Il a été prouvé dans le chapitre précédent que la technique de précodage appliquée à l'OFDM permet d'obtenir en simple lien un meilleur débit comparé à l'OFDM. On peut s'attendre à ce que la simple application du précodage sur le canal équivalent CeO apporte une augmentation des débits. Le débit multicast LP-LCG s'écrit

$$\mathcal{R}_{\text{LP-LCG}} = K \sum_{b=1}^B L \log_2 \left(1 + \frac{1}{\Gamma} \frac{L}{\sum_{n=1}^N \frac{d_{b,n}}{|H_n^{\text{eq}}|^2}} \frac{E}{N_0} \right), \quad (4.11)$$

où $d_{b,n}$ sont les éléments d'une matrice de décision D de taille $B \times N$. Cette matrice de décision D détermine la répartition des N sous-porteuses dans les B blocs, et D satisfait les contraintes suivantes

$$d_{b,n} = \begin{cases} 1 & \text{si } n \in S_b, \\ 0 & \text{sinon,} \end{cases} \quad \text{et } \forall n, \sum_{b=1}^B d_{b,n} = 1. \quad (4.12)$$

Du fait de l'utilisation de la détection ZF, la matrice D optimale doit être définie de sorte à minimiser les distorsions entre les sous-porteuses du canal équivalent dans chaque bloc. Les blocs sont alors composés de sous-porteuses adjacentes après un réarrangement des sous-porteuses du canal CeO dans l'ordre décroissant [69]. Soit O le vecteur des indices des sous-porteuses du canal CeO réarrangées dans l'ordre décroissant. La matrice de décision D s'écrit alors

$$d_{b,n} = \begin{cases} 1 & \text{si } n \in \{O_j \mid (b-1)L + 1 \leq j \leq bL\}, \\ 0 & \text{sinon.} \end{cases} \quad (4.13)$$

Voici un exemple avec $N = 8$, $L = 4$, $B = 2$ et $O = [5 \ 1 \ 8 \ 7 \ 4 \ 3 \ 2 \ 6]$. Il s'en suit que

$$D = \begin{pmatrix} 1 & 0 & 0 & 0 & 1 & 0 & 1 & 1 \\ 0 & 1 & 1 & 1 & 0 & 1 & 0 & 0 \end{pmatrix}. \quad (4.14)$$

4.2.3 Allocation des ressources en LP-OFDM multicast

Pour augmenter le débit des systèmes multicast à débit identique pour tous les utilisateurs, nous proposons d'exploiter la technique de précodage linéaire au contexte multicast monodébit afin de mieux tirer profit de la diversité des liens de transmission des utilisateurs. Dans le cadre du système LP-OFDM multicast, le débit réalisé dans le bloc b en multicast sera le plus faible débit des utilisateurs sur ce bloc. Ainsi, le canal équivalent LP-OFDM (CeL) sur le bloc b , pour une matrice D donnée, est défini par

$$|H_b^{\text{eq}}(D)|^2 = \min_k \frac{L}{\sum_{n=1}^N \frac{d_{b,n}}{|H_{k,n}|^2}}. \quad (4.15)$$

Ainsi, le débit LP-OFDM multicast en fonction de la définition de la matrice D s'écrit pour ce bloc

$$\mathcal{R}_b^{\text{LP}}(D) = L \log_2 \left(1 + \frac{E}{\Gamma N_0} |H_b^{\text{eq}}(D)|^2 \right). \quad (4.16)$$

En considérant que la longueur L des séquences d'étalement est constante d'un bloc

de sous-porteuses à un autre et d'un utilisateur à un autre, la seule donnée paramétrable est la matrice de décision D . Notre but étant de maximiser le débit multicast total, il s'agira donc de maximiser la somme des plus faibles débits des utilisateurs sur les différents blocs. Pour cela, il faut trouver une stratégie de regroupement des sous-porteuses dans les différents blocs qui maximise les plus faibles débits des blocs. Le problème d'optimisation pour un système de B blocs s'écrit

$$\max_D \sum_{b=1}^B \mathcal{R}_b^{\text{LP}} = \max_D \sum_{b=1}^B \left(\min_k \mathcal{R}_{k,b} \right), \quad (4.17)$$

sous les contraintes

$$\left\{ \begin{array}{l} d_{b,n} = \begin{cases} 1 & \text{si } n \in S_b \\ 0 & \text{sinon} \end{cases} \quad (C1) \\ \forall n, \sum_{b=1}^B d_{b,n} = 1 \quad (C2) \\ \forall b, \sum_{n=1}^N d_{b,n} = L \quad (C3) \end{array} \right. \quad (4.18)$$

4.2.3.1 Solution combinatoire

Le problème (4.17) est un problème d'optimisation combinatoire et la méthode basique de sa résolution est de tester toutes les définitions possibles de la matrice D et de choisir la meilleure définition. Pour chaque ligne de la matrice, il faut mettre L colonnes à « 1 » et les autres à « 0 » conformément à la contrainte (C3) dans (4.18). Il est clair qu'un tel algorithme serait fini mais le problème correspondant est de complexité non polynomiale. Ainsi, pour un grand nombre N de sous-porteuses, la recherche de la solution optimale n'est plus faisable. Dans le calcul du débit en LP-OFDM, l'ordre des sous-porteuses dans un bloc et l'ordre des blocs n'influent pas sur le résultat final. Ainsi, le nombre total de définitions possibles de la matrice D est donnée par

$$\frac{C_N^L \times C_{N-L}^L \times C_{N-2L}^L \times \dots \times C_{2L}^L}{\left(\frac{N}{L}\right)!} = \frac{N!}{(L!)^{\left(\frac{N}{L}\right)} \left(\frac{N}{L}\right)!}, \quad (4.19)$$

où N est le nombre total de sous-porteuses et est proportionnel à la taille L des blocs. Ce nombre représente le nombre total de combinaisons de L parmi N , de L parmi $N - L$ et ainsi de suite. Il est ensuite divisé par le nombre total d'arrangement des blocs. Le tableau 4.1 donne le nombre de possibilités pour $N = 128$ pour différentes valeurs de L . Cette solution donne le débit optimal en LP-OFDM multicast mais sa mise en œuvre devient vite irréalisable lorsque le nombre de sous-porteuses croît.

TABLE 4.1 – Nombre total de définitions possibles de D pour $N = 128$.

Longueur L	Nombre de possibilités
4	1×10^{136}
8	378×10^{126}
16	260×10^{102}
32	3×10^{72}
64	12×10^{36}

Le temps de calcul de cet algorithme peut être réduit en utilisant l'algorithme dit de séparation et d'évaluation, également appelé selon le terme anglo-saxon *branch and bound*. Cet algorithme fait une énumération implicite de toutes les solutions candidates tout en évitant d'énumérer de larges classes de mauvaises solutions, selon l'analyse des propriétés du problème. Des outils logiciels basés sur cet algorithme peuvent être utilisés pour résoudre ce problème combinatoire. Néanmoins, le temps de calcul reste élevé pour un grand nombre de sous-porteuses, notamment dans un contexte CPL. Dans la suite, nous proposons une solution simple au problème d'optimisation représenté par l'expression (4.17). Cette solution tire profit à la fois du précodage linéaire et de la diversité des liens de transmission des utilisateurs.

4.2.3.2 Solution LBCG (*low block channel gain*)

La méthode LP-LCG (cf. paragraphe 4.2.2.2) exploite la dimension du précodage linéaire sur le canal équivalent CeO mais ne tire pas profit de la diversité des liens de transmission des utilisateurs sur les blocs. En effet, de part la construction du canal CeO, la moyenne harmonique des sous-porteuses d'un utilisateur quelconque sur un bloc est meilleure que la moyenne harmonique des sous-porteuses du canal CeO sur ce bloc, c'est-à-dire,

$$\forall k, b, \frac{L}{\sum_{n=1}^N \frac{d_{b,n}}{|H_{k,n}|^2}} \geq \frac{L}{\sum_{n=1}^N \frac{d_{b,n}}{|H_n^{\text{eq}}|^2}}, \quad (4.20)$$

puisque $|H_{k,n}|^2 \geq |H_n^{\text{eq}}|^2$. Avec la méthode LBCG, on définit tout d'abord la matrice D de répartition des blocs à partir du canal équivalent OFDM CeO (4.2). Ensuite, le débit multicast sur chaque bloc est calculé avec le canal équivalent LP-OFDM CeL (4.15).

Le débit multicast offert par la méthode LBCG est meilleur que le débit multicast

offert par la méthode LP-LCG. En effet, à partir de (4.20), il s'en suit que

$$\forall k, b, \mathcal{R}_{k,b} \geq L \log_2 \left(1 + \frac{E}{\Gamma N_0} \frac{L}{\sum_{n=1}^N \frac{d_{b,n}}{|H_n^{\text{eq}}|^2}} \right) \quad (4.21)$$

et ainsi

$$\sum_{b=1}^B \left(\min_k \mathcal{R}_{k,b} \right) \geq L \sum_{b=1}^B \log_2 \left(1 + \frac{E}{\Gamma N_0} \frac{L}{\sum_{n=1}^N \frac{d_{b,n}}{|H_n^{\text{eq}}|^2}} \right). \quad (4.22)$$

A partir de (4.22), on en déduit que le débit LP-OFDM multicast possède une borne inférieure. Selon (4.17), on cherche à maximiser le membre de gauche de l'inégalité (4.22). Maximiser le membre de droite de cette inégalité permet d'accroître le membre de gauche. Or, nous avons vu que la définition de D dans (4.13), maximise le membre de droite qui correspondrait alors au débit multicast obtenu avec la méthode LP-LCG. On montre ainsi que le débit multicast offert par la méthode LBCG est meilleur que le débit multicast offert par la méthode LP-LCG mais aussi que cette méthode LBCG n'est pas optimale.

A partir de (4.16), le débit multicast LBCG s'écrit

$$\mathcal{R}_{\text{LBCG}} = K \sum_{b=1}^B L \log_2 \left(1 + \frac{E}{\Gamma N_0} |H_b^{\text{eq}}(D)|^2 \right). \quad (4.23)$$

L'algorithme 3 décrit la méthode de calcul du débit multicast pour la solution LBCG. Le résultat pour la méthode classique LCG s'obtient en fixant L à 1.

4.2.4 Comportement asymptotique du débit multicast total

Dans ce paragraphe, nous étudions le comportement asymptotique du débit multicast total pour un grand nombre d'utilisateurs. Contrairement à la méthode LCG où la limite asymptotique a été calculée pour K tendant vers l'infini (cf. (4.10)), le développement des équations n'a pas permis de trouver une formule analytique pour les méthodes basées sur le précodage linéaire. L'évolution à l'infini du débit multicast total sera donc déterminée par simulation. Afin d'étudier la solution optimale, différents canaux indépendants et identiquement distribués sont utilisés avec un nombre réduit de sous-porteuses. Le nombre total de sous-porteuses est ainsi fixé à $N = 12$ et la longueur des séquences de précodage à $L = 4$. La figure 4.4 présente la moyenne des débit multicast totaux en bit par sous-porteuse pour les différentes méthodes d'allocation des ressources. Ces débits

Algorithme 3 : Méthode LBCG**Entrées** : $N, K, B, L, \forall k, n |H_{k,n}|^2$ **Sorties** : débit par utilisateur R **début** $R \leftarrow 0;$ **pour** toute sous-porteuse $n, n \in [1; N]$ **faire** calculer le canal CeO $|H_n^{\text{eq}}|^2$ à partir de (4.2); trier $|H_n^{\text{eq}}|^2$ par ordre décroissant, soit O l'ensemble des indices ainsi triés; définir D comme dans (4.13);**pour** tout bloc $b, b \in [1; B]$ **faire** calculer le canal CeL $|H_b^{\text{eq}}|^2(D)$ à partir de (4.15); calculer $\mathcal{R}_b^{\text{LP}}$ à partir de (4.16); $R \leftarrow R + \mathcal{R}_b^{\text{LP}};$ **fin**

sont calculés à la fois pour les modulations non contraintes (c'est-à-dire des débits réels, dans $\mathbb{R}_+$) et pour les modulations contraintes (c'est-à-dire des débits entiers, dans $\mathbb{N}$). La limite à l'infini du débit multicast total offert par la méthode LCG est de 144 bits par sous-porteuse et cette limite est atteinte pour un très grand nombre d'utilisateurs (plus de 10000). Ces résultats montrent que le débit multicast total obtenu avec les méthodes basées sur la technique de précodage linéaire (optimale, LP-LCG et LBCG) semble aussi borné. Comme attendu, la méthode combinatoire dite optimale donne le meilleur débit multicast total. L'application directe du LP-OFDM sur le canal équivalent CeO (méthode LP-LCG) est moins performante que la méthode LCG dans $\mathbb{R}_+$ mais donne un meilleur résultat dans $\mathbb{N}$. Cela confirme les résultats obtenus lors de la comparaison faite au chapitre précédent entre l'OFDM et le LP-OFDM. Contrairement à la méthode LP-LCG, la méthode LBCG offre toujours un débit multicast supérieur à celui offert par la méthode LCG avec un ratio de plus de 2. De plus, cette méthode offre des performances très proches de celles de la méthode optimale. Néanmoins ce dernier résultat ne peut être généralisé pour tout nombre N de sous-porteuses car la méthode optimale ne peut être exécutée pour N supérieur à 12.

Dans la pratique, le nombre maximum de bits par sous-porteuse est limité et la granularité est de 1 bit. La figure 4.4 montre aussi que les débits dans $\mathbb{N}$ (c'est-à-dire pour les modulations contraintes) deviennent nuls lorsque le nombre d'utilisateurs K augmente. Nous allons montrer que lorsque le nombre d'utilisateurs tend vers l'infini, le débit multicast tend vers zéro. Soient $X \in \{\text{LCG, LP-LCG, LBCG, optimale}\}$ et α_X la limite du débit multicast total offert par la méthode X dans $\mathbb{R}_+$. En utilisant la fonction

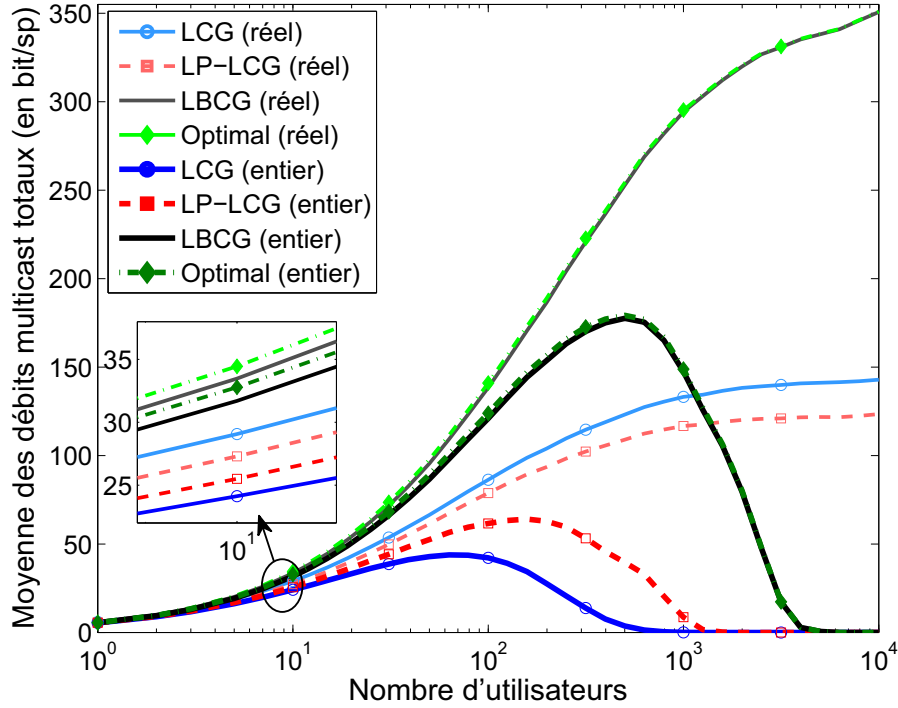


FIGURE 4.4 – Moyenne des débits multicast totaux en bit par sous-porteuse en fonction du nombre d'utilisateurs avec $\frac{E}{\Gamma N_0} = 20$ dB et $\sigma = 1$.

« partie entière », la limite dans $\mathbb{N}$ s'écrit

$$\lim_{K \rightarrow +\infty} K \left\lfloor \frac{\alpha_X}{K} \right\rfloor \quad (4.24)$$

et pour $K > \alpha_X$, $\left\lfloor \frac{\alpha_X}{K} \right\rfloor = 0$. Il s'en suit que la limite dans $\mathbb{N}$ de ces quatre méthodes est égale à zéro.

4.2.5 Application de l'algorithme du *bit-loading* avec le critère EQM

Puisque la méthode LBCG présente le meilleur débit multicast, nous appliquerons l'algorithme du *bit-loading* avec le critère EQM uniquement pour cette méthode. À partir du canal équivalent LP-OFDM CeL, nous avons vu que le débit multicast offert par la méthode LBCG sur le bloc b avec un détecteur de type ZF s'écrit

$$\mathcal{R}_b^{\text{LBCG}}(D) = L \log_2 \left(1 + \frac{E}{\Gamma N_0} |H_b^{\text{eq}}(D)|^2 \right), \quad (4.25)$$

avec la matrice D définie dans (4.13). Cette méthode sera considérée dans la suite comme la méthode « LBCG (ZF) » en référence au détecteur utilisé. Ainsi, la nouvelle méthode,

qui elle, est basée sur la technique de détection EQM, sera considérée comme la méthode « LBCG (EQM) ». Nous avons vu que l'application sur un bloc de l'algorithme du *bit-loading* avec le critère EQM nécessite la connaissance d'une répartition initiale des bits dans le bloc. La répartition des bits dans le bloc b , s'écrit donc (cf. paragraphe 3.3.3.1)

$$\begin{cases} \forall l \in [1, \dots, n_b] & \dot{r}_{u,b} = \lfloor \mathcal{R}_b^{\text{LBCG}}/L \rfloor + 1, \\ \forall l \in [n_b + 1, \dots, L] & \dot{r}_{u,b} = \lfloor \mathcal{R}_b^{\text{LBCG}}/L \rfloor, \\ \text{avec, } n_b = \lfloor L \left(2^{\mathcal{R}_b^{\text{LBCG}}/L} - \lfloor \mathcal{R}_b^{\text{LBCG}}/L \rfloor - 1 \right) \rfloor \end{cases} \quad (4.26)$$

À partir de ce résultat, la répartition des bits pour les systèmes LP-OFDM multicast avec le critère EQM peut être déterminée en utilisant l'algorithme 3.6 proposé au chapitre 3, appliqué sur le canal équivalent CeL. Les simulations sur les canaux CPL seront réalisées au paragraphe 4.4.

4.3 Allocation des ressources en multicast multidébit

Dans les méthodes d'allocation des ressources en multicast précédemment présentées, tous les utilisateurs partagent les mêmes sous-porteuses ou blocs de sous-porteuses, ce qui conduit à un même débit pour tous les utilisateurs. On a vu que lorsque les liens de transmission sont très différents, le débit multicast est limité par les utilisateurs qui ont les plus mauvais canaux. Dans ce contexte, il est justifiable d'allouer plus de ressources à certains utilisateurs qu'à d'autres. Afin d'augmenter le débit total en multicast et de mieux prendre en compte les conditions des liens, il a été proposé de passer d'un système multicast monodébit à un système multicast multidébit [63]. Les utilisateurs, dans un tel système, sont alors regroupés en sous-groupes et les destinataires dans chaque sous-groupe reçoivent un service en fonction de l'état de leurs canaux de transmission. Les sous-groupes peuvent être définis dans le domaine fréquentiel [65, 70] ou dans le domaine temporel [62]. Dans le domaine fréquentiel, chaque sous-porteuse est allouée à un sous-groupe alors que dans le domaine temporel, chaque intervalle de temps est alloué à un sous-groupe. En s'adaptant aux exigences des différents destinataires et à l'hétérogénéité du réseau, le système multicast multidébit présente un avantage par rapport au système multicast monodébit [64]. Pour réaliser un tel système multicast multidébit, plusieurs flux de données à débit variable sont générés à partir d'une même donnée source.

Dans cette partie, nous développons des méthodes d'allocation de ressources en LP-OFDM multicast pour le contexte multicast multidébit. Les problèmes d'optimisation du débit total traités dans la littérature pour le système OFDM multicast multidébit en technologie sans-fil seront reformulés et adaptés au contexte CPL avant d'être étendus aux systèmes LP-OFDM. Nous proposons également une solution empirique de définition

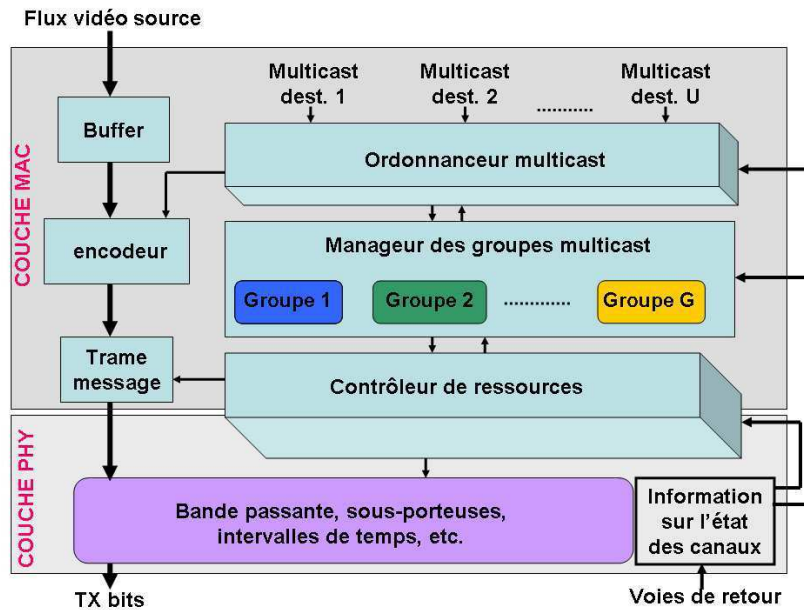


FIGURE 4.5 – Module inter-couche PHY-MAC du côté de l'émetteur.

des sous-groupes d'utilisateurs dans le domaine temporel.

4.3.1 Description du système multicast multidébit

En supposant que l'émetteur peut fournir plusieurs flux à débits différents pour une même donnée, les destinataires peuvent alors être séparés en sous-groupes et chaque sous-groupe reçoit un même service en fonction de l'état des canaux de transmission. La figure 4.5 présente le module inter-couche PHY-MAC du côté de l'émetteur pour la transmission de données dans un système multicast à débits hétérogènes. L'information sur l'état des canaux de transmission est utilisée par l'ordonnanceur, le manageur des groupes multicast et le contrôleur des ressources. Le manageur des groupes multicast regroupe les utilisateurs en sous-groupes et les sous-groupes sont définis dans le domaine temporel (sur chaque intervalle de temps) ou dans le domaine fréquentiel (sur chaque sous-porteuse). De plus, le manageur délivre des informations telles que les débits à l'ordonnanceur. Le module d'ordonnancement multicast détermine la quantité de données dans chaque trame. Le contrôleur des ressources quant à lui réalise l'allocation des sous-porteuses et des intervalles de temps, et détermine l'ordre de modulation sur chaque sous-porteuse. L'encodeur code les données du flux source en fonction des différents débits et des rendements de codage. Les approches possibles pour l'encodage à débits variables des données multicast sont décrites dans la section 4.3.1.1. Nous supposons que la taille du *buffer* est suffisamment grande pour stocker les données en temps réel. Dans la trame message, les bits sont regroupés à partir de la combinaison des bits des flux

de données encodés et des informations de management des sous-groupes [64]. Suivant les différentes fonctions de la couche physique (modulation OFDM, intervalle de garde, conversion, etc.), les bits TX sont transmis aux différents destinataires. Il est à noter que les paramètres du signal (la bande passante, le nombre de sous-porteuses, l'espacement inter-porteuse, etc.) sont identiques pour tous les utilisateurs quelle que soit la méthode choisie.

4.3.1.1 Réalisations possibles de l'encodage à débits variables

La scalabilité ou la capacité à être échelonné est la propriété d'un flux vidéo de permettre une transmission et un décodage partiels à un débit évoluant dans un intervalle donné et permettant la restitution d'une résolution ou d'une qualité variables en fonction de ce débit. La figure 4.6 présente quelques méthodes pour l'encodage des données sources en vue d'une transmission dans les systèmes multicast à débits variables. Les approches suivantes peuvent être utilisées pour obtenir la scalabilité en débit des données à transmettre :

- Utilisation de plusieurs encodeurs où la sortie de chaque encodeur est à débit prédéfini (constant ou variable) comme dans la figure 4.6a. Les encodeurs créent des flux à différents débits. La scalabilité est alors obtenue en commutant entre les différents flux encodés disponibles [71] ;
- Utilisation tout d'abord d'un encodeur dont la seule sortie est à débit prédéfini (constant ou variable). Ensuite plusieurs scaleurs sont utilisés pour créer des flux à différents débits (voir figure 4.6b). Un scaleur est un système qui permet de mettre à l'échelle une image suivant plusieurs résolutions. Différents types de scaleurs peuvent être utilisés, entre autres la re-quantification des coefficients DCT (*discrete cosine transform*) ou la sélection de certaines trames à supprimer (c'est la technique du *frame-dropping*, en anglais). La scalabilité est alors obtenue en sélectionnant les flux scalés disponibles [72] ;
- Utilisation d'un encodeur qui code un flux avec plusieurs niveaux hiérarchiques de qualité (voir figure 4.6c). Un flux de base et plusieurs améliorations sont disponibles. Pour décoder le niveau hiérarchique courant, il est nécessaire de recevoir le niveau précédent. Plus on ajoute de niveaux de qualité au flux de base, plus la qualité du flux reçu est meilleure. La scalabilité est alors obtenue en ajoutant des niveaux de qualité au flux de base [72] ;
- Utilisation d'un encodeur où on peut dynamiquement ajuster le débit de sortie pour obtenir la scalabilité en débit (voir figure 4.6d). Ce système sera moins performant lorsque plusieurs utilisateurs requièrent des débits très différents [72].

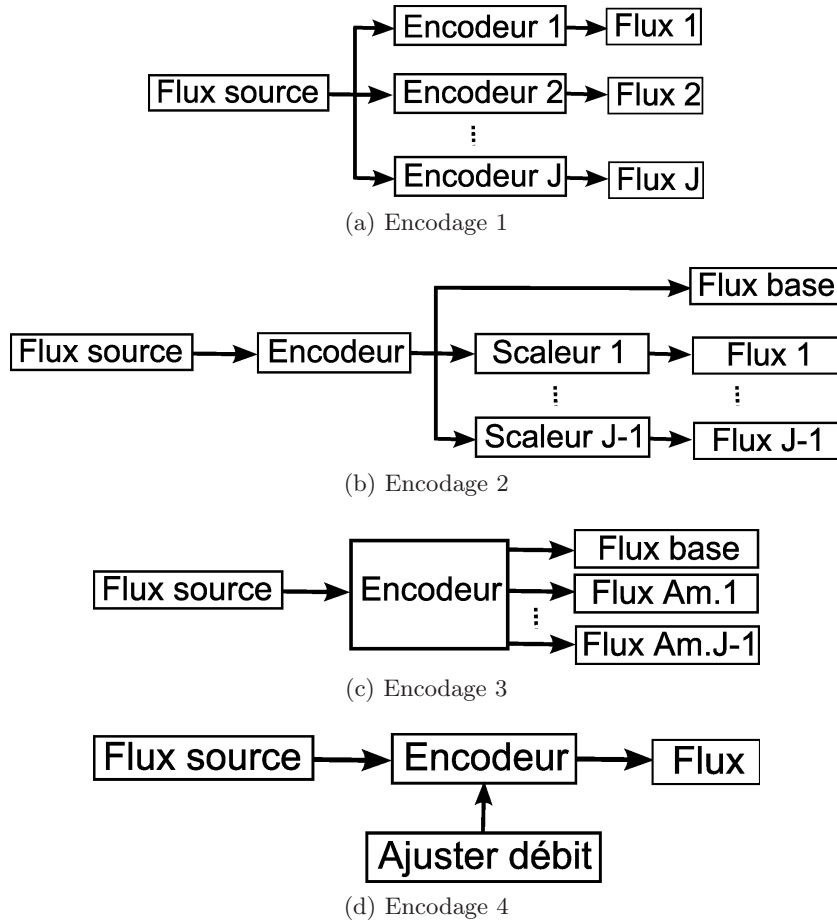


FIGURE 4.6 – Types d'encodage des données sources

4.3.1.2 Mesure de l'équité

Nous avons vu que pour les systèmes multicast à débit fixe, tous les utilisateurs partageaient les mêmes ressources ce qui avait pour conséquence l'allocation d'un même débit à tous les utilisateurs. Dans un contexte où les liens de transmission des utilisateurs sont très différents, ce débit est limité par les utilisateurs ayant les plus mauvais canaux. Dans la littérature, l'équité implique directement l'allocation égale des ressources entre les différents éléments en jeu, bien qu'il y ait peu d'accord entre les chercheurs sur ce qui devrait être réalisé. Dans les réseaux informatiques, certains recherchent un temps de retard égal, d'autres recherchent un débit égal, et encore d'autres recherchent une puissance égale pour tous les utilisateurs partageant une ressource. Dans certaines applications, des désaccords peuvent exister au sujet de la bonne métrique pour juger une allocation [73]. Une étude de quelques métriques de mesure d'équité a été faite dans [74]. L'équité ne signifie pas nécessairement une répartition égale des ressources. Dans certains cas, il est justifié de donner davantage de ressources à certains utilisateurs qu'à

d'autres. Dans de tels cas, la métrique de performance appropriée serait le ratio de la ressource allouée et le droit d'attribution de cette ressource [75]. A partir de la métrique d'équité de Jain [75], l'indice d'équité est alors défini par

$$\text{FI} = \frac{\left(\sum_{k=1}^K x_k \right)^2}{K \left(\sum_{k=1}^K x_k^2 \right)}, \quad (4.27)$$

où x_k est le rapport entre la quantité de ressources allouées à l'utilisateur k et son droit d'attribution de cette ressource.

Exemple 1 Une seule ressource doit être distribuée entre deux consommateurs A et B. Le consommateur A paie deux fois plus que le consommateur B. L'algorithme d'allocation alloue x part de la ressource à A et $1 - x$ part à B. L'indice d'équité s'écrit

$$\text{FI} = \frac{\left[\frac{x}{2} + (1 - x) \right]^2}{2 \left[\left(\frac{x}{2} \right)^2 + (1 - x)^2 \right]} \quad (4.28)$$

et vaut 1 que lorsque le consommateur A obtient les deux-tiers (deux fois plus que le consommateur B).

Dans le cas multicast, le but étant de satisfaire tous les utilisateurs, nous considérons que tous les utilisateurs ont le même droit d'attribution des ressources, soit $x_k = R_k$. Cet indice est continu en ce sens que tout changement dans l'allocation fait varier l'équité, alors qu'avec l'indice d'équité « min-max », seul un changement sur le débit minimum ou maximum peut faire varier l'indice d'équité. L'indice d'équité (4.27) peut être réécrit de la façon suivante [75]

$$\text{FI} = \frac{1}{K} \sum_{k=1}^K \frac{x_k}{x_f} \quad (4.29)$$

où x_f est la marque d'allocation équitable calculée comme suit

$$x_f = \frac{\sum_{k=1}^K x_k^2}{\sum_{k=1}^K x_k}. \quad (4.30)$$

Ainsi, chaque utilisateur compare son allocation x_k avec la quantité x_f et trouve l'algorithme d'allocation équitable, ou pas équitable, selon que le rapport x_k/x_f est plus ou

moins grand. Ce rapport x_k/x_f est l'indice d'équité pour l'utilisateur k . L'indice d'équité globale est la moyenne de l'équité perçue par tous les K utilisateurs.

Exemple 2 On souhaite distribuer la somme de 20\$ à 100 personnes. En se basant sur l'âge de ces personnes, on suppose que 10 personnes reçoivent chacune 2\$ et les 90 autres ne reçoivent rien. Dans ce contexte, la marque d'allocation équitable est

$$x_f = \frac{\sum_{k=1}^K x_k^2}{\sum_{k=1}^K x_k} = \frac{\sum_{k=1}^{10} 2^2 + \sum_{k=11}^{100} 0^2}{\sum_{k=1}^{10} 2 + \sum_{k=11}^{100} 0} = 2. \quad (4.31)$$

Ainsi, les 10 utilisateurs qui ont reçu 2\$ trouvent la méthode d'allocation 100% équitable alors que les 90 autres qui n'ont rien reçu trouvent cette allocation 0% équitable. L'équité globale est de 10%. Intuitivement, on peut déduire que l'indice d'équité globale donne le pourcentage d'utilisateurs satisfaits.

4.3.2 Séparation des destinataires dans le domaine fréquentiel

Pour augmenter le débit multicast total, une première approche consiste à séparer les destinataires dans le domaine fréquentiel [62, 65]. En effet, chaque sous-porteuse est allouée à un sous-groupe d'utilisateurs qui recevront le même symbole de données sur cette sous-porteuse. La méthode LCG est alors utilisée pour l'allocation des ressources dans le sous-groupe, c'est-à-dire que l'ordre de modulation sur chaque sous-porteuse est déterminé par l'utilisateur, dans le sous-groupe, qui présente la plus faible amplitude sur la sous-porteuse. Il a été prouvé que cette méthode apporte un gain significatif en terme de débit total comparée à la méthode LCG classique dans un contexte de communication sans fil [62, 65]. Dans ce paragraphe, nous essayerons d'adapter le problème d'optimisation à un contexte CPL et d'y ajouter par la suite la technique de précodage linéaire. Cette approche sera nommée dans la suite MDHF (système multicast à débits hétérogènes dans le domaine fréquentiel).

Soit $r_{k,n}$ le nombre de bits, dans $\mathbb{N}$, que peut charger l'utilisateur k , sur la sous-porteuse n dans le cadre de l'OFDM, ou sur la séquence n dans le cadre du LP-OFDM. En posant R_k comme le débit multicast alloué à l'utilisateur et r_n le nombre de bits chargés sur la sous-porteuse ou la séquence n , il s'en suit

$$\begin{cases} R_k = \sum_{n=1}^N r_n \rho_{k,n} \\ \rho_{k,n} = 1 \text{ si } r_{k,n} \geq r_n \text{ sinon } 0. \end{cases} \quad (4.32)$$

La variable binaire $\rho_{k,n}$ indique si l'élément n est alloué à l'utilisateur k ou non, et peut être définie par

$$\rho_{k,n} = \mathcal{U}(r_{k,n} - r_n), \quad (4.33)$$

où $\mathcal{U}$ est la fonction échelon d'Heaviside définie par

$$\mathcal{U}(x) = \begin{cases} 1 & \text{si } x \geq 0 \\ 0 & \text{si } x < 0. \end{cases} \quad (4.34)$$

4.3.2.1 Formulation du problème

Le but ici est de maximiser la somme des débits alloués aux différents utilisateurs. Le problème d'optimisation s'écrit

$$R_{\text{MDHF}} = \max_{r_n} \sum_{k=1}^K R_k = \max_{r_n} \sum_{k=1}^K \sum_{n=1}^N r_n \mathcal{U}(r_{k,n} - r_n). \quad (4.35)$$

En raison de la contrainte de DSP sur chaque sous-porteuse, la décision de l'ordre de modulation et du groupe d'utilisateurs sur chaque sous-porteuse est indépendante. En communication sans fil par contre, due à la contrainte de puissance totale, les décisions sont liées [65]. Le problème (4.35) avec une contrainte de puissance crête se réduit à

$$\max_{r_n} r_n \sum_{k=1}^K \mathcal{U}(r_{k,n} - r_n). \quad (4.36)$$

Il est évident que $R_{\text{MDHF}} \geq R_{\text{LCG}}$. En effet, pour chaque sous-porteuse ou séquence n , on a

$$\begin{aligned} \max_{r_n} r_n \sum_{k=1}^K \mathcal{U}(r_{k,n} - r_n) &\geq (\min_k r_{k,n}) \underbrace{\sum_{k=1}^K \mathcal{U}(r_{k,n} - (\min_k r_{k,n}))}_{=K} \\ &\geq K (\min_k r_{k,n}) \end{aligned} \quad (4.37)$$

4.3.2.2 Solution au problème d'optimisation

Comme la fonction $\mathcal{U}$ considérée n'est pas continue en zéro, la détermination d'un maximum local en utilisant le test classique de la dérivée première ou de la dérivée seconde s'avère être impossible. L'utilisation des expressions de la dérivation au sens des distribution pour la fonction échelon n'aboutit pas à un résultat. Une solution au

problème (4.36) s'écrit

$$r_n = r_{k^*,n} \quad (4.38)$$

où

$$k^* = \arg \max_k r_{k,n} \sum_{j=1}^K \mathcal{U}(r_{j,n} - r_{k,n}). \quad (4.39)$$

4.3.2.3 Prise en compte de l'équité entre les utilisateurs

L'approche MDHF ne considère pas l'équité entre les différents destinataires dans le processus d'allocation des ressources. En cherchant à maximiser le débit multicast total, certains utilisateurs peuvent se retrouver sans aucun débit alloué, ce qui conduit à l'insatisfaction de la qualité de service pour ces utilisateurs. Afin de respecter un minimum d'équité entre les utilisateurs tout en minimisant la dégradation du débit multicast total, il a été proposé une méthode d'allocation des ressources basée sur l'équité proportionnelle (*proportional fairness (PF), en anglais*) [65]. Cette méthode adapte le débit multicast alloué à chaque utilisateur pendant chaque intervalle de temps en fonction du débit reçu lors de l'allocation précédente. Le système multicast résultant peut alors être considéré comme un système multicast à débits hétérogènes dans le domaine temps-fréquence. Ici, nous allons adapter les équations déjà développées dans [65] pour les communications sans fils, au contexte CPL avec une extension au système LP-OFDM. Dans la suite, cette méthode sera nommée MDHF-PF.

Soit $R_k(t)$ le débit alloué à l'utilisateur k et $r_n(t)$ le nombre de bits chargés sur la sous-porteuse ou la séquence n sur l'intervalle de temps t . Pour réduire la complexité des calculs, les auteurs dans [76] ont proposé un algorithme simplifié basé sur l'équité proportionnelle en utilisant le débit moyen suivant

$$R_k(t) = (1 - \frac{1}{T_W})R_k(t-1) + \frac{1}{T_W} \sum_{n=1}^N r_n(t) \mathcal{U}(r_{k,n}(t) - r_n(t)), \quad (4.40)$$

où T_W indique la taille de la fenêtre de moyennage et t l'indice des intervalles de temps. Dans un contexte CPL, le problème d'optimisation s'écrit

$$\max_{r_n} \sum_{k=1}^K R_k(t) = \max_{r_n} \prod_{k=1}^K R_k(t) \quad (4.41)$$

sous contrainte de connaître $R_k(t-1)$. Ce problème (4.41) est asymptotiquement équi-

valent au problème suivant, lorsque T_W est grand,

$$\max_{r_n} \sum_{k=1}^K \sum_{n=1}^N \frac{r_n(t) \mathcal{U}(r_{k,n}(t) - r_n(t))}{R_k(t-1)} \quad (4.42)$$

sous contrainte de connaître $R_k(t-1)$. En effet, on a

$$\begin{aligned} \prod_{k=1}^K R_k(t) &= \underbrace{\left[\left(1 - \frac{1}{T_W}\right)^K \prod_{k=1}^K R_k(t-1) \right]}_{=c_k \text{ (constante)}} \times \prod_{k=1}^K \left(1 + \frac{\sum_n r_n(t) \mathcal{U}(r_{k,n}(t) - r_n(t))}{(T_W - 1) R_k(t-1)} \right) \\ &= c_k \left[1 + \frac{1}{T_W - 1} \sum_{k=1}^K \sum_{n=1}^N \frac{r_n(t) \mathcal{U}(r_{k,n}(t) - r_n(t))}{R_k(t-1)} \right. \\ &\quad + \left(\frac{1}{T_W - 1} \right)^2 \sum_{i \neq j} \frac{\sum_n r_n(t) \mathcal{U}(r_{i,n}(t) - r_n(t)) \sum_m r_m(t) \mathcal{U}(r_{j,m}(t) - r_m(t))}{R_i(t-1) R_j(t-1)} \\ &\quad \left. + \dots \right]. \end{aligned} \quad (4.43)$$

Pour T_W grand, les termes d'ordres supérieurs ou égaux à deux peuvent être négligés. De ce fait, on déduit que

$$\max_{r_n} \prod_{k=1}^K R_k(t) \Leftrightarrow \max_{r_n} \sum_{k=1}^K \sum_{n=1}^N \frac{r_n(t) \mathcal{U}(r_{k,n}(t) - r_n(t))}{R_k(t-1)} \quad (4.44)$$

sous contrainte de connaître $R_k(t-1)$. Spécialement pour $T_W = 1000$ proposé dans [76], le problème simplifié (4.42) est quasiment identique au problème d'origine (4.41). Toutefois, une grande valeur de T_W pourrait augmenter la latence du système et spécialement lorsque l'état de canal des utilisateurs varie rapidement, comme en communication sans fil [65]. En outre, le temps de latence peut être un élément critique en particulier pour les services multimédias. Par conséquent, il faudra examiner les performances de la méthode en prenant en compte des petites valeurs de T_W . Dans le cadre des communications sans fil, il a été prouvé que la dégradation de performance pour une grande valeur de T_W est négligeable par rapport à la performance pour une petite valeur de T_W [65].

En raison de la contrainte de DSP sur chaque sous-porteuse, la décision de l'ordre de modulation et du groupe d'utilisateurs sur chaque sous-porteuse est indépendante. Le problème (4.42), pour la sous-porteuse ou la séquence n , se réduit à

$$\max_{r_n} r_n(t) \sum_{k=1}^K \frac{\mathcal{U}(r_{k,n}(t) - r_n(t))}{R_k(t-1)}, \quad (4.45)$$

sous contrainte de connaître $R_k(t-1)$. Comme pour le cas sans équité, une solution au

problème (4.45) s'écrit

$$r_n(t) = r_{k^*,n}(t) \quad (4.46)$$

où

$$k^* = \arg \max_k r_{k,n}(t) \sum_{j=1}^K \frac{\mathcal{U}(r_{j,n}(t) - r_{k,n}(t))}{R_j(t-1)}. \quad (4.47)$$

4.3.3 Séparation des destinataires dans le domaine temporel

Pour augmenter le débit multicast total et pour réduire l'impact des mauvais liens de transmission sur le débit multicast, nous proposons une seconde approche qui consiste à séparer les destinataires dans le domaine temporel [15]. Ainsi les utilisateurs sont regroupés en sous-groupes en fonction de leur lien de transmission. Ici, un intervalle de temps est alloué à chaque sous-groupe d'utilisateurs contrairement aux méthodes vues dans les systèmes multicast à débits homogènes (LCG, LP-LCG, LBCG). Dans ces méthodes, tous les utilisateurs partagent les mêmes intervalles de temps. Ces méthodes seront utilisées pour l'allocation des ressources dans chaque sous-groupe et il est à noter qu'un utilisateur fait partie d'un seul sous-groupe. Cette approche sera nommée dans la suite MDHT (système multicast à débits hétérogènes dans le domaine temporel).

Le débit multicast total est défini comme étant la moyenne des débits sur tous les intervalles de temps utilisés et s'écrit

$$\mathcal{R}_{\text{MDHT}} = \frac{1}{G} \sum_{g=1}^G \mathcal{R}_{G_g}, \quad (4.48)$$

où G est le nombre total de sous-groupes, G_g est le g^e sous-groupe et $\mathcal{R}_{G_g}$ est le débit multicast du sous-groupe G_g . Ainsi,

$$\bigcup_g G_g = \bigcup_g \{\text{canaux de classe } C_g \text{ à canaux de classe } C_{g+1} - 1\} = \{1, 2, \dots, 9\} \quad (4.49)$$

et

$$\sum_{g=1}^G \text{card}(G_g) = K. \quad (4.50)$$

Le système multicast à débits fixes est obtenu pour $G = 1$. Le problème d'optimisation

TABLE 4.2 – Liste exhaustive des répartitions possibles des différents utilisateurs en 2 sous-groupes.

Répartition	$G_1 \cup G_2$	Répartition	$G_1 \cup G_2$
1	$\{1\} \cup \{2, \dots, 9\}$	5	$\{1, \dots, 5\} \cup \{6, \dots, 9\}$
2	$\{1, 2\} \cup \{3, \dots, 9\}$	6	$\{1, \dots, 6\} \cup \{7, 8, 9\}$
3	$\{1, 2, 3\} \cup \{4, \dots, 9\}$	7	$\{1, \dots, 7\} \cup \{8, 9\}$
4	$\{1, \dots, 4\} \cup \{5, \dots, 9\}$	8	$\{1, \dots, 8\} \cup \{9\}$

du débit multicast MDHT s'écrit

$$\max_{G_g} \mathcal{R}_{\text{MDHT}} = \max_{G_g} \frac{1}{G} \sum_{g=1}^G \mathcal{R}_{G_g}. \quad (4.51)$$

et consiste à trouver le nombre optimal G de sous-groupes et la répartition optimale des utilisateurs dans ces sous-groupes. Étant donné l'absence d'une expression donnant une relation directe entre les sous-groupes et le débit multicast réalisé, il est difficile de trouver une solution analytique au problème d'optimisation (4.51). Nous nous contenterons donc de trouver des méthodes empiriques de répartition des différents utilisateurs. Dans cette optique, nous fixons le nombre de sous-groupes et nous faisons une étude statistique du débit multicast calculé à partir de l'exploitation exhaustive des répartitions possibles. En se basant sur le nombre de classes de canaux de transmission (9 classes), deux différentes valeurs de G (nombre de sous-groupes) sont alors analysées. Les modes 2 et 3 représentent respectivement les cas où les utilisateurs sont regroupés en 2 et 3 sous-groupes. Les tableaux 4.2 et 4.3 donnent respectivement la liste exhaustive des répartitions possibles des différents utilisateurs en 2 et 3 sous-groupes.

Pour simplifier les simulations de l'étude statistique du débit multicast, nous utiliserons la capacité des canaux de transmission et on fait l'hypothèse suivante

$$C_{\text{LCG}} = C_{\min}, \quad (4.52)$$

où C_{LCG} et $C_{\min}$ sont respectivement la capacité du système avec la méthode LCG et la capacité minimale des canaux des utilisateurs. Sous l'hypothèse (4.52), la capacité du système multicast avec l'approche MDHT s'écrit alors

$$C_{\text{MDHT}} = \max_{G_g} \frac{1}{G} \sum_{g=1}^G \text{card}(G_g) C_{\min}^{G_g}. \quad (4.53)$$

La figure 4.7 présente les fonctions de répartition de la capacité du système MDHT à 9 utilisateurs et où chaque utilisateur a un canal de transmission appartenant à une

TABLE 4.3 – Liste exhaustive des répartitions possibles des différents utilisateurs en 3 sous-groupes.

Répartition	$G_1 \cup G_2 \cup G_3$	Répartition	$G_1 \cup G_2 \cup G_3$
1	$\{1\} \cup \{2\} \cup \{3, \dots, 9\}$	15	$\{1, 2, 3\} \cup \{4, 5\} \cup \{6, \dots, 9\}$
2	$\{1\} \cup \{2, 3\} \cup \{4, \dots, 9\}$	16	$\{1, 2, 3\} \cup \{4, 5, 6\} \cup \{7, 8, 9\}$
3	$\{1\} \cup \{2, 3, 4\} \cup \{5, \dots, 9\}$	17	$\{1, 2, 3\} \cup \{4, \dots, 7\} \cup \{8, 9\}$
4	$\{1\} \cup \{2, \dots, 5\} \cup \{6, \dots, 9\}$	18	$\{1, 2, 3\} \cup \{4, \dots, 8\} \cup \{9\}$
5	$\{1\} \cup \{2, \dots, 6\} \cup \{7, 8, 9\}$	19	$\{1, \dots, 4\} \cup \{5\} \cup \{6, \dots, 9\}$
6	$\{1\} \cup \{2, \dots, 7\} \cup \{8, 9\}$	20	$\{1, \dots, 4\} \cup \{5, 6\} \cup \{7, 8, 9\}$
7	$\{1\} \cup \{2, \dots, 8\} \cup \{9\}$	21	$\{1, \dots, 4\} \cup \{5, 6, 7\} \cup \{8, 9\}$
8	$\{1, 2\} \cup \{3\} \cup \{4, \dots, 9\}$	22	$\{1, \dots, 4\} \cup \{5, \dots, 8\} \cup \{9\}$
9	$\{1, 2\} \cup \{3, 4\} \cup \{5, \dots, 9\}$	23	$\{1, \dots, 5\} \cup \{6\} \cup \{7, 8, 9\}$
10	$\{1, 2\} \cup \{3, 4, 5\} \cup \{6, \dots, 9\}$	24	$\{1, \dots, 5\} \cup \{6, 7\} \cup \{8, 9\}$
11	$\{1, 2\} \cup \{3, \dots, 6\} \cup \{7, 8, 9\}$	25	$\{1, \dots, 5\} \cup \{6, 7, 8\} \cup \{9\}$
12	$\{1, 2\} \cup \{3, \dots, 7\} \cup \{8, 9\}$	26	$\{1, \dots, 6\} \cup \{7\} \cup \{8, 9\}$
13	$\{1, 2\} \cup \{3, \dots, 8\} \cup \{9\}$	27	$\{1, \dots, 6\} \cup \{7, 8\} \cup \{9\}$
14	$\{1, 2, 3\} \cup \{4\} \cup \{5, \dots, 9\}$	28	$\{1, \dots, 7\} \cup \{8\} \cup \{9\}$

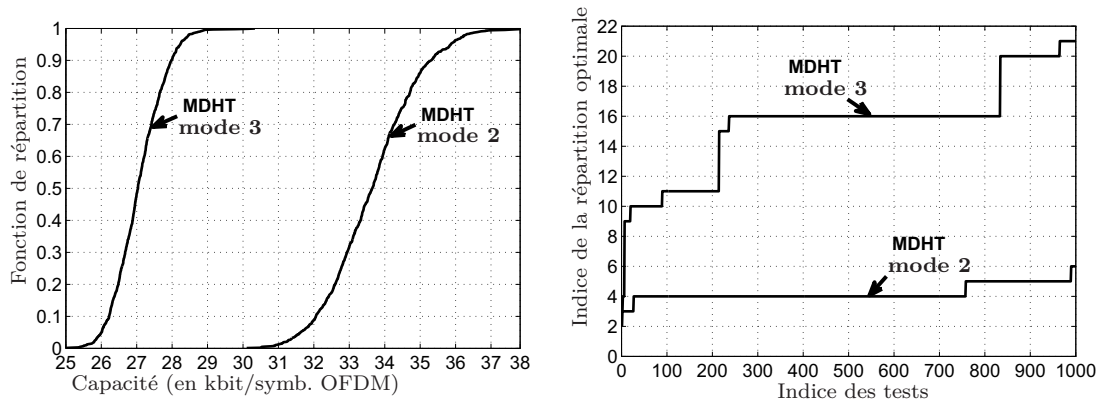


FIGURE 4.7 – (a) Fonction de répartition de la capacité du système MDHT, (b) indice de la répartition optimale des utilisateurs sous l'hypothèse (4.52).

classe différente de canaux. Ces résultats montrent que la capacité du système MDHT diminue avec le nombre de sous-groupes. Ceci peut s'expliquer par le fait que plus il y a de sous-groupes et moins il y a d'intervalles de temps alloués aux sous-groupes. Le tableau 4.4 est un exemple de l'allocation des intervalles de temps aux différents sous-groupes sur 12 intervalles de temps. Les utilisateurs dans le mode 2 reçoivent 3/2 fois plus d'intervalles de temps que les utilisateurs dans le mode 3.

La figure 4.7 présente aussi l'indice des répartitions optimales des utilisateurs dans les sous-groupes en fonction des différents modes (voir tableaux 4.2 et 4.3). Dans le mode 2, la répartition n° 4 (class $\{1, \dots, 4\} \cup \{5, \dots, 9\}$) est optimale dans plus de 73% des cas

TABLE 4.4 – Allocation des intervalles de temps (IT) aux différents sous-groupes sur 12 intervalles de temps pour les mode 2 et 3.

IT	1	2	3	4	5	6	7	8	9	10	11	12
Mode 2	G21	G22	G21	G22	G21	G22	G21	G22	G21	G22	G21	G22
Mode 3	G31	G32	G33	G31	G32	G33	G31	G32	G33	G31	G32	G33

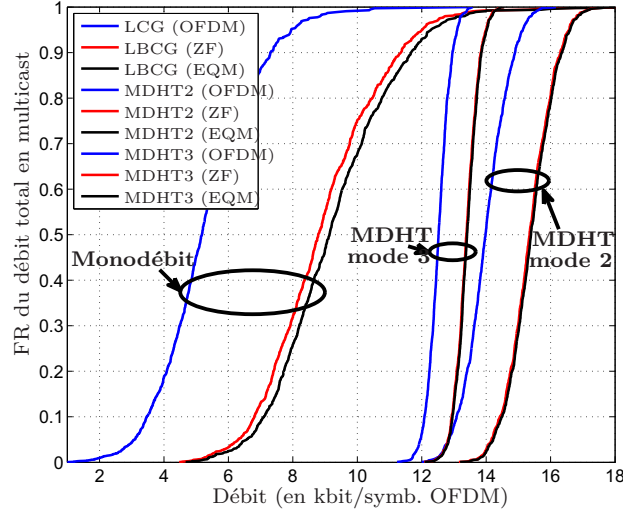
testés. Dans le mode 3, c'est la répartition n° 16 (class $\{1, 2, 3\} \cup \{4, 5, 6\} \cup \{7, 8, 9\}$) qui est optimale dans plus de 59% des cas testés. A partir de ces observations, nous utiliserons ces répartitions dites « optimales » comme répartitions empiriques des utilisateurs en fonction de chaque mode.

4.4 Simulations et performances

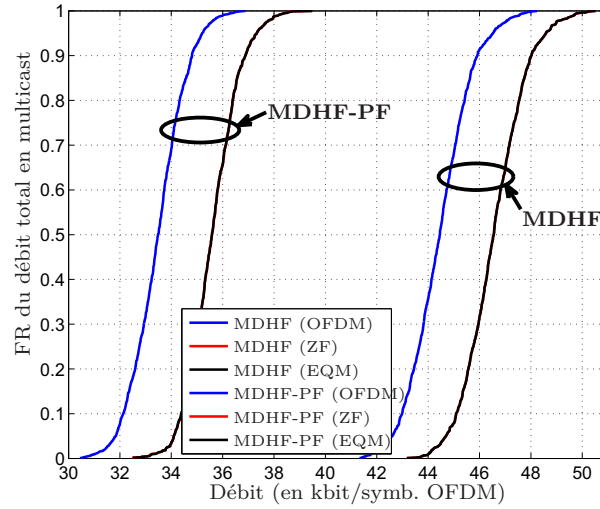
Le but de ce paragraphe est alors d'évaluer précisément les performances des algorithmes présentés, et particulièrement de comparer les résultats obtenus pour les systèmes LP-OFDM avec ceux obtenus avec les méthodes d'allocation classique en multicast. L'ensemble des simulations est mené à partir des modèles de canaux CPL présentés dans le chapitre 1. Ces canaux ont été répartis dans neuf classes différentes et un modèle de fonction de transfert est associé à chacune des classes. Comme nous l'avons vu au chapitre 1, dans une même classe, les capacités des canaux sont variables et peuvent atteindre des différences de l'ordre de 400 Mb/s. Il devient alors insignifiant de chercher à caractériser les performances des algorithmes sur un seul cas, ce qui traduirait assez mal la réalité des choses. Pour tenir compte de ces variations, on préfère utiliser la fonction de répartition (FR) des débits maximaux atteignables en multicast qui exprime la probabilité qu'un débit réalisable soit inférieur à une certaine valeur. Afin d'évaluer les différentes fonctions de répartitions, les simulations seront réalisées sur un grand nombre de réalisations de ces canaux CPL. Pour cela, nous nous limitons à un nombre plus faible de sous-porteuses par rapport aux chapitres précédents à savoir 1024 sous-porteuses prises dans la bande [2; 27] MHz. Les simulations effectuées dans un contexte similaire à celui des chapitres précédents (bande de fréquence étendue à 100 MHz, prise en compte du masque de puissance) montrent des tendances identiques à celles présentées dans ce paragraphe [15]. Les différents paramètres de simulation sont alors résumés dans le tableau 4.5 et serviront également pour le dernier chapitre. Par ailleurs, la probabilité d'erreur symbole cible sera fixée à 10^{-3} en sortie de l'égaliseur. Enfin, rappelons qu'on se place ici sous l'hypothèse d'une connaissance parfaite du canal de transmission à l'émission. Le système est aussi considéré comme étant parfaitement synchronisé.

TABLE 4.5 – Principaux paramètres de simulation.

Bande passante	Δ_f	N_{utile}	E_{DSP}	N_0	Bitcap
[2; 27] MHz	24,414 kHz	1024	-50 dBm/Hz	-110 dBm/Hz	14



(a) Domane temporel



(b) Domane fréquentiel

FIGURE 4.8 – Fonction de répartition du débit total en multicast multidébit, (a) dans le domaine temporel et (b) dans le domaine fréquentiel pour $L = 32$.

4.4.1 Étude d'un système multicast à 9 utilisateurs

Dans ce paragraphe, nous considérons un système multicast comprenant 9 utilisateurs et chaque utilisateur possède un canal de transmission pris dans une classe différente.

Dans un premier temps, la longueur des séquences de précodage est fixée à 32. La figure 4.8 donne les fonctions de répartition obtenues pour les débits totaux en multicast avec les différents algorithmes. Pour une question de lisibilité des courbes, ces fonctions de répartition ont été regroupées sur deux figures. La figure 4.8a présente les résultats dans un contexte multicast multidébit dans le domaine temporel. Il faut noter que les méthodes (LCG, LBCG (ZF/EQM)) développées dans le contexte multicast monodébit sont ici perçues comme des cas particuliers du contexte MDHT avec un seul groupe d'utilisateurs. La figure 4.8b présente les résultats dans un contexte multicast multidébit dans le domaine fréquentiel. L'analyse de ces figures fait ressortir trois points.

Premièrement, comme dans le contexte mono-utilisateur, les résultats obtenus pour le LP-OFDM, quel que soit le critère d'égalisation choisi, sont supérieurs à ceux obtenus pour l'OFDM. À la valeur 0,5 de la fonction de répartition, on relève un gain en débit d'au moins 70% pour les méthodes basées sur la technique de précodage (LBCG) par rapport à la méthode classique basée sur l'OFDM (LCG). Ce gain est moins important (10% au maximum) dans un contexte multicast multidébit.

Deuxièmement, l'utilisation de la technique d'égalisation suivant le critère EQM offre de meilleures performances par rapport au critère ZF. L'amélioration du débit total atteint jusqu'à 8% de gain dans le contexte monodébit, ce qui est meilleur que les gains obtenus en mono-utilisateur. Cette amélioration est par contre faible dans le contexte multicast multidébit.

Troisièmement, on note également que le passage d'un système multicast à débit fixe à un système multicast à débit variable permet d'accroître le débit total réalisable. Malgré le fort gain apporté par la technique de précodage dans le contexte multicast monodébit, le débit final reste limité par l'état du canal des mauvais utilisateurs.

- En séparant les utilisateurs dans le domaine fréquentiel, on réalise un gain en débit d'au moins 270% avec la méthode MDHF-PF. Il faut noter que cette méthode (MDHF-PF) prend en compte l'équité entre les utilisateurs en adaptant les débits reçus par chaque utilisateur à chaque intervalle de temps. Par conséquent, le débit total réalisable est inférieur de 10 kbit/symbole OFDM, soit 35% de perte, par rapport à la méthode MDHF simple (sans prise en compte de l'équité). Dans notre contexte de simulation où les canaux des utilisateurs sont très différents, cette dernière approche offre le meilleur débit total, sans allouer de ressources à certains utilisateurs [15]. En effet, en s'intéressant au débit minimum reçu par les utilisateurs, on remarque que le débit minimum est presque nul avec la méthode MDHF (cf. figure 4.9). Par conséquent, les exigences de qualité de service de tous les utilisateurs ne sont pas garanties et le service multicast n'est plus assuré. Par contre, la méthode MDHF-PF, en adaptant les débits, améliore le niveau du débit minimum et par la suite l'équité entre les utilisateurs.

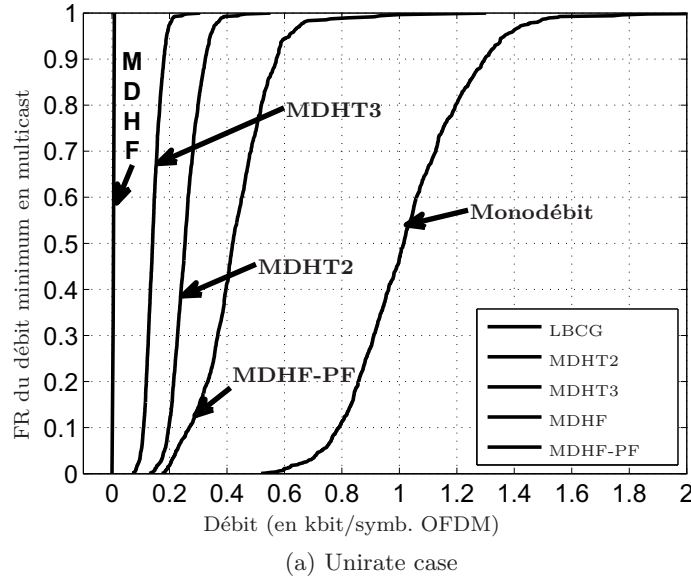


FIGURE 4.9 – Fonction de répartition du débit minimum en multicast multidébit pour $L = 32$.

- En séparant par contre les utilisateurs dans le domaine temporel, le gain par rapport au système multicast à débit fixe est moins important. On remarque même que dans moins de 1% des cas, le débit réalisé avec la méthode LBCG est supérieur aux débits des méthodes MDHT. Le débit réalisé dans le mode 2 (2 sous-groupes d'utilisateurs) est plus important que celui réalisé dans le mode 3 (3 sous-groupes d'utilisateurs). Cela s'explique par le fait que les utilisateurs reçoivent les données dans plus d'intervalles de temps dans le mode 2 que dans le mode 3.

La figure 4.9 montre la fonction de répartition du débit minimum alloué aux utilisateurs. Ces fonctions de répartition sont tracées à partir des résultats obtenus pour le LP-OFDM avec le critère EQM. On note que le système multicast monodébit assure le meilleur débit minimum aux utilisateurs (qui est aussi le débit reçu par tous les utilisateurs). La méthode MDHF-PF assure un bon niveau de débit minimum aux utilisateurs et aussi un bon niveau d'équité comme nous le verrons dans les discussions sur l'indice d'équité.

4.4.1.1 Influence de la longueur des séquences de précodage

Nous allons à présent mettre en évidence le gain en débit apporté par la composante de précodage linéaire. Pour cela nous nous intéressons aux résultats obtenus lorsque L varie, où L est un diviseur de N . On rappelle que la recherche de la longueur optimale des séquences de précodage conduit à un problème d'optimisation combinatoire complexe qui

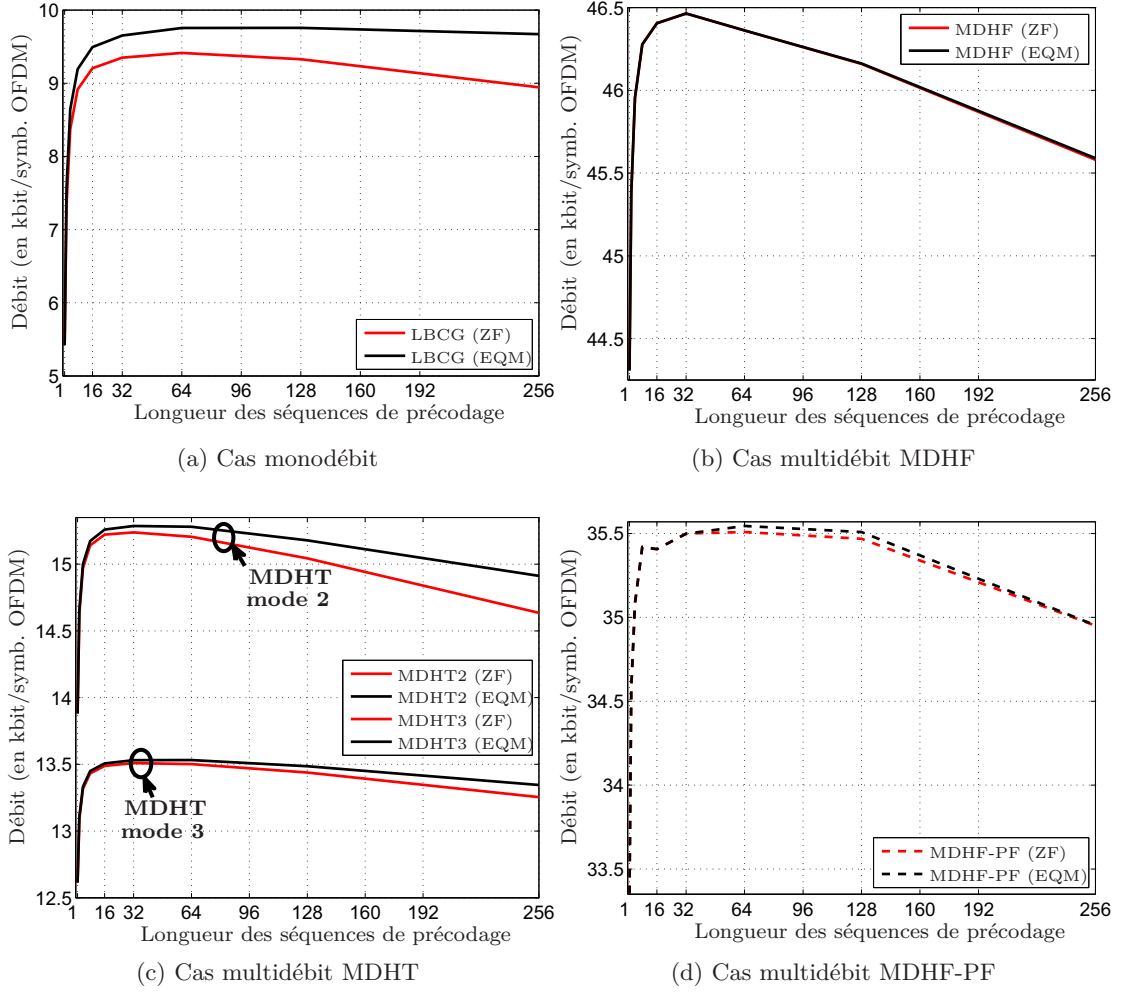


FIGURE 4.10 – Débit total en kbit par symbole OFDM versus la longueur des séquences de précodage.

ne peut être réduit à un problème convexe équivalent. Ainsi, il n'existe pas de solution analytique donnant la longueur optimale, cette dernière ne pouvant être obtenue que par une recherche exhaustive. La figure 4.10 présente l'évolution du débit total en fonction de la longueur des séquences de précodage. Le débit total obtenu avec les méthodes basées sur l'OFDM est donné pour $L = 1$.

Il est clair que le débit total réalisable avec les méthodes basées sur le LP-OFDM, quel que soit le critère d'égalisation, est amélioré lorsque la longueur des séquences de précodage est supérieure à 1. Ce débit total atteint un certain palier pour $L \geq 32$ sauf dans un contexte multicast multidébit avec séparation des utilisateurs dans le domaine fréquentiel où le débit commence à diminuer. Ces courbes confirment ainsi les analyses faites précédemment sur les performances des détecteurs ZF et EQM. Sur la base de ces

résultats, nous pouvons affirmer que l'utilisation de la technique de précodage linéaire augmente le débit des systèmes multicast. La raison de la meilleure performance des systèmes LP-OFDM par rapport aux systèmes OFDM est l'utilisation efficace de la limite de DSP. La composante de précodage accumule les énergies résiduelles d'un bloc donné de sous-porteuses pour transmettre des bits supplémentaires. Une tendance claire quant à l'évolution du débit total avec la longueur des séquences ne peut pas être dégagée. Néanmoins, on constate que pour $L \geq 64$, le débit total commence à diminuer.

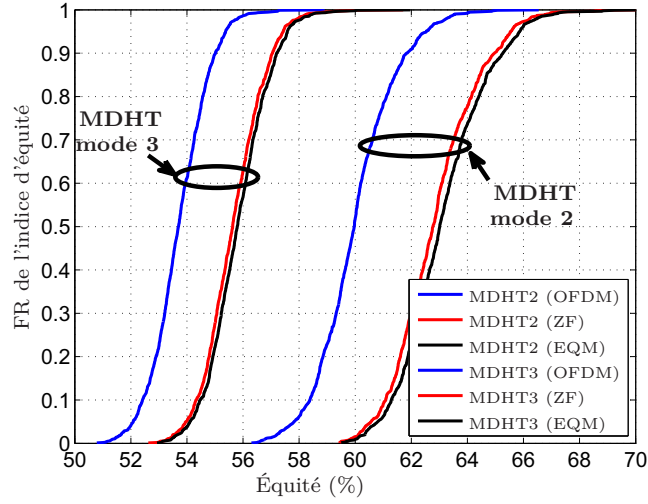
4.4.1.2 Étude de l'équité entre les utilisateurs

Les méthodes d'allocation des ressources dans le contexte multicast monodébit allouent de façon équitable les ressources et tous les utilisateurs reçoivent *in fine* le même débit. Lorsque les états des canaux des utilisateurs sont très différents, il est justifiable d'allouer davantage de ressources à certains utilisateurs qu'à d'autres [75]. C'est cette idée qui sous-tend la thèse de séparation des utilisateurs en fonction de leurs canaux de transmission. Les algorithmes proposés dans un tel contexte ont déjà démontré leurs meilleures performances en termes de débit total en multicast multidébit. Il est également nécessaire de mesurer leur performance en termes d'équité entre les utilisateurs. Comme un indicateur de performance, l'indice d'équité défini à l'équation (4.27) sera utilisé. Cet indice mesure l'égalité de la répartition des ressources entre les utilisateurs.

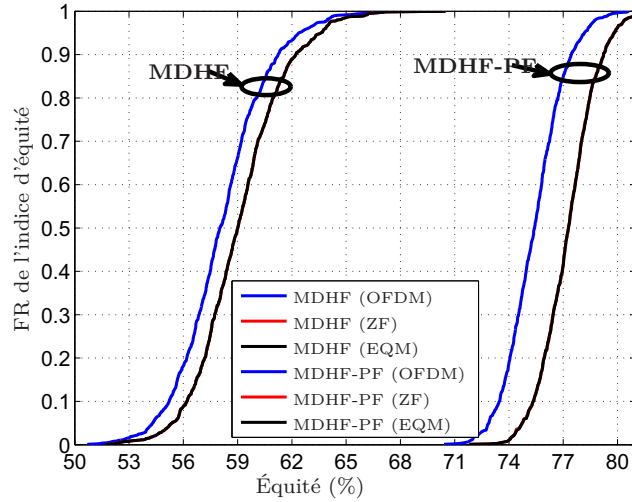
La figure 4.11 montre la fonction de répartition de l'indice d'équité des utilisateurs en multicast multidébit. Cet indice est égal à 1 pour les méthodes proposées dans le contexte multicast monodébit et n'est donc pas représenté. Les méthodes basées sur la technique de précodage, en plus d'apporter un gain en débit total, améliore l'équité entre les utilisateurs. Cette amélioration est plus importante avec les méthodes MDHT (allant jusqu'à 3%) qu'avec les méthodes MDHF (inférieure à 2%). L'ordre d'équité des méthodes est la suivante : MDHF-PF, MDHT2, MDHF et MDHT3. Lorsque les utilisateurs sont séparés dans le domaine fréquentiel, la perte de débit (environ 35%) engendrée par la méthode MDHF-PF (avec prise en compte de l'équité) par rapport à la méthode MDHF (sans prise en compte de l'équité) est compensée par l'indice d'équité. En effet, on passe d'une équité minimale de 50,7% pour la méthode MDHF à 70,4% pour la méthode MDHF-PF. En outre, la séparation en temps des utilisateurs en deux sous-groupes reste plus équitable que la séparation en trois sous-groupes.

4.4.2 Évolution du débit avec le nombre d'utilisateurs

Dans cette partie, nous étudions les débits réalisables en multicast en fonction du nombre d'utilisateurs. Chaque utilisateur utilise aléatoirement un canal parmi les 9 classes de canaux. Les résultats de simulations sont donnés uniquement pour le sys-



(a) Domaine temporel



(b) Domaine fréquentiel

FIGURE 4.11 – Fonction de répartition de l'indice d'équité entre les utilisateurs, (a) dans le domaine temporel et (b) dans le domaine fréquentiel pour $L = 32$.

tème LP-OFDM avec une égalisation de type ZF. Au vu des résultats obtenus plus haut, les commentaires qui seront faits restent valables pour les autres systèmes. La figure 4.12 donne un exemple de débits totaux et de débits minimums réalisables avec les différentes solutions proposées en fonction du nombre d'utilisateurs. Les observations faites pour le système multicast à 9 utilisateurs sont encore vraies lorsque le nombre d'utilisateurs augmente. De plus, le débit total pour le système multicast monodébit stagne voire décroît avec le nombre d'utilisateurs. Ceci confirme les résultats obtenus au paragraphe 4.2.4 sur le comportement asymptotique du débit total. En outre, la méthode MDHF offre le

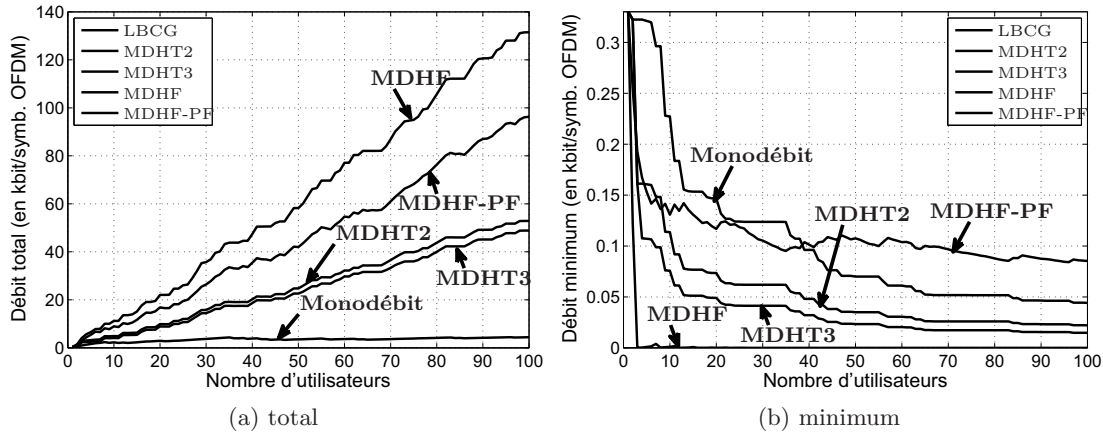


FIGURE 4.12 – Évolution des débits (débit (a) total et (b) minimum) en fonction du nombre d'utilisateurs pour $L = 32$.

meilleur débit total mais le débit minimum devient très vite nul. La méthode MDHF-PF devient plus performante que la méthode LBCG à la fois en termes de débit total et de débit minimum pour $K \geq 39$ utilisateurs.

4.4.3 Discussion sur l'*overhead* des solutions proposées

Pour terminer cette partie, il est intéressant de s'intéresser à la complexité de mise en œuvre des algorithmes proposés pour les systèmes multicast. Outre la comparaison des performances en termes de débit et d'indice d'équité, la surcharge de signalisation en liaison descendante est aussi prise en compte dans ce paragraphe. Dans les systèmes multicast à débit fixe, seule l'information sur l'ordre de modulation sur chaque sous-porteuse ou chaque séquence de précodage doit être transmise aux utilisateurs. L'utilisation de la technique de précodage engendre de plus la signalisation de la répartition des sous-porteuses dans les différents blocs. Quant aux systèmes multicast à débits hétérogènes, en plus des informations sur l'ordre de modulation et de la répartition éventuelle des sous-porteuses, l'information sur les différents sous-groupes d'utilisateurs doit être transmise. Avec les méthodes MDHF, les sous-groupes d'utilisateurs ne sont pas les mêmes d'une sous-porteuse à une autre. Ainsi, on en déduit que la surcharge de signalisation avec ces méthodes est plus élevée que pour les autres méthodes.

En outre, sous l'hypothèse que toute combinaison de flux de données peut être décodée au niveau du récepteur, un algorithme de *mapping* intelligent permettant de récupérer efficacement les données d'origine est nécessaire [65]. Cela pourrait donc apporter des surcharges de signalisation supplémentaires.

4.5 Conclusion

Dans ce chapitre, nous avons traité le problème de l'allocation des ressources dans un contexte multicast. Les problèmes de maximisation du débit multicast dans les contextes monodébit et multidébit ont été abordés sous la contrainte d'une limitation de DSP. Nous avons montré que le débit offert par la méthode classique d'allocation des ressources en OFDM multicast est limité par le débit du plus mauvais utilisateur. Nous avons aussi rappelé qu'afin de mieux correspondre aux conditions des liens et d'augmenter le débit total en multicast, il a été proposé d'utiliser le concept de multicast multidébit pour les systèmes OFDM. L'utilisation de l'approche multicast multidébit a permis d'augmenter significativement le débit multicast total par rapport à la méthode classique mais au prix d'une dégradation de l'équité entre les utilisateurs. Nous avons donc proposé de tirer profit des bonnes performances de la technique de précodage linéaire dans un contexte multicast afin d'augmenter les débits des utilisateurs et d'améliorer l'indice d'équité. Pour exploiter la diversité des liens de transmission des utilisateurs, nous avons proposé de nouvelles méthodes basées sur la technique de modulation LP-OFDM. Les algorithmes ont été développés pour un système LP-OFDM avec une mise en œuvre de l'égalisation suivant les critères ZF et EQM.

Lors des simulations, nous avons mis en évidence le fait que la mise en œuvre de la technique de précodage linéaire permet à la fois d'augmenter les débits totaux et d'améliorer l'équité entre les utilisateurs. Le gain en débit a été évalué à environ 70% et 10% respectivement pour les systèmes multicast monodébit et multidébit par rapport aux méthodes basées sur l'OFDM.

On a montré aussi que l'utilisation de la technique d'égalisation suivant le critère EQM offrait de meilleures performances que le critère ZF. L'amélioration du débit total pouvant atteindre jusqu'à 8% de gain dans le contexte monodébit, reste faible dans le contexte multicast multidébit. Cependant, du fait de la complexité des expressions de débit avec le critère EQM, on peut se limiter à mettre en œuvre les procédures d'allocation des ressources avec un détecteur de type ZF.

L'ensemble de ce travail a fait l'objet de contributions dans des conférences internationales [77, 78] et nationales [79, 80], et d'une soumission d'article de revue [81]. Un rapport technique dans le cadre du projet OMEGA a également été publié sur ce sujet [15].

Chapitre 5

Allocation des ressources dans un contexte d'accès

Ce dernier chapitre concerne le problème d'allocation des ressources dans un contexte d'accès où plusieurs utilisateurs partagent le même support physique de transmission. Contrairement au cas multicast vu au chapitre précédant, les données des utilisateurs sont différentes. Au cours de ces dernières années, plusieurs activités de recherche liées aux problèmes d'allocation des ressources pour les systèmes à porteuses multiples dans un contexte multi-utilisateur ont été menées. Pourtant, c'est un sujet qui continue de générer de nombreuses publications scientifiques consacrées à la résolution d'une grande variété de problèmes. Le problème d'allocation des ressources pour les systèmes à porteuses multiples, ou la façon de répartir les ressources temps-fréquences, les bits et les puissances, peut être défini de différentes manières selon les exigences du système mis en œuvre. Des problèmes d'allocation des ressources existent aussi bien pour les communications en voie montante ou descendante avec des contraintes de puissance, des contraintes de débits, des contraintes de qualité de services ou d'autres contraintes supplémentaires. Fondamentalement, il existe deux principaux objectifs d'optimisation qui sont poursuivis par la plupart des chercheurs, auxquels sont associées des contraintes liées à la demande spécifique ou au matériel utilisé : l'allocation des ressources en vue de la minimisation de la puissance et l'allocation des ressources en vue de la maximisation du débit. Ces grands objectifs sont habituellement mis en œuvre conjointement avec des contraintes supplémentaires telles que les limites de puissances maximales imposées par sous-canal (masques de densité spectrale de puissance). En outre, il est quelquefois préférable d'allouer les sous-canaux par groupe et non individuellement aux utilisateurs.

Une fois le problème formulé, l'algorithme de résolution doit aussi répondre à certaines exigences de complexité en fonction du système où il sera implanté. Ainsi, l'algorithme peut être centralisé ou distribué. Rappelons qu'un environnement centralisé

est un environnement dans lequel sont disponibles des informations globales sur tous les utilisateurs dans le système. Les utilisateurs sont administrés par une unité centrale de traitement, qui collecte ces informations. Par contre, dans un environnement décentralisé, il n'y a aucune unité centrale de traitement, les utilisateurs administrent le réseau par eux-mêmes. Les algorithmes de communication dans ce cas sont dit distribués.

Les systèmes actuels de communication par courant porteur sont caractérisés par des procédés d'accès multiple quelque peu aléatoires et traités par les protocoles de couches supérieures. Ces techniques sont plus ou moins liées à un accès multiple en temps en ce sens que les différents utilisateurs transmettent leurs signaux dans des intervalles de temps distincts. Les normes actuelles mettent en œuvre la technique TDMA (*time domain multiple access*) pour les communications nécessitant une garantie de qualité de service et le protocole CSMA/CA (*carrier sense multiple access collision avoidance*) pour les autres types de communications. Ces méthodes sont assez simples à implémenter mais les débits obtenus ne sont pas optimaux [82]. On peut s'attendre à avoir de meilleures performances si les utilisateurs sont autorisés à transmettre simultanément sur un même support physique. Dans ce chapitre, nous nous intéressons dans un premier temps à l'étude des algorithmes d'allocation de ressources suivant l'approche centralisée dans un réseau CPL. Nous considérons alors un réseau de communications CPL avec une structure centralisée. Dans cette structure, la répartition des ressources entre les utilisateurs du réseau est réalisée par un nœud central. Ainsi, ce nœud central, doté d'une capacité de calcul élevée, est en mesure d'allouer les différentes ressources. Nous ferons une synthèse des algorithmes centralisés en OFDMA (orthogonal frequency division multiple access) avant de proposer de nouveaux algorithmes d'allocations de ressources. Ces nouveaux algorithmes utilisent conjointement les politiques d'allocation en multi-utilisateur et une adaptation des algorithmes déjà développés au chapitre 3 dans le contexte mono-utilisateur.

Dans un second temps, nous étudions l'approche distribuée de l'allocation des ressources. Nous introduisons succinctement quelques concepts de la théorie des jeux avant d'étudier le problème de minimisation de la puissance de transmission sous les contraintes de débits minimums et de DSP.

5.1 Allocation centralisée des ressources

En télécommunication, un système centralisé est un système dans lequel la plupart des communications sont administrées par un ou plusieurs grands contrôleurs centraux. Un tel système permet la concentration de certaines fonctions (exemple, les processus d'allocation des ressources) au niveau des nœuds centraux déchargeant ainsi les nœuds terminaux de ces fonctions. Dans cette première partie, nous nous intéressons

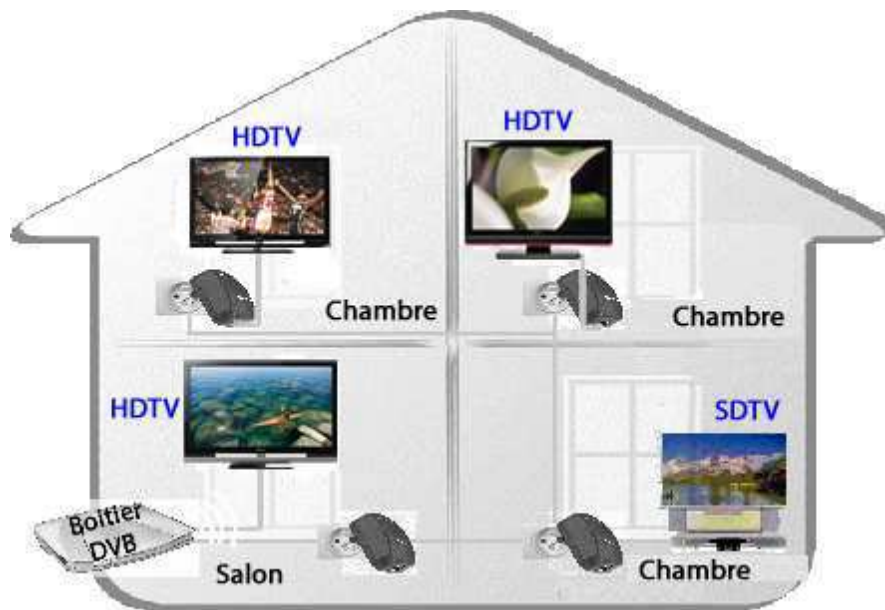


FIGURE 5.1 – Un exemple de communications centralisées dans un réseau CPL.

aux problèmes d'allocation centralisée des ressources pour les communications en voie descendante dans un réseau CPL. Le problème est identique pour les communications en voie montante ou multipoint-à-multipoint à partir du moment où les communications partagent les mêmes ressources. Un scénario de communications centralisées peut être illustré par le cas *multiroom* vidéo donné figure 5.1. Dans ce contexte, le modem CPL central (au salon) diffuse les différents flux vidéos aux autres modems (dans les chambres) selon leurs demandes. Il faut noter que le modem central a, avec les produits CPL actuels, une connaissance de l'état de chaque lien de transmission entre lui et les autres modems. Comme annoncé plus haut, le modem central utilise un accès multiple de type TDMA. Ici, nous étudions le cas d'un accès multiple de type LP-OFDMA, combinant l'OFDMA et la technique de précodage linéaire.

5.1.1 Problèmes d'allocation des ressources en OFDMA : état de l'art

L'OFDMA est l'extension à plusieurs utilisateurs du système de modulation numérique OFDM. L'accès multiple est réalisé en attribuant différents sous-ensembles de sous-canaux disjoints à chaque utilisateur. Contrairement au cas multicast étudié dans le chapitre précédent, un sous-canal ne peut être alloué qu'à un seul utilisateur. Comme les fonctions de transfert sont différentes pour chaque utilisateur, il peut arriver que des sous-canaux soient très atténués pour un utilisateur et très bons pour les autres utilisateurs. Chaque utilisateur est alors autorisé à transmettre ses données en utilisant certains sous-canaux en fonction de l'état de son canal de transmission et aussi des contraintes

du système. Ainsi, la distribution des sous-canaux entre les utilisateurs est une question clé dans le processus d'allocation des ressources. Après la répartition des sous-canaux, les algorithmes de *bit-loading* déterminent les meilleures allocations des bits et des puissances sur les sous-ensembles de sous-canaux de chaque utilisateur. Le débit total du système s'exprime alors comme la somme des débits R_k relatifs à chaque utilisateur k , et s'écrit

$$R_{\text{Total}} = \sum_{k=1}^K R_k = \sum_{k=1}^K \sum_{i \in \Omega_k} r_{k,i}, \quad (5.1)$$

où Ω_k indique le sous-ensemble de sous-canaux attribué à l'utilisateur k et $r_{k,i}$, défini à l'équation (2.15), est le nombre de bits chargé sur le sous-canal i . On note I le nombre total de sous-canaux.

5.1.1.1 Critères d'optimisation

Dans la littérature, l'allocation des ressources a été généralement étudiée suivant deux différentes approches qui mènent à deux familles d'algorithmes.

1. Allocation des ressources en vue de la minimisation de la puissance : un problème qui met l'accent sur la réduction de la puissance d'émission pour un débit donné. Ce problème peut être vu comme un problème de maximisation de la marge du système sous contrainte d'un débit cible. Dans certains systèmes, les utilisateurs ont besoin de transmettre un débit précis à un certain niveau de probabilité d'erreur. Les algorithmes d'allocation répartissent alors les bits et les puissances de manière à minimiser la puissance utilisée [83–88] ;
2. Allocation des ressources en vue de la maximisation du débit : un problème d'optimisation dont l'objectif est la maximisation du débit global (ou des débits individuels). Les contraintes sont placées sur les puissances d'émission qui peuvent être considérées par utilisateur (pour la liaison montante) ou pour tous les utilisateurs (pour la liaison descendante) [33, 89–92].

D'autres critères d'optimisation qui découlent de ces deux principaux problèmes existent aussi dans la littérature. On peut citer, par exemple, la minimisation de la probabilité de coupure, qui consiste à minimiser le nombre d'utilisateurs insatisfaits [93, 94]. Dans ce document, nous allons nous intéresser au problème de maximisation du débit du système sous les contraintes de DSP et de qualité de service (probabilité d'erreur, débit minimum).

5.1.1.2 Maximisation du débit sous contrainte en OFDMA

Sous les contraintes de DSP et de QoS, les sous-canaux, les bits et les puissances doivent être répartis entre les différents utilisateurs pour maximiser le débit du système.

Le tableau 5.1 donne les différentes stratégies de maximisation du débit dans un contexte multi-utilisateur.

TABLE 5.1 – Stratégies de maximisation du débit.

	Objectif	Avantage	Désavantage
Max somme des débits	$\max \sum_{k=1}^K R_k$	Meilleur débit total	Aucune équité entre les utilisateurs
Max somme des débits pondérés	$\max \sum_{k=1}^K w_k R_k$	Répartition ajustable des débits en fonction des poids	Pas de garantie de satisfaction pour tous les utilisateurs
Max débit minimum	$\max \min_k R_k$	Débit équivalent pour tous les utilisateurs	Pas de flexibilité dans la répartition des débits

La première stratégie consiste à maximiser la somme des débits et est mise en œuvre par les algorithmes dits opportunistes. Le principe d'allocation de ces algorithmes est d'attribuer chaque sous-canal à l'utilisateur qui y présente le meilleur gain. *In fine*, l'utilisateur qui voit un meilleur canal de transmission est favorisé par rapport aux autres utilisateurs sans tenir compte de leurs exigences en terme de débit [95]. Cette stratégie est appropriée pour les services de type « *best effort* » où aucune qualité de service n'est garantie aux utilisateurs. Pour les applications nécessitant un certain niveau de qualité de service, par exemple un débit minimum, les algorithmes doivent faire un compromis entre la maximisation du débit total et la satisfaction des exigences des utilisateurs. La deuxième stratégie consiste donc à maximiser la somme pondérée des débits, où les poids sont définis en fonction des exigences de chaque utilisateur. La répartition des débits en vue de la maximisation du débit total est alors ajustable en fonction des poids alloués aux utilisateurs [33]. Néanmoins, cette stratégie ne garantit pas à tous les utilisateurs la satisfaction de leurs exigences comme nous allons le voir dans les résultats de simulation. Nous détaillerons les procédures d'optimisation pour une telle stratégie dans le paragraphe suivant. Quant à la troisième stratégie, elle consiste à maximiser le débit minimum des utilisateurs. Cette stratégie permet de traiter les différents utilisateurs de manière équitable mais reste moins flexible quant à la répartition des débits. L'idée générale de l'algorithme d'allocation est d'attribuer les différents sous-canaux du spectre un par un aux différents utilisateurs, de manière à ce que chacun puisse tour à tour accroître son débit [90, 96].

5.1.1.3 Maximisation du débit total sous contraintes de débits minimums

Dans ce paragraphe, nous nous intéressons plus particulièrement au problème de maximisation du débit total sous les contraintes de débits minimums et de densité spectrale de puissance. Dans le contexte CPL avec une limite de DSP E_{DSP} , le problème d'optimisation peut s'écrire

$$\max_{P_{k,i}, \Omega_k} \sum_{k=1}^K \sum_{i \in \Omega_k} \log_2 \left(1 + \frac{1}{I} P_{k,i} \phi_{k,i} \right), \quad (5.2)$$

sous les contraintes de DSP $P_{k,i} \leq E_{\text{DSP}}$ et de débits minimums $R_k \geq R_{\min}^k$. La plupart des solutions algorithmiques proposées pour la résolution de ce problème comprennent deux étapes. La première étape consiste à déterminer les sous-ensembles de sous-canaux alloués à chaque utilisateur. La seconde étape consiste à choisir la répartition des bits et des puissances sur ces sous-ensembles de sous-canaux. Dans un contexte de limite de puissance crête par sous-canal, la solution optimale consiste à allouer la puissance maximale autorisée sur chaque sous-canal. La répartition des bits découle de l'allocation des puissances, puisque les débits individuels transmis sur chaque sous-canal sont reliés de manière univoque à la répartition des puissances. Nous nous limitons alors à présenter les algorithmes de la première étape. Plusieurs algorithmes existent dans la littérature plus ou moins similaires [97–99]. Nous développons les algorithmes proposés dans [92, 100] pour les système OFDM qui serviront de base pour la suite de l'étude avec les systèmes LP-OFDM.

5.1.1.4 Détermination des sous-ensembles de sous-canaux

La détermination des sous-ensembles de sous-canaux consiste dans un premier lieu à déterminer le nombre de sous-canaux à attribuer à chaque utilisateur et ensuite à affecter les sous-canaux spécifiques aux utilisateurs. Soit I_k le nombre de sous-canaux alloués à l'utilisateur k tel que $\text{card}(\Omega_k) = I_k$.

Phase 1 : Allocation du nombre de sous-canaux par utilisateur

Suivant les algorithmes, la détermination de I_k peut être ou non indépendante de la phase d'affectation des sous-canaux. Dans [92], on parle alors d'estimation souple lorsque $\sum_k I_k \leq I$ et d'estimation rigide lorsque $\sum_k I_k = I$. Dans le cas de l'estimation souple, le reste $I^r = I - \sum_{k=1}^K I_k$ des sous-canaux est alloué à la phase d'affectation des sous-canaux. On parle par contre d'estimation rigide lorsque la répartition $\{I_k\}_{1 \leq k \leq K}$ obtenue est utilisée sans modifications dans la phase d'affectation des sous-canaux. Le

nombre initial de sous-canaux alloués à l'utilisateur k est

$$I_k = \lfloor w_k I \rfloor , \quad (5.3)$$

où w_k est un poids attribué à l'utilisateur k . Une modification apportée à l'algorithme BARE (*bandwidth allocation on rate estimation*) proposé dans [92] permet de fixer w_k en fonction des débits minimums requis R_{est}^k et de l'état des canaux de transmission. À partir du rapport signal à bruit moyen $\overline{\text{rsb}}_k$ sur tout le canal de transmission, l'algorithme calcule un débit estimé par sous-porteuse comme suit

$$R_{\text{est}}^k = \log_2 \left(1 + \frac{\overline{\text{rsb}}_k}{\Gamma} \right) . \quad (5.4)$$

Le poids est alors donné par

$$w_k = \frac{R_{\text{est}}^k}{R_{\text{min}}^k} . \quad (5.5)$$

L'algorithme cherche à minimiser la différence entre le débit estimé et le débit minimal pour chaque utilisateur. Cette différence est donnée par :

$$\text{gap}_k = I_k R_{\text{est}}^k - R_{\text{min}}^k . \quad (5.6)$$

Ainsi, tant qu'il y a des sous-canaux à allouer, l'utilisateur qui présente le plus faible gap_k reçoit un sous-canal supplémentaire. À la fin de cet algorithme, si tous les gap_k sont positifs, on dit que l'algorithme s'est déroulé avec succès sinon c'est un échec c'est-à-dire certains utilisateurs n'arrivent pas à satisfaire leur débit minimum. Dans ce cas, il faudrait baisser les débits minimums ou augmenter la puissance de transmission. Cette augmentation de puissance n'est pas envisageable dans un contexte CPL. Il est à noter que le déroulement avec succès de l'algorithme ne signifie pas que tous les utilisateurs atteindront à la fin du processus d'allocation leur débit minimum. En effet, le nombre de sous-canaux étant estimé à partir du RSB moyen, il n'est pas garanti que les sous-canaux alloués à un utilisateur conduisent à la satisfaction de son débit minimum.

Dans l'algorithme WONG, du nom de l'auteur, le poids w_k est tel que $\sum_{k=1}^K w_k = 1$ et s'écrit [100]

$$w_k = \frac{R_{\text{min}}^k}{\sum_{j=1}^K R_{\text{min}}^j} . \quad (5.7)$$

Après le calcul de I_k pour chaque utilisateur (5.3), il reste I' sous-canaux qui seront attribués lors de la phase d'affectation des sous-canaux.

Phase 2 : Affectation des sous-canaux

L'affectation des sous-canaux est résolue de façon optimale par l'algorithme « hongrois » pour un ensemble $\{I_k\}_{k=1:K}$ fixé. Cet algorithme est utilisé dans de nombreux domaines car il résout les problèmes d'affectation avec minimisation d'un coût. Le problème peut s'exprimer sous une forme matricielle où chaque terme $c(k, i)$ contient le coût associé à l'affectation du sous-canal i à l'utilisateur k . La fonction de coût choisie pour la maximisation du débit est $c(k, i) = -r_{k,i}$. Le fonctionnement de l'algorithme « hongrois » est décrit dans [101]. L'algorithme « hongrois » est réputé lent pour une résolution temps réel. Plusieurs heuristiques sont proposées dans la littérature pour améliorer le temps d'exécution de l'affectation des sous-canaux [97, 98]. Nous présenterons les algorithmes proposés dans [92, 100] qui présentent des performances proches de l'algorithme « hongrois ».

5.1.1.4.1 Algorithme RPO (rate profit optimisation) Cet algorithme propose une règle de résolution de conflit plus juste entre les différents utilisateurs et qui pourrait être plus profitable pour le débit global du système [92]. Il y a conflit entre des utilisateurs lorsqu'ils désirent le même meilleur sous-canal. Dans la plupart des algorithmes d'allocation, ce problème est résolu soit par un ordre de traitement arbitraire des utilisateurs ou par l'allocation du sous-canal à l'utilisateur qui y présente le meilleur RSB. Ces politiques ne sont pas sans conséquence pour un utilisateur qui se verrait retirer son seul meilleur sous-canal au profit d'un autre, qui lui, aurait plusieurs meilleurs sous-canaux. Dans cet algorithme, à chaque itération, les utilisateurs, qui n'ont pas encore atteint leur nombre I_k de sous-canaux, choisissent leurs meilleurs sous-canaux. L'utilisateur qui n'est pas en conflit avec un autre sur son meilleur sous-canal choisi, se voit attribuer ce dernier. Les autres, qui sont en conflit, essaient de le résoudre sur un critère de profit. Le profit de l'utilisateur k représente le gain du débit global si k reçoit le sous-canal en conflit à la place de l'utilisateur qui y présente le meilleur gain. Pour calculer le profit, chaque utilisateur en conflit sur le sous-canal, choisit son second meilleur sous-canal après celui-ci. On appelle $\text{rateGap}(k)$ l'écart entre le débit implémentable sur le meilleur sous-canal (en conflit) de k et le débit implémentable sur son second meilleur sous-canal. Si k^* est l'utilisateur qui présente le meilleur RSB sur le sous-canal en conflit alors le profit est donné par

$$\text{profit}(k) = \text{rateGap}(k) - \text{rateGap}(k^*) \quad (5.8)$$

Par construction, $\text{profit}(k^*) = 0$. L'utilisateur qui présente le meilleur profit se voit allouer le sous-canal conflictuel.

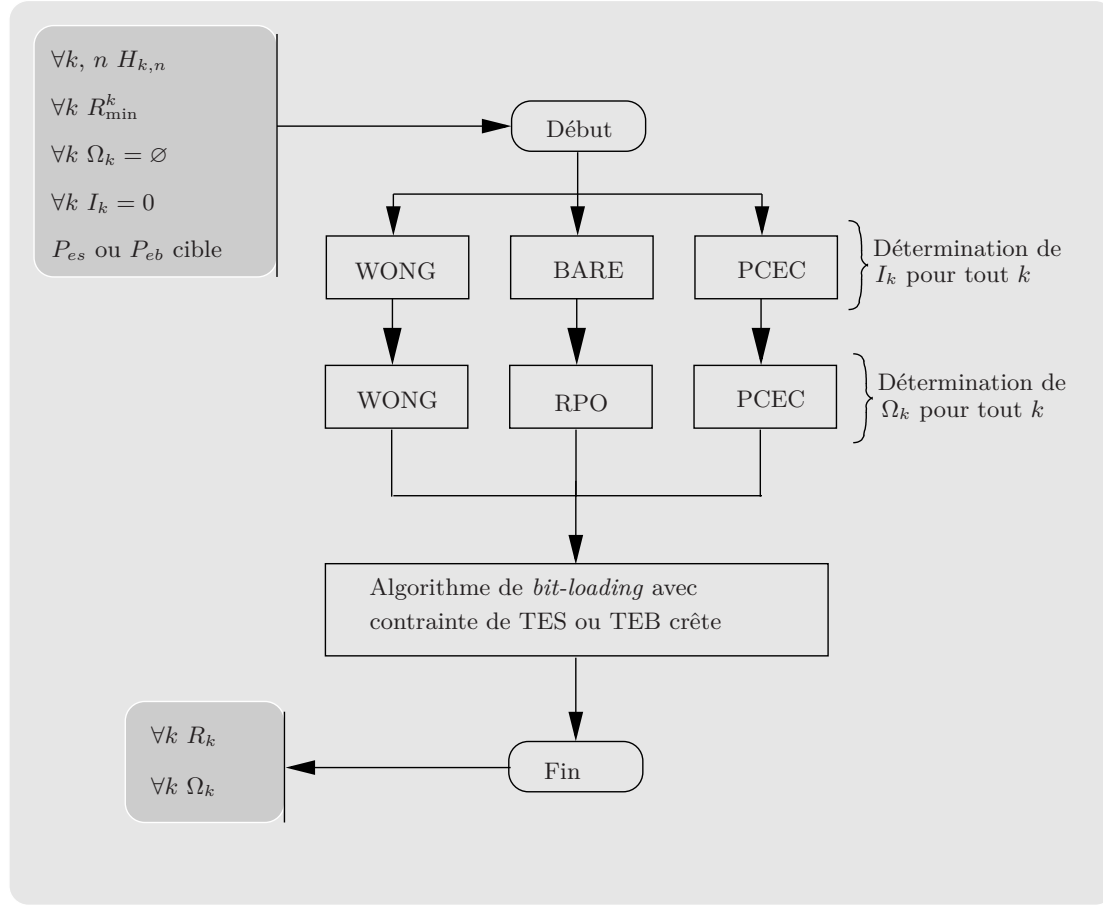


FIGURE 5.2 – Schéma global des procédures d'allocation des ressources.

5.1.1.4.2 Algorithme WONG Cet algorithme alloue, dans un ordre arbitraire, à chaque utilisateur son meilleur sous-canal parmi ceux disponibles [100]. À chaque itération, l'utilisateur qui n'a pas atteint son nombre I_k de sous-canaux et qui a son débit normalisé (par w_k) le plus faible reçoit un sous-canal supplémentaire. L'allocation des I^r sous-canaux restants se fait selon l'approche opportuniste à savoir l'utilisateur qui présente le meilleur gain sur le sous-canal traité garde ce dernier. Lorsque les utilisateurs ont les mêmes contraintes en débit, le principe d'allocation des sous-canaux rejoint celui de l'algorithme « max-min » [90].

5.1.2 Allocation des ressources en LP-OFDMA

Dans cette partie, nous proposons de nouveaux algorithmes d'allocation qui utilisent conjointement les politiques d'allocation en multi-utilisateur et une adaptation des algorithmes déjà développés au chapitre 3 dans le contexte mono-utilisateur. Ainsi, le nouveau système LP-OFDMA combinant la technique de précodage linéaire et l'OFDMA

est analysé avec les différentes contraintes de probabilités d'erreur. Le schéma global des procédures d'allocation des ressources est donné à la figure 5.2. À partir des données d'entrée du problème, nous commençons par déterminer les différents nombres de sous-canaux (qui peuvent être des sous-porteuses, des blocs de sous-porteuses ou des séquences de précodage) attribués aux utilisateurs avant de réaliser l'affectation de ces sous-canaux. Dans un souci de comparaison, les algorithmes WONG, BARE et RPO permettant de déterminer les nombres I_k et les sous-ensembles Ω_k de sous-canaux alloués aux différents utilisateurs sont considérés. L'affectation des sous-canaux permet alors de déterminer la répartition des bits et des puissances sur chaque sous-canal à partir des algorithmes de *bit-loading*, avec prise en compte d'une contrainte de TES ou TEB crête telle que décrite au chapitre 3.

Nous proposons ici un nouvel algorithme d'allocation des ressources dans un contexte multi-utilisateur pour le système LP-OFDMA avec prise en compte de l'état des canaux (PCEC). Dans l'algorithme WONG, le poids attribué à l'utilisateur k est la proportion de son débit dans le débit total requis par tous les utilisateurs. Ce poids ne prend pas en compte l'état de canal des utilisateurs. Un utilisateur k avec un mauvais canal peut avoir besoin de plus de I_k sous-canaux pour atteindre son débit minimum ou inversement un utilisateur j avec un très bon canal atteint son débit minimum avec moins de I_j sous-canaux. Par conséquent, l'algorithme WONG ne garantit pas aux utilisateurs la satisfaction de leur exigence en terme de débit minimum. L'algorithme que nous proposons prend en compte l'état des canaux de transmission lors de la phase d'attribution des sous-canaux. Ainsi, l'utilisateur qui atteint son débit minimum requis ne reçoit plus de sous-canaux et les sous-canaux restants sont redistribués aux utilisateurs insatisfaits. De plus, lors de la redistribution des sous-canaux restants, les utilisateurs qui sont proches d'être satisfaits sont privilégiés. Pour cela, on estime le nombre I_k^r de sous-canaux dont l'utilisateur insatisfait k a *a priori* besoin pour atteindre son débit minimum. Le nombre I_k^r s'écrit

$$I_k^r = \left\lfloor \left(R_{\min}^k - R_k \right) / R_{\text{est}}^k \right\rfloor. \quad (5.9)$$

L'algorithme 4 réalise la répartition des ressources entre les différents utilisateurs afin de satisfaire leur débit minimum. Contrairement à l'algorithme WONG, chaque utilisateur reçoit un sous-canal à l'étape d'initialisation suivant un ordre de priorité fonction de l'état de son canal. Plus l'utilisateur a un bon canal de transmission, moins il est prioritaire.

5.1.3 Simulation et performances

Le but de ce paragraphe est d'évaluer les performances des algorithmes d'allocation des ressources présentés dans un contexte centralisé, et particulièrement d'analyser les résultats obtenus avec le nouvel algorithme pour les systèmes LP-OFDMA. Comme pour

Algorithme 4 : Algorithme d'allocation des ressources PCEC.

```

1: Initialiser  $R_k = 0$ ,  $i_k = 0$ ,  $\mathcal{K} = \{1, 2, \dots, K\}$ 
2: pour chaque utilisateur  $k$  faire
3:   Calculer  $R_{\text{est}}^k$ 
4: fin pour
5: Ranger les utilisateurs par ordre croissant des  $R_{\text{est}}^k$ 
6: pour chaque utilisateur  $k$  faire
7:   Choisir son meilleur sous-canal  $i^*$ 
8:    $R_k = R_k + R_u^*$ ,  $i_k = i_k + 1$ 
9:   si  $i_k = I_k$  ou  $R_k \geq R_{\min}^k$  alors
10:     $\mathcal{K} = \mathcal{K} - \{k\}$ 
11:     $I_k = i_k$ 
12:   fin si
13: fin pour
14: tant que  $\mathcal{K} \neq \emptyset$  faire
15:   Trouver  $k^* = \arg \min_{k \in \mathcal{K}} R_k / w_k$ 
16:   Choisir le meilleur sous-canal de  $k^*$  parmi ceux disponibles
17:    $R_{k^*} = R_{k^*} + r_{k^*, i^*}$ ,  $i_{k^*} = i_{k^*} + 1$ 
18:   si  $i_{k^*} = I_{k^*}$  ou  $R_{k^*} \geq R_{\min}^{k^*}$  alors
19:     $\mathcal{K} = \mathcal{K} - \{k^*\}$ 
20:     $I_{k^*} = i_{k^*}$ 
21:   fin si
22: fin tant que
23: //redistribution des sous-canaux restants
24: pour chaque utilisateur  $k$  appartenant à  $\mathcal{K}$  faire
25:   Calculer  $I_k^r$ 
26: fin pour
27: Ranger les utilisateurs par ordre croissant des  $I_k^r$ 
28: tant que il y a des sous-canaux faire
29:   Chaque utilisateur reçoit  $I_k^r$  sous-canaux
30: fin tant que

```

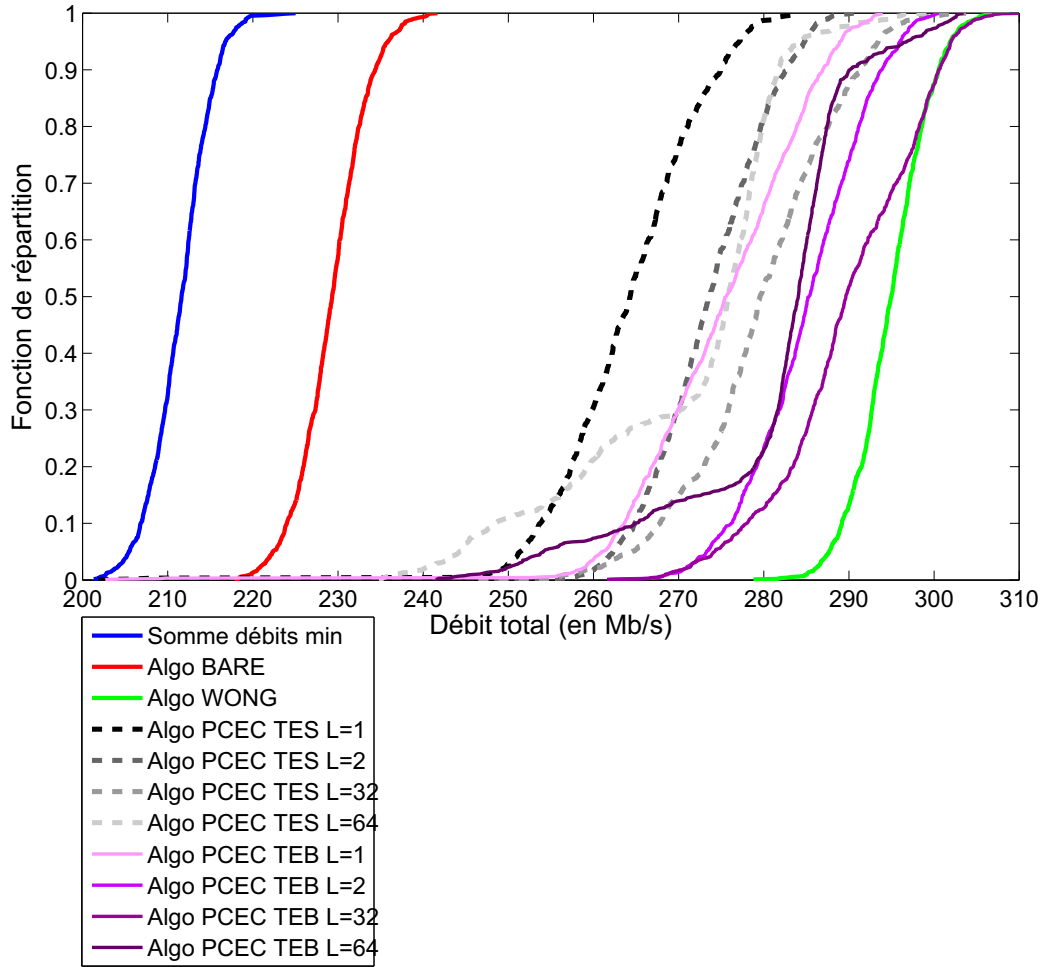


FIGURE 5.3 – Fonctions de répartition des débits totaux réalisables avec les différents algorithmes pour 9 utilisateurs.

les chapitres précédents, l'ensemble des simulations est mené à partir des modèles de canaux CPL présentés dans le chapitre 1. Ces canaux ont été répartis dans neuf classes différentes et un modèle de fonction de transfert est associé à chacune des classes. Les paramètres de simulations (densité spectrale de puissance du signal et du bruit, nombre de sous-porteuses, bande de fréquence et probabilité d'erreur) sont les mêmes que ceux du chapitre précédent, et qui sont donnés dans le tableau 4.5. Le débit minimum requis par chaque utilisateur est par hypothèse donné par le rapport entre son débit lorsqu'il est considéré tout seul et le nombre d'utilisateurs dans le système. Ainsi, la courbe qui donne la somme des débits minimaux représente aussi le débit total obtenu pour un système TDMA. Rappelons qu'un utilisateur est dit insatisfait lorsqu'il n'atteint pas son débit minimum après le processus d'allocation des ressources.

Dans un premier temps, nous étudions les performances des différents algorithmes

TABLE 5.2 – Pourcentage de cas testés où respectivement 0, 1, 2, 3 et 4 utilisateurs sur 9 sont insatisfaits.

Algorithmes	Nombres d'utilisateurs insatisfaits				
	0	1	2	3	4
BARE	51,2 %	44,4 %	4,2 %	0,2 %	0 %
WONG	9,7 %	0 %	0 %	3,1 %	87,2 %
PCEC TES $L = 1$	80,6 %	19 %	0 %	0,1 %	0,3 %
PCEC TES $L = 2$	98,9 %	1,1 %	0 %	0 %	0 %
PCEC TES $L = 8$	96,8 %	3,2 %	0 %	0 %	0 %
PCEC TES $L = 16$	66,6 %	33,2 %	0,2 %	0 %	0 %
PCEC TES $L = 32$	19,3 %	69,6 %	10,8 %	0,3 %	0 %
PCEC TES $L = 64$	9,7 %	11 %	76,6 %	2,7 %	0 %
PCEC TEB $L = 1$	98,3 %	1,3 %	0 %	0,3 %	0,1 %
PCEC TEB $L = 2$	99,9 %	0,1 %	0 %	0 %	0 %
PCEC TEB $L = 8$	100 %	0 %	0 %	0 %	0 %
PCEC TEB $L = 16$	88,8 %	11,2 %	0 %	0 %	0 %
PCEC TEB $L = 32$	27,8 %	67,4 %	4,8 %	0 %	0 %
PCEC TEB $L = 64$	9,7 %	6,8 %	81,2 %	2,3 %	0 %

pour un système avec 9 utilisateurs où chaque utilisateur possède un canal de transmission pris dans une classe différente. La figure 5.3 donne les fonctions de répartition des débits totaux maximaux avec les différents algorithmes dans un contexte multi-utilisateur. Tout d'abord, on note que le fait de permettre aux utilisateurs de transmettre simultanément dans une même trame (système OFDMA) offre un meilleur débit total que si les utilisateurs transmettaient dans des intervalles de temps distincts (système TDMA, résultat donné par la courbe représentant la somme des débits minimums). On note que tous les algorithmes présentés offrent des débits totaux supérieurs à la somme des débits requis par les utilisateurs. Néanmoins, en considérant chaque utilisateur séparément, nous verrons que la répartition des débits ne garantit pas à tous les utilisateurs leur débit minimum. L'algorithme WONG présente globalement le meilleur débit total et l'algorithme BARE présente le plus faible débit total. Le nouvel algorithme proposé pour le système LP-OFDMA montre un débit total croissant en fonction de la longueur des séquences de précodage avec $L < 64$. Par contre, une grande longueur ($L = 64$) des séquences dégrade le débit total et aussi augmente le nombre d'utilisateurs insatisfaits comme nous le verrons dans le tableau 5.2. Ce résultat peut s'expliquer par le fait que le nombre de sous-canaux à répartir entre les utilisateurs diminue en fonction de

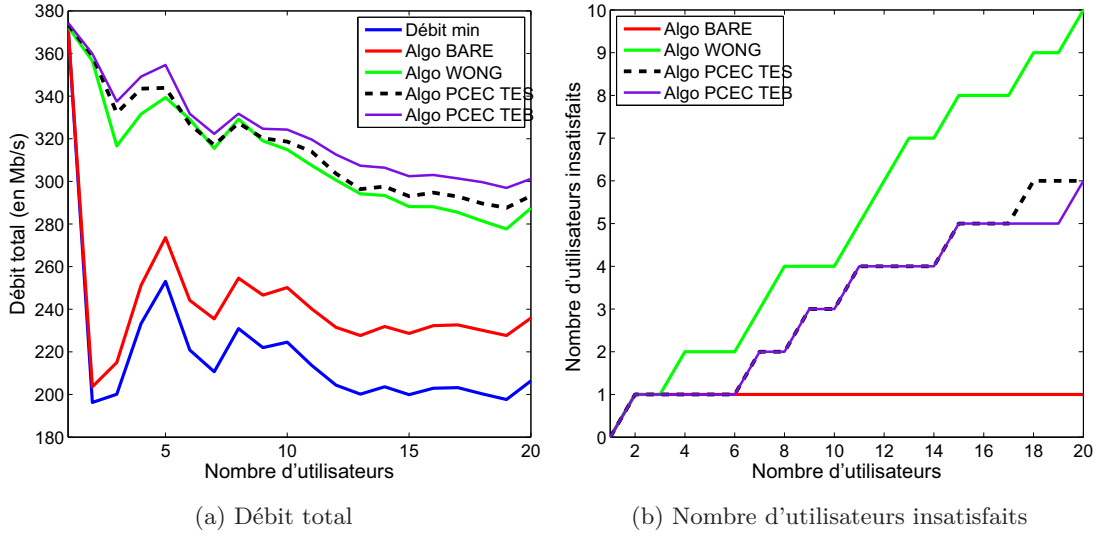


FIGURE 5.4 – Évolution du débit total (a) et du nombre d'utilisateurs insatisfaits en fonction du nombre d'utilisateurs du système.

la longueur des séquences de précodage. Par conséquent, la flexibilité à redistribuer les sous-canaux se retrouve réduite.

Le tableau 5.2 présente les pourcentages de cas testés dans lesquels respectivement 0, 1, 2, 3 et 4 utilisateurs sont insatisfaits. Ce tableau indique par exemple que dans 51,2% des systèmes à 9 utilisateurs testés, l'algorithme BARE aboutit à 0 utilisateur insatisfait c'est-à-dire tous les utilisateurs sont satisfaits. Aussi, l'algorithme WONG conduit à 4 utilisateurs insatisfaits dans 87,2% des cas testés. Pour 0 utilisateur insatisfait, le nouvel algorithme, quelle que soit la contrainte de TEB ou de TES choisie, offre le meilleur pourcentage comparé aux autres algorithmes. Ces résultats confirment également les résultats obtenus pour les algorithmes d'allocation sous contrainte de TEB et TES crête dans un contexte mono-utilisateur (cf. chapitre 3). La figure 5.3 montre également que le nouvel algorithme avec prise en compte d'une contrainte de TEB crête (courbes violettes en trait plein) est plus performant que l'algorithme avec prise en compte d'une contrainte de TES crête (courbes noires/grises en trait pointillé). En revanche, l'algorithme WONG, avec un meilleur débit total, présente le plus faible pourcentage d'utilisateurs satisfaits et le plus fort pourcentage dans le cas de 4 utilisateurs insatisfaits.

Dans un second temps, nous étudions les performances des algorithmes en faisant varier le nombre d'utilisateurs. Chaque utilisateur possède de manière aléatoire un canal parmi les 9 classes de canaux. Les résultats obtenus sont des moyennes sur plusieurs simulations effectuées avec le même jeu de canaux choisis. La figure 5.4 présente l'évolution des débits totaux et des nombres d'utilisateurs insatisfaits avec les différentes solutions

présentées en fonction du nombre d'utilisateurs. On note une tendance globale à la baisse, liée au jeu de canaux choisis, pour les débits totaux réalisables, mais une tendance à la hausse pour les nombres d'utilisateurs insatisfaits. Ces résultats montrent que le nouvel algorithme proposé, quelle que soit la contrainte de probabilité d'erreur choisie, est plus performant que l'algorithme WONG en terme de débit total et de nombre d'utilisateurs insatisfaits. L'algorithme BARE, quant à lui, permet d'assurer les débits minimums avec au maximum un seul utilisateur insatisfait.

5.1.4 Conclusion

Dans cette première partie, nous avons étudié le problème d'allocation centralisée des ressources pour les systèmes adaptatifs OFDMA. Nous avons passé en revue les aspects fondamentaux de ce type de problème d'allocation ainsi que quelques algorithmes d'allocation de ressources permettant de maximiser le débit total du système sous la contrainte de débit minimum pour chaque utilisateur. L'analyse de ces algorithmes nous a permis de développer un nouvel algorithme d'allocation pour le système LP-OFDMA. Cet algorithme permet à la fois d'augmenter le débit total du système et de réduire le nombre d'utilisateurs insatisfaits. De plus, les résultats de simulation ont permis de mettre en exergue l'augmentation en débit d'un système centralisé OFDMA par rapport au système TDMA actuellement mis en œuvre dans les modems CPL.

5.2 Allocation décentralisée des ressources

Nous avons vu que dans le contexte de communications centralisées, le nœud central disposait de l'information sur l'état des canaux de transmission de tous les utilisateurs. Le passage à un contexte de communications distribuées permet de limiter les échanges entre les différentes entités (utilisateurs et nœud central) et aussi d'adapter rapidement le réseau aux variations des paramètres, particulièrement pour un système avec un grand nombre d'utilisateurs et de sous-porteuses. Les réseaux actuels CPL *indoor* sont organisés de manière décentralisée en ce sens que les communications entre les utilisateurs ne passent pas par un nœud central. Néanmoins, un coordinateur central, CCo (*central coordinator*), contrôle les activités du réseau et envoie périodiquement des trames balises contenant des informations de planification allouant ainsi du temps à l'accès multiple TDMA ou CSMA [7]. Le CPL utilise actuellement une technique d'accès multiple décentralisée de type CSMA-CA. Le CSMA est un protocole d'accès multiple aléatoire dans lequel un utilisateur qui veut transmettre des données commence par vérifier l'absence de trafic des autres utilisateurs sur le support de transmission partagé avant de démarrer la transmission de ses données. L'utilisation du réseau avec ce protocole est inférieure à 50% [8]. Le protocole CSMA/CA utilise un mécanisme d'esquive de collision basé sur un

principe d'accusé de réception réciproque entre l'émetteur et le récepteur. L'utilisateur voulant transmettre écoute le support. Si ce dernier est encombré, la transmission est différée. Dans le cas contraire, si le support est libre pendant un temps donné, alors l'utilisateur peut émettre. L'utilisateur transmet alors un message contenant des informations sur le volume des données qu'il souhaite émettre et son débit de transmission. Le récepteur répond par un message signifiant que le champ est libre pour émettre, puis l'émetteur commence l'émission des données. À réception de toutes les données émises par l'utilisateur, le récepteur envoie un accusé de réception [102].

Dans cette partie, nous étudions un cas de transmission simultanée de deux utilisateurs dans un réseau CPL, comme indiqué sur la figure 5.5. Les signaux de l'émetteur E1 sont considérés comme des interférences pour la communication entre E2 et R2. Les canaux mis en jeu sont $\mathbf{H}_{11}$ et $\mathbf{H}_{22}$ pour les principaux trajets, et $\mathbf{H}_{12}$ et $\mathbf{H}_{21}$ pour les interférences. Plusieurs études dans la littérature cherchent alors à maximiser les débits (le débit total ou les débits individuels) sous la contrainte de puissance totale [103–106] ou à minimiser la puissance totale sous contrainte de débit minimum [107]. Ces problèmes ont été modélisés par un jeu non coopératif ne nécessitant pas de coopération entre les utilisateurs. Les travaux existants ont permis d'obtenir des conditions suffisantes qui garantissent l'existence et l'unicité de l'équilibre de Nash. Rappelons que ces études font les hypothèses suivantes : l'état du canal de chaque source à son propre destinataire est supposé connu de la source correspondante mais pas des autres sources ; chaque destinataire est supposé mesurer parfaitement la densité spectrale de puissance du bruit ainsi que les perturbations dues aux autres liens. Sur la base de ces informations, chaque destinataire calcule la signalisation nécessaire pour son propre lien et le renvoie à son émetteur par un canal de retour. De plus, toutes les transmissions sont synchronisées avec un retard de propagation inférieur ou égal à la taille de l'intervalle de garde. Cela impose une taille maximale de l'intervalle de garde qui dépend de l'étalement maximal des échos du système.

Dans un contexte CPL avec une limitation de densité spectrale de puissance, nous étudions le problème de minimisation de la puissance sous les contraintes de débit minimum et de DSP. Sous les hypothèses faites plus haut, la maximisation des débits individuels sous la contrainte de DSP conduirait, comme dans le cas mono-utilisateur, à l'allocation de la puissance maximale autorisée sur chaque sous-porteuse.

5.2.1 Quelques éléments de la théorie des jeux

La théorie des jeux est la discipline mathématique qui étudie les situations où le sort de chaque participant dépend non seulement des décisions qu'il prend mais également des décisions prises par d'autres participants. En conséquence, le choix « optimal » pour un participant dépend généralement de ce que font les autres. Parce que chacun n'est

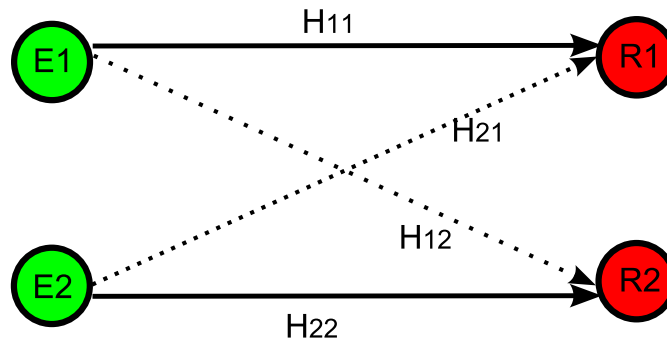


FIGURE 5.5 – Système distribué avec 2 utilisateurs.

pas totalement maître de son sort, on dit que les participants se trouvent en situation d'interaction stratégique [108]. Tout d'abord, les joueurs se connaissent (ils savent combien il y a de participants et qui ils sont). Ensuite, ils ne peuvent pas se contenter de choisir leurs propres plans d'actions, en négligeant ce que font les autres. Ils doivent au contraire se faire une idée aussi précise que possible des plans choisis par les autres. Pour cela, la théorie admet :

1. que chaque joueur s'efforce de prendre les meilleures décisions pour lui-même et sait que les autres joueurs font de même
2. que chacun sait qu'il en va de même pour tous les autres et ainsi de suite.

Il existe deux grandes familles de jeux qui sont les jeux coopératifs et les jeux non coopératifs. Un jeu est coopératif lorsque des joueurs peuvent passer entre eux des accords qui les lient de manière contraignante (par exemple, sous la forme d'un contrat qui prévoit une sanction légale dans le cas du non respect de l'accord). On dit alors qu'ils forment une coalition dont les membres agissent de concert. Dans le cas contraire, c'est-à-dire lorsque les joueurs n'ont pas la possibilité de former des coalitions, le jeu est non coopératif. Par définition, dans un jeu non coopératif on spécifie toutes les options stratégiques offertes aux joueurs, alors que les contrats qui sous-tendent les coalitions dans un jeu coopératif ne sont pas décrits. Un tel jeu peut être défini comme stratégique. Un jeu en forme stratégique est une collection de stratégies décrivant les actions de chaque joueur dans toutes les situations concevables du jeu, ainsi que les gains (*payoffs*) que chacun obtient lorsque les stratégies de tous les joueurs sont connues.

5.2.1.1 Définition d'un jeu en forme stratégique

Les éléments constitutifs d'un jeu G en forme stratégique sont les suivants :

1. On suppose que les joueurs sont en nombre fini. Un joueur quelconque est désigné par l'indice k . L'extension au cas d'une infinité de joueurs ne pose pas de problèmes conceptuels particuliers.

2. s_k désigne une stratégie du joueur $k \in \{1, \dots, K\}$. Une stratégie décrit de manière précise tout ce qu'un joueur fait. Remarquons que s_k n'est pas nécessairement un nombre ; ce peut être aussi un vecteur ou une fonction.
3. $\mathbf{S}_k$ est l'ensemble des stratégies du joueur $k \in \{1, \dots, K\}$. Cet ensemble décrit toutes les stratégies disponibles pour le joueur k .
4. $\mathbf{s} = (s_1, \dots, s_k, \dots, s_K) \in \mathbf{S}_1 \times \dots \times \mathbf{S}_k \times \dots \times \mathbf{S}_K \equiv \mathbf{S}$ est une issue du jeu, c'est-à-dire une combinaison de stratégies à raison d'une stratégie par joueur. On désigne par $\mathbf{s}_{-k} \in \mathbf{S}_{-k}$ toutes les stratégies choisies sauf celle du joueur k .
5. $u_k(\mathbf{s}) \in \mathbb{R}$ est la fonction de gain du joueur k . Autrement dit, la « fonction d'objectif » du joueur k dépend non seulement de sa stratégie s_k , mais aussi de celles des autres joueurs résumées dans $\mathbf{s}_{-k}$. Le joueur k préfère strictement l'issue $\mathbf{s}$ à l'issue $\mathbf{s}'$ si $u_k(\mathbf{s}) > u_k(\mathbf{s}')$. Si $u_k(\mathbf{s}) = u_k(\mathbf{s}')$, le joueur est indifférent entre les deux issues.
6. Chaque joueur connaît les ensembles de stratégies et les fonctions de gains de tous les joueurs, y compris donc les siennes.

Du fait de cette dernière hypothèse, on dit que le jeu est en information complète. Dans le cas contraire, le jeu est dit en information incomplète.

5.2.1.2 Équilibre de Nash

On dit qu'une combinaison de stratégies $\mathbf{s}^*$ est un équilibre de Nash (ou un équilibre non coopératif) si l'inégalité suivante est satisfaite pour chaque joueur $k = 1, 2, \dots, K$:

$$u_k(s_k^*, \mathbf{s}_{-k}^*) \geq u_k(s_k, \mathbf{s}_{-k}^*) \quad (5.10)$$

pour tout $s_k \in \mathbf{S}_k$, c'est-à-dire que chaque joueur maximise ses gains compte tenu de l'action supposée de l'autre. A cet équilibre, aucun joueur n'a intérêt à changer de stratégie si tous les autres gardent la même stratégie car sa fonction d'utilité sera dégradée [104]. L'existence d'un équilibre n'implique pas que celui-ci soit nécessairement optimal. Il peut en effet exister d'autres choix des joueurs qui conduisent, pour chacun, à un gain supérieur.

Exemple du dilemme du prisonnier On suppose que deux suspects sont interrogés séparément par la police pour une action délictueuse grave. La police ne dispose pas d'éléments de preuve suffisants pour obtenir la condamnation des prévenus pour l'acte dont ils sont accusés. L'aveu d'au moins l'un des deux est donc indispensable. La police propose à chaque accusé d'avouer, dans quel cas il sera relâché. S'il n'avoue pas mais que l'autre le fait, il écope d'une peine de prison de 15 ans. Si les deux avouent, ils peuvent

TABLE 5.3 – La matrice des gains du dilemme du prisonnier

	Avoue	N'avoue pas
Avoue	$(-8; -8)$	$(0; -15)$
N'avoue pas	$(-15; 0)$	$(-1; -1)$

espérer bénéficier de circonstances atténuantes et recevoir une peine de 8 ans chacun. Enfin, si aucun des deux n'avoue, ils seront condamnés pour des délits mineurs à 1 an de prison. Le tableau 5.3 donne la matrice des gains en fonction des décisions des deux prisonniers. L'équilibre de Nash est atteint lorsque les deux prisonniers décident d'avouer $(-8; -8)$ car dans cette situation, aucun d'eux n'a intérêt à changer de stratégie. Par contre, le gain optimal $(-1; -1)$ est obtenu lorsqu'aucun des deux prisonniers n'avoue l'action délictueuse.

5.2.1.3 Existence et unicité de l'équilibre de Nash

Certains jeux ne possèdent pas nécessairement un équilibre de Nash alors que d'autres peuvent en posséder plusieurs. Dans ce paragraphe, nous donnerons quelques conditions suffisantes pour l'existence d'un équilibre de Nash. On définit la meilleure réponse (*best reply*) du joueur k par la stratégie qui lui procure le meilleur gain quelle que soit la stratégie des autres joueurs. Cette meilleure réponse est définie par la correspondance de $\mathbf{S}_{-k}$ vers $\mathbf{S}_k$:

$$\forall \mathbf{s}_{-k} \in \mathbf{S}_{-k}, \quad r_k(\mathbf{s}_{-k}) = \arg \max u_k(s_k, \mathbf{s}_{-k}). \quad (5.11)$$

En considérant l'application $r(\mathbf{s}) = (r_1(\mathbf{s}_{-1}), \dots, r_K(\mathbf{s}_{-K}))$ de $\mathbf{S}$ vers $\mathbf{S}$, on a le lemme suivant [108] :

Lemme 1. *Si l'application $r(\mathbf{s})$ possède un point fixe, ce point est un équilibre de Nash du jeu, et réciproquement.*

Un ensemble de conditions suffisantes pour qu'un point fixe existe est donné par le théorème de Brouwer qui s'énonce de la manière suivante [109] : soit f une application de S vers S où S est un sous-ensemble de $\mathbb{R}^K$. Si S est compact et convexe et si f est continu, alors f possède un point fixe. Il s'en suit alors que si les ensembles de stratégie sont des sous-ensembles compacts et convexes de $\mathbb{R}^K$, et si la fonction de gain u_k est continue en $\mathbf{s}$ et strictement quasi concave en s_k pour chaque joueur, alors le jeu non coopératif admet un équilibre de Nash.

Théorème 1. *Si $r(\mathbf{s})$ est une contraction, alors l'équilibre de Nash est unique.*

Sachant qu'une fonction f de $S \in \mathbb{R}^K$ vers S est une contraction s'il existe $\lambda \in]0, 1[$ tel que, pour tout x' et x'' , on ait $d[f(x'), f(x'')] \leq \lambda d(x', x'')$, où d est la fonction distance.

5.2.2 Minimisation de la puissance sous contraintes de débit et de DSP

Dans le contexte CPL avec contrainte de densité spectrale de puissance, une première approche pour concevoir des algorithmes distribués est de maximiser les débits individuels sous la contrainte de DSP. En considérant les signaux des autres utilisateurs comme des interférences, la solution au problème de maximisation du débit, pour chaque utilisateur, consiste à allouer la puissance maximale autorisée sur chaque sous-porteuse, ce qui pénalise fortement les autres utilisateurs. Comme la plupart des communications en CPL nécessite un certain niveau de qualité de service en terme de débit, on peut bien imaginer un problème d'optimisation distribuée où l'on minimise la puissance totale sous les contraintes de débit minimum et de DSP. On considère dans cette étude un système distribué à 2 utilisateurs. Le débit de l'utilisateur 1 s'écrit

$$R_1 = \sum_{n=1}^N \log_2 \left(1 + \frac{P_{1,n} |H_{1,1,n}|^2}{\Gamma (\sigma_{1,n} + P_{2,n} |H_{2,1,n}|^2)} \right) \quad (5.12)$$

où $P_{1,n}$ et $\sigma_{1,n}$ sont respectivement la puissance utile et la puissance de bruit pour la sous-porteuse n de l'utilisateur 1. La même écriture peut être faite pour l'utilisateur 2. En posant

$$\beta_{1,n} = \frac{\Gamma \sigma_{1,n}}{|H_{1,1,n}|^2} \text{ et } \alpha_{2,n} = \frac{\Gamma |H_{2,1,n}|^2}{|H_{1,1,n}|^2}, \quad (5.13)$$

il vient que pour la transmission 1, le débit de l'utilisateur 1 sachant l'allocation de l'utilisateur 2 s'écrit

$$R_1(\mathbf{P}_1|\mathbf{P}_2) = \sum_{n=1}^N \log_2 \left(1 + \frac{P_{1,n}}{\beta_{1,n} + \alpha_{2,n} P_{2,n}} \right) \quad (5.14)$$

où $\mathbf{P}_1 = (P_{1,n})_{n=1}^N$ est la stratégie d'allocation de puissance de l'utilisateur 1, $\mathbf{P}_2$ est la stratégie de l'utilisateur 2 et $P_{k,n} \leq E_{\text{DSP}}$, $k = 1, 2$.

On considère un jeu stratégique non coopératif dans lequel chaque joueur (utilisateur) est en compétition avec les autres pour choisir l'allocation de puissance qui minimise sa puissance d'émission sous la contrainte d'un débit minimum réalisable sur son lien de transmission. L'équilibre de Nash d'un tel jeu est atteint lorsque chaque utilisateur, étant donné la stratégie des autres, ne peut atteindre une puissance totale plus faible en modifiant sa stratégie tout en garantissant la contrainte de débit [110]. À l'équilibre de Nash, la stratégie de chaque utilisateur est alors la réponse optimale à la stratégie des autres utilisateurs. Notons $\mathcal{P}_1(\mathbf{P}_2) \subset \mathbb{R}_+^N$ l'ensemble des stratégies $\mathbf{P}_1$ admissibles de l'utilisateur 1 sachant la stratégie de l'utilisateur 2 et $\mathcal{P} = (\mathcal{P}_1(\mathbf{P}_2), \mathcal{P}_2(\mathbf{P}_1))$ l'ensemble

des stratégies des deux utilisateurs. On a

$$\mathcal{P}_1(\mathbf{P}_2) = \{\mathbf{P}_1 \in \mathbb{R}_+^N | R_1(\mathbf{P}_1 | \mathbf{P}_2) \geq R_{\min}^1\}. \quad (5.15)$$

Connaissant l'allocation de puissance $P_{2,n}$ de l'utilisateur 2, l'allocation optimale $P_{1,n}$ pour l'utilisateur 1 est la solution au problème

$$\min \sum_{n=1}^N P_{1,n} \quad (5.16)$$

sous les contraintes

$$R_1(\mathbf{P}_1 | \mathbf{P}_2) \geq R_{\min}^1 \text{ et } 0 \leq P_{1,n} \leq E_{\text{DSP}} \quad (5.17)$$

Le lagrangien s'écrit

$$\begin{aligned} \mathcal{L}(P_{1,n}, \lambda, \gamma, \delta) = & \sum_{n=1}^N P_{1,n} + \lambda \left[R_{\min}^1 - \sum_{n=1}^N \log_2 \left(1 + \frac{P_{1,n}}{\beta_{1,n} + \alpha_{2,n} P_{2,n}} \right) \right] \\ & + \sum_{n=1}^N \gamma_n (P_{1,n} - E_{\text{DSP}}) + \sum_{n=1}^N \delta_n (-P_{1,n}) \end{aligned} \quad (5.18)$$

Les conditions de Karush-Kuhn-Tucker [111] s'écrivent

$$\begin{aligned} \forall n, \quad \frac{\partial \mathcal{L}}{\partial P_{1,n}} = 1 - \frac{\lambda}{\ln(2)} \frac{1}{P_{1,n} + \beta_{1,n} + \alpha_{2,n} P_{2,n}} + \gamma_n - \delta_n = 0 \\ \Rightarrow P_{1,n} = \frac{\lambda}{\ln(2) (1 + \gamma_n - \delta_n)} - (\beta_{1,n} + \alpha_{2,n} P_{2,n}) \\ \text{et aussi } \begin{cases} R_{\min}^1 = \sum_{n=1}^N \log_2 \left(1 + \frac{P_{1,n}}{\beta_{1,n} + \alpha_{2,n} P_{2,n}} \right) \\ \forall n, \gamma_n (P_{1,n} - E_{\text{DSP}}) = 0 \\ \forall n, \delta_n (-P_{1,n}) = 0 \end{cases} \end{aligned} \quad (5.19)$$

La solution analytique au problème est

$$P_{1,n}^* = \begin{cases} 0 & \text{si } \mu_1 < \beta_{1,n} + \alpha_{2,n} P_{2,n}, \\ E_{\text{DSP}} & \text{si } \mu_1 > E_{\text{DSP}} + \beta_{1,n} + \alpha_{2,n} P_{2,n}, \\ \mu_1 - \beta_{1,n} - \alpha_{2,n} P_{2,n} & \text{si } \beta_{1,n} + \alpha_{2,n} P_{2,n} \leq \mu_1 \leq E_{\text{DSP}} + \beta_{1,n} + \alpha_{2,n} P_{2,n}. \end{cases} \quad (5.20)$$

pour tout $n \in [1, N]$ et en choisissant μ_1 de sorte à vérifier la contrainte de débit $R_1(\mathbf{P}_1^* | \mathbf{P}_2, \mu_1) = R_{\min}^1$. Il est à noter que $\mu_1 > 0$, sinon il advient que $P_{1,n}^* = 0$ pour tout n , ce qui contredit la contrainte de débit. La variable μ_1 est la solution de l'équation

$R_1(\mathbf{P}_1^*|\mathbf{P}_2, \mu_1) - R_{\min}^1 = 0$. Sous l'hypothèse de l'existence et de la positivité de μ_1 , la stratégie d'allocation de puissance $\mathbf{P}_1^*$ est la meilleure allocation de l'utilisateur 1 sachant l'allocation de 2, c'est-à-dire

$$\forall \mathbf{P}_1 \in \mathcal{P}_1(\mathbf{P}_2^*), \quad \sum_{n=1}^N P_{1,n}^* \leq \sum_{n=1}^N P_{1,n}. \quad (5.21)$$

Il en est de même pour l'utilisateur 2. D'après la définition de l'équilibre de Nash (cf. 5.2.1.2), la combinaison $\{\mathbf{P}_1^*, \mathbf{P}_2^*\}$ est donc un équilibre de Nash.

Détermination de μ_1

L'expression de la solution est proche de celle de l'algorithme du *water-filling* [112] à la seule différence de la contrainte supplémentaire de DSP. Pour déterminer μ_1 , nous utiliserons la méthode de la sécante qui est une méthode algorithmique de recherche de zéros d'une fonction. La méthode de la sécante nécessite deux points de départ x_1 et x_2 tels que $f(x_1) \cdot f(x_2) < 0$, et une certaine tolérance ϵ . L'idée est de remplacer localement la fonction f par la droite qui passe par les deux points $(x_1, f(x_1)), (x_2, f(x_2))$. Cette méthode est donnée par l'algorithme suivant, pour une fonction f donnée

1. $i = 0, y_0 = f(x_1)$;
- 2.

$$x_0 = \frac{x_1 f(x_2) - x_2 f(x_1)}{f(x_2) - f(x_1)}, \quad y_1 = f(x_0);$$

3. Tant que $|y_{i+1} - y_i| > \epsilon$ faire

si $y_{i+1} < 0$ alors $x_1 = x_0$ sinon $x_2 = x_0$

$$x_0 = \frac{x_1 f(x_2) - x_2 f(x_1)}{f(x_2) - f(x_1)}, \quad y_{i+1} = f(x_0)$$

$i = i + 1$;

Soient les ensembles de sous-porteuses $\Omega_1 = \{n | 0 < P_{1,n}^* < E_{\text{DSP}}\}$, $\Omega_2 = \{n | P_{1,n}^* = E_{\text{DSP}}\}$ et $\Omega_3 = \{n | P_{1,n}^* = 0\}$ avec $\text{card}(\Omega_1) = N_1$, $\text{card}(\Omega_2) = N_2$ et $\text{card}(\Omega_3) = N_3$. On peut donc écrire

$$\begin{aligned} R_1(\mathbf{P}_1^*|\mathbf{P}_2, \mu_1) = & N_1 \log_2(\mu_1) - \sum_{n \in \Omega_1} \log_2(\beta_{1,n} + \alpha_{2,n} P_{2,n}) \\ & + \sum_{n \in \Omega_2} \log_2 \left(1 + \frac{E_{\text{DSP}}}{\beta_{1,n} + \alpha_{2,n} P_{2,n}} \right). \end{aligned} \quad (5.22)$$

La fonction f dont on cherche le zéro est alors donnée par

$$\begin{aligned} f(\mu_1) = & R_{\min}^1 - N_1 \log_2(\mu_1) + \sum_{n \in \Omega_1} \log_2(\beta_{1,n} + \alpha_{2,n} P_{2,n}) \\ & - \sum_{n \in \Omega_2} \log_2 \left(1 + \frac{E_{DSP}}{\beta_{1,n} + \alpha_{2,n} P_{2,n}} \right). \end{aligned} \quad (5.23)$$

5.2.2.1 Existence et unicité de l'équilibre de Nash

On note $\text{MR}_1(\mathbf{P}_2)$ la meilleure réponse de l'utilisateur 1 quelle que soit l'allocation de puissance de l'utilisateur 2 telle que $[\text{MR}_1(\mathbf{P}_2)]_n = P_{1,n}^*$, l'allocation optimale de puissance. Considérons l'application $\text{MR} = (\text{MR}_1(\mathbf{P}_2), \text{MR}_2(\mathbf{P}_1))$ de $\mathcal{P}$ vers $\mathcal{P}$. On montre aisément que l'ensemble $\mathcal{P}$ est fermé et borné et donc est compact. De plus

$$\forall P_{k,n}, P'_{k,1} \in \mathcal{P}_k, aP_{k,n} + (1-a)P'_{k,n} \in \mathcal{P}_k \text{ avec } a \in [0, 1] \quad (5.24)$$

d'où $\mathcal{P}$ est convexe. D'après le lemme 1, l'existence d'un équilibre de Nash dépend de l'existence d'un point fixe pour l'application MR . Pour cela, nous allons chercher à montrer que l'application MR est continue. En utilisant le théorème de Brouwer qui s'énonce de la manière suivante [109] : soit f une application de $\mathcal{P}$ vers $\mathcal{P}$ où $\mathcal{P}$ est un sous-ensemble de $\mathcal{R}^n$. Si $\mathcal{P}$ est compact et convexe et si f est continu, alors f admet un point fixe.

Théorème 2. *L'application MR_k , ($k = 1, 2$) est continue, de plus si $\max_n \alpha_{jn} < 1$, ($j = 1, 2$; $j \neq k$) alors MR_1 est une contraction.*

Preuve. voir Annexe A. □

Sous l'hypothèse de l'existence et de la positivité de μ_1 , on déduit alors que le problème de minimisation de la puissance sous les contraintes de débit et de DSP admet un unique équilibre de Nash.

5.2.2.2 Mise en œuvre de l'allocation décentralisée des ressources

Nous avons montré, sous certaines conditions, que le jeu stratégique en vue de la minimisation de la puissance d'émission admet un unique équilibre de Nash. À cet équilibre, chaque utilisateur atteint son débit minimum avec le minimum de puissance de transmission. Dans ce paragraphe, nous nous intéressons à la mise en place d'un tel jeu non coopératif avec l'absence de signalisation entre les différents utilisateurs. Dans la littérature, les algorithmes distribués d'allocation des ressources sont proposés suivant que le jeu est synchrone ou asynchrone. Dans un jeu synchrone, les joueurs décident de leur coup simultanément, sans savoir ce que les autres jouent. Dans un jeu asynchrone,

ils jouent les uns après les autres, en disposant à chaque fois de l'information sur le coup de l'adversaire [113]. La convergence de ces algorithmes a été démontré dans [103, 107] sous la condition de l'existence et de l'unicité de l'équilibre de Nash.

- Dans l'algorithme asynchrone, chaque utilisateur, à tour de rôle, met à jour séquentiellement son allocation de puissance en résolvant le problème (5.16). Un tel algorithme peut être mis en œuvre d'une manière distribuée puisque chaque utilisateur, pour déterminer son allocation optimale de puissance, a besoin de connaître localement pour chaque sous-canal les termes de bruit et d'interférences. Par contre, cet algorithme présente une convergence lente lorsque le nombre d'utilisateurs dans le système est grand [107]. En effet, chaque utilisateur, avant de choisir sa propre stratégie d'allocation, est contraint d'attendre son tour selon l'ordonnancement fixé. Force est de constater qu'il faudrait une instance supérieure centrale ou un minimum de coopération entre les utilisateurs pour gérer l'ordonnancement. De ce fait, l'algorithme séquentiel n'a pas le mérite d'être totalement distribué car, en principe, chaque utilisateur doit pouvoir choisir sa stratégie à n'importe quel moment.
- Pour remédier au problème d'ordonnancement, l'algorithme synchrone permet aux utilisateurs de mettre simultanément à jour, à chaque itération, leur allocation de puissance en considérant les interférences générées à l'itération précédente. Il a été prouvé que cet algorithme converge plus rapidement vers l'équilibre de Nash, en particulier lorsque le nombre d'utilisateurs actifs dans le réseau est grand [107].

5.2.2.3 Simulation

Dans ce paragraphe, nous examinons les performances de l'algorithme d'allocation de puissance pour deux utilisateurs. Les simulations sont réalisées sur des canaux OMEGA avec les paramètres donnés dans le tableau 4.5. Comme annoncé plus haut, nous nous plaçons dans l'hypothèse d'une connaissance parfaite de l'état des canaux de transmission au niveau de chaque émetteur. À partir de ces informations et à chaque itération, l'algorithme d'allocation des puissances est appliqué pour un seul utilisateur en considérant que l'allocation de l'autre utilisateur est fixe et connue. Les canaux directs et interférents sont choisis de sorte à respecter les conditions du théorème 2. De plus, les canaux sont normalisés pour s'approcher au mieux des conditions limites, c'est-à-dire $0.9 < \max_n \alpha_j n < 1$, ($j = 1, 2; j \neq k$). Dans ce contexte, l'existence de l'équilibre de Nash dépend de l'existence et de la positivité de μ_k en fonction des débits minimums requis par les utilisateurs.

Nous nous intéressons à des niveaux de débits minimums permettant d'atteindre un équilibre. La figure 5.6 présente le taux d'existence de l'équilibre de Nash en fonction du rapport entre la capacité des canaux et les débits minimum requis. Ce taux croît lorsque

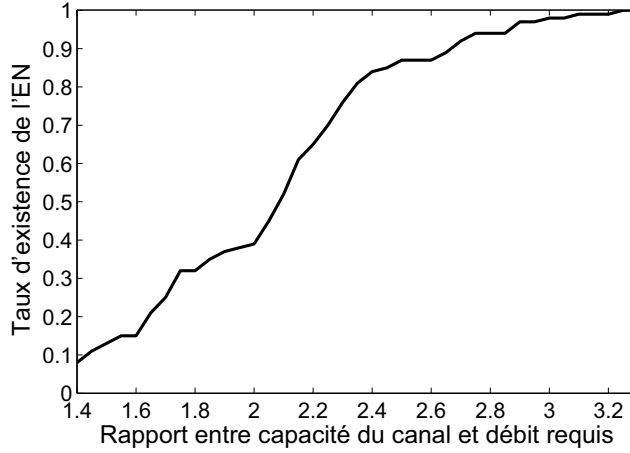


FIGURE 5.6 – Taux d'existence d'un équilibre de Nash versus le rapport entre la capacité du canal des utilisateurs et le débit requis.

le débit requis s'éloigne de la capacité du canal. Lorsque les contraintes de débits sont fixées à la moitié de la capacité des canaux des utilisateurs, l'algorithme converge vers un équilibre dans 39% des cas testés. Pour un système donné, il n'existera pas d'équilibre de Nash lorsque certaines contraintes de débits ne seront plus atteignables. En effet, dans un tel jeu, chaque utilisateur augmente sa puissance de transmission afin de satisfaire sa propre contrainte de débit, augmentant ainsi l'interférence générée pour les autres utilisateurs. L'augmentation des puissances de transmission de tous les utilisateurs ne garantit donc pas la satisfaction des contraintes de débit de tous les utilisateurs et donc l'existence d'un équilibre.

La figure 5.7 présente l'évolution de la somme des puissances allouées aux différentes sous-porteuses, lorsque l'équilibre de Nash existe. Les contraintes de débits pour les utilisateurs 1 et 2 sont respectivement $R_{\min}^1 = 8,1$ kb/symbole OFDM et $R_{\min}^2 = 6,3$ kb/symbole OFDM. À chaque itération et sous l'hypothèse de la connaissance de l'allocation de puissance de l'autre utilisateur, chaque utilisateur augmente sa puissance de transmission de sorte à respecter son débit minimum requis. Si les débits requis sont élevés, chaque utilisateur augmentera sa puissance jusqu'à atteindre la DSP (10^{-8} W/Hz). La figure 5.8 présente l'évolution de l'allocation de puissance par sous-porteuse lorsque les contraintes de débits sont fixées respectivement à $R_{\min}^1 = 8,1$ kb/symbole OFDM et $R_{\min}^2 = 6,9$ kb/symbole OFDM. Après 6 itérations, l'utilisateur 1 doit allouer la DSP par sous-porteuse pour atteindre son débit maximum, qui lui même est inférieur au débit requis. En réponse, l'utilisateur 2 doit aussi allouer la DSP par sous-porteuse. Au final, les contraintes de débits ne sont plus respectées.

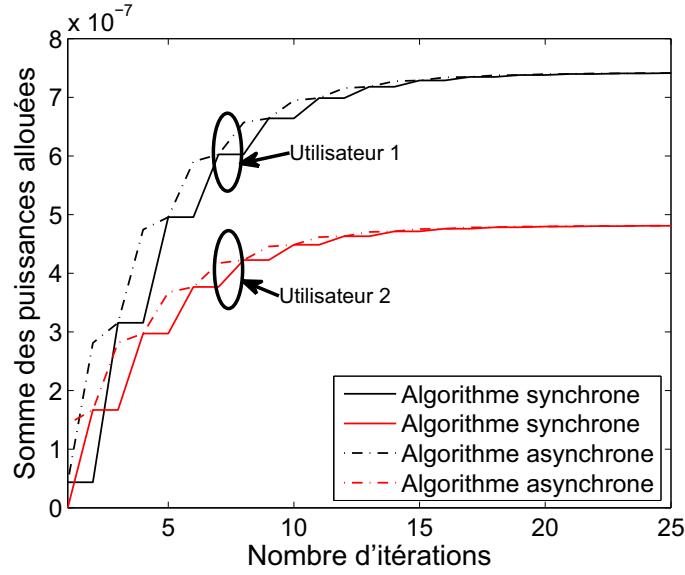


FIGURE 5.7 – Évolution de la somme des puissances allouées aux différentes sous-porteuses avec $R_{\min}^1 = 8,1$ kb/symbole OFDM et $R_{\min}^2 = 6,3$ kb/symbole OFDM.

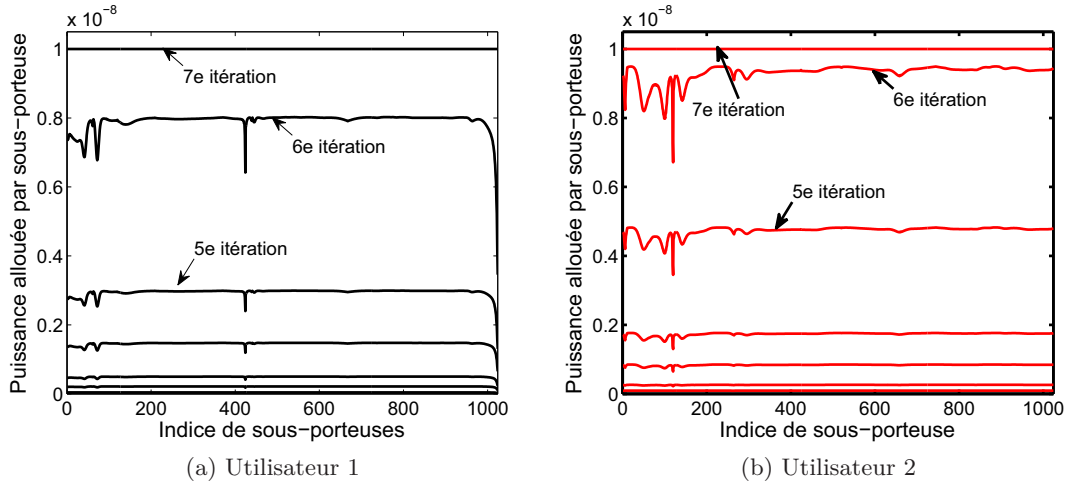


FIGURE 5.8 – Évolution de la répartition de la puissance par sous-porteuse avec $R_{\min}^1 = 8,1$ kb/symbole OFDM et $R_{\min}^2 = 6,9$ kb/symbole OFDM.

5.2.3 Conclusion

Dans cette seconde partie, nous avons étudié le problème d'allocation décentralisée des ressources dans un réseau CPL. L'idée étant de permettre plusieurs communications simultanées sur le support, nous avons étudié le problème de minimisation des puissances de transmission sous les contraintes de débits minimums pour les utilisateurs. Cette première étude a été menée sous forme de jeu non coopératif pour le cas simple de deux

utilisateurs. Les formules analytiques permettant d'obtenir les puissances minimales ont été développées. Lorsque les débits minimums requis par les utilisateurs sont atteignables, le système possède un unique équilibre de Nash où les différents utilisateurs peuvent communiquer simultanément sur le support.

5.3 Conclusion du chapitre

Dans ce chapitre, nous avons étudié dans une première approche, la possibilité de plusieurs communications simultanées dans un réseau CPL *indoor* sur le même support physique. Le problème d'allocation des ressources dans une telle éventualité a été analysé dans les contextes centralisé et décentralisé. Nous avons, tout d'abord, analysé le problème d'allocation des ressources dans un contexte centralisé et proposé un nouvel algorithme d'allocation des ressources pour les systèmes LP-OFDMA. Ensuite, le problème de minimisation des puissances de transmission sous les contraintes de débits minimums et de densité spectrale de puissance a été étudié dans un contexte décentralisé. Les premiers résultats de simulation ont permis de mettre en évidence un certain nombre de points tout à fait pertinents. L'utilisation d'un dispositif centralisé permet d'augmenter les débits de transmission dans le réseau CPL. La mise en place d'un tel dispositif nécessite néanmoins la connaissance des conditions de canal de tous les utilisateurs au niveau du nœud central. Par contre, l'utilisation d'un dispositif décentralisé permet de limiter les échanges entre les différents nœuds en se basant sur les informations connues localement par chaque nœud. Les communications simultanées sont sans collision (sans perte de données) après un certain temps de convergence. Une comparaison intéressante des deux dispositifs consisterait à mesurer le temps nécessaire au nœud central pour collecter toutes les informations des différents nœuds et le temps de convergence vers un équilibre dans le cas décentralisé. Ce qui permettrait d'évaluer la pertinence des deux approches.

Conclusion générale

Le travail mené durant cette thèse a permis d'étudier et de proposer différentes stratégies d'allocation des ressources dans des contextes mono et multi-utilisateurs pour des communications à très haut débit sur lignes d'énergie.

Les deux premiers chapitres de ce document ont été consacrés à la présentation du contexte de l'étude et des spécifications du système mis en œuvre. La technologie CPL a été décrite au travers d'une synthèse bibliographique qui a également portée sur les principales caractéristiques du canal de transmission. Après avoir passé en revue les techniques de modulations multiporteuses et d'étalement de spectre, le système LP-OFDM, né de la combinaison de la technique de précodage linéaire et de l'OFDM, a été décrit en détails avec ses paramètres. À la fin du second chapitre, nous avons introduit de façon générale le problème d'allocation des ressources en présentant les grands principes, avec en particulier les politiques d'optimisation.

Nous nous sommes intéressés dans le troisième chapitre au problème de maximisation du débit dans un contexte mono-utilisateur. De nouveaux algorithmes d'allocation des bits et des puissances ont été développés et ont servi de base pour le contexte multi-utilisateur. Un nouvel algorithme d'allocation des ressources a été proposé pour le système LP-OFDM avec une mise en œuvre de l'égalisation suivant le critère de l'erreur quadratique moyenne au lieu du critère de distorsion crête qui est pris en référence. Les résultats obtenus montrent que l'utilisation d'un égaliseur basé sur le critère de l'erreur quadratique moyenne offre de meilleures performances comparée à l'utilisation du critère de distorsion crête dans le contexte mono-utilisateur. Cependant, les gains restent relativement faibles du fait de la faible influence des distorsions apportées par le canal de transmission. En outre, les simulations réalisées sur les canaux CPL confirment le fait que le système OFDM n'avait pas la capacité d'exploiter efficacement la ressource en puissance disponible, contrairement au nouveau système LP-OFDM, quel que soit le critère d'égalisation utilisé. Les résultats présentés montrent que la composante de précodage linéaire permet de mutualiser les puissances disponibles sur les différentes sous-porteuses d'un même bloc, et de les exploiter collectivement pour augmenter le débit du système. Dans la dernière partie de ce chapitre, nous avons proposé deux nouveaux algorithmes

de *bit-loading*, de faible complexité, qui cherchent à maximiser le débit du système, tout en respectant une contrainte de taux d'erreur binaire crête ou moyen. Les résultats de simulations montrent que le passage d'une contrainte de TES crête à une contrainte de taux d'erreur binaire crête ou moyen permet d'augmenter le débit des systèmes OFDM.

Dans le quatrième chapitre, nous avons traité le problème de l'allocation des ressources dans un contexte multicast. Les problèmes de maximisation du débit multicast dans les contextes monodébit et multidébit ont été abordés sous la contrainte d'une limitation de DSP. Nous avons proposé un nouveau système LP-OFDM multicast avec des solutions permettant d'accroître les débits de transmissions des différents utilisateurs. Lors des simulations, nous avons mis en évidence le fait que la mise en œuvre de la technique de précodage linéaire permet à la fois d'augmenter les débits totaux et d'améliorer l'équité entre les utilisateurs.

Le dernier chapitre a été consacré à l'étude de cas de communications simultanées sur un même support physique dans le réseau CPL. Nous avons proposé de nouveaux algorithmes d'allocation des ressources entre plusieurs utilisateurs voulant accéder simultanément au même support physique, dans un contexte centralisé et décentralisé. Dans un premier temps, une nouvelle technique d'accès multiple LP-OFDMA combinant la technique de précodage linéaire et l'OFDMA a été proposée dans un contexte centralisé. Les résultats obtenus mettent à nouveau en évidence les bonnes performances de la solution LP-OFDM comparée à la solution OFDM. Dans un second temps, le problème de la minimisation des puissances de transmission sous contrainte de débit minimum a été traité sous forme de jeu non coopératif. Les premiers résultats de simulation montrent qu'il est possible, sous certaines hypothèses, d'avoir des communications simultanées sur le support physique du réseau sans perte de données.

Perspectives

Les perspectives à cette étude sont nombreuses et se situent dans le prolongement direct ou non des travaux déjà menés.

En premier lieu, notons que les performances des solutions proposées dans cette thèse ont été obtenues en considérant la seule présence du bruit stationnaire. Les communications CPL sont affectées par de nombreux bruits et brouilleurs en raison de la grande variété d'appareils connectés au réseau et de la multiplicité des perturbations captées par rayonnement. Les solutions proposées devront être testées et optimisées en présence de bruits impulsifs et de brouilleurs. Des algorithmes spécifiques de réjection des bruits impulsifs, et de gestion des brouilleurs peuvent également être exploités, comme ceux proposés dans [38, 114].

Ces dernières années, les techniques MIMO (multiple input multiple output) se sont

affirmées comme une nouvelle voie très prometteuse pour améliorer la robustesse ou l'efficacité spectrale des systèmes hertziens à haut débit en exploitant la dimension spatiale. Des études [115] montrent que des communications MIMO sont possibles sur les lignes d'énergie. Une perspective est d'examiner et d'étudier les possibilités offertes par les techniques MIMO pour augmenter toujours plus les débits transmis par les communications CPL indoor ou améliorer leur fiabilité.

Concernant les communications multicast, et particulièrement dans le contexte multicast multidébit, une étude approfondie est nécessaire pour comprendre comment mettre en œuvre un tel système.

Enfin, la modélisation, en jeu non coopératif, du problème de la minimisation des puissances de transmissions sous contraintes de débits minimums a permis de montrer l'existence d'un équilibre de Nash. La prochaine étape est de rechercher d'autres types d'équilibre et de savoir si l'équilibre de Nash dans notre cas est optimal. De plus, il serait également intéressant de comparer le temps de convergence d'un tel dispositif décentralisé avec le temps nécessaire à un nœud central pour collecter toutes les informations sur les états des canaux de transmissions dans un dispositif centralisé.

Annexe A

Preuve de l'existence et de l'unicité de l'équilibre de Nash

Pour montrer l'existence et l'unicité de l'équilibre de Nash, pour le problème de minimisation de la puissance sous les contraintes de débit et de DSP, nous montrons que l'application $\text{MR} = (\text{MR}_1(\mathbf{P}_2), \text{MR}_2(\mathbf{P}_1))$ est continue et contractante. Pour cela, montrons que $\text{MR}_1(\mathbf{P}_2)$ est lipschitzienne. Rappelons que $[\text{MR}_1(\mathbf{P}_2)]_n = P_{1,n}^*$ avec

$$P_{1,n}^* = \min \left(E_{\text{DSP}}, [\mu_1 - (\beta_{1,n} + \alpha_{2,n}P_{2,n})]^+ \right), \quad (\text{A.1})$$

où $(x)^+ = \max(0, x)$. Pour tout $\mathbf{P}_2$ et $\mathbf{P}_2'$ dans $\mathcal{P}_2$, on a

$$[\text{MR}_1(\mathbf{P}_2) - \text{MR}_1(\mathbf{P}_2')]_n = \begin{cases} 0, & \text{si } \mu_1 < V_1 \text{ et } \mu_1 < V_1' \\ -E_{\text{DSP}}, & \text{si } \mu_1 < V_1 \text{ et } \mu_1 > V_2' \quad (*) \\ -(\mu_1 - (\beta_{1,n} + \alpha_{2,n}P_{2,n})), & \text{si } \mu_1 < V_1 \text{ et } V_1' \leq \mu_1 \leq V_2' \quad (**) \\ E_{\text{DSP}}, & \text{si } \mu_1 > V_2 \text{ et } \mu_1 < V_1' \\ 0, & \text{si } \mu_1 > V_2 \text{ et } \mu_1 > V_2' \\ E_{\text{DSP}} - (\mu_1 - (\beta_{1,n} + \alpha_{2,n}P_{2,n})), & \text{si } \mu_1 > V_2 \text{ et } V_1' \leq \mu_1 \leq V_2' \\ \mu_1 - (\beta_{1,n} + \alpha_{2,n}P_{2,n}), & \text{si } V_1 \leq \mu_1 \leq V_2 \text{ et } \mu_1 < V_1' \\ \mu_1 - (\beta_{1,n} + \alpha_{2,n}P_{2,n}) - E_{\text{DSP}}, & \text{si } V_1 \leq \mu_1 \leq V_2 \text{ et } \mu_1 > V_2' \\ -\alpha_{2,n}(P_{2,n} - P_{2,n}'), & \text{si } V_1 \leq \mu_1 \leq V_2 \text{ et } V_1' \leq \mu_1 \leq V_2' \end{cases} \quad (\text{A.2})$$

avec

$$\begin{aligned} V_1 &= \beta_{1,n} + \alpha_{2,n}P_{2,n}, \quad V_2 = E_{\text{DSP}} + \beta_{1,n} + \alpha_{2,n}P_{2,n}, \\ V_1' &= \beta_{1,n} + \alpha_{2,n}P_{2,n}', \quad V_2' = E_{\text{DSP}} + \beta_{1,n} + \alpha_{2,n}P_{2,n}'. \end{aligned} \quad (\text{A.3})$$

Pour le cas (*), on a

$$\begin{cases} \mu_1 < V_1 \\ -\mu_1 < V_2' \end{cases} \Rightarrow 0 < E_{\text{DSP}} < \alpha_{2,n}(P_{2,n} - P_{2,n}'),$$

et pour le cas (**), on a

$$\begin{cases} \mu_1 < V_1 \\ V_1' \leq \mu_1 \leq V_2' \end{cases} \Rightarrow 0 > -(\mu_1 - (\beta_{1,n} + \alpha_{2,n}P_{2,n})) > -\alpha_{2,n}(P_{2,n} - P_{2,n}').$$

Le même raisonnement peut être fait pour tous les cas. On en déduit d'abord

$$\left| [\text{MR}_1(\mathbf{P}_2) - \text{MR}_1(\mathbf{P}_2')]_n \right| \leq \alpha_{2,n} \left| (P_{2,n} - P_{2,n}') \right|, \quad (\text{A.4})$$

et ensuite

$$\left\| \text{MR}_1(\mathbf{P}_2) - \text{MR}_1(\mathbf{P}_2') \right\| = \sqrt{\sum_n \left| [\text{MR}_1(\mathbf{P}_2) - \text{MR}_1(\mathbf{P}_2')]_n \right|^2} \leq \max_n \alpha_{2,n} \left\| (\mathbf{P}_2 - \mathbf{P}_2') \right\|. \quad (\text{A.5})$$

L'application MR_1 est lipschitzienne et donc continue et par conséquent admet un point fixe qui est un équilibre de Nash d'après le lemme 1. De plus MR_1 est une contraction si $\max(\alpha_{1,n}) < 1$. Dans ce cas, l'équilibre de Nash est unique.

Bibliographie

- [1] M. Tlich, P. PAGANI, G. AVRIL, F. GAUTHIER, A. ZEDDAM, A. KARTIT, O. ISSON, A. TONELLO, F. PECILE, S. D’ALESSANDRO, T. ZHENG, M. BIONDI, G. MIJIC, K. KRIZNAR, J.-Y. BAUDAIS et A. MAIGA, « PLC channel characterization and modelling », rap. tech., projet OMEGA, déc. 2008.
- [2] P. SIOHAN, A. ZEDDAM, G. AVRIL, P. PAGANI, S. PERSON, M. LE BOT, E. CHEVREAU, O. ISSON, F. ONADO, X. MONGABOURE, F. PECILE, A. TONELLO, S. D’ALESSANDRO, S. DRAKUL, M. VUKSIC, J.-Y. BAUDAIS, A. MAIGA et J.-F. HÉLARD, « State of the art, application scenario and specific requirements for PLC », rap. tech., projet OMEGA, avril 2008.
- [3] M. CRUSSIÈRE, *Etude et optimisation de communications à haut-débit sur lignes d’énergie : exploitation de la combinaison OFDM/CDMA*. Thèse de doctorat, INSA de Rennes, France, nov. 2005.
- [4] F. MUHAMMAD, *Various resource allocation and optimization strategies for high bit rate communications on power lines*. Thèse de doctorat, INSA de Rennes, France, mars 2010.
- [5] J.-P. JAVAUDIN, M. BELLEC et D. V. V. SURACI, « OMEGA ICT project: Towards convergent gigabit home networks », in *IEEE International Symposium on Personal, Indoor and Mobile Radio Communications*, (Cannes, France), p. 1–5, sept. 2008.
- [6] J.-Y. BAUDAIS, *Etude des modulations à porteuses multiples et à spectre étalé : analyse et optimisation*. Thèse de doctorat, INSA de Rennes, France, mai 2001.
- [7] HOMEPLUG POWERLINE ALLIANCE, « Homeplug AV white paper », rap. tech., HomePlug Powerline Alliance, Inc, 2005.
- [8] H. HRASNICA, A. HAIDINE et R. LEHNERT, *Broadband powerline communications: network design*. John Wiley & Sons Ltd, 2004.
- [9] « <http://www.cclab.com/fcc-part-15.htm> ».

- [10] R. RAZAFFERSON, P. PAGANI, A. ZEDDAM, J.-Y. BAUDAIS, A. MAIGA et O. ISSON, « Report on electromagnetic compatibility of power line communications », rap. tech., projet OMEGA, déc. 2009.
- [11] « <http://www.homeplug.org/> ».
- [12] « <http://www.upaplc.org> ».
- [13] J.-P. JAVAUDIN et M. BELLEC, « Gigabit home networks OMEGA ICT project », in *ICT MobileSummit*, (Stockholm, Sweden), juin 2008.
- [14] A. TONELLO, S. D’ALESSANDRO, M. ANTONIALI, M. BIONDI, F. VERSOLATTO, A. MAIGA, J.-Y. BAUDAIS, F. S. MUHAMMAD, P. SIOHAN, M. L. BOT, H. LIN, P. ACHAICHA, G. NDO, G. MIJIC, B. CERATO, S. DRAKUL, E. VITERBO et O. ISSON, « Performance report of optimized PHY algorithms », rap. tech., projet OMEGA, juin 2010.
- [15] A. MAIGA, J.-Y. BAUDAIS, O. ISSON, A. TONELLO et S. D’ALESSANDRO, « Optimized MAC algorithms and performance report », rap. tech., projet OMEGA, mai 2010.
- [16] H. PHILIPPS, « Modeling of power line communications channels », in *IEEE International Symposium on Power Line Communications and Its Applications*, (Lancaster, U.K.), p. 14–21, avril 1999.
- [17] M. ZIMMERMANN et K. DOSTERT, « A multipath model for the power line channel », *IEEE Transactions on Communications*, vol. 50, p. 553–559, avril 2002.
- [18] O. HOOIJEN, « On the relation between network topology and power line signal attenuation », in *IEEE International Symposium on Power Line Communications and Its Applications*, (Tokyo, Japan), p. 45–56, mars 1998.
- [19] S. GALLI et T. BANWELL, « The indoor power line channel: New results and modem design considerations », in *IEEE Consumer Communications and Networking Conference*, (Las Vegas, USA), p. 5–8, jan. 2004.
- [20] H. MENG, S. CHEN, Y. GUAN, C. LAW, P. SO, E. GUNAWAN et T. LIE, « Modeling of transfer characteristics for the broadband power line communication channel », *IEEE Transactions on Power Delivery*, vol. 19, p. 1057–1064, juil. 2004.
- [21] M. KUHN, S. BERGER, I. HAMMERSTROM et A. WITTNEBEN, « Power line enhanced cooperative wireless communications », *IEEE Journal on Selected Areas in Communications*, vol. 24, p. 1401–1410, juil. 2006.
- [22] J. CORTES, F. CANETE, L. DIEZ et J. ENTRAMBASAGUAS, « Characterization of the cyclic short-time variation of indoor power-line channels response », in *IEEE International Symposium on Power Line Communications and Its Applications*, (Vancouver, Canada), p. 326–330, avril 2005.

- [23] A. TONELLO, J. CORTES et S. D'ALESSANDRO, « Optimal time slot design in an OFDM-TDMA system over power-line time-variant channels », in *IEEE International Symposium on Power Line Communications and Its Applications*, p. 41–46, mars 2009.
- [24] M. ZIMMERMANN et K. DOSTERT, « Analysis and modeling of impulsive noise in broad-band powerline communications », *IEEE Transactions on Electromagnetic Compatibility*, vol. 44, p. 249–258, fév. 2002.
- [25] G. AVRIL, *Étude et Optimisation des systèmes à courant porteurs domestiques face aux perturbations du réseau électrique*. Thèse de doctorat, INSA de Rennes, France, oct. 2008.
- [26] D. LIU, E. FLINT, B. GAUCHER et Y. KWARK, « Wide band AC power line characterization », *IEEE Transactions on Consumer Electronics*, vol. 45, p. 1087–1097, nov. 1999.
- [27] V. DEGARDIN, M. LIENARD, P. DEGAUQUE, A. ZEDDAM et F. GAUTHIER, « An analysis of the broadband noise scenario in powerline networks », in *IEEE International Symposium on Electromagnetic Compatibility*, (Istanbul, Turkey), p. 131–138, mai 2003.
- [28] V. B. BALAKIRSKY et A. J. H. VINCK, « Potential limits on power-line communication over impulsive noise channels », in *IEEE International Symposium on Power Line Communications and Its Applications*, (Kyoto, Japan), p. 32–37, mars 2003.
- [29] R. CHANG et R. GIBBY, « A theoretical study of performance of an orthogonal multiplexing data transmission scheme », *IEEE Transactions on Communication Technology*, vol. 16, p. 529–540, août 1968.
- [30] A. STEPHAN, *Resource Allocation Strategies and Linear Precoded OFDM Optimization for Ultra-Wideband Communications*. Thèse de doctorat, INSA de Rennes, France, déc. 2008.
- [31] O. ISSON, J.-M. BROSSIER et D. MESTDAGH, « Multi-carrier bit-rate improvement by carrier merging », *Electronics Letters*, vol. 38, no. 19, p. 1134–1135, 2002.
- [32] R. C. DIXON, *Spread spectrum systems*. John Wiley & Sons, 2nd éd., 1984.
- [33] C. Y. WONG, R. S. CHENG, K. B. LETAIEF et R. D. MURCH, « Multiuser OFDM with adaptive subcarrier, bit, and power allocation », *IEEE Journal on Selected Areas of Communications*, vol. 17, p. 1747–1758, oct. 1999.
- [34] K. WESOŁOWSKI, *Introduction to digital communication systems*. John Wiley & Sons Ltd, 2009.
- [35] L. CARIOU et J.-F. HÉLARD, « A simple and efficient channel estimation for MIMO OFDM code division multiplexing uplink systems », in *IEEE Workshop on Signal Processing Advances in Wireless Communications*, p. 176–180, juin 2005.

- [36] J.-Y. BAUDAIS et M. CRUSSIÈRE, « Allocation MC-CDMA : augmentation des débits sur les lignes de transmission », in *Colloque GRETSI*, (Louvain-la-Neuve, Belgique), p. 735–738, sept. 2005.
- [37] O. MACCHI, « L'égalisation numérique en communications », *Annales des télécommunications*, vol. 53, no. 1–2, p. 39–58, 1998.
- [38] E. GUEGUEN, *Etude et optimisation des techniques UWB haut débit multibandes OFDM*. Thèse de doctorat, INSA de Rennes, France, jan. 2009.
- [39] N. YEE et J.-P. LINNARTZ, « Controlled equalization of multi-carrier CDMA in an indoor Rician fading channel », in *IEEE Vehicular Technology Conference*, vol. 3, (Stockholm, Sweden), p. 1665–1669, juin 1994.
- [40] S. B. SLIMANE, « Partial equalization of multi-carrier CDMA in frequency selective fading channels », in *IEEE International Conference on Communications*, vol. 1, p. 26–30, août 2000.
- [41] J. CIOFFI, « A multicarrier primer », committee contribution, ANSI T1E1.4/91157, nov. 1991.
- [42] K. CHO et D. YOO, « On the general BER expression of one- and two-dimensional amplitude modulations », *IEEE Transactions on Communications*, vol. 50, p. 1074–1080, juil. 2002.
- [43] T. ZOGAKIS, J. T. ASLANIS et J. CIOFFI, « A coded and shaped discrete multitone system », *IEEE Transactions on Communications*, vol. 43, p. 2941–2949, déc. 1995.
- [44] A. WYGLINSKI, *Physical Layer Loading Algorithms for Indoor Wireless Multicarrier Systems*. Thèse de doctorat, McGill University, Montreal, Canada, nov. 2004.
- [45] W. YU et J. M. CIOFFI, « On constant power water-filling », in *IEEE International Conference on Communications*, vol. 6, (Helsinki, Finland), p. 1665–1669, juin 2001.
- [46] M. COLLIN, *Etude de l'optimisation d'un système DMT-ADSL : application à la transmission vidéo MPEG-2 en mode hiérarchique*. Thèse de doctorat, Université de Valenciennes et du Hainaut Cambrésis, jan. 1999.
- [47] J. CAMPELLO, *Discrete bit loading for multicarrier modulation systems*. Thèse de doctorat, Stanford University, USA, mai 1999.
- [48] P. CHOW, J. CIOFFI et J. BINGHAM, « A practical discrete multitone transceiver loading algorithm for data transmission over spectrally shaped channels », *IEEE Transactions on Communications*, vol. 43, p. 773–775, fév. 1995.
- [49] A. CZYLWIK, « Adaptive OFDM for wideband radio channels », in *IEEE Global Telecommunications Conference*, vol. 1, p. 713–718, nov. 1996.

- [50] M. CRUSSIÈRE, J.-Y. BAUDAIS et J.-F. HÉLARD, « Adaptive spread-spectrum multicarrier multiple-access over wirelines », *IEEE Journal on Selected Areas in Communications*, vol. 24, p. 1377–1388, juil. 2006.
- [51] D. GERAKOULIS et S. GHASSEMZADEH, « Extended orthogonal code designs with applications in CDMA », in *IEEE International Symposium on Spread Spectrum Techniques and Applications*, vol. 2, p. 657–661, 2000.
- [52] A. S. HEDAYAT, N. J. A. SLOANE et J. STUFKEN, *Orthogonal Arrays: Theory and Applications*, vol. 7. New York : Springer-Verlag, 1999.
- [53] T. M. COVER et J. A. THOMAS, *Elements of Information Theory*. New York : John Wiley & Sons, 2006.
- [54] D. GUININ et B. JOPPIN, *Mathématique — Exercices — MPSI*. Bréal, 1 éd., nov. 2003.
- [55] A. WYGLINSKI, F. LABEAU et P. KABAL, « Bit loading with BER-constraint for multicarrier systems », *IEEE Transactions on Wireless Communication*, vol. 4, p. 1383–1387, juil. 2005.
- [56] F. MUHAMMAD, J.-Y. BAUDAIS, J.-F. HÉLARD et M. CRUSSIÈRE, « Coded adaptive linear precoded discrete multitone over PLC channel », in *IEEE International Symposium on Power Line Communications and Its Applications*, (Jeju city, Brazil), p. 123–128, avril 2008.
- [57] E. GUERRINI, G. DELL’AMICO, P. BISAGLIA et L. GUERRIERI, « Bit-loading algorithms and SNR estimate for homeplug AV », in *IEEE International Symposium on Power Line Communications and Its Applications*, (Pisa, Italy), p. 419–424, mars 2007.
- [58] Y. GEORGE et O. AMRANI, « Bit loading algorithms for OFDM », in *IEEE International Symposium on Information Theory*, (Chicago, USA), p. 391, juin 2004.
- [59] B. WANG et J. HOU, « Multicast routing and its QoS extension: problems, algorithms, and protocols », *IEEE Journal of Network*, vol. 14, p. 22–36, jan. 2000.
- [60] L. SAHASRABUDDHE et B. MUKHERJEE, « Multicast routing algorithms and protocols: a tutorial », *IEEE Journal of Network*, vol. 14, p. 90–102, jan. 2000.
- [61] N. MIR, « A survey of data multicast techniques, architectures, and algorithms », *IEEE Communications Magazine*, vol. 39, p. 164–170, sept. 2001.
- [62] C. SUH et C.-S. HWANG, « Dynamic subchannel and bit allocation multicast OFDM systems », in *IEEE International Symposium on Personal, Indoor and Mobile Radio Communications*, vol. 3, p. 2102–2106, sept. 2004.

- [63] J. SCHMITT, F. ZDARSKY, M. KARSTEN et R. STEINMETZ, « Heterogeneous multicast in heterogeneous QoS networks », in *IEEE International Conference on Networks*, (Bangkok, Thailand), p. 349–354, oct. 2001.
- [64] A. MOHAMED et H. ALNUWEIRI, « Cross-layer optimal rate allocation for heterogeneous wireless multicast », *EURASIP Journal on Wireless Communication Networks*, vol. 2009, p. 1–16, jan. 2009.
- [65] C. SUH et J. MO, « Resource allocation for multicast services in multicarrier wireless communications », in *IEEE International Conference on Computer Communications*, (Barcelona, Spain), p. 1–12, avril 2006.
- [66] J. PROAKIS, *Digital Communications*. McGraw-Hill Science, 4 éd., août 2000.
- [67] H. A. DAVID et H. N. NAGARAJA, *Order Statistics*. Wiley InterScience, 3 éd., 2003.
- [68] I. S. GRADSHTEYN et I. M. RYZHIK, *Table of Integrals, Series, and Products*. Academic Press, 7 éd., jan. 2007.
- [69] J.-Y. BAUDAIS et M. CRUSSIÈRE, « Resource allocation with adaptive spread spectrum OFDM using 2D spreading for power line communications », *EURASIP Journal on Applied Signal Processing*, vol. 2007, p. 1–13, 2007.
- [70] C.-S. HWANG et Y. KIM, « An adaptive modulation method for multicast communications of hierarchical data in wireless networks », in *IEEE International Conference on Communications*, (New York, USA), p. 896–900, avril 2002.
- [71] X. SUN, S. LI, F. WU, G. SHEN et W. GAO, « Efficient and flexible drift-free video bitstream switching at predictive frames », in *IEEE International Conference on Multimedia and Expo*, vol. 1, p. 541–544, août 2002.
- [72] R. WIJNHOFEN, E. JASPERS et P. de WIT, « Multi-channel video streaming server for surveillance systems », in *IEEE International Symposium on Consumer Electronics*, (Reading, UK), p. 353–358, sept. 2004.
- [73] X. BAI, A. SHAMI et C. ASSI, « On the fairness of dynamic bandwidth allocation schemes in ethernet passive optical networks », *Journal of Computer Communications*, vol. 29, no. 11, p. 2123–2135, 2006.
- [74] J.-Y. LE BOUDEC, « Rate adaptation, congestion control and fairness: a tutorial ». Note de cours, École Polytechnique Fédérale de Lausanne, déc. 2008.
- [75] D.-M. CHIU et R. JAIN, « Analysis of the increase and decrease algorithms for congestion avoidance in computer networks », *Computer Networks and ISDN Systems*, vol. 17, no. 1, p. 1–14, 1989.
- [76] A. JALALI, R. PADOVANI et R. PANKAJ, « Data throughput of CDMA-HDR a high efficiency-high data rate personal communication wireless system », in *IEEE*

- Vehicular Technology Conference Spring*, vol. 3, (Tokyo, Japan), p. 1854–1858, mai 2000.
- [77] A. MAIGA, J.-Y. BAUDAIS et J.-F. HÉLARD, « Subcarrier, bit and time slot allocation for multicast precoded OFDM systems », in *IEEE International Conference on Communications*, (Cap Town, South Africa), p. 1–6, mai 2010.
- [78] A. MAIGA, J.-Y. BAUDAIS et J.-F. HÉLARD, « Increase in multicast OFDM data rate in PLC network using adaptive LP-OFDM », in *International Conference on Adaptive Science Technology*, (Accra, Ghana), p. 384–389, déc. 2009.
- [79] A. MAIGA, J.-Y. BAUDAIS et J.-F. HÉLARD, « Allocation des ressources basée sur le précodage linéaire pour les systèmes OFDM multicast », in *Colloque GRETSI*, (Dijon, France), sept. 2009.
- [80] A. MAIGA, J.-Y. BAUDAIS et J.-F. HÉLARD, « Maximisation du débit des systèmes OFDM multicast dans un contexte de courant porteur en ligne », in *Colloque MajecSTIC*, (Avignon, France), nov. 2009.
- [81] A. MAIGA, J.-Y. BAUDAIS et J.-F. HÉLARD, « Bit rate optimization with MMSE detector for multicast LP-OFDM systems ». soumis à *IEEE transactions on power delivery*.
- [82] T. SARTENAER, *Multiuser communications over frequency selective wired channels and applications to the powerline access network*. Thèse de doctorat, Université Catholique de Louvain, Louvain, Belgium, sept. 2004.
- [83] E. BAKHTIARI et B. KHALAJ, « A new joint power and subcarrier allocation scheme for multiuser OFDM systems », in *IEEE Personal, Indoor and Mobile Radio Communications*, vol. 2, p. 1959–1963, sept. 2003.
- [84] D. KIVANC et H. LIU, « Subcarrier allocation and power control for OFDMA », in *Asilomar Conference on Signals, Systems and Computers*, vol. 1, (Pacific Grove, CA), p. 147–151, oct. 2000.
- [85] L. ZHEN, Z. GEQING, W. WEIHUA et S. JUNDE, « Improved algorithm of multiuser dynamic subcarrier allocation in OFDM system », in *International Conference on Communication Technology*, vol. 2, p. 1144–1147, avril 2003.
- [86] M. CHO, W. SEO, Y. KIM et D. HONG, « A joint feedback reduction scheme using delta modulation for dynamic channel allocation in OFDMA systems », in *IEEE International Symposium on Personal, Indoor and Mobile Radio Communications*, vol. 4, (Berlin, Germany), p. 2747–2750, sept. 2005.
- [87] G. ZHANG, « Subcarrier and bit allocation for real-time services in multiuser OFDM systems », in *IEEE International Conference on Communications*, vol. 5, (Paris, France), p. 2985–2989, juin 2004.

- [88] G. YU, Z. ZHANG, Y. CHEN, J. SHI et P. QIU, « A novel resource allocation algorithm for real-time services in multiuser OFDM systems », in *IEEE Vehicular Technology Conference Spring*, vol. 3, (Melbourne, Australia), p. 1156–1160, mai 2006.
- [89] W. WANG, T. OTTOSSON, M. STERNADT, A. AHLEN et A. SVENSSON, « Impact of multiuser diversity and channel variability on adaptive OFDM », in *IEEE Vehicular Technology Conference Fall*, vol. 1, p. 547–551, oct. 2003.
- [90] W. RHEE et J. CIOFFI, « Increase in capacity of multiuser OFDM system using dynamic subchannel allocation », in *IEEE Vehicular Technology Conference Spring*, vol. 2, (Tokyo, Japan), p. 1085–1089, mai 2000.
- [91] J. JANG et K. B. LEE, « Transmit power adaptation for multiuser OFDM systems », *IEEE Journal on Selected Areas in Communications*, vol. 21, p. 171–178, fév. 2003.
- [92] C. LENGOUNBI, *Accès multiple OFDMA pour les systèmes cellulaires post 3G : allocation de ressources et ordonnancement*. Thèse de doctorat, TELECOM ParisTech (ENST), Paris, France, mars 2008.
- [93] H. KWON, W.-I. LEE et B. G. LEE, « Low-overhead resource allocation with load balancing in multi-cell OFDMA systems », in *IEEE Vehicular Technology Conference Spring*, vol. 5, (Stockholm, Sweden), p. 3063–3067, mai 2005.
- [94] S. HAMOUDA, P. GODLEWSKI et S. TABBANE, « Enhanced capacity for multi-cell OFDMA systems with efficient power control and reuse partitioning », in *IEEE International Conference on Communication systems*, (Singapore), p. 1–5, oct. 2006.
- [95] K. KIM, H. KIM et Y. HAN, « Subcarrier and power allocation in OFDMA systems », in *IEEE Vehicular Technology Conference Fall*, vol. 2, p. 1058–1062, sept. 2004.
- [96] Y. PENG, A. DOUFEXI, S. ARMOUR et J. MCGEEHAN, « An investigation of dynamic sub-carrier allocation in OFDMA systems », in *IEEE Vehicular Technology Conference Spring*, vol. 3, p. 1808–1811, mai 2005.
- [97] D. KIVANC, G. LI et H. LIU, « Computationally efficient bandwidth allocation and power control for OFDMA », *IEEE Transactions on Wireless Communications*, vol. 2, p. 1150–1158, nov. 2003.
- [98] J. GROSS, H. KARL, F. FITZEK et A. WOLISZ, « Comparison of heuristic and optimal subcarrier assignment algorithms », in *International Conference on Wireless Networks*, (Las Vegas, USA), p. 249–255, juin 2003.

- [99] S. KO, J. HEO et K. CHANG, « Aggressive subchannel allocation algorithm for efficient dynamic channel allocation in multi-user OFDMA system », in *IEEE International Symposium on Personal, Indoor and Mobile Radio Communications*, (Helsinki, Finland), p. 1–5, sept. 2006.
- [100] I. WONG, Z. SHEN, B. EVANS et J. ANDREWS, « A low complexity algorithm for proportional resource allocation in OFDMA systems », in *IEEE Workshop on Signal Processing Systems*, p. 1–6, oct. 2004.
- [101] M. JIMENEZ, « Algorithme hongrois ». Cours, École polytechnique de Montréal, automne 2001.
- [102] M. NATKANIEC et A. R. PACH, « An analysis of the modified backoff mechanism for IEEE 802.11 networks », in *First Polish-German Teletraffic Symposium*, (Dresden, Germany), p. 24–26, sept. 2000.
- [103] W. YU, G. GINIS et J. CIOFFI, « Distributed multiuser power control for digital subscriber lines », *IEEE Journal on Selected Areas in Communications*, vol. 20, p. 1105–1115, juin 2002.
- [104] F. MESHKATI, M. CHIANG, H. POOR et S. SCHWARTZ, « A game-theoretic approach to energy-efficient power control in multicarrier CDMA systems », *IEEE Journal on Selected Areas in Communications*, vol. 24, p. 1115–1129, juin 2006.
- [105] S. T. CHUNG, S. J. KIM, J. LEE et J. CIOFFI, « A game-theoretic approach to power allocation in frequency-selective gaussian interference channels », in *IEEE International Symposium on Information Theory*, (Kanagawa, Japan), p. 316–316, juin 2003.
- [106] P. V. WRYCZA, M. R. B. SHANKAR, M. BENGTSSON et B. OTTERSTEN, « Spectrum allocation for decentralized transmission strategies: properties of Nash equilibria », *EURASIP Journal on Advance Signal Process*, vol. 2009, p. 1–11, 2009.
- [107] J.-S. PANG, G. SCUTARI, F. FACCHINEI et C. WANG, « Distributed power allocation with rate constraints in gaussian parallel interference channels », *IEEE Transactions on Information Theory*, vol. 54, p. 3471–3489, août 2008.
- [108] J.-F. THISSE, « Théorie des jeux : une introduction ». Cours, Université catholique de Louvain, oct. 2003.
- [109] S. KAKUTANI, « A generalization of Brouwer's fixed point theorem », *Duke Mathematical Journal*, vol. 8, p. 457–459, jan. 1941.
- [110] J. B. ROSEN, « Existence and uniqueness of equilibrium points for concave n-person games », *Econometrica*, vol. 33, p. 520–534, juil. 1965.
- [111] S. BOYD et L. VANDENBERGHE, « Convex optimization », in *Cambridge University Press*, (Cambridge, U.K.), mars 2004.

- [112] D. PALOMAR et J. FONOLLOSA, « Practical algorithms for a family of waterfilling solutions », *IEEE Transactions on Signal Processing*, vol. 53, p. 686– 695, fév. 2005.
- [113] C. CAMERER, *Behavioral Game Theory: Experiments in Strategic Interaction*. New Jersey, USA : Princeton University Press, mars 2003.
- [114] V. DÉGARDIN, *Analyse de la faisabilité d'une transmission de données haut-débit sur le réseau électrique basse tension*. Thèse de doctorat, Université des sciences et technologies de Lille, France, déc. 2002.
- [115] R. HASMAT, P. PAGANI, A. ZEDDAM et T. CHONAVEL, « MIMO communications for inhome PLC networks: Measurement and results up to 100 MHz », in *IEEE International Symposium on Power Line Communications and Its Applications*, (Rio de Janeiro, Brésil), p. 120–124, mars 2010.

Au cours des dernières années, la demande de services sur les réseaux CPL (courant porteur en ligne) a connu une forte augmentation du fait de la disponibilité des infrastructures et du faible coût de déploiement. Ce type de réseau supporte plusieurs trafics à haut débit avec une facilité de connexion entre plusieurs nœuds et points d'accès. Dans ce contexte, une gestion efficace des ressources disponibles, par l'intermédiaire de politiques d'allocation, s'avère indispensable pour satisfaire les contraintes de qualité de service. Ces politiques consistent à répartir efficacement les ressources dans le but d'optimiser les débits de transmission dans le réseau CPL. La présente thèse propose des stratégies d'allocation des ressources permettant d'augmenter les débits de transmission dans les contextes mono et multi-utilisateurs. Cependant, le canal de propagation CPL est peu favorable à la transmission de données à haut débit, puisqu'il n'a pas, à l'origine, été conçu dans ce but. Afin d'exploiter ce canal difficile, les données sont transmises via une forme d'onde combinant la technique de précodage linéaire aux modulations à porteuses multiples de type OFDM et conduisant à la solution LP-OFDM (linear precoded OFDM). Sous l'hypothèse d'une connaissance parfaite de la réponse du canal, cette combinaison permet une exploitation plus efficace de la puissance disponible. Les débits de transmission sont alors augmentés en adaptant les ordres de modulation, les niveaux de puissance et la répartition des ressources temps-fréquences aux conditions des liens de transmission.

L'objectif principal de cette thèse est d'étudier et d'optimiser les stratégies de distribution :

- des sous-canaux (qui peuvent être les sous-porteuses du système OFDM ou les séquences de précodage du système LP-OFDM) ;
- des bits et puissances attribués à ces sous-canaux.

Dans un premier temps, le problème de maximisation du débit dans un contexte mono-utilisateur est étudié et sert de base pour le contexte multi-utilisateur. Un nouvel algorithme d'allocation des ressources pour le système LP-OFDM avec une mise en œuvre de l'égalisation suivant le critère de l'erreur quadratique moyenne est proposé, ce qui constitue une première contribution originale. De plus, deux autres nouveaux algorithmes d'adaptation des ordres de modulation et des niveaux de puissance, de faible complexité, sont proposés pour maximiser le débit total tout en respectant une contrainte de taux d'erreur binaire.

Dans un second temps, une approche couche physique de la communication multicast est étudiée pour le système LP-OFDM. Les méthodes proposées permettent de mieux exploiter la diversité des liens de transmission pour augmenter les débits des utilisateurs. Comparés à la méthode classique d'allocation des ressources en OFDM multicast, les résultats de simulations montrent des gains en débit pouvant atteindre 70%.

Enfin, la possibilité pour plusieurs utilisateurs de transmettre simultanément des données dans un réseau CPL est analysée. Les systèmes actuels CPL sont caractérisés par des procédés d'accès multiple où les différents utilisateurs transmettent leurs signaux dans des intervalles de temps distincts. De nouveaux algorithmes d'allocation des ressources, dans un contexte centralisé ou décentralisé, sont alors proposés et analysés pour une transmission simultanée dans le réseau. Les résultats obtenus mettent à nouveau en évidence l'intérêt de la solution LP-OFDM.

The demand for broadband services over PLC (powerline communications) networks has been growing rapidly during the recent past due to the availability of infrastructures and low deployment costs. This kind of network supports multiple high-speed traffics with seamless connectivity among multiple nodes and access points. In this context, efficient management of available resources through allocation policies is needed to satisfy the quality of service requirements. These policies consist in determining efficient rules for allocating resources in order to optimize the transmission rates in PLC network. This thesis proposes resource allocation strategies to increase transmission rates in single and multi-user contexts. However, the transmission characteristics of the powerline channel are less favorable for data transfer, since it was originally not designed for that purpose. To exploit these difficult channel conditions, data is transmitted via a waveform combining the linear precoding technique and the multicarrier OFDM modulation scheme, leading to the LP-OFDM (linear precoding OFDM) solution. Assuming a perfect knowledge of the channel conditions at the transmitter side, this combination allows a more efficient utilization of the available transmission power. The achieved data rates are then increased by adapting modulation orders, transmitted power levels and the distribution of time-frequency resources, to the channel conditions.

The main objective of this thesis is to study and optimize distribution strategies, for one or more users, of different subchannels (subcarriers in OFDM case and precoding sequences in LP-OFDM case) of the multicarrier systems and the bits and powers allocated to these subchannels. First, the problem of maximizing the bit rate is studied in a single user context and uses as the basis for multi-user context. A new resource allocation algorithm for LP-OFDM systems with minimum mean square error equalizer is proposed and constitutes the first original contribution. In addition, two novel bit and power allocation algorithms, with low complexity, are proposed to maximize the total bit rate while satisfying a bit error rate constraint.

Then, a physical layer approach of multicast communications is addressed for LP-OFDM systems. The proposed methods better exploit the diversities of transmission links to increase the users' bit rates. Compared to the conventional resource allocation method in multicast OFDM systems, simulation results show bit rate gains up to 70% with linear precoding based methods.

Finally, the possibility for several users to simultaneously access to the same physical medium is analysed for PLC networks. Current powerline communication systems are characterized by multiple access methods where different users transmit their signals in separate time intervals. New resource allocation algorithms are then proposed and analysed for simultaneous transmission over the same physical medium, in centralized and decentralized manner. Again, the results show the interest of the LP-OFDM solution.

Utiliser la police Arial Taille 9 dans les champs texte « résumé » et « abstract » - Texte justifié -
Ne pas dépasser le nombre de caractères des cadres de texte ci-dessus.
Ne pas modifier la taille des cadres de texte