



Accès multiple OFDMA pour les systèmes cellulaires post 3G : allocation de ressources et ordonnancement

Carle Lengoumbi

► To cite this version:

Carle Lengoumbi. Accès multiple OFDMA pour les systèmes cellulaires post 3G : allocation de ressources et ordonnancement. domain_other. Télécom ParisTech, 2008. English. NNT: . pastel-00003747

HAL Id: pastel-00003747

<https://pastel.hal.science/pastel-00003747>

Submitted on 10 Apr 2009

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 présentée par
Carle Lengoumbi

pour obtenir le grade de DOCTEUR
de l'ECOLE NATIONALE SUPERIEURE DES TELECOMMUNICATIONS DE PARIS
Spécialité "Informatique et Réseaux"



Accès multiple OFDMA pour les systèmes cellulaires post 3G :
allocation de ressources et ordonnancement

Soutenue le 14 mars 2008 devant le jury composé de :

Président

Xavier Lagrange, Télécom Bretagne

Rapporteurs :

Mongi Marzoug, Orange, Issy-les-Moulineaux

Sami Tabbane, SupCom, Tunis

Examineurs :

Corinne Berland, ESIEE Paris

Jerome Brouet, Alcatel

Walid Hachem, Télécom ParisTech

Directeurs de thèse :

Philippe Martins, Télécom ParisTech

Philippe Godlewski, Télécom ParisTech

A mes parents,
Victor et Florence Lengoumbi
Avec tout mon amour,

A mes tantes,
Emérentienne Ngombo
et Marie-Hélène Mbomba

Remerciements

Je tiens tout d'abord à remercier mes deux directeurs de thèse Philippe Martins et Philippe Godlewski sans lesquels cette thèse n'aurait pas vu le jour.

Je remercie chaleureusement Mr Mongi Marzoug et Mr Sami Tabbane d'avoir accepté de relire ma thèse. Je remercie tous les membres du jury pour leur présence à la soutenance et leurs précieux commentaires.

Mes remerciements vont ensuite à plusieurs professeurs et maîtres de conférences du département INFRES, pour leur disponibilité : Laurent Decreusefond, Nicolas Puech, Marceau Coupechoux, et bien d'autres ont souvent pris le temps de répondre à mes nombreuses questions théoriques.

Je ne saurais oublier de remercier de nombreux collègues doctorants pour leur aide pratique et leur bonne humeur. Nicolas Dailly et Kinda Khawam en font partie. Le secrétariat INFRES, notamment Hayette Soussou, Celine Bizart et Sophie Bérenger ne sont pas en reste.

Grâce à l'implication et l'exemplarité de mes parents Victor et Florence Lengoumbi, j'ai obtenu un diplôme d'ingénieur et un doctorat en Télécommunications. Le départ brutal de mon père, peu avant la soutenance, m'empêche de partager avec lui ces précieux moments et tous ceux à venir. Je remercie toute ma famille et plus précisément ma tante Solange Kouya et mon frère Samson Batolo pour leur soutien inconditionnel. Je remercie tous mes amis pour leurs encouragements.

Enfin, je remercie Camille Tacail, celui qui a partagé toutes les épreuves.

Résumé

Cette thèse s'intéresse à l'allocation de ressources des réseaux WMAN (*Wireless Metropolitan Area Networks*) utilisant l'OFDMA (*Orthogonal Frequency Division Multiple Access*). Les applications récentes destinées aux réseaux sans fils requièrent des débits importants. L'OFDMA, technique d'accès multiple basée sur l'OFDM (*Orthogonal Frequency Division Multiplexing*), permet d'obtenir des débits élevés en tirant avantage de la diversité multiutilisateur.

Dans une première partie, l'allocation de ressources est considérée au niveau de la couche physique (affectation des fréquences et de la puissance). Dans un contexte monocellulaire, un algorithme est proposé afin de maximiser le débit de la cellule tout en assurant des débits individuels aux utilisateurs (problème RA, *Rate Adaptive optimization*). L'impact de la sous canalisation sur le débit global est analysé. Puis, un algorithme est proposé dans un contexte multicellulaire pour adapter dynamiquement le facteur de réutilisation fréquentiel.

Dans une deuxième partie, l'ordonnancement est traité. L'objectif est de garantir les délais du trafic temps réel, de maximiser le débit du trafic non temps réel tout en assurant une équité proportionnelle entre les flux. Pour cela, les extensions du GPS (*Global Processor Sharing*) sont étudiées. Deux algorithmes, extensions du WFS (*Wireless Fair Service*) pour l'OFDMA, sont proposés. Leurs performances sont comparées au WFS et aux algorithmes existants en OFDMA. L'une des propositions, l'OWFS (*Opportunist Wireless Fair Service*) est particulièrement adaptée pour maximiser le débit du trafic non temps réel et comporte un paramètre, le poids de délai, qui permet de maintenir un taux de pertes acceptable pour le trafic temps réel.

Mots clés : OFDMA, problème RA, débits minimaux, allocation de bande, affectation de sous porteuses, facteur de réutilisation fréquentiel, ordonnancement, trafic temps réel, trafic non temps réel, équité proportionnelle, WFS.

Abstract

This thesis deals with radio resource allocation in WMAN (*Wireless Metropolitan Area Networks*) networks based on OFDMA (*Orthogonal Frequency Division Multiple Access*). Several applications for wireless networks demand high data rates. OFDMA is a multiple access technique based on OFDM (*Orthogonal Frequency Division Multiplexing*) ; it takes advantage of multiuser diversity and enables high data rates.

In the first part of the thesis, resource allocation is considered at the physical layer (sub-carrier and power allocation). In a single cell network, an algorithm is proposed to maximize the cell data rate while ensuring individual rates to users (*Rate Adaptive optimization*). Data rates obtained under subchannelization are compared to data rates obtained under individual subcarrier allocation. In a multicell network, a new algorithm using a dynamic frequency reuse factor is described.

In a second part, scheduling is studied. The main goal is to provide delay guaranties to real time flows, maximize throughput of non real time flows while ensuring proportional fairness to the flows. GPS (*Global Processor Sharing*) extensions are examined. Two algorithms derived from the WFS (*Wireless Fair Service*) are defined. Performances are characterized and compared to existing algorithms in OFDMA. One of them, the OWFS (*Opportunist Wireless Fair Service*) algorithm is shown to maximize the data rate of non real time flows. Besides, thanks to the delay weight parameter, it is also possible to maintain a satisfying packet drop rate for real time flows.

Keywords : OFMA, RA, bandwidth allocation, subcarrier allocation, frequency reuse factor, scheduling, real time traffic, non real time traffic, proportional fairness, WFS.

Table des matières

Remerciements.....	5
Résumé.....	7
Abstract.....	8
Index des tables.....	13
Index des illustrations.....	14
Chapitre 1.Introduction générale.....	17
1. L'OFDMA, une technique d'accès multiple avantageuse.....	17
1.1.Des services exigeants et un canal radio contraignant.....	17
1.2.Systèmes multiporteuses et accès multiple.....	17
2. L'OFDMA et ses enjeux.....	19
2.1.Diversité multiutilisateur et diversité fréquentielle.....	19
2.2.L'ordonnancement.....	20
3. Problématiques et spécificités de la thèse.....	20
3.1.Problématiques.....	20
3.2.Démarche adoptée.....	21
3.3.Hypothèses de travail.....	21
4. Plan de la thèse (résumé par chapitre).....	22
Partie I : Allocation de ressources radio en OFDMA (sous porteuses et puissance).....	25
Acronymes (partie 1).....	26
Notations (partie 1).....	27
Chapitre 2.Synthèse des problèmes d'allocation de ressources en OFDMA.	29
1. Introduction.....	29
2. Problèmes d'optimisation avec contraintes en OFDMA.....	30
2.1.Minimiser la puissance émise (problème MA).....	30
2.2.Maximiser le débit (problème RA).....	30
2.3.Optimiser « l'équité »	30
2.4.Minimiser « l'outage ».....	31
2.5.Positionnement des principales références.....	31
3. Classification des algorithmes proposés.....	32
3.1.Enjeux.....	32
3.2.Deux approches : points communs et différences.....	32
3.3.Approche A et terminologie.....	33
4. Optimalités dans le contexte monocellulaire.....	34
4.1.Optimalité globale.....	34
4.2.Optimalité des étapes.....	34
4.3.Synthèse sur les cas d'optimalité.....	36
5. Principaux algorithmes dans le contexte monocellulaire.....	36

5.1. Attribution de sous porteuses dans l'approche A.....	37
5.2. Allocation de puissance dans l'approche A.....	42
5.3. Allocation conjointe de sous porteuses et de MCS dans l'approche B.....	43
5.4. Amélioration de la puissance transmise dans l'approche B.....	44
6. Conclusions.....	45

Chapitre 3. Allocation de ressources en OFDMA : contribution dans le contexte monocellulaire..... 47

1. Introduction.....	47
2. BARE et RPO : deux nouveaux algorithmes en RA.....	47
2.1. Motivations et contexte de notre contribution.....	47
2.2. Le BARE : un algorithme pour l'allocation de largeur de bande.....	48
2.3. Le RPO : un algorithme pour l'affectation des sous porteuses.....	49
3. Méta-Problème RA.....	53
3.1. Réduction proportionnelle.....	53
3.2. Réduction par décrétement.....	53
4. Simulations et résultats.....	54
4.1. Objectifs.....	54
4.2. Conditions de simulation.....	54
4.3. Débit par sous porteuse.....	55
4.4. Facteur d'équité.....	58
4.5. Probabilité d'échec (outage).....	60
4.6. Évaluation de la complexité.....	62
5. Conclusions.....	63

Chapitre 4. Modes de sous canalisation en WiMAX (IEEE 802.16)..... 65

1. Introduction.....	65
1.1. Algorithmes d'affectation de sous porteuses et conditions pratiques.....	65
1.2. Objectifs du chapitre.....	65
2. Le standard IEEE 802.16.....	66
2.1. Généralités.....	66
2.2. Quelques définitions.....	67
2.3. Structure de la trame TDD.....	67
3. Modes de sous canalisation en IEEE 802.16.....	69
3.1. Mode sélectif en fréquence : l'AMC.....	69
3.2. Modes de diversité : description et caractérisation des collisions.....	70
4. Performance des modes de sous canalisation et des algorithmes d'affectation des sous porteuses.....	79
4.1. Méthodes de sous canalisation comparées.....	79
4.2. Conditions de simulation.....	79
4.3. Résultats de simulation.....	82
5. Conclusions.....	86
5.1. La sous canalisation en WiMax.....	86
5.2. Résultats.....	87

Chapitre 5. Allocation de ressources en OFDMA : contribution dans le contexte multicellulaire.....	89
1. Introduction.....	89
2. Travaux existants.....	90
2.1. Aperçu des travaux par type d'optimisation.....	90
2.2. Travaux d'importance en RA.....	92
3. Un nouvel algorithme en RA : l'OSA-IL.....	95
3.1. Formulation du problème.....	95
3.2. Motivations.....	95
3.3. Objectifs et «philosophie» de l'algorithme.....	96
3.4. Description de l'algorithme.....	96
3.5. Illustration sur un cas d'école.....	98
4. Simulations et résultats.....	99
4.1. Algorithmes comparés.....	99
4.2. Modèle canal et Paramètres de simulation.....	100
4.3. Débit moyen d'une cellule.....	101
4.4. Probabilité d'échec (outage).....	103
4.5. Évaluation de la complexité.....	104
5. Conclusions.....	105
5.1. Résultats.....	105
5.2. Améliorations.....	106
5.3. Évolutions du contexte d'étude.....	106
Partie II: Ordonnancement en OFDMA.....	107
Acronymes (partie 2).....	108
Notations (partie 2).....	109
Chapitre 6. Contribution pour l'ordonnancement en OFDMA.....	111
1. Introduction.....	111
1.1. Présentation des objectifs.....	111
1.2. État de l'art : ordonnancement dans les réseaux avec et sans fils.....	112
1.3. État de l'art : ordonnancement en OFDMA.....	116
2. Le WFS et ses évolutions.....	120
2.1. Algorithme initial : le Wireless Fair Service	120
2.2. Notre contribution : évolutions du WFS pour l'OFDMA.....	123
2.3. Propriétés.....	124
2.4. Caractérisation du WFS, EWFS et OWFS.....	129
3. Simulations : ordonnancer deux classes de service en OFDMA.....	138
3.1. Modèle de canal et paramètres de simulation.....	139
3.2. Comparaison entre EWFS et OWFS.....	139
3.3. Comparaison avec les algorithmes existants en OFDMA.....	142
4. Conclusions.....	144
4.1. Travail effectué et principaux résultats.....	144
4.2. Limitations.....	145
4.3. Améliorations éventuelles.....	146

Chapitre 7. Conclusion générale.....	147
1. Travaux réalisés.....	147
2. Travaux futurs.....	148
Annexes.....	151
Annexe 1.A Le canal de transmission.....	152
Annexe 1.B L'OFDM.....	154
Annexe 1.C Le waterfilling.....	155
Annexe 2.A Puissance d'émission et SNR gap.....	158
Annexe 2.B Relation entre le débit et le SNR.....	159
Annexe 2.C Optimisation MA, RA et astuces	160
Annexe 2.D L'algorithme Hongrois.....	161
Annexe 4.A Règles de permutation en FUSC.....	163
Annexe 4.B Règles de permutation en PUSC et collisions	165
Annexe 5.A Détails sur deux contributions (contexte multicellulaire).....	168
Annexe 5.B Précisions sur l'OSA-IL (contexte multicellulaire).....	170
Annexe 6.A L'ordonnancement dans les réseaux filaires	172
Annexe 6.B L'ordonnancement dans les réseaux sans fils (canal ON-OFF).....	174
Annexe 6.C L'ordonnancement dans les réseaux sans fils (canal multidébit).....	176
Annexe 6.D L'ordonnancement en OFDMA.....	179
Annexe 6.E Preuves sur la région schedulable en EWFS et OWFS.....	181
Annexe 6.F Preuves sur les garanties de service en EWFS et OWFS.....	184
Annexe 6.G Preuves sur les garanties d'équité en EWFS et en OWFS.....	188
Annexe 6.H Preuves sur les garanties de délai en EWFS et en OWFS.....	190
Références	191

Index des tables

Tableau 2.1 : Problème d'optimisation dans les références monocellulaires.....	31
Tableau 2.2 : Problème d'optimisation dans quelques références multicellulaires.....	32
Tableau 2.3 : Résumé sur l'optimalité des algorithmes.....	36
Tableau 2.4 : Initialisation et conditions d'arrêt des algorithmes d'allocation de bande.....	39
Tableau 2.5 : Formules pour l'allocation de bande.....	40
Tableau 2.6 : Approches des références dans le contexte monocellulaire.....	46
Tableau 2.7 : Approches de quelques références dans le contexte multicellulaire.....	46
Tableau 3.1 : Paramètres de simulation, contexte monocellulaire.....	55
Tableau 4.1 : Nombre de sous canaux en FUSC en fonction de la bande totale.....	70
Tableau 4.2 : Étude de la probabilité de collision (FUSC, sous canalisation aléatoire).....	75
Tableau 4.3 : Formation des groupes majeurs en PUSC, [IEEE 802.16e].....	76
Tableau 4.4 : Constitution d'un slot en fonction des modes.....	78
Tableau 4.5 : MCS et seuils de SINR ([IEEE 802.16]).....	81
Tableau 4.6 : Paramètres de simulation, impact de la sous canalisation.....	81
Tableau 4.7 : Débit global par modes, $d = 500$ m, $dBS-BS = 2.5$ Km, $\alpha = 3.5$, $PBS = 43$ dBm.....	82
Tableau 4.8 : Facteur de charge et débit, $d = 500$ m, $dBS-BS = 2.5$ km, $\alpha = 3.5$, $PBS = 43$ dBm.....	85
Tableau 5.1 : Paramètres de simulation, allocation multicellulaire.....	101
Tableau 6.1 : Paramètres de modélisation du canal.....	130
Tableau 6.2 : Découplage débit/délai, WFS et extensions pour $(M = 1, S = 1)$	131
Tableau 6.3 : Découplage débit/délai, EWFS pour $(M = 5, S = 2)$	131
Tableau 6.4 : Découplage débit/délai, OWFS pour $(M = 5, S = 2)$	131
Tableau 6.5 : Comparaison entre délai moyen et délai maximum en OWFS.....	134
Tableau 6.6 : Paramètres de simulation en ordonnancement	139
Tableau 6.7 : Performance des algorithmes d'ordonnancement pour deux classes de service.....	144

Index des illustrations

Figure 1.1: Réponse du canal en fonction de la fréquence.....	19
Figure 2.1 : Structures des algorithmes d'allocation des ressources.....	33
Figure 2.2 : Détail sur l'attribution de sous porteuses.....	34
Figure 2.3: Répartition de puissance dans le waterfilling.....	35
Figure 3.1 : Comparaison des méthodes d'allocation de bande.....	56
Figure 3.2 : Comparaison de débit entre les algorithmes d'affectation de sous porteuses.....	57
Figure 3.3 : Comparaison de débit entre algorithmes d'affectation de sous porteuses, zoom et intervalles de confiance.....	58
Figure 3.4 : Comparaison d'un facteur d'équité entre algorithmes d'affectation de sous porteuses...	59
Figure 3.5 : Comparaison d'un facteur d'équité entre algorithmes d'affectation de sous porteuses, zoom et intervalles de confiance.....	59
Figure 3.6 : Débit minimal commun (après réduction par décrétement) en fonction de la puissance..	60
Figure 3.7 : Procédé de calcul de l'outage.....	61
Figure 3.8 : Comparaison du taux d'insatisfaction des utilisateurs entre les algorithmes d'affectation de sous porteuses.....	62
Figure 3.9 : Comparaison du temps d'exécution des algorithmes d'affectation de sous porteuses....	63
Figure 4.1 : Structure de la trame TDD, [IEEE 802.16].....	68
Figure 4.2 : Subdivision de la trame TDD en zones, [IEEE 802.16e].....	68
Figure 4.3 : Types de sous canaux AMC.....	70
Figure 4.4 : Illustration de la sous canalisation FUSC dans une cellule.....	71
Figure 4.5 : Illustration de la sous canalisation FUSC dans différentes cellules.....	72
Figure 4.6 : Illustration de la sous canalisation FUSC dans différentes cellules, zoom.....	72
Figure 4.7 : Densité de probabilité de collisions en FUSC, facteur de charge 1/16.....	73
Figure 4.8 : Densité de probabilité de collisions en FUSC, facteur de charge 2/16.....	74
Figure 4.9 : Position des fréquences pilotes dans un cluster en PUSC, [IEEE 802.16].....	77
Figure 4.10 : Schéma global de la construction des sous canaux en PUSC.....	77
Figure 4.11 : Position des fréquences pilotes dans un tile en PUSC, [IEEE 802.16].....	78
Figure 4.12 : Configuration cellulaire.....	80
Figure 4.13 : Débit en fonction de la distance inter-BS, $d = 500$ m, $\alpha = 3.5$, $PBS = 43$ dBm	83
Figure 4.14 : Débit en fonction de la distance inter-BS, zoom et intervalles de confiance.....	83
Figure 4.15 : Débit en fonction du coefficient d'atténuation ($d = 500$ m, $dBS-BS = 2.5$ km, $PBS = 43$ dBm).....	84
Figure 4.16 : Débit en fonction du coefficient d'atténuation, zoom et intervalles de confiance.....	84
Figure 4.17 : Comparaison des facteurs de charge en FUSC ($dBS-BS = 2.5$ Km, $\alpha = 3.5$, $PBS = 43$ dBm).....	85
Figure 4.18 : Comparaison des facteurs de charge en FUSC, zoom et intervalles de confiance.....	86
Figure 5.1 : Illustration des interférences dans le calcul du SIF de [Koutsopoulos_02].....	93
Figure 5.2: Exemple de configuration cellulaire avant l' application de l'OSA-IL.....	98
Figure 5.3 : Débit moyen d'une cellule en fonction du débit minimal.....	102
Figure 5.4 : Évolution du FRF moyen d'une sous porteuse en fonction du débit minimal requis...	103
Figure 5.5 : Comparaison de la probabilité d'échec sur tous les utilisateurs.....	103
Figure 5.6 : Comparaison de la probabilité d'échec sur les utilisateurs lointains.....	104
Figure 5.7 : Comparaison des temps d'exécution entre l'heuristique B et de l'OSA-IL.....	105
Figure 6.1 : Partage de la trame en PLFS.....	117
Figure 6.2 : Service reçu en EWFS.....	132

Figure 6.3 : Différence de service entre les flux en EWFS.....	133
Figure 6.4 : Différence de service entre les flux en OWFS.....	133
Figure 6.5 : Dégradation exponentielle de service WFS (1).....	134
Figure 6.6 : Dégradation exponentielle de service WFS (2)	135
Figure 6.7 : DES en EWFS et OWFS (1).....	136
Figure 6.8 : DES en EWFS (2).....	137
Figure 6.9 : DES en OWFS (2).....	137
Figure 6.10 : Débit de sortie par classe de service en fonction de la charge d'entrée.....	140
Figure 6.11 : Délais en EWFS et OWFS en fonction de la charge d'entrée.....	141
Figure 6.12 : Taux de pertes des paquets RT (EWFS, OWFS) en fonction de la charge d'entrée...	141
Figure 6.13 : Taux de pertes RT en fonction du rapport entre poids de délai RT et NRT.....	142

Chapitre 1. Introduction générale

1. L'OFDMA, une technique d'accès multiple avantageuse

1.1. Des services exigeants et un canal radio contraignant

De nos jours, les services multimédia et l'accès à Internet avec ou sans fils connaissent un essor croissant. Les réseaux sans fils doivent s'adapter pour supporter des débits toujours plus élevés. Or le canal de transmission radio, constitué de trajets multiples, est très contraignant. Un débit élevé en modulation monoporteuse provoque une durée symbole faible. L'étalement en temps du canal provoque des problèmes d'interférences entre symboles. De plus, lorsque le canal est sélectif en fréquence, le signal subit des atténuations qui varient avec la fréquence (cf. annexe 1.A). Pour contourner ces difficultés, les modulations multiporteuses ont été introduites. Il s'agit de techniques de multiplexage en fréquence qui modulent le signal sur un grand nombre de porteuses à la fois. Ces techniques sont intéressantes pour des transmissions haut débit sur un canal multitrajet et sélectif en fréquence. Sur chaque sous porteuse, la transmission est ainsi bas débit. L'espace inter-fréquence est inférieure à la bande de cohérence du canal.

L'OFDM (*Orthogonal Frequency Division Multiplexing*) est une modulation multiporteuse très utilisée, on peut citer les standards DAB (*Digital Audio Broadcasting*), DVB-T (*Digital Video Broadcasting Terrestrial*) et les versions *a* puis *g* de IEEE 802.11. L'OFDM autorise un fort recouvrement spectral entre sous porteuses. Le canal large bande est transformé en un ensemble de N canaux à bande étroite avec une orthogonalité entre canaux. Chaque symbole M -QAM est transmis sur un canal, simplement caractérisé par un gain complexe sur la sous porteuse correspondante (cf. annexe 1.B). On peut tirer profit de la diversité fréquentielle en privilégiant les bonnes sous porteuses. Le principe du *waterfilling* (ou *power loading*) est alors utilisé : les sous porteuses qui ont un gain trop faible ne reçoivent pas de puissance (cf. annexe 1.C).

1.2. Systèmes multiporteuses et accès multiple

Les techniques d'accès multiple permettent aux utilisateurs de partager le médium de transmission. Chaque utilisateur reçoit une fraction des ressources disponibles. Lorsque l'on considère des systèmes multiporteuses, les principales techniques à accès multiple sont l'OFDM-TDMA ([RohlGrund_96]), l'OFDMA ([Andrews&_07]) et le MC-CDMA ([Mourad_06]).

L'OFDM-TDMA est une technique hybride entre l'OFDM et le TDMA (*Time Division Multiple Access*). Les utilisateurs transmettent tour à tour, la modulation OFDM est appliquée sur toutes les sous porteuses.

L'OFDMA est une technique hybride entre l'OFDM, le TDMA et le FDMA (*Frequency Division Multiple Access*). Dans un même symbole OFDM, plusieurs utilisateurs reçoivent des parties distinctes de la bande fréquentielle.

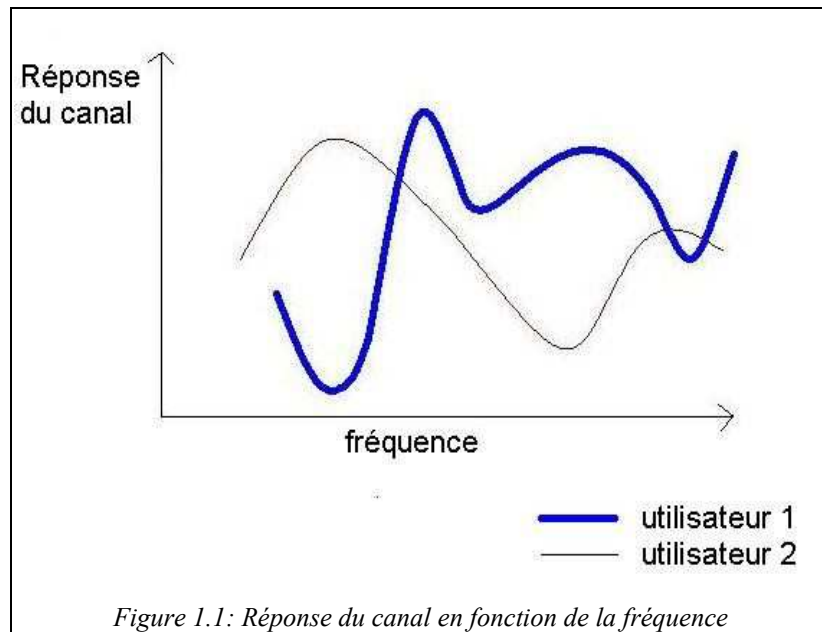
Le MC-CDMA signifie *Multi Carrier Code Division Multiple Access*. Les bits sont étalés grâce à un code (propre à chaque utilisateur) pour obtenir des chips. Les chips sont envoyés grâce à la modulation OFDM appliquée sur toutes les sous porteuses.

Les performances de ces techniques sont comparées dans [RohlGrund_97]. Les techniques OFDM-TDMA et OFDMA sont plus adaptées lorsque le canal est connu à l'émetteur tandis que le MC-CDMA est plus adapté lorsque le canal est inconnu à l'émetteur ([Ciblat_05]). Dans cette thèse, le MC-CDMA ne sera pas traité ; nous nous intéressons au cas où le canal est connu à l'émetteur.

Les avantages des techniques basées sur l'OFDM par rapport au TDMA et au FDMA sont la diversité fréquentielle et la robustesse aux multitrajets. Dans ces techniques, on peut faire de l'adaptation de lien (ou modulation adaptative) : on adapte le MCS (*Modulation and Coding Scheme*) en fonction du gain du canal ([Czylwik_96]). Le *bit loading* est un algorithme qui détermine la QAM pour minimiser la puissance transmise sur une sous porteuse. Le débit utile est élevé sur les bons canaux et le codage correcteur d'erreurs est robuste sur les mauvais canaux.

Le *waterfilling* permet de maximiser le débit d'un utilisateur avec un budget de puissance fixé (cf. annexe 1.C). En OFDM-TDMA, grâce au *waterfilling*, un utilisateur peut ignorer les sous porteuses qui ont un gain de canal très faible. Durant le *slot* d'allocation, ces sous porteuses ne sont pas utilisées du tout. Cette situation apparaît moins en OFDMA. Pour une sous porteuse fixée, plus le nombre d'utilisateurs est élevé, plus il y a de chances de trouver un utilisateur qui réalise un gain élevé. C'est la diversité multiutilisateur ([Viswanath_05], [Andrews&_07]). En OFDMA, profiter pleinement de cette diversité permet d'utiliser efficacement toutes les fréquences de la bande.

Sur la figure 1.1, la réponse du canal est représentée en fonction de la fréquence pour deux utilisateurs différents; cela permet de visualiser le profit que l'on peut tirer de la diversité multiutilisateur. Dans la suite, on utilisera la notion de CgNR pour *Channel gain to Noise Ratio* pour désigner la réponse du canal. Le CgNR de l'utilisateur u sur la sous porteuse n est noté $\phi_{u,n}$ et s'exprime comme le rapport entre le carré du module du gain du canal $|g_{u,n}|^2$ et la puissance du bruit sur l'espacement Δf entre deux sous porteuses; c'est à dire $\phi_{u,n} = |g_{u,n}|^2 / N_0 \Delta f$. Si on note p la puissance sur la sous porteuse n alors le SNR (*Signal to Noise Ratio*) est égal au produit entre la puissance p et le CgNR.



2. L'OFDMA et ses enjeux

Dans cette thèse, on étudie l'OFDMA (*Orthogonal Frequency Division Multiplexing Access*) dans le sens descendant. L'OFDMA est une candidate prometteuse pour les réseaux d'accès large bande post 3G : IEEE 802.16 ou WiMax (*Worldwide Interoperability for Microwave Access*), LTE (*Long Term Evolution*)... Cette technique supporte un grand nombre d'utilisateurs aux caractéristiques variables (QoS, débits).

Au niveau d'un utilisateur, on peut exploiter la diversité fréquentielle. Au niveau de la cellule, on peut exploiter la diversité multiutilisateur. L'allocation de ressources en OFDMA peut être orientée selon ces deux types de diversité.

2.1. Diversité multiutilisateur et diversité fréquentielle

2.1.1. Sous canalisation

En OFDMA, on attribue aux utilisateurs des fréquences distinctes. La sous canalisation consiste à regrouper les sous porteuses. La sous canalisation est par exemple utilisée en IEEE 802.16.

Un sous canal peut être constitué de sous porteuses adjacentes ou non.

Lorsque les sous porteuses sont adjacentes et que la largeur du sous canal est inférieure à la bande de cohérence du canal, les sous porteuses d'un sous canal ont un CgNR similaire. Dans la norme IEEE 802.16, on parle de sous canalisation AMC (*Adaptive Modulation and Coding*). Ce mode nécessite la connaissance de l'état du canal.

Lorsque les sous porteuses ne sont pas adjacentes, un sous canal peut être constitué de sous porteuses réparties de façon homogène sur la bande fréquentielle. La norme IEEE 802.16 appelle ce type de sous canalisation FUSC et PUSC (*Full Usage of SubChannels, Partial Usage of SubChannels*). Ce sont des modes robustes qui demandent peu de connaissances sur l'état du canal.

2.1.2. *Opportunisme grâce la diversité multiutilisateur*

Lorsqu'on profite de la diversité multiutilisateur, on affecte les sous porteuses aux différents utilisateurs de façon à maximiser le débit sur chacune d'entre elles : c'est une forme d'opportunisme. En théorie, on affecte les sous porteuses une par une. En pratique, on affecte plutôt des sous canaux composés de sous porteuses adjacentes (sous canaux AMC en IEEE 802.16). Le *waterfilling* et l'adaptation de lien peuvent être utilisés pour déterminer la puissance et la modulation sur une sous porteuse (ou un sous canal). La connaissance de l'état du canal est alors critique.

2.1.3. *Moyenner les interférences grâce à la diversité fréquentielle*

Lorsque les sous canaux sont composés de sous porteuses distribuées sur la bande fréquentielle (exemple des sous canaux FUSC, PUSC en IEEE 802.16), ils ont une qualité moyenne grâce à la diversité fréquentielle. En effet, pour un utilisateur donné, un sous canal comporte à la fois des sous porteuses de bonne et de mauvaise qualité. Dans un contexte multicellulaire, cela permet de moyenner les interférences entre les cellules ([FuSheen_07]).

2.2. L'ordonnancement

En OFDMA, l'ordonnancement de niveau MAC (*Medium Access Control*) est indissociable de la couche physique. Ces problématiques sont souvent énoncées sous l'expression *cross layer*. Désormais, en plus de l'état du trafic (considéré dans les techniques classiques), la sélection des utilisateurs doit tenir compte des conditions radio pour garantir un bon *throughput* (bonne utilisation des ressources). Les décisions doivent cependant tenir compte de l'équité. Un utilisateur pourrait être longtemps privé de transmissions si l'on ne considèrait que les conditions radio. De nouveaux mécanismes sont donc à étudier pour assurer efficacité et équité de façon conjointe (exemple de l'algorithme *proportional fair* en CDMA, [Jalali&_00]).

3. Problématiques et spécificités de la thèse

3.1. Problématiques

La première partie de la thèse est intitulée ***Allocation de ressources radio en OFDMA (sous porteuses et puissance)***. Dans cette partie, nous cherchons à maximiser le débit global d'une cellule tout en assurant un débit minimal aux utilisateurs (problème RA, *Rate Adaptive Optimization*). Pour cela, nous cherchons la meilleure instance d'affectation en sous porteuses et en puissance.

La seconde partie est intitulée ***Ordonnancement en OFDMA***. L'objectif du problème d'ordonnancement est de garantir les délais d'un trafic temps réel, maximiser le débit d'un trafic non temps réel tout en réalisant une équité proportionnelle entre les flux. Dans cette partie, il faut sélectionner les flux qui vont transmettre à un instant donné puis se poser, comme dans la première partie, le problème des ressources à leur affecter.

3.2. Démarche adoptée

Le problème d'allocation de ressources est un problème complexe, tant les variables sont nombreuses. Les utilisateurs sont soumis à des contraintes de QoS vis à vis de leurs applications. Ils peuvent être mobiles et leurs conditions radio peuvent donc changer rapidement. Le réseau doit, au sein d'une cellule, satisfaire tous les utilisateurs y compris ceux qui subissent de mauvaises conditions radio. Le réseau doit coordonner les différentes cellules, faire des choix judicieux concernant la réutilisation des ressources. Il faut choisir entre un facteur de réutilisation fréquentiel fixe ou dynamique. Le réseau doit respecter les contraintes internes des cellules et maîtriser les interférences pour que les cellules ne se gênent pas.

Pour éviter de traiter tous ces problèmes en même temps, nous adoptons une démarche progressive. Chaque hypothèse simplificatrice sert de préambule à des études plus complexes.

3.3. Hypothèses de travail

Les études se restreignent à la voie descendante.

3.3.1. *Première partie: Allocation de ressources radio en OFDMA*

L'allocation faite à un instant est indépendante des allocations aux instants précédents. Les sous porteuses sont modélisées individuellement selon un modèle classique : atténuation de parcours, effet de masque et évanouissement rapide. Nous considérons des *snapshots* du système où ces trois composantes sont tirées aléatoirement. Les performances en terme de débit sont moyennées sur plusieurs *snapshots*.

Dans cette partie, des algorithmes sont proposés pour le problème RA dans le contexte monocellulaire puis multicellulaire. Dans le problème d'optimisation, les utilisateurs doivent transmettre un débit minimal par unité de temps. Les utilisateurs ont toujours des données à transmettre et les envoient bit par bit. Ces algorithmes affectent les sous porteuses individuellement et exploitent la diversité multiutilisateur. Un utilisateur reçoit les sous porteuses sur lesquelles son SNR (*Signal Noise Ratio*) est élevé. Le gain du *waterfilling* étant significatif pour de faibles valeurs de SNR ([Martinian_04], [Viswanath_05]), nous adoptons une puissance constante par sous porteuse.

3.3.2. *Deuxième partie: Ordonnancement en OFDMA*

Des modèles de trafic paquet sont considérés : un trafic temps réel et non temps réel. Les queues des utilisateurs ne sont pas forcément pleines. En plus d'allouer les ressources radio, il faut au préalable sélectionner les utilisateurs qui vont transmettre. Nous considérons des politiques non stationnaires : chaque décision dépend de l'état des queues et des conditions radio des utilisateurs. L'état des queues dépend des décisions antérieures.

Pour simplifier le problème d'optimisation au cours de l'ordonnancement, nous considérons des sous canaux au lieu des sous porteuses individuelles. Ces sous canaux sont constitués de fréquences adjacentes et sont indépendants entre eux. L'évolution dans le temps de l'état d'un sous canal est modélisée par une chaîne de Markov ([Liu&_04]). La puissance par sous porteuse est toujours constante.

4. Plan de la thèse (résumé par chapitre)

La première partie comprend les chapitres 2 à 5.

Dans le **chapitre 2**, nous réalisons une synthèse bibliographique. Nous exprimons le problème d'allocation de ressources en terme de problème d'optimisation. Il s'agit soit de minimiser la puissance totale de transmission avec une contrainte de débit (ce qui convient aux applications à débit fixe de type voix), soit de maximiser le débit global sous une contrainte de puissance (ce qui convient aux applications à débit sporadique). Nous proposons une classification des algorithmes existants dans le contexte monocellulaire. Les sous porteuses sont affectées individuellement avec une connaissance parfaite de l'état du canal à l'émetteur. Ces algorithmes constituent une borne supérieure au regard des performances dans le contexte monocellulaire et servent de base pour étudier le contexte multicellulaire et l'ordonnancement.

Dans le **chapitre 3**, nous résolvons le problème RA de maximisation de débit avec contrainte en puissance et en débits minimaux. Le contexte est celui de la voie descendante d'une cellule radio isolée de toutes interférences extérieures. Les utilisateurs ni contraintes de taille de *buffer* ni de contraintes de temps. Comme plusieurs travaux antérieurs, nous considérons une puissance uniforme sur les sous porteuses. Nous proposons un algorithme appelé BARE (*Bandwidth Allocation based on Rate Estimation*, [LengGod&_06]) qui estime le nombre de sous porteuses nécessaire pour satisfaire le débit minimal de chaque utilisateur ; peu de propositions existent dans le contexte RA. Nous proposons le RPO (*Rate Profit Optimization*, [LengGod&_06]) pour l'affectation de sous porteuses, cet algorithme résout les conflits entre utilisateurs de façon à maximiser le débit global de la cellule. Dans ces algorithmes, les débits minimaux peuvent être quelconques. Par soucis d'équité entre les utilisateurs, nous nous intéressons à la recherche du plus grand débit minimal commun que l'on peut garantir aux utilisateurs. Cette recherche est appelée *méta problème RA*.

Dans le **chapitre 4**, nous comparons les modes de sous canalisation décrits dans la norme IEEE 802.16 avec les algorithmes présentés dans les chapitres 2 et 3 ([LengGodM_07], [LengGod&_07]). Cela permet d'évaluer l'impact de la sous canalisation sur le débit. De plus, cela permet de comparer les stratégies qui utilisent la diversité multiutilisateur et celles qui profitent de la diversité fréquentielle. La bande de cohérence correspond à la largeur d'un sous canal. Dans un contexte multicellulaire, nous évaluons la densité de probabilité de collisions des modes FUSC et PUSC ; on montre qu'elle se rapproche de celle d'une sous canalisation aléatoire.

Dans le **chapitre 5**, nous progressons en introduisons les interférences extérieures dans le problème d'optimisation RA. Après une courte synthèse des algorithmes d'allocation de sous porteuses dans le contexte multicellulaire, nous proposons l'OSA-IL (*Opportunist Subcarrier Allocation with Interference Limitation*, [LengGodM_06]). L'allocation multicellulaire est centralisée dans un équipement qui gère les stations de base (par exemple un ASN-GW en WiMax, *Access Service Network Gateway*). La puissance transmise par chaque station de base est fixe et la répartition est uniforme sur les sous porteuses. Le facteur de réutilisation fréquentiel varie entre les sous porteuses ;

il est modifié en fonction des utilisateurs qui n'ont pas encore atteint leur débit minimal. On s'intéresse au débit par cellule et à la proportion d'utilisateurs qui ne satisfont pas leur débit minimum.

La deuxième partie de la thèse comprend le chapitre 6 et les annexes 6.A à 6.H.

Dans le **chapitre 6**, nous nous intéressons aux politiques d'ordonnancement. Pour traiter l'ordonnancement en OFDMA, plusieurs stratégies existantes mêlent les techniques classiques d'ordonnancement aux algorithmes d'optimisation présentés dans le chapitre 2. Nous adoptons ici une approche différente. Partant du WFS (*Wireless Fair Service*), une politique de partage proportionnel des ressources qui supporte des types de trafic différents, nous proposons et caractérisons deux algorithmes spécifiques à l'OFDMA. On s'intéresse notamment au débit global des flux non temps réel et au taux de pertes des flux temps réel.

-

-

Partie I : Allocation de ressources radio en OFDMA (sous porteuses et puissance)

Dans cette partie, l'allocation est réalisée en fonction des conditions radio et non des décisions antérieures. Le trafic est illimité et sans contraintes de temps.

Acronymes (partie 1)

ACG	<i>Amplitude Craving Greedy algorithm</i>
AMC	<i>Adaptive Modulation and Coding</i>
BABS	<i>Bandwidth Assignment Based on SNR</i>
BARE	<i>Bandwidth Allocation based on Rate Estimation</i>
bDA	<i>basic Dynamic Assignment</i>
BER	<i>Bit Error Rate</i>
BS	<i>Base Station</i>
CgINR	<i>Channel gain Interference to Noise Ratio</i>
CgNR	<i>Channel gain to Noise Ratio</i>
CQI	<i>Channel Quality Indicator</i>
DCD	<i>Downlink Channel Descriptor</i>
DIUC	<i>Downlink Interval Usage Code</i>
DL-MAP	<i>Downlink map</i>
FCH	<i>Frame Control Header</i>
FRF	Facteur de Réutilisation Fréquentiel ou <i>Frequency Reuse Factor</i>
FUSC	<i>Full Usage of Subchannels</i>
MA	<i>Margin Adaptive optimization</i>
MCS	<i>Modulation and Coding Scheme</i>
OFDM	<i>Orthogonal Frequency Division Multiplexing</i>
OFDMA	<i>Orthogonal Frequency Division Multiple Access</i>
OSA-IL	<i>Opportunist Subcarrier Allocation with Interference Limitation</i>
PUSC	<i>Partial Usage of Subchannels</i>
QAM	<i>Quadrature Amplitude Modulation</i>
RA	<i>Rate Adaptive optimization</i>
RPO	<i>Rate Profit Optimization</i>
RRV	<i>Rate Requirement Violation ratio</i>
RTG	<i>Receive Transition Gap</i>
SINR	<i>Signal Interference to Noise Ratio</i>
SNR	<i>Signal to Noise Ratio</i>
SER	<i>Symbol Error Rate</i>
TTG	<i>Transmit Transition Gap</i>

UCD	<i>Uplink Channel Descriptor</i>
UIUC	<i>Uplink Interval Usage Code</i>
UL-MAP	<i>Uplink Map</i>

Notations (partie 1)

b	indice de station de base
$b_{u,n}$	nombre de bits par symbole sur la porteuse n de l'utilisateur u (paramètre de bit loading, $b_{u,n}$ est un entier)
B	nombre de stations de bases
$\mathbf{B}^{(n)}$	ensemble des cellules où la sous porteuse n est utilisée
$g_{u,n}$	gain de canal (complexe) sur la sous porteuse n vu par l'utilisateur u
$G_{b,u}^{(n)}$	gain de la sous porteuse n entre l'utilisateur u et la station de base b
$I_{b,u}^{(n)}$	interférences sur n pour l'utilisateur u de la station de base b
n, m	indices de sous porteuses
N	nombre total de sous porteuses (utiles)
N_{CHAN}	nombre de sous canaux (en IEEE 802.16)
N_{data}	nombre de sous porteuses de données (en IEEE 802.16)
N_{FFT}	taille de la FFT (en IEEE 802.16)
N_0	densité de puissance du bruit blanc gaussien
$N_{scPerChan}$	nombre de sous porteuses par sous canal (en IEEE 802.16)
N_u	nombre total de sous porteuses attribuées à u
$P_{T,max}$	budget de puissance total
P_u	budget de puissance alloué à l'utilisateur u
$P_{u,n}$	puissance affectée à la porteuse n de l'utilisateur u
$P_{b,u}^{(n)}$	puissance affectée à la porteuse n de l'utilisateur u
$r_{u,n}$	débit sur la porteuse n de l'utilisateur u ($r_{u,n} \in \Re$)
r_u	débit de l'utilisateur u sur Ω_u
r_u°	débit minimal de u (contrainte du problème d'optimisation)
$r_{u,b}^\circ$	débit minimal de l'utilisateur u de la cellule b
R_{max}	<i>nombre maximal de bits par symbole sur une sous porteuse</i>

u, v	indices sur l'utilisateur u
U	nombre d'utilisateurs par cellule
U_{total}	nombre total d'utilisateurs
V_b	ensemble des utilisateurs à satisfaire dans la cellule b dans OSA-IL
W	bande de fréquence totale
$\gamma_{u,n}$	rapport signal à bruit de l'utilisateur u sur la porteuse n
$\gamma^{(n)}_{b,u}$	SINR de l'utilisateur u de la station de base b sur la sous porteuse n
Δf	écart entre deux sous porteuses consécutives
η	seuil utilisé pour le RRV dans l'OSA-IL
Π_n	ensemble des utilisateurs qui utilisent la sous porteuse n
σ_{sh}	écart type de l'effet de masque (<i>shadowing</i>)
$\phi_{u,n}$	CgNR de la sous porteuse n de l'utilisateur u ($ g_{u,n} ^2 / N_0 \Delta f$)
$\phi^{(n)}_{b,u}$	CgINR de l'utilisateur u de la station de base b sur la sous porteuse n
Ω_u	ensemble des sous porteuses affectées à l'utilisateur u
Ω_{ub}	ensemble des sous porteuses affectées à l'utilisateur u de la cellule b

Chapitre 2. Synthèse des problèmes d'allocation de ressources en OFDMA

1. Introduction

Dans ce chapitre, nous présentons les problèmes d'optimisation qui peuvent se poser lorsqu'on effectue l'allocation de ressources sur la voie descendante en OFDMA. Il s'agit principalement de répartir les sous porteuses et la puissance disponible entre U utilisateurs. On considère que chaque utilisateur a toujours des données à transmettre.

Les utilisateurs sont caractérisés par le niveau de leur gain de parcours. Le bruit subi dans la bande est blanc. Selon le contexte considéré, monocellulaire ou multicellulaire, nous utilisons le CgNR (*Channel gain to Noise Ratio*) ou le CgINR (*Channel gain to Interference and Noise Ratio*) pour caractériser la qualité des sous porteuses et ce, pour chaque utilisateur. Dans le contexte monocellulaire, l'ensemble des CgNR $\{\varphi_{u,n}\}_{1 \leq u \leq U, 1 \leq n \leq N}$ est supposé parfaitement connu lors de l'allocation des ressources. Dans le contexte multicellulaire, on connaît l'ensemble des gains $\{G^{(n)}_{b,u}\}_{1 \leq b \leq B, 1 \leq u \leq U, 1 \leq n \leq N}$. En pratique cette connaissance de l'état du canal sur la voie descendante peut être obtenue de la station de base grâce à des canaux logiques de type CQICH (*Channel Quality Indicator Channel*) ; les utilisateurs y transmettent leurs mesures de la voie descendante.

Les contraintes du problème d'optimisation s'expriment pour chaque utilisateur u en terme de débit minimal à satisfaire r_u° (ou $r_{u,b}^\circ$). Pour une sous porteuse n , le SINR (*Signal to Interference and Noise Ratio*) final dépend de l'attribution de sous porteuses et de l'allocation de puissance.

L'objectif du problème d'optimisation peut être de minimiser la puissance transmise, de maximiser le débit total des utilisateurs, d'optimiser une équité (par exemple maximiser le débit minimum parmi les utilisateurs) ou de minimiser le nombre d'utilisateurs qui n'atteignent par leur débit minimal (r_u° ou $r_{u,b}^\circ$). Dans l'annexe 2.A (resp. 2.B), nous rappelons la (les) relation(s) entre la puissance (resp. le débit) à optimiser et la qualité radio du canal.

Pour résoudre ces problèmes d'optimisation, plusieurs approches ont été proposées. Dans la section 3, nous proposons une classification de ces approches. Dans la section 4, nous proposons une synthèse de ces travaux dans le contexte monocellulaire. Les principales approches du contexte monocellulaire sont importantes car elles sont souvent réutilisées dans le contexte multicellulaire (cf. chapitre 5) et en ordonnancement (cf. chapitre 6).

2. Problèmes d'optimisation avec contraintes en OFDMA

Nous utilisons la distinction MA (*Margin Adaptive optimization*) / RA (*Rate Adaptive optimization*) présente dans la littérature. Le premier type de problème cherche à minimiser la puissance transmise tandis que le second entend maximiser le débit global des utilisateurs. Dans ces deux types de problème, il y a généralement une contrainte sur les débits individuels. En l'absence de cette contrainte, la résolution analytique devient aisée ([WangNiu_05] en MA, [LiLiu_05] en RA).

2.1. Minimiser la puissance émise (problème MA)

Un objectif peut être de garantir un débit $r_{u,b}^o$ à tous les utilisateurs ($1 \leq u \leq U$, $1 \leq b \leq B$), tout en minimisant la puissance totale émise $P_T(\underline{r}) = \sum_{b=1}^B \sum_{u=1}^U P_{u,b}(r_{u,b})$.

On écrit :

$$\min_{\underline{r}} P_T(\underline{r})$$

avec $\underline{r} \geq \underline{r}^o$

où $\underline{r} = (r_{11}, \dots, r_{u,b}, \dots, r_{U,B})$ et $\underline{r}^o = (r_{11}^o, \dots, r_{u,b}^o, \dots, r_{U,B}^o)$

Ce problème est connu sous le nom de *Margin Adaptive optimization* (MA). La plupart des références qui apportent une contribution à ce problème sont citées dans les tableaux 2.1 et 2.2. Les auteurs de [WangNiu_05] considèrent plutôt une contrainte sur le débit global à transmettre.

2.2. Maximiser le débit (problème RA)

Un autre objectif peut être de maximiser le débit total des utilisateurs avec comme contraintes un budget de puissance fixé et éventuellement des débits minimaux $r_{u,b}^o$ pour chaque utilisateur u . L'expression du problème est alors :

$$\max \sum_{b=1}^B \sum_{u=1}^U r_{u,b}$$

avec $P_T^{(b)} < P_{T,max}$, $1 \leq b \leq B$ (1)

et $\underline{r} \geq \underline{r}^o$ (2)

Ce problème est connu sous le nom de *Rate Adaptive optimization* (RA). Par rapport au problème MA, le budget de puissance $P_{T,max}$ est une entrée supplémentaire du problème (en plus de $\underline{r}^o$ et $\{G_{b,u}^{(n)}\}_{1 \leq b \leq B, 1 \leq u \leq U, 1 \leq n \leq N}$). La contrainte (2) peut ne pas être considérée ([JangLee_03], [KimHan&_04], [LiLiu_05], [Koutsopoulos_02]). Le problème RA en est alors grandement simplifié. La contrainte (2) peut exprimer une proportionnalité entre les débits finaux des utilisateurs ([Wong&_04]). Quelque soit sa forme, la présence de la contrainte (2) baisse le débit global pour pouvoir garantir un débit minimal aux utilisateurs qui ont de mauvaises conditions radio.

2.3. Optimiser « l'équité »

Lorsqu'on cherche à réaliser une équité entre les utilisateurs, on peut penser à l'équité en bande qui consiste à attribuer aux utilisateurs une quantité égale de ressources fréquentielles. Etant donnée la disparité des conditions radio des utilisateurs, ce type d'équité ne garantit rien aux utilisateurs en terme de débit. Il est peut être plus

intéressant, pour les utilisateurs, de réaliser une équité en débit. Les utilisateurs peuvent ainsi espérer des débits similaires indépendamment de leurs conditions radio.

En MA, une équité en débit peut se traduire par $r_{u,b}^o = r^o$ pour tout u et b dans la contrainte sur le vecteur de débit. En RA, on peut de même fixer un débit commun $r_{u,b}^o = r^o$ pour tout u et b dans la contrainte (2). On peut ensuite chercher le plus grand débit commun réalisable : on qualifiera cette recherche de méta-problème RA.

Pour réaliser, en RA, une équité *totale* en débit, on peut maximiser le plus faible débit parmi les utilisateurs. Ce type de problème n'a été rencontré que dans le contexte monocellulaire. Cela revient à résoudre le problème suivant ([RheeCioffi_00]) :

$$\max_A (\min_u (r_{u,A}))$$

$$\text{avec } P_T < P_{T,max}$$

où A est une instance d'allocation en sous porteuses et puissance.

2.4. Minimiser « l'outage »

Pour garantir un débit minimal aux utilisateurs, un critère d'optimisation peut être de minimiser le nombre d'utilisateurs insatisfaits. Un utilisateur insatisfait est un utilisateur qui n'atteint pas le débit minimal $r_{u,b}^o$; cela constitue un *outage*. On appelle probabilité d'échec (ou d'*outage*) le rapport entre le nombre d'utilisateurs qui n'atteignent pas leur débit minimal et le nombre d'utilisateurs total. Ce calcul peut être effectué par le biais de *snapshots* de l'état du canal. On a rencontré ce type de problème dans le contexte multicellulaire ([Kwon_05], [Hamouda_06]). La probabilité d'échec peut être globale ou calculée par cellule.

2.5. Positionnement des principales références

Les tableaux 2.1 et 2.2 classent les principales références monocellulaires et multicellulaires en fonction du type de problème traité. Les contributions des références du tableau 2.1 sont décrites dans la section 5 ; celles du tableau 2.2 sont décrites dans le chapitre 6.

Tableau 2.1 : Problème d'optimisation dans les références monocellulaires

Contexte monocellulaire				
MA		RA		Équité
[Acena_05]	[Chen_04]	[Anas_04]	[ChangKuo_04]	[DoufexiArmour_05]
[ChenKron_07]	[Cho_05]	[Chee_04]	[Gross_03]	[MorettiMorelli_06]
[Inyoung_01]	[KimKwak_04]	[HuiZhou_06]	[Inyoung_01]	[RheeCioffi_00]
[Kivanc_03]	[Kivanc_00]	[JangLee_03]	[KimHan_04]	
[Pfltschinger_02]	[PietrykJan_02]	[Ko_06]	[MaoWang_06]	
[WangNiu_05]	[Wong_99]	[LiLiu_05]	[NguyenHan_06]	
[WongTsui_99]	[Yu_06]	[Wong_04]	[YinLiu_00]	
[Zhang_04]	[Zhen_03]			

Tableau 2.2 : Problème d'optimisation dans quelques références multicellulaires

Contexte multicellulaire		
MA	RA	Minimiser l'outage
[Gault&_05] , [Junqiang&_03] [PietrzykJanssen_03]	[Han_03], [Kim&_04] [Koutsopoulos_02] [Yan&_03]	[Hamouda&_06] [Kwon&_05]

3. Classification des algorithmes proposés

3.1. Enjeux

Un problème majeur des algorithmes concerne leur temps d'exécution ; ce dernier permet d'évaluer leur complexité. Sans simplifications préalables, les problèmes d'affectation conjointe (optimale) de sous porteuses et de puissance (avec contraintes) sont réputés difficiles et qualifiés de NP complets dans le contexte monocellulaire ([Chen&_04],[Anas_04]).

Un autre problème majeur concerne l'efficacité des algorithmes, c'est à dire leur écart avec la solution optimale.

La littérature aborde ces deux problèmes en tentant d'attribuer les ressources de façon pertinente et cela en faisant face, en temps réel, aux variations du canal.

Nous proposons une classification des algorithmes heuristiques qui exploitent la diversité multiutilisateur dans la littérature. Cette classification est commune aux différents types d'optimisation. Nous avons répertorié deux façons d'aborder le problème général d'optimisation.

3.2. Deux approches : points communs et différences

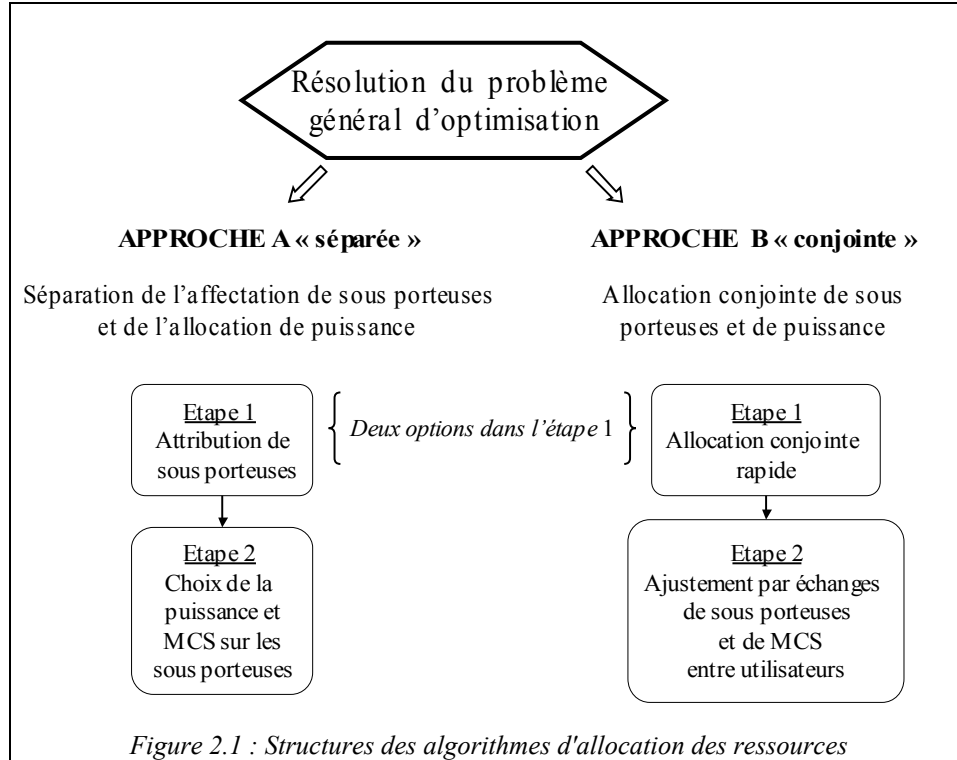
Les deux approches que nous avons rencontrées sont résumées dans la Fig. 2.1.

L'*approche A* considère séparément l'attribution des sous porteuses et l'allocation de la puissance (et/ou des MCS). Deux étapes en découlent naturellement : la première étape réalise l'attribution des sous porteuses et la deuxième étape effectue l'allocation de la puissance (et/ou des MCS).

L'*approche B* réalise de façon conjointe l'attribution des sous porteuses et l'allocation de la puissance (et/ou des MCS). Cependant, on y distingue aussi deux étapes. La première étape est une allocation conjointe rapide et la deuxième étape réalise des échanges entre utilisateurs pour améliorer un critère qui dépend du type d'optimisation.

La plupart des algorithmes (tous types d'optimisation confondus) que nous avons rencontrés sont de type *approche A*. Tous les algorithmes de type *approche B* que nous avons rencontrés résolvent un problème MA.

La différence entre les deux approches est que dans l'*approche B*, l'étape 2 est facultative. Rappelons qu'à cause du contexte cellulaire, cette étape n'est effectuée que si le temps de cohérence du canal est suffisamment grand. Dans l'*approche A*, l'étape 2 ne peut être ignorée.



Dans les deux approches, l'étape 1 a pour point commun d'attribuer les sous porteuses. De plus, l'étape 1 comporte deux options dans les approches A et B. Dans l'option 1, il n'y a pas de connaissance a priori du nombre de sous porteuses par utilisateur. Dans l'option 2, une première tâche doit évaluer le nombre de sous porteuses par utilisateur avant de poursuivre. Nous appelons **allocation de bande** le calcul du nombre de sous porteuses par utilisateur.

3.3. Approche A et terminologie

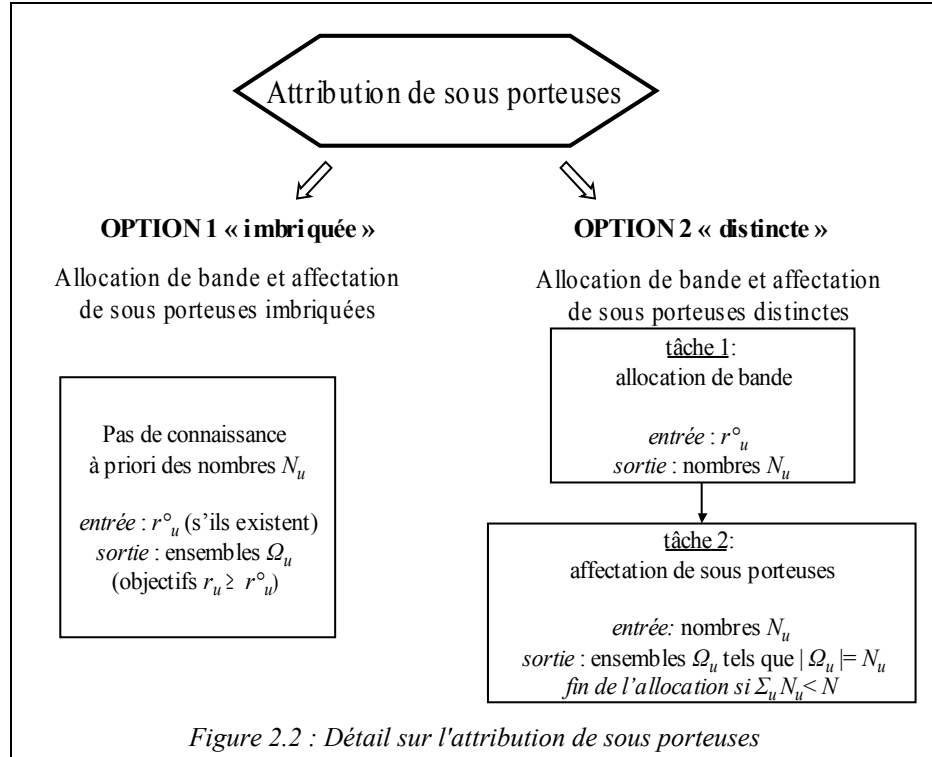
On appelle **attribution de sous porteuses** l'ensemble de l'étape 1 qui est constituée de **l'allocation de bande** et de **l'affectation de sous porteuses** ; que ces deux tâches soient imbriquées (option 1) ou séparées (option 2). Cette terminologie est utile pour la première partie de la thèse.

La figure 2.2 détaille l'attribution de sous porteuses (étape 1) dans le cas de l'approche A qui est l'approche la plus répandue.

Lorsque l'allocation de bande est séparée de l'affectation de sous porteuses, on note N_u le nombre de sous porteuses nécessaires pour chaque utilisateur u ; ce nombre dépend du débit minimal r_u^o à satisfaire et des CgNR. On parlera d'**estimation souple** lorsque

$$\sum_{u=1}^N N_u \leq N \text{ et d'estimation rigide lorsque } \sum_{u=1}^N N_u = N .$$

Dans l'affectation de sous porteuses, on note Ω_u l'ensemble des sous porteuses affectées à l'utilisateur u . A la fin de l'affectation de sous porteuses, les ensembles Ω_u vérifient $|\Omega_u| = N_u$. En cas d'estimation souple, l'affectation continue (au delà de $|\Omega_u| = N_u$ pour tout u) jusqu'à l'utilisation de toutes les sous porteuses.



4. Optimalités dans le contexte monocellulaire

4.1. Optimalité globale

Dans le cas du problème RA sans contraintes sur les débits individuels, la solution optimale a été établie en séparant attribution de sous porteuses et allocation de puissance (approche A). Dans [JangLee_03] et [LiLiu_05], il est montré que chaque sous porteuse doit être attribuée à son meilleur utilisateur et que l'allocation de puissance sur les sous porteuses d'un utilisateur doit être réalisée par le *waterfilling* (cf. section 4.2).

Dans le cas des problèmes MA et RA avec contraintes sur les débits individuels, le problème peut être exprimé sous une forme linéaire ([Inyoung&_01], [MaoWang_06]) et résolu de façon optimale par la programmation entière (résolution conjointe : approche B). Une autre technique consiste à rendre le problème convexe en relâchant la contrainte sur le partage des sous porteuses dans le temps ([Wong&_99]). L'annexe 2.C revient brièvement sur ces deux techniques. La complexité de ce type d'approches a conduit les auteurs, dans la littérature, à proposer des heuristiques (cf. section 5).

4.2. Optimalité des étapes

Dans le cas des problèmes MA et RA avec contraintes sur les débits individuels, la résolution optimale n'est pas évidente. La plupart des auteurs sépare l'attribution de sous porteuses et l'allocation de puissance (approche A).

Lorsque l'allocation de bande et l'affectation de sous porteuses sont séparées (cf. Fig. 2.2 : option 2), l'affectation de sous porteuses peut être résolue de façon optimale par l'algorithme Hongrois. Cet algorithme décrit dans l'annexe 2.D considère l'allocation de N ressources à N utilisateurs; chaque allocation est associée à un coût. Grâce à une écriture matricielle, l'algorithme parvient à minimiser le coût global. Concernant l'affectation de sous porteuses, l'optimalité de l'algorithme Hongrois est relative l'ensemble $\{N_u\}_{1 \leq u \leq U}$ calculé pendant l'allocation de bande.

Une fois l'affectation de sous porteuses réalisée (ensembles $\{\Omega_u\}_{1 \leq u \leq U}$), la solution optimale de l'allocation de puissance (étape 2) est le soit *waterfilling*, soit le *bit loading*; cela dépend du type d'optimisation considéré. Nous rappelons ci après le principe de ces algorithmes.

Principe de l'algorithme de bit loading en MA ([Campello_98])

Chaque utilisateur u est traité séparément. Sur chaque sous porteuse n de Ω_u , on calcule l'augmentation de puissance nécessaire pour transmettre le MCS immédiatement supérieur. Le nombre de bits $b_{u,n}$ est incrémenté sur la sous porteuse qui minimise l'augmentation de puissance transmise. Le procédé s'arrête lorsque la somme des bits $b_{u,n}$ des sous porteuses n de Ω_u atteint le débit cible r_u^o .

Principe du waterfilling en RA ([CoverThomas_91])

Le *waterfilling* très utilisé en théorie de l'information a été redécouvert pour des applications pratiques. Il consiste à maximiser le débit en répartissant la puissance en fonction du CgNR du canal et d'un «niveau d'eau» qui dépend de la contrainte de puissance globale. Il est montré (cf. annexe 1.C) que la puissance sur n vaut $p_n(v) = (v - \varphi_n^{-1})^+$ où v est le niveau d'eau, φ_n est le CgNR de la sous porteuse n et $x^+ = 0$ si $x < 0$. Finalement, plus le CgNR d'une sous porteuse est élevé (φ_n^{-1} est faible), plus elle recevra de puissance (cf. Fig.2.3). Dans le cadre multiutilisateur, le *waterfilling* peut être appliqué pour chaque utilisateur u sur l'ensemble Ω_u qui lui est attribué.

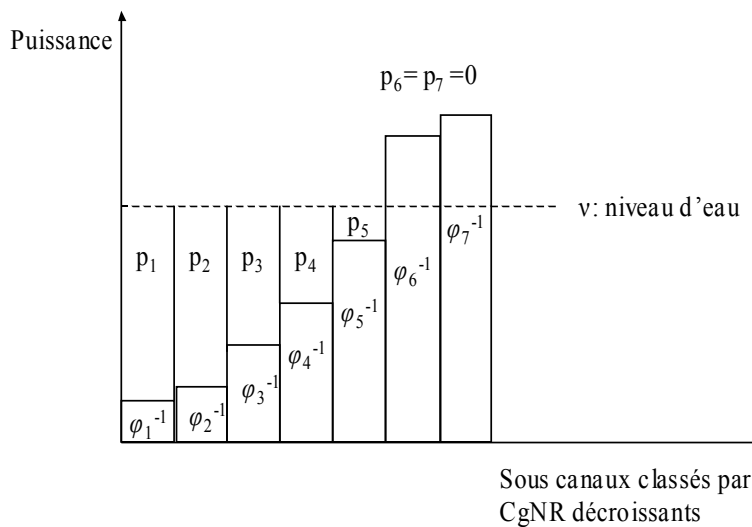


Figure 2.3: Répartition de puissance dans le waterfilling

4.3. Synthèse sur les cas d'optimalité

L'allocation de ressources dans le cas général est réputé NP-complet. Différentes répartitions $\{N_u\}_{1 \leq u \leq U}$ engendrent des ensembles Ω_u différents. Ces derniers créent des allocations de puissance différentes. On voit que le problème est très sensible à l'allocation de bande. Lorsque le canal est plat, un utilisateur voit toutes les sous porteuses avec le même CgNR. L'attribution de sous porteuses se réduit à l'allocation de bande. On peut dans ce cas trouver une solution optimale à l'allocation de bande, c'est à dire la répartition $\{N_u\}_{1 \leq u \leq U}$ qui maximise le débit (ou minimise la puissance). Lorsque le canal est sélectif en fréquence, on ne peut qu'évaluer la répartition $\{N_u\}_{1 \leq u \leq U}$.

Tableau 2.3 : Résumé sur l'optimalité des algorithmes

		Problème RA sans débits minimaux	Problème MA ou RA avec débits minimaux
Approche A	Étape 1	affecter chaque sous porteuse au meilleur utilisateur	Tâche 1 : pas de solution optimale Tâche 2 : algorithme hongrois
	Étape 2	<i>waterfilling</i>	<i>waterfilling</i> ou <i>bit loading</i>
	Optimalité globale (étape 1 + étape 2)	oui	non
Approche B		programmation entière	

5. Principaux algorithmes dans le contexte monocellulaire

Nous présentons ici les principales heuristiques que nous avons rencontrées, tous types d'optimisation confondus, en fonction de la classification proposée dans la section 3. Bien que notre contribution (cf. chapitre 3) ne porte que sur un problème RA, la description des mécanismes MA est utile pour avoir une compréhension globale de l'allocation de ressources en OFDMA. Dans tous les algorithmes la station de base utilise la connaissance des CgNR pour l'affectation des ressources.

Simplification fréquente de l'allocation de puissance en RA

Dans les chapitres 3 et 5, nous nous intéressons particulièrement au problème RA avec débits minimaux. En RA, les auteurs de [JangLee_03] montrent que le fait d'appliquer une puissance uniforme sur les sous porteuses ne dégrade pas beaucoup les performances par rapport à la solution optimale du *waterfilling*. Plusieurs travaux ont utilisé ce résultat (entre autres [RheeCioffi_00], [Gross&_03], [HuiZhou_06]). Cela simplifie non seulement l'étape d'allocation de puissance mais aussi l'étape d'attribution de sous porteuses en fournissant une connaissance des SNR. Dans notre contribution ([LengGod&_06], cf. chap. 3), nous utilisons aussi une allocation de puissance uniforme sur les sous porteuses. Nous nous intéressons principalement à l'attribution des sous porteuses, c'est pourquoi nous ne décrivons pas les contributions qui portent exclusivement sur l'allocation de puissance.

5.1. Attribution de sous porteuses dans l'approche A

Dans l'approche A (« séparée »), l'étape d'attribution de sous porteuses (étape 1) est suivie d'une étape d'allocation de puissance (étape 2) que l'on étudie dans la section 5.2.

5.1.1. Option 1 « imbriquée »

Dans l'étape d'attribution de sous porteuses, l'option 1 consiste à ne pas déterminer à priori le nombre de sous porteuses à affecter à un utilisateur (cf. Fig 2.2). On rencontre souvent ce cas de figure en l'absence de contraintes sur le débit minimal des utilisateurs.

Problème opportuniste

Les auteurs de [JangLee_03] et [KimHan&_04] tentent de maximiser le débit total sans imposer de contrainte de débit minimal. Il s'agit du problème opportuniste. La solution optimale consiste, comme on l'a vu, à attribuer chaque sous porteuse à celui qui y présente le meilleur CgNR. Le nombre de sous porteuses par utilisateur n'est pas calculé à priori et un utilisateur peut ne rien recevoir. L'opposé du problème opportuniste est le problème d'équité en débit.

Problème d'équité en débit

Dans [RheeCioffi_00], les auteurs tentent de maximiser le plus faible débit parmi les utilisateurs. Ce problème est appelé problème d'équité en débit ou *problème «max-min»*. Les auteurs proposent un algorithme heuristique, initialisé en affectant à chaque utilisateur la sous porteuse sur laquelle il a le meilleur CgNR (l'ordre des utilisateurs est arbitraire). La répartition de puissance est uniforme sur les sous porteuses. Les débits r_u sont calculés en utilisant la formule de Shannon. Puis, tant qu'il reste des sous porteuses, l'utilisateur de plus faible débit r_u reçoit sa meilleure sous porteuse (dans l'ensemble H des sous porteuses disponibles) ; les débits r_u sont ensuite actualisés. Un algorithme similaire est rencontré dans [DoufexiArmour_05] , la métrique de l'utilisateur u n'est pas le débit r_u mais la somme des $|g_{u,n}|^2$ sur Ω_u . Dans [Anas_04], l'auteur propose une variante de l'algorithme de [RheeCioffi_00] pour traiter deux classes d'utilisateurs¹.

Problème RA avec débits minimaux

Ci après, nous décrivons les travaux [Ko&_06] et [NguyenHan_06]. Contrairement à notre contribution ([LengGod&_06], publiée à la même conférence), l'allocation de bande et l'affectation de sous porteuses y sont imbriquées.

¹ Dans [Anas_04], l'auteur considère une classe GP (*Guaranteed Performance*) et une classe BE (*Best Effort*). Les utilisateurs GP ont un débit minimal à satisfaire tandis que les utilisateurs BE n'en ont pas. L'algorithme de [RheeCioffi_00] est utilisé jusqu'à ce que chaque utilisateur GP vérifie sa contrainte de débit. Le nombre de sous porteuses restantes est affecté aux utilisateurs de la classe BE avec le même algorithme.

Dans [Ko&_06], l'utilisateur v qui reçoit une sous porteuse supplémentaire est le plus éloigné de son débit minimal ($v = \arg \min_{1 \leq u \leq U} \left(\sum_{n \in \Omega_u} r_{u,n} \right) / r_u^\circ$ où $r_{u,n}$ est le débit réalisé sur la sous porteuse n de Ω_u et r_u° est le débit minimal à atteindre). Puis la meilleure sous porteuse m de l'utilisateur v (parmi l'ensemble H des sous porteuses disponibles) lui est affectée. Si plusieurs sous porteuses permettent à l'utilisateur v d'atteindre son débit le plus élevé, on note ce sous ensemble H_v . Dans ce cas, la BS affecte à l'utilisateur v la sous porteuse qui vérifie $\min_{n \in H_v, u, u \neq v} r_{u,n}$. L'algorithme s'arrête quand l'ensemble H des sous porteuses disponibles est vide.

L'auteur de [NguyenHan_06], dans une première phase, détermine l'ordre dans lequel les utilisateurs vont recevoir des sous porteuses grâce à leur débit moyen (sur une fenêtre temporelle). Dans une deuxième phase, chaque utilisateur u reçoit des sous porteuses jusqu'à satisfaction de son débit minimal. La sous porteuse m qui minimise $|r_u + r_{u,n} - r_u^\circ|$ (parmi l'ensemble H des sous porteuses disponibles) est affectée à l'utilisateur dont c'est le tour. Dans une dernière phase les sous porteuses restantes sont affectées, chaque utilisateur (toujours par ordre de priorité) reçoit sa meilleure sous porteuse.

5.1.2. Option 2 « distincte »

Pour réaliser l'attribution de sous porteuses (étape 1), on effectue l'allocation de bande (tâche 1 : déterminer le nombre de sous porteuses à affecter à un utilisateur), puis on affecte les sous porteuses aux utilisateurs (tâche 2). Nous n'avons rencontré ce cas de figure que lorsqu'il y a un débit minimal r_u° à atteindre pour chaque utilisateur u .

a) Tâche 1 : allocation de bande

Lorsque le MCS est unique, il y a une relation directe entre le débit minimal et le nombre de sous porteuses nécessaires pour réaliser un débit minimum ([WongTsui&_99], [PietrykJan_02]). La tâche 1 est alors simple et le résultat obtenu est exact.

En modulation adaptative (possibilité de plusieurs MCS), le lien est moins direct et des solutions plus ou moins élaborées sont proposées pour évaluer le nombre de sous porteuses par utilisateur. On parle d'*estimation rigide* lorsque la répartition $\{N_u\}_{1 \leq u \leq U}$ obtenue est utilisée sans modifications dans la tâche 2 $\left(\sum_{u=1}^N N_u = N \right)$. C'est le cas

dans [Kivanc&_00], [Kivanc&_03] et [YinLiu_00]. D'autres auteurs utilisent une formule que l'on qualifiera d'*estimation souple* comme [Wong&_04] et [HuiZhou_06].

Dans ce cas $\left(\sum_{u=1}^N N_u \leq N \right)$, les nombres $\{N_u\}_{1 \leq u \leq U}$ peuvent augmenter dans la tâche 2.

Estimation rigide en modulation adaptative

Le BABS (*Bandwidth Assignment Based on SNR*, [Kivanc&_00], [Kivanc&_03]), proposé en MA) utilise le vecteur des débits minimaux $\underline{r}^o = (r^o_1, \dots, r^o_u, \dots, r^o_U)$ et le vecteur des CgNR moyens $\underline{\varphi} = (\varphi_1, \dots, \varphi_u, \dots, \varphi_U)$. L'hypothèse de travail est un canal plat : un utilisateur u voit toutes les sous porteuses avec le même CgNR. Il s'agit alors du CgNR moyen φ_u : $\varphi_u = \frac{1}{N} \sum_{n=1}^N \varphi_{u,n}$. Le nombre de sous porteuses d'un utilisateur est déterminé selon la routine ci-après.

Routine de réduction de puissance

Soit $i(u)$ le nombre de sous porteuses attribué à l'utilisateur u à un stade donné. La boucle décrite ci-dessous porte sur le vecteur $\underline{i}$. Les actions ci-après sont effectuées de façon itérative :

- pour chaque u , connaissant φ_u , on calcule $P(i(u))$, la puissance¹ nécessaire pour transmettre r^o_u bits avec $i(u)$ sous porteuses,
- de même, on calcule la puissance ($P(i(u)+1)$) nécessaire pour transmettre r^o_u bits si l'utilisateur avait $i(u)+1$ sous porteuses,
- on incrémente $i(u)$ pour l'utilisateur qui vérifie le plus grand $\Delta(u) = P(i(u)) - P(i(u)+1)$.

A chaque étape, on augmente donc le nombre de sous porteuses de l'utilisateur u qui réalise la plus grande économie de puissance.

Le principe de réduction de puissance est aussi utilisé en RA. Le tableau 2.4 présente les conditions d'initialisation et d'arrêt des contributions utilisant le principe de réduction de puissance.

Tableau 2.4 : Initialisation et conditions d'arrêt des algorithmes d'allocation de bande

Publication	Problème d'optimisation	Contrainte	Initialisation	Condition d'arrêt
[Kivanc&_00] [Kivanc&_03]	MA	$r \geq r^o$	$i(u) = \lceil r^o_u / R_{max} \rceil$	$\sum_u i(u) = N$
[YinLiu_00]	RA	$r \geq r^o$	$i(u) = 1$	$\frac{P'_T}{N'_a} \leq \frac{P_{T,max}}{N}$

Dans le BABS, R_{max} est le nombre de bits utiles maximal dans un MCS. Si à l'initialisation, le nombre total de sous porteuses dépasse N , les utilisateurs de plus faible débit minimal sont éliminés². La répartition $\{N_u\}_{1 \leq u \leq U}$ obtenue par le BABS est optimale lorsque le canal est plat.

- 1 Pour calculer $P(i(u))$, on considère la fonction f qui relie le SNR γ au débit r : $\gamma = f(r)$ (la fonction est différente selon le BER exigé). L'expression de la puissance nécessaire à u est: $P(i(u)) = (i(u)/\varphi_u) \times f(r^o_u / i(u))$; cf annexe 2.B pour un exemple de f . Pour f croissante $f(r^o_u / i(u))$ décroît quand $i(u)$ augmente ($\Delta(u) \geq 0$); entre deux utilisateurs u_1 et u_2 , à r^o_u et $i(u)$ identiques, u_1 reçoit la sous porteuse supplémentaire s'il vérifie $\varphi_{u1} < \varphi_{u2}$.
- 2 Au contraire, l'auteur de [Nakad_03] propose d'éliminer les utilisateurs qui ont un grand débit minimal pour servir un plus grand nombre d'utilisateurs.

Dans la condition d'arrêt de [YinLiu_00], P'_T est la somme des puissances allouées à un instant donné, N'_a est le nombre total de sous porteuses attribuées aux utilisateurs à une étape donnée. Dans [YinLiu_00], l'auteur utilise le principe de réduction de puissance alors qu'il résout un problème RA. Notre contribution [LengGod&_06] (cf. chap. 3) utilise un critère plus adapté au problème RA.

A l'arrêt des algorithmes, la valeur finale du nombre de sous porteuses de l'utilisateur u est notée N_u (à distinguer avec les valeurs temporaires $i(u)$ au cours de l'algorithme).

Estimation souple en modulation adaptative

Dans certaines contributions, les auteurs proposent une estimation souple du nombre de sous porteuses, $i(u) = \lceil r_u^\circ / R_{max} \rceil$, plutôt qu'un nombre rigide N_u déterminant le cardinal de Ω_u pendant la tâche 2. Dans [Wong&_04] et [HuiZhou_06], le problème comporte une contrainte proportionnelle (où équité pondérée) sur le débit : le vecteur $\underline{r}$ doit vérifier pour tout couple d'utilisateurs (u, v) $r_u/r_v = \beta_u/\beta_v$ où le vecteur $\underline{\beta}$ est donné en entrée. On voit dans le tableau 2.5 que ce vecteur influe directement sur l'estimation du nombre de sous porteuses.

Tableau 2.5 : Formules pour l'allocation de bande

Publication	Problème d'optimisation	Contrainte	Estimation souple
[KimKwak_04]	MA	$\underline{r} \geq \underline{r}^\circ$	$i(u) = \lceil r_u^\circ / R_{max} \rceil$
[Wong&_04] [HuiZhou_06]	RA	$\forall (u, v), \frac{r_u}{r_v} = \frac{\beta_u}{\beta_v}$	$i(u) = \lfloor \beta_u \cdot N \rfloor$

b) Tâche 2 : affectation de sous porteuses

La tâche 2 affecte les sous porteuses avec les contraintes $|\Omega_u| = N_u$ pour tout u (cf. Fig. 2.2). Quand il s'agit d'une estimation souple, l'affectation continue jusqu'à l'utilisation de toutes les sous porteuses disponibles.

Solution optimale

On a vu que la solution optimale est obtenue par l'algorithme Hongrois par rapport à un ensemble $\{N_u\}_{1 \leq u \leq U}$ fixé (pour cet algorithme l'estimation doit être rigide). L'algorithme Hongrois est réputé lent pour une résolution temps réel (donné en $\mathcal{O}(N^3)$ dans [PapaSteig_82]). Plusieurs heuristiques sont proposées dans la littérature pour améliorer le temps d'exécution de la phase d'affectation de sous porteuses.

Approches sous optimales orientées utilisateurs

L'auteur de [Pfltschinger&_02] traite un problème MA avec débits minimaux et estimation rigide des nombres $\{N_u\}_{1 \leq u \leq U}$. Les auteurs introduisent des variables *accessoires* nommées priorités. L'utilisateur qui a la plus grande priorité reçoit sa meilleure sous porteuse. Les priorités des utilisateurs sont mises à jour à chaque fois

qu'un utilisateur reçoit une sous porteuse. La priorité initiale de chaque utilisateur est N_u/N . A chaque étape, la priorité de l'utilisateur u est la différence entre $N'_u/(\sum_u N'_u)$ et sa priorité initiale, où N'_u est le nombre de sous porteuses restant à lui attribuer.

Les auteurs de [Gross&_03] traite un problème RA avec débits minimaux et estimation rigide des $\{N_u\}_{1 \leq u \leq U}$. Les auteurs attribuent la totalité des N_u meilleures sous porteuses de chaque utilisateur u quand celui-ci est pris en compte dans le processus d'affectation. L'algorithme est appelé bDA (*basic Dynamic Assignment*). La puissance est uniforme sur les sous porteuses. Pour assurer une sorte d'équité entre les utilisateurs, les auteurs considèrent des cycles de plusieurs affectations. Tous les S symboles OFDM, chaque utilisateur voit sa priorité s'affaiblir excepté la priorité la plus faible qui devient la priorité la plus forte.

Les deux approches précédentes sont orientées utilisateurs (en terme d'implémentation), les utilisateurs ordonnés selon certains critères reçoivent des sous porteuses. Ci-après nous examinons des approches orientées sous porteuses : pour une sous porteuse déterminée, un utilisateur est sélectionné.

Approches sous optimales orientées sous porteuses

Nous présentons maintenant les algorithmes ACG et ses dérivés. Le problème résolu est de type MA, avec débits minimaux et estimation rigide des $\{N_u\}_{1 \leq u \leq U}$.

Les auteurs de [Kivanc&_00] proposent l'algorithme ACG (*Amplitude Craving Greedy algorithm*). C'est un algorithme opportuniste dont le nombre de sous porteuses de chaque utilisateur est limité à N_u . A chaque étape une sous porteuse n est fixée, l'algorithme alloue la sous porteuse n à l'utilisateur qui a le meilleur CgNR et qui n'a pas encore atteint son nombre de sous porteuses N_u . L'algorithme s'arrête lorsqu'il n'y a plus de sous porteuses à attribuer. La principale limitation de cet algorithme est l'ordre de traitement des sous porteuses. Elles sont traitées par index croissant ([Kivanc&_00]) ou dans un ordre aléatoire ([Kivanc&_03]). L'ordre de traitement a pourtant un impact sur les performances.

Pour améliorer l'ordre de traitement des sous porteuses, les auteurs de [Zhen&_03] proposent un algorithme où la sous porteuse et l'utilisateur sont choisis conjointement (par la BS) de la façon suivante : $(u^*, n^*) = \arg \max \varphi_{u,n}$. Ainsi, à chaque étape on est assuré de réunir le couple qui réalise les meilleures performances. Dans la suite, l'algorithme proposé par [Zhen&_03] sera appelé «improved ACG».

Une autre modification de l'ACG est exposée dans [Cho&_05], les auteurs ordonnent les sous porteuses selon le CgNR minimum ($\varphi_{min,n} = \min_{1 \leq u \leq U} \varphi_{u,n}$). On procède par $\varphi_{min,n}$ croissants. Le principe d'affectation de l'ACG est conservé : allocation opportuniste en tenant compte de N_u . En affectant très tôt dans le déroulement de l'algorithme une sous porteuse de faible $\varphi_{min,n}$ à son meilleur utilisateur, on améliore l'utilisation de cette sous porteuse par rapport à l'ACG. L'algorithme est appelé «modified ACG».

L'algorithme ACG et ses dérivés sont proposés dans un contexte MA (avec débits minimaux), mais ils sont utilisables en RA (avec débits minimaux) pourvu que l'allocation de bande qui les précède soit effectuée dans un contexte RA. Notre

contribution répondant à cette exigence, nous nous intéressons particulièrement aux performances des algorithmes ACG, «modified ACG» et «improved ACG » dans le chapitre 3.

Approches sous optimales mixtes (orientées utilisateurs puis sous porteuses)

Dans [KimKwak&_04], les auteurs traitent un problème MA avec débits minimaux et estimation souple. Une première phase affecte aux utilisateurs u , dans un ordre arbitraire, leur meilleure sous porteuse $m(u)$. Une deuxième phase se déroule jusqu'à ce que chaque utilisateur reçoive $\lceil r_u^\circ / R_{max} \rceil$ sous porteuses. Le principe de réduction de puissance est généralisé, on pose : $\Delta(u, n) = P(i(u), \Omega_u) - P(i(u)+1, \Omega_u \cup \{n\})$. Il s'agit de la réduction de puissance obtenue si n est affectée à u ¹. A chaque étape, l'utilisateur qui maximise $\Delta(u, m(u))$ se voit attribuer sa meilleure sous porteuse $m(u)$. Dans la dernière phase, chaque sous porteuse m restante est affectée à celui qui maximise $\Delta(u, m)$.

Enfin, nous présentons un algorithme qui résout le problème RA avec une contrainte d'équité pondérée (cf. tableau 2.5). Si le vecteur $\underline{\beta}$ est le vecteur unité (i.e. $\mathbf{1}$) dans [Wong&_04], on obtient un problème similaire au «max-min». Trois phases constituent l'algorithme de [Wong&_04]. La première phase est identique à [KimKwak&_04]. Une deuxième phase se déroule jusqu'à ce que chaque utilisateur reçoive $\lfloor \beta_u \cdot N \rfloor$ sous porteuses. Pour cela, le principe d'attribution de [RheeCioffi_00] est généralisé : l'utilisateur qui minimise r_u / β_u reçoit sa meilleure sous porteuse. S'il reste des sous porteuses non affectées, une dernière phase les affecte à l'utilisateur qui y présente le meilleur CgNR. On peut noter que [HuiZhou_06] propose un algorithme très largement inspiré de [Wong&_04]. La seule différence est l'introduction de *priorités* entre les utilisateurs au moment de la première phase (pour pallier un ordre arbitraire). Les deux dernières phases sont inchangées.

5.2. Allocation de puissance dans l'approche A

L'étape 2 de l'approche A n'est réalisée qu'une fois que les ensembles $\{\Omega_u\}_{1 \leq u \leq U}$ sont formés (ensembles des sous porteuses affectées à chaque utilisateur u). L'étape 2 détermine le MCS ou le niveau de puissance sur chaque sous porteuse. Nous ne décrivons pas les différentes heuristiques proposées pour cette étape car elles sont complètement décorréées de l'attribution des sous porteuses. Dans cette thèse, nous nous intéressons principalement à l'attribution de sous porteuses car il a été montré qu'une répartition uniforme de puissance est presque aussi compétitive que le *waterfilling*.

Problème RA

Le lecteur intéressé par des algorithmes d'allocation de puissance en RA peut se référer à : [Shen&_03], [Chee&_04], [Anas_04], [KimHan&_04], [LiLiu_05].

¹ Le CgNR moyen ϕ_u est calculé sur les sous porteuses que possède l'utilisateur. Dans le BABS, la moyenne était faite sur l'ensemble des sous porteuses.

Problème MA

Le lecteur intéressé pourra se référer à [WangNiu_05] et [ErmoMak_05].

5.3. Allocation conjointe de sous porteuses et de MCS dans l'approche B

L'approche B consiste à réaliser conjointement l'attribution de sous porteuses et le choix des MCS sur les sous porteuses. Les auteurs proposent des heuristiques pour éviter la résolution par programmation entière et améliorer ainsi le temps d'exécution. Toutes les heuristiques rencontrées sont MA.

La première étape, que l'on traite ici, réalise une affectation concernant à la fois les sous porteuses et le choix de MCS. Souvent, il existe une deuxième étape (cf. section 5.4) d'amélioration dont l'objectif est de réduire la puissance transmise. L'intérêt de cette étape est de pouvoir affiner l'allocation quand le temps de cohérence du canal le permet.

5.3.1. Option 1 : les nombres de sous porteuses N_u ne sont pas connus a priori

Approche par résolution de conflits

Dans l'initialisation de [Zhang_04], pour chaque utilisateur u la BS construit une liste des meilleures sous porteuses pour atteindre le débit r_u° ; le nombre de bits $b_{u,n}$ prévu sur chacune d'entre elles est calculé par l'algorithme du *bit loading*. Les sous porteuses n'apparaissant que dans une seule liste sont affectées ; les autres sous porteuses sont dites « conflictuelles » et sont traitées par CgNR moyen décroissant. Soit n une sous porteuse « conflictuelle », pour chaque utilisateur v en conflit, on calcule le coût¹ en puissance $D_{v,n}$ s'il ne reçoit pas la sous porteuse n . Pour chaque u , $Dtot_{u,n} = \sum_{v \neq u} D_{v,n}$ représente l'augmentation de la puissance transmise si u recevait la sous porteuse n . L'utilisateur qui minimise cette dépense globale ($u^* = \arg \min_u Dtot_{u,n}$) reçoit la sous porteuse conflictuelle.

Allocation progressive du spectre

Dans l'initialisation de l'algorithme de [Yu&_06], chaque utilisateur u reçoit tour à tour (l'ordre est arbitraire) la meilleure sous porteuse (sur l'ensemble disponible) ; il y place son débit minimal de r_u° bits. Ensuite, à chaque nouvelle allocation, la sous porteuse n et l'utilisateur u auquel elle sera affectée sont déterminés conjointement. Comme dans [KimKwak&_04], l'utilisateur u qui maximise la réduction de puissance $\Delta(u, m(u))$ se voit attribuer sa meilleure sous porteuse $m(u)$. La différence avec [KimKwak&_04] vient d'un calcul plus précis de $\Delta(u, m(u))$; de plus la répartition de bits sur le nouvel ensemble Ω_u est réalisée par une alternative au *bit loading* qu'on ne détaillera pas ici.

Dans cet algorithme, les conflits sont automatiquement réglés car si deux utilisateurs ont la même meilleure sous porteuse, c'est celui qui maximise $\Delta(u, m(u))$ qui la reçoit.

Dans [Zhang_04] et [Yu&_06], il n'y a pas d'étape d'amélioration de la puissance transmise.

1 Pour calculer $D_{u,n}$, l'auteur réaffecte les $b_{u,n}$ bits qui étaient affectés à n sur le reste des sous porteuses (à l'exclusion des anciennes sous porteuses conflictuelles) grâce à l'algorithme du *bit loading*.

5.3.2. Option 2 : les nombres de sous porteuses N_u doivent être connus

Dans [WongTsui&_99], [PietrykJan_02] et [Chen&_04], un processus similaire est utilisé pour la phase d'allocation grossière (une fois les N_u connus). Il s'agit pour chaque utilisateur de classer les sous porteuses par CgNR décroissant ([WongTsui&_99], [Chen&_04]) ou par puissance croissante en cas de MCS fixe ([PietrykJan_02]). L'ordre de traitement des utilisateurs est arbitraire¹.

Soit une matrice $N \times U$, nommée $\mathbf{A}$, comportant des index de sous porteuses : $a_{n,u} \in \{1 \dots N\}$ avec $1 \leq n \leq N$ et $1 \leq u \leq U$. Pour chaque u , la colonne correspondante est ordonnée de la meilleure sous porteuse à la moins bonne. Chaque ligne n est parcourue de 1 à U : une sous porteuse disponible $a_{n,u}$ est affectée à u si N_u n'est pas atteint. Le premier utilisateur est privilégié par rapport au dernier : une même meilleure sous porteuse sera allouée au premier utilisateur plutôt qu'au dernier (à l'index 1 plutôt qu'à l'index U). Une fois la matrice parcourue ligne par ligne, une phase d'amélioration a pour objectif de minimiser la puissance émise en échangeant les sous porteuses entre les utilisateurs et, si le MCS n'est pas fixe, en modifiant les MCS sur les sous porteuses. A notre connaissance, l'initiateur du procédé d'échange de sous porteuses est [WongTsui&_99], cette idée est ensuite reprise dans plusieurs articles ([PietrykJan_02], [Chen&_04] et dans des références récentes [ChenKron_07]).

5.4. Amélioration de la puissance transmise dans l'approche B

Cette phase constitue la deuxième étape de l'approche B.

Le procédé d'amélioration de la puissance transmise dans [WongTsui&_99] est le suivant. Pour chaque couple (u,v) , on calcule $p_{u,v}$ le facteur de réduction de puissance : $p_{u,v} = \delta p_{u,v} + \delta p_{v,u}$. La quantité $\delta p_{u,v}$ est la plus grande économie de puissance calculée en faisant successivement l'hypothèse que chacune des sous porteuses de l'utilisateur u est réallouée à l'utilisateur v . On nomme $n_{u,v}$ l'index de la sous porteuse de u sur laquelle le maximum $\delta p_{u,v}$ est réalisé. Si $\max(p_{u,v}) > 0$ (maximum pris sur l'ensemble des couples (u,v)), on réalise l'échange : l'utilisateur v reçoit la sous porteuse $n_{u,v}$ et l'utilisateur u reçoit la sous porteuse $n_{v,u}$. Pratiquement, cette tâche peut prendre du temps. Pendant la durée de l'allocation impartie, on la poursuit tant que $\max(p_{u,v}) > 0$. Quand $\max(p_{u,v}) < 0$, alors aucune amélioration n'est possible. Dans [Chen&_04] et [ChenKron_07] il y a en plus des échanges, la possibilité de réaffecter sans contrepartie une sous porteuse : dans ce cas l'utilisateur u_1 reçoit une sous porteuse provenant de l'utilisateur u_2 et l'utilisateur u_2 perd une sous porteuse.

¹ Par index croissant pour une simplifier la lecture

6. Conclusions

Nous avons réalisé une étude bibliographique sur l'allocation de ressources en OFDMA. Deux principaux problèmes, MA et RA, sont traités dans la littérature. Le problème MA minimise la puissance totale transmise tandis que le problème RA maximise le débit global de la cellule.

Dans ces deux types de problèmes, l'optimisation conjointe en sous porteuses et en puissance est coûteuse en temps d'exécution. Cette optimisation est simplifiée lorsqu'on distingue l'attribution de sous porteuses et l'allocation de puissance. L'attribution de sous porteuses est souvent réalisée en deux tâches successives : l'allocation de largeur de bande et l'affectation spécifique des sous porteuses. Ces tâches, respectivement appelées tâche 1 et tâche 2, sont aussi bien considérées en MA qu'en RA.

Pour l'allocation de bande, il n'existe pas d'algorithme optimal dans le cas général (sélectivité en fréquence et modulation adaptative). Le nom de l'algorithme à retenir pour cette phase est le BABS ([Kivanc&_00], cf. section 5.1.2.a), il a été proposé en MA.

Concernant la phase d'affectation de sous porteuses, on peut garder en mémoire l'algorithme ACG et ses algorithmes dérivés ([Kivanc&_00], [Cho&_05], [Zhen&_03], cf. section 5.1.2.b). Ces algorithmes proposés en MA peuvent être utilisés en RA pourvu qu'un algorithme soit proposé pour l'allocation de bande. De plus, on retrouve ces algorithmes dans le contexte multicellulaire (cf. chap. 5).

Dans les problèmes MA, certains auteurs réalisent l'attribution de sous porteuses et le choix des MCS de façon conjointe. Dans ce cas, ils considèrent une phase d'allocation grossière et une phase d'amélioration. Dans la phase d'amélioration, pour diminuer la puissance transmise, les sous porteuses peuvent être échangées entre deux utilisateurs distincts et les MCS peuvent être «échangés» sur les sous porteuses d'un utilisateur. Un algorithme que l'on peut retenir dans cette catégorie est la proposition de [WongTsui&_99] (cf. sections 5.3.2 et 5.4). Cet algorithme est parfois réutilisé en ordonnancement (cf. chap. 6).

La classification que nous avons proposé est assez efficace comme on peut le voir dans les tableaux 2.6 et 2.7. Le premier classe les approches dans le contexte monocellulaire. Le second classe quelques références dans le contexte multicellulaire ; elles sont décrites dans le chapitre 5.

Enfin, on peut remarquer que, dans le contexte monocellulaire, peu de références intéressantes ont été proposées en RA avec contraintes en débit minimaux (sans équité pondérée). Il n'existe pas d'algorithmes spécifiques pour l'allocation de bande (calcul du nombre de sous porteuses). Concernant l'affectation de sous porteuses, c'est en MA que l'on trouve des idées intéressantes sur la résolution de conflits ([Zhang_04], [Yu&_06], cf. section 5.3.1). Une sous porteuse est attribuée sur un critère précis de coût et non d'après un choix arbitraire sur l'ordre des utilisateurs ([Gross&_03], cf. section 5.1.2.b). C'est en se fondant sur ce constat que nous proposons des contributions en RA dans le chapitre 3.

Tableau 2.6 : Approches des références dans le contexte monocellulaire

	Séparation de l'attribution de sous porteuses et de l'allocation de puissance (Approche A)	Attribution de sous porteuses et allocation de puissance imbriquée (Approche B)
Allocation de bande et affectation de sous porteuses imbriquées (Option 1)	[Anas_04] [DoufexiArmour_05] [JangLee_03] [Ko&_06] [KimHan&_04] [NguyenHan_06] [RheeCioffi_00]	[Yu&_06] [Zhang_04]
Allocation de bande et affectation de sous porteuses séparées (Option 2)	[Chen&_04] [Cho&_05] [Gross&_03] [HuiZhou_06] [KimKwak&_04] [Kivanc&_00] [Kivanc&_03] [Pfletschinger&_02] [Wong&_04] [YinLiu_00] [Zhen&_03]	[Chen&_04] [ChenKron_07] [PietrykJan_02] [WongTsui&_99]

Tableau 2.7 : Approches de quelques références dans le contexte multicellulaire

	Séparation de l'attribution de sous porteuses et de l'allocation de puissance (Approche A)	Attribution de sous porteuses et allocation de puissance imbriquée (Approche B)
Allocation de bande et affectation de sous porteuses imbriquées (Option 1)	[Yan&_03] [Han_03]	[Koutsopoulos_02]
Allocation de bande et affectation de sous porteuses séparées (Option 2)	[PietrzykJanssen_03] [Kwon&_05] [Hamouda&_06]	[Junqiang&_03]

Chapitre 3. Allocation de ressources en OFDMA : contribution dans le contexte monocellulaire

1. Introduction

Dans ce chapitre, nous étudions les algorithmes d'allocation de ressources radio dans un contexte monocellulaire. Nous nous intéressons au problème RA sur la voie descendante. Les utilisateurs doivent transmettre un débit minimal r_u^o . Dans la section 2, nous expliquons les motivations de notre contribution puis nous la décrivons. Nous comparons ensuite plusieurs approches (décrites dans le chapitre 2) entre elles et avec de nouvelles propositions.

Une recherche de contrainte en débits réalisables est présentée dans la section 3 sous le nom de méta problème RA. Il s'agit ici de trouver un débit minimal commun qui peut être garanti à tous les utilisateurs. La section 4 synthétise les résultats de simulations où les principaux algorithmes sont comparés.

2. BARE et RPO : deux nouveaux algorithmes en RA

2.1. Motivations et contexte de notre contribution

Notre contribution s'applique au problème RA qui semble assez adapté aux problématiques des réseaux haut débits actuels. Nous rappelons ici la formulation de ce problème :

$$\begin{aligned} &\text{Maximiser } \sum_u r_u \\ &\text{avec } P_T < P_{T,Max}, \end{aligned} \quad (1)$$

$$\text{pour tout } u, r_u \geq r_u^o \quad (2)$$

Comme nous l'avons vu, la contrainte sur le vecteur de débit complique la résolution du problème. Dans le cadre de l'approche A¹, nous travaillons avec une allocation uniforme de puissance sur les sous porteuses (choix similaire dans [RheeCioffi_00], [JangLee_03], [Gross&_03], [HuiZhou_06] et [NguyenHan_06]). L'attribution des sous porteuses est séparée en deux tâches : allocation de bande et affectation de sous porteuses.

¹ À notre connaissance, cette approche a été unanimement adoptée dans les problèmes RA.

Deux constats ont motivé notre contribution. Premièrement, les travaux existants en allocation de bande (tâche 1) ne sont adaptés pas au problème RA. La proposition de [YinLiu_00] reprend le principe de réduction de puissance utilisé dans le BABS. Or le principe de réduction de puissance a pour objectif de minimiser la puissance transmise. En RA, on peut utiliser la totalité de la puissance et cette dernière est souvent uniformément répartie. Dans ce contexte, utiliser le principe de réduction de puissance est contreproductif pour maximiser le débit. Nous proposons un algorithme appelé BARE (*Bandwidth Allocation based on Rate Estimation*, [LengGod&_06]) et adapté aux contraintes du problème RA.

Deuxièmement, les travaux existants en affectation de sous porteuses (tâche 2) et traitant directement les conflits¹ résolvent un problème MA ([Zhang_04], [Yu&_06]). Nous avons proposé le RPO (*Rate Profit Optimization*, [LengGod&_06]) qui apporte une solution à la résolution de conflits en RA.

Remarques pour la mise en oeuvre

Les deux algorithmes proposés s'utilisent l'un à la suite de l'autre pour déterminer l'attribution en sous porteuses sur une période temporelle où la réponse du canal est supposée invariante. Dans un contexte TDD, par exemple, la trame montante remonte les mesures des terminaux sur la qualité du canal et en début de trame descendante la station de base informe les terminaux du résultat de l'allocation. Les deux algorithmes proposés seraient donc appliqués avant le début de chaque trame descendante.

2.2. Le BARE : un algorithme pour l'allocation de largeur de bande

Nous proposons un algorithme pour déterminer le nombre de sous porteuses (ou largeur de bande) par utilisateur en RA (cf. Fig. 2.2: option 2, tâche 1). L'expression de la contrainte sur les débits est $\underline{r} \geq \underline{r}^o$ où $\underline{r}^o$ est un vecteur quelconque.

En MA, l'algorithme BABS augmente le nombre de sous porteuses d'un utilisateur sur un critère de puissance. L'algorithme BARE (*Bandwidth Allocation based on Rate Estimation*) augmente le nombre de sous porteuses d'un utilisateur sur un critère de débit.

Soit $\underline{\varphi} = (\varphi_1, \dots, \varphi_u, \dots, \varphi_U)$, le vecteur de CgNR moyens. On considère un canal plat : un utilisateur u voit toutes les sous porteuses avec le même CgNR moyen φ_u . L'allocation de puissance est supposée uniforme. Grâce à ces deux hypothèses, on estime le débit des utilisateurs. On appelle h la fonction qui donne le débit en fonction du SNR.

Lorsqu'un utilisateur u reçoit une largeur de bande de M sous porteuses, son débit est estimé à $r_{est}(u, M) = M \times h(p \varphi_u)$ avec $h(p \varphi_u)$ le débit estimé d'une sous porteuse de u et $p = P_{T,max}/N$.

1 Cas où plusieurs utilisateurs ont la même meilleure sous porteuse.

Principe de l'algorithme

L'algorithme cherche à minimiser la différence entre le débit estimé et le débit minimal de chaque utilisateur. Pour cela, l'algorithme modifie la largeur de bande $i(u)$ qui est attribuée à un utilisateur u .

La différence entre le débit estimé et le débit minimal recherché est appelée *gap*. Le vecteur $\underline{i} = (i(u))_u$ représente les nombres de sous porteuses (ou largeur de bande) des utilisateurs au cours de l'algorithme BARE. Pour un utilisateur u , la différence entre le débit estimé et le débit minimal recherché dépend de la largeur de bande $i(u)$ reçue :

$$gap(u, i(u)) = r_{est}(u, i(u)) - r_u^o.$$

Initialisation

$$i(u) = \lfloor N/U \rfloor$$

Phase post_initialisation

On s'assure que toutes les sous porteuses seront utilisées en augmentant le nombre de sous porteuses des utilisateurs les plus faibles. L'utilisateur le plus éloigné de son débit minimal cible reçoit une sous porteuse supplémentaire. Cette phase peut être décrite par le pseudo code suivant :

```
while  $\sum_u i(u) < N$ 
   $u^* := \arg \min_u gap(u, i(u))$ 
   $i(u^*) := i(u^*) + 1$ 
end while
```

Phase principale de l'algorithme BARE

S'il existe un utilisateur u^* qui ne satisfait pas son objectif et un utilisateur u' qui satisfait son objectif malgré le retrait d'une sous porteuse, on retire une sous porteuse à u' et on l'ajoute à u^* . L'algorithme s'arrête dans deux cas. Dans le premier cas tous les utilisateurs satisfont leurs objectifs. Dans le deuxième cas aucun des utilisateurs «satisfaits» ne peut donner une sous porteuse et rester «satisfait». Cette phase peut être décrite par le pseudo code suivant :

```
while  $\min_u (gap(u, i(u)) < 0 \ \& \ \max_u gap(u, i(u) - 1) > 0,$ 
   $u^* := \arg \min_u gap_u(u, i(u))$ 
   $u^o := \arg \max_u gap_u(u, i(u) - 1)$ 
   $i(u^*) := i(u^*) + 1$ 
   $i(u^o) := i(u^o) - 1$ 
end while
```

2.3. Le RPO : un algorithme pour l'affectation des sous porteuses

Nous avons proposé le RPO (*Rate Profit Optimization*, [LengGod&_06]) qui est une nouvelle alternative aux heuristiques décrites dans le chapitre 2 (cf. Fig. 2.2: option 2, tâche 2). Les performances de cet algorithme s'approchent de celles de l'algorithme Hongrois.

2.3.1. Notion de profit et principe de l'algorithme

Il arrive souvent que des utilisateurs différents aient la même meilleure sous porteuse. Comment résoudre le conflit? Des priorités arbitraires ([Gross&_03]) peuvent être utilisées entre les utilisateurs (sans garantie sur l'évolution du débit total). Très souvent, l'utilisateur qui a le meilleur CgNR reçoit la sous porteuse ([Zhen&_03]). Cela peut être injuste si l'on prive un utilisateur d'une unique très bonne sous porteuse pour un utilisateur qui a plusieurs sous porteuses dont le CgNR est élevé. On a vu dans [Zhang_04], un critère de résolution de collisions axé sur la minimisation de la puissance transmise. Ici, l'algorithme RPO propose une règle de résolution de conflits axée sur l'amélioration du débit global. On considère *la situation de référence* dans laquelle l'utilisateur u° qui a le meilleur CgNR reçoit la sous porteuse en conflit. Pour maîtriser l'évolution du débit global, il faut vérifier si un autre utilisateur ne ferait pas un meilleur usage de la sous porteuse.

On définit la notion de **profit** de l'utilisateur u . Il s'agit de l'impact sur le débit global dans l'hypothèse où l'utilisateur u reçoit la sous porteuse en conflit. Si le profit de l'utilisateur u est positif, affecter la sous porteuse à u plutôt qu'à u° (situation de référence) augmente le débit global. Le profit est calculé pour les utilisateurs qui ont la même meilleure sous porteuse et dépend du niveau de leurs secondes meilleures sous porteuses. Dans [Zhang_04] (problème MA), l'utilisateur qui minimise le coût en puissance reçoit la sous porteuse en conflit. Dans le RPO ([LengGod&_06], problème RA), l'utilisateur qui maximise le profit reçoit la sous porteuse en conflit.

Remarque : Dans [Ko&_06], publié à la même période que l'algorithme RPO, lorsqu'un utilisateur a deux sous porteuses de même qualité (même débit), il reçoit celle qui intéresse le moins les autres utilisateurs. On voit que le RPO a généralisé l'idée : avant d'affecter une sous porteuse on s'intéresse systématiquement aux autres utilisateurs.

Le **principe de l'algorithme** est le suivant. La meilleure sous porteuse des utilisateurs est déterminée. En l'absence de conflit, les sous porteuses sont affectées. En cas de conflit, la sous porteuse est affectée à celui qui y présente le meilleur profit. Un utilisateur qui a reçu N_u sous porteuses n'est plus en compétition pour recevoir des sous porteuses. L'algorithme se termine lorsque toutes les sous porteuses sont affectées.

2.3.2. Notations et calcul du profit

Au cours de l'algorithme, on note Π l'ensemble des utilisateurs non satisfaits ($|\Omega_u| < N_u$) et H l'ensemble des sous porteuses encore disponibles.

A chaque itération de l'algorithme, la BS détermine la meilleure sous porteuse des utilisateurs. L'ensemble C comporte les sous porteuses désignées pour des utilisateurs différents dans Π . Soit n une sous porteuse conflictuelle ($n \in C$), on note Λ_n l'ensemble des utilisateurs qui ont n comme meilleure porteuse. C'est uniquement parmi ces utilisateurs que l'on calcule le profit : celui qui a le meilleur profit reçoit la sous porteuse n . On note $u^\circ(n)$, l'utilisateur de Λ_n qui a le meilleur CgNR (pour n fixée, on peut simplifier la notation à u°).

Calcul du profit

Lorsqu'à une étape donnée, un utilisateur u ne reçoit pas sa meilleure sous porteuse, il est candidat pour sa seconde meilleure sous porteuse à l'étape suivante.

La seconde meilleure sous porteuse $n''(u)$ de chaque utilisateur u de Λ_n est donc déterminée et on appelle $rateGap(u)$ le débit perdu par l'utilisateur u s'il recevait sa seconde meilleure sous porteuse au lieu de sa meilleure sous porteuse. Alors $rateGap(u)$ est l'écart entre le débit de u sur sa meilleure sous porteuse $n'(u)$ et sur sa seconde meilleure sous porteuse $n''(u)$:

$$rateGap(u) = r_{u, n'(u)} - r_{u, n''(u)}$$

Si on attribue la sous porteuse n à u plutôt qu'à u° , l'utilisateur u gagne $rateGap(u)$ bits et l'utilisateur u° perd $rateGap(u^\circ)$ bits. Finalement, si on attribue la sous porteuse n à u plutôt qu'à u° , le débit global varie de $rateGap(u) - rateGap(u^\circ)$. Cette quantité représente le profit de l'utilisateur u :

$$profit(u) = rateGap(u) - rateGap(u^\circ)$$

Le profit de l'utilisateur u° est par définition nul. Comparé à u° , un utilisateur qui a un plus grand écart entre sa meilleure et sa seconde meilleure sous porteuse a un profit positif. Affecter la sous porteuse à un tel utilisateur (plutôt qu'à u°) améliore le débit global. L'utilisateur de Λ_n qui reçoit la sous porteuse n maximise le profit. Par conséquent l'utilisateur u° ne reçoit la sous porteuse que lorsque les autres ont un profit négatif.

2.3.3. Pseudo code de l'algorithme

```

H := {1..N}           # all subcarriers
Π := {1..U}           # all users

while |H| > 1

    for each u ∈ Π
        n'(u) := arg max_n (φu,n)
    end for
    Let C be the set of subcarriers such that n'(u) == n'(v) and u ≠ v ;

    for each u ∈ Π           # Simple cases resolution
        if n'(u) ∉ C
            Ωu := Ωu + {n'(u)}
            H := H - {n'(u)}
            if |Ωu| == Nu
                Π := Π - {u}
            end if
        end if
    end for

    while |C| ≠ 0, pick n in C   # Collision cases resolution
        u° = argmaxu in Π (φu,n)
        Λn is formed

```

```

    for each  $u \in \Lambda_n$  (begin with  $u^\circ$ )
         $n''(u) := \arg \max_{m \neq n'}(u) (\varphi_{u,m})$ 
         $r_{u,1} := h(\gamma_{u,n})$ 
         $r_{u,2} := h(\gamma_{u,n''(u)})$ 
         $\text{rateGap}(u) := r_{u,1} - r_{u,2}$ 
         $\text{profit}(u) := \text{rateGap}(u) - \text{rateGap}(u^\circ)$ 
    end for

     $u' := \arg \max_{u \in \Lambda_n} (\text{profit}(u))$ 

     $\Omega_{u'} := \Omega_{u'} + \{n\}$ 

     $C := C - \{n\}$ 
     $H := H - \{n\}$ 
end while

end while

The last subcarrier (if any), is assigned to the last user of  $\Pi$ 

```

3. Méta-Problème RA

Lorsqu'on résout un problème RA avec une contrainte sur les débits individuels, le problème peut ne pas avoir de solutions. C'est le cas lorsque, compte tenu des conditions du canal, aucune instance d'affectation de sous porteuses n'assure aux utilisateurs les débits minimaux recherchés. Le «méta-problème RA» est alors la recherche d'un vecteur de débits minimaux réalisable ; c'est à dire un vecteur pour lequel on obtient une solution. Nous proposons d'utiliser l'algorithme BARE pour effectuer cette recherche. Lorsque le vecteur de débits minimaux est quelconque, on peut utiliser une réduction proportionnelle. Si le débit minimal est commun à tous, on peut utiliser une réduction par décrémentation.

3.1. Réduction proportionnelle

Le vecteur $\underline{r}^\circ$ est quelconque. Si les valeurs du vecteur $\underline{r}^\circ$ sont trop grandes, le problème n'a pas de solutions. Dans ce cas, l'algorithme BARE proposé pour l'allocation de bande n'aboutit pas (`success_BARE==0`) : le débit estimé de certains utilisateurs est inférieur au débit minimal.

Si l'important est de trouver un vecteur de débit minimal réalisable (et non des valeurs figées du vecteur), on peut relancer l'algorithme BARE avec $\underline{r}^\circ = \lambda' \underline{r}^\circ$ avec $\lambda' < 1$. Le facteur λ' permet de respecter la structure initiale du vecteur $\underline{r}^\circ$ (c'est à dire les rapports de proportionnalité entre les débits minimaux des utilisateurs) et de réduire les valeurs des débits minimaux à satisfaire.

Le pseudo code est le suivant :

```
while success_BARE==0 & (minu( $\underline{r}^\circ > r^\circ_{\min}$ ))  
     $\underline{r}^\circ = \lambda' \underline{r}^\circ$   
    BARE  
end
```

3.2. Réduction par décrémentation

La contrainte est commune à tous les utilisateurs : le vecteur de débit minimal est $\underline{r}^\circ = r^\circ \mathbf{1}$. Cela correspond à une équité en débit minimal.

Pour ce problème, l'algorithme BARE est appliqué avec $r^\circ_u = r^\circ$ pour tous les utilisateurs. Si le débit estimé de certains utilisateurs est inférieur au débit minimal on peut baisser la contrainte r° d'un facteur λ ($\lambda \in \mathbb{R}^*$).

Le pseudo code suivant :

```
while success_BARE==0 & ( $r^\circ > \lambda$ )  
     $r^\circ = r^\circ - \lambda$   
    BARE  
end
```

4. Simulations et résultats

4.1. Objectifs

Dans les simulations, nous nous intéressons au problème d'équité en débit minimal. De ce fait, la contrainte sur les débits individuels est commune à tous. Dans ce contexte nous résolvons le méta problème RA : nous recherchons le plus grand débit commun réalisable avant d'effectuer l'affectation des sous porteuses.

Dans la section 4.3.a, nous montrons que l'algorithme BARE est plus efficace dans la résolution du méta problème RA que le BABS. Après cette section, le BARE est adopté pour l'allocation de bande. Les algorithmes BARE et RPO sont complètement indépendants. L'algorithme BARE peut être utilisé à l'entrée d'autres algorithmes.

Les algorithmes ainsi comparés au RPO (*Rate Profit Optimization*) sont les algorithmes Hongrois, ACG (*Amplitude Craving Greedy algorithm*), «improved ACG», «modified ACG» et bDA (*basic Dynamic Assignment*). Nous traçons aussi les performances de l'algorithme opportuniste ([JangLee_03]) et de l'heuristique équitable ([RheeCioffi_00], problème «max-min») qui résolvent deux problèmes RA extrêmes.

On trace le débit moyen par sous porteuse pour évaluer les performances en débit. Concernant les contraintes sur le débit minimal, on s'intéresse à une probabilité d'échec (ou *outage*) par rapport au débit minimal à réaliser. Il s'agit du taux d'utilisateurs qui n'atteignent pas le débit minimal. On définit un facteur d'équité qui correspond au rapport maximal entre les débits finaux des utilisateurs. Enfin pour évaluer la complexité, on compare le temps d'exécution des différents algorithmes.

4.2. Conditions de simulation

Nous considérons une cellule constituée d'une station de base et de U utilisateurs uniformément répartis. Le modèle de canal adopté est constitué de N sous porteuses parallèles sur la bande B . Aucune corrélation n'est considérée entre les sous porteuses.

4.2.1. Modèle de canal

Le carré du module du gain du canal de l'utilisateur u sur la sous porteuse n est : $|g_{u,n}|^2 = K d(u)^{-\alpha} a_{sh}(u) a_f(u,n)$. Dans cette formule, K est l'atténuation à la distance de référence 1 m ; $d(u)$ est la distance de l'utilisateur ; α est le coefficient d'atténuation ; $a_{sh}(u)$ est une variable lognormale qui traduit l'effet de masque subi par toutes les sous porteuses de l'utilisateur u ; enfin, $a_f(u,n)$ est une variable exponentielle pour modéliser l'évanouissement rapide (*fast fading*) de la sous porteuse n .

Les sous porteuses subissent par ailleurs un bruit blanc gaussien de variance $\sigma^2 = N_0 W / N$. Ainsi le CgNR est donné par $\phi_{u,n} = |g_{u,n}|^2 / \sigma^2$ et le SNR (*Signal to Noise Ratio*) s'en déduit simplement : $\gamma_{u,n} = P_{u,n} \phi_{u,n}$. Rappelons que nous adoptons une allocation de puissance uniforme $P_{u,n} = p$ sur chaque sous porteuse. Pour calculer le débit à partir du SNR, nous utilisons comme fonction $h(\gamma) = \min(R_{max}, f^{-1}(\gamma))$ avec la fonction f proposée par [Kivanc&_03] : $\gamma = f(r) = 0.06 r^3$.

4.2.2. Déroulement de la simulation

La puissance p par sous porteuse varie de 1 à 35 mW (0 à 15 dBm). Pour $N = 128$, la puissance émise par la BS varie alors de 21 à 36 dBm. Pour chaque valeur de p , les résultats sont moyennés sur 10000 *snapshots*. Les intervalles de confiances sont pris à 95 %. Pour chaque *snapshot*, on tire de façon aléatoire la distance des utilisateurs (uniformément répartis dans la cellule), l'effet de masque et l'évanouissement rapide sur tous les couples (u, n) . Les paramètres de simulation sont résumés dans le tableau 3.1.

Tableau 3.1 : Paramètres de simulation, contexte monocellulaire

Paramètres	Valeurs
W (MHz)	1
U	8
N	128
$R_{\max}$ (bits)	4.5
K : constante de l'atténuation de parcours α : exposant de l'atténuation de parcours	10^{-4} 2.8
$D_{\max}$ (km)	5
σ_{sh} (dB)	8
Moyenne de la variable exponentielle : a_f Ecart type de la variable exponentielle : a_f	1 1
Puissance par sous porteuse (mW)	1...35
N_0 : densité de bruit blanc (dBm/Hz)	-174

4.3. Débit par sous porteuse

a) Tâche 1

L'algorithme BARE est un algorithme mieux adapté au cas RA ; il utilise les N sous porteuses et le budget $P_{T, \max}$ disponible. Pour le montrer, on compare le BARE et le BABS en utilisant les ensembles $\{N_u\}_{1 \leq u \leq U}$ qu'ils produisent dans l'algorithme Hongrois. En effet, l'algorithme Hongrois apporte la solution optimale à la tâche 2 pour un ensemble $\{N_u\}_{1 \leq u \leq U}$ fixé. Le meilleur algorithme de la tâche 1 peut donc être désigné en comparant le débit par sous porteuse obtenu en utilisant les combinaisons BARE + Hongrois et BABS + Hongrois.

Pour résoudre le méta problème RA avec équité en débit minimal, on utilise la réduction par décrétement. Si une contrainte est trop élevée, le problème se pose en terme de dépassement du budget de puissance dans le BABS, tandis que dans le BARE il se pose en terme d'*outage* (nombre d'utilisateurs insatisfaits).

Dans la figure 3.1, le débit obtenu par l'algorithme BARE est plus élevé que celui du BABS. Lors de l'ajustement du débit minimal, l'algorithme BARE trouve un débit minimal réalisable plus faible que celui du BABS. La différence est au maximum de 2 bits par temps symbole OFDM (soit 20 kbps pour un temps symbole OFDM de 100 μ s) pour chaque valeur de p . Grâce au BARE, ces ressources sont utilisées pour maximiser le débit global.

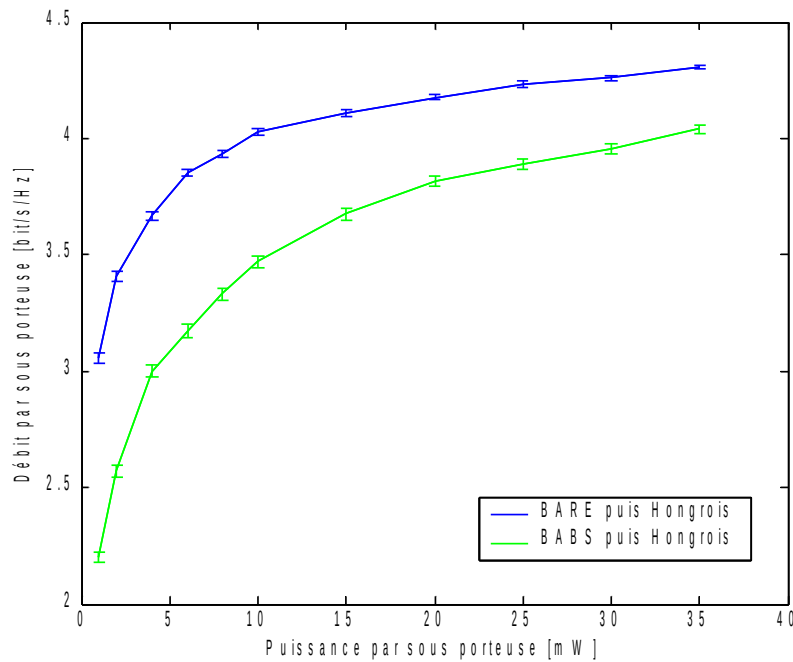


Figure 3.1 : Comparaison des méthodes d'allocation de bande

b) Tâche 2

Impact des stratégies d'équité

Garantir à chaque utilisateur u un débit minimal r_u° ou maximiser le débit minimal d'un utilisateur provoque une baisse de la capacité globale par rapport à une stratégie purement opportuniste. La différence de débit entre l'algorithme opportuniste et l'algorithme Hongrois (précédé du BARE) représente la perte de capacité provoquée par la garantie d'un débit minimal. La différence de débit entre l'algorithme opportuniste et l'heuristique équitale ([RheeCioffi_00]) représente la perte de capacité induite par la stratégie «max-min» (recherche d'une équité totale en débit).

Nous avons cherché à garantir un débit minimal r° commun à chaque utilisateur. Nous qualifions cela d'équité en débit minimal. Sur la figure 3.2, on voit que les stratégies d'équité pénalisent moins la capacité lorsque la puissance augmente. La perte de capacité due à l'équité en débit minimal diminue de 30 % à 4 % tandis que celle due à l'équité totale diminue de 46 % à 7 % quand p varie de 1 à 15 mW par sous-porteuse.

Comparaison : algorithme Hongrois, RPO, ACG, «modified ACG», «improved ACG»

L'algorithme Hongrois donne la solution optimale par rapport à la répartition $\{N_u\}_{1 \leq u \leq U}$ indiquée par le BARE. Il présente donc le meilleur débit par sous-porteuse. L'algorithme RPO présente de bonnes performances car il s'approche de l'algorithme Hongrois (moins d'1 % de différence). La différence du RPO avec l'algorithme Hongrois est stable : elle ne dépend ni de la puissance ni des paramètres du modèle (K, α, σ_{sh} , etc...).

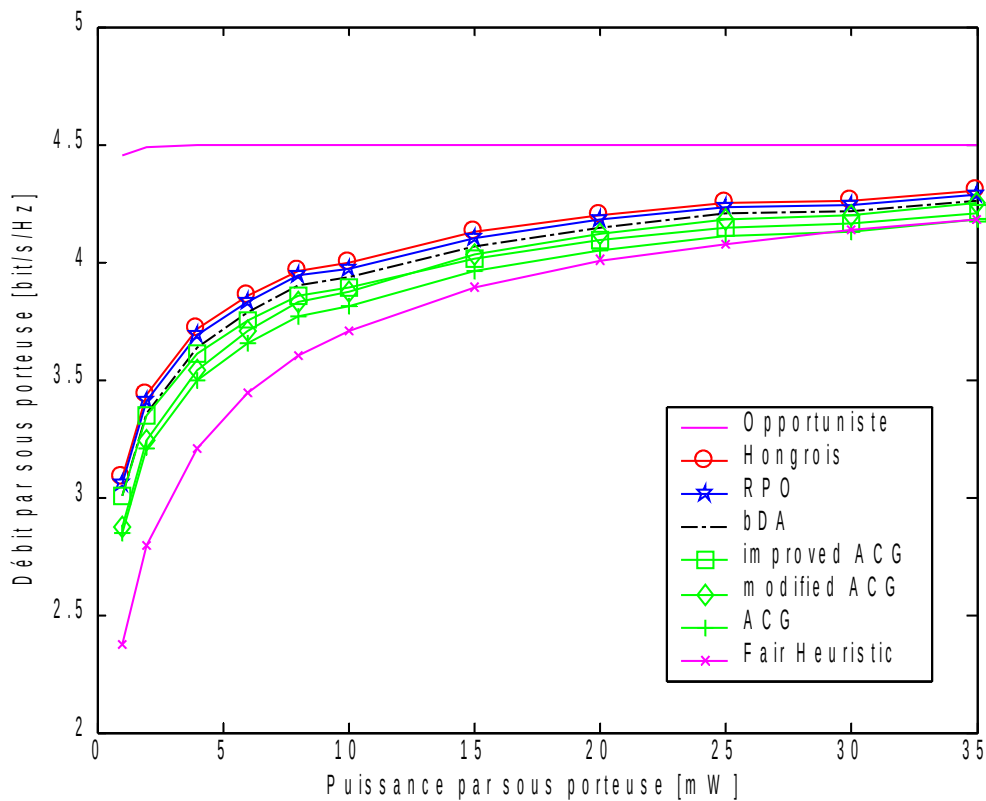


Figure 3.2 : Comparaison de débit entre les algorithmes d'affectation de sous porteuses

Le bDA et l'algorithme «improved ACG» suivent avec respectivement 2 % et 3 % de différence avec l'algorithme Hongrois.

Les performances de l'algorithme «modified ACG» sont assez variables et dépendent beaucoup des paramètres du modèle de propagation. L'écart entre cet algorithme et l'algorithme Hongrois baisse de 8 % à 3 % avec la puissance. A partir de $p = 10$ mW et $\alpha = 2.8$, le débit de «modified ACG» dépasse celui de «improved ACG». On a remarqué que pour $\alpha > 3$, le débit de «modified ACG» reste inférieur à celui de «improved ACG». L'algorithme «modified ACG» n'est pas insensible aux paramètres du modèle canal.

L'algorithme ACG présente les performances les plus faibles à cause de l'impact de l'ordre de traitement des sous porteuses.

La figure 3.3 permet de visualiser les intervalles de confiances des résultats sur le débit ; il s'agit d'un *zoom* de la figure 3.2. On voit que les intervalles de confiance ne se chevauchent pas. Au vue des marges d'erreurs, le RPO reste l'algorithme le plus proche de l'algorithme Hongrois.

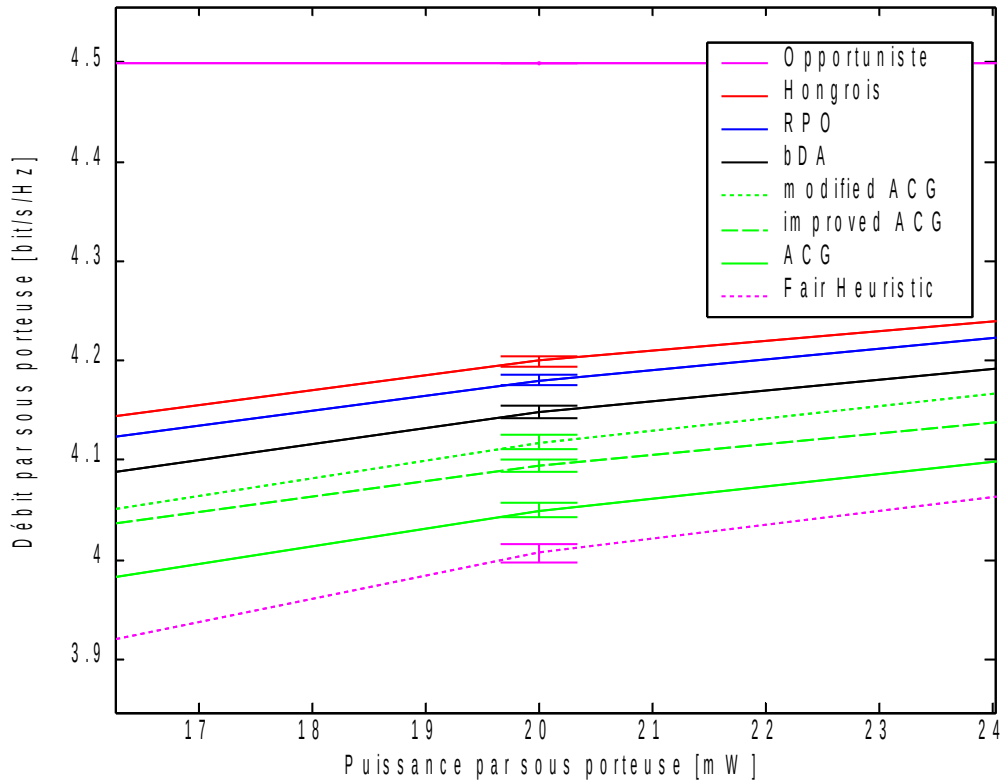


Figure 3.3 : Comparaison de débit entre algorithmes d'affectation de sous porteuses, zoom et intervalles de confiance

4.4. Facteur d'équité

Dans le cadre de l'équité en débit minimal, on définit un facteur d'équité noté F qui est le rapport entre le débit le plus faible et le débit le plus élevé rencontrés dans la cellule : $F = \min_u r_u / \max_u r_u$. Une équité totale en débit est caractérisée par un facteur F de 1.

Sur la figure 3.4, l'équité de l'heuristique de [RheeCioffi_00] (stratégie «max-min») est confirmée par un facteur de 0.9 tandis que «l'injustice» de l'algorithme opportuniste est soulignée par un facteur F nul. L'algorithme «modified ACG» traite les mauvaises sous porteuses en premier afin qu'elles ne soient pas allouées à leur pire utilisateur, cet ordre de traitement permet à l'algorithme de réaliser un facteur d'équité très proche de l'algorithme Hongrois. Le RPO s'assure qu'un utilisateur ne perde pas sa meilleure sous porteuse lorsque sa seconde meilleure sous porteuse est mauvaise. Cette règle assure au RPO un facteur F proche de ceux de l'algorithme Hongrois et du «modified ACG». Sur la figure 3.5, seuls les intervalles de confiance des algorithmes RPO et «modified ACG» se chevauchent. Le RPO réussit à donc concurrencer «modified ACG» du point de l'équité.

Le facteur d'équité du RPO dépasse ceux des algorithmes bDA, ACG et «improved ACG».

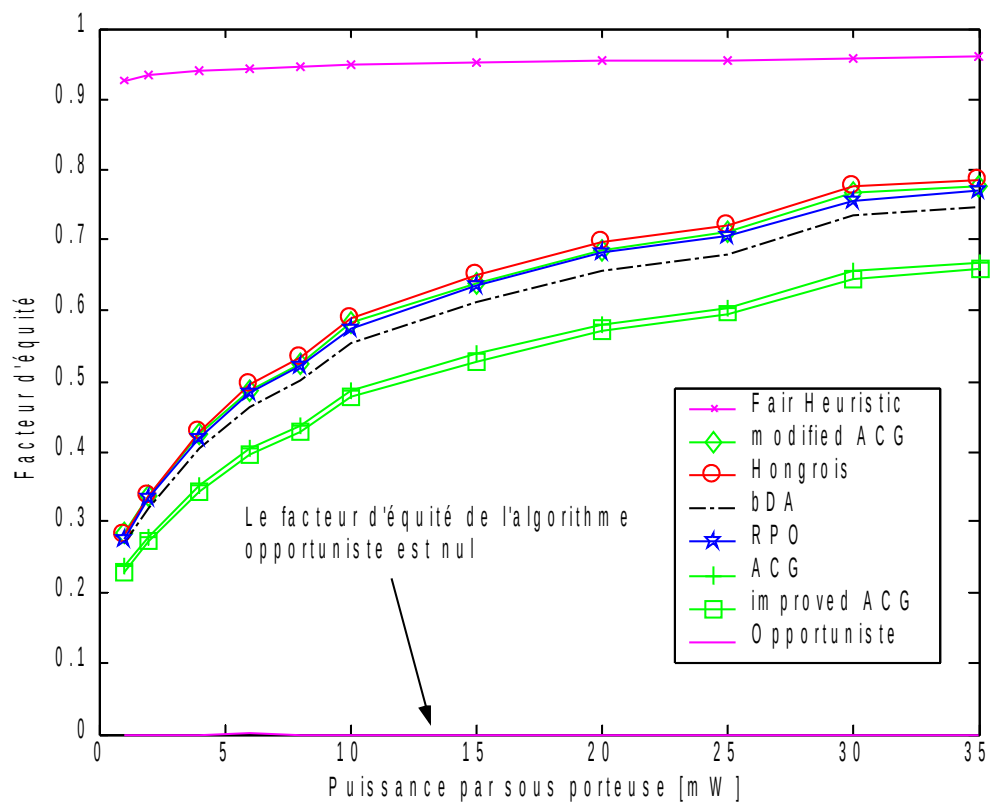


Figure 3.4 : Comparaison d'un facteur d'équité entre algorithmes d'affectation de sous porteuses

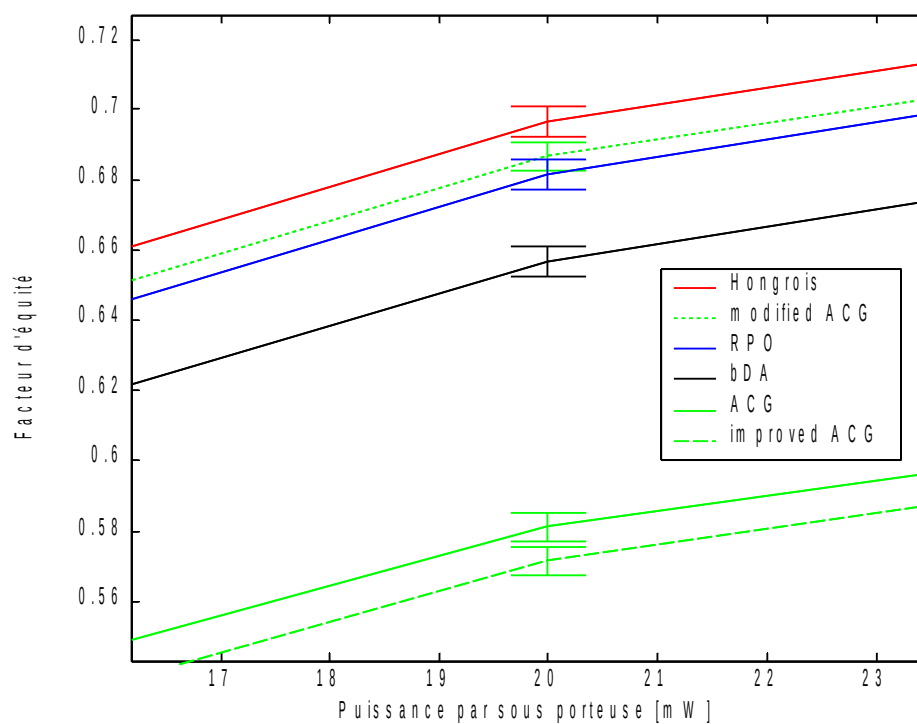


Figure 3.5 : Comparaison d'un facteur d'équité entre algorithmes d'affectation de sous porteuses, zoom et intervalles de confiance

4.5. Probabilité d'échec (*outage*)

Pendant la résolution du méta problème RA, nous avons cherché un débit minimal commun réalisable. Le terme *réalisable* signifie que la probabilité qu'un utilisateur n'atteigne pas ce débit après l'affectation de sous porteuses doit être faible.

La figure 3.6 donne l'évolution, en fonction de la puissance, de ce débit minimal commun r^o . Ce débit minimal est exprimé en nombre de bits par temps symbole OFDM. On s'intéresse à un débit minimal dit "long terme" et un débit minimal dit "court terme". Le débit minimal court terme est moyenné sur plusieurs *snapshots* : la distance et l'effet de masque sont constants, seul l'évanouissement rapide varie. Le débit minimal long terme est moyenné sur plusieurs *snapshots* où les trois paramètres varient : distance, effet de masque et évanouissement rapide. Le débit minimal long terme a une grande variance comparée au débit minimal court terme.

On voit que la différence entre le débit minimal long terme et le débit minimal court terme peut atteindre 30 bits par symbole OFDM (soit 300 kbps pour un temps symbole OFDM de 100 μ s).

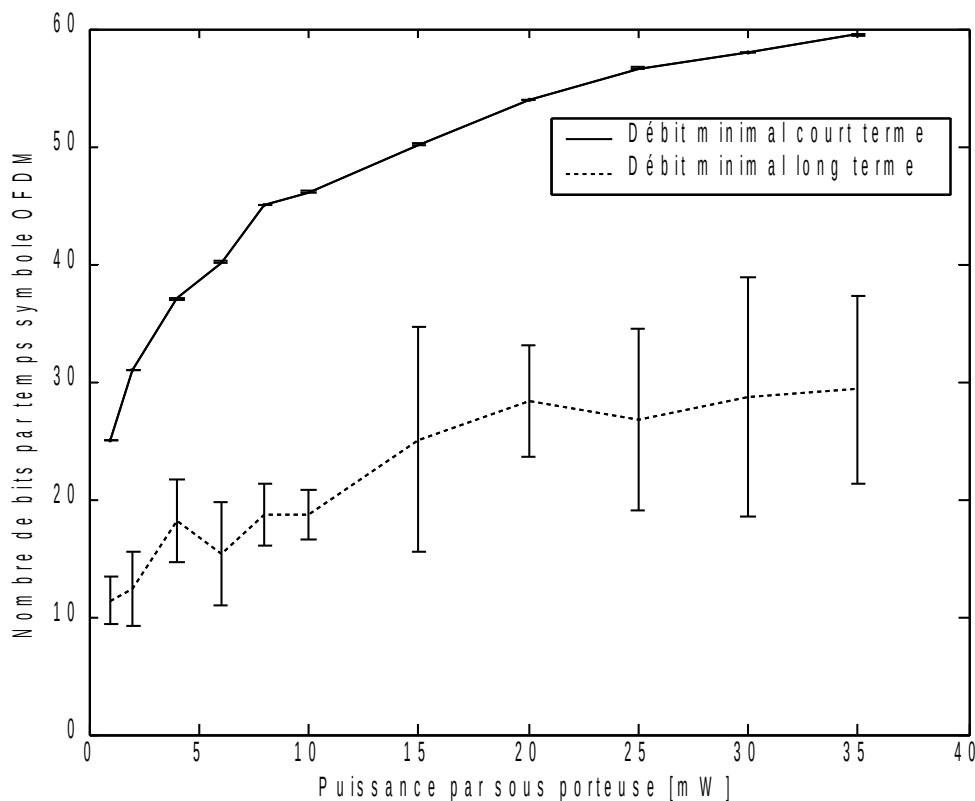
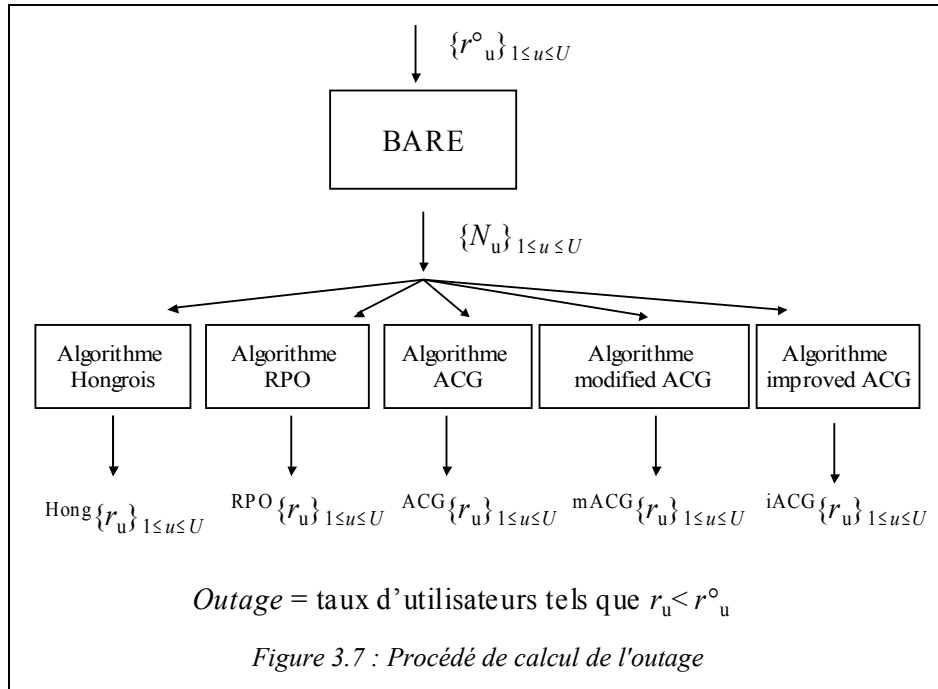


Figure 3.6 : Débit minimal commun (après réduction par décrément) en fonction de la puissance

Dans la suite, on veut évaluer la probabilité d'échec (ou *outage*) c'est à dire la probabilité qu'un utilisateur n'atteigne pas le débit minimal commun r^o long terme. La probabilité d'échec est calculée de façon équitable pour tous les algorithmes (cf. Fig. 3.7). Les algorithmes d'affectation de sous porteuses (tâche 2) reçoivent en entrée une même répartition $\{N_u\}_{1 \leq u \leq U}$. En sortie, on calcule le taux d'utilisateurs qui n'atteignent pas leur débit minimal.



Remarque: La répartition calculée par le BARE (plus adapté que le BABS en RA) est indépendante de l'algorithme utilisé pour la tâche 2. Rappelons que les hypothèses utilisées par le BARE sont simplement une puissance égale sur les sous porteuses et un canal plat ; le canal est caractérisé pour chaque utilisateur par un CgNR moyen (cette dernière hypothèse est également faite dans le BABS). Ces hypothèses ne favorisent aucun algorithme par rapport à un autre.

Pour la figure 3.8, les résultats ont été moyennés sur 20000 simulations pour chaque valeur de la puissance p . Aucun *outage* n'a été recensé pour l'algorithme Hongrois; c'est pourquoi il n'apparaît pas en Fig. 3.8. La probabilité d'échec de l'algorithme Hongrois est donc inférieure à 10^{-6} . Un faible *outage* de l'algorithme Hongrois valide les répartitions $\{N_u\}_{1 \leq u \leq U}$ obtenues par le BARE. L'algorithme Hongrois donne en effet la solution optimale correspondant à une répartition $\{N_u\}_{1 \leq u \leq U}$ fixée. Lorsque l'*outage* de l'algorithme Hongrois est élevé, on peut en déduire que la répartition $\{N_u\}_{1 \leq u \leq U}$ d'entrée ne permet pas de garantir les débits minimaux de manière satisfaisante.

Sur un total de 200 000 itérations, seuls deux cas d'*outage* ont été recensés pour l'algorithme RPO ; ceci donne un ordre de grandeur de 10^{-5} pour la probabilité d'échec de l'algorithme RPO. Il est meilleur que les algorithmes «modified ACG» (dont la probabilité d'échec est de 10^{-4}), bDA (dont la probabilité d'échec est de 10^{-3}) ainsi que les algorithmes ACG et «improved ACG» (dont la probabilité d'échec est de 10^{-2}).

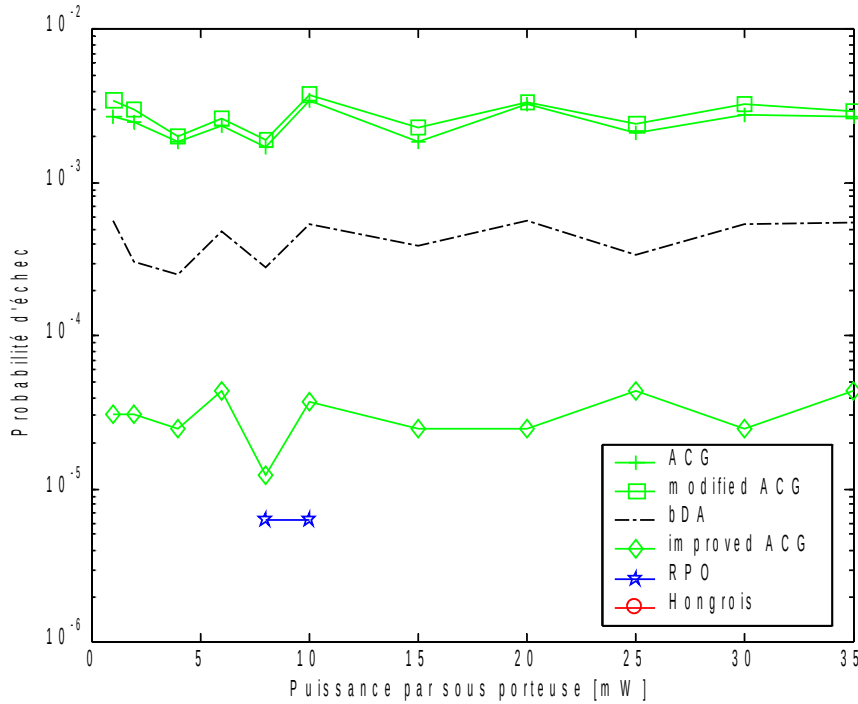


Figure 3.8 : Comparaison du taux d'insatisfaction des utilisateurs entre les algorithmes d'affectation de sous porteuses

4.6. Évaluation de la complexité

La complexité des algorithmes est un paramètre important pour analyser leur «faisabilité». Le temps de cohérence du canal est très court (de l'ordre de la milliseconde), l'allocation doit donc être assez rapide.

L'algorithme Hongrois est donné en $\mathcal{O}(N^3)$ dans [PapaSteig_82]. L'implémentation de l'algorithme n'est pas envisageable lorsque N augmente. C'est pourquoi, des algorithmes sous optimaux sont proposés dans la littérature. La complexité de l'ACG est donnée en $\mathcal{O}(N \times U)$ par les auteurs de [Kivanc&_00]. La complexité du bDA est donnée en $\mathcal{O}(N \times U \log(N))$ par les auteurs de [Gross&_03]. Les auteurs de [Zhen&_03] ne donnent pas la complexité de l'algorithme «improved ACG» mais on peut l'évaluer. A chaque itération, U utilisateurs choisissent leur meilleure sous porteuse parmi N puis on choisit le meilleur utilisateur parmi U . Chaque itération est donc en $\mathcal{O}(N+U)$. Il y a N itérations donc la complexité de l'«improvedACG» est en $\mathcal{O}(N^2)$.

Dans le RPO, U utilisateurs choisissent leur meilleure sous porteuse parmi N et en cas de conflit, ils doivent recommencer pour déterminer leur seconde meilleure sous porteuse. Il y a entre N/U et N itérations. Donc dans le pire cas, le RPO est en $\mathcal{O}(N^2)$.

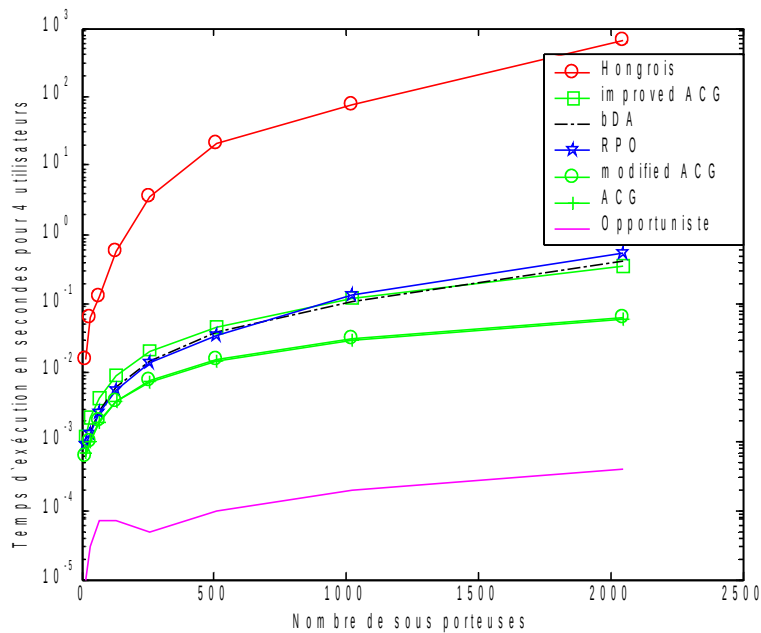


Figure 3.9 : Comparaison du temps d'exécution des algorithmes d'affectation de sous porteuses

La figure 3.9 trace le temps d'exécution (en secondes) des différents algorithmes lorsqu'ils affectent N sous porteuses à quatre utilisateurs; N varie entre 16 et 2048. L'algorithme opportuniste est le plus rapide car le plus simple. Les algorithmes heuristiques sont jusqu'à mille fois plus rapide que l'algorithme Hongrois. Les algorithmes ACG et «modified ACG» ont des temps d'exécution les plus faibles au dépend des performances en débit. Un peu plus lents, les algorithmes bDA, RPO et «improved ACG» ont des temps d'exécution similaires. On voit que jusqu'à $N = 1000$, le RPO est légèrement plus rapide que l'«improved ACG» puis la situation s'inverse.

Pour des temps d'exécution comparables, le RPO obtient de meilleures performances que l'«improved ACG» et le bDA en terme de débit, de facteur d'équité et de probabilité d'échec.

5. Conclusions

Nous avons proposé des algorithmes pour la résolution du problème RA. Nous avons utilisé lors de cette résolution :

- les tâches 1 et 2 lors de l'attribution des sous porteuses,
- une allocation uniforme de puissance sur les sous porteuses car une étude montre que les performances rivalisent avec le *waterfilling* ([JangLee_03]).

Concernant le calcul du nombre de sous porteuses par utilisateur (allocation de largeur de bande, tâche 1), nous avons proposé le BARE (*Bandwidth Allocation based on Rate Estimation*) qui, comme nous l'avons montré, est plus efficace dans le contexte RA que le BABS (*Bandwidth Assignment Based on SNR*). Au cours de cette tâche, nous avons

cherché à assurer une équité en débit minimal. Néanmoins, le BARE est également adapté à une contrainte quelconque sur le vecteur de débit.

Concernant l'affectation spécifique de sous porteuses (tâche 2), nous proposons le RPO (*Rate Profit Optimization*). Son but est d'éviter l'allocation «aveugle» d'une sous porteuse à son meilleur utilisateur. Lorsque plusieurs utilisateurs ont une meilleure sous porteuse commune, la seconde meilleure sous porteuse de chacun est évaluée. Ainsi le débit global n'est pas pénalisé par une affectation systématique au meilleur utilisateur de la sous porteuse considérée. Pour comparer les performances des heuristiques, on utilise l'algorithme Hongrois qui est l'algorithme optimal pour la tâche 2.

D'après les simulations, le RPO réalise l'efficacité spectrale la plus proche de l'algorithme Hongrois. Le facteur d'équité du RPO est de même très proche de celui de l'algorithme Hongrois. La probabilité d'échec du RPO, c'est-à-dire la probabilité qu'un utilisateur n'atteigne pas le débit minimal, est inférieure à celle de toutes les heuristiques rencontrées. Enfin, la complexité du RPO est comparable aux heuristiques existantes. Ainsi, pour une complexité comparable, le RPO est meilleur que les heuristiques existantes au regard de l'efficacité spectrale, du facteur d'équité et de la probabilité d'échec de la tâche 2 (par rapport aux objectifs définis par la tâche 1).

Dans les systèmes, il peut être difficile d'affecter les sous porteuses une par une à cause notamment des problèmes de signalisation. En pratique, la sous canalisation (c'est à dire le regroupement des sous porteuses) est souvent considérée ([IEEE 802.16]). Dans le chapitre suivant, nous comparons ces deux approches : l'approche théorique de l'affectation individuelle et l'approche pratique de la sous canalisation.

Chapitre 4. Modes de sous canalisation en WiMAX (IEEE 802.16)

1. Introduction

Dans ce chapitre, nous souhaitons comparer les algorithmes d'affectation de sous porteuses proposés dans la littérature avec les modes de sous canalisation utilisés en pratique.

1.1. Algorithmes d'affectation de sous porteuses et conditions pratiques

Dans le chapitre 3, nous avons présenté des algorithmes d'affectation de sous porteuses et nous avons évalué leur performance en terme de débit par sous porteuse. Dans ce chapitre, nous avons travaillé dans un contexte idéal sur trois points. Premièrement, la connaissance du canal est parfaite. Deuxièmement, les algorithmes traités affectent les sous porteuses individuellement, cela suppose une signalisation par sous porteuse. Enfin, l'adaptation de lien est réalisée sous porteuse par sous porteuse : chaque sous porteuse transporte un MCS adapté à son propre SNR.

Grâce à ces trois hypothèses idéales, les algorithmes présentés dans le chapitre 3 donnent une borne supérieure de la performance en débit global.

La sous canalisation permet de réduire la signalisation nécessaire car elle regroupe un nombre fixe de sous porteuses pour former un sous canal. Les sous porteuses peuvent être soit adjacentes, soit issues de toute la bande fréquentielle. Dans le premier cas, on peut mettre en place des stratégies opportunistes d'allocation de sous canaux pour profiter d'une diversité multiutilisateur. Dans le deuxième cas, on peut attribuer un sous canal à n'importe quel utilisateur car le niveau des sous canaux est supposé uniforme. Cela permet de décorréler les sous canaux par rapport au *fading* sélectif en fréquence.

La signalisation est utilisée pour l'adaptation de lien. Cette dernière est réalisée par sous canal : les sous porteuses d'un même sous canal transportent un MCS commun.

La sous canalisation et l'adaptation de lien par sous canal sont utilisés dans le réseau IEEE 802.16 pour résoudre des contraintes évoquées dans la suite. Il est intéressant d'étudier l'impact de ces solutions pratiques sur le débit global.

1.2. Objectifs du chapitre

Dans ce chapitre, un premier objectif est de comprendre les principaux modes de sous canalisation proposés dans la norme IEEE 802.16.

Le standard définit des modes dits de diversité : un sous canal est composé de sous porteuses issues de toute la bande fréquentielle. Les modes de diversité sont définis pour utiliser la bande uniformément et moyenner les interférences entre les cellules.

Un deuxième objectif est d'établir, dans les modes de diversité, le taux de collision entre les sous canaux pour évaluer les interférences entre deux cellules.

Un dernier objectif du chapitre est de comparer les performances obtenues en débit global par les modes de sous canalisation avec les performances idéales obtenues dans le chapitre 3. Cette comparaison est réalisée dans un contexte multicellulaire où les fréquences sont réutilisées dans toutes les cellules (le facteur de réutilisation fréquentiel vaut 1). Dans ce chapitre, pour réduire la durée des simulations, nous utilisons le RPO comme borne supérieure sur le débit global (on a vu que c'est le meilleur algorithme après l'algorithme Hongrois). Pour comparer les performances des modes de sous canalisation et du RPO, nous utilisons un nombre de fixe de sous porteuses par utilisateur pour l'allocation de bande (tâche 1 de l'étape 1 définie dans le chapitre 2) et une puissance uniforme sur les sous porteuses (étape 2 définie dans le chapitre 2).

2. Le standard IEEE 802.16

2.1. Généralités

Le standard IEEE 802.16 2004 définit des réseaux métropolitains sans fils (MAN : *Metropolitan Area Network*) dans le cas d'une vision directe (LOS : *Line Of Sight*) dans la bande des 10-60 GHz et dans le cas d'une vision non directe (NLOS : *Non Line Of Sight*) dans la bande des 2-11GHz. Le terme Wimax (*Worldwide Interoperability for Microwave Access*) désigne le label commercial délivré par le Wimax Forum aux équipements conformes à la norme 802.16 afin de garantir une inter-opérabilité entre les équipements de différents constructeurs.

Le réseau défini par le standard 802.16 peut, dans les zones peu denses, servir d'alternative aux technologies filaires classiques (lignes xDSL, câble, liaisons spécialisées...) au niveau du dernier kilomètre ([CCM_Wimax]). Une autre application consiste à utiliser un tel réseau pour relier des réseaux locaux sans fils entre eux ([PaD_Wimax]). L'architecture du réseau est semblable à une architecture cellulaire classique : des stations de base utilise une architecture PMP (point multipoint) pour servir les utilisateurs.

Un amendement au standard définit les caractéristiques nécessaires pour gérer la mobilité des utilisateurs : il s'agit de l'IEEE 802.16e. Les systèmes basés sur le standard 802.16 2004 sont couramment appelés «Wimax fixe» (*Fixed WiMAX*) et ceux qui sont basés sur le standard 802.16e sont appelés «Wimax mobile» (*Mobile Wimax*).

L'OFDMA est un mode d'accès obligatoire dans le 802.16e tandis qu'il est optionnel dans le 802.16 2004. En 802.16, les modulations possibles sont les modulations à monoporteuse et l'OFDM à 256 sous porteuses. L'OFDMA à 2048 sous porteuses est optionnelle. En 802.16e, la bande peut varier de 1.25 MHz à 20 MHz. La taille N_{FFT} de la FFT (*Fast Fourier Transform*) varie de manière à fixer l'espace inter-fréquence $\Delta f = 10.93$ kHz ([WimaxEval_07], p19). C'est le principe du SOFDMA : *Scalable OFDMA*.

Le standard supporte les modes duplex TDD (*Time Division Duplexing*) et FDD (*Frequency Division Duplexing*) et *half duplex* FDD. La *release* actuelle de certification des profils de 802.16e ne considère que le mode TDD.

Le mode TDD permet de supporter des terminaux *half duplex*. D'autre part, il permet de gérer une asymétrie des trafics des voies montante et descendante. Enfin, pendant les procédés d'adaptation de lien, on peut faire l'hypothèse de la réciprocité du canal.

2.2. Quelques définitions

Chaque standard comporte sa propre terminologie, nous rappelons ici les principales définitions en 802.16. Attention car les significations usuelles de certains termes sont assez différentes.

Un **mode de sous canalisation** (appelé aussi mode de permutation) est un ensemble de règles définissant la construction d'un sous canal. Nous étudierons dans la section 3, les principaux modes de sous canalisation.

Un **slot** est une unité d'allocation minimal : un utilisateur reçoit au minimum un *slot*. Un *slot* est une parcelle ou un ensemble de parcelles dans le plan temps / fréquence. La constitution du *slot* dépend du mode de sous canalisation. Les différentes possibilités pour un *slot* seront récapitulées en tableau 4.4.

Un **burst** est un ensemble de *slots*.

Une **région** est un groupe de sous canaux consécutifs (par leurs numéros logiques) sur un nombre de symboles OFDM consécutifs.

Un **segment** est une sous division de l'ensemble des sous canaux disponibles. Lors d'un facteur de réutilisation fractionnel, les secteurs d'une cellule reçoivent des segments différents.

Une **zone de permutation** est un ensemble de symboles OFDMA qui utilisent le même mode de sous canalisation.

2.3. Structure de la trame TDD

Une trame montante (cf. Fig. 4.1) est séparée de la trame descendante précédente par un intervalle de temps appelé TTG (*Transmit Transition Gap*) ; elle est séparée de la trame descendante suivante par le RTG (*Receive Transition Gap*). La trame descendante commence par un préambule, un FCH (*Frame Control Header*), une DL-MAP (*downlink map*) et une UL-MAP (*uplink map*).

Le FCH spécifie le profil de burst et la longueur des premiers messages : DL-MAP, UL-MAP, UCD (*Uplink Channel Descriptor*), DCD (*Downlink Channel Descriptor*). On appelle profil de burst l'association d'une modulation et d'un taux de codage correcteur d'erreur.

La DL-MAP indique les transitions physiques dans la trame descendante courante c'est à dire les changements de profils de burst. Elle indique aussi les changements de zones de permutation (PUSC, FUSC, AMC, AAS *Advanced Antenna Systems etc*), cf. Fig. 4.2. Une zone de permutation est caractérisée par un paramètre DL_PermBase compris entre 0 et 31 et indiqué par la DL_MAP. Les trames montantes et descendantes commencent toujours par une zone PUSC ([Yaghoobi_04]) ; dans cette zone dite

zone 1, le paramètre DL_PermBase est égal au paramètre IDCell¹ de la cellule (connu dans le préambule, [IEEE 802.16e] p535).

La UL-MAP indique les allocations de ressources pour la prochaine trame montante. Cette carte contient les codes UIUC (*Uplink Interval Usage Code*) désignant les profils de burst avec lesquels les utilisateurs transmettront.

Si des messages UCD et DCD sont envoyés, ils doivent immédiatement suivre les messages UL-MAP et DL-MAP. Ces messages associent chaque profil de burst à un code DIUC (*Downlink Interval Usage Code*) et UIUC (*Uplink Interval Usage Code*).

La trame montante commence par une région de *ranging* et une région d'acquiescement et d'information canal (CQI, *channel quality indicator*). Les utilisateurs envoient, durant les *slots* indiqués par la UL-MAP, le CINR (*Channel Interference to Noise Ratio*) du préambule de la trame descendante pour les modes de diversité (PUSC, FUSC) ou les CINR des cinq meilleures bandes (définies en 3.1) en AMC. La procédure de *ranging* est nécessaire pour toutes les futures transmissions d'un terminal : il obtient des informations sur sa puissance d'émission et l'avance en temps.

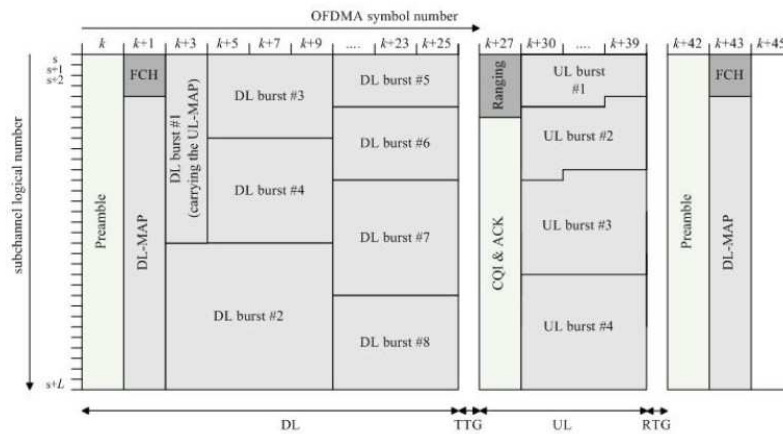


Figure 4.1 : Structure de la trame TDD, [IEEE 802.16]

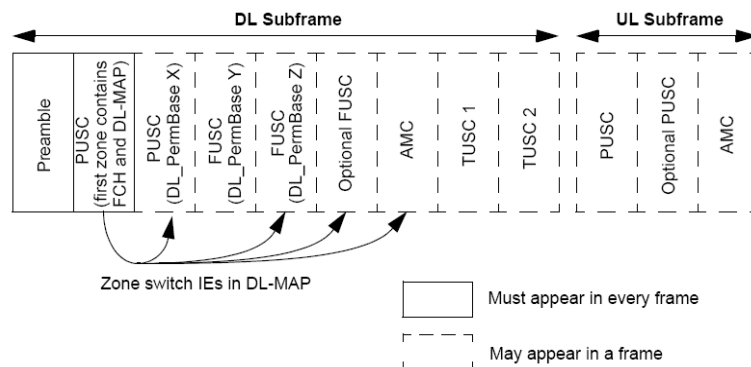


Figure 4.2 : Subdivision de la trame TDD en zones, [IEEE 802.16e]

¹ IDCell vaut entre 0 et 31 d'après [WimaxEval_07] p 122

3. Modes de sous canalisation en IEEE 802.16

Dans l'introduction, nous avons vu que la sous canalisation et l'adaptation de lien par sous canal sont utilisés dans le réseau IEEE 802.16. On distingue le mode sélectif en fréquence et appelé AMC (*Adaptive Modulation and Coding*) des modes dits de diversité qui utilisent uniformément la bande. Dans les modes de diversité, le mode PUSC (*Partial Usage of SubChannels*) permet, à l'opposé du FUSC (*Full Usage of SubChannels*), d'utiliser un facteur de réutilisation fréquentiel différent de 1. Pour faciliter les mesures physiques, des structures de sous porteuses adjacentes sont définies (tels que le *tile* en PUSC voie montante et le *bin* en AMC).

3.1. Mode sélectif en fréquence : l'AMC

Il n'est pas vraiment envisageable d'affecter les sous porteuses individuellement, comme dans le chapitre 3, pour des raisons de signalisation. Les algorithmes décrits dans le chapitre 3 tentent de profiter de la diversité multiutilisateur. Le mode AMC est le mode qui s'en rapproche le plus dans la norme.

Les sous canaux AMC sont constitués de sous porteuses consécutives. En faisant des hypothèses sur la bande de cohérence, on peut supposer que les sous porteuses d'un sous canal ont une qualité homogène. On peut ainsi allouer un sous canal à l'utilisateur qui y présente le meilleur CgNR moyen.

Un sous canal est construit de la même façon en voie montante et descendante. On définit les structures de *bin* et de bande. Un *bin* est un ensemble de 9 sous porteuses consécutives sur l'axe des fréquences et une bande est constituée de 4 *bins* consécutifs.

Un *slot* AMC s'étale sur J temps symboles consécutifs et comprend I *bins* voisins. Les entiers I et J doivent vérifier $I \times J = 6$. Les différentes «factorisations» de 6 sont illustrées dans la figure 4.3 ; la combinaison 6×1 semble ne pas en faire parti ([IEEE 802.16e], p 353). Dans un *bin*, il y a toujours une fréquence pilote. Un *slot* comporte donc 48 ($6 \times (9-1)$) tons de données (*data subcarriers*). Nous introduisons la définition d'un ton comme étant une fréquence sur un temps symbole : (f_1, t_1) et (f_1, t_2) constituent deux tons différents.

Dans le sens descendant, la BS décide quel mode est le mieux adapté au terminal. Un terminal reçoit des sous canaux AMC si pendant un nombre¹ fixé de trames TDMA :

- la moyenne des mesures de CINR sur les cinq meilleures bandes est supérieure au seuil «*band AMC entry average CINR*»,
- l'écart type du CINR est inférieur au seuil «*band AMC Alloc threshold*».

¹ Ce nombre est indiqué par le message UCD, cf 2.3

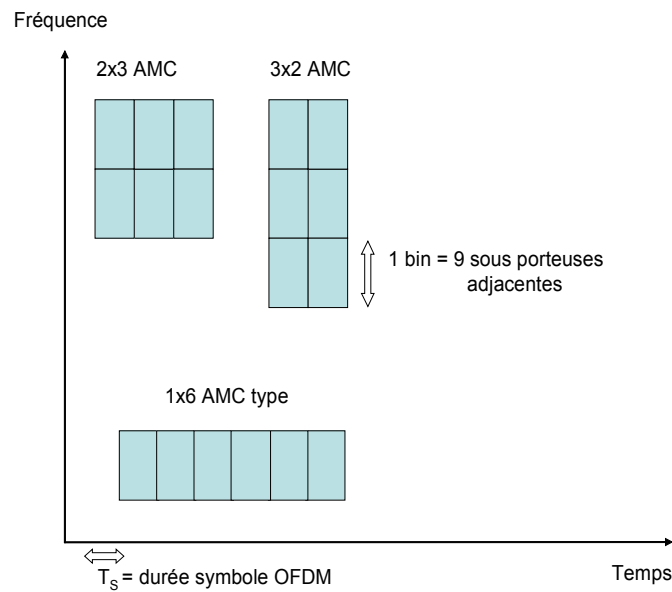


Figure 4.3 : Types de sous canaux AMC

3.2. Modes de diversité : description et caractérisation des collisions

3.2.1. FUSC : Full Usage of Subchannels

Ce mode est utilisé uniquement en voie descendante. Nous nous intéressons particulièrement à ce mode dans nos simulations. Un sous canal comporte $N_{scPerChan} = 48$ fréquences de données. Ce nombre ne change pas avec la bande du système SOFDMA ; seul change le nombre total de sous canaux appelé N_{CHAN} (cf. tableau 4.1). Un *slot* FUSC est composé d'un sous canal sur un temps symbole OFDM.

Tableau 4.1 : Nombre de sous canaux en FUSC en fonction de la bande totale

W (bande en MHz)	1.25	5	10	20
N_{FFT} (taille FFT)	128	512	1024	2048
N_{data} (nombre total de sous porteuses de données) en FUSC	96	384	768	1536
N_{CHAN} (nombre total de sous canaux)	2	8	16	32

a) Construction des sous canaux

On considère l'ensemble de la bande fréquentielle. Premièrement, les fréquences pilotes sont désignées ; elles seront partagées par l'ensemble des sous canaux.

Les fréquences restantes sont renumérotées et divisées en groupes de sous porteuses consécutives. Pour constituer un sous canal, on prend une sous porteuse appartenant à chaque groupe. Ainsi, N_{CHAN} est à la fois le nombre de sous canaux obtenus et le nombre de sous porteuses consécutives par groupe. Le nombre $N_{scPerChan}$ est à la fois le nombre de

sous porteuses par sous canal (invariable) et le nombre de groupes. Si par exemple, on se place dans une bande de 20 MHz, il y a 1536 sous porteuses de données réparties en 48 ($N_{scPerChan}$) groupes de 32 (N_{CHAN}) sous porteuses.

La façon de choisir une sous porteuse dans chacun des 48 groupes pour former un sous canal est régie par une formule dite de *permutation*. Elle détermine l'index global de la sous porteuse à partir de l'index du sous canal (entre 0 et $N_{CHAN}-1$), de l'index de la sous porteuse dans le sous canal (entre 0 et $N_{scPerChan}-1$) et enfin d'un paramètre ID_{Cell} caractérisant la cellule (entre 0 et 31) ; la formule est décrite dans l'annexe 4.A.

Grâce à la formule de permutation, les sous porteuses d'un sous canal ne sont pas consécutives ; elles sont distribuées sur toute la bande disponible comme l'illustre la figure 4.4. Sur cette figure, l'axe des ordonnées représente l'index absolu des sous porteuses de données (il varie entre 0 et 1535), l'axe des abscisses représente l'index relatif d'une sous porteuse de données dans le sous canal auquel elle appartient (il varie entre 0 et 47). Les traits rouges délimitent chacun des 48 groupes de 32 sous porteuses consécutives sur la bande disponible. Pour un paramètre ID_{Cell} fixe, on marque les 48 sous porteuses constituant les sous canaux d'index 0 et 1 avec respectivement des croix et des ronds. On voit qu'un sous canal n'est constitué que d'une sous porteuse par groupe (il n'y a qu'un point et/ou une croix entre deux traits rouges).

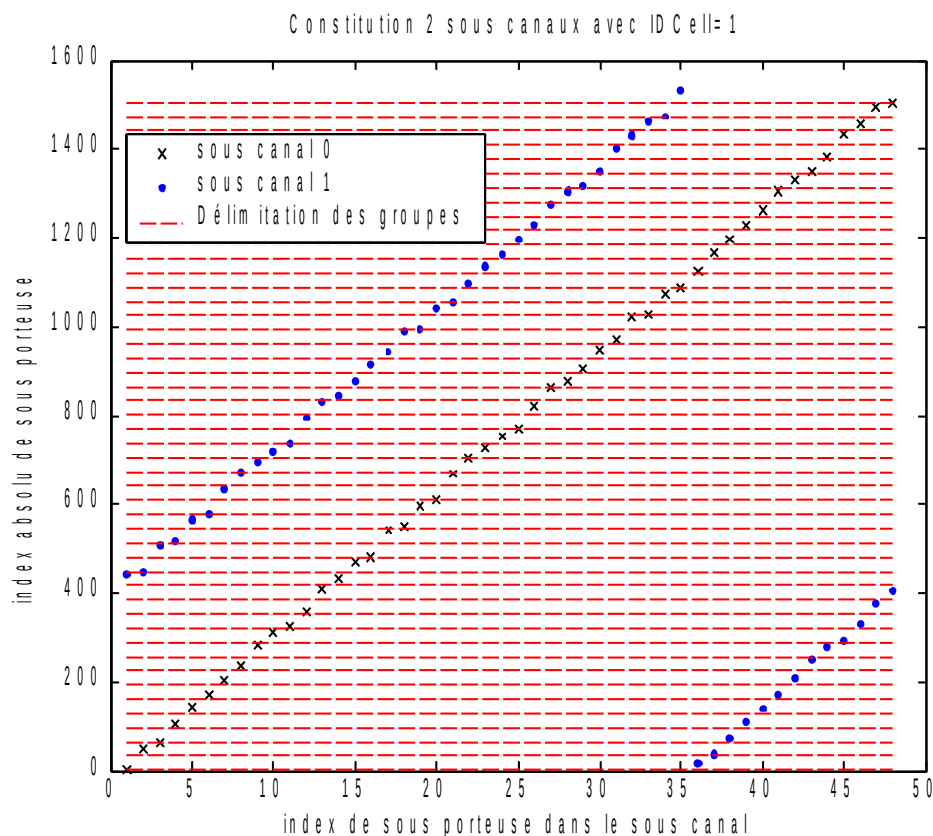


Figure 4.4 : Illustration de la sous canalisation FUSC dans une cellule

Mode FUSC et contexte multicellulaire

Les figures 4.5. et 4.6 illustrent la formule de permutation pour des sous canaux de même index dans des cellules différentes (paramètre ID_{Cell} distincts). On peut remarquer que quelque soit l' ID_{Cell} , la $k^{ième}$ sous porteuse d'un sous canal d'un même index est choisie dans le même groupe (cf. Fig. 4.5). Si chaque cellule utilise un seul sous canal de même index, on observe qu'il n'y a pas de collisions (cf. Fig. 4.6) grâce à la formule de permutation.

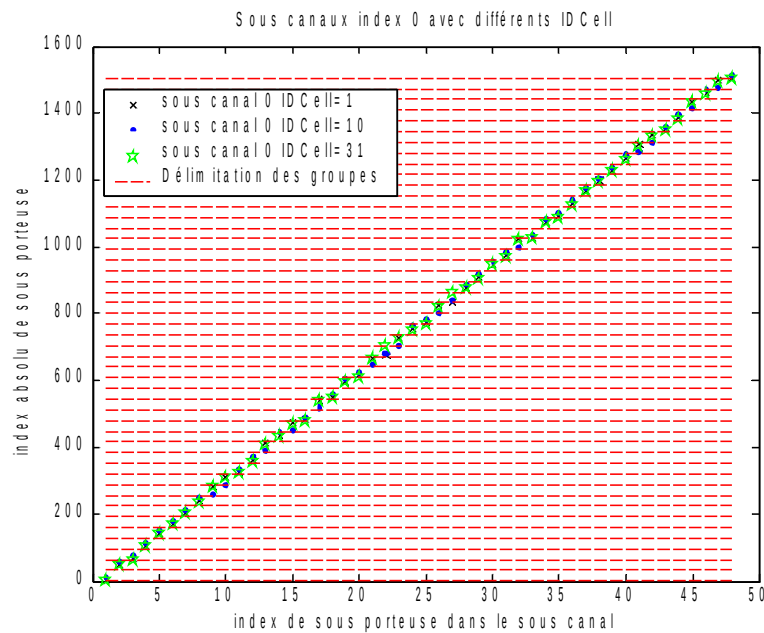


Figure 4.5 : Illustration de la sous canalisation FUSC dans différentes cellules

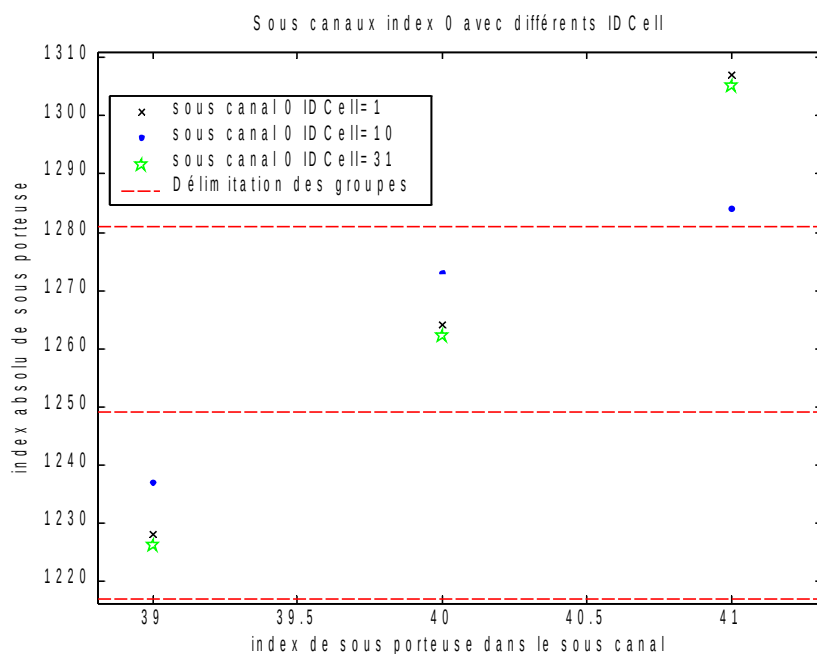


Figure 4.6 : Illustration de la sous canalisation FUSC dans différentes cellules, zoom

b) Analyse des collisions

On veut maintenant déterminer le nombre de collisions entre les sous canaux de deux cellules distinctes. On considère que chaque cellule utilise le même nombre L de sous canaux. On définit alors le facteur de charge L / N_{CHAN} . Pour l'illustration, on se place dans le cas $W = 10$ MHz, alors $N_{CHAN} = 16$.

Facteur de charge de 1/16 ($= L / N_{CHAN}$)

On s'intéresse à la densité de probabilité du nombre de collisions en FUSC entre deux sous canaux utilisés dans deux cellules différentes.

On la compare à une sous canalisation aléatoire, où 48 fréquences sont choisies aléatoirement parmi les N_{data} sous porteuses de données. Soit χ la variable aléatoire qui compte le nombre de collisions entre deux sous canaux. La probabilité d'avoir c collisions entre deux sous canaux vaut dans ce cas $p_{\chi}(c) = \frac{C_{N_{data}-48}^{48-c} C_{48}^c}{C_{N_{data}}^{48}}$. Un résultat

similaire est développé dans [Lerbour&_06].

La figure 4.7 compare les densités de probabilité du nombre de collisions quand chaque cellule utilise $L = 1$ sous canal sur 16. Pour obtenir la densité de probabilité en FUSC, on traite de façon exhaustive¹ tous les couples d'*IDCell*, le nombre de collisions sur tous les couples d'index de sous canaux est déterminé. On voit sur la figure 4.7 que le nombre de collisions en sous canalisation FUSC, pour cette charge de 1/16, ne peut prendre que certaines valeurs : $\{0, 3, 6, 48\}$.

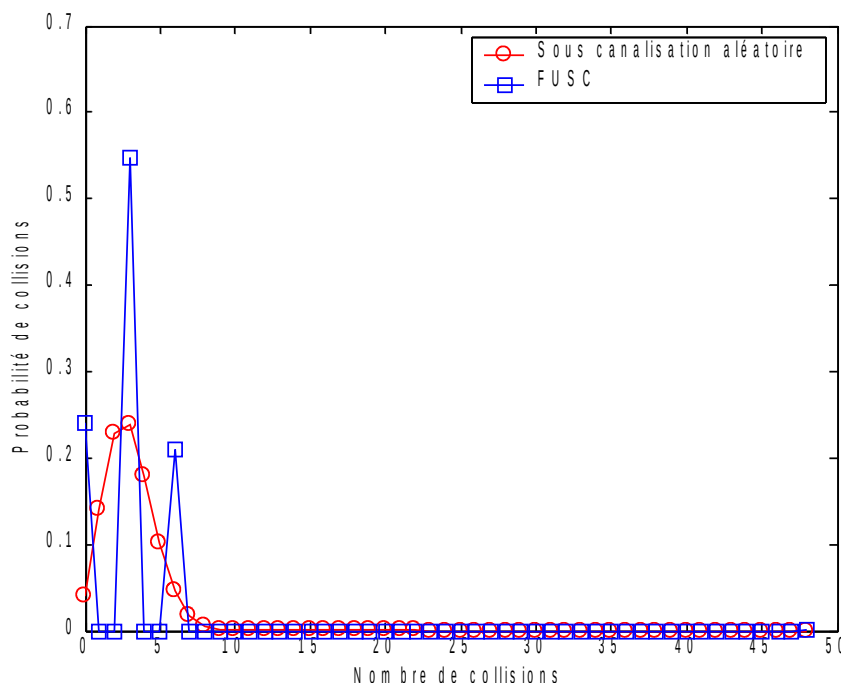


Figure 4.7 : Densité de probabilité de collisions en FUSC, facteur de charge 1/16

¹ Il y a 496 couples à traiter et pour chacun $16 \times 16 = 256$ couples de sous canaux à parcourir.

Facteur de charge de 2/16 ou plus

Pour un facteur de charge de 2/16, chaque cellule utilise deux sous canaux de 48 sous porteuses. Le nombre de collisions varie entre 0 et 2×48 comme on le voit sur la figure 4.8. Pour l'obtenir, on considère la moitié des couples d'*IDCell* (en fait 248) ; pour chaque couple, on traite la totalité des quadruplets (distincts) d'index de sous canaux (c'est à dire 120).

On voit que le nombre de collisions en sous canalisation FUSC ne peut prendre que certaines valeurs : $\{0, 3, 6, 9, 12, 15, 18, 21, 24, 48, 96\}$ pour la charge de 2/16. La densité de probabilité du nombre de collisions en sous canalisation aléatoire n'a pas la même forme. Malgré cette différence de forme, les deux densités présentent le même nombre moyen de collisions.

On a poursuivi cette étude en réalisant le «comptage» des collisions pour des facteurs de charge plus élevés. Pour chaque facteur de charge, on traite une quinzaine de n -uplets d'index de sous canaux pour une centaine de couples d'*IDCell*. On a ainsi obtenu le nombre moyen de collisions en fonction du facteur de charge en FUSC. En divisant ce nombre moyen de collisions par le nombre total de sous porteuses utilisées dans la cellule, on a vérifié que la probabilité de collision sur une sous porteuse est égal au facteur de charge (cf tableau 4.2).

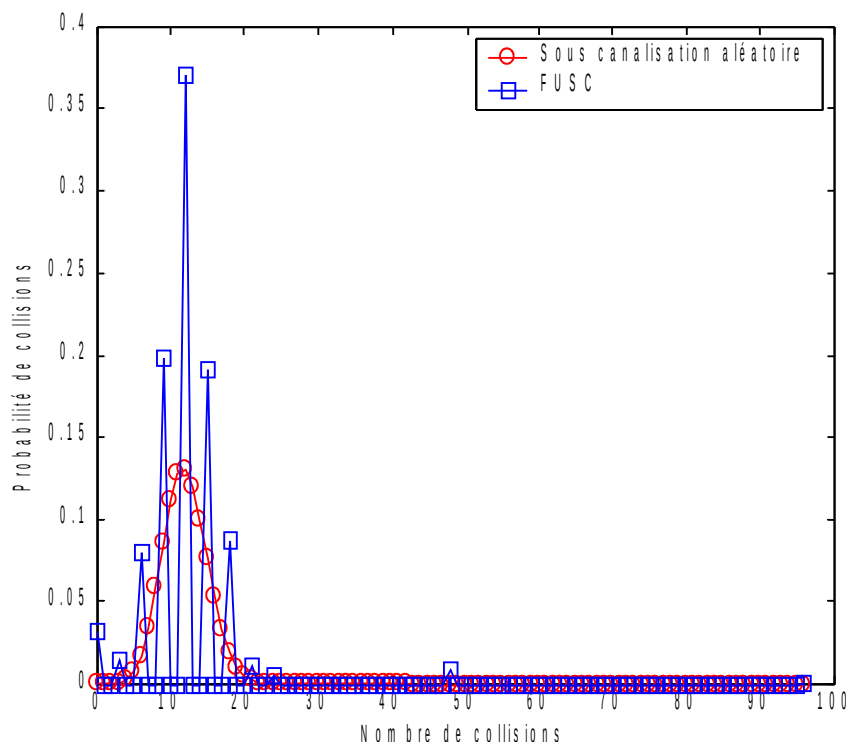


Figure 4.8 : Densité de probabilité de collisions en FUSC, facteur de charge 2/16

Tableau 4.2 : Étude de la probabilité de collision (FUSC, sous canalisation aléatoire)

Nombre de sous canaux utilisés	Nombre de sous porteuses utilisées	Facteur de charge	Nombre moyen de collisions en FUSC	Nombre moyen de collisions en sous canalisation aléatoire	Probabilité de collision pour une sous porteuse en FUSC
1	48	0.062	3	3	0.062
2	96	0.125	11.98	12	0.125
3	144	0.187	26.88	27	0.186
6	288	0.375	108.17	108	0.375
10	480	0.625	299.65	300	0.624
13	624	0.812	504.87	507	0.809

Collisions en PUSC

Dans [WimaxEval_07], une analyse sur la densité de probabilité du nombre de collisions est menée en PUSC. Dans l'annexe 4.B, nous menons une analyse complémentaire sur la comparaison avec la sous canalisation aléatoire. La section 3.2.2 décrit le mode de PUSC.

3.2.2. PUSC : Partial Usage of Subchannels

Ce mode est utilisé en voie montante et en voie descendante. C'est un mode qui garde les mêmes propriétés que le FUSC (vis à vis de l'utilisation uniforme de la bande) et qui permet de varier le facteur de réutilisation fréquentiel. Pour cela, des groupes majeurs sont définis en PUSC et contiennent des fréquences issues de l'ensemble de la bande. Ils pourront être affectés à des secteurs ou cellules différents. Dans le sens descendant, les pilotes sont liés au groupe majeur et dans le sens montant les pilotes sont liés au sous canal. Ci-après on décrit la sous canalisation en voie descendante et en voie montante.

a) Voie descendante

Un *slot* PUSC, unité d'allocation minimale, s'étend sur deux temps symboles OFDM. Dans chaque temps symbole, le sous canal comporte 24 fréquences de données ; cela fait 48 tons¹ de données dans un *slot*.

Principe de construction des sous canaux

En PUSC, on divise la bande en sous ensembles appelés groupes majeurs (*major groups*). Chaque sous canal est construit à l'intérieur d'un seul groupe majeur et partage les fréquences pilotes de ce même groupe majeur. Cela permet d'appliquer un facteur de réutilisation fréquentiel (FRF) différent de 1 en affectant à différents secteurs (ou cellules) des groupes majeurs distincts.

Dans chaque groupe majeur, on détermine les fréquences pilotes ; elles seront partagées par les sous canaux du groupe majeur. Puis, on divise les sous porteuses restantes en 24 groupes de sous porteuses de données consécutives. Enfin, un sous canal est composé

¹ Un ton est une fréquence sur un temps symbole

d'une fréquence de chaque groupe grâce à une formule de permutation similaire à celle utilisée en FUSC. On voit par analogie avec le FUSC, que ce qui était réalisé sur l'ensemble de la bande est fait en PUSC, à l'intérieur de chaque groupe majeur. On peut noter que les noms des modes laissaient présager ce phénomène : *Full Usage* (de FUSC) signifie usage total de la bande et *Partial Usage* (de PUSC) signifie usage partiel de la bande.

Détails de construction des groupes majeurs

Nous avons vu le principe de fonctionnement global du PUSC, mais pour le lecteur curieux de la terminologie utilisée dans la norme, nous pouvons expliquer les principaux points de détails illustrés par la figure 4.10.

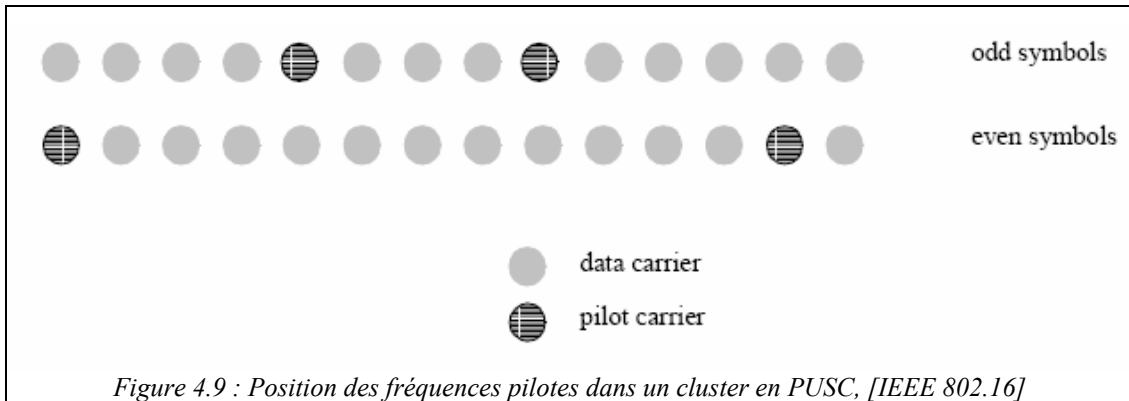
Le mode PUSC définit des *clusters* qui sont des ensembles de 14 sous porteuses adjacentes. Le nombre de *clusters* $N_{clusters}$ varie avec la taille de la bande (cf. tableau 4.3). Les *clusters* ont initialement un numéro physique. Après une renumérotation appelée *outer permutation*, chaque cluster possède un numéro logique. Les groupes majeurs sont composés d'un certain nombre de *clusters* identifiés par leur numéro logique. Notons que grâce à la renumérotation (ou *outer permutation*), un groupe majeur contient des *clusters* issus de toute la bande. Ainsi, un sous canal bien que construit à l'intérieur d'un seul groupe majeur peut contenir des fréquences sur toute la bande fréquentielle initiale.

Les groupes majeurs d'index pairs contiennent plus de *clusters* que les groupes majeurs d'index impairs (cf. tableau 4.3). La conséquence directe est que le nombre de sous canaux obtenu est plus faible dans un groupe majeur impair. Quelque soit la parité du groupe majeur, les sous canaux contiennent 24 sous porteuses de données sur un temps symbole OFDM.

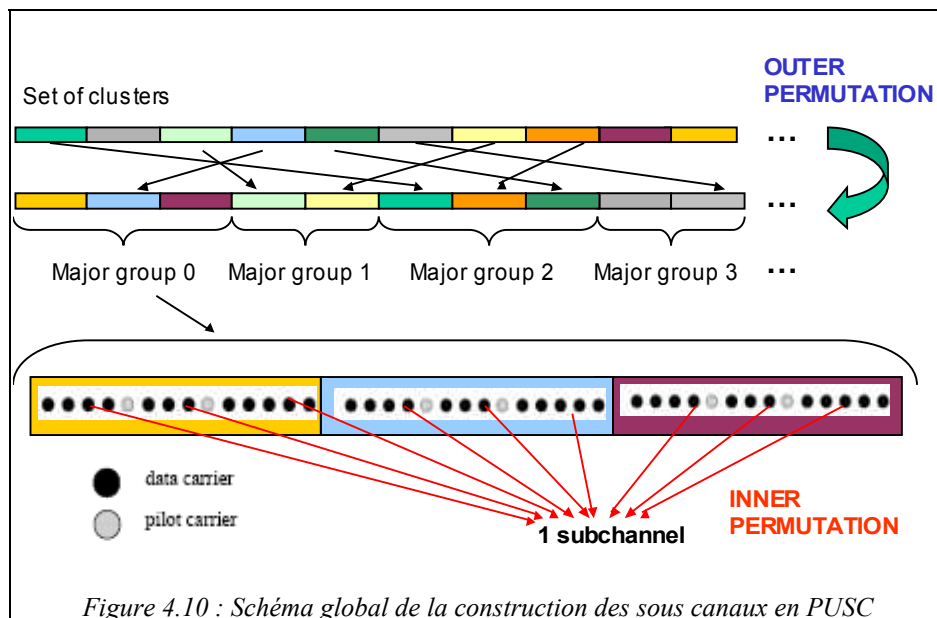
Dans chaque groupe majeur, on commence par identifier les fréquences pilotes grâce à une position prédéfinie dans les *clusters* (cf. Fig. 4.9). Les sous porteuses restantes sont réparties en 24 groupes de sous porteuses consécutives (les groupes sont plus petits dans les groupes majeurs impairs). Un sous canal est constitué d'une sous porteuse de chaque groupe, la formule utilisée et similaire au FUSC, est ici appelée *inner permutation*.

Tableau 4.3 : Formation des groupes majeurs en PUSC, [IEEE 802.16e]

Clusters	$N_{clusters}$	Group 0	Group 1	Group 2	Group 3	Group 4	Group 5
FFT 2048	120	0-23	24-39	40-63	64-79	80-103	104-119
FFT 1024	60	0-11	12-19	20-31	32-39	40-51	52-59
FFT 512	30	0-9		10-19		20-29	
FFT 128	6	0-1		2-3		4-5	



Note : L'intérêt d'un *cluster* est le positionnement des pilotes qui seront partagés par les sous canaux du groupe majeur correspondant. Les fréquences d'un *cluster* sont ensuite attribuées à différents sous canaux avec une formule similaire au FUSC.



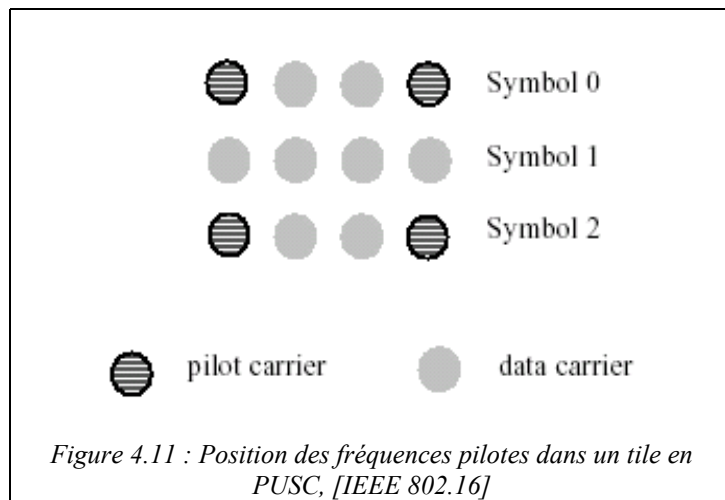
b) Voie montante

Un *slot* PUSC *uplink* s'étend sur trois temps symboles : un sous canal est constitué de 12 fréquences de données sur le premier et le dernier temps symbole et de 24 fréquences de données sur le deuxième temps symbole. Un *slot* contient donc 48 tons¹ données. En voie montante, chaque sous canal comporte ses propres fréquences pilotes.

On appelle *tile* un ensemble de quatre sous porteuses adjacentes. Sur chaque temps symbole, la bande est divisée en six groupes de *tiles*. Par exemple, sur 20 MHz, un groupe comporte 70 *tiles*. Sur un temps symbole, chaque sous canal est constitué d'un *tile* de chaque groupe grâce à une formule similaire à celle du FUSC. Sur un temps symbole, un sous canal comporte 6*4 sous porteuses. Dans la figure 4.11, la position des fréquences pilotes dans un *tile* dépend du numéro du temps symbole. D'après cette

¹ Un ton est une fréquence sur temps symbole: la même fréquence sur 2 symboles différents représentent 2 tons

figure, 12 des 24 fréquences du premier et du dernier temps symbole sont des pilotes ; les 24 fréquences du temps symbole du milieu sont des fréquences de données.



Un *tile* est un bloc compact temps / fréquence. On voit qu'un *tile* est différent d'un *cluster* car un *tile* est envoyé tout entier dans un sous canal : c'est une fraction de sous canal, les pilotes compris. On a vu qu'au contraire, les fréquences d'un *cluster* étaient envoyées dans différents sous canaux. En voie descendante, grâce à la structure de *cluster*, pour chaque fréquence d'un sous canal il y a toujours deux pilotes du groupe majeur situés à une distance minimale d_{DL} . En voie montante, la distance minimale d_{UL} est plus faible grâce à la structure de *tile*.

3.2.3. Résumé sur la sous canalisation en IEEE 802.16

Le tableau 4.4 résume la constitution des *slots* pour les différents modes.

Tableau 4.4 : Constitution d'un slot en fonction des modes

	FUSC	PUSC voie descendante	PUSC voie montante	AMC
Nombre de temps symboles OFDM	1	2	3 (Symbole : 0/1/2)	Cas 1 : $J = 2$ Cas 2 : $J = 3$ Cas 3 : $J = 6$
Nombre de fréquences de données sur un temps symbole OFDM	48	24	Symbole 0 : 12 Symbole 1 : 24 Symbole 2 : 12	Cas 1 : 24 Cas 2 : 16 Cas 3 : 8
Nombre de fréquences pilotes (non partagées) sur un temps symbole OFDM	0	0	Symbole 0 : 12 Symbole 1 : 0 Symbole 2 : 12	Cas 1 : 3 Cas 2 : 2 Cas 3 : 1

4. Performance des modes de sous canalisation et des algorithmes d'affectation des sous porteuses

Un des objectifs de ce chapitre est de comparer les performances obtenues en débit global par les modes de sous canalisation avec les performances idéales obtenues dans le chapitre 3. Les modes de sous canalisation traités ici sont le mode AMC et un mode dit de diversité. Les résultats présentés dans ce chapitre ont été publiés dans [LengGod&_07]¹.

4.1. Méthodes de sous canalisation comparées

Comme algorithme d'affectation de sous porteuses, nous choisissons l'algorithme que nous avons proposé le RPO (*Rate Profit Optimization*) et le bDA (*basic Dynamic Assignment*). Pour pouvoir les comparer avec la sous canalisation définie en WiMax, nous effectuons une allocation de bande simplifiée qui fixe le nombre de sous porteuses par utilisateur à 48.

Rappelons le fonctionnement du RPO : les utilisateurs n'ayant pas reçu 48 sous porteuses reçoivent leur meilleure sous porteuse. En l'absence de conflit, les utilisateurs reçoivent la sous porteuse choisie et recommencent. En cas de conflit, la sous porteuse concernée est affectée à l'utilisateur présentant le meilleur profit (cf. chap. 3, section 3.3).

Dans le bDA, les utilisateurs reçoivent selon un ordre arbitraire leur 48 meilleures sous porteuses. Le dernier utilisateur reçoit les 48 dernières sous porteuses.

Nous comparons ces deux algorithmes au mode FUSC. Les sous canaux FUSC sont construits conformément à la norme. Une fois construits, ils sont affectés aux utilisateurs sans considération sur le SINR.

Concernant le mode AMC, les sous canaux construits sont du type 6x1 c'est à dire six *bins* sur un temps symbole². Un sous canal contient 48 sous porteuses de données. Nous considérons deux possibilités : les sous canaux sont alloués avec ou sans considération du SINR. Lorsque le SINR est considéré, les utilisateurs pris dans un ordre arbitraire reçoivent le sous canal présentant le meilleur SINR moyen.

4.2. Conditions de simulation

Nous nous intéressons aux performances, sur la voie descendante, d'une cellule centrale constituée d'une station de base et de U utilisateurs. Cette cellule est entourée de 18 cellules interférentes réparties sur deux anneaux. Pour simplifier, les antennes sont omni-directionnelles et il n'y a pas de sectorisation.

¹ Ce papier a reçu le Best Award Paper à la conférence WiMob 2007

² Ce type n'est pas considéré dans la norme, mais il est pratique pour la comparaison

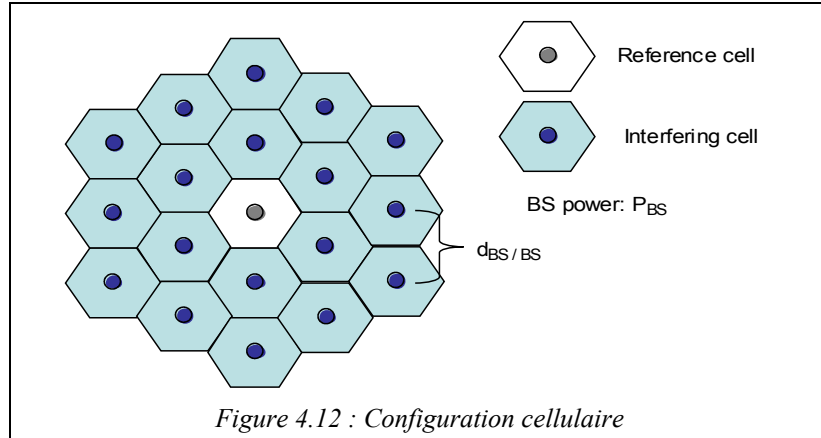


Figure 4.12 : Configuration cellulaire

4.2.1. Modèle de canal

Le carré du module du gain du canal de l'utilisateur u sur la sous porteuse n est :

$$|g_{u,n}|^2 = K d(u)^{-\alpha} a_{sh}(u) a_f(u,n)$$

Dans cette formule, K est l'atténuation à la distance de référence 1 m ; $d(u)$ est la distance de l'utilisateur ; α est le coefficient d'atténuation ; $a_{sh}(u)$ est une variable lognormale qui traduit l'effet de masque subi par toutes les sous porteuses de l'utilisateur u ; enfin $a_f(u,n)$ est une variable exponentielle pour modéliser l'évanouissement rapide (*fast fading*) de la sous porteuse n . Nous considérons une corrélation entre l'évanouissement rapide subie par deux sous porteuses consécutives. La bande de cohérence du canal est donnée par $B_c = 1/2\pi\sigma_{rms}$. Pour $\sigma_{rms} = 290$ ns on a $B_c = 539$ kHz, cela correspond à peu près à la largeur d'un sous canal AMC lorsque $N_{scPerChan} = 48$ et $\Delta f = 10.93$ kHz : $N_{scPerChan} \Delta f = 525$ kHz.

Les sous porteuses subissent par ailleurs un bruit blanc gaussien de variance $\sigma^2 = N_0 B/N$. Nous adoptons une allocation de puissance uniforme $P_{u,n} = p$ sur chaque sous porteuse. Nous pouvons évaluer le niveau d'interférence subi par chaque sous porteuse n . Le SINR (*Signal to Interference and Noise Ratio*) s'exprime : $\gamma_{u,n} = p g_{u,n} / (\sigma^2 + I(u,n))$ avec $I(u,n) = \sum_{b=1 \dots B} K d_b(u)^{-\alpha} a_{sh}^b(u) a_f^b(u,n) \delta_{b,n}$ où $\delta_{b,n} = 1$ si la sous porteuse n est utilisée dans la cellule b et 0 sinon. Le facteur de réutilisation fréquentiel (FRF) vaut donc 1 lorsque le facteur de charge est égal à 1 (c'est à dire que tous les sous canaux sont utilisés). Dans ce cas, le paramètre $\delta_{b,n}$ vaut 1 pour tout b et n .

Les sous porteuses d'un sous canal utilisent le même MCS. Pour le déterminer on calcule le SINR effectif du sous canal (cf. [WimaxEval_07]) : $2^{\text{MIC}} - 1$ où MIC est le *mean instantaneous capacity* qui est la moyenne de la capacité de $\log_2(1 + \gamma_{u,n})$ sur les 48 sous porteuses du sous canal alloué à u . La connaissance du canal est supposée parfaite. Une fois le SINR effectif calculé, le MCS est déterminé selon le tableau 4.5 ([IEEE 802.16], p 621). Le MCS est unique pour le sous canal. Dans le chapitre 3, le MCS était individuel (par sous porteuse) et déterminé par une fonction continue proposée dans la littérature ([Kivanc&_03]).

Tableau 4.5 : MCS et seuils de SINR ([IEEE 802.16])

Modulation et codage	Nbre de bits utiles/s/Hz	SINR requis (dB)
QPSK 1/2	1	6
QPSK 3/4	1.5	9
16 QAM 1/2	2	12
16 QAM 3/4	3	15
64 QAM 1/2	4	18
64 QAM 3/4	4.5	21

4.2.2. Déroulement de la simulation

La puissance P_{BS} transmise par toutes les BS vaut 43 dBm. Les utilisateurs sont situés à distance fixe de leur station de base de référence. Il y a autant d'utilisateurs que de sous canaux disponibles. On réalise une simulation à base de *snapshots* : on réalise 7500 *snapshots* et les résultats sont moyennés. Les intervalles de confiance sont pris à 95%. Pour chaque *snapshot*, on tire de façon aléatoire l'effet de masque et l'évanouissement rapide corrélé sur tous les triplets (u, n, b) .

Tableau 4.6 : Paramètres de simulation, impact de la sous canalisation

Paramètres	Valeurs
W (MHz)	10
U	16
S : nombre de sous canaux dans une cellule	16
N_{data}	768
F : fréquence centrale (GHz)	3.3
Δf (kHz)	10.93
K : constante de l'atténuation de parcours α : exposant de l'atténuation de parcours	$1.4 \cdot 10^{-4}$ variable
σ_{sh} (dB)	8.9
σ_{rms} (ns)	295
P_{BS} : puissance (dBm)	43
N_0 : Densité de bruit blanc (dBm/Hz)	-174
D_{BS-BS} : distance inter-BS	variable
d : distance entre utilisateur et BS	variable

Dans un premier temps, la distance entre deux stations de base voisines vaut 2.5 km. La distance entre un utilisateur et sa station de base vaut 500 m. On s'intéresse au débit global des différentes méthodes de sous canalisation. Puis on fait successivement varier la distance inter-BS (entre 1 et 5 km) et l'exposant d'atténuation α . On s'intéresse alors à l'évolution du débit par utilisateur et par sous porteuse. Enfin, la distance entre un utilisateur et sa BS augmente entre 100 m et 1 km pour une distance inter-BS de 2.5 km. Les paramètres de simulation sont résumés en tableau 4.6.

4.3. Résultats de simulation

4.3.1. Débit global

On compare le débit global moyen de la cellule centrale pour les différentes méthodes de sous canalisation. Le débit global est la somme des débits individuels et une moyenne est réalisée sur plusieurs *snapshots*. La puissance P_{BS} émise par BS vaut 43 dBm. La distance inter-BS d_{BS-BS} vaut 2.5 km. La distance entre les utilisateurs et leur BS vaut 500 m. L'exposant α d'atténuation vaut 3.5. La charge vaut 1 dans toutes les cellules et chaque utilisateur reçoit un sous canal. Le FRF vaut 1. Les débits obtenus sont résumés en tableau 4.7.

Tableau 4.7 : Débit global par modes, $d = 500$ m, $d_{BS-BS} = 2.5$ Km, $\alpha = 3.5$, $P_{BS} = 43$ dBm

Modes	RPO	bDA	AMC 6x1 avec SINR	AMC 6x1 sans SINR	FUSC
Débit global en Mbits/s	18.06	17.45	16.59	12.75	12.86

Dans le RPO, chaque utilisateur reçoit 48 sous porteuses sur lesquelles il réalise le meilleur profit. On voit que cela permet d'obtenir le débit global moyen le plus élevé. On peut cependant reprocher au RPO une durée d'exécution légèrement plus élevée que les autres méthodes et la nécessité d'une signalisation par sous porteuse.

On peut noter qu'à la différence du chapitre 3, les résultats obtenus ici par le RPO sont obtenus avec un MCS commun aux 48 sous porteuses attribuées à un utilisateur. Dans le chapitre 3, le MCS de chaque sous porteuse était individuel c'est à dire adapté à son SNR. La différence en efficacité spectrale du RPO entre le chapitre 3 et chapitre 4 est de 2 bits/s/Hz.

Le mode AMC regroupe des sous porteuses consécutives. Le regroupement en sous canaux de sous porteuses consécutives baisse le débit global de 8 %. Mais cela réduit grandement la charge de signalisation. Le fait de ne pas prendre en compte le SINR dans l'affectation de sous canaux AMC provoque une réduction de 23 % sur le débit global (AMC avec SINR $\rightarrow$ AMC sans SINR). L'utilisation de sous canaux AMC sans considération des SINR s'apparente, au regard du débit global moyen, à l'utilisation de sous canaux FUSC.

L'utilisation de sous canaux AMC sans considération des SINR peut représenter le cas de figure où les informations sur la qualité du canal sont *périmées* ce qui peut être le cas lors d'une grande mobilité du terminal. Dans ce cas, on voit que l'AMC ne présente pas d'avantage sur le mode FUSC.

4.3.2. Allègement des interférences

Dans cette section, nous évaluons l'impact des paramètres d_{BS-BS} et α sur les performances en débit. La distance inter-BS d_{BS-BS} et l'exposant α d'atténuation influent directement sur le niveau des interférences reçues par les cellules voisines.

Dans la figure 4.13, la distance inter-BS d_{BS-BS} varie tandis que les autres paramètres sont fixés : la distance entre utilisateurs et BS vaut 500 m, l'exposant α d'atténuation vaut 3.5 et la puissance émise par les BS vaut 43 dBm. Le facteur de charge et le FRF valent 1. La figure 4.14 nous permet de visualiser les intervalles de confiance.

Comme on pouvait s'y attendre, le débit par sous canal s'améliore lorsque la distance inter-BS augmente. On observe les mêmes positions relatives des modes de sous canalisation que dans la section précédente. Cette courbe montre que lorsque $P_{BS} = 43$ dBm et $\alpha = 3.5$, les cellules sont isolées à une distance d_{BS-BS} de 5 km car on retrouve les performances d'une cellule sans interférences extérieures¹

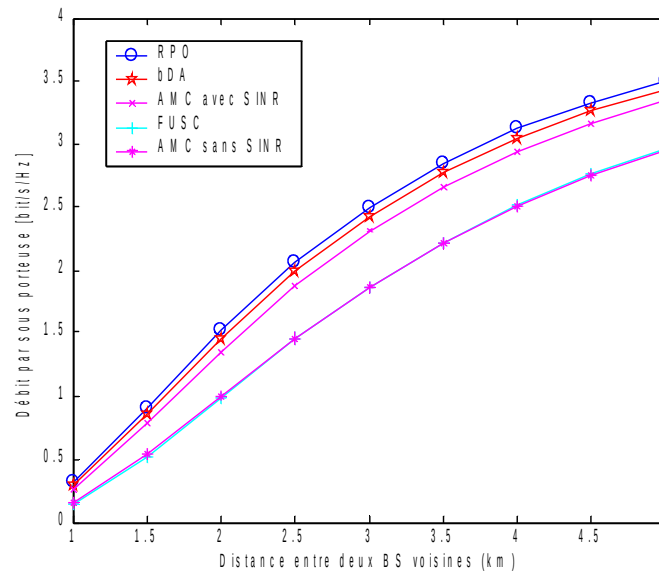


Figure 4.13 : Débit en fonction de la distance inter-BS, $d = 500$ m, $\alpha = 3.5$, $P_{BS} = 43$ dBm

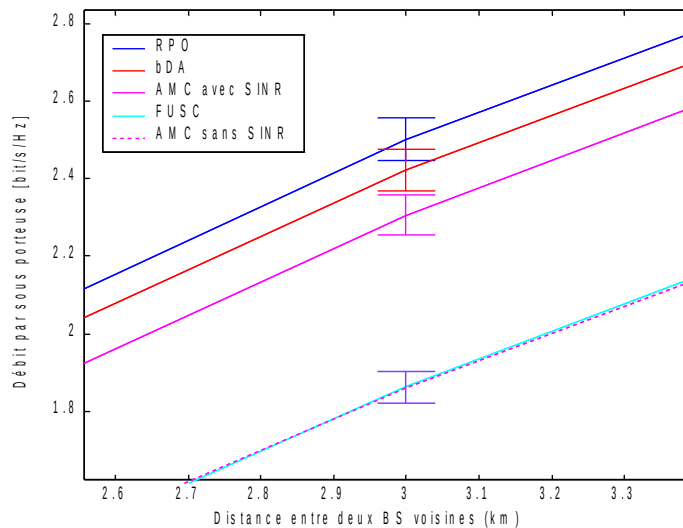


Figure 4.14 : Débit en fonction de la distance inter-BS, zoom et intervalles de confiance

¹ Résultat montré dans [LengGodM_07].

Dans la figure 4.15, l'exposant α varie tandis que les autres paramètres sont fixés : la distance inter-BS vaut 2.5 km, la distance entre utilisateurs et BS vaut 500 m la puissance émise par BS vaut 43 dBm. Le facteur de charge et le FRF valent 1.

Cette courbe montre que lorsque $P_{BS}=43$ dBm et $\alpha=2$, l'efficacité spectrale est très faible (inférieure à 0.5) et ce quelque soit le mode de sous canalisation. Lorsque le paramètre α augmente, l'atténuation du signal est plus rapide et les cellules subissent moins d'interférences en provenance de leurs voisines. Les meilleures performances sont observées lorsque $\alpha=3.5$. Pour $\alpha=4$, on observe une légère baisse, cela s'explique par fait que la forte atténuation du signal utile dans la cellule est plus significative que la baisse du niveau des interférences.

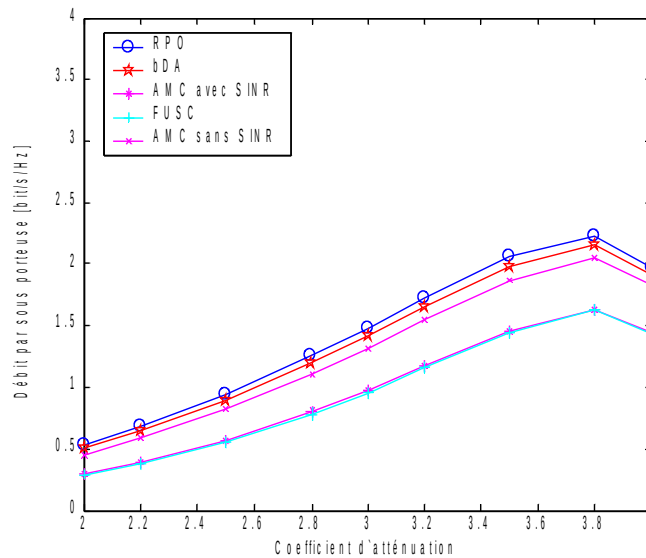


Figure 4.15 : Débit en fonction du coefficient d'atténuation ($d = 500$ m, $d_{BS-BS} = 2.5$ km, $P_{BS} = 43$ dBm)

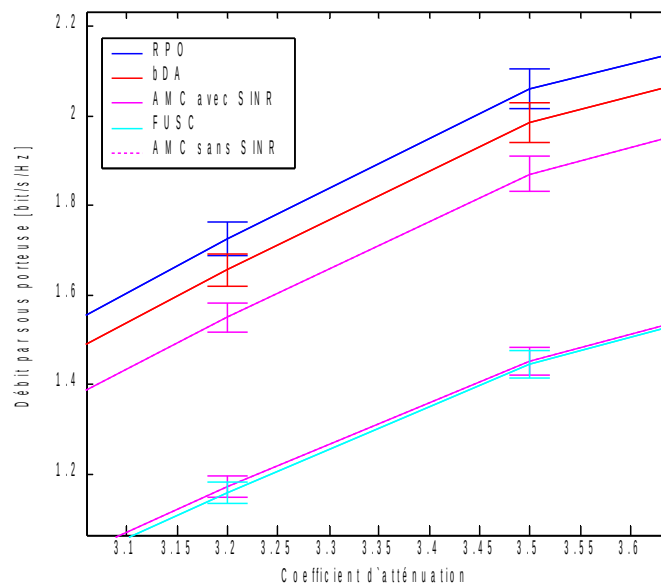


Figure 4.16 : Débit en fonction du coefficient d'atténuation, zoom et intervalles de confiance

Sur le terrain, on ne peut agir sur l'exposant d'atténuation. Lors du déploiement, on décide de la distance inter-BS mais une fois le réseau déployé, on a besoin d'autres leviers pour maîtriser le niveau des interférences. C'est l'objet de la section suivante.

4.3.3. Effet de la charge en FUSC

Dans cette section on ne s'intéresse qu'au mode de sous canalisation FUSC. Pour alléger le niveau des interférences, les cellules peuvent limiter la charge en utilisant seulement une partie des sous canaux disponibles. Le facteur de charge est alors différent de 1.

Dans le tableau 4.8, on teste différents facteurs de charge et on compare le débit par sous canal et le débit global. Le débit par sous canal est améliorée de 61 % quand on passe d'un facteur de 1 à 0.4 ; mais le débit global baisse de 40 % puisqu'il y a moins de sous canaux utilisés.

Tableau 4.8 : Facteur de charge et débit, $d = 500$ m, $d_{BS-BS} = 2.5$ km, $\alpha = 3.5$, $P_{BS} = 43$ dBm

Facteur de charge	0.4	0.6	0.8	1
Débit global (Mbits/s)	7.49	9.25	10.46	12.35
Débit par sous canal FUSC (Mbits/s)	1.25	1.03	0.87	0.77

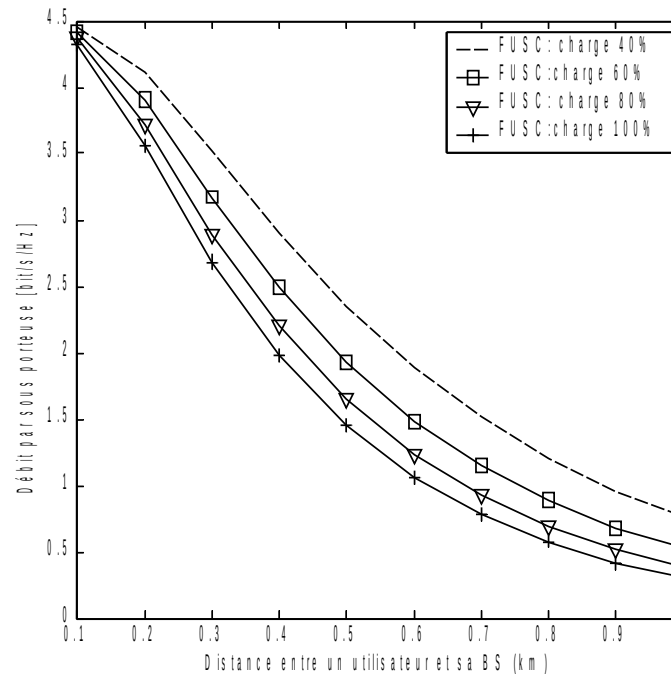


Figure 4.17 : Comparaison des facteurs de charge en FUSC
($d_{BS-BS} = 2.5$ Km, $\alpha = 3.5$, $P_{BS} = 43$ dBm)

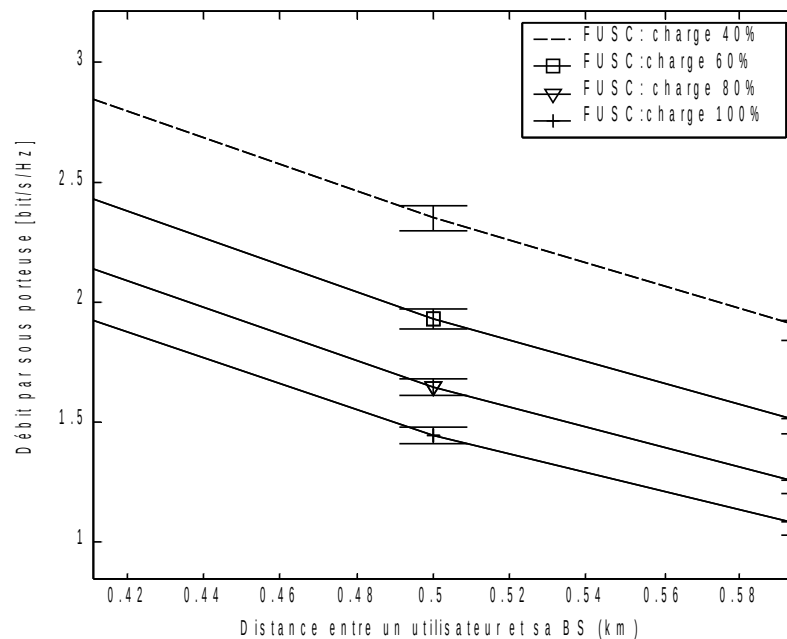


Figure 4.18 : Comparaison des facteurs de charge en FUSC, zoom et intervalles de confiance

Dans la figure 4.17, on trace l'efficacité spectrale en fonction de la distance entre utilisateurs et BS. Les autres paramètres sont fixés : la distance inter-BS vaut 2.5 km, l'exposant α vaut 3.5, la puissance émise par BS vaut 43 dBm et le FRF vaut 1. Quatre courbes correspondant à quatre facteurs de charge différents sont tracées. La figure 4.18 permet de visualiser les intervalles de confiance.

A 500 m de la BS (sachant que l'espacement entre BS est de 2.5 km) un utilisateur peut espérer 706 kbps avec un sous canal en pleine charge. Avec une charge de 0.4, ce débit devient 1.2 Mbps. A 1 km de la BS (sachant que l'espacement entre BS est de 2.5 km), un utilisateur peut espérer 240 kbps avec un sous canal en pleine charge. Avec une charge de 0.4, ce débit devient 336 kbps. Pour un utilisateur situé entre 300 et 700 m de la BS, utiliser un facteur de charge de 0.6 au lieu de 1 permet à un utilisateur d'utiliser le MCS immédiatement supérieur.

5. Conclusions

5.1. La sous canalisation en WiMax

Le chapitre 4 présente la structure de trame en IEEE 802.16 et ses principaux modes de sous canalisation.

Le mode AMC permet d'utiliser des stratégies opportunistes lors de l'allocation des sous canaux car les fréquences y sont adjacentes. Les pilotes sont, en voie montante et descendante, affectés spécifiquement à un sous canal.

Le mode FUSC permet de construire des sous canaux uniformes. Le facteur de réutilisation fréquentiel est forcément de 1. Il est possible cependant de faire varier le facteur de charge c'est à dire le nombre de sous canaux utilisés dans chaque cellule. Ce mode n'est utilisé que sur la voie descendante.

Le mode PUSC construit des sous canaux uniformes avec la possibilité d'affecter des parties distinctes de la bande à des cellules ou secteurs différents. Cela permet différentes possibilités pour le facteur de réutilisation fréquentiel. Les pilotes sont en voie montante affectés spécifiquement à un sous canal.

5.2. Résultats

Nous avons évalué, en FUSC et en PUSC, le nombre de collisions entre les sous canaux utilisés dans deux cellules distinctes. Nous avons comparé les densités de probabilité de nombre de collisions en PUSC et en FUSC avec celle de la sous canalisation aléatoire. Cette comparaison a montré que le nombre moyen de collisions en PUSC et en FUSC est le même que celui de la sous canalisation aléatoire. En PUSC et en FUSC, la probabilité de collision sur une sous porteuse est égale au facteur de charge.

Dans ce chapitre, nous avons pu comparer les performances en débit idéales (obtenues par le RPO) et les performances en débit que l'on peut obtenir en sens descendant avec les modes de sous canalisation (et une adaptation de lien par sous canal). La connaissance du canal est supposée parfaite. Nous avons principalement travaillé avec un FRF de 1 et une charge totale.

Pour une puissance par BS de 43 dBm, une distance inter-BS de 2.5 km, une distance entre utilisateurs et BS de 500 m et un exposant α d'atténuation de 3.5, le débit global obtenu par le RPO est de 18.06 Mbps sur une bande de 10 MHz. La baisse de performances par rapport au chapitre 3 est principalement due au MCS commun sur les sous porteuses d'un utilisateur.

Regrouper les sous porteuses consécutives et affecter les sous canaux selon le SINR moyen baisse le débit globale de 8 % par rapport aux performances idéales du RPO. Le mode FUSC présente des performances inférieures de 23 % à celle du mode AMC (avec prise en compte du SINR). Cependant ce résultat dépend fortement de l'hypothèse sur la connaissance du canal.

Lorsque le canal varie vite ou si le terminal est très mobile, les mesures arrivant à la station de base peuvent être périmées. Les performances du mode AMC peuvent être sérieusement dégradées. Dans [Riati&_07], l'impact de la mobilité est étudié en 802.16. Il y est montré que le mode AMC présente des bénéfices sur les modes de diversités (notamment PUSC) pour des vitesses inférieures à 8 km/h et des types de trafic tolérants au délai.

Les performances des modes de diversités sont celles que l'on peut garantir en cas de grande mobilité ou avec une connaissance pauvre du canal. Pour un FRF de 1 et en pleine charge, le débit global par sous canal en FUSC peut être assez faible. Nous avons vu qu'à 1 km de la BS, un utilisateur peut espérer 240 kbps. Nous avons donc envisagé d'autres facteurs de charge. Pour des distances entre 300 et 700 m de la BS, avec nos paramètres de simulation, on peut utiliser un facteur de charge de 0.6 au lieu de 1. Cela permet à un utilisateur d'utiliser le MCS immédiatement supérieur.

Chapitre 5. Allocation de ressources en OFDMA : contribution dans le contexte multicellulaire

1. Introduction

Nous abordons l'attribution de sous porteuses dans un contexte multicellulaire. Le contexte considéré peut être qualifié d'idéal compte tenu des deux hypothèses qui suivent. La connaissance du canal est supposée parfaite. L'attribution et l'adaptation de lien se fait sous porteuse par sous porteuse.

On retrouve la distinction entre les optimisations MA (minimiser la puissance transmise dans la cellule) et RA (maximiser le débit total de la cellule). Nous cherchons à résoudre un problème de type RA avec une contrainte en débit associée à chaque utilisateur.

Ce contexte est similaire à celui du chapitre 3 où nous avons étudié l'affectation de sous porteuses sans interférences co-canal. La principale difficulté réside dans la résolution des conflits ; ces derniers surviennent lorsque deux utilisateurs (ou plus) ont la même sous porteuse.

Lorsque l'on se place dans un contexte multicellulaire, il faut ajouter aux difficultés précédemment évoquées, la prise en compte des interférences dues à la réutilisation d'une fréquence dans les cellules voisines. Résoudre le problème RA dans un contexte multicellulaire revient à déterminer, pour chaque sous porteuse, l'ensemble des cellules dans lesquelles elle sera utilisée (notion de FRF – facteur de réutilisation fréquentiel ou encore *frequency reuse factor*) ainsi que les utilisateurs qui vont la partager. Nous adoptons la convention suivante concernant le facteur de réutilisation fréquentiel :

- $FRF = 1$: la (les) sous porteuse(s) sont utilisées dans toutes les cellules
- $FRF = 1/k$: la (les) sous porteuse(s) sont utilisées dans une cellule sur k .

Dans ce chapitre, le facteur de réutilisation n'est pas uniforme. Les hypothèses de ce chapitre diffèrent donc de celles du chapitre 4 : il y avait de la sous canalisation et le facteur de réutilisation fréquentiel était uniforme et fixé à 1.

Nous considérons la voie descendante. Par soucis de simplicité, les cellules ne sont pas sectorisées . Un modèle en tore est considéré pour simuler les interférences des cellules extérieures. Nous nous intéresserons au cas particulier dans lequel la puissance allouée à chaque sous porteuse est uniforme. Cette hypothèse peut sembler préjudiciable pour les utilisateurs qui subissent de mauvaises conditions de propagation (faible niveau du signal utile et niveau élevé des interférences) par rapport aux utilisateurs dont les

conditions de propagation sont favorables. Pour pallier les inégalités entre utilisateurs, nous diminuons le facteur de réutilisation fréquentiel (FRF) des sous porteuses affectées aux utilisateurs défavorisés afin de baisser le niveau des interférences. Cette solution est préférée à l'augmentation de la puissance du signal utile dans chaque cellule. Elle a l'avantage d'éviter la surenchère en puissance.

Nous présentons dans la section 2, quelques algorithmes proposés pour l'allocation de ressources dans un contexte multicellulaire en RA. Nous décrivons, dans la section 3, un nouvel algorithme : l'OSA-IL (*Opportunist Subcarrier Allocation with Interference Limitation*, [LengGodM_06]). Cet algorithme est centralisé et procède par sous porteuse. Dans cet algorithme, on construit pour chaque sous porteuse, l'ensemble des utilisateurs qui vont la partager. Les utilisateurs d'une même cellule ne doivent pas partager une même sous porteuse. L'algorithme adapte le FRF de façon à maintenir un SINR seuil (*Signal to Interference and Noise Ratio*) sur les utilisateurs qui partagent une même sous porteuse. Dans la section 4, les résultats de cet algorithme sont comparés à ceux des algorithmes existants (adaptés au contexte de la puissance uniforme entre les sous porteuses).

2. Travaux existants

2.1. Aperçu des travaux par type d'optimisation

2.1.1. Travaux MA

La classe des travaux MA (*Margin Adaptive optimization*) est celle des travaux où l'on minimise la puissance transmise afin de réaliser un débit minimal (r_u°) pour chaque utilisateur u . Parmi les travaux MA, on peut citer [Junqiang&_03], [PietrzykJanssen_03] et [Gault&_05].

Dans [PietrzykJanssen_03], on accepte un nouvel entrant dans le système multicellulaire seulement s'il ne dégrade pas la QoS des utilisateurs existants (débit minimal r_u°). Dans ce cas, les auteurs utilisent d'abord le BABS (*Bandwidth Assignment based on SNR*, cf. chap. 2) pour déterminer le nombre N_u de sous porteuses par utilisateur. Puis l'algorithme ACG (*Amplitude Craving Greedy algorithm* avec utilisation des CgINR¹, cf. chap. 2) pour la phase d'allocation spécifique de sous porteuses. Enfin, l'algorithme du *bit loading* (cf. chap. 2) est appliqué sur les sous porteuses.

Dans [Junqiang&_03], on considère des cellules virtuelles ; il s'agit de la réunion de trois secteurs interférents (i.e. issus de trois cellules adjacentes). Les auteurs utilisent le BABS, puis l'algorithme «modified ACG» (avec utilisation des CgINR). Le nombre de bits initial sur une sous porteuse est r_u° / N_u . Des transferts de capacité (en bits) entre sous porteuses sont réalisés s'ils améliorent la puissance transmise.

¹ La notion de CgNR est remplacée par le CgINR (*Channel gain Interference to Noise Ratio*) pour adapter l'algorithme au contexte multicellulaire. Le meilleur utilisateur d'une sous porteuse est donc celui qui présente le meilleur CgINR.

Les auteurs de [Gault&_05] effectuent l'allocation des ressources en se basant sur une connaissance statistique des gains des utilisateurs (en particulier la variance). Dans la plupart des travaux existants, les gains d'un utilisateur, sur toutes les sous porteuses, sont connus par la station de base (BS). Ici, un déploiement linéique de cellules est considéré. Les auteurs utilisent le saut de fréquence. Un algorithme d'allocation de puissance augmente la puissance d'émission de chaque BS pour combattre les interférences. Les auteurs montrent que cet algorithme converge si la charge (débit moyen par km) requise par les utilisateurs est inférieur à un seuil.

2.1.2. Travaux RA

La classe des travaux RA (*Rate Adaptive optimization*) est celle des travaux où l'on maximise le débit global de la cellule avec une contrainte sur le budget de puissance.

Le contexte de [YihGerianotis_00] est l'OFDM-TDMA. On cite ce travail car il réalise une adaptation du *waterfilling* au cas multicellulaire. Dans chaque cellule, pour l'utilisateur en cours de transmission, l'ensemble des sous porteuses actives est privé de la sous porteuse subissant le plus d'interférences jusqu'à ce que le *waterfilling* (appliqué avec les CgINR) réalise sur chaque sous porteuse active un SINR supérieur à un seuil prédéterminé. Dans [Yan&_03], ce principe est repris et adapté à l'OFDMA. Une sous porteuse n'est affectée à un utilisateur que si elle présente un SINR¹ supérieur à un seuil fixé. Le *waterfilling* est appliqué pour chaque utilisateur u sur l'ensemble des sous porteuses Ω_u qui lui est attribué.

Dans [Kim&_04], on essaie de satisfaire des débits minimaux pour les utilisateurs lors de la maximisation du débit global. Les auteurs considèrent K facteurs de réutilisation préfixés. Le nombre de BS qui utilisent un sous canal identique dépend de la valeur du FRF du sous canal. Les sous porteuses d'un sous canal sont entrelacées sur toute la bande et subissent le même facteur de réutilisation (FRF). Trois phases constituent l'allocation : (i) préférences intra-BS, (ii) l'harmonisation inter-BS et (iii) l'allocation intra-BS. Dans la première phase, chaque BS calcule, en fonction des conditions radio de ses utilisateurs, le nombre de sous canaux souhaités par valeur de FRF. La deuxième phase détermine les nombres définitifs, communs à toutes les BS, de sous canaux utilisés par valeur de FRF. Enfin, de façon indépendante dans chaque BS, les utilisateurs reçoivent des sous canaux selon leurs demandes dans la première phase (cf. annexe 5.A pour le détail de chaque phase). La dernière phase manque de précision dans les cas où, dans une BS, l'offre en sous canaux est inférieur à la demande.

Nous détaillons les travaux [Koutsopoulos_02] et [Han_03] que nous étudions particulièrement en 2.2. Notre proposition en RA sera comparée à ces deux contributions en fin de chapitre.

2.1.3. Autres Travaux

Certains travaux ne se situent dans aucune des classes MA et RA.

Dans [Kwon&_05], les auteurs minimisent l'*outage* maximal dans les pseudo cellules. Une pseudo cellule est définie comme la réunion de k secteurs interférents ($k = 3$ pour des cellules hexagonales et $k = 4$ pour des cellules carrées). L'*outage* est la proportion d'utilisateurs ne satisfaisant pas leur débit minimal. La puissance transmise sur les sous

1 La puissance utilisée pour effectuer ce calcul préalable du SINR n'est pas précisée par les auteurs.

porteuses est uniforme. Deux FRF sont considérés : 1 et $1/k$. Pour un FRF de $1/k$, chaque SG¹ (*Subcarrier Group*) est divisée en k portions. Une portion de SG est exclusivement utilisée dans un des k secteurs d'une pseudo cellule. Les utilisateurs d'une pseudo cellule sont classés dans deux ensembles distincts selon le FRF qui leur convient le mieux (cf. annexe 5.A). Grâce à ce classement et à une phase d'harmonisation entre les secteurs, le nombre de sous porteuses allouées avec chaque FRF est calculé (cf. annexe 5.A). Finalement, l'allocation se fait de façon indépendante dans chaque secteur : l'utilisateur le plus loin de son objectif (en débit minimal) reçoit sa meilleure sous bande. Si l'*outage* dépasse le seuil, l'utilisateur le plus éloigné de son débit minimal est abandonné et ses ressources sont redistribuées.

Dans [Hamouda&_06], on trouve un travail similaire avec contrôle de puissance.

2.2. Travaux d'importance en RA

Dans [Koutsopoulos_02] et [Han_03], le FRF est «adaptatif». Cela signifie que le FRF d'une sous porteuse n'est pas uniforme et peut changer d'une sous porteuse à l'autre.

2.2.1. Problème RA sans contraintes sur le débit minimal ([Koutsopoulos_02])

Dans [Koutsopoulos_02], l'auteur veut maximiser le débit lorsqu'il n'y a pas de contraintes en débit². L'auteur de [Koutsopoulos_02] considère la notion d'utilisateur «co-canal» : il s'agit de l'ensemble Π_n des utilisateurs qui vont partager une sous porteuse n ; ces utilisateurs doivent provenir de cellules différentes. Chaque sous porteuse n est traitée successivement (et indépendamment). Au début de la construction de l'ensemble Π_n des utilisateurs «co-canal», tous les utilisateurs du système sont candidats. Quand un utilisateur est sélectionné, les utilisateurs appartenant à la même cellule sont disqualifiés pour la suite de l'affectation de la sous porteuse n . Concernant les critères de sélection d'un utilisateur, on distingue les cas avec et sans contrôle de puissance.

a) Allocation sans contrôle de puissance

La sélection d'un utilisateur u appartenant à une cellule b dépend des notions de SIF – *Signal Interference Factor* –, d'IRF – *Incremental Rate Factor* – et d'APF – *Assignment Preference Factor*.

Notion de SIF

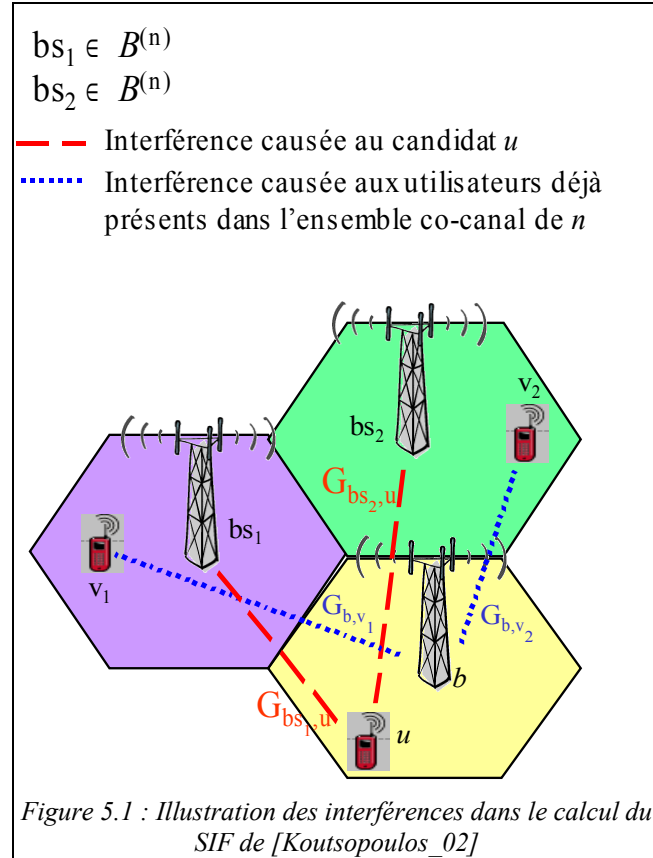
La BS teste si l'utilisateur u de la cellule b peut entrer dans l'ensemble Π_n . On illustre la situation en Fig.5.1, on appelle $B^{(n)}$ l'ensemble des BS qui utilisent la sous porteuse n . L'auteur distingue le niveau cumulé d'interférence que l'utilisateur u causerait aux utilisateurs déjà présents dans l'ensemble Π_n $\left(\sum_{bs \in B^{(n)}} G_{bs,u}^{(n)} \right)$ et l'interférence que les utilisateurs déjà présents dans l'ensemble Π_n causeraient à l'utilisateur u $\left(\sum_{v \in \Pi_n} G_{b,v}^{(n)} \right)$. Le maximum de ces deux interférences est utilisé pour calculer le SIF :

1 Ensemble de sous porteuses consécutives appartenant à la même bande de cohérence.

2 Lorsqu'il existe des contraintes en débit fixées, l'auteur minimise le nombre de sous porteuses nécessaires.

$$SIF = \frac{G_{b,u}^{(n)}}{\max(\sum_{v \in \Pi_n} G_{b,v}^{(n)}, \sum_{bs \in B^{(n)}} G_{bs,u}^{(n)})}$$

car les BS émettent toutes la même puissance ; de plus, le bruit est supposé négligeable.



Notion d'IRF

L'IRF permet d'évaluer les conséquences de l'introduction d'un utilisateur sur le débit de la sous porteuse. Chaque MCS est associée à une plage de SIR. L'introduction d'un nouvel utilisateur dans l'ensemble co-canal Π_n augmente le niveau d'interférence des utilisateurs déjà présents : le MCS diminue et doit être mis à jour. L'IRF est la différence entre le nombre de bits utiles perdus par les utilisateurs déjà présents dans Π_n et le nombre MCS de l'éventuel nouvel entrant u de la cellule b .

Éligibilité d'un utilisateur et notion d'APF

Seuls les candidats dont l'IRF est positif sont éligibles. Si tous les IRF sont négatifs alors aucun utilisateur n'est capable d'améliorer le débit de la sous porteuse ; l'ensemble Π_n n'accepte plus de nouvel entrant. On introduit un nouveau facteur l'APF pour sélectionner un utilisateur parmi les utilisateurs éligibles. L'APF est le produit du SIF et de l'IRF. Parmi les éligibles, l'utilisateur sélectionné présente le plus grand APF.

Procédure d'affectation

Une fois les trois notions précédentes définies, la procédure d'affectation en découle simplement. Pour une sous porteuse, tant qu'il existe des candidats dont l'APF est positif, l'ensemble co-canal est construit en sélectionnant l'utilisateur présentant le plus grand APF. L'ensemble des candidats est mis à jour en retirant l'utilisateur sélectionné et

les utilisateurs venant de la même cellule, puis les APF sont recalculés. Une fois l'allocation d'une sous porteuse terminée, on passe à la suivante. Soit M le nombre de MCS, la complexité annoncée par sous porteuse est $O(M U_{total} B^2)$ avec U_{total} le nombre total d'utilisateurs et B le nombre de BS.

b) Intoduction du contrôle de puissance

Dans [Koutsopoulos_02], le contrôle de puissance peut être activé lorsque l'IRF de tous les candidats est négatif (même avec la plus faible modulation pour tous les utilisateurs de l'ensemble co-canal).

Dès lors, il s'agit de déterminer un vecteur de puissance (puissance P_b émise par chaque BS) qui permette l'entrée d'un nouvel utilisateur tout en augmentant le débit sur la sous porteuse i.e $IRF > 0$. Les seuils de SIR que doivent respecter les utilisateurs de l'ensemble co-canal (y compris le nouvel entrant potentiel) sont notés $(\gamma_j)_{1 \leq j \leq B}$. On doit avoir

$$\frac{G_{jj} P_j}{\sum_{i=1, i \neq j, i \in B^{(n)}}^B G_{ij} P_i} \geq \gamma_j . \text{ L'existence d'une solution est garantie lorsque la plus grande}$$

valeur propre de la matrice $\frac{\gamma_j G_{i,j}}{(1 + \gamma_j) G_{j,j}}$ est inférieure à 1 ; le vecteur propre $(P_j)_{1 \leq j \leq B}$ associé est le vecteur de puissance solution.

L'algorithme se déroule comme suit, une fois le contrôle de puissance activé¹, on vérifie l'existence d'un vecteur de puissance pour chaque vecteur de SIR². Pour chaque candidat, les MCS sont testés du plus fort au plus faible jusqu'à ce qu'un IRF positif soit trouvé (en l'absence de solution $IRF = -\infty$). Le SIF est alors

$$SIF = \frac{P_b^{(n)} G_{b,u}^{(n)}}{\max(\sum_{v \in \Pi_n} P_b^{(n)} G_{b,v}^{(n)}, \sum_{bs \in B^{(n)}} P_{bs}^{(n)} G_{bs,u}^{(n)})} . \text{ L'APF est calculé et le nouvel entrant est}$$

désigné (celui qui présente le plus grand APF). La complexité par sous porteuse est $O(M U_{total} B^5)$.

2.2.2. Résolution avec contraintes sur le débit minimal ([Han_03])

Dans [Han_03], on essaie de satisfaire un vecteur de débit minimaux lors de la maximisation du débit global. L'auteur préconise deux phases : une phase d'allocation initiale et une phase d'amélioration de débit réalisée par un algorithme itératif.

La phase d'amélioration de débit (calculs de gradients sur la capacité) est susceptible d'augmenter l'*outage* (proportion d'utilisateurs ne satisfaisant pas leur débit minimal). Nous ignorons cette phase et ne décrivons dans la suite que la phase d'allocation initiale. En l'absence de précisions de la part de l'auteur, nous considérons une puissance uniforme sur les sous porteuses.

1 Il nous paraît logique qu'une fois, le contrôle de puissance activé, il le reste pour tous les futurs entrants de l'ensemble co-canal.

2 Par soucis de simplicité, lorsque le contrôle de puissance est activé, tous les utilisateurs de l'ensemble co-canal (y compris le candidat) utilisent la même modulation

Parmi les propositions faites pour la phase d'allocation initiale nous nous intéressons à celle qui réalise le principe de *maximum packing*. Le *maximum packing* était déjà présent dans [Koutsopoulos_02] : pour chaque n , l'ensemble des utilisateurs «co-canal» Π_n est optimisé afin de maximiser le débit de la sous porteuse n .

A la différence de [Koutsopoulos_02], l'ordre de traitement des sous porteuses n'est pas prédéterminé dans [Han_03]. Nous décrivons ci après le processus d'affectation.

A la fin de l'allocation d'une sous porteuse, on cherche le couple (n', u') qui présente le meilleur rapport signal sur interférence parmi l'ensemble des sous porteuses disponibles et l'ensemble V des utilisateurs non satisfaits ($r_u < r_{min,u}$). La sous porteuse n' est alors sélectionnée et l'utilisateur u' est le premier utilisateur de $\Pi_{n'}$. Tant que la somme des débits des utilisateurs de $\Pi_{n'}$ augmente, l'utilisateur présentant le meilleur rapport sur interférence (dans l'ensemble V) est ajouté à $\Pi_{n'}$ (sachant que les utilisateurs de l'ensemble «co-canal» ne doivent pas être de la même cellule).

Si l'ensemble V devient vide (tous les utilisateurs sont satisfaits) et qu'il reste des sous porteuses disponibles, la sélection des prochains utilisateurs est désormais réalisée sur tous les utilisateurs. C'est une phase opportuniste. Les meilleurs utilisateurs de chaque cellule peuvent recevoir toutes les sous porteuses restantes. Le lecteur intéressé par la phase d'amélioration qui suit trouvera une description dans [Han_03].

3. Un nouvel algorithme en RA : l'OSA-IL

3.1. Formulation du problème

Le problème que l'on cherche à résoudre est la maximisation du débit global (problème RA) de la cellule sous contrainte de débit minimaux pour les utilisateurs, c'est à dire :

$$\max_{u,b} \sum_{b=1}^B \sum_{u=1}^U \sum_{n \in \Omega_{u,b}} \log_2(1 + \gamma_{u,b}^{(n)})$$

$$r_{u,b} \geq r_{u,b}^o \quad \text{avec } 1 \leq u \leq U \text{ et } 1 \leq b \leq B$$

On utilise une fonction $r = h(\gamma)$ de type formule de *Shannon*. Dans chacune des B cellules, il y a U utilisateurs. Le SINR (*Signal to Interference and Noise Ratio*) perçu par l'utilisateur u de la cellule b sur la sous porteuse n est noté $\gamma_{u,b}^{(n)}$. L'ensemble des sous porteuses affectées à l'utilisateur u de la cellule b est appelé $\Omega_{u,b}$. Enfin, le débit minimal d'un utilisateur u de la cellule b est noté $r_{u,b}^o$.

3.2. Motivations

Notre contribution adapte le FRF de façon dynamique selon la sous porteuse traitée. Parmi les algorithmes étudiés dans la littérature, seuls les travaux de [Han_03] résolvent le problème formulé ci dessus en adoptant un FRF adaptatif.

L'algorithme proposé par [Han_03] fournit de bons résultats mais cet algorithme est coûteux en temps d'exécution car il effectue des recherches exhaustives. A chaque étape, on doit déterminer la sous porteuse dont on va construire l'ensemble co-canal. Pour cela, il faut faire la recherche du meilleur couple (n', u') sur l'ensemble H des sous porteuses disponibles et l'ensemble V des utilisateurs non satisfaits. De plus, une fois la

sous porteuse n' fixée, l'ensemble co-canal se construit en cherchant l'utilisateur présentant le meilleur rapport sur interférence dans l'ensemble V (utilisateurs non satisfaits de toutes les cellules). On peut éviter ces recherches exhaustives.

Notre motivation est de proposer un algorithme heuristique utilisant un FRF adaptatif et présentant une complexité inférieure à l'algorithme de [Han_03].

3.3. Objectifs et «philosophie» de l'algorithme

Pour résoudre le problème ci-dessus, notre objectif est d'optimiser la réutilisation des ressources tout en assurant un débit minimal aux utilisateurs. Les sous porteuses peuvent avoir des FRF différents.

Pour maximiser la réutilisation des ressources, nous tentons tout d'abord d'utiliser un FRF unité. Si ce FRF ne réalise pas les débits minimaux, l'algorithme modifie le niveau d'interférences en cherchant un FRF adapté. L'algorithme s'appelle OSA-IL pour *Opportunist Subcarrier Allocation with Interference Limitation* ([LengGodM_06]).

Une allocation initiale alloue chaque sous porteuse avec un FRF unité au meilleur utilisateur insatisfait de chaque cellule. Si à la fin de cette phase, le seuil de violation des débits minimaux est dépassé, il est nécessaire de mieux gérer les interférences. Une phase de limitation d'interférences est alors lancée. Son objectif est d'adapter le niveau d'interférence à l'utilisateur traité afin d'optimiser la réutilisation des ressources. En effet, tous les utilisateurs ne sont pas en mesure de supporter le même FRF¹. L'algorithme parcourt les utilisateurs selon la qualité de leurs conditions radio. Les utilisateurs ayant de bonnes conditions de propagation sont plus faciles à satisfaire, ils toléreront un FRF élevé. Pour les autres utilisateurs, le FRF le plus adapté est déterminé de façon itérative. Le critère utilisé est le maintien d'un SINR seuil sur l'ensemble «co-canal» d'une sous porteuse.

3.4. Description de l'algorithme

L'ordre de traitement des sous porteuses est prédéterminé. Par simplicité on considère l'ordonnancement par index croissant. Une fois une sous porteuse n choisie, on construit l'ensemble co-canal de cette sous porteuse.

L'ordre de traitement pourrait être amélioré mais il est évident qu'il existe un compromis entre la complexité de l'algorithme et les performances obtenues.

3.4.1. Phase d'allocation opportuniste avec débits minimaux

Une sous porteuse est allouée dans toutes les cellules. Cela signifie que le facteur de réutilisation est égal à 1 pour toutes les sous porteuses. Dans une cellule, la règle d'allocation attribue la sous porteuse en cours au meilleur utilisateur insatisfait. S'il n'existe pas d'utilisateur insatisfait, la sous porteuse est allouée au meilleur utilisateur de la cellule. Lorsque toutes les sous porteuses ont été traitées, le taux de violation de

¹ Par exemple, le FRF supporté par un utilisateur de faible débit minimal avec de bonnes conditions de propagation est plus élevé que celui d'un utilisateur ayant un débit minimal élevé et de mauvaises conditions de propagation.

débats minimaux (RRV pour *Rate Requirement Violation ratio*) est calculé sur l'ensemble des utilisateurs du système. S'il dépasse un seuil η , la phase de limitation d'interférence est déclenchée sinon l'allocation est terminée.

3.4.2. Phase avec limitation d'interférence

Allocation intra-cellulaire

Pendant le traitement d'une sous porteuse n , lorsqu'une cellule b est considérée, le principe d'allocation consiste à l'attribuer au meilleur utilisateur insatisfait. S'il n'existe pas d'utilisateur insatisfait, la sous porteuse est affectée au meilleur utilisateur de la cellule.

Dans cet algorithme, on voit que les utilisateurs dont les conditions de propagation sont bonnes sont traités avant les utilisateurs dont les conditions de propagation sont médiocres. C'est une remarque importante car elle est à l'origine du choix du FRF utilisé pour l'affectation d'une sous porteuse.

Allocation multicellulaire : principes

Tant qu'il reste des utilisateurs insatisfaits (ensemble V) l'algorithme modifie (si besoin) le FRF qui caractérise l'affectation d'une sous porteuse n . Le nombre de cellules qui utilise la sous porteuse dépend du FRF en cours. La sélection des cellules est faite en donnant la priorité aux cellules qui ont les meilleurs utilisateurs non satisfaits (cf. détails en annexe 5.B).

Lorsque tous les utilisateurs sont satisfaits, toutes les sous porteuses qui ne sont pas encore attribuées sont toutes affectées avec un FRF unité.

Allocation multicellulaire : critère de choix du FRF

Les valeurs prédéterminées de FRF sont classées dans un vecteur par ordre décroissant. Une sous porteuse n est affectée avec le FRF en cours. Il commence à 1 et peut diminuer tant qu'il y a des utilisateurs insatisfaits ($V \neq \emptyset$).

Le changement du FRF courant (son maintien ou sa réduction) dépend du SINR des utilisateurs des ensemble co-canaux des sous porteuses précédentes. On considère que le FRF courant est inadapté lorsque le SINR moyen (des utilisateurs des ensembles co-canaux) de C_{max} sous porteuses consécutives est inférieur à un seuil δ . Comme vu précédemment, les meilleurs utilisateurs sont traités en premier. Ainsi, les conditions radio des utilisateurs de l'ensemble co-canal I_{n-1} sont à priori meilleures que celles des conditions radio des utilisateurs qu'il reste à satisfaire. Lorsque le FRF courant est jugé inadapté, il est diminué (valeur suivante dans le vecteur des FRF) afin de baisser le niveau d'interférences sur les sous porteuses qui ne sont pas encore attribuées. Cela permet aux utilisateurs insatisfaits de bénéficier de meilleures conditions sur les sous porteuses restantes.

Un cas exceptionnel se produit quand le SINR des utilisateurs d'une sous porteuse reste inférieur à ε bien que le FRF courant soit le plus faible. Dans ce cas, la sous porteuse peut exceptionnellement être utilisée dans toutes les cellules.

Le lecteur intéressé par l'implémentation de cet algorithme trouvera un *flow chart* dans l'annexe 5.B.

3.5. Illustration sur un cas d'école

On considère la situation suivante, on dispose d'une bande constituée de 10 sous porteuses F_i $1 \leq i \leq 10$, que l'on peut affecter avec les FRF 1, 2/3 et 1/3. Six utilisateurs $\omega_{u,b}$, $1 \leq u \leq 2$, $1 \leq b \leq 3$ ont des débits minimaux différents à satisfaire. On se situe dans le cas où l'affectation de toutes les sous porteuses avec le FRF 1 ne permet pas d'obtenir un taux de satisfaction suffisant ($RRV < \eta$).

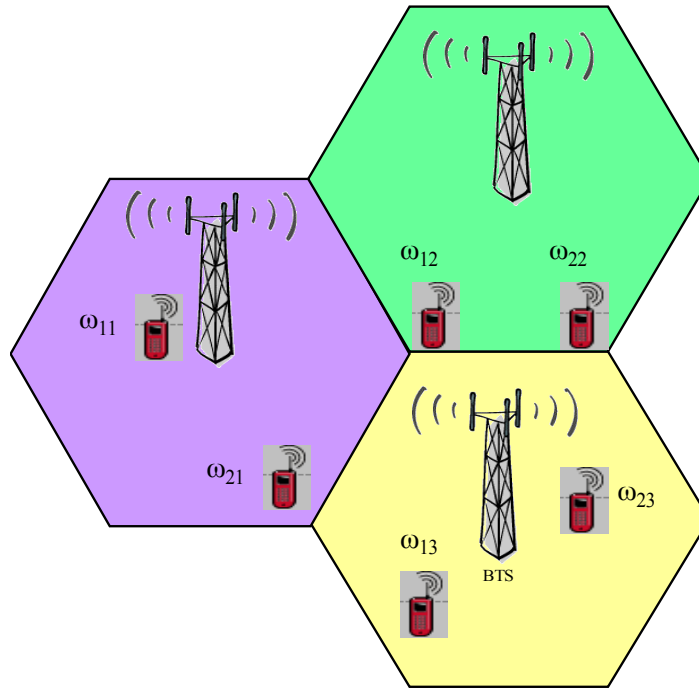


Figure 5.2: Exemple de configuration cellulaire avant l'application de l'OSA-IL

Application de la phase de limitation d'interférence

La fréquence F_1 est affectée avec le FRF 1. Elle sera donc utilisée dans toutes les cellules. Dans chaque cellule, l'utilisateur qui a le meilleur SINR sur F_1 parmi les insatisfaits (c'est à dire ceux qui n'ont pas encore atteint leur débit minimal) est sélectionné. On a alors $\Pi_1 = \{\omega_{11}, \omega_{12}, \omega_{23}\}$. Le SINR moyen de cet ensemble est supérieur au seuil δ .

La fréquence F_2 est donc affectée avec le même FRF c'est à dire le FRF 1. Les utilisateurs qui ont le meilleur SINR sur F_2 parmi les insatisfaits sont $\Pi_2 = \{\omega_{11}, \omega_{22}, \omega_{23}\}$. Le SINR moyen de cet ensemble est inférieur au seuil δ . On doit changer de FRF (dans cet exemple $C_{max} = 1$).

La fréquence F_3 est alors affectée avec le FRF suivant c'est à dire le FRF 2/3. Elle ne sera utilisée que sur 2 cellules. Dans chaque cellule, l'utilisateur qui a le meilleur SINR sur F_3 parmi les insatisfaits est candidat : les deux meilleurs sont choisis et cela désigne

automatiquement les cellules qui utilisent F_3 . On a $\Pi_3 = \{\omega_{11}, \omega_{23}\}$. Le SINR moyen de cet ensemble est supérieur au seuil δ .

La fréquence F_4 est affectée avec le même FRF c'est à dire le FRF 2/3. On a $\Pi_4 = \{\omega_{11}, \omega_{13}\}$. A la suite de cette allocation, l'utilisateur ω_{11} atteint son débit minimal. Tant qu'il y a des insatisfaits dans la cellule 1, ω_{11} est écarté de l'allocation. Le SINR moyen de Π_4 est supérieur au seuil δ .

La fréquence F_5 est affectée avec le FRF 2/3. On a $\Pi_5 = \{\omega_{21}, \omega_{23}\}$. A la suite de cette allocation, l'utilisateur ω_{23} atteint son débit minimal. Tant qu'il y a des insatisfaits dans la cellule 3, ω_{23} est écarté de l'allocation. Le SINR moyen de Π_5 est supérieur au seuil δ .

La fréquence F_6 est affectée avec le FRF 2/3. On a $\Pi_6 = \{\omega_{12}, \omega_{13}\}$. A la suite de cette allocation, l'utilisateur ω_{13} atteint son débit minimal. Il n'y a plus d'insatisfaits dans la cellule 3. Le SINR moyen de Π_5 est supérieur au seuil δ .

La fréquence F_7 est affectée avec le FRF 2/3. On a $\Pi_7 = \{\omega_{21}, \omega_{22}\}$. A la suite de cette allocation, l'utilisateur ω_{21} atteint son débit minimal. Il n'y a plus d'insatisfaits dans la cellule 1. Le SINR moyen de Π_7 est inférieur au seuil δ .

La fréquence F_8 est alors affectée avec le FRF suivant c'est à dire le FRF 1/3. Elle ne sera utilisée que dans une cellule. On l'utilise dans la cellule 2 car c'est la seule où il y a des utilisateurs insatisfaits. On a $\Pi_8 = \{\omega_{12}\}$. Le SINR moyen de Π_8 est supérieur au seuil ε (sinon on pourrait utiliser pour cette sous porteuse le FRF 1 pour augmenter le débit des autres utilisateurs). A la suite de cette allocation, l'utilisateur ω_{12} est satisfait.

La fréquence F_9 est affectée avec le FRF 1/3. On a $\Pi_9 = \{\omega_{22}\}$. A la suite de cette allocation, l'utilisateur ω_{12} est satisfait.

Finalement, pour chaque utilisateur les ensembles de sous porteuses reçues sont : $\mathcal{Q}_{11} = \{F_1 F_2 F_3 F_4\}$; $\mathcal{Q}_{21} = \{F_5 F_7\}$; $\mathcal{Q}_{12} = \{F_1 F_6 F_8\}$; $\mathcal{Q}_{22} = \{F_2 F_7 F_9\}$; $\mathcal{Q}_{13} = \{F_4 F_6\}$ et $\mathcal{Q}_{23} = \{F_1 F_2 F_3 F_5\}$.

Remarque importante

L'algorithme OSA-IL s'applique sans modifications lorsque les cellules n'ont pas le même nombre d'utilisateurs. Lorsqu'une cellule a tous ses utilisateurs satisfaits, elle peut encore recevoir des sous porteuses pour améliorer le débit global mais elle n'est pas prioritaire par rapport aux cellules dont les utilisateurs sont insatisfaits.

4. Simulations et résultats

4.1. Algorithmes comparés

Nous avons implémenté l'algorithme proposé dans [Koutsopoulos_02], il est noté *Heuristic A*. Cet algorithme est complètement opportuniste car il ne considère pas de débits minimaux pour les utilisateurs, il nous fournit une borne maximale sur le débit.

L'algorithme proposé dans [Han_03] nous sert de référence pour comparer les performances de notre algorithme, il est appelé *Heuristic B*. Nous ne considérons que la version *maximum packing* de la phase d'allocation initiale(cf. 2.2.2).

Une version de l'algorithme OSA-IL ([LengGodM_06]) sans phase de limitation d'interférences sera notée OSA. Comparer l'OSA et l'OSA-IL permet de voir l'intérêt de la phase de limitation d'interférences.

Un dernier algorithme, appelé *Fair Heuristic*, est examiné. Cet algorithme cherche à réaliser une équité en débit parmi les utilisateurs, il fournit une borne minimale sur le débit moyen d'une cellule. Il s'agit d'une extension de l'algorithme proposé dans [RheeCioffi_00] pour un contexte monocellulaire. Au cours de l'algorithme, l'utilisateur au plus faible débit reçoit sa meilleure sous porteuse. Pour l'étendre au contexte multicellulaire, nous appliquons l'algorithme dans chaque cellule en considérant un $\text{FRF} = 1$ pour toutes les sous porteuses.

4.2. Modèle canal et Paramètres de simulation

Nous considérons B cellules carrées suivant un modèle en tore (pour éviter les effets de bords). Dans chaque cellule, il y a U utilisateurs. Un utilisateur ω est caractérisé dans le système par ses coordonnées (u, b) avec $1 \leq u \leq U$ et $1 \leq b \leq B$. Ainsi, u est l'index local de l'utilisateur dans sa cellule b . Les utilisateurs sont uniformément répartis dans les cellules. Un utilisateur situé au delà d'une distance $D_{\max}$ de sa BS est appelé utilisateur lointain. Nous accordons un souci particulier à la probabilité d'échec (ou *outage*) de ce type d'utilisateur.

Le modèle canal adopté est constitué de N sous porteuses parallèles sur la bande W . Aucune corrélation n'est considérée entre les sous porteuses. Le carré du module du gain d'une sous porteuse n entre l'utilisateur ω et la BS i est déterminé suivant un modèle classique : $|G_{i,\omega}^{(n)}|^2 = K d(\omega, i)^{-\alpha} a_{sh}(\omega) a_f(\omega, n)$. Dans cette formule (selon des notations similaires au chapitre 3), K est l'atténuation à la distance de référence 1 m, $d(\omega, i)$ est la distance entre l'utilisateur ω et la BS i , α est le coefficient d'atténuation, $a_{sh}(\omega)$ est une variable lognormale qui traduit l'effet de masque subi par toutes les sous porteuses de l'utilisateur ω , enfin, $a_f(\omega, n)$ est une variable exponentielle pour modéliser l'évanouissement rapide (*fast fading*) de la sous porteuse n .

Le niveau d'interférence reçu sur la sous porteuse n par l'utilisateur u de la cellule b est la somme des signaux émis par les autres cellules utilisant cette sous porteuse :

$I_{u,b}^{(n)} = p \sum_{i \neq b} G_{i,\omega}^{(n)}$ où p est la puissance par sous porteuse et $\omega = (u, b)$. Le bruit blanc gaussien a pour variance $\sigma^2 = N_0 W / N$. Ainsi le SINR perçu par l'utilisateur u de la cellule b sur la sous porteuse n est donné par $\gamma_{u,b}^{(n)} = G_{b,\omega}^{(n)} / (I_{u,b}^{(n)} + \sigma^2)$.

Le problème que l'on cherche à résoudre est la maximisation du débit global de la cellule avec des contraintes de débit minimal pour chaque utilisateur. Pour simplifier la présentation des résultats, les simulations sont effectuées avec un débit minimal r° commun à tous les utilisateurs.

Le débit minimal r° varie de 16 kbps à 2 Mbps. On moyenne les résultats obtenus (débit moyen d'une cellule, taux d'utilisateurs insatisfaits) sur 1000 *snapshots* pour chaque valeur de r° .

Tableau 5.1 : Paramètres de simulation, allocation multicellulaire

Paramètres	Valeur
U : nombre d'utilisateurs par cellule	10
B : nombre de cellules	9
Taille d'une cellule(km)	2
D_{max} : distance limite entre BS et un utilisateur proche (km)	1
η : seuil de RRV dans l'OSA-IL	0.001
δ : seuil sur le SINR moyen des utilisateurs co-canaux dans l'OSA-IL	3
ε : seuil sur le SINR minimum des utilisateurs co-canaux dans l'OSA-IL	-3
C_{max} : nombre limite de sous porteuses de SINR inférieur à δ dans l'OSA-IL	5
Fréquence centrale (GHz)	3.3
W (MHz)	5
N : nombre de sous porteuses	512
f_{rf} : liste de FRF utilisés dans l'OSA-IL	[1 8/9 7/9 6/9 5/9 4/9 1/3]
r° : débit minimal(kbps)	16...2000
α : facteur d'atténuation	4
σ_{sh} : écart type de la variable log normale	6
Puissance par sous porteuse (dBm)	15
N_0 (dBm/Hz)	-174

4.3. Débit moyen d'une cellule

Sur la figure 5.3, on trace le débit moyen d'une cellule en fonction du débit minimal r° . Sur la figure 5.4, on suit l'évolution du FRF moyen d'une sous porteuse en fonction du débit minimal.

Les performances des algorithmes *Heuristic B* et OSA et OSA-IL sont encadrées par les cas extrêmes *Heuristic A* et *Fair Heuristic*. Ils présentent respectivement la meilleure et la moins bonne performance vis à vis du débit moyen d'une cellule. Leurs résultats sont constants car ces deux algorithmes ne considèrent pas de contraintes sur le débit des utilisateurs.

Lorsqu'on compare les algorithmes OSA et OSA-IL, on peut voir que la phase de limitation d'interférences améliore le débit moyen d'une cellule. On voit que l'*Heuristic B* présente un meilleur débit que l'OSA-IL (au prix d'une plus grande complexité, cf section 4.5).

Pour ces trois algorithmes, le débit d'une cellule diminue avec la contrainte r° jusqu'à un seuil que l'on note r_l° (cf. Fig. 5.3). Garantir un débit r° de plus en plus grand provoque naturellement une baisse du débit de la cellule. Les utilisateurs qui ont de faibles conditions radio ont besoin soit de plus de sous porteuses soit de sous porteuses de meilleure qualité. On voit, Fig. 5.4, que jusqu'à r_l° , le FRF moyen d'une sous porteuse baisse (donc la qualité s'améliore) en OSA-IL et dans l'*Heuristic B*. C'est une condition nécessaire pour réaliser un bon débit avec des utilisateurs qui ont de faibles conditions radio.

Le seuil r_l° dépend de l'algorithme. En fait, r_l° représente le plus grand débit minimal que ces algorithmes peuvent garantir aux utilisateurs. Au delà de ce seuil, l'*outage* devient important. Pour un *outage* de 2 %, on observe $r_l^\circ = 1,2$ Mbps pour *Heuristic B* et $r_l^\circ = 800$ kbps pour l'OSA et l'OSA-IL.

Au delà du seuil r_l° , le débit moyen de la cellule se remet à augmenter. Il y a une correspondance entre la courbe de débit (cf. Fig. 5.3) et celle de FRF moyen (cf. Fig. 5.4) pour l'OSA_IL et l'*Heuristic B*. Comme la contrainte est élevée, les meilleurs utilisateurs insatisfaits sont toujours les mêmes. Dans ce cas, le débit de la cellule augmente car seuls les meilleurs utilisateurs reçoivent des sous porteuses. Au delà de $r^\circ = 1.2$ Mbps, les intervalles de confiance à 95% (non représentés ici) de l'OSA, l'OSA_IL et l'*Heuristic B* se chevauchent : pour des contraintes trop élevées les algorithmes se valent.

Enfin, on peut comparer l'OSA et le *Fair Heuristic*. Les deux algorithmes appliquent un FRF de 1 à toutes les sous porteuses. La comparaison de ces deux algorithmes permet d'observer le coût en débit global causé par la recherche d'une équité complète en débit.

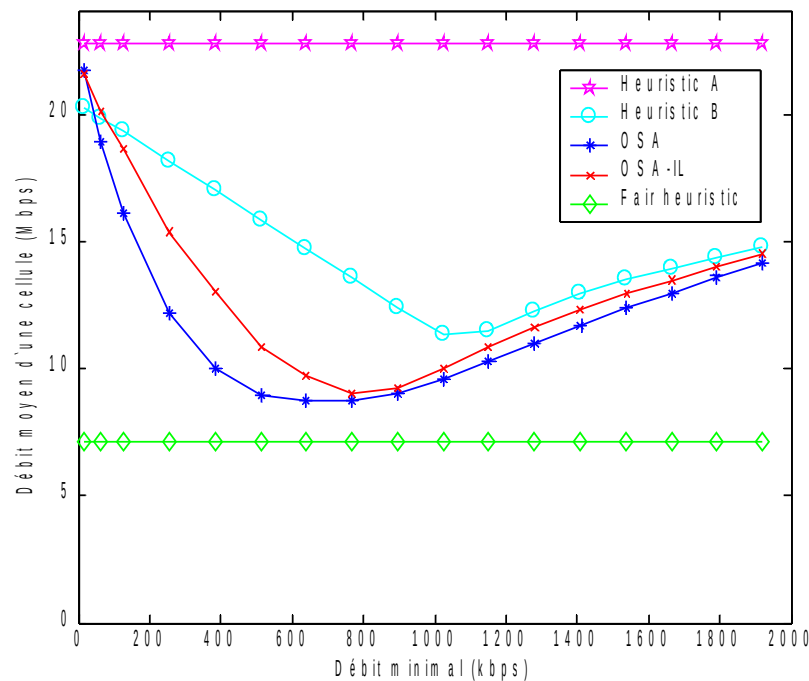


Figure 5.3 : Débit moyen d'une cellule en fonction du débit minimal

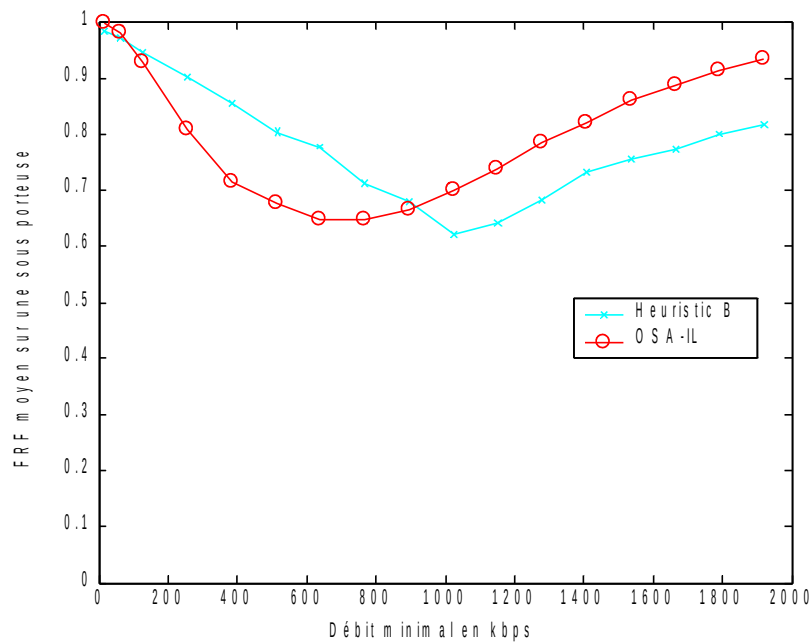


Figure 5.4 : Évolution du FRF moyen d'une sous porteuse en fonction du débit minimal requis

4.4. Probabilité d'échec (outage)

On calcule la proportion d'utilisateurs qui ne sont pas satisfaits sur l'ensemble des utilisateurs du système (cf. Fig. 5.5). La comparaison entre l'OSA et l'OSA-IL montre que la proportion d'utilisateurs non satisfaits est améliorée grâce à la phase de limitation d'interférences. Cela permet à l'OSA-IL de se rapprocher des performances de l'Heuristic B, du point de vue de l'outage, avec une complexité inférieure (cf. section 4.5).

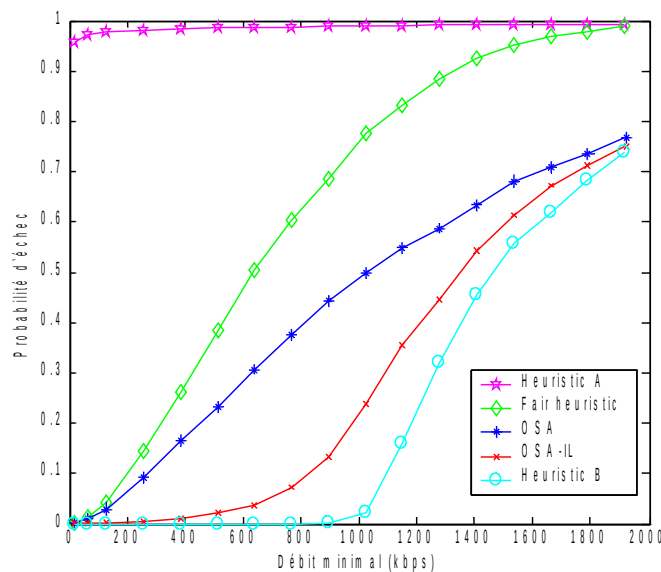


Figure 5.5 : Comparaison de la probabilité d'échec sur tous les utilisateurs

On s'intéresse aux utilisateurs dits lointains qui sont bien sûr défavorisés. Les valeurs d'*outage* sont plus élevées (que sur la Fig. 5.5) car les conditions de propagation des utilisateurs lointains sont médiocres. La proportion des utilisateurs non satisfaits parmi ce type d'utilisateurs (cf. Fig. 5.6) est améliorée grâce à la phase de limitation d'interférences (comparaison entre OSA et OSA-IL).

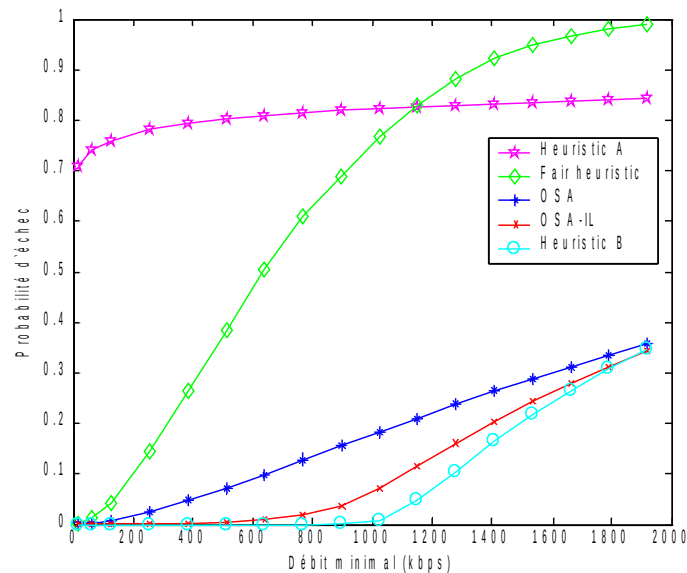


Figure 5.6 : Comparaison de la probabilité d'échec sur les utilisateurs lointains

4.5. Évaluation de la complexité

Pour évaluer la complexité, on compare les temps d'exécution des algorithmes *Heuristic B* et OSA-IL. Sur la figure 5.7, on trace le rapport entre le temps d'exécution de ces deux algorithmes. On fixe le nombre d'utilisateurs par cellule à 12 et on fait varier le nombre de sous porteuses. On voit que l'*Heuristic B* a un temps d'exécution plus important que l'OSA-IL. On a pu établir par ailleurs que le temps d'exécution est plus sensible au nombre de sous porteuses qu'au nombre des utilisateurs (cf. [LengGodM_06]).

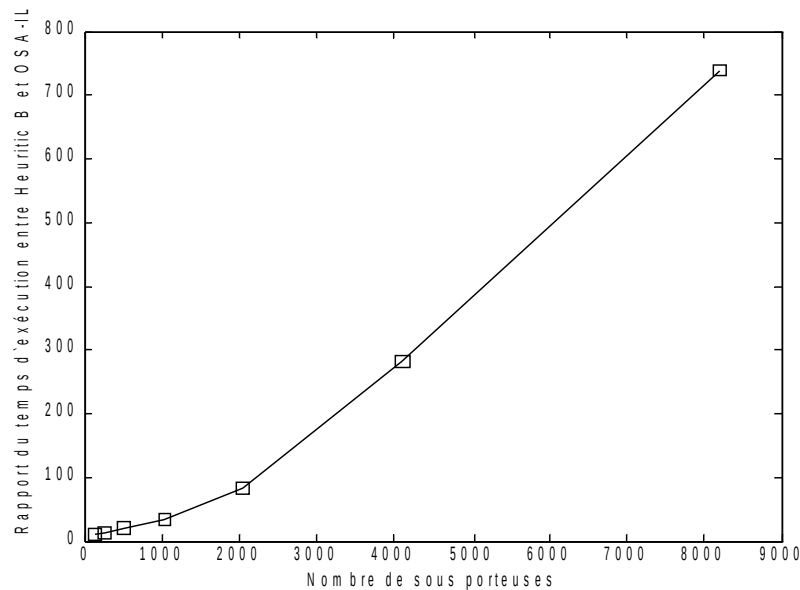


Figure 5.7 : Comparaison des temps d'exécution entre l'heuristique B et de l'OSA-IL

5. Conclusions

5.1. Résultats

Nous avons proposé un algorithme l'OSA-IL pour l'affectation des sous porteuses sur la voie descendante d'un système multicellulaire. Notre algorithme s'approche des performances de l'algorithme décrit dans [Han_03] (*Heuristic B*) en terme de probabilité d'*outage* avec une complexité inférieure. Pour $N = 2048$ (et 12 utilisateurs dans chacune des 9 cellules), le temps d'exécution est de l'*Heuristic B* est 100 fois plus élevé que celui de l'OSA-IL.

Le débit moyen d'une cellule de l'OSA-IL est inférieur à celui obtenu par l'*Heuristic B*. Sur une bande de 5 MHz, l'*Heuristic B* garantit un débit individuel de 500 kbps aux dix utilisateurs de chacune des neuf cellules ; le débit de chaque cellule étant de 16 Mbps. Dans les mêmes conditions, le débit moyen par cellule obtenu par l'OSA-IL est de 11 Mbps.

Lorsqu'on compare l'OSA-IL à l'OSA, on voit que le débit moyen par cellules est meilleur mais c'est surtout la probabilité d'*outage* qui est améliorée. On en déduit qu'à puissance constante sur les sous porteuses, la phase de limitation d'interférences agit surtout sur la probabilité d'*outage*.

Tant que le débit minimum peut être réalisé avec une probabilité d'*outage* satisfaisante, le FRF moyen d'une sous porteuse baisse lorsque le débit minimum augmente. Dans notre contexte de simulation, un nombre moyen de 45 sous porteuses de FRF moyen de $2/3$ permet de réaliser un débit individuel entre 800 kbps et 1 Mbps.

5.2. Améliorations

Dans le sens descendant d'un système multicellulaire, un travail intéressant serait de mettre en place une version avec contrôle de puissance pour les utilisateurs en bord de cellules. Cela permettrait éventuellement à la phase de limitation d'interférences, en plus de la probabilité d'*outage*, d'améliorer significativement le débit moyen par cellule.

5.3. Évolutions du contexte d'étude

Dans cette première partie de la thèse, nous avons essayé de maximiser le débit d'une cellule en assurant aux utilisateurs un débit minimal. L'allocation des ressources a été considérée comme indépendante d'un instant à l'autre. Les résultats ont été moyennés sur des *snapshots*. De plus, les queues de utilisateurs ont été considérées comme infinies et les données ont été envoyées sans contraintes de temps.

Dans la deuxième partie de la thèse, nous faisons évoluer ces conditions. L'allocation de ressources sera considérée dans le temps. Les queues des utilisateurs sont désormais finies et les utilisateurs sont regroupés en deux classes. Une classe temps-réel réunit les utilisateurs dont les paquets doivent être transmis dans un certain délai. Ces utilisateurs sont aussi caractérisés par un débit source faible. Une classe non temps réel, réunit des utilisateurs sans contraintes temporelles et dont le débit source est élevé. Pour ces deux classes, on désire affecter les ressources simultanément. La décision prise à un instant dépend désormais des décisions précédentes.

Partie II: Ordonnancement en OFDMA

Dans cette partie, les ressources sont attribuées en fonction de l'état du trafic et des conditions radio. L'état du trafic résulte du débit des sources et des décisions antérieures.

Acronymes (partie 2)

AMC	<i>Adaptive Modulation and Coding</i>
BS	<i>Base Station</i>
CAFQ	<i>Channel Adaptive Fair Queuing</i>
CIF-Q	<i>Channel Independent Packet Fair Queuing</i>
CPLD-PGPS	<i>Channel-Condition Packet Length Dependent Packet Generalized Processor Sharing</i>
CQICH	<i>Channel Quality Indicator Channel</i>
CSDPS	<i>Channel State Dependent Packet Scheduling</i>
CS-WFQ	<i>Channel State Independent Weighted Fair Queuing</i>
EDF	<i>Earliest Deadline First</i>
EWFS	<i>Enhanced Wireless Fair Service</i>
FSMC	<i>Finite State Markov Chain</i>
GPF	<i>Generalized Proportional Fair</i>
GPS	<i>Generalized Processor Sharing</i>
IWFQ	<i>Idealized Wireless Fair Queuing</i>
MCS	<i>Modulation and Coding Scheme</i>
MR-FQ	<i>MultiRate Fair Queuing</i>
MWFS	<i>MultiRate Wireless Fair Scheduling</i>
NRT	<i>Non Real Time</i>
OFDMA	<i>Orthogonal Frequency Division Multiple Access</i>
ORCA_SRT	<i>Optimal Radio Channel Allocation Single Rate Transmission</i>
ORCA_MRT	<i>Optimal Radio Channel Allocation Multi Rate Transmission</i>
OWFQ	<i>Opportunistic Weighted Fair Queuing</i>
OWFS	<i>Opportunistic Wireless Fair Service</i>
PF	<i>Proportional Fair</i>
PGPS	<i>Packetized Generalized Processor Sharing</i>
PLFS	<i>Packet Loss Fair Scheduling</i>
RR	<i>Round Robin</i>
RT	<i>Real Time</i>
SCFQ	<i>Self Clock Fair Queuing</i>
STFQ	<i>Start-Time Fair Queuing</i>

TDD	<i>Time Division Duplexing</i>
TGS	<i>Throughput Guaranteed Scheduling</i>
VC	<i>Virtual Clock</i>
WFQ	<i>Weighted Fair Queuing</i>
W2FQ	<i>Worst case WFQ</i>
WFS	<i>Wireless Fair Scheduling ou Wireless Fair Service</i>

Notations (partie 2)

$A(t)$	ensemble des flux actifs à l'instant t
b	indice sur les stations de base
B	nombre de stations de base
c_m	capacité en bit/s/Hz du mode m
$C(t)$	capacité du canal en bit/s à l'instant t
$D(p_i^k)$	délai d'attente du k ème paquet du flux i jusqu'à transmission complète
$D_{max,i}$	délai maximal d'attente d'un paquet du flux i
$E(i)$	compteur du flux i , spécifie l'avance ou le retard du flux
$E_{min,i}$	limite inférieure du compteur de retard du flux i
$E_{max,i}$	limite supérieure du compteur d'avance d'un flux i
$F(p_i^k)$	étiquette de fin de service du k ème paquet du flux i
i,j	indice sur les flux
H	ensemble des sous porteuses disponibles
k	indice sur le rang d'un paquet
$L(p_i^k)$	longueur du k ème paquet du flux i
L_{max}	longueur maximale d'un paquet du flux i
n	indice sur les sous porteuses
m	indice sur les modes de transmissions (MCS)
$m_{i,n}$	mode de transmission du flux i sur la sous porteuse n
M	nombre de modes de transmissions (MCS)
N	nombre de sous porteuses
N_u	nombre de sous porteuses attribuées à l'utilisateur u
N_{sc}	nombre de sous porteuses dans un sous canal
p_i^k	k ème paquet du flux i

p_i^{HOL}	paquet de tête de queue du flux i
QL_i	longueur de la queue de paquets du flux i
r_i	poids de débit du flux i
R_u°	nombre de bits à recevoir au minimum par l'utilisateur u
R_u	nombre de bits reçus par l'utilisateur u
$R_{u,n}(t)$	nombre de bits reçus par l'utilisateur u sur la sous porteuse n à l'instant t
$R_{u,max}(t)$	nombre de bits maximum de l'utilisateur u sur une sous porteuse
S	nombre de sous canaux
$S(p_i^k)$	étiquette de début de service du k ème paquet du flux i
t_i^k	heure réelle d'arrivée du k ème paquet du flux i
u, v	indices sur les utilisateurs
U	nombre d'utilisateurs
$V(t)$	temps virtuel correspondant au l'heure réelle t
$W_i(t_1, t_2)$	service reçu en nombre de bits par le flux i pendant l'intervalle $[t_1, t_2]$
$\gamma_{m-1,m}^{th}$	SNR seuil entre les modes de transmission $m-1$ et m
$\mu_u(t)$	métrique d'ordonnancement de l'utilisateur u à l'instant t
ρ	paramètre <i>lookahead</i> dans le <i>Wireless Fair Service</i>
Φ_i	poids de délai du flux i
Ψ_{RT}	ensemble des flux temps réel
Ψ_{NRT}	ensemble des flux non temps réel
Ω_u	ensemble des sous porteuses attribuées à l'utilisateur u

Chapitre 6. Contribution pour l'ordonnancement en OFDMA

1. Introduction

1.1. Présentation des objectifs

1.1.1. Contexte d'étude

Dans ce chapitre, nous considérons les techniques d'ordonnancement et d'allocation de ressources en OFDMA. Les problèmes d'équité sont abordés sur un intervalle temporel : les politiques d'allocation dépendent non seulement de l'état du canal mais aussi de l'instant t . On se situe sur la voie descendante d'une cellule unique. Nous travaillons avec des sous canaux (de sous porteuses adjacentes). La répartition de puissance est uniforme sur les sous porteuses. Les utilisateurs génèrent des paquets et les queues sont de taille finie. Il peut y avoir plus d'utilisateurs actifs (i.e. dont les queues sont non vides) que de sous canaux disponibles. Nous considérons deux types de trafic différents : un trafic temps réel et un trafic non temps réel. L'objectif est de transmettre les paquets du trafic temps réel en respectant leurs contraintes, de maximiser le débit global du trafic non temps réel tout en respectant une équité proportionnelle entre flux.

Ce contexte est plus large que celui étudié dans la première partie où on considérait U utilisateurs qui avaient des données à transmettre. Cela représente deux cas de figure : soit il s'agit toujours des mêmes utilisateurs qui transmettent sans interruption soit une sélection a été réalisée en amont pour choisir ceux qui vont transmettre. Nous nous étions placés dans la situation où le nombre N de sous porteuses est nettement supérieur au nombre U d'utilisateurs actifs. Nous avons également considéré des *snapshots* du système : l'allocation réalisée dans un *snapshot* est indépendante des autres instants. On peut dire que les politiques d'allocation que l'on a étudié auparavant sont stationnaires, ne dépendant que de l'état du canal et non de l'instant t .

1.1.2. Garanties des techniques d'ordonnancement

Nous étudions les techniques d'ordonnancement pour réaliser un partage des ressources dans le temps. Les techniques seront présentées selon les critères ci-après.

a) Équité

Réaliser l'équité proportionnelle sur un lien radio est plus difficile que sur un lien filaire. On distinguera l'équité à long terme et l'équité à court terme.

Équité à long terme

Lorsqu'un flux ne peut pas transmettre à cause de mauvaises conditions radio, il est souvent remplacé par un flux bénéficiant de conditions radio plus favorables. Pour réaliser une équité à long terme, on doit garder en mémoire ces échanges. On mémorise les opportunités de transmission qui ont été cédées par un flux (ce dernier a accumulé une sorte de retard) ainsi que celles qui ont été reçues par les autres flux (qui ont accumulé une sorte d'avance). L'équité à long terme assure aux flux en retard qu'ils recouvriront le service perdu, pourvu que la période de temps soit suffisamment longue. La façon dont les flux en avance cèdent leur avance afin que les flux en retard rattrapent leur retard constitue le modèle de compensation. Les algorithmes proposent différents modèles de compensation.

Équité à court terme

Pour réaliser une équité à court terme, la différence de service entre deux flux actifs doit être bornée sur tout intervalle de temps. L'équité à court terme est fortement compromise par la nature du lien radio. Une période (plus ou moins longue) de non transmission (ou de «famine»), due à un mauvais canal ou au fait que d'autres flux rattrapent leur retard, est un risque pour l'équité à court terme. La «dégradation gracieuse de service» (en abrégé DGS) assure que les flux en avance ne connaîtront pas une période de famine pour que d'autres flux rattrapent intégralement leur retard. Selon le principe DGS, un flux en avance doit, dans un intervalle de temps, recevoir une fraction de service correspondant au modèle de service sans erreurs.

b) Garanties de service

Outre les problèmes d'équité, on souhaite traiter des flux de caractéristiques différentes. Les algorithmes doivent fournir des garanties en terme de délai pour les applications temps réels.

1.1.3. *Organisation du chapitre*

Pour comprendre les motivations de nos contributions en OFDMA, nous présentons un état de l'art des techniques d'ordonnancement dans le domaine filaire et sans fils en section 1.2. Dans la section 1.3. nous présentons les travaux d'ordonnancement existants en OFDMA. Certains algorithmes sont décrits en annexe pour ne pas alourdir l'exposé. Enfin, notre contribution est présentée dans la section 2 et ses performances sont évaluées dans la section 3.

1.2. État de l'art : ordonnancement dans les réseaux avec et sans fils

1.2.1. *Ordonnancement dans les réseaux filaires*

Les techniques d'ordonnancement à partage de bande tentent d'assurer un accès équitable aux ressources aux différents flux. En effet, en cas de pénurie de ressources, ces techniques empêchent un flux de transmettre plus que le taux qui lui est alloué. On distingue les algorithmes GPS (*Generalized Processor Sharing*, [Jeinroek_76]), WFQ (*Weighted Fair Queuing*, [ParekhGallager_93]), RR (*Round Robin*, [Nagle_87]), VC (*Virtual Clock*, [Zhang_90]), SCFQ (*Self Clock Fair Queuing*, [Golestani_94]), W2FQ (*Worst case WFQ*, [ZhangBennett_96]) et STFQ (*Start-Time Fair Queuing*,

[Goyal&_97]).

Le GPS (*Generalized Processor Sharing*, [Jainroek_76]) assure une équité proportionnelle grâce à la notion de poids. Soit r_i le poids du flux i , et $W_i(t_1, t_2)$ le service reçu par le flux i dans l'intervalle temporel $[t_1, t_2]$. Le GPS garantit à deux flux i et j continuellement actifs dans l'intervalle $[t_1, t_2]$ une équité proportionnelle à leur poids, à savoir :

$$\frac{W_i(t_1, t_2)}{W_j(t_1, t_2)} \geq \frac{r_i}{r_j} \quad (6.1)$$

Pour chaque flux, la garantie de service est indépendante des autres flux. Dans le GPS, les flux peuvent être servis simultanément ce qui n'est pas le cas dans un système réel.

Le WFQ est une version paquet du GPS, il introduit la notion de temps virtuel. Ce dernier compte le nombre de tours de *scheduling* écoulés. Durant un tour, chaque flux actif reçoit r_i bits (cf. détails en 2.1.1). Le temps virtuel $V(t)$ est nul lorsque tous les flux sont inactifs. Lorsqu'au moins un flux est actif, l'avancée de $V(t)$ est régie par la formule :

$$\frac{dV(t)}{dt} = \frac{C_{server}}{\sum_{j \in A(t)} r_j} \quad (6.2)$$

Les étiquettes de début et de fin de service représentent respectivement le temps virtuel auquel (i.e. le nombre de tours de *scheduling*) la transmission d'un paquet commence et se termine en GPS. En WFQ, les paquets sont transmis par étiquette de fin de service (*virtual finish tag*) croissante. Les étiquettes de début et de fin de service du $k^{ième}$ paquet du flux i sont respectivement notées $S(p_i^k)$ et $F(p_i^k)$ (*virtual start / finish tag*) :

$$S(p_i^k) = \max(V(t_i^k), F(p_i^{k-1})) \quad (6.3)$$

$$F(p_i^k) = S(p_i^k) + \frac{L(p_i^k)}{r_i} \quad (6.4)$$

Dans l'équation (6.3), t_i^k est l'heure réelle d'arrivée du paquet p_i^k .

Les garanties en débit et en délai sont fortement liées dans le WFQ. Le PWFQ (*Priority-based Weighted Fair Queueing*, [Wang&_02]) et le (m, k) WFQ ([Koubaa_04]) tentent de découpler ces garanties (cf. annexe 6 A).

D'autres versions paquet du GPS existent à savoir le VC, le SCFQ, le W2FQ et le STFQ. Ces algorithmes ne découplent pas les garanties en débit et en délai. Le VC et le SCFQ modifient la définition du temps virtuel tandis que le W2FQ et le STFQ ordonnancent en fonction de l'étiquette de début de service (cf. annexe 6A). Dans le STFQ et le SCFQ, la différence de service reçue par deux flux vérifie :

$$\left| \frac{W_i(t_1, t_2)}{r_i} - \frac{W_j(t_1, t_2)}{r_j} \right| \leq \frac{L_{i,max}}{r_i} + \frac{L_{j,max}}{r_j} \quad (6.5)$$

1.2.2. Ordonnancement dans les réseaux sans fils

Les algorithmes précédents (WFQ, SCFQ, STFQ ...) ne peuvent être appliqués directement au canal radio car il n'est pas aussi fiable que le lien filaire. Donner l'accès au canal à un utilisateur qui a de mauvaises conditions radio revient à gâcher les ressources. Des algorithmes ont été introduits pour émuler la politique GPS dans un contexte radio. Ces algorithmes s'inspirent de ceux définis dans le domaine filaire et introduisent un modèle dit de compensation. Le but du mécanisme de compensation est de contrebalancer les inégalités introduites par le lien radio. Le modèle de compensation gère le transfert entre les opportunités de transmission d'un terminal percevant un mauvais canal aux terminaux qui perçoivent un bon canal. Certains flux accumulent une sorte d'avance tandis que d'autres accumulent un retard. Lorsqu'un terminal recouvre un bon canal, le modèle de compensation lui transfère des opportunités de transmission pour compenser celles qu'il a cédées. Les modalités exactes de transfert d'opportunités de transmission diffèrent selon les algorithmes et le modèle de canal.

Généralement, les algorithmes adaptés aux réseaux sans fils sont composés :

- d'un modèle de service sans erreurs (*error free service model*) qui définit le flux autorisé à transmettre si le canal est bon,
- d'un modèle d'avance et de retard (*lead and lag model*) qui définit les différentes catégories de flux,
- d'un modèle de compensation qui définit le flux de remplacement si le canal du premier flux est mauvais,
- d'un modèle d'écoute et de prédiction du canal.

a) Modèle du canal de transmission

Le rapport entre le niveau du signal utile et le bruit caractérise l'état du canal radio. Il s'agit du SNR (*Signal To Noise Ratio*). Si on considère un taux d'erreurs bit fixe, on peut avoir un débit utile plus ou moins important selon le SNR. La plage de variation du SNR est généralement divisée en intervalles consécutifs séparés par des seuils (cf. tableau 4.5). Dans chaque intervalle, un MCS (caractérisé par une modulation et un taux de codage) peut être utilisé. Lorsque $\gamma \in [\gamma_{m-1,m}, \gamma_{m,m+1}]$ alors le mode m est choisi. On considère généralement $m \in \{0 \dots M\}$ avec $M = 5$. Pour $m = 0$, il est impossible de transmettre et pour $m = M$, le débit utile maximal est atteint. Certains considèrent un canal tel que $M = 1$ c'est à dire qu'il y a un seul mode qui permet de transmettre. Le canal est dit ON-OFF car soit $m = 0$ soit $m = 1$.

b) Ordonnancement sur modèle de canal ON-OFF

Dans le modèle ON-OFF : soit le canal est bon et on peut transmettre avec un débit unique soit il est mauvais et toute transmission est perdue.

Parmi ces algorithmes, on a retenu le WFS (*Wireless Fair Service*, [Lu&_98]) qui découple les garanties en délai de celles en débit. Pour cela, le modèle sans erreurs est une sorte de WFQ dans lequel on introduit un poids de délai dans l'étiquette de fin de

service. Le WFS est détaillé dans la section 2.1 car il est à la base de notre contribution. Le WFS garantit des bornes de délai, des débits minimaux, une équité à court et long terme et une dégradation gracieuse de service.

D'autres algorithmes utilisent un modèle ON-OFF. On distingue le CSDPS (*Channel State Dependent Packet Scheduling*, [Bhagwat&_96]), le CIF-Q (*Channel Independent Packet Fair Queuing*, [Eugene&_98]), l'IWFQ (*Idealized Wireless Fair Queuing*, [Lu&_99]), l'ORCA_SRT (*Optimal Radio Channel Allocation Single Rate Transmission*, [Issaiyakul_99]). De brèves descriptions sont données dans l'annexe 6. B.

L'hypothèse d'un canal ON-OFF est une sévère limitation. Cette hypothèse est pessimiste et empêche d'exploiter pleinement les possibilités du canal de transmission. Les algorithmes présentés ci-après prennent en compte l'aspect multidébit du canal de transmission.

c) Ordonnancement sur canal multidébit

Grâce aux différents MCS, on peut transmettre sur un canal radio variable et assurer un taux d'erreurs fixe. La possibilité, pour chaque utilisateur, d'utiliser différents modes de transmission a d'importantes conséquences sur les débits utiles et remet en cause la notion d'équité. Malgré une équité temporelle entre flux, l'usage de MCS différents, peut aboutir à d'importantes disparités en terme de débit. Les algorithmes prenant en compte l'aspect multidébit sont : l'ORCA-MRT (*Optimal Radio Channel Allocation Multi Rate Transmission*, [Issaiyakul_99]), le CS-WFQ (*Channel State Independent Wireless Fair Queuing*, [Lin&_00]), le CAFQ (*Channel Adaptive Fair Queuing*, [Wang&_04]), le MR-FQ (*MultiRate Fair Queueing*, [Wang&_05]), le MWFS (*Multirate Wireless Fair Scheduling*, [Yuan&_05]) et l'OWFQ (*Opportunistic Weighted Fair Queuing*, [Khawam_06]).

Tous ces algorithmes sont décrits dans l'annexe 6.C, on peut noter ici que le point commun de ces algorithmes est d'insérer l'état du canal dans la décision d'ordonnancement. Voyons quelques exemples.

Dans le CS-WFQ, le poids r'_i utilisé dans l'étiquette de fin de service (équation (6.4)) diffère du poids de débit initial r_i . Ce poids r'_i évolue en fonction de l'état du canal :

$$r'_i = \frac{r_i}{\max(W_i^{th}, W_i(t))} \quad (6.6) \text{ où } W_i(t) = W_i(0, t) \text{ correspond au débit du flux } i. \text{ Si un}$$

flux transmet peu par rapport à son poids de débit, c'est qu'il est victime d'un mauvais canal. Dans ce cas, il va obtenir momentanément un poids r'_i supérieur à r_i et va recevoir une sorte de compensation immédiate ; W_i^{th} est le seuil en dessous duquel un tel flux n'est plus favorisé. Le seuil W_i^{th} évite qu'un flux qui présente un mauvais canal trop longtemps ne préempte toutes les opportunités de transmission.

Dans le CAFQ, le flux choisi possède le paramètre v_i^N le plus faible :

$$v_i^N = v_i^N + \frac{L}{r_i f(m)^\eta} \quad (6.7)$$

La fonction $f(m)$ augmente avec le mode de transmission m (il évolue de 0 à 1 quand le mode augmente de 0 à M). Le paramètre η est fixe et positif. Dans l'équation (6.7), plus

le canal est mauvais, plus v_i^N est grand et moins le flux a de chances d'être sélectionné au prochain tour.

Dans le MWFS, l'étiquette de début de service est donnée par l'équation (6.3). Les paquets sont transmis par étiquette de début de service croissante. L'étiquette de fin de service devient :

$$F(p_i^k) = S(p_i^k) + \frac{L(p_i^k)}{r_i c_{m_i(t)}} \quad (6.8) \text{ où } m_i(t) \text{ est le débit du mode de transmission du flux } i$$

au temps t . Comme dans l'équation (6.7), l'étiquette est d'autant plus faible que le mode est élevé ; cela favorise les flux qui ont un bon canal.

Enfin l'OWFQ définit une étiquette de fin de service similaire à l'équation (6.8) mais transmet les paquets par étiquette de fin de service croissante.

1.3. État de l'art : ordonnancement en OFDMA

En OFDMA, on distingue deux stratégies d'ordonnancement : un ordonnancement par sous canal et un ordonnancement global.

Dans la première stratégie, les auteurs proposent une règle de priorité qui est appliquée indépendamment sur chaque sous canal. L'utilisateur maximisant la métrique est sélectionné et peut transmettre sur le sous canal. A notre connaissance, les algorithmes de ce type traitent principalement le trafic non temps réel : [KiKim_05], [Wengerter&_05]. Les règles de décisions proposées (cf. annexe 6. D) sont inspirées de celle de l'ordonnancement PF (*Proportional Fair*, [Jalali&_00]).

Nous nous intéressons à la seconde stratégie consistant en un ordonnancement global. Nous avons constaté que les algorithmes qui supportent à la fois un trafic temps réel et non temps réel, utilisent cette stratégie. Dans une première étape, un ordonnancement global sélectionne les utilisateurs qui vont transmettre. On retrouve dans cette étape certains algorithmes présentés en 1.2.1 et 1.2.2. Puis, une seconde étape décide quelles ressources allouer aux utilisateur.

Dans certaines références ([ShinRyu_04], [Shin&_05]), la seconde étape est imbriquée dans la première. On dira alors qu'il s'agit d'un *ordonnancement avec allocation de ressources « simplifiée »*. Dans d'autres références ([ZhangLetaief_04], [Diao&_04] et [Jeong&_06]), la seconde partie est clairement séparée et on y retrouve des algorithmes d'optimisation présentés dans la première partie de la thèse (chapitre 2). On dira dans ce cas qu'il s'agit d'un *ordonnancement avec allocation de ressources « optimisée »*. Cette stratégie peut être plus coûteuse en temps que la première ; la contrepartie serait alors l'obtention d'une meilleure utilisation des ressources grâce à la phase d'optimisation.

1.3.1. Ordonnancement avec allocation de ressources simplifiée

Les auteurs de [ShinRyu_04] proposent le PLFS (*Packet Loss Fair Scheduling*) pour un trafic temps réel. Dans [Shin&_05], le PLFS est élargi pour traiter deux types de trafic temps réel (ou RT pour *real time*) et non temps réel (ou NRT pour *non real time*). Le

PLFS est l'une des références choisies pour évaluer les performances de notre proposition en OFDMA (cf. section 3.3) ; C'est pourquoi nous le détaillons ci-après.

PLFS (Packet Loss Fair Scheduling, [Shin&_05])

La trame est divisée en trois parties (cf. Fig. 6.1) : la première est réservée au trafic RT et la dernière au trafic NRT ; l'affectation de la partie du milieu est variable : elle est en priorité affectée au trafic RT pour respecter les échéances temporelles. Dans chaque région de la trame, les S utilisateurs de plus haute priorité sont sélectionnés pour la transmission et se voient affecter un sous canal en fonction de leur SNR¹. Dans le PLFS, deux règles distinctes sont appliquées pour chaque type de trafic :

$$\mu_u(t) = \frac{R_{u,n(u)}(t)}{\bar{R}_u(t)} \frac{PLR_u(t)}{PLR_{req,u} D_{max,u}} \quad (6.9)$$

$$\mu_u(t) = QL_u(t) D(p_u^{HOL}) \frac{R_{u,n(u)}(t)}{\bar{R}_u(t)} \quad (6.10)$$

La région RT est affectée en utilisant la métrique (6.9) où $PLR_u(t)$ est le taux de paquets perdus par l'utilisateur u à l'instant t , $PLR_{req,u}$ est le taux de pertes maximum autorisé et $D_{max,u}$ est le délai maximum d'un paquet. Pour la région NRT, on utilise la métrique (6.10) où $QL_u(t)$ est la longueur de la queue de l'utilisateur u tandis que $D(p_u^{HOL})$ est le délai du paquet de tête de la queue de l'utilisateur u . Une implémentation possible pour ces deux équations est la suivante : $\bar{R}_u(t)$ est un débit moyenné sur une fenêtre glissante et sur tous les sous canaux ; $R_{u,n(u)}(t)$ est le débit sur le meilleur sous canal de u noté $n(u)$.

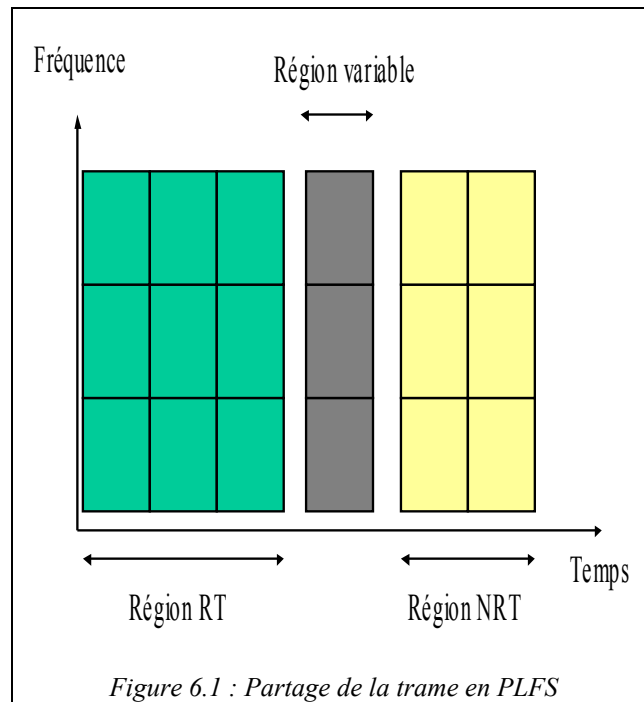


Figure 6.1 : Partage de la trame en PLFS

¹ Le mécanisme précis n'est pas explicité dans [Shin&_05], mais dans [ShinRyu_04] il mentionné que les utilisateurs sélectionnent leur meilleur sous canal tour à tour et ce en fonction de la priorité obtenue par la métrique.

1.3.2. Ordonnancement avec allocation de ressources optimisée

Dans cette catégorie, [ZhangLetaief_04] traitent les flux NRT : les auteurs appliquent le *Virtual Clock* suivi d'une optimisation MA avec l'algorithme de [WongTsui&_99] (cf. annexe 6. D).

Les contributions [Diao&_04] et [Jeong&_06] traitent les deux types de flux. Ces travaux nous servent de références pour évaluer les performances de notre proposition en OFDMA (cf. section 3.3). C'est pourquoi nous les détaillons ci-après.

CPLD-PGPS ([Diao&_04])

Les auteurs de [Diao&_04] proposent le CPLD-PGPS (*Channel-Condition Packet Length Dependent Packet Generalized Processor Sharing*). Les hypothèses sont les suivantes : modèle d'erreur canal ON-OFF, modulation fixe et puissance également répartie sur les sous porteuses. Les auteurs généralisent le calcul du temps virtuel¹ au contexte de l'OFDM. Les paquets sont marqués avec une étiquette virtuelle de fin de service comme en WFQ. Au moment de l'ordonnancement, les paquets de tête de queue sont parcourus par étiquette de fin de service croissante. Soit H l'ensemble des sous porteuses disponibles, cet ensemble décroît à chaque fois qu'un paquet est sélectionné. Pour chaque paquet, soit :

- $N_u(p_u^{HOL})$ le nombre de sous porteuses nécessaires pour traiter le paquet (ce nombre est calculé grâce à la taille du paquet et la connaissance d'un MCS fixe),
- $N_{good,u}$ le nombre de bonnes sous porteuses (une sous porteuse est soit bonne soit mauvaise à cause du modèle d'erreur ON-OFF).

Lors du parcours des paquets, on incrémente le facteur d'équité² $G(i)$ d'un flux i dont le paquet p n'est pas choisi, s'il existe des paquets d'autres flux qui sont choisis avec des étiquettes de fin de service supérieures à celle de p . La règle de sélection des paquets est la suivante :

- soit $N_u(p_u^{HOL}) < |H|$ (il y a assez de sous porteuses disponibles pour traiter le paquet) et $G(u) > G_{threshold}$ (le flux a été différé plusieurs fois),
- soit $N_u(p_u^{HOL}) < |H|$ et $N_{good,u} / N_u(p_u^{HOL}) > seuil$ (un flux a beaucoup de bonnes sous porteuses, cela facilite la résolution des conflits pendant l'étape d'attribution des sous porteuses).

Une fois les paquets choisis, les auteurs appliquent l'algorithme de [WongTsui&_99] (cf. chap. 2).

Practical Cross Layer Resource Management ([Jeong&_06])

Dans [Jeong&_06], la phase d'ordonnancement détermine les paquets qui seront émis dans la prochaine trame. Pour la transmission des paquets, une phase d'optimisation affecte des unités d'allocation (UA) ; ces dernières sont composées d'un sous canal (sur l'axe fréquentiel) et d'un *slot* (sur l'axe temporel). L'ordonnancement (décrit ci-après) sélectionne les flux et calcule les débits minimaux à respecter pendant l'optimisation.

1 On reviendra sur ce calcul en 2.2

2 Ce facteur ressemble fort à un retard.

Pour profiter de la diversité multiutilisateur, les auteurs imposent un nombre minimal K d'utilisateurs à choisir pour chaque trame ; S/K est la limite initiale¹ sur le nombre d'UA qu'un utilisateur peut recevoir. Dans la première phase, la priorité est donnée aux paquets RT urgents définis par $\frac{D(p_u^k)}{D_{max,u}} \geq \delta$ (6.11) ; $D_{max,u}$ est la contrainte de délai de l'utilisateur u , $D(p_u^k)$ est le délai du $k^{ième}$ paquet de u et δ est un seuil fixe. Les utilisateurs qui possèdent des paquets urgents sont considérés par ordre décroissant du rapport $\frac{\max_k(D(p_u^k))}{D_{max,u}}$. Pour chacun, on note $\sum_p L_{u,p}$ le volume de données urgentes.

Soit N_u^1 , le nombre d'UA de l'utilisateur u à la première phase :

$$N_u^1 = \min\left(E\left(\frac{\sum_p L_{u,p}}{R_{u,max}}\right) + 1, \frac{S}{K}, |H|\right) \quad (6.12) \text{ où } E(x) \text{ est la partie entière de } x, R_{u,max} \text{ est}$$

le nombre maximum de bits que l'utilisateur u peut écouler sur une UA et H est l'ensemble des UA disponibles. S'il reste des UA ($|H| > 0$) après le traitement des données urgentes, une deuxième phase considère les utilisateurs RT et NRT simultanément. Ils sont traités par CgNR moyen décroissant, le CgNR moyen étant calculé sur l'ensemble des UA. Un utilisateur reçoit N_u^2 UA supplémentaires². Enfin, s'il reste encore des UA, les utilisateurs qui ont été limités à la quantité S/K peuvent recevoir N_u^3 UA supplémentaires selon leurs besoins. Le nombre final d'UA de l'utilisateur u est $N_u = N_u^1 + N_u^2 + N_u^3$. Le débit minimal de u pendant la phase d'optimisation vaut $R_u^o = N_u R_{u,max}$. Pendant cette phase, les UA sont affectées à leur meilleur utilisateur puis une réaffectation a lieu pour respecter les débits R_u^o de chacun.

c) Motivations pour une contribution en OFDMA

Dans le contexte de l'ordonnancement PF (trafic NRT), les auteurs de [Anchun&_03] montrent qu'un ordonnancement global vaut mieux que plusieurs indépendants (i.e. un ordonnancement par sous canal) au regard du débit global (*throughput*). De plus, l'ordonnancement global permet de supporter deux types de trafic. Partant de ces constats, nous proposons une contribution avec un ordonnancement global.

A notre connaissance, il n'y a pas d'algorithmes approximant le modèle GPS qui ait été proposé dans le contexte de l'OFDMA avec modulation adaptative. Il y a certes la contribution de [Diao&_04] mais la modulation est fixe et il n'y a pas de découplage entre débit et délai.

Nous modifions le WFS pour tenir compte d'un grand nombre de canaux multidébits. Nous choisissons le WFS car il peut traiter des flux de QoS différentes grâce à l'existence d'un poids de débit et d'un poids de délai et à la différenciation entre queue de *slots logiques* et queue de paquets (cf. section 2.1.2). Deux algorithmes sont proposés l'EWFS (*Enhanced Wireless Fair Service*) et l'OWFS (*Opportunist Wireless Fair Service*). Notre contribution entre dans la catégorie *ordonnancement avec allocation de ressources simplifiée*.

¹ L'algorithme comporte trois phases, la limite sur le nombre d'UA est utilisée sur les deux premières phases.

² N_u^2 est calculé grâce à (23) sachant que dans l'étape 2, $\sum L_u$ est le volume total de données à transmettre

2. Le WFS et ses évolutions

Nous étudions en 2.1, l'algorithme initial proposé dans [Lu&_98]. L'algorithme WFS considère un seul canal de transmission qui peut être soit bon soit mauvais. Comme on l'a vu, le modèle d'erreurs du canal peut être affiné : selon l'état du canal, on peut utiliser différentes modulations et différents taux de codage. En 2.2, nous présentons deux algorithmes, extensions du WFS, qui considèrent plusieurs canaux multidébits. En 2.3, les théorèmes et propriétés de ces algorithmes sont énoncés. Les principales propriétés de ces algorithmes sont illustrées en 2.4.

2.1. Algorithme initial : le *Wireless Fair Service*

2.1.1. Le temps virtuel

Soient $\{r_i\}_i$, les poids de débit normalisés des utilisateurs. En GPS (*General Processor Sharing*), on définit un tour d'ordonnancement (*scheduling*), comme le temps passé à servir tous les flux actifs. Pendant un tour, un flux actif i reçoit r_i bits. Plus il y a de flux actifs, plus le tour d'ordonnancement est lent. La durée de transmission d'un paquet dépend du nombre de flux actifs. Celle du paquet le plus ancien varie au gré du nombre de flux actifs.

Par contre, le nombre de tours (L/r_i) nécessaires à la transmission d'un paquet (du flux i et de taille L) est fixe et indépendant du nombre de flux actifs. Considérons une horloge $V(t)$, on définit une étiquette de fin de service qui s'écrit $F(p_i^k) = V(t) + L(p_i^k)/r_i$, alors $F(p_i^k)$ est indépendant du nombre de flux actifs.

Ainsi, le temps virtuel $V(t)$ est introduit pour disposer d'étiquettes de début et de fin de service invariantes. Pour cela, le temps virtuel compte le nombre de tours d'ordonnancement effectués depuis le début d'une période active. Une période active est une période où il y a toujours au moins un flux actif. Lorsqu'aucun flux n'est actif, le temps virtuel est remis à zéro.

Le temps virtuel est mis à jour à chaque événement : une arrivée ou un départ de paquet. Un tel événement peut en effet changer le nombre de flux actifs et affecter la vitesse d'exécution d'un tour d'ordonnancement. On considère les notations suivantes :

- t_j l'heure (réelle) du $j^{\text{ème}}$ événement en secondes,
- A_j est l'ensemble des flux actifs entre $[t_{j-1}, t_j[$,
- $C(t_{j-1})$ est la capacité du serveur constante sur l'intervalle $[t_{j-1}, t_j[$ en bits/s.

Alors, le temps virtuel est mis à jour comme suit ([TsengChang_98]):

$$V(t_j) = V(t_{j-1}) + \frac{(t_j - t_{j-1}) C(t_{j-1})}{\sum_{i \in A_j} r_i} \quad (6.13)$$

Le nombre de bits reçus par le flux i pendant l'intervalle $[t_{j-1}, t_j[$ est :

$$\Delta(i, t_j) = r_i \frac{(t_j - t_{j-1}) C(t_{j-1})}{\sum_{i \in A_j} r_i} \quad (6.14)$$

$$\text{ou } \Delta(i, t_j) = r_i(V(t_j) - V(t_{j-1})) \quad (6.15)$$

En supposant qu'il n'y ait pas d'arrivée après le temps t_j , la prochaine heure (réelle) de fin de service d'un paquet est donnée par :

$$t = t_j + \frac{(F_{\min} - V(t_j)) \sum_{i \in A_j} r_i}{C(t_j)} \quad (6.16)$$

En effet :

- $F_{\min}$ est la plus petite étiquette de fin de service, $(F_{\min} - V(t_j))$ représente le nombre de tours qui s'écoulent entre t_j et l'instant de fin de service.
- $(F_{\min} - V(t_j)) \sum_{i \in A_j} r_i$ représente le nombre de bits traités par le serveur et étant donné sa capacité $C(t_j)$ en bits par secondes, on peut obtenir la durée écoulée entre t_j et l'instant de fin de service

2.1.2. Le modèle sans erreurs

Le modèle sans erreurs est inspiré de celui du WFQ. Il diffère par l'existence d'un deuxième poids Φ_i appelé poids de délai. Il permet de découpler les garanties en bande des garanties en délai. Ces deux notions sont en effet intimement liées en WFQ. Un paramètre appelé *lookahead* et noté ρ spécifie «une région de sélection» en définissant la condition sur les étiquettes de début de service :

$$S(p_i^k) < V(t) + \rho \quad (6.17)$$

$$S(p_i^k) = \max(V(t_i^k), S(p_i^{k-1}) + \frac{L(p_i^{k-1})}{r_i}) \quad (6.18)$$

$$F(p_i^k) = S(p_i^k) + \frac{L(p_i^k)}{\Phi_i} \quad (6.19)$$

Les paquets sont sélectionnés par étiquette de fin service croissante parmi les paquets vérifiant l'équation (6.17).

2.1.3. Queue de «jetons» et queue de paquets

Le WFS propose de gérer une queue de paquets et une queue de «jetons». En anglais le terme *slot queue* est utilisé. Pour les auteurs les *slots* représentent plutôt des *slots logiques* et sont des opportunités de transmission ; nous utiliserons le terme «jeton» dans le reste du document. Nous réservons le terme *slot* au *slots* physiques, intervalles temporels destinés à la transmission.

Les queues de jetons et de paquets sont d'égale longueur. Lorsqu'un paquet arrive, il est ajouté à la fin de la queue de paquets et un jeton est ajouté à la fin de la queue de jetons. Les étiquettes de début et de fin de service sont associées au jeton et non au paquet. Lorsqu'un flux est désigné pour transmettre, il transmet le paquet de tête de la queue de paquets et détruit le jeton de tête de la queue de jetons. Lorsqu'un paquet est supprimé pour expiration d'échéance (*timer*) ou excès de retransmissions, le jeton de tête de la queue de jetons est inchangé¹. Ainsi, la politique d'un flux concernant la gestion des

¹ Pour équilibrer la taille des deux queues, on peut supprimer le dernier jeton

paquets n'a pas d'impact sur ses opportunités de transmission. On peut alors gérer simultanément des flux sensibles au délai et des flux sensibles aux pertes.

2.1.4. *Modèle d'avance et de retard*

Ce paragraphe décrit le modèle d'avance et de retard et le modèle de compensation. Il existe trois catégories de flux, en avance, en retard et en équilibre. Des compteurs sont maintenus au niveau de chaque flux. La catégorie d'un flux dépend de la valeur des compteurs. Un flux i en avance a un compteur d'avance positif $E(i)$, un flux i en retard a un compteur de retard $G(i)$ positif et un flux en équilibre a les deux compteurs précédents nuls¹. Le compteur d'avance $E(i)$ est incrémenté lorsque $E(i)$ est inférieur à la limite $E_{max}(i)$ et que le flux i ayant un bon canal reçoit un jeton venant d'un autre flux. Un flux cède son service pour deux raisons : soit il ne peut transmettre à cause d'un mauvais canal, soit il s'agit d'un flux en avance. Les flux j en avance ont en effet le devoir de relâcher une proportion de service variable: $E(j)/E_{max}(j)$. Grâce à cette proportion variable, la dégradation de service est exponentielle.

Les compteurs du flux initial ne sont modifiés que si un autre flux a pu transmettre à sa place. Dans ce cas, le flux initial devrait laisser sa queue de paquet inchangée et supprimer son jeton de tête². S'il ne supprimait pas son jeton de tête, il aurait, comme en IWFQ, priorité sur les autres flux dès qu'il retrouve un bon canal et jusqu'à ce qu'il rattrape son retard.

2.1.5. *Modèle de compensation*

Un ordre de priorité est fixé afin de déterminer quel flux peut transmettre à la place d'un flux dont le canal est mauvais. On cherche un flux qui a un bon canal parmi : 1) les flux en retard, puis 2) parmi les flux en avance et qui vérifient $E(i) < E_{max}(i)$, puis 3) parmi les flux synchrones et enfin 4) parmi les flux en avance et qui vérifient $E(i) = E_{max}(i)$. Si un flux en avance relâche une opportunité de transmission et qu'aucun flux en retard ne peut transmettre, alors on redonne la main au flux initial. S'il a un bon canal, il transmet sinon la recherche d'un flux reprend à partir de l'étape 2 (les flux en retard constituent la première étape et ont déjà été parcourus).

2.1.6. *Écoute et prédiction du canal*

Dans la définition initiale du WFS, le modèle d'écoute du canal est le CSMA (*Carrier Sense Multiple Access*). Si un acquittement n'est pas reçu, le canal est considéré comme mauvais. La norme prévue pour WFS est en effet l'IEEE 802.11. Le mécanisme de prédiction est le suivant : l'état du canal durant un *slot* (physique) est supposé identique à celui du *slot* précédent. Dans la suite nous adapterons le modèle d'écoute et de prédiction du canal à la norme considérée.

1 On peut sans pertes d'information avoir un unique compteur qui serait positif, négatif ou nul selon la catégorie avance retard ou équilibre

2 Pour conserver la taille de la queue de jetons, il peut recréer un jeton à la fin de la queue de jetons

2.2. Notre contribution : évolutions du WFS pour l'OFDMA

2.2.1. Modèle de service sans erreurs

Enhanced Wireless Fair Service (EWFS)

L'expression de l'étiquette de début de service est inchangée par rapport au WFS. Cependant, les étiquettes de fin de service ne sont désormais calculées que pour les paquets de tête de queue. Cette nouvelle définition prend en compte l'état des sous canaux du flux à l'instant de la sélection des paquets.

$$F(p_i^k) = S(p_i^k) + \frac{L(p_i^k)}{m_{i,best}(t)\Phi_i} \quad (6.20)$$

Le paramètre introduit, $m_{i,best}(t)$, représente le meilleur mode (en modulation adaptative) sur tous les sous canaux encore disponibles à l'instant t . Dans le contexte OFDMA, cela permet au moment de la transmission de favoriser les flux qui ont de bonnes conditions radio. Un flux sélectionné reçoit son meilleur sous canal parmi ceux qui sont encore disponibles. Le poids Φ_i représente le poids de délai.

Opportunist Wireless Fair Service (OWFS)

L'étiquette de début de service diffère du WFS :

$$S(p_i^k) = \max(V(t_i^k), S(p_i^{k-1}) + \frac{L(p_i^{k-1})}{m_{i,max}(t_i^k)r_i}) \quad (6.21)$$

Le paramètre introduit, $m_{i,max}(t_i^k)$, représente le meilleur mode (en modulation adaptative) sur tous les sous canaux à l'instant t_i^k d'arrivée du paquet. Lorsqu'un paquet arrive, les sous canaux sont tous disponibles pour les prochains instants de «sélection».

La définition des étiquettes de fin de service est similaire à celle de l'EWFS. L'OWFS est un algorithme plus opportuniste que l'EWFS. Dès l'arrivée d'un paquet, un bon canal diminue l'étiquette de début de service. Si le canal est toujours bon à l'instant de sélection l'étiquette de fin de service est alors plus faible que celle qui aurait été obtenue en EWFS. Si un flux a un bon canal, le délai de ses paquets est ainsi diminué. La prise en compte du canal dans l'étiquette de début de service permet de transmettre le plus tôt possible. Cela augmente les chances de transmettre tant que le canal est bon.

2.2.2. Modèle de compensation

Le modèle de compensation de l'EWFS est identique à celui de l'OWFS. Comparé au WFS, l'ordre de priorité entre les différentes catégories de flux est scrupuleusement le même (on cherche parmi les flux en retard, puis en avance, et parmi les flux synchrones). Cependant, une différence apparaît entre les flux d'une même catégorie. Lorsqu'on cherche un flux synchrone (respectivement en avance) on prendra celui qui a le meilleur sous canal disponible. Concernant les flux en retard, le WFS considère un WRR (*weighted round robin*) où les poids sont égaux aux compteurs de retard. Pour l'EWFS et OWFS, nous supprimons le WRR et nous considérons simplement le flux en retard qui a le meilleur sous canal disponible. Nous pensons que la meilleure façon de réduire le retard d'un flux est la transmission dans des conditions radio privilégiées.

Choix parmi des flux en retard dont le mode de transmission est identique

En cas d'égalité, c'est à dire que le meilleur mode est identique pour plusieurs flux en retard, nous distinguons trois cas. Premièrement, il n'y a que des flux non temps réel ; le flux avec le plus grand retard transmet. Deuxième cas, il n'y a que des flux temps réel ; le flux avec le retard le plus faible transmet. Dernier cas : des flux de type différents sont à égalité ; si le flux de retard maximum est non temps réel, il est choisi sinon le flux temps réel de retard minimum est choisi.

Nos motivations sont les suivantes. Concernant les flux non temps réel, nous essayons d'éviter les interactions malencontreuses entre TCP et le canal radio ; on veut donc éviter que les flux non temps réel aient des retards qui déclenchent les retransmissions TCP. Concernant les flux temps réel, étant donné que les connexions qui ont un accumulé un retard important sont susceptibles d'être supprimées, on préfère privilégier les connexions qui ont un faible retard pour leur éviter une dégradation de QoS. Choisir les flux temps réel qui ont le plus grand retard, c'est utiliser des ressources pour des connexions qui sont peut être condamnées et laisser pendant ce temps se dégrader des connexions qui auraient pu être sauvées.

Pour montrer l'intérêt de ses règles dans le modèle de compensation, il faudrait les comparer aux règles du WFS dans une simulation où la couche TCP est modélisée. Faute de temps, cette simulation n'a pas été réalisée au cours de cette thèse. C'est un travail très intéressant pour la suite de la caractérisation de l'EWFS et de l'OWFS.

Valeur des compteurs

On considère $E_{min,i}$ la limite inférieure (en valeur absolue) du retard d'un flux et $E_{max,i}$ la limite supérieure de l'avance du flux i . La modification des compteurs lors d'un échange de jetons entre i et j n'a lieu que si $|E(i)| < |E_{min,i}|$ et $E(j) < E_{max,j}$. Ainsi la somme des compteurs sera toujours nulle. Lorsqu'un échange ne donne pas lieu à une modification de compteurs, cela signifie que le flux qui cède son jeton ne recevra pas de compensation pour ce jeton. Si un flux i a de mauvaises conditions radio pendant une longue période, il ne peut espérer recevoir plus de jetons que $E_{min,i}$ en guise de compensation.

2.2.3. *Écoute et prédiction du canal*

Nous considérons en simulation une connaissance parfaite de l'état de chaque sous canaux et ce pour tous les utilisateurs. Concernant les valeurs système, nous nous basons sur la norme IEEE 802.16. En pratique, sur réception du signal de la BS, les utilisateurs envoient leurs cinq meilleurs modes à la BS sur un canal prévu à cet effet (CQICH, *Channel Quality Indicator Channel*). En mode TDD, on peut supposer que le canal est symétrique.

2.3. Propriétés

Toutes les propriétés énoncées dans cette section sont démontrées dans les annexes 6.E- 6.H. Les démonstrations se basent sur des arguments similaires à ceux utilisés dans [Lu&_00] pour caractériser le WFS. On notera U le nombre total de flux.

2.3.1. Région schedulable en temps virtuel

Définition 6.1 : Un vecteur de délai $(D_1, D_2, \dots, D_U)$ est dit *schedulable* en temps virtuel si tout paquet p_i^k du flux i est transmis avant le temps virtuel $S(p_i^k) + D_i$. La *région schedulable* est l'ensemble des vecteurs *schedulables*.

Lemme 6.1

Considérons $\rho \in [0, \infty[$. Soit $D_1 \leq D_2 \leq \dots \leq D_U$ où $D_i = L_{max} / \Phi_i$. Les poids de débit et de délai (respectivement r_i et Φ_i) sont normalisés. En EWFS et OWFS, si le vecteur $(D_1, \dots, D_U)$ est *schedulable* en temps virtuel, alors pour tout intervalle en temps virtuel $[v_o, v(t)]$,

$$L_{max} \leq D_1 \quad (6.22)$$

$$\begin{aligned} &\text{Pour } L_{max} \leq v(t) - v_o < D_U \\ &\sum_{i=1}^U [L_{max} + r_i \min(A, B)] * I(\min(A, B)) + L_{max} \leq v(t) - v_o \\ &\text{avec } A = v(t) - v_o - L_{max} + \rho ; B = v(t) - v_o - D_i \\ &\text{et } I(x) = 1 \text{ si } x \geq 0 \text{ sinon } I(x) = 0 \end{aligned} \quad (6.23)$$

$$\begin{aligned} &\text{Pour } v(t) - v_o \geq D_U, \sum_{i=1}^U [L_{max} + r_i \min(A, B)] \leq v(t) - v_o \\ &\text{avec } A = v(t) - v_o - L_{max} + \rho ; B = v(t) - v_o - D_i \end{aligned} \quad (6.24)$$

Preuve: cf. annexe 6.E

Ce lemme est identique à celui du WFS car pour l'EWFS et l'OWFS on considère un vecteur de délai similaire à celui du WFS. En effet, la démonstration se base sur l'attente d'un paquet : $L_{max} / m_{i,best} \Phi_i$, c'est à dire la différence entre l'étiquette de fin de service (*finish tag*) et celle de début de service (*start tag*). Pour que le vecteur de délai soit *schedulable* quelque soit l'évolution des conditions radio des sous canaux, on tient compte du cas $m_{i,best} = 1$; $D_i = L_{max} / \Phi_i$ est le nombre de tours qu'un paquet doit attendre au maximum avant d'être servi. On retrouve donc un vecteur de délai identique au WFS.

Lemme 6.2 : Tout vecteur $(D_1, D_2, \dots, D_U)$ avec $D_1 \leq D_2 \leq \dots \leq D_U$ et $D_i = L_{max} / \Phi_i$ qui vérifie les contraintes du lemme est *schedulable* en temps virtuel en EWFS et en OWFS.

Preuve: cf. annexe 6.E

Grâce aux lemmes 6.1 et 6.2, la région *schedulable* en temps virtuel est définie. Le théorème 6.1, caractérise cette région *schedulable*.

Théorème 6.1 : La région *schedulable* des algorithmes EWFS et OWFS est, en temps virtuel, l'ensemble des vecteurs vérifiant $L_{max} \leq D_1 \leq D_2 \leq \dots \leq D_U$ de la forme $D_i = L_{max} / \Phi_i$ et qui satisfont les contraintes :

$$1 \leq j \leq U-1, \quad (j+1)L_{max} \leq D_j - \sum_{i=1}^{j-1} r_i \min(D_j - L_{max} + \rho, D_j - D_i) \quad (6.25)$$

$$UL_{max} \leq D_U - \sum_{i=1}^{U-1} r_i \min(D_U - L_{max} + \rho, D_U - D_i) \quad (6.26)$$

Preuve: cf. annexe 6.E

Corollaire 6.1 : Comme pour tout i $L_{max} \leq D_i$ et $\rho \geq 0$, la condition de *schedulabilité* de l'EWFS et de l'OWFS se simplifie et devient :

$$1 \leq j \leq U-1, \quad (j+1)L_{max} \leq D_j \left(1 - \sum_{i=1}^{j-1} r_i\right) + \sum_{i=1}^{j-1} r_i D_i \quad (6.27)$$

$$UL_{max} \leq D_U \left(1 - \sum_{i=1}^{U-1} r_i\right) + \sum_{i=1}^{U-1} r_i D_i \quad (6.28)$$

Preuve: cf. annexe 6. E

La région *schedulable* en temps virtuel de l'EWFS et de l'OWFS est la même que celle du WFS. Elle ne dépend pas du paramètre *lookahead*. Elle est identique à la région *schedulable* en temps réel de l'EDF.

2.3.2. Garanties de service

Lemme 6.3 : En EWFS et en OWFS, un flux actif dont la queue est non vide durant l'intervalle (en temps virtuel) $[v_1, v_2]$ reçoit une quantité de service $W_i(v_1, v_2)$ qui vérifie :

- si le premier paquet de la queue du flux i arrive à v_1 , alors

$$W_i(v_1, v_2) \geq r_i(v_2 - v_1) - L_{max} - 2L_{max} \frac{r_i}{\phi_i} \quad (6.29)$$

- si un paquet du flux i est servi à $v_1 - \varepsilon$ (avec ε arbitrairement petit), alors

$$W_i(v_1, v_2) \geq r_i(v_2 - v_1) - L_{max} - L_{max} \frac{r_i}{\phi_i} \quad (6.30)$$

Preuve: cf. annexe 6. F

On souhaite à partir du lemme 6.3 déterminer le service minimum reçu en temps réel. Pour cela, étant donné que la capacité du canal est variable, on travaille avec la capacité vue par les utilisateurs avec une probabilité β donnée. On considère donc la capacité $C_{gar,\beta}$ telle que $P(C \geq C_{gar,\beta}) \geq \beta$.

Théorème 6.2 : En EWFS et en OWFS, si un flux i est continuellement actif sur un intervalle réel $[t_1, t_2]$, alors le service reçu par le flux i vérifie :

$$P(W_i(t_1, t_2) \geq r_i C_{gar,\beta}(t_2 - t_1) - L_{max} - \alpha L_{max} \frac{r_i}{\phi_i}) \geq \beta \quad (6.31)$$

avec β et $C_{gar,\beta}$ tels que $P(C \geq C_{gar,\beta}) \geq \beta$

$\alpha = 2$, si le flux i reçoit son service à t_1

$\alpha = 1$, si le flux i reçoit son service à $t_1 - \varepsilon$

Preuve : cf. annexe 6. F.

Dans le lemme suivant, on revient dans le temps virtuel pour majorer la quantité de service reçue sur un intervalle où le flux est continuellement actif.

Lemme 6.4 : En EWFS, si tous les flux sont continuellement actifs sur l'intervalle $[v_1, v_2]$ en temps virtuel alors le service qu'un flux i peut recevoir est majoré comme suit :

$$W_i(v_1, v_2) \leq r_i(v_2 - v_1) + L_{max} \quad (6.32)$$

En OWFS, si tous les flux sont continuellement actifs sur l'intervalle $[v_1, v_2]$ alors le service qu'un flux i peut recevoir est majoré comme suit :

$$W_i(v_1, v_2) \leq r_i m_{max}(v_2 - v_1) + L_{max} \quad (6.33)$$

Preuve : cf. annexe 6.F.

2.3.3. Garanties d'équité

Théorème 6.3 : En EWFS, pour un flux continuellement actif, l'équité réalisée à long terme vérifie :

$$\lim_{v \rightarrow \infty} \frac{W_i(0, v)}{v} = r_i \quad (6.34)$$

En OWFS, pour un flux continuellement actif, l'équité réalisée à long terme vérifie :

$$r_i \leq \lim_{v \rightarrow \infty} \frac{W_i(0, v)}{v} \leq r_i m_{max} \quad (6.35)$$

Preuve : cf. annexe 6. G.

Corollaire 6.2 : En EWFS, pour tout intervalle en temps virtuel $[v_1, v_2]$, dans lequel deux flux i et j sont continuellement actifs, la différence de service reçue entre les deux flux vérifie :

$$\left| \frac{W_i(v_1, v_2)}{r_i} - \frac{W_j(v_1, v_2)}{r_j} \right| \leq \frac{L_{max}}{r_i} + \frac{L_{max}}{r_j} + \frac{\alpha L_{max}}{\min(\phi_i, \phi_j)} \quad (6.36)$$

En OWFS, pour tout intervalle en temps virtuel $[v_1, v_2]$, dans lequel deux flux i et j sont continuellement actifs, la différence de service reçue entre ces flux vérifie :

$$\left| \frac{W_i(v_1, v_2)}{r_i} - \frac{W_j(v_1, v_2)}{r_j} \right| \leq (v_2 - v_1)(m_{max} - 1) + \frac{L_{max}}{r_i} + \frac{L_{max}}{r_j} + \frac{\alpha L_{max}}{\min(\phi_i, \phi_j)} \quad (6.37)$$

$\alpha = 2$, si le flux i reçoit son service à v_1

$\alpha = 1$, si le flux i reçoit son service à $v_1 - \varepsilon$

Preuve : cf. annexe 6. G.

Le théorème 6.3 et le corollaire 6.2 illustrent une différence fondamentale entre l'EWFS et l'OWFS, la différence de service entre deux flux n'est pas indépendante de la longueur de l'intervalle en OWFS. A cause de l'aspect opportuniste de l'OWFS, la différence de service entre deux flux peut augmenter avec la taille de l'intervalle.

2.3.4. Garanties de délai

Théorème 6.4 : En EWFS et en OWFS, le délai d'un paquet de tête du flux i vérifie :

$$P\left(d_i \leq \frac{L_{max}}{\phi_i C_{gar, \beta}} + \sum_{j \in F} \frac{L_{max}}{C_{gar, \beta}}\right) \geq \beta \quad (6.38)$$

Si de plus les facteurs r_i et Φ_i sont invariants, on a :

$$P\left(d_i \leq \frac{L_{max}}{\phi_i C_{gar, \beta}} + \frac{L_{max}}{C_{gar, \beta}}\right) \geq \beta \quad (6.39)$$

avec β et $C_{gar, \beta}$ tels que $P(C \geq C_{gar, \beta}) \geq \beta$

Preuve : cf. annexe 6.H.

Définition 6.2

Pour un paquet arbitraire du flux i , le délai est défini par rapport à l'heure d'arrivée attendue ou *Expected Arrival Time* (EAT). L'EAT est une estimation de l'heure à laquelle un paquet quelconque devient le paquet de tête de queue. La définition de l'EAT est une translation de l'étiquette de début de service dans le temps réel, on utilise la même définition pour l'EWFS et l'OWFS :

$$EAT(p_i^{k+1}) = \max\left\{t_i^k, EAT(p_i^k) + \frac{L(p_i^k)}{r_i C_{gar, \beta}}\right\} \quad (6.40)$$

Théorème 6.5 : En EWFS, l'heure de départ du paquet p_i^k (*Packet Departure Time, PDT*) du flux i vérifie :

$$P\left(PDT(p_i^k) \leq EAT(p_i^k) + \frac{L_{max}}{\phi_i C_{gar,\beta}} + \sum_{j \in F} \frac{L_{max}}{C_{gar,\beta}}\right) \geq \beta \quad (6.41)$$

Si de plus les facteurs r_i et Φ_i sont invariants, on a :

$$P\left(PDT(p_i^k) \leq EAT(p_i^k) + \frac{L_{max}}{\phi_i C_{gar,\beta}} + \frac{L_{max}}{C_{gar,\beta}}\right) \geq \beta \quad (6.42)$$

avec β et $C_{gar,\beta}$ tels que $P(C \geq C_{gar,\beta}) \geq \beta$.

En OWFS, l'heure de départ d'un paquet vérifie les relations (6.41) et (6.42).

Preuve : cf. annexe 6.H.

2.3.5. Modèle de canal et garanties de service ou de délai

Dans les théorèmes 6.2, 6.4 et 6.5, à β donné on considère $C_{gar,\beta}$ tels que $P(C \geq C_{gar,\beta}) \geq \beta$. Dans cette section, nous explicitons quelques couples $(\beta, C_{gar,\beta})$ pour un modèle de canal donné.

Le canal est constitué de S sous canaux. Sur chaque sous canal, le mode m a une capacité c_m et peut être choisi avec la probabilité π_m . Le mode $m = 0$ est exclu car le service est sans erreurs. La capacité d'un sous canal est une variable aléatoire dont la moyenne est égale à $c_{avg} = \sum_{m=1}^M \pi_m c_m$.

Pour toute variable aléatoire X , l'inégalité de Bienaymé-Tchebitcheff donne :

$P(|X - E(X)| \geq k \sigma) \leq 1/k^2$ où σ est l'écart type de la variable aléatoire X . On a alors $P(C \geq (c_{avg} - k \sigma)) \geq (k^2 - 1)/k^2$. La probabilité β vaut $(k^2 - 1)/k^2$. Pour $k = 2$, la probabilité β vaut 0.75 ; pour $k = 5$, la probabilité β vaut 0.95. La capacité sur un sous canal est minorée par c_1 .

Pour la capacité globale du canal, on peut finalement écrire : $C_{gar,0.75} = S \max(c_{avg} - 2\sigma, c_1)$, $C_{gar,0.95} = S \max(c_{avg} - 5\sigma, c_1)$ et $C_{gar,1} = S c_1$.

2.4. Caractérisation du WFS, EWFS et OWFS

2.4.1. Découplage entre le débit et le délai

En présence d'un seul poids de débit, le débit d'un flux et les délais subis par ses paquets sont fortement imbriqués. Ce n'est pas le cas en WFS, où l'on peut découpler ces deux caractéristiques. Dans [Lu&_98], un scénario illustre cette propriété. Nous étendons ce scénario au contexte OFDMA pour montrer que l'EWFS et OWFS ont conservé le découplage entre le débit et le délai.

Scénario initial ([Lu&_98])

On considère trois flux dont les poids de débit sont 0.11, 0.44, 0.44 et les poids de délai sont respectivement 0.9, 0.09, 0.009. On considère un service sans erreurs. Les sources sont Poisson avec des taux d'arrivées proportionnels aux poids de débit¹ (le vecteur des taux d'arrivées est noté λ). Le *lookahead* est infini. On simule des *slots* temporels de 1 ms et un paquet peut être envoyé par *slot*.

Extension du scénario pour l'EWFS et l'OWFS

On considère M modes et S sous canaux indépendants entre eux. On considère un service sans erreurs donc le mode $m = 0$ est exclu. L'évolution de l'état d'un sous canal s d'un *slot* temporel à l'autre est modélisée par une chaîne de Markov à M états. On note π le vecteur de probabilité stationnaire de la chaîne (c'est à dire la probabilité stationnaire d'être dans un mode donné). Le vecteur c donne le nombre de paquets que l'on peut envoyer dans chaque mode. Soit λ' le vecteur déterminant les taux d'arrivées des sources, $\lambda' = S(c.\pi)\lambda$. Dans cette expression, on distingue :

- S le nombre de sous canaux,
- λ le taux d'arrivée des sources sachant qu'un paquet est transmis par *slot*,
- le produit scalaire de π et c qui représente le nombre moyen de paquets écoulés par un sous canal sur un *slot*.

Tableau 6.1 : Paramètres de modélisation du canal

$M = 2$	$M = 5$
$c = [1 \ 2]$ $\pi = [0.5 \ 0.5]$	$c = [1 \ 2 \ 3 \ 6 \ 9]$ $\pi = [0.36 \ 0.37 \ 0.18 \ 0.06 \ 0.01]$

Dans les tableaux 6.2 à 6.4, on distingue *simulation entière* (50000 *slots*, moyenne sur 25 simulations) et *petites fenêtres* (5 fenêtres de 200 *slots* sur 5 simulations). On s'intéresse à :

- r : rapport entre le nombre de paquets transmis par une source et le nombre total de paquets transmis,
- D_{max} : délai² maximum recensé parmi les paquets transmis,
- $D_{average}$: délai moyen des paquets transmis,
- σ_D : écart type du délai
- d^{nq} : délai maximum d'un paquet de tête de queue.

Le tableau 6.2 donne la valeur des paramètres pour le WFS. On trouve des valeurs identiques pour l'OWFS et l'EWFS lorsque $M = 1$ et $S = 1$. Cela confirme que l'OWFS et l'EWFS se réduisent au WFS sur un unique canal ON-OFF.

1 Le taux d'arrivée des paquets de la source 1 est de 100 paquets par seconde, celui de la source de 2 et 3 est de 400 paquets par seconde.

2 Résultats donnés en nombre de *slots*

Dans le tableau 6.2, on voit que la proportion de paquets transmis est le reflet des poids de débit. Bien que le flux 1 transmette le moins, ses paquets subissent le délai moyen le plus faible. Il transmet peu car il a un poids de débit faible et ses paquets attendent peu car il a un poids de délai élevé. En l'absence de poids de débit, un poids de débit faible aboutit à un délai de paquets élevé. D'un autre côté, le flux 3 transmet quatre fois plus que le flux 1 en accord avec les poids de débit. Le fait de transmettre une grande proportion des paquets de ce flux n'a pas diminué leur délai moyen. En effet, le flux 3 a un poids de délai très faible et son poids de débit élevé n'influence pas le délai.

L'existence du poids de délai permet de découpler le débit d'un flux et le délai de ses paquets. Ce constat est vérifié aussi bien sur une longue période simulation que sur de courtes fenêtres. L'équité proportionnelle est donc vérifiée à court et à long terme.

Tableau 6.2 : Découplage débit/délai, WFS et extensions pour ($M = 1, S = 1$)

	Flux	Poids débit	Poids délai	r	$D_{average}$	D_{max}	d^{nq}	σ_D
Simulation entière	1	0.11	0.9	0.11	0.6	9.3	1.2	0.5
	2	0.44	0.09	0.44	1.1	8.8	4.5	1
	3	0.44	0.009	0.44	10.1	77.7	28.7	10.7
Petites fenêtres	1	0.11	0.9	0.11	0.6	1.2	1	0.3
	2	0.44	0.09	0.43	1	3.7	2.2	0.8
	3	0.44	0.009	0.45	8.4	21	6.6	8.8

Les tableaux 6.3 et 6.4 montrent la valeur des paramètres de l'EWFS et de l'OWFS dans un contexte OFDMA ($M > 1, S > 1$). Les taux d'arrivée des paquets ont été augmenté pour maintenir une charge constante (environ 0.9) par rapport au tableau 6.2. Le découplage entre débit et délai moyen est vérifié aussi bien pour l'EWFS que l'OWFS. En ce qui concerne la comparaison entre OWFS et EWFS, on observe des débits moyens similaires. Par contre les délais maximaux et l'écart type du délai sont moins élevés en OWFS qu'en EWFS.

Tableau 6.3 : Découplage débit/délai, EWFS pour ($M = 5, S = 2$)

	Flux	Poids débit	Poids délai	r	$D_{average}$	D_{max}	d^{nq}	σ_D
Simulation entière	1	0.11	0.9	0.11	1.3	36.5	7.1	2.7
	2	0.44	0.09	0.44	2.7	28.1	6.7	3.2
	3	0.44	0.009	0.44	12	52.1	12.6	7.4
Petites fenêtres	1	0.11	0.9	0.11	1.3	6.5	2.8	2.3
	2	0.44	0.09	0.43	2.5	8.4	2.7	2.4
	3	0.44	0.009	0.45	10.5	20.9	4.2	6.6

Tableau 6.4 : Découplage débit/délai, OWFS pour ($M = 5, S = 2$)

	Flux	Poids débit	Poids délai	r	$D_{average}$	D_{max}	d^{nq}	σ_D
Simulation entière	1	0.11	0.9	0.11	1.2	15.5	10	1.4
	2	0.44	0.09	0.44	2	14.9	7	1.8
	3	0.44	0.009	0.44	9.2	29.9	14	4.7
Petites fenêtres	1	0.11	0.9	0.11	1.1	5	3.8	1
	2	0.44	0.09	0.43	1.8	7	3.2	1.6
	3	0.44	0.009	0.45	8.7	18	5.5	4.6

2.4.2. Garanties de service et de délai

Nous souhaitons illustrer les propriétés de la section 2.4 : le théorème 6.2 sur les garanties de service, le corollaire 6.2 sur la différence de service reçue ainsi que les théorèmes 6.4 et 6.5 sur les garanties de délai. Les illustrations de ces théorèmes ne sont pas réalisées dans [Lu&_98] et [Lu&_00].

Trois flux sont considérés : les poids de débit et de délai et valent (0.11, 0.44, 0.44). Les résultats présentés dans cette section sont réalisés avec les paramètres du tableau 6.1 avec $M = 5$ et $S = 2$.

La figure 6.2 illustre le théorème 6.2 sur les garanties de service. On peut voir que le service reçu par un flux est bien supérieur au service minimum garanti par le théorème. Ici, on a utilisé¹ $C_{gar,0.75} = S \max(c_{avg} - 2\sigma, c_1)$. Le service maximum qu'on a tracé est déduit du lemme 6.4 : il s'agit de $W_i(t) \leq r_i C_{gar,0.75} t + L_{max}$. Le service reçu se rapproche de cette borne maximale lorsque la charge d'entrée ($\lambda' = S(c \cdot \pi) \lambda$, cf. section 2.4.1) est élevée. On établit le même type de figure en OWFS.

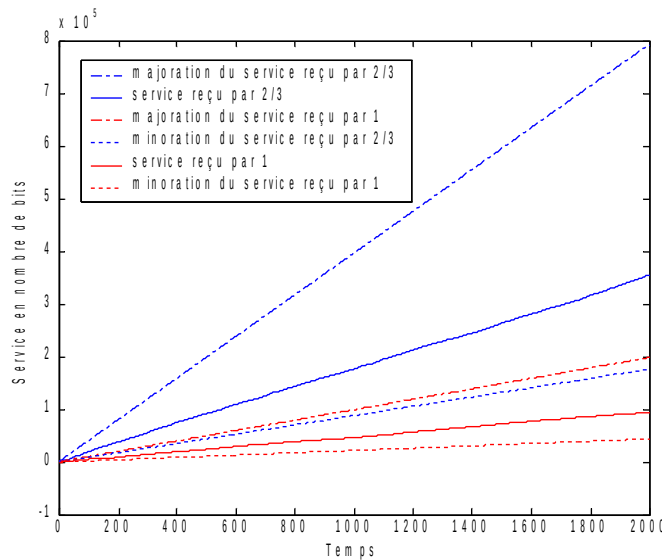


Figure 6.2 : Service reçu en EWFS

Les figures 6.3 et 6.4 illustrent la différence de service entre deux flux en EWFS et en OWFS. Le corollaire 6.2 majore cette différence de service en temps virtuel. Ici, on trace la différence de service en temps réel. On voit sur la figure 6.3 que la différence de service en temps réel peut dépasser la majoration définie en temps virtuel. En effet, en temps réel un flux peut transmettre plusieurs paquets simultanément et ce n'est pas le cas en temps virtuel. En temps réel, la différence de service ne respecte pas exactement la majoration stricte de l'EWFS. En OWFS, cette majoration est moins stricte car contrairement à l'EWFS, elle dépend de l'intervalle et de sa taille. On voit sur la figure 6.4 que pour les flux 1 et 2, on passe d'une majoration de 3000 bits en EWFS à près de 7000 bits en OWFS sur le même intervalle. La différence de service en temps réel en OWFS respecte la majoration définie en temps virtuel car elle est peu stricte.

¹ Pour $c_{avg} - 5\sigma < 0$ dans cette simulation.

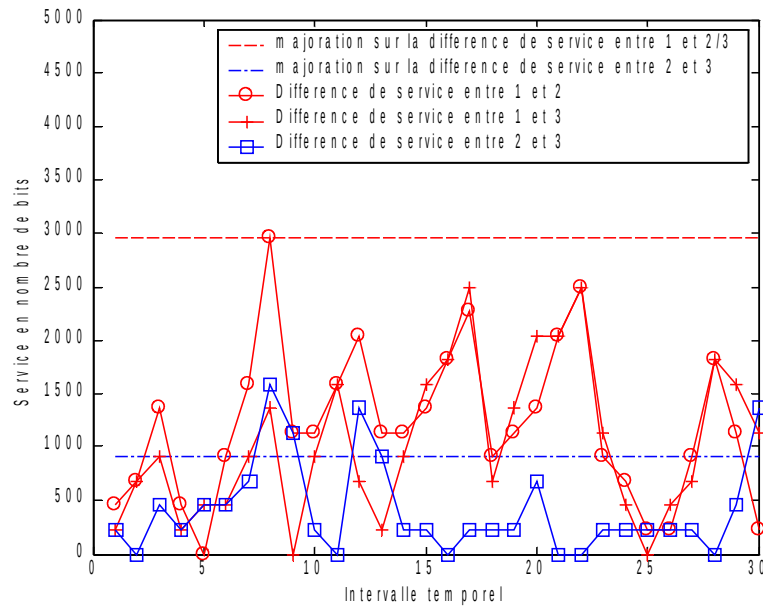


Figure 6.3 : Différence de service entre les flux en EWFS

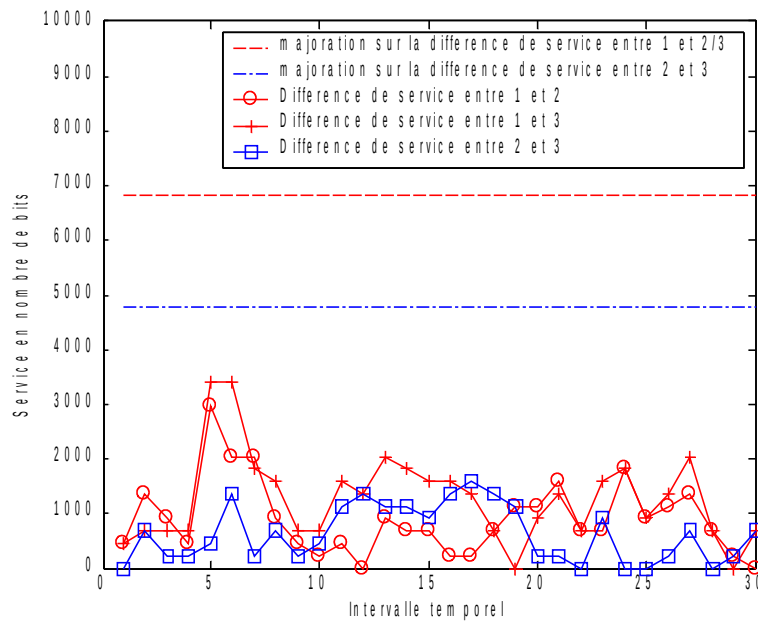


Figure 6.4 : Différence de service entre les flux en OWFS

Le tableau 6.5 illustre les théorèmes 6.4 et 6.5 sur les garanties de délai avec le tableau ci-après. Lorsque les majorations sont calculées avec la capacité $C_{gar,1}$, elles sont toujours vraies mais on voit qu'elles sont très larges : pour le flux 1, le délai moyen maximal est de 151 ms pour un délai observé de 3.1 ms. Avec $C_{gar,0.75} = S \max(C_{avg} - 2\sigma, c_1)$ la majoration devient 12 ms, elle n'est vraie qu'avec une probabilité de 0.75 mais on voit qu'elle est beaucoup plus proche du délai observé.

Tableau 6.5 : Comparaison entre délai moyen et délai maximum en OWFS

Délai moyen (ms)	Borne délai moyen (ms) calculé avec $C_{gar,1} / C_{gar,0.75}$	Délai moyen HOL (ms)	Borne délai moyen HOL (ms) calculé avec $C_{gar,1} / C_{gar,0.75}$
3.4	151 / 12	2.3	50.5 / 4.9
2.2	124 / 8.5	0.6	16.4 / 1.6
2.2	123 / 8.4	0.6	16.4 / 1.6

2.4.3. Dégradation gracieuse de service

Dans cette section on illustre le mécanisme de compensation. Nous reprenons le scénario de [Lu&_98] pour illustrer cette propriété dans le WFS. Nous l'appliquons dans le contexte de l'OFDMA pour observer les caractéristiques de l'EWFS et l'OWFS.

Scénario initial ([Lu&_98])

On considère trois flux constamment actifs (en pratique la charge d'entrée est très élevée). Leur poids de débit et de délai sont identiques (tous à 0.33). Le *lookahead* est infini. Le flux 1 a un mauvais canal de $t = 0$ à $t = 100$ puis un bon canal jusqu'à la fin. Les flux 2 et 3 ont toujours un bon canal. Les valeurs des limites des compteurs $E_{min} = -50$ et $E_{max} = 50$. On observe l'évolution du nombre de paquets transmis en fonction du temps, on simule sur 2000 trames.

Résultats du scénario initial

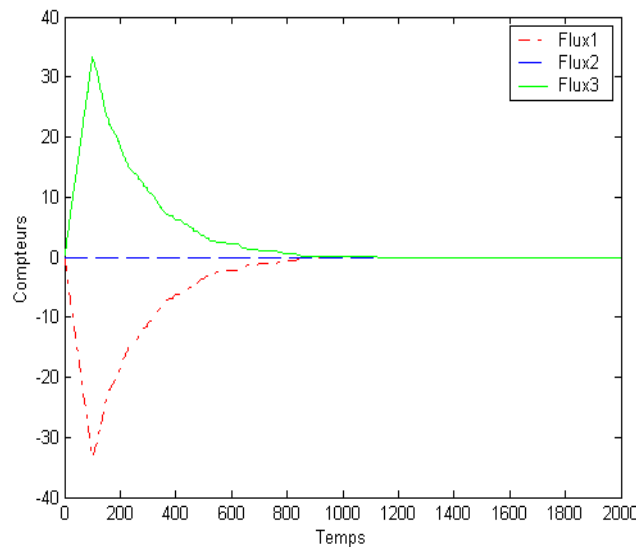


Figure 6.5 : Dégradation exponentielle de service WFS (1)

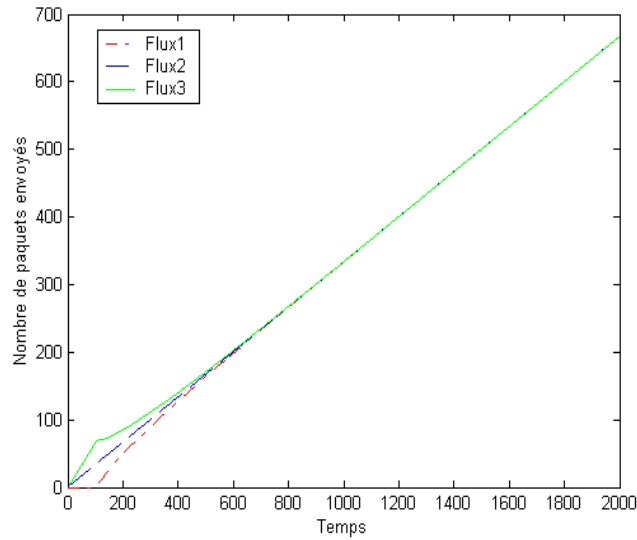


Figure 6.6 : Dégradation exponentielle de service WFS (2)

Les flux sont ordonnancés tour à tour à la manière d'un *round robin* car tous les poids sont égaux. Avant $t = 100$, lorsque le flux 1 est choisi, il doit céder son jeton car son canal est mauvais. La première fois, il y a deux flux synchrones, l'un des deux est choisi arbitrairement pour recevoir le jeton cédé. Dès lors, il y a un flux en retard (ici le flux 1), un flux synchrone (ici le flux 2) et un flux en avance (ici le flux 3). A chaque fois que le flux 1 est ordonnancé, il cède son jeton et le supprime de sa tête de queue de jetons (il en recrée un à la fin de la queue de jetons pour ajuster le nombre de paquets au nombre de jetons). Remarque: S'il ne supprimait pas le jeton cédé, c'est le cas en IWFQ, ce serait toujours le même flux (ici le flux 1 en retard) qui gagnerait l'accès au canal car son jeton de tête aurait la plus faible étiquette de fin de service.

Quand le flux 1 cède un jeton, c'est toujours le flux en avance (le flux 3) qui le reçoit, car on évite au maximum de perturber les flux synchrones. Quand le flux 3 reçoit un jeton du flux 1, il garde son jeton de tête de queue (et supprime son dernier jeton de la queue de jetons pour ajuster le nombre de paquets au nombre de jetons). Garder son jeton de tête de queue, lui permet de garder sa prochaine opportunité de transmission.

Sur la figure 6.6, entre $t = 0$ et $t = 100$, la pente de la courbe du flux 3 est deux fois plus grande que celle du flux 2 car il bénéficie de ses propres opportunités de transmission (il n'a pas supprimé ses jetons de tête de queue) et de toutes celles cédées par le flux 1 qui a un mauvais canal. La pente du flux 2 est la pente qu'auraient tous les flux si le flux 1 n'avait pas eu un mauvais canal. Remarque: En IWFQ, le flux 1 ne supprime pas les jetons cédés du début de sa queue de jeton. A chaque fois, c'est le flux 1 qui est choisi (plus faible étiquette de fin de service) et il cède toujours la transmission au flux 3. Le flux 2 ne peut donc pas émettre entre $t = 0$ et $t = 100$; sa pente serait nulle.

Le flux 2 reste synchrone pendant toute la simulation. Il n'est jamais perturbé. A partir de $t = 100$, le flux 3 cède son avance au flux en retard de façon exponentielle car il cède une proportion $E(3)/E_{max}(3)$ de ses opportunités de transmission. La dégradation exponentielle de service (notée DES dans la suite) se voit particulièrement sur l'évolution de compteurs des différents flux sur la figure 6.5.

Pour $M = 1$ et $S = 1$, l'OWFS et l'EWFS se comportent comme le WFS au niveau de la compensation. On obtient exactement les figures 6.5 et 6.6. Voyons ce qu'il en est lorsque M et S augmentent.

Extension du scénario pour l'EWFS et l'OWFS

On considère $M+1$ modes et S sous canaux, chaque sous canal est représenté par une chaîne de Markov (les valeurs du tableau 2.1 sont appliquées). Le flux 1 voit le mode 0 avec une probabilité 1 sur tous les sous canaux durant les 100 premiers *slots*. Puis, à partir de $t=100$, le flux 1 voit le mode 0 avec une probabilité nulle quelque soit les sous canaux. Concernant les flux 2 et 3, ils voient le mode 0 avec une probabilité nulle sur toute la simulation. On simule sur 2000 trames et on fait la moyenne sur 5 simulations. On utilise $E'_{min} = S E'_{min}$ et $E'_{max} = S E_{max}$. Plus il y a de sous canaux, plus il y a d'opportunités de transmissions échangées : pour harmoniser les conditions de test, nous multiplions les limites de compteurs par le nombre de sous canaux.

Résultats du scénario OFDMA

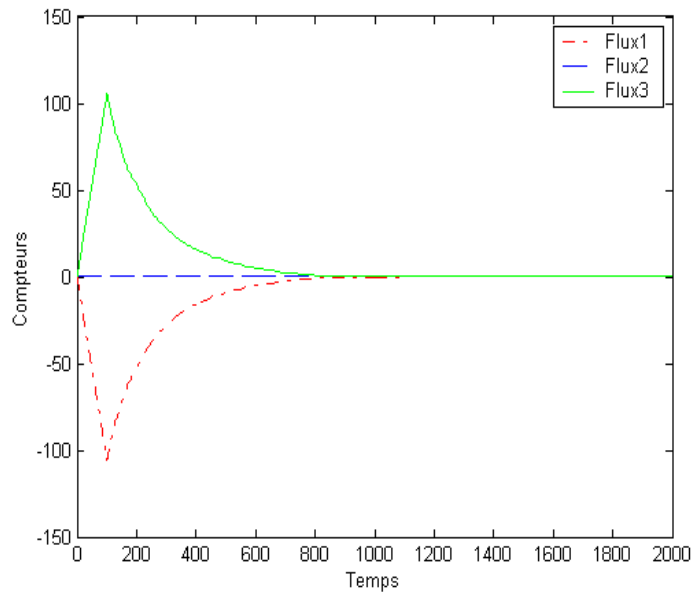


Figure 6.7 : DES en EWFS et OWFS (1)

Le mécanisme de compensation est réalisé en EWFS et en OWFS. On voit en effet que le flux 3 cède exponentiellement son avance au flux 1. L'équité temporelle est atteinte à $t = 800$. On remarque que, à $t = 100$, la valeur de l'avance (respectivement du retard) est plus élevée que dans le cas $M = 1$, $S = 1$; c'était prévisible le nombre d'opportunités de transmission augmente avec M et S et donc le nombre d'échanges aussi.

Les figures 6.8 et 6.9 illustrent, respectivement en EWFS et en OWFS, l'évolution du nombre de paquets transmis par chaque flux lorsque $M = 5$ et $S = 2$.

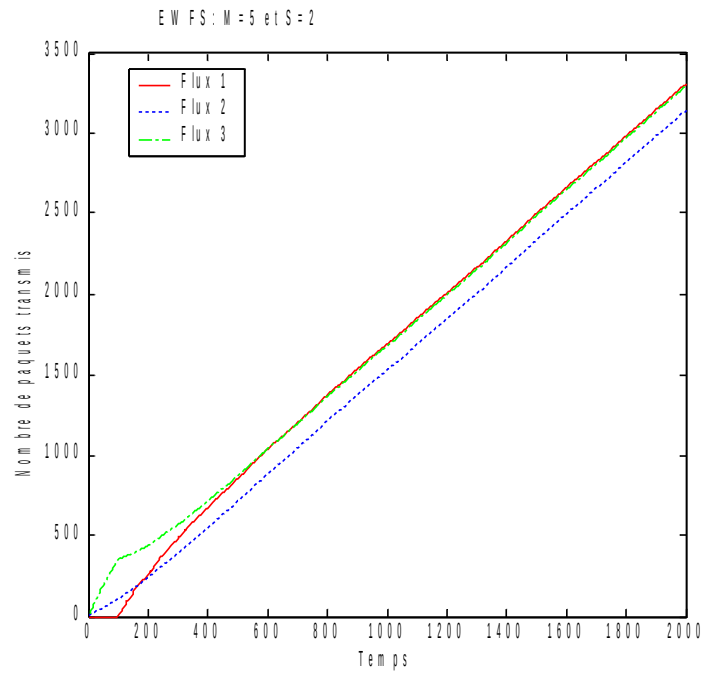


Figure 6.8 : DES en EWFS (2)

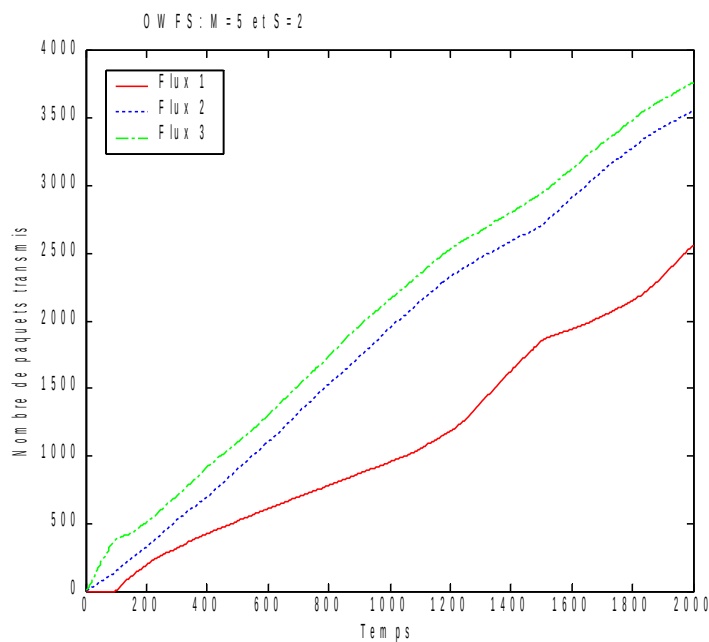


Figure 6.9 : DES en OWFS (2)

La Fig. 6.7 nous a montré que le flux 1 rattrape son retard en terme de jetons c'est à dire d'opportunités de transmissions aux environs de $t = 800$. Concernant le nombre de paquets transmis, on voit que cela n'aboutit pas du tout au même résultat en EWFS et en OWFS. En EWFS, l'équité temporelle est proche d'une équité en débit (cf. Fig 6.8) tandis que ce n'est pas le cas en OWFS (cf. Fig 6.9).

En OWFS, malgré la récupération des jetons cédés, le flux 1 ne parvient pas à transmettre autant que les autres flux. Ce phénomène s'observe à forte charge, il est causé par la définition des étiquettes de début de service en OWFS.

De $t = 0$ à $t = 100$, le flux 1 a un mauvais canal, les étiquettes de début de service des jetons sont calculées avec L/r_i et les étiquettes de fin de service sont calculées avec L/Φ_i . Pour les flux 2 et 3 ces mêmes calculs sont faits avec $L/(m_i r_i)$ et $L/(m_i \Phi_i)$ où $m_i > 1$ (ces flux voient un bon canal). Les étiquettes du flux 1 sont donc plus élevées que celles des flux 2 et 3 et cela favorise ces derniers.

A forte charge, les flux 2 et 3 sont favorisés bien au delà de la période de $t = 0$ à $t = 100$. Les étiquettes des nouveaux jetons sont calculées en fonction de celles des jetons déjà présents : $S(p_i^k) \approx S(p_i^{k-1}) + L / (m_i r_i)$. Donc même si le canal redevient bon (et donc $L/(m_i r_i)$ devient petit), le fait que $S(p_i^{k-1})$ soit élevé pénalise les étiquettes des paquets entrants. On peut considérer cela comme un «effet de mémoire du mauvais canal». Ainsi malgré la restitution des jetons cédés, le retard du flux 1 s'auto-entretient jusqu'à ce que la queue se vide. Ici, elle ne se vide jamais car la charge est trop élevée.

A charge modérée l'«effet mémoire du mauvais canal» s'efface rapidement quand la queue de paquets se vide.

2.4.4. Conclusions sur les caractéristiques des algorithmes

On a pu illustrer qu'en WFS, l'équité proportionnelle est réalisée à court et long terme grâce aux poids de débit. On peut noter que les rapports entre les charges d'entrée doivent vérifier les proportions des poids de débit. On a vu que les poids de délai influencent les délais moyens et délais maximaux des paquets et cela indépendamment des poids de débits.

Les algorithmes EWFS et OWFS maintiennent ces caractéristiques dans le contexte de l'OFDMA. On a vu que l'EWFS présente des délais maximaux plus élevés qu'en OWFS. A très forte charge, un flux longuement pénalisé au niveau radio peut, malgré la compensation, subir une période encore plus longue de «famine» en OWFS.

3. Simulations : ordonnancer deux classes de service en OFDMA

Dans cette section, nous caractérisons nos deux algorithmes EWFS et OWFS par rapport au support de deux classes de service. Une classe dite temps réel (ou RT pour *Real Time*) et une classe de service non temps réel (ou NRT pour *Non Real Time*) sont supportées. Les paquets de la classe RT doivent être transmis avant un délai limite au delà duquel ils sont supprimés. Le débit d'entrée de la classe NRT est supérieur à celui de la classe RT. En 3.2, les algorithmes OWFS et EWFS sont comparés entre eux. En 3.3, les algorithmes OWFS et EWFS sont comparés à d'autres algorithmes proposés en OFDMA et présentés dans la section 1.

3.1. Modèle de canal et paramètres de simulation

Nous considérons un système constitué de S sous canaux indépendants. Chaque sous canal est modélisé par une chaîne de Markov à états finis. Chaque état représente un mode de transmission et correspond à une plage de SNR. Nous reprenons la méthode proposée dans [Liu&_04] ; les auteurs définissent le calcul de la matrice de transition de chaque utilisateur en fonction des paramètres physiques (distance, puissance d'émission, fréquence Doppler etc...). Les sous canaux vus par un utilisateur sont soumis à la même matrice de transition mais le mode de transmission initial est différent. Les utilisateurs sont répartis uniformément dans la cellule. Nous nous intéressons à la voie descendante. Nous considérons une trame TDD (*Time Division Duplexing*) de 5 ms, la sous trame DL (*downlink*) dure 3 ms. Les sous canaux sont alloués à un utilisateur pour la durée de la trame DL. La puissance par sous canaux est fixe. La durée de la simulation est 10000 trames soit 50 s, les résultats sont moyennés sur 15 simulations.

Le trafic d'entrée est Poisson. Le débit du trafic RT est 64 kbps tandis que celui trafic NRT est de 384 kbps. Les paquets RT dont l'attente excède 100 ms sont supprimés.

Tableau 6.6 : Paramètres de simulation en ordonnancement

Paramètres	Valeurs
Nombre de modes : M	5
Nombre de sous canaux indépendants : S	8
Nombre de sous porteuses (de données)	512 (384)
Bande : W	5 MHz
Espace entre 2 sous porteuses	10.93 kHz
Puissance de la BS : P_{BS}	43 dBm

3.2. Comparaison entre EWFS et OWFS

Les résultats présentés dans cette section ont fait l'objet d'une publication ([LengMar&_07]).

3.2.1. Impact de la charge d'entrée

Le nombre de flux RT est fixé à 14. On fait varier la charge d'entrée en terme de nombre flux NRT ; ce nombre varie entre 10 et 20. Le vecteur de charge normalisée correspondant à ce vecteur d'entrée est [0.9 1 1.2 1.4 1.5 1.7]. Il s'agit d'une charge moyenne, la véritable charge est certainement inférieure. Les poids de délai RT et NRT sont identiques.

La figure 6.10 montre que les flux NRT atteignent un meilleur débit en OWFS qu'en EWFS. L'OWFS, en favorisant les flux qui ont les meilleures conditions radio dès l'étiquette de début de service, permet d'augmenter le nombre de paquets transmis par les sources qui ont un débit d'entrée élevé. En fait, en OWFS la charge d'entrée influence beaucoup l'ordre de transmission. On voit qu'à forte charge ($U_{NRT} = 20$) le nombre de flux RT envoyés en OWFS est en baisse : il y a beaucoup de paquets NRT, il y en a toujours un qui a de bonnes conditions radio et les paquets RT sont moins transmis.

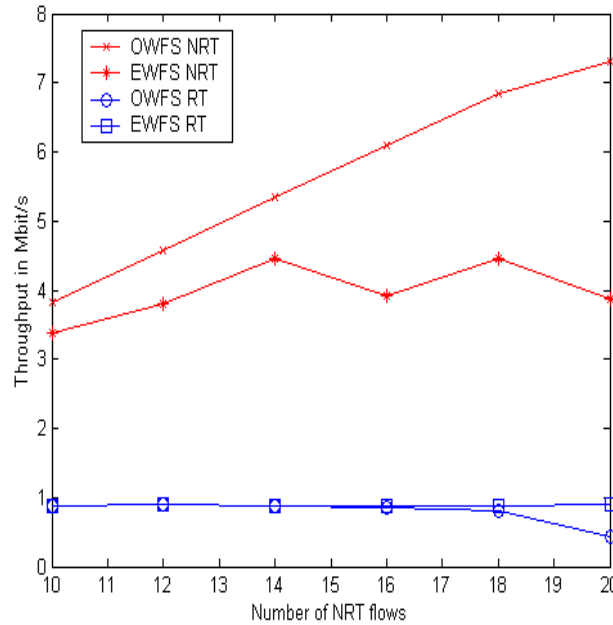


Figure 6.10 : Débit de sortie par classe de service en fonction de la charge d'entrée

La figure 6.11 trace les délais moyens des paquets en fonction de la charge en nombre de flux NRT. En EWFS, le délai des paquets NRT est beaucoup plus élevé qu'en OWFS. En OWFS, le délai NRT est toujours inférieure à 50 ms même à forte charge tandis qu'il est compris entre 500 ms et 1,5 s en EWFS. En EWFS, la charge d'entrée n'influence pas les étiquettes de début de service : les flux dont le débit d'entrée est élevé ont des queues pleines et les délais d'attente des paquets NRT s'en ressentent.

A l'opposé, les paquets RT ont un délai plus faible en EWFS qu'en OWFS. Ici, le poids de délai est identique pour les flux RT et NRT. C'est donc le faible débit d'entrée des flux RT qui les favorise en EWFS : les queues des flux RT sont moins pleines et les paquets attendent moins longtemps.

Sur la figure 6.12, on s'intéresse au nombre de paquets RT qui ont été supprimés : c'est à dire ceux dont le délai d'attente a atteint 100 ms. L'OWFS commence à perdre des paquets RT à partir de $U_{NRT}=16$ tandis que l'EWFS résiste à la charge et ne perd pas de paquets. Au delà de $U_{NRT}=16$, le taux de pertes obtenu en OWFS (supérieur à 10 %) n'est pas tolérable dans un système.

On se place au point ($U_{RT}=14$, $U_{NRT}=16$), le taux de pertes des paquets RT en OWFS est de l'ordre de 3 % lorsque le poids de délai RT est identique à celui des flux NRT. L'OWFS est très intéressant du point de vue du nombre de paquets NRT transmis. Dans la section 3.3, on étudie l'influence du poids de délai sur le taux de pertes des paquets RT en OWFS.

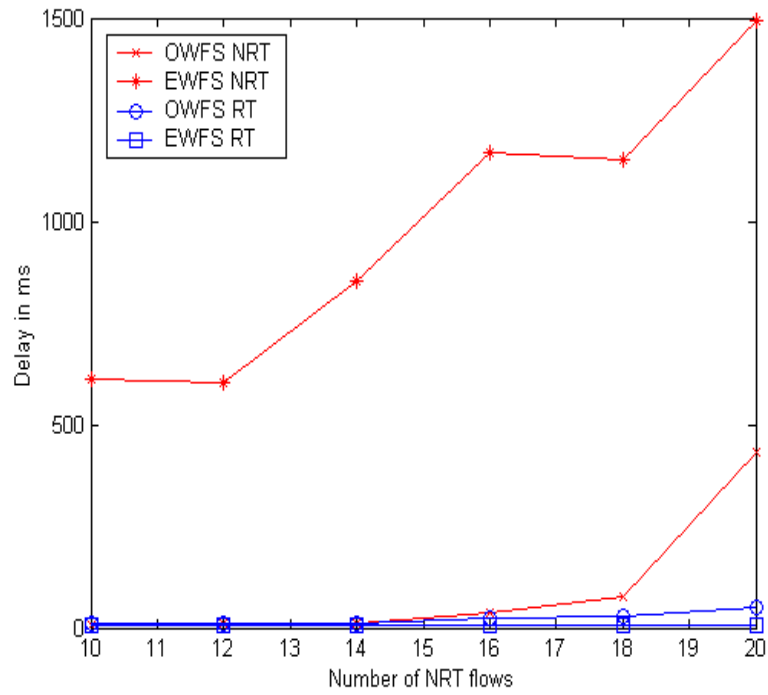


Figure 6.11 : Délais en EWFS et OWFS en fonction de la charge d'entrée

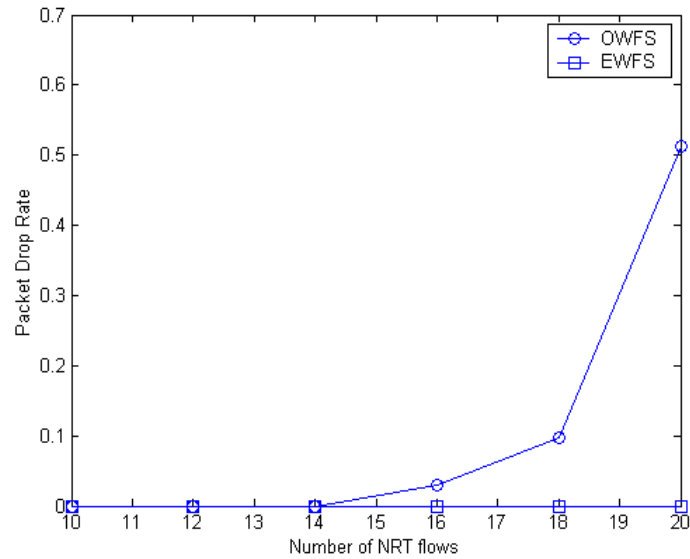


Figure 6.12 : Taux de pertes des paquets RT (EWFS, OWFS) en fonction de la charge d'entrée

3.2.2. Impact du poids de délai

Dans la figure 6.13, en abscisse, le rapport Φ_{RT}/Φ_{NRT} entre le poids de délai RT et le poids délai NRT varie entre 1 et 20. On voit que le taux de pertes RT en OWFS baisse considérablement. A partir de $\Phi_{RT}/\Phi_{NRT}=8$, le taux de pertes RT est meilleur en OWFS.

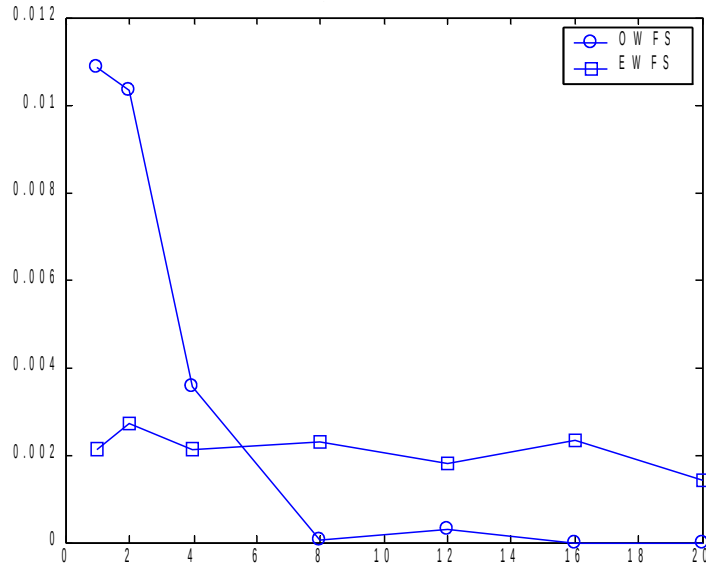


Figure 6.13 : Taux de pertes RT en fonction du rapport entre poids de délai RT et NRT

3.3. Comparaison avec les algorithmes existants en OFDMA

Nous avons vu en 3.2 que l'OWFS permet d'obtenir des débits NRT importants et, dans certaines conditions, de maintenir un taux de pertes RT satisfaisant. Dans cette section, nous montrons que l'OWFS est intéressant y compris par rapport aux algorithmes récents proposés en OFDMA dans la catégorie *ordonnancement global*.

Le nombre de flux RT est 14 et le nombre de flux NRT est 16. Le débit du trafic RT est 64 kbps tandis que celui trafic NRT est de 384 kbps. Les paquets RT dont l'attente excède $D_{max} = 100$ ms sont supprimés. Dans cette configuration nous avons vu qu'un poids de délai RT de 8 et un poids de délai NRT de 1 permettent d'obtenir des taux de pertes satisfaisant en OWFS. C'est ce que nous utilisons dans la suite. La durée de la simulation est toujours de 10000 trames.

3.3.1. Algorithmes comparés

L'EWFS et l'OWFS sont comparés avec les trois algorithmes suivants.

Algorithme proposé par [Jeong&_06] (CLRM)

Nous appelons cet algorithme CLRM pour *Cross Layer Resource Management*. Pour traiter deux classes de flux, le CLRM déclare un paquet RT comme urgent lorsque son attente a dépassé δD_{max} (nous prenons $\delta=0.6$). Cette méthode est comparée à l'OWFS qui utilise plutôt un poids de délai.

Dans notre simulation, une unité d'allocation (UA) s'étend sur une sous trame DL dans l'axe temporel et un sous canal sur l'axe fréquentiel. Trois étapes constituent l'algorithme (cf. section 1.3.2 b) : les utilisateurs reçoivent un nombre d'UA selon l'urgence des paquets RT (étape 1) et les conditions canal (étape 2). A la fin des trois étapes, une phase d'optimisation décide de la répartition des UA. On applique ici une version simplifiée : dans chaque étape, l'utilisateur sélectionné reçoit la meilleure UA parmi celles disponibles.

Extension du PGPS (EPGPS)

Pour évaluer l'intérêt du poids de délai introduit par le WFS, nous souhaitons comparer l'OWFS à une extension du PGPS à l'OFDMA. Le PGPS n'utilise qu'un poids de débit. On pourrait simuler le CPLD-PGPS mais c'est un algorithme avec *allocation optimisée*. En effet, selon les étiquettes de fin de service et le nombre de sous porteuses de bonne qualité, un nombre de sous porteuses est déterminé pour envoyer le paquet (cf. section 1.3.2 b). Une phase d'optimisation répartit les bonnes sous porteuses.

Nous simulons ici un algorithme avec allocation simplifiée que l'on peut appeler EPGPS (pour *Extended Packetized General Processor Sharing*) : les paquets sont sélectionnés par étiquette de fin de service croissante. Si le flux i peut transmettre alors il reçoit son meilleur sous canal. Si le flux i ne peut pas transmettre alors un autre flux j est sélectionné. Le flux initial voit son compteur G incrémenté. Le flux j sélectionné présente le meilleur compromis entre qualité du sous canal et valeur du compteur.

Pour favoriser le trafic RT, le poids de délai RT est plus élevé que celui du trafic NRT.

PLFS

Nous implémentons une adaptation du PLFS (cf. section 1.3.2 a). Les auteurs considèrent une trame temporelle qui est divisée en régions RT et NRT. Dans une région, tous les sous canaux sont alloués à une seule classe. Dans notre adaptation, la sous trame n'est pas divisée en région cependant une proportion des sous canaux est réservée aux flux RT (60 %). Tant que cette proportion n'est pas allouée, la métrique d'ordonnancement RT du PLFS sélectionne l'utilisateur RT ayant la plus grande priorité. L'utilisateur reçoit le meilleur sous canal disponible. Une fois la proportion de RT allouée, la métrique NRT est appliquée jusqu'à l'allocation de tous les sous canaux restants. L'allocation est maintenue pendant toute la durée de la sous trame DL.

Note : Dans [LengMarG_07], le PLFS est comparé à l'OWFS pour un nombre variable de flux NRT. Les trafics RT et NRT sont de type ON-OFF. Les périodes ON et OFF du trafic RT ont des durées exponentielles de moyennes 1 s et 1.35 s. Le trafic NRT correspond à des sessions FTP et sont simulées d'après [WimaxEval_07].

3.3.2. Résultats

Nous nous intéressons à W_{NRT} le taux de paquets NRT transmis, au débit sortant d'un flux NRT et au délai moyen d'attente d'un paquet NRT. La taille des queues correspond à 30 paquets de 1500 octets. Le délai moyen d'un paquet RT est observé et le taux de paquets perdus pour cause de dépassement du délai d'attente est calculé.

Dans le tableau 6.7, on voit que l'OWFS permet d'obtenir le meilleur débit en observant le débit sortant du trafic NRT. L'OWFS est suivi par le CLRM, puis l'EWFS, le PLFS et enfin l'EPGPS réalise le moins bon débit sortant NRT.

Concernant le trafic RT, on n'a pas observé de pertes pour l'EPGPS et le CLRM. Le PLFS suit avec une perte de l'ordre de 10^{-5} , puis l'OWFS avec une perte de 10^{-4} et l'EWFS avec une perte de l'ordre de 10^{-3} .

Grâce aux poids de délai choisis, l'OWFS est meilleur que l'EWFS à la fois vis à vis du débit sortant NRT (de presque 100 kbps de plus) et par rapport au taux de pertes RT.

Ne disposant pas de poids de délai, l'EPGPS sacrifie le débit du trafic NRT pour assurer au trafic RT un bon taux de pertes. L'EPGPS n'est donc pas compétitif face à l'OWFS : la différence de débit NRT s'élève à 240 kbps.

Le PLFS assure un taux de pertes RT satisfaisant (de l'ordre de 10^{-5}) avec une réservation de 60 % des ressources RT. Cependant le débit NRT est moins bon que celui de l'OWFS d'environ 150 kbps.

Des quatre algorithmes comparés, c'est l'algorithme CLRM qui est l'algorithme le plus compétitif par rapport à l'OWFS. Aucune pertes RT n'ont été enregistrées contre un taux de pertes de l'ordre de 10^{-4} pour l'OWFS. Toutefois le débit NRT du CLRM est inférieur de 50 kbps à celui de l'OWFS. Un autre avantage de l'OWFS concerne le délai moyen d'attente des paquets : un paquet NRT attend en moyenne 640 ms en CLRM contre 20 ms en OWFS, tandis qu'un paquet RT attend en moyenne 22 ms en CLRM contre 34 ms en OWFS.

Tableau 6.7 : Performance des algorithmes d'ordonnancement pour deux classes de service

	OWFS	EWFS	EPGPS	CLRM	PLFS
W_{NRT}	0.87	0.83	0.74	0.85	0.81
Débit moyen sortant d'un flux NRT (kb/s)	382	286	156	325.5	235
Délai NRT en ms	20	1200	2400	640	1400
Intervalle de confiance à 95 % sur D_{NRT} (ms)	[13, 27]	[1040, 1360]	[1980, 2820]	[498, 782]	[1300, 1500]
Taux de pertes RT	$1.3 \cdot 10^{-4}$	$2.6 \cdot 10^{-3}$	0	0	10^{-5}
Délai moyen RT en ms	22	5.8	2.5	34	6.8
Intervalle de confiance à 95 % sur D_{RT} (ms)	[14.6, 29.4]	[2.7, 8.9]	[1.8, 3.2]	[23.5, 44.3]	[2.8, 10.8]

4. Conclusions

4.1. Travail effectué et principaux résultats

Dans la partie bibliographique de ce chapitre, nous nous sommes intéressés à des algorithmes inspirés du PGPS (*Packetized General Processor Sharing*) car ils fournissent aux différents flux une équité proportionnelle. Nous avons revus les extensions de ce type d'algorithmes au domaine sans fils ; ces algorithmes tentent de maintenir une équité sur le long terme à cause de la variabilité des conditions radio. Un algorithme, le WFS ([Lu&_00]), a retenu notre attention car il introduit un poids de délai en plus du poids de débit. Grâce à ce poids de délai, plusieurs classes de service peuvent être ordonnancées de façon spécifique, tout en respectant l'équité proportionnelle.

Nous avons vu que peu d'algorithmes inspirés du PGPS ont été proposés dans le contexte de l'OFDMA avec modulation adaptative. Le CPLD-PGPS ([Diao&_04]) ne considère pas de modulation adaptative.

D'autre part, en OFDMA, quelques algorithmes ont été proposés. Nous avons distingué l'ordonnancement indépendant par sous canaux et l'ordonnancement global. Certains algorithmes considèrent une métrique appliquée indépendamment sur les sous canaux ([Wengerter&_05]) ; seul le trafic NRT est alors considéré. D'autres algorithmes, plus intéressants, proposent un ordonnancement global pour plusieurs types de trafic ([Shin&_05],[Jeong&_06]).

Nous nous sommes intéressés à un algorithme d'ordonnancement global capable de traiter simultanément plusieurs types de trafic et basé sur le PGPS. C'est pourquoi nous avons défini l'EWFS et l'OWFS, deux extensions du WFS en OFDMA. L'EWFS introduit la qualité du meilleur canal vu par un utilisateur dans l'étiquette de fin de service pour favoriser les flux qui ont de bonnes conditions radio. L'OWFS généralise le concept en introduisant aussi l'état du canal dans l'étiquette de début de service. Pour ces deux algorithmes, nous avons généralisé la définition du temps virtuel à l'OFDMA où la capacité est variable dans le temps et d'un sous canal à l'autre.

Dans ce chapitre, nous avons caractérisé les propriétés de l'EWFS et l'OWFS dans l'hypothèse du modèle de service sans erreurs. Nous avons démontré des garanties de service minimum et des garanties de délai pour les paquets (cf. section 2.4). L'équité proportionnelle est garantie à long terme en EWFS et plutôt à court terme en OWFS. Nous avons illustré ces propriétés fondamentales en 2.5.1 et 2.5.2.

Nous avons défini un modèle de compensation commun à l'OWFS et l'EWFS. Mais ce modèle de compensation reste à caractériser du point de vue des garanties de service et de délai notamment pour les flux en retard. Dans ce chapitre, les démonstrations n'ont porté que sur le modèle sans erreurs.

En comparant l'OWFS et l'EWFS aux principaux algorithmes existants en OFDMA, nous avons montré que l'OWFS est prometteur. Il peut gérer des classes de trafic différentes avec efficacité. Pour le trafic NRT, il réalise des débits beaucoup plus importants que les autres algorithmes. Pour le trafic RT, un taux de pertes satisfaisant peut être réalisé par un choix adéquat du poids de délai. On peut noter que l'algorithme proposé par [Jeong&_06] est assez compétitif par rapport à l'OWFS. Mais l'OWFS réalise tout de même des délais moyens inférieurs dans les deux classes de trafic.

L'OWFS traite simultanément différentes classes de trafic en OFDMA sans réservation de ressources préalables.

4.2. Limitations

On peut citer une faiblesse relative de l'OWFS, héritée du WFS et déjà citée par [Goyal&_97], il s'agit du calcul du temps virtuel. Par soucis de précision, il faut parallèlement évaluer les départs des paquets en GPS. Mais cette faiblesse n'est que relative car les calculs en jeu ne sont comportent que des opérations élémentaires (cf. section 2.1.1). De plus dans un système, nous proposons d'évaluer la capacité résultante du canal grâce aux mesures des terminaux, effectuées sur les pilotes et envoyées à la BS sur un canal logique tel que le CQICH.

Le modèle de compensation a été peu caractérisé mais un travail préliminaire a été

réalisé. Par un scénario défini dans [Lu&_00], nous avons illustré la compensation des flux en retard en EWFS et OWFS. On pressent qu'à forte charge, l'OWFS peut pénaliser les flux dont les conditions radio sont médiocres longtemps. En EWFS, l'équité temporelle se traduit par une équité en débit, ce qui n'est pas le cas à forte charge en OWFS. En OWFS, selon la charge, l'équité proportionnelle peut être mise à mal par la variation des conditions radio.

Enfin, on a pu constater en simulation que l'OWFS et l'EWFS sont sensibles à la charge d'entrée et notamment au rapport de charge entre les classes. Il n'existe pas, pour l'instant, de forme close permettant d'obtenir la valeur adéquate du poids de délai à partir des charges d'entrée et du taux de pertes à réaliser.

4.3. Améliorations éventuelles

Nous avons fait le choix d'un ordonnancement avec *allocation simplifiée* pour l'OWFS et l'EWFS. Le fonctionnement est le suivant : tant qu'il y a des sous canaux libres, le flux dont le paquet de tête présente la plus faible étiquette de fin de service est sélectionné. Le flux sélectionne son meilleur sous canal et transmet autant de paquets que les conditions radio le permettent. Ces paquets ne correspondent pas forcément à ceux qui auraient été choisis (tous flux confondus) au strict regard des étiquettes de fin de service. Cela crée une différence entre le service reçu en temps réel et celui reçu en temps virtuel, comme nous l'avons illustré en section 2.5.2. On pourrait reprocher à ce type d'allocation d'augmenter la différence de service entre deux flux et de dégrader l'équité proportionnelle à court terme.

Un travail intéressant serait de réaliser un ordonnancement OWFS avec *allocation optimisée* (comme le CPLD-PGPS). Une première phase serait alors de sélectionner un certain nombre de paquets strictement en fonction de l'étiquette de fin de service. Cela pourrait éviter de dégrader l'équité proportionnelle. Une deuxième phase serait une phase d'optimisation par rapport à l'affectation des sous canaux (ou des sous porteuses). Nous avons présenté dans le chapitre 2 plusieurs algorithmes pour réaliser cette phase selon les problèmes MA ou RA. Mais deux questions importantes sont à résoudre pour réaliser un ordonnancement avec *allocation optimisée*. Premièrement, comment prendre en compte les conditions radio dans les étiquettes de fin de service (et éventuellement de début de service) ? En effet, on peut pas savoir ce que va recevoir l'utilisateur : l'affectation des ressources n'a lieu que plus tard et plusieurs utilisateurs peuvent réclamer la même ressource. Alors quel état du canal introduire dans les étiquettes de début et de fin de service ? Deuxièmement, combien de paquets sélectionne t'on pour la phase d'optimisation ? C'est d'autant plus difficile à prévoir que la modulation est adaptative et que l'affectation des ressources n'est pas encore réalisée. Le CPLD-PGPS a simplifié le problème en prenant une modulation est fixe.

Une fois ce travail réalisé, on pourrait comparer l'équité à court terme d'un algorithme OWFS avec *allocation simplifiée* et celle d'un algorithme OWFS avec *allocation optimisée*. On verrait alors si la différence de service en temps réel est améliorée par l'*allocation optimisée*.

Chapitre 7. Conclusion générale

Nous l'observons chaque jour, les réseaux sans fils doivent répondre à des demandes croissantes en terme de débit et de qualité de service.

Les techniques d'accès multiple jouent un grand rôle sur les performances obtenues. L'OFDMA est une technique d'accès multiple hybride entre le TDMA, le FDMA (*Frequency Division Multiple Access*) et la modulation OFDM. C'est une candidate prometteuse pour de nouvelles normes cellulaires et de réseaux d'accès large bande. Elle a été choisie pour le sens descendant de la norme LTE (*Long Term Evolution*) et pour les sens montant et descendant de la norme IEEE 802.16. Grâce à cette technique, les interférences inter symboles dues aux multitrajets sont diminuées, des schémas variés de modulation et codage correcteurs d'erreurs sont supportés ; des débits élevés sont possibles, jusqu'à 80 Mbps dans 20 MHz.

1. Travaux réalisés

En OFDMA, il faut répartir la puissance et les ressources fréquentielles à chaque *slot* temporel. Cela doit être fait en fonction de la charge des cellules, du trafic des utilisateurs et de leurs conditions radio. Dans cette thèse nous avons étudié ces problèmes d'ordonnancement et d'allocation de ressources radio de manière progressive.

Dans la première partie, nous avons résolu le problème d'allocation de ressources radio en faisant l'hypothèse que les utilisateurs avaient toujours des données à transmettre et qu'il n'existait pas de contraintes temporelles. Dans le contexte monocellulaire, nous avons présenté les algorithmes existants qui exploitent la diversité multiutilisateur. Nous avons montré leurs limites et proposé un algorithme le RPO (*Rate Profit Optimization*) dont les performances s'approchent de l'algorithme optimal Hongrois. Cet algorithme s'inscrit dans le cadre de la maximisation du débit de la cellule sous contraintes de puissance et de débits individuels (problème RA). Concernant les débits individuels, nous avons recherché le plus grand débit minimal commun réalisable. Dans le contexte multicellulaire, nous avons proposé l'OSA-IL (*Opportunist Subcarrier Allocation with Interference limitation*) pour résoudre le problème RA. Dans cet algorithme, le facteur de réutilisation fréquentiel (FRF) des sous porteuses varie en fonction des conditions radio des utilisateurs qu'il reste à satisfaire. D'après les simulations, un utilisateur peut obtenir un débit minimum de 800 kbps sur une bande accumulée de 500 kHz et un FRF moyen de $2/3$ (c'est à dire une sous porteuse utilisée dans deux cellules sur trois).

Le RPO et l'OSA-IL ont pour point commun d'affecter les sous porteuses individuellement. En pratique, la sous canalisation est souvent adoptée. Nous avons étudié l'impact de la sous canalisation sur les performances en débit. Nous nous sommes situés dans le contexte de l'IEEE 802.16 où une sous canalisation de sous porteuses adjacentes et une sous canalisation dite de diversité sont définies. Nous avons observé une différence de 8 % entre les performances en débit du RPO et celles de la sous canalisation adjacente (dite AMC, *Adaptive Modulation and Coding*). Cette dernière réalise un débit supérieur de 23 % à celui de la sous canalisation de diversité. La sous canalisation de diversité présente l'intérêt d'être peu sensible à la connaissance de l'état du canal et donc à la mobilité des terminaux.

Dans une deuxième partie, nous avons considéré deux classes de trafic : une classe RT (*Real Time*) et une classe NRT (*Non Real Time*). Notre objectif a été de maintenir un taux de pertes satisfaisant pour le trafic RT, de maximiser le débit du trafic NRT tout en maintenant une équité proportionnelle entre les flux. Deux algorithmes ont été proposés : l'EWFS (*Enhanced Wireless Fair Service*) et l'OWFS (*Opportunistic Wireless Fair Service*). Il s'agit de deux extensions du WFS (*Wireless Fair Service*) ; ce dernier découple le débit et le délai des flux grâce à un poids de délai.

Nous avons actualisé la définition du temps virtuel en présence de plusieurs sous canaux de capacité variable. Les deux algorithmes sélectionnent les utilisateurs par étiquette de fin de service croissante. L'étiquette de fin de service représente le temps virtuel auquel le paquet aurait terminé son service dans le GPS (*Global Processor Sharing*). Dans ces deux algorithmes, meilleures sont les conditions radio d'un utilisateur, plus faibles sont les étiquettes de fin de service de ses paquets. L'OWFS, à la différence de l'EWFS, prend en compte l'état du canal dans l'étiquette de début de service. Cette particularité lui permet de maximiser le débit des flux NRT. Un poids de délai judicieux permet en OWFS de maintenir un taux de pertes acceptable pour le trafic RT. Nous abordons le problème du choix du poids de délai dans la section suivante.

2. Travaux futurs

Dans la deuxième partie, nous avons montré que l'OWFS est un algorithme prometteur. Il reste cependant du travail pour faciliter sa mise en oeuvre dans un système.

Pour généraliser le contexte d'étude, les simulations présentées en section 3.3 (chapitre 6) pourraient être réalisées avec les modèles de trafic définis dans [WimaxEval_07] à la place d'un trafic de type Poisson. Nous avons utilisé ce type de modèle dans [LengMarG_07] ; grâce à ces travaux préliminaires, on peut conjecturer la généralisation des résultats de la section 3.3 (chapitre 6) à un modèle de trafic *bursty*.

Le modèle de compensation choisi pour l'OWFS et l'EWFS maintient une équité temporelle. Pour des flux de conditions radio similaires, l'équité en débit peut être atteinte. Durant les simulations, le modèle de compensation a été peu sollicité. Une raison simple explique cela : la présence de nombreux sous canaux a évité aux utilisateurs de se retrouver dans l'incapacité de transmettre. Certaines règles de priorité

ont été définies selon la catégorie (RT ou NRT) des flux en retard. Pour montrer l'intérêt de ces règles, il faudrait les comparer aux règles du WFS dans une simulation où la couche TCP est modélisée. Faute de temps, cette simulation et les démonstrations associées n'ont pas été réalisées au cours de la thèse.

Enfin, la région *schedulable* de l'algorithme a été démontrée en temps virtuel. Dans cette région, nous avons démontré des propriétés très intéressantes (service minimum, garanties d'équité, garanties de délai) dans le modèle sans erreurs. Des propriétés similaires doivent être démontrées pour le modèle de compensation.

Il faudrait définir, en temps réel, la zone de charge dans laquelle ces propriétés sont réalisées. La résolution de ce problème est une clé pour l'utilisation de l'OWFS. On pourrait alors obtenir une forme close de l'expression du poids de délai d'un flux en fonction du délai maximum et de la charge d'entrée. Nos simulations nous ont permis de déterminer le poids de délai adéquat dans une situation donnée ; une forme close permettrait d'utiliser l'OWFS dans un système réel.

Annexes

Les annexes sont numérotées en fonction du chapitre auquel elles se rapportent ; les notations adoptées correspondent donc à celles de la partie qui contient le chapitre.

Annexe 1.A Le canal de transmission

Composantes du signal reçu

La puissance du signal reçu est la résultante de trois composantes : atténuation, effet de masque et évanouissement rapide. Le signal émis subit une atténuation moyenne qui agit sur une distance d de plusieurs kilomètres. L'atténuation est proportionnelle à $K.d^\alpha$ où α est compris entre 2 et 4 et K est une constante dépendant de l'environnement.

Le signal reçu est la somme de différentes versions atténuées et retardées du signal émis à cause de divers phénomènes (réflexions, diffractions, réfractions). C'est pourquoi, autour de l'atténuation moyenne, la puissance du signal oscille suivant des variations d'une période de quelques kilomètres ; ces variations sont modélisées par une loi lognormale. Autour de cette variation lente, une variation plus rapide (sur quelques mètres) peut par exemple être identifiée par une loi de Rayleigh. La loi de Rayleigh est utilisée lorsque les différents trajets ont une importance égale.

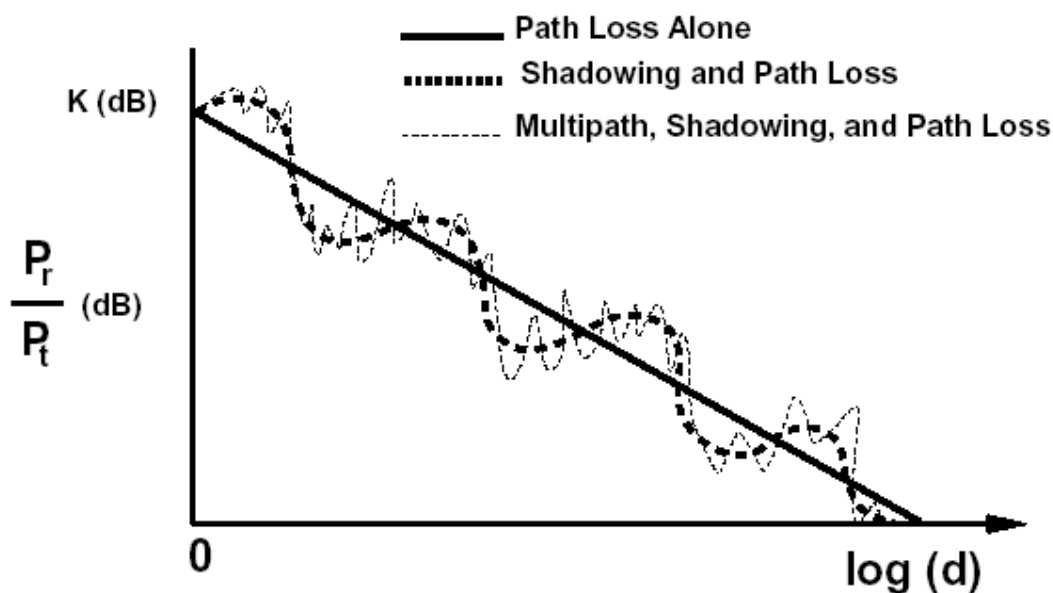


Illustration 1: Variations du signal reçu en fonction de la distance, [Goldsmith]

Propriétés dispersives du canal en fréquence

Le *fading* a des effets sur l'amplitude du signal reçu mais aussi sur la forme de l'impulsion envoyée. Cette dernière peut être élargie car constituée par la réception de plusieurs impulsions décalées en temps. Le temps moyen τ mis par une impulsion pour traverser le canal est calculé en fonction de la puissance et du retard de chaque trajet ([Shankar_01]). Le délai d'étalement σ_{rms} ($\sigma_{rms} = \sqrt{(\langle \tau^2 \rangle - \langle \tau \rangle^2)}$ où rms signifie *root mean square*) détermine si l'impulsion est déformée ou non à la réception. Si σ_{rms} est

nul, l'impulsion n'est pas étalée ; plus σ_{rms} est grand et plus l'impulsion est élargie à la réception. Il est possible de quantifier l'élargissement de l'impulsion en définissant un filtre passe bas de bande B_c représentant le canal. Soit B la bande du signal, le signal n'est pas modifié si $B < B_c$. On parle alors de *flat fading*. Dans le cas contraire, on parle d'un canal sélectif en fréquence.

Le caractère plat ou sélectif en fréquence d'un même canal (même σ_{rms}) dépend du débit des données (en effet, le débit conditionne la bande du signal).

La bande B_c est appelée bande de cohérence du canal ; c'est la bande sur laquelle la corrélation entre les enveloppes de deux fréquences atteint une certaine valeur. Pour un coefficient de corrélation de 0.5 on a $B_c = \frac{1}{2\pi\sigma_{rms}}$; pour un coefficient de corrélation de 0.9, on a $B_c = \frac{1}{50\sigma_{rms}}$ ([Sklar_01]).

Les deux effets, niveau aléatoire de l'enveloppe du signal reçu et sélectivité en fréquence du canal, sont des manifestations séparées du *fading* et peuvent exister individuellement ou conjointement.

Propriétés dispersives du canal en temps

A cause de la mobilité d'un terminal, la fréquence du signal reçu est modifiée et sa variation f_d est la suivante : $f_d = f_0 \frac{v}{c}$. Si le nombre de trajets est suffisamment grand l'enveloppe du signal suit toujours une loi de Rayleigh.

Les changements sur le canal introduits par la mobilité ont lieu autour de f_d Hz. On considère la largeur de l'impulsion, si elle est faible, la bande signal évaluée à l'inverse de cette pulsation sera plus grande que f_d et il y aura peu d'impact : le canal ne sera pas sélectif en temps. La largeur d'impulsion critique qui constitue la limite de la sélectivité en temps est appelée temps de cohérence.

L'expression du temps de cohérence est liée au coefficient de corrélation d'enveloppe dans le domaine temporel. Pour un coefficient de 0.5 on a $T_c = \frac{9}{16\pi f_d}$. Une valeur souvent utilisée résulte de la moyenne géométrique entre l'équation précédente et l'équation simplifiée $T_c = \frac{1}{f_d}$, ce qui donne $T_c = \sqrt{\frac{9}{16\pi f_d}}$ ([Sklar_01]). Si la largeur de l'impulsion est supérieure à ce temps de cohérence, le canal est sélectif en temps.

Pour conclure, B_c et T_c sont la bande et la durée sur lesquelles le canal est supposé constant.

Annexe 1.B L'OFDM

Le lecteur intéressé peut se référer à [ShinsukePrasad_03], nous faisons ici de brefs rappels. La différence fondamentale entre les différentes techniques classiques de modulations multiporteuses et l'OFDM est que cette dernière autorise un fort recouvrement spectral entre les porteuses. Cela permet d'augmenter sensiblement le nombre de porteuses et d'améliorer l'efficacité spectrale.

Le signal OFDM peut s'écrire :

$$s(t) = \sum_k \sum_{n=0}^{N-1} c_{n,k} e^{2i\pi f_n t} g(t - kT_s)$$

Dans cette équation, $c_{n,k}$ est le symbole M-QAM à émettre sur la porteuse f_n d'indice n , T_s est la durée d'un symbole OFDM, $g(t)$ est la forme d'onde de la modulation. Le train binaire initial subit une conversion série-parallèle, sur chaque fréquence les symboles subissent une transformée de Fourier inverse. Après cela, les données parallèles sont sommées et transmises.

Interférences entre symboles et interférence entre porteuses

Pour éviter les interférences entre symboles générées par la dispersion du canal, on doit introduire un intervalle de garde entre les symboles OFDM : $T_s = T_u + T_g$. Cet intervalle T_g doit être plus grand que l'étalement du canal. On reviendra sur le choix du contenu de l'intervalle de garde.

Les porteuses doivent respecter une contrainte d'orthogonalité afin de garantir une récupération simple des données. Pour avoir une orthogonalité fréquentielle, l'espacement entre les porteuses est égal à $1/T_u$ où T_u est la partie utile du symbole OFDM. Ce résultat est lié au choix de la forme d'onde $g(t)$ qui est une fonction porte qui

vaut, par normalisation, $\frac{1}{\sqrt{T_u}}$ sur $[0, T_u]$ et 0 ailleurs. La transformée de Fourier de cette fonction est un sinus cardinal. En choisissant $\Delta f = 1/T_u$, le spectre d'une fréquence est maximal lorsque ceux de toutes les autres sont nuls d'où l'orthogonalité fréquentielle.

Choix de l'intervalle de garde

Le traitement en fréquence est facilité par l'introduction du préfixe cyclique dans l'intervalle de garde : le signal émis durant le temps de garde est la copie des derniers échantillons du paquet. Il en découle que la convolution linéaire entre le filtre du canal et le signal des données peut être remplacée par une convolution circulaire. L'effet du canal pouvant être vu comme une convolution circulaire, la matrice du canal est diagonalisable dans l'espace de Fourier ([Ciblat_04], la transformée de Fourier discrète transforme une convolution en multiplication). Chaque sous-canal est alors caractérisé simplement par un coefficient complexe (chaque sous canal est non sélectif en fréquence), ce qui simplifie l'égalisation au récepteur.

Annexe 1.C Le waterfilling

On considère un unique utilisateur qui doit répartir la puissance entre les différentes sous porteuses afin de maximiser son débit (exprimé avec la formule de Shannon). Le *waterfilling* détermine un *waterlevel* ou niveau d'eau.

On note φ le CgNR (*Channel gain to Noise Ratio*) ; il est égal au rapport entre le carré du module du gain g du canal (résultat de l'atténuation, de l'effet de masque et l'évanouissement) et le bruit présent sur la bande d'une sous porteuse (Δf) :

$$\varphi = \frac{|g|^2}{N_0 \Delta f}$$

Si l'inverse du CgNR d'une sous porteuse est supérieur au niveau d'eau alors la sous porteuse ne reçoit pas de puissance. En fait, plus le CgNR est bas, moins la sous porteuse reçoit de puissance.

Expression du niveau d'eau et du vecteur de puissance

Le problème à résoudre est le suivant :

$$\begin{aligned} \max_{p_n} \sum_{n=1}^N \log_2(1 + p_n \varphi_n) \\ \sum_{n=1}^N p_n = P_{T, \max} \\ h_n(p) \leq 0, \quad 1 \leq n \leq N \quad \text{avec } h_n(p) = -p_n \end{aligned}$$

Maximiser $\sum_{n=1}^N \log_2(1 + p_n \varphi_n)$ revient à minimiser l'opposé. On exprime le lagrangien comme la somme de la fonction à minimiser et des fonctions contraintes associées à leurs multiplicateurs de Lagrange respectifs :

$$L(p, \lambda, \mu) = - \sum_{n=1}^N \log_2(1 + p_n \varphi_n) + \lambda \left(\sum_{n=1}^N p_n - P_{T, \max} \right) + \sum_{n=1}^N \mu_n (-p_n)$$

Note : p et μ sont des vecteurs de N éléments, ils ne seront pas soulignés pour alléger les notations ; λ est un scalaire.

Les conditions nécessaires de Kuhn et Tucker sont les suivantes ([Bertsekas_82], p71) :

$$1 \leq n \leq N, \quad \frac{\partial L(\tilde{p}, \tilde{\lambda}, \tilde{\mu})}{\partial p_n} = 0 \quad (1)$$

$$1 \leq n \leq N, \quad \tilde{\mu}_n \geq 0 \quad \text{et} \quad \tilde{\mu}_n h_n(\tilde{p}) = 0 \quad (2)$$

L'expression des dérivées partielles du lagrangien suivant la variable p est :

$$\frac{\partial L(\tilde{p}, \tilde{\lambda}, \tilde{\mu})}{\partial p_n} = \frac{\varphi_n}{\ln(2)(1 + p_n \varphi_n)} + (\tilde{\lambda} - \tilde{\mu}_n)$$

La condition (1) donne l'expression de la puissance p_n du canal n :

$$\tilde{p}_n = \frac{1}{\ln(2)(\tilde{\lambda} - \tilde{\mu}_n)} - \frac{1}{\varphi_n}$$

La condition (2) nous dit que si $\tilde{p}_n > 0$ alors $\tilde{\mu}_n = 0$. Cela donne l'intuition de la solution ci dessous :

$$\tilde{p}_n = I\left(\frac{1}{\ln(2)\tilde{\lambda}} - \frac{1}{\varphi_n}\right) \text{ avec } I(t) = t \text{ si } t > 0 \text{ sinon } I(t) = 0.$$

$$\tilde{\mu}_n = \tilde{\lambda} - \frac{\varphi_n}{\ln(2)} \text{ si } \left(\frac{1}{\ln(2)\tilde{\lambda}} - \frac{1}{\varphi_n}\right) \leq 0$$

sinon $\tilde{\mu}_n = 0$

Cette solution vérifie les conditions de Kuhn et Tucker. Ces conditions sont suffisantes car la fonction à minimiser est convexe et les contraintes sont convexes (car linéaires). Les vecteurs p et μ exprimés correspondent donc aux vecteurs optimaux $\tilde{p}, \tilde{\mu}$.

Adoptons les notations suivantes :

Soit $v = \frac{1}{\ln(2)\tilde{\lambda}}$; on a alors $\tilde{p}_n = I\left(v - \frac{1}{\varphi_n}\right)$ et v est le *waterlevel* ou niveau d'eau.

On note Ω l'ensemble des canaux utilisés c'est à dire ceux tels que $\tilde{p} \neq 0$ et on note $|\Omega| = \tilde{N}$.

Il reste à expliciter v pour que l'expression des vecteurs $\tilde{p}$ et $\tilde{\mu}$ soit complète. Il suffit d'utiliser la première contrainte, on a (en insérant l'expression de $\tilde{p}$ puis en développant) :

$$\sum_{n=1}^N \tilde{p}_n - P_{T,max} = \sum_{n=1}^N I\left(v - \frac{1}{\varphi_n}\right) - P_{T,max}$$

$$\sum_{n=1}^N \tilde{p}_n - P_{T,max} = \tilde{N} v - \sum_{n \in \Omega} \frac{1}{\varphi_n} - P_{T,max}$$

Finalement, $\sum_{n=1}^N \tilde{p}_n - P_{T,max} = 0 \Rightarrow v = \frac{P_{T,max} + \sum_{n \in \Omega} \frac{1}{\varphi_n}}{\tilde{N}}$

Un algorithme de programmation du waterfilling

Principe

La tâche principale est la détermination du niveau d'eau noté v . Plusieurs algorithmes sont proposés dans la littérature, un algorithme itératif particulièrement simple est proposé dans [PerezFon_05]. Pour calculer le niveau d'eau il faut trouver le zéro de la

fonction $g(v) = \sum_{n=1}^N p_n(v) - P_{T,max}$. Étant donné l'expression de $p_n(v)$,

$$g(v) = Nv - \sum_{n=1}^N \frac{1}{\varphi_n} - P_{T,max}. \text{ Tant que } g(v) > 0, \text{ la contrainte en puissance n'est pas}$$

respectée. Certains canaux sont tellement mauvais que le budget de puissance n'est pas suffisant : ils doivent être abandonnés. Ainsi, à chaque étape lorsque $g(v) > 0$, on abandonne le canal qui présente le plus mauvais CgNR. Dès que $g(v) < 0$ on a atteint le nombre $\tilde{N}$ de canaux (et surtout l'identification de ses sous canaux) pour lesquels la contrainte de puissance est respectée.

Déroulement

L'algorithme commence par classer les canaux en fonction de leur CgNR. Appelons N_{sc} le nombre de sous canaux utilisables à chaque étape de l'algorithme. N_{sc} est initialisé à N . L'estimation de v à chaque étape est : $v_{est} = 1/\varphi_{N_{sc}}$. On calcule

$$g(v_{est}) = N_{sc} v_{est} - \sum_{n=1}^{N_{sc}} \frac{1}{\varphi_n} - P_{T,max} \text{ on a donc } g(v_{est}) = \frac{N_{sc}}{\varphi_{N_{sc}}} - \sum_{n=1}^{N_{sc}} \frac{1}{\varphi_n} - P_{T,max} .$$

Si $g(v_{est})$ est positif, alors la puissance qui est allouée à ce stade dépasse la contrainte, un sous canal est abandonné et $N_{sc} = N_{sc} - 1$. Lorsque $g(v_{est})$ est négatif, l'allocation est réalisable. On pose $\tilde{N} = N_{sc}$ et par conséquent v et p sont complètement déterminés. Dans l'illustration 2, les canaux 6 et 7 ont été abandonnés.

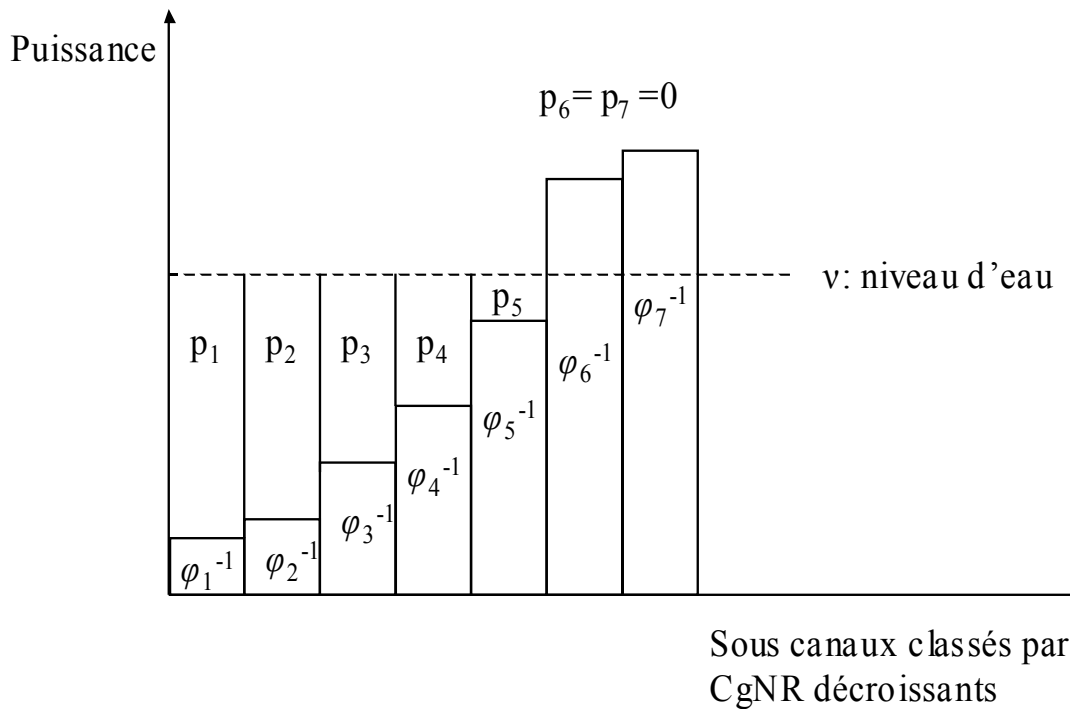


Illustration 2: Waterfilling et niveau d'eau

La détermination du niveau peut être fastidieuse surtout si elle répétée plusieurs fois (c'est le cas en présence de plusieurs utilisateurs) et que le nombre de sous porteuses est élevé.

Annexe 2.A Puissance d'émission et SNR gap

La puissance nécessaire pour transmettre un débit r sur la sous porteuse n est :

$$P(r) = \frac{\Gamma(2^r - 1)}{\varphi} \quad (1)$$

Dans cette formule, φ est le CgNR (*Channel gain to Noise Ratio*) : $\varphi = \frac{|g|^2}{N_0 \Delta f}$ (2), ([YinLiu_00],[Wong&_04]).

Le paramètre Γ (connu sous le nom de SNR-gap où SNR signifie Signal Noise To Ratio) est la marge de puissance qu'il faut concéder pour assurer un SER (*Symbol Error Rate*) donné.

Expression de Γ

Exprimons la probabilité d'erreur symbole d'une M-QAM à $|g|$ fixé (où $|g|$ suit une loi de Rayleigh) ([Lamy_00], p.126) :

$$SER \approx 4Q\left(\sqrt{\frac{3\log_2(M)|g|^2 E_b}{(M-1)N_0}}\right) \quad (3) \text{ avec } Q(x) = \frac{1}{\sqrt{2\pi}} \int_x^\infty e^{-\frac{t^2}{2}} dt$$

Cette formule de la probabilité d'erreur d'une M-QAM s'établit en considérant une M-QAM comme deux $\sqrt{M}$ -PAM (*Pulse Amplitude Modulation*) indépendantes en quadrature ([Benedetto&_99], p231).

On a $E_b = P(r) T_b$ où la puissance d'émission d'un symbole QAM de $r = \log_2(M)$ bits et $T_b = \frac{1}{\log_2(M) \Delta f}$. Dans (3) nous remplaçons donc E_b par $\frac{P(r)}{\log_2(M) \Delta f}$ et M

par 2^r . D'où $SER \approx 4Q\left(\frac{3P(r)|g|^2}{(2^r-1)N_0 \Delta f}\right)$. D'après (2) $SER \approx 4Q\left(\frac{3P(r)\varphi}{(2^r-1)}\right)$. On en

déduit que $P(r) \approx \frac{2^r-1}{\varphi} \frac{Q^{-1}\left(\frac{SER}{4}\right)}{3}$ (4).

Par identification entre (1) et (4), on obtient que $\Gamma = \frac{Q^{-1}\left(\frac{SER}{4}\right)}{3}$

([Pfletschinger&_02]).

Annexe 2.B Relation entre le débit et le SNR

Le calcul du débit d'un utilisateur s'obtient par une fonction $r = h(\gamma)$ qui lie le débit r avec le SNR γ . Le tableau suivant résume quelques possibilités pour la fonction $h(\cdot)$. Dans ce tableau, Γ désigne le *SNR gap* (il correspond à une marge de puissance, cf. annexe 2.A). La fonction $h(\cdot)$ communément adoptée est la formule de *Shannon*. Cette dernière peut être bornée par R_{max} le débit du meilleur MCS (*Modulation and Coding Scheme* ou schémas de modulation et de codage) que l'on peut utiliser sur une sous-porteuse. Il suffit de prendre $R_{max} = \infty$ pour supprimer cette borne.

Fonction puissance-débit : $r = h(\gamma)$
$h(\gamma) = \min (R_{max}, \log_2(1+\gamma/\Gamma))$
$h(\gamma) = \min(R_{max}, f^{-1}(\gamma))$ où $\gamma = f(r)$

Un exemple de la fonction est donné dans [Kivanc&_03], $f(r) = 0.6 r^3$; la fonction est illustrée ci-dessous, les auteurs tentent de prendre en compte les MCS utilisés par le système considéré.

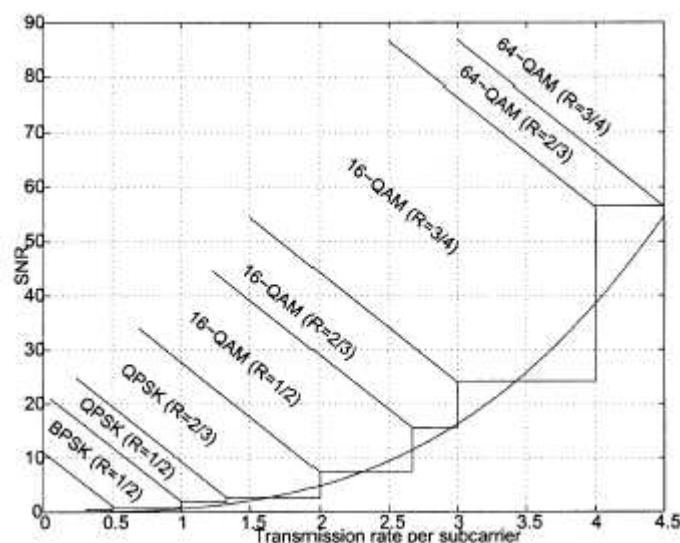


Illustration 3: Exemple de fonction débit-puissance

Annexe 2.C Optimisation MA, RA et astuces

1) Programmation entière

Pour appliquer la programmation entière, il faut reformuler les problèmes MA et RA sous une forme linéaire ([Inyoung&_01] et [MaoWang_06]).

Prenons l'exemple du problème MA

On considère la fonction utilisée telle que $\gamma = f(r)$. Si l'utilisateur u reçoit la sous porteuse n , on définit $\rho_{u,n} = 1$ sinon $\rho_{u,n} = 0$.

$$\min_{r_{u,n}, \rho_{u,n}} \sum_{n=1}^N \sum_{u=1}^U \rho_{u,n} \frac{f(r_{u,n})}{|g_{u,n}|^2}$$

$$\forall u, r_u \geq \sum_{n=1}^N r_{u,n}$$

Ce problème est non linéaire car la fonction f est non linéaire (on a vu des exemples dans les annexes 2.A et 2.B).

Pour linéariser le problème, on considère que l'ensemble des débits possibles sur une sous porteuse est un ensemble discrétisé $r_{u,n} \in \{c_1, c_2, \dots, c_M\}$. On a alors $f(r_{u,n}) \in \{f(c_1), f(c_2), \dots, f(c_M)\}$ et ces valeurs peuvent être précalculées. Un nouvel indicateur est alors défini : $\delta_{u,n,m} = 1$ si $r_{u,n} = c_m$ et $\delta_{u,n,m} = 0$ sinon. A la suite de cette nouvelle définition, on a $\delta_{u,n,m} \rho_{u,n} = \delta_{u,n,m}$

Le problème se reformule alors comme suit :

$$\min_{\delta_{u,n,m}} \sum_{n=1}^N \sum_{u=1}^U \sum_{m=1}^M \delta_{u,n,m} \frac{f(c_m)}{|g_{u,n}|^2}$$

$$\forall u, r_u \geq \sum_{n=1}^N \sum_{m=1}^M \delta_{u,n,m} c_m \quad \text{et} \quad \forall n, \sum_{u=1}^U \sum_{m=1}^M \delta_{u,n,m} \in \{0, 1\}$$

Ce problème linéaire peut désormais être résolu par programmation entière. La reformulation se fait de façon similaire pour le problème RA.

2) Relaxation de la contrainte sur le partage des sous porteuses

Dans la section précédente, on a défini $\rho_{u,n}$ tel que $\rho_{u,n} = 1$ si l'utilisateur u reçoit la sous porteuse n et $\rho_{u,n} = 0$ sinon. Certains auteurs considèrent $\rho \in [0, 1]$ et interprètent cela comme un partage de la sous porteuse dans le temps. Si la fonction f est convexe cela contribue à rendre le problème ci dessus (avant linéarisation) convexe. Le problème peut alors être résolu grâce à la méthode des multiplicateurs de Lagrange. Dans [Wong&_99], cette technique est utilisée. L'algorithme proposé est critiqué pour sa complexité et la lenteur de sa convergence ([Pfetschinger&_02], [PietrykJan_02]).

Annexe 2.D L'algorithme Hongrois

Utilisation de l'algorithme pour l'affectation de sous porteuses

L'affectation de sous porteuses est résolue de façon optimale par l'algorithme Hongrois par rapport à un ensemble $\{N_u\}_{1 \leq u \leq U}$ fixé (Approche A, option 2). 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 $c(u,n)$ contient le coût associé à l'affectation de la sous porteuse n à l'utilisateur u . La matrice de coût doit être carrée. Généralement, on dispose de N sous porteuses et U utilisateurs avec $N > U$. Pour appliquer l'algorithme Hongrois, il suffit dupliquer chaque utilisateur en N_u utilisateurs fictifs où N_u est le nombre de sous porteuses qu'un utilisateur va recevoir. Chaque utilisateur fictif a les mêmes coûts que l'utilisateur initial. Les utilisateurs fictifs reçoivent une unique sous porteuse à la fin du processus d'affectation. Pour obtenir la solution du problème initial, on regroupe les sous porteuses des utilisateurs fictifs issus d'un même utilisateur initial. L'ensemble $\{N_u\}_{1 \leq u \leq U}$ est déterminant pour la matrice de coût et la répartition de sous porteuse résultante.

Choix de la fonction de coût de l'algorithme Hongrois

L'optimalité de l'algorithme Hongrois par rapport au problème initial dépend des nombres N_u et de l'expression de la fonction de coût. L'algorithme Hongrois donne la solution optimale par rapport à la matrice de coût, mais cette dernière est plus ou moins pertinente au regard du problème initial. Comme le coût total doit être minimisé, les possibilités de fonctions de coût sont les suivantes :

- $c(u,n) = -\varphi_{u,n}$ (1)

- $c(u,n) = -10 \log_{10}(\varphi_{u,n})$ ([YinLiu_00]) (2)

Les fonctions de coût (1) et (2) sont optimales lorsqu'on résout respectivement les problèmes $\sum_{u,n} \varphi_{u,n}$ et $\max_{u,n} \sum_{u,n} 10 \log_{10}(\varphi_{u,n})$. Elles peuvent être utilisées aussi bien en MA qu'en RA.

- $c(u,n) = -\log_2(1 + P_{u,n} \varphi_{u,n})$ (3)

- $c(u,n) = P_{u,n}$ ([PietrykJan_02]) (4)

La fonction de coût (3) peut être utilisée en RA quand on cherche à maximiser le débit global ($\sum_{u,n} \log_2(1 + P_{u,n} \varphi_{u,n})$). La fonction de coût (4) peut être utilisée en MA quand on cherche à minimiser la puissance transmise. Un problème se pose lorsqu'on considère les fonctions de coût (3) et (4). Pour exprimer l'une ou l'autre de ces fonctions de coût, il faut connaître l'allocation de puissance. Donc pour effectuer la tâche 2 de façon optimale, il faut connaître l'allocation de puissance. Or dans l'approche A, l'allocation de puissance (ou le choix des MCS) n'est possible qu'après de la tâche 2. Le problème «se mord la queue». Pratiquement, pour exprimer la fonction de coût (3) (respectivement (4)) et réaliser la tâche 2, on fait l'hypothèse d'allocation uniforme de puissance sur les sous porteuses (respectivement MCS fixe).

En conclusion, si l'ensemble $\{N_u\}_{1 \leq u \leq U}$ ne se résume pas à une évaluation et que la

stratégie d'allocation de puissance respecte l'hypothèse utilisée pour exprimer la fonction de coût, l'affectation de sous porteuses proposée par l'algorithme Hongrois est optimale au regard du problème initial.

Quelques éléments théoriques

Le problème qui se pose est l'affectation de chaque sous porteuse à un utilisateur unique. La liaison entre un utilisateur et une sous porteuse est associée à un coût. Il s'agit de minimiser un coût global lors de l'affectation. Lorsqu'une formulation équivalente du problème comporte un nombre N d'utilisateurs égal à celui des sous porteuses, le problème revient à déterminer un couplage parfait et de poids minimal dans un graphe biparti ([Gondran&_79]). Un graphe biparti est un graphe dont l'ensemble de sommets peut être divisé en deux ensembles distincts A et B tels que les arêtes du graphe relient un noeud de A à un noeud de B . Un couplage du graphe est un sous ensemble d'arêtes qui n'ont aucun sommet en commun. Le couplage est dit parfait lorsque tous les sommets du graphe initial sont impliqués dans le couplage, il n'est possible que lorsque $|A| = |B|$. Autrement dit chaque élément de A est relié à un élément de B de manière unique. Le poids du couplage est égal à la somme du poids des arêtes.

Dans une autre formulation du problème, on recherche une permutation σ qui minimise $\sum c_{i\sigma(i)}$ parmi les $N!$ permutations possibles.

La solution optimale à ce problème est fournie par l'algorithme Hongrois créé par H. W. Kuhn en 1955 en hommage aux mathématiciens J. Egervary et D. König. C'est l'algorithme le plus efficace ($O(N^3)$) et de ce fait l'algorithme de référence pour la résolution optimale du problème formulé ([PapaSteig_82]).

Le principe de l'algorithme est le suivant : la (les) solution(s) optimale(s) sont inchangées en cas d'augmentation ou de diminution d'un même facteur sur tous les éléments d'une même ligne ou d'une même colonne de la matrice de coût. L'objectif de ces opérations est l'apparition de N zéros : un seul zéro par ligne et par colonne qui indiquent l'affectation optimale.

Déroulement de l'algorithme :

Considérons une matrice de coûts.

- 1) On soustrait le coût minimal de chaque ligne à tous les autres éléments de la ligne correspondante.
- 2) Dans les colonnes où il n'y a pas de zéros, on soustrait le coût minimal de la colonne à tous les autres éléments de la colonne.
- 3) On procède au marquage des zéros i.e. on tire un trait sur les lignes et les colonnes où il y a des zéros. Si le nombre minimal de traits nécessaires pour que chaque zéro soit marqué est égal à la dimension de la matrice on s'arrête. Sinon on va à l'étape 4.
- 4) Soustraire le coût du plus faible élément non marqué à tous les éléments non marqués et le rajouter aux éléments marqués deux fois (i.e. intersections d'une ligne rayée d'un trait et d'une colonne rayée d'un trait) puis on retourne à l'étape 3. A l'arrêt, les zéros indiquent l'affectation optimale.

Pour une illustration de l'application de l'algorithme, on peut se référer à [Issaiyakul_99].

Annexe 4.A Règles de permutation en FUSC

Nous considérons $N_{FFT} = 2048$ (p566 de [IEEE 802.16]), le raisonnement est ensuite valable pour $N_{FFT} = 1024, 512$ et 128 en appliquant les paramètres des tables 311a, 311b et 311c définies p379-381 de [IEEE 802.16e].

Sur 1703 fréquences utilisées, 1536 servent pour les données. Les fréquences pilotes sont déterminées avant de construire les sous canaux. On va construire $N_{CHAN} = 32$ sous canaux de $N_{scPerCHAN} = 48$ sous porteuses. Pour cela, on forme 48 groupes de 32 sous porteuses consécutives. Un sous canal est constitué d'une fréquence de chaque groupe.

La formule¹ exacte qui permet d'aboutir à ce résultat est la suivante :

$$\text{subcarrier}(k, s) = N_{CHAN} \cdot n(k, s) + SS(k, DL_PermBase)$$

Cette formule donne l'index absolu (entre 0 et $N_{data}-1$) de la fréquence d'index local k du sous canal d'indice s (s varie entre 0 et $N_{CHAN}-1$ et k varie entre 0 et $N_{scPerCHAN}-1$).

Principe général de la formule :

- Le produit $N_{CHAN} \cdot n(k, s)$ est un multiple de N_{CHAN} et pointe donc sur la première fréquence d'un groupe. Ce produit permet de choisir un groupe.
- $SS(k, DL_PermBase)$ est compris entre 0 et $N_{CHAN}-1$. A partir de la première fréquence d'un groupe (désignée par $N_{CHAN} \cdot n(k, s)$), on peut obtenir toutes les fréquences du groupe. On appelle cette fonction SS pour *Subcarrier Selection*.

Détails de chaque terme :

- $N_{CHAN} = 32$
- $n(k, s) = (k + 13s) \bmod N_{scPerCHAN}$. Ainsi $n(k, s)$ vaut entre 0 et 47. La dépendance par rapport à l'indice s assure que la fréquence d'index k de deux sous canaux distincts ne sera pas choisie dans le même groupe.
- $SS(k, DL_PermBase) = (p_s[n(k, s) \bmod N_{CHAN}] + DL_PermBase) \bmod N_{CHAN}$
 - Le tableau 311 ([IEEE 802.16]) donne la liste de permutation (*permutation base*) p_0 : 3, 18, 2, 8, 16, 10, 11, 15, 26, 22, 6, 9, 27, 20, 25, 1, 29, 7, 21, 5, 28, 31, 23, 17, 4, 24, 0, 13, 12, 19, 14, 3, 0. Si on construit le sous canal 0, on utilise cette liste. Par contre, si on construit le sous canal s , on utilise p_s qui est obtenue en décalant s fois p_0 vers la gauche. Si $s=10$, on obtient p_{10} = 6, 9, 27, 20, 25, 1, 29, 7, 21, 5, 28, 31, 23, 17, 4, 24, 0, 13, 12, 19, 14, 3, 0, 3, 18, 2, 8, 16, 10, 11, 15, 26, 22.
 - Quelle que soit la valeur de s , $p_s(x)$ est compris entre 0 et 31 ($N_{CHAN}-1$) et identifie une fréquence à l'intérieur d'un groupe. Le nombre x doit être compris entre 1 et 32 (N_{CHAN}). Le nombre $n(k, s)$ est compris entre 0 et 47

¹ Formule 111 p 535 dans [IEEE 802.16e]

$(N_{scPerCHAN} - 1)$, on lui applique un modulo N_{CHAN} .

- $DL_PermBase$ intervient dans le calcul pour créer des sous canaux différents. La fréquence d'index k du sous canal s avec $DL_PermBase$ i et la fréquence d'index k du sous canal s avec $DL_PermBase$ $i+1$ sont consécutives .
- Enfin, on applique un modulo N_{CHAN} à $(p_s(x) + DL_PermBase)$ pour pouvoir obtenir un index de fréquence à l'intérieur d'un groupe.

Exemple de sous canal :

Ci dessous, on donne les index absolus des fréquences contenues dans le sous canal d'index 1 pour $DL_PermBase = 1$

Pour k variant de 0 à 47, $subcarrier(k,1) = [442, 450, 510, 520, 566, 582, 637, 640, 696, 722, 741, 793, 801, 846, 877, 916, 943, 991, 996, 1043, 1059, 1097, 1137, 163, 1196, 1232, 1275, 1303, 1319, 354, 1404, 1429, 1466, 1474, 1534, 19, 35, 73, 113, 139, 172, 208, 251, 279, 295, 330, 380, 405]$.

Annexe 4.B Règles de permutation en PUSC et collisions

En PUSC, sur un temps symbole, la bande est subdivisée en $N_{clusters}$ clusters de 14 fréquences consécutives. Ces *clusters* ont un numéro physique i_{PHY} . Après une renumérotation, il comporteront des numéros logiques. Des groupes majeurs de *clusters* de numéros logiques consécutifs sont alors formés. Ces groupes majeurs ont des fréquences issues de différentes parties de la bande grâce à la renumérotation.

La formule de renumérotation (ou *outer permutation*) est la suivante :

– $RenumberingSequence(i_{PHY})$ si la zone vaut 1 ou l'indicateur $Use\ All\ SC = 0$

– $\{RenumberingSequence(i_{PHY}) + 13 * DL_PermBase\} \bmod N_{clusters}$

où $RenumberingSequence$ est une liste de $N_{clusters}$ entiers entre 0 et $N_{clusters}-1$, donnée par les tables 310, 310a, 310b et 310c selon la taille de la FFT (p528, [IEEE 802.16e]) et l'indicateur $Use\ All\ SC$ est un indicateur de broadcast¹.

Les groupes majeurs sont formés suivant le tableau 4.3. Les sous porteuses sont affectées aux sous canaux suivant la même formule que le FUSC (avec les paramètres des tables 310, 310a, 310b et 310c selon la taille de la FFT), c'est l'*inner permutation*.

Exemple de construction d'un sous canal

Pour $N_{FFT} = 1024$, le nombre de sous porteuses utilisées est 841. On a $N_{clusters} = 60$ clusters de 14 fréquences consécutives.

La $RenumberingSequence$ vaut [6, 48, 37, 21, 31, 40, 42, 56, 32, 47, 30, 33, 54, 18, 10, 15, 50, 51, 58, 46, 23, 45, 16, 57, 39, 35, 7, 55, 25, 59, 53, 11, 22, 38, 28, 19, 17, 3, 27, 12, 29, 26, 5, 41, 49, 44, 9, 8, 1, 13, 36, 14, 43, 2, 20, 24, 52, 4, 34, 0].

Exemple : dans la zone 1, un *cluster* d'index physique 3 a un index logique de 37.

Une fois les *clusters* renumérotés, on construit les groupes majeurs. D'après le tableau 4.3, le groupe majeur 0 est constitué des 12 *clusters* dont les numéros logiques sont 0-11

Exemple : dans la zone 1, les numéros physiques des *clusters* du groupe majeur 0 sont [60, 49, 54, 38, 58, 43, 1, 27, 48, 47, 15, 32].

Ensuite, on place les pilotes dans chaque *cluster* en fonction du temps symbole. Il reste dans le groupe majeur 0 qui est pair, $12 * (14 - 2)$ fréquences de données à répartir en $N_{CHAN} = 6$ sous canaux. On construit 24 groupes de 6 sous porteuses et on applique la même formule que le PUSC avec $p_s = [3, 2, 0, 4, 5, 1]$.

Exemple de sous canal obtenu : on donne les index absolus des fréquences contenues dans le sous canal d'index 1 pour $DL_PermBase = 1$ et zone = 1. Pour k variant de 0 à 23, $subcarrier(k, 1) = [597, 790, 791, 450, 460, 661, 667, 426, 427, 464, 474, 87, 93, 678, 679, 520, 530, 297, 303, 440, 441, 562, 572, 591]$.

¹ Voir p.437 dans [IEEE 802.16e]

Densité de probabilité des collisions en PUSC

On s'est placé dans le cas où on utilise tous les groupes majeurs. Pour $N_{FFT}=1024$ et $N_{data}=720$, on peut construire en tout $N_{CHAN}=30$ sous canaux (de 24 fréquences par temps symbole). Pour les facteurs de charge 1/30 et 2/30, on trace les densités de probabilité du nombre de collisions ci-après. Le comportement est différent selon la zone où on se situe.

La première conclusion est la suivante : en zone 1, chaque secteur (ou cellule) doit utiliser une partie des groupes majeurs et éviter ainsi les collisions entre secteurs (ou cellules).

En zone 1, l'utilisation de tous les groupes majeurs dans chaque secteur (ou cellule) cause une probabilité de collision très élevée. En effet, en zone 1, la renumérotation des *clusters* est identique dans chaque cellule, les groupes majeurs sont donc identiques. Si on utilise tous les groupes majeurs dans chaque secteur (ou cellule), les sous canaux subissent plus de collisions que lorsque la zone est différente de 1.

En zone 1, il faut faire de la sectorisation et affecter des groupes majeurs distincts à chaque secteur (ou cellule). Les groupes majeurs sont orthogonaux et le nombre de collisions entre les secteurs concernés (ou cellules) est alors nul.

La deuxième conclusion est la suivante : dans une zone différente de 1, la sous canalisation PUSC a la même densité de probabilité de collision que la sous canalisation aléatoire (choix aléatoire de 24 fréquences parmi 720).

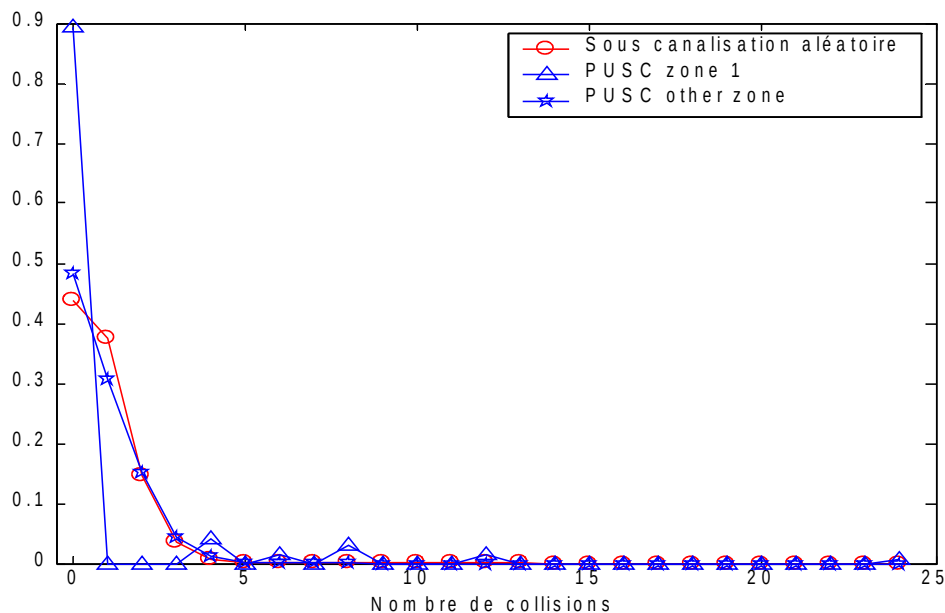


Illustration 4: Densité de probabilité du nombre de collisions en PUSC, facteur de charge 1/30

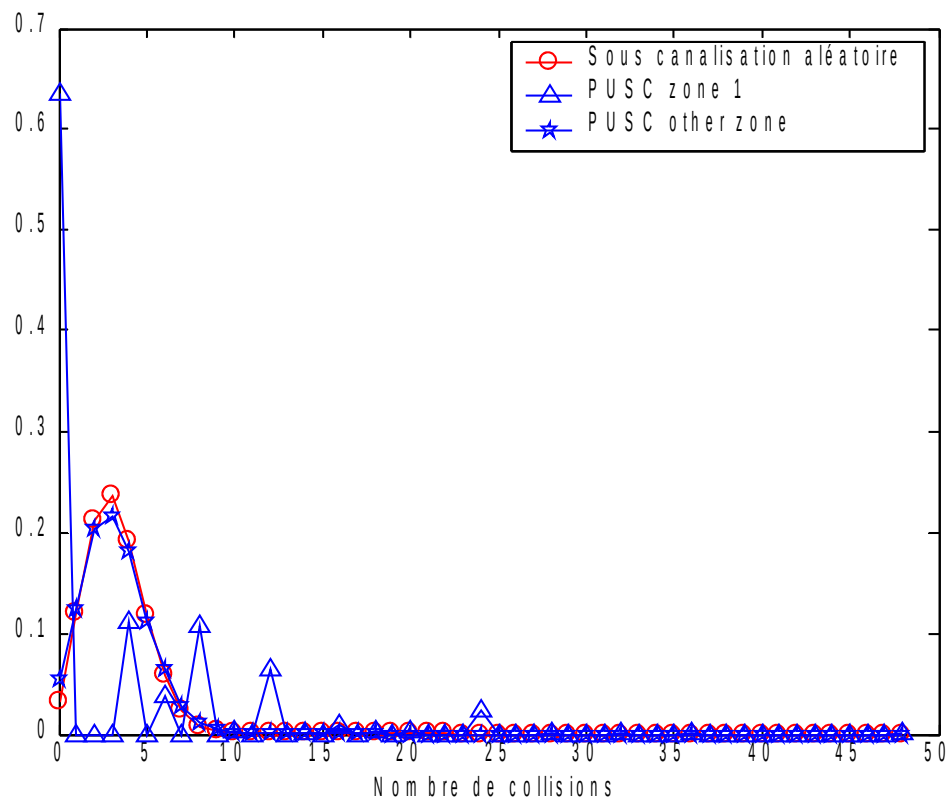


Illustration 5: Densité de probabilité du nombre de collisions en PUSC, facteur de charge 2/30

Annexe 5.A Détails sur deux contributions (contexte multicellulaire)

Détails sur le fonctionnement des trois phases dans [Kim&_04]

Les sous canaux (constitués de sous porteuses distribuées sur la bande) sont identiques dans toutes les cellules. L'algorithme détermine combien de fois un sous canal est utilisé ainsi que les utilisateurs correspondants.

Phase 1 : Préférences intra-cellulaires

Lors de la première phase, les utilisateurs soumettent leur choix $1/k_{best}$ de FRF à leur station de base ainsi que le débit associé $R_{k_{best}}$: $k_{best} = \arg \max_k \left(\frac{R_k}{k} \right)$, où R_k est le débit que l'utilisateur réalise avec un sous canal de FRF $1/k$. Chaque BS peut alors calculer le nombre de sous canaux dont elle a besoin pour chaque FRF :

$$N_k^b = \sum_{u=1}^U \frac{R_{min}}{R_{k_{best}}^u}.$$

Phase 2 : Harmonisation inter-cellulaire

La deuxième phase vise à déterminer le nombre définitif de sous canaux utilisés pour chaque FRF. Ces nombres seront communs à toutes les BS. Les auteurs procèdent du FRF le plus faible au plus élevé pour calculer le nombre de sous canaux nécessaires (par exemple si k peut valoir 1, 3 et 7 alors l'ordre de traitement est 1/7, 1/3, 1).

Pour le FRF $1/k$, le nombre de sous canaux nécessaires est $N_k^{(final)} = \max_b N_k^b$ où N_k^b est la demande en sous canaux de FRF $1/k$ de la BS b .

Avant de procéder au traitement du FRF suivant $1/k'$ (où $k' < k$), les demandes des cellules b telles que $N_k^{(final)} \geq N_k^b$, en sous canaux de FRF $1/k'$, sont mises à jour. Voyons pourquoi. Les cellules telles que $N_k^{(final)} \geq N_k^b$, ont demandé N_k^b sous canaux de FRF $1/k$ et en reçoivent $N_k^{(final)}$. Ces cellules ont donc $S_k^b = N_k^{(final)} - N_k^b$ canaux supplémentaires de FRF $1/k$. Les sous canaux de FRF $1/k$ sont de meilleure qualité que ceux de FRF $1/k'$ (où $k' < k$) : à gain égal, ils subissent moins d'interférences. Les S_k^b canaux supplémentaires de FRF $1/k$ peuvent être utilisés pour satisfaire les demandes en FRF $1/k'$. Finalement, pour les cellules b telles que $N_k^{(final)} \geq N_k^b$, la demande en canaux en FRF $1/k'$ est mise à jour : $new N_{k'}^b = \max(0, N_k^b - S_k^b)$.

Exemple numérique : soit deux FRF 1/7 et 1/3, on a $N_7^1 = 1$ et $N_3^1 = 6$. Si $N_7^{(final)} = 3$ alors la BS 1 a $S_7^1 = N_7^{(final)} - N_7^1$ sous canaux supplémentaires pour le FRF 1/7. Elle peut les attribuer à des utilisateurs qui ont demandé des sous canaux de FRF 1/3. Alors $new N_3^1 = N_3^1 - S_7^1 = 6 - 2 = 4$, la demande de la cellule 1 en FRF 1/3 est désormais de 4.

Phase 3 : Allocation intra-cellulaire

Après la phase d'harmonisation, l'allocation se termine de façon indépendante dans chaque BS. Chaque utilisateur reçoit des sous canaux selon le nombre et le FRF demandé. S'il reste des sous canaux disponibles, ils sont alloués de façon opportuniste (i.e. l'utilisateur avec les meilleures conditions de propagation reçoit le sous canal). La

phase d'allocation intra-BS est peu décrite et aucune stratégie n'est proposée en cas de pénurie de sous canaux (lorsque la demande est supérieure au nombre total de sous canaux).

Détails sur l'algorithme proposé dans [Kwon&_05]

[Kwon&_05] construit des pseudo cellules de k secteurs interférents et s'intéresse à l'allocation d'un débit minimal spécifique à chaque utilisateur d'une pseudo cellule. Deux facteurs de réutilisation fréquentielle sont considérés : 1 et k . La totalité de la bande est divisée en sous bandes de fréquences adjacentes. Lorsque le FRF vaut 1 la sous bande est utilisée en entier dans chaque secteur ; lorsque le FRF vaut $1/k$, la sous bande est divisée en k parties réservées de façon exclusive à chaque secteur. Les auteurs minimisent l'*outage* maximal sur l'ensemble des secteurs. Deux phases sont considérées : (i) le calcul du nombre de sous bandes associées aux deux FRF et identification de ces sous bandes, (ii) l'allocation des ressources à l'intérieur de chaque secteur. Pendant la première phase, chaque secteur construit l'ensemble des utilisateurs (V_1 et V_k) préférant chaque FRF : pour chacun le débit moyen sur l'ensemble des S sous bandes avec les FRF 1 et k sont calculés, $\bar{r}_u(FRF) = \frac{1}{S} \sum_1^S r_{u,s}(FRF)$. Pour classer

les utilisateurs, le rapport $\frac{\bar{r}_u(1)}{\bar{r}_u(k)}$ est comparé à un seuil. Grâce à ce classement,

chaque secteur calcule le nombre moyen de sous porteuses nécessaires pour chaque type de FRF ρ_1^{secteur} et ρ_k^{secteur} . Après une phase d'harmonisation (transfert entre les ensembles V_1 et V_k dans les secteurs concernés) pour que chaque secteur ait les mêmes besoins moyens, les nombres N_1 et N_k (nombre de sous bandes nécessaires pour chaque FRF) sont calculés :

$$N_1 = S \frac{1/k \sum_{\text{secteur}=1}^k \rho_1^{\text{secteur}}}{1/k \sum_{\text{secteur}=1}^k \rho_1^{\text{secteur}} + \sum_{\text{secteur}=1}^k \rho_k^{\text{secteur}}}$$

Pendant la seconde phase, l'allocation est indépendante dans chaque secteur. L'utilisateur le plus loin de son objectif en terme de débit minimal reçoit sa meilleure sous bande. Cela constitue une adaptation de [RheeCioffi_00] quand il y a des débits minimaux ; en effet [RheeCioffi_00], qui ne considère pas de contraintes sur le débit minimal, alloue à l'utilisateur qui a le plus faible débit sa meilleure sous porteuse. Enfin dans [Kwon&_05], lorsque l'*outage* dépasse un seuil, l'utilisateur le plus mal loti est abandonné et ses ressources redistribuées.

Annexe 5.B Précisions sur l'OSA-IL (contexte multicellulaire)

Dans cette annexe, nous présentons le *flow chart* de la phase de limitation d'interférences (cf. illustration 4). Pour en faciliter la lecture nous l'accompagnons du texte ci dessous.

Rappelons quelques notations : B est le nombre de cellules (avec b l'indice sur les cellules), N est le nombre de sous porteuses (avec n l'indice sur les sous porteuses). Onnote Θ l'ensemble des cellules qui ont tous leurs utilisateurs satisfaits.

Cas le plus simple : toutes les cellules sont satisfaites (cf. bloc A, illustration 4). Dans ce cas la sous porteuse est attribuée dans chaque cellule à l'utilisateur qui a le meilleur SINR.

Cas contraire : il reste des cellules à satisfaire. Dans ce cas la sous porteuse est affectée avec le FRF courant.

- Changement de FRF (cf. transitions B1-B3, illustration 4)

Le changement de FRF est réalisé lorsqu'un nombre C_{max} sous porteuses voient le SINR moyen de leur ensemble co-canal Π_n inférieur au seuil δ et que le FRF n'a pas atteint la valeur minimale. Le vecteur fif contient les valeurs de FRF, i est l'index sur ce vecteur et i_{max} la valeur maximum de cet index. Le FRF courant détermine le nombre x de cellules qui ont le droit d'utiliser la sous porteuse.

- Sélection des cellules et des utilisateurs (cf. bloc C-E, illustration 4)

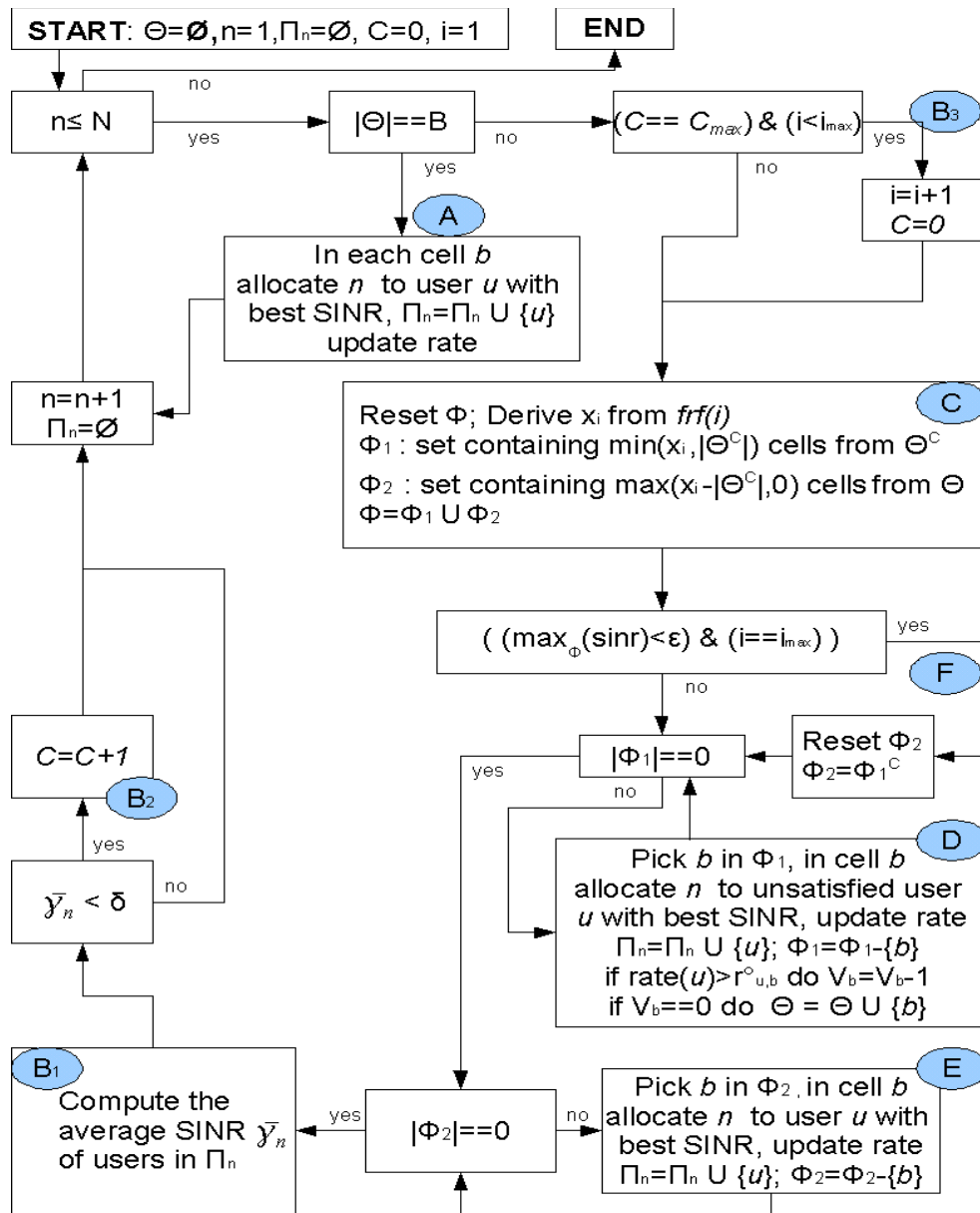
Les cellules de l'ensemble $|\Theta|^c$ ont la priorité pour recevoir la sous porteuse. On choisit un nombre de cellules non satisfaites égal à $\min(x, |\Theta|^c)$; elles sont stockées dans un ensemble Φ_1 (cf. bloc C). On complète éventuellement par un nombre $\max(x - |\Theta|^c, 0)$ de cellules satisfaites ; elles sont stockées dans un ensemble Φ_2 (cf. bloc C).

Choix des cellules non satisfaites : dans chacune de ces cellules on calcule le CgINR du meilleur utilisateur insatisfait (c'est celui qui reçoit la sous porteuse si la cellule est choisie, cf. bloc D). Les $|\Phi_1|$ meilleurs CgINR déterminent les cellules qui vont utiliser la sous porteuse.

Choix des cellules satisfaites : dans chacune de ces cellules on calcule le CgINR du meilleur utilisateur (c'est celui qui reçoit la sous porteuse si la cellule est choisie, cf. bloc E). Les $|\Phi_2|$ meilleurs CgINR déterminent les cellules qui vont utiliser la sous porteuse.

- Décision exceptionnelle pour un FRF = 1 (cf. transition F, illustration 4)

Si malgré l'utilisation du plus faible FRF (du vecteur fif), le meilleur SINR estimé sur l'ensemble «co-canal» de la sous porteuse est inférieur à un seuil ε , la sous porteuse est exceptionnellement utilisée dans toutes les cellules (l'ensemble Φ_2 devient le complémentaire de l'ensemble Φ_1). Cela arrive quand les utilisateurs restants sont très durs à satisfaire : le FRF le plus faible n'améliore pas leur débit sur la sous porteuse. Dans un tel cas, on préfère permettre à toutes les cellules d'utiliser la sous porteuse pour optimiser le débit global.



Annexe 6.A L'ordonnancement dans les réseaux filaires

Chapitre 6, section 1.2.1

Round Robin ([Nagle_87])

Le RR sert successivement chaque flux en ignorant les flux inactifs (i.e. flux sans paquets en attente de transmission). Les paquets doivent être de même taille pour assurer une équité entre les flux. De plus, des besoins différents en terme de bande passante ne peuvent être gérés.

Weighted Fair Queuing ([ParekhGallager_93])

Le principe du WFQ est décrit en 1.2.1. Le retard du WFQ par rapport au GPS est borné que ce soit en terme de fin de service dans le temps virtuel (1) et de service reçu (2).

$$|F_{WFQ}(p_i^k) - F_{GPS}(p_i^k)| \leq \frac{L_{max}}{C_{server}} \quad (1)$$

$$|W_{i,GPS}(0,t) - W_{i,WFQ}(0,t)| \leq L_{max} \quad (2)$$

Dans la suite, on préférera l'appellation WFQ à celle du PGPS.

Priority-based Weighted Fair Queueing ([Wang&_02])

Le PWFQ introduit des priorités dans le WFQ pour découpler délai et bande passante. Les auteurs considèrent une fenêtre coulissante sur l'axe du temps virtuel. Dans une telle fenêtre, on transmet le paquet de plus haute priorité et ayant la plus faible étiquette de fin de service. Les auteurs proposent une heuristique pour calculer les priorités en fonction des délais requis et des poids en WFQ. La limitation de cet algorithme réside dans le calcul de la taille de la fenêtre. Ce dernier serait au pire en $O(U^2)$. Il dépend des priorités calculées et semble devoir être renouvelé lors d'un changement important des caractéristiques des utilisateurs.

(m,k) WFQ

Dans [Koubaa_04], l'auteur propose d'intégrer des contraintes temporelles (m,k) firm dans le WFQ afin de fournir simultanément des garanties en terme de débit et de délai. Le principe (m,k) firm consiste à respecter les échéances de m tâches (ou paquets) sur n'importe quel intervalle de k tâches (ou paquets) consécutives. Les paquets sont classés en paquets optionnels ou critiques ; les échéances temporelles de ces derniers doivent absolument être respectées. Sur k paquets, il y en a m critiques. L'algorithme (m,k) WFQ est le suivant : s'il existe des paquets critiques, on transmet le paquet critique qui a la plus petite étiquette de fin de service sinon on transmet le paquet optionnel qui a la plus petite étiquette de fin de service. Un paquet optionnel n'est transmis que si l'échéance n'est pas dépassée, dans le cas contraire il est supprimé afin de diminuer le délai des paquets critiques.

Cet algorithme se base sur l'assertion selon laquelle plusieurs applications temps réel peuvent tolérer le non respect de quelques échéances. Les applications temps-réel seraient principalement sensibles à la perte de messages consécutifs. Spécifier m et k ($k > m$) pour un flux permettrait de limiter l'impact de la perte des paquets sur la QoS.

Virtual Clock ([Zhang_90])

Le VC transmet les paquets par estampille temporelle croissante. La formule (6.3) est utilisée en remplaçant le temps virtuel par le temps réel. La formule (6.4) est inchangée. Les estampilles sont plus simples à calculer que pour le WFQ. Cependant, le VC ne partage pas équitablement la bande parmi les utilisateurs lorsque tous les flux ne sont pas actifs ([Zhang_95]).

Self Clock Fair Queuing ([Golestani_94])

Le SCFQ définit le temps virtuel comme étant l'étiquette de fin de service (*virtual finish tag*) du paquet en cours de transmission. La formule (6.3) est utilisée avec cette modification. La formule (6.4) est inchangée. Les écarts entre le SCFQ et le GPS sont plus grands que ceux entre le WFQ et le GPS à cause de la nouvelle définition du temps virtuel.

Worst Case Fair Queuing ([ZhangBennett_96])

Dans le W2FQ, l'étiquette de début de service (*virtual start tag*) intervient dans la sélection d'un paquet. Le W2FQ transmet le paquet qui a la plus faible étiquette de fin de service (*virtual finish tag*) parmi ceux qui auraient commencé leur service dans le GPS (c'est à dire ceux qui vérifient $S(p_i^k) < V(t)$).

Le W2FQ vérifie (3) et réalise une approximation plus fine du GPS que le WFQ :

$$|W_{i,GPS}(0,t) - W_{i,W2FQ}(0,t)| \leq (1 - \frac{r_i}{C_{server}}) L_{i,max} \quad (3)$$

En revanche la complexité du W2FQ est supérieure à celle du WFQ à cause de la vérification supplémentaire sur l'étiquette de début de service.

Start-Time Fair Queuing ([Goyal&_97])

Les auteurs de [Goyal&_97] montrent que le WFQ n'émule pas correctement le GPS quand la capacité C_{server} est variable. Dans le WFQ, le VC et le SCFQ, les paquets étaient transmis par étiquette de fin de service croissante (*virtual finish tag*). Dans le STFQ, les paquets sont transmis par étiquette de début de service croissante (*virtual start tag*) . Le temps virtuel correspond à l'étiquette de début de service du paquet en cours de transmission.

Annexe 6.B L'ordonnancement dans les réseaux sans fils (canal ON-OFF)

(Chapitre 6, section 1.2.2)

On distingue le CSDPS (*Channel State Dependent Packet Scheduling*), le CIF-Q (*Channel Independent Packet Fair Queuing*), l'IWFQ (*Idealized Wireless Fair Queuing*) et le WFS (*Wireless Fair Scheduling*), l'ORCA_SRT (*Optimal Radio Channel Allocation Single Rate Transmission*).

Channel State Dependent Packet Scheduling ([Bhagwat&_96])

Dans le CSDPS, un flux est marqué lorsqu'un acquittement n'est pas reçu. Lors du choix du paquet à transmettre, on transmet prioritairement les paquets des flux non marqués. Pour ordonnancer les paquets des flux non marqués, la politique est à définir ; par exemple si une politique de type *round robin* est utilisée, on parle de CSDPS-RR. Quand toutes les queues non marquées sont vides, alors un paquet d'un flux marqué peut être envoyé. Un flux reste marqué tant que l'estimation de son canal est mauvaise. L'estimation du canal est un élément clé pour l'efficacité de l'algorithme. Il n'y a pas de mécanisme garantissant la bande passante aux flux servis. L'excès de service reçu par un flux n'est pas borné.

Idealized Wireless Fair Queuing ([Lu&_99])

Le modèle de service sans erreurs est le WFQ. Une queue virtuelle et une queue réelle sont parallèlement remplies lors de l'arrivée des paquets. Le WFQ est appliqué. Si le flux désigné bénéficie d'un bon canal, le paquet est envoyé et retiré des deux queues. Dans le cas contraire, il n'est retiré que de la queue virtuelle¹.

Le modèle d'avance et de retard consiste à comparer les tailles des queues réelle et virtuelle d'un même flux. Un flux est en avance (resp. en retard) si sa queue réelle a moins (resp. plus) de paquets que sa queue virtuelle. Lorsqu'un flux a un mauvais canal, il cède son opportunité de transmission au flux ayant la plus petite étiquette de fin de service et un bon canal. Un flux qui n'a pas transmis depuis longtemps est sûr de transmettre dès qu'il a un bon canal. Il a en effet une étiquette de fin de service plus faible que celles des autres flux. Dans cette situation, le flux retardé capture toute la bande tant qu'il n'a pas rattrapé son retard. Le retard global des flux en retard est borné à la valeur $E_{min,TOT}$. Le retard d'un flux ne doit pas dépasser le seuil $E_{min,i}$ où

$$E_{min,i} = \frac{r_i E_{min,TOT}}{L \sum_{j \in F} r_j} \text{ avec } F \text{ l'ensemble des flux et } L \text{ la taille d'un paquet. Pour cela, à}$$

tout moment, un flux doit avoir au maximum B_i paquets tels que $F(p_i^k) \leq V(t)$. S'il y en a plus, le flux i garde les B_i plus faibles étiquettes de fin de service, et supprime les autres. Une queue de *slots logiques* (ou jetons) est gérée parallèlement à la queue de

¹ La queue virtuelle est le reflet de la discipline choisie pour le modèle sans erreurs (ici le WFQ), c'est à dire l'état de la queue s'il n'y avait jamais d'erreurs sur le canal.

paquets pour découpler la politique de gestion des paquets et les opportunités de transmission (cf. WFS en 2.1). En IWFQ, l'excès de service n'est pas borné, des bornes de délai strictes ne peuvent être assurées.

Channel-condition Independent packet Fair Queuing ([Eugene&_98])

L' algorithme, en abrégé le CIF-Q, utilise le STFQ (*Start Time Fair Queuing*) comme modèle de service sans erreurs. Le CIF-Q présente plusieurs propriétés intéressantes : bornes de délai et garanties de débit, équité à court et long terme, dégradation gracieuse de service. Rappelons que la dernière propriété stipule qu'un flux en avance ne restitue pas l'avance reçue d'une seule traite.

On considère pour chaque flux, une queue réelle et une queue virtuelle¹. Le temps virtuel (par flux) n'est calculé que pour les queues réelles. Il est mis à jour lorsqu'un paquet ne peut être transmis à cause de l'état du canal.

Le paquet sélectionné est celui qui a l'étiquette de début de service (*virtual start tag*) le plus faible, à moins que 1) le canal soit mauvais ou que 2) le flux soit en avance et son débit minimum soit dépassé. La dégradation gracieuse de service est assurée par un paramètre δ . Les flux en avance relâchent une opportunité de transmission avec la probabilité constante $(1-\delta)$. Les flux en retard reçoivent en priorité les opportunités de transmission cédées (pour les raisons 1 et 2) et cela proportionnellement à leur poids. Un flux en retard ne peut capturer tout le canal comme en IWFQ car il ne rattrape son retard que dans les cas 1 et 2. Si aucun flux en retard ne peut transmettre, l'excès de service est ré-attribué aux autres flux proportionnellement à leur poids.

Optimal Radio Channel Allocation Single Rate Transmission (ORCA_SRT)

Comme alternative aux algorithmes décrits précédemment et obtenus de manière «heuristique», [Issaiyakul_99] propose un algorithme d'ordonnancement équitable basé sur l'optimisation ; pour cela l'algorithme Hongrois est utilisé. Le problème d'ordonnancement est formulé comme un problème d'affectation. Un flux i est dupliqué en r_i flux fictifs. Un flux fictif reçoit exactement une opportunité de transmission sur une trame. Un flux réel recevra donc r_i opportunités de transmission sur une trame. A partir du résultat de l'algorithme Hongrois, le modèle de compensation est utilisé pour éviter les transmissions sur de mauvais canaux. Le modèle de compensation correspond à celui du WFS .

¹ La queue virtuelle est le reflet de la discipline choisie pour le modèle sans erreurs (ici le STFQ), c'est à dire l'état de la queue s'il n'y avait jamais d'erreurs sur le canal.

Annexe 6.C L'ordonnancement dans les réseaux sans fils (canal multidébit)

Chapitre 6, section 1.2.2

Optimal Radio Channel Allocation Multi Rate Transmission ([Issaiyakul_99])

L'algorithme ORCA_MRT est basé sur l'optimisation par l'algorithme Hongrois. La matrice de coût est modifiée en fonction de l'état du canal et des compteurs de retard. Les flux en retard sont favorisés proportionnellement à leur compteur de retard. Après la transmission, les compteurs de retard¹ sont mis à jour comme suit :

$$E(i) = -\left[\left(\max_k \frac{W_k(t_1, t)}{r_k}\right) - \frac{W_i(t_1, t)}{r_i}\right]$$

Channel State Independent Wireless Fair Queuing ([Lin&_00])

L'algorithme CS-WFQ utilise le STFQ (*start-time fair queuing*) comme modèle de service sans erreurs. Les paquets sont transmis par étiquette de début de service (6.3) croissante. En CS-WFQ, le poids r'_i utilisé dans l'étiquette de fin de service (6.4) diffère du poids de débit initial r_i . Ce poids r'_i évolue en fonction de l'état du canal :

$$r'_i = \frac{r_i}{\max(W_i^{th}, W_i(t))} \quad \text{où } W_i(t) = W_i(0, t) \text{ correspond au débit du flux } i \text{ et } W_i^{th} \text{ est le}$$

seuil en dessous duquel un flux qui a eu un mauvais canal n'est plus favorisé. Les auteurs opposent leur «compensation immédiate» à la «compensation différée» de la plupart des algorithmes. Cependant, à l'opposé des algorithmes à «compensation différée», on ne vérifie pas que les conditions canal d'un flux en retard se sont améliorées. Un flux «en retard» est donc susceptible de gâcher la ressource parce que la compensation a été trop rapide. Notons que les performances de cet algorithme dépendent fortement des seuils W_i^{th} .

Channel Adaptive Fair Queuing ([Wang&_04])

L'algorithme CAFQ utilise le STFQ (*start-time fair queuing*) comme modèle de service sans erreurs. Un flux est dit en avance s'il a reçu plus d'opportunités dans le service réel (sur canal radio) que dans le service sans erreurs. Le compteur E représente ici la différence entre le service réel et le service sans erreurs. Si le compteur $E(i)$ est positif, le flux est en avance sinon le flux est dit normal.

Un flux i normal est sélectionné en fonction d'un paramètre v_i^N . Dans le service réel, on sélectionne en priorité les flux normaux ; le choix se porte sur celui qui a le plus faible v_i^N . Les flux en avance sont caractérisés par un paramètre v_i^A . Si aucun flux normal n'est actif, le flux en avance qui a le plus faible v_i^A est sélectionné. Lorsqu'un flux transmet, son paramètre v_i^N (respectivement v_i^A) est mis à jour comme suit :

¹ Nous adoptons une convention selon laquelle les compteurs des flux en retard sont négatifs.

$v_i^N = v_i^N + \frac{L}{r_i f(m)^\eta}$; $f(m)^\eta$ augmente avec le mode de transmission m (il évolue de 0 à 1 quand le mode augmente de 0 à M).

Un flux virtuel, dit de compensation, reçoit un poids de débit et entre en compétition avec tous les flux actifs durant l'ordonnancement du modèle sans erreurs. Les flux constituant ce flux virtuel sont les flux normaux ($E(i) \leq 0$) qui ont le meilleur état de canal ($m_i = M$) pour effectuer une compensation efficace. Ils sont classés dans une queue, le flux de tête a le plus grand retard ($\max |E(i)|$). Lorsque ce flux virtuel est choisi par le modèle sans erreurs (STFQ), le flux i placé en tête de queue du flux virtuel de compensation transmet et on applique $E(i) = E(i) + L$. Sinon, le service réel choisi le plus faible v_i^N parmi les flux normaux ou si besoin le plus faible v_i^A parmi les flux en avance. Dans cet algorithme, les flux en avance peuvent connaître de longue période de «famine» car aucune dégradation gracieuse de service ne leur est accordée.

MultiRate Fair Queueing ([Wang&_05])

Dans le MR-FQ, chaque flux possède un temps virtuel v_i qui lui est propre. Le flux sélectionné présente le plus faible temps virtuel v_i . Le MR-FQ possède trois modules : un module de *sélection de débit*, un module de *dégradation de service* (pour les flux en avance) et un module de *compensation*. Les flux sont classés en deux catégories (notées Ψ_{RT} et Ψ_{NRT}) : temps réel (Ψ_{RT} pour *real time*) et non temps réel (Ψ_{NRT} pour *non real time*). Ces deux classes interviennent dans le module de compensation et de dégradation de service.

Le flux i sélectionné initialement peut transmettre si :

- le module de sélection de débit le permet (le débit c retourné est positif)
- le module de dégradation de service ne l'oblige pas à abandonner son service.

Lorsque le flux i sélectionné initialement transmet, son temps virtuel v_i est mis à jour :

$$v_i = v_i + \frac{L}{r_i} \times \frac{c_{max}}{c} \quad (1) \text{ où } L \text{ est la longueur du paquet, } c_{max} \text{ est le débit maximum de}$$

transmission sur le canal, c est la sortie du module de sélection de débit et enfin r_i est le poids de débit du flux. Chaque flux a un compteur E . Les flux en retard (respectivement en avance) ont un compteur positif (respectivement négatif). Un flux dont le compteur est nul est dit satisfait.

Dans le module de sélection de débit, le mode de transmission qu'un flux peut utiliser dépend des conditions radio et surtout de son «retard relatif» : E_i / r_i .

Le module fonctionne avec des seuils ($\delta_1 \geq \dots \delta_{M-1} \geq \delta_M$)¹ associés aux modes de transmission utilisables ($c_1 \leq \dots c_{M-1} \leq c_M$). Si $\delta_M \leq \frac{lag_i}{r_i} \leq \delta_{M-1}$ (c'est à dire un faible retard relatif) et que le meilleur mode de transmission possible pour le flux i vaut c_M alors le module retourne $c = c_M$, dans le cas contraire (c'est à dire des conditions radio moyennes voire faibles) le débit retourné est négatif. Cela signifie qu'un flux n'a le droit

¹ Nos conventions concernant la numérotation des modes sont opposées à celles des auteurs, ce qui inverse le sens des inégalités. Le sens de leur proposition est évidemment préservé.

de transmettre avec un débit moyen ($c \leq c_M$) que s'il a un «retard relatif» important. Lorsque le débit indiqué par le module de sélection de débit est négatif, le module de compensation est lancé.

Dans le module de compensation, la classe Ψ_{RT} est choisie pour recevoir l'opportunité de transmission cédée si $v'_{RT} \leq v'_{NRT}$; ces paramètres sont des temps virtuels qui représentent la compensation reçue par chacune des classes. Si la classe Ψ_{RT} (resp. Ψ_{NRT}) est choisie, v'_{RT} (resp. v'_{NRT}) est incrémenté de L/r'_{RT} (resp. L/r'_{NRT}). Le poids r'_{RT} est inférieur au poids r'_{NRT} pour favoriser la classe temps réel. Quand un flux j est choisi en compensation (quelque soit sa classe), il présente le temps virtuel de compensation $v_{comp,j}$ le plus faible ; ce dernier est mis à jour par une formule similaire à l'équation (1). Les compteurs deviennent : $E(i) = E(i) + L$ et $E(j) = E(j) - L$. Le temps virtuel du flux initial i est actualisé : $v_i = v_i + \frac{L}{r_i}$.

Dans le module de dégradation de service, un flux en avance peut recevoir une quantité de service supplémentaire proportionnelle (facteur α) à son quantité de service en cours. Comme dans le module de compensation, deux facteurs α_{RT} et α_{NRT} distinguent les classes temps réel et non temps réel.

Certaines décisions du MR-FQ rappellent le WFS notamment la mise à jour des compteurs lorsqu'un flux en retard devient inactif (il n'a plus de paquets à transmettre) ou la reprise d'un flux en avance après le module de dégradation gracieuse de service lorsqu'il n'a pu être remplacé par un flux en retard.

En MR-FQ, les quantités de service en temps et en débit sont bornées sur tout intervalle de temps. L'intervalle de temps au bout duquel un flux rattrape son retard est majoré. Cependant, les expressions des bornes sont complexes. De plus, le MR-FQ maintient jusqu'à cinq temps virtuels différents pour distinguer les quantités de service d'un flux reçues directement, en compensation, en service additionnel, en dégradation de service etc... Par ailleurs, les règles du module de sélection de débit interdisent à un flux en avance d'utiliser de faibles débits de transmission ; le module de compensation est alors appelé. Cela peut se comprendre s'il y a vraiment un flux en retard qui peut transmettre, mais cette condition n'est pas assurée. Dans l'hypothèse où aucun flux en retard ne peut transmettre, l'opportunité de transmission est simplement perdue.

Opportunistic Weighed Fair Queuing ([Khawam_06])

Dans l'OWFQ, l'état du canal est introduit dans l'étiquette de fin de service pour favoriser les utilisateurs ayant un bon canal. L'étiquette de fin de service n'est calculée que pour les paquets de tête de queue :

$$F(p_i^k) = S(p_i^k) + \frac{L(p_i^k)}{r_i c_{m_i}} \quad \text{où } m_i \text{ est le mode de transmission adapté au flux } i \text{ à l'instant}$$

T d'ordonnancement et c_M est la capacité du mode m . L'étiquette de début de service est inchangée. Les garanties en terme d'équité sont données par :

$$\left| \frac{W_i(t_1, t_2)}{r_i} - \frac{W_j(t_1, t_2)}{r_j} \right| \leq L_{max} \left(\frac{1}{r_i} + \frac{1}{r_j} + \frac{1}{c_{min}} \max\left(\frac{1}{r_i}, \frac{1}{r_j}\right) \right)$$

Annexe 6.D L'ordonnancement en OFDMA

Ordonnancement indépendant par sous canal ; trafic NRT

Chapitre 6, section 1.3

Rappel sur l'ordonnancement PF ([Chaponniere&_02], [Jalali&_00])

Dans le PF, l'utilisateur qui maximise le rapport entre le débit instantané et le débit moyenné jusqu'à l'instant t est sélectionné. Cette technique d'ordonnancement proposée dans le cadre du CDMA, peut être appliquée en OFDMA sur chaque sous canal. Alors, la métrique s'exprime de la sorte :

$$\mu_{u,n}(t) = \frac{R_{u,n}(t)}{\overline{R_{u,n}}(t)} \quad (1)$$

$$\overline{R_{u,n}}(t+1) = \left(1 - \frac{1}{T_c}\right) \overline{R_{u,n}}(t) + \frac{1}{T_c} x_{u,n} R_{u,n}(t) \quad (2) \text{ où } r_{u,n}(t) \text{ est le débit instantané de}$$

l'utilisateur u sur le sous canal n et la variable $x_{u,n}$ vaut 1 si l'utilisateur u a reçu le sous canal n et 0 sinon ; T_c est une constante qui est de l'ordre de 1000 trames.

Variations de l'ordonnancement PF en OFDMA

Les auteurs de [Wengerter&_05] proposent le GPF (*Generalized Proportional Fair*). Ils introduisent des exposants au numérateur et au dénominateur de la métrique (1) afin de pouvoir moduler leur importance relative. La métrique de l'utilisateur u sur le sous canal¹ n s'écrit alors :

$$\mu_{u,n}(t) = \frac{R_{u,n}(t)^\alpha}{\overline{R_{u,n}}(t)^\beta} \quad . \text{ Lorsque } \alpha > \beta \text{ la tendance de la nouvelle métrique est plus}$$

opportuniste et dans le cas contraire, elle est plus équitable.

Les auteurs de [KiKim_05] proposent deux contributions qui conservent le rapport initial (1) de la métrique PF et se distinguent par un facteur supplémentaire. Leur première proposition, le WFS (*Weighted Fair Scheduling*), cherche à garantir une équité en débit entre les utilisateurs. La métrique (1) est modifiée comme suit :

$$\mu_{u,n}(t) = \frac{R_{u,n}(t)}{\overline{R_{u,n}}(t)} \frac{\frac{1}{U} \sum_{v=1}^U R_v(t)}{R_u(t)} \quad ; r_u \text{ le débit de l'utilisateur } u. \text{ Les utilisateurs qui sont}$$

favorisés sont ceux qui présentent de faibles débits par rapport au débit moyenné sur l'ensemble des utilisateurs.

La seconde proposition, le TGS (*Throughput Guarantee Scheduling*), cherche à garantir un débit requis par utilisateur. La métrique (1) est inchangée pour les utilisateurs u qui ont un débit supérieur au débit requis R_u^o . Pour ceux qui n'ont pas atteint leur débit requis la métrique est la suivante :

¹ Dans [Wengerter&_05], les sous porteuses d'un sous canal sont adjacentes.

$\mu_{u,n}(t) = \frac{R_{u,n}(t)}{\overline{R_{u,n}}(t)} \left(\frac{\max_v(\overline{R_{v,n}}(t))}{\overline{R_{u,n}}(t)} \right)^{R_u(t)/R_u^\circ}$. Le rapport $\frac{\max_v(\overline{R_{v,n}}(t))}{\overline{R_{u,n}}(t)}$ favorise les utilisateurs qui ont un faible débit moyenné sur le sous canal. L'exposant $R_u(t)/R_u^\circ$ augmente le rapport $\frac{\max_v(\overline{R_{v,n}}(t))}{\overline{R_{u,n}}(t)}$ pour les utilisateurs qui sont proches de leur débit minimum (car ils sont supposés plus faciles à satisfaire). Ces utilisateurs sont traités en priorité pour minimiser le nombre total d'insatisfaits.

Ordonnancement global en OFDMA ;trafic NRT

(chapitre 6, section 1.3.2)

Virtual Clock + optimisation MA

Les auteurs de [ZhangLetaief_04] utilisent l'ordonnancement *Virtual Clock*. Le canal est considéré sans erreurs. Après une estimation du nombre de paquets maximum que l'on peut envoyer sur une trame, les paquets (de taille fixe) présentant les plus faibles estampilles temporelles sont choisis. Ensuite, une optimisation MA est réalisée en donnant à chaque utilisateur un nombre de sous porteuses proportionnel au nombre de bits qu'il a émettre. Le but principal de l'optimisation est d'assurer un PER très faible. Cela permet de se rapprocher de l'hypothèse d'un canal sans erreurs faite au moment de l'ordonnancement. L'utilisateur ayant les pires conditions canal peut être écarté si la phase d'optimisation de l'allocation de ressources n'aboutit pas. Pour réaliser l'optimisation, les auteurs évoquent l'algorithme de [WongTsui&_99].

Annexe 6.E Preuves sur la région *schedulable* en EWFS et OWFS

Preuve du lemme 6.1

Cette preuve reprend les arguments de [Georgiadis&_97] p.1526, où les auteurs démontrent un lemme similaire dans le cadre de l'EDF (*Earliest Deadline First*) non préemptif et en temps réel. Dans [Lu&_00], les auteurs se contentent de citer la référence sans transposer la démonstration dans le temps virtuel ni introduire l'effet du paramètre *lookahead*.

On se place dans le pire cas où tous les flux sont actifs. En EWFS et en OWFS, la différence entre l'étiquette de fin de service (*finish tag*) et celle de début de service (*start tag*) d'un paquet est égale à $L_{max} / m_{i,best} \Phi_i$. C'est le nombre de tours que le paquet de tête doit attendre au maximum avant d'être servi. Pour que le vecteur de délai soit *schedulable* quelque soient les conditions radio, on pose pour tout i : $D_i = L_{max} / \Phi_i$.

L'équation (6.22) assure qu'un paquet de taille maximale peut être envoyé en respectant la plus forte contrainte de délai. En effet, le *scheduleur* doit pouvoir attribuer au flux 1,

L_{max} bits en D_1 tours. La quantité $D_1 \sum_{i=1}^U r_i$ est le nombre de bits traités par le

scheduleur en D_1 tours. Les poids sont normalisés, alors $\sum_{i=1}^U r_i = 1$. Finalement, D_1 doit être supérieur ou égal à L_{max} .

Pour établir l'équation (6.23), on considère maintenant le cas $L_{max} \leq v(t) - v_0 < D_U$

La quantité $v(t) - v_0$ représente à la fois le nombre de tours de scheduling réalisés pendant l'intervalle $[v_0, v(t)]$ et (pour les raisons de normalisation que l'on a déjà cité) le nombre de bits distribués par le *scheduleur* pendant cet intervalle. A v_0 , un paquet est en train d'être servi, il reste $(v(t) - v_0 - L_{max})$ bits à distribuer pendant $(v(t) - v_0)$ tours. Considérons donc ce qui se passe durant $(v(t) - v_0)$ tours pour chaque flux i :

Un paquet peut être sélectionné en avance si $S(p_i^k) < v(t) + \rho$. Si tel est le cas, il fini avec ρ tours d'avance. Si $v(t) - v_0 > L_{max} - \rho$, alors un paquet peut être envoyé avec ρ tours d'avance.

Si $v(t) - v_0 > D_i$, alors un paquet du flux i doit être envoyé (à cause de la contrainte de *schedulabilité* le service d'un paquet arrivé au temps virtuel v doit être terminé à $v + D_i$).

Finalement, si $v(t) - v_0 > \max(L_{max} - \rho, D_i)$, un paquet doit être envoyé, ce qui fait L_{max} bits. Dans ce cas, le nombre de tours restants pour transmettre d'autres bits du flux i vaut : $\min(v(t) - v_0 - L_{max} + \rho, v(t) - v_0 - D_i)$. On multiplie ce nombre de tours par r_i pour obtenir le nombre de bits restant à transmettre pendant l'intervalle $[v_0 \dots v(t)]$.

On définit $I(x) = x$ si $x > 0$ et 0 sinon.

Le nombre de bits à émettre entre v_0 et $v(t)$ est :

$$\sum_{i=1}^U [L_{max} + r_i \min(A, B)] * I(\min(A, B))$$

avec $A = v(t) - v_o - L_{max} + \rho$; $B = v(t) - v_o - D_i$

Le nombre de bits disponibles vaut $(v(t) - v_o) \sum_{i=1}^U r_i - L_{max} = v(t) - v_o - L_{max}$

On a donc:

$$\sum_{i=1}^U [L_{max} + r_i \min(A, B)] * I(\min(A, B)) \leq v(t) - v_o - L_{max} \quad \text{d'où le résultat (6.23)}$$

avec $A = v(t) - v_o - L_{max} + \rho$; $B = v(t) - v_o - D_i$

Pour montrer le résultat (6.24), on se place dans le cas $v(t) - v_o \geq D_U$

On a pour tout i , $v(t) - v_o \geq D_i$ et à fortiori $v(t) - v_o \geq L_{max}$.

Ainsi, $I(\min(v(t) - v_o - L_{max} + \rho, v(t) - v_o - D_i))$ vaut toujours 1. Le nombre de bits à émettre entre v_o et $v(t)$ est alors :

$$\sum_{i=1}^U [L_{max} + r_i \min(v(t) - v_o - L_{max} + \rho, v(t) - v_o - D_i)]$$

Le nombre de bits à distribuer vaut $v(t) - v_o$. On obtient ainsi le résultat (6.24).

Preuve du lemme 6.2

On appelle $W^R(v, d)$ la quantité de travail qui, au temps virtuel v , a une date limite de traitement de d . On veut montrer que si un vecteur de délai vérifie le lemme, alors pour tout v , $W^R(v, v) = 0$; ce qui implique que le vecteur est *schedulable*.

Soit $[v_0, v)$ l'intervalle où le serveur est occupé à servir des paquets de date limite au plus v . Soit P , l'ensemble des paquets servis dans $[v_0, v)$ ou présents à v . Soit p_e le paquet dont l'étiquette de début de service est la plus faible dans l'ensemble P .

La quantité $(v - v_0) + W^R(v, v)$ est le nombre de bits à distribuer par le serveur (on considère le pire cas où tous les flux sont actifs).

On a $v - S(p_e) \geq F(p_e) - S(p_e)$ et $F(p_e) - S(p_e) \geq D_1$ d'après la définition des étiquettes de début et de fin de service,

Comme le vecteur de délai vérifie (6.22), on a $v - S(p_e) \geq L_{max}$

Pour montrer que $W^R(v, v) = 0$, on traite deux cas.

Cas 1: $S(p_e) = v_0$

$v - v_0 \geq L_{max}$ donc on pourra utiliser les équations (6.23) et (6.24).

Comme on a l'a vu dans la preuve du lemme 6.1, la majoration du nombre de bits à émettre

est donnée par $\sum_{i=1}^U [L_{max} + r_i \min(A, B)] * I(\min(A, B))$ et on a donc

avec $A = v(t) - v_o - L_{max} + \rho$; $B = v(t) - v_o - D_i$

$$\sum_{i=1}^U [L_{max} + r_i \min(A, B)] * I(\min(A, B)) \geq v - v_o + W^R(v, v)$$

$$\text{avec } A = v(t) - v_o - L_{max} + \rho ; B = v(t) - v_o - D_i$$

Si $v - v_o < D_U$, le résultat (6.23) et la majoration ci-dessus impliquent que $W^R(v, v) \leq -L_{max}$

Si $v - v_o \geq D_U$, le résultat (6.24) et la majoration ci-dessus impliquent que $W^R(v, v) \leq 0$

On voit quelque soit la situation de $v - v_o$ par rapport à D_U , $W^R(v, v) = 0$ car c'est une quantité qui ne peut être que positive ou nulle.

Cas 1: $S(p_e) < v_o$

Soit q le paquet en service ou terminant son service à v_o ; par définition de v_o le paquet q n'est pas l'ensemble P , la date limite est donc supérieure à v . Le paquet q part avant le paquet p_e donc (le service étant non préemptif), on a $S(p_e) > S(q)$.

On a $v - S(q) < F(q) - S(q)$ et $F(q) - S(q) \leq D_U$ d'après la définition des étiquettes de début et de fin de service. On en déduit que $v - S(p_e) < v - S(q) < D_U$. On pourra utiliser le résultat (6.23) dans l'intervalle $[S(p_e), v]$

En majorant le nombre de bits à émettre pendant $[v_o, v]$ on a

$$\sum_{i=1}^U [L_{max} + r_i \min(A, B)] * I(\min(A, B)) \geq v - v_o + W^R(v, v)$$

$$\text{avec } A = v(t) - S(p_e) - L_{max} + \rho ; B = v(t) - S(p_e) - D_i$$

ou encore

$$\sum_{i=1}^U [L_{max} + r_i \min(A, B)] * I(\min(A, B)) \geq (v - S(p_e)) - (v_o - S(p_e)) + W^R(v, v)$$

$$\text{avec } A = v(t) - S(p_e) - L_{max} + \rho ; B = v(t) - S(p_e) - D_i$$

Le paquet p_e arrive avant que le paquet q ne soit envoyé alors $v_o - S(p_e) \leq L_{max}$

$$\sum_{i=1}^U [L_{max} + r_i \min(A, B)] * I(\min(A, B)) \geq (v - S(p_e)) - L_{max} + W^R(v, v)$$

$$\text{avec } A = v(t) - S(p_e) - L_{max} + \rho ; B = v(t) - S(p_e) - D_i$$

La relation (6.23) donne donc $W^R(v, v) \leq 0$ d'où $W^R(v, v) = 0$. Le lemme 6.2 est ainsi démontré dans les cas 1 et 2.

Preuve du théorème 6.1

Il découle directement des lemmes 6.1 et 6.2 en posant $v(t) - v_o = D_j$ pour $1 \leq j \leq U - 1$. L'équation (6.23) devient :

$$\sum_{i=1}^j [L_{max} + r_i \min(D_j - L_{max} + \rho, D_j - D_i)] + L_{max} \leq D_j \text{ d'où (6.25).}$$

Preuve du corollaire 6.1

Dans le théorème 6.1, les équations se simplifient avec : $L_{max} \leq D_i \leq D_k$

On alors $\min(D_k - L_{max} + \rho, D_k - D_i) = D_k - D_i$. D'où le corollaire.

Annexe 6.F Preuves sur les garanties de service en EWFS et OWFS

Preuve du lemme 6.3

On commence par prouver la partie (6.30). Si $v_2 < v_1 + \frac{L_{max}}{r_i} + \frac{L_{max}}{\phi_i}$ alors

$$r_i(v_2 - v_1) - L_{max} - L_{max} \frac{r_i}{\phi_i} < 0 \text{ et comme } W_i(v_1, v_2) \geq 0 \text{ l'inégalité (6.30) est vérifiée.}$$

On considère le cas non trivial $v_2 \geq v_1 + \frac{L_{max}}{r_i} + \frac{L_{max}}{\phi_i}$. Montrons qu'un paquet du flux i est servi dans l'intervalle $[v_1, v_2]$. On note p_i^{k-1} le paquet du flux i qui est servi à $v_1 - \varepsilon$, ce paquet vérifie $S(p_i^{k-1}) \leq v_1$ et $F(p_i^{k-1}) \geq v_1$. Selon l'étiquetage de début et de fin de service :

$$S(p_i^k) = S(p_i^{k-1}) + \frac{L(p_i^{k-1})}{r_i m_{i,max}(t_i^k)} \text{ (OWFS) ou } S(p_i^k) = S(p_i^{k-1}) + \frac{L(p_i^{k-1})}{r_i} \text{ (EWFS)}$$

$$\text{et } F(p_i^k) = S(p_i^k) + \frac{L(p_i^k)}{\phi_i m_{i,best}} \text{ (dans les deux cas).}$$

En majorant l'expression de l'étiquette de début de service (de l'OWFS), on a

$$S(p_i^k) \leq v_1 + \frac{L_{max}}{r_i m_{i,max}(t_i^k)} \text{ que l'on peut encore majorer ainsi (majoration obtenue}$$

directement pour EWFS) : $S(p_i^k) \leq v_1 + \frac{L_{max}}{r_i}$. D'après, l'hypothèse de travail

$$v_1 + \frac{L_{max}}{r_i} + \frac{L_{max}}{\phi_i} \leq v_2, \text{ on obtient alors (pour EWFS et OWFS) } S(p_i^k) < v_2.$$

En majorant l'expression de l'étiquette de fin de service (commune à l' EWFS et l'OWFS), on a $F(p_i^k) \leq S(p_i^k) + \frac{L_{max}}{\phi_i m_{i,best}}$ que l'on peut encore majorer ainsi :

$$F(p_i^k) \leq S(p_i^k) + \frac{L_{max}}{\phi_i}. \text{ Comme } S(p_i^k) \leq v_1 + \frac{L_{max}}{r_i}, \text{ on a}$$

$$F(p_i^k) \leq v_1 + \frac{L_{max}}{r_i} + \frac{L_{max}}{\phi_i}. \text{ D'après, l'hypothèse de travail } v_1 + \frac{L_{max}}{r_i} + \frac{L_{max}}{\phi_i} \leq v_2, \text{ on}$$

obtient alors $F(p_i^k) \leq v_2$.

On a montré que $S(p_i^k) < v_2$ et $F(p_i^k) \leq v_2$, ce qui prouve que le paquet p_i^k est servi dans l'intervalle $[v_1, v_2]$.

Soit p_i^{k+n+1} le dernier paquet du flux i à recevoir son service dans l'intervalle $[v_1, v_2]$, ce paquet vérifie $F(p_i^{k+n+1}) \geq v_2$.

En EWFS on peut dire que :

$$F(p_i^{k+n+1}) - \frac{L(p_i^{k+n+1})}{\phi_i m_{i,best}(T)} - \frac{L(p_i^{k+n})}{r_i} \geq v_2 - \frac{L(p_i^{k+n+1})}{\phi_i} - \frac{L(p_i^{k+n})}{r_i}$$

de façon similaire en OWFS on peut dire que :

$$F(p_i^{k+n+1}) - \frac{L(p_i^{k+n+1})}{\phi_i m_{i,best}(T)} - \frac{L(p_i^{k+n})}{r_i m_{i,max}(t_i^{k+n})} \geq v_2 - \frac{L(p_i^{k+n+1})}{\phi_i} - \frac{L(p_i^{k+n})}{r_i}$$

Puis, en identifiant à gauche et en majorant à droite on a (inégalité commune à l'EWFS et à l'OWFS) $S(p_i^{k+n}) \geq v_2 - \frac{L_{max}}{\phi_i} - \frac{L(p_i^{k+n})}{r_i}$.

D'autre part (en OWFS) $S(p_i^{k+n}) - S(p_i^k) = \sum_{j=0}^{n-1} \frac{L(p_i^{k+j})}{r_i m_{i,max}(t_i^{k+j})}$ donc $\sum_{j=0}^{n-1} \frac{L(p_i^{k+j})}{r_i} \geq S(p_i^{k+n}) - S(p_i^k)$ (cette inégalité en OWFS est une égalité en EWFS).

On obtient $\sum_{j=0}^{n-1} \frac{L(p_i^{k+j})}{r_i} \geq v_2 - \frac{L_{max}}{\phi_i} - \frac{L(p_i^{k+n})}{r_i} - S(p_i^k)$ ou encore $\sum_{j=0}^{n-1} \frac{L(p_i^{k+j})}{r_i} \geq v_2 - \frac{L_{max}}{\phi_i} - \frac{L(p_i^{k+n})}{r_i} - (v_1 + \frac{L_{max}}{r_i})$ (car $-S(p_i^k) \geq -(v_1 + \frac{L_{max}}{r_i})$).

On a $\sum_{j=0}^n \frac{L(p_i^{k+j})}{r_i} \geq v_2 - v_1 - \frac{L_{max}}{\phi_i} - \frac{L_{max}}{r_i}$.

Comme $W_i(v_1, v_2) \geq \sum_{j=0}^n L(p_i^{k+j})$, $W_i(v_1, v_2) \geq r_i(v_2 - v_1 - \frac{L_{max}}{\phi_i} - \frac{L_{max}}{r_i})$: c'est l'inégalité (6.30) qui est valable en OWFS et en EWFS.

Pour montrer l'inégalité (6.29), il faut considérer le pire cas suivant : le premier paquet du flux i arrive et que tous les autres flux sont et demeurent actifs. La démonstration est identique en OWFS et EWFS.

$F(p_i^{k-1}) = V(t_i^{k-1}) + \frac{L(p_i^{k-1})}{\phi_i m_{i,best}} \leq v_1 + \frac{L(p_i^{k-1})}{\phi_i}$, le paquet attend au maximum $\frac{L(p_i^{k-1})}{\phi_i}$ tours, le service perdu dans l'intervalle $[v_i, v_i + \frac{L(p_i^{k-1})}{\phi_i}]$ est au maximum $r_i \frac{L_{max}}{\phi_i}$. Une fois que le premier paquet a été servi le raisonnement de l'inégalité (6.30) tient à condition de retirer le service perdu par l'attente du premier paquet. Cela donne $W_i(v_1, v_2) \geq r_i(v_2 - v_1 - \frac{L_{max}}{\phi_i} - \frac{L_{max}}{r_i}) - r_i \frac{L_{max}}{\phi_i}$.
Finalement, $W_i(v_1, v_2) \geq r_i(v_2 - v_1) - \frac{L_{max}}{r_i} - 2r_i \frac{L_{max}}{\phi_i}$.

Preuve du théorème 6.2

Pour démontrer le théorème, on se place sur le sous ensemble des trajectoires dans lequel $C > C_{gar,\beta}$.

Soit v_1 tel que $V(t_1) = v_1$ et $V(t_2) = v_2$. On décompose l'intervalle $[v_1, v_2]$ en sous intervalles $[v^{j-1}, v^j]$ où l'ensemble A_j des flux actifs est fixe, sur ces intervalles, sachant que, $V(t^{j-1}) = v^{j-1}$ et $V(t^j) = v^j$, on peut écrire d'après (6.13)

$$v^j - v^{j-1} = \frac{(t^j - t^{j-1}) C_{gar,\beta}}{\sum_{i \in A_j} r_i}$$

D'après notre décomposition, $v_2 - v_1 = \sum_{j=2}^J v^j - v^{j-1}$ avec $v^J = v_2$ et $v^1 = v_1$

$$\text{D'où } v_2 - v_1 = C_{gar,\beta} \sum_{j=2}^J \frac{(t^j - t^{j-1})}{\sum_{i \in A_j} r_i}.$$

Comme $\sum_{i \in A_j} r_i \leq 1$, on a $v_2 - v_1 \geq C_{gar,\beta} \sum_{j=2}^J (t^j - t^{j-1})$; c'est à dire $v_2 - v_1 \geq C_{gar,\beta} (t_2 - t_1)$.

On a $W(t_1, t_2) = W(v_1, v_2)$ et d'après le lemme 6.3 :

$$W_i(v_1, v_2) \geq r_i(v_2 - v_1) - \frac{L_{max}}{r_i} - \alpha r_i \frac{L_{max}}{\phi_i}.$$

$$\text{On en déduit: } W_i(t_1, t_2) \geq r_i C_{gar,\beta} (t_2 - t_1) - \frac{L_{max}}{r_i} - \alpha r_i \frac{L_{max}}{\phi_i}.$$

Cette relation est vérifiée sur le sous ensemble où $C > C_{gar,\beta}$.

Comme $P(C > C_{gar,\beta}) \geq \beta$, on en déduit (6.31).

Preuve du lemme 6.4

- En EWFS

Soit A l'ensemble des paquets k tels que $v_1 \leq S(p_i^k)$, $S(p_i^k) < v_2$ et $F(p_i^k) < v_2$. On note p_i^k et p_i^{k+n} le premier et le dernier de ces paquets.

Le paquet p_i^{k+n+1} vérifie $S(p_i^{k+n+1}) \leq v_2$ et donc $S(p_i^{k+n+1}) - S(p_i^k) \leq v_2 - v_1$.

D'autre part, on a $S(p_i^{k+n+1}) - S(p_i^k) = \sum_{j=0}^n \frac{L(p_i^{k+j})}{r_i}$. On en déduit

$$\sum_{j=0}^n L(p_i^{k+j}) \leq r_i(v_2 - v_1). \text{ Dans l'ensemble A } W_i(v_1, v_2) = \sum_{j=0}^n L(p_i^{k+j}) \text{ on obtient : } W_i(v_1, v_2) \leq r_i(v_2 - v_1).$$

L'ensemble B des paquets tels que $S(p_i^k) \leq v_2$ et $F(p_i^k) > v_2$ ne contient qu'un paquet. Finalement, on a (6.32) c'est à dire $W_i(v_1, v_2) \leq r_i(v_2 - v_1) + L_{max}$.

- En OWFS

L'ensemble A est défini comme ci-dessus.

Par contre $S(p_i^{k+n+1}) - S(p_i^k) \leq v_2 - v_1$ et $S(p_i^{k+n+1}) - S(p_i^k) \leq \sum_{j=0}^n \frac{L(p_i^{k+j})}{r_i}$: on ne peut pas conclure tel quel.

On utilise plutôt $\sum_{j=0}^n \frac{L(p_i^{k+j})}{r_i m_{\max}} \leq S(p_i^{k+n+1}) - S(p_i^k)$. Avec

$S(p_i^{k+n+1}) - S(p_i^k) \leq v_2 - v_1$ on a alors $\sum_{j=0}^n L(p_i^{k+j}) \leq r_i m_{\max} (v_2 - v_1)$. Dans

l'ensemble A $W_i(v_1, v_2) = \sum_{j=0}^n L(p_i^{k+j})$, on obtient : $W_i(v_1, v_2) \leq r_i m_{\max} (v_2 - v_1)$.

L'ensemble B des paquets tels que $S(p_i^k) \leq v_2$ et $F(p_i^k) > v_2$ ne contient qu'un paquet. Finalement, on a (6.33) c'est à dire $W_i(v_1, v_2) \leq r_i m_{\max} (v_2 - v_1) + L_{\max}$.

Annexe 6.G Preuves sur les garanties d'équité en EWFS et en OWFS

Preuve du théorème 6.3

● en EWFS

D'après les lemmes 6.3 et 6.4, on peut encadrer la quantité de service reçue entre 0 et v

$$r_i v - L_{\max} - \alpha L_{\max} \frac{r_i}{\phi_i} \leq W_i(0, v) \leq r_i v + L_{\max} . \text{ En divisant tout par } v, \text{ on a}$$

$$r_i - \frac{L_{\max}}{v} - \alpha \frac{L_{\max}}{v} \frac{r_i}{\phi_i} \leq \frac{W_i(0, v)}{v} \leq r_i + \frac{L_{\max}}{v} . \text{ En faisant tendre } v \text{ vers l'infini, on obtient}$$

(6.34).

● en OWFS

D'après les lemmes 6.3 et 6.4, on peut encadrer la quantité de service reçue entre 0 et v

$$r_i v - L_{\max} - \alpha L_{\max} \frac{r_i}{\phi_i} \leq W_i(0, v) \leq r_i m_{\max} v + L_{\max} . \text{ En divisant tout par } v, \text{ on a}$$

$$r_i - \frac{L_{\max}}{v} - \alpha \frac{L_{\max}}{v} \frac{r_i}{\phi_i} \leq \frac{W_i(0, v)}{v} \leq r_i m_{\max} + \frac{L_{\max}}{v} . \text{ En faisant tendre } v \text{ vers l'infini, on}$$

obtient (6.35).

Preuve du Corollaire 6.2

● en EWFS

D'après les lemmes 6.3 et 6.4, on peut encadrer la quantité de service reçue entre v_1 et v_2 puis en divisant par r_i (respectivement r_j) on a

$$(v_2 - v_1) - \frac{L_{\max}}{r_i} - \alpha \frac{L_{\max}}{\phi_i} \leq \frac{W_i(v_1, v_2)}{r_i} \leq (v_2 - v_1) + \frac{L_{\max}}{r_i}$$

$$\rightarrow \text{ et } -(v_2 - v_1) - \frac{L_{\max}}{r_j} \leq \frac{-W_j(v_1, v_2)}{r_j} \leq -(v_2 - v_1) + \frac{L_{\max}}{r_j} + \alpha \frac{L_{\max}}{r_j}$$

En additionnant terme à terme :

$$-\frac{L_{\max}}{r_j} - \frac{L_{\max}}{r_i} - \alpha \frac{L_{\max}}{\phi_i} \leq \frac{W_i(v_1, v_2)}{r_i} - \frac{W_j(v_1, v_2)}{r_j} \leq \frac{L_{\max}}{r_j} + \frac{L_{\max}}{r_i} + \alpha \frac{L_{\max}}{\phi_j}$$

Comme $-\alpha \frac{L_{\max}}{\min(\phi_i, \phi_j)} \leq -\alpha \frac{L_{\max}}{\phi_i}$ et $\alpha \frac{L_{\max}}{\phi_j} \leq \alpha \frac{L_{\max}}{\min(\phi_i, \phi_j)}$

On remplace à gauche et à droite et on obtient :

$$-\frac{L_{\max}}{r_j} - \frac{L_{\max}}{r_i} - \alpha \frac{L_{\max}}{\min(\phi_i, \phi_j)} \leq \frac{W_i(v_1, v_2)}{r_i} - \frac{W_j(v_1, v_2)}{r_j} \leq \frac{L_{\max}}{r_j} + \frac{L_{\max}}{r_i} + \alpha \frac{L_{\max}}{\min(\phi_i, \phi_j)}$$

On a de la même façon :

$$\frac{-L_{\max}}{r_j} - \frac{L_{\max}}{r_i} - \alpha \frac{L_{\max}}{\min(\phi_i, \phi_j)} \leq \frac{W_j(v_1, v_2)}{r_j} - \frac{W_i(v_1, v_2)}{r_i} \leq \frac{L_{\max}}{r_j} + \frac{L_{\max}}{r_i} + \alpha \frac{L_{\max}}{\min(\phi_i, \phi_j)}$$

D'où (6.36), $\left| \frac{W_i(v_1, v_2)}{r_i} - \frac{W_j(v_1, v_2)}{r_j} \right| \leq \frac{L_{\max}}{r_i} + \frac{L_{\max}}{r_j} + \frac{\alpha L_{\max}}{\min(\phi_i, \phi_j)}$

● en OWFS

D'après les lemmes 6.3 et 6.4, on peut encadrer la quantité de service reçue entre v_1 et v_2 puis en divisant par r_i (respectivement r_j) on a :

$$(v_2 - v_1) - \frac{L_{\max}}{r_i} - \alpha \frac{L_{\max}}{\phi_i} \leq \frac{W_i(v_1, v_2)}{r_i} \leq m_{\max}(v_2 - v_1) + \frac{L_{\max}}{r_i}$$

et $-m_{\max}(v_2 - v_1) - \frac{L_{\max}}{r_j} \leq \frac{W_j(v_1, v_2)}{r_j} \leq -(v_2 - v_1) + \frac{L_{\max}}{r_j} + \alpha \frac{L_{\max}}{r_j}$

En additionnant terme à terme :

$$(v_2 - v_1)(1 - m_{\max}) - \frac{L_{\max}}{r_j} - \frac{L_{\max}}{r_i} - \alpha \frac{L_{\max}}{\phi_i} \leq \frac{W_i(v_1, v_2)}{r_i} - \frac{W_j(v_1, v_2)}{r_j}$$

et $\frac{W_i(v_1, v_2)}{r_i} - \frac{W_j(v_1, v_2)}{r_j} \leq (v_2 - v_1)(m_{\max} - 1) + \frac{L_{\max}}{r_j} + \frac{L_{\max}}{r_i} + \alpha \frac{L_{\max}}{\phi_j}$

D'après les mêmes arguments que précédemment :

$$(v_2 - v_1)(1 - m_{\max}) - \frac{L_{\max}}{r_j} - \frac{L_{\max}}{r_i} - \alpha \frac{L_{\max}}{\min(\phi_i, \phi_j)} \leq \frac{W_i(v_1, v_2)}{r_i} - \frac{W_j(v_1, v_2)}{r_j}$$

et $\frac{W_i(v_1, v_2)}{r_i} - \frac{W_j(v_1, v_2)}{r_j} \leq (v_2 - v_1)(m_{\max} - 1) + \frac{L_{\max}}{r_j} + \frac{L_{\max}}{r_i} + \alpha \frac{L_{\max}}{\min(\phi_i, \phi_j)}$

on peut interchanger i et j et d'où (6.37) en OWFS :

$$\left| \frac{W_i(v_1, v_2)}{r_i} - \frac{W_j(v_1, v_2)}{r_j} \right| \leq (v_2 - v_1)(m_{\max} - 1) + \frac{L_{\max}}{r_i} + \frac{L_{\max}}{r_j} + \frac{\alpha L_{\max}}{\min(\phi_i, \phi_j)}$$

Annexe 6.H Preuves sur les garanties de délai en EWFS et en OWFS

Preuve du théorème 6.4

Le pire cas se produit lorsqu' à l'arrivée du paquet, les autres flux sont et demeurent actifs. Le délai maximal du flux i est en nombre de tours L_{max}/Φ_i . Chaque flux $j \neq i$ reçoit au maximum pendant cette période $r_j L_{max}/\Phi_i + L_{max}$ (où on prend en compte au maximum un paquet additionnel).

La quantité totale de service reçue par les flux qui sont servis avant le flux i est alors :

$$\sum_{j \in F}^{j \neq i} \left(r_j \frac{L_{max}}{\Phi_i} + L_{max} \right)$$

Pour démontrer le théorème, on se place sur le sous ensemble des trajectoires dans lequel $C > C_{gar,\beta}$. Étant donné la capacité garantie du serveur $C_{gar,\beta}$; la relation entre l'heure de départ et heure d'arrivée du paquet de tête (*Head of Line*, HOL) vaut :

$$T_i^{HOL} \leq A_i^{HOL} + \frac{1}{C_{gar,\beta}} \left\{ \sum_{j \in F}^{j \neq i} \left(r_j \frac{L_{max}}{\Phi_i} + L_{max} \right) + L_{max} \right\}$$

Comme $\sum_{j \in F}^{j \neq i} r_j \leq 1$ on a $T_i^{HOL} - A_i^{HOL} \leq \frac{L_{max}}{C_{gar,\beta} \Phi_i} + \sum_{i \in F} \frac{L_{max}}{C_{gar,\beta}}$ d'où le résultat (6.38).

Pendant la période d'attente L_{max}/Φ_i du flux i , si on ne prends pas en compte le service d'un paquet additionnel par flux $j \neq i$ (cela correspond au cas des facteurs r_i et Φ_i invariants), on a plutôt :

$$T_i^{HOL} \leq A_i^{HOL} + \frac{1}{C_{gar,\beta}} \left\{ \sum_{j \in F}^{j \neq i} \left(r_j \frac{L_{max}}{\Phi_i} \right) + L_{max} \right\}$$

Comme $\sum_{j \in F}^{j \neq i} r_j \leq 1$, cela donne $T_i^{HOL} - A_i^{HOL} \leq \frac{L_{max}}{C_{gar,\beta} \Phi_i} + \frac{L_{max}}{C_{gar,\beta}}$.

Cette relation est vérifiée sur le sous ensemble dans lequel $C > C_{gar,\beta}$.

Comme $P(C > C_{gar,\beta}) \geq \beta$, on en déduit (6.39).

Preuve du théorème 6.5

A partir du moment où un paquet arbitraire devient un paquet de tête, on peut appliquer le théorème précédent sur l'EWFS et l'OWFS. L'heure de départ d'un paquet est donc majorée par l'EAT du paquet (l'heure maximale à partir de laquelle il est en tête de queue) auquel on ajoute la borne de délai d'un paquet de tête (6.38 ou 6.39).

Références

- [Acena&_05] M. Acena, S. Pflersinger. "A Spectrally Efficient Method for Subcarrier and Bit Allocation in OFDMA". IEEE 61st Vehicular Technology Conference, VTC -Spring, 30 May-1 June 2005 Volume 3, Page(s) : 1773 - 1777 Vol. 3.
- [Anas_04] Mohammad Anas. PHD Thesis, "QoS aware Subcarrier and Power Allocation in OFDMA Systems for Broadband Wireless Applications". Advisor : Kiseon Kim. 2004.
- [Anchun&_03] Wang Anchun, Xiao Liang, Zhou Shidong, Xu Xibin, Yao Yan. "Dynamic resource management in the fourth generation wireless systems". ICCT 2003. International Conference on Communication Technology Proceedings, 2003. Publication Date : 9-11 April 2003 Volume : 2, on page(s) : 1095- 1098 vol.2
- [Andrews&_07] J. Andrews, A. Ghosh, R. Muhamed, "Fundamentals of WiMax", chapter 6: OFDMA, collection: communications engineering and emerging technologies series. Prentice Hall. 23/07/2007.
- [AzginKrunz_03] Azgin, A. Krunz, M. "Scheduling in wireless cellular networks under probabilistic channel information" The 12th International Conference on Computer Communications and Networks, ICCCN 2003. Proceedings. Publication Date : 20-22 Oct. 2003. On page(s) : 89- 94.
- [Bhagwat&_96] P. Bhagwat, A. Krishna, and S. Tripathi, Enhancing throughput over wireless lanss using channel state dependent packet scheduling, in Proceedings of INFOCOMS96 (1996), 1133-1140.
- [Benedetto&_99] Sergio Benedetto, Ezio Biglieri. « Principles of Digital transmission with wireless Applications». 1999.
- [Bertsekas_82] Dimitri P. Bertsekas. Computer science and applied mathematics : "Constrained optimization and Lagrange multiplier methods". Academic Press, Inc. 1982.
- [Campello_98] Jorge Campello «Optimal discrete Bit loading for Multicarrier Modulation Systems». 1998
- [CCM_Wimax] Wimax-802.16-Worldwide Interoperability for Microwave Access. Issu de "Comment ça marche", sous les termes de la license Creative Commons ; www.commentcamarche.net/wimax/wimax-intro.php3
- [Chaponniere&_02] E.F. Chaponniere, P.J. Black, J.M. Holtzman, and D.N.C. Tse, 'Transmitter directed code division multiple access system using path diversity to equitably maximize throughput, US Patent 6 (2002).
- [Ciblat_04] Philippe Ciblat, «Introduction à l'OFDM» ENST, 2004.
- [Ciblat_05] Philippe Ciblat, « Systèmes multiporteuses à accès multiples », ENST, 2005
- [ChangKuo_04] Min-Kuan Chang, Jay Kuo. Power Control, Adaptive modulation and Subchannel Allocation for Multiuser downlink OFDM, 2004.
- [Chee&_04] Tommy K. Chee, Cheng-Chew Lim, Jinho Choi. "Sub-optimal Power Allocation for downlink OFDMA Systems". IEEE 60th VTC2004, 26-29 Sept. 2004, Page(s) : 2015- 2019 Vol. 3.
- [Chen&_04] Y. Chen, J. Chen, C. Li. "A Fast Suboptimal Subcarrier, Bit, and Power Allocation Algorithm for Multiuser OFDM-based Systems", IEEE International Conference on Communications, 20-24 June 2004, Vol : 6, pp. 3212- 3216.

- [ChenKron_07] Liang Chen, Brian Krongold «An Iterative resource Allocation Algorithm for Multiuser OFDM with Fairness and QoS Constraints». Vehicular Technology Conference, VTC 2007-Spring. Dublin Ireland 22-25 April 2007.
- [Cho&_05] M. Cho, W. Seo, Y. Kim, D.Hong. "A Joint Feedback Reduction Scheme Using Delta Modulation for Dynamic Channel Allocation in OFDMA Systems". The 16th IEEE International Symposium on Personal, Indoor and Mobile Radio Communications PIMRC 2005.
- [CoverThomas_91] T. M. Cover and J. A. Thomas. "Elements of Information Theory". John Wiley and sons, Inc.1991.
- [Czylwik_96] A. Czylwik. "Adaptive OFDM for wideband radio channels". Global Telecommunications Conference, 1996. GLOBECOM '96. 'Communications : The Key to Global Prosperity Volume 1,18-22 Nov. 1996. Page(s) :713 - 718 vol.1.
- [Diao&_04] Z. Diao, D. Shen Li, V.O.K."CPLD-PGPS scheduling algorithm in wireless OFDM systems». Global Telecommunications Conference, GLOBECOM '04. Date : 29 Nov.-3 Dec. 2004. Vol : 6, On page(s) : 3732- 3736 Vol.6.
- [DoufexiArmour_05] A. Doufexi, S. Armour. « Design considerations and Physical Layer Performance Results for a 4G OFDMA System Employing Dynamic Subcarrier Allocation ». The 16th IEEE International Symposium on Personal, Indoor and Mobile Radio Communications, PIMRC 2005.
- [ErmoMak_05] N.Y. Ermolova and B.Makarevitch. "Power and Subcarrier Allocation Algorithms for OFDMA Systems". The 16th IEEE International Symposium on Personal, Indoor and Mobile Radio Communications, PIMRC 2005.
- [Eugene&_98] T. S. Eugene Ng, I. Stoica, and H. Zhang, « Packet fair queueing algorithms for wireless networks with location-dependent errors » In Proceedings of IEEE Infocom'98 (1998).
- [FuSheen_07] I-K. Fu, W. Sheen. "An Analysis on Downlink Capacity of Multi-Cell OFDMA Systems under Randomized Inter-cell/sector Interference". IEEE 64th Vehicular Technology Conference, 2007. VTC -Spring. 22-25 april 2007.
- [Gault&_05] Sophie Gault, Walid Hachem, Philippe Ciblat, "Performance Analysis of an OFDMA Transmission System in a Multi-Cell Environment". Submitted to IEEE Transactions on Communication in july 2005.
- [Georgiadis&_97] Leonidas Georgiadis, Roch Guérin, Abhay Parekh. "Optimal Multiplexing on a single link : Delay and Buffer Requirements". IEEE Transactions on Information Theory, Vol 43, No. 5, september 1997.
- [Gross&_03] James Gross, Holger Karl, Frank Fitzek, Adam Wolisz. "Comparison of heuristic and optimal subcarrier assignment algorithms", Proceedings of IEEE International Conference on Wireless Networks (ICWN) 2003, pp. 249-255.
- [Goldsmith] Andrea Goldsmith, "Wireless Communications", Stanford University.
- [Golestani_94] S. J. Golestani, "A Self-Clocked Fair Queueing Scheme for Broadband Applications," Proceedings of IEEE INFOCOM, Toronto, Canada, Jun. 1994.
- [Gondran&_79] Michel Gondran, Michel Minoux. «Graphes et algorithmes». Eyrolles. 1979.
- [Goyal&_97] P. Goyal, H.M. Vin, and H. Cheng, "Start-Time Fair Queueing : A Scheduling Algorithm for Integrated Services Packet Switching Networks," IEEE/ACM Transactions on Networking, Vol. 5, No. 5, , Oct. 1997, pp. 690-704.
- [Hamouda&_06] Hamouda, S.; Godlewski, P.; Tabbane, S.; "Enhanced Capacity for Multi-Cell OFDMA Systems with Efficient Power Control and Reuse Partitioning". International Conference on Communication systems. ICCS 2006. 10th IEEE SingaporeOct. 2006 Page(s):1 - 5
- [Han_03] Zhu Han. Directed by K.J.Ray.Liu. PHD Thesis :«An optimization Theoretical Framework for resource allocation over Wireless Networks», 2003.
- [HuiZhou_06] Juhong Hui, Yongxing Zhou. "Enhanced Rate Adaptive Resource Allocation Scheme in Downlink OFDMA System". Vehicular Technology Conference, VTC2006-Spring. 7-10 may 2006.
- [IEEE 802.16] 802.16-2004. «Part 16 : Air Interface for Fixed Broadband Wireless Access Systems" IEEE Computer Society and the IEEE Microwave Theory and Techniques Society .

- [IEEE 802.16e] Part 16 : Air Interface for Fixed and Mobile Broadband Wireless Access Systems . Amendment for Physical and Medium Access Control Layers for Combined Fixed and Mobile Operation in Licensed Bands 18 février 2005
- [Inyoung&_01] Inhyoung Kim, Hae Leem Lee, Beomsup Kim, Yong H. Lee. “On the use of linear programming for dynamic subchannel and bit allocation in multiuser OFDM”, IEEE Global Telecommunications Conference, no. 1, Nov 2001, pp. 3648 – 3652.
- [Issaiyakul_99] Teerawat Issaiyakul, « Scheduling and Error Control in Cellular and Multi-hop Wireless Networks : Analysis and Optimization ». PHD Thesis directed by Ekram Hossein, Dpt of Electrical engineering. University of Manitoba.
- [Jalali&_00] R. Jalali, A Padovani, R Pankaj, “Data throughput of CDMA-HDR a high efficiency-high data rate personal communication wireless system” IEEE 51st Vehicular Technology Conference Proceedings, VTC 2000-Spring Tokyo. Volume : 3, On page(s) : 1854-1858 vol.3
- [JangLee_03] Jiho Jang, Kwang Bok Lee. “Transmit Power Adaptation for Multiuser OFDM Systems”, IEEE Journal on Selected Areas in Communications, vol. 21, no. 2, pp. 171-178, Feb. 2003.
- [Jainroek_76]•L. K Jainroek, Queueing Systems Vol. 2: Computer Applications. NewYork: Wiley, 1976.
- [Jeong&_06] S. Jeong ; D. Jeong ; W. Jeon.“Cross-layer Design of Packet Scheduling and Resource Allocation in OFDMA Wireless Multimedia Networks”. 63rd Vehicular Technology Conference, VTC 2006-Spring. Vol : 1,On page(s) : 309- 313
- [Junqiang&_03] J. Li, H. Kim, Y. Lee, Y. Kim “A novel Broadband Wireless OFDMA Scheme for Downlink in Cellular Communication”.Wireless Communications and Networking, WCNC 2003. IEEE 16-20 March 2003. 1907- 1911 vol.3.
- [Khawam_06] K. Khawam, « Ordonnancement opportuniste dans les réseaux mobiles de nouvelle génération ». 14 novembre 2006. Thèse de doctorat, dir: daniel Kofman. Ecole Nationale supérieure de Télécommunications, ENST-Paris.
- [Kim&_04] H. Kim, Y. Han and J. Koo. “Optimal Subchannel Allocation Scheme in Multicell OFDMA Systems”. Vehicular Technology Conference, 2004. VTC 2004-Spring. IEEE 59th-19 May 2004. Vol : 3, On page(s) : 1821- 1825 Vol.3.
- [KiKim_05] Young Min KI, Dong Ku KIM, « Packet Scheduling Algorithms for Throughput fairness and Coverage Enhancement in TDD-OFDMA Downlink Network ». IEICE Transactions On Communications : 2005, Volume 88, Numéro 11, Pages 4402-4405.
- [KimHan&_04] Keunyoung Kim, Hoon Kim, Younghan Han. “Subcarrier and Power Allocation in OFDMA Systems”. IEEE 60th Vehicular Technology Conference,VTC-Fall. 26-29 Sept. 2004, Volume 2, Page(s) :1058 - 1062 Vol.2.
- [KimKwak&_04] Ho Seok Kim, Jin Sam Kwak, Jung Min Choi, and Jae Hong Lee. “Efficient Subcarrier and Bit Allocation Algorithm for OFDMA System with Adaptive Modulation”, in Proc. IEEE VTC Spring, May 2004.
- [Kivanc&_00] Didem Kivanc, Hui Liu. “Subcarrier Allocation and Power Control for OFDMA”,Conference Record of the Thirty-Fourth Asilomar Conference on Signals, Systems and Computers, 2000, pp. 147-151 vol.1.
- [Kivanc&_03] Didem Kivanc, Guoqing Li, Hui Liu. “Computationally Efficient Bandwidth Allocation and Power Control for OFDMA”, IEEE Transactions on Wireless Communications, Nov 2003, vol. 2, no. 6, pp. 1150 – 1158.
- [Ko&_06] SanJun Ko, Joo Heo, KyungHi Chang. «Aggressive subchannel allocation algorithm for efficient dynamic channel allocation in multi-user OFDMA system». The 17th IEEE International Symposium on Personal, Indoor and Mobile Radio Communications, PIMRC 2006 : 11-14 September 2006, Helsinki, Finland.
- [Koubaa_04] Anis Koubaa, « Gestion de la qualité de Service temporelle selon la contrainte (m,k)-firm dans les réseaux à commutation de paquets ». Thèse de doctorat sous la direction de Jean Pierre Thomesse, 27 octobre 2004. Institut National Polytechnique de Lorraine.
- [Koutsopoulos_02] Iordanis Koutsopoulos. Advisor : Leandros Tassiulas. PHD Thesis :“Resource Allocation issues in broadband wireless networks with OFDM signalling”, 2002. University of Maryland.

- [Kwon&_05] H. J. Kwon ; W. I. Lee ; B. G. Lee, "Low-Overhead resource allocation with load balancing in multi-cell OFDMA systems", IEEE Vehicular Technology Conference, IEEE VTC 2005-Spring, Stockholm, Sweden, Mar. 2005.
- [Lamy_00] C. Lamy, «Communications à grande efficacité spectrale sur le canal à évanouissements», 2000, Thèse ENST.
- [Lawrey_99] E. Lawrey "Multiuser OFDM". International Symposium on Signal Processing and Its Applications, ISSPA 22-25 Aug '99. Proceedings of the Fifth Volume 2, Page(s) :761 - 764 vol.2.
- [LengGod&_06] C. Lengoumbi, P. Godlewski, P. Martins, «An Efficient Subcarrier Assignment Algorithm for Downlink OFDMA» . The 64th IEEE Semiannual Vehicular Technology Conference, 25-28 september 2006, Montreal, Canada.
- [LengGodM_06] C. Lengoumbi, P. Godlewski, P. Martins, "Dynamic Subcarrier Reuse with Rate Guaranty in a Downlink Multicell OFDMA System". The 17th IEEE International Symposium on Personal, Indoor and Mobile Radio Communications, PIMRC 2006 : 11-14 September 2006, Helsinki, Finland.
- [LengGod&_07] C. Lengoumbi, P. Godlewski, P. Martins "Subchannelization Performance for the Downlink of a Multi-Cell OFDMA System". Third IEEE International Conference on Wireless and Mobile Computing, Networking and Communications, WiMob'07. New York, USA 8-10 October 2007.
- [LengGodM_07] C. Lengoumbi, P. Godlewski and P. Martins, "Comparison of Different Subchannelization Modes for OFDMA". 18th IEEE Personal Indoor Mobile Radio Conference, PIMRC'07. Athens, Greece 3-6 september 2007.
- [LengMar&_07] C. Lengoumbi, P. Martins and P. Godlewski. "Characterization of Wireless Fair Service Extensions for Packet Scheduling in OFDMA". 12th International OFDM-workshop 2007. Hamburg, Germany 29-30 August 2007.
- [LengMarG_07] C. Lengoumbi, P. Martins and P. Godlewski. «An Opportunist extension of Wireless Fair Service for Packet Scheduling in OFDMA» Vehicular Technology Conference, 2007. VTC2007-Spring. IEEE 65th Publication Date : 22-25 April 2007. On page(s) : 3001-3005.
- [Lerbour&_06] Lerbour, T. Kurt, Y. Le Helloco, B. Breton. « Broadband Wireless Access Interference and capacity Estimation. ». IEEE 64th Vehicular Technology Conference, Montreal, Canada, 25-28 sept 2006.
- [LiLiu_05] Guoqing Li, Hui Liu. « On the Optimality of the OFDMA Network ». 2005
- [Lin&_00] P. Lin, B. Bensaou, Q.L. Ding, K.C. Chua. "A wireless fair scheduling algorithm for error-prone wireless channels". International Workshop on Wireless Mobile Multimedia. Proceedings of the 3rd ACM international workshop in Wireless mobile multimedia. Year of Publication :2000. Pages : 11 – 20.
- [Liu&_04] Q. Liu, S. Zhou, G. B. Giannakis, "Cross-Layer Modeling of Adaptive Wireless Links for QoS Support in Heterogeneous Wired-Wireless Networks". Wireless Networks, Volume 12, Issue 4, Pages : 427 - 437.
- [Lu&_98] Lu S. , Nandagopal, T., and Bharghavan, V. "A wireless fair service algorithm for packet cellular networks". In Proceedings of the 4th Annual ACM/IEEE international Conference on Mobile Computing and Networking (Dallas, Texas, United States, October 25 - 30, 1998). W. P. Osborne and D. Moghe, Eds. MobiCom '98.
- [Lu&_99] S. Lu and V. Bharghavan, R. Srikant, « Fair scheduling in wireless packet networks », IEEE/ACM Trans. Networking 7(4) (1999), 473-489.
- [Lu&_00] Songwu Lu, Thyagarajan Nandagopal, Vaduvur Bharghavan. "Design and analysis for fair service in error prone wireless channels". Wireless Networks, 2000.
- [MaoWang_06] Z. Mao, X. Wang, « Branch and Bound Approach to OFDMA Radio ressource Allocation ». Vehicular Technology Conference, 2006. VTC 2006-Fall, The 64th IEEE Semiannual, 25-28 september 2006, Montreal, Canada.
- [Martinian_04] E. Martinian, "Waterfilling gains at most $O(1/\text{SNR})$ at high SNR." unpublished noted, Feb. 2004, available at <http://www.csua.berkeley.edu/~emin/research/wfill.pdf>.
- [MorettiMorelli_06] M. Moretti, M. Morelli, «A novel dynamic subcarrier assignment scheme for multiuser OFDMA

systems» 2006.

[Mourad_06] Abdel-Majid Mourad “On the System Level Performance of MC-CDMA Systems in the Downlink”. PHD thesis, Ramesh PYNDIAH (Dir.), Ecole Nationale Supérieure des Télécommunications de Bretagne ENST-B, 2006.

[Nagle_87] J. Nagle, “On Packet Switches with Infinite Storage,” (RFC 970) IEEE Transactions on Communications, Vol. COM-35, No. 4, April 1987, pp. 435-438.

[Nakad_03] Jawad Nakad. “Allocations de ressources radio dans un réseau local sans fil (WLAN) de type OFDM. Mémoire de DEA. 2003. Faculté de Genie du Liban.

[NguyenHan_06] T. Nguyen, Y. Han. «A Dynamic Channel Assignment Algorithm for OFDMA Systems». Vehicular Technology Conference, 2006. The 64th IEEE Semiannual, 25-28 september 2006, Montreal, Canada.

[PapaSteig_82] C.H Papadimitiou, K.Steiglitz. «Combinatorial Optimization, algorithms and complexity». 1982 Prentice Hall.

[PaD_Wimax] Découverte du WiMax www.touslesreseaux.com/test-dossier/Partez-a-la-decouverte-du-WIMAX.html

[ParekhGallager_93] K. Parekh, R. G. Gallager, “A Generalized Processor Sharing Approach to Flow Control in Integrated Services Networks : The single-node case” IEEE/ACM Tr. on Networking, Vol. 1, No. 3, June 1993, pp. 344—357.

[Park&_04] H Park, JS Kim, S Yoon, T Kim, J Lee. "Scheduling algorithm in the multiuser OFDM wireless mobile communication systems". International Conference on Telecommunications ICT, 2004.

[PerezFon_05] Daniel Perez Palomar, Javier R. Fonollosa. “Practical Algorithms for a Family of Waterfilling Solutions”, IEEE Trans. on Signal Processing, Vol. 53, No. 2, pp. 686-695, Feb. 2005.

[Pfletschinger&_02] S. Pfletschinger, J. Speidel, G. Münz. “Efficient Subcarrier Allocation for Multiple Access in OFDM Systems”, 7th International OFDM-Workshop 2002 (InOWo'02), Hamburg, Germany, September 2002, pp. 21-25, 10-11.

[PietrykJan_02] Slawomir Pietrzyk, Gerard J.M. Janssen. “Multiuser Subcarrier Allocation for QoS provision in the OFDMA Systems”, IEEE 56th Vehicular Technology Conference Proceedings. VTC 2002-Fall. pp. i- xl.

[PietrzykJanssen_03] S.Pietrzyk, G.J.M.Janssen, "Subcarrier allocation and power control for QoS provision in the presence of CCI for the downlink of cellular OFDMA systems". Vehicular Technology Conference, 2003. VTC 2003-Spring. The 57th IEEE Semiannual, 22-25 April 2003, Volume : 4, On page(s) : 2221- 2225 vol.4.

[RheeCioffi_00] W. Rhee, J. M. Cioffi. “Increase in Capacity of Multiuser OFDM System Using Dynamic Subchannel Allocation”, IEEE, Vehicular Technology Conference Proceedings, VTC 2000-Spring Tokyo. 2000 pp.1085-1089 vol.2.

[Riato&_07] N. Riato, S. Sorrentino, D. Franco, Carlo Masseroni, M. Rastelli, R. Trivisonno. “Impact of Mobility on physical and MAC layer algorithms performance in Wimax System”. The 18th IEEE International Symposium on Personal, Indoor and Mobile Radio Communications, PIMRC 2007 : 3-7 September 2007, Athens, Greece.

[RohlGrund_96] H. Rohling, R. Grunheid: “Performance of an OFDMTDMA Mobile Communication System”, Proc. Vehicular Technology Conference 1996, Atlanta, pp. 1589-1593.

[RohlGrund_97] H. Rohling, R. Grundheid, “Performance comparison of different multiple access schemes for the downlink of an OFDM communication system”. IEEE 47th Vehicular Technology Conference, 4-7 May 1997, pp 1365-1369 vol.3.

[Shankar_01] P. Mohana Shankar, «Introduction to wireless systems» John Wiley and sons, INC, 2001.

[Sklar_01] Bernard Sklar , «Digital Communications : fundamentals and applications», Prentice Hall Communications, 2001.

[Shen&_03] Zukang Shen, Jeffrey G. Andrews, Brian L. Evans. “Optimal Power Allocation in Multiuser OFDM Systems”, *Proc. IEEE Global Comm. Conf*, Dec. 1-5, 2003, vol. 1, pp. 337-341, San Francisco, CA USA.

[Shin&_05] Seokjoo Shin, Seungjae Bahng, Insoo Koo and Kiseon Kim. “QoS-Oriented Packet Scheduling Schemes for

- Multimedia Traffics in OFDMA Systems”. Networking - ICN 2005 : 4th International Conference on Networking, Reunion Island, France, April 17-21, 2005, Proceedings, Part I.
- [ShinRyu_04] Seokjoo Shin and Byung-Han Ryu. “Packet loss fair scheduling scheme for real-time traffic in ofdma systems”. ETRI Journal, vol.26, no.5, Oct. 2004, pp.391-396.
- [ShinsukePrasad_03] H. Shinsuke, R. Prasad. “Multimedia Techniques for 4G mobile communications”. Artech House, 2003.
- [TsengChang_98] Yang-Chung Tseng and Cheng-Sang Chang. « PGPS servers with Time-Varying Capacities ». IEE Communications Letters, Vol.2, No 9, september 1998.
- [Viswanath_05] P. Viswanath, chapter 1: Opportunistic Communication: A System View, “Space-Time Wireless Systems: From Array Processing to MIMO Systems” Cambridge University Press, 2005
- [Wang&_02] Wang, S., Wang, Y., Lin, K., “ Integrating Priority with Share in the Priority-Based Weighted Fair Queueing Scheduler for Real-Time Networks,” Journal of Real Time Systems, Vol. 22, 2002, pp.119-149.
- [Wang&_04] Wang L ; Kwok Y-K ; Lau W-C ; Lau V.K.N. ” Efficient Packet Scheduling Using Channel Adaptive Fair Queueing in Distributed Mobile Computing Systems”. Mobile Networks and Applications, Volume 9, Number 4, August 2004, pp. 297-309.
- [Wang&_05] You-Chiun Wang, Yu-Chee Tseng, Wen-Tsuen Chen, Kun-Cheng Tsai. “MR-FQ : a fair scheduling algorithm for wireless networks with variable transmission rates”. 3rd International Conference on Information Technology : Research and Education, 2005. 27-30 June 2005. On page(s) : 250- 254
- [WangNiu_05] L.Wang, Z. Niu. «An efficient Rate and Power Allocation Algorithm for Multiuser OFDM Systems». 2005.
- [Wengerter&_05] Christian Wengerter, Jan Ohlhorst, Alexander Golitschek Edler von Elbwart. “Fairness and Throughput Analysis for Generalized Proportional Fair Frequency Scheduling in OFDMA” VTC spring 2005.
- [WimaxEval_07] WiMAX System Evaluation Methodology. V1.0 30/01/2007. Wimax Forum.
- [Wong&_99] Cheong Yui Wong, Roger S. Chen, Khaled Ben Letaief, Ross D. Murch. “Multiuser OFDM with Adaptive Subcarrier, Bit, and Power Allocation”, IEEE Journal on Selected Areas in Communications, vol. 17, no. 10, October 1999, pp. 1747 – 1758.
- [WongTsui&_99] C. Wong, C. Tsui, S. Chen, Khaled Ben Letaief. “A Real-time Sub-carrier Allocation Scheme for Multiple Access Downlink OFDM Transmission”, Vehicular Technology Conference, 1999. Fall. 50th pp. 1124-1128 vol.2.
- [Wong&_04] Ian C. Wong, Zukang Shen, Jeffrey G. Andrews, Brian L. Evans. “A Low Complexity Algorithm for proportional Resource Allocation in OFDMA Systems”, 2004.
- [Yaghoobi_04] Hassan Yaghoobi, « Scalable OFDMA Physical Layer in IEEE 802.16 WirelessMAN » Intel Technology Journal, Volume 8, Issue 3, 2004.
- [Yan&_03] L. Yan ; Z. Wenan ; S. Junde. “An adaptive subcarrier, bit and power allocation algorithm for multicell OFDM systems”. IEEE Electrical and Computer Engineering, CCECE 2003. Volume 3, 4-7 May 2003 Page(s) :1531 - 1534 vol.3
- [YihGerianotis_00] Chi-Hsiao Yih and Evaggelos Geraniotis. “Adaptive Modulation, Power Allocation and Control for OFDM Wireless Networks”. The 11th IEEE International Symposium on Personal, Indoor and Mobile Radio Communications, PIMRC 2000. Publication Date : 2000, Volume : 2, On page(s) : 809-813 vol.2
- [YinLiu_00] Hujun Yin, Hui Liu, “An Efficient multiuser Loading Algorithm for OFDM-based Broadband Wireless Systems”, GLOBECOM 2000 - IEEE Global Telecommunications Conference, no. 1, Nov 2000, pp. 103 – 107.
- [Yu&_06] Gunading Yu, Zhaoyang Zhang, Yan Chen, Ying Shi and Peiliang Qiu. “A novel Resource Allocation Algorithm for Real-time Services in Multiuser OFDM Systems”. Vehicular Technology Conference, VTC2006-Spring. 7-10 may 2006.
- [Yuan&_05] Yuan, Y ; Gu, D ; Arbaugh, W ; Zhang, J., "Achieving Fair Scheduling Over Error-Prone Channels in Multirate

WLANs", International Conference on Wireless Networks, Communications and Mobile Computing (*WIRELESSCOM*), Vol. 1, pp. 698-703, June 2005

[Zhang_90] L. Zhang, "Virtual Clock : a New Traffic Control Algorithm for Packet Switching Networks," Proceedings of ACM SIGCOMM Symposium on Communications Architectures and Protocols, Philadelphia , Sep.1990, pp. 19-29.

[Zhang_95] H. Zhang, "Service disciplines for guaranteed performance service in packet-switching networks," in *Proceedings of IEEE* 1995, vol. 83(10), pp. 1374-1399, October 1995.

[Zhang_04] G. Zhang, "Subcarrier and bit allocation for real-time services in multiuser OFDM systems" IEEE ICC 04 Paris France, June 20-24, 2004, pp 2985_2989.

[ZhangBennett_96] H. Zhang, J. C. R. Bennett, "W2FQ : Worst-Case Fair Weighted Fair Queueing," Proceedings of IEEE INFOCOM'96, 1996.

[ZhangLetaief_04] Y. Zhang ; Letaief, K.B. "Adaptive resource allocation and scheduling for multiuser packet-based OFDM networks". IEEE International Conference on Communications, Vol 5, 20-24 June 2004 Page(s) :2949 - 2953

[Zhen&_03] L. Zhen Z. Geqing, W. Weihua, S. Junde. "Improved algorithm of multiuser dynamic subcarrier allocation in OFDM system". International Conference on Communication Technology Proceedings, ICCT: 9-11 April 2003, page(s) : 1144- 1147 vol.2.