



Utilisation de la Radio Intelligente pour un Réseau Mobile à Faible Consommation d'Energie

Rémi Bonnefoi

► To cite this version:

Rémi Bonnefoi. Utilisation de la Radio Intelligente pour un Réseau Mobile à Faible Consommation d'Energie. Traitement du signal et de l'image [eess.SP]. CentraleSupelec, 2018. Français. NNT : . tel-02430551

HAL Id: tel-02430551

<https://hal.science/tel-02430551>

Submitted on 7 Jan 2020

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 DE DOCTORAT DE

CENTRALESUPELEC RENNES

COMUE UNIVERSITE BRETAGNE LOIRE

ECOLE DOCTORALE N° 601

*Mathématiques et Sciences et Technologies
de l'Information et de la Communication*

Spécialité : Télécommunications

Par

Rémi BONNEFOI

Utilisation de la Radio Intelligente pour un Réseau Mobile à Faible Consommation d'Energie

Thèse présentée et soutenue à Rennes, le 15 Octobre 2018

Unité de recherche : UMR 6164 - IETR (Equipe SCEE)

Thèse N° : 2018-05-TH

Rapporteurs avant soutenance :

Claire GOURSAUD Maître de conférences à INSA Lyon

Samson LASAULCE Directeur de recherche au CNRS

Composition du Jury :

Loutfi NUAYMI Professeur à l'IMT Atlantique

Président

Antonio DE DOMENICO Ingénieur de recherche au CEA

Estêvão FERNANDES Professeur à l'Université Fédérale
de Céara

Yves LOUËT Professeur à CentraleSupélec

Directeur de thèse

Christophe MOY Professeur à l'Université de
Rennes 1

Encadrant

Jacques PALICOT Professeur à CentraleSupélec

A Hortense et Lucile.

Remerciements

Je voudrais, tout d'abord, remercier mes deux encadrants Christophe Moy et Jacques Palicot qui m'ont guidé et accompagné pendant toute la durée de ce doctorat. Je remercie également très sincèrement l'ensemble de l'équipe SCEE dans laquelle j'ai effectué ma thèse. Je remercie l'ensemble des permanents de l'équipe, et en particulier Yves Louët mon directeur de thèse. J'ai aussi une pensée particulière pour Haïfa et Amor avec qui j'ai eu l'occasion de travailler ; sans oublier Carlos, Georgios et Pascal. Je tiens également à remercier l'ensemble des doctorants et stagiaires que j'ai eu l'occasion de croiser, je pense en particulier à ceux qui m'ont accompagné pendant une grande partie de ce doctorat, Muhammad et Rami qui m'ont accompagné pendant ces trois années, Quentin, Lilian avec qui j'ai eu la chance de travailler, Ali, Ali, Vincent, Vincent, Marwa, Esteban, Raphaël, Pierre-Samuel, Bastien, Nabil, Adrien, Eloïse, Majed et tous ceux que j'ai oublié.

Je remercie tout particulièrement Claire Goursaud et Samson Lasaulce qui ont accepté d'être rapporteurs de mon manuscrit de thèse. Je remercie également Estêvão R. Fernandes, Antonio de Domenico et Loutfi Nuaymi qui ont accepté de faire partie de mon jury de thèse.

Je remercie l'ensemble des étudiants que j'ai eu la chance d'encadrer et tout particulièrement Jihane, Salma, Théo et Clément qui m'ont supporté pendant un an. Je remercie également Aymeric Sanchez qui a modernisé le site internet de l'équipe ainsi que Pierre Bélis qui a réalisé un superbe stage au sein de SCEE.

Je remercie le laboratoire GTEL qui m'a accueilli pendant un mois au Brésil, en particulier, Tarcisio Maciel et Estêvão qui m'ont encadré pendant cette mobilité. Je remercie également l'ensemble des étudiants de l'UFC qui ont facilité mon séjour. Merci à Raul et Juliana.

Je tiens à remercier le personnel du campus de Rennes pour sa sympathie. Je remercie Jeannine Hardy et Karine Bernard qui ont facilité mes démarches dans la jungle administrative de CentraleSupélec. Je remercie également l'ensemble du service technique et en particulier Stéphane et Maud pour leur sympathie. Je remercie également le personnel de la Sodexo qui a égayé tous mes déjeuners. Merci encore à tout ce qui ont été présents pendant les 6 années passées dans cet établissement, Ophélie, Martine, Philippe, Gabriel et tous ceux que j'oublie.

J'ai appris beaucoup de choses pendant ces trois années de thèse. Par exemple, à jouer au foot. Et cela n'aurait pas été possible si je n'avais pas eu l'occasion de jouer avec Florian, Greg, Martin, Pape, ...

Je tiens à dédier quelques mots à ma famille qui, sans trop comprendre ce que je fais, m'a supporté pendant ces trois années. Merci infiniment Lucile, merci Hortense de m'avoir écouté lorsque je lisais ma thèse. Merci papa, maman, Anne, Simon et Jérôme. Merci Françoise et Dédé, et plus généralement, merci les Carnac d'Aveyron, de Lozère et mamie Huguette.

Enfin, voici une liste de personnes qui n'ont pas trop œuvré pour cette thèse mais que j'ai envie de citer, merci : Philippe, Gwen, Charley, Inès, Gurban, Johnny, Micka, Tex et Ella, Julien, Pauline et Mila, Emilien, Thomas, les Fahrenheit, Mathilde et Vincent, Jessy.

Table des matières

Introduction	1
1 Cadre de l'étude	7
1.1 La radio intelligente	8
1.1.1 Présentation	8
1.1.2 Le cycle intelligent	8
1.1.3 L'accès opportuniste au spectre	9
1.2 L'architecture HDCRAM	10
1.3 Le réseau mobile considéré	12
1.3.1 Introduction	13
1.3.2 Description HDCRAM du réseau mobile efficace en énergie	14
1.3.3 Position dans le réseau électrique	15
1.3.4 Description complète	17
1.4 Études menées dans cette thèse	18
I Economie d'énergie dans les réseaux cellulaires	21
2 Allocation de ressources et de puissance pour des réseaux mobiles sobres en énergie	23
2.1 Les réseaux mobiles	24
2.2 L'accès multiple	25
2.3 Formulation du problème d'allocation de ressources et de puissance en OFDMA	27
2.3.1 Définitions et vocabulaire	27
2.3.2 Maximisation de la capacité	28
2.3.3 Prise en compte de la consommation énergétique	30
2.4 Résolution du problème d'optimisation	31

2.4.1	Allocation de ressources	32
2.4.2	Allocation de puissance	33
2.5	Évolutions récentes des réseaux cellulaires	36
2.6	Consommation d'énergie des réseaux mobiles	37
2.7	Mise en veille des stations de base et transmission discontinue	38
2.7.1	Mise en veille prolongée	38
2.7.2	Transmission discontinue	39
2.8	Allocation de puissance et de ressources avec du Cell DTx	40
2.8.1	Allocation de puissance à une seule station de base	40
2.8.2	Cell DTx dans des réseaux à plusieurs cellules	43
2.9	Conclusions	44
3	Allocation de puissance et transmission discontinue en TDMA	45
3.1	Problème d'optimisation et reformulation	46
3.2	Résolution du problème	48
3.2.1	Méthodologie	48
3.2.2	Première étape : gestion des contraintes sur le temps de service minimum	49
3.2.3	Seconde étape : gestion de la contrainte sur le temps total	52
3.2.4	Algorithme final	56
3.3	Résultats numériques	57
3.3.1	Pour une macro-cellule	57
3.3.2	Pour une femto-cellule	61
3.4	Dans un réseau composé de plusieurs cellules	62
3.4.1	Introduction	62
3.4.2	Modèle d'interférence	62
3.4.3	Paramètres de simulation	63
3.4.4	Comportement de notre solution dans un réseau à plusieurs cellules .	64
3.4.5	Stratégie optimale irréalisable	66
3.4.6	Evaluation de la méthode proposée Algorithme 3	69
3.4.7	Réduction des interférences pour améliorer les performances	70
3.4.8	Puissance d'émission optimale	71
3.5	Conclusions	75

4	Allocation de ressources et de puissance avec de la transmission discontinue en OFDMA	77
4.1	Introduction	77
4.2	Dans des canaux plats	78
4.2.1	Formulation du problème	78
4.2.2	Réécriture du problème	79
4.3	Dans un canal sélectif en fréquence	81
4.3.1	Position du problème	81
4.3.2	Allocation de puissance	82
4.3.3	Allocation de ressources	87
4.3.4	Résultats numériques	91
4.4	Conclusions	93
II	Amélioration des communications du réseau électrique intelligent	95
5	Communications dans le réseau électrique intelligent	97
5.1	Introduction	97
5.2	Le réseau électrique intelligent	98
5.3	Les communications du réseau électrique intelligent	99
5.4	Rappels sur l'architecture HDCRAM	100
5.5	HDCRAM pour la gestion de la production et de la consommation	102
5.6	Description HDCRAM des mécanismes de gestion du réseau	105
5.6.1	HDCRAM pour la gestion du réseau de transport	106
5.6.2	HDCRAM pour la gestion du réseau de distribution	107
5.7	Conclusion	108
6	Algorithmes de bandits pour les réseaux d'objets connectés	109
6.1	Introduction	110
6.2	Les <i>Low Power Wide Area Networks</i>	110
6.3	Probabilité de collision dans les réseaux LPWAN	111
6.3.1	Réseau ALOHA non-slotté avec le mécanisme d'acquittement du standard LoRaWAN	113
6.3.2	Réseau ALOHA slotté	116
6.4	Latence dans les réseaux LPWAN	119

6.4.1	Sélection aléatoire du canal	121
6.4.2	Transmission uniquement dans le meilleur canal	122
6.5	Algorithmes de bandits multibras	123
6.5.1	L'approche fréquentiste : l'algorithme UCB	124
6.5.2	L'approche bayésienne : l'algorithme Thompson-Sampling	125
6.6	Application aux réseaux d'objets connectés	126
6.7	Validation expérimentale sur Plateforme USRP	128
6.8	Évaluation avec un grand nombre d'objets dynamiques	131
6.8.1	Modèle	132
6.8.2	Stratégies de référence	132
6.8.3	Évaluation des performances	134
6.9	Conclusion	137
Conclusion et perspectives		139
A Quelques preuves relatives au chapitre 3		145
A.1	Dérivation de l'équation (3.37)	145
A.2	Convexité du problème défini par l'équation (3.40)	147
B Quelques preuves relatives au Chapitre 4		151
B.1	Non-convexité du problème défini équation (4.5)	151
B.2	Allocation de puissance optimale en OFDMA avec du Cell DTx	153
C Quelques preuves relatives au Chapitre 6		155
C.1	Calcul des probabilités de collision dans un réseau ALOHA	155
C.2	Une approximation pour la probabilité de collision dans un réseau ALOHA slotté	163
Liste des publications		165
Table des figures		167
Liste des tableaux		173
Bibliographie		175

Introduction Générale

Contexte de l'étude

L'homme a un impact de plus en plus visible sur son environnement. Mais parmi les conséquences invisibles de l'activité humaine, l'augmentation de la concentration en dioxyde de carbone dans l'atmosphère est l'une des plus préoccupantes. En effet, cette évolution a pour effet une augmentation de la température qui pourrait avoir des répercussions graves au niveau climatique engendrant de nombreuses répercussions écologiques et économiques, donc sur les conditions de vie humaines, animales et végétales. C'est pourquoi, il est aujourd'hui important de repenser la façon dont nous produisons et consommons de l'énergie. Autrement dit, il faut, à la fois, réduire notre consommation d'énergie et augmenter la part d'énergie renouvelable dans la production.

Dans cette transformation, les Technologies Numériques de l'Information et de la Communication (TNIC) jouent un rôle côté production et côté consommation. Si on s'intéresse à la consommation d'énergie, les TNIC sont responsables de plus de 2% des émissions de gaz à effet de serre [1]. Même si la majorité de ces émissions sont à imputer aux datacenters, les points d'accès radio (tels que les stations de base du réseau mobile, point d'accès Wifi ou autres) ont un rôle non négligeable. On estime aujourd'hui qu'environ 50% des émissions de gaz à effet de serre provoqués par les réseaux mobiles est due aux stations de base [2]. De plus, avec la densification des réseaux mobiles et le déploiement de nouveaux réseaux, par exemple pour l'Internet des Objets (IdO), le nombre de points d'accès va augmenter. C'est pourquoi, il est aujourd'hui nécessaire d'améliorer la gestion de l'énergie des points d'accès du réseau mobile. C'est-à-dire qu'il faut réduire leur consommation d'énergie et de les alimenter par des sources d'énergie renouvelables.

Mais le rôle des réseaux de communication sans fil dans la transformation écologique ne se limite pas à la réduction de sa propre consommation. L'insertion de sources d'énergie renouvelables intermittentes, telles que l'énergie solaire ou éolienne, rend, aujourd'hui, plus difficile la gestion du réseau électrique et en particulier l'adaptation de la production électrique à la consommation. Ceci augmente le risque de panne du réseau. C'est pourquoi, les opérateurs

du réseau électrique sont en train de le transformer pour que celui-ci puisse continuer à fonctionner alors que la part de renouvelable dans la production d'électricité augmente. Ceci transforme le réseau électrique en réseau électrique intelligent (*Smart Grid*). Dans un réseau électrique intelligent, l'opérateur du réseau électrique doit avoir une estimation fine et quasiment en temps réel de la production et de la consommation électrique. Pour cela, il collecte de l'information venant de capteurs disséminés tout au long du réseau électrique. Cette collecte d'information ne peut se faire sans qu'un réseau de communication ne transmette l'information mesurée par les capteurs à l'opérateur du réseau électrique. C'est pourquoi, les réseaux de communication ont un rôle clé dans l'insertion d'énergie renouvelable dans le réseau électrique et, par extension, dans la transformation de nos sources d'électricité.

Les principes de la radio intelligente peuvent être utilisés pour que les TNIC jouent pleinement leur rôle dans la transition écologique. Une radio intelligente est un appareil capable de se reconfigurer dynamiquement afin d'adapter ses paramètres à son environnement [3]. La radio intelligente peut servir à la fois pour la réduction de la consommation du réseau de communication mobile et pour les communications du réseau électrique intelligent. Dans un premier temps, rendre les équipements radio flexibles et capables de s'adapter à leur environnement permet de réduire la consommation des réseaux mobiles [4]. Avec cette flexibilité, les stations de base du réseau peuvent adapter leurs paramètres de transmission dans le but de réduire leur consommation. Par exemple, le réseau peut adapter son fonctionnement au nombre d'utilisateurs actifs. Il peut fonctionner à plein régime pendant les heures d'utilisations pleines et avoir un fonctionnement moins gourmand en énergie pendant les heures creuses, par exemple la nuit [5].

Les principes de la radio intelligente peuvent aussi s'appliquer aux communications du réseau électrique intelligent. En effet, il est extrêmement probable que la majorité des communications du réseau électrique intelligent se fassent par le biais de réseaux d'objets connectés [6]. Ces réseaux seront partagés par un grand nombre d'appareils ayant des contraintes différentes en débit, latence ou consommation d'énergie. Ce grand nombre d'objets risque de mener à une congestion du spectre fréquentiel. Les principes de la radio intelligente peuvent s'appliquer dans ce contexte. En effet, la congestion du spectre peut être évitée en rendant plus flexible le comportement des objets [7].

De plus, si les stations de base du réseau cellulaire sont associées à des sources d'énergie renouvelables (éoliennes ou panneaux solaires), comme montré sur la Figure 1, le réseau mobile et le réseau de communications du réseau électrique intelligent sont deux réseaux qui composent un même système. Dans ce cas, le réseau mobile produit et consomme de l'énergie qu'il peut acheter ou vendre à l'opérateur du réseau électrique. Le réseau mobile va donc vouloir réduire sa consommation d'énergie pour diminuer son coût de fonctionnement. De plus, en tant que producteur et consommateur d'énergie, il doit échanger de l'information, avec le réseau électrique, à propos de sa production et de sa consommation d'énergie. L'information envoyée

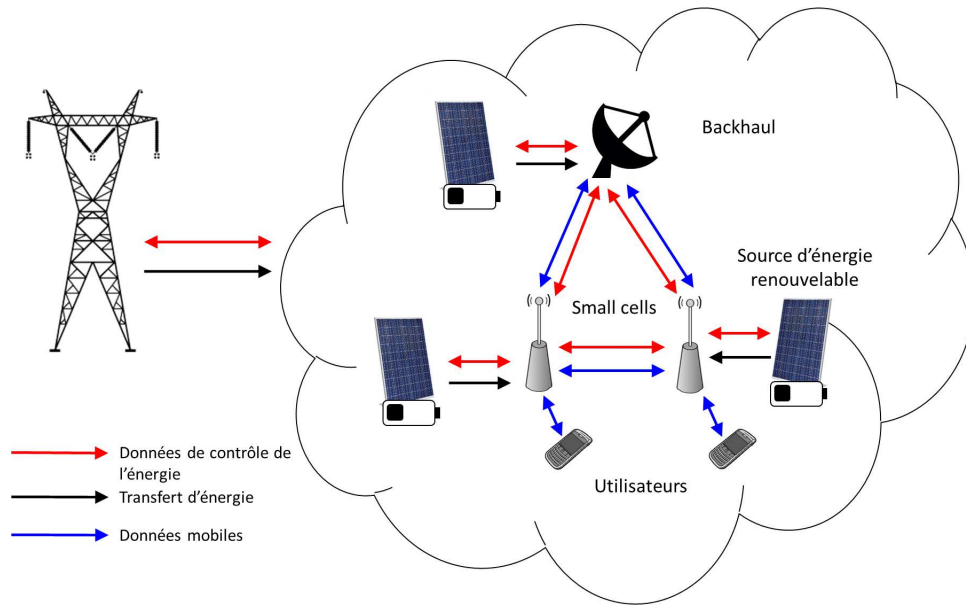


Figure 1 – Dans le système étudié nous avons différents réseaux de communication : un réseau de communication mobile et un réseau de communication pour le réseau électrique intelligent.

par les différents producteurs et consommateurs permet à l'opérateur du réseau électrique d'estimer la production et la consommation d'énergie dans celui-ci. Il peut alors envoyer une consigne aux producteurs et consommateurs afin que ceux-ci adaptent leur production et leur consommation. En résumé, réseaux électriques et réseaux de communication sont de plus en plus intriqués, et c'est dans ce contexte que se situe cette thèse.

Objectifs et contributions

L'objectif de cette thèse est de réduire l'empreinte carbone du réseau mobile présenté sur la Figure 1. Pour cela, nous appliquons les principes de la radio intelligente afin de :

1. réduire la consommation d'énergie des stations de base du réseau mobile,
2. améliorer les performances du réseau électrique, et en particulier la latence du réseau de communication utilisé pour les communications entre les producteurs/consommateurs d'énergie électrique et l'opérateur du réseau électrique.

La transmission discontinue [8] peut être utilisée pour réduire la consommation d'énergie des stations de base du réseau mobile. Avec cette solution, la station de base peut être mise en veille pendant des temps très courts. L'utilisation de cette technologie influe sur la façon dont les ressources radio et la puissance d'émission sont réparties entre les utilisateurs. Le premier objectif de cette thèse est de proposer des algorithmes pour l'allocation de ressources et de puissance avec la transmission discontinue.

Le second objectif de cette thèse est de proposer des algorithmes pour améliorer les performances des communications du réseau électrique intelligent. Pour cela, il faut identifier les protocoles de communication qui peuvent être utilisés pour le transfert d'information dans un réseau électrique. Ensuite, il faut identifier et étudier les indicateurs de performance de ces protocoles, et enfin proposer des solutions pour améliorer ces indicateurs.

Organisation du document

Cette thèse est organisée en deux parties et six chapitres.

Le **Chapitre 1** décrit le cadre de l'étude et les objectifs de cette thèse. Dans ce chapitre, nous commençons par introduire les principes de la radio intelligente et plus généralement la prise de décision dans un système intelligent. Ensuite, nous décrivons l'architecture *Hierarchical and Distributed Cognitive Radio Architecture Management* (HDCRAM) qui est un outil de description que nous utiliserons pour présenter le réseau mobile sobre en énergie qui sera le sujet de cette thèse. Nous verrons, dans ce premier chapitre, que pour réduire l'empreinte carbone de ce réseau, il est nécessaire de proposer à la fois des mécanismes pour réduire la consommation d'électricité du réseau et des mécanismes pour améliorer les communications du réseau électrique.

Première partie

L'objectif de la **Première Partie** est de présenter des algorithmes d'allocation de puissance efficaces lorsque la transmission discontinue est utilisée. Nous montrons que ces algorithmes permettent de réduire significativement la consommation d'énergie des stations de base.

Dans le **Chapitre 2**, nous présentons le problème d'allocation de ressources et de puissance dans un réseau mobile sobre en énergie. Après avoir présenté brièvement les principes de fonctionnement d'un réseau mobile, nous présenterons un état de l'art du problème d'allocation de ressources et de puissance dans les réseaux mobiles. En particulier, nous nous intéresserons au problème d'allocation de ressources et de puissance pour la réduction de la consommation d'énergie d'une station de base. Cette présentation nous permettra d'introduire les conditions de Karush-Kuhn et Tucker (KKT) qui seront utilisées dans les Chapitres 3 et 4. Ensuite, nous présenterons les mécanismes de mise en veille qui permettent de réduire davantage l'énergie consommée par les stations de base. Nous verrons que l'utilisation de ces mécanismes de mise en veille et en particulier l'utilisation la transmission discontinue change l'allocation de puissance. C'est pourquoi, dans les chapitres suivants, nous nous efforcerons de proposer

des algorithmes qui permettent de minimiser l'énergie consommée par une station de base lorsque de la transmission discontinue est utilisée.

Dans le **Chapitre 3** nous proposons un algorithme d'allocation de puissance pour minimiser la consommation d'une station de base qui utilise du *Time Division Multiple Access* (TDMA) et de la transmission discontinue. Nous commençons par formuler le problème d'allocation de puissance. Ensuite, nous reformulons ce problème sous-forme convexe. Nous proposons, par la suite, un algorithme efficace qui permet de calculer l'allocation de puissance optimale. Ensuite, nous évaluons les performances de cette solution dans un réseau composé de plusieurs stations de base. Nous verrons alors que la solution proposée peut être améliorée en réduisant la puissance d'émission maximale de la station de base. Cette résolution dans un scénario TDMA nous permettra de bien comprendre le problème d'allocation de ressources et de puissance avec de la transmission discontinue. Cependant, dans les réseaux actuels, de l'*Orthogonal Frequency Division Multiple Access* (OFDMA) est utilisée pour l'accès multiple.

Par conséquent, dans le **Chapitre 4** nous traitons le problème d'allocation de ressources et de puissance en OFDMA lorsque de la transmission discontinue est utilisée. Nous commençons par supposer que le canal de tous les utilisateurs est plat et nous montrerons que dans ce cas là, on peut utiliser les résultats et dérivations menées dans le Chapitre 3. Ensuite, nous nous passerons de cette hypothèse, et nous traiterons le problème d'allocation de ressources et de puissance en OFDMA lorsque le canal des utilisateurs est à évanouissement. Nous formulerons d'abord le problème d'allocation de puissance pour montrer qu'il est convexe et pour ainsi pouvoir proposer un algorithme pour le résoudre de manière optimale. Une fois le problème d'allocation de puissance résolu, nous étudierons le problème d'allocation de ressources.

Deuxième partie

Nous avons vu dans cette introduction que pour réduire l'empreinte carbone des réseaux mobiles il n'est pas suffisant de réduire leur consommation d'énergie, il est aussi possible d'améliorer la façon dont l'électricité est produite, délivrée et consommée par les différents consommateurs et en particulier par le réseau mobile. C'est pourquoi, dans la **Deuxième Partie**, nous proposons des solutions pour améliorer les réseaux de communication qui sont utilisés pour la gestion d'un réseau électrique intelligent.

Dans le **Chapitre 5**, nous présentons les différents mécanismes qui gèrent un réseau électrique intelligent. Pour cela, nous utilisons l'architecture HDCRAM qui nous permet de les présenter de manière unifiée. Cette présentation nous permettra d'identifier les technologies et standards de communication sans fil qui peuvent être employés pour le réseau électrique. Cet état de l'art nous permettra de nous rendre compte que les *Low Power Wide Area Networks* (LPWAN) sont des bons candidats pour les communications à longues portées utilisées pour

la gestion de la production et de la consommation d'électricité. C'est pourquoi ces réseaux seront le sujet des études menées dans le Chapitre 6.

Dans le **Chapitre 6**, nous proposons d'améliorer les communications du réseau électrique intelligent qui se font via des LPWAN. Nous commençons par décrire le fonctionnement des LPWAN avant d'expliquer que la probabilité de collision est un indicateur de performance dans ces réseaux. Nous verrons que cette probabilité a un impact sur d'autres mesures telles que la latence. Ensuite, nous montrerons que des algorithmes de bandits multibras peuvent être utilisés pour réduire les collisions et ainsi améliorer les différents indicateurs de performance du réseau.

Chapitre 1

Cadre de l'étude

Sommaire

1.1	La radio intelligente	8
1.1.1	Présentation	8
1.1.2	Le cycle intelligent	8
1.1.3	L'accès opportuniste au spectre	9
1.2	L'architecture HDCRAM	10
1.3	Le réseau mobile considéré	12
1.3.1	Introduction	13
1.3.2	Description HDCRAM du réseau mobile efficace en énergie	14
1.3.3	Position dans le réseau électrique	15
1.3.4	Description complète	17
1.4	Études menées dans cette thèse	18

Dans ce chapitre, nous introduisons la notion de radio intelligente. Ensuite, nous décrivons l'architecture HDCRAM. C'est une architecture hiérarchique et distribuée qui a été proposée pour la gestion d'équipements de radio intelligente [9]. Cet outil de modélisation de haut niveau est ensuite utilisé pour décrire la prise de décision dans un réseau mobile dans lequel les stations de base sont associées à des sources d'énergie renouvelables. L'étude de ce réseau sera l'objet de la suite de ce manuscrit. Nous verrons qu'il peut être à la fois, vu comme un système complet mais aussi comme un élément dans des mécanismes plus larges dans la gestion du réseau électrique intelligent. Cette description globale nous permet de positionner plus précisément les études réalisées pendant cette thèse.

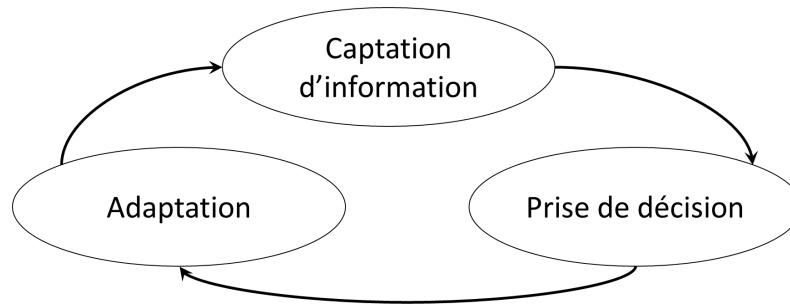


Figure 1.1 – Version simplifiée du cycle intelligent.

1.1 La radio intelligente

1.1.1 Présentation

L'objectif d'un système de radiocommunication est de transférer de l'information entre deux équipements distants. Cette communication va se faire dans un environnement donné en utilisant des paramètres spécifiques. Par exemple, si on considère une communication entre un téléphone mobile et une station de base, cette communication peut se faire en milieu rural, en milieu urbain, à l'intérieur d'un bâtiment, à l'arrêt ou pendant un déplacement ... Elle va donc utiliser différents paramètres qui vont dépendre du type de communication et de cet environnement. Parmi ces paramètres, on va retrouver la puissance de transmission, la largeur de bande, ou encore la modulation.

Dans la radio intelligente [3], au lieu d'utiliser des paramètres de communication fixes, ces équipements radio, qu'ils soient émetteurs ou récepteurs, ont la possibilité de se reconfigurer dans le but de s'adapter à un changement d'environnement ou à un changement des paramètres de communication. Le passage d'une communication inflexible à une communication dans laquelle les paramètres peuvent être modifiés est facilitée par la numérisation d'une grande partie des fonctions radios. C'est à dire le passage à une radio logicielle [10] dans laquelle la conversion analogique numérique doit se faire au plus près de l'antenne.

Les performances des Convertisseurs Analogique-Numérique (CAN) et Numérique-Analogique (CNA) ne permettent pas de numériser une bande passante infinie, néanmoins, la radio intelligente peut être utilisée pour de nombreuses applications et dans différents contextes où les contraintes sont un peu restreintes.

1.1.2 Le cycle intelligent

Dans une radio intelligente, les appareils radios doivent prendre des décisions. Comme dans n'importe quel système intelligent, cette prise de décision se fait en suivant le cycle intelligent [3] dont une version simplifiée est décrite sur la Figure 1.1. Ce cycle intelligent peut

être décomposé en trois étapes principales : la captation d'information, la prise de décision et l'adaptation.

Le rôle de chacune de ces étapes est le suivant :

- La captation de l'environnement consiste à récolter de l'information.
- Pendant la prise de décision, l'appareil décide quelle action il va faire. Cette décision s'appuie sur les informations relevées par les capteurs lors des précédentes mesures. Cette étape de prise de décision est au cœur des études menées pendant cette thèse. En particulier, on la retrouvera dans le Chapitre 3 pour la minimisation de la puissance consommée par les stations de base et dans le Chapitre 6 pour la sélection des bandes de fréquence dans les réseaux d'objets connectés.
- Enfin on a l'étape d'adaptation ou de reconfiguration. Dans cette étape, l'appareil applique la décision précédemment prise.

Pour mieux comprendre le principe du cycle intelligent dans le domaine des communications sans fil, nous donnons un exemple d'utilisation pour l'Accès Opportuniste au Spectre (AOS) [11] qui est l'application la plus répandue de la radio intelligente.

1.1.3 L'accès opportuniste au spectre

L'ensemble des fréquences pouvant être utilisées pour les communications sont aujourd'hui attribuées. En effet, en France, l'Agence Nationale des Fréquences (ANFR) qui définit les applications des différentes portions du spectre a planifié l'usage de toutes les fréquences exploitables avec les technologies actuelles. C'est-à-dire, les portions du spectre entre 8.3 kHz et 275 GHz [12].

Si un nouveau service, un nouvel usage ou un nouveau standard de communication apparaît, on va devoir l'insérer dans ce spectre. Pour cela, on peut envisager trois possibilités. Il peut être inséré dans les bandes non-licenciées qui sont partagées par divers services, mais, dans ce cas là, il devra cohabiter avec d'autres usages, par exemple avec le Wifi dans la bande de fréquence autour de 2.4 GHz ; ou pourra être limité par la régulation qui régit ces bandes. Par exemple, l'usage des bandes de fréquence autour de 433 et 868 MHz est limité par le rapport cyclique des communications. Une seconde solution consiste à lui réserver une bande de fréquence au détriment d'un autre service qui n'est plus d'actualité. Par exemple, les bandes allouées pour la Télévision Numérique Terrestre (TNT) ont permis de libérer de la place dans l'ancienne bande de la TV analogique autour de 800 MHz pour laisser place à des nouveaux standards de téléphonie mobile de 4G.

La dernière des possibilités consiste à utiliser un accès dynamique au spectre dont l'un des exemples est l'AOS. Les bandes réservées n'étant pas toujours utilisées, à un instant donné, on peut avoir une bande réservée par un standard mais laissée libre de façon intermittente

dans le temps, mais aussi dans l'espace. C'est ce que l'on appelle un trou dans le spectre. On peut donc insérer un nouveau service qui n'utilisera que ces trous. L'insertion de ce nouveau service doit se faire sans perturber l'utilisateur principal du service qui a payé pour une utilisation de la bande. Par exemple, cela peut se faire dans les réseaux de communication de la TNT, dans lequel, deux zones de couverture voisines n'utilisent pas les mêmes bandes de fréquence pour limiter les interférences. On a donc des bandes de fréquence non-utilisées dans chaque zone. Celles-ci peuvent être utilisés pour d'autres communications [13] tant que les interférences générées sont supportées par les récepteurs TV. Il s'agit ici d'un cas statique, les bandes libres à un endroit donné le restent tout le temps.

Dans l'AOS, on appelle utilisateur primaire celui qui est licencié dans une bande de fréquence (par exemple, le réseau de télévision) et utilisateur secondaire celui qui va utiliser une partie du spectre lorsque l'utilisateur primaire ne le fait pas. L'usage du spectre par l'utilisateur secondaire doit se faire sans perturber les communications de l'utilisateur primaire. Si, à un instant donné et à un endroit donné, un utilisateur secondaire veut communiquer à une fréquence donnée, il va donc devoir :

- Évaluer la présence ou l'absence de l'utilisateur primaire. Cette opération passe par l'utilisation de *spectrum sensing* [14].
- Décider s'il communique ou choisit d'aller ou non explorer d'autres fréquences.
- En fonction de cette décision, il va changer sa configuration en modifiant les paramètres de communication et en particulier sa fréquence de communication.

Les trois étapes de prise de décision décrites ici sont exactement les trois étapes qui composent le cycle intelligent (captation, décision, reconfiguration).

1.2 L'architecture HDCRAM

Le mécanisme simplifié en trois étapes du cycle intelligent est inhérent à n'importe quelle prise de décision. Mais la modélisation de la prise de décision par ce cycle intelligent n'est plus suffisante dans des systèmes plus complexes, et par exemple dans un appareil de radio intelligente. Dans ce cas là, on a des décisions qui peuvent impacter des fonctions différentes du système, et d'autres qui peuvent avoir un impact à plus grande échelle. On a alors une prise de décision hiérarchique et distribuée. L'objectif de l'architecture *Hierarchical and Distributed Cognitive Radio Architecture Management* (HDCRAM) [9] est d'organiser cette prise de décision dans des systèmes électroniques et en particulier pour la radio intelligente, qui est le contexte dans lequel elle a été initialement développée.

L'architecture HDCRAM est un outil de modélisation haut niveau qui permet de structurer les décisions dans un système complexe. En d'autres termes, c'est un ensemble de règles de

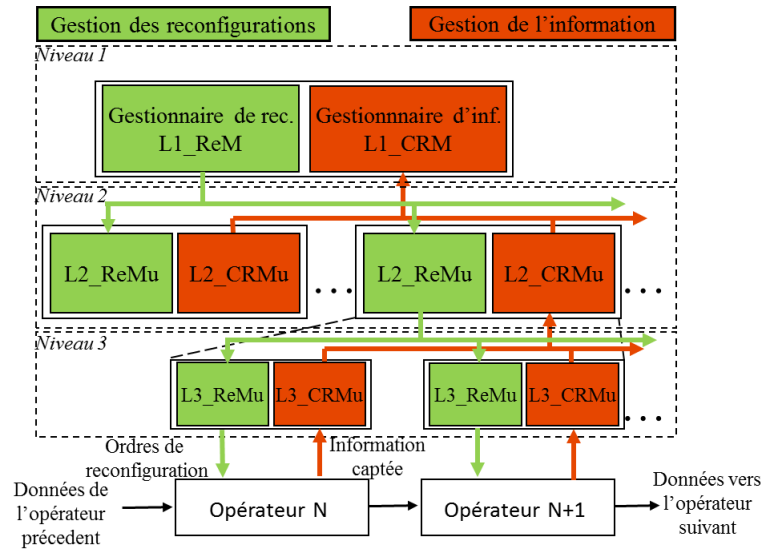


Figure 1.2 – L'architecture hiérarchique HDCRAM avec ses deux branches de gestion de la reconfiguration et de l'information.

modélisation qui permettent d'organiser la prise de décision lorsque celle-ci concerne différents éléments du système [15]. HDCRAM est présentée sur la Figure 1.2. Cette architecture à trois niveaux se décompose en deux parties distinctes. L'une pour le traitement de l'information et la prise de décision (en rouge sur la Figure 1.2) et la seconde pour les ordres de reconfiguration (en vert sur cette figure).

La gestion de l'information relevée par les capteurs se fait par le biais d'unités de traitement ou *Cognitive Radio Management units* (CRMu). Lorsqu'une de ces unités reçoit de l'information captée, elle va l'utiliser afin de prendre des décisions. Elle peut alors ordonner une reconfiguration d'un ou plusieurs opérateurs, ou encore, décider de transmettre une partie de l'information reçue à d'autres unités de traitement.

Chacune des unités de traitement est associée à une unité de reconfiguration ou *Reconfiguration Management unit* (ReMu) dont le rôle est d'interpréter (d'après des règles pré-définies) et d'appliquer les ordres envoyés par les unités de traitement ou par les unités de reconfiguration d'un niveau hiérarchique supérieur.

HDCRAM est composée de trois niveaux hiérarchiques. Cette distribution de la prise de décision permet de gérer les contraintes temps réel tout en réduisant la complexité. Au niveau 3, on peut avoir un grand nombre d'unités de gestion (un couple CRMu-ReMu par opérateur) leur rôle est de traiter l'information captée dans un temps court et de prendre des décisions simples en temps réel. Elles peuvent choisir de transmettre une partie de l'information envoyée par le capteur au niveau hiérarchique supérieur. Celui-ci, (le niveau 2) pourra utiliser cette information pour prendre une décision et la faire appliquer ou décidera d'informer le niveau

1. Lorsque la décision est prise au niveau 2, elle peut impliquer une sous-partie du système et donc plusieurs opérateurs. Si une décision est prise par la L2_CRMu, elle va la transmettre à la L2_ReMu correspondante. Celle-ci va transmettre l'ordre de reconfiguration aux L3_ReMu concernées qui vont appliquer le changement aux opérateurs qu'elles contrôlent. Finalement, au niveau 1, on a une unique unité de gestion. c'est le contrôleur central du système. Son rôle est de prendre des décisions complexes qui concernent l'ensemble du système.

En résumé, le niveau 3 est utilisé pour prendre des décisions rapides et n'impliquant qu'un seul opérateur. Par exemple, si on suppose un détecteur d'énergie dans un canal, la L3_CRMu associée peut ajuster le seuil de détection sans avoir à informer les autres opérateurs. Nous verrons, dans le paragraphe suivant, qu'un changement de modulation peut-être appliqué au niveau 2 de l'architecture. Le niveau 1 va prendre des décisions plus rares et ayant un impact plus global. Par exemple, si l'équipement radio veut changer de standard de communication, une grande partie des opérateurs radio vont être impliqués et le niveau 1 sera le seul capable de prendre une telle décision.

Afin de clarifier les explications présentées ci-dessus, nous donnons un exemple simple dans lequel on adapte la modulation au Rapport Signal à Bruit (RSB). Cet exemple est illustré sur la Figure 1.3. Sur celle-ci, nous n'avons représenté que les niveaux 2 et 3 de l'architecture. Supposons qu'une modulation est utilisée par l'équipement radio pour communiquer (par exemple une *Quadrature Phase-Shift Keying* (QPSK)). Le RSB est mesuré par le capteur de RSB. La L3_CRMu associé à ce capteur transmet cette information à l'unité de niveau 2 en charge de la modulation. Si le RSB a changé, et est, par exemple, passé au-dessus d'un seuil pré-défini, la L2_CRMu prend la décision de changer la modulation et d'utiliser une 16-*Quadrature Amplitude Modulation* (16-QAM) afin d'augmenter le débit. Cette décision est transmise à la L2_ReMu qui la transmet à la L3_ReMu en charge du mapping. Cette dernière applique la reconfiguration et l'opérateur est modifiée.

L'architecture HDCRAM est un métamodèle. C'est à dire qu'on peut l'utiliser pour obtenir un squelette de code. Elle a été proposée pour la gestion de la reconfiguration partielle de FPGA [16] et a aussi permis de structurer la prise de décision dans un démonstrateur pour l'AOS [17].

Tout comme le cycle intelligent, HDCRAM peut être utilisé pour la modélisation de nombreux systèmes et en particulier en dehors du contexte de la radio intelligente. Par exemple, elle a été proposée pour la gestion du réseau électrique intelligent [18].

1.3 Le réseau mobile considéré

Dans cette section, nous décrivons le réseau de communication mobile qui sera étudié dans la suite de cette thèse. Après l'avoir brièvement décrit, nous montrerons que HDCRAM

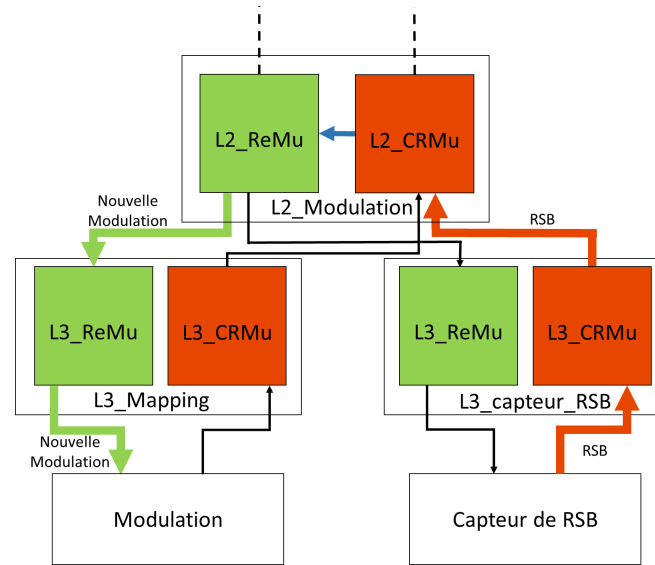


Figure 1.3 – Application de l’architecture HDCRAM pour l’adaptation de la modulation.

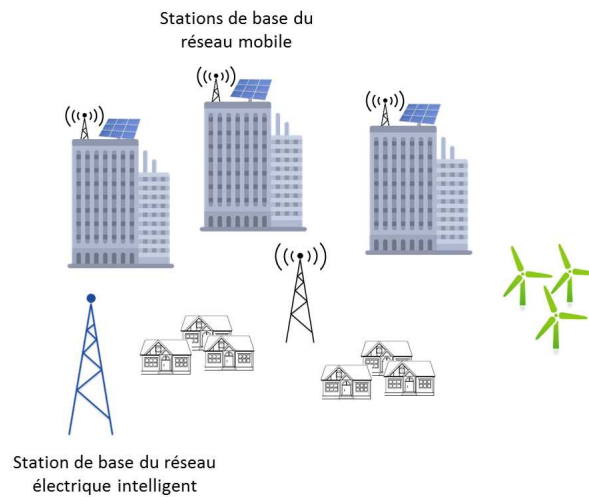


Figure 1.4 – Le système étudié dans lequel les stations de base sont associées à des sources d’énergie renouvelables et intermittentes.

peut être utilisée pour décrire la constitution de ce réseau. Cette description de l’architecture du système va ensuite être utilisée pour détailler les différents aspects étudiés pendant cette thèse.

1.3.1 Introduction

On considère le système de la Figure 1.4. Cette figure représente une aire urbaine dans laquelle nous avons un réseau mobile composé d’une station de base principale ayant une grande couverture (macro-cellule) et de petites stations de base ayant une plus petite

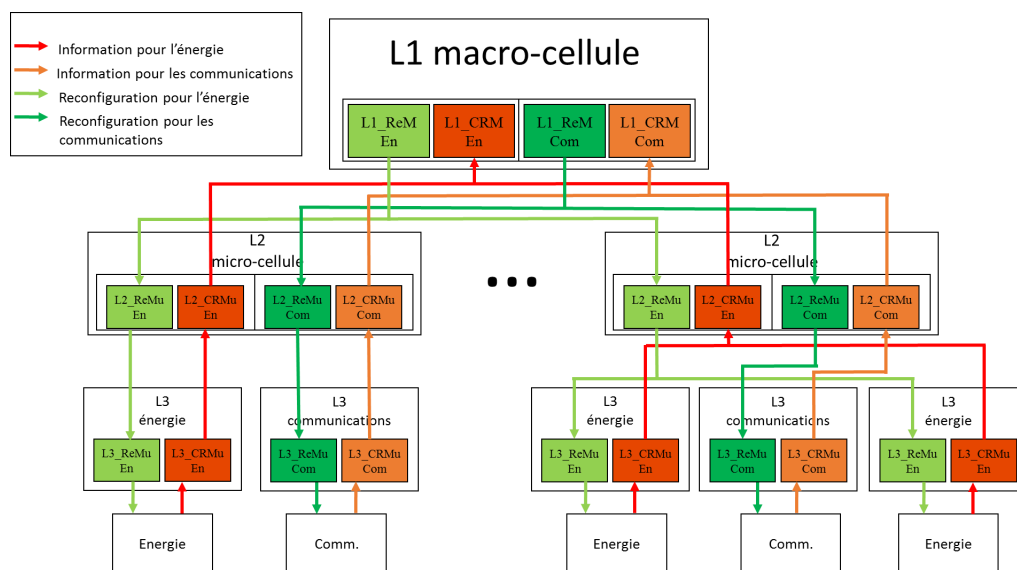


Figure 1.5 – Description HDCRAM du réseau mobile étudié. Dans cette description, on utilise deux HDCRAM. La première pour l'énergie et la seconde pour les communications.

couverture. Chacune d'entre elles peut être alimentée par des sources d'énergie renouvelables et reliée à une batterie. Par conséquent, ce système "réseau mobile" est à la fois producteur et consommateur d'énergie. Dans un réseau mobile de ce type, la gestion de l'énergie devient aussi importante que la gestion des communications mobiles. En d'autres termes, la gestion de l'énergie peut se faire en parallèle de la gestion des communications mobiles par le biais des mécanismes de communication du système, lui-même.

On s'appuie sur HDCRAM pour modéliser les architectures des deux systèmes de gestion. On a, par conséquent, une gestion divisée en deux parties : la première pour les communications mobiles et la seconde pour gérer l'énergie produite et consommée. Ces deux gestions ne sont pas forcément assurées par le même gestionnaire mais sont tout de même liées. La description HDCRAM de notre réseau mobile ne doit pas se faire avec une seule architecture mais via deux architectures HDCRAM interconnectées.

1.3.2 Description HDCRAM du réseau mobile efficace en énergie

La description HDCRAM du réseau mobile étudié est présentée sur la Figure 1.5. Elle comprend deux architectures HDCRAM en parallèle. La première sert pour la gestion des communications et la seconde pour la gestion de l'énergie du système.

Dans la description présentée, chacun des opérateur du réseau est relié à une unité de gestion de niveau 3. Les opérateurs énergétiques (panneaux solaires, batteries ou autres) sont reliés au niveau 3 du gestionnaire d'énergie et les opérateurs de radiocommunication (éléments des chaînes RF) sont reliés au niveau 3 du gestionnaire des communications.

Chacun de ces opérateurs est un élément d'un ensemble composé d'une station de base et d'éléments de stockage et de production d'énergie. Pour chacun de ces ensembles on a un gestionnaire de niveau 2 qui est divisé en deux gestionnaires séparés : le premier pour l'énergie et le second pour les communications. Ces deux gestionnaires peuvent prendre des décisions séparément mais ont aussi la possibilité d'échanger de l'information. Par exemple, le gestionnaire d'énergie de niveau 2 peut communiquer l'état de charge de la batterie au gestionnaire des communications qui peut utiliser cette information pour ajuster le service fourni aux utilisateurs.

Au niveau 1, on a aussi deux gestionnaires. Leur rôle est de gérer l'ensemble du réseau. Ils sont situés dans la macro-cellule qui est l'élément central du réseau. Ces deux gestionnaires ont deux rôles principaux. Le premier est de prendre des décisions globales, par exemple, le gestionnaire des communications pourra choisir d'allumer ou d'éteindre certaines micro-cellules en fonction de la densité d'utilisateurs dans le réseau [5]. Le second rôle de ces gestionnaires globaux est de gérer le lien avec le reste de l'infrastructure. Le gestionnaire des communications gère le réseau de *backhaul* qui fait le lien avec les serveurs centraux de l'architecture réseau. De son côté, le gestionnaire électronique gère le lien avec le reste du réseau électrique.

Ce dernier lien a une grande importance vu que le réseau de communication mobile considéré est un élément au sein d'un ensemble de producteurs et consommateurs qui partagent la même infrastructure électrique. Leur fonctionnement est régi par le biais d'un mécanisme appelé *Advanced Metering Infrastructure* (AMI) qui est l'une des fonctions les plus importantes du réseau électrique intelligent que nous allons étudier particulièrement dans ce manuscrit.

1.3.3 Position dans le réseau électrique

Dans un réseau électrique intelligent ou l'électricité produite provient, en partie, de sources d'énergie renouvelables et de producteurs indépendants, le gestionnaire du réseau électrique doit avoir une connaissance quasiment en temps réel (la remontée d'information peut se faire toutes les heures, ou plus fréquemment) de la production et de la consommation électrique. Il utilise cette connaissance pour s'assurer que l'équilibre qui existe entre production et consommation se fait sans risquer l'écroulement du réseau électrique. En cas de changement de la quantité d'électricité produite par les sources d'énergies renouvelables intermittentes, ou de la consommation, il va pouvoir modifier la quantité d'énergie qu'il produit (cet ajustement est possible dans des centrales nucléaires ou à charbon ou en délestant une partie de cette électricité). Il peut aussi envoyer des commandes incitatives par exemple en ajustant le prix de l'électricité pour influencer sur la consommation [19].

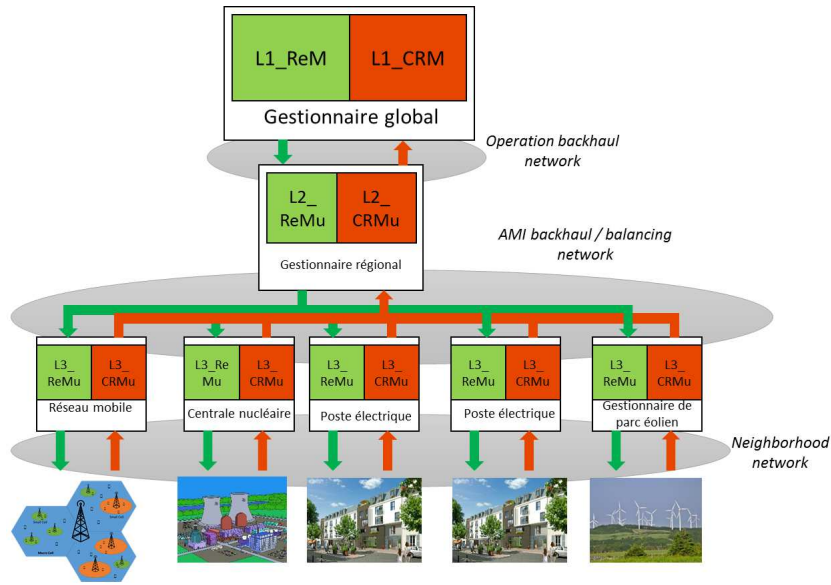


Figure 1.6 – Description HDCRAM de la gestion de la production et de la consommation dans un réseau électrique.

Les producteurs et consommateurs d'électricité et en particulier le réseau mobile présenté dans la section précédente vont adapter leur consommation d'électricité à ce prix. Si le prix est faible, ils vont consommer l'électricité fournie par le réseau, par exemple en rechargeant les batteries, et si le prix est élevé, ils peuvent consommer l'énergie qui a été stockée dans les batteries ou passer dans un régime de fonctionnement moins énergivore. Le choix de l'achat ou non d'électricité se fait au niveau L1 de l'architecture présenté sur la Figure 1.5.

Du point de vue de l'opérateur du réseau électrique, le gestionnaire d'énergie du réseau mobile est un producteur et consommateur d'électricité comme un autre et a le même rôle que tout autre client ou fournisseur. La gestion du prix de l'électricité est souvent décrite en trois étapes [20]. Nous pouvons donc la décrire grâce à l'architecture HDCRAM [21]. La description HDCRAM de la gestion de la production et de la consommation d'énergie dans un réseau électrique est présentée sur la Figure 1.6.

Dans cette architecture, la consommation et production d'électricité de chacun des agents du réseau électrique est transmise par les L3_CRMu vers le gestionnaire régional (L2_CRMu) qui les transfèrent vers le gestionnaire global. Celui-ci, va utiliser cette information pour envoyer des ordres et par exemple le nouveau prix de l'électricité par le biais du L1_ReMu. Ce nouveau prix va être transmis aux L3_ReMu et utilisé par le gestionnaire de chacun des réseaux pour décider de l'achat ou non d'électricité.

Il est important de noter que dans la description HDCRAM du réseau électrique intelligent, il y a de grandes distances entre les gestionnaires des différents niveaux. C'est pourquoi il est primordial de s'intéresser aux réseaux de communications qui existent entre eux.

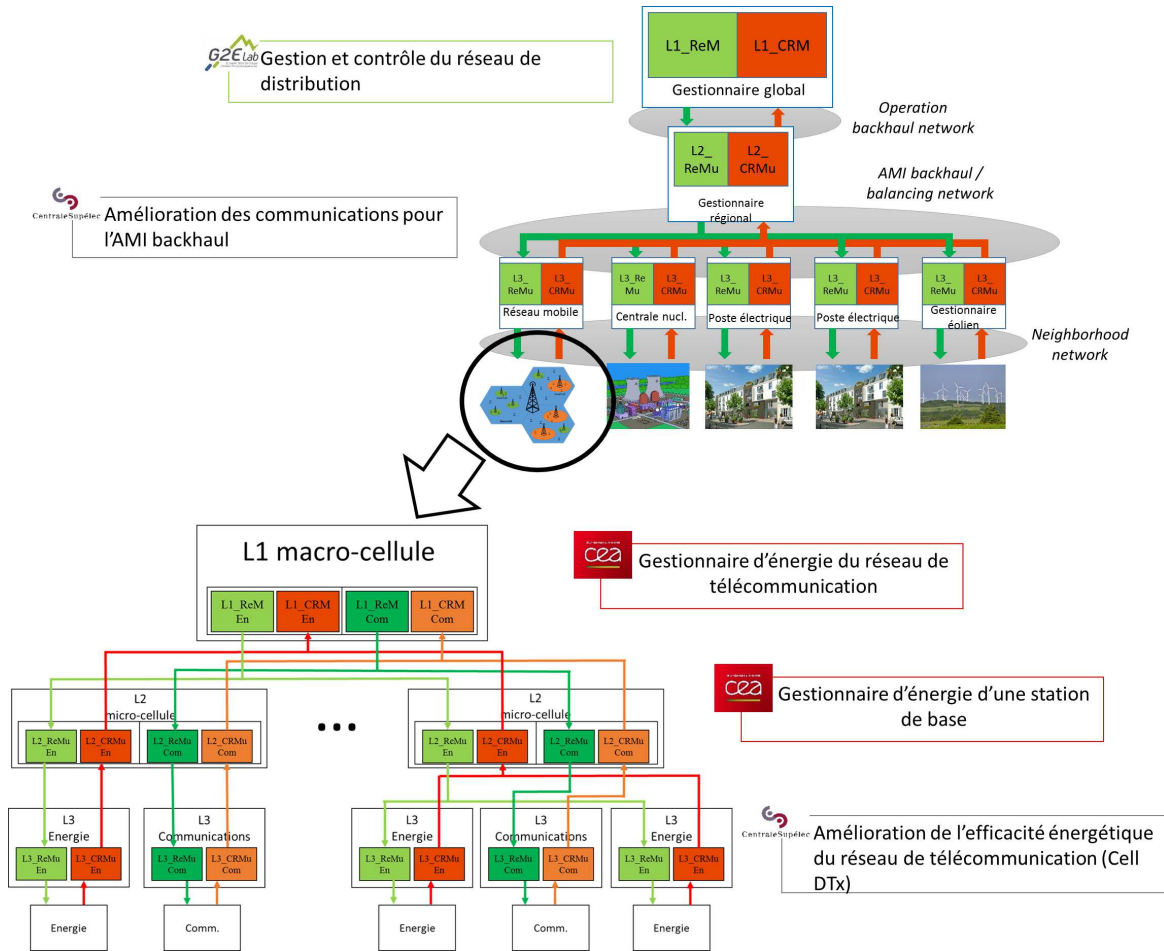


Figure 1.7 – Description complète du système considéré.

Dans la Figure 1.6, on peut voir trois réseaux de communications [22]. Le premier pour les communications locales entre le gestionnaire de niveau 3 et l'agent auquel il est connecté. Le second réseau est le réseau est l'AMI backhaul. Il fait le lien entre les gestionnaires locaux et le gestionnaire régional de niveau 2. Enfin, on a un dernier réseau qui fait le lien entre les gestionnaire régionaux et le centre de contrôle de l'opérateur du réseau électrique. Nous proposons une analyse plus complète des mécanismes de gestion du réseau électrique intelligent et des réseaux de communication associés dans le Chapitre 5.

1.3.4 Description complète

Pour avoir une vue complète du système étudié, on peut placer sur une même figure la description HDCRAM du réseau électrique considéré de la Figure 1.5 et la représentation HDCRAM de la gestion de la production et consommation de la Figure 1.6. C'est ce qui est fait sur la Figure 1.7.

Sur la Figure 1.7, le réseau de communication mobile sobre en énergie est à la fois vu comme un système intelligent qui a ses propres mécanismes de gestion mais aussi comme un élément d'un système plus global qu'est le réseau électrique intelligent. C'est pourquoi, on a un grand nombre d'unités de traitement et de reconfiguration. La prise de décision dans chacun des gestionnaires pourrait être optimisée. L'étude complète d'un tel système dépasse largement le cadre de cette thèse et demande de nombreuses études : pour la gestion de l'énergie, pour la gestion des communications mobiles sobres en énergie, pour les interactions entre les différents gestionnaires ou pour l'étude des réseaux de communication pour le réseau électrique intelligent.

Ce système a été étudié dans le projet ANR SOGREEN (*Smart Power Grid for Energy Efficient small cell Networks*) qui a impliqué trois partenaires :

- Le CEA-LETI (Laboratoire d'Electronique et de Technologie de l'Information) qui a travaillé sur la gestion de l'énergie du réseau mobile.
- Le G2ELab (Laboratoire de Génie Electrique de Grenoble) qui a travaillé sur la gestion du réseau électrique intelligent.
- CentraleSupélec qui a travaillé sur les communications du réseau mobile et du réseau électrique intelligent.

Sur la Figure 1.7, nous avons placé les contributions des trois partenaires.

Cette thèse s'inscrit dans le cadre de la contribution de CentraleSupélec au projet SOGREEN et s'est, par conséquent, focalisée sur l'étude des réseaux de communication. Nous avons donc une étude en deux parties :

- Une partie pour l'étude des communications du réseau mobile qui doivent être le plus sobre possible en énergie.
- Une second pour les réseaux de communication pour le réseau électrique intelligent.

1.4 Études menées dans cette thèse

Dans la suite de ce manuscrit, nous étudions séparément les deux réseaux cités dans la section précédente. L'étude de chacun de ces réseaux se fera dans une partie dédiée.

Dans un premier temps, nous nous intéressons à la réduction de la consommation d'énergie dans les réseaux mobiles et en particulier à la transmission discontinue [8]. Ce mécanisme permet de réduire la consommation d'énergie de chaque cellule indépendamment des cellules voisines. Dans la Figure 1.7, ce mécanisme opère dans un gestionnaire de niveau 3 du réseau mobile associé aux communications (L3 Communication).

Ensuite, dans une seconde partie, nous étudions les réseaux de communication pour le réseau électrique intelligent. En particulier, nous nous intéressons aux communications de l'*AMI backhaul*. C'est-à-dire aux communications entre les gestionnaires de niveau 2 et de niveau 3 de la Figure 1.6. Nous verrons que des solutions qui ont été présentées pour l'AOS peuvent améliorer les communications du réseau électrique intelligent.

Première partie

Economie d'énergie dans les
réseaux cellulaires

Chapitre 2

Allocation de ressources et de puissance pour des réseaux mobiles sobres en énergie

Sommaire

2.1	Les réseaux mobiles	24
2.2	L'accès multiple	25
2.3	Formulation du problème d'allocation de ressources et de puissance en OFDMA	27
2.3.1	Définitions et vocabulaire	27
2.3.2	Maximisation de la capacité	28
2.3.3	Prise en compte de la consommation énergétique	30
2.4	Résolution du problème d'optimisation	31
2.4.1	Allocation de ressources	32
2.4.2	Allocation de puissance	33
2.5	Évolutions récentes des réseaux cellulaires	36
2.6	Consommation d'énergie des réseaux mobiles	37
2.7	Mise en veille des stations de base et transmission discontinue .	38
2.7.1	Mise en veille prolongée	38
2.7.2	Transmission discontinue	39
2.8	Allocation de puissance et de ressources avec du Cell DTx . . .	40
2.8.1	Allocation de puissance à une seule station de base	40
2.8.2	Cell DTx dans des réseaux à plusieurs cellules	43
2.9	Conclusions	44

L'objectif de ce chapitre est de présenter un état de l'art du problème d'allocation de ressources et de puissance dans un réseau cellulaire sobre en énergie. Après avoir décrit

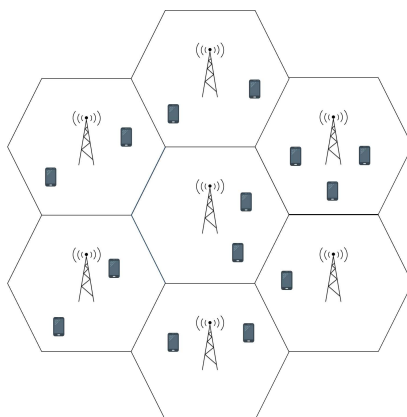


Figure 2.1 – Représentation schématique d'un réseau cellulaire composé de plusieurs stations de base chacune servant différents utilisateurs.

brèvement le fonctionnement d'un réseau mobile, nous présentons le problème d'allocation de ressources et de puissance. Nous verrons que, dans un réseau à faible consommation d'énergie, ce problème a été formulé de sorte à minimiser la puissance d'émission tout en fournissant aux utilisateurs une certaine qualité de service. Nous décrivons, ensuite, la solution de ce problème telle qu'elle a été proposée dans la littérature. Ensuite, nous verrons qu'avec l'augmentation du nombre de stations de base, il n'est plus suffisant de réduire la puissance d'émission et qu'il faut faire appel à des mécanismes de mise en veille pour réduire davantage l'énergie consommée par la station de base. Nous décrirons finalement l'état de l'art du problème d'allocation de ressources et de puissance avec des transmissions discontinues qui sera le sujet des Chapitres 3 et 4.

2.1 Les réseaux mobiles

Comme illustré sur la Figure 2.1, un réseau mobile est composé d'un ensemble de stations de base dont l'objectif est d'échanger des données avec les utilisateurs qui se trouvent dans sa zone de couverture, aussi appelée cellule. Une cellule est souvent symbolisée de forme hexagonale.

Dans un réseau mobile, on a différents types de communication. On a d'abord des échanges de données entre les stations de base pour gérer la mobilité des utilisateurs. Il est important de noter que lorsque la taille des cellules diminue, la quantité de données à échanger entre les stations de base augmente.

Les communications entre les stations de base et les utilisateurs sont les communications principales du réseau. Dans celles-ci, on a deux types de données :

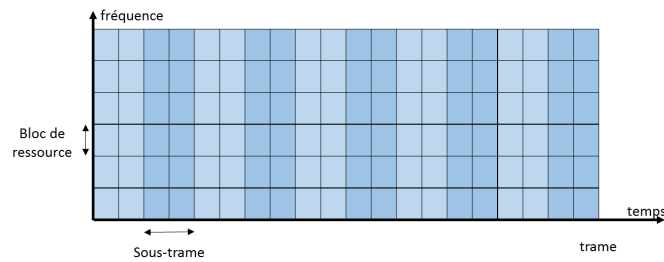


Figure 2.2 – Une trame LTE est divisée en sous-trames temporelles et en blocs de ressources fréquentielles.

- Les données utiles. C'est à dire l'information que l'on veut échanger et qui va être utilisée pour un autre usage.
- Les données de contrôle. Ce sont des données qui contiennent de l'information relative à la communication.

De plus, lorsqu'un utilisateur et une station de base communiquent, on a un échange de données de l'utilisateur vers la station de base (voie montante) et de la station de base vers l'utilisateur (voie descendante). Dans une grande partie des applications mobiles actuelles (*streaming* vidéo, réseaux sociaux, écoute de musique en continu) la majorité des données mobiles sont envoyées sur la voie descendante, c'est à dire, vers le mobile. La voie montante étant principalement utilisée pour des données de contrôle. Quel que soit le type de données échangées, les données binaires à transmettre sont mises en trame, codées, et modulées. Les symboles modulés sont alors mis sur la bonne fréquence porteuse et transmis.

Avant d'être reçu par l'utilisateur, le signal transféré va passer à travers un canal. Celui-ci va avoir plusieurs effets sur le signal. En particulier ce canal va atténuer le signal émis. Finalement, l'utilisateur va recevoir un signal atténué et bruité. Il va pouvoir utiliser les données reçues une fois qu'il aura réussi à les démoduler. Il est important de noter que l'effet du canal sur les symboles émis n'est pas forcément constant sur l'ensemble de la bande de fréquence.

Pour transférer des données aux utilisateurs sur la voie descendante, la station de base dispose d'un espace temps-fréquence limité. Elle va, donc, devoir répartir cet espace entre les différents utilisateurs. Cette répartition s'appelle l'accès multiple ou *multiple access*.

2.2 L'accès multiple

Pour bien comprendre ce problème, nous représentons sur la Figure 2.2 une trame du standard *Long Term Evolution* (LTE) utilisé dans la quatrième génération de téléphonie mobile (4G) [23].

Dans ce standard, le temps est divisé en trames. Une trame dure 10 ms et est divisée en dix sous-trames de 1 ms. La largeur de la bande de fréquence utilisée pour servir les utilisateurs peut varier entre 1,08 MHz et 18 MHz. Cette bande est divisée en blocs de ressources ayant une largeur de bande de 180 kHz.

Ce découpage suggère deux stratégies différentes pour l'accès multiple : le *Time Division Multiple Access* (TDMA) et le *Frequency Division Multiple Access* (FDMA)¹. En TDMA, les utilisateurs sont séparés temporellement et servis les uns après les autres. On alloue un certain nombre de sous-trames à chaque utilisateur. Au contraire, en FDMA, ils sont séparés en fréquence, c'est à dire que l'on alloue à chaque utilisateur un certain nombre de blocs de ressources. Cette seconde solution a deux avantages comparé au TDMA :

- Pour un utilisateur donné, on peut avoir des portions de bande dans lesquelles le canal de propagation qui le sépare de la station de base est très mauvais. On dira alors que le canal est sélectif en fréquence. En FDMA, contrairement au TDMA, on peut répartir les utilisateurs de sorte que chacun des blocs ressources soit alloué à un utilisateur qui a un bon canal de propagation dans celui-ci.
- La puissance d'émission instantanée de la station de base est limitée par des contraintes électroniques et de régulation. Par exemple, dans une macro-cellule, la puissance d'émission ne peut pas dépasser 20 W. Cette contrainte de puissance restreint davantage les performances d'une station de base TDMA que celles d'une station de base FDMA [24].

Dans le standard 4G, l'allocation FDMA est utilisée. Plus précisément, dans ce standard une modulation *Orthogonal Frequency Division Multiplexing* (OFDM) est utilisée et les blocs de ressources sont orthogonaux. Dans ce cas là, on parle de *Orthogonal Frequency Division Multiple Access* (OFDMA).

Les approches TDMA et FDMA ne sont pas les seules possibles. Une autre possibilité consiste à servir plusieurs utilisateurs en même temps et dans les mêmes blocs de ressources et à la séparer par d'autres solutions. Les différentes technologies d'accès multiple sont les suivantes :

- Le *Code Division Multiple Access* (CDMA) [25], dans lequel les données modulées sont multipliées par des séquences orthogonales. Chacune des séquences code les données pour un utilisateur qui pourra décoder l'information qui lui est destinée en multipliant par la même séquence. Cette solution est, par exemple, utilisée dans la troisième génération de téléphonie mobile (3G).

1. Il est envisageable d'utiliser à la fois du TDMA et du FDMA. Pour des raisons de clarté, dans ce chapitre, nous présentons le TDMA et le FDMA comme deux stratégies séparées.

- Le *Spatial Division Multiple Access* (SDMA) [26]. Une station de base ayant plusieurs antennes d'émission peut les utiliser pour donner une direction aux signaux envoyés. Elle peut ainsi séparer spatialement les données transmises aux différents utilisateurs.
- Le *Non-Orthogonal Multiple Access* (NOMA) [27]. En NOMA les utilisateurs sont servis sans être séparés orthogonalement en temps ou en fréquence. Par exemple en *Power Domain NOMA* (PD-NOMA), ils sont servis avec des puissances d'émission très différentes. Grâce à cette grande différence, on peut successivement identifier le signal qui a la plus forte puissance et le soustraire au signal reçu. Ainsi tous les utilisateurs peuvent avoir accès à toutes les données et chaque utilisateur peut récupérer les données qui lui sont dédiées. Cette stratégie d'accès est très récente comparé à celles présentées ci-dessus.

Aujourd'hui la technologie OFDMA est utilisée dans les réseaux 4G. C'est aussi celle qui a suscité le plus d'intérêt dans la communauté scientifique. C'est pourquoi, dans les Sections 2.3 et 2.4, nous présentons en détail le problème d'allocation de ressources et de puissance en OFDMA. Plus précisément, dans la Section 2.3, nous présentons les formulations possibles de ce problème. Ensuite, dans la Section 2.4, nous présentons un état de l'art des solutions proposées pour sa résolution.

2.3 Formulation du problème d'allocation de ressources et de puissance en OFDMA

2.3.1 Définitions et vocabulaire

Avant de présenter le problème d'allocation de ressources et de puissance. Il est nécessaire de préciser le vocabulaire qui est utilisé dans cette thèse.

Dans la suite, on décompose le problème complet d'**allocation de ressources et de puissance** en un problème d'allocation de puissance et un problème d'allocation de ressources. Ces deux termes ne vont pas avoir la même signification lorsque le canal est plat et lorsqu'il varie sur la largeur de bande.

Si le canal d'au moins un des utilisateurs servis a de l'évanouissement (*fading*), c'est à dire que son canal n'est pas plat :

- On résout d'abord le problème d'**allocation de ressources** qui consiste à déterminer l'ensemble des blocs de ressources qui sont utilisés pour servir chaque utilisateur. On doit donc déterminer le nombre de blocs de ressources utilisés pour chacun ainsi que la position de ces blocs de ressources.

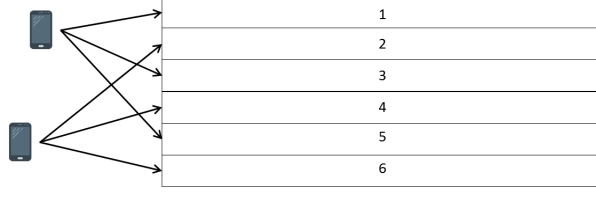


Figure 2.3 – Allocation de ressources en OFDMA. Dans cet exemple, la station de base dispose de $N_c = 6$ blocs de ressources pour servir $N_u = 2$ utilisateurs.

- Une fois cette allocation de ressources faite, résoudre le problème d'**allocation de puissance** revient à distribuer la puissance d'émission entre les blocs de ressources.

Lorsque le canal de tous les utilisateurs est plat, le problème d'allocation de ressources et de puissance se limite au problème d'allocation de puissance. En effet, vu que le coefficient d'atténuation du canal de chaque utilisateur ne change pas d'un bloc de ressources à l'autre, en réalisant l'allocation de puissance, on détermine aussi le nombre de blocs de ressources à allouer à chaque utilisateur. En d'autres termes, résoudre le problème d'allocation de puissance revient à déterminer la puissance utilisée pour chaque utilisateur ainsi que le nombre de blocs de ressources dans lesquels il sera servi. Tous les canaux étant identiques, la répartition des utilisateurs dans les blocs de ressources n'a pas d'impact sur la métrique optimisée.

De plus, dans la suite de ce document, nous utilisons le terme de **contrôle de puissance** pour désigner le fait de servir un utilisateur avec une puissance d'émission inférieure à la puissance d'émission maximale de la station de base. Ainsi, avec une stratégie sans contrôle de puissance, la puissance d'émission de la station de base est toujours la plus élevée possible.

2.3.2 Maximisation de la capacité

On considère une station de base qui doit servir N_u utilisateurs. Pour les servir, elle dispose de N_c blocs de ressources, comme indiqué sur la Figure 2.3. Si les blocs de ressources sont suffisamment fins, on peut supposer que l'atténuation dans le canal de propagation de chaque utilisateur est constant sur la largeur d'un bloc de ressources. Dans la suite on note $h_{k,n}$, $k, n \in \llbracket 1; N_u \rrbracket \times \llbracket 1; N_c \rrbracket$, le coefficient d'atténuation complexe du canal de l'utilisateur k dans le bloc de ressources n . De plus, le carré de son module est noté $g_{k,n} = |h_{k,n}|^2$. Si on suppose ces coefficients connus par la station de base, celle-ci va vouloir répartir les utilisateurs dans les différents blocs de ressources en fonction des modules $g_{k,n}$. En d'autres termes, chaque utilisateur se voit allouer un ensemble $\mathcal{S}_k$ de blocs de ressources. Ensuite, la station de base alloue une certaine puissance d'émission à chaque utilisateur dans chaque canal. La puissance d'émission de l'utilisateur k dans le bloc de ressources n est noté $P_{T_X}^{k,n}$.

Cette allocation des utilisateurs dans les différents blocs de ressources et l'allocation de puissance associée ne sont pas faites de façon anodine mais pour optimiser une certaine métrique. Par exemple, cette allocation de ressources et de puissance peut être faite pour maximiser la somme des capacités de canal des utilisateurs [28]. La capacité de canal est le débit théorique maximum qu'un utilisateur peut avoir dans un canal additif gaussien. Cette métrique a deux avantages. Le premier est qu'elle est indépendante de la modulation choisie (contrairement au taux d'erreur binaire, par exemple). De plus, lorsque le *Modulation and Coding Scheme* (MCS) est bien adapté, le débit d'une communication LTE est proportionnel à la capacité de Shannon [29].

La maximisation de cette capacité doit se faire en considérant que la puissance d'émission instantannée ne doit pas dépasser une certaine valeur maximale notée $P_{\max}$ (cette valeur maximale peut être la puissance maximale des composants de la station de base, la puissance maximale de sortie de l'amplificateur de puissance ou encore une contrainte de régulation). On a, finalement, le problème d'optimisation suivant, qui a, par exemple, été posé en 2000 par Rhee et al. dans [30] :

$$\max_{\mathcal{S}_k, P_{T_X}^{k,n}} B_c \sum_{k=1}^{N_u} \sum_{n \in \mathcal{S}_k} \log_2 \left(1 + \frac{P_{T_X}^{k,n} g_{k,n}}{N} \right), \quad (2.1a)$$

$$\text{S.à; } \sum_{k=1}^{N_u} \sum_{n \in \mathcal{S}_k} P_{T_X}^{k,n} \leq P_{\max}, \quad (2.1b)$$

$$P_{T_X}^{k,n} \geq 0, \forall k, n \in \llbracket 1; N_u \rrbracket \times \llbracket 1; N_c \rrbracket, \quad (2.1c)$$

$$B_c \sum_{n \in \mathcal{S}_k} \log_2 \left(1 + \frac{P_{T_X}^{k,n} g_{k,n}}{N} \right) \geq C_k, \forall k \in \llbracket 1; N_u \rrbracket, \quad (2.1d)$$

$$\mathcal{S}_k \cap \mathcal{S}_{k'} = \emptyset, k \neq k'. \quad (2.1e)$$

Dans l'équation (2.1), S. à, est l'abréviation de Sujet à, qui permet de désigner les contraintes du problème. Dans cette formulation du problème, B_c est la largeur de bande d'un bloc de ressources et $N = k_B T B_c$ est la puissance du bruit qui est le produit de la constante de Boltzmann k_B , la température en Kelvin et de la largeur de bande d'un bloc de ressources.

Dans la formulation du problème d'allocation de ressources et de puissance donné par l'équation (2.1), on veut maximiser la capacité totale qui est donnée par l'équation (2.1a) et est la somme de la capacité de chaque utilisateur. Cette maximisation doit se faire en satisfaisant certaines contraintes. La première est que l'on ne doit pas dépasser la puissance d'émission maximale de la station de base (équation (2.1b)). Ensuite, la puissance d'émission dans chaque bloc de ressource doit être positive (équation (2.1c)). On a aussi la contrainte de l'équation (2.1d) qui nous force à fournir à chaque utilisateur un débit minimal C_k qui

dépend du service demandé par l'utilisateur k . Avec cette contrainte, on évite que, dans le cas où un utilisateur a le meilleur coefficient de canal dans chaque bloc de ressources, toutes les ressources lui soient allouées. Enfin, la dernière contrainte de ce système d'équation assure que chaque bloc de ressources n'est alloué qu'à un seul utilisateur.

2.3.3 Prise en compte de la consommation énergétique

Dans la section précédente, nous avons vu comment formuler le problème d'allocation de ressources et de puissance en OFDMA pour maximiser la capacité. En maximisant la capacité, on va utiliser toute la puissance disponible au niveau de la station de base. Cette approche est donc énergivore et ne peut pas être utilisée dans un réseau sobre en énergie. Dans la littérature, on retrouve principalement deux objectifs pour l'allocation de puissance dans les réseaux sobres en énergie :

- la maximisation de l'efficacité énergétique (EE) [31, 32],
- la minimisation de la puissance consommée [33].

L'EE est définie comme le rapport entre la capacité de Shannon et la puissance consommée par la station de base qui est une fonction des puissances d'émission des utilisateurs. Dans le cas où la station de base ne sert qu'un seul utilisateur à travers un canal plat (le module du canal est constant sur la largeur de bande), son expression est :

$$EE = \frac{C}{P(P_{T_X})} = \frac{B \log_2 \left(1 + \frac{P_{T_X} g}{N} \right)}{P(P_{T_X})}. \quad (2.2)$$

Dans l'équation 2.2, g est le module du canal, P_{T_X} est la puissance d'émission et B est la largeur de bande et P est la puissance consommée par la station de base. De plus, dans un système OFDMA, l'efficacité énergétique peut s'écrire :

$$EE = \frac{B_c \sum_{k=1}^{N_u} \sum_{n \in \mathcal{S}_k} \log_2 \left(1 + \frac{P_{T_X}^{k,n} g_{k,n}}{N} \right)}{P(P_{T_X}^{1,1}, \dots, P_{T_X}^{N_u, N_c})} \quad (2.3)$$

Bien que l'efficacité énergétique soit souvent utilisée dans des problèmes d'optimisation, nous pensons, pour plusieurs raisons que cette métrique n'est pas la meilleure. Voici une liste non-exhaustive des raisons qui nous font penser cela :

- Le ratio $\frac{C}{P}$ crée une asymétrie entre puissance consommée et capacité. En effet, à capacité donnée, une diminution de 20 % de la puissance consommée augmente l'EE de 25 % alors qu'une augmentation de 20 % de la capacité à puissance donnée va augmenter l'EE de 20 %.

- Si on a un réseau composé de plusieurs stations de base, la capacité totale est la somme des capacités des stations de base, la puissance consommée totale est la somme des puissances consommées alors que l'EE totale n'est pas la somme des EE. Ce qui pose un problème de passage à l'échelle.
- Nous verrons dans la suite de ce chapitre que la puissance consommée par une station de base est égale à une puissance qui dépend de la puissance d'émission à laquelle on ajoute une puissance qui ne dépend pas directement de la puissance d'émission. Cette dernière, que nous appellerons puissance statique va avoir un effet important sur l'efficacité énergétique alors que l'allocation de puissance n'aura aucun impact sur celle-ci. Par exemple, si la puissance statique est égale à 0, l'EE est maximisée en mettant à 0 la puissance d'émission et si elle est très élevée, on a intérêt à maximiser la puissance d'émission.

En plus des éléments, cités ci-dessus, il nous semble, que d'un point de vue opérateur, il est plus raisonnable de poser le problème de minimisation de la puissance d'émission pour une qualité de service donnée, et par exemple, pour une capacité donnée. On retrouve alors le problème d'allocation de ressources et de puissance tel qu'il a été posé dans [34] :

$$\min_{\mathcal{S}_k, P_{T_X}^{k,n}} \sum_{k=1}^{N_u} \sum_{n \in \mathcal{S}_k} P_{T_X}^{k,n}, \quad (2.4a)$$

$$\text{S.à } C_k - B_c \sum_{n \in \mathcal{S}_k} \log_2 \left(1 + \frac{P_{T_X}^{k,n} g_{k,n}}{N} \right) \leq 0, \forall k \in \llbracket 1; N_u \rrbracket, \quad (2.4b)$$

$$P_{T_X}^{k,n} \geq 0, \forall k, n \in \llbracket 1; N_u \rrbracket \times \llbracket 1; N_c \rrbracket, \quad (2.4c)$$

$$\mathcal{S}_k \cap \mathcal{S}_{k'} = \emptyset, k \neq k'. \quad (2.4d)$$

Dans ce problème d'optimisation, C_k est la contrainte de capacité de l'utilisateur k . Dans la suite de ce chapitre, nous expliquons comment résoudre le problème d'allocation de ressources et de puissance tel qu'il est posé dans l'équation (2.4).

2.4 Résolution du problème d'optimisation

La résolution du problème d'optimisation de l'équation (2.4) permet d'allouer les blocs de ressources aux utilisateurs et de leur fournir une certaine puissance d'émission. Comme expliqué dans la Section 2.3.1, lorsqu'il y a du *fading*, ce problème est résolu en deux étapes :

- On résout d'abord le problème d'**allocation de ressources**, ce qui consiste à choisir, pour chaque utilisateur, le jeu de blocs de ressources qui lui est alloué.

- Ensuite on s'intéresse au problème d'**allocation de puissance**. C'est-à-dire que, une fois que les blocs de ressources ont été alloués aux utilisateurs, on calcule la puissance que l'on doit allouer à chaque utilisateur dans chaque bloc de ressources.

2.4.1 Allocation de ressources

On commence par traiter le problème d'allocation de ressources. Dans ce problème, on veut allouer un certain nombre de blocs de ressources à chacun des utilisateurs servis par la station de base. Avec N_u utilisateurs et $N_c \geq N_u$ blocs de ressources, le nombre d'associations possibles est égal à [35] :

$$N_{as} = \sum_{i=0}^{N_u} (-1)^i \binom{N_u}{i} (N_u - i)^{N_c} = \mathcal{O}(N_u^{N_c}) \quad (2.5)$$

Ce nombre prend en compte le fait que chaque utilisateur doit être servi dans au moins un bloc de ressources. N_{as} est bien trop grand pour espérer pouvoir faire une recherche exhaustive. Par exemple, si on a 4 utilisateurs et 25 blocs de ressources, le nombre d'allocations de ressources possibles est supérieur à 10^{15} . On doit donc trouver des algorithmes sous-optimaux efficaces pour résoudre ce problème. Un grand nombre de solutions ont été proposées dans la littérature pour résoudre ce problème [36]. Parmi les nombreuses stratégies qui permettent de trouver une solution sous-optimale, les plus remarquables sont les suivantes :

1. Une première solution pour obtenir un algorithme sous-optimal est la relaxation de contraintes. C'est à dire, la transformation des contraintes discrètes en contraintes continues. Cette relaxation transforme le problème combinatoire en un problème à variables continues qui est plus facile à résoudre. Enfin, on discrétise le résultat de cette première optimisation pour obtenir la solution espérée [30, 33].
2. Pour résoudre le problème d'allocation de ressources, on peut aussi utiliser le problème dual [37]. Pour cela il faut d'abord écrire le Lagrangien du problème complet (allocation de ressources et de puissance). Ce Lagrangien dépend de $x_1, \dots, x_p$, les variables du problème d'optimisation, et des multiplicateurs de Lagrange $\lambda_1, \dots, \lambda_n$. Pour résoudre notre problème d'optimisation, on va d'abord exprimer les valeurs des variables qui minimisent ce Lagrangien. On obtient alors une fonction $h(\lambda_1, \dots, \lambda_n) = \min_{x_i} \mathcal{L}$. Cette première étape peut être réalisée en utilisant la théorie de la décomposition [34, 38]. Ensuite, on cherche de, manière itérative, les valeurs de $\lambda_1, \dots, \lambda_n$ qui maximisent ce Lagrangien. Une fois que ce problème a été résolu, on obtient une solution sous-optimale.
3. Enfin, l'allocation de ressources peut se faire de façon heuristique. C'est à dire en s'appuyant sur des intuitions et non sur des théories mathématiques. C'est le cas de l'algorithme proposé dans [39]. Dans cet algorithme, le nombre de blocs de ressources

N_k est calculé en fonction du coefficient de canal moyen de l'utilisateur. Ensuite, pour l'allocation des blocs de ressources, pour chaque bloc, on va choisir, parmi les utilisateurs qui n'ont pas encore leurs N_k blocs de ressources, celui qui a le meilleur coefficient d'atténuation du canal. Bien que très simple (en terme de complexité informatique et de compréhension), cet algorithme a de très bonnes performances.

Nous verrons dans le Chapitre 4 que ces solutions, et en particulier celle proposée dans [39], ont de bonnes performances lorsqu'elles sont utilisées avec des mécanismes de mise en veille dynamique de la station de base.

Nous venons de présenter un bref état de l'art du problème d'allocation de ressources. Dans la section suivante, nous présentons la méthode de résolution du problème d'allocation de puissance.

2.4.2 Allocation de puissance

Une fois que l'allocation de ressources a été réalisée, on doit résoudre le problème d'allocation de puissance. Ce problème étant convexe, il peut être résolu de manière optimale. L'objectif de cette section est donc de présenter l'algorithme de résolution optimal. Nous commençons par présenter des généralités sur la convexité utiles pour comprendre cet algorithme de résolution optimal. Les notions introduites ici seront utilisées dans les Chapitres 3 et 4.

Généralités sur la convexité

Un problème de la forme :

$$\min_x f_0(x_1, \dots, x_p), \quad (2.6a)$$

$$\text{S.à } f_i(x_1, \dots, x_p) \leq 0, \forall i \in \llbracket 1; m \rrbracket, \quad (2.6b)$$

$$g_j(x_1, \dots, x_p) = 0, \forall j \in \llbracket 1; n \rrbracket, \quad (2.6c)$$

est convexe si et seulement si [40] :

- La fonction objectif f_0 est convexe.
- Les m fonctions f_i qui définissent les contraintes d'inégalité sont convexes.
- Les n fonctions g_j qui définissent les contraintes d'égalité sont affines.

De plus, une fonction dont la dérivée seconde est continue, est convexe si et seulement si sa matrice Hessienne (la matrice des dérivées partielles secondes) est définie positive. Pour les fonctions à une seule variable, cette propriété revient à prouver que la dérivée seconde est positive. Pour prouver qu'une fonction est convexe on peut aussi utiliser les propriétés de

conservation de la convexité par des opérations élémentaires. Par exemple, la somme de deux fonctions convexes est convexe et si f est concave, $-f$ est convexe [40].

La convexité d'un problème d'optimisation implique que, si le minimum existe, on peut trouver son optimum global en appliquant les conditions de Karush-Kuhn et Tucker (KKT) [41]. Ces conditions nous disent que si le problème n'a pas de contraintes, on peut trouver le minimum de la fonction en trouvant le point auquel la dérivée de la fonction objectif s'annule. Si le problème n'a que des contraintes d'égalité de la forme $g_j(x_1, \dots, x_p) = 0$, le minimum de la fonction objectif est trouvé au point d'annulation de la dérivée du Lagrangien dont l'expression est :

$$\mathcal{L} = f_0(x_1, \dots, x_p) + \sum_{j=1}^n \lambda_j g_j(x_1, \dots, x_p). \quad (2.7)$$

Dans ce Lagrangien, les λ_j sont les multiplicateurs de Lagrange. La valeur de ces multiplicateurs doit être calculée pour qu'ils satisfassent les contraintes d'égalité du problème.

On se place maintenant dans le cas général. Dans ce cas là, on a des contraintes d'inégalité. Pour chacune d'elles, on a deux possibilités :

- Soit l'optimum est atteint lorsque la contrainte est une inégalité stricte. Dans ce cas là, la fonction f_i correspondante ne doit pas être considérée dans le Lagrangien.
- Soit la contrainte est une égalité. C'est à dire que la contrainte doit être prise en compte dans le Lagrangien.

Pour mieux comprendre cet aspect, nous l'illustrons par un exemple simple. Supposons le problème d'optimisation suivant :

$$\min_x f_0(x), \quad (2.8a)$$

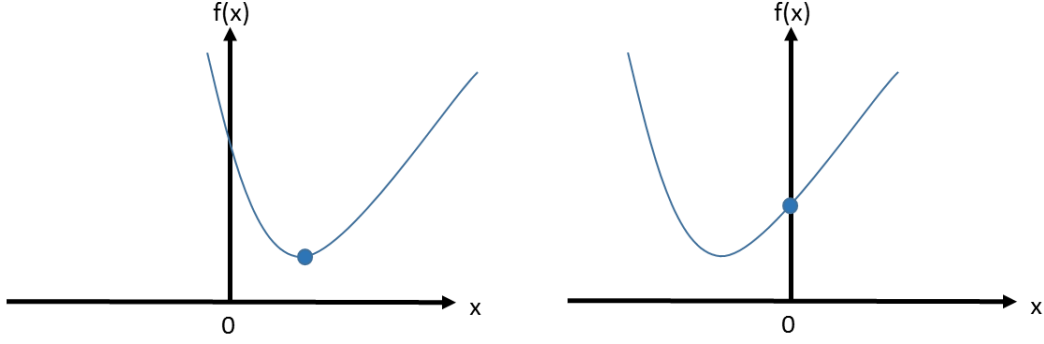
$$\text{S.à } x \geq 0. \quad (2.8b)$$

Dans ce problème la fonction f_0 est une fonction convexe. Pour ce problème, on a une contrainte d'inégalité. On peut donc avoir deux cas différents comme illustré sur la Figure 2.4.

Dans le premier cas, la fonction f_0 est minimum pour une valeur de x positive. La contrainte n'intervient pas dans la minimisation. La solution du problème est donnée par le minimum de la fonction f_0 . C'est à dire, au point x^* qui vérifie $f'_0(x^*) = 0$.

Dans le second cas, le minimum de f_0 est atteint pour un x inférieur à 0. Les conditions de KKT nous disent que dans ce cas, la solution au problème (2.8) est obtenue lorsque la contrainte est saturée, c'est à dire pour $x^* = 0$. Cette valeur x^* minimise le Lagrangien qui est égal à :

$$\mathcal{L} = f_0(x) - \lambda x \quad (2.9)$$



(a) Premier cas : la contrainte n'est pas saturée. (b) Second cas : la contrainte est saturée.

Figure 2.4 – Illustrations des deux cas possibles pour l'exemple donné par l'équation (2.8).

On voit dans cet exemple que pour une contrainte d'inégalité on a deux Lagrangiens possibles. Pour N contraintes d'inégalité on aura 2^N Lagrangiens possibles. Il est évidemment trop long de considérer l'ensemble des Lagrangiens lors de la résolution d'un problème. Heureusement, dans plusieurs problèmes d'optimisation, en particulier lorsque toutes les contraintes sont similaires (par exemple lorsqu'on a une contrainte de positivité pour chacune des variables), on peut réduire la difficulté de la résolution en analysant l'ordre dans lequel les contraintes vont être saturées. En d'autres termes, on va d'abord résoudre le problème sans considérer les contraintes d'inégalités et on va ensuite vérifier si elles sont satisfaites. Si elles ne le sont pas, on va devoir considérer comme saturées certaines de ces contraintes. On obtient alors un nouvel optimum potentiel, et ainsi de suite. Le processus itératif s'arrête lorsqu'on obtient un optimum potentiel pour lequel l'ensemble des contraintes sont satisfaites. Cet optimum sera alors l'optimum global de notre problème.

Application au problème d'allocation de puissance

Le procédé de résolution décrit ci-dessus permet de trouver la solution optimale du problème d'allocation de puissance suivant :

$$\min_{P_{T_X}^{k,n}} \sum_{k=1}^{N_u} \sum_{n \in \mathcal{S}_k} P_{T_X}^{k,n}, \quad (2.10a)$$

$$\text{S.à } C_k - B_c \sum_{n \in \mathcal{S}_k} \log_2 \left(1 + \frac{P_{T_X}^{k,n} g_{k,n}}{N} \right) = 0, \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (2.10b)$$

$$P_{T_X}^{k,n} \geq 0, \quad \forall n \in \llbracket 1; N_c \rrbracket. \quad (2.10c)$$

Pour le problème décrit dans l'équation (2.10), le jeu de blocs de ressources qui est attribué à chaque utilisateur a été déterminé lors de l'allocation de ressources. On cherche donc à déterminer la valeur de la puissance d'émission dans chaque bloc de ressources.

On remarque que, dans le problème de l'équation (2.10), la contrainte de capacité est une égalité alors que c'était une inégalité dans le problème défini équation (2.4). Cette inégalité a été transformée en égalité car il est évident que la puissance consommée est minimale lorsque la capacité de chaque utilisateur est égale à la capacité qu'il demande. On consommerait un surplus d'électricité en fournissant à un utilisateur plus de capacité que nécessaire.

Ce problème est convexe. On peut donc le résoudre en appliquant les conditions de KKT et la procédure décrite précédemment. On commence donc par résoudre le problème sans considérer les contraintes de positivité des puissances d'émission. Dans ce cas là, on dérive la valeur optimale de $P_{T_X}^{k,n}$ en trouvant la valeur pour laquelle la dérivée du Lagrangien est nulle :

$$P_{T_X}^{k,n} = N \left(\frac{2^{\frac{C_k}{N_k B_c}}}{\left(\prod_{n=1}^{N_k} g_{k,n} \right)^{\frac{1}{N_k}}} - \frac{1}{g_{k,n}} \right), \forall n \in \llbracket 1; N_c \rrbracket. \quad (2.11)$$

Dans cette équation N_k désigne le cardinal de $\mathcal{S}_k$. On a maintenant des valeurs de $P_{T_X}^{k,n}$ qui sont potentiellement optimales. En effet, ces valeurs ne sont optimales que si elles satisfont les contraintes de l'équation (2.10c). C'est à dire si tous les $P_{T_X}^{k,n}$ sont positifs. Si ce n'est pas le cas, on doit mettre à 0 les valeurs de $P_{T_X}^{k,n}$ qui sont négatives et recalculer les nouvelles valeurs. Ce processus itératif s'arrête lorsque tous les $P_{T_X}^{k,n}$ sont positifs. Finalement la valeur optimale de $P_{T_X}^{k,n}$ peut s'écrire :

$$P_{T_X}^{k,n} = N \left(\frac{2^{\frac{C_k}{N'_k B_c}}}{\left(\prod_{n=1}^{N'_k} g_{k,n} \right)^{\frac{1}{N'_k}}} - \frac{1}{g_{k,n}} \right)^+, \forall n \in \mathcal{S}_k. \quad (2.12)$$

Dans cette équation, $(.)^+ = \max(., 0)$ et N'_k désigne le nombre de blocs de ressources dans lesquels la puissance allouée à l'utilisateur k est non-nulle. L'algorithme présenté ici s'appelle algorithme de water-filling [28, 34].

2.5 Évolutions récentes des réseaux cellulaires

Dans la section précédente, nous avons présenté le problème d'allocation de puissance dans un scénario à une seule cellule et lorsque la station de base n'a qu'une seule antenne d'émission. Cependant, pour satisfaire la demande toujours plus grande des utilisateurs, les opérateurs sont amenés à déployer un grand nombre de petites stations de base [42]. Ce

déploiement permet de réduire la distance entre les utilisateurs et la station de base et donc d'augmenter le débit.

Dans un réseau dense, c'est à dire avec une grande densité de stations de base, on a des stations de base qui vont ré-utiliser les mêmes fréquences pour communiquer. C'est pourquoi, les interférences entre-cellules sont un facteur limitant dans ce type de réseaux. Par conséquent, l'allocation de ressources doit être réalisée en les prenant en compte. Celle-ci peut-être réalisée de manière coordonnée ou non-coordonnée [43]. Dans la plupart des travaux sur l'allocation de ressources et de puissance, on suppose que les stations de base sont coordonnées [44]. C'est à dire qu'elles échangent de l'information pour servir les utilisateurs, un réseau dans lequel les stations se coordonnent pour échanger de l'information s'appelle un *Coordinated Multipoint* (COMP) [45].

Une autre solution pour satisfaire la demande toujours croissante en données mobiles est d'utiliser des systèmes multi-antennes ou *Multiple-Input-Multiple-Output* (MIMO). Il a été prouvé qu'augmenter le nombre d'antennes d'émission et de réception permet d'augmenter le débit d'une communication entre une station de base et un utilisateur [46]. L'utilisation d'un grand nombre d'antennes transforme évidemment le problème d'allocation de ressources et de puissance [47].

2.6 Consommation d'énergie des réseaux mobiles

Lorsqu'on augmente le nombre de stations de base et d'antennes d'émission, il n'est plus suffisant de minimiser la puissance d'émission en résolvant le problème posé dans l'équation (2.4). En effet, pour réduire de façon significative la consommation du réseau mobile il est nécessaire de réduire à la fois :

- L'énergie dynamique de la station de base. C'est une fonction de la puissance d'émission.
- L'énergie statique de la station de base qui n'est pas directement dépendante de la puissance d'émission.

La puissance dynamique d'une station de base peut être réduite en ajustant la puissance d'émission. Au contraire, réduire la puissance statique requiert l'utilisation de mécanismes de mise en veille [5].

Pour pouvoir proposer des mécanismes de mise en veille, il est nécessaire d'avoir un modèle de consommation global de la station de base. La puissance instantanée consommée par une station de base est généralement supposée être une fonction affine de la puissance d'émission :

$$\begin{cases} P_{BS} = P_0 + m_p P_{TX}, & \text{si la station de base est active, avec } 0 < P_{TX} \leq P_{max}, \\ P_{BS} = P_s, & \text{en veille.} \end{cases} \quad (2.13)$$

Cette modélisation a d'abord été proposée dans [48] et a été vérifiée expérimentalement dans [49]. Dans l'équation (2.13), P_0 désigne la puissance statique consommée par la station lorsque celle-ci est en train de servir les utilisateurs. P_{Tx} est la puissance d'émission de la station de base et m_p est le coefficient de proportionnalité qui relie la puissance d'émission et la puissance consommée. c'est l'inverse du rendement de la chaîne RF. En d'autres termes, pour une puissance d'émission égale à 1 W, la station de base consomme $P_0 + m_p$ W. P_s est la puissance consommée par la station de base en veille. P_s est inférieure à P_0 . Nous rappelons que $P_{\max}$ désigne la puissance d'émission maximale de la station de base.

Dans [49], les auteurs donnent des valeurs réalistes pour P_0 , P_s , m_p et $P_{\max}$ pour différents types de stations de base :

Tableau 2.1 – Puissance consommée par les différents types de stations de base et dans les différentes phases. P_s en veille, P_0 en statique et $P_{\max}$ ($\geq P_{Tx}$)

type de SB	$P_{\max}$ (W)	P_0 (W)	m_p	P_s (W)
Macro	20.0	130.0	4.7	75.0
Micro	6.3	56.0	2.6	39.0
Pico	0.13	6.8	4.0	4.3
Femto	0.05	4.8	8.0	2.9

2.7 Mise en veille des stations de base et transmission discontinue

Il existe deux types de veille différents qui permettent de réduire la consommation des stations de base. La première consiste à mettre en veille pendant des temps longs la station de base. La seconde est la transmission discontinue dans laquelle certains éléments de la station de base sont mis en veille pendant des temps très courts [50].

2.7.1 Mise en veille prolongée

Avec la première solution, on met en veille prolongée certaines stations de base [5, 50]. Le choix des stations de base à éteindre peut se faire de façon centralisée [51], c'est à dire qu'un contrôleur central décide des stations de base à éteindre. Une seconde solution revient à éteindre les stations de base de façon décentralisée [52]. Dans ce cas là, chaque station de base décide si elle doit être active ou mise en veille. Le choix des stations de base à allumer ou éteindre peut se faire pour minimiser l'énergie consommée ou maximiser l'efficacité énergétique [53]. Pour que le réseau ait toujours la même couverture et puisse toujours servir

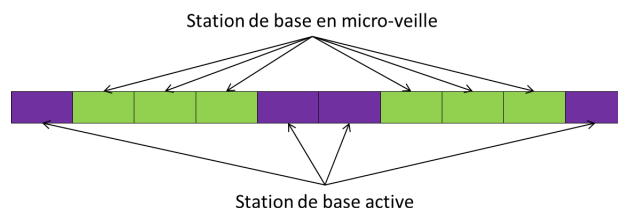


Figure 2.5 – Dans le standard LTE, la station de base peut être mise en veille pendant six des dix sous-trames.

tous les utilisateurs, les stations de base actives doivent ajuster leur zone de couverture afin de compenser la mise en veille de leurs voisines. C'est ce qu'on appelle le *cell-breathing* [54].

Mettre en veille prolongée une station de base demande du temps. C'est pourquoi, la mise en veille prolongée des stations de base ne peut pas être instantanée. Le choix du nombre de stations de base à allumer peut se faire en fonction des évolutions moyennes journalières du trafic [55] et sur des échelle temporelles longues (de l'ordre de la minute). En d'autres termes, on ne peut pas mettre en veille prolongée une station de base de manière immédiate.

2.7.2 Transmission discontinue

Le principe de la transmission discontinue, ou *Discontinuous Transmission* (Cell DTx) [8] est de mettre en veille une partie de la station de base pendant des temps très courts. Pendant ces micro-veilles, certains éléments de la chaîne RF tels que l'amplificateur de puissance sont éteints [56]. Une station de base qui utilise de la transmission discontinue alterne entre des périodes pendant lesquelles elle sert des utilisateurs et d'autres pendant lesquelles elle est en veille.

Afin de mieux comprendre cette solution, nous détaillons, dans ce paragraphe, son utilisation possible dans le standard LTE [57]. Dans ce standard, certaines des sous-trames d'une trame LTE peuvent être des trames *Multicast-Broadcast Single-Frequency Network* (MBSFN). L'une des particularités de ces sous-trames est que les symboles pilotes sont envoyés en début de sous-trame. On n'a donc aucun symbole pilote transmis pendant la quasi-totalité de la sous-trame. Il a donc été envisagé dans [8] de mettre en veille l'amplificateur de puissance pendant cette période. Une trame LTE est composée de dix sous-trames et jusqu'à six de ces trames peuvent être des sous-trames MBSFN. On peut donc, si la charge de la station de base n'est pas trop élevée, comme indiqué sur la Figure 2.5, mettre la station de base en micro-veille pendant six sous-trames.

Le Cell DTx ne peut être utilisé que si le temps d'allumage et d'éteignage de l'amplificateur de puissance est court devant la durée d'une sous-trame et si le pic de consommation à l'allumage n'est pas trop important. D'après [56], le temps de mise en route et d'arrêt d'un

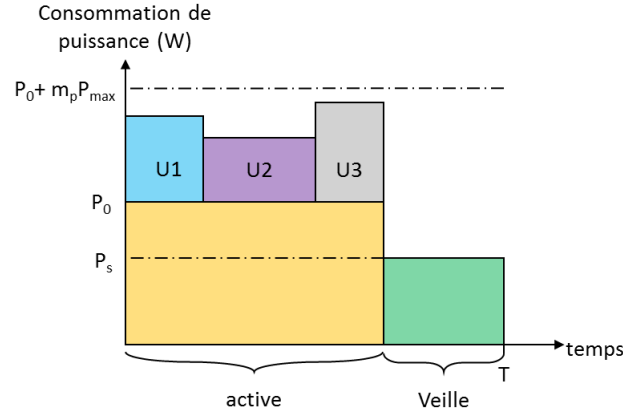


Figure 2.6 – Modélisation du problème d'allocation de puissance en TDMA avec de la transmission discontinue. Dans cet exemple, la station de base sert trois utilisateurs $\{U1, U2, U3\}$ avant de passer en veille.

amplificateur de puissance sont respectivement égaux à 65 et 25 μs . On en déduit que le temps d'allumage et d'éteignage de l'amplificateur n'est pas limitant pour le standard LTE dans lequel les sous-trames durent 1 ms. Nous n'avons pas trouvé de chiffres permettant d'assurer que le pic de consommation à l'allumage n'est pas trop important. Cependant, la prise en compte du Cell DTx dans le très récent standard 5G-NR (5G-*New Radio*) [58] nous permet de dire que le Cell DTx est utilisable en pratique et que le pic de consommation à l'allumage pourra être négligé dans nos travaux.

Comme cela a été fait dans [59], il est tout à fait envisageable de combiner à la fois des veilles prolongées et du Cell DTx dans des réseaux cellulaires.

2.8 Allocation de puissance et de ressources avec du Cell DTx

Si l'on regarde la Figure 2.5, on se rend compte qu'utiliser de la transmission discontinue change l'allocation de ressources et de puissance. En effet, en plus de se poser les questions habituelles de répartition des utilisateurs et d'allocation de puissance, avec du Cell DTx, on doit aussi se poser la question du nombre de sous-trames pendant lesquelles la station de base doit être active.

2.8.1 Allocation de puissance à une seule station de base

La question de l'allocation de ressources et de puissance avec du Cell DTx a d'abord été posée dans [60]. Dans cet article, les auteurs se focalisent sur le problème d'allocation de puissance qu'ils posent en TDMA. Ils utilisent le modèle présenté sur la Figure 2.6.

On voit sur la Figure 2.6, que la trame de durée T est divisée en deux parties. Une première pendant laquelle les utilisateurs sont servis et une seconde pendant laquelle la station de base est en veille. On voit sur cette figure, qu'avec du Cell DTx, résoudre le problème d'allocation de puissance revient à déterminer, à la fois :

- La puissance instantanée utilisée pour servir l'utilisateur.
- Le temps pendant lequel l'utilisateur est servi, que nous appellerons temps de service de l'utilisateur.

En effet, lorsque du Cell DTx est utilisé, la puissance d'émission et le temps de service sont liés et ne peuvent pas être optimisés séparément.

Dans [60], les auteurs ont posé le problème d'allocation de puissance pour minimiser l'énergie consommée par la station de base. Ce qui revient à minimiser la puissance moyenne consommée pendant une trame. Pour poser ce problème, on va supposer qu'on a une connaissance parfaite du canal. On note respectivement P_{TX}^k et t_k la puissance d'émission et le temps de service de l'utilisateur k . De plus, $\mu_k = \frac{t_k}{T}$ désigne la proportion de trame² utilisée pour servir l'utilisateur k . Avec du Cell DTx, on peut supposer que le débit d'un utilisateur est égal à la capacité de canal lorsqu'il est servi, et égal à 0 le reste du temps. Dans ce cas là, ce débit peut s'écrire :

$$C_k = B\mu_k \log_2 \left(1 + \frac{P_{TX}^k g_k}{N} \right), \quad \forall k \in \llbracket 1; N_u \rrbracket. \quad (2.14)$$

Dans cette équation, B désigne la largeur de bande qui est identique pour tous les utilisateurs. Le problème d'allocation de puissance pour minimiser la puissance moyenne de la station de base peut donc s'écrire :

$$\min_{P_{TX}^k, \mu_k} \left[P_s \left(1 - \sum_{k=1}^{N_u} \mu_k \right) + P_0 \sum_{k=1}^{N_u} \mu_k + m_p \sum_{k=1}^{N_u} \mu_k P_{TX}^k \right], \quad (2.15a)$$

$$\text{S.à } C_k = B\mu_k \log_2 \left(1 + \frac{P_{TX}^k g_k}{N} \right), \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (2.15b)$$

$$0 \leq P_{TX}^k \leq P_{\max}, \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (2.15c)$$

$$\mu_k \geq 0, \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (2.15d)$$

$$\sum_{k=1}^{N_u} \mu_k \leq 1. \quad (2.15e)$$

La dernière contrainte de ce système d'équations permet d'assurer que le temps de service de tous les utilisateurs n'est pas plus long que la durée de chaque trame. Sous cette forme

2. Dans la suite, le terme temps de service sera aussi utilisé pour désigner μ_k

le problème de (2.15) n'est pas convexe (en particulier la fonction objectif est une somme de produits de variables, et donc n'est pas convexe). Dans [60], les auteurs ont utilisé la contrainte de capacité pour réécrire la fonction objectif uniquement en fonction des μ_k . Pour cela, on utilise l'expression de (2.15b) pour exprimer $P_{T_X}^k$ en fonction de μ_k :

$$P_{T_X}^k = \frac{N}{g_k} \left(2^{\frac{C_k}{B\mu_k}} - 1 \right). \quad (2.16)$$

Cette expression permet de retirer les variables $P_{T_X}^k$ du problème d'optimisation. Le problème devient alors :

$$\min_{\mu_k} P_m = \left(\sum_{k=1}^{N_u} \mu_k \right) P_0 + \left(1 - \sum_{k=1}^{N_u} \mu_k \right) P_s + m_p N \sum_{k=1}^{N_u} \frac{\mu_k}{g_k} \left(2^{\frac{C_k}{B\mu_k}} - 1 \right), \quad (2.17a)$$

$$\text{S.à } \frac{N}{g_k} \left(2^{\frac{C_k}{B\mu_k}} - 1 \right) \leq P_{\max}, \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (2.17b)$$

$$\mu_k \geq 0, \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (2.17c)$$

$$\sum_{k=1}^{N_u} \mu_k \leq 1. \quad (2.17d)$$

Il a été prouvé dans [60] que ce problème est convexe. Mais les auteurs de ce papier n'ont pas proposé d'algorithme spécifique ou de formes analytiques pour sa résolution. Il est intéressant de noter qu'une solution analytique partielle a été proposée dans [61]. En effet, lorsqu'aucune contrainte n'est saturée, la valeur optimale de μ_k s'écrit :

$$\mu_k = \frac{C_k \ln(2)}{B} \frac{1}{\mathcal{W} \left(e^{-1} \left[\frac{g_k}{N} \frac{P_0 - P_s}{m_p} - 1 \right] \right) + 1}. \quad (2.18)$$

Dans cette équation, $\mathcal{W}$ désigne la fonction $\mathcal{W}$ de Lambert [62], qui est la bijection réciproque de la fonction $x \mapsto xe^x$.

Si on s'intéresse maintenant au problème d'allocation de ressources et de puissance en OFDMA, le problème a principalement été traité dans deux articles :

- Dans [63], les auteurs proposent une solution pour la résolution du problème d'allocation de ressources et de puissance en OFDMA avec du Cell DTx. Dans leur algorithme, ils utilisent la solution du problème TDMA (2.17) pour calculer le nombre de slots temps-fréquence alloués à chaque utilisateur (un slot a la bande passante d'un bloc de ressources et dure une demi sous-trame comme indiqué sur la Figure 2.2). Ensuite ils appliquent la méthode proposée dans [39] pour la répartition des utilisateurs dans les blocs de ressources. L'inconvénient de cette solution est que le nombre de sous-frames pendant lesquelles la station de base est active est calculé de façon sous-optimale.

- Dans [64], les auteurs proposent de résoudre le problème d'allocation de ressources et de puissance en OFDMA, sans passer par la résolution TDMA. La modélisation proposée dans l'article est intéressante, mais les justifications mathématiques qui mènent les auteurs à une solution ne sont pas convaincantes.

On se rend compte ici qu'aucun algorithme suffisamment efficace pour être implémenté dans une station de base n'a été proposé dans la littérature pour résoudre le problème d'allocation de puissance en TDMA. De plus, aucun algorithme n'a été proposé pour résoudre le problème d'allocation de ressources et de puissance avec du Cell DTx en OFDMA de manière optimale.

C'est pour cette raison que nous proposerons des algorithmes pour la résolution des problèmes d'allocation de ressources et de puissance en TDMA en supposant que le canal est plat, et en OFDMA. En particulier, nous proposerons un algorithme pour résoudre le problème d'allocation de puissance avec du Cell DTx en TDMA dans le Chapitre 3 et un algorithme pour résoudre le problème d'allocation de ressources et de puissance avec du Cell DTx en OFDMA dans le Chapitre 4.

2.8.2 Cell DTx dans des réseaux à plusieurs cellules

La transmission discontinue a aussi été étudiée dans des réseaux denses dans lesquels des stations de base voisines utilisent les mêmes bandes de fréquence. Le potentiel du Cell DTx dans ce type de réseaux a été étudié grâce à de la géométrie stochastique dans [65, 66]. Dans ces deux articles, les auteurs considèrent que l'utilisation du DTx est aléatoire et que la puissance d'émission de chaque station de base est constante. Dans les articles [67, 68], les auteurs évaluent la densité des stations de base qui permet de maximiser l'efficacité énergétique du réseau lorsque de la transmission discontinue est utilisée.

Si on se focalise maintenant sur le problème d'allocation de ressources et de puissance dans des réseaux à plusieurs cellules, le problème d'allocation de ressources est traité de manière heuristique et sans regarder le problème d'allocation de puissance dans [69, 70]. Le problème d'allocation de puissance, avec du Cell DTx, a été traité dans [71, 72]. Dans ces deux articles, l'allocation de puissance est réalisée par chaque station de base de sorte à minimiser sa propre consommation énergétique. Pour l'allocation de ressources, deux solutions différentes sont proposées. Dans [71], les auteurs proposent de former des coalitions entre les stations de base adjacentes. Une fois les coalitions formées, ces stations de base peuvent se coordonner pour limiter les interférences. Dans [72], les auteurs formulent le problème d'allocation de ressources comme un problème de théorie des jeux non-coopératif.

Dans ce manuscrit, nous étudierons aussi le problème d'allocation de puissance en TDMA avec du Cell DTx dans un réseau composé de plusieurs cellules. Cette étude est menée dans le Chapitre 3

2.9 Conclusions

Dans ce chapitre, nous avons proposé un état de l'art pour le problème d'allocation de ressources et de puissance dans un réseau mobile sobre en énergie. Nous montrons comment a été posé et résolu le problème d'allocation de ressources et de puissance en OFDMA. Ensuite, nous avons vu que, lorsque le nombre de stations de base augmente, il n'est pas suffisant de minimiser la puissance d'émission pour minimiser la consommation des réseaux mobiles. Dans ce cas là, il faut utiliser des mécanismes de mise en veille. Nous avons ensuite présenté l'état de l'art du problème d'allocation de ressources et de puissance avec du Cell DTx.

Cette présentation de l'état de l'art nous montre qu'aucun algorithme permettant de résoudre efficacement le problème d'allocation de ressources et de puissance avec Cell DTx en TDMA n'a été proposé dans la littérature. Nous allons proposer un algorithme efficace pour résoudre ce problème dans le Chapitre 3. Dans le Chapitre 4, nous proposerons un algorithme pour résoudre le problème d'allocation de ressources et de puissance en OFDMA avec du Cell DTx.

Chapitre 3

Allocation de puissance et transmission discontinue en TDMA

Sommaire

3.1	Problème d'optimisation et reformulation	46
3.2	Résolution du problème	48
3.2.1	Méthodologie	48
3.2.2	Première étape : gestion des contraintes sur le temps de service minimum	49
3.2.3	Seconde étape : gestion de la contrainte sur le temps total	52
3.2.4	Algorithme final	56
3.3	Résultats numériques	57
3.3.1	Pour une macro-cellule	57
3.3.2	Pour une femto-cellule	61
3.4	Dans un réseau composé de plusieurs cellules	62
3.4.1	Introduction	62
3.4.2	Modèle d'interférence	62
3.4.3	Paramètres de simulation	63
3.4.4	Comportement de notre solution dans un réseau à plusieurs cellules	64
3.4.5	Stratégie optimale irréalisable	66
3.4.6	Evaluation de la méthode proposée Algorithme 3	69
3.4.7	Réduction des interférences pour améliorer les performances	70
3.4.8	Puissance d'émission optimale	71
3.5	Conclusions	75

L'objectif de ce chapitre est de proposer un algorithme efficace pour résoudre le problème d'allocation de ressource et de puissance avec Cell DTx en TDMA. Dans tout ce chapitre, nous supposons que le canal de tous les utilisateurs est plat. Par conséquent, comme expliqué

Section 2.3.1, il est suffisant de traiter le problème d'allocation de puissance. Nous nous plaçons d'abord dans un scénario à une seule cellule. Dans ce contexte, nous proposons un algorithme permettant de résoudre le problème posé de manière optimale. Dans un second temps, nous étudierons le comportement de cette solution dans un réseau composé de plusieurs cellules.

Les travaux présentés dans ce chapitre ont fait l'objet de deux publications dans des conférences internationales [73, 74], d'une publication dans une conférence nationale [75] et d'un article de revue [76].

3.1 Problème d'optimisation et reformulation

On considère une station de base qui sert plusieurs utilisateurs. Pour traiter le problème d'allocation de puissance en TDMA avec du Cell DTx, on se place dans le même scénario que [60]. C'est à dire que l'on suppose une station de base qui sert N_u utilisateurs indexés par l'indice k ; le canal est supposé plat et la station de base a une connaissance des coefficients des canaux des utilisateurs notés g_k . La station de base va ajuster la puissance d'émission et le temps de service des utilisateurs afin de minimiser sa puissance consommée tout en leur fournissant une capacité donnée. Dans ce cas là, on se retrouve avec le problème introduit dans le chapitre précédent équation (2.17). Nous rappelons ici son expression :

$$\min_{\mu_k} P_m = \left(\sum_{k=1}^{N_u} \mu_k \right) P_0 + \left(1 - \sum_{k=1}^{N_u} \mu_k \right) P_s + m_p N \sum_{k=1}^{N_u} \frac{\mu_k}{g_k} \left(2^{\frac{C_k}{B\mu_k}} - 1 \right), \quad (3.1a)$$

$$\text{S.à } \frac{N}{g_k} \left(2^{\frac{C_k}{B\mu_k}} - 1 \right) \leq P_{\max}, \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (3.1b)$$

$$\mu_k \geq 0, \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (3.1c)$$

$$\sum_{k=1}^{N_u} \mu_k \leq 1. \quad (3.1d)$$

Pour faciliter la résolution du problème nous allons introduire les variables $\mu_k^{\min}, \forall k \in \llbracket 1; N_u \rrbracket$. $\mu_k^{\min}$ désigne le temps minimum qu'il faut à la station de base pour servir l'utilisateur k . En d'autres termes, c'est le temps qu'il faut à la station de base pour servir l'utilisateur k lorsque la puissance d'émission est maximale :

$$\mu_k^{\min} = \mu^{\min}(C_k, g_k) = \frac{C_k}{B \log_2 \left(1 + \frac{P_{\max} g_k}{N} \right)}. \quad (3.2)$$

Les variables $\mu_k^{\min}$ ont deux intérêts. Le premier est qu'elles permettent de vérifier rapidement si les utilisateurs peuvent être servis ou non pendant la durée de la trame. En

effet, les $\mu_k^{\min}$ peuvent être évalués rapidement et, dans le cas où :

$$\sum_{k=1}^{N_u} \mu_k^{\min} \geq 1, \quad (3.3)$$

il n'est pas possible de servir les utilisateurs pendant cette trame, la station de base devra, soit réduire leur qualité de service, soit refuser de servir certains d'entre eux.

Le second avantage des variables $\mu_k^{\min}$ est qu'elles permettent de réduire le nombre de contraintes du problème de l'équation (3.1). En réalité, pour s'assurer que la puissance d'émission de l'utilisateur k soit plus faible que $P_{\max}$, il suffit de s'assurer que son temps de service est plus long que $\mu_k^{\min}$. En effet, servir un utilisateur pendant un temps plus court que $\mu_k^{\min}$ nécessite d'utiliser une puissance d'émission plus élevée que $P_{\max}$, et vice versa. On obtient alors le problème d'optimisation suivant :

$$\min_{\mu_k} P_m = P_s + \left(\sum_{k=1}^{N_u} \mu_k \right) (P_0 - P_s) + m_p N \sum_{k=1}^{N_u} \frac{\mu_k}{g_k} \left(2^{\frac{C_k}{B\mu_k}} - 1 \right), \quad (3.4a)$$

$$\text{S.à } \mu_k \geq \mu_k^{\min} = \frac{C_k}{B \log_2 \left(1 + \frac{P_{\max} g_k}{N} \right)}, \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (3.4b)$$

$$\sum_{k=1}^{N_u} \mu_k \leq 1. \quad (3.4c)$$

Grâce à notre réécriture, le problème de l'équation (3.4) n'a plus que $N_u + 1$ contraintes au lieu de $2N_u + 1$.

Par ailleurs, la fonction objectif donnée par l'équation (3.4a) est la somme de deux termes :

- La puissance statique moyenne qui est égale à $P_s + \left(\sum_{k=1}^{N_u} \mu_k \right) (P_0 - P_s)$. Cette puissance ne dépend pas directement de la puissance d'émission et est une fonction croissante du temps de service total $\sum_{k=1}^{N_u} \mu_k$. Pour minimiser la puissance statique moyenne, il faut servir chaque utilisateur pendant un temps le plus court possible, c'est-à-dire égal à $\mu_k^{\min}$.
- La puissance dynamique moyenne, qui est égale à $m_p N \sum_{k=1}^{N_u} \frac{\mu_k}{g_k} \left(2^{\frac{C_k}{B\mu_k}} - 1 \right)$. Elle décroît lorsque les valeurs de μ_k augmentent. Pour minimiser la puissance dynamique moyenne, il faut servir les utilisateurs pendant toute la trame, sans passer en veille.

Par conséquent, trouver l'allocation de puissance optimale revient à trouver le compromis entre ces deux puissances moyennes.

3.2 Résolution du problème

Dans cette section, nous proposons un algorithme de résolution efficace pour le problème d'allocation de puissance en TDMA. Cet algorithme opère en deux étapes. Dans la suite, nous commencerons par exprimer notre méthodologie de résolution avant de détailler les deux étapes.

3.2.1 Méthodologie

L'algorithme de résolution proposé s'appuie sur la convexité du problème d'allocation de puissance. Pour prouver cette convexité, il suffit de prouver que la matrice Hessienne de la fonction objectif est définie positive. On a :

$$\frac{\partial^2 P_m}{\partial \mu_k \partial \mu_l} = 0, \quad \forall k \neq l, \quad (3.5)$$

et,

$$\frac{\partial^2 P_m}{\partial^2 \mu_k} = \frac{m_p N}{g_k} 2^{\frac{C_k}{B \mu_k}} \left(\frac{C_k}{B} \right)^2 \frac{\ln(2)^2}{\mu_k^3} \geq 0 \quad \mu_k \in [\mu_k^{\min}; 1]. \quad (3.6)$$

La matrice Hessienne de la fonction objectif définie par (3.4) est diagonale, à éléments diagonaux positifs. Elle est donc définie positive. De plus, les contraintes du problème d'optimisation étudié sont des fonctions linéaires. Par conséquent, le problème étudié est convexe. On peut donc le résoudre en utilisant les conditions de KKT.

Comme détaillé dans le chapitre précédent, avec ces conditions, chaque contrainte d'inégalité peut soit être une inégalité stricte (dans ce cas là elle n'est pas considérée dans la résolution), soit être une contrainte d'égalité qui doit être prise en compte dans la résolution du problème. Pour gérer les contraintes, on procède donc comme suit :

- Une fois la solution sans contrainte explicitée, on vérifie les contraintes sur les $\mu_k^{\min}$ données par l'équation (3.4b).
- Après avoir pris en compte les contraintes sur les $\mu_k^{\min}$, il faut vérifier si le temps de service total ne dépasse pas la durée de la trame. C'est-à-dire vérifier la contrainte de l'équation (3.4c).

La méthode de résolution est résumée sur la Figure 3.1.

Dans la suite de ce chapitre, nous détaillons les deux étapes de notre algorithme de résolution.

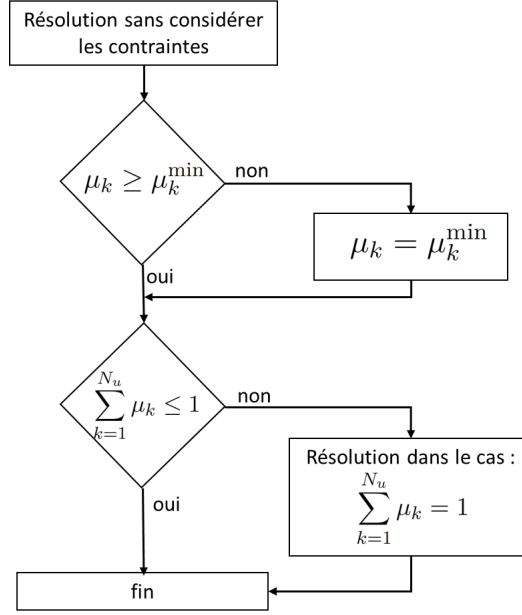


Figure 3.1 – Schéma de résolution du problème posé.

3.2.2 Première étape : gestion des contraintes sur le temps de service minimum

Pour cette première étape, on applique les conditions de KKT. C'est-à-dire qu'on commence par résoudre le problème sans considérer les contraintes avant de regarder, parmi les contraintes de l'équation (3.4b), lesquelles sont saturées.

Pour la résolution sans contrainte, on calcule le point d'annulation du gradient de la fonction objectif. On obtient une valeur de μ_k potentiellement optimale notée μ_k^{opt1} [61] :

$$\mu_k^{\text{opt1}} = \mu^{\text{opt1}}(C_k, g_k) = \frac{C_k \ln(2)}{B} \frac{1}{\mathcal{W}\left(e^{-1} \left[\frac{g_k}{N} \frac{P_0 - P_s}{m_p} - 1 \right] \right) + 1}. \quad (3.7)$$

Dans cette équation, la fonction $(C, g) \mapsto \mu^{\text{opt1}}(C, g)$ est la fonction qui, à une valeur de C et g données fait correspondre le temps de service optimal sans contrainte.

Une fois que l'on a obtenu cette expression, on doit vérifier quelles contraintes sur les temps de service minimums sont saturées. En d'autres termes, on doit vérifier pour chaque utilisateur s'il doit être servi pendant un temps minimum μ_k^{min} avec une puissance d'émission égale à P_{max} , ou s'il doit être servi pendant un temps plus long égal à μ_k^{opt1} avec une puissance d'émission moins élevée. Pour cela, on introduit une première fonction $g \mapsto \rho^M(g)$:

$$\rho^M(g) = \frac{g P_{\text{max}}}{N}, \quad g \geq 0. \quad (3.8)$$

On peut maintenant définir $\rho_k^M = \rho^M(g_k)$ qui est le RSB qu'aurait l'utilisateur k s'il était servi avec toute la puissance d'émission. De plus on définit le rapport r :

$$r = \frac{P_0 - P_s}{m_p P_{\max}} \quad (3.9)$$

r est le rapport entre la plage de variation de la puissance statique et la plage de variation de la puissance dynamique. La puissance statique de la station de base est soit égale à P_0 , soit à P_s . C'est pourquoi, le numérateur de r est égal à $P_0 - P_s$. Pour le dénominateur, la puissance d'émission varie dans l'intervalle $[0; P_{\max}]$. La puissance dynamique a donc une plage de variation égale à $m_p P_{\max}$.

Nous verrons dans la suite de cette section que ce coefficient r permet de caractériser le type de station de base. C'est-à-dire que l'on pourra trier les stations de base en fonction de sa valeur.

On remarque que l'on peut écrire la fonction μ^{opt1} , comme une fonction de C , r et ρ^M et la fonction $\mu^{\min}$ comme une fonction de C et ρ^M :

$$\mu^{\min}(C, \rho^M) = \frac{C}{B \log_2(1 + \rho^M)}, \quad (3.10)$$

$$\mu^{\text{opt1}}(C, \rho^M, r) = \frac{C \ln(2)}{B} \frac{1}{\mathcal{W}(e^{-1} [\rho^M r - 1]) + 1}. \quad (3.11)$$

Dans la suite, on note $\rho_{\lim}^M$ la valeur de ρ^M qui vérifie :

$$\mu^{\text{opt1}}(C, \rho^M, r) = \mu^{\min}(C, \rho^M). \quad (3.12)$$

On a la proposition suivante :

Proposition 3.1. *La fonction $f : r \mapsto \rho_{\lim}^M$ est bijective et strictement croissante sur $\mathbb{R}_+^*$*

Démonstration. Pour prouver que la fonction $f : r \mapsto \rho_{\lim}^M$ est bijective et strictement croissante, on s'intéresse à la fonction $h : \rho_{\lim}^M \mapsto r$. Pour $\rho_{\lim}^M > 0$, on dérive de l'équation (3.12) :

$$r = h(\rho_{\lim}^M) = \left[1 + \frac{1}{\rho_{\lim}^M} \right] \ln(1 + \rho_{\lim}^M) - 1. \quad (3.13)$$

La dérivée de cette équation par rapport à $\rho_{\lim}^M$ est égale à :

$$h'(\rho_{\lim}^M) = -\frac{1}{(\rho_{\lim}^M)^2} \ln(1 + \rho_{\lim}^M) + \frac{1}{\rho_{\lim}^M}. \quad (3.14)$$

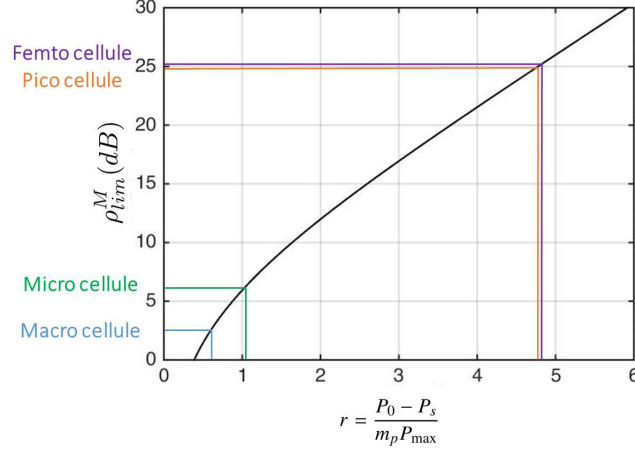


Figure 3.2 – Fonction $\rho_{\text{lim}}^M(r)$ qui permet, pour une station de base, de savoir si un utilisateur doit être servi pendant $\mu_k^{\min}$ ou pendant μ_k^{opt1} . Cette figure présente un nouveau mode de classification des stations de base spécifique au problème d'allocation de puissance avec Cell DTx.

En appliquant l'inégalité de concavité du logarithme :

$$\ln(1+x) < x \quad \forall x > 0, \quad (3.15)$$

on prouve que $h : \rho_{\text{lim}}^M \mapsto r$ est une fonction strictement croissante sur $]0; +\infty[$. Elle est donc bijective et a une bijection réciproque. Celle-ci, $f : r \mapsto \rho_{\text{lim}}^M$ est aussi bijective et strictement croissante.

□

La proposition 3.1 nous dit que pour chaque type de station de base, c'est-à-dire pour chaque valeur de r , on a une valeur de ρ^M appelée ρ_{lim}^M qui permet de choisir entre $\mu^{\min}$ et μ^{opt1} . En effet, pour chaque utilisateur, en comparant ρ_k^M avec ρ_{lim}^M , on peut savoir s'il doit être servi avec une puissance d'émission maximale égale à P_{max} pendant $\mu_k^{\min}$ ou avec une puissance plus faible pendant un temps plus long égal à μ^{opt1} . Si $\rho_k^M < \rho_{\text{lim}}^M$, l'utilisateur est servi pendant $\mu_k^{\min}$ et dans le cas contraire, pendant μ_k^{opt1} .

On trace, sur la Figure 3.2, la fonction $\rho_{\text{lim}}^M(r)$. Sur cette courbe, on ajoute la valeur de r des différentes stations de base décrites dans le Tableau 2.1.

Pour chacune des stations de base listées dans la Table 2.1, on peut calculer r en fonction des paramètres P_0 , P_s , m_p et P_{max} . Une fois r calculé, on peut calculer la valeur de ρ_{lim}^M de la station de base. Nous avons listé ces valeurs dans le Tableau 3.1.

Le nouveau mode de représentation proposé par la Figure 3.2 permet de classer les stations de base. En effet, une analyse de la Figure 3.2 et du Tableau 3.1 nous montre que l'on peut séparer les stations de base en deux catégories. On a d'un côté les pico-cellules et femto-cellules,

Tableau 3.1 – Valeur de ρ_{lim}^M pour les différents types de stations de base.

Type de SB	r	ρ_{lim}^M (dB)
Macro	0.58	2.6
Micro	1.04	6.2
Pico	4.80	25.1
Femto	4.75	24.9

ces stations de base ont une puissance d'émission faible et donc une valeur de r élevée (aux alentours de 5). Pour ces stations de base, la puissance statique est dominante. Pour minimiser la consommation de la station de base, il est suffisant de minimiser la puissance statique, c'est-à-dire, de servir tous les utilisateurs pendant $\mu_k^{\min}$ avec une puissance d'émission égale à $P_{\max}$.

On s'intéresse maintenant aux stations de base à grande couverture pour lesquelles la valeur de ρ_{lim}^M est faible. Cet ensemble contient les macro-cellules et micro-cellules. Pour ces stations de base, la puissance d'émission maximale $P_{\max}$ est élevée. Par conséquent, la puissance dynamique est du même ordre de grandeur que la puissance statique. Comme r est proche de 1, on doit à la fois ajuster la puissance statique et la puissance dynamique de la station de base. Pour cela, on doit ajuster la puissance d'émission (c'est-à-dire faire du contrôle de puissance) pour ajuster à la fois la puissance statique et la puissance dynamique. Pour ce type de stations de base, la quasi-totalité des utilisateurs doit être servi pendant μ_k^{opt1} (tous ceux ayant une valeur de ρ_k^M supérieure à ρ_{lim}^M).

Finalement, quel que soit le type de station de base, pour calculer le temps de service de chaque utilisateur, lorsqu'on connaît la valeur de g_k , on procède comme indiqué dans l'Algorithme 1.

Algorithme 1

```

if  $\rho_k^M > \rho_{\text{lim}}^M$  then
   $\mu_k = \mu_k^{\text{opt1}}$  (calculée avec (3.7))
else
   $\mu_k = \mu_k^{\min}$  (calculée avec (3.2))
end if

```

3.2.3 Seconde étape : gestion de la contrainte sur le temps total

Une fois que l'étape décrite dans la section précédente est réalisée, il nous faut vérifier si la contrainte de l'équation (3.4c) est saturée. C'est-à-dire si, après la première étape, le temps de service n'est pas plus long que la durée de la trame.

Si la contrainte est satisfaite, on peut servir les utilisateurs pendant les temps de service calculés précédemment. Dans le cas contraire, et toujours d'après les conditions de KKT, le problème peut être réécrit sous la forme suivante :

$$\min_{\mu_k} P_m = P_0 + m_p N \sum_{k=1}^{N_u} \frac{\mu_k}{g_k} \left(2^{\frac{C_k}{B\mu_k}} - 1 \right), \quad (3.16a)$$

$$\text{S.à } \mu_k \geq \mu_k^{\min}, \forall k \in \llbracket 1; N_u \rrbracket, \quad (3.16b)$$

$$\sum_{k=1}^{N_u} \mu_k = 1. \quad (3.16c)$$

Dans ce cas là, à la fin de la première étape de l'algorithme, on a obtenu des temps de service qui sont trop longs et qui ne vérifient pas la contrainte de l'équation (3.16b). On va donc devoir servir les utilisateurs pendant des temps plus courts et qui vérifient $\sum_{k=1}^{N_u} \mu_k = 1$.

Si on ne prend pas en compte les contraintes sur les temps de service minimum de l'équation (3.16b), le Lagrangien du problème s'écrit :

$$\mathcal{L} = P_0 + \sum_{k=1}^{N_u} \mu_k \frac{m_p N}{g_k} \left(2^{\frac{C_k}{B\mu_k}} - 1 \right) + \lambda \left(\sum_{k=1}^{N_u} \mu_k - 1 \right). \quad (3.17)$$

Dans cette équation, λ désigne le multiplicateur de Lagrange. En dérivant ce Lagrangien par rapport à μ_k , on obtient :

$$\mu_k^{\text{opt2}}(\lambda) = \mu^{\text{opt2}}(\lambda, C_k, g_k) = \frac{C_k \ln(2)}{B} \frac{1}{\mathcal{W} \left(e^{-1} \left[\frac{g_k}{N} \frac{\lambda}{m_p} - 1 \right] \right) + 1}. \quad (3.18)$$

En comparant l'équation (3.18) avec l'équation (3.7), le multiplieur de Lagrange λ peut prendre plusieurs valeurs. Lorsque le problème est résolu (et que les contraintes sont satisfaites), il est égal à la valeur théorique de $P_0 - P_s$ qui réduit suffisamment le temps de service total pour que la somme des μ_k soit égale à 1. Par conséquent, plus λ est élevé, plus les temps de service sont courts.

Pour $\lambda = P_0 - P_s$, on a $\mu_k^{\text{opt2}}(P_0 - P_s) = \mu_k^{\text{opt1}}$, pour cette valeur de λ , $\sum_{k=1}^{N_u} \mu_k > 1$. Dans ce cas là les temps de service des utilisateurs sont trop longs. Pour avoir des valeurs des μ_k^{opt2} suffisamment petites, la valeur de λ doit donc être plus grande que $P_0 - P_s$. En augmentant la valeur de λ , on ne va pas pouvoir réduire le temps de service de tous les utilisateurs. En effet, le temps de service de chaque utilisateur ne peut pas être plus court que $\mu_{\min}^k$. Par conséquent, on ne pourra pas réduire le temps de service des utilisateurs qui étaient servis pendant $\mu_{\min}^k$ à la fin de la première étape. Leur temps de service sera donc toujours égal à $\mu_{\min}^k$ à la fin de la seconde étape de l'algorithme. On notera N_m le nombre

d'utilisateurs servis pendant $\mu_k^{\min}$ à la fin de la première étape. De plus, nous noterons λ_{opt} la valeur optimale de λ qui vérifie $\sum_{k=1}^{N_u} \mu_k = 1$.

En augmentant la valeur de λ , on va avoir d'autres utilisateurs qui seront servis pendant $\mu_k^{\min}$. On propose une résolution en deux étapes :

- On va d'abord identifier l'ensemble des utilisateurs qui doivent être servis pendant $\mu_k^{\min}$.
- Une fois que c'est fait, on se retrouve avec deux catégories d'utilisateurs. Les premiers sont servis pendant $\mu_k^{\min}$. Les seconds sont servis pendant $\mu_k^{\text{opt2}}(\lambda_{\text{opt}})$. Pour calculer le temps de service de ces derniers, il nous faut donc calculer la valeur λ_{opt} qui permet de satisfaire la contrainte de l'équation (3.16c).

Tout d'abord, lorsque λ augmente, la valeur de chacun des μ_k^{opt2} diminue. Pour chaque utilisateur, cette valeur va diminuer jusqu'à atteindre $\mu_k^{\min}$. On a donc une valeur de λ à partir de laquelle l'utilisateur k est servi pendant $\mu_k^{\min}$. Dans la suite on note λ_k la valeur de λ qui vérifie $\mu_k^{\text{opt2}} = \mu_k^{\min}$:

$$\lambda_k = \lambda_m(g_k) = \left[\frac{m_p N}{g_k} + m_p P_{\max} \right] \ln \left(1 + \frac{P_{\max} g_k}{N} \right) - m_p P_{\max}. \quad (3.19)$$

Dans l'équation (3.19), $\lambda_m(g)$ est la fonction qui vérifie $\mu^{\min}(C, g) = \mu^{\text{opt2}}(C, g, \lambda)$. L'équation (3.19) nous donne la valeur de λ à partir de laquelle l'utilisateur k doit être servi pendant $\mu_k^{\min}$. En d'autres termes, si $\lambda_k \leq \lambda_{\text{opt}}$, l'utilisateur k doit être servi pendant $\mu_k^{\min}$. Dans le cas contraire, il doit être servi pendant μ_k^{opt2} . Si on étudie la fonction $\lambda_m(g)$, elle a les propriétés suivantes :

- λ_m est une fonction croissante de g (ceci peut être prouvé par une analyse similaire à celle faite dans la preuve de la Proposition 3.1).
- μ^{opt2} et par conséquent le temps de service total sont des fonctions décroissantes de λ_m .

Ces deux points nous permettent de dire que si un utilisateur est servi pendant $\mu_k^{\min}$, c'est à dire que $\lambda_k < \lambda_{\text{opt}}$, tous les utilisateurs avec un coefficient de canal inférieur devront être servis pendant $\mu_k^{\min}$. On peut donc identifier l'ensemble des utilisateurs à servir pendant $\mu_k^{\min}$ en trouvant celui qui a le plus grand coefficient de canal (g_k) qui est servi pendant $\mu_k^{\min}$. On désignera par l'indice $N_{\min}$ cet utilisateur.

Pour cet utilisateur, $\sum_{k=1}^{N_u} \mu_k(\lambda_{N_{\min}}) \geq 1$ et pour l'utilisateur ayant un coefficient d'atténuation directement supérieur (l'utilisateur $N_{\min} + 1$), $\sum_{k=1}^{N_u} \mu_k(\lambda_{N_{\min}+1}) \leq 1$.

On peut donc identifier l'ensemble des utilisateurs servis pendant le temps de service minimum en triant les utilisateurs par ordre croissant de g_k et en calculant successivement la somme $\sum_{k=1}^{N_u} \mu_k(\lambda_{N_{\min}})$. Ceci nous permet d'identifier l'utilisateur $N_{\min}$ qui est le dernier pour lequel cette somme est supérieure à 1. L'Algorithme 2 détaille ces étapes.

Algorithme 2 Utilisateurs servis pendant $\mu_{k \min}$

```

1: Trier les utilisateurs par ordre croissant de  $g_k$ 
2:  $\lambda = P_0 - P_s$  et  $p := N_m$ 
3: while  $\sum_{k=1}^{N_u} \mu_k(\lambda) > 1$  do
4:    $p = p + 1$ 
5:    $\lambda = \lambda_p$  (calculé avec l'équation (3.19))
6:   for  $k$  de 1 à  $N_u$  do
7:      $\mu_k(\lambda) = \max(\mu_k^{\min}, \mu_k^{\text{opt2}}(\lambda))$ 
8:   end for
9: end while
10:  $N_{\min} = p - 1$ 

```

Une fois que l'on a identifié l'ensemble des utilisateurs qui doivent être servis pendant $\mu_k^{\min}$, on doit calculer la valeur λ_{opt} qui permet de satisfaire la contrainte de l'équation (3.16b). C'est à dire qui vérifie :

$$\sum_{k=1}^{N_{\min}} \mu_k^{\min} + \sum_{k=N_{\min}+1}^{N_u} \mu_k^{\text{opt2}}(\lambda_{\text{opt}}) - 1 = 0. \quad (3.20)$$

Après l'étape précédente, on sait que $\sum_{k=1}^{N_u} \mu_k(\lambda_{N_{\min}}) \geq 1$ et que $\sum_{k=1}^{N_u} \mu_k(\lambda_{N_{\min}+1}) \leq 1$. On en déduit que λ_{opt} appartient à l'intervalle $[\lambda_{N_{\min}}; \lambda_{N_{\min}+1}]$. On peut trouver la valeur optimale λ_{opt} en utilisant une méthode de recherche de 0 à une seule variable telle que la méthode de la bisection, Regula Falsi, la méthode de Brent, la méthode du point fixe ou la méthode de Newton [77]. Dans la liste des méthodes cités ci-dessus, la méthode de Newton est celle qui garantit la vitesse de convergence la plus rapide. Cependant, cette méthode ne converge pas toujours.

Si on note $\lambda_{N_{\min}}$ et $\lambda_{N_{\min}+1}$ les bornes de l'intervalle de recherche, λ_0 le point de départ de l'algorithme et f la fonction dont on cherche la racine, on peut s'assurer de la convergence de la méthode de Newton en vérifiant les propriétés suivantes :

- f n'a qu'une seule racine et $f(\lambda_0) > 0$,
- $f \in \mathcal{C}^2 [\lambda_{N_{\min}}; \lambda_{N_{\min}+1}]$, $f'(\lambda) \neq 0$ and $f''(\lambda) > 0$
- $f(\lambda_{N_{\min}})f(\lambda_{N_{\min}+1}) < 0$

Proposition 3.2. Si $\lambda_0 = \lambda_{N_{\min}}$, la méthode de Newton appliquée à la fonction :

$$f : \lambda \mapsto \sum_{i=1}^{N_{\min}} \mu_{i \min} + \sum_{i=N_{\min}+1}^{N_u} \mu_{i \text{opt2}}(\lambda) - 1 \quad (3.21)$$

converge vers son unique racine.

Démonstration. $f \in \mathcal{C}^2 [\lambda_{N_{min}}; \lambda_{N_{min}+1}]$, de plus, les résultats précédemment obtenus montrent que, $f(\lambda_{N_{min}+1}) < 0$ et $f(\lambda_0) = f(\lambda_{N_{min}}) > 0$. De plus,

$$\mu'_{k_{opt2}}(\lambda) = -\frac{e^{-1}C_k \ln(2)g_k}{NBm_p} \frac{e^{-\mathcal{W}\left(e^{-1}\left[\frac{g_k}{N}\frac{\lambda}{m_p}-1\right]\right)}}{\left(\mathcal{W}\left(e^{-1}\left[\frac{g_k}{N}\frac{\lambda}{m_p}-1\right]\right)+1\right)^3} < 0 \quad (3.22)$$

f' est donc la somme de fonctions strictement négatives. C'est donc une fonction strictement négative. De plus, le théorème de la bijection nous dit que f n'a qu'une seule racine.

Finalement,

$$\mu''_{k_{opt2}}(\lambda) = \frac{e^{-2}C_k \ln(2)g_k^2}{N^2Bm_p^2} \frac{e^{-2\mathcal{W}\left(e^{-1}\left[\frac{g_k}{N}\frac{\lambda}{m_p}-1\right]\right)}}{\left(\mathcal{W}\left(e^{-1}\left[\frac{g_k}{N}\frac{\lambda}{m_p}-1\right]\right)+1\right)^5} \left(\mathcal{W}\left(e^{-1}\left[\frac{g_k}{N}\frac{\lambda}{m_p}-1\right]\right)+4\right) > 0 \quad (3.23)$$

Par conséquent, f'' est strictement positive. Nous avons vérifié les trois conditions suffisantes de convergence et, par conséquent, prouvé la convergence de la méthode de Newton. $\square$

Une fois que la méthode de Newton a été appliquée, on connaît les valeurs des μ_k optimales. Pour obtenir les puissances d'émissions correspondantes, on utilise l'équation (2.16).

3.2.4 Algorithme final

Afin d'améliorer les explications proposées ci-dessus, nous proposons d'en faire un résumé.

Nous avons une station de base qui sert N_u utilisateurs. Elle a connaissance du coefficient d'atténuation du canal g de chacun de ces utilisateurs. Sans aucune restriction, nous pouvons indexer les utilisateurs par ordre croissant de g .

$$g_1 \leq g_2 \leq \dots \leq g_{N_u}. \quad (3.24)$$

Pendant la première étape (Algorithme 1), on va séparer ces utilisateurs en deux catégories. Les N_m qui ont le coefficient de canal le plus faible sont servis pendant $\mu_k^{\min}$, les autres pendant μ_k^{opt1} . La séparation des utilisateurs se faisant en comparant le RSB de chaque utilisateur avec un seuil qui dépend des caractéristiques de la station de base.

Une fois cette première étape terminée, on peut se retrouver dans deux situations. Soit le temps de service total est inférieur à la durée de la trame. Dans ce cas là on pourra servir tous les utilisateurs. Soit on a un temps de service plus long. Et dans ce cas, on doit appliquer la deuxième partie de l'algorithme.

Dans cette deuxième partie, on va d'abord identifier les $N_{\min}$ utilisateurs qui sont servis pendant $\mu_k^{\min}$ et ensuite calculer λ_{opt} , la valeur du multiplicateur de Lagrange qui permet de satisfaire la contrainte :

$$\sum_{k=1}^{N_{\min}} \mu_k^{\min} + \sum_{k=N_{\min}+1}^{N_u} \mu_k^{\text{opt2}}(\lambda_{\text{opt}}) - 1 = 0. \quad (3.25)$$

Pour trouver la valeur de $N_{\min}$, on utilise le fait que pour chaque utilisateur, on a une valeur λ_k qui vérifie $\mu_k^{\min} = \mu_k^{\text{opt2}}(\lambda_k)$. On a :

$$\underbrace{\lambda_1 \leq \dots \leq \lambda_{N_{\min}}}_{\text{Utilisateurs servis pendant } \mu^{\min}} \leq \lambda_{\text{opt}} \leq \underbrace{\lambda_{N_{\min}+1} \leq \dots \leq \lambda_{N_u}}_{\text{Utilisateurs servis pendant } \mu^{\text{opt2}}}. \quad (3.26)$$

Vu que l'on ne connaît pas λ_{opt} à priori, on applique la somme des μ_k à ces inégalités :

$$\underbrace{\sum_{k=1}^{N_u} \mu_k(\lambda_1) \leq \dots \leq \sum_{k=1}^{N_u} \mu_k(\lambda_{N_{\min}})}_{\text{Utilisateurs servis pendant } \mu^{\min}} \leq 1 \leq \underbrace{\sum_{k=1}^{N_u} \mu_k(\lambda_{N_{\min}+1}) \leq \dots \leq \sum_{k=1}^{N_u} \mu_k(\lambda_{N_u})}_{\text{Utilisateurs servis pendant } \mu^{\text{opt2}}}. \quad (3.27)$$

On peut donc identifier les $N_{\min}$ utilisateurs servis pendant $\mu_k^{\min}$ en comparant $\sum_{k=1}^{N_u} \mu_k(\lambda_k)$ et 1. Une fois ces utilisateurs identifiés, on peut calculer λ_{opt} en utilisant un algorithme de recherche de zéro. L'Algorithme 3 décrit l'algorithme complet.

3.3 Résultats numériques

Nous évaluons maintenant les performances de l'algorithme de résolution optimal proposé. La courbe 3.2 nous montre que l'on peut classer les stations de base en deux catégories en fonction de la valeur de r . Nous conduisons donc des simulations numériques pour deux types de cellule différents : les macro-cellules qui ont une faible valeur de r et les femto-cellules qui ont une valeur de r élevée.

3.3.1 Pour une macro-cellule

On suppose une macro-cellule qui sert des utilisateurs répartis dans l'espace suivant un processus de Poisson. La distance entre la station de base et les utilisateurs varie entre 300 et 1500 m. En moyenne, on a 10 utilisateurs actifs dans la zone de couverture de la station de base. On suppose que la largeur de bande est égale à 10 MHz avec une fréquence centrale égale à 2 GHz. Le facteur de bruit et le gain d'antenne sont considérés égaux à 0 dans nos

Algorithme 3 Calcul du temps de service des utilisateurs

```
1: Trier les utilisateurs par ordre croissant de  $g_k$ 
2: for  $k$  de 1 à  $N_u$  do
3:   if  $\rho_k^M > \rho_{\text{lim}}^M$  then
4:      $\mu_k = \mu_k^{\text{opt1}}$  (calculée avec (3.7))
5:   else
6:      $\mu_k = \mu_k^{\text{min}}$  (calculée avec (3.2))
7:   end if
8: end for
9: if  $\sum_{k=1}^{N_u} \mu_k > 1$  then
10:   $\lambda = P_0 - P_s$  et  $p := N_m$ 
11:  while  $\sum_{k=1}^{N_u} \mu_k(\lambda) > 1$  do
12:     $p := p + 1$ 
13:     $\lambda = \lambda_p$  (calculé avec l'équation (3.19))
14:    for  $k$  de 1 à  $N_u$  do
15:       $\mu_k(\lambda) = \max\{\mu_k^{\text{min}}, \mu_k^{\text{opt2}}(\lambda)\}$ 
16:    end for
17:  end while
18:   $N_{\text{min}} = p - 1$ 
19:  Calculer  $\lambda_{\text{opt}}$  avec la méthode de Newton
20:  for  $k$  de  $N_{\text{min}} + 1$  à  $N_u$  do
21:     $\mu_k = \mu_k^{\text{opt2}}(\lambda_{\text{opt}})$ 
22:  end for
23: end if
```

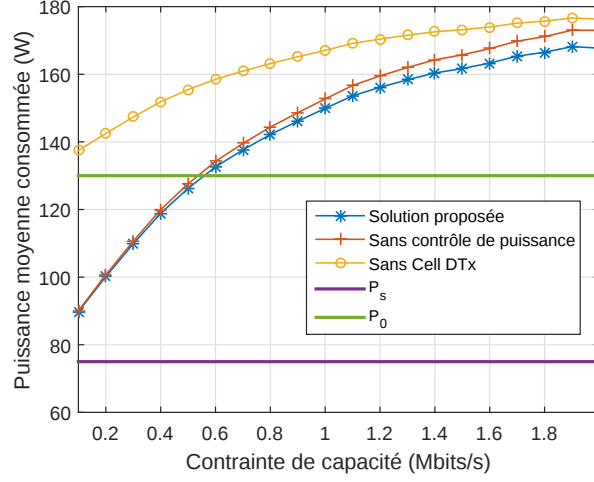


Figure 3.3 – Évaluation des performances de l'algorithme proposé lorsque une macro cellule sert en moyenne 10 utilisateurs, en termes de puissance moyenne consommée.

simulations. Le *path-loss* est calculé avec le modèle Winner II *suburban macro-cell* (C1) [78]. On suppose que tous les utilisateurs ont la même contrainte de capacité.

Aucun algorithme n'a été proposé dans la littérature pour la résolution du problème d'allocation de puissance avec du Cell DTx. C'est pourquoi, afin d'évaluer les performances de la solution optimale, on compare seulement les stratégies suivantes :

- La stratégie optimale proposée
- Une stratégie dans laquelle le Cell DTx est utilisé sans contrôle de puissance. C'est-à-dire dans laquelle tous les utilisateurs sont servis pendant $\mu_k^{\min}$ avec une puissance d'émission égale à $P_{\max}$.
- Une stratégie dans laquelle le Cell DTx n'est pas utilisé et dans laquelle la puissance d'émission est allouée de manière optimale.

On peut rappeler que la stratégie sans contrôle de puissance minimise la puissance statique moyenne consommée par la station de base, la stratégie sans Cell DTx minimise la puissance dynamique moyenne consommée par la station de base alors que la stratégie optimale proposée est un compromis entre ces deux puissances moyennes.

On affiche sur la Figure 3.3 la puissance moyenne consommée par la station de base avec les différentes stratégies comparées en fonction de la contrainte de capacité. On remarque que, comme attendu, la solution proposée offre les meilleures performances. Pour le cas des macro-cellules, lorsque la contrainte de capacité est faible, la puissance dynamique est faible comparée à la puissance statique. C'est pourquoi l'utilisation du Cell DTx réduit fortement la puissance moyenne consommée par rapport à une solution sans Cell DTx. Dans ce cas là,

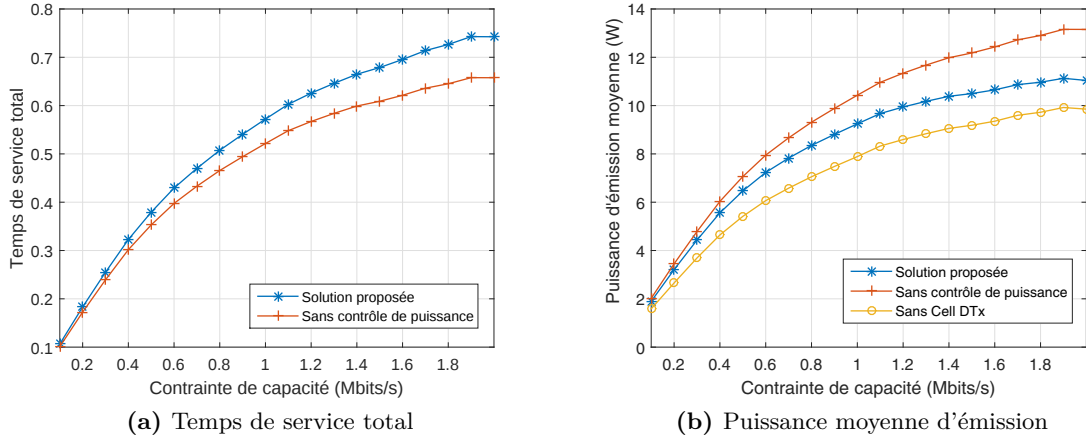


Figure 3.4 – (a) Temps total pendant lequel la station de base est active. On voit que ce temps est plus long lorsque le contrôle de puissance est utilisé. (b) Puissance d'émission moyenne de la station de base pendant une trame.

on a un gain en puissance moyenne qui représente 35 % de la puissance totale consommée de la station de base.

Lorsque la contrainte de capacité augmente, la puissance dynamique augmente et la solution optimale s'éloigne de la solution sans contrôle de puissance (qui minimise la puissance statique moyenne) et se rapproche de la solution sans Cell DTx (qui minimise la puissance dynamique moyenne). Dans le scénario simulé, la solution proposée apporte un gain en consommation qui représente jusqu'à 3,5 % de la puissance moyenne totale de la station de base.

On s'intéresse maintenant au temps pendant lequel la station de base est active et à la puissance moyenne d'émission qui sont représentés sur la Figure 3.4. On peut voir qu'avec la solution proposée, le temps de service est plus long qu'avec la solution sans contrôle de puissance. Au contraire, la puissance d'émission moyenne de la station de base est plus faible avec la solution proposée comparée à celle sans contrôle de puissance.

On voit, sur la Figure 3.4b, que la puissance d'émission moyenne avec la solution optimale proposée est comprise entre la puissance d'émission moyenne avec la solution sans Cell DTx et la puissance d'émission moyenne sans contrôle de puissance. Par conséquent, dans un réseau dans lequel plusieurs cellules adjacentes utilisent la même bande de fréquence, en moyenne, la solution proposée va générer moins d'interférence sur les cellules voisines que la solution sans contrôle de puissance mais plus que la solution sans Cell DTx.

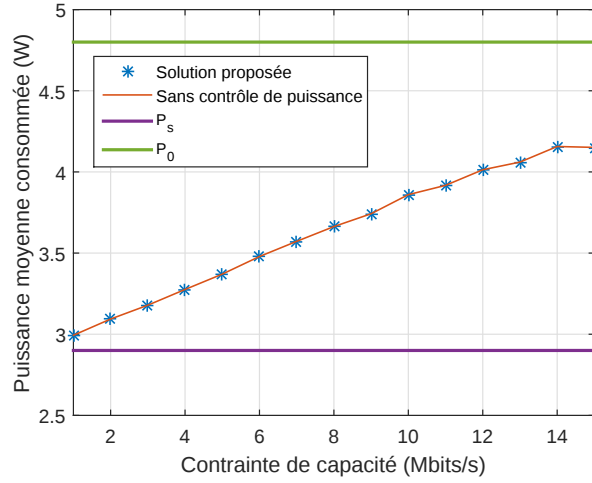


Figure 3.5 – Puissance moyenne consommée par une femto-cellule. On remarque que pour ce type de station de base le contrôle de puissance n’a aucun effet.

3.3.2 Pour une femto-cellule

Pour l’étude des stations de base à faible zone de couverture, on considère une femto-cellule qui sert en moyenne 4 utilisateurs. Les utilisateurs sont répartis dans l’espace suivant un processus de Poisson et la distance entre un utilisateur et la station de base varie entre 0 et 50 m. Les *path-losses* sont calculés avec le modèle *ITU for indoor propagation* [79]. La fréquence centrale est de 2 GHz, la largeur de bande est égale à 10 MHz et le nombre d’étages entre l’utilisateur et la station de base est uniformément réparti entre 0 et 5. Le facteur de bruit et le gain d’antenne sont supposés égaux à 0.

On compare la puissance moyenne consommée par la station de base avec la solution optimale proposée et avec la solution sans contrôle de puissance. Les résultats sont affichés sur la Figure 3.5. Comme anticipé dans la Section 3.2.2, on voit sur cette courbe que la solution sans contrôle de puissance et la solution optimale consomment la même énergie. Ce résultat s’explique par le fait que, pour les stations de base considérées, la puissance statique est bien plus grande que la puissance dynamique. Par conséquent, il suffit de minimiser la puissance statique moyenne pour minimiser la puissance moyenne consommée par la station de base.

Sur la courbe 3.5, nous n’avons pas comparé l’algorithme proposé avec d’autres algorithmes car aucun algorithme n’a été proposé dans la littérature pour résoudre le problème d’allocation de puissance avec du Cell DTx.

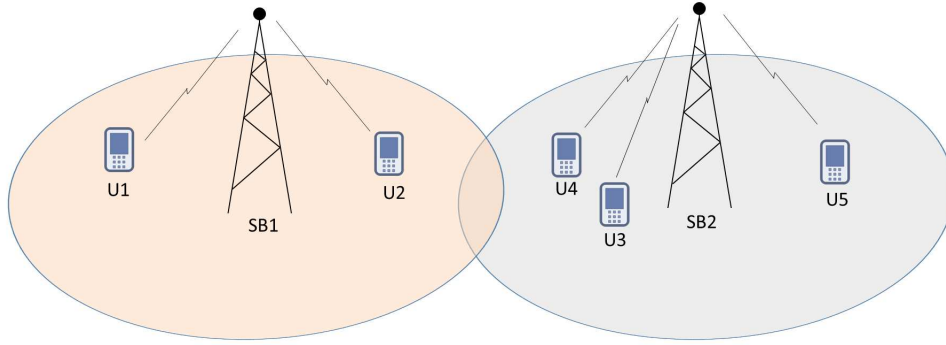


Figure 3.6 – Un exemple de réseau mobile dans lequel on a $N_{BS} = 2$ stations de base et $N_u = 5$ utilisateurs. Dans ce cas là, $S_1 = \{1, 2\}$ et $S_2 = \{3, 4, 5\}$.

3.4 Dans un réseau composé de plusieurs cellules

3.4.1 Introduction

L'objectif de cette section est d'analyser les performances de la solution proposée dans les sections précédentes dans un réseau dense dans lequel plusieurs cellules proches utilisent la même bande de fréquence. Dans un tel contexte, les interférences entre cellules voisines diminuent la qualité de service des utilisateurs. Il est donc nécessaire d'adapter la solution proposée pour les prendre en compte.

On se focalise sur les stations de base à large couverture, pour lesquelles il y a un intérêt à utiliser du contrôle de puissance, c'est-à-dire pour lesquelles il est préférable de servir les utilisateurs avec une puissance d'émission inférieure à la puissance d'émission maximale $P_{\max}$. En particulier, dans nos simulations nous considérerons une micro-cellule dont les caractéristiques sont celles données dans le Tableau 2.1.

3.4.2 Modèle d'interférence

On considère N_{BS} stations de base de micro-cellules adjacentes qui servent leurs utilisateurs en TDMA. Chacun des N_u utilisateurs dans le réseau est servi par la station de base dont il est le plus proche. Les utilisateurs sont indexés par l'indice k et les stations de base par l'indice n . Comme illustré sur la Figure 3.6, on dénote S_n l'ensemble des utilisateurs servis par la station de base n et N_n dénote le cardinal de S_n . De plus, μ_{nk} et P_{TX}^{nk} désignent respectivement la proportion de la trame utilisée pour émettre les données et la puissance d'émission utilisée par la station de base n pour l'utilisateur k . On note g_{nk} le coefficient d'atténuation du canal entre la station n et l'utilisateur k .

Le niveau d'interférence subi par un utilisateur dépend de la puissance d'émission des stations de base voisines, du coefficient du canal g_{mk} entre cet utilisateur et les stations de base voisines et de l'échelonnement des utilisateurs par ces dernières (c'est-à-dire de la répartition des utilisateurs sur la durée de la trame).

Le problème de répartition des utilisateurs a été étudié dans divers articles incluant [69]. Dans nos travaux, nous nous concentrons sur le problème d'allocation de puissance et ne regardons pas l'échelonnement des utilisateurs. Pour cela, tout comme cela a été fait dans [67, 68], nous considérons que chaque utilisateur subit les interférences générées par la puissance d'émission moyenne des stations de base voisines. Par conséquent, les interférences subies par l'utilisateur k , servi par la station de base n , ont pour expression :

$$I_k = \sum_{m=1, m \neq n}^{N_{BS}} I_{m \rightarrow k} = \sum_{m=1, m \neq n}^{N_{BS}} g_{mk} \underbrace{\sum_{l \in S_m} \mu_{ml} P_{TX}^{ml}}_{\text{Puissance d'émission moyenne de la station de base } m}. \quad (3.28)$$

Dans cette équation I_k désigne les interférences subies par l'utilisateur k . Ces interférences sont causées par les stations de base voisines indexées par m . Avec cette formule pour les interférences, le débit d'un utilisateur s'écrit :

$$C_k = B \mu_{nk} \log_2 \left(1 + \frac{P_{TX}^{nk} g_{nk}}{N + \sum_{m=1, m \neq n}^{N_{BS}} \sum_{l \in S_m} g_{mk} \mu_{ml} P_{TX}^{ml}} \right), \quad \forall k \in \llbracket 1; N_u \rrbracket. \quad (3.29)$$

Si la station de base a connaissance de la valeur du coefficient $\frac{g_{nk}}{N + I_k}$, elle peut utiliser l'Algorithme 3. La seule différence entre l'application de cet algorithme dans un réseau mono-cellulaire et dans un réseau à plusieurs cellules est que, dans un réseau à plusieurs cellules, les interférences s'ajoutent au bruit. C'est-à-dire qu'on remplace le RSB par le Rapport Signal à Interférence plus Bruit (RSIB).

L'utilisation de notre solution dans un réseau à plusieurs cellules peut se faire sans qu'aucune information ne soit échangée entre les stations de base, c'est-à-dire de manière non-coordonnée. Pour évaluer ses performances on suppose donc que les stations de base sont parfaitement synchronisées mais n'échangent pas d'information entre elles pour l'allocation de puissance.

3.4.3 Paramètres de simulation

Pour évaluer les performances de notre solution dans un réseau composé de plusieurs cellules, nous considérons deux réseaux différents :

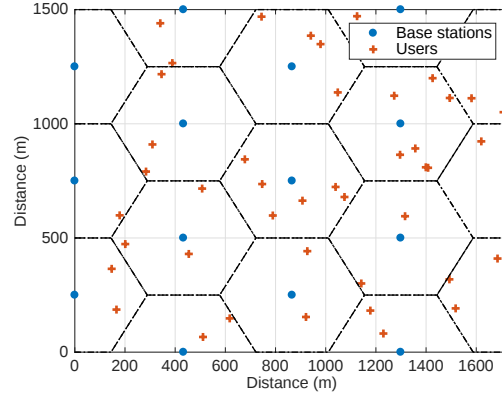


Figure 3.7 – Le réseau 2 qui est composé de 17 cellules hexagonales.

1. Le réseau 1 est uniquement composé de deux micro-cellules. Dans ce réseau les utilisateurs servis sont répartis dans un rectangle de 1000×500 m. Les coordonnées des stations de base dans ce réseau sont $(250, 250)$ et $(750, 250)$. Dans ce réseau, les utilisateurs sont distribués selon un processus de Poisson, avec en moyenne 20 utilisateurs à servir.
2. Le réseau 2 est composé de 17 cellules hexagonales comme illustré sur la Figure 3.7. Dans celui-ci la distance entre deux cellules est de 500 m. Dans ce réseau on a en moyenne 20 utilisateurs par km^2 .

Dans ces deux réseaux, du TDMA est utilisé et les utilisateurs sont servis dans une bande de 10 MHz avec une fréquence centrale de 2 GHz. Les *path-losses* sont calculés avec le modèle de Winner II *Urban micro-cell* (B1) [78].

3.4.4 Comportement de notre solution dans un réseau à plusieurs cellules

Avec la solution proposée pour l'allocation de puissance, on n'a aucune communication entre les cellules voisines. Sans coordination, on peut avoir des phases de transition pendant lesquelles toutes les stations de base mettent à jour la puissance d'émission et le temps pendant lequel elles servent les utilisateurs.

Par exemple, si un nouvel utilisateur se connecte au réseau, la station de base qui le sert va augmenter les interférences qu'elle va générer sur les stations de base voisines. Toutes ces stations de base vont donc mettre à jour leur puissance d'émission pour que la contrainte de capacité des utilisateurs qu'elles servent soit toujours satisfaite. Par conséquent, ces stations de base vont, à leur tour, augmenter les interférences qu'elles génèrent sur les stations de base voisines qui vont alors augmenter les puissances d'émission et les temps de service et ainsi de suite. On a ici un processus itératif qu'il est nécessaire d'analyser. On peut avoir deux cas :

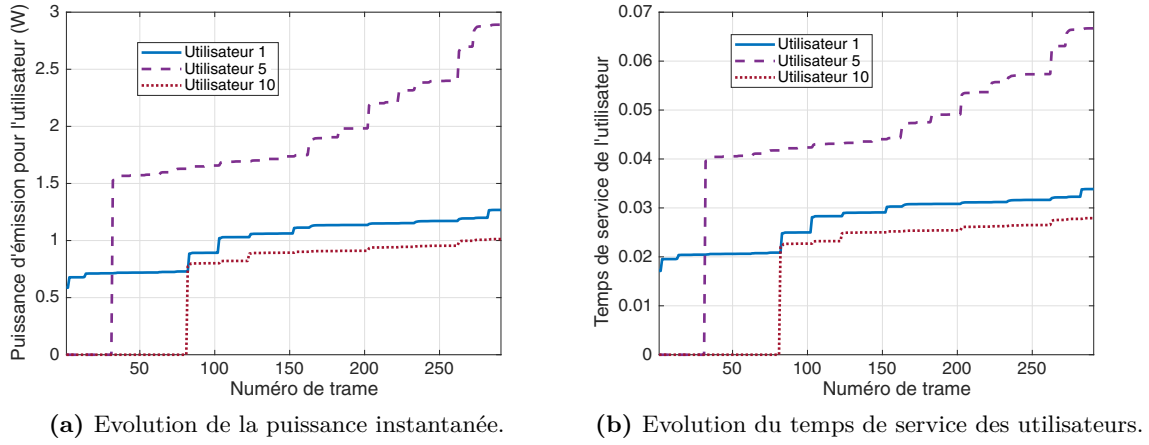


Figure 3.8 – Evolution des temps de service et des puissances instatannées lorsque de nouveaux utilisateurs se connectent au réseau.

- Si ce processus itératif converge rapidement, on peut utiliser la solution proposée dans la section 3.2 dans un réseau à plusieurs cellules.
- Si ce processus itératif ne converge pas rapidement, on va devoir ajouter de la coordination entre les cellules ou développer de nouvelles stratégies d'allocation de puissance.

Nous avons étudié numériquement l'évolution des puissances d'émission et des temps de service dans le réseau 2 (composé de 17 cellules). Dans ce réseau, on suppose que 30 utilisateurs se connectent les uns après les autres toutes les 10 trames entre 0 et 300. On indexe les utilisateurs par ordre d'arrivée. Chacun de ces utilisateurs a une contrainte de capacité égale à 3 Mbits/s. Entre chaque connexion d'un nouvel utilisateur, on a 10 trames (c'est-à-dire 10 mises à jours des puissances d'émission). Sur les Figures 3.8a et 3.8b, on affiche respectivement l'évolution de la puissance d'émission et du temps de service de certains utilisateurs. On observe sur ces figures que à chaque nouvelle connexion, on a une augmentation nette de la puissance d'émission et du temps de service avant d'avoir une stabilisation rapide de ces deux valeurs (ces valeurs peuvent être considérées comme constantes après moins de 5 trames).

Cette stabilisation rapide nous permet de conclure que lorsque la solution proposée est utilisée dans un réseau à plusieurs cellules, les stations de base adaptent rapidement la puissance d'émission et les temps de service. On peut donc utiliser notre solution dans un réseau à plusieurs cellules. De plus, les transitions sont courtes. On peut donc négliger leur durée et évaluer les performances de la solution lorsque tous les utilisateurs sont connectés.

Nous venons de voir qu'il est possible d'utiliser notre solution dans un réseau composé de plusieurs cellules. Nous devons maintenant évaluer ses performances. Pour cela, il nous faut définir des stratégies de référence. On compare notre stratégie avec :

- la stratégie sans contrôle de puissance et la stratégie sans Cell DTx qui ont été introduites précédemment. Pour avoir un intérêt, notre solution doit surpasser ces solutions naïves.
- Une stratégie optimale irréalisable. En comparant notre solution avec cet optimum idéal, on voit s'il y a un intérêt ou non à développer des stratégies plus élaborées. Cette stratégie est décrite dans la section suivante.

3.4.5 Stratégie optimale irréalisable

Dans cette section, on s'intéresse à cette stratégie optimale irréalisable. Pour que l'on puisse utiliser cette stratégie dans un réseau cellulaire, il faudrait que le réseau soit au courant de tous les *path-losses* entre chaque station de base et chaque utilisateur. Ce qui est impossible.

Pour l'étude de cet optimum irréalisable, on se place dans un réseau qui, comme le réseau 1 présenté Section 3.4.3, est composé de 2 cellules. Dans ce réseau, si on suppose que l'on a connaissance de tous les coefficients des canaux, on peut écrire le problème de minimisation de la puissance moyenne consommée par le réseau sous forme convexe. Une fois que cette convexité est établie, on peut résoudre numériquement notre problème. On obtient alors la puissance minimale nécessaire au réseau pour servir les utilisateurs. Cette puissance minimale est inaccessible mais permet d'évaluer l'optimalité de notre solution.

Dans un réseau composé de deux cellules, le problème de minimisation de la puissance moyenne consommée s'écrit :

$$\min_{\mu_{nk}, P_{TX}^{nk}} 2P_s + \sum_{n=1}^2 \left[(P_0 - P_s) \sum_{k \in S_n} \mu_{nk} + m_p \sum_{k \in S_n} \mu_{nk} P_{TX}^{nk} \right], \quad (3.30a)$$

$$\text{S.à } C_k = B\mu_{nk} \log_2 \left(1 + \frac{P_{TX}^{nk} g_{nk}}{N + g_{mk} \sum_{l \in S_m} \mu_{ml} P_{TX}^{ml}} \right), \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (3.30b)$$

$$P_{TX}^{nk} \in]0; P_{\max}], \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (3.30c)$$

$$\mu_{nk} > 0, \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (3.30d)$$

$$\sum_{k \in S_n} \mu_{nk} \leq 1, \quad \forall n \in \{1; 2\}. \quad (3.30e)$$

Dans ce problème $m = n \bmod (2) + 1$, c'est-à-dire égal à 2 lorsque n est égal à 1 et vice-versa. L'équation (3.30a) est la puissance moyenne consommée par le réseau pendant une trame. La minimisation de la puissance moyenne consommée par le réseau doit se faire en satisfaisant les contraintes sur la capacité (équation (3.30b)), sur la puissance d'émission maximale (équation (3.30c)), sur la positivité des temps de service (équation (3.30d)) et sur le temps de service total ((3.30e)).

Sous cette forme, ce problème n'est pas convexe. On peut, cependant, le réécrire sous forme convexe. Pour cela, on introduit les variables :

$$x_k = \mu_{nk} P_{TX}^{nk}, \quad \forall k \in \llbracket 1; N_u \rrbracket. \quad (3.31)$$

x_k est le produit entre la proportion de trame et la puissance d'émission moyenne utilisée pour l'utilisateur k pendant chaque trame. En introduisant x_k dans l'équation (3.30b), on obtient l'équation suivante :

$$C_k = B \mu_{nk} \log_2 \left(1 + \frac{P_{TX}^{nk} g_{nk}}{N + g_{mk} \sum_{l \in S_m} x_l} \right), \quad \forall k \in \llbracket 1; N_u \rrbracket. \quad (3.32)$$

Cette équation nous permet d'écrire :

$$P_{TX}^{nk} = \left(2^{\frac{C_k}{B \mu_{nk}}} - 1 \right) \left(\frac{N}{g_{nk}} + \frac{g_{mk}}{g_{nk}} \sum_{l \in S_m} x_l \right), \quad \forall k \in \llbracket 1; N_u \rrbracket. \quad (3.33)$$

En multipliant la $k^{\text{ième}}$ équation par μ_{nk} , on obtient :

$$x_k - \alpha_k \sum_{l \in S_m} x_l = \beta_k, \quad \forall k \in \llbracket 1; N_u \rrbracket. \quad (3.34)$$

Dans cette équation :

$$\alpha_k = \frac{g_{mk}}{g_{nk}} \mu_{nk} \left(2^{\frac{C_k}{B \mu_{nk}}} - 1 \right), \quad (3.35)$$

et,

$$\beta_k = \frac{N \mu_{nk}}{g_{nk}} \left(2^{\frac{C_k}{B \mu_{nk}}} - 1 \right). \quad (3.36)$$

L'équation (3.34) définit un système de N_u équations linéaires en x_k . En résolvant ce système, on obtient l'expression des x_k en fonction des μ_{ml} :

$$x_k = \underbrace{\beta_k}_{\text{Valeur de } x_k \text{ lorsque les interférences sont nulles}} + \underbrace{\alpha_k \left[\frac{\sum_{i \in S_n} \sum_{j \in S_m} \alpha_j \beta_i + \sum_{j \in S_m} \beta_j}{1 - \sum_{i \in S_n} \sum_{j \in S_m} \alpha_i \alpha_j} \right]}_{\text{Puissance supplémentaire liée aux interférences}}. \quad (3.37)$$

Le détail des calculs est donné Annexe A.1. Dans cette expression, l'utilisateur k est servi par la station de base n et reçoit des interférences de la station de base m . Une fois cette expression trouvée, on peut réécrire le problème de l'équation (3.30) uniquement en fonction des temps de service. Dans notre reformulation du problème, on modifie aussi les contraintes. Pour cela on utilise la propriété suivante :

Proposition 3.3. Si $\mu_{nk} > 0$, $\forall k \in \llbracket 1; N_u \rrbracket$, on a l'équivalence suivante :

$$P_{T_X}^{nk} > 0, \quad \forall k \in \llbracket 1; N_u \rrbracket \Leftrightarrow \sum_{i \in S_1} \sum_{j \in S_2} \alpha_i \alpha_j < 1.$$

L'expression des α_i étant donnée par l'équation (3.35)

Démonstration. Dans cette preuve, on prouve séparément les deux implications.

$\Rightarrow$ Si $\mu_{nk} > 0$, $\forall k \in \llbracket 1; N_u \rrbracket$ et $P_{T_X}^{nk} > 0$, $\forall k \in \llbracket 1; N_u \rrbracket$ alors $x_k > 0$, $\forall k \in \llbracket 1; N_u \rrbracket$.

Ce qui implique que $\sum_{k=1}^{N_u} x_k > 0$. La somme des x_k est égale à :

$$\sum_{k=1}^{N_u} x_k = \frac{\sum_{i \in S_1} \beta_i + \sum_{i \in S_1} \sum_{j \in S_2} \alpha_i \beta_j + \sum_{i \in S_1} \sum_{j \in S_2} \alpha_j \beta_i + \sum_{j \in S_2} \beta_j}{1 - \sum_{i \in S_1} \sum_{j \in S_2} \alpha_i \alpha_j} > 0. \quad (3.38)$$

L'équation (3.38) implique :

$$\sum_{i \in S_1} \sum_{j \in S_2} \alpha_i \alpha_j < 1. \quad (3.39)$$

$\Leftarrow$ La réciproque est immédiate, si $\sum_{i \in S_1} \sum_{j \in S_2} \alpha_i \alpha_j < 1$, et $\mu_{nk} > 0$, $\forall k \in \llbracket 1; N_u \rrbracket$ alors $x_k > 0$, $\forall k \in \llbracket 1; N_u \rrbracket$. Par conséquent, $P_{T_X}^{nk} > 0$, $\forall k \in \llbracket 1; N_u \rrbracket$. □

On a finalement le problème d'optimisation suivant qui ne dépend plus que des variables sur les temps de service :

$$\min_{\mu_{nk}} 2P_s + (P_0 - P_s) \sum_{n=1}^2 \sum_{k \in S_n} \mu_{nk} + m_p \sum_{n=1}^2 \sum_{k \in S_n} x_k, \quad (3.40a)$$

$$\text{S.à } x_k \leq \mu_{nk} P_{\max}, \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (3.40b)$$

$$\sum_{i \in S_1} \sum_{j \in S_2} \alpha_i \alpha_j < 1, \quad (3.40c)$$

$$\mu_{nk} > 0, \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (3.40d)$$

$$\sum_{k \in S_n} \mu_{nk} \leq 1, \quad \forall n \in \llbracket 1; N_{BS} \rrbracket. \quad (3.40e)$$

On a la proposition suivante :

Proposition 3.4. Le problème d'optimisation défini par l'équation (3.40) est convexe.

La démonstration est donnée dans l'Annexe A.2. Une fois cette convexité établie, on peut utiliser une méthode de résolution numérique telle qu'une descente de gradient ou une méthode de points intérieurs [40] pour trouver cet optimum.

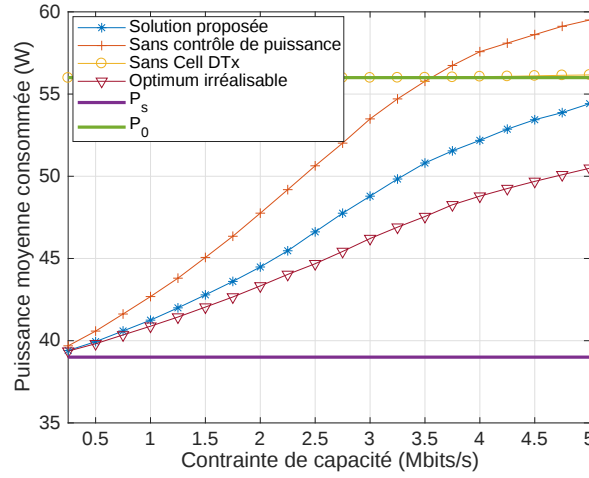


Figure 3.9 – Puissance moyenne consommée par chacune des stations de base du réseau avec les différentes stratégies comparés.

On sait maintenant que l'on peut calculer la consommation du réseau lorsque la solution optimale irréalisable est utilisée. Bien que cet optimum ne soit pas atteignable, il va nous permettre de jauger les performances de notre solution. En effet, en comparant notre solution avec cet optimum, on pourra évaluer les gains potentiels que pourraient apporter des améliorations.

3.4.6 Évaluation de la méthode proposée Algorithme 3

On a maintenant suffisamment de stratégies de comparaison pour évaluer notre solution. Pour cela, on se place dans le réseau 1, présenté dans la Section 3.4.3, et on compare :

- La solution proposée.
- La solution sans contrôle de puissance dans laquelle tous les utilisateurs sont servis avec une puissance d'émission égale à $P_{\max}$.
- La solution sans Cell DTx.
- L'optimum irréalisable qui a été présenté dans la section précédente.

Le réseau 1 étant symétrique, on peut se contenter de l'évaluation de la consommation d'une seule des deux stations de base. L'énergie totale consommée étant la somme des énergies consommées par chaque station de base, il est équivalent de comparer l'énergie consommée par le réseau et celle consommée par une seule station de base. On trace sur la Figure 3.9 l'évolution de la puissance consommée par une des stations de base en fonction de la contrainte de capacité de chaque utilisateur.

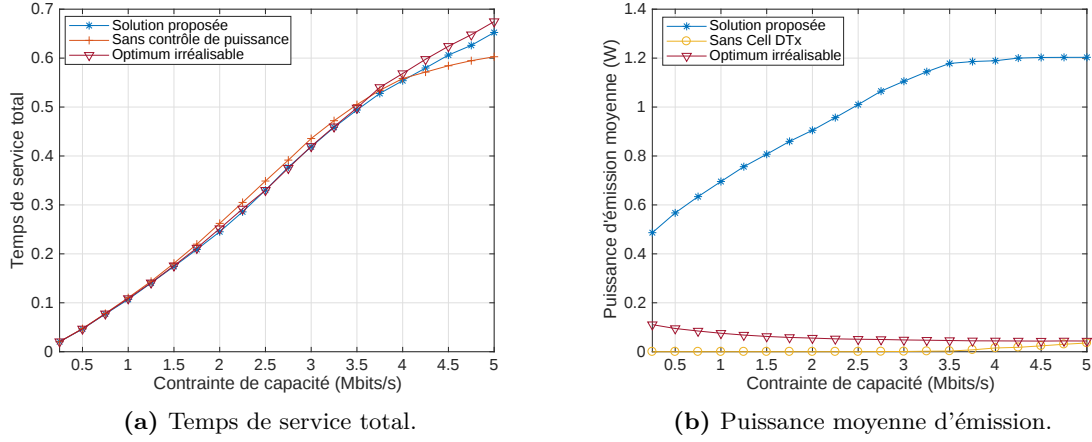


Figure 3.10 – (a) Temps moyen pendant lequel une station de base est active. (b) Puissance d'émission moyenne de la station de base lorsqu'elle est active.

Si on compare, dans un premier temps, la solution proposée et les stratégies naïves (sans contrôle de puissance ou sans Cell DTx), on remarque que, notre solution consomme toujours moins de puissance que ces deux solutions. La solution proposée, permet de gagner jusqu'à 30% de puissance moyenne consommée par rapport à une solution sans Cell DTx et jusqu'à 10% par rapport à une solution sans contrôle de puissance.

On regarde maintenant l'écart entre notre solution et la référence qu'est l'optimum irréalisable. Bien entendu, notre solution a de moins bonnes performances que la solution optimale irréalisable. L'écart entre les deux stratégies est faible lorsque la contrainte de capacité est faible et augmente avec celle-ci pour atteindre jusqu'à 8% de la puissance moyenne consommée par la station de base. Afin de mieux comprendre cet écart, on trace sur la Figure 3.10, le temps moyen pendant lequel une station de base est active ainsi que la puissance moyenne d'émission lorsque la station de base est active.

Si on compare notre solution à l'optimum irréalisable, on voit sur la Figure 3.10a que les temps de service sont quasiment identiques, alors que, la Figure 3.10b nous montre que la puissance d'émission moyenne est bien plus élevée avec notre solution. Cette observation suggère que pour réduire la puissance consommée par le réseau, il faut réduire la puissance d'émission des stations de base.

3.4.7 Réduction des interférences pour améliorer les performances

Un bon moyen de réduire la puissance d'émission des stations de base est de réduire $P_{\max}$, la puissance d'émission maximale d'une station de base. En réduisant $P_{\max}$, on réduit les interférences et donc la consommation du réseau.

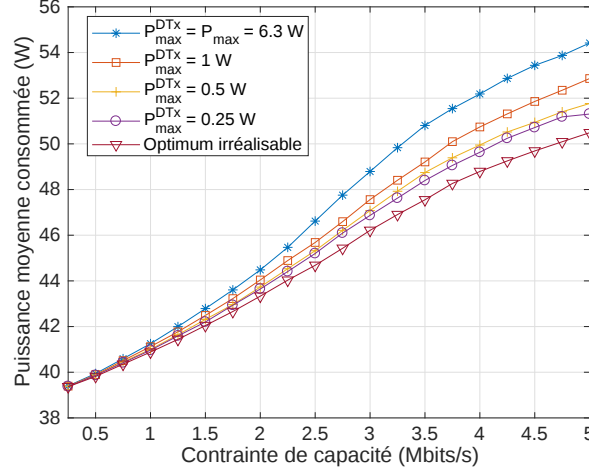


Figure 3.11 – Puissance moyenne consommée par la station de base pour différentes valeurs de $P_{\max}^{\text{DTx}}$.

On propose ici d'utiliser la solution introduite dans la Section 3.2 avec une puissance maximale plus faible. On appellera $P_{\max}^{\text{DTx}} \leq P_{\max}$ cette puissance d'émission maximale. La puissance d'émission ne pourra donc pas dépasser $P_{\max}^{\text{DTx}}$. Exception faite pour le cas où la station de base ne passe pas en mode veille pendant la trame. Dans ce cas là, la puissance d'émission maximale reste égale à $P_{\max}$. Cette exception évite d'augmenter le nombre de coupures ou *outages* dans le réseau.

Afin de valider cette solution, on se place dans le réseau 1 décrit dans la Section 3.4.3. On trace sur la Figure 3.11 la puissance moyenne consommée par la station de base en fonction de la contrainte de capacité.

On remarque qu'en diminuant la puissance d'émission maximale, on réduit la puissance consommée par la station de base. Pour les valeurs choisies dans cet exemple, on voit que la puissance moyenne consommée est d'autant plus faible que $P_{\max}^{\text{DTx}}$ est faible.

3.4.8 Puissance d'émission optimale

Nous venons de voir que réduire la puissance d'émission maximale de la station de base permet de réduire la puissance consommée par celle-ci. Dans cette section, on s'intéresse à la valeur de $P_{\max}^{\text{DTx}}$ qui permet de réduire au maximum la puissance consommée. Pour cela, nous conduirons d'abord quelques simulations avant de faire une étude analytique à faible RSIB.

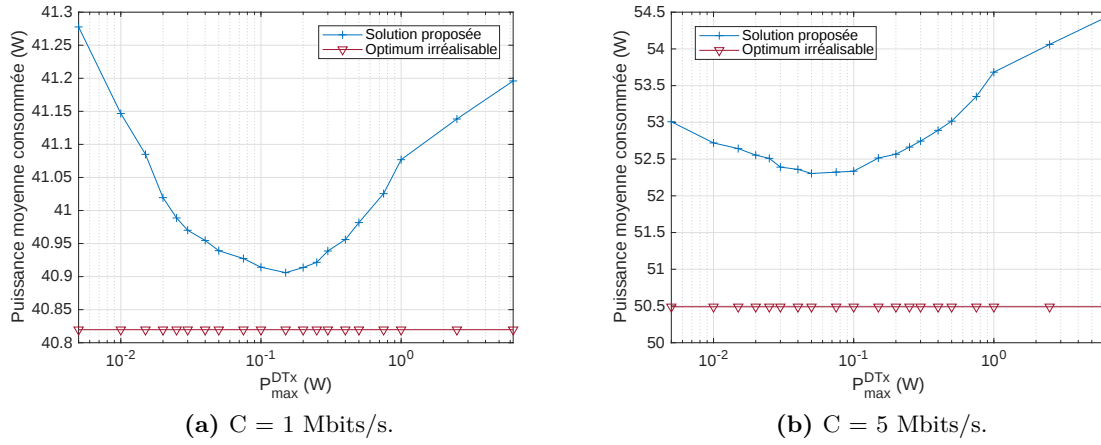


Figure 3.12 – Évolution de la puissance consommée en fonction de $P_{\max}^{\text{DTx}}$ pour (a) une contrainte de capacité de 1 Mbits/s, (b) une contrainte de capacité de 5 Mbits/s.

Résultats de simulation

On se place tout d'abord dans le réseau 1. Et on évalue la puissance consommée en moyenne par une des stations de base en fonction de $P_{\max}^{\text{DTx}}$. Les résultats sont affichés sur la Figure 3.12.

Lorsque $P_{\max}^{\text{DTx}}$ est égal à $P_{\max}$, on est loin de l'optimum. En réduisant la valeur de $P_{\max}^{\text{DTx}}$, on se rapproche de l'optimum. Pour une contrainte de capacité de 1 Mbits/s on obtient un minimum lorsque $P_{\max}^{\text{DTx}}$ est égal à 150 mW et pour une contrainte de 5 Mbits/s, lorsque $P_{\max}^{\text{DTx}} = 50$ mW. Lorsque $P_{\max}^{\text{DTx}}$ devient plus faible, la puissance consommée augmente, en effet, lorsque la puissance d'émission est trop faible, les utilisateurs sont alors servis pendant des temps longs et les stations de base ne passent que très rarement en veille pour économiser de l'énergie. Pour une contrainte de capacité égale à 5 Mbits/s, la réduction de $P_{\max}^{\text{DTx}}$ réduit de 50 % l'écart entre notre solution et la solution optimale.

On réalise maintenant des simulations dans le réseau 2, présenté Section 3.4.3, qui est composé de 17 cellules afin de confirmer nos observations. Dans ce réseau, on ne peut pas comparer notre solution à l'optimum. On trace l'évolution de la puissance consommée par la station de base au centre du réseau en fonction de $P_{\max}^{\text{DTx}}$ pour différentes valeurs de la contrainte de capacité. Les résultats sont affichés sur la Figure 3.13. On voit sur cette figure, que quelle que soit la contrainte de capacité, la puissance consommée est élevée lorsque $P_{\max}^{\text{DTx}}$ est égal à 6.3 W. Cette puissance diminue jusqu'à atteindre un minimum qui se trouve entre 15 et 40 mW. Pour une contrainte de capacité de 5 Mbits/s, en ajustant la puissance d'émission, on réduit de 8 % la puissance consommée par la station de base.

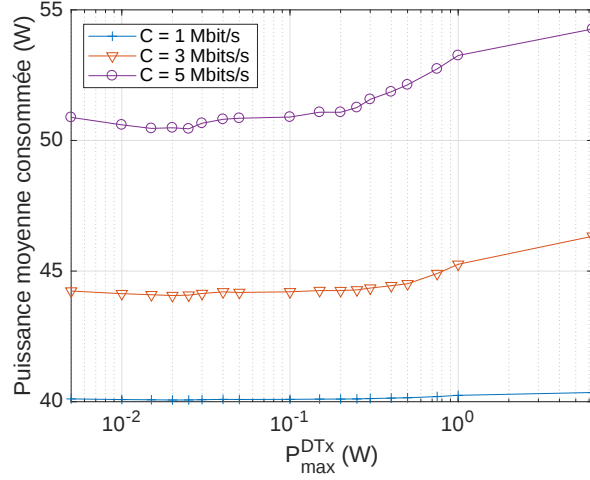


Figure 3.13 – Évolution de la puissance consommée par une station de base du réseau 2 en fonction de $P_{\max}^{\text{DTx}}$.

On remarque que nos résultats de simulation montrent que notre solution fournit les meilleures performances lorsque la puissance d'émission maximale $P_{\max}^{\text{DTx}}$ est faible (de l'ordre de la centaine de milliwatts). D'après la Figure 3.2, pour des valeurs aussi faibles, la quasi-totalité des utilisateurs sont servis avec une puissance d'émission égale à $P_{\max}^{\text{DTx}}$.

Analyse théorique à faible RSIB

Les résultats exposés ci-dessus montrent deux choses :

- Dans un réseau composé de plusieurs cellules, on peut améliorer significativement notre solution en ajustant correctement la valeur de la puissance d'émission maximale.
- Les meilleures performances sont atteintes lorsque la puissance d'émission maximale de la station de base est faible (de l'ordre de la centaine de milliwatts).

Afin de confirmer ce second point, nous réalisons une analyse analytique à faible RSIB. Si on suppose que la contrainte sur le temps de service total d'une station de base n'est saturé pour aucune station de base, tous les utilisateurs sont servis avec une puissance d'émission égale à $P_{\max}^{\text{DTx}}$ pendant un temps de service égal à $\mu_{nk}^{\min}$ égal à :

$$\mu_{nk}^{\min} = \frac{C_k}{B \log_2 \left(1 + \frac{P_{\max}^{\text{DTx}} g_{nk}}{N + I_k} \right)} = \frac{C_k}{B \log_2 (1 + \rho_{nk}^M)}. \quad (3.41)$$

Afin de trouver la valeur de $P_{\max}^{\text{DTx}}$ optimale, nous utilisons les développements limités des fonctions $f(x) = \log(1 + x)$ and $f(x) = \frac{1}{1-x}$ lorsque $x \approx 0$. Ceux-ci nous permettent de dériver une approximation de $\mu_{nk}^{\min}$ à faible RSIB :

$$\mu_{nk}^{\min}(\rho_{nk}^M) = \frac{C_k}{B \log_2(1 + \rho_{nk}^M)} \approx \frac{\ln(2)C_k}{B\rho_{nk}^M} \frac{1}{1 - \frac{\rho_{nk}^M}{2}} \approx \frac{C_k \ln(2)}{B} \left[\frac{1}{\rho_{nk}^M} + \frac{1}{2} \right]. \quad (3.42)$$

De plus, avec les hypothèses considérées ici, la puissance moyenne consommée par le réseau mobile peut s'écrire :

$$\mathbb{E}[P_{\text{net}}] = N_{BS}P_s + \left(P_0 - P_s + m_p P_{\text{max}}^{\text{DTx}}\right) \sum_{n=1}^{N_{BS}} \mathbb{E}[\Sigma_n]. \quad (3.43)$$

Dans cette équation $\Sigma_n = \sum_{k \in S_n} \mu_{nk}^{\min}$. De plus, en utilisant les développements limités, on peut réécrire Σ_n :

$$\Sigma_n = \frac{\gamma_n}{P_{\text{max}}^{\text{DTx}}} + \sum_{m \neq n} \Sigma_m \delta_{nm} + \epsilon_n, \quad \forall n \in \llbracket 1; N_{BS} \rrbracket. \quad (3.44)$$

Dans cette équation $\gamma_n = \frac{N \ln(2)}{B} \sum_{k \in S_n} \frac{C_k}{g_{nk}}$, $\delta_{nm} = \frac{\ln(2)}{B} \sum_{k \in S_n} \frac{C_k g_{mk}}{g_{nk}}$, et, $\epsilon_n = \frac{\ln(2)}{2B} \sum_{k \in S_n} C_k$. L'équation (3.44) nous donne un système de N_{BS} équations en Σ_n . La matrice de ce système s'écrit :

$$\Delta \Sigma = \frac{1}{P_{\text{max}}^{\text{DTx}}} \gamma + \epsilon. \quad (3.45)$$

Dans cette équation, $\Sigma = (\Sigma_1, \dots, \Sigma_{N_{BS}})^T$, $\gamma = (\gamma_1, \dots, \gamma_{N_{BS}})^T$, $\epsilon = (\epsilon_1, \dots, \epsilon_{N_{BS}})^T$ et, $\Delta_{nm} = -\delta_{nm}$ si $m \neq n$ et $\Delta_{nn} = 1$, $\forall m \in \llbracket 1; N_{BS} \rrbracket$. En résolvant ce système, on peut réécrire la puissance totale consommée par le réseau :

$$\mathbb{E}[P_{\text{net}}] = N_{BS}P_s + \left(P_0 - P_s + m_p P_{\text{max}}^{\text{DTx}}\right) \sum_{n=1}^{N_{BS}} \left(\frac{1}{P_{\text{max}}^{\text{DTx}}} \mathbb{E}[(\Delta^{-1} \gamma)_n] + \mathbb{E}[(\Delta^{-1} \epsilon)_n] \right). \quad (3.46)$$

L'équation (3.46) est convexe en $P_{\text{max}}^{\text{DTx}}$. On peut donc trouver la valeur optimale de $P_{\text{max}}^{\text{DTx}}$ en trouvant le point d'annulation de sa dérivée par rapport à $P_{\text{max}}^{\text{DTx}}$:

$$P_{\text{max}}^{\text{DTx}*} = \sqrt{\frac{(P_0 - P_s) \sum_{n=1}^{N_{BS}} \mathbb{E}[(\Delta^{-1} \gamma)_n]}{m_p \sum_{n=1}^{N_{BS}} \mathbb{E}[(\Delta^{-1} \epsilon)_n]}} \quad (3.47)$$

Dans cette équation, $P_{\text{max}}^{\text{DTx}*}$ désigne la valeur optimale de $P_{\text{max}}^{\text{DTx}}$. De plus, d'après l'équation (3.47), à faible RSIB, dans le réseau 1, la valeur optimale de $P_{\text{max}}^{\text{DTx}}$ est environ égale à 25 mW et dans le réseau 2, la valeur de $P_{\text{max}}^{\text{DTx}}$ optimale est égale à 11 mW.

Cette analyse à faible RSIB nous a donc permis de confirmer nos résultats de simulations qui montrent que notre solution apporte de meilleurs résultats lorsque $P_{\text{max}}^{\text{DTx}}$ est faible.

3.5 Conclusions

Dans ce chapitre, nous avons proposé un algorithme optimal pour la résolution du problème d'allocation de puissance en TDMA avec de la transmission discontinue. Ensuite, nous avons montré que cet algorithme peut être utilisé pour l'allocation de puissance dans des réseaux composés de plusieurs cellules. Dans ces réseaux, les performances de notre solution peuvent être améliorées en réduisant la puissance d'émission maximale de la station de base $P_{\max}^{\text{DTx}}$. Enfin, nous avons montré que, les performances de notre solution sont optimales lorsque $P_{\max}^{\text{DTx}}$ est faible, c'est-à-dire de l'ordre de la centaine de milliwatts. Pour de si faibles valeurs de $P_{\max}^{\text{DTx}}$, la quasi-totalité des utilisateurs doit être servie avec une puissance d'émission égale à $P_{\max}^{\text{DTx}}$ sans faire de contrôle de puissance.

Dans tout ce chapitre, nous avons supposé que le canal de tous les utilisateurs était plat. Lorsque ce n'est pas le cas, de l'OFDMA est généralement utilisé. Dans le prochain chapitre, nous nous intéressons au problème d'allocation de ressources et de puissance avec du Cell DTx pour une station de base OFDMA.

Chapitre 4

Allocation de ressources et de puissance avec de la transmission discontinue en OFDMA

Sommaire

4.1	Introduction	77
4.2	Dans des canaux plats	78
4.2.1	Formulation du problème	78
4.2.2	Réécriture du problème	79
4.3	Dans un canal sélectif en fréquence	81
4.3.1	Position du problème	81
4.3.2	Allocation de puissance	82
4.3.3	Allocation de ressources	87
4.3.4	Résultats numériques	91
4.4	Conclusions	93

4.1 Introduction

Dans ce chapitre, nous supposons la voie descendante d'une station de base OFDMA et nous proposons un algorithme pour la résolution du problème d'allocation de ressources et de puissance.

Dans un premier temps, nous supposons que tous les utilisateurs ont un canal plat. Comme expliqué Section 2.3.1, dans ce cas là, il suffit de résoudre le problème d'allocation de puissance. Nous verrons que ce problème peut être ramené à un problème très semblable au problème TDMA qui a été traité dans le chapitre précédent.

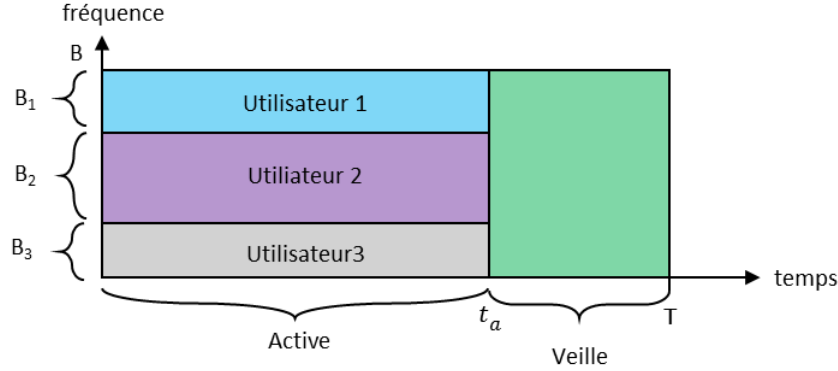


Figure 4.1 – Pendant chaque trame, la station de base est active pendant t_a et ensuite est mise en veille. Lorsque la station de base est active, de l'OFDMA est utilisé pour l'accès multiple.

Dans un second temps, nous verrons comment résoudre le problème d'allocation de ressources et de puissance lorsque le canal des utilisateurs est sélectif en fréquence. Dans cette résolution, nous verrons comment résoudre le problème d'allocation de puissance de façon optimale. Ensuite, nous regardons le problème d'allocation de ressources. Ces contributions ont fait l'objet d'une publication dans une conférence internationale [80].

4.2 Dans des canaux plats

4.2.1 Formulation du problème

Dans cette section nous supposons que le canal de tous les utilisateurs est plat et nous utilisons des notations analogues à celles du chapitre précédent. On va supposer que les trames sont découpées comme indiqué sur la Figure 4.1.

Durant chaque trame, la station de base est active pendant une portion de la trame de durée égale à $\mu_a = \frac{t_a}{T}$, ensuite elle est mise en veille. Lorsque la station de base est active, elle sert N_u utilisateurs en OFDMA. Elle alloue une portion de la bande B à chacun des utilisateurs. On note B_k la portion de bande utilisée pour l'utilisateur k et N_0 désigne la densité spectrale de puissance du bruit thermique (égale à la constante de Boltzmann multipliée par la température). Dans cette portion de bande, la station de base sert l'utilisateur k avec une puissance d'émission notée P_{TX}^k .

Pour minimiser l'énergie consommée par la station de base, on minimise la puissance d'émission moyenne consommée par celle-ci pendant une trame. On a le problème d'optimisation suivant :

$$\min_{P_{TX}^k, \mu_a} P_s (1 - \mu_a) + \mu_a \left(P_0 + m_p \sum_{k=1}^{N_u} P_{TX}^k \right), \quad (4.1a)$$

$$\text{S.à } C_k = B_k \mu_a \log_2 \left(1 + \frac{P_{TX}^k g_k}{N_0 B_k} \right), \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (4.1b)$$

$$0 \leq P_{TX}^k, \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (4.1c)$$

$$\sum_{k=1}^{N_u} P_{TX}^k \leq P_{\max}, \quad (4.1d)$$

$$0 \leq \mu_a \leq 1, \quad (4.1e)$$

$$B_k \geq 0, \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (4.1f)$$

$$\sum_{k=1}^{N_u} B_k = B. \quad (4.1g)$$

Dans le problème défini dans (4.1), la puissance moyenne consommée par la station de base pendant une trame est donnée par l'équation (4.1a). Avec l'équation (4.1b), on a la contrainte de débit de chaque utilisateur. Les quatre dernières équations donnent respectivement les contraintes sur la puissance d'émission, sur le temps de service qui doit être positif et plus court que la durée de la trame ainsi que sur la largeur de bande de chaque utilisateur qui doit être positive. Enfin, la somme des largeurs de bande de chaque utilisateur doit être égale à la largeur de bande totale.

4.2.2 Réécriture du problème

Formulé ainsi, le problème d'allocation de puissance OFDMA semble plus difficile à résoudre que son pendant TDMA (décrit par l'équation (2.15)). En effet, alors qu'en TDMA on ne s'intéresse qu'au temps de service et à la puissance allouée à chaque utilisateur, en OFDMA, on doit calculer la largeur de bande et la puissance d'émission de chaque utilisateur et le temps de service total μ_a . Afin de résoudre ce problème, on va le transformer pour qu'il s'apparente au problème TDMA. Pour cela, on pose :

$$\mu'_k = \frac{\mu_a B_k}{B} \quad (4.2)$$

$$P_{TX}^{k'} = \frac{P_{TX}^k B}{B_k} \quad (4.3)$$

L'introduction de ces variables nous permet d'écrire le problème sous la forme :

$$\min_{P_{TX}^{k'}, \mu'_k} P_s \left(1 - \sum_{k=1}^{N_u} \mu'_k \right) + \left(\sum_{k=1}^{N_u} \mu'_k \right) P_0 + m_p \sum_{k=1}^{N_u} \mu'_k P_{TX}^{k'}, \quad (4.4a)$$

$$\text{S.à } C_k = B \mu'_k \log_2 \left(1 + \frac{P_{TX}^{k'} g_k}{N_0 B} \right), \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (4.4b)$$

$$0 \leq P_{TX}^{k'}, \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (4.4c)$$

$$\sum_{k=1}^{N_u} \mu'_k P_{TX}^{k'} \leq \left(\sum_{k=1}^{N_u} \mu'_k \right) P_{\max}, \quad (4.4d)$$

$$\mu'_k \geq 0, \quad \forall k \in \llbracket 1; N_u \rrbracket, \quad (4.4e)$$

$$\sum_{k=1}^{N_u} \mu'_k \leq 1. \quad (4.4f)$$

Grâce à cette réécriture, on se retrouve avec un problème très semblable au problème TDMA introduit dans l'équation (2.15). La seule différence vient de la contrainte sur la puissance d'émission totale de l'équation (4.4d). En effet, alors qu'en TDMA on avait une contrainte sur la puissance maximale par utilisateur en OFDMA, on n'a plus qu'une seule contrainte qui concerne la somme des puissances d'émission.

La réécriture faite ici nous permet de proposer une méthode de résolution très proche de celle utilisée dans le Chapitre 3. Pour résoudre ce problème on va donc pouvoir :

- Utiliser la contrainte de capacité de l'équation (4.4b) pour réécrire les problèmes en fonction des μ'_k .
- Montrer que le problème est convexe et donc utiliser les conditions de KKT pour le résoudre.
- Montrer que lorsqu'aucune des contraintes n'est saturée, la solution est donnée par l'équation (3.7).
- Si la contrainte de l'équation (4.4d) ou celle de l'équation (4.4f) est saturée, on peut utiliser une méthode de résolution numérique pour trouver la valeur optimale des μ'_k .
- Enfin, on revient aux variables du problème initial par le biais des équations (4.2) et (4.3).

Cette résolution est très similaire à celle présentée dans le chapitre précédent, c'est pourquoi nous ne la décrivons pas davantage. Dans la suite de ce chapitre, nous nous concentrons sur le cas où au moins un des utilisateurs a un canal à évanouissements.

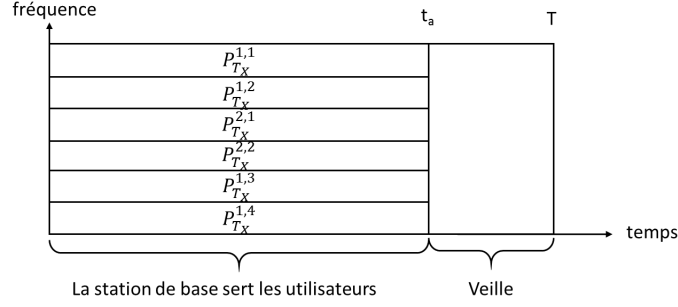


Figure 4.2 – La largeur de bande est divisée en N_c blocs de ressources. Dans cet exemple, on a 6 blocs de ressources qui sont répartis entre 2 utilisateurs.

4.3 Dans un canal sélectif en fréquence

4.3.1 Position du problème

On se place dans la situation où les utilisateurs ont des canaux sélectifs en fréquence. Dans ce cas là, la largeur de bande B est généralement divisée en N_c blocs de ressources de largeur fixe. La largeur de bande B_c de chacun de ces blocs de ressources est choisie de sorte qu'ils soient assez fins pour que l'on puisse considérer que le canal ne varie pas sur la largeur du bloc. De plus, $N = N_0 B_c$ désigne la puissance du bruit.

On suppose le modèle décrit par la Figure 4.2. On a N_c blocs de ressources pour servir N_u utilisateurs. On indexe les utilisateurs par la lettre k et les blocs de ressources par la lettre n . On note $g_{k,n}$ le coefficient du canal de l'utilisateur k dans le bloc de ressources n et $P_{T_X}^{k,n}$ la puissance d'émission utilisée pour cet utilisateur dans ce bloc de ressources. On rappelle que pour résoudre le problème d'allocation de ressources et de puissance avec du Cell DTx, la station de base va devoir :

- Résoudre le problème d'**allocation de ressources**, c'est-à-dire allouer à chaque utilisateur un jeu de blocs de ressources. Dans la suite de ce chapitre, on notera S_k le jeu de blocs de ressources dédié à l'utilisateur k et N_k son cardinal.
- Résoudre le problème d'**allocation de puissance** qui consiste à calculer la puissance d'émission de chaque utilisateur dans chaque bloc de ressources ainsi que le temps μ_a pendant lequel la station de base est active.

Dans la suite de ce chapitre, nous décrirons la résolution du problème dans l'ordre inverse. Nous commencerons par décrire la résolution du problème d'allocation de puissance. Pour cela, nous supposons qu'une allocation de ressources à été réalisée (quelle qu'elle soit) et nous verrons comment doit être allouée la puissance pour minimiser la consommation de la station de base. Ensuite, nous étudierons le problème d'allocation de ressources.

Bien que cet ordre de présentation soit inverse à celui dans lequel les deux allocations sont réalisées, il permet de mieux les décrire. En effet, la stratégie d'allocation de puissance ne dépend pas de l'allocation de ressources. Elle peut donc être décrite en premier. L'allocation de ressources, quand à elle, dépend de l'allocation de puissance. Cet ordre de présentation est souvent adopté dans la littérature [30, 34].

4.3.2 Allocation de puissance

On se concentre sur le problème d'allocation de puissance. On suppose que l'allocation de ressources a été réalisée et que l'on a choisi S_k pour chaque utilisateur. On cherche, maintenant, à répartir la puissance d'émission entre les utilisateurs afin de minimiser la puissance moyenne consommée par la station de base. Ce problème d'allocation de puissance s'écrit :

$$\min_{\mu_a, P_{T_X}} \quad P_m = (1 - \mu_a)P_s + \left(P_0 + m_p \sum_{k=1}^{N_u} \sum_{n=1}^{N_k} P_{T_X}^{k,n} \right) \mu_a \quad (4.5a)$$

$$\text{S.t.} \quad \mu_a \in [0; 1] \quad (4.5b)$$

$$C_k - B_c \mu_a \sum_{n=1}^{N_k} \log_2 \left(1 + \frac{g_{k,n} P_{T_X}^{k,n}}{N} \right) = 0, \quad \forall k \in \llbracket 1; N_u \rrbracket \quad (4.5c)$$

$$\sum_{k=1}^{N_u} \sum_{n=1}^{N_k} P_{T_X}^{k,n} \leq P_{max} \quad (4.5d)$$

$$P_{T_X}^{k,n} \geq 0, \quad \forall k, n \in \llbracket 1; N_u \rrbracket \times \llbracket 1; N_k \rrbracket \quad (4.5e)$$

Dans le problème décrit par le système d'équation (4.5), l'expression de la puissance moyenne consommée par la station de base pendant une trame est donnée par l'équation (4.5a). Les autres équations donnent respectivement les contraintes sur le temps de service (équation (4.5b)), les contraintes de capacité des utilisateurs (equation (4.5c)) ainsi que les contraintes sur les puissances d'émission.

On prouve dans l'Annexe B.1 que le problème formulé par l'équation (4.5) n'est pas convexe au sens donné dans [40].

Reformulation du problème

Pour pouvoir traiter le problème d'allocation de puissance, on doit le réécrire sous forme convexe. Pour cela, on introduit la variable $C_{k,n}$ qui est la capacité de l'utilisateur k dans le bloc de ressources n :

$$C_{k,n} = B_c \mu_a \log_2 \left(1 + \frac{g_{k,n} P_{T_X}^{k,n}}{N} \right). \quad (4.6)$$

L'introduction des $C_{k,n}$ permet de réécrire la fonction objectif :

$$(1 - \mu_a)P_s + \left(P_0 + m_p \sum_{k=1}^{N_u} \sum_{n=1}^{N_k} \frac{N}{g_{k,n}} \left[2^{\frac{C_{k,n}}{B_c \mu_a}} - 1 \right] \right) \mu_a. \quad (4.7)$$

On peut alors montrer que :

Proposition 4.1. *Si $C_{k,n} \geq 0$ pour tout k et tout n , et $\mu_a > 0$, la fonction P_m définie par l'équation (4.7) est convexe.*

Démonstration. La matrice Hessienne de la fonction $g_{k,n} : (\mu_a, C_{k,n}) \mapsto \frac{m_p N \mu_a}{g_{k,n}} \left(2^{\frac{C_{k,n}}{B_c \mu_a}} - 1 \right)$ est égale à :

$$H_{k,n} = \frac{m_p N \ln(2)^2}{g_{k,n} B_c^2 \mu_a} 2^{\frac{C_{k,n}}{B_c \mu_a}} \begin{pmatrix} \frac{C_{k,n}^2}{\mu_a^2} & -\frac{C_{k,n}}{\mu_a} \\ -\frac{C_{k,n}}{\mu_a} & 1 \end{pmatrix}. \quad (4.8)$$

Les valeurs propres de cette matrice Hessienne sont $\lambda_1 = 0$ et $\lambda_2 = \frac{m_p N \ln(2)^2}{g_{k,n} B_c^2 \mu_a} \left(1 + \frac{C_{k,n}^2}{\mu_a^2} \right) 2^{\frac{C_{k,n}}{B_c \mu_a}}$. Si $C_{k,n} \geq 0$ pour tout k et tout n , et $\mu_a > 0$, ces deux valeurs propres sont positives ou nulles. La fonction objectif définie par l'équation (4.7) est convexe. $\square$

L'introduction des $C_{k,n}$ permet aussi de réécrire la contrainte de l'équation (4.5d). En effet, cette contrainte peut être réécrite :

$$\mu_a \sum_{k=1}^{N_u} \sum_{n=1}^{N_k} \frac{N}{g_{k,n}} \left[2^{\frac{C_{k,n}}{B_c \mu_a}} - 1 \right] \leq \mu_a P_{max}. \quad (4.9)$$

Cette contrainte est aussi convexe. La preuve est similaire à celle faite ci-dessus. Finalement, on obtient la formulation suivante pour le problème d'optimisation :

$$\min_{\mu_a, C_{k,n}} (1 - \mu_a)P_s + \left(P_0 + m_p \sum_{k=1}^{N_u} \sum_{n=1}^{N_k} \frac{N}{g_{k,n}} \left[2^{\frac{C_{k,n}}{B_c \mu_a}} - 1 \right] \right) \mu_a \quad (4.10a)$$

$$\text{S. à } \mu_a \in [0; 1] \quad (4.10b)$$

$$C_k - \sum_{n=1}^{N_k} C_{k,n} = 0, \quad \forall k \quad (4.10c)$$

$$\mu_a \sum_{k=1}^{N_u} \sum_{n=1}^{N_k} \frac{N}{g_{k,n}} \left[2^{\frac{C_{k,n}}{B_c \mu_a}} - 1 \right] \leq \mu_a P_{max} \quad (4.10d)$$

$$C_{k,n} \geq 0, \quad \forall k, n \quad (4.10e)$$

Nous avons montré que la fonction objectif est convexe, tout comme les différentes contraintes. C'est suffisant pour montrer que le problème est convexe.

Résolution du problème d'allocation de puissance

Le problème étant convexe, on peut le résoudre en utilisant les conditions de Karush-Kuhn et Tucker. La méthode employée est donc similaire à celle utilisée dans le Chapitre 3. Dans cette section, nous énonçons seulement les principaux résultats. Le lecteur intéressé par les détails mathématiques peut consulter l'Annexe B.2.

Dans un premier temps, en appelant μ_{opt} la valeur optimale de μ_a et en dérivant le Lagrangien de notre problème, on obtient l'expression de la valeur optimale des $C_{k,n}$ comme la solution d'un algorithme de water-filling :

$$C_{k,n} = \left(\frac{C_k}{N'_k} + B_c \mu_{opt} \log_2 \left(\frac{g_{k,n}}{\left(\prod_{n=1}^{N'_k} g_{k,n} \right)^{\frac{1}{N'_k}}} \right) \right)^+ \quad (4.11)$$

Dans cette équation, $(.)^+ = \max(., 0)$ (cette notation a été introduite dans le Chapitre 2 lors de l'introduction de l'algorithme du water-filling). De plus, N'_k désigne le nombre de blocs de ressources alloués à l'utilisateur k dans lesquels la puissance d'émission est non-nulle. Une fois que cette expression est établie, on peut utiliser l'équation (4.6) pour déduire l'expression de la puissance d'émission optimale pour chaque utilisateur :

$$P_{TX}^{k,n} = N \left(\frac{2^{\frac{C_k}{N'_k B_c \mu_{opt}}}}{\left(\prod_{n=1}^{N'_k} g_{k,n} \right)^{\frac{1}{N'_k}}} - \frac{1}{g_{k,n}} \right)^+ \quad (4.12)$$

L'allocation de puissance optimale donnée par l'équation (4.12) est très proche de celle sans Cell DTx, qui a déjà été calculée dans la littérature [34] et que nous avons rappelée dans le Chapitre 2, équation (2.12). A la différence de l'allocation de puissance sans Cell DTx, avec l'allocation de puissance avec Cell DTx, la valeur de μ_{opt} apparait dans le calcul des puissances.

Cette valeur μ_{opt} doit être calculée de façon temporaire à chaque étape de l'algorithme de water-filling. on rappelle que, après chaque itération de l'algorithme du water-filling, on va successivement mettre à 0 la puissance d'émission dans les blocs de ressources dans lesquels elle est négative et recalculer les puissances d'émission. Ceci jusqu'à ce qu'elles soient toutes positives. Lorsque le Cell DTx est utilisé, avant de recalculer la puissance d'émission, on doit calculer la valeur de μ_{opt} dont va dépendre la valeur des $P_{TX}^{k,n}$.

Cette valeur de μ_{opt} va être calculée de différentes manières en fonction des contraintes qui sont saturées ou non. On a trois cas différents :

1. Soit $\mu_{\text{opt}} = 1$, c'est-à-dire que la contrainte de l'équation (4.10b) est saturée.
2. Soit la contrainte de l'équation (4.10d) est saturée. Dans ce cas là, on notera μ_{min} la valeur de μ_{opt} .
3. Soit aucune des deux contraintes énoncées ci-dessus n'est saturée. Dans ce cas là, on notera $\mu_{\text{opt}1}$ la valeur de μ_{opt} .

Les trois cas possibles sont présentés sur la Figure 4.3. Pour bien les comprendre, il faut comparer les puissances statique et dynamique de la station de base. Pour rappel, la puissance dynamique est proportionnelle à la puissance d'émission alors que la puissance statique n'en est pas directement dépendante. Pour minimiser la puissance statique, il faut que la station de base soit en veille le plus longtemps possible alors que pour minimiser la puissance dynamique, il faut réduire au maximum la puissance d'émission, c'est-à-dire ne pas passer en mode veille.

Si la puissance dynamique de la station de base est dominante devant la puissance statique, on gagne peu de puissance en passant en veille, dans ce cas là, la contrainte de l'équation (4.10b) est saturée (cas 1). Au contraire, si la puissance statique est dominante, on a un grand intérêt à passer en veille et peu d'intérêt à faire du contrôle de puissance. Dans ce cas là, la contrainte de l'équation (4.10d) est saturée (cas 2). Si les deux puissances sont du même ordre de grandeur, on va se retrouver dans le dernier cas, c'est-à-dire qu'aucune des deux contraintes n'est saturée.

Les observations faites ci-dessus nous montrent que l'on ne pourra pas avoir les deux contraintes saturées en même temps. On détaille maintenant le calcul de μ_{opt} dans les trois différents cas.

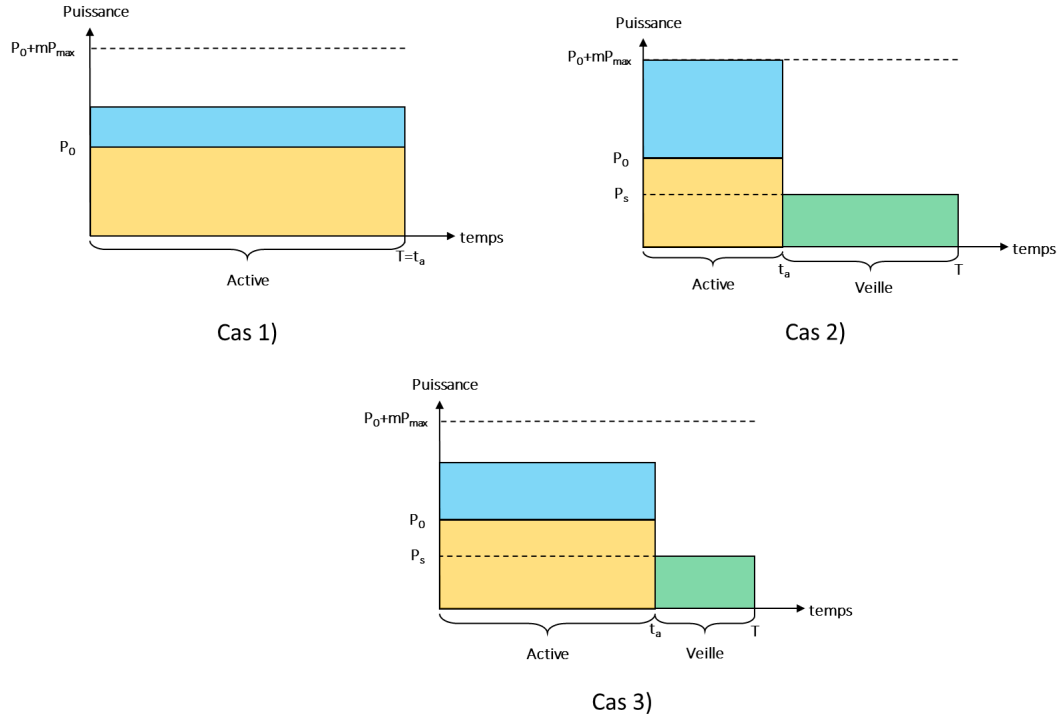


Figure 4.3 – Les trois situations possibles pour le calcul de μ_{opt} . Dans le premier cas, la contrainte de temps de l'équation (4.10b) est saturée. Dans le second, la contrainte de puissance de l'équation (4.10d) est saturée. Dans le dernier, aucune de ces contraintes n'est saturée.

Cas 1 : Si la contrainte sur le temps de service total est saturée, la station de base ne se met jamais en mode veille. On a alors $\mu_{\text{opt}} = 1$. La puissance d'émission de chaque utilisateur dans chaque bloc de ressources est donc égale à :

$$P_{T_X}^{k,n} = N \left(\frac{2^{\frac{C_k}{N'_k B_c}}}{\left(\prod_{n=1}^{N'_k} g_{k,n} \right)^{\frac{1}{N'_k}}} - \frac{1}{g_{k,n}} \right)^+ . \quad (4.13)$$

On retrouve alors l'allocation de puissance optimale sans mise en veille [34].

Cas 2 : Lorsque la contrainte sur la puissance maximale (équation (4.10d)) est saturée, toute la puissance d'émission est utilisée pour servir les utilisateurs. Par conséquent, le temps de service optimal est égal au temps de service minimum permettant de servir les utilisateurs. On note $\mu_{\min}$ ce temps de service qui peut être calculé en résolvant l'équation suivante grâce à un algorithme de recherche de racines à une seule variable :

$$N \sum_{k=1}^{N_u} \left(\frac{N'_k 2^{\frac{C_k}{N'_k B_c \mu_{\min}}}}{\left(\prod_{n=1}^{N'_k} g_{k,n} \right)^{\frac{1}{N'_k}}} - \sum_{n=1}^{N'_k} \frac{1}{g_{k,n}} \right) = P_{\max}. \quad (4.14)$$

Cas 3 : Dans ce cas là, aucune des contraintes mentionnées ci-dessus n'est saturée. On appelle μ_{opt1} la valeur optimale de μ_a qui est la solution de :

$$\sum_{k=1}^{N_u} \frac{2^{\frac{C_k}{N'_k B_c \mu_{\text{opt1}}}}}{\left(\prod_{n=1}^{N'_k} g_{k,n} \right)^{\frac{1}{N'_k}}} \left[\frac{C_k \ln(2)}{B_c \mu_{\text{opt1}}} - N'_k \right] - \frac{P_0 - P_s}{m_p N} + \sum_{k=1}^{N_u} \sum_{n=1}^{N'_k} \frac{1}{g_{k,n}} = 0. \quad (4.15)$$

La preuve est donnée en Annexe B.2.

On peut remarquer que, si la station de base ne sert qu'un seul utilisateur, le temps de service de l'utilisateur est égal à :

$$\mu_{\text{opt}} = \frac{C \ln(2)}{N'_c B_c} \frac{1}{\mathcal{W} \left(e^{-1} \left(\prod_{n=1}^{N'_c} g_n \right)^{\frac{1}{N'_c}} \left(\frac{P_0 - P_s}{m_p N'_c N} - \frac{1}{N'_c} \sum_{n=1}^{N'_c} \frac{1}{g_n} \right) \right) + 1}. \quad (4.16)$$

Dans cette équation N'_c désigne le nombre de blocs de ressources dans lesquels la puissance d'émission est non-nulle. On remarque que cette formule est très similaire à celle dérivée équation (3.7) dans le cas TDMA.

Finalement, l'Algorithme 4 détaille le calcul des puissances d'émission et du temps de service μ_{opt} .

Nous venons de résoudre le problème d'allocation de puissance avec du Cell DTx de manière optimale. Dans la suite nous nous intéressons au problème d'allocation de ressources.

4.3.3 Allocation de ressources

Nous avons vu dans la section précédente que la solution du problème d'allocation de puissance avec Cell DTx est similaire à celle sans Cell DTx. Ainsi, avant de proposer des solutions spécifiques pour l'allocation de ressources avec de la transmission discontinue, il peut être judicieux d'analyser les performances des solutions proposées pour l'allocation de ressources sans Cell DTx.

L'objectif des solutions proposées pour l'allocation de ressources en OFDMA sans Cell DTx est de s'approcher de l'allocation de ressources optimale sans Cell DTx. Ces solutions ne seront efficaces avec du Cell DTx que si l'allocation de ressources optimale sans Cell DTx

Algorithme 4 Calcul des puissances d'émission et du temps de service μ_{opt}

```

 $P_{T_X}^{k,n} = -1, \forall k, n$ 
Soit  $g^{\text{temp}}$  la liste des  $g_{k,n}$ .
while  $\min P_{T_X}^{k,n} < 0$  do
     $P_{T_X}^{k,n} = 0, \forall k, n$ 
    Calculer  $\mu_{\text{opt1}}$  avec l'équation (4.15) en considérant uniquement les  $g_{k,n}$  de  $g^{\text{temp}}$ .
    Calculer  $\mu_{\text{min}}$  avec l'équation (4.14) en considérant uniquement les  $g_{k,n}$  de  $g^{\text{temp}}$ .
    if  $\mu_{\text{opt1}} > 1$  then
         $\mu_{\text{opt}} = 1$ 
    else if  $\mu_{\text{opt1}} < \mu_{\text{min}}$  then
         $\mu_{\text{opt}} = \mu_{\text{min}}$ 
    else
         $\mu_{\text{opt}} = \mu_{\text{opt1}}$ 
    end if
    for  $n$  de 1 à  $N_c$  do
        if  $g_{k,n} \in g^{\text{temp}}$  then
            
$$P_{T_X}^{k,n} = N \left( \frac{2^{\frac{C_k}{N'_k B_c}}}{\left( \prod_{n=1}^{N'_k} g_{k,n} \right)^{\frac{1}{N'_k}}} - \frac{1}{g_{k,n}} \right).$$

            if  $P_{T_X}^{k,n} < 0$  then
                 $g^{\text{temp}} = g^{\text{temp}} \setminus g_{k,n}$ 
            end if
        end if
    end for
end while

```

est proche de l'allocation de ressources optimale avec du Cell DTx. On va donc comparer la puissance consommée par la station de base dans deux cas différents :

- Avec l'allocation de ressources et de puissance optimale avec Cell DTx (cette stratégie sera nommée (r1) dans la suite de cette section),
- Avec l'allocation de ressources optimale sans Cell DTx et l'allocation de puissance optimale avec Cell DTx (cette stratégie sera nommée (r2) dans la suite de cette section).

Si le nombre de blocs de ressources et d'utilisateurs ne sont pas trop importants, les allocations de ressources optimales peuvent être trouvées grâce à une recherche exhaustive. C'est-à-dire en essayant toutes les combinaisons de ressources possibles et en choisissant celle pour laquelle la consommation d'énergie est la plus faible.

Pour trouver l'allocation de ressources optimale avec Cell DTx (r1), on va réaliser l'allocation de puissance présentée dans la section précédente pour chaque combinaison de blocs de ressources et choisir celle pour laquelle la puissance consommée est la plus faible. Pour calculer la puissance consommée par la station de base avec l'allocation de ressources

optimale sans Cell DTx (r2), on utilise l'équation (4.13) pour réaliser l'allocation de puissance pour toutes les combinaisons possibles et on choisit celle qui consomme le moins de puissance. Ensuite, on utilise l'algorithme 4 pour réaliser l'allocation de puissance avec Cell DTx. Il est à noter que cette deuxième allocation de ressources est sous-optimale comparée à la première.

Si ces deux solutions apportent des performance similaires, on pourra utiliser les allocations de ressources proposées sans prendre en compte le Cell DTX. Au contraire, si la seconde solution a de biens moins bonnes performances, il faudra trouver des allocations de ressources spécifiques à l'utilisation du Cell DTx.

Pour comparer les deux allocations de ressources mentionnées ici, on suppose que deux utilisateurs sont servis par la station de base. Tout comme dans le standard LTE, on suppose que la bande est divisée en blocs de ressources de 180 kHz. La station de base sert les utilisateurs dans 6 ou 15 blocs de ressources (ces valeurs correspondent respectivement à la bande LTE de 1,4 MHz et à celle de 3 MHz). Le *fading* est généré en utilisant le modèle *ETU Typical Urban* [81].

Pour chaque utilisateur, on va fixer ses conditions de propagation moyennes. Pour cela, on va dénoter ρ_i^M le RSB moyen qu'aurait l'utilisateur i s'il était servi avec la moitié de la puissance d'émission sur la moitié de la bande. On va considérer deux cas différents :

- Dans le premier cas, les deux utilisateurs ont le même RSB égal à 15 dB.
- Dans le second, le premier utilisateur a un RSB égal à 5 dB et le second a un RSB égal à 25 dB.

On considèrera aussi deux types de stations de base différents, une femto-cellule et une macro-cellule. On a finalement 8 cas différents. Pour chacun d'eux, on fait varier la contrainte de capacité des utilisateurs et on regarde la perte liée à l'utilisation de l'allocation de ressources optimale sans Cell DTx (qui est sous-optimale quand on utilise du DTx). Pour évaluer cette perte, on la compare au gain apporté par le Cell DTX, c'est à dire que l'on calcule :

$$\Delta = \frac{P_{DTx}^{r2} - P_{DTx}^{r1}}{P_{noDTx} - P_{DTx}^{r1}}. \quad (4.17)$$

Dans cette équation P_{DTx}^{r2} est la puissance moyenne consommée par la station de base avec l'allocation de ressources r2 (optimale si on ne fait pas de Cell DTx), P_{DTx}^{r1} est la puissance moyenne consommée par la solution optimale avec Cell DTx (r1) et P_{noDTx} est la puissance moyenne consommée par la station de base avec l'allocation de ressources et de puissance sans Cell DTx. Les résultats obtenus sont détaillés dans le Tableau 4.1.

On peut voir dans le Tableau 4.1, que la perte liée à l'utilisation de l'allocation de ressources (r2) est très faible. Elle représente moins de 10 % du gain qu'apporte le Cell DTx. Afin de visualiser cette faible perte, on trace sur la Figure 4.4 la puissance consommée par la

Tableau 4.1 – Consommation supplémentaire liée à l'utilisation de l'allocation de ressources optimale sans Cell DTx.

type de station de base	ρ_1^M (dB)	ρ_2^M (dB)	N_c	Perte (%)	Scenario
Macro	15	15	6	1.1	#1
			15	1.7	#2
	5	25	6	10.3	#3
			15	2.5	#4
Femto	15	15	6	0.9	#5
			15	1.3	#6
	5	25	6	6.8	#7
			15	1.2	#8

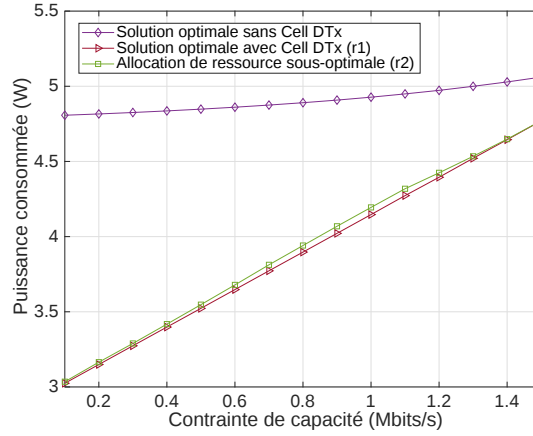


Figure 4.4 – Puissance moyenne consommée par la station de base avec l'allocation de puissance optimale appliquée avec soit l'allocation de ressources optimale (r1), soit avec l'allocation de ressources qui ne considère pas le Cell DTx (r2).

station de base avec la stratégie (r1) et avec la stratégie (r2) dans le scénario #7 (un des pires cas).

On voit bien sur la Figure 4.4 que l'utilisation de r1 et r2 donne quasiment les mêmes résultats. Par conséquent, l'allocation de ressources optimale sans Cell DTx et l'allocation de ressources optimale avec Cell DTx ont quasiment les mêmes performances. Par conséquent, on peut utiliser les solutions proposées pour l'allocation de ressources sans prendre en compte le Cell DTx lorsque du DTx est utilisé. Par exemple, on peut utiliser l'une des solutions proposées dans [30, 33]. Sans restriction, dans la section suivante, nous utiliserons les algorithmes *Bandwidth Assignment Based on SNR* (BABS) et *Amplitude Creeving Greedy* (ACG) pour l'allocation de ressources qui ont été proposés par Kivanc et al. en 2003 [39]. Cette solution a l'avantage d'avoir une faible complexité.

4.3.4 Résultats numériques

Suite aux conclusions de la section précédente, nous proposons d'utiliser une solution qui utilise l'algorithme BABS+ACG pour l'allocation de ressources avec l'allocation de puissance présentée par l'Algorithme 4.

L'algorithme BABS permet de calculer le nombre N_k de blocs de ressources alloués à chaque utilisateur simplement en calculant le nombre de blocs de ressources qu'il faudrait pour servir l'utilisateur si son coefficient d'atténuation du canal était identique dans tous les canaux et égal au coefficient de canal moyen de l'utilisateur. Ensuite, l'algorithme ACG est un algorithme glouton qui alloue à chaque utilisateur N_k canaux. Pour cela, il parcourt les canaux et les alloue successivement à l'utilisateur qui a le meilleur coefficient d'atténuation dans ce canal. Ceci jusqu'à ce que le nombre de canaux alloués à chaque utilisateur soit égal à N_k .

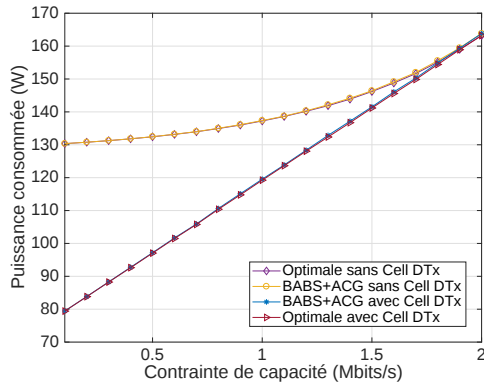
On commence par évaluer les performances de la solution proposée pour l'allocation de ressources dans des cas simples. Pour cela, on considère les scénarios #1, #3, #5, #7 décrits dans le Tableau 4.1. Pour chacun de ces scénarios, on compare la consommation de la station de base avec l'allocation de ressources optimale et avec l'algorithme BABS+ACG.

On peut voir sur la Figure 4.5 que l'on a trois cas différents :

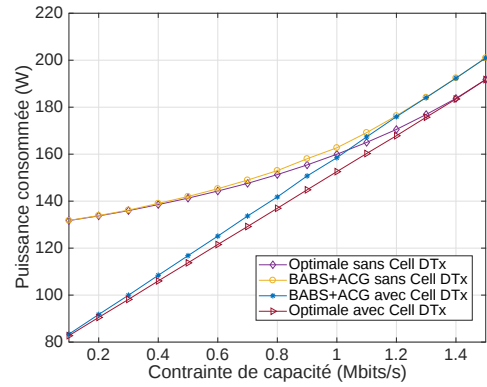
1. On peut voir sur la Figure 4.5a et sur la Figure 4.5c que lorsque les deux utilisateurs ont le même RSB moyen, l'algorithme BABS+ACG donne des résultats optimaux.
2. On voit sur la Figure 4.5b que, pour une macro-cellule, l'erreur liée à l'utilisation de l'algorithme BABS+ACG est identique avec ou sans Cell DTx. Par conséquent, on pourrait améliorer les performances de l'allocation de ressources en utilisant une allocation de ressources plus performante mais plus complexe, qui a déjà été proposée dans la littérature sans prendre en compte le Cell DTx. Par exemple, celle proposée dans [30].
3. Finalement, on peut voir sur la Figure 4.5d, que, pour une femto-cellule, lorsque les utilisateurs ont un RSB très différent, la perte causée par l'allocation de ressources est plus importante si on applique le Cell DTx.

Les deux derniers cas s'expliquent en regardant la puissance statique moyenne et la puissance dynamique moyenne des stations de base. En effet, l'allocation de ressources utilisée a été conçue pour minimiser la puissance d'émission de la station de base. Elle a donc pour but de minimiser la puissance dynamique.

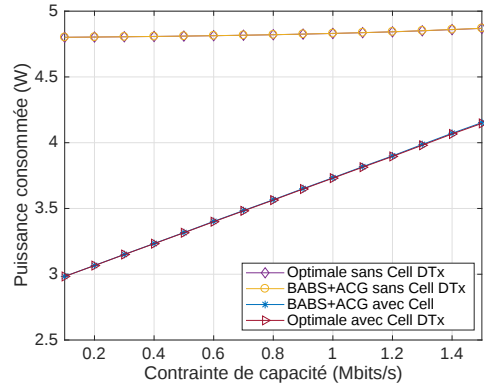
Les macro-cellules ont une puissance dynamique élevée. C'est pourquoi, pour ces stations de base, l'utilisation d'une allocation de ressources qui minimise la puissance dynamique a de bonnes performances. Au contraire, la puissance dynamique d'une femto-cellule est faible. Par conséquent, l'allocation de ressources proposée aura des performances un peu moins bonnes.



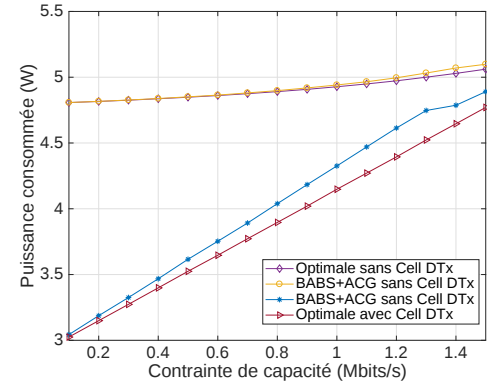
(a) Macro-cellule, $\rho_1^M = \rho_2^M = 15$ dB



(b) Macro-cellule, $\rho_1^M = 5$ dB, $\rho_{max}^2 = 25$ dB



(c) Femto-cellule, $\rho_1^M = \rho_2^M = 15$ dB



(d) Femto-cellule, $\rho_1^M = 5$ dB, $\rho_2^M = 25$ dB

Figure 4.5 – Comparaison de la solution proposée (utilisation de l'algorithme BABS+ACG pour l'allocation de ressources) avec une solution optimale (recherche exhaustive).

Avec l'allocation de ressources utilisée, la puissance consommée par la station de base est, au pire, 5% supérieure à celle consommée avec une solution optimale (qui demande une recherche exhaustive).

Nous évaluons maintenant les performances de l'allocation de puissance optimale dans un scénario réaliste. On suppose une macro-cellule qui sert des utilisateurs répartis suivant un processus de Poisson. On a en moyenne 10 utilisateurs dans la zone de couverture et la distance entre les utilisateurs et la station de base varie entre 300 et 1500 m. On suppose que le facteur de bruit du récepteur est égal à 2 dB et le gain d'antenne est de 10 dBi. La bande de 20 MHz est centrée autour de 2 GHz et est divisée en 100 canaux. Les *path-losses* sont calculés avec le modèle de Winner II *suburban macrocell* et le *fading* est généré via le modèle ETU *Extended Pedestrian A*. On suppose que tous les utilisateurs ont la même contrainte de capacité.

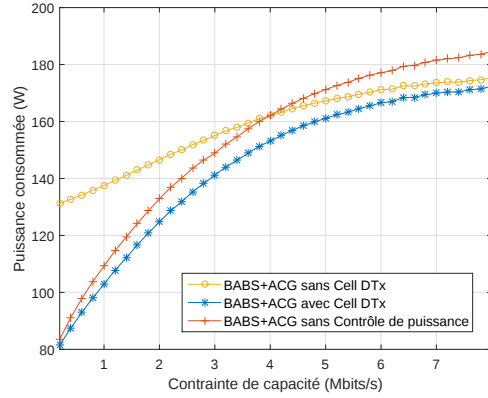


Figure 4.6 – Puissance moyenne consommée par la station de base en utilisant l’algorithme BABS+ACG pour l’allocation de ressources et avec différentes allocations de puissance.

De manière analogue à ce qui a été fait dans le Chapitre 3, on compare trois stratégies d’allocation de puissance différentes :

- L’allocation de puissance optimale proposée.
- Une stratégie sans contrôle de puissance. C’est-à-dire, pour laquelle la somme des puissances d’émission est toujours égale à $P_{\max}$. Avec cette stratégie, le temps de service (pendant lequel la station de base est active) est calculé avec l’équation (4.14).
- L’allocation de puissance optimale sans Cell DTx.

Pour chacune des stratégies comparées, on réalise d’abord l’allocation de ressources avec l’algorithme BABS+ACG. Ensuite, on réalise l’allocation de puissance désirée. Sur la Figure 4.6, on montre l’évolution de la puissance consommée par la station de base en fonction de la contrainte de capacité des utilisateurs.

On peut voir sur cette figure que, comme attendu, l’allocation de puissance optimale a les meilleures performances. Lorsque les utilisateurs ont des contraintes de capacité faible, cette solution a un fort gain par rapport à la solution sans Cell DTx. Notre solution permet de réduire la puissance consommée de 38 % par rapport à la solution sans Cell DTx. Si on compare maintenant la solution optimale avec la solution sans contrôle de puissance, on voit que l’écart entre les deux courbes est maximum lorsque la contrainte de capacité est élevée. Le gain apporté par l’allocation de puissance optimale représente jusqu’à 8 % de la puissance consommée par la station de base.

4.4 Conclusions

Dans ce chapitre nous nous sommes intéressés au problème d’allocation de ressources et de puissance en OFDMA. Nous avons d’abord supposé que le canal de tous les utilisateurs

était plat. Nous avons montré que, dans ce cas là, ce problème peut-être reformulé pour donner un problème similaire à celui traité en TDMA dans le Chapitre 3.

Ensuite, nous nous sommes intéressés au cas où le canal des utilisateurs n'est pas plat. Nous avons proposé un algorithme pour réaliser l'allocation de puissance optimale. Nous avons ensuite vu que les solutions proposées pour l'allocation de ressources sans Cell DTx peuvent être utilisées avec du Cell DTx.

Nous venons de voir, dans les Chapitres 2, 3 et 4 que l'utilisation du Cell DTx avec une allocation de ressources et de puissance efficace permet de réduire l'énergie consommée par les réseaux mobiles. Cependant, le potentiel des réseaux de communication dans la réduction de l'empreinte écologique de l'activité humaine ne s'arrête pas à la réduction de la consommation des réseaux mobiles. En effet, les réseaux de communication, et en particulier les réseaux d'objets connectés, peuvent aider à mieux gérer des systèmes distribués complexes tels que le réseau électrique intelligent. Ceci permet, par exemple, d'intégrer plus facilement les énergies renouvelables dans la production d'électricité et dans le réseau cellulaire. Ceci permet, donc, de mieux gérer l'énergie du réseau mobile.

Dans la Partie 2, nous étudions les réseaux de communication pour le réseau électrique intelligent.

Deuxième partie

Amélioration des communications du réseau électrique intelligent

Chapitre 5

Communications dans le réseau électrique intelligent

Sommaire

5.1	Introduction	97
5.2	Le réseau électrique intelligent	98
5.3	Les communications du réseau électrique intelligent	99
5.4	Rappels sur l'architecture HDCRAM	100
5.5	HDCRAM pour la gestion de la production et de la consommation	102
5.6	Description HDCRAM des mécanismes de gestion du réseau	105
5.6.1	HDCRAM pour la gestion du réseau de transport	106
5.6.2	HDCRAM pour la gestion du réseau de distribution	107
5.7	Conclusion	108

5.1 Introduction

L'intégration d'énergies renouvelables dans la production d'électricité permet de réduire les émissions de gaz carbonique. Cependant, certaines de ces sources ont une production d'électricité intermittente (c'est le cas du solaire et de l'éolien). C'est pourquoi leur insertion dans le réseau électrique réduit sa stabilité.

Pour palier ce problème, il faut mettre en place des mécanismes de gestion du réseau électrique afin de le transformer en réseau électrique intelligent. La mise en place de ces mécanismes se fait en répartissant des capteurs et actionneurs tout au long du réseau électrique. Ceux-ci étant distribués dans l'espace, leur mise en place nécessite l'utilisation de réseaux de communication pour transmettre les relevés et les informations de contrôle.

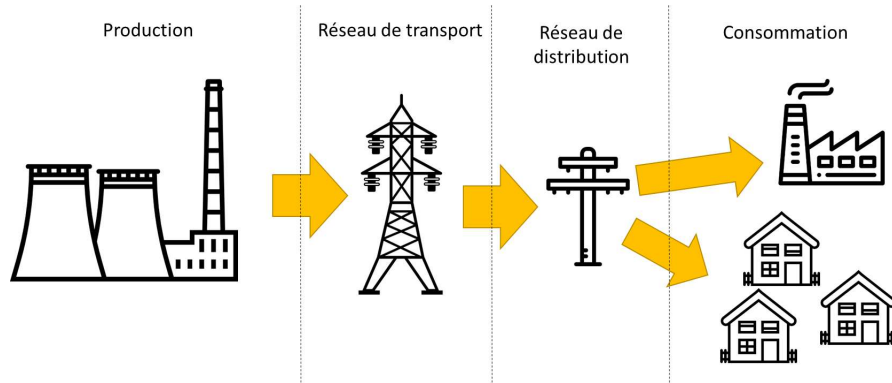


Figure 5.1 – Le réseau électrique sous sa forme historique. On a un découpage en quatre parties : la production, le réseau de transport, le réseau de distribution et la consommation.

L'objectif de ce chapitre est de dresser un état de l'art des réseaux de communication utilisés et envisagés pour le réseau électrique intelligent. Pour cela, nous présentons d'abord les différents réseaux de communication qui sont nécessaires au bon fonctionnement du réseau électrique. Ensuite, nous utilisons l'architecture HDCRAM qui a été introduite dans le Chapitre 1 pour décrire plus précisément les mécanismes de gestion du réseau électrique intelligent. L'état de l'art présenté dans cette section a fait l'objet d'une publication dans une conférence internationale [21].

5.2 Le réseau électrique intelligent

Historiquement, le réseau électrique est divisé en quatre parties distinctes :

- La partie production qui comprend l'ensemble des centrales électriques, qu'elles soient hydrauliques, nucléaires, à charbon ou autres. Traditionnellement, la majorité de l'électricité est produite par des centrales qui ont une forte capacité de production. Ces centrales sont reliées au réseau de transport.
- Le réseau de transport est un réseau haute tension qui fait le lien entre les centrales électriques et le réseau de distribution (basse tension).
- Le réseau de distribution est un réseau basse tension qui a pour but d'amener l'électricité jusqu'aux consommateurs finaux.
- Enfin, les consommateurs reliés au réseau de distribution consomment l'électricité.

Le réseau électrique traditionnel est schématisé sur la Figure 5.1.

Dans un réseau électrique, la puissance produite doit toujours être égale à la puissance consommée. Si ce n'est pas le cas, la fréquence du courant électrique dans le réseau (qui a une valeur nominale de 50 Hz) va être modifiée. Cette variation peut endommager le réseau et causer des pannes. C'est pourquoi, il est nécessaire de maintenir sans-cesse l'équilibre

entre production et consommation pour éviter que le réseau électrique sorte de son régime de fonctionnement et devienne instable.

Aujourd'hui, chaque consommateur peut devenir un producteur d'électricité. Par exemple, chacun peut mettre des panneaux solaires sur son toit. Par conséquent, une part de l'électricité produite est maintenant directement injectée dans le réseau de distribution. Cette production étant intermittente, il est nécessaire d'améliorer la gestion du réseau électrique, et en particulier celle du réseau de distribution, pour éviter que celui-ci ne devienne instable. Ceci se fait en utilisant des technologies de l'information et de la communication. En effet, le réseau électrique étant un système distribué à grande échelle, les réseaux de communication sont un élément clé dans la gestion du réseau électrique [82]. Plus globalement, en plus de permettre une bonne intégration des énergies renouvelables, l'utilisation de ces technologies optimise le fonctionnement du réseau et améliore sa fiabilité et sa sécurité [20].

Dans la section suivante, nous présentons les différents réseaux de communication qui vont permettre de mieux gérer les quatre parties du réseau électrique.

5.3 Les communications du réseau électrique intelligent

Pour l'analyse des réseaux de communication on s'appuie principalement sur les documents produits par le *Smart Grid Coordination Group* qui est un comité de normalisation qui réunit le Comité Européen de Normalisation (CEN), le Comité Européen de Normalisation en Électronique et Électrotechnique (CENELEC) et l'European Telecommunications Standards Institute (ETSI). Dans [83, 84], ce consortium décrit les réseaux de communication qui sont utiles à la gestion du réseau électrique. Dans le Tableau 5.1, nous avons classé ces réseaux en fonction de deux critères :

- Leur étendue : un réseau peut-être à faible étendue (à l'échelle d'une maison ou d'un bâtiment), d'étendue moyenne (à l'échelle d'un quartier ou d'un campus) ou de grande échelle (à l'échelle d'une ville, d'un pays, voire à l'échelle internationale).
- La partie du réseau électrique pour laquelle ils sont utilisés : un réseau de communication peut être utilisé pour gérer la production, la consommation, ou alors pour la gestion du réseau (surveillance des lignes et autres éléments). Cet aspect dimensionne les contraintes du système, c'est-à-dire le débit, la latence, ou le taux de perte de paquets [85].

Dans le Tableau 5.1, nous avons aussi listé les standards de communication sans fil envisagés pour chaque réseau de communication. On remarque dans ce tableau que les mêmes standards peuvent être utilisés pour la gestion de la production et de la consommation. En effet, les mécanismes de gestion de la production et de la consommation sont similaires et ont les mêmes contraintes de communication (en particulier pour la latence). C'est pourquoi, il a

Tableau 5.1 – Réseaux de communication utilisés pour la gestion du réseau électrique intelligent [83, 84]

Partie gérée	Etendue	Réseaux	Technologies et standards envisagés
Consommation	Faible portée	<i>Subscriber access network, Home automation</i>	Zigbee, Bluetooth, Wifi
	Portée moyenne	<i>Neighborhood access network, Backbone network</i>	Zigbee, Wifi, WirelessHART
	Longue portée	<i>AMI backhaul</i>	GSM, GPRS, UMTS, LTE, WiMAX, LoRaWAN, NB-IoT, Sigfox
Production	Faible portée	<i>Industrial fieldbus</i>	Zigbee, Bluetooth, Wifi
	Longue portée	<i>Balancing network</i>	GSM, GPRS, UMTS, LTE, WiMAX, LoRaWAN, NB-IoT, Sigfox
Réseau	Faible portée	<i>Intra-substation, Low-end intra-substation</i>	La plupart des protocoles envisagés pour ces communications sont filaires. Comme par exemple le standard IEC 61850
	Longue portée	<i>Inter-substation</i>	LTE, LTE-M, NB-IoT, 5G

été proposé dans le document du *US Department of Energy* [85] de présenter la gestion de la production et de la consommation comme un seul mécanisme. C'est ce que nous ferons dans la suite de ce chapitre.

Afin de mieux comprendre le rôle de chacun de ces réseaux, dans la suite, nous utilisons l'architecture HDCRAM qui a été présentée dans le Chapitre 1. Nous commencerons par rappeler le principe de cette architecture. Ensuite, nous l'appliquerons pour la gestion du réseau électrique intelligent. Ceci nous permettra de mieux comprendre les mécanismes de gestion et d'expliquer plus précisément le rôle et les contraintes de chacun des réseaux.

5.4 Rappels sur l'architecture HDCRAM

L'architecture HDCRAM est schématisée sur la Figure 5.2 [9]. Cette architecture définit un ensemble de règles qui servent à organiser la prise de décision dans un système intelligent. Grâce à HDCRAM, on va donc pouvoir définir à la fois le rôle et la position, dans le système global, des différents éléments qui composent la prise de décision dans un réseau

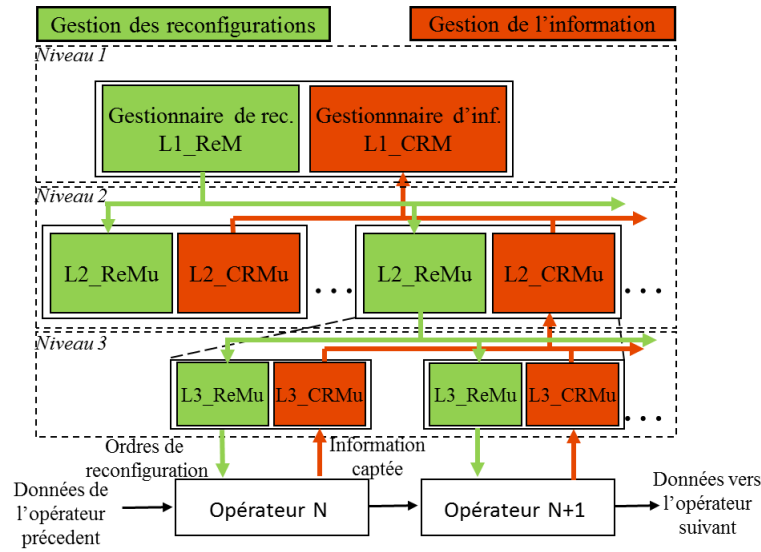


Figure 5.2 – L'architecture HDCRAM est une architecture de gestion hiérarchique des systèmes distribués intelligents.

électrique intelligent et définir les différents réseaux de communication qui existent entre ces différents éléments. De plus, en utilisant HDCRAM, on peut présenter de manière homogène les différents mécanismes qui gèrent le réseau électrique. Pour rappel, cette architecture est composée de deux parties distinctes : la première pour la gestion de l'information, la seconde pour la transmission des ordres de reconfiguration. Dans HDCRAM la prise de décision peut se faire à trois niveaux. On peut donc avoir des décisions locales, qui ne concernent qu'un seul élément du système géré, des décisions qui concernent plusieurs éléments et des décisions globales qui concernent l'ensemble du système. A chaque niveau, la prise de décision se fait dans des unités qui se divisent en deux parties : la première pour le traitement de l'information (CRMu) et la seconde pour la transmission des ordres (ReMu).

Par ailleurs, l'architecture HDCRAM est indépendante du mécanisme de gestion utilisé. On peut donc actualiser un algorithme de gestion sans modifier l'organisation hiérarchique de la prise de décision.

Nous allons utiliser cette architecture à trois niveaux pour décrire les mécanismes de gestion du réseau électrique. Des architectures de gestion à trois niveaux ont déjà été proposées dans la littérature pour la gestion hiérarchique de certaines parties du réseau électrique. Par exemple, pour la gestion de la consommation [86, 87], pour la gestion d'un réseau électrique de petite taille (*microgrid*) [88] ou encore pour la gestion de la consommation d'un bâtiment intelligent [89]. Dans tous ces articles, les auteurs proposent des architectures spécifiques pour un mécanisme de gestion. L'architecture HDCRAM a l'avantage de pouvoir être utilisée pour décrire de la même manière les différents mécanismes de gestion du réseau électrique.

Il aurait aussi été possible de présenter les différents mécanismes de gestion du réseau électrique et les réseaux de communication associés en utilisant la *Smart Grid Architecture Model* (SGAM) qui est l'architecture proposée par le *Smart Grid Coordination Group* [83, 84]. Cependant, cette architecture a été conçue pour décrire l'interopérabilité entre les différents éléments qui composent le réseau électrique, alors qu'avec l'architecture HDCRAM, on décrit les mécanismes de gestion du réseau.

5.5 HDCRAM pour la gestion de la production et de la consommation

Tout d'abord, pour la gestion de la consommation et de la production, (cinq premières lignes du tableau 5.1), on a deux mécanismes qui sont l'*Advanced Metering Infrastructure* et la réponse à la demande. Avec l'AMI, l'opérateur du réseau électrique relève à intervalles réguliers la production et la consommation des différents utilisateurs [90]. Ensuite, en fonction de ce relevé, il peut envoyer des consignes aux utilisateurs afin que ceux-ci adaptent leur comportement. Par exemple, l'opérateur du réseau peut calculer un nouveau prix pour l'électricité. Celui-ci est retransmis aux usagers qui le considèrent dans leur façon de produire et consommer de l'électricité. C'est ce que l'on appelle la réponse à la demande [85]. Les deux mécanismes décrits dans ce paragraphe sont des composants essentiels du réseau électrique intelligent.

Les deux mécanismes de gestion de la production et de la consommation cités ci-dessus peuvent être représentés à deux échelles géographiques différentes : soit à l'échelle locale, soit à l'échelle globale.

La représentation HDCRAM de la gestion de la production et de la consommation à l'échelle locale est visible sur la Figure 5.3. A l'échelle locale, ou d'un quartier, on a différents producteurs ou consommateurs (entreprises, bâtiments administratifs, champs de panneaux solaires ou maisons). Chacun d'entre eux gère sa propre électricité. Par exemple, à l'intérieur de la maison, cela concerne les appareils électriques, et en particulier ceux dont la consommation peut être ajustée (climatisation), ceux dont le démarrage peut être différé (machine à laver), les sources d'énergie (panneaux solaires) et les batteries (dans un véhicule électrique). Chacun de ces appareils est donc géré par une unité de gestion de niveau 3. Le gestionnaire de la maison, qui peut être un compteur intelligent, est au niveau 2 de l'architecture. Il a deux rôles. Le premier est qu'il doit gérer le fonctionnement des différents appareils du bâtiment. Le second est qu'il doit échanger des informations avec le reste du réseau électrique. En particulier, il peut recevoir des consignes, ou encore, le prix de l'électricité et transmettre la consommation électrique du bâtiment ou d'autres informations utiles pour la gestion. Cet échange d'information se fait avec le gestionnaire de niveau 1 qui se trouve généralement dans

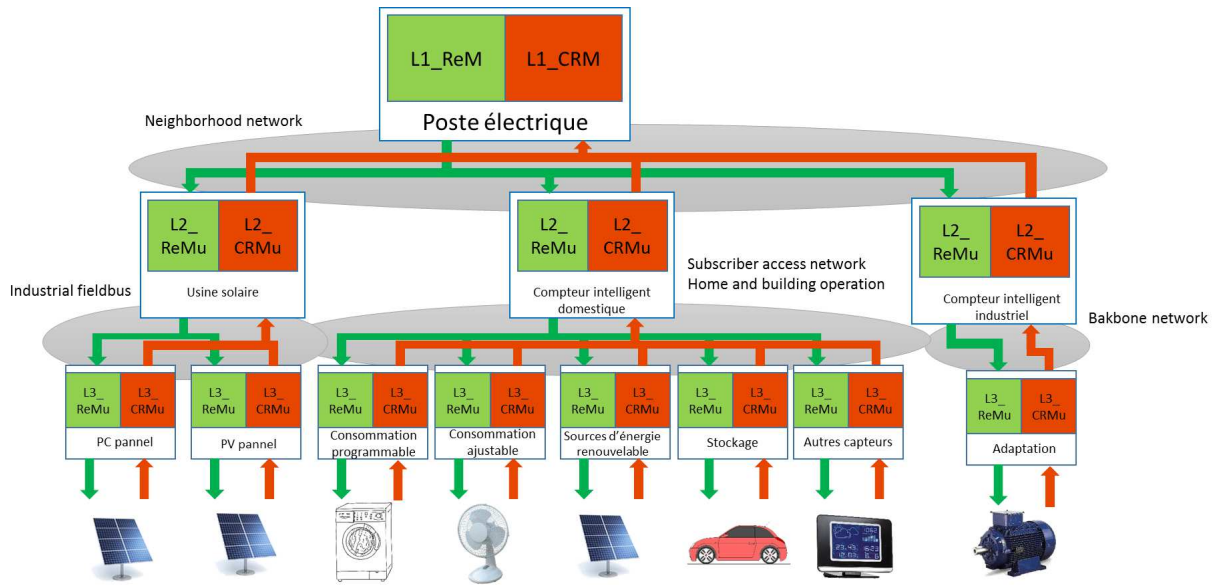


Figure 5.3 – Représentation HDCRAM pour la gestion de la production et de la consommation à l'échelle locale.

le poste électrique. Par exemple, dans l'AMI, on a un agrégateur au sein du poste électrique qui récupère les données envoyées par les compteurs intelligents, qui les agrège et les transmet au centre de contrôle.

La Figure 5.4 montre la représentation HDCRAM du mécanisme de gestion de la consommation et de la production d'électricité à l'échelle globale. C'est la vue qu'ont les opérateurs du réseau électrique. Au niveau 3 on a les postes de gestions qui sont reliés aux gestionnaires régionaux (niveau 2) qui sont eux même reliés au gestionnaire global de la production et de la consommation. La consommation d'électricité mesurée est remontée du gestionnaire de niveau 3 vers celui de niveau 1 alors que les ordres (le prix) va être transmis depuis le niveau 1 vers l'unité de gestion de niveau 3.

Les communications de l'AMI et de la réponse à la demande ne sont pas très contraintes. D'après le *US Department of Energy* [85], la latence des communications de ces deux mécanismes peut-être comprise entre quelques secondes (AMI) et aller jusqu'à plusieurs minutes (réponse à la demande). Pour ces deux mécanismes, on a plusieurs réseaux de communication qui peuvent être classés en fonction de leur étendue. Pour chaque réseau, on peut soit utiliser des communications filaires, soit utiliser des Courants Porteurs en Ligne (CPL), soit utiliser des technologies sans fil. Dans la suite, pour des raisons de clarté, nous ne présentons que les technologies sans fil qui peuvent être utilisées dans chaque réseau.

Tout d'abord, les réseaux de faible portée tel que le *Subscriber access network*, l'*Home automation* et l'*Industrial fieldbus*, sont utilisés entre les unités de niveau 3 et celles de niveau 2. Ces communications sont internes aux bâtiments (maisons, immeubles, usines) et

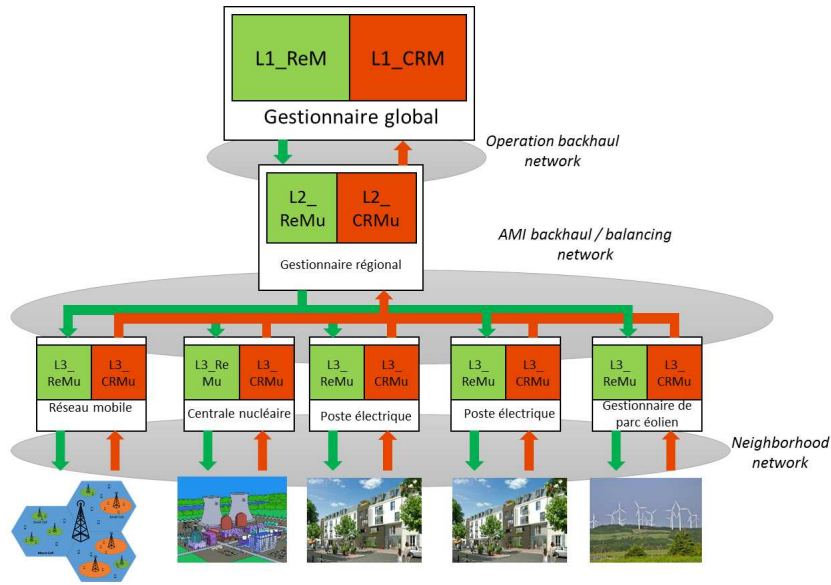


Figure 5.4 – Représentation HDCRAM pour la gestion de la production et de la consommation à l'échelle globale.

permettent de relier les différents appareils au gestionnaire de l'énergie de la maison, qui peut être un compteur intelligent. Pour ces réseaux, on va pouvoir utiliser des technologies déjà existantes dans la quasi-totalité des foyers (Wifi), des standards développés pour l'Internet des Objets (IdO) à faible portée tels que le Zigbee ou le Bluetooth [91].

On a ensuite les réseaux de moyenne portée (*Neighborhood access network*, *Backbone network*) qui permettent d'échanger des informations à l'échelle d'un quartier, d'un campus ou d'un ensemble de bâtiments. Ces communications peuvent être utilisés entre les unités de niveau 3 et celles de niveau 2 ou entre les unités de niveau 2 et celle de niveau 1 de l'architecture HDCRAM locale. Pour ces communications, on peut utiliser des réseaux Wifi ou Zigbee qui ont une portée moyenne. En plus de ces deux technologies génériques (c'est à dire qui peuvent être utilisées pour différents usages), on peut aussi utiliser des standards spécifiques tels que le WirelessHART qui est spécifiquement conçu pour les communications entre compteurs intelligents et agrégateurs de données [92].

Enfin, on a les réseaux longue portée tels que l'*AMI backhaul* qui est utilisé pour transmettre la consommation électrique depuis les agrégateurs de données vers le centre de contrôle du réseau électrique. Pour ces communications longue distance on a plusieurs possibilités. La première est d'utiliser les réseaux mobiles déjà existants. Par exemple, le réseau *Global System for Mobile communications* (GSM) ou son extension le *General Packet Radio Service* (GPRS) qui sont utilisés pour la 2G, sont aujourd'hui sous-utilisés par les utilisateurs mobiles. C'est pourquoi le standard *Extended Coverage-GSM* (EC-GSM) adapte ces technologies pour

les objets connectés [93]. Les auteurs de [94] proposent de faire évoluer les standards GSM et GPRS spécifiquement pour le réseau électrique intelligent.

Pour les communications entre agrégateurs et les centres de contrôle régionaux, on peut aussi utiliser des standards *Low Power Wide Area Networks* (LPWAN). Ces standards ont récemment été conçus pour les réseaux IdO de longue portée [6]. Les standards LPWAN ont plusieurs particularités. La première c'est qu'ils utilisent des protocoles radio très simples souvent basés sur un schéma d'accès aux fréquences radio de type ALOHA [95]. Ces réseaux ont aussi une longue portée (de l'ordre de la dizaine de kilomètres) et une très faible efficacité spectrale. Les solutions les plus connues sont la technologie *Ultra Narrow Band* développée par Sigfox [96] et la technologie LoRaWAN [97] développée par la LoRa Alliance. On peut aussi citer la technologie Weightless [98] développée pour les communications dans les bandes TV laissées libres et le standard Ingenu [6].

Les communications de chaque agrégateur de l'*AMI backhaul* sont des communications à faible débit et peu contraintes en latence. Elles peuvent donc être faites par le biais de réseaux LPWAN [99]. C'est pourquoi ces réseaux seront au centre des études menées dans le Chapitre 6.

Pour donner un exemple concret de standards utilisés pour la gestion de la consommation, en France, Enedis, l'opérateur du réseau de distribution historique, est en train de déployer 35 millions de compteurs intelligents (Linky). Ces compteurs transmettent les données collectées à un des 700000 agrégateurs de données situés dans les postes électriques en utilisant des CPL. Ces derniers utilisent alors le réseau GPRS pour transmettre l'information au centre de contrôle [100].

5.6 Description HDCRAM des mécanismes de gestion du réseau

Les réseaux de communications présentés dans la section précédente seront le sujet du Chapitre 6. Ce ne sont, cependant, pas les seuls réseaux de communications nécessaires pour le fonctionnement d'un réseau électrique intelligent. Certains réseaux de communications doivent être utilisés pour la maintenance des réseaux de transport et de distribution. En d'autres termes, en plus de gérer la production et la consommation, il faut aussi gérer les lignes électriques, les relais et les autres éléments qui composent le réseau électrique. Cette gestion requiert aussi l'utilisation de réseaux de communication. Dans cette section, nous présentons les mécanismes de gestion des réseaux de transport et de distribution et nous présentons les différents réseaux de communication nécessaires pour la mise en place de ces mécanismes.

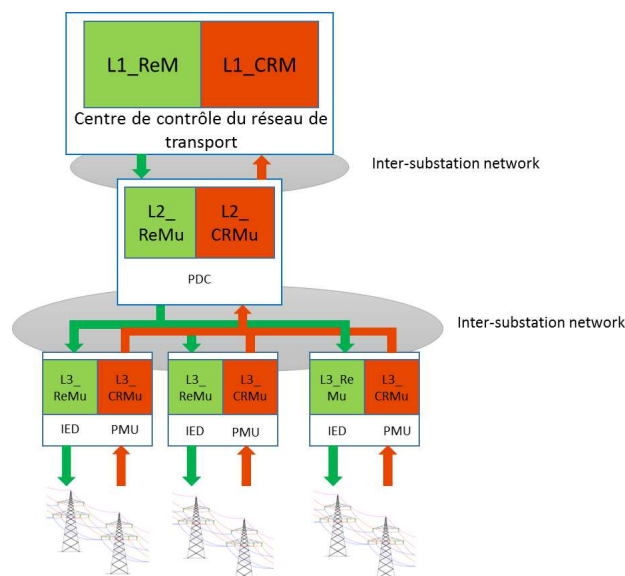


Figure 5.5 – Représentation HDCRAM pour la gestion du réseau de transport dans un réseau électrique intelligent.

5.6.1 HDCRAM pour la gestion du réseau de transport

La gestion du réseau de transport se fait par le biais d'*Intelligent Electronic Devices* (IED) et de *Phasor Measurement Units* (PMU) [101]. Le rôle des PMU est de mesurer le courant et la tension du réseau, tout en étant synchronisés par un signal GPS. Ces PMU envoient leurs informations vers le *Phasor Data Concentrator* (PDC) qui récupère les données de plusieurs PMU, les agrège et les transfère au centre de contrôle. L'opérateur du réseau de transport peut ainsi avoir une vue d'ensemble. C'est ce qu'on appelle le *Wide Area Situational Awareness* (WASA) [102]. Grâce à cette vue globale du réseau de transport, l'opérateur peut prendre des décisions et les appliquer via les actionneurs que sont les IED. Ceci lui permet d'amortir les oscillations électriques qui peuvent endommager le réseau et de mieux gérer les lignes et les installations de secours [103].

La Figure 5.5 montre la représentation HDCRAM de la gestion du réseau de transport. Au niveau 3 de l'architecture, on a les PMU qui ont le rôle de capteur et les IED qui ont le rôle d'actionneur. Au niveau 2, on a les PDC qui servent d'intermédiaires entre les PMU et IED et le centre de contrôle, qui est au niveau 1 de l'architecture de gestion. Dans le réseau de transport, la majorité des décisions demandent une vue globale du réseau et sont donc prises au niveau 1.

La gestion du réseau de transport se fait par le biais d'un unique réseau de communication qui est le réseau *Inter-substation*. Ce réseau est très contraint en terme de latence. En effet, dans [103], les auteurs listent les différents mécanismes qui utilisent les PMU pour la gestion

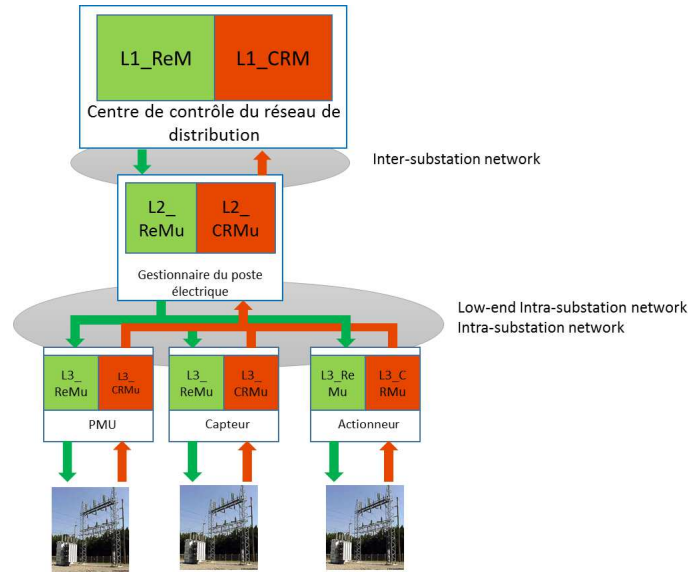


Figure 5.6 – Représentation HDCRAM pour la gestion du réseau de distribution dans un réseau électrique intelligent.

du réseau de transport. Pour la plupart de ces mécanismes la latence doit être très faible, c'est-à-dire en deçà de 50 ms. Seuls les standards de communication sans fil les plus récents peuvent satisfaire une telle contrainte. Pour le réseau *Inter-substation*, on peut donc utiliser la technologie 4G (LTE), ou ses variantes développées pour les objets connectés telles que le *NarrowBand-IoT* (NB-IoT) [104]. De plus, la gestion du réseau de transport est une application typique pour les communications très fiable et à faible latence (*ultra-reliable and low latency*) qui seront introduites avec la cinquième génération de téléphonie mobile (5G) [105].

5.6.2 HDCRAM pour la gestion du réseau de distribution

Le réseau de distribution est celui qui va connaître les plus grands changements lors de la transformation du réseau électrique en réseau électrique intelligent. En effet, avec le développement des énergies renouvelables, la multiplication des sources de stockage et le développement des *microgrids* [106], ce réseau demande aujourd'hui une gestion similaire à celle faite pour le réseau de transport.

En d'autres termes, il faut remplacer le système *Supervisory Control and Data Acquisition* (SCADA) utilisé pour la gestion du réseau de distribution par des mécanismes flexibles capables de se reconfigurer rapidement [107].

Comme montré sur la Figure 5.6, la gestion agile du réseau de distribution va se faire en déployant des capteurs et actionneurs dans les postes électriques. On a donc deux types de réseaux de communication différents pour la gestion du réseau de distribution : ceux qui sont

internes aux postes électriques et ceux qui permettent de communiquer avec le gestionnaire du réseau de distribution.

A l'intérieur des postes électriques on a les réseaux *Intra-substation*, *Low-end intra-substation* qui permettent la communication entre les différents capteurs et actionneurs du poste électrique (niveau 3 de l'architecture) et le gestionnaire de celui-ci (niveau 2 de l'architecture). Pour ces réseaux, les interférences causées par le réseau électrique dans le spectre des ondes radio empêchent l'utilisation de technologies sans-fils. C'est pourquoi des standards de communication filaires sont préférables.

En ce qui concerne le réseau *Inter-substation* qui relie le gestionnaire du poste électrique (niveau 2 de l'architecture) et le centre de contrôle du réseau de distribution (niveau 1 de l'architecture), celui-ci a des contraintes extrêmement strictes. En effet, contrairement au réseau de transport, la stabilité du réseau de distribution n'est pas aidée par des sources d'énergie stables. Ce réseau est donc plus difficile à maintenir à son régime de fonctionnement. Dans [108], les auteurs estiment que la latence de ce réseau doit être inférieure à 10 ms. Une telle contrainte peut être satisfaite avec le standard LTE [109] ou avec une technologie *ultra-reliable and low latency* de la 5G [105].

5.7 Conclusion

Dans ce chapitre, nous avons utilisé l'architecture HDCRAM pour décrire un réseau électrique intelligent. Cette architecture nous a permis de présenter de manière unifiée les différents mécanismes qui sont nécessaires au fonctionnement d'un tel réseau électrique. Une fois ces mécanismes identifiés, nous avons pu positionner sur cette architecture les différents réseaux de communications nécessaires pour l'échange d'information dans un réseau électrique. Nous avons, enfin, discuté des standards de communication qui peuvent être utilisés dans chaque réseau.

Nous avons vu que les standards développées spécifiquement pour les objets connectés jouent un rôle important dans les communications du réseau électrique intelligent. En particulier, les LPWANs sont des réseaux particulièrement intéressants pour l'*AMI backhaul*. Dans le prochain chapitre, nous montrons comment des algorithmes d'apprentissage permettent d'améliorer les communications dans les LPWAN.

Chapitre 6

Algorithmes de bandits pour les réseaux d'objets connectés

Sommaire

6.1	Introduction	110
6.2	Les <i>Low Power Wide Area Networks</i>	110
6.3	Probabilité de collision dans les réseaux LPWAN	111
6.3.1	Réseau ALOHA non-slotté avec le mécanisme d'acquittement du standard LoRaWAN	113
6.3.2	Réseau ALOHA slotté	116
6.4	Latence dans les réseaux LPWAN	119
6.4.1	Sélection aléatoire du canal	121
6.4.2	Transmission uniquement dans le meilleur canal	122
6.5	Algorithmes de bandits multibras	123
6.5.1	L'approche fréquentiste : l'algorithme UCB	124
6.5.2	L'approche bayésienne : l'algorithme Thompson-Sampling	125
6.6	Application aux réseaux d'objets connectés	126
6.7	Validation expérimentale sur Plateforme USRP	128
6.8	Évaluation avec un grand nombre d'objets dynamiques	131
6.8.1	Modèle	132
6.8.2	Stratégies de référence	132
6.8.3	Évaluation des performances	134
6.9	Conclusion	137

6.1 Introduction

Nous avons vu dans le Chapitre 5 que les réseaux d'objets connectés peuvent jouer un rôle important dans la gestion du réseau électrique. En particulier, l'*AMI backhaul* est une application idéale pour les LPWAN.

Dans ce chapitre nous montrons que des algorithmes de bandits multibras [110] peuvent être utilisés pour améliorer les performances des réseaux d'objets connectés et en particulier des LPWAN. Pour cela, nous commençons par décrire les indicateurs de performance d'un réseau LPWAN. Nous verrons que la probabilité de collision est une métrique importante dans ces réseaux. Nous la calculerons dans différents réseaux et verrons son impact sur la latence. Une fois les métriques identifiées et calculées, nous verrons comment les algorithmes de bandit multibras peuvent les améliorer.

Les travaux présentés dans ce chapitre ont donné lieu à un article de revue [111] et à une publication dans une conférence internationale [112].

6.2 Les *Low Power Wide Area Networks*

Les LPWAN sont des standards de communication récemment développés pour l'internet des Objets (IdO) [6]. Plus précisément, ces standards ont pour objectif de connecter n'importe quel objet à internet en passant par une station de base (parfois appelée *gateway*). Ces standards ont deux contraintes :

1. Avoir un faible coût de déploiement pour l'opérateur.
2. La consommation d'énergie des objets qui utilisent ces réseaux doit être faible.

Pour satisfaire la première de ces contraintes, il faut que le nombre de stations de base déployées soit faible. Pour cela, la portée des communications doit être importante. Par conséquent, les protocoles LPWAN sont déployés dans des bandes libres à basse fréquence (en dessous du gigahertz) et utilisent des modulations qui permettent de démoduler même lorsque la puissance du signal reçu est faible. Par exemple, la société Sigfox transmet sur des bandes très fines pour limiter la puissance du bruit de réception, ce qui lui permet de communiquer à plus grande distance [96]. Au contraire, la modulation LoRa utilisée avec le standard LoRaWAN utilise de l'étalement de spectre [113], qui consiste à prendre le problème à l'inverse en bénéficiant d'un gain d'étalement.

La seconde contrainte a deux impacts sur les caractéristiques des communications LPWAN. Le premier est que toute la complexité de la communication doit être concentrée au niveau de

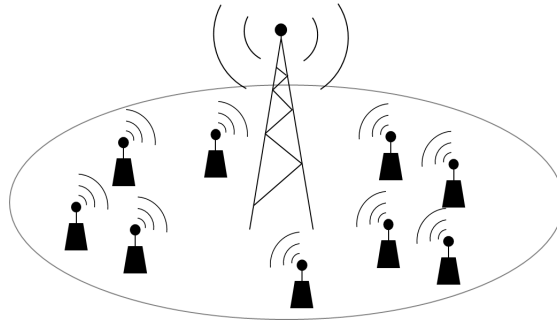


Figure 6.1 – Le modèle considéré dans lequel un grand nombre d’objets connectés envoie des paquets vers une station de base.

la station de base. Les objets ne doivent pas avoir besoin de réaliser des opérations complexes pour communiquer. La seconde est que le protocole utilisé pour communiquer doit être simple et ne demander que très peu de signalisation. On doit donc utiliser des protocoles basés sur de l’ALOHA [95] comme dans les protocoles utilisés par Sigfox ou LoRaWAN.

On peut noter que dans les standards LPWAN, cités ci-dessus et dans le Chapitre 5, aucun ne se base sur de l’ALOHA slotté [114]. Cependant, cette technologie pourrait être utilisée pour ce type de standards et est aujourd’hui utilisée dans des protocoles tels que le Zigbee qui peuvent aussi être utilisés pour le réseau électrique intelligent [115] mais pas pour l’AMI *backhaul* (voir tableau 5.1).

Dans la section suivante, nous étudions la probabilité de collision dans un réseau LPWAN. Nous verrons que c’est une métrique qui permet d’évaluer les performances de ces réseaux.

6.3 Probabilité de collision dans les réseaux LPWAN

Afin de calculer la probabilité de collision dans un réseau LPWAN, on va supposer le système de la Figure 6.1. On a une station de base, par exemple d’un réseau LPWAN, qui échange des informations avec des objets qui peuvent, par exemple, être des agrégateurs du réseau électrique intelligent (AMI *backhaul*). La plupart des communications sont en *uplink*, c’est-à-dire, des objets vers la station de base et aucun échelonnement des paquets n’est effectué par la station de base.

Dans un réseau ALOHA, chacun des objets du réseau envoie ses données par paquets. Chaque objet envoie ses paquets quand il le veut, et, en cas de perte du paquet (par exemple en raison d’une collision avec un autre paquet), le retransmet après un délai aléatoire. Dans le cas où plusieurs canaux sont disponibles, pour chaque retransmission, l’objet peut choisir aléatoirement un des canaux pour chacune de ses communications.

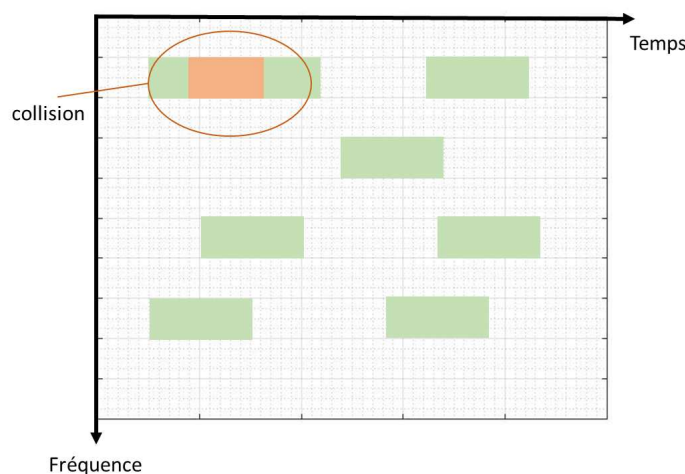


Figure 6.2 – On suppose que, dès que deux paquets se retrouvent au même moment dans un canal, on a une collision et aucun des deux paquets ne peut être décodé.

Dans les LPWAN, les stations de base ont une grande zone de couverture. Par conséquent, le nombre d'objets gérés par chaque station de base est élevé. Dans ce cas là, il est fréquent que la station de base reçoive deux paquets en même temps dans le même canal fréquentiel, il sera alors difficile de les décoder.

On suppose que, comme montré sur la Figure 6.2, dès que deux paquets se chevauchent dans un canal (le premier paquet se termine après que le second ait commencé), les deux paquets ne peuvent être décodés. C'est ce qu'on appelle une collision. Dans la suite, nous allons calculer la probabilité de collision dans certains réseaux d'objets connectés et nous utiliserons cette hypothèse. Bien que simpliste, cette hypothèse est commune dans la littérature [116, 117]. Elle n'est cependant pas toujours utilisée et certains travaux tels que [118] s'en passent. S'en passer permet d'obtenir des résultats plus réalistes. Cependant, cela augmente la difficulté des calculs et les expressions obtenus sont plus complexes et difficiles à interpréter.

Nous venons de décrire les réseaux ALOHA non-slottés ou purs. Cependant, dans cette section, nous nous intéresserons aussi aux réseaux ALOHA slottés. La seule différence entre un réseau ALOHA pur et un réseau ALOHA slotté est que, dans ce dernier, le temps est divisé en slots.

Dans un réseau ALOHA ou dans un réseau ALOHA slotté partagé par un grand nombre d'objets, on peut supposer que, la plupart des pertes de paquets sont dues aux collisions. Dans ce cas là, les différentes métriques d'évaluation du réseau et en particulier la latence et le débit, sont des fonctions de cette probabilité de collision.

Nous préférons ne pas utiliser le débit comme métrique d'évaluation du réseau. Pour justifier ce choix, on considère un réseau ALOHA composé d'un seul canal et dans lequel le

trafic généré par les objets suit une loi de Poisson d'intensité λ sans acquittement de paquets [95], le débit dans un tel réseau est égal à :

$$D = \lambda T_m e^{-2\lambda T_m} \quad (6.1)$$

Où T_m est la durée de chaque paquet. Ce débit est maximum lorsque $\lambda T_m = \frac{1}{2}$. Pour cette valeur de λ , la probabilité de collision est élevée (60 %). Autrement dit, lorsque le débit est maximum dans le réseau, le nombre de collisions est élevé. Ce grand nombre de collisions est désavantageux pour les objets qui doivent transmettre plusieurs fois un paquet avant de réussir à le transmettre correctement. En d'autres termes, le régime de fonctionnement qui maximise le débit total n'est pas bénéfique pour les objets pris individuellement. C'est pour cette raison que dans la suite de ce chapitre, nous utiliserons la probabilité de collision et la latence comme métriques d'évaluation du réseau LPWAN.

Pour l'analyse des performances des algorithmes de bandit multibras, nous considérerons soit des réseaux ALOHA non-slotté (ou ALOHA pur), soit des réseaux ALOHA slottés. Afin d'évaluer les performances de ces algorithmes dans ces réseaux, nous allons commencer par calculer la probabilité de collision. On considèrera en particulier deux réseaux spécifiques :

- On considèrera d'abord un réseau ALOHA non-slotté dans lequel l'acquittement est similaire à celui utilisé dans le standard LoRaWAN [97].
- On considèrera ensuite un réseau ALOHA slotté.

Ensuite, nous donnerons l'expression de la latence dans un réseau ALOHA en fonction de la probabilité de collision et de la stratégie d'accès utilisée par chaque objet. Une fois que ces deux métriques ont été établies, nous les utiliserons pour évaluer les performances des algorithmes de bandits multibras dans des réseaux d'objets connectés.

6.3.1 Réseau ALOHA non-slotté avec le mécanisme d'acquittement du standard LoRaWAN

Nous considérons un réseau LPWAN utilisant un mécanisme d'acquittement similaire à celui utilisé dans le standard LoRaWAN. Dans ce standard, avec les spécifications européennes [119], une fois que la station de base reçoit un paquet, elle envoie un acquittement dans le même canal après un intervalle de durée fixe¹.

On suppose un réseau composé de N_c canaux. Pour l'analyse des collisions, nous faisons les hypothèses suivantes :

1. Dans le standard LoRaWAN, un acquittement peut aussi être envoyé par la station de base dans un second canal dédié aux acquittements, ce canal n'est pas considéré ici.

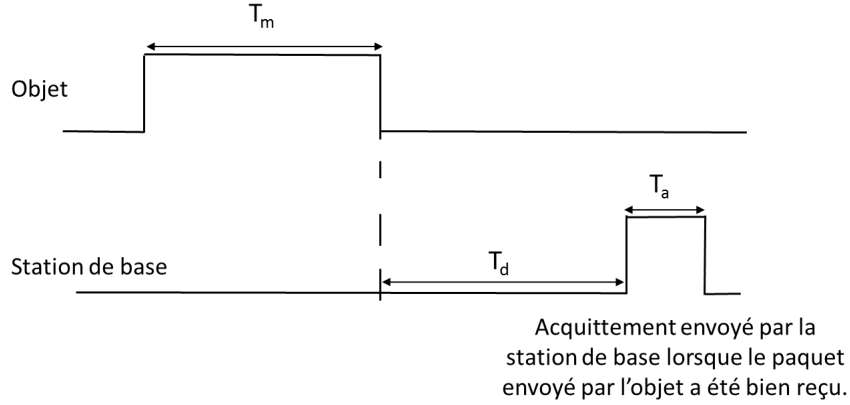


Figure 6.3 – Lorsque la station de base reçoit un paquet envoyé par un objet, elle attend pendant un intervalle de durée T_d et transmet ensuite un acquittement si le canal est vide.

- La probabilité qu'un paquet soit envoyé, à un instant donné, par un des objets dans le canal j suit une loi de Poisson d'intensité λ_j . Il est important de noter ici que le trafic n'est pas forcément identique dans tous les canaux.
- Pour l'analyse des collisions, on suppose que tous les paquets ont la même durée notée T_m . Comme illustré sur la Figure 6.3, on appelle T_d la durée de l'intervalle entre la fin du paquet envoyé par l'objet et le début de la transmission de l'acquittement et T_a la durée de l'acquittement qui est plus courte que T_m .
- On suppose que la station de base a une connaissance parfaite de la présence ou de l'absence d'un signal dans un canal et qu'elle ne transmet un acquittement que si le canal est vide. C'est-à-dire qu'elle applique du *Listen Before Talk* (LBT) qui a deux avantages. Le premier est qu'il permet de réduire le nombre de collisions dans le réseau. De plus, bien que cette contrainte ne soit pas considérée ici, les bandes utilisées par le standard LoRaWAN en Europe sont limitées par un *duty cycle* qui impose une durée d'attente longue entre deux communications [97]. Cependant, en combinant du LBT et de l'*Adaptive Frequency Agility* (AFA), la station de base peut passer outre les restrictions du *duty cycle*. En effet, d'après [120], si en plus de faire du LBT, la station de base a la possibilité de choisir entre plusieurs canaux pour transmettre l'acquittement (AFA), elle n'a plus besoin d'attendre entre deux transmissions dans le même canal.
- La station de base acquitte chacun des paquets reçus. A condition que le canal soit libre au moment d'envoyer cet acquittement.

On peut utiliser deux probabilités de collision pour évaluer les performances d'un canal. La première est la probabilité d'avoir une transmission sans collision sur la voie montante dans le canal j . On notera cette probabilité $P^j(\text{su})$. On peut aussi calculer la probabilité que l'objet reçoive l'acquittement envoyé par la station de base dans le canal j . Cette seconde

probabilité est la probabilité que ni le paquet envoyé sur la voie montante, ni l'acquittement ne subissent de collision. Elle est notée $P^j(sd)$.

Nous dérivons, ici, l'expression de ces deux probabilités. Le détail des calculs menés est donné dans l'annexe C.1. On note que l'on a deux expressions différentes pour $P^j(su)$ et $P^j(sd)$ selon si $T_d \leq T_m$ ou si $T_d \geq T_m$.

Proposition 6.1. *Si $T_d \leq T_m$, la probabilité d'avoir une transmission sans collision sur la voie montante est égale à :*

$$P^j(su) = \frac{e^{-2\lambda_j T_m}}{1 + e^{-\lambda_j(T_d+T_m)} - e^{-\lambda_j(T_d+T_m+T_a)}}. \quad (6.2)$$

La probabilité de recevoir un acquittement sans collision est égale à :

$$P^j(sd) = \frac{e^{-\lambda_j(2T_m+T_d+T_a)}}{1 + e^{-\lambda_j(T_d+T_m)} - e^{-\lambda_j(T_d+T_m+T_a)}}. \quad (6.3)$$

Proposition 6.2. *Si $T_d \geq T_m$, la probabilité d'avoir une transmission sans collision sur la voie montante est égale à :*

$$P^j(su) = \frac{e^{-2\lambda_j T_m}}{1 + f(\lambda_j, T_m, T_d, T_a)}. \quad (6.4)$$

la probabilité de recevoir un acquittement sans collision est égale à :

$$P^j(sd) = \frac{e^{-\lambda_j(3T_m+T_a)}}{1 + f(\lambda_j, T_m, T_d, T_a)}. \quad (6.5)$$

Dans ces deux équations,

$$f(\lambda_j, T_m, T_d, T_a) = \left(e^{-\lambda_j T_m} - e^{-\lambda_j(T_m+T_a)} \right) \left[e^{-\lambda_j T_d} + \frac{1}{\lambda_j T_a} \left(e^{-\lambda_j T_m} - e^{-\lambda_j(T_m+T_a)} - e^{-\lambda_j T_d} + e^{-\lambda_j(T_d+T_a)} \right) \right]. \quad (6.6)$$

Les expressions données ci-dessus peuvent être vérifiées par simulation. Pour cela, on compare la probabilité $P^j(sd)$ approchée avec des simulations de Monte-Carlo et calculée avec les formules proposées ci-dessus. Pour chaque simulation, on suppose que $\lambda_j = \frac{10^{-4}}{T_m} \text{ s}^{-1}$. De plus, $T_d = 1 \text{ s}$ et les valeurs de T_m et T_a sont choisies pour être en adéquation avec le standard LoRaWAN [97]. Les résultats sont montrés sur la Figure 6.4. On voit sur cette figure que les résultats théoriques et numériques coïncident.

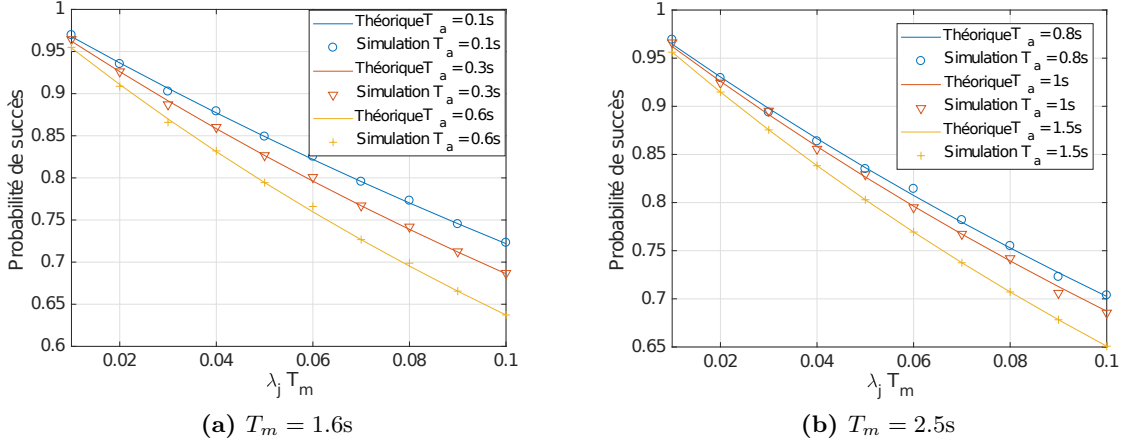


Figure 6.4 – Probabilité d’avoir une bonne transmission à la fois sur la voie montante et la voie descendante ($P^j(sd)$) en fonction de $\lambda_j T_m$. Pour ces simulations on considère que $T_d = 1$ s et les valeurs de T_a et T_m sont en adéquation avec le standard LoRaWAN.

6.3.2 Réseau ALOHA slotté

Le protocole ALOHA slotté peut aussi être utilisé dans les réseaux LPWAN. De plus, ce protocole sera utilisé pour l’évaluation des performances des algorithmes de bandits multibras qui seront proposés dans la suite de ce chapitre. C’est pourquoi, nous étudions, dans cette section, la probabilité de collision dans un tel réseau. On considère un réseau ALOHA slotté utilisé par un grand nombre d’objets qui envoient des paquets de données vers une station de base. Tout comme avec le protocole ALOHA présenté dans la section précédente, chaque objet va choisir l’instant auquel il transmet ses données. A la différence de l’ALOHA non-slotté, ou ALOHA pur, dans un réseau ALOHA slotté, le temps est divisé en slots. Chaque objet choisit les slots pendant lesquels il communique. On notera que, comme illustré sur la Figure 6.5, on suppose que chaque slot est découpé en deux parties : une première pour la voie montante, la seconde pour la voie descendante. Cette voie descendante permet, par exemple, d’acquitter les messages envoyés. Avec un tel découpage, si tous les objets utilisent le même standard, on ne peut avoir des collisions que sur la voie montante. Dans ce cas là, $P^j(su) = P^j(sd)$.

Si deux objets envoient un paquet vers la station de base dans le même slot, on a une collision. Dans ce cas là, la station de base ne pourra pas décoder les paquets et ils seront perdus. La probabilité de collision dans un réseau ALOHA slotté dépend des hypothèses faites et du comportement supposé des objets. Par exemple, si on suppose que dans le canal j on a N_j objets et que chacun envoie un paquet avec une probabilité p suivant un processus de Bernoulli, la probabilité qu’un paquet envoyé ne subisse pas de collision s’écrit [121] :

$$P^j(su) = P^j(sd) = (1 - p)^{N_j - 1} \underset{N_j \rightarrow \infty}{\approx} e^{-N_j p} \quad (6.7)$$

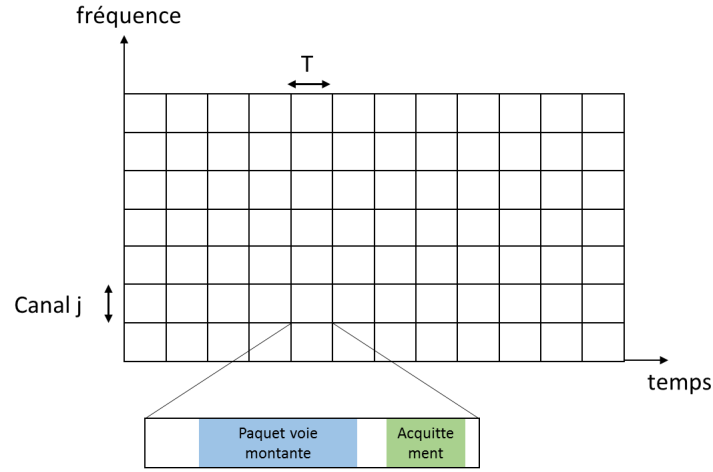


Figure 6.5 – Dans le réseau ALOHA slotté considéré, le temps est divisé en slots de même durée. On suppose que chaque slot se décompose en deux parties. La première est utilisée pour la voie montante, la seconde pour l’acquittement envoyé par la station de base.

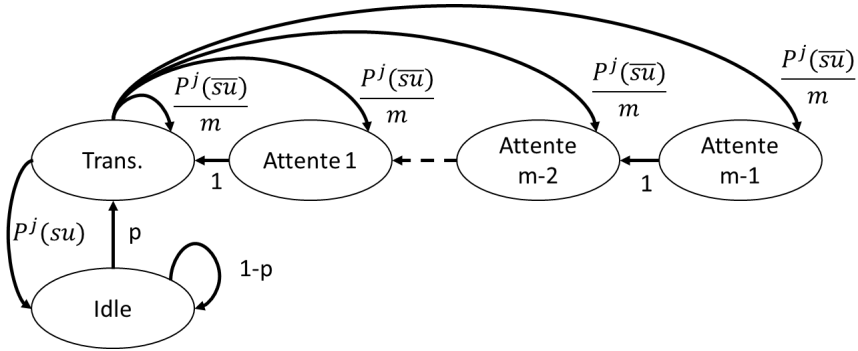


Figure 6.6 – On modélise par une chaîne de Markov le comportement de chacun des objets qui communiquent dans le canal j .

Si on suppose maintenant que, si le message n’est pas acquitté par la station de base, l’objet le retransmet après un délai (nombre de slots) uniformément distribué dans un intervalle de *back-off* prédéfini dans le même canal, on peut modéliser le comportement de chacun des N_j objets par une chaîne de Markov [117, 122] qui est illustrée sur la Figure 6.6. On peut voir sur cette figure que, lorsqu’il est en veille, un objet a une probabilité p d’envoyer un paquet dans un slot temporel donné. Ce paquet sera bien reçu avec une probabilité $P^j(su)$ et aura une probabilité $P^j(\overline{su}) = 1 - P^j(su)$ de subir une collision. Dans ce second cas, le paquet sera retransmis après un nombre de slots tiré aléatoirement dans $\llbracket 0; m - 1 \rrbracket$ (intervalle de *back-off*) suivant une loi uniforme.

Pour étudier la probabilité de collision avec un tel modèle, on calcule d’abord la probabilité que l’objet soit en train d’émettre dans un slot temporel donné. En supposant que la chaîne

de Markov est stationnaire, on peut calculer cette probabilité en dérivant l'expression du vecteur $\mathbf{x}$:

$$\mathbf{x} = \begin{bmatrix} x_v & x_t & x_{a_{m-1}} & \dots & x_{a_1} \end{bmatrix}^T. \quad (6.8)$$

Dans cette équation x_v est la probabilité que l'appareil soit en veille, x_t celle qu'il soit en train de transmettre et x_{a_i} est la probabilité qu'il soit en train d'attendre dans l'état d'attente i avant de retransmettre. On appelle $\mathbf{M}$ la matrice de transition de la chaîne de Markov donnée sur la Figure 6.6 qui a pour expression :

$$\mathbf{M} = \begin{bmatrix} 1-p & 1-P^j(\overline{\text{su}}) & 0 & \dots & 0 & 0 \\ p & \frac{P^j(\overline{\text{su}})}{m} & 0 & \dots & 0 & 1 \\ 0 & \frac{P^j(\overline{\text{su}})}{m} & 0 & \dots & 0 & 0 \\ \vdots & \vdots & 1 & \dots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & \frac{P^j(\overline{\text{su}})}{m} & 0 & \dots & 1 & 0 \end{bmatrix}, \quad (6.9)$$

$\mathbf{x}$ vérifie :

$$\mathbf{M}\mathbf{x} = \mathbf{x}. \quad (6.10)$$

En résolvant ce système d'équations linéaires et en utilisant le fait que la somme des éléments de $\mathbf{x}$ est égale à 1, on obtient l'expression suivante pour x_t , la probabilité que l'objet soit en train de transmettre dans un slot donné :

$$x_t = \frac{1}{1 + \frac{P^j(\text{su})}{p} + (1 - P^j(\text{su}))\frac{(m-1)}{2}}. \quad (6.11)$$

De plus, si l'intervalle de *back-off* est suffisamment grand et qu'on a un grand nombre d'objets, on peut supposer que la probabilité de collision suit une loi de Bernoulli. Dans ce cas là :

$$P^j(\text{su}) = (1 - x_t)^{N_j-1} \underset{N_j \gg 1}{\approx} e^{-N_j x_t}. \quad (6.12)$$

Comme montré par l'équation (6.11), x_t est une fonction de $P^j(\text{su})$. Par conséquent, l'équation (6.12) est une expression implicite de $P^j(\text{su})$. On peut cependant proposer une approximation de $P^j(\text{su})$ dans le cas où le réseau est saturé et le nombre de collisions élevé.

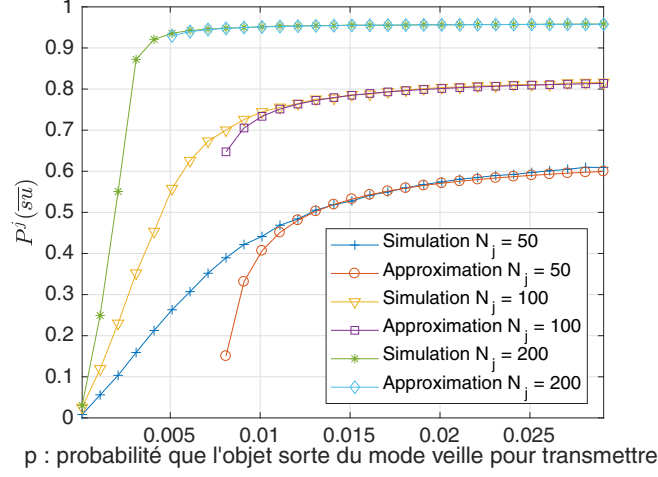


Figure 6.7 – Approximation de la probabilité de collision $P^j(su)$ dans un réseau ALOHA slotté.

Proposition 6.3. *Lorsque la probabilité de collision dans un réseau ALOHA slotté est élevée, elle peut être approximée par :*

$$P^j(\overline{su}) = 1 - P^j(su) \approx 1 + \frac{\mathcal{W}_0 \left(-\frac{4N_j}{(m+1)^2} \left(\frac{1}{p} - \frac{m-1}{2} \right) e^{-\frac{2N_j}{m+1}} \right)}{\left(\frac{1}{p} - \frac{m-1}{2} \right) \frac{4N_j}{(m+1)^2}}. \quad (6.13)$$

La preuve est fournie en Annexe C.2.

On peut vérifier par simulation l'approximation proposée ici. Pour cela, on suppose que $m = 128$, et on trace la probabilité de collision $P^j(\overline{su}) = 1 - P^j(su)$ en fonction de p , la probabilité qu'un objet veuille envoyer un paquet et pour différentes valeurs de N_j , le nombre d'objets dans le canal. Les résultats sont donnés sur la Figure 6.7. On peut voir sur cette figure que, l'approximation proposée est valide lorsque le nombre de collisions est élevé.

6.4 Latence dans les réseaux LPWAN

Une fois la probabilité de collision calculée, on peut s'intéresser aux métriques qui en dépendent. En particulier, on peut regarder l'impact de cette probabilité de collision sur la latence. En effet, plus on a de collisions, plus un objet va devoir retransmettre un paquet avant que celui-ci ne soit bien reçu par la station de base. L'augmentation du nombre de la probabilité de collision augmente, donc, la latence dans le réseau. Pour l'analyse de la latence on considère un objet, par exemple un agrégateur de données du réseau électrique intelligent, qui veut communiquer avec une station de base dans un réseau LPWAN comme présenté sur la Figure 6.8. Ce réseau de type ALOHA (slotté ou non) est utilisé par un grand nombre

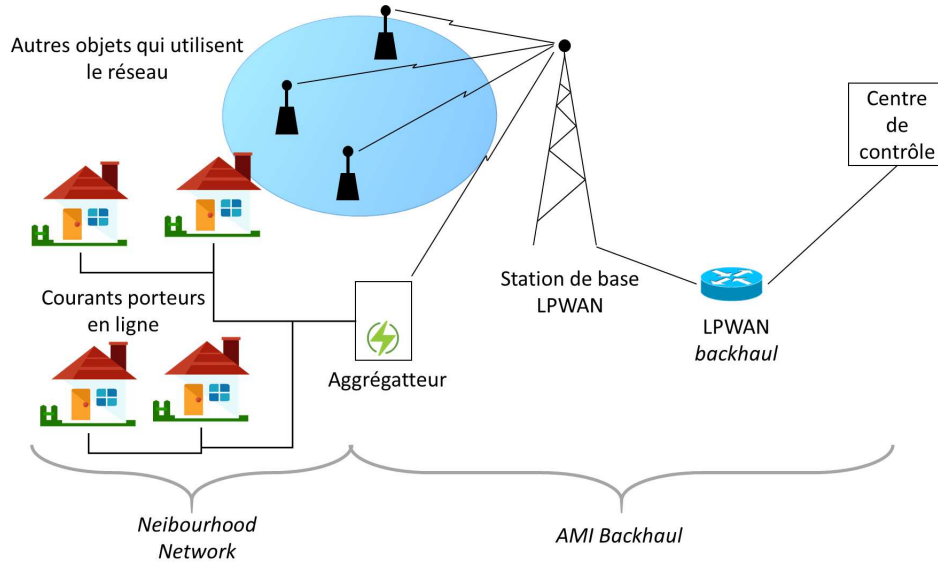


Figure 6.8 – On s'intéresse aux communications d'un agrégateur de données qui communique avec un centre de contrôle par le biais d'un réseau LPWAN.

d'objets et la bande utilisée est découpée en canaux. Dans chacun de ces canaux, on a une probabilité de collision sur la voie montante notée $P^j(\overline{s}u)$ et une probabilité de collision sur la voie descendante notée $P^j(\overline{s}d)$. L'expression de ces probabilités dépend du type de réseau, du comportement des objets et de l'environnement des communications. On s'intéresse, dans cette section à la latence dans un réseau LPWAN de type ALOHA dans lequel le protocole utilisé peut être, par exemple, l'un de ceux décrits dans les Sections 6.3.1 et 6.3.2.

Dans tous les cas, la probabilité de collision a un impact sur la latence des communications de l'agrégateur. Pour l'étude de cette latence, nous faisons les hypothèses suivantes :

- La durée qui sépare l'envoi de deux paquets indépendants (qui contiennent des données issues de mesures différentes) est grande devant la durée qui sépare deux retransmissions d'un même paquet. Par exemple, la première est de l'ordre de quelques dizaines de minutes voire d'heures alors que la seconde est de l'ordre de quelques secondes.
- Si le paquet envoyé par l'agrégateur n'est pas acquitté, il le retransmet après un temps aléatoire uniformément distribué dans l'intervalle de *back-off* de longueur T_{bo} . On note que, dans un réseau ALOHA slotté, la durée de l'intervalle de *back-off* est égale à $T_{bo} = mT_s$, où T_s est la durée d'un slot et m est le nombre de slots qui composent l'intervalle de *back-off*.
- La probabilité de collision dans un canal ne dépend pas de l'index de retransmission.

La latence dépend de la stratégie d'accès au spectre utilisée par notre agrégateur de données. On peut comparer deux stratégies différentes :

1. L'agrégateur choisit aléatoirement un des N_c canaux disponibles pour ses transmissions.

2. L'aggrégateur ne transmet que dans le meilleur canal.

6.4.1 Sélection aléatoire du canal

Quelle que soit la stratégie d'accès au spectre utilisée par l'aggrégateur, on peut utiliser le théorème de l'espérance totale pour réécrire l'espérance de la latence $\mathbb{E}[\mathcal{L}]$ des communications :

$$\mathbb{E}[\mathcal{L}] = \sum_{i=1}^M P(N_t = i) \mathbb{E}[\mathcal{L} | N_t = i]. \quad (6.14)$$

Dans cette équation, N_t est le nombre de fois qu'un paquet a été transmis avant d'être reçu et correctement décodé par la station de base. M désigne le nombre de transmissions maximum pour un paquet. C'est à dire que si la station de base ne reçoit pas le paquet après M transmissions, il n'est pas renvoyé par l'aggrégateur. De plus, avec les stratégies d'accès considérées, la latence sachant que l'on a i transmissions, $\mathbb{E}[\mathcal{L} | N_t = i]$, ne dépend pas de la stratégie d'accès de l'aggrégateur et est égale à :

$$\mathbb{E}[\mathcal{L} | N_{ret} = i] = (i - 1) (T_m + T_c + \mathbb{E}[T_r]) + T_m \quad (6.15)$$

$$= (i - 1) \left(T_m + T_c + \frac{T_{bo}}{2} \right) + T_m. \quad (6.16)$$

Dans cette équation, T_c est le temps pendant lequel l'aggrégateur attend son acquittement. Si après T_c l'acquittement n'est pas reçu, il va calculer un temps aléatoire T_r et ensuite retransmettre le paquet. De plus, vu que nous supposons que la probabilité de collision ne dépend pas de l'indice i de transmission, si l'aggrégateur choisit aléatoirement son canal, la probabilité que l'aggrégateur envoie avec succès un paquet est égale à :

$$P(\text{su}) = \frac{1}{N_c} \sum_{j=1}^{N_c} P^j(\text{su}) = P_{\text{moy}}(\text{su}), \quad (6.17)$$

où $P_{\text{moy}}(\text{su})$ est la probabilité de succès moyenne dans les canaux. La probabilité d'avoir N_t transmissions avant d'en avoir une réussie s'écrit donc :

$$P(N_t = i) = (1 - P_{\text{moy}}(\text{su}))^{i-1} P_{\text{moy}}(\text{su}). \quad (6.18)$$

Cette équation nous permet d'écrire la latence :

$$\mathbb{E}[\mathcal{L}] = P_{\text{moy}}(\text{su})T_m \sum_{i=1}^M (1 - P_{\text{moy}}(\text{su}))^{i-1} + P_{\text{moy}}(\text{su}) \left(T_m + T_c + \frac{T_{bo}}{2} \right) \sum_{i=2}^M (i-1) (1 - P_{\text{moy}}(\text{su}))^{i-1}. \quad (6.19)$$

On utilise maintenant l'expression de la somme de la dérivée d'une série géométrique pour déduire l'expression de la latence lorsque M est grand. Dans ce cas, on peut approximer la latence par sa limite lorsque M tend vers l'infini :

$$\mathbb{E}[\mathcal{L}] \xrightarrow{M \rightarrow \infty} \left(T_m + T_c + \frac{T_{bo}}{2} \right) \frac{1 - P_{\text{moy}}(\text{su})}{P_{\text{moy}}(\text{su})} + T_m. \quad (6.20)$$

6.4.2 Transmission uniquement dans le meilleur canal

On suppose maintenant que, au lieu de choisir aléatoirement un canal pour chacune de ses transmissions, l'aggrégateur ne transmet que dans le canal le plus libre. Dans ce cas, la probabilité que la transmission d'un paquet soit réussie est égale à $P^{j*}(\text{su})$, la probabilité de succès dans le meilleur canal. En réitérant les dérivations faites dans la section précédente, on dérive l'expression de la latence lorsque M est grand est égale à :

$$\mathbb{E}[\mathcal{L}] \xrightarrow{M \rightarrow \infty} \left(T_m + T_c + \frac{T_{bo}}{2} \right) \frac{1 - P^{j*}(\text{su})}{P^{j*}(\text{su})} + T_m. \quad (6.21)$$

En comparant l'expression donnée par l'équation (6.20) et celle de (6.21), on se rend compte que, bien évidemment, il y a un intérêt à choisir le meilleur canal pour chacune des transmissions. En effet, en calculant la différence entre les deux expressions données ci-dessus on obtient :

$$\mathbb{E}[\mathcal{L}]_{\text{rand}} - \mathbb{E}[\mathcal{L}]_{\text{MC}} = \left(T_m + T_c + \frac{T_{bo}}{2} \right) \left(\frac{1}{P_{\text{moy}}(\text{su})} - \frac{1}{P^{j*}(\text{su})} \right). \quad (6.22)$$

Dans cette équation, $\mathbb{E}[\mathcal{L}]_{\text{rand}}$ est la latence avec une sélection aléatoire du canal et $\mathbb{E}[\mathcal{L}]_{\text{MC}}$ est la latence lorsque l'aggrégateur n'émet que dans le meilleur canal.

Pour améliorer la latence, un aggrégateur doit donc s'insérer dans le meilleur canal. L'aggrégateur ne peut pas, à priori, savoir quel canal est le meilleur. Pour acquérir cette connaissance il faut :

1. Soit que la station de base lui transmette cette information. Le problème de cette solution est qu'elle rajoute une surcharge de signalisation et que l'information envoyée par la station de base peut être inexacte vu que l'objet peut subir des interférences que la station de base ne voit pas.
2. Une seconde solution consiste à utiliser des algorithmes d'apprentissage pour que l'aggrégateur apprenne quel est le meilleur canal.

Dans la suite de ce chapitre, nous présentons les algorithmes de bandits multibras qui permettent à l'aggrégateur et à n'importe quel objet d'utiliser quasi-exclusivement le meilleur canal et ainsi de réduire la probabilité d'avoir des collisions et donc la latence de ses communications.

6.5 Algorithmes de bandits multibras

Les algorithmes de bandits multibras [110, 123] sont des algorithmes de prise de décision qui s'appliquent dans une situation où un agent fait face à N possibilités. Par exemple, il peut s'agir d'un joueur de casino qui fait face à N machines à sous. Toutes ces possibilités lui apportent une même récompense, mais la probabilité d'obtenir cette récompense change d'une possibilité à l'autre (d'une machine à sous à l'autre). Lorsqu'il joue, l'utilisateur obtient de l'information à propos de la machine qu'il vient d'utiliser. Grâce à cette information il peut apprendre au fil du temps quelle machine à sous permet de gagner le plus souvent et ainsi n'utiliser que celle là.

Ces algorithmes ont déjà été proposés pour résoudre des problématiques radio. Par exemple, les auteurs de [124] les utilisent pour améliorer la qualité des communications dans un scénario d'accès opportuniste au spectre. Ces algorithmes ont aussi été utilisés dans des réseaux d'objets connectés pour la sélection des stations de base dans [125].

Dans cette thèse, nous utilisons les algorithmes de bandit multibras pour l'accès fréquentiel dans des réseaux d'objets connectés. On suppose un objet (un aggrégateur de données du réseau électrique) qui peut envoyer ses paquets vers une station de base d'un réseau LPWAN dans N_c canaux. Du point de vue de l'aggrégateur, la transmission est réussie s'il reçoit l'acquittement que lui envoie la station de base. La probabilité de recevoir l'acquittement n'est pas forcément la même dans tous les canaux. Dans le canal j , on a une probabilité $P^j(\text{sd})$ que l'aggrégateur reçoive l'acquittement. La bonne ou mauvaise réception de cet acquittement apporte de l'information à l'objet. Cet acquittement est, donc, la récompense que peut obtenir l'aggrégateur lorsqu'il transmet dans le canal j . Dans la suite, on notera $r_t(j)$ la récompense reçue par l'aggrégateur après la transmission t qui s'est faite dans le canal j ; celle-ci est égale à :

$$r_t(j) = \begin{cases} 1 & \text{si l'aggrégateur reçoit l'acquittement,} \\ 0 & \text{sinon.} \end{cases} \quad (6.23)$$

Avec ces algorithmes d'apprentissage, on a un compromis entre exploitation et exploration. En effet, l'aggrégateur va devoir explorer suffisamment tous les canaux pour en avoir une connaissance fine et ainsi choisir le meilleur d'entre eux. Il va aussi devoir exploiter un maximum le meilleur canal afin de maximiser sa récompense. Plusieurs algorithmes ont été proposés dans la littérature pour résoudre ce compromis. Ces algorithmes peuvent être classés dans deux familles [126]. Les premiers sont les algorithmes fréquentistes avec lesquels le canal est choisi de manière déterministe en fonction des récompenses précédentes. Les seconds sont les algorithmes bayésiens, avec lesquels le choix du canal se fait en tirant un échantillon d'une distribution de probabilité dont les paramètres dépendent des récompenses reçues précédemment. Dans les sections suivantes, nous décrivons et évaluons les performances de deux algorithmes, chacun appartenant à une famille différente. Le premier est l'algorithme *Upper Confidence Bound* (UCB) [127], qui est un algorithme fréquentiste. Le second est l'algorithme *Thompson Sampling* (TS) qui est un algorithme bayésien [128, 129]. Il a été prouvé dans la littérature que ces algorithmes ont des performances asymptotiques optimales.

6.5.1 L'approche fréquentiste : l'algorithme UCB

Avec l'algorithme UCB, avant chaque nouvelle transmission (transmission t), l'aggrégateur calcule la moyenne arithmétique des récompenses reçues dans chaque canal :

$$\bar{X}_j(t) = \frac{1}{T_j(t)} \sum_{l=0}^{t-1} r_l(j) \mathbf{1}(a_l = j). \quad (6.24)$$

Dans cette équation $T_j(t)$ est le nombre de transmissions qui ont eu lieu dans le canal j et $\mathbf{1}(a_l = j)$ désigne la fonction indicatrice qui est égale à 1 si la $l^{\text{ième}}$ transmission a été faite dans le canal j . Cette moyenne est ensuite utilisée pour calculer l'index UCB du canal qui est égal à :

$$B_j(t) = \bar{X}_j(t) + A_j(t). \quad (6.25)$$

Dans cette équation, $A_j(t)$ est un biais d'exploration qui, pour l'UCB₁, version la plus courant de l'algorithme UCB, est égal à :

$$A_j(t) = \sqrt{\frac{\alpha \ln t}{T_j(t)}}. \quad (6.26)$$

Dans l'équation (6.26), α est le coefficient d'exploration. Si ce coefficient est élevé, l'algorithme va être exploratoire et avoir tendance à utiliser plus souvent tous les canaux pour parfaire la connaissance qu'il en a. Si, au contraire, ce coefficient est faible, l'algorithme va avoir tendance à se focaliser rapidement dans le canal qui lui a apporté la meilleure récompense grâce à la connaissance qu'il a acquise jusqu'à présent. Les performances de l'algorithme présenté ici sont théoriquement d'ordre optimal pour une valeur de $\alpha > 0.5$ [127]. Cependant, il a été montré expérimentalement dans [130], que cet algorithme est efficace même pour des valeurs de $\alpha \leq 0.5$. Avant chaque transmission, l'agrégateur calcule l'index UCB de tous les canaux avec l'équation (6.25). Sa prochaine transmission se fera dans le canal qui a l'index le plus élevé :

$$a_t = \underset{j}{\operatorname{argmax}}(B_j(t)). \quad (6.27)$$

On peut noter que l'algorithme UCB₁ est un algorithme simple, que ce soit en terme de compréhension, de facilité d'implémentation ou de complexité de mise en œuvre tant au sens informatique du terme qu'au sens électronique du terme. En effet, seules quelques opérations et quelques espaces mémoires sont nécessaires à cet algorithme. Cette simplicité lui permet d'être implémentable dans n'importe quel objet connecté, aussi fortes soient ses contraintes de consommation et de complexité, et en particulier dans un agrégateur de données du réseau électrique intelligent.

6.5.2 L'approche bayésienne : l'algorithme Thompson-Sampling

Comme nous l'avons déjà expliqué, contrairement à l'algorithme UCB qui est un algorithme fréquentiste, l'algorithme TS est un algorithme bayésien. Avec cet algorithme, les index sont des réalisations d'une loi bêta dont les paramètres dépendent de l'expérience acquise lors des précédentes émissions. On note :

$$S_j(t) = \sum_{l=0}^{t-1} r_l(j) \mathbb{1}_{\{a_l=j\}}, \quad (6.28)$$

la somme des récompenses acquises à l'instant t dans le canal j , et :

$$F_j(t) = T_j(t) - S_j(t), \quad (6.29)$$

est le nombre de transmissions qui ont échouées dans ce même canal. Avec l'algorithme TS, l'index du canal j à l'instant t est tiré suivant la loi bêta suivante :

$$B_j(t) \sim \beta(1 + S_j(t), 1 + F_j(t)). \quad (6.30)$$

Ensuite, tout comme avec l'algorithme UCB, le canal avec l'index le plus élevé est choisi pour la $t^{\text{ième}}$ transmission. Avant la première transmission, l'algorithme n'a aucune information sur les canaux. Les index sont donc uniformément répartis dans $[0; 1]$. Au fur et à mesure que l'algorithme apprend, la distribution de la loi beta dans le canal j , $\beta(1 + S_j(t), 1 + F_j(t))$ devient de plus en plus fine (sa variance diminue) et se centre autour de $P^j(\text{sd})$. Ainsi, au cours de l'apprentissage l'algorithme va de plus en plus utiliser le canal dans lequel $P^j(\text{sd})$ est la plus élevée (meilleur canal).

6.6 Application aux réseaux d'objets connectés

Nous proposons, dans cette section, d'évaluer l'impact des algorithmes d'apprentissage, présentés dans la section précédente, sur la latence des communications dans un réseau d'objets connectés. Les algorithmes proposés ici pourraient être utilisés par n'importe quel objet d'un réseau LPWAN. Nous nous intéresserons en particulier, mais sans aucune restriction, au cas où les objets qui utilisent les algorithmes proposés sont des agrégateurs du réseau électrique intelligent.

Pour nos simulations, on suppose un réseau LPWAN dans lequel l'accès au spectre se fait suivant un protocole ALOHA non-slotté et dans lequel la station de base acquitte chaque paquet. Le mécanisme d'acquiescement est similaire à celui du standard LoRaWAN [97] qui a été présenté dans la Section 6.3.1. Avec ce mécanisme, l'acquiescement est envoyé dans le canal qui a été utilisé pour le paquet envoyé par l'objet sur la voie montante. Tous les paquets émis sur la voie montante ont la même durée égale à 0.7 ms, si le paquet est reçu par la station de base sans subir de collision, la station de base attend pendant un intervalle de durée $T_d = 1$ s et, si le canal est libre, envoie un acquiescement de durée $T_a = 0.3$ s. Un paquet est au maximum transmis 5 fois sur la voie montante. On suppose que l'intervalle de *back-off* est de 10 s. On note que les durées considérées ici sont en accord avec le standard LoRaWAN [131]. Par ailleurs, la bande passante utilisée par le réseau LPWAN considéré est divisée en 10 canaux.

Dans ce réseau on a deux types d'objets. On a d'abord les objets dit statiques². Ces objets n'utilisent qu'un seul canal. Chacun de ces objets envoie un paquet vers la station de base suivant un processus de Poisson. L'intensité de ce processus est telle que, en moyenne, chacun d'eux envoie un paquet toutes les deux heures (sans compter les retransmissions). Dans

2. Dans cette section et jusqu'à la fin de ce chapitre, nous utiliserons le terme statique pour désigner un objet qui ne peut utiliser qu'un seul canal et le terme dynamique pour parler désigner un objet qui peut utiliser plusieurs canaux

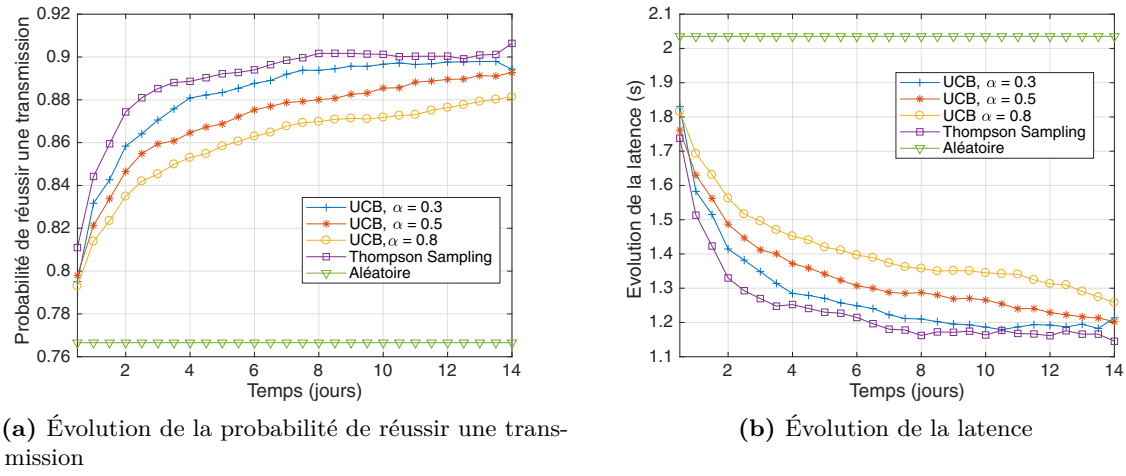


Figure 6.9 – Évolution de la probabilité de réussir une transmission et de la latence avec les algorithmes UCB et TS.

des bandes non-licenciées, le trafic interférant peut être généré par des objets qui utilisent différents standards ou qui n'ont pas tous le même fabricant. On peut donc avoir un trafic interférant inégalement réparti entre les canaux. Pour le générer, on suppose que l'on a 1000 objets statiques dans le premier canal, 900 dans le deuxième, et ainsi de suite jusqu'à 100 dans le dernier canal.

Ces objets statiques partagent le réseau LPWAN avec 50 agrégateurs du réseau électrique qui intègrent les algorithmes proposés ici. Chacun d'eux envoie des paquets suivant un processus de Poisson avec en moyenne un paquet envoyé toutes les 30 minutes. Ces agrégateurs ont un comportement dynamique qui leur permet de sélectionner leur canal pour chacune de leurs transmissions. De leur point de vue, le trafic généré par les objets statiques est interférant. Ce trafic étant inégalement réparti dans les différents canaux, les agrégateurs peuvent utiliser des algorithmes d'apprentissage afin d'éviter les canaux les plus engorgés. Afin de montrer l'effet des algorithmes UCB et TS, on simule le réseau considéré pendant 14 jours et on regarde l'évolution de la latence des communications des agrégateurs et de la probabilité qu'un agrégateur réussisse une transmission.

Sur la Figure 6.9, on compare les performances (en terme de latence et de probabilité de succès) obtenues avec les algorithmes TS et UCB avec celles que l'on obtient si les agrégateurs choisissent aléatoirement le canal qu'ils utilisent pour chaque transmission. Avec les algorithmes d'apprentissage, pendant les premières transmissions, les agrégateurs ont peu de connaissances des canaux. Ils les essayent donc tous. C'est pourquoi, au début de l'apprentissage, les algorithmes proposés ont des performances similaires à celles de la stratégie aléatoire. Très rapidement, au fur et à mesure des transmissions, les agrégateurs acquièrent de la connaissance sur les différents canaux. Ils l'utilisent, ensuite, pour se focaliser

dans les canaux dans lesquels la probabilité de succès est la plus élevée. Ainsi, après 14 jours d'exploitation, les agrégateurs ne communiquent avec la station de base que dans les meilleurs canaux (canaux les plus libres). Ils augmentent donc leur probabilité de succès et réduisent ainsi la latence des communications. Dans le scénario simulé, les algorithmes d'apprentissage permettent d'augmenter de 14 % la probabilité de succès et ainsi de réduire de 40 % la latence par rapport à une sélection aléatoire. On voit ici qu'un gain apparemment faible sur la probabilité de succès a un impact important sur la latence.

6.7 Validation expérimentale sur Plateforme USRP

Nos résultats de simulation montrent que les algorithmes de bandits multibras améliorent les performances des communications dans des réseaux IdO. Afin de les vérifier, nous avons réalisé, avec Lilian Besson, doctorant de l'équipe SCEE, une implémentation sur plateforme *Universal Software Radio Peripheral* (USRP), de *National Instruments*, de ces algorithmes d'apprentissage³. Cette implémentation c'est faite sur le testbed mis en place par notre équipe pour ses expérimentations (SCEE TestBed) [132].

Les cartes USRP sont des plateformes de radio logicielle sur lesquelles nous pouvons coder nos propres fonctions radio. Chacune de ces cartes radio logicielle permet d'émettre et de recevoir des données sur différentes bandes de fréquence.

Pour cette démonstration, nous avons utilisé trois plateformes USRP N210 :

- La première USRP a le rôle de station de base du réseau LPWAN. Elle acquitte les paquets envoyés par l'objet dynamique (agrégateur).
- La seconde USRP joue le rôle de l'agrégateur. C'est donc un objet dynamique dont l'objectif est d'envoyer des paquets sur la voie montante vers la station de base. Un algorithme d'apprentissage est implémenté dans cette plateforme (UCB). Elle peut donc apprendre à n'utiliser que les canaux les plus libres.
- Une troisième USRP génère un trafic interférant qui serait celui issu d'un grand nombre d'objets à portée de la station de base et va parfois interférer avec les communications entre les deux premières USRP.

Cette réalisation expérimentale a deux objectifs :

- Elle permet d'abord de souligner la facilité d'implémentation des algorithmes d'apprentissage présenté dans ce chapitre. Elle montre aussi la possibilité de les faire fonctionner en temps réel.

3. Cette démonstration a été présentée à la conférence ICT 2018 à Saint-Malo.

- Ensuite, elle nous permet de valider expérimentalement le fonctionnement de ces algorithmes sur des systèmes réels et donc dans lesquels nous pouvons avoir des erreurs et des dysfonctionnements.

L'objectif de cette démonstration étant d'évaluer des algorithmes d'accès, nous utilisons une couche physique simple. Comme illustré sur la Figure 6.10, les horloges des USRP sont synchronisées par une octoclock de chez National Instrument. Cette synchronisation nous permet de nous passer de certaines étapes de traitement du signal de synchronisation.

On pourrait utiliser la couche physique de n'importe quel standard IdO. Mais pour éviter les développements couteux en temps, nous avons conçu notre propre couche physique que voici. La station de base écoute dans quatre canaux de même largeur de bande. Les données envoyées dans les différents canaux sont modulées par une QPSK. L'agrégateur (première USRP) envoie des paquets vers la station de base dans ces canaux. Les paquets envoyés sont composés de deux parties : un préambule qui sert à la synchronisation et un index qui est répété plusieurs fois (cet index est un symbole QPSK qui peut soit être égal à $1 + j$, soit égal à $-1 + j$)⁴. Lorsque ces paquets sont reçus par la station de base, elle vérifie qu'elle retrouve bien les index. Si c'est le cas, elle va attendre pendant une durée fixe puis acquitter le paquet reçu dans le canal utilisé pour l'émission. Cet acquittement est aussi découpé en deux parties : on a un préambule qui est suivi par le conjugué de l'index de l'objet à acquitter. Ainsi, lorsque l'objet reçoit l'acquittement et le démodule, il vérifie si l'index décodé est bien le conjugué de son index. Si c'est le cas, il sait que le paquet qu'il a transmis a bien été reçu par la station de base. Il peut donc ajouter 1 à la somme des récompenses du canal dans l'algorithme UCB.

Afin de rendre la démo plus visuelle, les paquets envoyés par l'USRP 1 et les acquittements durent une seconde. La troisième USRP, elle, envoie des paquets interférants qui sont plus courts, mais suffisamment longs pour gêner la bonne réception des paquets envoyés par l'agrégateur. L'intensité du trafic généré dans les différents canaux est réglable et nous permet donc d'évaluer les performances de l'algorithme UCB pour différents taux d'utilisation des canaux par les objets aux alentours.

Sur la Figure 6.11, nous montrons une vue temps/fréquence des quatre canaux qui composent le réseau tels qu'ils sont vus par la station de base (USRP 2). Dans ces canaux, on retrouve les paquets envoyés par l'USRP 1, les acquittements et le trafic interférant généré dans tous les canaux par l'USRP 3.

Pour notre expérimentation, on génère, avec l'USRP 3, un trafic interférant qui est tel que : le canal 1 est utilisé 20 % du temps, le canal 2, 10 % du temps, le canal 3, 5 % du temps et le dernier canal 25 % du temps. On peut mesurer expérimentalement la probabilité

4. j désigne ici le nombre imaginaire donc le carré est égal à -1

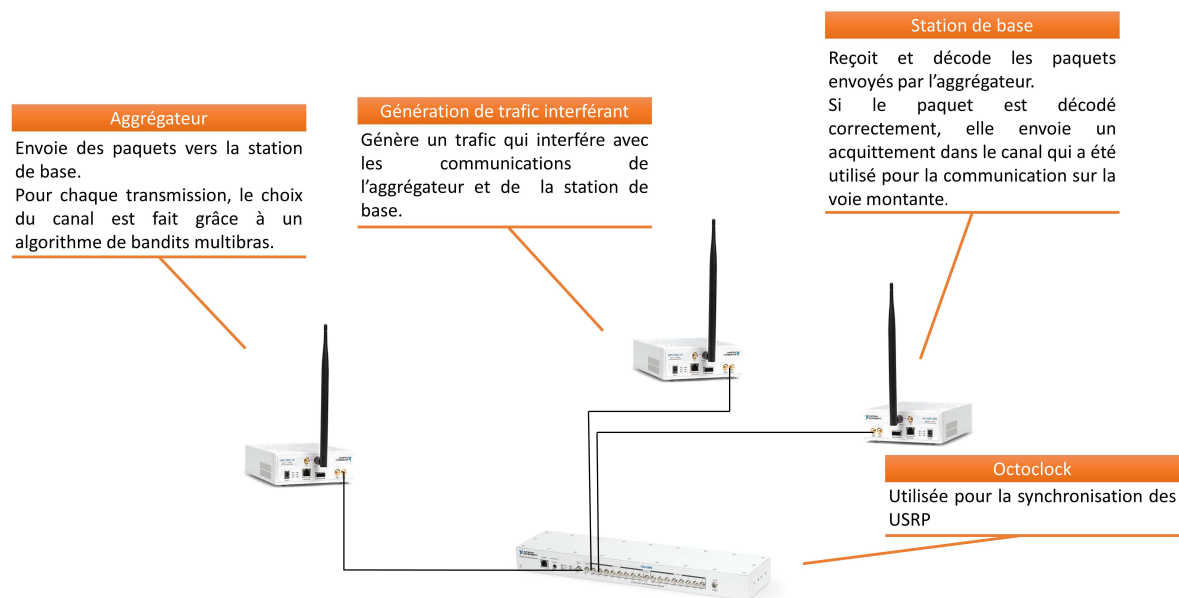


Figure 6.10 – Schéma de la démonstration réalisée. On utilise trois USRP synchronisées par une octoclock. La première USRP joue le rôle d'aggrégateur qui transmet régulièrement des paquets vers une seconde USRP qui joue le rôle de station de base et qui acquitte les paquets qu'elle reçoit correctement. Enfin une troisième USRP génère un trafic interférant (généré par les objets à portée).

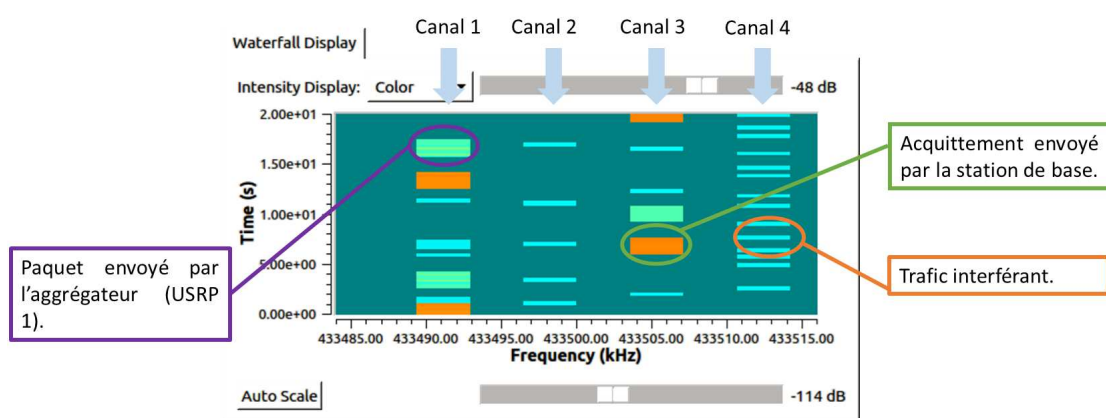


Figure 6.11 – Les quatre canaux tels qu'ils sont vus par la station de base.

que l'acquittement soit bien reçu par l'objet dans chaque canal, avec ces taux d'occupation, elle est égale à $[0.69 ; 0.80 ; 0.91 ; 0.58]$.

On montre, sur la Figure 6.12, le taux d'occupation de chaque canal, avec l'algorithme UCB et avec une sélection aléatoire après plus de 200 transmissions. On voit que, comme attendu, l'algorithme UCB permet à l'aggrégateur de transmettre la majorité de ses paquets dans le canal le plus libre (canal 3) qui est occupé à 5 % par le trafic interférant. Il améliore, ainsi, la qualité de ses communications. Avec l'algorithme UCB on a, après environ 200

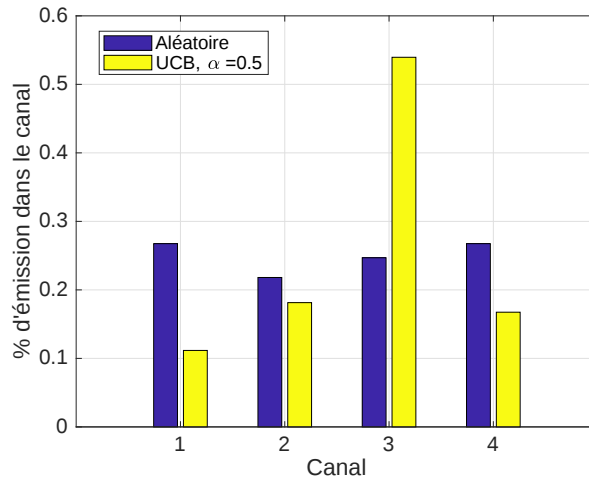


Figure 6.12 – pourcentage d’utilisation de chaque canal par l’aggrégateur avec l’algorithme UCB et avec une sélection aléatoire du canal. On voit que l’aggrégateur a concentré ses émissions sur le canal ayant le plus faible taux d’occupation par les objets aux alentours.

transmissions, un taux de succès de 80 %, alors que, sur la même durée, avec une sélection aléatoire, le taux de succès est de 74 %. En d’autres termes, on a une réduction de 23 % du taux de collisions.

En continuant d’apprendre, l’algorithme UCB aurait permis à l’aggrégateur de transmettre de plus en plus de paquets dans le meilleur canal et ainsi d’augmenter encore son taux de succès.

6.8 Évaluation avec un grand nombre d’objets dynamiques

Nous venons de montrer que les algorithmes de bandits multibras permettent d’augmenter le taux de transmissions réussies par les aggrégateurs du réseau électrique, et, par conséquent, de réduire la latence des communications l’*AMI Backhaul*.

Cependant, dans toutes les simulations et expériences réalisées, nous avons soit considéré un seul aggrégateur (Section 6.7), soit considéré des scénarios dans lesquels le trafic généré par les aggrégateurs est faible devant le trafic interférant (Section 6.6). Dans aucun des scénarios précédemment considérés, nous n’avons observé ce qui se passe si beaucoup d’objets (aggrégateurs) utilisent des algorithmes de bandits.

En effet, il a été prouvé dans la littérature [110, 126, 127] que les performances des algorithmes de bandits multibras, présentés dans la Section 6.5, sont d’ordre optimal lorsque la probabilité de collision dans chaque canal suit une loi de Bernoulli. Dans nos expériences précédentes, la majorité du trafic était généré par les objets statiques dont le comportement

était indépendant de celui des agrégateurs. Ainsi l'hypothèse Bernoulli était valide. Au contraire, si le trafic généré par des agrégateurs qui utilisent des algorithmes d'apprentissage est du même ordre de grandeur, ou plus important, que le trafic généré par les objets statiques, le trafic interférant vu par chaque objet dynamique (agrégateurs qui utilise l'apprentissage) ne suit plus une loi de Bernoulli et va même varier au cours du temps. Son comportement n'est plus stochastique.

L'objectif de cette section est d'évaluer les performances des algorithmes de bandits multibras lorsque le nombre d'objets connectés qui utilisent ces algorithmes (le nombre d'agrégateurs) est important. Dans la suite de ce chapitre, nous présentons d'abord le modèle ainsi que quelques stratégies de référence. Ensuite, nous utiliserons des simulations numériques pour comparer les performances des algorithmes de bandits avec celles-ci.

6.8.1 Modèle

Nous supposons un réseau ALOHA slotté qui utilise la structure de trame présentée par la Figure 6.5. Dans ce réseau, la bande est divisée en N_c canaux. L'unique station de base du réseau reçoit, décode et acquitte les paquets envoyés par les objets environnants. Deux types d'objets communiquent avec cette station de base :

- On a d'abord S objets statiques qui n'utilisent qu'un seul canal. On a donc S_i objets statiques dans le canal $i \in \llbracket 1; N_c \rrbracket$.
- Ensuite, on a D objets dynamiques (agrégateurs) qui peuvent utiliser tous les canaux et, donc, utiliser des algorithmes d'apprentissage.

Tous les objets du réseau transmettent vers la station de base suivant une loi de Bernoulli de probabilité p . En d'autres termes, chaque objet à une probabilité p de transmettre à un instant donné (slot donné).

6.8.2 Stratégies de référence

Afin d'évaluer les algorithmes de bandits multibras étudiés ici, on compare leurs performances avec celles des stratégies de référence présentées ci-dessous.

Approche aléatoire

Comme nous l'avons fait dans les sections précédentes, nous comparons les performances des algorithmes de bandits avec celles d'une stratégie dans laquelle chaque objet dynamique choisit aléatoirement son canal pour chaque transmission. Pour qu'il y ait un intérêt à utiliser des algorithmes de bandits, il faut que leurs performances soient meilleures que celles de cette solution simple.

Approche optimale

Afin d'étudier l'optimalité des algorithmes de bandits, on doit comparer leurs performances avec celles d'une stratégie optimale. C'est pourquoi, on analyse, dans cette section, la répartition des agrégateurs dans les canaux qui maximise la probabilité de réussir une transmission.

Cette stratégie optimale est une bonne stratégie de référence, cependant, elle est difficile à mettre en place sur un système réel. En effet, pour que cette stratégie optimale puisse être utilisée en pratique, il faudrait que la station de base transmette à chaque objet l'index de son canal. Ceci ajouterait de la signalisation, ce qui n'est pas souhaitable dans l'IdO, fortement contraint en énergie. De plus, la station de base ne peut pas connaître, à priori les conditions de propagation de chaque objet dans chaque canal.

On appelle D_i le nombre d'agrégateurs qui sont insérés dans le canal i . Une fois que ces objets sont fixés dans un canal, la probabilité qu'un objet dynamique réussisse une transmission est égale à la moyenne pondérée des probabilités de succès dans chaque canal :

$$P(\text{sd}) = \frac{1}{D} \sum_{i=1}^{N_c} D_i (1-p)^{D_i-1} (1-p)^{S_i} \quad (6.31)$$

Ainsi, trouver la répartition des agrégateurs qui maximise la probabilité de succès revient à résoudre le problème suivant :

$$\arg \max_{D_1, \dots, D_{N_c}} \sum_{i=1}^{N_c} D_i (1-p)^{S_i+D_i-1}, \quad (6.32a)$$

$$\text{S. à } \sum_{i=1}^{N_c} D_i = D, \quad (6.32b)$$

$$D_i \geq 0 \quad \forall i \in \llbracket 1; N_c \rrbracket. \quad (6.32c)$$

Dans ce problème, les variables D_i sont des variables discrètes. Afin de trouver l'optimum, on va d'abord résoudre le problème en considérant les D_i comme des variables réelles pour ensuite discrétiser le résultat obtenu.

Une analyse de la dérivée seconde par rapport à D_i de l'équation (6.32a) montre que, si $D_i \leq \frac{-2}{\ln(1-p)} \approx \frac{2}{p}$, $\forall D_i \in \llbracket 1; N_c \rrbracket$ le problème défini par (6.32) est concave. On remarque que, pour $D_i = \frac{-2}{\ln(1-p)}$, la probabilité d'avoir une collision avec un autre objet dynamique qui

utilise le même canal est de 87 %. Dans ce cas là, le réseau est sur-utilisé. On pourra donc, dans la suite, supposer que l'hypothèse $D_i \leq \frac{-2}{\ln(1-p)}$ est valide⁵.

En appliquant les conditions de KKT, on trouve la solution de ce problème. On dérive :

$$D_i^*(\lambda) = \left(\frac{1}{\log(1-p)} \left[\mathcal{W} \left(\frac{\lambda e}{(1-p)^{S_i-1}} \right) - 1 \right] \right)^+, \quad (6.33)$$

où, λ vérifie :

$$\sum_{i=1}^{N_c} D_i^*(\lambda) = D. \quad (6.34)$$

Dans l'équation (6.33), $(a)^+ = \max(a, 0)$. Pour calculer D_i^* , on peut utiliser une méthode numérique de recherche de 0 à une seule variable pour résoudre (6.34) et trouver la valeur de λ et ainsi trouver la valeur des D_i^* . Une fois qu'on a trouvé les valeurs réelles de D_i^* , on les arrondit pour obtenir la valeur discrète optimale.

Une approche sous-optimale intuitive

Enfin, comme stratégie de référence, nous utilisons aussi une stratégie sous-optimale efficace. Comme la stratégie optimale présentée ci-dessus, la stratégie présentée dans cette section n'est pas utilisable en pratique. Avec celle-ci, les objets dynamiques (aggrégateurs) sont répartis comme s'ils avaient été insérés les uns après les autres dans le canal le plus libre, c'est-à-dire celui qui contient le moins d'objets (statiques et aggrégateurs). En d'autres termes, avec cette stratégie, les objets dynamiques sont placés comme si on avait rempli le canal qui contient le moins d'objets jusqu'à ce que le nombre d'objets dans celui-ci soit égal à celui dans le deuxième canal le plus libre. Ensuite, les aggrégateurs sont insérés dans ces deux canaux jusqu'à ce que le nombre d'objets dans chacun d'eux soit égal à celui dans le troisième canal le plus libre, et ainsi de suite jusqu'à ce que l'on ait placé tous les aggrégateurs.

6.8.3 Évaluation des performances

On compare, maintenant, les performances des algorithmes de bandits multibras avec celles des stratégies de référence présentées ci-dessus. Pour cela, on considère une station de base qui écoute en continu dans 10 canaux de même largeur. On suppose qu'on a un total de $S + D = 2000$ objets qui communiquent avec la station de base. Parmi ces objets, η % sont des objets dynamiques (aggrégateurs du réseau électrique intelligent) qui peuvent choisir leur

5. La fonction objectif définie par l'équation (6.32a) n'est ni concave, ni quasi-concave sur $\mathbb{R}^+$ [133, 134] c'est pourquoi, nous sommes obligés de restreindre l'intervalle d'étude de la concavité.

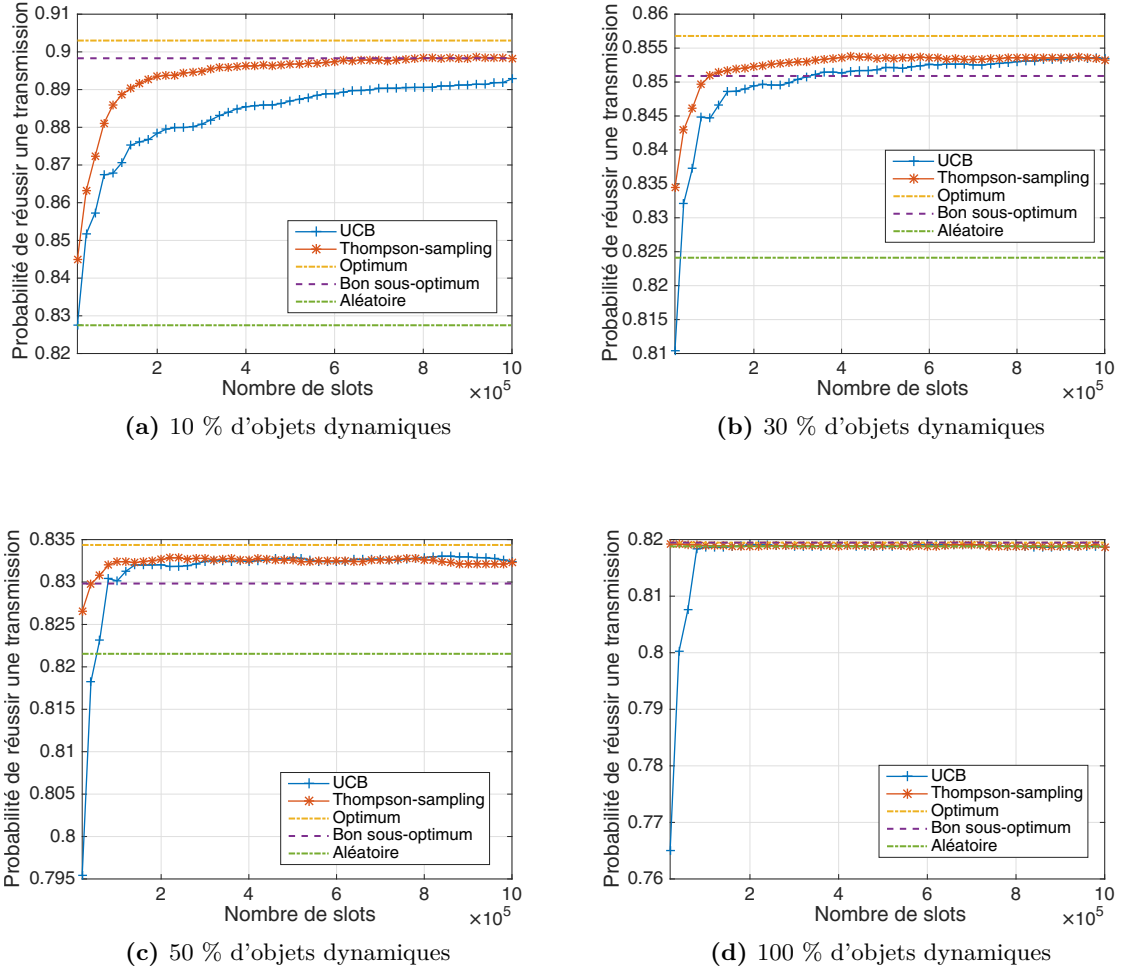


Figure 6.13 – Performances des algorithmes de bandits multibras en fonction de la proportion d'objets dynamiques (aggrégateurs).

canal selon une stratégie définie. Si $\eta = 0$, tous les objets sont statiques et si $\eta = 100$, tous les objets sont dynamiques. On suppose que le nombre d'objets statiques dans chaque canal est égal à $(S_1, \dots, S_{N_c}) = S \times (0.3, 0.2, 0.1, 0.1, 0.05, 0.05, 0.02, 0.08, 0.01, 0.09)$. De plus, la probabilité qu'un objet envoie un paquet à un instant donné est égale à $p = 10^{-3}$.

On compare sur la Figure 6.13 les performances des algorithmes d'apprentissage décrits dans la Section 6.5 (UCB avec $\alpha = 0.5$ et TS) avec les stratégies de référence décrites dans la section précédente pour $\eta = 10\%$, $\eta = 30\%$, $\eta = 50\%$ et $\eta = 100\%$.

On remarque que, quelle que soit la proportion d'objets statiques et d'aggrégateurs, les performances des algorithmes d'apprentissage s'améliorent avec le temps jusqu'à s'approcher de l'optimum. Par exemple, comme on peut le voir sur la Figure 6.13a, lorsque 10 % des

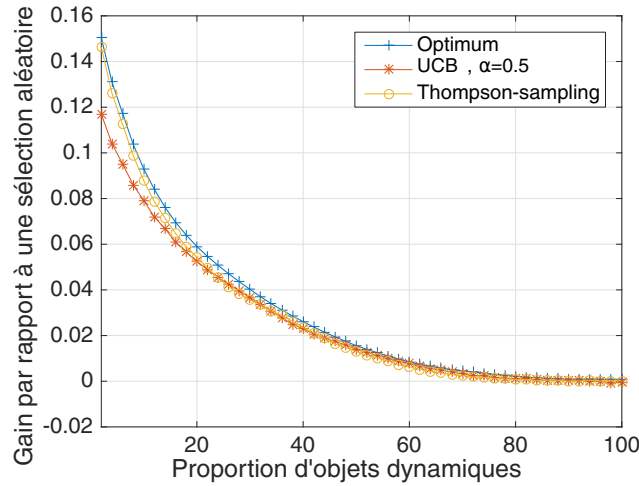


Figure 6.14 – Gain apporté par l'apprentissage comparé à une stratégie avec laquelle le canal est choisi aléatoirement. Ces résultats sont obtenus après 10^6 slots. Soit en moyenne après que chaque objet ait transmis, en moyenne, 1000 paquets.

objets sont dynamiques, on a un taux de succès qui atteint 90 % avec l'algorithme TS et 89% avec l'algorithme UCB alors qu'il n'est que de 82 % avec la solution aléatoire.

On peut voir sur la Figure 6.13 que plus la proportion d'objets dynamiques est importante, plus l'approche aléatoire est proche de l'optimum. En effet, lorsque le nombre d'objets statiques est faible, on a peu de différence entre les canaux. Dans ce cas là, l'apprentissage apporte peu de gain.

Afin de mieux comparer les performances des algorithmes d'apprentissages proposés dans ce chapitre, on évalue sur la Figure 6.14, le gain, en probabilité de transmission réussie, qu'apporte l'apprentissage. Les résultats sont obtenus après 10^6 slots temporel soit, en moyenne, après que chaque objet dynamique ait transmis 1000 paquets.

On peut voir sur la Figure 6.14 que, plus on a d'objets statiques, plus on a d'intérêt à utiliser de l'apprentissage. En effet, si tous les canaux sont identiques, on n'a aucun intérêt à en sélectionner un et plus on a de différence entre les canaux, plus le choix du meilleur canal apporte du gain.

On remarque aussi, sur la Figure 6.14, que quel que soit la proportion d'objets statiques, l'utilisation des algorithmes TS et UCB est toujours quasi-optimale. Ce résultat est logique lorsque le nombre d'objets dynamiques est faible. En effet, dans ce cas là, le trafic interférant observé par chaque objet dynamique peut être approximé par une loi de Bernoulli. Dans ce cas là, les résultats obtenus par simulation sont corroborés par des analyses théoriques [127, 126]. La quasi-optimalité des algorithmes UCB et TS est plus surprenante, en effet, aucune preuve théorique n'appuie ces résultats de simulation. Après avoir mené une analyse simulatoire plus poussée, les auteurs de [135] se sont rendus compte qu'il existait une probabilité très

faible que les algorithmes UCB et TS ait de très mauvaises performances lorsque le nombre d'objets qui les utilise augmente.

Nous venons de voir que l'apprentissage est toujours meilleur qu'une sélection aléatoire du canal. Par ailleurs, le coût des algorithmes d'apprentissages étant quasiment nul (très faible), on en déduit qu'il est toujours avantageux d'utiliser des algorithmes d'apprentissage pour la sélection des canaux. Jamais ils ne dégraderont les performances des communications.

6.9 Conclusion

Dans ce chapitre, nous avons utilisé des algorithmes de bandits multibras pour améliorer la qualité des communications du réseau électrique intelligent. En particulier lorsque les communications de l'*AMI Backhaul* se font via un LPWAN, ces algorithmes peuvent être utilisés pour la sélection fréquentielle. Nous avons réalisé des simulations numériques et mis en place un démonstrateur afin d'éprouver ces algorithmes. Grâce aux algorithmes de bandits, les agrégateurs du réseau électrique peuvent éviter les canaux engorgés et ainsi améliorer leur probabilité de réussir une transmission. Ils réduisent ainsi la latence des communications. Nous avons montré par simulation que ces résultats restent valables lorsque le nombre d'objets qui utilisent ces algorithmes augmente.

Le gain en latence obtenu grâce aux algorithmes proposés dans ce chapitre permet de mieux gérer le réseau électrique intelligent. Grâce à cette meilleure gestion du réseau, les consommateurs (et en particulier le réseau mobile) vont pouvoir être associés à des sources d'énergies renouvelables sans que ce soit préjudiciable pour le réseau électrique. On réduit ainsi l'empreinte carbone de ces consommateurs.

Conclusion et Perspectives

Conclusion

L'objectif de ce travail de thèse était de réduire l'empreinte carbone des réseaux de communication mobile. Pour cela, il faut réduire l'empreinte carbone des stations de base du réseaux qui sont de grosses consommatrices d'électricité. Pour ce faire, nous avons deux possibilités. Une première piste consiste à réduire la quantité d'énergie consommée par ces stations de base. Un second axe de recherche consiste à améliorer les réseaux de communication utilisés pour la gestion du réseau électrique. Ceci permet de mieux intégrer des sources d'énergie renouvelables dans la production et en particulier d'associer les stations de base à des sources d'énergie renouvelables (solaires ou éoliennes). Dans nos travaux, nous avons exploré ces deux pistes séparément.

Résumé des contributions

Nous avons d'abord proposé des solutions pour réduire la consommation d'énergie du réseau mobile. Pour cela, nous nous sommes intéressés au problème d'allocation de puissance lorsque de la transmission discontinue est utilisée. Cette solution étant standardisée dans le récent standard 5G-NR [58]. Nous avons résolu ce problème de manière optimale sous-diverses hypothèses. Nous avons d'abord supposé qu'un ensemble d'utilisateurs était servi par une station de base et que le canal de chacun d'eux était plat. Sous cette hypothèse, nous avons proposé un algorithme pour l'allocation de puissance en TDMA et en OFDMA avec de la transmission discontinue. Une fois ces problèmes résolus, nous avons amélioré notre modèle en considérant le cas où les utilisateurs ont un canal à évanouissement. Nous avons, alors, proposé un algorithme pour résoudre le problème d'allocation de ressources et de puissance en OFDMA avec un canal à évanouissement.

Ensuite, nous avons analysé les performances des algorithmes proposés en TDMA sous l'hypothèse d'un canal plat, dans un réseau composé de plusieurs cellules. Nous avons comparé les performances de notre solution avec celles d'une solution optimale irréalisable. Cette

comparaison nous a permis de montrer que notre solution a de bonnes performances qui peuvent être améliorées en réduisant la puissance d'émission maximale de la station de base.

Après avoir proposé des solutions pour réduire la consommation d'énergie des stations de base, nous nous sommes intéressés aux communications du réseau électrique intelligent. Améliorer ces communications permet de mieux gérer le réseau électrique et ainsi mieux intégrer des sources d'énergie renouvelables et de pouvoir, ainsi, les associer aux stations de base du réseau mobile. Nous avons d'abord utilisé l'architecture HDCRAM pour présenter de manière unifiée les mécanismes de gestion du réseau électrique. Cette représentation unifiée nous a permis, entre autres, d'identifier les standards de communication qui peuvent être utilisés pour cette gestion. Nous avons ainsi pu voir que les réseaux LPWAN sont parfaitement dimensionnés pour l'*AMI Backhaul*.

Nous avons ensuite montré que la probabilité de réussir une transmission (ou réciproquement de subir une collision) et la latence sont deux métriques importantes pour le dimensionnement de ces réseaux. Nous avons, proposé une expression analytique pour la probabilité de collision dans un réseau LPWAN ALOHA non-slotté qui utilise un mécanisme d'acquiescement similaire à celui proposé dans le standard LoRaWAN. Ensuite, nous avons dérivé l'expression de la latence dans un tel réseau comme une fonction de la probabilité de collision.

Enfin, nous avons montré que les algorithmes de bandits multibras peuvent être utilisés pour améliorer ces métriques lorsque le trafic interférant est inégalement réparti dans les canaux disponibles, cette hypothèse a de fortes chances de se vérifier dans les réseaux IdO tels que LoRaWAN vu que ceux-ci opèrent dans des bandes libres. En particulier, nous avons montré par des simulations numériques que ces algorithmes restent efficaces lorsqu'ils sont utilisés par plusieurs objets en même temps. Nous avons, finalement, réalisé un démonstrateur sur des plateformes de radio-logicielle pour valider les résultats obtenus via des simulations numériques. Ce démonstrateur reste, à notre connaissance, la seule implémentation temps-réel d'apprentissage pour l'IdO.

Perspectives

Les travaux présentés dans cette thèse ouvrent, bien entendu, plusieurs perspectives qui sont présentés dans cette section. Nous présenterons plus en détails les perspectives qui ont le plus attiré notre attention. Nous avons commencé à travailler sur certaines d'entre elles.

Réduction de la consommation d'énergie des stations de base

Amélioration du modèle : les algorithmes d'allocation de ressources et de puissance avec transmission discontinue présentés dans cette thèse ont été conçus en prenant en compte certaines hypothèses (connaissance parfaite du canal par la station de base, canal plat, absence de mobilité). Il est donc possible d'améliorer ces algorithmes en se passant de certaines hypothèses. Par exemple, nous aurions pu traiter ce problème d'allocation de ressources et de puissance avec de la transmission discontinue avec une connaissance imparfaite du canal.

Allocation de puissance adaptée aux réseaux denses : dans le Chapitre 3, nous avons choisi d'étendre l'utilisation de la solution proposée, pour une seule cellule, à un réseau dense. Il serait donc intéressant de proposer une solution pour l'allocation de puissance qui n'est pas une simple extension de la solution proposée dans un réseau à une seule cellule.

Allocation de ressources dans des réseaux denses : dans le Chapitre 3, nous avons étudié les performances de notre algorithme d'allocation de puissance dans un réseau dense. Dans ce chapitre, nous n'avons pas considéré le problème d'allocation de ressources (échelonnement des utilisateurs) et considéré les interférences générées par la puissance d'émission moyenne des stations de base environnantes. Il reste donc à proposer des algorithmes d'allocation de ressources et de puissance avec de la transmission discontinue dans des réseaux denses.

MIMO : Les systèmes de communication MIMO utilisent un grand nombre d'antennes pour augmenter le débit des utilisateurs. L'utilisation de ces antennes augmente la consommation statique des stations de base, il est donc intéressant d'en éteindre certaines pour diminuer la consommation des stations de base MIMO. Le pendant MIMO de la transmission discontinue est l'*Antenna Selection* [136]. Nous avons commencé à travailler sur ce sujet et publié un article de conférence à ce propos [137]. Le travail publié est un travail préliminaire et peut être largement perfectionné. Par exemple, en étudiant ce problème dans un scénario dans lequel du SDMA est utilisé.

NOMA : Le NOMA est une stratégie pour l'accès multiple qui a été proposée récemment et qui concentre, donc, beaucoup d'attention. Nous avons commencé à travailler sur ce sujet et en particulier sur les problématiques d'allocation de ressources et de puissance en NOMA avec du Cell DTX. Ce problème est intéressant, mais, en réalité, les premiers résultats que nous avons obtenus montrent que sa résolution est assez similaire à la résolution du même problème en OFDMA.

Amélioration des communications du réseau électrique intelligent

Amélioration du modèle : lors des dérivations des probabilités de collision, nous avons supposé que dès que deux paquets se chevauchent dans un canal, il y a collision. Ce qui est

une hypothèse simplificatrice car pessimiste. De plus, nous avons considéré que l'ensemble des pertes paquets été causées par des collisions et n'avons pas considéré les pertes liées aux conditions de propagation. Il serait intéressant de dériver une probabilité de perte paquet plus réaliste. Une des difficulté que nous avons rencontrée dans l'analyse des réseaux d'objets connectés (réseaux ALOHA) a été de trouver une modélisation réaliste mais qui n'est pas trop complexe et permet de comprendre et analyser l'impact des différents paramètres.

Optimisation des réseaux LoRaWAN : les réseaux LPWAN sont aujourd'hui en plein essor et peuvent être utilisés pour de nombreuses applications industrielles. L'un des standards les plus en vogue est le standard LoRaWAN. Il serait intéressant d'optimiser les performances d'un tel réseau au delà de ce que nous avons proposé dans cette thèse. Par exemple, dans cette thèse, nous avons utilisé l'acquiescement qui est disponible dans ce standard pour utiliser des algorithmes d'apprentissage. L'intérêt de cet acquiescement, tel qu'il est proposé dans ce standard, est aujourd'hui questionné [138]. Il serait donc intéressant de voir comment le faire évoluer pour améliorer les communications dans des réseaux LoRaWAN.

Algorithmes de bandits multibras dans des environnements non-stochastiques : pour que les algorithmes de bandits multibras puissent être utilisés en pratique, il faut d'abord qu'ils puissent s'adapter aux variations potentielles du trafic interférant. Le développement et l'évaluation d'algorithmes de bandits qui ont de bonnes performances dans des environnements dont les statistiques varient dans le temps est la priorité pour le développement d'algorithmes de bandits dans des réseaux d'objets connectés.

Bilan personnel

L'objectif de cette thèse et plus largement du réseau SOGREEN était de limiter l'empreinte carbone des réseaux mobiles et par extension de l'activité humaine en améliorant la gestion de l'électricité des stations de base. Cet objectif était bien entendu trop ambitieux pour pouvoir être atteint en 3 ans et résumé en moins de 200 pages. Pour l'atteindre, il est nécessaire d'optimiser le fonctionnement de tous les systèmes conçus et utilisés par l'homme. C'est dans ce cadre que s'inscrit cette thèse. En effet, il est évident que l'apport des algorithmes présentés dans ce manuscrit n'est pas suffisant pour réduire significativement l'impact de l'activité humaine sur l'environnement. Cependant, on peut obtenir des résultats significatifs si des démarches similaires à celles proposées dans cette thèse sont appliquées pour optimiser un grand nombre de systèmes.

Dans cette thèse, en plus d'avoir proposé des algorithmes qui permettent de réduire la consommation du réseau mobile, nous avons proposé des algorithmes utiles pour mieux gérer le réseau électrique. Le réseau électrique intelligent et, par extension, les réseaux de communication qu'il utilise sont un élément majeur de la transition écologique actuelle. La

transformation du réseau électrique en réseau électrique intelligent a un impact bénéfique sur l'environnement bien plus significatif que la réduction de la consommation d'énergie des réseaux mobiles. C'est pourquoi, l'amélioration des communications du réseau électrique intelligent doit devenir un axe majeur dans la recherche en éco-radio.

Par ailleurs, la problématique de réduction de l'empreinte carbone des réseaux mobiles est un sujet vaste qui, par conséquent, m'a permis de découvrir divers domaines (réseaux mobiles, objets connectés, réseau électrique) et d'acquérir des compétences sur divers sujets (probabilité, optimisation, apprentissage par renforcement).

Enfin, comme d'autres travaux dans le domaine des télécommunications, certains algorithmes proposés dans cette thèse ne sont pas utilisables en l'état sur des systèmes réels. En effet, les hypothèses formulées dans cette thèse pour poser le problème de minimisation de la puissance consommée ne sont pas suffisamment réalistes. Il est évidemment difficile d'évaluer à priori la marche qu'il reste à franchir pour pouvoir implémenter ces algorithmes. Dans cette thèse, nous nous sommes efforcé lorsque c'était possible d'implémenter les algorithmes que nous proposons afin de les confronter à une certaine réalité. Même si elle n'est pas aussi complète qu'un système déployé en vraie grandeur, le pas le plus important a été franchi en passant des études théoriques et par simulations aux expérimentations en condition de laboratoire. C'est ce que nous avons fait avec les algorithmes de bandits multibras.

Annexe A

Quelques preuves relatives au chapitre 3

A.1 Dérivation de l'équation (3.37)

Dans cette preuve, sans aucune perte de généralité, on suppose que $S_1 = \llbracket 1; N_u \rrbracket$ et que $S_2 = \llbracket N_1 + 1; N_u \rrbracket$. Pour exprimer les x_k uniquement comme une fonction des temps de service, on résout le système d'équations linéaires défini par (3.34). Ce système peut s'écrire sous forme matricielle :

$$\mathbf{M}\mathbf{X} = \boldsymbol{\beta}. \quad (\text{A.1})$$

Dans cette équation matricielle, $\boldsymbol{\beta} = (\beta_1, \dots, \beta_{N_u})^T$, $\mathbf{X} = (x_1, \dots, x_{N_u})^T$ et la matrice $\mathbf{M}$ peut s'écrire comme la matrice par blocs suivante :

$$\mathbf{M} = \left(\begin{array}{c|c} \mathbf{I}_{N_1} & \mathbf{\Gamma}_1 \\ \hline \mathbf{\Gamma}_2 & \mathbf{I}_{N_2} \end{array} \right). \quad (\text{A.2})$$

Dans cette équation, N_1 et N_2 désignent respectivement le nombre d'utilisateurs servis par la station de base 1 et 2. $\mathbf{I}_{N_1}$ et $\mathbf{I}_{N_2}$ sont les matrices identités de taille N_1 et N_2 . De plus, d'après l'équation (3.34), les matrices $\mathbf{\Gamma}_1$ et $\mathbf{\Gamma}_2$ ont pour expression :

$$\mathbf{\Gamma}_1 = \begin{pmatrix} -\alpha_1 & \cdots & -\alpha_1 \\ \vdots & & \vdots \\ -\alpha_{N_1} & \cdots & -\alpha_{N_1} \end{pmatrix}, \quad (\text{A.3})$$

et,

$$\mathbf{\Gamma}_2 = \begin{pmatrix} -\alpha_{N_1+1} & \cdots & -\alpha_{N_1+1} \\ \vdots & & \vdots \\ -\alpha_{N_u} & \cdots & -\alpha_{N_u} \end{pmatrix}. \quad (\text{A.4})$$

On rappelle que la définition de α_k est donnée par l'équation (3.35). On remarque que les matrices $\mathbf{\Gamma}_1 \in M_{N_1 \times N_2}(\mathbb{R})$ et $\mathbf{\Gamma}_2 \in M_{N_2 \times N_1}(\mathbb{R})$ ne sont pas des matrices carrées. Dans la suite, on sera amené à utiliser les matrices $\mathbf{\Gamma}'_1 \in M_{N_1}(\mathbb{R})$ et $\mathbf{\Gamma}'_2 \in M_{N_2}(\mathbb{R})$ qui sont carrées et dont les colonnes sont respectivement égales à celles de $\mathbf{\Gamma}_1$ et $\mathbf{\Gamma}_2$. Pour inverser la matrice $\mathbf{M}$, on utilise la formule d'inversion d'une matrice par bloc [139] :

$$\begin{pmatrix} \mathbf{A} & \mathbf{B} \\ \mathbf{C} & \mathbf{D} \end{pmatrix}^{-1} = \begin{pmatrix} \mathbf{A}^{-1} + \mathbf{A}^{-1}\mathbf{B}(\mathbf{D} - \mathbf{C}\mathbf{A}^{-1}\mathbf{B})^{-1}\mathbf{C}\mathbf{A}^{-1} & -\mathbf{A}^{-1}\mathbf{B}(\mathbf{D} - \mathbf{C}\mathbf{A}^{-1}\mathbf{B})^{-1} \\ -(\mathbf{D} - \mathbf{C}\mathbf{A}^{-1}\mathbf{B})^{-1}\mathbf{C}\mathbf{A}^{-1} & (\mathbf{D} - \mathbf{C}\mathbf{A}^{-1}\mathbf{B})^{-1} \end{pmatrix}. \quad (\text{A.5})$$

L'équation (A.5) nous permet de déduire que :

$$\mathbf{M}^{-1} = \begin{pmatrix} \mathbf{I}_{N_1} + \mathbf{\Gamma}_1(\mathbf{I}_{N_2} - \mathbf{\Gamma}_2\mathbf{\Gamma}_1)^{-1}\mathbf{\Gamma}_2 & -\mathbf{\Gamma}_1(\mathbf{I}_{N_2} - \mathbf{\Gamma}_2\mathbf{\Gamma}_1)^{-1} \\ -(\mathbf{I}_{N_2} - \mathbf{\Gamma}_2\mathbf{\Gamma}_1)^{-1}\mathbf{\Gamma}_2 & (\mathbf{I}_{N_2} - \mathbf{\Gamma}_2\mathbf{\Gamma}_1)^{-1} \end{pmatrix}. \quad (\text{A.6})$$

On doit maintenant calculer l'inverse de $\mathbf{I}_{N_2} - \mathbf{\Gamma}_2\mathbf{\Gamma}_1$. Les expressions de $\mathbf{\Gamma}_1$ et $\mathbf{\Gamma}_2$ nous permettent d'écrire :

$$\mathbf{\Gamma}_2\mathbf{\Gamma}_1 = \text{Tr}(\mathbf{\Gamma}'_1)\mathbf{\Gamma}'_2. \quad (\text{A.7})$$

Dans cette équation $\text{Tr}(\cdot)$ désigne l'opérateur trace. On en déduit :

$$(\mathbf{\Gamma}_2\mathbf{\Gamma}_1)^2 = \text{Tr}(\mathbf{\Gamma}_2\mathbf{\Gamma}_1)\mathbf{\Gamma}_2\mathbf{\Gamma}_1. \quad (\text{A.8})$$

Cette dernière expression nous permet de montrer que :

$$\begin{aligned} (\mathbf{I}_{N_2} - \mathbf{\Gamma}_2\mathbf{\Gamma}_1) \left(\mathbf{I}_{N_2} + \frac{1}{1 - \text{Tr}(\mathbf{\Gamma}_2\mathbf{\Gamma}_1)} \mathbf{\Gamma}_2\mathbf{\Gamma}_1 \right) = \\ \left(\mathbf{I}_{N_2} + \frac{1}{1 - \text{Tr}(\mathbf{\Gamma}_2\mathbf{\Gamma}_1)} \mathbf{\Gamma}_2\mathbf{\Gamma}_1 \right) (\mathbf{I}_{N_2} - \mathbf{\Gamma}_2\mathbf{\Gamma}_1) = \mathbf{I}_{N_2}. \end{aligned} \quad (\text{A.9})$$

Par conséquent :

$$(\mathbf{I}_{N_2} - \mathbf{\Gamma}_2\mathbf{\Gamma}_1)^{-1} = \mathbf{I}_{N_2} + \frac{\text{Tr}(\mathbf{\Gamma}'_1)}{1 - \text{Tr}(\mathbf{\Gamma}'_2) \text{Tr}(\mathbf{\Gamma}'_1)} \mathbf{\Gamma}'_2. \quad (\text{A.10})$$

En insérant cette expression dans l'équation (A.6), on déduit :

$$\mathbf{M}^{-1} = \begin{pmatrix} \mathbf{I}_{N_1} + \frac{\text{Tr}(\mathbf{\Gamma}'_2)}{1 - \text{Tr}(\mathbf{\Gamma}'_2) \text{Tr}(\mathbf{\Gamma}'_1)} \mathbf{\Gamma}'_1 & \frac{-1}{1 - \text{Tr}(\mathbf{\Gamma}'_2) \text{Tr}(\mathbf{\Gamma}'_1)} \mathbf{\Gamma}_1 \\ \frac{-1}{1 - \text{Tr}(\mathbf{\Gamma}'_2) \text{Tr}(\mathbf{\Gamma}'_1)} \mathbf{\Gamma}_2 & \mathbf{I}_{N_2} + \frac{\text{Tr}(\mathbf{\Gamma}'_1)}{1 - \text{Tr}(\mathbf{\Gamma}'_2) \text{Tr}(\mathbf{\Gamma}'_1)} \mathbf{\Gamma}'_2 \end{pmatrix}. \quad (\text{A.11})$$

On peut finalement résoudre le problème d'équations linéaires introduit équation (3.34) :

$$\mathbf{X} = \mathbf{M}^{-1}\boldsymbol{\beta}. \quad (\text{A.12})$$

Pour $k \in \llbracket 1; N_u \rrbracket$, on a finalement :

$$x_k = \beta_k + \alpha_k \left[\frac{\sum_{i \in S_n} \sum_{j \in S_m} \alpha_j \beta_i + \sum_{j \in S_m} \beta_j}{1 - \sum_{i \in S_n} \sum_{j \in S_m} \alpha_i \alpha_j} \right]. \quad (\text{A.13})$$

Dans cette équation $m = n \bmod (2) + 1$. Cette équation pouvant être réécrite :

$$x_k = \frac{N\mu_{nk}}{g_{nk}} \left(2^{\frac{C_k}{B\mu_{nk}}} - 1 \right) \times \left[1 + \frac{g_{mk} \sum_{j \in S_m} \frac{\mu_{mj}}{g_{mj}} \left(2^{\frac{C_j}{B\mu_{mj}}} - 1 \right) \left(1 + g_{nj} \sum_{i \in S_n} \frac{\mu_{ni}}{g_{ni}} \left(2^{\frac{C_i}{B\mu_{ni}}} - 1 \right) \right)}{1 - \sum_{i \in S_n} \sum_{j \in S_m} \frac{\mu_{ni} g_{mi}}{g_{ni}} \left(2^{\frac{C_i}{B\mu_{ni}}} - 1 \right) \frac{\mu_{mj} g_{nj}}{g_{mj}} \left(2^{\frac{C_j}{B\mu_{mj}}} - 1 \right)} \right]. \quad (\text{A.14})$$

Nous venons d'exprimer les variables x_k uniquement en fonction des, temps de service des utilisateurs.

A.2 Convexité du problème défini par l'équation (3.40)

Avant d'entrer dans le coeur de la démonstration, il est important de noter que les contraintes données par les équations (3.40c) et (3.40d) sont des contraintes d'inégalité stricte. D'après [40], un problème convexe n'a que des inégalités larges. Le problème ne peut donc pas être rigoureusement convexe. Vu qu'on ne s'intéresse qu'à l'optimum de ce problème, on peut introduire un $\epsilon \ll 1$ qui change ces inégalités strictes en inégalités large. On a alors les inégalités suivantes :

$$\sum_{i \in S_1} \sum_{j \in S_2} \alpha_i \alpha_j \leq 1 - \epsilon, \quad (\text{A.15})$$

$$\mu_{nk} \geq \epsilon, \quad \forall k \in \llbracket 1; N_u \rrbracket. \quad (\text{A.16})$$

Si ϵ est choisi suffisamment petit, son introduction ne change rien à la convexité du problème.

On entre maintenant dans le coeur de la démonstration. Dans cette preuve, sans aucune perte de généralité, on suppose que $S_1 = \llbracket 1; N_1 \rrbracket$ et $S_2 = \llbracket N_1 + 1; N_u \rrbracket$. Pour prouver la convexité du problème défini par l'équation (3.40), on prouve la convexité de la fonction objectif et celle de l'ensemble délimité par les contraintes. Pour cela, il suffit de montrer que

les x_k sont des fonctions convexe de μ_{nl} pour $\mu_{nl} \in]0; 1]$, $\forall l \in \llbracket 1; N_u \rrbracket$. Notre preuve s'appuie sur le lemme suivant :

Lemme A.1. *la fonction $f(x) = \ln(x(e^{1/x} - 1))$ est convexe pour tout $x \in \mathbb{R}_+^*$.*

Démonstration. Pour $x > 0$, la dérivée de f est égale à :

$$f'(x) = -\frac{1}{x} \left(\frac{1}{x(1 - e^{-\frac{1}{x}})} - 1 \right) = -\frac{1}{x} g(x). \quad (\text{A.17})$$

Pour prouver la convexité de f , on va montrer que sa dérivée est croissante. Pour cela, on s'intéresse à la fonction g . Tout d'abord, l'inégalité :

$$e^x > 1 + x, \quad \forall x \in \mathbb{R}^*, \quad (\text{A.18})$$

prouve que $g(x) > 0$, $\forall x > 0$. De plus, la dérivée de $x(1 - e^{-\frac{1}{x}})$ est égale à :

$$\left(x(1 - e^{-\frac{1}{x}}) \right)' = 1 - e^{-\frac{1}{x}} - \frac{e^{-\frac{1}{x}}}{x}. \quad (\text{A.19})$$

L'inégalité (A.18) nous montre que cette dérivée est positive. Donc g est décroissante car c'est l'inverse d'une fonction croissante. Ensuite, $\frac{g(x)}{x}$ est décroissante sur $\mathbb{R}_+^*$ comme produit de deux fonctions décroissantes à valeur positive. Enfin, f' est croissante car c'est l'opposée d'une fonction croissante. Finalement, f est convexe.

□

Le Lemme A.1 permet de prouver les fonctions α_k et β_k sont des fonctions log-convexes pour tout k . Par conséquent, chaque produit de ces fonctions est convexe [40]. De plus, comme $\sum_{i=1}^{N_1} \sum_{j=N_1+1}^{N_u} \alpha_i \alpha_j \in]0; 1[$, on peut utiliser le développement en série entière suivant :

$$\frac{1}{1-x} = \sum_{n=0}^{+\infty} x^n, \quad \forall x \in [0; 1[. \quad (\text{A.20})$$

Dans la suite de cette preuve, pour améliorer sa lisibilité, on notera :

$$\sum_{l=1}^{N_1 N_2} \gamma_l = \sum_{i=1}^{N_1} \sum_{j=N_1+1}^{N_u} \alpha_i \alpha_j. \quad (\text{A.21})$$

Dans cette équation, $\gamma_l = \alpha_{(l-1)(\text{mod } N_1)+1} \times \alpha_{N_1 + \lfloor \frac{l-1}{N_1} \rfloor + 1}$, et $\lfloor \cdot \rfloor$ désigne la fonction partie entière. Chacune des fonctions γ_l est log-convexe comme produit de fonctions log-convexes. En appliquant le développement en série entière de l'équation (A.20), x_k peut s'écrire :

$$x_k = \beta_k + \alpha_k \left[\sum_{i=1}^{N_1} \sum_{j=N_1+1}^{N_u} \alpha_j \beta_i + \sum_{j=N_1+1}^{N_u} \beta_j \right] \left[\sum_{n=0}^{+\infty} \left(\sum_{l=1}^{N_1 N_2} \gamma_l \right)^n \right]. \quad (\text{A.22})$$

Les coefficients multinomiaux [140] permettent de développer cette expression, on obtient alors :

$$x_k = \beta_k + \sum_{n=0}^{+\infty} \left[\left(\sum_{i=1}^{N_1} \sum_{j=N_1+1}^{N_u} \alpha_j \beta_i \alpha_k + \sum_{j=N_1+1}^{N_u} \alpha_k \beta_j \right) \left(\sum_{\sum n_l = n} \binom{n}{n_1, \dots, n_{N_1 N_2}} \prod_{l=1}^{N_1 N_2} \gamma_l^{n_l} \right) \right]. \quad (\text{A.23})$$

Enfin, chaque produit $\prod_{l=1}^{N_1 N_2} \gamma_l^{n_l}$ est log-convexe comme produit de fonctions log-convexes. Ainsi, en développant les termes dans la somme infinie, on exprime les termes de cette somme comme des sommes de produits de fonctions log-convexes. Ces produits sont des fonctions log-convexes. Finalement, une somme infinie de fonctions convexes est convexe [40]. Les x_k sont donc convexes. Le problème étudié est donc un problème convexe.

Annexe B

Quelques preuves relatives au Chapitre 4

B.1 Non-convexité du problème définit équation (4.5)

On a deux façons de prouver que le problème définit par l'équation (4.5) n'est pas convexe. On peut soit choisir de prouver que la fonction :

$$\left(\mu_a, P_{TX}^{k,1}, \dots, P_{TX}^{k,N_k}\right) \mapsto C_k - B_c \mu_a \sum_{n=1}^N \log_2 \left(1 + \frac{g_{k,n} P_{TX}^{k,n}}{N}\right), \quad (\text{B.1})$$

n'est pas convexe, soit prouver que la fonction objectif donnée par l'équation (4.5a) n'est pas convexe. Dans cette annexe, nous prouvons que la fonction définie par l'équation (B.1) n'est pas convexe. Pour cela, on note :

$$\alpha_n = \frac{B_c}{N \ln(2)} \frac{g_{k,n}}{1 + \frac{g_{k,n} P_{TX}^{k,n}}{N}} \geq 0 \quad (\text{B.2})$$

$$\beta_n = \frac{g_{k,n}^2 B_c \mu_a}{N^2 \ln(2)} \frac{1}{\left(1 + \frac{g_{k,n} P_{TX}^{k,n}}{N}\right)^2} \geq 0 \quad (\text{B.3})$$

Pour $k \in \llbracket 1; N_u \rrbracket$, la matrice Hessienne de la fonction étudiée s'écrit :

$$H_k = \begin{pmatrix} 0 & -\alpha_1 & \cdots & -\alpha_{N_k} \\ -\alpha_1 & \beta_1 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ -\alpha_{N_k} & 0 & \cdots & \beta_{N_k} \end{pmatrix} \quad (\text{B.4})$$

On déduit assez facilement, que pour des valeurs de μ_a et $P_{TX}^{k,n}$ positives, la trace de cette matrice est positive. Ceci nous permet de dire qu'au moins une des valeurs propres de cette matrice est positive.

On cherche maintenant à calculer le déterminant de cette Hessienne. En développant suivant la première ligne, on obtient :

$$\det(H_k) = \alpha_1 \begin{vmatrix} -\alpha_1 & 0 & \cdots & 0 \\ \vdots & \beta_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ -\alpha_{N_k} & 0 & \cdots & \beta_{N_k} \end{vmatrix} + \cdots + (-1)^{N_k+1} \alpha_{N_k} \begin{vmatrix} -\alpha_1 & \beta_1 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ \vdots & 0 & \cdots & \beta_{N_k} \\ -\alpha_{N_k} & 0 & \cdots & 0 \end{vmatrix} \quad (\text{B.5})$$

On vient d'exprimer le déterminant de la Hessienne comme la somme de N_k déterminants. Pour les calculer, on va effectuer $n-1$ opérations élémentaires sur le $n^{\text{ième}}$ de ces déterminants. Afin de faire remonter successivement la $n^{\text{ième}}$ ligne jusqu'à la première position :

$$\det(H_k) = \alpha_1 \begin{vmatrix} -\alpha_1 & 0 & \cdots & 0 \\ -\alpha_2 & \beta_2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ -\alpha_{N_k} & 0 & \cdots & \beta_{N_k} \end{vmatrix} + \cdots + \alpha_{N_k} \begin{vmatrix} -\alpha_{N_k} & 0 & \cdots & 0 \\ -\alpha_1 & \beta_1 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ -\alpha_{N_k-1} & 0 & \cdots & \beta_{N_k-1} \end{vmatrix} \quad (\text{B.6})$$

Les déterminants obtenus sont triangulaires inférieurs, on peut finalement les calculer :

$$\det(H_k) = - \sum_{n=1}^{N_k} \alpha_n^2 \prod_{j=1, j \neq n}^{N_k} \beta_j < 0 \quad (\text{B.7})$$

Donc le produit des valeurs propres de H_k est négatif. Ce qui prouve qu'au moins une valeur propre de H_k est négative. H_k a donc au moins une valeur propre positive et une valeur propre négative. La contrainte étudiée n'est donc pas convexe. Ceci est suffisant pour dire que, sous cette forme, le problème donné par l'équation (4.5a) n'est pas convexe.

B.2 Allocation de puissance optimale en OFDMA avec du Cell DTx

Comme expliqué dans la Section 4.3.2, et d'après les conditions de KKT, pour ce problème, on peut avoir trois cas différents. Dans cette preuve on se limite au cas où ni la contrainte de l'équation (4.10d), ni celle de l'équation (4.10b) ne sont saturées. Dans ce cas là, le Lagrangien du problème s'écrit :

$$\mathcal{L} = P_s + \mu_a(P_0 - P_s) + m_p\mu_a N \sum_{k=1}^{N_u} \sum_{n=1}^{N_k} \frac{1}{g_{k,n}} \left(2^{\frac{C_{k,n}}{B_c\mu_a}} - 1 \right) - \sum_{k=1}^{N_u} \lambda_k \left(\sum_{n=1}^{N_k} C_{k,n} - C_k \right). \quad (\text{B.8})$$

La dérivée de ce Lagrangien para rapport à $C_{k,n}$ est égale à :

$$\frac{\partial \mathcal{L}}{\partial C_{k,n}} = \frac{m_p N \ln(2)}{g_{k,n} B_c} 2^{\frac{C_{k,n}}{B_c\mu_a}} - \lambda_k. \quad (\text{B.9})$$

Au point minimum, cette dérivée est égale à 0. On en déduit :

$$C_{k,n} = B_c\mu_a \log_2 \left(\frac{B_c\lambda_k g_{k,n}}{m_p N \ln(2)} \right). \quad (\text{B.10})$$

Vu que $\sum_{n=1}^{N_k} C_{k,n} = C_k$, on a :

$$B_c\mu_a \log_2 \left(\frac{\lambda_k B_c}{m_p N \ln(2)} \right) = \frac{C_k}{N_k} - \frac{B_c\mu_a}{N_k} \log_2 \left(\prod_{n=1}^{N_k} g_{k,n} \right). \quad (\text{B.11})$$

Par conséquent,

$$C_{k,n} = \frac{C_k}{N_k} + B_c\mu_a \log_2 \left(\frac{g_{k,n}}{\left(\prod_{n=1}^{N_k} g_{k,n} \right)^{\frac{1}{N_k}}} \right). \quad (\text{B.12})$$

On note que l'expression (B.12) ne dépend pas de la façon dont on calcule la valeur de μ_a . On va maintenant exprimer μ_{opt} , la valeur optimale de μ_a comme la solution d'une équation à une seule variable. Pour cela, on écrit la dérivée du Lagrangien par rapport à μ_a :

$$\frac{\partial \mathcal{L}}{\partial \mu_a} = \sum_{k=1}^{N_u} \sum_{n=1}^{N_k} \left[\frac{m_p N}{g_{k,n}} \left(2^{\frac{C_{k,n}}{B_c\mu_a}} - 1 \right) - \frac{m_p N C_{k,n} \ln(2)}{B_c g_{k,n} \mu_a} 2^{\frac{C_{k,n}}{B_c\mu_a}} \right] + (P_0 - P_s). \quad (\text{B.13})$$

En remplaçant $C_{k,n}$ dans l'équation (B.13) par son expression, donnée par l'équation (B.12), l'équation $\frac{\partial \mathcal{L}}{\partial \mu_a} = 0$ nous donne μ_{opt} comme la solution d'une équation à une seule variable :

$$\sum_{k=1}^{N_u} \frac{2^{\frac{C_k}{N_k B_c \mu_{\text{opt}}}}}{\left(\prod_{n=1}^{N_k} g_{k,n}\right)^{\frac{1}{N_k}}} \left[\frac{C_k \ln(2)}{B_c \mu_{\text{opt}}} - N_k \right] - \frac{P_0 - P_s}{m_p N} + \sum_{k=1}^{N_u} \sum_{n=1}^{N_k} \frac{1}{g_{k,n}} = 0. \quad (\text{B.14})$$

Pour trouver l'allocation de puissance optimale, on va devoir résoudre l'équation (B.14) et ensuite calculer la capacité par canal grâce à l'équation (B.12). Une fois que c'est fait, on doit vérifier si les $C_{k,n}$ sont tous positifs. Si ce n'est pas le cas, on va mettre à 0 la capacité dans les blocs de ressource dans lesquels elle est négative et recalculer μ_{opt} ainsi que les valeurs des $C_{k,n}$. On a donc un procédé itératif de type water-filling, comme décrit Section 2.4.2.

Annexe C

Quelques preuves relatives au Chapitre 6

C.1 Calcul des probabilités de collision dans un réseau ALOHA

Dans cette Annexe, nous prouvons les Propositions 6.1 et 6.2. L'objectif de cette annexe est donc de calculer la probabilité d'avoir une transmission sans collision dans un canal. Pour cela, on utilisera les notations du Tableau C.1

Pour le calcul de la probabilité d'avoir une transmission réussie dans le canal j , on va supposer qu'un des objets a envoyé un paquet dans ce canal. On appellera paquet 1 ce paquet et on calcule la probabilité qu'il arrive sans subir de collisions. Comme indiqué dans le Chapitre 6, on a deux cas selon si $T_d \leq T_m$ ou si $T_d \geq T_m$.

Cas 1 : $T_d \leq T_m$

On commence par calculer $P^j(su)$, la probabilité que le paquet 1 soit bien reçu et décodé par la station de base. Afin de simplifier les notations, dans la suite on ne mentionnera pas l'indice du canal et on notera $P(su)$ la probabilité qu'un paquet envoyé sur la voie montante soit bien reçu par la station de base.

Le paquet 1 est transmis avec succès (su) si et seulement si il n'y a pas de collision entre ce paquet et un autre paquet, que celui-ci soit envoyé avant ou après le paquet 1 :

$$P(su) = P(\overline{cb} \cap \overline{ca}) = P(\overline{cb})P(\overline{ca}). \quad (\text{C.1})$$

Tableau C.1 – Évènements dont la probabilité est utilisée dans le calcul des probabilités d'avoir une bonne transmission.

Notation	Évènement
su	<i>Successful Uplink</i> : il n'y pas eu de collision entre le paquet envoyé sur la voie montante et d'autres paquets. Il a été décodé par la station de base.
sd	<i>Successful Downlink</i> : L'objet a bien reçu l'acquittement et l'a décodé.
sa	<i>Successful Acknowledgement</i> : la transmission sur la voie descendante s'est passée sans que l'acquittement n'ait de collision avec un autre paquet.
ca	<i>Collision After</i> : La transmission du paquet sur la voie montante a échoué à cause d'une collision avec un paquet envoyé après le paquet considéré.
cb	<i>Collision Before</i> : La transmission du paquet sur la voie montante a échoué suite à une collision avec un paquet envoyé avant. Celui-ci pouvant être un paquet envoyé sur la voie montante ou sur la voie descendante.
cub	<i>Collision Uplink Before</i> : La transmission du paquet sur la voie montante a échoué en raison d'une collision avec un paquet envoyé par un autre objet avant le paquet considéré.
cd	<i>Collision Downlink</i> : Le paquet envoyé sur la voie montante a eu une collision avec un acquittement.
pss	<i>Packet Successfully Sent</i> : Un paquet a été envoyé avec succès par un objet dans l'intervalle de temps $[-T_d - T_m - T_a; -T_d - T_m]$.
pb	<i>Packet Between</i> : Un paquet est envoyé entre le paquet envoyé sur la voie montante et l'acquittement sans gêner leur transmission.

Dans cette équation, et en accord avec le Tableau C.1, $P(cb)$ et $P(ca)$ sont respectivement la probabilité d'avoir une collision avec un paquet envoyé avant et après le paquet 1. Vu que le trafic sur la voie montante suit une loi de Poisson, ces deux probabilités sont indépendantes.

On commence par calculer $P(cb)$. Pour cela, on utilise la formule des probabilités totales et la formule de Bayes pour l'écrire comme la somme de deux probabilités. La première est la probabilité d'avoir une collision avec un paquet envoyé par un objet sur la voie montante avant le paquet 1 (cub). La seconde est la probabilité d'avoir une collision avec un paquet envoyé avant, sans collision avec un paquet envoyé sur la voie montante. Dans ce cas là, on aura une collision avec un acquittement envoyé par la station de base. On a donc la formule suivante :

$$P(cb) = P(cb|cub)P(cub) + P(cb|\overline{cub})P(\overline{cub}) = P(cub) + P(cd, \overline{cub}). \quad (C.2)$$

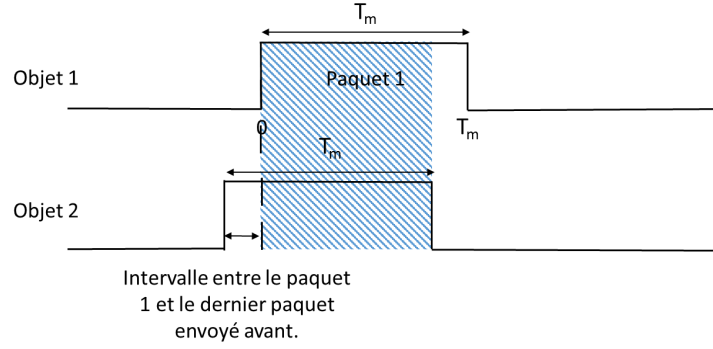


Figure C.1 – Collision entre le paquet 1 et un paquet envoyé avant sur la voie montante.

Dans cette équation, $P(cd)$ est la probabilité d'avoir une collision avec un acquittement envoyé par la station de base avant le paquet 1.

On représente sur la Figure C.1 une collision entre le paquet 1 et un paquet envoyé avant sur la voie montante (cub). Comme on peut le voir sur cette figure, on a une collision entre le paquet 1 et un paquet envoyé avant si et seulement si l'intervalle entre le paquet 1 et le dernier paquet envoyé avant est inférieur à T_m . Par conséquent :

$$P(cub) = 1 - e^{-\lambda_j T_m}. \quad (C.3)$$

Ensuite, comme montré sur la Figure C.2, dans le cas où $T_d \leq T_m$ et sachant que la station de base n'envoie un acquittement que si le canal est libre, on a une collision avec un acquittement sans collision avec un paquet envoyé sur la voie que si le dernier paquet envoyé sur la voie montante avant le paquet considéré est envoyé dans l'intervalle de temps $I_a = [-T_m - T_d - T_a; -T_d - T_m]$ et est reçu et décodé avec succès par la station de base. La probabilité $P(cd, \overline{cub})$ s'écrit donc :

$$P(cd, \overline{cub}) = \underbrace{\left(e^{-\lambda_j(T_d+T_m)} - e^{-\lambda_j(T_d+T_m+T_a)} \right)}_{\text{L'intervalle de temps entre le paquet 1 et le dernier paquet envoyé avant est dans } I_a} \times \underbrace{P(\overline{cb})}_{\substack{\text{Le paquet envoyé dans } I_a \\ \text{n'a pas eu de collision} \\ \text{avec un paquet envoyé avant lui}}} . \quad (C.4)$$

En remplaçant $P(cd, \overline{cub})$ et $P(cub)$ par leurs expressions données ci-dessus, on peut calculer $P(cb)$ la probabilité de ne pas de collision avec un paquet envoyé avant paquet 1 :

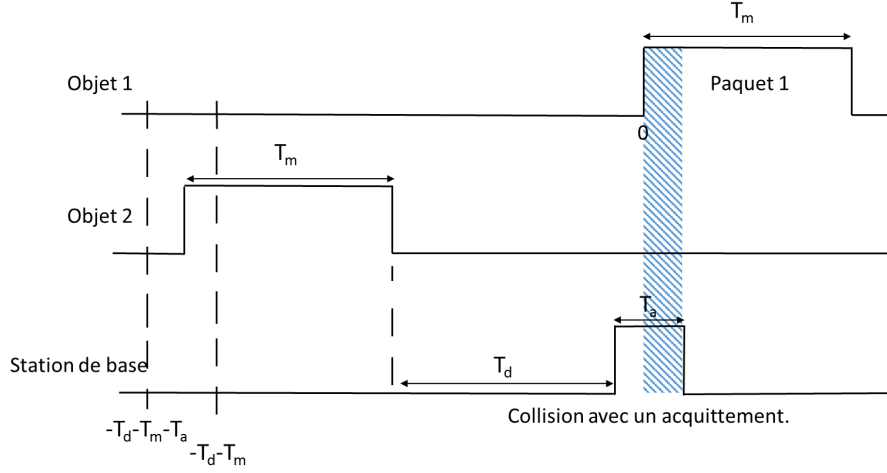


Figure C.2 – Collision entre le paquet 1 et un acquittement, sans avoir de collision avec un paquet envoyé sur la voie montante avant le paquet 1.

$$\begin{aligned}
 P(\overline{cb}) &= 1 - P(cb) \\
 &= \frac{e^{-\lambda_j T_m}}{1 + e^{-\lambda_j (T_d + T_m)} - e^{-\lambda_j (T_d + T_m + T_a)}}.
 \end{aligned} \tag{C.5}$$

Ensuite, la probabilité de ne pas avoir une collision avec un paquet envoyé après le paquet 1 s'écrit :

$$P(\overline{ca}) = e^{-\lambda_j T_m} \tag{C.6}$$

On a maintenant tous les éléments qui permettent de calculer $P(su)$ à partir de l'équation 1 :

$$\begin{aligned}
 P(su) &= P(\overline{cb})P(\overline{ca}) \\
 &= \frac{e^{-2\lambda_j T_m}}{1 + e^{-\lambda_j (T_d + T_m)} - e^{-\lambda_j (T_d + T_m + T_a)}}.
 \end{aligned} \tag{C.7}$$

On doit maintenant regarder la probabilité $P(sd)$, que l'acquittement soit bien reçu. Si on sait que le paquet envoyé sur la voie montante a bien été décodé par la station de base, on a une bonne transmission sur la voie descendante si et seulement si l'acquittement ne subit pas de collisions. C'est à dire :

$$P(sd) = P(su)P(sa|su). \quad (C.8)$$

Dans le cas où $T_d \leq T_m$,

$$P(sa|su) = e^{-\lambda_j(T_d+T_a)}. \quad (C.9)$$

Ceci nous permet de conclure que :

$$P(sd) = \frac{e^{-\lambda_j(2T_m+T_d+T_a)}}{1 + e^{-\lambda_j(T_d+T_m)} - e^{-\lambda_j(T_d+T_m+T_a)}}. \quad (C.10)$$

Nous venons donc de prouver la Proposition 6.1.

Cas 2 : $T_d \geq T_m$

On prouve maintenant la Proposition 6.2, pour cela, on utilise un raisonnement très similaire à celui utilisé dans la section précédente. On commence par calculer $P(cb)$ que l'on décompose suivant l'équation (C.2). Une fois cette décomposition faite, on calcule $P(cd, \overline{cub})$.

Dans le cas où $T_d \geq T_m$, on peut avoir une collision avec un acquittement dans deux situations différentes qui sont illustrées sur la Figure C.3. Afin de bien comprendre ces deux cas, on va numéroter les différents paquets comme ceci :

- On appelle toujours paquet 1 le paquet que l'on considère.
- Paquet 2 désigne un paquet qui a été envoyé par l'objet 2 dans l'intervalle de temps $[-T_d - T_m - T_a; -T_d - T_m]$ et dont l'acquittement interfère avec le paquet 1.
- On appelle paquet 3 le dernier paquet envoyé sur la voie montante avant l'acquittement de paquet 2.

Les deux situations décrites par la Figure C.3 étant incompatibles, on peut écrire $P(cd, \overline{cub})$ comme la somme des probabilités de ces deux situations :

$$P(cd, \overline{cub}) = P(cd, \overline{cub}, \overline{pb}) + P(cd, \overline{cub}, pb). \quad (C.11)$$

Où pb désigne le fait d'avoir un paquet envoyé sur la voie montante entre paquet 2 et son acquittement. On a déjà calculé le premier terme de cette somme. En effet, si aucun paquet n'est envoyé entre paquet 2 et son acquittement, on se retrouve dans la situation étudiée dans le cas $T_d \leq T_m$. En d'autres termes, l'expression de $P(cd, \overline{cub}, \overline{pb})$ dans le cas où $T_d \geq T_m$ est la même que celle de $P(cd, \overline{cub})$ dans le cas où $T_d \leq T_m$ et est donnée par l'équation (C.4).

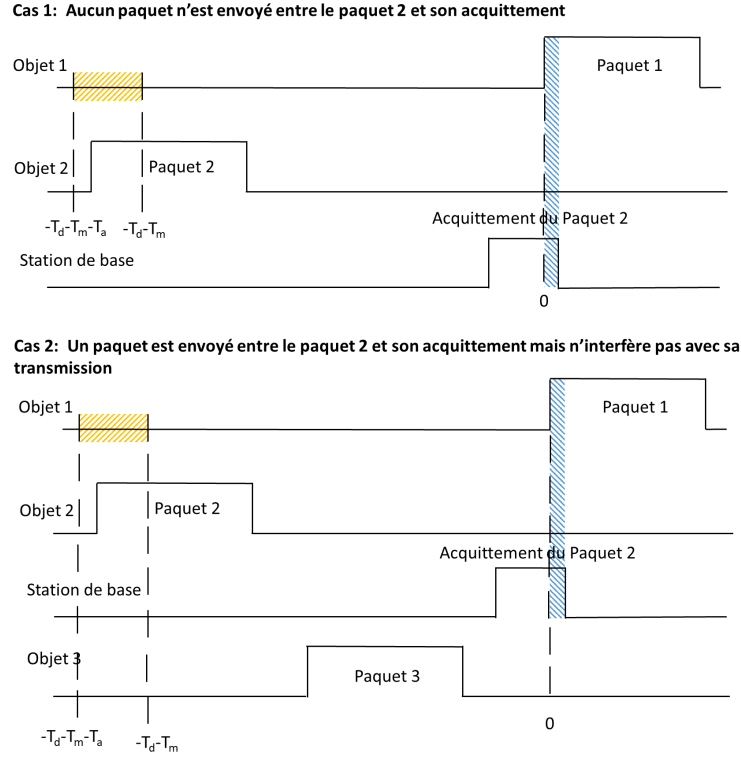


Figure C.3 – Lorsqu'on a une collision avec un acquittement, on peut avoir deux cas possibles. Soit on a des paquets envoyés entre le paquet 2 et son acquittement, soit aucun paquet n'est envoyé dans cet intervalle de temps.

On peut maintenant calculer $P(cd, \overline{cub}, pb)$. On a l'évènement $cd \cap \overline{cub} \cap pb$ si et seulement si :

- Le paquet 2 est bien reçu par la station de base.
- Le paquet 3, qui est le dernier paquet envoyé sur la voie montante avant l'acquittement du paquet 2 est transmis entre le paquet 2 et son acquittement sans gêner la réception de ce dernier (sans collision).

Ces deux évènements sont indépendants. On peut donc écrire $P(cd, \overline{cub}, pb)$ comme le produit de deux probabilités :

$$P(cd, \overline{cub}, pb) = P(pss)P(T_c + T_m \leq T_1 \leq T_c + T_d) \quad (C.12)$$

La première étant la probabilité de pss , c'est à dire la probabilité que le paquet 2 ne subisse pas de collisions, elle s'écrit :

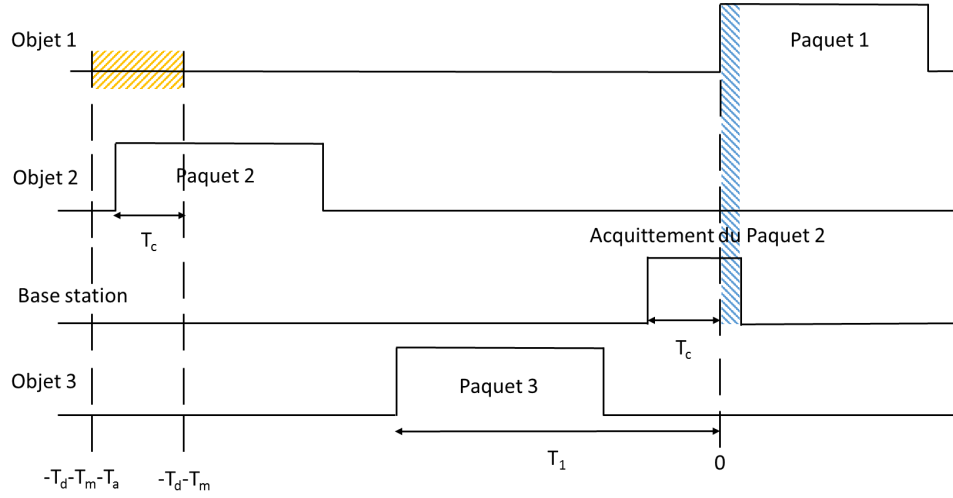


Figure C.4 – T_1 est la durée de l'intervalle entre le paquet 1 et le paquet 3 et T_c est l'intervalle entre le paquet 1 et l'acquittement du paquet 2.

$$\begin{aligned}
 P(pss) &= \underbrace{(1 - e^{-\lambda_j T_a})}_{\text{Proba. qu'un paquet soit envoyé dans } I_a} \times \underbrace{e^{-\lambda_j T_m}}_{\text{Proba. de ne pas avoir de collision avec un paquet envoyé après.}} \times \underbrace{(1 - P(cb))}_{\text{Proba. de ne pas avoir de collision avec un paquet envoyé avant}} \\
 &= (e^{-\lambda_j T_m} - e^{-\lambda_j (T_m + T_a)}) (1 - P(cb)). \tag{C.13}
 \end{aligned}$$

On doit maintenant calculer la seconde probabilité qui est la probabilité que le paquet 3 (dernier paquet envoyé avant l'acquittement) ne gêne pas la transmission de cet acquittement. Pour cela, on va regarder la durée T_1 de l'intervalle entre le paquet 1 et le paquet 3 et la durée T_c de l'intervalle entre le paquet 1 et l'acquittement du paquet 2. Ces deux intervalles sont illustrés sur la Figure C.4. Comme on peut le voir sur cette figure, la probabilité que le paquet 3 soit envoyé après le paquet 2 et ne gêne pas l'acquittement est donnée par :

$$P(T_c + T_m \leq T_1 \leq T_c + T_d) = P(T_m \leq T_1 - T_c \leq T_d). \tag{C.14}$$

Pour le calcul de cette probabilité, on doit calculer la densité de probabilité de $T_1 - T_c$ que l'on notera $f_{T_1 - T_c}$. Vu que le paquet 2 a été transmis sans collision, on sait qu'un seul paquet a été envoyé sur la voie montante dans l'intervalle I_a . On en déduit que T_c suit une loi uniforme sur $[0; T_a]$ [141]. De plus, T_1 suit une loi exponentielle. On peut donc calculer $f_{T_1 - T_c}$ en calculant la convolution des densités de probabilité de T_1 et $-T_c$. On obtient :

$$f_{T_1-T_c}(\tau) = \begin{cases} 0 & \text{si } \tau < -T_a \\ \frac{1}{T_a}(1 - e^{-\lambda_j(\tau+T_a)}) & \text{si } \tau \in [-T_a; 0] \\ \frac{1}{T_a}(e^{-\lambda_j\tau} - e^{-\lambda_j(\tau+T_a)}) & \text{si } 0 < \tau \end{cases} \quad (\text{C.15})$$

On peut, ensuite, en déduire, l'expression de $P(T_m \leq T_1 - T_c \leq T_d)$:

$$P(T_m \leq T_1 - T_c \leq T_d) = \frac{1}{\lambda_j T_a} \left(e^{-\lambda_j T_m} - e^{-\lambda_j(T_m+T_a)} - e^{-\lambda_j T_d} + e^{-\lambda_j(T_d+T_a)} \right). \quad (\text{C.16})$$

On peut, maintenant, réécrire l'ensemble $P(cb)$ en fonction des différentes probabilités calculés précédemment :

$$P(cb) = P(cub) + P(cd, \overline{cub}, \overline{pb}) + P(pss)P(T_m \leq T_1 - T_c \leq T_d). \quad (\text{C.17})$$

Ceci nous permet d'obtenir l'expression de $P(\overline{cb})$ en fonction des paramètres des communications :

$$P(\overline{cb}) = 1 - P(cb) = \frac{e^{-\lambda_j T_m}}{1 + f(\lambda_j, T_m, T_d, T_a)}. \quad (\text{C.18})$$

Dans cette équation,

$$f(\lambda_j, T_m, T_d, T_a) = \left(e^{-\lambda_j T_m} - e^{-\lambda_j(T_m+T_a)} \right) \times \left[e^{-\lambda_j T_d} + \frac{1}{\lambda_j T_a} \left(e^{-\lambda_j T_m} - e^{-\lambda_j(T_m+T_a)} - e^{-\lambda_j T_d} + e^{-\lambda_j(T_d+T_a)} \right) \right]. \quad (\text{C.19})$$

Ensuite, on peut utiliser l'équation (C.1) pour trouver l'expression de $P(su)$. Enfin, dans le cas où $T_d \geq T_m$, $P(ca)$ est égale à :

$$P(sa|su) = e^{-\lambda_j(T_m+T_a)}. \quad (\text{C.20})$$

Ce qui nous permet de conclure que :

$$P(sd) = \frac{e^{-\lambda_T(3T_m+T_a)}}{1 + f(\lambda_T, T_m, T_d, T_a)}. \quad (\text{C.21})$$

Nous venons de prouver la Proposition 6.2.

C.2 Une approximation pour la probabilité de collision dans un réseau ALOHA slotté

Dans cette Annexe, nous prouvons la Proposition 6.3. Tout d'abord, x_t , la probabilité qu'un objet soit en train d'émettre peut être réécrite :

$$x_t = \frac{2}{m+1} \frac{1}{1 + \frac{2P^j(su)}{m+1} \left(\frac{1}{p} - \frac{(m-1)}{2} \right)}. \quad (\text{C.22})$$

Lorsque la probabilité de collision $P^j(\overline{su})$ est élevée et, par conséquent, la probabilité de succès $P^j(su)$ est faible ou alors lorsque p est approximativement égal à on peut utiliser le développement limité $\frac{1}{1+x} \underset{x \approx 0}{\approx} 1 - x$:

$$x_t \approx \frac{2}{m+1} \left[1 - \frac{2P^j(su)}{m+1} \left(\frac{1}{p} - \frac{(m-1)}{2} \right) \right]. \quad (\text{C.23})$$

En remplaçant x_t par son expression dans (6.12), on obtient :

$$P^j(su) \approx e^{\frac{-2N}{m+1} \left[1 - \frac{2P^j(su)}{m+1} \left(\frac{1}{p} - \frac{(m-1)}{2} \right) \right]}. \quad (\text{C.24})$$

En multipliant par la même quantité des deux cotés de l'équation, cette équation peut se réécrire :

$$\frac{-4NP^j(su)}{(m+1)^2} \left(\frac{1}{p} - \frac{m-1}{2} \right) e^{\frac{-4N}{(m+1)^2} P^j(su) \left(\frac{1}{p} - \frac{m-1}{2} \right)} \approx \frac{-4N}{(m+1)^2} \left(\frac{1}{p} - \frac{m-1}{2} \right) e^{\frac{-2N}{m+1}}. \quad (\text{C.25})$$

De plus, dans le cas où :

$$p \geq \frac{1}{e^{\frac{2N}{m+1}-1} \frac{(m+1)^2}{4N} + \frac{m-1}{2}}, \quad (\text{C.26})$$

on peut exprimer la solution de (C.25) sous forme analytique grâce à la fonction $\mathcal{W}$ de Lambert [62]. Afin d'obtenir des valeurs de $P^j(su)$ cohérentes, c'est-à-dire qui appartiennent à l'intervalle $[0; 1]$, on applique la branche principale de cette fonction. On obtient finalement :

$$P^j(\overline{su}) = 1 - P^j(su) \approx 1 + \frac{\mathcal{W}_0 \left(-\frac{4N_j}{(m+1)^2} \left(\frac{1}{p} - \frac{m-1}{2} \right) e^{-\frac{2N_j}{m+1}} \right)}{\left(\frac{1}{p} - \frac{m-1}{2} \right) \frac{4N_j}{(m+1)^2}}. \quad (\text{C.27})$$

Ce qui prouve la Proposition 6.3.

Liste des publications

Revue internationale

- [RI-4] **Bonnefoi, R.**, Moy, C., Palicot, J., “Power Control and Cell Discontinuous Transmission used as a means of decreasing Small-Cell Networks’ Energy Consumption,” in *IEEE Transactions on Green Communications and Networking*, 2018
- [RI-3] **Bonnefoi, R.**, Moy, C., Palicot, J., “Improvement of the LPWAN AMI Backhaul’s Latency thanks to Reinforcement Learning Algorithms”, *Eurasip Journal on Wireless Communications and Networking*, Feb. 2018
- [RI-2] Nafkha, A. and **Bonnefoi, R.**, “Upper and lower bounds for the ergodic capacity of MIMO Jacobi fading channels”, *Opt. Express*, volume 25, pp.12144-12151, 2017
- [RI-1] Palicot, J., Moy, C., Résimond, B., **Bonnefoi, R.**, “Application of Hierarchical and Distributed Cognitive Radio Architecture Management for the Smart Grid”, *Ad hoc networks*, volume 41, pp. 86-98, May 2016

Conférences internationales avec actes

- [CI-9] **Bonnefoi, R.**, Moy, C., Palicot, J., Nafkha A. “Low-Complexity Antenna Selection for Minimizing the Power Consumption of a MIMO Base Station”, *AICT 2018*, July 2018
- [CI-8] **Bonnefoi, R.**, Tarcisio, M., C. Estêvão Fernandes “Latency Efficient Request Access Rate for Congestion Reduction in LTE MTC”, *ICT*, June 2018
- [CI-7] **Bonnefoi, R.**, Besson, L., Moy, C., Kaufmann, E., Palicot, J. “Multi-Armed Bandit Learning in IoT Networks : Learning helps even in non-stationary settings”, *CROWN-COM*, September 2017
- [CI-6] **Bonnefoi, R.**, Moy, C., Palicot, J., “Framework for Hierarchical and Distributed Smart Grid Management”, *URSI General Assembly*, August 2017
- [CI-5] **Bonnefoi, R.**, Moy, C., Farès, H., Palicot, J., “Power Allocation for Minimizing Energy Consumption of OFDMA Downlink with Cell DTx”, *ICT, 2017 IEEE*, May 2017

- [CI-4] **Bonnefoi, R.**, Nafkha, A., “ A New Lower Bound on the Ergodic Capacity of Optical MIMO Channels”, *ICC, 2017 IEEE*, May 2017
- [CI-3] **Bonnefoi, R.**, Moy, C., Palicot, J., “New Macrocell Downlink Energy Consumption Minimization with Cell DTx and Power Control”, *ICC, 2017 IEEE*, May 2017
- [CI-2] **Bonnefoi, R.**, Moy, C., Palicot, J., “Advanced Metering Infrastructure Backhaul Reliability Improvement with Cognitive Radio”, *SmartGridComm, 2017 IEEE*, November 2016
- [CI-1] **Bonnefoi, R.**, Moy, C., Palicot, J., “Dynamic Sleep Mode for Minimizing a Femtocell Power Consumption”, *CROWNCOM*, May 2016

Revues nationales

- [RN-1] A. De Domenico, **R. Bonnefoi**, M. Mendil, C. Gavriluta, J. Palicot, C. Moy, V. Heiries “ Une architecture intelligente pour l’amélioration de l’efficacité énergétique du réseau cellulaire 5G”, *Revue de l’Electricité et de l’Electronique*, December 2016

Conférences nationales avec actes

- [CN-2] **Bonnefoi, R.**, Moy, C., Palicot, J., “Mises en veille dynamique pour Minimiser la Consommation d’Energie d’une Station de Base”, *Colloque GRETSI*, Septembre 2017
- [CN-1] **Bonnefoi, R.**, Mendil, M., Gavritula, C., Moy, C., Palicot, J., Heiries, V. De Domenico, A., Caire, R. Hadjsaid, N., “A Low Energy Consumption Wireless Network”, *URSI Scientific Days*, March 2016

Communications sans actes

- [CSA-2] **Bonnefoi, R.**, “Power Allocation with Cell Discontinuous Transmission (DTx)”, Fortaleza, October 2017
- [CSA-1] **Bonnefoi, R.**, “Power Allocation with Cell Discontinuous Transmission (DTx)”, Séminaire SCEE, March 2017

Démonstrations

- [D-1] **R. Bonnefoi**, L. Besson and C. Moy “ Multi-Armed Bandit Learning in IoT Networks (MALIN)”, *ICT*, June 2018

Table des figures

1	Dans le système étudié nous avons différents réseaux de communication : un réseau de communication mobile et un réseau de communication pour le réseau électrique intelligent.	3
1.1	Version simplifiée du cycle intelligent.	8
1.2	L'architecture hiérarchique HDCRAM avec ses deux branches de gestion de la reconfiguration et de l'information.	11
1.3	Application de l'architecture HDCRAM pour l'adaptation de la modulation. .	13
1.4	Le système étudié dans lequel les stations de base sont associées à des sources d'énergie renouvelables et intermittentes.	13
1.5	Description HDCRAM du réseau mobile étudié. Dans cette description, on utilise deux HDCRAM. La première pour l'énergie et la seconde pour les communications.	14
1.6	Description HDCRAM de la gestion de la production et de la consommation dans un réseau électrique.	16
1.7	Description complète du système considéré.	17
2.1	Représentation schématique d'un réseau cellulaire composé de plusieurs stations de base chacune servant différents utilisateurs.	24
2.2	Une trame LTE est divisée en sous-frames temporelles et en blocs de ressources fréquentielles.	25
2.3	Allocation de ressources en OFDMA. Dans cet exemple, la station de base dispose de $N_c = 6$ blocs de ressources pour servir $N_u = 2$ utilisateurs.	28
2.4	Illustrations des deux cas possibles pour l'exemple donné par l'équation (2.8).	35
2.5	Dans le standard LTE, la station de base peut être mise en veille pendant six des dix sous-frames.	39

2.6	Modélisation du problème d'allocation de puissance en TDMA avec de la transmission discontinue. Dans cet exemple, la station de base sert trois utilisateurs $\{U1, U2, U3\}$ avant de passer en veille.	40
3.1	Schéma de résolution du problème posé.	49
3.2	Fonction $\rho_{\text{lim}}^M(r)$ qui permet, pour une station de base, de savoir si un utilisateur doit être servi pendant μ_k^{min} ou pendant μ_k^{opt1} . Cette figure présente un nouveau mode de classification des stations de base spécifique au problème d'allocation de puissance avec Cell DTx.	51
3.3	Évaluation des performances de l'algorithme proposé lorsque une macro cellule sert en moyenne 10 utilisateurs, en termes de puissance moyenne consommée.	59
3.4	(a) Temps total pendant lequel la station de base est active. On voit que ce temps est plus long lorsque le contrôle de puissance est utilisé. (b) Puissance d'émission moyenne de la station de base pendant une trame.	60
3.5	Puissance moyenne consommée par une femto-cellule. On remarque que pour ce type de station de base le contrôle de puissance n'a aucun effet.	61
3.6	Un exemple de réseau mobile dans lequel on a $N_{BS} = 2$ stations de base et $N_u = 5$ utilisateurs. Dans ce cas là, $S_1 = \{1, 2\}$ et $S_2 = \{3, 4, 5\}$	62
3.7	Le réseau 2 qui est composé de 17 cellules hexagonales.	64
3.8	Evolution des temps de service et des puissances instatannées lorsque de nouveaux utilisateurs se connectent au réseau.	65
3.9	Puissance moyenne consommée par chacune des stations de base du réseau avec les différentes stratégies comparés.	69
3.10	(a) Temps moyen pendant lequel une station de base est active. (b) Puissance d'émission moyenne de la station de base lorsqu'elle est active.	70
3.11	Puissance moyenne consommée par la station de base pour différentes valeurs de $P_{\text{max}}^{\text{DTx}}$	71
3.12	Évolution de la puissance consommée en fonction de $P_{\text{max}}^{\text{DTx}}$ pour (a) une contrainte de capacité de 1 Mbits/s, (b) une contrainte de capacité de 5 Mbits/s.	72
3.13	Évolution de la puissance consommée par une station de base du réseau 2 en fonction de $P_{\text{max}}^{\text{DTx}}$	73
4.1	Pendant chaque trame, la station de base est active pendant t_a et ensuite est mise en veille. Lorsque la station de base est active, de l'OFDMA est utilisé pour l'accès multiple.	78

4.2	La largeur de bande est divisée en N_c blocs de ressources. Dans cet exemple, on a 6 blocs de ressources qui sont répartis entre 2 utilisateurs.	81
4.3	Les trois situations possibles pour le calcul de μ_{opt} . Dans le premier cas, la contrainte de temps de l'équation (4.10b) est saturée. Dans le second, la contrainte de puissance de l'équation (4.10d) est saturée. Dans le dernier, aucune de ces contraintes n'est saturée.	86
4.4	Puissance moyenne consommée par la station de base avec l'allocation de puissance optimale appliquée avec soit l'allocation de ressources optimale (r1), soit avec l'allocation de ressources qui ne considère pas le Cell DTx (r2). . . .	90
4.5	Comparaison de la solution proposée (utilisation de l'algorithme BABS+ACG pour l'allocation de ressources) avec une solution optimale (recherche exhaustive). 92	
4.6	Puissance moyenne consommée par la station de base en utilisant l'algorithme BABS+ACG pour l'allocation de ressources et avec différentes allocations de puissance.	93
5.1	Le réseau électrique sous sa forme historique. On a un découpage en quatre parties : la production, le réseau de transport, le réseau de distribution et la consommation.	98
5.2	L'architecture HDCRAM est une architecture de gestion hiérarchique des systèmes distribués intelligents.	101
5.3	Représentation HDCRAM pour la gestion de la production et de la consommation à l'échelle locale.	103
5.4	Représentation HDCRAM pour la gestion de la production et de la consommation à l'échelle globale.	104
5.5	Représentation HDCRAM pour la gestion du réseau de transport dans un réseau électrique intelligent.	106
5.6	Représentation HDCRAM pour la gestion du réseau de distribution dans un réseau électrique intelligent.	107
6.1	Le modèle considéré dans lequel un grand nombre d'objets connectés envoie des paquets vers une station de base.	111
6.2	On suppose que, dès que deux paquets se retrouvent au même moment dans un canal, on a une collision et aucun des deux paquets ne peut être décodé. .	112

6.3	Lorsque la station de base reçoit un paquet envoyé par un objet, elle attend pendant un intervalle de durée T_d et transmet ensuite un acquittement si le canal est vide.	114
6.4	Probabilité d'avoir une bonne transmission à la fois sur la voie montante et la voie descendante ($P^j(sd)$) en fonction de $\lambda_j T_m$. Pour ces simulations on considère que $T_d = 1$ s et les valeurs de T_a et T_m sont en adéquation avec le standard LoRaWAN.	116
6.5	Dans le réseau ALOHA slotté considéré, le temps est divisé en slots de même durée. On suppose que chaque slot se décompose en deux parties. La première est utilisée pour la voie montante, la seconde pour l'acquittement envoyé par la station de base.	117
6.6	On modélise par une chaîne de Markov le comportement de chacun des objets qui communiquent dans le canal j	117
6.7	Approximation de la probabilité de collision $P^j(su)$ dans un réseau ALOHA slotté.	119
6.8	On s'intéresse aux communications d'un agrégateur de données qui communique avec un centre de contrôle par le biais d'un réseau LPWAN.	120
6.9	Évolution de la probabilité de réussir une transmission et de la latence avec les algorithmes UCB et TS.	127
6.10	Schéma de la démonstration réalisée. On utilise trois USRP synchronisées par une octoclock. La première USRP joue le rôle d'agrégateur qui transmet régulièrement des paquets vers une seconde USRP qui joue le rôle de station de base et qui acquitte les paquets qu'elle reçoit correctement. Enfin une troisième USRP génère un trafic interférant (généré par les objets à portée).	130
6.11	Les quatre canaux tels qu'ils sont vus par la station de base.	130
6.12	pourcentage d'utilisation de chaque canal par l'agrégateur avec l'algorithme UCB et avec une sélection aléatoire du canal. On voit que l'agrégateur a concentré ses émissions sur le canal ayant le plus faible taux d'occupation par les objets aux alentours.	131
6.13	Performances des algorithmes de bandits multibras en fonction de la proportion d'objets dynamiques (agrégateurs).	135
6.14	Gain apporté par l'apprentissage comparé à une stratégie avec laquelle le canal est choisi aléatoirement. Ces résultats sont obtenus après 10^6 slots. Soit en moyenne après que chaque objet ait transmis, en moyenne, 1000 paquets. .	136

C.1	Collision entre le paquet 1 et un paquet envoyé avant sur la voie montante. .	157
C.2	Collision entre le paquet 1 et un acquittement, sans avoir de collision avec un paquet envoyé sur la voie montante avant le paquet 1.	158
C.3	Lorsqu'on a une collision avec un acquittement, on peut avoir deux cas possibles. Soit on a des paquets envoyés entre le paquet 2 et son acquittement, soit aucun paquet n'est envoyé dans cet intervalle de temps.	160
C.4	T_1 est la durée de l'intervalle entre le paquet 1 et le paquet 3 et T_c est l'intervalle entre le paquet 1 et l'acquittement du paquet 2.	161

Liste des tableaux

2.1	Puissance consommée par les différents types de stations de base et dans les différentes phases. P_s en veille, P_0 en statique et $P_{\max}$ ($\geq P_{Tx}$)	38
3.1	Valeur de ρ_{lim}^M pour les différents types de stations de base.	52
4.1	Consommation supplémentaire liée à l'utilisation de l'allocation de ressources optimale sans Cell DTx.	90
5.1	Réseaux de communication utilisés pour la gestion du réseau électrique intelligent [83, 84]	100
C.1	Évènements dont la probabilité est utilisée dans le calcul des probabilités d'avoir une bonne transmission.	156

Bibliographie

- [1] S. Mingay, “Green IT : A New Industry Shock Wave,” tech. rep., Gartner, 01 2007.
- [2] A. Fehske, G. Fettweis, J. Malmudin, and G. Biczok, “The global footprint of mobile communications : The ecological and economic perspective,” *IEEE Communications Magazine*, vol. 49, pp. 55–62, August 2011.
- [3] J. Mitola and G. Q. Maguire, “Cognitive radio : making software radios more personal,” *IEEE Personal Communications*, vol. 6, pp. 13–18, Aug 1999.
- [4] G. Gur and F. Alagoz, “Green wireless communications via cognitive dimension : an overview,” *IEEE Network*, vol. 25, pp. 50–56, March 2011.
- [5] I. Ashraf, F. Boccardi, and L. Ho, “SLEEP mode techniques for small cell deployments,” *IEEE Communications Magazine*, vol. 49, pp. 72–79, August 2011.
- [6] U. Raza, P. Kulkarni, and M. Sooriyabandara, “Low power wide area networks : An overview,” *IEEE Communications Surveys Tutorials*, vol. 19, pp. 855–873, Secondquarter 2017.
- [7] M. Hasan, E. Hossain, and D. Niyato, “Random access for machine-to-machine communication in lte-advanced networks : issues and approaches,” *IEEE Communications Magazine*, vol. 51, pp. 86–93, June 2013.
- [8] P. Frenger, P. Moberg, J. Malmudin, Y. Jading, and I. Godor, “Reducing Energy Consumption in LTE with Cell DTX,” in *2011 IEEE 73rd Vehicular Technology Conference (VTC Spring)*, pp. 1–5, May 2011.
- [9] L. Godard, C. Moy, and J. Palicot, “From a configuration management to a cognitive radio management of sdr systems,” in *2006 1st International Conference on Cognitive Radio Oriented Wireless Networks and Communications*, pp. 1–5, June 2006.
- [10] J. Mitola, “The software radio architecture,” *IEEE Communications Magazine*, vol. 33, pp. 26–38, May 1995.
- [11] S. Haykin, “Cognitive radio : brain-empowered wireless communications,” *IEEE Journal on Selected Areas in Communications*, vol. 23, pp. 201–220, Feb 2005.
- [12] “Tableau National de Répartition des Bandes de Fréquences,” tech. rep., Agence Nationale des Fréquences, Dec 2017.

- [13] M. Nekovee, “Quantifying the Availability of TV White Spaces for Cognitive Radio Operation in the UK,” in *2009 IEEE International Conference on Communications Workshops*, pp. 1–5, June 2009.
- [14] T. Yucek and H. Arslan, “A survey of spectrum sensing algorithms for cognitive radio applications,” *IEEE Communications Surveys Tutorials*, vol. 11, pp. 116–130, First 2009.
- [15] J. Palicot, *Radio Engineering : From Software Radio to Cognitive Radio*. Wiley-ISTE, Aug. 2011.
- [16] X. Wu, *Hierarchical Reconfiguration Management for Heterogeneous Cognitive Green Radio Equipments*. PhD thesis, 2016.
- [17] O. Lazrak, P. Leray, and C. Moy, “Hdcram proof-of-concept for opportunistic spectrum access,” in *2012 15th Euromicro Conference on Digital System Design*, pp. 453–458, Sept 2012.
- [18] J. Palicot, C. Moy, B. Résimont, and R. Bonnefoi, “Application of hierarchical and distributed cognitive architecture management for the smart grid,” *Ad Hoc Networks*, vol. 41, pp. 86 – 98, 2016. Cognitive Radio Based Smart Grid The Future of the Traditional Electrical Grid.
- [19] P. Samadi, A. H. Mohsenian-Rad, R. Schober, V. W. S. Wong, and J. Jatskevich, “Optimal Real-Time Pricing Algorithm Based on Utility Maximization for Smart Grid,” in *2010 First IEEE International Conference on Smart Grid Communications*, pp. 415–420, Oct 2010.
- [20] K. Moslehi and R. Kumar, “A reliability perspective of the smart grid,” *IEEE Transactions on Smart Grid*, vol. 1, pp. 57–64, June 2010.
- [21] R. Bonnefoi, C. Moy, and J. Palicot, “Framework for hierarchical and distributed smart grid management,” in *2017 XXXIInd General Assembly and Scientific Symposium of the International Union of Radio Science (URSI GASS)*, pp. 1–4, Aug 2017.
- [22] “Smart Grid Reference Architecture,” tech. rep., CEN-CENELEC-ETSI Smart Grid Coordination Group, November 2012.
- [23] “3GPP Long Term Evolution : System Overview, product Development and Test Challenges,” tech. rep., Keysight Technologies, August 2014.
- [24] G. Caire, D. Tuninetti, and S. Verdu, “Suboptimality of TDMA in the low-power regime,” *IEEE Transactions on Information Theory*, vol. 50, pp. 608–620, April 2004.
- [25] A. J. Viterbi, *CDMA : Principles of Spread Spectrum Communication*. Redwood City, CA, USA : Addison Wesley Longman Publishing Co., Inc., 1995.
- [26] H. Yin and H. Liu, “Performance of space-division multiple-access (SDMA) with scheduling,” *IEEE Transactions on Wireless Communications*, vol. 1, pp. 611–618, Oct 2002.

-
- [27] A. Benjebbour, Y. Saito, Y. Kishiyama, A. Li, A. Harada, and T. Nakamura, "Concept and practical considerations of non-orthogonal multiple access (noma) for future radio access," in *2013 International Symposium on Intelligent Signal Processing and Communication Systems*, pp. 770–774, Nov 2013.
 - [28] J. Jang and K. B. Lee, "Transmit power adaptation for multiuser OFDM systems," *IEEE Journal on Selected Areas in Communications*, vol. 21, pp. 171–178, Feb 2003.
 - [29] A. Serra Pages, *Link Level Performance Evaluation and Link Abstraction for LTE/LTE-Advanced Downlink*. PhD thesis, Universitat Politècnica de Catalunya, 2016.
 - [30] W. Rhee and J. M. Cioffi, "Increase in capacity of multiuser OFDM system using dynamic subchannel allocation," in *VTC2000-Spring. 2000 IEEE 51st Vehicular Technology Conference Proceedings (Cat. No.00CH37026)*, vol. 2, pp. 1085–1089 vol.2, 2000.
 - [31] G. Miao, N. Himayat, and G. Y. Li, "Energy-efficient link adaptation in frequency-selective channels," *IEEE Transactions on Communications*, vol. 58, pp. 545–554, February 2010.
 - [32] L. Sboui, Z. Rezki, and M. S. Alouini, "Energy-efficient power control for ofdma cellular networks," in *2016 IEEE 27th Annual International Symposium on Personal, Indoor, and Mobile Radio Communications (PIMRC)*, pp. 1–6, Sept 2016.
 - [33] Y. J. Zhang and K. B. Letaief, "Multiuser adaptive subcarrier-and-bit allocation with adaptive cell selection for ofdm systems," *IEEE Transactions on Wireless Communications*, vol. 3, pp. 1566–1575, Sept 2004.
 - [34] K. Seong, M. Mohseni, and J. M. Cioffi, "Optimal Resource Allocation for OFDMA Downlink Systems," in *2006 IEEE International Symposium on Information Theory*, pp. 1394–1398, July 2006.
 - [35] N. Loehr, *Bijective Combinatorics*. Chapman & Hall/CRC, 1st ed., 2011.
 - [36] S. Sadr, A. Anpalagan, and K. Raahemifar, "Radio resource allocation algorithms for the downlink of multiuser ofdm communication systems," *IEEE Communications Surveys Tutorials*, vol. 11, pp. 92–106, rd 2009.
 - [37] W. Yu and R. Lui, "Dual methods for nonconvex spectrum optimization of multicarrier systems," *IEEE Transactions on Communications*, vol. 54, pp. 1310–1322, July 2006.
 - [38] D. P. Palomar and M. Chiang, "A tutorial on decomposition methods for network utility maximization," *IEEE Journal on Selected Areas in Communications*, vol. 24, pp. 1439–1451, Aug 2006.
 - [39] D. Kivanc, G. Li, and H. Liu, "Computationally efficient bandwidth allocation and power control for OFDMA," *IEEE Transactions on Wireless Communications*, vol. 2, pp. 1150–1158, Nov 2003.
 - [40] S. Boyd and L. Vandenberghe, *Convex Optimization*. New York, NY, USA : Cambridge University Press, 2004.

- [41] M. Luptacik, *Mathematical Optimization and Economic Analysis*. Springer, 2010.
- [42] J. Hoadley and P. Maveddat, “Enabling small cell deployment with HetNet,” *IEEE Wireless Communications*, vol. 19, pp. 4–5, April 2012.
- [43] I. Hwang, B. Song, and S. S. Soliman, “A holistic view on hyper-dense heterogeneous and small cell networks,” *IEEE Communications Magazine*, vol. 51, pp. 20–27, June 2013.
- [44] D. Lopez-Perez, X. Chu, A. V. Vasilakos, and H. Claussen, “Power Minimization Based Resource Allocation for Interference Mitigation in OFDMA Femtocell Networks,” *IEEE Journal on Selected Areas in Communications*, vol. 32, pp. 333–344, February 2014.
- [45] R. Irmer, H. Droste, P. Marsch, M. Grieger, G. Fettweis, S. Brueck, H. P. Mayer, L. Thiele, and V. Jungnickel, “Coordinated multipoint : Concepts, performance, and field trial results,” *IEEE Communications Magazine*, vol. 49, pp. 102–111, February 2011.
- [46] E. Telatar, “Capacity of Multi-antenna Gaussian Channels,” *European Transactions on Telecommunications*, vol. 10, no. 6, pp. 585–595, 1999.
- [47] E. S. Lo, P. W. C. Chan, V. K. N. Lau, R. S. Cheng, K. B. Letaief, R. D. Murch, and W. H. Mow, “Adaptive Resource Allocation and Capacity Comparison of Downlink Multiuser MIMO-MC-CDMA and MIMO-OFDMA,” *IEEE Transactions on Wireless Communications*, vol. 6, pp. 1083–1093, March 2007.
- [48] O. Arnold, F. Richter, G. Fettweis, and O. Blume, “Power consumption modeling of different base station types in heterogeneous cellular networks,” in *2010 Future Network Mobile Summit*, pp. 1–8, June 2010.
- [49] G. Auer, V. Giannini, C. Desset, I. Godor, P. Skillermark, M. Olsson, M. A. Imran, D. Sabella, M. J. Gonzalez, O. Blume, and A. Fehske, “How much energy is needed to run a wireless network?,” *IEEE Wireless Communications*, vol. 18, pp. 40–49, October 2011.
- [50] J. Wu, Y. Zhang, M. Zukerman, and E. K. N. Yung, “Energy-Efficient Base-Stations Sleep-Mode Techniques in Green Cellular Networks : A Survey,” *IEEE Communications Surveys Tutorials*, vol. 17, pp. 803–826, Secondquarter 2015.
- [51] M. Peng, Y. Li, J. Jiang, J. Li, and C. Wang, “Heterogeneous cloud radio access networks : a new perspective for enhancing spectral and energy efficiencies,” *IEEE Wireless Communications*, vol. 21, pp. 126–135, December 2014.
- [52] S. Samarakoon, M. Bennis, W. Saad, and M. Latva-aho, “Opportunistic sleep mode strategies in wireless small cell networks,” in *2014 IEEE International Conference on Communications (ICC)*, pp. 2707–2712, June 2014.

-
- [53] F. Han, S. Zhao, L. Zhang, and J. Wu, "Survey of strategies for switching off base stations in heterogeneous networks for greener 5g systems," *IEEE Access*, vol. 4, pp. 4959–4973, 2016.
 - [54] L. Suarez and L. Nuaymi, "Multi-Size Cell Expansion for Energy-Efficient Cell Breathing in Green Wireless Networks," in *2015 IEEE 82nd Vehicular Technology Conference (VTC2015-Fall)*, pp. 1–5, Sept 2015.
 - [55] E. Oh, K. Son, and B. Krishnamachari, "Dynamic Base Station Switching-On/Off Strategies for Green Cellular Networks," *IEEE Transactions on Wireless Communications*, vol. 12, pp. 2126–2136, May 2013.
 - [56] A. Chatzipapas, S. Alouf, and V. Mancuso, "On the minimization of power consumption in base stations using on/off power amplifiers," in *2011 IEEE Online Conference on Green Communications*, pp. 18–23, Sept 2011.
 - [57] T. Chen, Y. Yang, H. Zhang, H. Kim, and K. Horneman, "Network energy saving technologies for green wireless access networks," *IEEE Wireless Communications*, vol. 18, pp. 30–38, October 2011.
 - [58] B. Soret, A. D. Domenico, S. Bazzi, N. H. Mahmood, and K. I. Pedersen, "Interference coordination for 5g new radio," *IEEE Wireless Communications*, vol. 25, pp. 131–137, JUNE 2018.
 - [59] P. Chang and G. Miao, "Joint optimization of base station deep-sleep and dtx micro-sleep," in *2016 IEEE Globecom Workshops (GC Wkshps)*, pp. 1–6, Dec 2016.
 - [60] H. Holtkamp, G. Auer, and H. Haas, "On Minimizing Base Station Power Consumption," in *2011 IEEE Vehicular Technology Conference (VTC Fall)*, pp. 1–5, Sept. 2011.
 - [61] R. Wang, J. S. Thompson, H. Haas, and P. M. Grant, "Sleep mode design for green base stations," *IET Communications*, vol. 5, pp. 2606–2616, Dec 2011.
 - [62] R. M. Corless, G. H. Gonnet, D. E. G. Hare, D. J. Jeffrey, and D. E. Knuth, "On the Lambert W Function," in *ADVANCES IN COMPUTATIONAL MATHEMATICS*, pp. 329–359, 1996.
 - [63] H. Holtkamp, G. Auer, S. Bazzi, and H. Haas, "Minimizing base station power consumption," *IEEE Journal on Selected Areas in Communications*, vol. 32, pp. 297–306, February 2014.
 - [64] H. B. Ren, M. Zhao, J. K. Zhu, and W. Y. Zhou, "Energy-efficient resource allocation for ofdma networks with sleep mode," *Electronics Letters*, vol. 49, pp. 111–113, January 2013.
 - [65] W. Nie, Y. Zhong, F. C. Zheng, W. Zhang, and T. O'Farrell, "HetNets With Random DTX Scheme : Local Delay and Energy Efficiency," *IEEE Transactions on Vehicular Technology*, vol. 65, pp. 6601–6613, Aug 2016.

- [66] P. Chang and G. Miao, "Energy and Spectral Efficiency of Cellular Networks with Discontinuous Transmission," *IEEE Transactions on Wireless Communications*, vol. PP, no. 99, pp. 1–1, 2017.
- [67] S. Tombaz, S. w. Han, K. W. Sung, and J. Zander, "Energy Efficient Network Deployment With Cell DTX," *IEEE Communications Letters*, vol. 18, pp. 977–980, June 2014.
- [68] G. Andersson, A. Vastberg, A. Devlic, and C. Cavdar, "Energy efficient heterogeneous network deployment with cell DTX," in *2016 IEEE International Conference on Communications (ICC)*, pp. 1–6, May 2016.
- [69] A. D. Domenico, R. Gupta, and E. C. Strinati, "Dynamic Traffic Management for Green Open Access Femtocell Networks," in *2012 IEEE 75th Vehicular Technology Conference (VTC Spring)*, pp. 1–6, May 2012.
- [70] R. Gupta, E. C. Strinati, and D. Ktenas, "Energy efficient joint DTX and MIMO in cloud Radio Access Networks," in *2012 IEEE 1st International Conference on Cloud Networking (CLOUDNET)*, pp. 191–196, Nov 2012.
- [71] S. Wu, F. Liu, Z. Zeng, and H. Xia, "Cooperative Sleep and Power Allocation for Energy Saving in Dense Small Cell Networks," *IEEE Access*, vol. 4, pp. 6993–7004, 2016.
- [72] A. Cheng, H. Jin, J. Li, Y. Yu, and M. Peng, "Joint discontinuous transmission and power control for high energy efficiency in heterogeneous small cell networks," in *2014 IEEE 25th Annual International Symposium on Personal, Indoor, and Mobile Radio Communication (PIMRC)*, pp. 970–975, Sept 2014.
- [73] R. Bonnefoi, C. Moy, and J. Palicot, "Dynamic sleep mode for minimizing a femto-cell power consumption," in *Cognitive Radio Oriented Wireless Networks* (D. Nogu  t, K. Moessner, and J. Palicot, eds.), (Cham), pp. 618–629, Springer International Publishing, 2016.
- [74] R. Bonnefoi, C. Moy, and J. Palicot, "New macrocell downlink energy consumption minimization with cell DTx and power control," in *2017 IEEE International Conference on Communications (ICC)*, pp. 1–7, May 2017.
- [75] R. Bonnefoi, C. Moy, and J. Palicot, "Mises en Veille Dynamiques pour Minimiser la Consommation d'  nergie d'une Station de Base," in *Colloque GRETSI*, pp. 1–4, Sept. 2017.
- [76] R. Bonnefoi, C. Moy, and J. Palicot, "Power control and cell discontinuous transmission used as a means of decreasing small-cell networks' energy consumption," *IEEE Transactions on Green Communications and Networking*, pp. 1–1, 2018.
- [77] J. D. Hoffman, *Numerical methods for engineers and scientists*. New York, NY : McGraw-Hill, 1992.

- [78] P. Kyosti and al., “Winner II Deliverable 1.1.2. : Winner II Channel Models,” tech. rep., September 2007.
- [79] “Propagation data and prediction models for the planning of indoor radiocommunication systems and radio local area networks in the frequency range 900 MHz to 100 GHz,” tech. rep., ITU, September 1997.
- [80] R. Bonnefoi, C. Moy, H. Farès, and J. Palicot, “Power allocation for minimizing energy consumption of OFDMA downlink with Cell DTx,” in *2017 24th International Conference on Telecommunications (ICT)*, pp. 1–6, May 2017.
- [81] “3rd Generation Partnership Project ; Technical Specification Group Radio Access Network ; Evolved Universal Terrestrial Radio Access (E-UTRA) ; Base Station (BS) radio transmission and reception (Release 8),” tech. rep., 3GPP, May 2008.
- [82] V. C. Gungor, D. Sahin, T. Kocak, S. Ergut, C. Buccella, C. Cecati, and G. P. Hancke, “Smart grid technologies : Communication technologies and standards,” *IEEE Transactions on Industrial Informatics*, vol. 7, pp. 529–539, Nov 2011.
- [83] “Smart Grid Reference Architecture,” tech. rep., CEN-CENELEC-ETSI Smart Grid Coordination Group, November 2012.
- [84] “Overview of SG-CG Methodologies,” tech. rep., CEN-CENELEC-ETSI Smart Grid Coordination Group, August 2014.
- [85] “Communications Requirements of Smart Grid Technologies,” tech. rep., Department of Energy of USA, October 2010.
- [86] J. Lloret, J. Tomas, A. Canovas, and L. Parra, “An integrated iot architecture for smart metering,” *IEEE Communications Magazine*, vol. 54, pp. 50–57, December 2016.
- [87] R. Verschae, H. Kawashima, T. Kato, and T. Matsuyama, “A distributed hierarchical architecture for community-based power balancing,” in *Smart Grid Communications (SmartGridComm), 2014 IEEE International Conference on*, pp. 163–169, Nov 2014.
- [88] C. Dou and B. Liu, “Multi-agent based hierarchical hybrid control for smart microgrid,” *Smart Grid, IEEE Transactions on*, vol. 4, pp. 771–778, June 2013.
- [89] D. Li and S. Jayaweera, “Distributed smart-home decision-making in a hierarchical interactive smart grid architecture,” *Parallel and Distributed Systems, IEEE Transactions on*, vol. 26, pp. 75–84, Jan 2015.
- [90] R. R. Mohassel, A. Fung, F. Mohammadi, and K. Raahemifar, “A survey on advanced metering infrastructure,” *International Journal of Electrical Power & Energy Systems*, vol. 63, pp. 473 – 484, 2014.
- [91] Z. M. Fadlullah, M. M. Fouda, N. Kato, A. Takeuchi, N. Iwasaki, and Y. Nozaki, “Toward intelligent machine-to-machine communications in smart grid,” *IEEE Communications Magazine*, vol. 49, pp. 60–65, April 2011.

- [92] D. Chen, M. Nixon, and A. Mok, *Why WirelessHART*, pp. 195–199. Boston, MA : Springer US, 2010.
- [93] “3rd Generation Partnership Project ; Technical Specification Group GSM/EDGE Radio Access Network ; Cellular System Support for Ultra Low Complexity and Low Throughput Internet of Things ; (Release 13),” tech. rep., 3GPP, August 2015.
- [94] G. C. Madueno, C. Stefanovic, and P. Popovski, “Reengineering gsm/gprs towards a dedicated network for massive smart metering,” in *2014 IEEE International Conference on Smart Grid Communications (SmartGridComm)*, pp. 338–343, Nov 2014.
- [95] N. Abramson, “The aloha system : Another alternative for computer communications,” in *Proceedings of the November 17-19, 1970, Fall Joint Computer Conference*, AFIPS ’70 (Fall), (New York, NY, USA), pp. 281–285, ACM, 1970.
- [96] M. T. Do, C. Goursaud, and J. M. Gorce, “On the benefits of random fdma schemes in ultra narrow band networks,” in *2014 12th International Symposium on Modeling and Optimization in Mobile, Ad Hoc, and Wireless Networks (WiOpt)*, pp. 672–677, May 2014.
- [97] N. Sornin and M. Luis and T. Eirich and T. Kramp and O. Hersent, “LoRaWAN Specification,” tech. rep., LoRa Alliance, Inc., July 2016.
- [98] W. Webb, *Understanding Weightless : Technology, Equipment, and Network Deployment for M2M Communications in White Space*. New York, NY, USA : Cambridge University Press, 1st ed., 2012.
- [99] M. Centenaro, L. Vangelista, A. Zanella, and M. Zorzi, “Long-Range Communications in Unlicensed Bands : the Rising Stars in the IoT and Smart City Scenarios,” *IEEE Wireless Communications*, vol. 23, October 2016.
- [100] X. Mamo, S. Mallet, T. Coste, and S. Grenard, “Distribution automation : The cornerstone for smart grid development strategy,” in *2009 IEEE Power Energy Society General Meeting*, pp. 1–6, July 2009.
- [101] D. Tholomier, H. Kang, and B. Cvorovic, “Phasor measurement units : Functionality and applications,” in *2009 Power Systems Conference*, pp. 1–12, March 2009.
- [102] S. Matsumoto, Y. Serizawa, F. Fujikawa, T. Shioyama, Y. Ishihara, S. Katayama, T. Kase, and A. Ishibashi, “Wide-area situational awareness (wasa) system based upon international standards,” *IET Conference Proceedings*, pp. 82–82(1), January 2012.
- [103] A. G. Phadke and J. S. Thorp, “Communication needs for wide area measurement applications,” in *2010 5th International Conference on Critical Infrastructure (CRIS)*, pp. 1–7, Sept 2010.
- [104] Y. P. E. Wang, X. Lin, A. Adhikary, A. Grovlen, Y. Sui, Y. Blankenship, J. Bergman, and H. S. Razaghi, “A primer on 3gpp narrowband internet of things,” *IEEE Communications Magazine*, vol. 55, pp. 117–123, March 2017.

- [105] N. A. Johansson, Y. P. E. Wang, E. Eriksson, and M. Hessler, "Radio access for ultra-reliable and low-latency 5g communications," in *2015 IEEE International Conference on Communication Workshop (ICCW)*, pp. 1184–1189, June 2015.
- [106] S. Chowdhury and P. Crossley, *Microgrids and active distribution networks*. IET renewable energy series, Stevenage : The Institution of Engineering and Technology, 2009.
- [107] R. E. Brown, "Impact of smart grid on distribution system design," in *2008 IEEE Power and Energy Society General Meeting - Conversion and Delivery of Electrical Energy in the 21st Century*, pp. 1–4, July 2008.
- [108] K. Budka, J. Deshpande, and M. Thottan, *Communication Networks for Smart Grids : Making Smart Grid Real*. Springer Publishing Company, Incorporated, 2014.
- [109] K. V. Katsaros, B. Yang, W. K. Chai, and G. Pavlou, "Low latency communication infrastructure for synchrophasor applications in distribution networks," in *2014 IEEE International Conference on Smart Grid Communications (SmartGridComm)*, pp. 392–397, Nov 2014.
- [110] S. Bubeck and N. Cesa-Bianchi, "Regret analysis of stochastic and nonstochastic multi-armed bandit problems," *Foundations and Trends® in Machine Learning*, vol. 5, no. 1, pp. 1–122, 2012.
- [111] R. Bonnefoi, C. Moy, and J. Palicot, "Improvement of the lpwan ami backhaul's latency thanks to reinforcement learning algorithms," *EURASIP Journal on Wireless Communications and Networking*, vol. 2018, p. 34, Feb 2018.
- [112] R. Bonnefoi, L. Besson, C. Moy, E. Kaufmann, and J. Palicot, "Multi-armed bandit learning in iot networks : Learning helps even in non-stationary settings," in *Cognitive Radio Oriented Wireless Networks* (P. Marques, A. Radwan, S. Mumtaz, D. Noguét, J. Rodriguez, and M. Gundlach, eds.), (Cham), pp. 173–185, Springer International Publishing, 2018.
- [113] C. Goursaud and J. M. Gorce, "Dedicated networks for iot : Phy / mac state of the art and challenges," *EAI Endorsed Transactions on Internet of Things*, vol. 15, 10 2015.
- [114] J. Arnbak and W. van Blitterswijk, "Capacity of slotted aloha in rayleigh-fading channels," *IEEE Journal on Selected Areas in Communications*, vol. 5, pp. 261–269, Feb 1987.
- [115] M. Vilgelm, M. Gürsu, S. Zoppi, and W. Kellerer, "Time slotted channel hopping for smart metering : Measurements and analysis of medium access," in *2016 IEEE International Conference on Smart Grid Communications (SmartGridComm)*, pp. 109–115, Nov 2016.

- [116] M. Anteur, V. Deslandes, N. Thomas, and A. L. Beylot, "Ultra narrow band technique for low power wide area communications," in *2015 IEEE Global Communications Conference (GLOBECOM)*, pp. 1–6, Dec 2015.
- [117] W. Zhan and L. Dai, "Throughput optimization for massive random access of m2m communications in lte networks," in *2017 IEEE International Conference on Communications (ICC)*, pp. 1–6, May 2017.
- [118] M. Z. Win, P. C. Pinto, and L. A. Shepp, "A mathematical theory of network interference and its applications," *Proceedings of the IEEE*, vol. 97, pp. 205–230, Feb 2009.
- [119] LoRa Alliance Technical committee, "LoRaWAN Regional Parameters," tech. rep., LoRa Alliance, Inc., July 2016.
- [120] ETSI, "ETSI - TR 102 313 - ELECTROMAGNETIC COMPATIBILITY AND RADIO SPECTRUM MATTERS (ERM); FREQUENCY-AGILE GENERIC SHORT RANGE DEVICES USING LISTEN-BEFORE-TRANSMIT (LBT); TECHNICAL REPORT," tech. rep., ETSI, July 2004.
- [121] L. Kleinrock and S. S. Lam, "Packet-switching in a slotted satellite channel," in *AFIPS National Computer Conference*, 1973.
- [122] X. Yang, A. Fapojuwo, and E. Egbogah, "Performance analysis and parameter optimization of random access backoff algorithm in lte," in *2012 IEEE Vehicular Technology Conference (VTC Fall)*, pp. 1–5, Sept 2012.
- [123] T. Lai and H. Robbins, "Asymptotically efficient adaptive allocation rules," *Adv. Appl. Math.*, vol. 6, pp. 4–22, Mar. 1985.
- [124] W. Jouini, D. Ernst, C. Moy, and J. Palicot, "Upper confidence bound based decision making strategies and dynamic spectrum access," in *2010 IEEE International Conference on Communications*, pp. 1–5, May 2010.
- [125] S. Maghsudi and E. Hossain, "Multi-armed bandits with application to 5g small cells," *IEEE Wireless Communications*, vol. 23, pp. 64–73, June 2016.
- [126] E. Kaufmann, O. Cappe, and A. Garivier, "On bayesian upper confidence bounds for bandit problems," in *Proceedings of the Fifteenth International Conference on Artificial Intelligence and Statistics* (N. D. Lawrence and M. Girolami, eds.), vol. 22 of *Proceedings of Machine Learning Research*, (La Palma, Canary Islands), pp. 592–600, PMLR, 21–23 Apr 2012.
- [127] P. Auer, N. Cesa-Bianchi, and P. Fischer, "Finite-time Analysis of the Multi-armed Bandit Problem," *Machine Learning*, vol. 47, no. 2, pp. 235–256, 2002.
- [128] W. R. Thompson, "On the likelihood that one unknown probability exceeds another in view of the evidence of two samples," *Biometrika*, vol. 25, 1933.
- [129] S. Agrawal and N. Goyal, "Analysis of Thompson sampling for the Multi-Armed Bandit problem," in *JMLR, Conference On Learning Theory*, pp. 39–1, 2012.

- [130] L. Melian-Gutierrez, N. Modi, C. Moy, F. Bader, I. Perez-Alvarez, and S. Zazo, "Hybrid ucb-hmm : A machine learning strategy for cognitive radio in hf band," *IEEE Transactions on Cognitive Communications and Networking*, vol. 1, pp. 347–358, Sept 2015.
- [131] K. Mikhaylov, . J. Petaejaevaervi, and T. Haenninen, "Analysis of capacity and scalability of the lora low power wide area network technology," in *European Wireless 2016 ; 22th European Wireless Conference*, pp. 1–6, May 2016.
- [132] Q. Bodinier, *Coexistence of Communication Systems Based on Enhanced Multi-Carrier Waveforms with Legacy OFDM Networks*. PhD thesis, Universite de Rennes 1, November 2017.
- [133] D. Luenberger, "Quasi-convex programming," *SIAM Journal on Applied Mathematics*, vol. 16, no. 5, pp. 1090–1095, 1968.
- [134] M. E. Yaari, "A note on separability and quasiconcavity," *Econometrica*, vol. 45, no. 5, pp. 1183–1186, 1977.
- [135] L. Besson and E. Kaufmann, "Multi-Player Bandits Revisited," in *Proceedings of Algorithmic Learning Theory* (F. Janoos, M. Mohri, and K. Sridharan, eds.), vol. 83 of *Proceedings of Machine Learning Research*, pp. 56–92, PMLR, 07–09 Apr 2018.
- [136] S. Sanayei and A. Nosratinia, "Antenna selection in MIMO systems," *IEEE Communications Magazine*, vol. 42, pp. 68–73, Oct 2004.
- [137] R. Bonnefoi, C. Moy, J. Palicot, and A. Nafkha, "Low-Complexity Antenna Selection for Minimizing the Power Consumption of a MIMO Base Station," *AICT 2018*, Jul 2018.
- [138] A. I. Pop, U. Raza, P. Kulkarni, and M. Sooriyabandara, "Does bidirectional traffic do more harm than good in lorawan based lpwa networks?," in *GLOBECOM 2017 - 2017 IEEE Global Communications Conference*, pp. 1–6, Dec 2017.
- [139] D. S. Bernstein, *Matrix Mathematics : Theory, Facts, and Formulas*. Princeton University Press, second ed., 2011.
- [140] "NIST Digital Library of Mathematical Functions." <http://dlmf.nist.gov/>, Release 1.0.17 of 2017-12-22. F. W. J. Olver, A. B. Olde Daalhuis, D. W. Lozier, B. I. Schneider, R. F. Boisvert, C. W. Clark, B. R. Miller and B. V. Saunders, eds.
- [141] R. Gallager, *Discrete Stochastic Processes*. The Springer International Series in Engineering and Computer Science, Springer US, 1995.

Titre : Utilisation de la Radio Intelligente pour un Réseau Mobile à Faible Consommation d'Energie

Mots clés : Radio Intelligente, Eco-radio, Réseaux Mobiles, Internet des Objets

Résumé : La réduction de l'empreinte carbone de l'activité humaine est aujourd'hui un enjeu économique et écologique majeur. Les réseaux de communication ont un double rôle à jouer dans cette réduction. En premier lieu, les réseaux mobiles, et en particulier les stations de base, sont un gros consommateur d'électricité. Il est donc nécessaire d'optimiser leur fonctionnement pour réduire leur empreinte carbone. Ensuite, des réseaux de communication sont désormais nécessaires pour mieux gérer la production d'électricité et ainsi pouvoir augmenter la proportion d'électricité produite par des sources d'énergie renouvelables.

Nous commençons par proposer une solution pour réduire la consommation d'énergie des réseaux mobiles. Pour cela, nous proposons des algorithmes pour optimiser l'allocation de puissance lorsque des mécanismes de mise en veille dynamique sont utilisés.

Dans un second temps, nous proposons une solution pour améliorer le fonctionnement des réseaux d'objets connectés utilisés pour la gestion de l'électricité. Plus précisément, nous rendons plus fiables ces communications grâce à l'utilisation d'algorithmes de bandit multibras pour l'accès fréquentiel.

Dans cette thèse, nous regardons ces deux aspects.

Title : Cognitive Radio for Green Cellular Networks

Keywords: Cognitive Radio, Green Radio, Mobile Networks, Internet of Things

Abstract: The reduction of the carbon footprint of human activities is one of the current major economic and ecological challenges. Communication networks have a dual role in this reduction. On one hand, mobile networks, and in particular the base stations, are nowadays an important energy consumer. It is, thus, necessary to optimize their behavior in order to reduce their carbon footprint. On the other hand, some communication networks are necessary to better manage the electrical grid. Thanks to this better management, it is possible to improve the proportion of electricity produced by renewable energy sources.

In this thesis, we look at both aspects. In a first step, we propose a solution to reduce the energy consumption of wireless mobile networks. For that purpose, we propose algorithms that optimize the power allocation when Cell Discontinuous Transmission is used by the base stations.

In a second step, we propose a solution in order to improve the performance of Internet of Things networks used for the electrical grid. More precisely, we use multi-armed bandit algorithm for channel selection in IoT networks as a means of increasing the reliability of communications.