

N° d'ordre 3725

THESE

En vue de l'obtention du : **DOCTORAT**

Structure de Recherche : Laboratoire de recherche en informatique et télécommunications (LRIT).

Discipline : Informatique.

Spécialité : Réseaux et télécommunications.

Présentée et soutenue le 17/12/2022
par :

EL METTITI Abderrahmane

Contributions à l'optimisation des ressources spectrales dans la 6G

JURY

Khalid MINAOUI	PES, UM5, Faculté des sciences, Rabat	Président
Jamal MHAMDI	PES, UM5, Faculté des sciences, Rabat	Rapporteur
Abdessamad EL RHARRAS	PH, Ecole hassania des ingénieurs, Casablanca	Rapporteur
Hicham LAANAYA	PH, UM5, Faculté des sciences, Rabat	Rapporteur
Rachid SAADANE	PES, Ecole hassania des ingénieurs, Casablanca	Examineur
Ouadoudi ZYTOUNE	PES, Uit, ENSA kénitra	Examineur
Mohammed OUMSIS	PES, UM5, Ecole supérieure de technologies, Salé	Directeur

Année Universitaire : 2022/2023

Dédicaces

Je dédie cette thèse,

À ma très chère mère, source de tendresse, de patience et de sacrifice. Tes prières et ta bénédiction sont d'un grand secours le long de ma vie. Quoi que je puisse dire et écrire, je ne pourrai jamais suffisamment exprimer ma grande affection et ma profonde gratitude. Puisse dieu tout-puissant, te préserver et t'accorder santé, longue vie et bonheur.

À mon très cher père, aucune dédicace ne saura exprimer mon respect, ma reconnaissance et mon profond amour. De tous les pères, tu es le meilleur. Tu as été et tu seras toujours un exemple pour moi, de par tes qualités humaines, ta persévérance et ton perfectionnisme. J'espère que tu trouveras dans ce travail l'aboutissement de tant d'années de sacrifices, de sollicitudes, d'encouragement et de prière. En ce jour, un de tes rêves deviendra – je l'espère – réalité et tu pourras en être fier. Puisse dieu te préserver et te procurer santé et bonheur.

À ma très chère femme. Depuis que je te connais, tu ne cesses de me soutenir et de m'épauler. Tu me veux toujours « le meilleur ». À travers ton amour, tu sais me procurer confiance et stabilité. Dans le meilleur comme dans le pire, tu es toujours à mes côtés. Aucun mot ne pourra exprimer ma gratitude, mon amour et mon respect. Puisse Dieu – qui a croisé nos chemins – nous procurer santé et longue vie.

À mes sœurs Nabila, Soumaya et Mariam, je vous dédie ce travail en témoignage de ma profonde affection et de notre indéfectible union. Puisse dieu vous protéger, garder et renforcer notre fraternité.

À tous les membres de ma famille, que dieu vous protège tous.

À mes amis que je ne peux passer sous silence, Abdessamad, Rachid, Zakaria, Hassan, Loutfi, Mohammed qui ont fait en sorte que ma thèse puisse prendre forme dans les conditions les plus favorables. Que toute personne ayant participé de près ou de loin au déroulement de ce travail trouve ici l'expression de mes sentiments les meilleurs.

À mon petit bout de chou qui vient de naître. J'espère qu'il viendra le jour où tu écriras de tes propres mains des lignes comme celles qu'écrit ici ton père qui t'aime. Que dieu te bénisse et veille sur toi ta vie entière

Remerciements

C'est grâce à l'aide de nombreuses personnes que j'ai pu mener cette thèse à son terme et il m'est très difficile de citer chacun pour dire ma gratitude.

Je voudrais tout d'abord remercier grandement mon directeur de thèse, Monsieur **Mohammed Oumsis**, professeur de l'enseignement supérieur à l'université Mohammed V de Rabat. Il m'a encadré tout au long de cette thèse et il a partagé ses brillantes intuitions avec moi. Qu'il soit aussi remercié pour sa gentillesse, ses conseils, sa permanente disponibilité et pour les nombreux encouragements qu'il m'a prodigués. Je suis très honoré de sa présence à mon jury de thèse et je tiens à le remercier vivement !

Monsieur **Khalid Minaoui**, directeur du Laboratoire de recherche en informatique et télécommunications – LRIT, et professeur de l'enseignement supérieur à l'université Mohammed V de Rabat, merci pour l'honneur qu'il me fait en acceptant d'être membre de mon jury de thèse. Je tiens à l'assurer de ma profonde reconnaissance pour l'intérêt qu'il porte à ce travail. C'est grâce à lui que j'ai pu concilier avec bonheur recherche théorique et recherche appliquée pendant cette thèse.

Monsieur **Rachid Saadane**, professeur de l'enseignement supérieur de l'École Hassania des Travaux Publics de Casablanca, qui m'a encadré et orienté dans la première année de cette thèse. Nous avons partagé les premières réflexions de cette recherche, il m'a aidé à créer le premier cadre conceptuel de cette thèse. Je le remercie pour ses conseils toujours très pertinents qui m'ont aidé à trouver mes marques dans la recherche scientifique et à bien débiter les premières années de ma formation doctorale.

Merci à Monsieur **Abdessamad El Rharras**, professeur de l'enseignement supérieur de l'École Hassania des Travaux Publics de Casablanca pour sa participation à mon jury de thèse en qualité de rapporteur de mon travail et pour toutes les remarques intéressantes qu'il m'a faites pendant la pré-soutenance de cette thèse.

Monsieur **Hicham Laanaya**, professeur de l'enseignement supérieur à l'université Mohammed V de Rabat, Monsieur **Jamal Mhamdi**, professeur de l'enseignement supérieur à l'université Mohammed V de Rabat et Monsieur **Ouadoudi Zytoune**, professeur de l'enseignement supérieur à l'École nationale des Sciences Appliquées de Kénitra, merci à eux pour l'honneur qu'ils me font en acceptant d'être rapporteur ou examinateur de cette thèse.

Mes vifs remerciements et sincères reconnaissances vont également à toutes les personnes qui ont porté un regard critique à mon travail, et à ceux qui m'ont accordé leurs précieux conseils et inconditionnel soutien pendant la réalisation de cette recherche, je pense à Madame Hanane, à messieurs Karim, Mouhi, Eddine.

Résumé

Notre travail propose des techniques contribuant à l'amélioration des services du futur réseau 6G. Au moyen d'un auto-encodeur empilé pour compresser les jeux de données et de perceptrons multicouches, nous proposons une technique d'apprentissage profond pour la prédiction de formation de faisceaux analogiques à ondes millimétriques. Nous présentons également un système hybride pour la gestion dynamique du spectre dans les environnements radio cognitive à l'aide des surfaces intelligentes reconfigurables pour un accès optimisé aux ressources spectrales. Notre démonstration repose sur un corpus de descriptions allant de la première génération à la 5G, complété par les apports de la 5G nouvelle radio, qui explique les limites technologiques désormais atteintes et que la 6G doit dépasser. Autant de défis identifiés que nous exposons en détail concernant différents domaines dans une vaste vision structurelle de la 6G ; s'il reste de nombreux défis, des avancées significatives sont discutées ici. Nos travaux ont fait l'objet de trois publications qui ont fait part de nos résultats encourageants.

Mots clefs : Machine learning, Deep learning, Auto-encodeur, SIR, GDS, RNA, Perceptron multicouche

Abstract

Our work proposes techniques that contribute to the improvement of services in the future 6G network. Using a stacked auto-encoder to compress datasets and multi-layer perceptrons, we propose a deep learning technique for the prediction of millimeter wave analog beamforms. We also present a hybrid system for dynamic spectrum management in cognitive radio environments using reconfigurable smart surfaces for optimized access to spectrum resources. Our demonstration is based on a corpus of descriptions from the first generation to 5G, completed by the contributions of the new radio 5G, which explains the technological limits now reached and that 6G must overcome. As many identified challenges that we detail concerning different domains in a broad structural vision of 6G; while many challenges remain, significant advances are discussed here. Our work has been published in three papers, and our results are encouraging.

Keywords: Machine learning, Deep learning, Auto-encoder, SIR, GDS, RNA, Multilayer perceptron

Liste des figures

Figure 1 : Architecture d'un réseau 1G (AMPS)	8
Figure 2 : Architecture d'un réseau 2G (GSM)	9
Figure 3 : Architecture d'un réseau 3G (UMTS)	10
Figure 4 : Architecture d'un réseau 4G (LTE)	12
Figure 5 : Évolution des débits selon les technologies cellulaires utilisées	13
Figure 6 : Architecture d'un réseau d'accès 5G (5G NG-RAN)	15
Figure 7 : Nouveaux services pris en charge par la technologie 5G	17
Figure 8 : Découpage du réseau 5G en tranches (<i>network slicing</i>)	20
Figure 9 : Architecture globale du réseau d'accès radio 5G RAN	28
Figure 10 : Nœud de connexion NG-RAN	29
Figure 11 : Options de découpage CU/DU ^[6]	29
Figure 12 : Transition des états RRC en 5G NR	31
Figure 13 : Structure d'une trame 5G NR ($\mu=0$)	33
Figure 14 : Structure d'une trame 5G NR	34
Figure 15 : Identification du <i>minislot</i>	34
Figure 16 : Duplexage FDD et TDD	35
Figure 17 : Exemple de latence d'un système de l'internet tactile ^[22]	37
Figure 18 : Récupération de l'AS et du NAS après l'échec du $k^{\text{ième}}$ <i>handover</i>	39
Figure 19 : Latence du plan utilisateur pour 2 symboles OFDM par <i>slot</i>	42
Figure 20 : Latence du plan utilisateur pour 4 symboles OFDM par <i>slot</i>	42
Figure 21 : Latence du plan utilisateur pour 7 symboles OFDM par <i>slot</i>	43
Figure 22 : Latence du plan utilisateur pour 14 symboles OFDM par <i>slot</i>	43
Figure 23 : Exigence des réseaux 6G	50
Figure 24 : Architecture des réseaux non terrestres	51
Figure 25 : Cartographie entre les applications 6G, les exigences et les technologies clés potentielles	70
Figure 26 : Relation entre IA, ML et DL	80
Figure 27 : Neurone biologique	81
Figure 28 : Neurone formel	82
Figure 29 : Modèle mathématique du neurone formel	82
Figure 30 : Perceptron multicouche	83
Figure 31 : Auto-encodeur pour le débruitage d'images	85
Figure 32 : Fonction d'un auto-encodeur convolutif	85
Figure 33 : Schéma de principe d'un auto-encodeur	86
Figure 34 : Modélisation de la compression de données d'un auto-encodeur sous-complet	88
Figure 35 : Schéma de principe d'un auto-encodeur épars	89
Figure 36 : Méthodes de réduction de la dimensionnalité	90
Figure 37 : Modèle de la rétropropagation	92
Figure 38 : Modèle d'un auto-encodeur débruiteur	92
Figure 39 : Modèle de l'auto-encodeur variationnel	93
Figure 40 : Construction des attributs latents par l'auto-encodeur variationnel	93
Figure 41 : Modèle de la reparamétrisation	94
Figure 42 : Échantillon de chiffres MNIST généré par l'entraînement d'un auto-encodeur variationnel	95

Figure 43 : Architecture de formation de faisceau analogique du récepteur et de l'émetteur (RF : radiofréquence, DAC : convertisseurs numérique-analogique, ADC : convertisseur analogique-numérique)	99
Figure 44 : Vue plongeante et vue de dessus du scénario O1 ^[26]	101
Figure 45 : Architecture des perceptrons multicouches	102
Figure 46 : Architecture de l'auto-encodeur	102
Figure 47 : Notre proposition d'entraînement de faisceau utilisant un auto-encodeur empilé et des perceptrons multicouches	103
Figure 48 : Architecture empilée de l'auto-encodeur proposé	105
Figure 49 : Historique du montage du PMC	107
Figure 50 : Historique d'ajustement de l'auto-encodeur empilé (pour l'extraction de caractéristiques + PMC)	108
Figure 51 : Erreur quadratique moyenne de notre modèle par rapport au modèle PMC uniquement	108
Figure 52 : Cas typique d'utilisation d'une SIR qui reçoit un signal de l'émetteur et le réémet en le focalisant vers le récepteur ; pour focaliser le faisceau dans la bonne direction, la SIR doit être correctement configurée ^[182]	111
Figure 53 : Illustration de la communication SISO considérée avec NLOS	113
Figure 54 : Architecture typique de SIR ^[185]	115
Figure 55 : Transmission supportée par les SIR	115
Figure 56 : Transmission supportée par le relais DF	117
Figure 57 : Architecture du modèle proposé	120
Figure 58 : Schéma d'allocation hybride du modèle proposé	120
Figure 59 : Utilisation du spectre dans le modèle OSA ^[48]	123
Figure 60 : Illustration du modèle de partage du spectre de la CSA	125
Figure 61 : Principales contraintes de puissance pour la protection des services UP	126
Figure 62 : Modèle de simulation proposé	128
Figure 63 : Gains typiques du canal en fonction de la distance incluant plusieurs gains d'antenne	129
Figure 64 : Puissance d'émission nécessaire pour atteindre le débit $R = 4$ bit/s/Hz en fonction de la distance d_{sr}	130
Figure 65 : Puissance d'émission nécessaire pour atteindre le débit $R = 8$ bit/s/Hz en fonction de la distance d_{sr}	130
Figure 66 : Puissance d'émission requise en fonction du taux réalisable R lorsque $d_1 = 100$ m	130
Figure 67 : Efficacité énergétique en fonction de la distance d_1	131
Figure 68 : Efficacité énergétique en fonction de la distance d_1 lorsque N est optimisé	132
Figure 69 : Efficacité énergétique en fonction du débit réalisable lorsque N est optimisé	133
Figure 70 : Courbes ROC pour les modèles SISO et SIR avec différentes configurations	134
Figure 71 : Courbes ROC pour les modèles SISO et SIR avec différentes configurations SIR lorsque $P_t = 0,001$ W	134
Figure 72 : Courbes ROC des modèles SISO, SIR et relais DF lorsque $N = 100$ et $P_t = 0,001$ W	135
Figure 73 : Diagramme du résultat de l'opération de détection du spectre	136
Figure 74 : Diagramme d'allocation initiale pour les utilisateurs secondaires dans chaque canal	136
Figure 75 : Diagramme de résultat de l'allocation des canaux par le modèle OSA	137
Figure 76 : Diagramme de résultat de l'allocation des canaux basée sur le modèle hybride OSA-CSA	137
Figure 77 : Mobilité des US après la mise en œuvre de l'algorithme hybride OSA-CSA proposé	138
Figure 78 : Efficacité énergétique totale dans chaque canal alloué quand $\bar{R} = 10$ bits / s / Hz et $N_c = 20$	139
Figure 79 : Efficacité énergétique totale pour chaque canal dans le cas où $\bar{R} = 14$ bits / s / Hz et $N_c = 20$	139
Figure 80 : Efficacité énergétique totale pour chaque canal dans le cas où $\bar{R} = 10$ bits / s / Hz et $N_c = 50$	139
Figure 81 : Nouvelle actualisation de l'état du spectre	140
Figure 82 : Allocation du spectre avec le modèle hybride OSA-CSA	140
Figure 83 : Diagramme de la mobilité des US après l'exécution de l'algorithme lorsque la contrainte d'interférence moyenne est adoptée	141

Liste des tableaux

Tableau 1 : Comparaison entre technologie 5G et technologie 4G (LTE-Advanced) ^[4]	18
Tableau 2 : Comparaison entre les performances NB-IoT et LTE-M	21
Tableau 3 : Classes de mobilité	25
Tableau 4 : Débits de données de liaison de canal de trafic normalisé par bande passante	25
Tableau 5 : Structure de la trame temporelle	35
Tableau 6 : Technologies clés potentielles des réseaux 6G	68
Tableau 7 : Fonctions d'activation usuelles	83
Tableau 8 : Variables hyperparamètres utilisées pour le modèle PMC	106

Table des matières

Résumé, <i>abstract</i>	III
<i>abstract</i>	II
Liste des figures	V
Liste des tableaux	VII
♦ Introduction générale	5
Chapitre I ♦ Évolution des réseaux mobiles de la première génération à la 5G	7
Introduction	7
1. Différentes générations de réseaux mobiles	8
1.1 Réseaux mobiles de première génération (1G)	8
1.2 Réseaux mobiles de deuxième génération (2G)	8
1.3 Réseaux mobiles de troisième génération (3G)	10
1.4 Réseaux mobiles de quatrième génération (4G)	11
1.5 Réseaux mobiles de cinquième génération (5G)	14
2. Catégories d'utilisation dans un réseau 5G	16
3. Technologies mises en place par la 5G	18
3.1 Architecture radio	18
3.1.1 Techniques de transmission optimisées	18
3.1.2 Nouvelles technologies d'antennes	18
3.2 Architecture réseau	19
4. 5G pour l'internet des objets (IoT)	20
5. Exigences minimales des performances techniques de l'interface radio 5G NR	21
5.1 Débit de données maximal	22
5.2 Latence	22
5.2.1 Latence du plan utilisateur	22
5.2.2 Latence du plan contrôle	23
5.2.3 Densité de connexion	23
5.2.4 Efficacité énergétique	23
5.2.5 Fiabilité	24
5.2.6 Mobilité	24
5.2.7 Temps d'interruption de mobilité (MIT, <i>mobility interruption time</i>)	25
5.2.8 Bande passante	25
Conclusion	26
Chapitre II ♦ Cinquième génération nouvelle radio de communication mobile (5G NR)	27
Introduction	27
1. Caractéristiques globales du réseau d'accès radio 5G (NG-RAN)	28
1.1 Cartographie techniques globale du Réseau 5G RAN (schéma technique général du réseau 5G)	28
1.1.1 Architecture de la station de base gNB	28
1.1.2 Canaux physiques	30
1.2.3 Contrôle des ressources radio dans 5G NR	30
1.2.4 Formes d'ondes pour la 5G NR	31
1.1.5 Techniques d'accès au réseau 5G	32
1.1.6 Numérologie pour la 5G NR	32
1.1.4 Spectres de fréquences pour 5G NR	36
1.2 Technologie 5G pour l'internet tactile	37
1.2.1 Internet tactile et la question du débit maximum réalisable	38
1.2.2 L'Internet tactile et la fiabilité maximale	38



1.3 Optimisation de la latence pour les applications uRLLC	41
1.3.1 Cas d'une première transmission réussie	41
1.3.2 Cas de demande de retransmission HARQ	41
Conclusion	44
Chapitre III ♦ Réseaux 6G : vision, exigences, architecture, technologies et enjeux	45
Introduction	45
1. État de l'art	45
2. Vision, exigences et scénarios d'application	47
2.1 Vision et exigences	47
2.2 Scénarios d'application	48
3. Architecture et défis du réseau 6G	50
3.1 Défis 6G	50
3.1.1 Réseau intégré espace-air-sol-mer	50
3.1.2 Réseaux intelligents	51
3.2 Architecture pour 6G	52
3.2.1 Architecture cyber-jumelle	52
3.2.2 Architecture RAN découplée	52
3.2.3 Architecture généralisée pour la 6G	53
3.2.4 Architecture de haut niveau de la 6G	53
3.2.5 Architecture multiniveau pour la 6G	53
3.2.6 Architecture 6G tridimensionnelle ^[4]	54
3.2.7 Architecture 6G basée sur les neurosciences	54
3.2.8 Architecture proposée	54
3.3 Tendances futures de la 6G	55
3.3.1 Localisation en 6G	55
3.3.2 6G de confiance	56
3.2.3 6G pour les objectifs de développement durable (ODD)	57
3.2.4 Comment la 6G réduit-elle la fracture numérique ?	59
3.2.5 Quelle est la portée de l'intelligence de bord en 6G ?	60
3.2.6 Comment l'apprentissage automatique régit la 6G ?	62
3.2.6.1 ML au MAC	63
3.2.6.2 ML à la couche d'application 6G	64
3.2.7 Modèle commercial pour la 6G	64
3.2.7.1 Tendances clés et incertitudes	65
3.2.7.2 Scénarios	65
3.2.8 Tendances supplémentaires	66
3.2.8.1 Accès sans cellulaire	66
3.2.8.2 Communication par rétrodiffusion (BSC)	66
3.2.8.3 Transfert de puissance sans fil	67
4. Activer les technologies clés	67
4.1 Nouvelle technologie de ressources spectrales	70
4.1.1 Communication térahertz	70
4.1.2 Communication par lumière visible (VLC)	71
4.1.3 Technologie de partage dynamique du spectre	71
4.2 Technologie d'accès sans fil efficace	71
4.2.1 Amélioration de la technologie traditionnelle de la couche physique	71
4.2.1.1 Nouvelle modulation de codage	71
4.2.1.2 Multiantenne à grande échelle	72
4.2.2 Moment angulaire orbital (OAM)	72
4.3 Technologie de base innovante	72
4.3.1 Chaîne de blocs	72
4.3.2 Nouveaux matériaux	72
4.3.3 Gestion de l'énergie	73
4.4 Nouvelle technologie de communication	73
4.4.1 Communication et informatique quantique	73
4.4.2 Communication moléculaire	73
5. Activités de recherche	73
5.1 Union européenne	74
5.2 Finlande	74
5.3 États-Unis	74
5.4 Japon et Corée du Sud	75



5.5 Chine	75
Conclusion	75
Chapitre IV ♦ Structure technique d'un réseau 6G exigeant	77
Introduction	77
1. Intelligence artificielle	78
1.1 Apprentissage automatique	80
1.2 Types d'apprentissage	80
2. Réseaux de neurones artificiels	81
2.1 Neurone biologique	81
2.2 Neurone artificiel (ou neurone formel)	82
2.3 Réseaux de neurones multicouches	83
2.4 Apprentissage	84
3. Auto-encodeurs	84
3.1 Qu'est-ce qu'un auto-encodeur ?	85
3.2 Architecture d'un auto-encodeur	85
3.3 Comment entraîner un auto-encodeur ?	87
3.4 5 types d'auto-encodeur	87
3.4.1 Auto-encodeur sous-complet	88
3.4.2 Auto-encodeur épars	89
3.4.3 Auto-encodeur contractif	90
3.4.4 Auto-encodeur débruiteur	92
3.4.5 Auto-encodeur variationnel	93
3.5 Applications d'un auto-encodeur	96
3.5.1 Réduction de la dimensionnalité	96
3.5.2 Débruitage d'image	96
3.5.3 Génération de données d'image et de séries temporelles	96
3.5.4 Détection d'anomalies	96
3.6 Synthèse	97
4. Travaux connexes	97
5. Préliminaires et présentation du jeu de données	98
5.1 Prédiction de la formation de faisceaux à ondes millimétriques	98
5.2 Présentation du jeu de données	100
5.3 Modèle étudié de <i>deep learning</i>	101
6. Approche proposée	103
6.1 Acquisition et génération de jeux de données	103
6.2 Extraction de caractéristiques	104
6.3 Processus de formation	105
7. Résultats expérimentaux	106
7.1 Indicateur de performances	106
7.2 Discussion des résultats	107
Conclusion	109
Chapitre V ♦ Communication sans fil soutenue par surfaces intelligentes reconfigurables	110
Introduction	110
1. Surfaces intelligentes reconfigurables pour communications sans fil	112
1.1 Modèle du système étudié	112
1.2 Transmission supportée par SIR	114
1.3 Communication sans fil assistée par relais	117
1.4 SIR pour une gestion dynamique du spectre	119
1.4.1 Modèle du système proposé	119
1.4.2 Modèle de détection du spectre	121
2. Modèle hybride pour une gestion dynamique du spectre	123
3. Résultats numériques	128
3.1 Comparaison numérique des performances	128
3.2 Réseau radio cognitif supporté par les surfaces intelligentes reconfigurables	133
4. Implémentation de l'algorithme hybride GDS et résultats	136
Conclusion	141
♦ Conclusion générale	142
Bibliographie	145



◆ Introduction générale

La cinquième génération de communication sans fil apporte une optimisation de puissance indéniable comparée aux générations précédentes. Pour autant, ses avancées technologiques et ses déclinaisons, dont la 5G NR que nous regardons de près dans ces pages, n'absorbe pas l'entièreté des besoins sinon actuels, assurément futurs. Car comme nous le verrons, les idées ne manquent pas concernant des applications susceptibles de fonctionner sur un réseau non seulement à la fois dense et épars, tridimensionnel, ultra-sécurisé mais par-dessus tout : offrant un volume de débit de données jamais connu à ce jour.

Pour le commun des mortels, ce jour aura lieu en 2030, soit après-demain. À l'échelle de la R&D il est déjà advenu, les travaux sont en cours comme en témoigne la littérature. En effet, dans différents domaines, la recherche s'emploie aujourd'hui à construire un réseau d'une fiabilité et d'une puissance nettement accrues, c'est pourquoi nous évoquons ici au présent de l'indicatif tous les aspects de la prochaine et sixième génération de communication sans fil, appelée sans surprise, 6G.

Quelles sont les technologies qui sont à même de soutenir une évolution telle qu'elle appelle un changement de paradigme dans la conception, à de nombreuses échelles, d'un réseau qui dépend du spectre de radio fréquence ?

Les limites de ce dernier motivent toutes les recherches pour exploiter au mieux cette ressource dorénavant difficilement dépassable sans des méthodes radicalement innovantes dans plusieurs domaines. Cette évolution, repose en bonne partie sur la désormais incontournable intelligence artificielle sur laquelle se bâtit en autres, notamment les surfaces intelligentes reconfigurables, un réseau de communication sans fil de sixième génération que d'aucuns prédisent spectaculaire tant il offrira des champs d'application bien plus larges que la 5G ne permet.

Dans le premier chapitre, nous retraçons les grandes étapes techniques de la première à la cinquième génération, notamment les apports de cette dernière sur les problèmes de latence. Quand bien même la 6G doit dépasser, et de loin, les améliorations substantielles de la 5G telles que la 5G NR, toutes ces précédentes versions sont le socle structurel de la technologie 6G qui se bâtit sur l'existant, normé par le 3GPP.



Nous présentons dans le deuxième chapitre la nouvelle interface radio conçue pour la technologie 5G (5G NR) et ses diverses caractéristiques (numérologie, modulation, spectre de fréquences, etc.). Nous examinons le problème de la latence aller-retour qui doit égaliser 1 ms ainsi que les défis qu'il reste à relever pour proposer des services très exigeants en termes de fiabilité, de sécurité et de débit de données.

Aussi, nous en venons dans le troisième chapitre au cœur de la 6G et considérons de près les réponses qu'elle doit apporter. Nous évoquons les enjeux et les défis que la technologie aide à surmonter pour obtenir un réseau sol-espace-mer intègre, à travers ce qui compose son architecture technique généralisée et multiniveau. Sans omettre son empreinte globale sur nos sociétés et plus spécifiquement avec l'internet des objets et l'internet tactile. Nous concluons ce chapitre en passant en revue l'encours des travaux dans diverses parties du monde.

Dans le quatrième chapitre nous mettons en exergue quelques technologies clés à même de propulser la 6G à un niveau d'exigence jamais atteint. L'intelligence artificielle occupe la première partie pour évoquer entre autres les réseaux de neurones artificiels. Nous parlons aussi des travaux en cours et de procédés prometteurs comme le perceptron multicouche et plus particulièrement des auto-encodeurs sur lesquels nous nous attardons plus en détail. En effet, l'apprentissage automatique (ou *deep learning* en anglais) est un procédé particulièrement performant pour entraîner un dispositif à former le meilleur faisceau et répondre ainsi au problème de perte de puissance dans les ondes millimétriques. Ici, nous proposons une approche peu coûteuse dont l'efficacité est comparée à la méthode de référence.

Les surfaces intelligentes reconfigurables (SIR) méritent une attention particulière de par leur apport significatif aux futures performances de la 6G en termes de fiabilité et de couverture de territoires (terrestres et non terrestres). Nous consacrons le cinquième chapitre aux SIR, d'abord pour en définir le contour technique puis pour évaluer leur performance et comprendre la formation de faisceau, la gestion des ressources ainsi que les applications de l'apprentissage automatique aux réseaux sans fil soutenus par SIR. Puis, en proposant un dispositif qui compare une transmission assistée par SIR avec un relais de type décode-et-transfert (relais DF). À cette fin, nous optimisons les deux technologies en calculant les puissances d'émission optimales et le nombre optimal d'éléments dans un SIR.

Dans notre conclusion nous faisons la synthèse de nos travaux et des résultats obtenus. À partir de ceux-ci, nous présentons aussi quelques perspectives de recherches.

Chapitre I ♦ Évolution des réseaux mobiles de la première génération à la 5G

Introduction

L'histoire des réseaux mobiles est divisée à ce jour en cinq phases principales, souvent appelées générations, par conséquent, ils sont désignés par 1G, 2G, 3G, 4G et 5G. La principale différence entre ces cinq générations est la technologie mise en œuvre pour accéder aux ressources radioélectriques. Se prépare aujourd'hui l'arrivée de la prochaine génération, soit la 6G, qui, quant à elle et la technologie qui l'accompagne, fera franchir tant à la société qu'aux professionnels une étape majeure. Le développement de ces technologies est guidé par la volonté d'augmenter la capacité et les débits que le système peut fournir dans des bandes de fréquences restreintes. En effet, les fréquences sont des ressources qui deviennent de plus en plus rare car elles sont fortement sollicitées pour un grand nombre d'applications différentes, notamment la télévision, la radio, les liaisons satellites, les faisceaux hertziens, les communications militaires et les réseaux privés.

Dans tous les pays, dans tous les milieux, le spectre de fréquences disponible est limité. Par ailleurs, le développement des réseaux mobiles est toujours conditionné à la capacité d'optimiser les ressources spectrales disponibles. À l'origine, la capacité d'un réseau de télécommunication mobile était définie comme le nombre maximal d'appels téléphoniques simultanés pouvant être effectués sous la couverture d'une même cellule. Aujourd'hui, à mesure que l'utilisation des services de données évolue, la capacité du réseau dépend davantage du nombre d'utilisateurs pouvant se connecter au service de données en même temps et de la vitesse moyenne de chaque utilisateur lors d'une session de transfert de données. Plus généralement, la capacité du réseau peut être représentée par le débit total maximal que peut transporter une cellule fortement chargée.

L'objectif de ce premier chapitre est de mettre en exergue les différentes générations et de passer en revue les évolutions majeures apportées par chacune des générations.

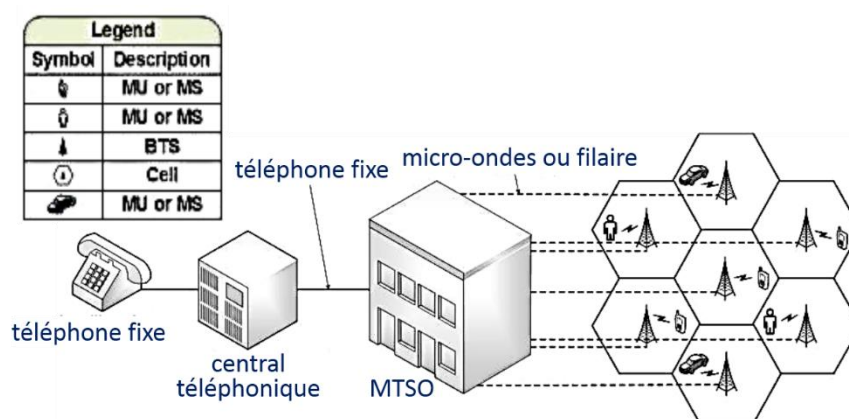


1. Différentes générations de réseaux mobiles

1.1 Réseaux mobiles de première génération (1G)

La première génération de réseaux mobile (1G) apparaît au cours des années 1980, elle est caractérisée par une multitude de technologies introduites en parallèle à travers le monde. La raison d'être de cette première génération est surtout d'offrir un service de téléphonie en situation de mobilité (figure 1). Cependant, ces technologies ne parviennent pas à franchir les frontières de leurs pays d'origine et à s'imposer au niveau international. Cela les rend incompatibles entre elles, non interopérables et par conséquent avec l'impossibilité d'offrir une itinérance internationale. Ce handicap enjoint les pays concernés à réfléchir sur une norme internationale en matière de téléphonie mobile.

Figure 1 : Architecture d'un réseau 1G (AMPS)

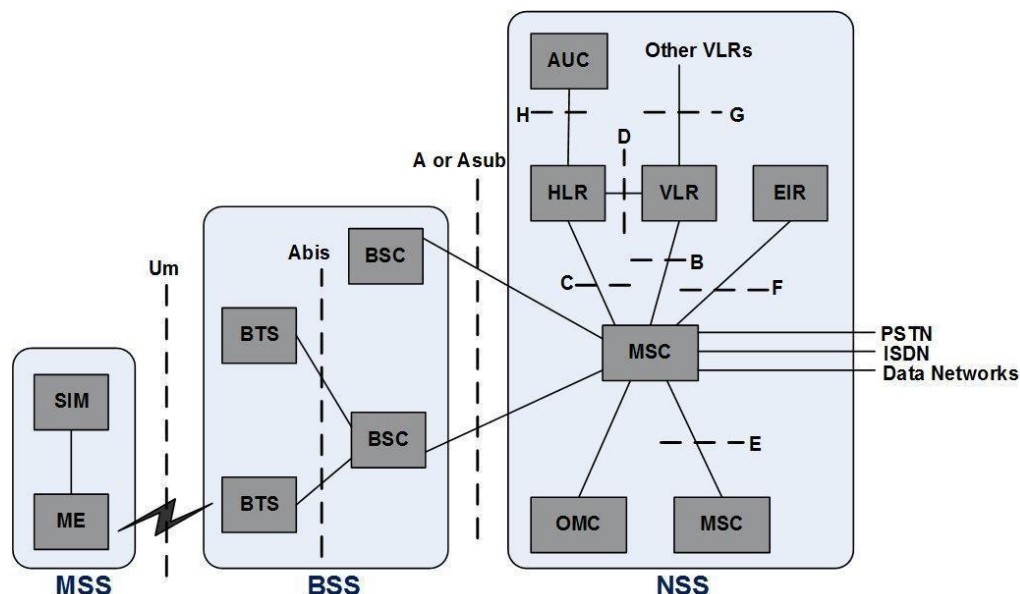


D'un point de vue technique, ces systèmes sont basés sur un codage et une modulation de type analogique. Le mode d'accès utilisé étant l'accès multiple par répartition en fréquences (FDMA, *frequency division multiple access*) il consiste à associer à chaque utilisateur une fréquence. La capacité de ces systèmes est très faible, elle permet seulement quelques appels-voix simultanés par cellule.

1.2 Réseaux mobiles de deuxième génération (2G)

Le réseau cellulaire de deuxième génération (2G), qui marque la transition des systèmes 1G analogiques vers les systèmes numériques, se caractérise également par le nombre de systèmes définis et déployés dans le monde (figure 2). Il existe le GSM (*global system for mobile communications*) en Europe, le PDC (*personal digital communication*) au Japon et l'IS-95 (Interim Standard 95) ou CDMAOne aux États-Unis. Dans une première version, ces systèmes permettent l'accès aux services vocaux en mobilité, ils permettent également l'accès aux SMS (*short message service*). Ces systèmes sont également capables de transmettre des données dans des modes à commutation de circuit à faible débit. De plus, la taille des équipements du terminal mobile est réduite par rapport à celle de la 1G, ce qui rend la mobilité encore plus confortable.

Figure 2 : Architecture d'un réseau 2G (GSM)



Il convient également de noter que les systèmes GSM sont les plus acceptés au niveau international. De nombreux pays décident alors de déployer à grande échelle cette technologie qui permet le *roaming* international. Il est standardisé au niveau européen au début des années 1990, il permet de réduire les coûts de production en tenant compte des économies d'échelle dues à la taille du marché. La réduction des coûts et la compatibilité incitent la plupart des pays à adopter cette technologie.

Les systèmes 2G utilisent tous des codages et des modulations numériques. Cela permet une augmentation significative de la capacité des réseaux de téléphonie mobile cellulaire. De plus, des techniques d'accès multiple plus complexes que FDMA sont employées. Les systèmes GSM et PDC utilisent la méthode d'accès FDMA combinée à une méthode de division temporelle TDMA (*time division multiple access*). Par ailleurs, et côté duplex du canal, ces systèmes utilisent un mode FDD (mode duplex FDD) qui utilise des fréquences différentes pour l'émission et la réception. D'autre part, le système IS-95 utilise une méthode d'accès basée sur une méthode de division par code appelée CDMA (*code division multiple access*).

Cependant, les systèmes 2G ont certaines limites, notamment sur leurs capacités. En effet, malgré la forte densité des réseaux, les appels sont souvent rejetés aux heures chargées. Une autre limitation est liée à la vitesse de transmission des données sur ce réseau car le cœur de réseau utilise la commutation de circuits. À cette fin, et pour surmonter ce problème, le réseau central de transmission par paquets appelé GPRS (*general packet radio service*) s'appuie sur les réseaux GSM pour prendre en charge les services de données et utiliser le multiplexage statistique. Ensuite, la technologie d'accès EDGE (*enhanced data rate for GSM evolution*), une évolution du GSM, développe le débit par cellule en améliorant la technologie d'accès aux canaux radio, en adoptant des schémas de modulation et de codage particulièrement améliorés. Cependant, à la fin des années 1990, le débit offert par le réseau 2G n'est pas encore suffisant pour permettre l'accès à des services de données plus évolués.

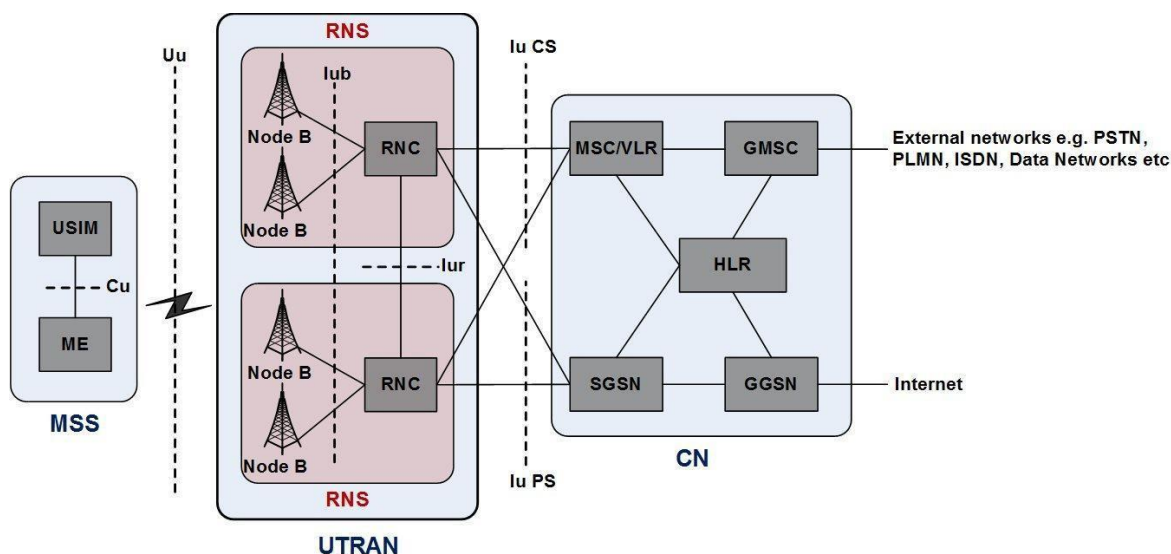
1.3 Réseaux mobiles de troisième génération (3G)

Le réseau cellulaire de troisième génération (3G) est une combinaison de deux familles de technologies qui connaissent un succès commercial : le système universel de télécommunications mobiles (UMTS) (figure 3). Cette génération est considérée comme une extension naturelle du GSM et est largement utilisée dans le monde. Et CDMA2000 qui est également une extension du système IS-95. Ce dernier est principalement utilisé en Asie et en Amérique du Nord. Les interfaces radio de ces deux systèmes sont principalement basées sur le schéma d'accès multiple par répartition en code (CDMA).

Guidés par la volonté de définir la 3G comme une norme internationale et universelle permettant l'itinérance internationale et des économies de coûts d'équipements et d'appareils, ces systèmes se regroupent à la fin des années 1990 à l'initiative de l'ETSI (European Telecommunications Standards Institute) au sein d'un consortium appelé 3GPP (3rd Generation Partnership Project). Cette approche conduit au développement de la norme UMTS à la fin des années 1990. La première version de cette norme est appelée Release 99. Dans cette première version, le réseau d'accès est développé au travers d'une interface avec le réseau cœur GPRS. L'objectif initial est d'augmenter la capacité du système, principalement pour les services voix et données, et de permettre les services multimédias.

La première version (Release 99) utilise la technologie Wideband CDMA (WCDMA). Cette technologie prend en charge le mode duplex FDD (avec bandes appariées) et TDD (sans bandes appariées). Cette technologie repose sur le principe de l'étalement de spectre, où les signaux utiles sont étalés sur une bande passante de 3,84 MHz avant d'être placés sur une porteuse qui occupe un canal de 5 MHz. Sur un même opérateur, les appels des utilisateurs sont distingués au sein d'une cellule par l'attribution de codes spécifiques connus de la station de base et du terminal. WCDMA permet de se connecter à plusieurs cellules en même temps, améliorant ainsi la qualité de la communication lors du changement de cellule en mobilité. La Release 99 est limitée à un débit maximum en amont et en aval de 384 kbps.

Figure 3 : Architecture d'un réseau 3G (UMTS)





Une variante de l'UMTS TDD appelée TDSCDMA (*time division synchronous CDMA*) est également normalisée par le 3GPP. Cette technologie fonctionne avec une bande passante de 1,28 MHz et est principalement déployée en Chine.

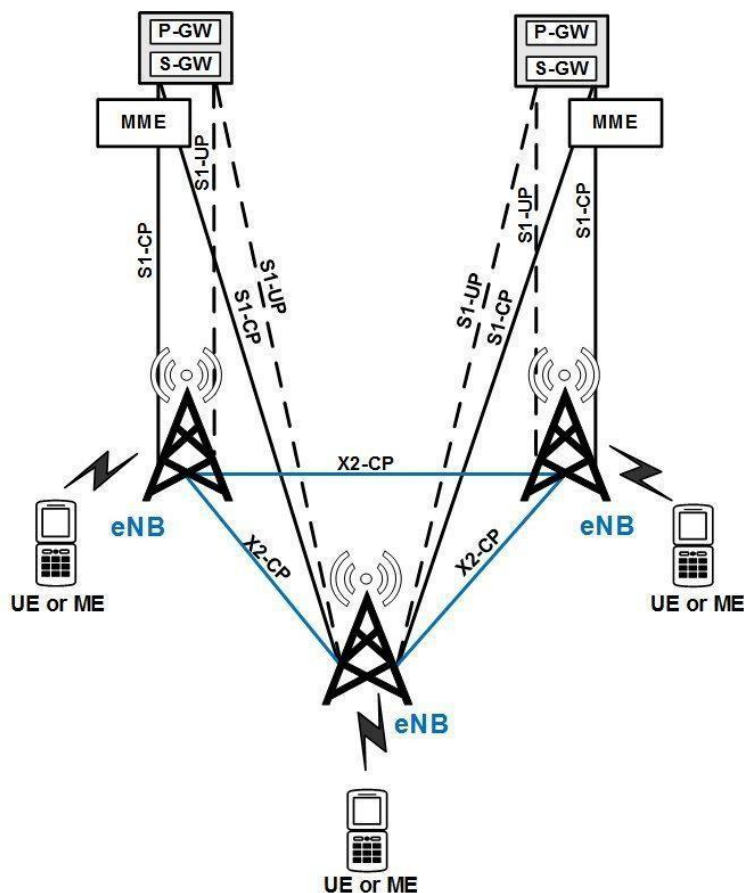
L'UMTS connaît deux évolutions majeures, le HSPA (*high speed packet access*) et le HSPA+, qui sont commercialement connus sous le nom de 3G+. La principale innovation de ces développements est le passage d'une commutation circuit sur l'interface radio, où des ressources radio sont réservées pour chaque équipement utilisateur (UE, *user equipment*) pendant la durée de l'appel, à une commutation de paquets, où la station de base (BS, *base station*) détermine dynamiquement le partage des ressources entre les UE actifs. L'allocation dynamique des ressources est effectuée par la fonction de planification (*scheduling*), en particulier en fonction de la qualité instantanée du canal radio de chaque UE, des limites de qualité de service et de l'efficacité globale du système. La commutation de paquets optimise ainsi l'utilisation des ressources radio pour les services de données. La modulation et le codage sont adaptatifs aux conditions radio de l'UE lors du service, et le débit instantané est augmenté en utilisant une modulation avec un nombre d'états plus élevé qu'en Release 99. C'est ainsi que la modulation 16QAM (*16 quadrature amplitude modulation*) est introduite pour la voie descendante en complément de la modulation QPSK (*quadrature phase shift keying*) utilisée dans la Release 99. De même, la modulation QPSK est introduite pour la voie montante en complément de la modulation BPSK (*binary phase shift keying*) utilisée en Release 99. Enfin, un nouveau mécanisme de retransmission rapide des paquets erronés, appelé HARQ (*hybrid automatic response request*), est défini entre l'UE et la station de base afin de réduire la latence du système en cas de perte de paquets. Ces évolutions offrent aux utilisateurs des débits maximaux de 14,4 Mb/s en voie descendante et de 5,8 Mb/s en voie montante, ainsi qu'une latence réduite.

1.4 Réseaux mobiles de quatrième génération (4G)

La quatrième génération de réseaux mobiles (4G) est envisagée en novembre 2004 comme l'évolution à long terme de l'UMTS (d'où le nom LTE, *long term evolution*), dans un atelier organisé par 3GPP appelé Future Evolution Workshop. Ce développement vise alors à maintenir la compétitivité de l'UMTS pendant au moins une décennie. C'est ainsi que LTE est adopté dans la Release 8 du consortium 3GPP (figure 4).



Figure 4 : Architecture d'un réseau 4G (LTE)



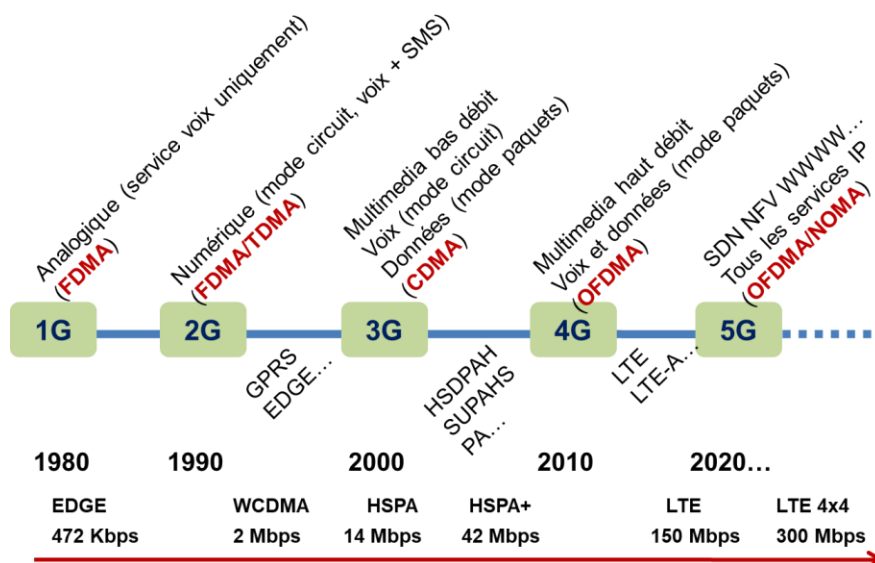
Les principales motivations et exigences de l'introduction de la 4G (LTE) (figure 5) sont les suivantes :

- La capacité : les gains associés aux évolutions HSPA et HSPA+ ont renforcé la capacité des réseaux par rapport à la Release 99 et au GSM. Toutefois, cette capacité reste faible quant aux besoins des opérateurs de réseaux mobiles. La mise sur le marché de terminaux tels que les téléphones intelligents (*smartphones*) ou les clés 3G+ contribue largement à entraîner l'explosion des demandes de services de données mobiles par les utilisateurs. L'utilisation de réseaux mobiles comme alternative aux réseaux de données résidentiels tels l'ADSL contribue également à cette forte croissance du trafic de données mobiles. Afin de satisfaire ces demandes croissantes, les opérateurs travaillent beaucoup sur la densification de leur réseau 3G et aussi sur les agrégations des porteuses. Cependant, ces fréquences sont des ressources très limitées et leur pénurie cause des rejets d'appels ou des réductions de débits au cours des appels. C'est ainsi qu'à la fin de l'année 2004, 3GPP discute l'introduction du LTE afin d'accroître la capacité des réseaux mobiles ;
- Les débits : La fourniture par les opérateurs des débits supérieurs à ceux offerts par les réseaux HSPA est l'objectif principal du déploiement du LTE.

Les exigences pour la technologie LTE, pour un UE de référence comprenant deux antennes en réception et une antenne en émission, sont les suivantes :

- 100 Mb/s en voie descendante pour une largeur de bande allouée de 20 MHz, soit une efficacité spectrale crête de 5 bit/s/Hz ;
- 50 Mb/s en voie montante pour une largeur de bande allouée de 20 MHz, soit une efficacité spectrale crête de 2,5 bit/s/Hz ;

Figure 5 : Évolution des débits selon les technologies cellulaires utilisées



- La latence : la latence d'un système est la mesure du délai introduit par ce système. On distingue deux types de latence, à savoir la latence du plan contrôle (temps nécessaire pour établir une connexion et accéder au service) et la latence du plan usager (délai de transmission d'un paquet au sein du réseau une fois la connexion établie). La latence traduit la capacité du système à traiter rapidement les demandes d'utilisateurs ou de service. Une latence forte limite l'interactivité d'un système et s'avère pénalisante pour l'utilisateur de certains services de données. L'UMTS et ses évolutions HSPA offrent une latence trop importante pour offrir des services tels que les jeux vidéo en ligne. L'amélioration de la latence est un des éléments qui concourt à la décision de définir un nouveau système ;

En termes d'exigence, l'objectif fixé pour le LTE est d'améliorer la latence du plan contrôle par rapport à l'UMTS via un temps de transition inférieur à 100 ms entre un état de veille de l'UE et un état actif autorisant l'établissement du plan usager. Pour la latence du plan usager, définie par le temps de transmission d'un paquet entre la couche IP de l'UE et la couche IP d'un nœud du réseau d'accès ou inversement, c'est-à-dire le délai de transmission d'un paquet IP au sein du réseau d'accès, il est visé une latence inférieure à 5 ms dans des conditions de faible charge du réseau et pour des paquets IP de petite taille.



- L'adaptation au spectre disponible : la technologie UMTS contraint les opérateurs à utiliser des canaux de 5 MHz. Cette limitation est pénalisante à deux titres :
 - les allocations spectrales dont la largeur est inférieure à 5 MHz ne peuvent pas être utilisées (sauf pour le TD-SCDMA), ce qui limite le spectre disponible ;
 - en cas de disponibilité de plusieurs bandes spectrales de largeur de 5 MHz, un opérateur est dans l'incapacité d'allouer simultanément plusieurs porteuses à un même UE. Cette contrainte limite le débit maximal potentiel du système ainsi que la flexibilité de l'allocation des ressources spectrales aux utilisateurs. Il faut noter que cette contrainte est partiellement levée en HSPA+ Release 8 avec la possibilité de servir un terminal UE sur deux porteuses de 5 MHz simultanément.

Ainsi, un consensus s'impose sur le besoin d'un système dit agile en fréquence, capable de s'adapter à des allocations spectrales variées. Cette agilité est un objectif de conception fort du LTE qui doit pouvoir opérer sur des porteuses de différentes largeurs afin de s'adapter à des allocations spectrales variées. Les largeurs de bandes initialement requises sont par la suite modifiées pour devenir les suivantes : 1,4 MHz ; 3 MHz ; 5 MHz ; 10 MHz ; 15 MHz et 20 MHz dans le sens montant et descendant. Les modes de duplexage FDD et TDD doivent être pris en charge pour toutes ces largeurs de bande.

L'émergence de l'OFDM : L'OFDM (*orthogonal frequency division multiplexing*) offre plusieurs avantages aux systèmes radio mobiles et bénéficie d'une grande immunité contre l'interférence entre symboles créée par les réflexions du signal sur les objets de l'environnement. En plus, l'OFDM permet de gérer simplement des largeurs de bande variables et potentiellement grandes, ce qui était une motivation de l'introduction du LTE ;

Enfin, l'OFDMA (*orthogonal frequency division multiple access*) assure le partage aisé des ressources fréquentielles entre un nombre variable d'utilisateurs bénéficiant de débits divers. Pour ces raisons, il est décidé de baser LTE sur l'OFDM, en rupture avec le CDMA.

Signalons enfin que la mobilité est une fonction clé pour un réseau mobile. Le LTE vise à rester fonctionnel pour des UE se déplaçant à des vitesses élevées (jusqu'à 350 km/h, et même 500 km/h).

Le LTE coexiste avec les autres technologies 3GPP : tout terminal UE qui met en œuvre les techniques GSM et UMTS en complément du LTE est capable d'effectuer les *handovers* en provenance et à destination des systèmes GSM et UMTS, ainsi que les mesures associées. Le temps d'interruption de service lors d'une procédure de *handover* entre le système LTE et les systèmes GSM ou UMTS doit rester inférieur à 300 ms pour les services en temps réel et inférieur à 500 ms pour les autres services.

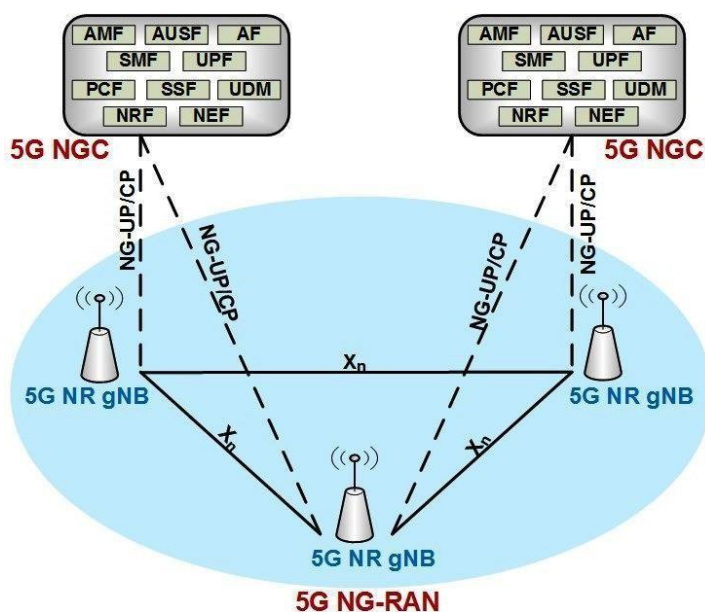
1.5 Réseaux mobiles de cinquième génération (5G)

Avec l'introduction de la 4G, la rupture avec les systèmes précédents n'est pas aussi franche. La 4G est plus focalisée sur les aspects services que sur les aspects techniques. Les objectifs principaux motivant l'introduction de cette nouvelle génération de système sont la réduction

des coûts des usagers et des opérateurs tout en multipliant les services, avec en parallèle, l'augmentation des débits et la réduction de la latence.

Selon les objectifs de l'Union Internationale des Télécommunications (UIT, IMT-2020), la technologie 5G doit offrir un débit utilisateur maximal et un débit maximal respectivement 10 et 20 fois supérieurs à ce qui est actuellement disponible avec la technologie 4G. La densité maximale de connexions est multipliée par 10 et la latence divisée par au moins 10 (la latence point à point cible est de 1 ms, contre 30 à 40 ms en 4G).

Figure 6 : Architecture d'un réseau d'accès 5G (5G NG-RAN)



Avec la technologie 5G, l'inquiétude n'est pas toujours centrée sur la seule augmentation du débit, à l'image des générations qui l'ont précédée, mais surtout intéressée à répondre à un certain nombre de contraintes et nécessités du nombre en croissance phénoménale d'utilisateurs (allant des téléphones mobiles aux différents objets connectés) et de leurs besoins en ce qui concerne la réduction de la latence, l'augmentation de la couverture intérieure en intégrant les petites cellules ainsi que la couverture extérieure, la réduction des coûts, l'augmentation de l'efficacité spectrale, les débits binaires, tout en garantissant une bonne qualité de service (QoS) et une bonne qualité d'expérience (QoE) et surtout avec des économies d'énergie. Cela facilite la mise en œuvre de villes intelligentes dans les meilleures conditions. En effet, une ville intelligente nécessite des connexions très importantes, rapides, efficaces et fiables avec des économies d'énergie considérables. La technologie 5G (figure 6) doit faire face à un grand nombre d'appareils connectés (ordinateurs, smartphones, tablettes...) qui atteignent, selon Cisco et Ericsson, près de 6,2 milliards fin 2021. L'efficacité énergétique bit / joule est un principe de conception central de la 5G.

Dans le livre blanc NGMN pour la 5G^[1], l'efficacité énergétique est identifiée comme un indicateur de performance clé (KPI) de la 5G et est définie comme le nombre de bits pouvant être transmis par joule d'énergie, où l'énergie est calculée sur le réseau entier, y compris les technologies cellulaires potentiellement héritées, l'accès radio, les réseaux cœurs et les centres



de données. Aussi, on peut citer la tendance des connexions du réseau 5G au réseau électrique, du type appelé *smart grid*. L'utilisation d'énergies renouvelables (notamment éolienne ou solaire), d'énergie propre, mais présentant en retour un nouveau problème de disponibilité non sécurisée, est une autre piste intéressante pour la 5G.

Selon l'UIT, dans sa recommandation UIT-R M.2083-0^[2], les systèmes IMT futurs (5G) doivent prendre en charge de nouveaux cas d'utilisation, y compris des applications qui nécessitent des communications à très haut débit de données, un grand nombre de dispositifs connectés, et des applications à temps de latence ultra-court et à haute fiabilité.

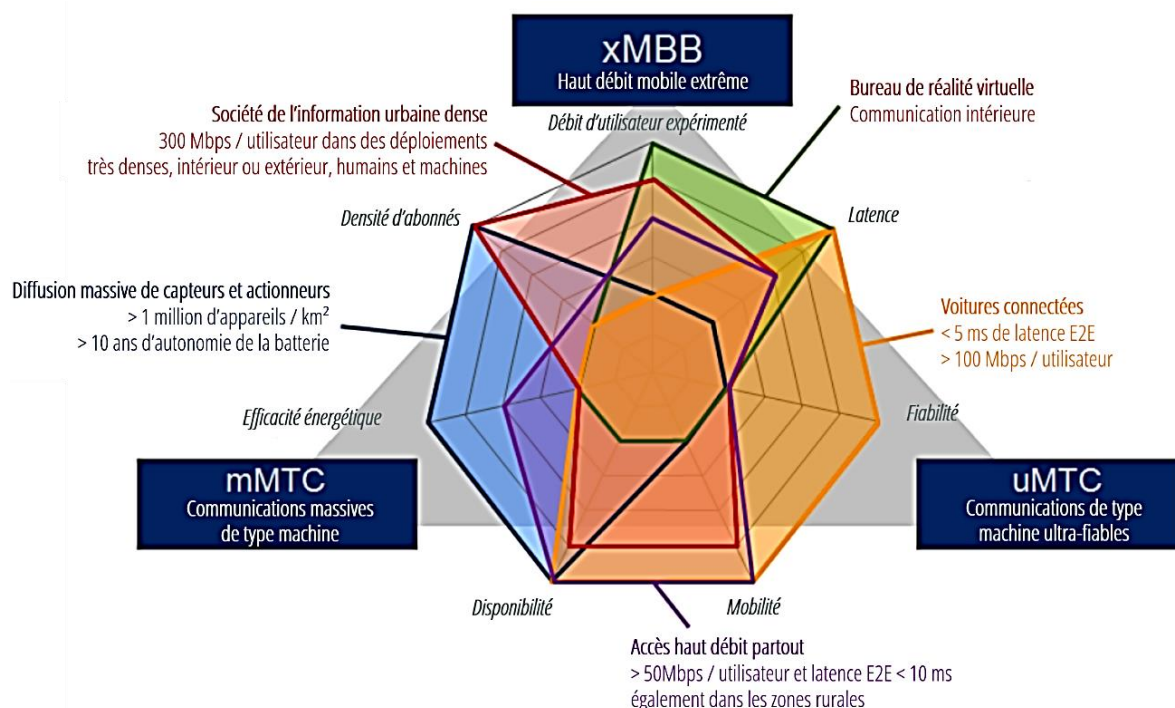
2. Catégories d'utilisation dans un réseau 5G

Pour la technologie 5G, l'Union internationale des télécommunications (UIT) identifie trois grandes catégories d'utilisation^[2], comme le montre la figure 7 :

- mMTC (*massive machine type communications*) : communications entre un grand nombre d'objets (appareils) avec des besoins de qualité de service différenciés. Le but de cette catégorie est de répondre à l'augmentation exponentielle du nombre d'objets connectés notamment les objets portables, les moniteurs cardiaques et les appareils connectés à un domicile intelligent et à une ville intelligente ainsi que les véhicules connectés ;
- eMBB (*enhanced mobile broadband*) : connexion ultra-haut débit aussi bien en extérieur (*outdoor*) qu'en intérieur (*indoor*) avec une qualité de service acceptable, même en bordure de cellule. Cette catégorie de services concerne notamment la réalité virtuelle et la réalité augmentée et la vidéo temps réel ;
- uRLLC (communications Ultra-fiables et *low latency*) : communications ultra-fiables pour les besoins critiques avec une très faible latence. De nombreux usages sont rendus possibles grâce à la fiabilité et à la réactivité du réseau 5G notamment les voitures autonomes, qui doivent réagir rapidement, soit en temps réel, aux situations rencontrées. La médecine (e-santé) et l'industrie sont également concernées par cette catégorie de services.

Il va sans dire qu'en considérant chaque type de service séparément et en construisant un réseau 5G en conséquence, on se retrouve avec des conceptions et des architectures de réseau d'accès radio (RAN) très différentes^[3].

Figure 7 : Nouveaux services pris en charge par la technologie 5G



L'objectif de la technologie 5G est également de fournir, entre autres :

- un réseau très fiable, avec des performances plus homogènes, quelle que soit la distance entre l'utilisateur et la station de base ;
- une connexion stable même en mobilité (avec des vitesses allant jusqu'à 500 km/h) ;
- une augmentation de l'efficacité énergétique (batteries jusqu'à 100 fois moins énergivores).

Pour mettre en œuvre ces trois catégories d'utilisation (mMTC, eMBB, uRLLC), l'UIT établit huit indicateurs clés de performance (KPI, *key performance indicator*) afin de spécifier et de mesurer les caractéristiques du système 5G (appelé IMT-2020).

Tableau 1 : Comparaison entre technologie 5G et technologie 4G (LTE-Advanced)^[4]

Performance / Generation	4G LTE-A (IMT-Advanced)	5G (IMT-2020)	Type of service
Peak data rate (Gb/s)	1	20 (DL), 10(UL)	eMBB
User experience data rate (Mb/s)	10	100 (DL), 50 (UL)	eMBB
Spectrum efficiency (bit/Hz)	1x	3x	eMBB
Device mobility (km/h)	350	500	uRLLC/eMBB
Latency (ms)	10	1	uRLLC (4ms for eMBB)
Connection density (number of connected objects/km ²)	10 ⁵	10 ⁶	mMTC
Network's energy efficiency	1x	100x	mMTC
Area traffic capacity (Mb/s/m ²)	0.1	10	eMBB

D'autres nouveaux indicateurs de performance clés peuvent être ajoutés à ces huit catégories : la fiabilité ; le temps d'interruption de la mobilité ; la bande passante ; l'efficacité maximale ; l'efficacité du spectre.

3. Technologies mises en place par la 5G

Les spécificités de ces différentes catégories de service exigent des compromis techniques complexes impliquant des décisions délicates en termes de conception. À cet effet, pour permettre la mise en place de ces différentes catégories de services, les technologies suivantes sont introduites dans l'architecture du réseau 5G, tant au niveau du réseau d'accès radio qu'au niveau du réseau cœur, dont les principales sont les suivantes :

3.1 Architecture radio

3.1.1 Techniques de transmission optimisées

La technologie 5G utilise un large spectre de fréquences et doit permettre la transmission d'un volume colossal de données, et doit donc gérer les interférences, elle utilise des techniques déjà disponibles dans la 4G mais avec une optimisation accrue et davantage de flexibilité et d'efficacité dans l'utilisation du spectre de fréquences.

3.1.2 Nouvelles technologies d'antennes

La technologie 5G utilise la technique des antennes MIMO déjà utilisée par la technologie 4G mais avec en plus un très nombre d'antennes nommées MIMO. Cette dernière est possible grâce au nouveau spectre de fréquence réservé à la 5G, contenant de très hautes fréquences et donc de très petites longueurs d'ondes (ondes millimétriques) et des antennes à très petites

dimensions. Lesquelles antennes permettent la focalisation des faisceaux (ou *beamforming*) pour mieux desservir les équipements utilisateurs et aussi économiser l'énergie en transmission de données.

3.2 Architecture réseau

La fourniture des différentes catégories de services (mMTC, eMBB et uRLLC) proposées par la 5G est rendue possible grâce à son architecture basée essentiellement sur la virtualisation de son réseau. Dans cette architecture, le réseau 5G est découpé en trois principales tranches ou couches (*network slicing*) (figure 8) et chacune des couches permet d'assurer une catégorie de services bien spécifique. La mise en place de ces différentes couches doit permettre d'adapter les services de connectivité du réseau aux besoins de services demandant chacun une qualité de service différente.

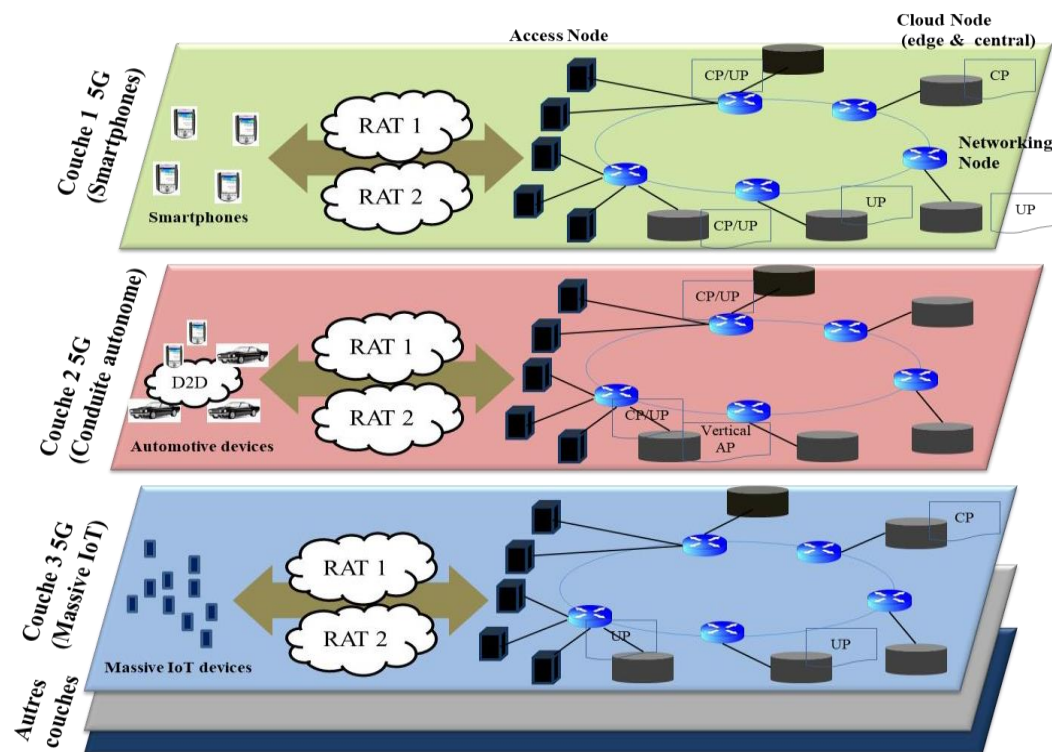
Chaque tranche (*slice*) englobe un ensemble de ressources logicielles et matérielles permettant au réseau de répondre à un besoin spécifique avec une qualité bien définie.

Architecture basée sur la virtualisation :

La gestion et le contrôle des infrastructures de la 5G sont devenus de plus en plus complexes du fait de la prolifération des catégories de services à fournir aux utilisateurs et également l'extension des bandes de fréquences allouées pour ces services. C'est ainsi que la 5G introduit dans son architecture réseau deux technologies principales à savoir la technologie *NFV* (*network function virtualisation*) et la technologie *SDN* (*soft defined networking*). Ces deux technologies sont décrites comme suit :

- la technologie *network function virtualization* (NFV) consiste à séparer les infrastructures matérielles (*hardware*) des applications fonctions logicielles (*software*). C'est ainsi que l'architecture se base sur des machines virtuelles par l'utilisation de logiciels qui peuvent être utilisés sur des équipements du réseau et aussi par l'utilisation des services du *cloud*. En adoptant la technologie NFV, la gestion et le contrôle du réseau 5G peuvent en plus se passer des différents équipements et des différentes technologies utilisés et donc assurer plus de flexibilité et d'interopérabilité. Cette technologie s'avère très bénéfique pour les opérateurs des télécommunications mobiles qui se voient offrir la possibilité de mieux gérer et exploiter les ressources physiques de leurs réseaux ;
- la technologie *software defined networking* (SDN), par le biais des applications logicielles utilisant des interfaces de programmation (API). La SDN assure la programmation du comportement du réseau de manière contrôlée et centralisée. La SDN consiste à séparer la partie contrôle de la partie transport des données. Le contrôle des différents équipements du réseau est centralisé.

Figure 8 : Découpage du réseau 5G en tranches (*network slicing*)



4. 5G pour l'internet des objets (IoT)

L'un des défis majeurs de la technologie 5G est la possibilité d'assurer une connectivité améliorée des objets connectés à l'internet (internet des objets, IoT). Le défi est donc de connecter des objets qui ont, en raison de leur nature, des besoins très différents et demandent des qualités de service différentes notamment en termes de débit et latence. On cite par exemple les cas suivants :

- les services nécessitant un bas débit comme le cas des capteurs de mesures permettant le suivi de consommation d'eau ou encore ceux pour la relève de température. C'est le cas également des capteurs servant à renseigner la localisation d'équipements se trouvant dans des endroits donnés (*tracking*). Généralement, ces objets connectés sont équipés de batteries à longue durée de vie et nécessitent donc une consommation de faible énergie ;
- les services nécessitant une transmission à haut débit, par exemple les équipements utilisés dans la télésurveillance qui nécessitent la transmission de données volumineuses et exigent une très faible latence ;
- les services de très fortes exigences en termes de fiabilité et de latence, par exemple les véhicules autonomes qui échangent un volume de données très important et exigent une grande fiabilité pour éviter les accidents ou le dérapage et aussi une communication en temps réel exigeant une latence très stricte.

Pour répondre à ces exigences, le consortium 3GPP introduit les technologies LTE-M et NB-IoT appartenant à la famille des LPWAN (*low power wide area networks*) déjà intégrées dans la technologie 4G. Les réseaux LPWAN sont effet conçus pour répondre aux besoins de communication des objets qui sont à faible consommation énergétique et aussi permettre une large couverture. Toutefois, ces besoins sont jusque-là assurés par les réseaux mobiles cellulaires 2G et 3G qui consomment beaucoup d'énergie mais permettent une transmission longue portée, ils sont aussi assurés par d'autres types de réseaux mobiles non cellulaires de la famille LPWAN notamment SigFox et LoRaWAN. En revanche, l'intégration par 3GPP des deux technologies LTE- M et NB-IoT permettent la gestion de l'inactivité des objets connectés et ainsi optimiser la consommation énergétique et également permettre le *roaming* et le *handover*. Le tableau 2 montre la comparaison des performances des deux technologies.

Tableau 2 : Comparaison entre les performances NB-IoT et LTE-M

NB-IoT (Narrow-Band IoT)	LTE-M
- Débits de 20 à 25 kbps - Communication bidirectionnelle - Latence de l'ordre de la dizaine de secondes - Bonne pénétration en intérieur (indoor) - Ne gère pas la voix en plus des données	- Débit maximal théorique de l'ordre de 1 Mb/s (ce qui est très supérieur à SigFox ou LoraWan, mais reste inférieur à la 4G LTE) - Communication bidirectionnelle - Latence < 100 ms - Gère la voix en plus des données - <i>Handover</i> et <i>roaming</i> LTE-M permet un meilleur débit que les autres LPWAN, mais consomme beaucoup plus.

Par ailleurs, et afin de répondre aux besoins en termes de latence critique pour les objets connectés (IoT), la 5G introduit une nouvelle architecture informatique (appelée Edge Computing), qui est en quelque sorte un petit centre de données à proximité de l'utilisateur, permettant, comme son nom l'indique, de traiter les données des utilisateurs à la périphérie du réseau. Le fait de déployer les serveurs à proximité des utilisateurs contribue largement à minimiser la latence et permet donc des traitements en temps réel, tels qu'exigés par exemple par les véhicules autonomes.

5 Exigences minimales des performances techniques de l'interface radio 5G NR

Dans son rapport ITU-R M.2410-0^[5], l'UIT définit les exigences minimales relatives aux performances techniques des interfaces radio de la technologie de cinquième génération (IMT-2020) avec l'objectif de rendre cette technologie plus flexible, fiable et sécurisée lors de la fourniture de divers services dans les trois catégories d'utilisation (eMBB, mMTC et uRLLC).

Dans cette section, nous citons les principales exigences relatives aux performances techniques minimales des technologies d'interface radio de la 5G telles qu'identifiées par l'UIT.



5.1 Débit de données maximal

Le débit de données maximal est le débit pouvant être atteint dans des conditions idéales (en bit/s), qui sont les bits de données reçus en supposant des conditions sans erreurs attribuables à une seule station mobile, lorsque toutes les ressources radio affectées pour la direction de liaison correspondante sont utilisées. Ceci correspond à l'exclusion des ressources radio utilisées pour la synchronisation de la couche physique, les signaux de référence (ou pilotes), les bandes de gardes et les temps de garde.

Le débit de données maximal est défini pour une seule station mobile. Dans une seule bande, il est lié à l'efficacité spectrale maximale dans cette bande. Soit W la largeur de bande du canal et SE_p le pic d'efficacité spectrale dans cette bande. Le débit de données maximal de l'utilisateur R_p est donné par :

$$R_p = W \cdot SE_p \quad (1)$$

L'efficacité spectrale maximale et la bande passante disponible peuvent avoir des valeurs différentes, et ce, dans différentes gammes de fréquences. Dans le cas où la bande passante est agrégée sur plusieurs bandes, le débit de données maximal sera additionné sur ces bandes. Par conséquent, si la bande passante est agrégée sur les bandes Q , le débit de données total maximal est donné par :

$$R = \sum_{i=1}^Q W_i \cdot SE_{pi} \quad (2)$$

où W_i et SE_{pi} ($i = 1, \dots, Q$) sont respectivement les largeurs de bande et les rendements spectraux des composants.

Cette exigence est définie à des fins d'évaluation dans le scénario d'utilisation eMBB. Les exigences minimales pour le débit de données maximal sont les suivantes :

- le débit de données maximal de la liaison descendante est de 20 Gbit / s ;
- le débit de données maximal de la liaison montante est de 10 Gbit / s.

5.2 Latence

5.2.1 Latence du plan utilisateur

La latence du plan utilisateur est la contribution du réseau d'accès radio au temps entre le moment où la source envoie un paquet et le moment où la destination le reçoit (en ms). Il est défini comme le temps unidirectionnel nécessaire pour transmettre avec succès un paquet / message de la couche application du point d'entrée SDU de la couche 2/3 du protocole radio au point de sortie SDU de la couche 2/3 du protocole radio de l'interface radio en liaison montante ou en liaison descendante dans le réseau pour un service donné et dans des conditions déchargées, en supposant que la station mobile est à l'état actif.

Cette exigence est définie à des fins d'évaluation dans les scénarios d'utilisation d'eMBB et d'URLLC.

Les exigences minimales pour la latence du plan utilisateur sont les suivantes :

- 4 ms pour eMBB ;
- 1 ms pour URLLC.

5.2.2 Latence du plan contrôle

La latence du plan de contrôle fait référence au temps de transition entre l'état le plus « efficace de la batterie » (par exemple, l'état de veille) et le début du transfert de données en continu (par exemple, le passage à l'état actif).

Cette exigence est définie à des fins d'évaluation dans les scénarios d'utilisation d'eMBB et d'URLLC.

La configuration minimale requise pour la latence du plan de contrôle est de 20 ms. Toutefois, il est préférable d'envisager une latence du plan de contrôle plus faible, par exemple aux alentours de 10 ms.

5.2.3 Densité de connexion

La densité de connexion est le nombre total d'appareils remplissant une qualité de service (QoS) spécifique par unité de surface (par km^2).

La densité de connexion doit être obtenue pour une bande passante et un nombre de TRxP limités. La qualité de service cible est de prendre en charge la remise d'un message d'une certaine taille dans un certain délai et avec une certaine probabilité de succès.

Cette exigence est définie à des fins d'évaluation dans le scénario d'utilisation mMTC. L'exigence minimale de densité de connexion est de un million d'appareils par km^2 .

5.2.4 Efficacité énergétique

L'efficacité énergétique du réseau est la capacité d'une technologie d'interface radio (TIR) à minimiser la consommation d'énergie du réseau d'accès radio par rapport à la capacité de trafic fournie. L'efficacité énergétique de l'appareil est la capacité du TIR à minimiser la puissance consommée par le modem de l'appareil par rapport aux caractéristiques du trafic.

L'efficacité énergétique du réseau et de l'appareil peut concerner la prise en charge des deux aspects suivants :

- la transmission de données efficace dans le cas où il est chargé ;
- la faible consommation d'énergie en l'absence de données.

Une transmission de données efficace dans un cas chargé est démontrée par l'efficacité spectrale moyenne. Une faible consommation d'énergie en l'absence de données peut être estimée par le taux de sommeil. Le taux de sommeil est la fraction de ressources temps inoccupés (pour le réseau) ou de temps de sommeil (pour l'appareil) dans une durée

correspondant au cycle de la signalisation de commande (pour le réseau) ou au cycle de réception discontinue (pour le périphérique) lorsqu'aucun transfert de données utilisateur n'a lieu. En outre, la durée du sommeil, c'est-à-dire la période continue sans transmission (pour le réseau et l'appareil) et la réception (pour l'appareil), doit être suffisamment longue.

Cette exigence est définie à des fins d'évaluation dans le scénario d'utilisation eMBB.

La technologie d'interface radio doit avoir la capacité de prendre en charge un taux de sommeil élevé et une longue durée de sommeil.

5.2.5 Fiabilité

La fiabilité concerne la capacité à transmettre une quantité donnée de trafic dans une durée prédéterminée avec une probabilité de succès élevée.

La fiabilité est la probabilité de succès de la transmission d'un paquet de couche 2/3 dans un temps maximum requis, qui est le temps nécessaire pour délivrer un petit paquet de données du point d'entrée SDU de la couche de protocole radio 2/3 à la couche de protocole radio 2/3 du point de sortie SDU de l'interface radio avec une certaine qualité de canal.

Cette exigence est définie à des fins d'évaluation dans le scénario d'utilisation d'uRLLC.

L'exigence minimale pour la fiabilité est la probabilité de succès ($1-10^{-5}$) de la transmission d'une PDU de couche 2 (unité de données de protocole) de 32 octets en 1 ms dans la qualité de canal au bord de couverture pour l'environnement de test Urbain Macro-uRLLC, en supposant de petits paquets de données d'application (par exemple 20 octets de données d'application + surcharge de protocole).

Toutefois, il est préférable de considérer des paquets de plus grandes tailles, par exemple une taille de PDU de couche 2 jusqu'à 100 octets.

5.2.6 Mobilité

La mobilité est la vitesse maximale de la station mobile à laquelle une QoS définie peut être atteinte (en km/h). Les classes de mobilité suivantes sont définies :

- Stationnaire : 0 km/h ;
- Piéton : de 0 km/h à 10 km/h ;
- Véhicule : de 10 km/h à 120 km/h ;
- Véhicule à grande vitesse : de 120 km / h à 500 km / h.

Les véhicules à grande vitesse (jusqu'à 500 km/h) sont principalement envisagés pour les trains à grande vitesse. Le tableau 3 définit les classes de mobilité qui doivent être prises en charge dans les environnements de test respectifs.



Tableau 3 : Classes de mobilité

ENVIRONNEMENT DE TEST POUR eMBB			
	Indoor Hotspot - eMBB	Urbain Dense - eMBB	Rural - eMBB
Classes de mobilité supportées	Stationnaire, Piéton	Stationnaire, Piéton, Véhicule (jusqu'à 30 km/h)	Piéton, Véhicule, Véhicule à grande vitesse

Une classe de mobilité est prise en charge si le débit de données de liaison de canal de trafic sur la liaison montante, normalisé par la bande passante, est comme indiqué dans le tableau 4. Cela suppose que l'utilisateur se déplace à la vitesse maximale dans cette classe de mobilité dans chacun des environnements de test.

Cette exigence est définie à des fins d'évaluation dans le scénario d'utilisation eMBB.

Tableau 4 : Débits de données de liaison de canal de trafic normalisé par bande passante

Environnement du test	Débit de données de liaison du canal de trafic normalisé (bit / s/ Hz)	Mobilité (km/h)
Indoor – Hotspot – eMBB	1,5	10
Urbain Dense – eMBB	1,12	30
Rural –eMBB	0,8	120
	0,45	500

5.2.7 Temps d'interruption de mobilité (MIT, *mobility interruption time*)

Le temps d'interruption de mobilité est la durée la plus courte prise en charge par le système pendant laquelle un terminal utilisateur ne peut échanger des paquets de plan utilisateur avec aucune station de base pendant les transitions.

Le temps d'interruption de mobilité comprend le temps nécessaire pour exécuter toute procédure de réseau d'accès radio, protocole de signalisation de commande de ressources radio ou autres échanges de messages entre la station mobile et le réseau d'accès radio.

Cette exigence est définie à des fins d'évaluation dans les scénarios d'utilisation d'eMBB et d'URLLC.

Le minimum requis pour le temps d'interruption de mobilité est de 0 ms.

5.2.8 Bande passante

La bande passante est la bande passante agrégée maximale du système. La bande passante peut être prise en charge par une ou plusieurs porteuses de radiofréquence (RF).

L'exigence de bande passante est d'au moins 100 MHz.

La technologie d'interface radio doit prendre en charge des largeurs de bande allant jusqu'à 1 GHz pour un fonctionnement dans des bandes de fréquences plus élevées (par exemple au-dessus de 6 GHz).

Conclusion

À travers ce chapitre, nous passons en revue les différentes technologies de la 1G à la 5G et ce, afin d'apprécier l'évolution des technologies en termes de capacité, de débit de latence et d'efficacité énergétique. Nous constatons que les générations qui ont précédé la 5G ont en commun le souci d'augmenter davantage les débits et la couverture pour satisfaire les demandes de trafic qui vont constamment croissantes. La 5G est plutôt considérée comme une révolution par rapport aux générations qui la précèdent. Elle est en mesure de fournir une très forte augmentation de débits et de réduction de latence. De plus, la 5G est capable d'assurer les besoins des services critiques comme dans le cas des communications en temps réel pour les véhicules autonomes dont les critères sont très stricts en ce qui concerne le débit, la latence et la fiabilité. Ces derniers étaient impossibles à réaliser par la 4G, même avec ses versions les plus évoluées.

Nous citons également différentes catégories de services qui doivent être assurées par la technologie 5G ainsi que les différentes architectures réseau mises en place pour assurer ces cas d'utilisation. Cette architecture est découpée en trois tranches (*slices*). Ce découpage en tranches s'appuie sur une plate-forme virtuelle et permet aux opérateurs de réseaux de télécommunication mobile de disposer d'un réseau pouvant supporter différentes qualités de service (latence, débit, fiabilité) et de répondre ainsi aux demandes de services selon les besoins de chacune. Nous présentons également les différentes techniques utilisées pour l'internet des objets. Enfin, nous citons les différentes exigences de l'UIT concernant les performances techniques de la technologie 5G, sur la base desquelles les différents acteurs concernés (instances de normalisation, équipementiers, opérateurs, chercheurs...) s'appuient pour concevoir les normes y répondant.

Chapitre II ♦ Cinquième génération nouvelle radio de communication mobile (5G NR)

Introduction

Le transfert des données issues de la couche IP relatives au service requis par l'utilisateur et qui se fait par voie des airs est assuré par l'interface radio. L'exigence de qualité du service est une nécessité à respecter en transfert au niveau de la latence et du débit malgré la variabilité extrême du médium, cependant cela ne peut se réaliser qu'à travers l'optimisation de l'accès aux ressources spectrales intrinsèquement limitées. Le spectre, aussi variable qu'il soit par région, demande une adaptabilité de nature à maîtriser les différentes bandes disponibles.

La performance connue de la technologie 4G sous ses versions les plus variées (LTE-A et LTE-A Pro) ne permet pas de satisfaire toutes les exigences technologiques de qualité qui lui sont assignées, néanmoins, le travail accompli ces dix dernières années sur la technologie LTE permet l'émergence de solutions techniques plus avancées, la 3GPP est le début d'une nouvelle technologie d'accès radio dénommée 5G NR.

Cette nouvelle génération radio 5G NR est la version 15 de la 3GPP, elle utilise des techniques de modulation ainsi que des formes d'ondes et d'accès qui permettent au système de satisfaire les besoins de services en haut débit, ainsi que les systèmes à faible débit en préservant la durée de vie de la batterie, elle apporte des améliorations importantes en termes de flexibilité, d'efficacité et de progrès au regard de la consommation d'énergie et d'occupation du spectre.

La 5G NR est la nouvelle génération réseau sans fil qui connecte mobiles, smartphones, voitures, compteurs et autres machines intelligentes avec les mêmes qualités de service qu'un système haut débit dont la réalisation est à moindre coûts.

Avec l'installation d'une latence de meilleur niveau, d'une bonne fiabilité et d'une meilleure sécurité, la 5G NR est un excellent moyen de connexion de l'internet des objets (IoT) permettant d'offrir de nouveaux types de services à cette fin.



Avec une large gamme de fréquence, la NR couvre une bande de moins de 1GHz à 100 GHz et déploie de nouveaux matériels, comme les pico, micro et macrocellules, c'est un défi relevé avec succès, le 3GPP introduit une couverture physique adaptée au NR, elle est par ce fait le moyen de transmission des signaux à travers les appareils technologiques de modulation de la technique d'accès.

Dans ce chapitre, nous discutons des caractéristiques les plus importantes de la couverture physique de 5G NR ainsi que des signaux et canaux déployés afin de répondre aux exigences de la 5G.

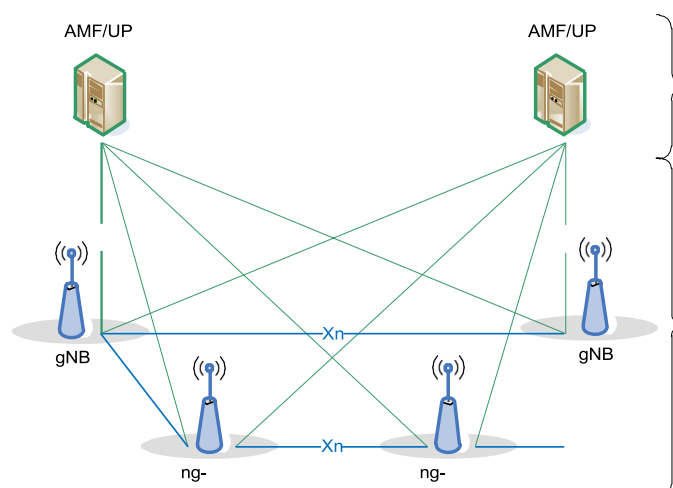
1. Caractéristiques globales du réseau d'accès radio 5G (NG RAN)

1.1 Cartographie techniques globale du Réseau 5G RAN (schéma technique général du réseau 5G)

Pour la connexion au réseau mobile 5G (5GC), les mobiles constituent l'outil de base permettant l'accès à la radio 5G à travers ses stations implantées, c'est pour cela que ces derniers constituent le moyen à travers lequel les terminaux mobiles communiquent via un lien radio 5G nommée gNB (*next génération node base station*) ou bien par un lien 4G nommé ng-eNB.

C'est grâce à une interface Xn que les stations de base gNB et ng-eNB sont reliées, c'est aussi à travers les interfaces NG et 5GC et plus particulièrement à l'AMF (*access and mobility function*), et aussi par l'interface NG-u et l'interface NG-C et la fonction UPF (figure 09).

Figure 9 : Architecture globale du réseau d'accès radio 5G RAN

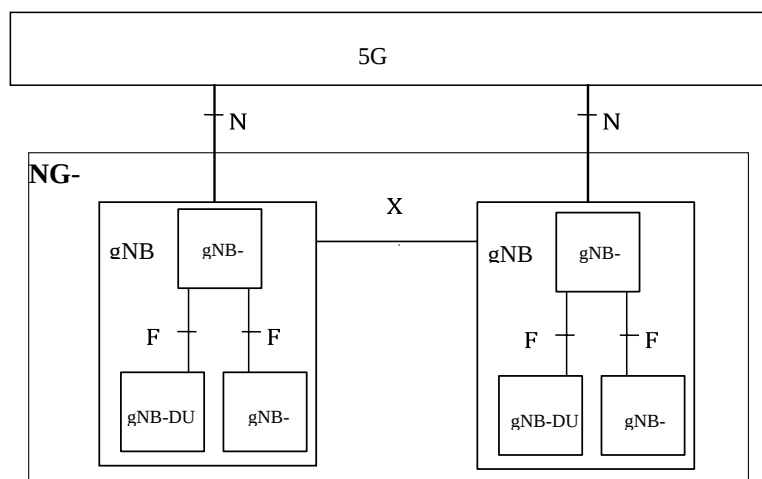


1.1.1 Architecture de la station de base gNB

On parle de la constitution du nœud de connexion NG-RAN (figure 10) par deux unités, une unité centralisée gNB-CN (*centralized unit*) et une unité distribuée gNB-DU (*distributed unit*) on parle ici d'une unité qui contrôle une ou une multitude d'unités gNB-DU et qui est l'unité gNB-CU alors que l'unité gNB-DU ne peut être en revanche contrôlée que par une seule unité gNB-CU.

La 3GPP permet de proposer plusieurs découpages fonctionnels entre entités gérées par l'unité gNB-CU et celles gérées par l'unité gNB-DU (figure 10).

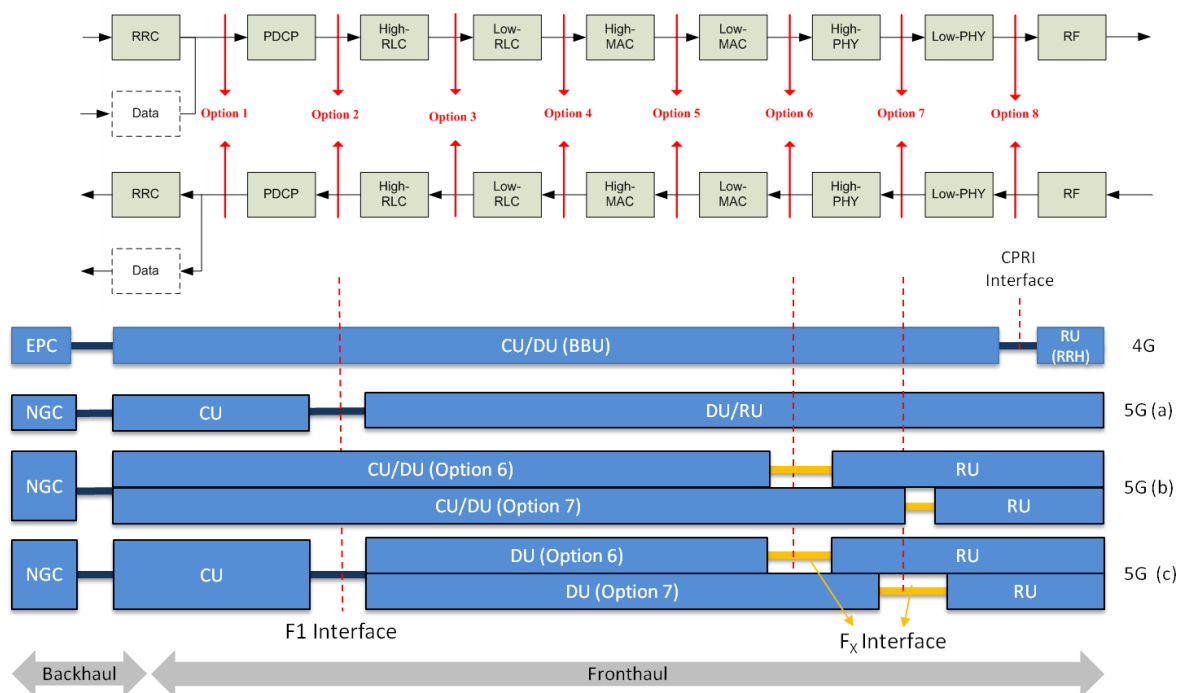
Figure 10 : Nœud de connexion NG-RAN



À travers le schéma suivant, nous pouvons apprécier le cheminement du signal entre les stations de base sans fil 4G et 5G et les points de partage en option, un lien nommé Midhaul permet la liaison entre les unités CU et DU.

Dans sa version 15 sortie récemment, la 3Gp l'a retenu comme option qui l'accompagne avec l'interface F1 entre les unités CU et DU.

Figure 11 : Options de découpage CU/DU [6]





1.1.2 Canaux physiques

Les canaux physiques de la technologie 5G NR sont les suivants :

- *Physical downlink shared channel* (PDSCH) ;
- *Physical downlink control channel* (PDCCH) ;
- *Physical broadcast Channel* (PBCH) ;
- *Physical uplink shared channel* (PUSCH) ;
- *Physical uplink control channel* (PUCCH) ;
- *Physical random-access channel* (PRACH).

1.2.3 Contrôle des ressources radio dans 5G NR

La couche radio ressource contrôle qui opère sur l'interférence radio 5G NR est connectée selon le schéma modélisant la structure d'interface radio, elle est de ce fait connectée aux quatre autres couches par le moyen d'accès du contrôle.

Au niveau du contrôle des couches de niveau 1 ou ce qu'on appelle les couches physiques, et le contrôle des couches de niveau 2 composées de MAC, RLL et PDCP, la couche RRC reste seule, et selon cette couche toujours que se passent les échanges d'information des unités distantes au cœur de l'équipement utilisateurs et la station de base.

Il est important de mentionner que les messages RRC sont traités selon les protocoles PDCP, RLC et MAC ainsi que la couche émettrice du signal avant qu'il soit transmis sur l'interface radio, ensuite reconstitué, vérifié et enfin interprété par l'entité distante RRC.

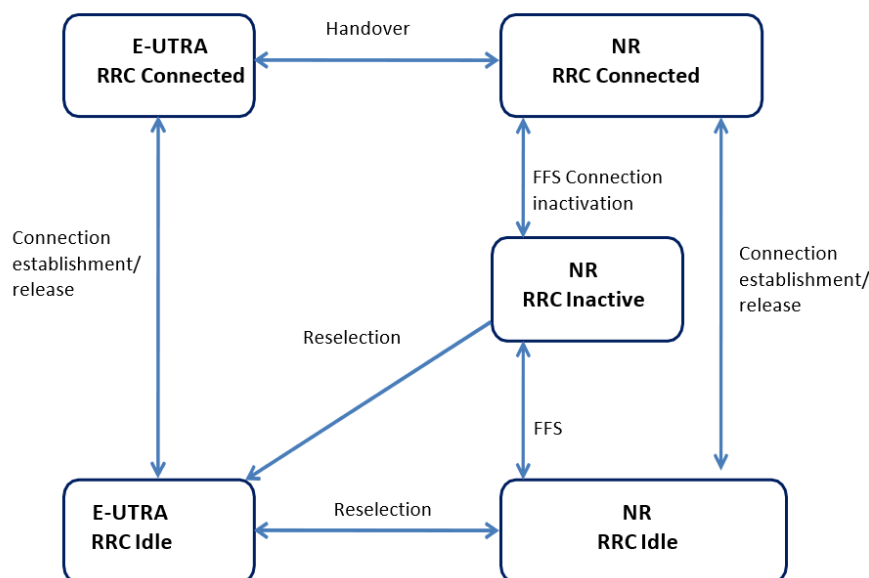
Lorsque la connexion entre la couche RRC active avec la station de base, la cellule utilisateur 5G NR se trouve coupée, cette dernière passe au mode RRC Idle mode, c'est-à-dire mode veille.

La cellule utilisateur décode régulièrement l'information système reçu de la station de base ainsi que les messages de paging, elle contrôle de ce fait de manière autonome sa mobilité.

Dans le cas contraire, le mode connecté (RRC *connected mode*), le protocole RRC est responsable de la gestion de la connexion active, ainsi que de la mobilité de l'équipement utilisateur, du transfert de la signalisation NAS, de la sécurité AS communément nommée gestion des clés et enfin des supports radio activés pour porter les données de service et / ou la signalisation (RC et NAS).

Le mode inactif est ajouté par la 5G comme troisième mode de connexion (RRC *inactive mode*) ou l'utilisateur passe au mode inactif au lieu de connecté après avoir stocké sa configuration au réseau de nouveau permet de déclencher la procédure de reprise, ce qui signifie la restauration de la configuration mis au stockage sans pour autant avoir besoin d'une signalisation étendue avec le réseau cœur.

Figure 12 : Transition des états RRC en 5G NR



Selon la figure 12, les technologies LTE (E-UTRA), UMTS (UTRA) et GSM/GPRS (GERAN), *new radio* (NR) censés interagir avec la technologie NR5G introduit alors la nouvelle condition RRC inactive pour être capable de satisfaire aux demandes du service 5G.

1.2.4 Formes d'ondes pour la 5G NR

Sensibles aux délais, les applications compatibles au fonctionnement en mode duplexage TDD passent en multiplexage en répartition orthogonale des fréquences (OFDM) qui reste la plus efficace pour la transmission multitrajet, ces derniers, sous forme de préfixe cyclique comme CP-OFDM, *cyclic prefix* OFDM, facilite son usage comme schéma d'accès pour le 5G NR et 3GPP.

Dans le cadre de la direction des liaisons montantes, la 5G NR est de ce fait capable de prendre en charge la CP-ORFM.

Cette dernière a pu transformé la porteuse relative à LTE (DFT-S-ORFDM) au multiplexage à division fréquentielle à une simple porteuse SC-FDMA *single carrier* FDMA pour permettre aux terminaux 5G des utilisateurs de choisir l'un de deux schémas à utiliser dans le cas des liaisons montantes.

Le préfixe cyclique (CP-OFDM) a plusieurs points forts :

- adaptabilité MIMO ;
- réduction du bruit de phase ;
- facilité de l'émetteur-récepteur ;
- efficacité du spectre élevée ;
- résistance à la sélectivité en temps et en fréquence de canal ;



- résistance aux erreurs de temporisation et aux interférences inter symboles.

L'équipement utilisateur peut prendre en charge les schémas de modulation suivants : 16QAM, 64QAM, QPSK et 256QAM, cependant le rapport constitué de la valeur de crête et de la valeur moyenne (PAPR, *peak-to-average power ratio*) est très élevé et constitue de ce fait un inconvénient majeur à cette technologie.

Le système de communication 5G devient plus flexible lorsque le schéma OFDM est combiné au schéma d'une seule porteuse de transmission DFT-S-OFDM, cette dernière est aussi connue sous le nom SC-FDMA permettant ainsi de moduler les sous-porteuses d'un émetteur avec les mêmes données de toutes les sous-porteuses, et permet aux utilisateurs de prendre en charge les schémas de modulateurs M/2- BPSK (16 QAM, 64 QAM et 256 QAM), sachant que le schéma M/2- BPSK nécessite une nouvelle adresse IP puisque c'est un nouveau schéma qui fait son entrée à la N. R.

1.1.5 Techniques d'accès au réseau 5G

Il est primordial pour la 5G de pouvoir maintenir l'orthogonalité maximale entre les utilisateurs dans les systèmes d'accès à travers des méthodes multiples adoptées :

- La technique d'accès OFDMA (*orthogonal frequency division multiple access*) ayant servi pour la 4G, elle exige une orthogonalité entre les sous-porteuses et l'usage d'un préfixe cyclique ;
- La nouvelle technique NOMA (*non orthogonal multiple access*) déjà employée par la 3GPP dans sa 16^e version, elle demeure une technique d'accès multiple, les utilisateurs sont capables alors d'utiliser toute la sous-porteuse à accès multiples sur la base d'allocation de puissance ;
- La technique SCMA (*sparse code multiple access*) est une combinaison de deux techniques OFDMA et CDMA permettant à un utilisateur donné d'être capable d'utiliser une seule ou plusieurs sous-porteuses dotées de code d'étalement utilisateur, applicable aux liaisons montantes et descendantes ;
- La technique non orthogonale MUSA opère dans deux domaines différents, celui du code et celui de la puissance.

1.1.6 Numérolgie pour la 5G NR

L'usage de la modulation OFDM comme schéma de transmission utilise le domaine de fréquence sous forme d'un nombre défini de sous-porteuse pour une durée de symbole donnée dans un domaine temporel.

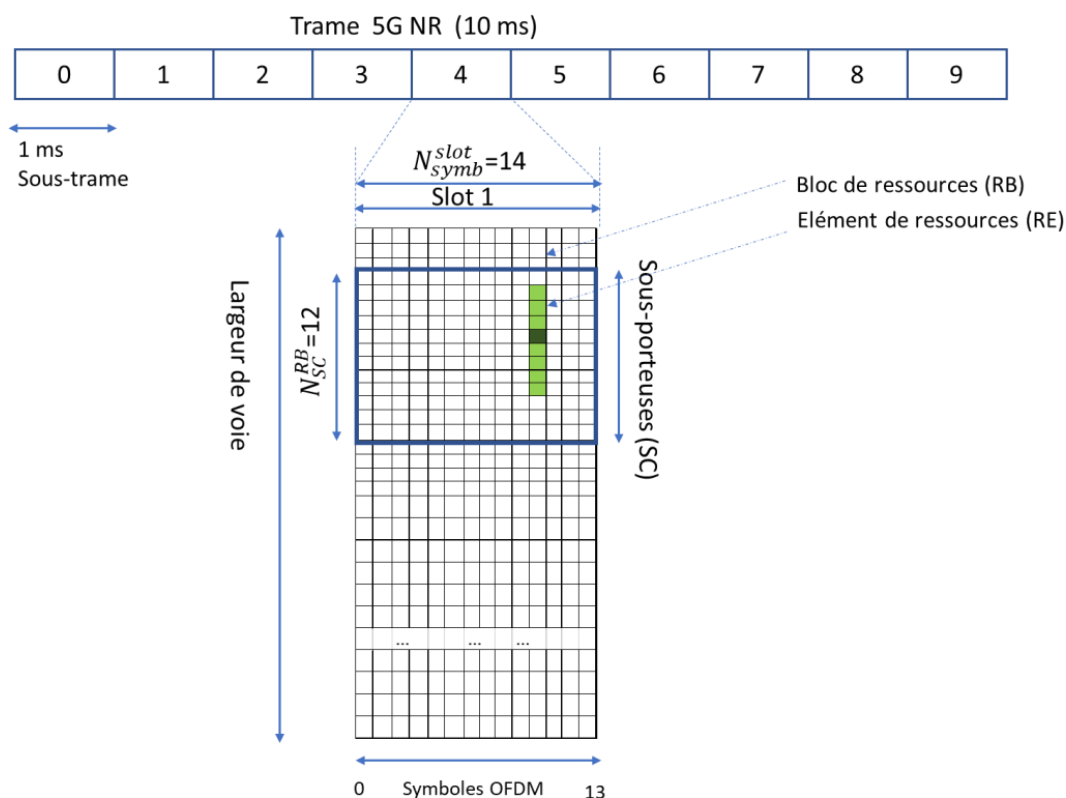
Cependant pour comprendre comment cela fonctionne pour la 5G on procède à la définition de ces éléments de schéma (figure 13) :

- un bloc de ressources RB (*resource bloc*) dans un domaine de fréquence

correspond à 12 sous-porteuses contiguës ;

- un *slot* dans un domaine temporel correspondant à 14 symboles consécutifs.

Figure 13 : Structure d'une trame 5G NR ($\mu=0$)



Un signal de référence est constitué d'une unité de ressources comme étant la plus petite pouvant être attribuée à un signal, c'est un symbole OFDM au niveau du domaine temporel et au niveau de domaine fréquentiel de la sous-porteuse.

Cet élément de ressource est identifié par les coordonnées (i, j) avec i coordonnée de sous-porteuse et j coordonnée du symbole OFDM ; au niveau de la grille temporelle par rapport à un point de référence relatif. Il est à noter par ailleurs qu'un bloc de ressources RB n'est qu'un alignement fréquentiel de douze sous-porteuses contiguës.

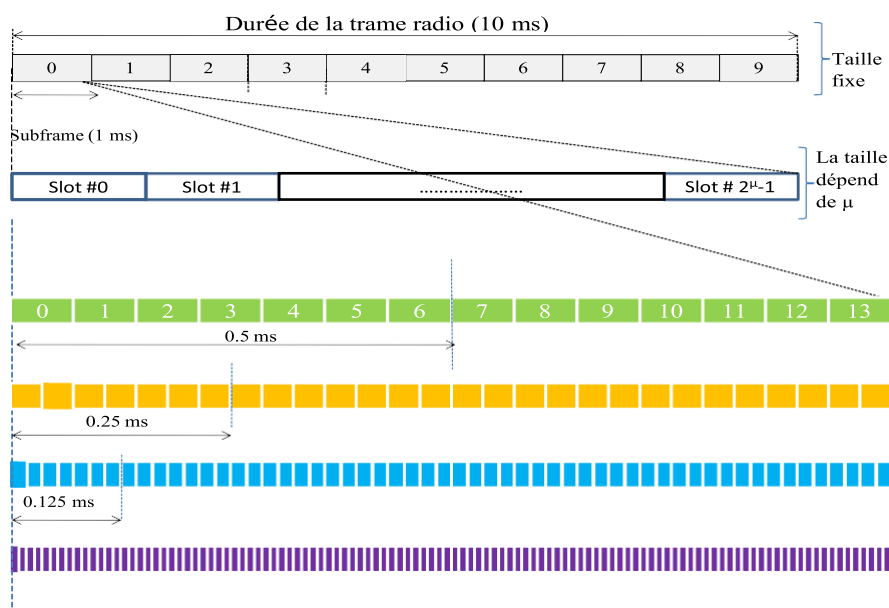
Les éléments ne peuvent être interprétés sans certaines explications complémentaires par rapport aux :

- espace plus large entre le porteur SCS ;
- taille du *slot* ;
- duplexage FDD et TDD ;
- espace plus large entre les porteuses (*high sub-carrier spacing, SCS*).

Le *slot* se définit comme une granularité minimale qui permet l'accès au canal à un utilisateur, chaque *slot* a une durée d'une milliseconde multipliée par dix et devient une trame composant la structure référencée à la 5G.

L'élargissement de la taille du *slot* permet d'accueillir plusieurs créneaux dans une seule milliseconde, c'est-à-dire qu'un créneau peut permettre la transmission de 15 KHz en SCS, et le double en deux créneaux soit 30 KHz, 60 KHz pour 4 créneaux et 120 KHz pour 8 créneaux, ce qui signifie que plusieurs opportunités sont possibles pour l'accès au canal lorsque la sous-trame contient plusieurs créneaux, chose primordiale à l'amélioration de la fiabilité.

Figure 14 : Structure d'une trame 5G NR



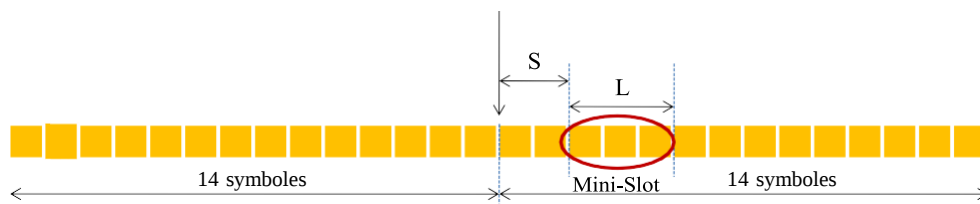
Taille de *slot* :

Le nombre de durée d'un *slot* configuré de manière standard avec 14 symboles dans la 5G comme dans la 4G, c'est-à-dire en répartition en petits *slots* de 2, 4 et 7 symboles (figure 15), ces derniers permettent le démarrage à n'importe quel symbole sans pour autant démarrer au début de la bordure du *slot*.

Plus courtes que ce qui est conventionnellement connu pour une durée de transmission, cette répartition par *slot* permet en même temps le démarrage instantané de transmission.

Un petit *slot* est reconnu par un identifiant du *slot* (K), un identifiant du symbole de départ par rapport au début du *slot* (S) et un identifiant du nombre de symboles consécutifs de S (L).

Figure 15 : Identification du *minislot*



Duplexage FDD et TDD plus flexible :

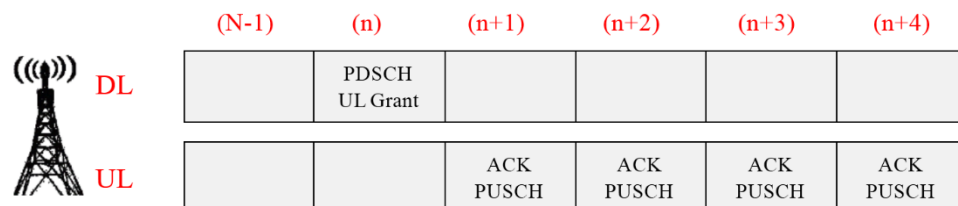
La question de l'accusé de réception qui reste requis dans la transmission d'un paquet est soulevée par la communication radio qui diffère entre la 4G et la 5G.

Optimisé en 5G par rapport au 4G, l'accusé de réception envoyé en sous-trame (n+4) après avoir transmis les informations dans un créneau n , alors que ce même accusé est transmis dans la sous-trame (n+1), (n+2), (n+3) ou au (n+4) dans les mêmes conditions d'envoi en 5G dans le créneau n , le calcul du dispositif mobile est la charge du réseau restent déterminants pour ce paramètre.

En 5G une trame est sensée se composer de 10 sous-trames, pour un total de 10 ms, c'est-à-dire une 1 ms par sous-trame, alors que chaque sous-châssis est constitué de 2 emplacements (figure 16) composés de 14 symboles OFDM chacun ou 12 symboles préfixes cycliques (CP étendu).

Ici le CP et la distance maximale intersites sont liés.

Figure 16 : Duplexage FDD et TDD



Cette dernière est expliquée sous sa forme (ISD) exposée dans le tableau suivant :

Tableau 5 : Structure de la trame temporelle

μ	Espacement de sous-porteuse	Type de préfixe cyclique (CP)	CP (μ s)	Nombre de slots par trame	Durée Slot (ms)	Distance intersite (ISD) (en km)
0	15	Normal	4,69	10	1	1,41
1	30	Normal	2,34	20	0,5	0,70
2	60	Normal	1,17	40	0,25	0,35
2	60	Étendu	4,16	40	0,25	1,25
3	120	Normal	0,59	80	0,125	0,18
4	240	Normal	0,29	160	0,0625	0,09

Le calcul de l'espacement de la sous-porteuse SCS (*subcarrier spacing*) respecte la formule :

$$SCS = 2^\mu * 15 \text{ kHz avec } \mu=0, 1, 2, 3, 4) \quad (3)$$

La puissance μ peut être égale à un nombre entier naturel de 0 à 4.



Les techniques déjà citées sont toutes des fonctionnalités importantes intégrées à la norme 5G, elles réduisent de ce fait la latence et adhèrent aux normes industrielles pour la version 4.0.

Par exemple, pour $\mu=0$ l'espacement est de 15 kHz et pour $\mu=1$ l'espacement est de 30 kHz.

1.1.4 Spectres de fréquences pour 5G NR

C'est par une gamme large et bien répartie que le spectre des fréquences 5G NR est composé, on peut distinguer :

- la gamme inférieure à 1GHz qui assure la diffusion du faisceau de couverture sur de grandes étendues territoriales que ce soit en intérieur ou en accès difficile, en zones urbaines ou en zones rurales ce qui permet une excellente couverture y compris au sein des bâtiments avec un nombre minimum d'antennes ;
- à un degré moindre, la gamme de 700 MHz n'offre pas une bonne compatibilité avec la nouvelle technologie 5G NR, surtout en *beamforming*, cependant leurs propriétés en IoT restent intéressantes avec un gain important en latence ;
- la gamme 3,5 GHz, avec un débit et une latence équilibrée et une fiabilité élevée, cette fréquence qui s'étend de 3,4 GHz à 3,8 GHz reste le meilleur compromis entre portée, fiabilité et débit ;
- l'ultra-haut débit ou la gamme des ondes millimétriques s'étend sur une fréquence de 26 GHz permettant ainsi l'exploitation d'une nouvelle large gamme jamais exploitée auparavant avec un nouveau type d'antenne et de nouvelles technologies de transmission accordant ainsi au débit une très haute performance.

Les fréquences sont utiles pour assurer une meilleure qualité du signal pour de bons services destinés à une forte agglomération gourmande en débit (internet des objets, réalité augmentée), néanmoins cette gamme manque en couverture souffrant de blocage devant les obstacles épais comme les murs et des rangées d'arbres, ou les perturbations climatiques comme la pluie, la densité du brouillard ou les vents, ce qui invalide l'idée selon laquelle l'usage de cette technologie serait urbain en premier lieu sauf si il y a multiplication d'installation d'antennes de petites tailles et de manière très rapprochées à travers des microantennes ou microcellules dont la liaison se fait moyennant les antennes macrocellules.

L'usage de haute fréquence signifie ainsi le recours à la courte longueur d'onde dont la portée est plus faible mais capable d'avoir un important débit, de ce fait nous concluons que chaque technique a des possibilités différentes et pour cela on distingue les fréquences par groupe pour les fréquences inférieures à 6 GHz des groupes d'ondes millimétriques.

L'usage de fréquences de 3,4 à 3,8 GHz rend possible le découpage de ces bandes en 50 et 100 MHz de large renforçant de ce fait l'accès à de nouvelles gammes de fréquence qui permettent d'avoir de grandes longueurs de bandes supérieures à celles existant actuellement.

Avec les ondes millimétriques on peut faire un découpage sur une largeur de 400 MHz, ce qui signifie un gain important en débit et par conséquent la nécessité d'adapter le matériel et l'équipement à ces nouvelles techniques.

1.2 Technologie 5G pour l'internet tactile

Les communications sans fil sont actuellement le moyen reliant les capteurs de tout type d'appareils, multimédias, médicales, la communication à travers les réseaux ne cesse d'avoir un débit qui croît de façon exponentielle grâce à l'innovation dans les domaines électroniques et la cybernétique.

Cependant, la latence reste ultra-faible^[9-12] pour les réseaux 5G face à la grande densité de connexion massive, la nécessité d'une fiabilité ultra-élevée et une grande capacité pour fournir la demande en débit.

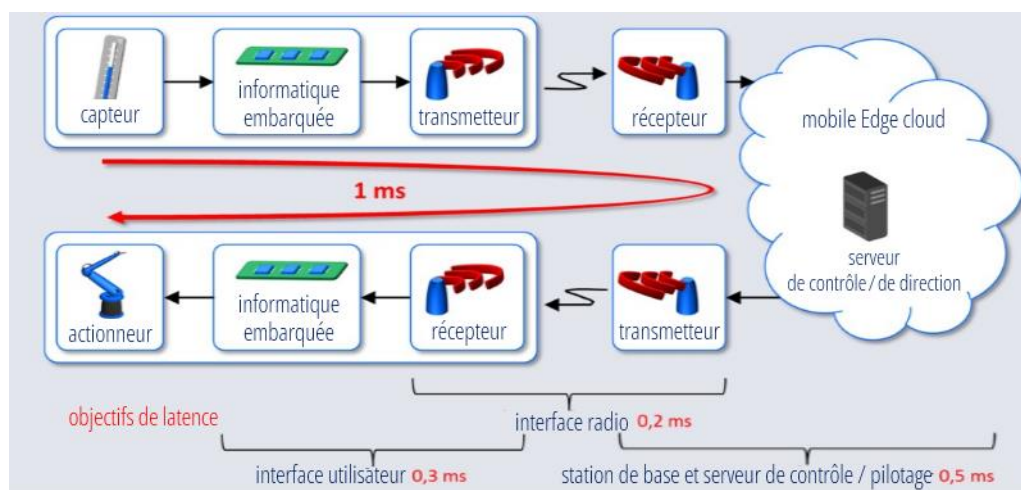
Alors que l'objectif de la technologie est de se doter d'une latence de plus en plus faible qui accompagne le changement de la façon avec laquelle communique l'humanité.

L'innovation technologique permet de rendre internet sous une forme tactile, une solution qui facilite la communication des domaines de santé, éducation, industrie et militaire, un moteur de développement et de croissance voire de l'innovation elle-même pour tous les secteurs y compris pour la technologie du *cloud* afin de la rapprocher de l'IoT réduisant ainsi la latence à travers ce qui s'appelle le Mobile Edge Computing^[13].

C'est une transformation immense que propose l'internet tactile en changeant le paradigme selon lequel la compétence et les main d'œuvre sont transférés numériquement partout dans le monde au lieu de diffusion du contenu existant^[14-17], ce qui intègre de nouveaux concepts d'interactions pour devenir capable de toucher et d'action à distance^[18-21].

L'Union internationale des télécommunications dans son rapport de 2004^[22] définit ce concept selon un schéma de capteurs et d'actionneurs permettant ainsi d'avoir la latence nécessaire de bout en bout par unité de milliseconde (figure 17).

Figure 17 : Exemple de latence d'un système de l'internet tactile^[22]



Cependant il est primordial de discuter des contraintes liées à l'évolution de l'internet tactile, pour évoquer la faible fiabilité et de la latence pour pouvoir caractériser les applications et progiciels industriels comme ceux de la surveillance de grandes structures avec des capteurs dispersés partout dépassant souvent les 100 à 1 000 connectés au réseau.



Selon les caractéristiques stochastiques du canal utilisé pour la connexion sans fil au service de la norme uRLLC^[23], il est capital de contrôler la variation de la force du signal traversant le canal afin de s'assurer de sa fiabilité, nous pouvons alors évoquer la diversité comme technique de réponse^[24].

1.2.1 Internet tactile et la question du débit maximum réalisable

L'ultra-fiabilité comme exigence stricte des services de la 5G uRLLC est lié à une latence faible qui pourvoit aux exigences de l'internet tactile, d'où la nécessité de procéder à la modification de la formule de Shannon, la capacité se voit alors modifiée selon^[25] :

$$C(P_e) = W \left[\log_2(1 + SINR) - \sqrt{\frac{V}{D_{tx}W}} Q^{-1}(P_e) \right] \text{ bit/s} \quad (4)$$

Avec $SINR = \frac{\alpha P_t g}{I + N_0 W}$ est le rapport signal sur bruit plus interférence, W est la bande passante, P_t est la puissance d'émission, α est le gain moyen du canal qui capture la perte de chemin et l'ombrage, g est le gain instantané normalisé du canal, N_0 est la densité spectrale de bruit d'un seul côté, I est l'interférence totale, Q^{-1} est l'inverse de la fonction Q-gaussienne et V est la dispersion de canal.

1.2.2 L'Internet tactile et la fiabilité maximale

L'UIT est défini comme étant « la capacité de transmettre une quantité donnée de trafic pendant une durée prédéterminée avec une probabilité de succès élevée »^[26].

La possibilité du succès d'une unité de donnée de protocole (PDU) de couche 2 de 32 octets demeure à 1 ms afin de respecter le minimum requis de fiabilité qui est de 99,999 % pour une qualité de canal du bord de couverture dans un environnement de test urbain Macro uRLLC, cette opération n'est réalisable qu'après validation des indicateurs de performance dans les conditions suivantes :

- trop de paquets sont perdus ;
- trop de paquets arrivent trop tard ;
- les paquets contiennent des erreurs.

Si nous notons R comme indicateur de fiabilité, il peut être exprimé comme suit :

$$R = (1 - \sum_{k=1}^{nHO} t_k)(1 - P_e)P_{wd} \quad (5)$$

où nHO est le nombre d'événements de *handover* (HO), t_k est le MIT pour le $k^{ième}$ événement HO, P_e est la probabilité d'erreur de transmission d'un paquet, P_{wd} est la probabilité de réception d'un paquet dans un certain délai.

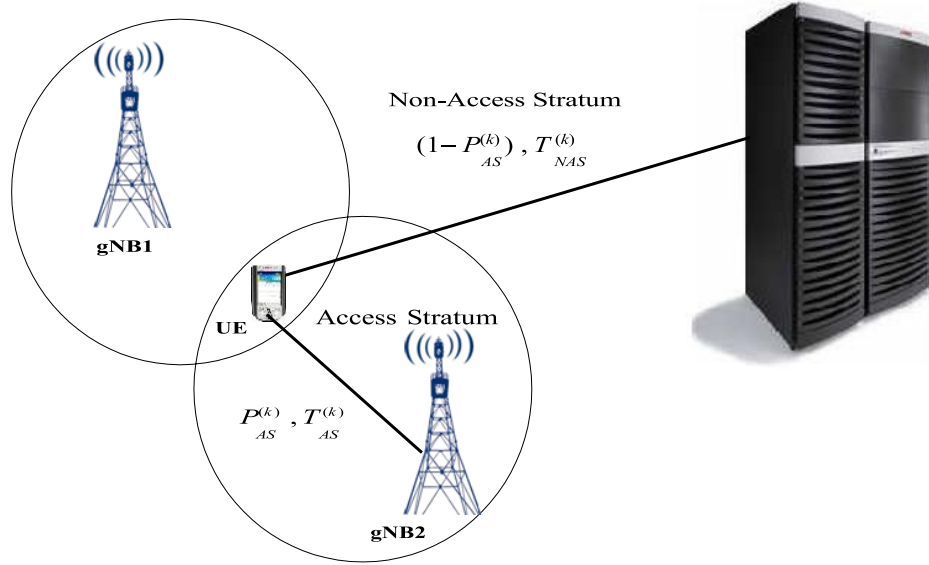
C'est alors que le MIT englobe le temps nécessaire pour exécuter toute procédure de réseau d'accès radio, protocole de signalisation de commande de ressources radio ou autres échanges



de messages entre la station mobile et le réseau d'accès radio, comme applicable aux technologies d'interface radio candidates.

Le rétablissement de la connexion se fait à travers les strates d'accès (AS, *access stratum*) comme premier moyen ou par les strates de nonaccès (NAS, *non-access stratum*) selon la figure 18^[29].

Figure 18 : Récupération de l'AS et du NAS après l'échec du $k^{\text{ième}}$ *handover*



Le MIT pour le $k^{\text{ième}}$ événement de *handover* peut être obtenu comme suit :

$$t_k = (1 - P_{HOF}^{(k)}) T_{SHO}^{(k)} + P_{HOF}^{(k)} T_{CR}^{(k)} \quad (6)$$

où $P_{HOF}^{(k)}$ est la probabilité d'échec du *handover* intercellulaire, $T_{SHO}^{(k)}$ est le MIT lors d'un *handover* intercellulaire réussi et $T_{CR}^{(k)}$ est le temps de récupération de la connexion après un échec de *handover* intercellulaire. Dans l'équation (6) $T_{CR}^{(k)}$ est calculé comme suit :

$$T_{CR}^{(k)} = P_{AS}^{(k)} T_{AS}^{(k)} + (1 - P_{AS}^{(k)}) T_{NAS}^{(k)} \quad (7)$$

$T_{AS}^{(k)}$ est le MIT dans une extraction de la strate d'accès (AS) après le $k^{\text{ième}}$ échec de *handover* intercellulaire, $P_{AS}^{(k)}$ est la probabilité de succès de la récupération de la strate d'accès après le $k^{\text{ième}}$ échec de *handover* intercellulaire et $T_{NAS}^{(k)}$ est le MIT dans un NAS récupération après le $k^{\text{ième}}$ échec de *handover*.

En fonction de la taille du paquet, de la qualité du canal et de la stratégie d'ordonnancement, la durée de transmission D_{tx} peut tout à fait varier d'un à plusieurs intervalles de temps de transmission (TTI). Pour les systèmes 4G, la signalisation de commande prend une grande partie de la latence de transmission, qui est généralement comprise entre 0,3 ms et 0,4 ms.

Ainsi, la conception d'un système de transmission par paquets de longueurs courtes et avec une latence de 0,5 ms peut entraîner un gaspillage de plus de 60% des ressources pour la



surcharge de contrôle. On suppose que le paquet transmis contient X bits, et le délai de transmission est D_{tx} . Par conséquent, en utilisant l'équation (1), nous avons :

$$X = D_{tx} C(P_e) \quad (bits) \quad (8)$$

$$\frac{X}{D_{tx}} = C(P_e) = W \left[\log_2(1 + SINR) - \sqrt{\frac{V}{D_{tx}W}} Q^{-1}(P_e) \right] \quad (bit/s) \quad (9)$$

Ensuite, la probabilité d'erreur de transmission d'un paquet de X bits est calculée comme suit :

$$P_e = Q \left[\sqrt{\frac{D_{tx}W}{V}} \left(\log_2(1 + SINR) - \frac{X}{D_{tx}W} \right) \right] \quad (10)$$

Nous désignons par $\rho = SINR$ et nous supposons que nous utilisons le mode FDD, dans cette technique l'émetteur et le récepteur fonctionnent à des fréquences porteuses différentes, donc il n'y a pas d'interférence entre les utilisateurs uRLLC et nous pouvons supposer que $I = 0$ dans ce cas. D'où,

$$\rho = \frac{\alpha P_e g}{N_0 W} \quad (11)$$

Nous posons également :

$$W = \frac{\xi}{\rho} \quad \text{avec} \quad \xi = \frac{\alpha P_e g}{N_0} \quad (12)$$

La dispersion de canal, V , est définie comme :

$$V = 1 - \frac{1}{(1+\rho)^2} = \frac{\rho(\rho+2)}{(\rho+1)^2} \quad (13)$$

Nous notons ici que la probabilité d'erreur de transmission (P_e) dépend étroitement de X et D_{tx} . Dans cette section, nous nous intéressons à l'étude de la variation de cette probabilité d'erreur de transmission en fonction de la variation de la longueur du paquet de données (X).

Cette probabilité est donnée par :

$$P_e = Q \left[\frac{\rho+1}{\rho\sqrt{\rho+2}} \sqrt{D_{tx}\xi} \left(\log_2(1 + \rho) - \frac{\rho}{D_{tx}\xi} X \right) \right] = Q(a - bX) \quad (14)$$

Et la dérivée de Q par rapport à X est donnée comme suit :

$$\frac{\partial P_e}{\partial X} = \frac{\rho+1}{\sqrt{2\pi(\rho+2)D_{tx}\xi}} e^{-\frac{(a-bX)^2}{2}} > 0 \quad (15)$$

Étant donné que cette dérivée est strictement positive, on en déduit donc que la probabilité d'erreur de transmission (P_e) est une fonction croissante avec X . Par conséquent, pour la transmission de paquets de données de très petites tailles, cette probabilité est également très faible ; et pour des paquets plus grands, cette probabilité devient également plus grande.

Par conséquent, dans le cas de l'internet tactile, qui nécessite une très grande fiabilité et un temps de latence très court et qui utilise donc des services uRLLC, les très petits paquets répondent parfaitement à la fiabilité de la transmission requise.



1.3 Optimisation de la latence pour les applications uRLLC

1.3.1 Cas d'une première transmission réussie

Le problème d'optimisation dans uRLLC est de minimiser la latence des utilisateurs uRLLC. La latence unidirectionnelle cible uRLLC, en supposant une première transmission réussie, notée D_I est calculée comme une combinaison du délai de mise en file d'attente noté D_q , du délai de traitement de la station de base noté D_{BSp} , le retard d'alignement de trame noté D_{Fa} , du délai de transmission noté D_{Tx} et du délai de traitement de l'équipement utilisateur noté D_{UEp} , comme suit :

$$D_I = D_q + D_{BSp} + D_{Fa} + D_{Tx} + D_{UEp} \quad (16)$$

L'arrivée des paquets uRLLC est modélisée comme un processus de Poisson et le modèle de mise en file d'attente est de type M/ M/1, avec une transmission uRLLC monocouche, est utilisée pour calculer la latence.

Il faut noter que le délai D_{Fa} est limité par la courte durée de l'intervalle de temps de transmission (TTI), qui est défini comme étant l'unité de temps minimum allouée pour toute transmission de liaison montante ou descendante et, par conséquent, sa durée a un impact sur la latence. D_{BSp} et D_{UEp} sont limités par la durée du symbole 3-OFDM [27], en raison des capacités de traitement améliorées pour la technologie 5G NR. Par conséquent, D_q et D_{Tx} deviennent le principal obstacle à l'atteinte du budget de latence pour le cas de services uRLLC. Certains paquets déclenchés par des événements générés par différents utilisateurs mobiles (UE) peuvent arriver intempestivement à une station de base (BS), et le temps entre les arrivées des paquets peut être plus court que le temps de transmission de chaque paquet, il y a un grand besoin de considérer le délai de mise en file d'attente. Selon [28-29], D_{Tx} dépend fortement de la qualité de la liaison radio, mesurée par le rapport SINR, pour les services uRLLC.

1.3.2 Cas de demande de retransmission HARQ.

Dans ce cas, la transmission et la retransmission HARQ entre l'équipement utilisateur (UE) et la station de base (BS), en termes de latence totale, D_{tot} , dans le plan utilisateur peuvent être modélisées comme suit :

$$D_{tot} = D_1 + p(D_2 + D_3) \quad (17)$$

où, p est la probabilité d'une retransmission, D_2 est la latence requise pour une requête HARQ et D_3 est la latence nécessaire pour retransmettre le contenu.

D_2 et D_3 sont calculés comme suit:

$$D_2 = D_{UEp1} + D_{Fa1} + D_{HARQ} + D_{BSp1} \quad (18)$$

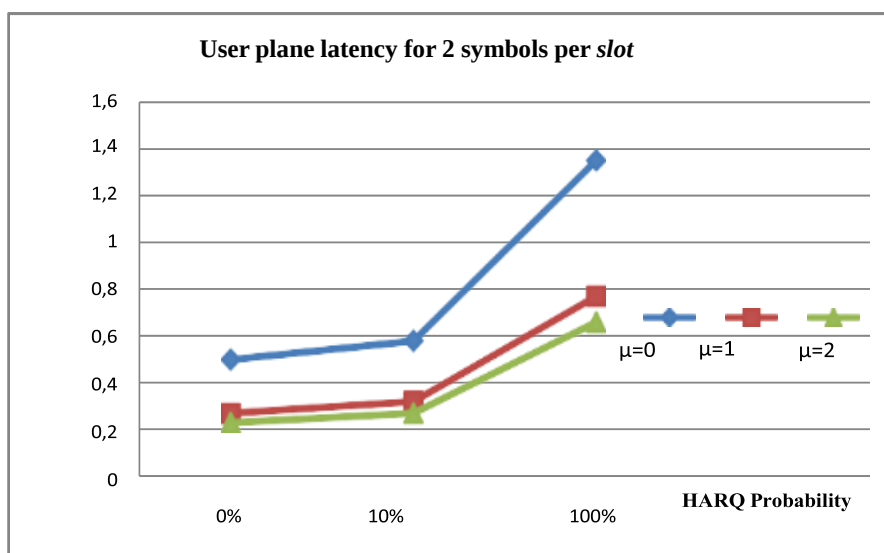
$$D_3 = D_{UEp2} + D_{Fa2} + D_{Tx2} + D_{BSp2} \quad (19)$$

où, D_{BSp_i} est le retard de traitement de la station de base i et D_{Fa_i} , D_{Tx_i} et D_{UEp_i} sont respectivement le retard d'alignement de trame, le retard de transmission et le retard

de traitement de l'équipement utilisateur i . D_{HARQ} peut être considéré comme égal à 1 symbole OFDM.

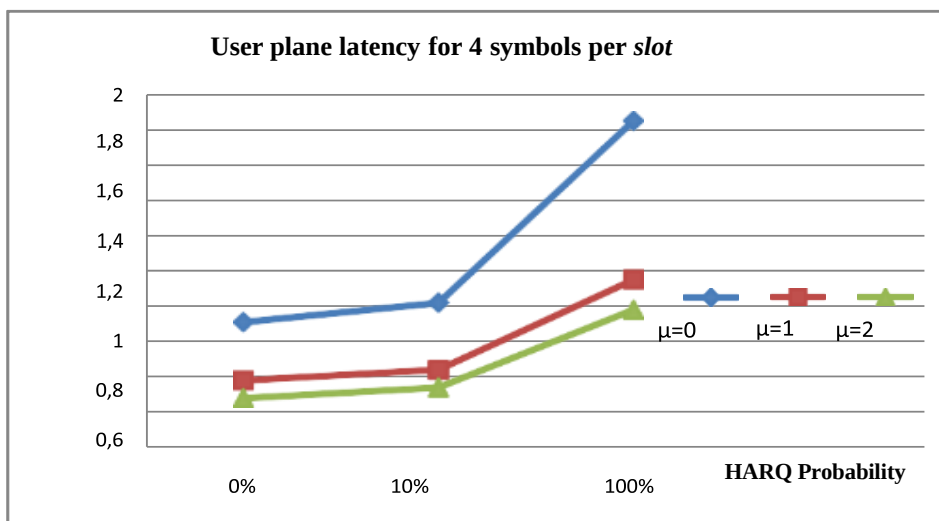
Dans ce paragraphe, nous donnons un exemple de simulation en considérant le cas d'un équipement utilisateur de capacité 2 et opérant en mode duplex FDD.

Figure 19 : Latence du plan utilisateur pour 2 symboles OFDM par *slot*



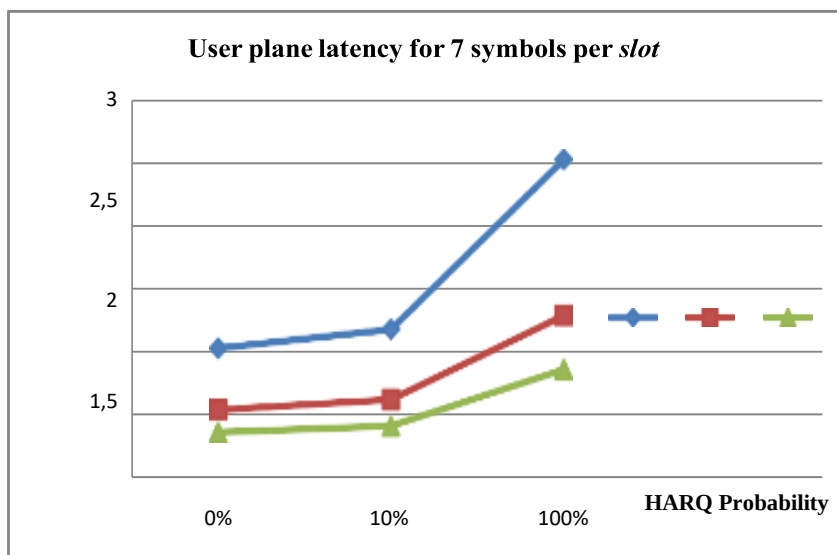
Sur la figure 19, nous pouvons voir que, dans le cas de 2 symboles par *slot*, la latence du plan utilisateur est tout à fait correcte en cas de transmission réussie (c'est-à-dire de probabilité de retransmission $p = 0$) et aussi avec $p = 10\%$. Par contre dans le cas de $p = 100\%$, la latence est légèrement plus élevée pour une séparation de sous-porteuses de 15 kHz. De sorte que 2 symboles par emplacement sont généralement utilisés dans uRLLC.

Figure 20 : Latence du plan utilisateur pour 4 symboles OFDM par *slot*



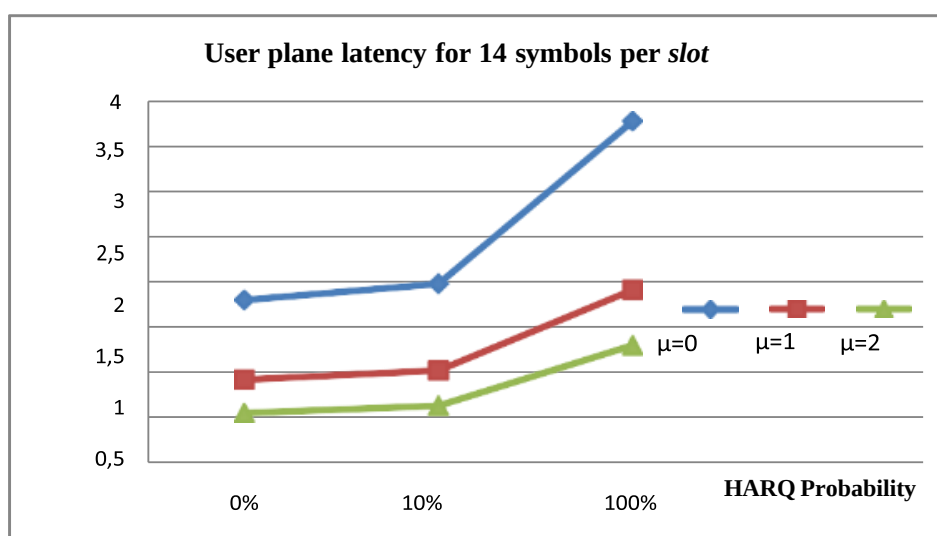
Sur la figure 20, nous pouvons voir que, dans le cas de 4 symboles par emplacement, la latence du plan utilisateur est également correcte en cas de transmission réussie (c.-à-d. la probabilité de retransmission $p = 0$) et aussi avec $p = 10\%$, mais avec une latence supérieure. Pour $p = 100\%$, la latence est plutôt plus importante pour une séparation de sous-porteuses de 15kHz . Par conséquent, nous pouvons également utiliser 4 symboles par emplacement dans uRLLC.

Figure 21 : Latence du plan utilisateur pour 7 symboles OFDM par slot



Sur la figure 21, on note que dans le cas de 7 symboles par slot , les latences du plan utilisateur commencent à dépasser la valeur attendue et surtout pour une séparation des sous-porteuses de 15kHz .

Figure 22 : Latence du plan utilisateur pour 14 symboles OFDM par slot





La figure 22 montre clairement que, dans le cas de 14 symboles par s/ot , comme dans la technologie 4G LTE, les latences dans le plan utilisateur deviennent plus importantes et dépassent le seuil attendu.

Conclusion

Nous présentons dans ce chapitre des définitions et présentations sur la nouvelle interface radio conçue pour la technologie 5G (5G NR) et ses diverses caractéristiques (numérologie, modulation, spectre de fréquences, etc.).

Ensuite, nous montrons que la technologie nécessaire pour garantir une latence aller-retour aussi faible que 1 ms est toujours à l'étude et n'est pas encore prête à relever les défis posés par des services très exigeants en termes de latence, de fiabilité et de qualité de service tel que l'internet tactile. À travers ce chapitre, nous montrons que la prise en compte de la transmission de petits paquets réduit également la probabilité d'erreur de transmission. Cette probabilité augmente à mesure que la longueur des paquets augmente. D'autre part, nous examinons également la variation de la probabilité d'erreur de transmission en fonction de la variation de la durée de transmission et nous constatons que cette probabilité décroît au-delà d'une certaine valeur de cette durée. La technologie 5G renforce donc ce concept d'internet tactile en utilisant son service uRLLC.

Dans ce chapitre, nous étudions également la latence dans le plan utilisateur entre la station de base et chaque utilisateur 5G NR, dans le cas où la transmission est réussie avec ou sans retransmission (HARQ). Ce chapitre montre que le concept de *minislot* introduit dans la 5G NR et l'espacement des sous-porteuses multiples (multiple de 15 kHz) vont contribuer à fournir des services exigeants en ce qui concerne la latence et la fiabilité, en particulier pour le cas d'utilisation URLL.

Chapitre III ♦ Réseaux 6G : vision, exigences, architecture, technologies et enjeux

Introduction

Notre société devient de plus en plus dépendante de tout ce qui est numérique. Par exemple, différents types d'objets physiques et virtuels sont connectés à l'internet des objets, tous les services sont numérisés, le nombre d'appareils connectés augmente, ce qui entraîne un échange de volumes de données énormes. Les réseaux de communication actuels 5G ne peuvent pas répondre aux demandes futures. Le besoin de communication mobile à très haut débit est donc vital pour mieux se préparer à l'arrivée de nouveaux services et d'applications émergentes. À savoir, la réalité étendue, la communication holographique, l'internet des sens, les jumeaux numériques humains, les villes intelligentes, l'industrie, etc. Ces nouveaux cas d'utilisation sont appliqués dans de nombreux domaines différents, par exemple la santé, le transport autonome, le climat, la cybersécurité, etc. Ainsi, les recherches sur les réseaux de nouvelle génération 6G sont lancées pour dépasser les limites de la 5G et relever les nouveaux défis. Ce chapitre présente l'étendue de la question à ce jour, l'état des lieux et le dessein en matière de réseaux de communication de sixième génération. Sont présentés d'abord la vision, les exigences et les scénarios d'application attendus des réseaux 6G ; sont ensuite décrites et analysées l'architecture intelligente et l'intégration des réseaux spatiaux, aériens, terrestres et maritimes, ainsi que les technologies clés potentielles les plus importantes nécessaires à la future sixième génération. Enfin, les principales activités de recherche lancées sont présentées.

1. État de l'art

Au cours des dernières années, l'évolution et l'émergence de la technologie ont changé et amélioré de nombreux systèmes dans divers domaines. Les systèmes de communication sans fil sont obligatoires et indispensables pour innover et conduire à de nouvelles technologies efficaces et faire évoluer celles qui existent déjà. Par conséquent, chercheurs et experts doivent s'interroger sur les besoins en systèmes de communication pour améliorer le domaine de la recherche technologique de la prochaine génération mobile^[1].

Avec la standardisation complète et le déploiement du système 5G dans plusieurs pays^[2], la recherche sur les systèmes de communication mobile 6G a commencé. Puisqu'une nouvelle génération de systèmes de communication mobile apparaît tous les dix ans environ, Cisco s'attend à ce que les systèmes 6G soient déployés commercialement avant 2030, il y aura alors 14,4 milliards d'appareils connectés^[3].

Il faut généralement plus de dix ans pour qu'une nouvelle technologie soit lancée à des fins commerciales. Par conséquent, comme nous l'avons mentionné précédemment, nous devons passer par la recherche en nouvelles technologies 6G et relever de nouveaux défis.

La force de la 5G vient de sa variété de services tels que le haut débit mobile amélioré (eMBB), les communications à faible latence ultra-fiables (URLLC), les communications massives de type machine (mMTC) qui fournissent respectivement des débits de données élevés de 1G bits/s pour les mobiles, connectivité, fiabilité de 99,99 %, millisecondes de latence de la communication et un million de connexions au Km²^[4].

Cependant, avec les nouvelles technologies émergentes, la demande en 2030 sera 100 fois supérieure au nombre actuel et aux applications à venir. De nouvelles technologies innovantes et des dispositifs internet des objets apparaissent, ils envoient et partagent une énorme quantité de données produites par la technologie hologramme, la réalité étendue, l'imagerie haute résolution, les véhicules autonomes, etc. Tout cela génère des volumes de données de plus en plus colossaux, ce qui oblige à envisager de nouveaux systèmes de communication dans divers domaines : santé, transports, climat, cybersécurité, secteurs civil et militaire, etc. Pour mettre en œuvre ces services IoT, les futurs réseaux de génération sans fil doivent fournir une vitesse de transmission élevée, une grande fiabilité et une faible latence que les systèmes 5G ne proposent pas.

Tous ces éléments limitent l'objectif initial de la 5G IoT. Par conséquent, il existe le besoin urgent d'un système sans fil 6G dont la conception s'adapte aux exigences de performance des applications IoT et aux tendances technologiques qui l'accompagnent.

Le système 6G est censé être une révolution exceptionnelle qui le distingue significativement des autres générations. En outre, on s'attend à ce que la 6G renforce les services d'intelligence artificielle (IA) omniprésents du cœur aux appareils finaux du réseau. Par conséquent, l'IA joue un rôle important dans la conception et l'optimisation des architectures, des protocoles et des opérations 6G^[5], ^[6].

Dans la littérature, divers articles discutent de réseau mobile 6G, en terme de vision, d'exigences, d'applications, de solutions potentielles et de défis en recherche et développement^{[6]-[10]}, les classes de service potentielles et l'architecture réseau de la 6G sont analysées dans [9], [11]–[14] et les technologies habilitantes clés dans [15] et [16].

En outre, plusieurs technologies et idées spécifiques à la 6G sont expliquées et discutées en fonction de différents aspects tels que la communication holographique, la communication intelligente, etc.

Par le prisme de la vision, des tendances, des exigences, des défis, des scénarios d'application, de l'architecture de réseau, des technologies clés potentielles de la 6G, 12 catalyseurs sont discutés et comparés entre l'Europe, les États-Unis, la Finlande, le Japon et la Chine.

Une première section présente ensuite la vision, les exigences et les scénarios d'application de la 6G, la deuxième section décrit l'architecture du réseau 6G, la troisième étudie les technologies clés potentielles de la 6G, et la quatrième et dernière section présente de brèves conclusions.

2. Vision, exigences et scénarios d'application

2.1 Vision et exigences

La 5G contribue au développement de l'internet mobile vers l'internet des objets. La 6G augmente considérablement la capacité et l'efficacité des réseaux de communication mobile, elle élargit et approfondit encore la portée et le champ des applications de l'internet des objets et est compatible avec les nouvelles technologies de l'information et de la communication (TIC) telles que l'intelligence artificielle et les mégadonnées. Il est envisagé que la nouvelle génération de communications mobiles 6G atteigne les caractéristiques suivantes.

a) De solides performances

En premier lieu, la réalisation de réseaux de communication 6G nécessite des bandes de fréquences plus élevées qui atteignent 1 à 3 térahertz (THz), ce qui peut prendre en charge les cas d'utilisation anticipés de la 6G. Puis, la latence de la 6G doit être aussi faible que 0,1 ms et tendre vers zéro dans certains cas, comme les futures applications de télé-présence. Enfin, la densité des objets connectés est supposée supporter 10 millions de connexions au km². Selon ces trois aspects, les performances de la 6G sont améliorées de 10 à 100 fois par rapport à la cinquième génération.

b) Plus intelligent

Pour être plus adaptable et robuste, la 6G doit s'autogérer. Par conséquent, l'entrée de technologies telles que l'intelligence artificielle, les techniques d'apprentissage automatique et le *big data* sont très pratiques pour permettre aux nœuds du réseau d'avoir des capacités intelligentes qui sont hors du contrôle humain et aussi pour parvenir à une auto-organisation complète, à l'auto-apprentissage, à l'auto-guérison, à l'auto-agrégation, à l'auto-protection et à l'auto-optimisation du réseau. De plus, l'IA aide la 6G à fournir des services intelligents connectés hautement personnalisés pour les particuliers, les industries et d'autres utilisateurs afin de répondre à des exigences raffinées.

c) Plus vert

Tout en améliorant les performances du réseau, le coût et la consommation d'énergie doivent être efficacement réduits, et l'efficacité énergétique du système par bit doit être multipliée par dix, voire cent. Ceci est censé être réalisé en soutenant le concept de développement vert et de récupération d'énergie.

d) Couverture plus large

La couverture du réseau 6G doit s'étendre des réseaux terrestres aux réseaux non terrestres (NTN) et comprendre des connexions air-espace-sol-mer très étroitement connectées pour atteindre une couverture mondiale profonde aussi prégnante que l'air ambiant.

e) Plus sûr

Grâce à la conception de signaux physiques, à la conception d'architecture, à la conception de protocoles et aux applications de nouvelles technologies telles que la *blockchain* et la communication quantique, la sécurité du réseau 6G est assurée et la fiabilité des communications comme la sécurité des informations sont améliorées.

f) Source ouverte

Le réseau 6G réalise la décentralisation et l'aplatissement, les équipements de réseau et les produits terminaux réalisent des plates-formes et des logiciels *open-source* et construisent un environnement écologique industriel plus ouvert et plus équitable.

2.2 Scénarios d'application

Avec la vision et les exigences du réseau 6G susmentionnées, les scénarios d'application de celui-ci seront plus étendus que la 5G. L'intégration de nouvelles technologies permet de réaliser de nouvelles applications et de nouveaux cas d'utilisation. Cette section présente les applications envisagées au cours du déploiement des réseaux 6G.

a) Communication holographique

L'une des applications les plus critiques de la 6G est la télé-présence holographique^[17]. Elle permet aux humains de se connecter à distance et de communiquer avec de vraies images-objets 3D. Les hologrammes sont utilisables dans de nombreux scénarios tels que les communications sociales, les jeux de divertissement, la conception de bureaux, la télémédecine, etc. Les réseaux 5G ne sont pas en mesure de réaliser de tels scénarios avec une grande fiabilité, surtout avec un débit de données élevé qui pourrait atteindre 4 Tbits par seconde et qui impose une latence ultra-faible.

b) Industrie 4.0

L'industrie 4.0 à l'ère de la 6G connaît une transformation à grande échelle. De plus en plus d'usines intelligentes intègrent le mode de fabrication intelligent et de collaboration entre l'homme, la machine et le matériel. Les robots intelligents remplacent les humains et les robots existants deviennent la force principale de la fabrication agile, la fabrication industrielle devient plus autonome et intelligente. Le développement de la nanotechnologie fournit une nouvelle façon de surveiller les aspects et les processus de production industrielle. Les nanorobots peuvent faire partie intégrante des produits et peuvent les contrôler tout au long de leur cycle

de vie. Les plans de production industrielle, de stockage et de vente sont analysés sur la base de données de marché dynamiques en temps réel, garantissant ainsi un maximum d'avantages pour la production industrielle. À cette fin, il est nécessaire que les réseaux 6G atteignent un haut niveau de fiabilité et une très faible latence.

c) Jumeaux numériques humains

En termes d'applications de jumeau numérique humain, la 5G met principalement en œuvre des fonctions telles que la surveillance de la santé humaine et la prévention primaire des maladies. Avec les progrès révolutionnaires de la théorie de la communication moléculaire et des technologies clés telles que les nanomatériaux et les capteurs, cette application réalise davantage la numérisation du corps humain et l'intelligence des traitements médicaux.

Par la reconstruction numérique du corps humain dans le monde réel, les jumeaux numériques humains construisent un « humain numérique » personnalisé dans le monde virtuel. Grâce à la surveillance de la santé et à la gestion des « humains numériques », une surveillance complète et précise des signes vitaux humains, une thérapie ciblée, des recherches pathologiques et la prédiction du risque de maladie grave peuvent être effectuées, protégeant ainsi des vies humaines et les maintenant en bonne santé.

Les jumeaux numériques humains peuvent devenir une réalité à l'ère de la 6G, compte tenu des exigences élevées en temps réel et de fiabilité, d'un accès internet haut débit dans les airs avec une mobilité élevée et des exigences de couverture complète.

d) Internet des sens

L'internet des sens est une sorte de réseau de transmission d'expériences qui relie les sens multidimensionnels pour réaliser une communication inter-sensorielle. Grâce à l'infrastructure interconnectée, les gens peuvent pleinement mobiliser la vision, l'ouïe, le toucher, l'odorat, le goût et même les émotions et réaliser la transmission et l'interaction à distance de ces sentiments importants. Peu importe où l'on est, on peut vivre une expérience immersive en musique, art, sport et d'autres compétences dans un environnement réel, on peut ressentir la vraie nourriture, l'expérience d'essai de soins de la peau sans consommer d'objets réels.

e) Réalité étendue

En particulier, la 6G serait très efficace pour les services XR (AR (réalité augmentée), VR (réalité virtuelle) et MR (réalité mixte)).

Avec les dernières technologies telles que AR, VR et MR, le contenu de grande capacité tel que la vidéo et 3DCG (*three-dimensional computer graphics*) devrait être fourni en tant que service, donc un trafic à grande vitesse et de grande capacité est nécessaire. Avec la réalité virtuelle et la RA sensible à l'espace, le contenu doit être modifié en fonction du mouvement de l'utilisateur, donc l'immédiateté est également requise.

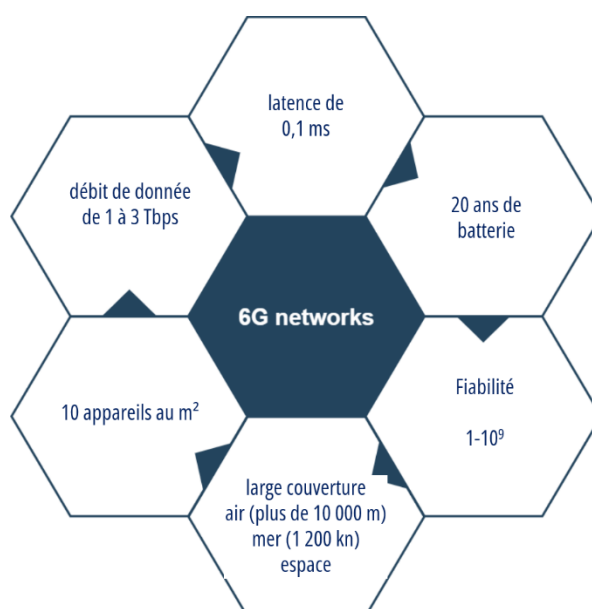
La prolifération d'applications XR épuisera le spectre 5G, en particulier pour certains cas d'utilisation sensibles tels que la chirurgie à distance. Il lui faut un système avec une très faible

latence inférieure à 0,1 ms et une grande capacité de transmission. AR / VR ne peut pas être compressé (le codage et le décodage est un processus qui prend du temps), par conséquent, le débit de données par utilisateur doit atteindre le gigabit par seconde, contrairement à l'objectif plus souple de 100 Mb/s de la 5G.

Cette variété de scénarios d'application est une caractéristique unique de la norme 6G, dont le plein potentiel ne peut être exposé qu'à travers une architecture réseau solide et des avancées technologiques révolutionnaires, comme décrit dans les sections suivantes.

3. Architecture et défis du réseau 6G

Figure 23 : Exigence des réseaux 6G



3.1 Défis 6G

Sur la base de la vision, des exigences et des scénarios d'application de la 6G, on s'attend à ce que les réseaux 6G réalisent la connexion sans fil transparente de la société à l'échelle mondiale, intégrant la communication, l'informatique, la navigation et la détection, ainsi que l'air, l'espace et le sol dans un aspect tridimensionnel avec des opérations et une maintenance intelligente et autonome. Les réseaux d'IA peuvent fournir une hyperconnectivité intelligente à très haute capacité, quasi instantanée, fiable et illimitée^[11].

3.1.1 Réseau intégré espace-air-sol-mer

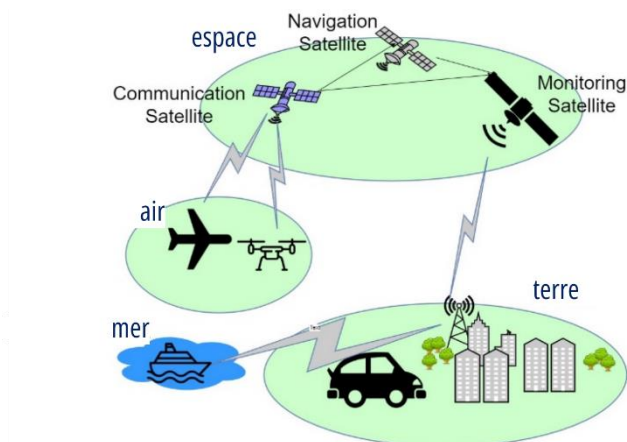
Le réseau terrestre actuel ne peut pas étendre la profondeur de la portée de communication. Dans le même temps, le coût des connexions à l'échelle mondiale est très élevé. Pour prendre en charge la couverture complète du réseau et le déplacement à grande vitesse des utilisateurs, la 6G optimise l'infrastructure des réseaux aériens, spatiaux, terrestres et maritimes, en

intégrant des réseaux terrestres et non terrestres pour fournir une couverture complète et illimitée.

Les réseaux spatiaux basés sur les communications par satellite fournissent une couverture sans fil pour les zones non desservies et non couvertes en déployant de manière dense des satellites en orbite (voir figure 24). Le réseau aérien à basse altitude peut être déployé plus rapidement et reconfiguré de manière flexible, pour être adapté à l'environnement de communication et révéler de meilleures performances pour les communications à courte distance.

Le réseau aérien de plates-formes à haute altitude peut être utilisé comme nœud relais dans les communications longue distance pour favoriser l'intégration de réseaux terrestres et non terrestres. Le réseau terrestre prend en charge la bande de fréquences térahertz et sa couverture réseau extrêmement réduite atteindra la limite de l'amélioration de la capacité du système. Les architectures réseaux « de-cellular » et réseau ultra-dense centré sur l'utilisateur^[6] émergeront au fur et à mesure que les temps l'exigeront. Les réseaux sous-marins fournissent une couverture et des services internet pour les activités en haute mer pour des applications militaires ou commerciales. Cependant, il existe une controverse quant à savoir si les réseaux sous-marins peuvent faire partie du futur réseau 6G.

Figure 24 : Architecture des réseaux non terrestres



3.1.2 Réseaux intelligents

Pour concrétiser la vision de systèmes sans fil 6G endogènes intelligents omniprésents, la conception de l'architecture 6G doit pleinement tenir compte de la possibilité d'une intelligence artificielle dans le réseau, ce qui en fait une caractéristique inhérente à la 6G.

Ces dernières années, l'IA et l'apprentissage automatique (ML, pour *machine learning*) suscitent une large attention dans l'industrie. L'intelligence artificielle est appliquée à de nombreux aspects des réseaux cellulaires 5G, y compris les applications de la couche physique (telles que le codage et l'estimation des canaux), les applications de la couche MAC (telles que l'accès multiple), les applications de la couche réseau (telles que l'allocation des ressources et la détection des erreurs). Cependant, l'application de l'IA dans les réseaux 5G est limitée à l'optimisation des architectures de réseau traditionnelles car ils n'ont pas pris en compte l'IA au

début de la conception de l'architecture. Il est donc difficile de réaliser pleinement le potentiel de l'IA à l'ère de la 5G.

L'intelligence artificielle est une réalisation de l'IA perceptive, incapable de répondre à des situations inattendues. Avec la diversification des besoins de service et la croissance explosive du nombre d'appareils connectés, le réseau est devenu un système hétérogène extrêmement complexe. Par conséquent, il existe le besoin urgent d'un nouveau type de réseau d'IA avec auto-agrégation, auto-organisation, auto-optimisation, auto-adaptation et auto-raisonnement. Il devra non seulement intégrer l'intelligence dans l'ensemble du réseau, mais devra également intégrer la logique de l'IA dans la structure du réseau, de sorte que la perception et le raisonnement interagissent systématiquement, et enfin tous les composants du réseau peuvent être connectés et contrôlés de manière autonome, et peuvent reconnaître et s'adapter aux accidents de façon autonome. Le déploiement de l'IA centralisée, de l'IA distribuée, de l'IA de périphérie, [<https://ieeexplore.ieee.org/ielx7/6287639/9312710/09335927.pdf>]

3.2 Architecture pour 6G

Alors que la feuille de route et les activités de normalisation doivent encore être lancées, quelques chercheurs proposent de nouvelles architectures de réseaux 6G. Dans cette sous-section, nous décrivons brièvement différentes architectures issues de la littérature. Dans la section précédente, nous avons décrit les caractéristiques critiques de deux architectures 6G basées sur l'IA. Parlons maintenant des autres architectures.

3.2.1 Architecture cyber-jumelle

Les auteurs proposent deux architectures pour les réseaux 6G, à savoir un modèle internet centré sur le *cloud* et un modèle RAN découplé^[104]. Dans le modèle internet centré sur le *cloud*, l'architecture IP existante est légèrement conservée avec certaines modifications. Tout d'abord, les utilisateurs sont connectés au RAN et les données du RAN entrent dans la couche *cloud* périphérique où les cyber-jumeaux acceptent les données. Ces cyber-jumeaux agissent comme des enregistreurs de données, des propriétaires d'assentiment ou une représentation virtuelle de l'utilisateur. Les *clouds* périphériques, à leur tour, sont connectés à la couche *cloud*, où plusieurs *clouds* sont interconnectés pour former le centre de l'architecture du réseau. La couche *cloud* comprend les fonctions de planificateur de ressources, d'orchestrateur, de communication, de calcul et de mise en cache. Cette couche *cloud* permet aux applications de fournir des services en périphérie et dans le *cloud* à un coût et une qualité de service raisonnables.

3.2.2 Architecture RAN découplée

Voici des points d'accès distincts pour gérer les données de contrôle et de plan utilisateur différentes des générations cellulaires précédentes. Le plan de contrôle BS (station de base) est une macro BS, à laquelle un utilisateur se connecte pour échanger les informations de contrôle. Ces BS de commande seront connectées aux BS du plan utilisateur (données) pour la signalisation de haut niveau. De plus, le trafic de liaison montante et de liaison descendante est

également découplé et géré par des BS distinctes pour les utilisateurs. Ici, la BS de liaison montante peut être un déploiement dense de microcellules à proximité de l'utilisateur pour collecter les données de l'utilisateur à faible puissance. Les stations de base de liaison montante et descendante auront une coordination interne pour une communication efficace^{[104], [105]}.

3.2.3 Architecture généralisée pour la 6G

Une architecture 6G généralisée pour l'IoT et les réseaux véhiculaires est présentée dans [139]. Il se compose de trois niveaux, à savoir le niveau utilisateur composé d'appareils intelligents avec capacité de mise en cache. Ces appareils envoient les données détectées au niveau de la station de base, où les stations de base ou les points d'accès disposent de serveurs périphériques pour effectuer la planification et l'allocation des ressources. Enfin, un serveur central effectuera les actions de découpage et de transfert au niveau du réseau. Cette architecture réduit le délai pour les services critiques (voitures autonomes) en traitant et en stockant au niveau de la périphérie. Cependant, en raison de son plan de contrôle partiellement centralisé, certaines tâches non critiques utilisent toujours le traitement dans le *cloud*.

3.2.4 Architecture de haut niveau de la 6G

Une architecture de réseau 6G à trois niveaux se compose d'un plan AI, d'un plan utilisateur et d'un plan de contrôle^[140]. Dans ce cas, le stockage, le calcul et la mise en réseau sont effectués au même niveau (dans le plan utilisateur) pour éliminer la hiérarchie. Le plan utilisateur est plat et défini entre un réseau d'accès et internet. Le plan de contrôle et les plans AI sont distribués et virtualisés pour divers services. De plus, le réseau de transport est virtualisé et isolé du reste par la virtualisation définie par logiciel. Les fonctions du réseau central sont réalisées sous forme de microservices et accessibles par des systèmes sans serveur.

3.2.5 Architecture multiniveau pour la 6G

Une architecture à trois niveaux de la 6G se compose d'utilisateurs intelligents, d'appareils périphériques et du *cloud*^[114]. Les utilisateurs mettent en œuvre des techniques de prise de décision intelligentes telles que des techniques basées sur des données ou des modèles pour prédire les modèles de mobilité et de trafic. Au niveau de l'intelligence de périphérie, les appareils mobiles de périphérie utilisent l'apprentissage par renforcement profond ou des réseaux de neurones profonds pour optimiser le planificateur pour l'allocation des ressources aux utilisateurs mobiles en fonction du CSI. De plus, au niveau du *cloud*, disposer d'un contrôle central de grande capacité peut entraîner le système avec une plate-forme numérique qui présente les échantillons étiquetés pour les algorithmes d'optimisation. Ceux-ci peuvent être utilisés pour entraîner des réseaux de neurones profonds et être ensuite implémentés dans le plan de contrôle.

3.2.6 Architecture 6G tridimensionnelle^[4]

Le FG-NET 2030 imagine une architecture 3D pour les réseaux 6G, couvrant trois aspects clés, à savoir la communication (vue infrastructure), l'intelligence (vue contrôle) et la gestion (vue réseau). Il dispose de différentes couches de communication au niveau de l'infrastructure, allant du terrestre, sous-marin, aérien et au satellite pour améliorer la gamme de communication dans les réseaux 6G. En outre, les auteurs de la vue de contrôle présentent comment les réseaux 6G incluent l'intelligence pour contrôler et optimiser les fonctions globales de détection, d'accès au spectre, de communication, de stockage et de traitement à l'aide de l'IA, de l'apprentissage en profondeur et du ML. Il recommande que l'intelligence soit répartie entre diverses entités du réseau. Enfin, dans la vue réseau, les fonctions du réseau global ont été divisées en sous-couches. Ils incluent la sous-couche application, routage, gestion, accès au spectre et support physique, qui ressemble à la pile IP en couches. Une version simplifiée de la même architecture pour la 6G est proposée dans [108].

3.2.7 Architecture 6G basée sur les neurosciences

Cette architecture est un cadre qui tente d'intégrer des signaux neuro pour émuler les signaux sans fil à appliquer dans les réseaux 6G^[136]. Ici, l'intelligence du cerveau humain et les propriétés de rayonnement des signaux sans fil seront intégrées pour permettre une communication intelligente entre l'humain et les ordinateurs. Les avancées récentes en matière de bio-IoT et d'appareils de communication implantables nous incitent à imaginer une communication réseau 6G à courte portée où des modules sans fil implantés à l'intérieur du cerveau humain acquièrent l'intelligence du cerveau et communiquent directement avec les appareils sans fil et la station de base du monde extérieur ou avec un autre être humain avec une capacité similaire.

3.2.8 Architecture proposée

L'architecture 6G proposée se compose d'appareils au niveau de la couche externe. Les appareils de cette couche collectent des données, des événements dans le réseau et communiquent avec le niveau suivant. Ici, les nœuds IoT, les utilisateurs cellulaires, les appareils intelligents, etc., qui ont été connectés au réseau d'accès radio peuvent avoir un certain niveau d'intelligence par défaut. Par conséquent, les interférences, la sélection des canaux, la détection sont mieux gérées. Nous décrivons également le paradigme du réseau sans cellule, des petites cellules et du MIMO massif activé par des surfaces intelligentes fonctionnant côte à côte, comme ce sera le cas dans la 6G. Les flèches dans cette couche indiquent les liaisons sans fil. De plus, dans la couche suivante, nous avons des appareils périphériques et un *cloud* qui utilisent l'IA pour l'allocation des ressources, la gestion des utilisateurs et d'autres tâches d'optimisation. Comme proposé dans [42], nous considérons une architecture intelligente de périphérie distribuée qui fournit des installations de processus, de stockage, de calcul et de prise de décision pour des réseaux physiques indépendants. Les agents intelligents de ces appareils périphériques apprennent des appareils pour aider à une meilleure gestion du réseau en dessous, au niveau de la couche 1. Par exemple, il peut s'agir de la sélection du canal sans fil et du réglage des paramètres en conséquence. Ici, un nœud périphérique centralisé fournit un service de liaison à tous les réseaux physiques via ces nuages périphériques, comme indiqué

par les lignes discrètes. Il comprend également des ressources cloud-RAN et de liaison pour divers cas d'utilisation. Cependant, toutes ces ressources sont virtualisées. De plus, nous avons divers nuages qui interconnectent les éléments du réseau. Enfin, au niveau de la couche la plus interne, nous avons les applications pour divers secteurs verticaux qui fournissent ou reçoivent des services des couches inférieures. Ainsi, l'architecture de réseau fournit la vue abstraite de la 6G. Cependant, plusieurs autres fonctionnalités telles que l'informatique quantique, les nœuds de brouillard, la sécurité, etc., ne sont pas montrées en exclusivité et sont les candidats potentiels de l'architecture 6G.

3.3 Tendances futures de la 6G

Le projet 6GFlagship identifie 11 domaines clés du projet pilote 6G et publie des livres blancs en juin 2020^{[36], [42], [44]–[55], [107]}. Dans cette section, nous résumons l'ambition globale de certains des domaines clés tels que la localisation et la détection, la confiance et la sécurité, les objectifs de développement durable des Nations unies, la connectivité rurale et la mise en réseau, l'intelligence de pointe, l'apprentissage automatique et la portée commerciale de la 6G.

3.3.1 Localisation en 6G

La détection et la localisation traditionnelles utilisant le GPS ou la coordination cellulaire deviendront obsolètes dans les futurs environnements urbains intelligents à haute densité. Dans les futurs réseaux de 2030, les KPI liés à la précision doivent inclure des informations sur l'environnement de l'utilisateur pour faciliter un positionnement précis. Par exemple, un cas d'utilisation des réseaux 6G tel que la conduite autonome où un positionnement haute résolution au niveau du centimètre est nécessaire le long de la voie de circulation ou des détails de la voie pour assurer la sécurité routière^[19]. De même, les applications mobiles basées sur des capteurs qui utilisent des informations de localisation doivent inclure la consommation d'énergie par unité de couverture de distance ; Les KPI liés à l'intégrité et à la confidentialité doivent inclure les détails de localisation, les nœuds défectueux et les alarmes^[53]. Ces exigences garantissent que les futures technologies ou dispositifs pour les réseaux 6G soient accompagnés d'informations supplémentaires telles que la consommation d'énergie et les types d'alarmes générées lors de toute panne de capteur pour une localisation exacte en mettant l'accent sur la confidentialité. En outre, la localisation et la détection d'information ont plusieurs applications à l'avenir, notamment la chirurgie robotique, la recherche de contacts pendant les pandémies, les jeux basés sur la réalité virtuelle, les réseaux sociaux et les applications de rencontre, la livraison de nourriture, le marketing contextuel, la conduite autonome, la navigation personnelle et le suivi d'animaux. Plus précisément, dans le cas de services contextuels, le réseau 6G apprend automatiquement les besoins des utilisateurs en détectant l'environnement ou l'emplacement et en fournissant la meilleure qualité de service et d'autres paramètres^[124]. Par conséquent, les systèmes 6G intègrent l'intelligence, la mise en réseau omniprésente, ainsi qu'une localisation de précision et une détection à haute résolution pour répondre aux cas d'utilisation spécifiques qui émergent en 2030, nécessitant un ajustement précis de la précision au niveau du centimètre. Cela permet une connectivité, une localisation et une détection transparentes pour les services contextuels^[19]. La 5G NR et les

ondes millimétriques offrent des fonctionnalités de localisation et de détection différentes de la génération précédente de réseaux cellulaires. Les réseaux 6G améliorent encore la précision de la localisation conduisant à une imagerie haute définition grâce à l'utilisation de la communication THz, d'antennes massives et des SIR^{[53], [91]}. Lorsque la largeur du faisceau diminue à des longueurs d'onde inférieures (ondes μm), la précision de positionnement augmente. Ensuite, la relation inverse entre la fréquence et la taille de l'appareil (antenne) permet un regroupement dense de plusieurs antennes, ce qui facilite une estimation précise de l'angle et de la direction. À cet égard, de meilleurs algorithmes d'estimation de canal sont nécessaires aux fréquences THz pour détecter et localiser. De plus, un débit de données plus élevé du réseau 6G permet le partage de cartes entre les appareils beaucoup plus facilement. Avec les grands SIR, l'onde réfléchie peut fournir un meilleur service de localisation en exploitant l'effet de champ proche et en analysant le front d'onde. De plus, cela supprime le besoin de synchronisation entre les stations de référence. Cependant, des méthodes plus précises sont nécessaires pour la localisation et le positionnement en temps réel si les surfaces intelligentes sont de plus petite taille^[91]. Il convient de noter que les techniques de localisation existantes utilisant le *machine Learning* (ML, en français : apprentissage automatique) dépendent fortement des méthodes d'empreinte digitale, de régression et de classification. Cependant, des méthodes de localisation et de cartographie plus intelligentes sont nécessaires pour le réseau 6G^[53]. Par des *frameworks* basés sur l'IA et le ML, le système de détection peut être formé pour apprendre à partir de données brutes, ou des informations vitales (temporelles et spatiales) présentes avec le signal radio bruyant, ou des données de capteur faibles peuvent être extraites pour analyser l'emplacement et les modèles de détection pour le système de réseau 6G. Cependant, la plupart des techniques basées sur ML et le *deep learning* (DL, en français : apprentissage profond) nécessitent un grand volume de données structurées pour la formation qui peuvent ne pas être disponibles (bruyantes, aléatoires) dans de minuscules applications basées sur des capteurs. Par conséquent, des méthodes d'analyse avancées doivent être introduites pour aider à la localisation formée. Comme mentionné précédemment, la localisation est également associée à un autre ensemble de tâches, à savoir l'imagerie et la détection qui impliquent des signaux à haute fréquence (60 GHz). Ces signaux lorsqu'ils sont réfléchis par les objets, la taille et la dimension des objets peuvent être mesurées avec précision^[127]. En outre, la localisation comprend également l'imagerie à l'aide de capteurs. Ici, des capteurs passifs captent les images et des capteurs actifs transmettent des signaux qui représentent la distance, l'angle et le *doppler* avec une résolution et une précision élevées. Par exemple, dans le cas des voitures autonomes, l'imagerie radar ajoute la distance et le *doppler* à l'image multidimensionnelle^[53].

3.3.2 6G de confiance

En quoi la sécurité dans la 6G est différente de la 5G ? le nombre d'appareils connectés (IoT, MTC et utilisateurs mobiles) dans les réseaux 6G augmentera à un rythme galopant et en termes de densité, de dix millions d'appareils dans une zone unitaire d'un km^2 ^[2]. Ces appareils hériteront de l'intelligence artificielle et s'étendront à des applications diversifiées telles que la banque, l'industrie, la santé, le gouvernement, les transports et bien d'autres. En raison de la dépendance à de nombreux secteurs verticaux, les réseaux 6G doivent appliquer un niveau de sécurité plus élevé à chaque étape de leur réseau pour empêcher toute forme d'attaque de

sécurité. Dans un autre cas, nous pouvons imaginer un monde hyperconnecté, où un compromis mineur dans les paramètres de sécurité peut entraîner la mort dans une industrie automatisée, une chirurgie robotique effectuant une incision incorrecte, etc.^{[55], [50]}. Actuellement, dans les réseaux 5G, la cryptographie traditionnelle répond à la plupart des besoins de sécurisation de la communication de données à travers les architectures *cloud* et Edge. Cependant, une nouvelle ère de sécurité commencera avec l'avènement de l'informatique quantique dans le réseau 6G. Par exemple, les systèmes QKD facilitent la détection des écoutes clandestines par rapport à la cryptographie classique. De plus, l'amélioration du niveau de sécurité est possible avec les communications directes sécurisées quantiques (QSDC), qui offrent une meilleure sécurité sur le canal quantique que le QKD ou les méthodes classiques^[120]. Dans une autre optique, l'automatisation des fonctions de sécurité avec des algorithmes ML à différents niveaux sera essentielle en raison de la densification du réseau. Cependant, il existe une menace d'attaque de sécurité inversée basée sur ML, car le réseau peut apprendre tous les paramètres à la minute près. Par conséquent, le réseau 6G nécessite une architecture de sécurité réseau holistique exclusive^[55].

Caractéristiques de sécurité requises :

Les fonctionnalités clés doivent inclure en particulier, les futures cartes SIM et leurs aspects de sécurité et l'utilisation de la cryptographie asymétrique, le transport *layer security* (TLS) utilisant la cryptographie à courbe elliptique (ECC), le logiciel et la sécurité basée sur l'IA, comme applicable au SDN et au NFV. Par exemple, le modèle de sécurité dans les réseaux 6G consiste en une passerelle VNF activée pour l'apprentissage en profondeur qui surveille le trafic d'entrée et de sortie et applique des politiques de filtrage pour détecter et prévenir les attaques de sécurité potentielles. En résumé, un modèle de confiance doit définir les règles pour collecter, traiter, distribuer et filtrer les données dans le réseau.

3.2.3 6G pour les objectifs de développement durable (ODD)

La numérisation due aux réseaux 5G répond actuellement à de nombreux problèmes sociaux tels que l'éducation, la protection de l'environnement et la faim. Néanmoins, il existe plusieurs problèmes qui nécessitent une urbanisation et une connectivité importantes. Étant donné que les réseaux 6G apportent une connectivité hyperdonnées basée sur l'expérience utilisateur, plusieurs problèmes mondiaux, en particulier en ce qui concerne les ODD des Nations unies, sont atténués^[86]. Le réseau 6G, grâce à sa capacité de communication à multiples facettes (longue et courte portée, connectivité 3D), a le potentiel de promouvoir la durabilité mondiale, ce qui est principalement dû à sa connectivité transparente et à la prise en charge d'une pléthore de services. Voyons quelques-unes des futures portées des réseaux 6G en ce qui concerne la promotion des objectifs de développement durable. Par exemple, les services bancaires en ligne permettront à quiconque d'accéder facilement aux services financiers, ce qui contribue à éradiquer la pauvreté. De même, les services de santé en ligne réduisent le temps de déplacement et atteignent facilement les nécessiteux chaque fois que nécessaire pour promouvoir une meilleure prestation des soins de santé. En disposant d'une connectivité réseau complète, les agriculteurs des zones rurales peuvent utiliser les transactions numériques pour leurs produits agricoles et domestiques. L'IdO est mis en œuvre dans la



production alimentaire pour améliorer le rendement et éliminer la faim^[93]. Avec l'aide des réseaux 6G, tous ces problèmes sont résolus à l'échelle mondiale. Lorsque nous considérons l'énergie zéro, le réseau 6G cible la miniaturisation des appareils pour augmenter l'efficacité énergétique et la récupération d'énergie, améliorant ainsi les performances environnementales. Par exemple, ces appareils peuvent maintenir l'énergie générée par les activités quotidiennes telles que la marche, le jogging, et les tâches ménagères de l'utilisateur de l'appareil. Cette énergie doit prendre en charge les dispositifs d'informations personnelles tels que les bracelets, les patchs intelligents qui surveillent de temps en temps les signes vitaux du corps d'une personne. De plus, dans un autre scénario, ceux-ci peuvent répondre aux besoins d'information et de divertissement grâce à une connectivité sur le dessus. Pour soutenir les ODD des Nations unies, le réseau 6G définit sa vision pour fournir une opportunité à la société et à l'économie mondiale d'accomplir ses rêves numériques^[44]. Il comprend la connectivité des données, une économie forte, facilite les soins de santé intelligents, l'efficacité énergétique pour n'en nommer que quelques-uns parmi les 17 ODD. La vision 6G cible trois piliers pour mettre en œuvre les ODD par le biais de divers services. Ce sont les suivants :

- Répondre aux problèmes des personnes et de la société en fournissant une infrastructure numérique appropriée pour les autonomiser.
- Sensibiliser au contexte basé sur l'environnement qui est possible grâce à des capacités de localisation et de détection très précises, à la surveillance en ligne des ressources, etc.
- Améliorer globalement l'écosystème des ODD. Outre ces exigences, la nécessité fondamentale pour atteindre l'objectif sera la connectivité complète et l'accès à internet.

Comment la 6G soutient un ODD :

La 6G couvre tous les secteurs de la vie, la vision des réseaux 6G doit être pluridisciplinaire. Par conséquent, une plate-forme ouverte et de co-innovation pour les activités de normalisation technologique de la 6G qui soutiennent la vie humaine et l'environnement est cruciale pour créer un écosystème qui profite à toutes les parties prenantes des ODD de l'ONU. L'écosystème 6G contribue au progrès et au bien-être de diverses communautés de la société. Par exemple, la 6G fournit un réseau sans fil pour connecter le patient, le médecin et l'équipement médical ensemble dans un système de santé intelligent à distance. Avec l'aide du réseau 6G, plusieurs services numériques peuvent être étendus tels que l'éducation, les services de distribution alimentaire, la banque et l'industrie pour progresser à grande échelle. Ceux-ci aboutissent finalement à l'autonomisation de la vie humaine vers la réalisation de meilleurs ODD. Cela dit, il y aura des préoccupations tout en ciblant les ODD tels que l'inclusion, la sécurité des données. Par exemple, en ce qui concerne la collecte de données, toutes les parties prenantes telles que les hôpitaux, la société et le gouvernement doivent gérer les données avec les règles de confidentialité les plus strictes pour l'intégrité du système. Désormais, les réseaux 6G jouent un rôle dans la protection de la confidentialité et de la sécurité de la communication des données. En résumé, les réseaux 6G doivent mettre en œuvre des actions qui favorisent la durabilité dans différents secteurs verticaux, encouragent la croissance avec une consommation d'énergie minimale et mettent en œuvre le zéro déchet lorsque les personnes, les machines et les



ressources sont interconnectées. En raison de la prévalence de l'intelligence dans le réseau 6G, il améliore l'efficacité et la durabilité environnementale en utilisant des technologies ou des dispositifs à faible consommation d'énergie qui tirent l'énergie des activités naturelles pour leur fonctionnement.

Opportunité pour la 6G :

Les cas d'utilisation des réseaux 6G doivent être alignés sur les intentions des ODD, en particulier les ODD qui peuvent être bien soutenus par l'infrastructure des TIC. Les aspects communs tels que les activités financières en ligne, y compris les services bancaires en ligne et les affaires en ligne, favorisent la croissance économique. En outre, la télémédecine, la cartographie à distance des ressources naturelles telles que les minéraux et les masses d'eau sont des cas d'utilisation qui améliorent la durabilité sociale et sanitaire. Lorsque nous examinons la gamme diversifiée de connectivité sans fil de la 6G, l'intelligence de pointe, la localisation, la détection et la consommation d'énergie ultra-faible soutiennent fortement la connectivité rurale, le contexte environnemental et les aspects énergétiques des ODD.

3.2.4 Comment la 6G réduit-elle la fracture numérique ?

Il existe plusieurs raisons au manque de connectivité dans de nombreuses régions éloignées du monde, telles que l'absence d'infrastructures, les faibles revenus, la géographie difficile et bien d'autres^[92]. Le réseau 6G a une vision forte de la réduction de la fracture numérique qui, à son tour, favorise la durabilité parallèlement à la connectivité. On estime que près de 3,7 milliards de la population mondiale est loin de portée à la connectivité internet^[44]. La fracture numérique donne lieu à plusieurs autres problèmes tels que le retard de la croissance économique, l'éducation et la sensibilisation, inhibe l'adoption de technologies de soins de santé à distance et avancées. Cela entraîne également une faible croissance dans les secteurs de l'industrie moderne et des transports du nouvel âge. La fracture numérique n'épargne aucune région géographique. En réalité, même si la plupart des pays dispose d'une connectivité internet mobile de haute qualité, il existe des régions éloignées des États-Unis et d'Europe qui connaissent un faible débit de données de 0,4 à 1 Mbps via une connectivité filaire / sans fil. La situation est bien pire en Afrique, au Brésil et en Inde^[52].

Portée de la 6G dans la connectivité rurale :

La connectivité sur des terrains difficiles, loin des principales villes, doit être simple, peu coûteuse et facilement réalisable. Par conséquent, le haut débit mobile est la meilleure solution par rapport à l'internet filaire^[48]. La fourniture de la connectivité du dernier kilomètre via le satellite LEO doit couvrir une vaste zone. Cependant, ce ne sera pas une solution raisonnable, compte tenu du débit, de la latence et du coût. Au lieu de cela, l'utilisation d'émetteurs de grande puissance (mégacellules) dans les zones rurales à faible densité est une solution viable. Pour atteindre cet objectif, de nouvelles règles de sécurité et de nouvelles réglementations énergétiques sont nécessaires. Parallèlement, les stations de base mobiles flottantes comme les drones ou les ballons favorisent les services de données mobiles tout en agissant comme des relais ou des caches de données^[52]. Étant donné que les régions éloignées ont des



contraintes financières, tout en offrant une bonne qualité de service, il est obligatoire de mettre l'accent sur l'accessibilité en termes de coûts. L'utilisation de sources d'énergie naturelles telles que l'énergie solaire et éolienne pour les réseaux électriques qui exploitent les liaisons de retour permet de réduire les coûts. De plus, la réduction des coûts et l'efficacité des ressources peuvent être obtenues en limitant le contenu à la région rurale générée au sein de la localité. Cela nécessite de mettre en cache le contenu généré localement dans les serveurs locaux eux-mêmes et de provisionner le découpage du réseau. La mise en cache locale permet de réduire les coûts en minimisant le besoin de connexions de liaison. De plus, plusieurs schémas de service de données peuvent être mis en image grâce au découpage du réseau pour améliorer la connectivité avec des ressources limitées^[76]. En ce qui concerne la connectivité *backhaul*, même si la lumière visible ou les micro-ondes semblent être les bons candidats, ils peuvent être confrontés à des défis dans les zones rurales. Donc, L'accès et le *backhaul* intégrés (IAB) semblent très prometteurs lorsque l'opérateur utilise la partie du spectre pour le *backhaul* et la partie restante pour l'accès cellulaire à un coût raisonnablement inférieur à celui des anciens candidats^[52]. Ainsi, les réseaux 6G doivent s'appuyer sur l'IAB pour réduire le coût de déploiement et améliorer l'utilisation des ressources^[107]. Avec IAB, seul un sous-ensemble de stations de base achemine le trafic de liaison via des liaisons connectées (fibre / câble) ; tandis que les stations de base restantes sont connectées via des liaisons sans fil pour transmettre le trafic de liaison via un relais multisaut. C'est tout à fait possible en 6G en raison de la disponibilité d'une large bande passante dans les bandes THz où le trafic d'accès et de liaison peut être multiplexé par des schémas de planification appropriés, réduisant ainsi le coût de l'infrastructure. L'IAB fournit un service en sous-6GHz et en ondes millimétriques. De la même manière, Les satellites en orbite basse, les drones et les ballons trouvent des solutions utiles pour étendre la portée de manière transparente. En fonction du besoin de communication, ces terminaux peuvent être rendus opérationnels (transréception discontinue) au fur et à mesure des besoins : réduisant ainsi la consommation électrique. Dans l'ensemble, la communication non terrestre joue un rôle clé dans les réseaux 6G pour réduire la fracture numérique^[52]. Dans les tendances récentes, l'informatique de périphérie et la mise en cache des données locales pour une utilisation locale, l'installation d'une micro-infrastructure de télécommunication et son intégration avec l'opérateur de réseau mobile, ainsi que leur maintenance ont pris de l'importance dans les régions rurales. Cependant, tous ceux-ci ont plusieurs défis, tels que les modèles financiers, les réglementations de fréquence, les retards de propagation et les problèmes de coexistence^[48]. En fonction du besoin de communication, ces terminaux peuvent être rendus opérationnels (transréception discontinue) au fur et à mesure des besoins, réduisant ainsi la consommation électrique. Dans l'ensemble, la communication non terrestre jouera un rôle clé dans les réseaux 6G pour réduire la fracture numérique^[52].

3.2.5 Quelle est la portée de l'intelligence de bord en 6G ?

L'essor de l'informatique de périphérie et du déploiement de nœuds ultra-dense ou d'appareils mobiles conduit à l'émergence de l'informatique de périphérie mobile pour optimiser les calculs et les performances du réseau. La nature même des appareils mobiles nécessite que les calculs soient distribués. Un facteur limitant du réseau 5G est le manque d'intelligence dans l'informatique de périphérie mobile, qui facilite autrement une transition vers de nouveaux marchés verticaux et services tels que l'industrie 4.0, les mégapoles intelligentes et les gadgets



intelligents. L'une des raisons de ce revers est que les algorithmes d'IA existants consomment énormément de temps et de ressources. De plus, la production de matériel et de solutions d'IA n'a pas encore atteint le synchronisme. Au total, lorsque le réseau 5G est considéré comme utilisant l'IA pilotée par le *cloud*, le réseau 6G utilise l'IA pilotée par les bords^{[80], [81]}. L'Edge Intelligence (EI) est une capacité extrême pour traiter les données à la périphérie. Aujourd'hui, les appareils mobiles disposent d'une capacité de calcul élevée, d'un espace de stockage et exécutent de petites applications d'IA. L'IA devient partie intégrante du réseau 6G avec une intelligence distribuée pour exécuter des opérations autonomes et automatisées remplaçant la plupart des interactions humaines. Par exemple, des secteurs verticaux tels que les voitures autonomes, les soins de santé ultra-intelligents et les industries modernes impliquent des tâches complexes en matière d'acquisition et de calcul de données. Par conséquent, une intelligence distribuée à différentes entités du réseau est nécessaire pour offrir des services optimisés. Cela constitue la base de l'internet 6G des objets intelligents^[86]. L'IE peut être vu sous différentes directions :

- le point de vue de l'analyse des données fait référence à l'analyse des données et au développement de solutions sur le site ou à proximité du site où les données sont générées ;
- la perspective réseau cible principalement les fonctions intelligentes qui sont déployées à la périphérie du réseau, comprenant l'utilisateur, le locataire ou la frontière du réseau^[42].

Besoin de logiciel pour AI Edge :

Actuellement, nous manquons de fluidité logicielle ; sans cela, la localisation précise de l'intelligence dans l'entité du réseau et le partage des connaissances est compromis. EI est utile dans plusieurs marchés verticaux pour offrir une meilleure qualité de service. Par exemple, lorsque nous considérons le scénario des véhicules autonomes, ces véhicules doivent calculer et communiquer de nombreux facteurs, tels que la distance inter-véhiculaire, les entités en bordure de route, les panneaux de signalisation et les véhicules dans les voies à proximité. Ces données sont communiquées à la station de base ou aux points d'accès compatibles EI situés le long de la route pour stocker et traiter rapidement (faible latence) les données et aider à décider de la charge sur les bandes de fréquences disponibles (spectre) pour faciliter le transfert de données à partir de ces véhicules dynamiquement. Notamment, ces unités EI ont également accès à la base de données *cloud* centrale qui conserve les données globales de plusieurs systèmes régionaux d'assurance-emploi. En outre, les applications de ville intelligente, la réalité étendue mobile (XR) qui nécessitent un traitement et des calculs de données volumineux, peuvent bénéficier de l'EI en raison du traitement local des données, d'un délai réduit, entraînant un débit de données, une tâche et un partage plus élevés, etc.^[42]. Néanmoins, la technologie matérielle permettant de prendre en charge les capacités de traitement intelligent sur un petit périphérique avec des contraintes de ressources n'est pas encore mûre. De même, les logiciels dotés de capacités de décision intelligentes, en temps réel et indépendantes en matière de formation et d'apprentissage à partir des données pour répondre aux exigences des réseaux 6G auront besoin de temps pour évoluer.

Défis :

Lorsque nous considérons un environnement distribué, fournir de l'intelligence à tous les appareils périphériques en rassemblant tous les paramètres liés à l'utilisation des ressources, à la formation et aux modèles d'apprentissage est un défi. En outre, le déploiement optimal des nœuds périphériques dans un réseau à haute densité et l'allocation des ressources sont complexes. Au niveau des appareils périphériques, l'action de l'IA est largement influencée par le type de précision des données d'entrée, en particulier lorsque les données disponibles sont volumineuses, vives et nécessitent une classification. Par conséquent, les modèles de formation doivent tenir compte de ces facteurs difficiles et s'appuyer sur des données synthétiques générées à partir de réseaux antagonistes généralisés pour gérer la cohérence des données. En ce qui concerne les appareils périphériques distribués qui gèrent l'IA, par exemple les téléphones mobiles, les algorithmes d'IA doivent être légers, et les applications doivent être partitionnées entre les appareils utilisateurs et la couche périphérique pour une gestion efficace des ressources. Les progiciels pour EI contiennent des applications de périphérie qui sont déployées sur les dispositifs de périphérie avec l'infrastructure de virtualisation. Dans ce contexte, la sélection des politiques systèmes, la facturation, etc., jouent un rôle clé. Dans les dispositifs EI, la rétroaction en temps réel est essentielle pour obtenir de meilleures performances. Par conséquent, il est nécessaire de redéfinir la rétroaction en temps réel en considérant les modèles d'apprentissage en ligne et préformés. Les algorithmes de formation doivent être distribués au lieu d'être centralisés et doivent adopter des modèles de formation en ligne en tenant compte des limites des dispositifs périphériques^{[61], [42], [80]}.

3.2.6 Comment l'apprentissage automatique régit la 6G ?

L'apprentissage automatique (ML) est une sous-branche de l'IA qui aide à la prédiction et à la classification des données en fonction des données formées. Dans le réseau 6G, ML trouve son application à différentes instances du réseau, en particulier au niveau des couches PHY, MAC, réseau et application. Nous discutons brièvement de ces applications dans cette section. D'après les discussions précédentes sur les cas d'utilisation à l'ère de la 6G, nous réalisons que les composants de réseaux intelligents interconnectés nécessitent le ML comme exigence fondamentale du réseau pour offrir une valeur ajoutée aux services, une optimisation sans contact et améliorer les performances. En effet, s'appuyer uniquement sur la modélisation mathématique du système sans fil rend l'ensemble du système complexe sur le plan informatique. Cependant, le ML peut modéliser efficacement le système sans fil qui ne pourrait pas être enveloppé dans des équations mathématiques pour répondre aux exigences des réseaux 6G. Fait intéressant, ML agit sur les vastes données en temps réel reçues de différents cas d'utilisation tels que les voitures autonomes, la télé-présence holographique et le contrôle intelligent du réseau 6G^[47].

Au niveau de la couche PHY, de nombreuses fonctionnalités peuvent être envisagées par des modèles mathématiques. Cependant, certains phénomènes non linéaires tels que la détection d'interférences et la prédiction de canal peuvent être efficacement traités avec le ML. De plus, les modules discrets tels que le modulateur, le démodulateur et le filtre optimisent les performances dans leur ensemble par des méthodes d'apprentissage appropriées en ML. Lorsque nous considérons le processus de codage de canal, l'émetteur et le récepteur doivent

mettre en œuvre certains schémas de codage pour protéger les données contre les disparités de canal. L'apprentissage en profondeur est l'une des méthodes prometteuses pour prédire la longueur de codage appropriée en apprenant la longueur du mot de code du canal. Ensuite, la synchronisation est un autre aspect qui dépend certainement des variations de canal, de la mobilité et des écarts de fréquence. À cet égard, les réseaux de neurones profonds se révèlent bénéfiques. Cependant, des recherches supplémentaires sont nécessaires pour commenter davantage à ce sujet. Un autre aspect du ML est l'utilisation de réseaux de neurones profonds pour le positionnement. Étant donné que les méthodes de positionnement existantes entraînent des problèmes de précision lorsqu'il n'y a pas de trajectoire de ligne de visée en raison de leur dépendance à la modélisation mathématique pure. Dans l'apprentissage en profondeur, la méthode des empreintes digitales utilise CSI et la force du signal reçu comme données d'apprentissage. Au total, plusieurs optimisations de couche physique telles que la formation de faisceaux, l'estimation de canal et la maximisation du débit ne sont pas convexes et ne donnent pas de bonnes performances avec les algorithmes heuristiques. Dès lors, le *deep learning* (CNN et DNN) apparaît comme une aubaine pour résoudre en temps réel les problématiques d'optimisation de la couche physique des réseaux 6G, tout en réalisant un bon compromis entre performance et temps de calcul.

3.2.6.1 ML au MAC

ML trouve son application dans des solutions optimales aux tâches de la couche MAC telles que l'allocation des ressources, la sélection du schéma de modulation et de codage (MCS), le transfert et le contrôle de la puissance de la liaison montante. Lorsque les conditions du canal varient en raison du mouvement de l'utilisateur, la variation du canal, les méthodes d'apprentissage par renforcement profond peuvent déterminer des solutions optimales en apprenant à partir des entrées variables.

- 1 - ML pour l'allocation des ressources : dans la plupart des scénarios IoT, la transmission de données est prévisible car elle suit un modèle spécifique. En conséquence, il est facile de prédire de tels modèles et de décider des ressources à allouer efficacement à l'aide du ML. Cela réduit la latence du réseau et l'accès aléatoire en liaison montante pour les appareils. De même, par allocation de temps, les ressources de fréquence pour les utilisateurs de cellules peuvent être optimisées avec ML. Il peut prédire les traces de données dans la cellule, charger le réseau et allouer des créneaux horaires de fréquence de manière flexible (*flexible duplex*) pour optimiser les ressources et éviter les interférences.
- 2 - Gestion de l'alimentation : les périphériques réseau 6G doivent disposer d'une capacité de gestion de batterie transparente. Étant donné que la couche MAC contrôle l'accès au support, la puissance radio est le facteur clé à gérer. ML peut analyser les modèles de trafic en temps réel à l'intérieur de la cellule et les données ciblées pour les cellules voisines afin d'ajuster la charge. Ainsi, le réseau peut programmer les utilisateurs entre les états de transmission et de veille de manière économe en énergie.
- 3 - Sécurité avec ML : en raison de la nature très médiatisée du réseau 6G, comme un débit de données élevé, une connectivité complète et une latence extrêmement faible,

la génération de données va monter en flèche. Par conséquent, pour faire face aux diverses menaces de sécurité qui surviendraient en raison des données massives, la sécurité du réseau doit être automatisée. Dans ce contexte, les fonctionnalités proactives et auto-adaptatives du ML sont des solutions prometteuses pour automatiser la sécurité.

3.2.6.2 ML à la couche d'application 6G

Les réseaux deviennent intelligents en intégrant la sensibilisation au contexte pour offrir de meilleurs services aux utilisateurs et l'apprentissage par renforcement multiagent pour améliorer l'efficacité globale. L'utilisation d'une approche basée sur des règles pour déterminer la configuration du contexte pour chaque classe de service sera fastidieuse. Cependant, ML peut prédire les configurations de contexte de manière significative à l'aide des choix de services passés et modifier la configuration pour s'adapter aux nouveaux services des applications. Par exemple, considérons une application qui surveille automatiquement les paramètres du réseau. Nous savons que les KPI doivent être maintenus au niveau du seuil. Par conséquent, ML peut être utilisé pour détecter des anomalies concernant la gestion du réseau, la sécurité et bien d'autres.

Les futures architectures de réseau doivent être des fournisseurs de services de bout en bout, où il y aura une coopération dynamique entre les segments de réseau et les entités de communication. Pour assurer une connectivité ubiquitaire complète, l'interopérabilité entre les réseaux hétérogènes est obligatoire. En outre, l'efficacité en termes de coût, de puissance, de performances du système, de communication et de connectivité, d'utilisation du spectre est nécessaire.

Le réseau central pour fournir le déploiement de services à la demande se compose de SBA, où les fonctions de réseau (NF) s'abonnent à la fonction de référentiel de réseau et offrent des services à un ou plusieurs NF. Chaque NF peut être indépendamment mis à niveau ou supprimé pour fournir un découpage du réseau.

3.2.7 Modèle commercial pour la 6G

Dans les réseaux 5G, les modèles commerciaux sont passés d'opérateurs de réseau monopolisés à des micro-opérateurs de périphérie locaux distribués, des fournisseurs de services de réseau virtualisés et logiciels et des gestionnaires de tranches. Contrairement au réseau 5G, le réseau 6G fournit une connectivité ultime entre le monde cyber-physique en temps réel et plusieurs KPI initiaux. Cela motive de nombreux verticaux, développeurs de technologies, développeurs de produits et d'applications pour la future société 5.0. Ces avancées du marché perturbent l'activité conventionnelle et invitent de nouveaux acteurs et investisseurs en ce qui concerne les applications de communication, de stockage et de contrôle de la 6G^[50]. En outre, le réseau 6G s'entremêle avec l'intelligence, la virtualisation, l'exploitation locale en périphérie, le partage du spectre et une variété de services de détection ; dans l'ensemble, cela influence nos vies personnelles et fait partie intégrante de la société en 2030. En conséquence, cela apporte plusieurs nouvelles formes de modèles commerciaux et d'opportunités pour générer des revenus^[16].

3.2.7.1 Tendances clés et incertitudes

Le système de réseau intelligent pour la distribution d'électricité, ou internet, prévoit une connectivité extrême et autonome. Compte tenu de ces applications, des modèles de revenus peuvent être planifiés avec la participation d'un partenariat public-privé. Ensuite, les entreprises de premier plan (OTT) offrent des services de stockage et cognitifs dans le *cloud* (IA, UAV en tant que service et services sensibles au contexte) ainsi que des fonctionnalités principales telles que les appels et la connectivité des données pour concurrencer les ORM traditionnels. Cette tendance stimule les modèles commerciaux innovants pour attirer les clients et améliorer leur profit. Les modèles commerciaux ainsi que l'augmentation des revenus doivent même viser des objectifs de durabilité au profit de l'humanité. Par exemple, la vision de la 6G comprend la fourniture d'une connectivité pour tous les humains ; dans ce contexte, au moins fournir l'internet mobile à tous dans les zones rurales devrait résoudre le problème de la fracture numérique. Lorsque tout le monde est connecté, cela favorise la croissance économique^[157]. À cet égard, les services de données attirent la monétisation. De plus, une autre opportunité commerciale intéressante est l'automatisation de l'industrie. Lorsque nous considérons l'industrie 5.0 et au-delà, tous les capteurs et machines qui communiquent entre eux ou sont exploités par des robots nécessitent un réseau de communication privé et sécurisé qui offre une fiabilité extrême, une assistance sans contact et fonctionne comme des réseaux autonomes. Plus important encore, le partage du spectre 6G et les réglementations politiques génèrent des revenus de premier ordre pour le gouvernement et les agences. En plus de cela, les jumeaux numériques et l'industrie FinTech vont révolutionner le modèle commercial avec les réseaux 6G. En effet, les réseaux 6G utilisent des technologies vertes, sans énergie et sans émission, telles que des surfaces réfléchissantes intelligentes et des nœuds sans fil à longue durée de vie, de manière à générer de nouvelles activités autour de la construction de ces technologies^[23].

3.2.7.2 Scénarios

Il est prévu que les scénarios suivants ouvrent de nouvelles opportunités commerciales pour la 6G^[50]. L'informatique de pointe, l'IoT, la robotique et l'IA deviennent les technologies courantes que chaque gadget électronique intègre et sont à un prix raisonnable pour les clients. Par conséquent, ces appareils pourraient être utilisés pour collecter une variété de données auprès des utilisateurs et les traiter localement avant la transmission à distance. Ensuite, le découpage du réseau va autoriser jusqu'à 10 000 tranches sous un opérateur de réseau. Cela accélère la génération de revenus^[23]. De plus, le développement de solutions durables leur permet de fonctionner de manière autonome et de simplifier la vie des gens. À cet égard, pour trouver localement des solutions durables et économiques, des projets de partenariat public-privé, des modèles de coopératives distribuées, et les modèles commerciaux *peer to peer* sont plus adaptés au marché monopolisé au niveau local^[50]. De plus, en raison de la faible diffusion de la technologie dans les zones rurales, la fabrication de solutions bon marché et peu coûteuses au niveau micro pourrait être une approche acceptable. Cependant, la 6G surmonte la fracture numérique pour promouvoir le partage de technologie dans lequel une conception 3D d'une machine peut être partagée en temps réel avec une autre ville sous-développée qui est privée d'installations de recherche pour la fabriquer en utilisant sa main-d'œuvre bon marché et son

marché local ainsi qu'au lieu d'origine à un prix raisonnable. Cela soutient une économie circulaire et crée une société meilleure à la fois dans les zones rurales et urbaines. La fabrication de solutions bon marché et peu coûteuses au niveau micro pourrait être une approche acceptable.

3.2.8 Tendances supplémentaires

3.2.8.1 Accès sans cellulaire

L'architecture sans cellule est un concept dans les futurs réseaux, où le paradigme existant de connexion d'un utilisateur à un point d'accès (gNB) confiné à une cellule n'est plus le cas. Au lieu de cela, un utilisateur cellulaire obtient la connectivité de plusieurs BS sans la restriction des limites des cellules. Dans une architecture cellulaire traditionnelle, les utilisateurs qui sont au bord de la cellule ont une mauvaise connectivité, qui est résolue dans l'architecture sans cellule. Cette fonctionnalité nécessite de peaufiner les procédures traditionnelles de recherche de cellules, de synchronisation, d'accès aléatoire et d'allocation des ressources, pour s'adapter au nouvel accès sans cellule^{[107], [112]}. De plus, en 6G en raison de la densité extrêmement élevée d'utilisateurs mobiles ou d'appareils intelligents qui nécessitent des débits de données élevés, la couverture des gNB est limitée car ils fonctionnent à des fréquences très élevées, ce qui réduit les interférences. Dans ce contexte, l'architecture sans cellule est très utile pour fournir une connectivité transparente. Le système sans cellule se compose de plusieurs gNB répartis dans une grande région géographique et fournit une connectivité à tous les utilisateurs situés en utilisant la même fréquence. Ces gNB sont connectés à une unité centrale de traitement (CPU) qui les contrôle. Normalement, un système *cellfree* est associé au MIMO pour améliorer la connectivité spatiale. L'architecture sans cellule présente plusieurs avantages, à savoir :

- un traitement de signal évolutif ;
- un contrôle de puissance évolutif ;
- une efficacité spectrale et énergétique ;
- et un traitement de signal simple et un système économique^[120].

3.2.8.2 Communication par rétrodiffusion (BsC)

C'est un paradigme intéressant pour le réseau 6G d'atteindre une efficacité énergétique extrêmement élevée (EHEE) du réseau lorsqu'il y a un nombre énorme d'appareils connectés. Pour une distribution massivement dense de nœuds sans fil, l'utilisation de la batterie comme source d'alimentation ne sera pas possible en raison des défis liés à la recharge ou au remplacement de la batterie. Dans ce scénario, les nœuds sans fil doivent recevoir passivement les signaux radio de l'environnement, ce qui conduit à un fonctionnement sans batterie^[123]. L'idée fondamentale est d'utiliser l'énergie RF incidente de l'ambiance et de la réfléchir après modulation. En fait, cette méthode réduit la charge sur le dispositif contraint et augmente l'efficacité du spectre. Il a des applications dans des scénarios de communication à courte portée et à faible puissance.

3.2.8.3 Transfert de puissance sans fil

Lors de la connexion d'appareils massifs, l'obtention des informations sur l'état du canal (CSI) pour connaître l'état de leur canal devient une tâche fastidieuse et consomme beaucoup d'énergie. Par conséquent, la qualité de la transmission peut être affectée si CSI est omis. Cependant, un compromis entre la communication de données et la puissance peut être atteint par le transfert de puissance sans fil (WPT) et les transmissions sans CSI. En outre, le transfert d'énergie sans fil est une méthode possible pour promouvoir la durabilité, les faibles émissions et la consommation d'énergie nulle parmi les appareils des réseaux 6G^[12]. Dans ce contexte, comme applicable aux nœuds IoT massifs, les nœuds peuvent être fabriqués en récupérant l'énergie de l'environnement. Il peut s'agir de l'utilisation de l'énergie solaire, éolienne, hydroélectrique ou de signaux RF, etc. Le transfert d'énergie sans fil à l'aide de RF implique d'alimenter (charger) à distance les appareils sans fil à l'aide du champ RF. Par exemple, lorsque la BS est transférée à un utilisateur, le voisin de cet utilisateur peut utiliser l'énergie du signal RF pour se recharger. Il utilise les propriétés de rayonnement en champ lointain des ondes électromagnétiques pour transmettre sur une courte distance. Cela présente plusieurs avantages, tels qu'une augmentation significative de la durabilité et de la durée de vie des nœuds, une empreinte énergétique réduite, la recharge automatique et la recharge sans contact.

Quelques solutions bien connues pour le transfert d'énergie sans fil impliquent :

- la formation de faisceaux d'énergie, où un groupe d'antennes à faisceau élevé est focalisé pour transmettre l'énergie à un ensemble spécifique de nœuds à l'aide de faisceaux étroits.
- des systèmes d'antennes distribuées pour reconstituer la perte d'énergie lors de la propagation de l'énergie vers les nœuds en utilisant des balises de puissance.

Cependant, il existe plusieurs défis impliqués dans le transfert d'énergie. Pour n'en citer que quelques-uns il s'agit :

- de l'inefficacité de la prise en charge du transfert d'énergie hors ligne de mire ;
- du transfert de puissance pour les utilisateurs non fixes (mobiles) ;
- des risques de rayonnement.

4. Activer les technologies clés

Grâce à des examens et des enquêtes 6G, des statistiques et des résumés des technologies clés potentielles, nous obtenons 12 technologies clés potentielles avec l'attention et la fréquence de mots les plus élevées. Les avantages et les défis correspondants de chaque technologie clé sont analysés et résumés. Voir le tableau 6 pour plus de détails.



Tableau 6 : Technologies clés potentielles des réseaux 6G

TECHNOLOGIE CLÉ	RÉFÉRENCES	AVANTAGES	DÉFIS
Nouvelle technologie de fonction réseau			
Intégration de l'air, de l'espace, du sol et de la mer	[6], [7], [13], [14], [19]–[28]	Un espace de couverture complète avec une capacité maximisée, des connexions omniprésentes denses et un spectre haute densité est formé	Scénarios nécessitant une faible latence et une grande fiabilité
Intelligence artificielle	[6], [7], [13], [14], [19]–[29]	Réaliser un nouveau réseau autonome et sans contact, en adaptant le réseau à des scénarios d'application omniprésents	Défis en temps réel, partage, efficacité énergétique, etc.
Nouvelle technologie de ressources spectrales			
Communication térahertz	[6], [7], [13], [19], [21]–[30]	Il peut répondre aux exigences de spectre du taux de transmission de données extrêmement élevé de la 6G et dispose de ressources spectrales plus riches	Caractéristiques de propagation du spectre THz et limitations de la maturité des dispositifs RF
Communication par lumière visible	[6], [7], [13], [14], [19], [20], [22], [24]–[26]	Il a une bande ultra-large et il est largement disponible	Besoin de réaliser des percées dans les puces et modules d'émetteurs-récepteurs de transmission de lumière visible ultra-rapides
Technologie dynamique de partage intelligent du spectre	[6], [7], [14], [22], [24], [25], [27]	L'utilisation dynamique du spectre est un moyen important d'améliorer efficacement l'efficacité de l'utilisation du spectre existant	Il doit être soutenu par un algorithme coordonné, rationnel et intelligent
Technologie d'accès sans fil efficace			
Amélioration de la technologie traditionnelle	[7], [14], [19], [20], [24]–[28]	Il satisfait plus de bande passante, une plus grande capacité, un débit de données ultra-élevé et un déploiement d'équipement de manière haute densité	La haute fréquence et la haute puissance font que la conception des émetteurs-récepteurs et des antennes, la conception de l'efficacité énergétique et d'autres technologies sont confrontées à des défis
OAM	[7], [24]	Multiplexage en parallèle pour atteindre une efficacité spectrale élevée	Le goulot d'étranglement de la fusion et de la séparation des faisceaux d'ondes radio

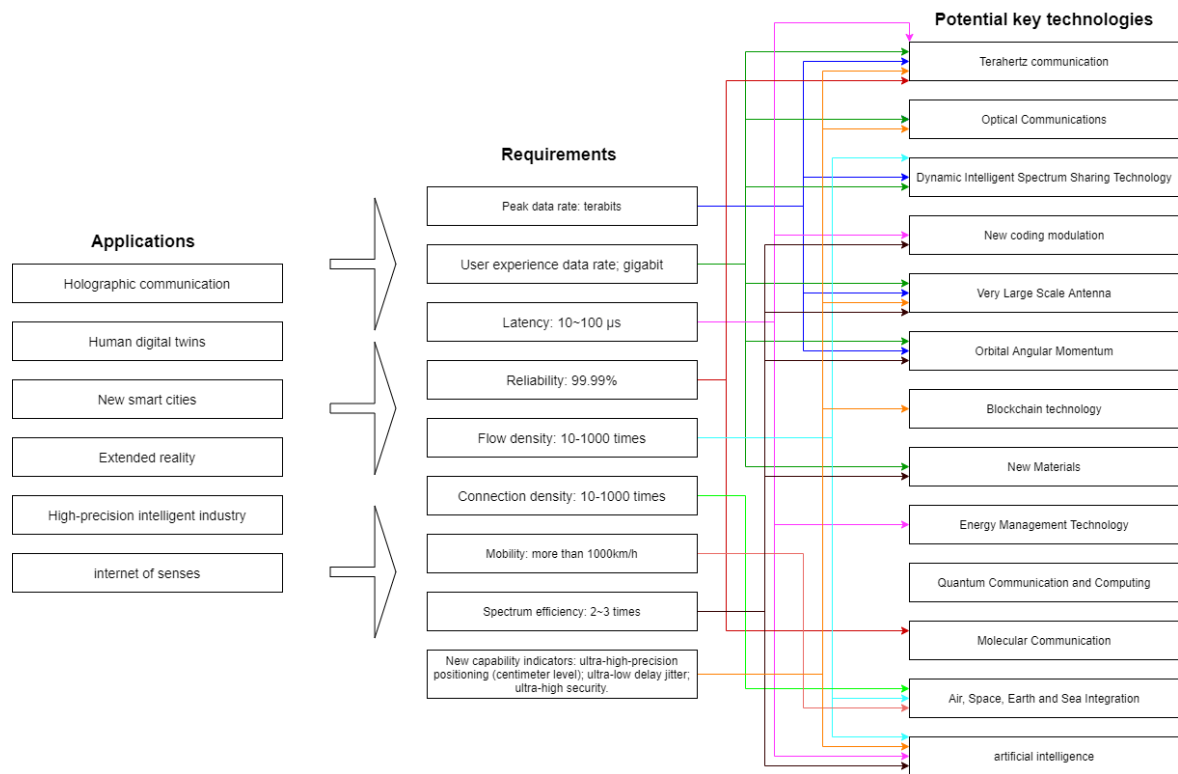


TECHNOLOGIE CLÉ	RÉFÉRENCES	AVANTAGES	DÉFIS
Technologie de base innovante			
Chaîne de blocs	[11], [22], [25], [26], [28], [30], [31]	Fournir une forte performance de sécurité	Issu de défis tels que la sécurité du système, la confidentialité des données, la supervision, l'évolutivité, etc.
Nouveaux matériaux	[14], [22], [25], [26]	Promouvoir l'innovation des batteries, appareils et antennes	Résoudre les performances des appareils clés
Gestion de l'énergie	[13], [14], [19], [20], [22], [25], [26], [28]	Réaliser des opérations de réseau efficaces et flexibles	Nécessité d'intégrer les caractéristiques énergétiques dans le système
Nouvelle technologie de communication			
Communication et informatique quantique	[6], [13], [14], [22], [26]	Améliorer l'efficacité informatique et fournir une sécurité renforcée pour la 6G	Problèmes de sécurité et problèmes de transmission longue distance dans des conditions réalistes
Communication moléculaire	[6], [14], [22], [26]	Réaliser une communication et une interconnexion à l'échelle nanométrique	Interface et assurance de la sécurité entre les domaines électriques et chimiques

D'après le nombre de technologies correspondantes dans la deuxième colonne du tableau, on peut voir que les trois technologies majeures d'intégration air-espace-sol-mer, l'intelligence artificielle et la communication térahertz sont presque les principales technologies clés de la 6G. Parmi les 18 documents de recherche que nous étudions, il y a 15 articles liés à ces trois technologies. On en déduit que cette technologie favorise principalement la direction du développement de la 6G. De plus, la communication par lumière visible, le partage intelligent dynamique du spectre, les améliorations de la couche physique, la *blockchain*, les technologies de gestion de l'énergie sont également très concernées et elles jouent un rôle non-clé au milieu de la 6G. Au lieu du moment cinétique orbital, les communications et l'informatique quantique, les nouveaux matériaux et la communication moléculaire sont relativement moins concernés, mais ils ont un grand potentiel après les percées dans les défis technologiques ultérieurs et l'émergence d'une demande accrue. Sur la base des technologies clés potentielles ci-dessus, les capacités du réseau 6G sont considérablement améliorées en offrant aux utilisateurs des applications et des services plus impressionnants. Après avoir analysé et résumé la relation entre les futurs scénarios d'application, les exigences de performance et les technologies clés potentielles, la relation de mappage mutuel est notée (voir figure 25). On peut voir que pour atteindre les exigences de débit de pointe au niveau du téraoctet, la prise en charge de technologies clés telles que la communication térahertz, la communication par lumière visible, la technologie de partage intelligent dynamique du spectre et les antennes à très grande échelle sont nécessaires.

Parmi les 12 technologies clés potentielles susmentionnées issues de la recherche et de l'analyse, les deux technologies d'intégration air-espace-sol-mer et d'intelligence artificielle ont été condensées en détail dans la deuxième section de l'architecture du réseau 6G. Ainsi, cette section résume les autres technologies clés.

Figure 25 : Cartographie entre les applications 6G, les exigences et les technologies clés potentielles



4.1 Nouvelle technologie de ressources spectrales

La théorie de l'information de Shannon est une base de modèle pertinente pour les réseaux 6G. Il propose deux façons principales d'augmenter la capacité du système : augmenter la bande passante du système et améliorer l'efficacité du spectre^[18]. À ce stade, la communication Hertz, la communication par lumière visible et le partage du spectre sont principalement des technologies cohérentes pour augmenter les ressources du spectre 6G.

4.1.1 Communication térahertz

La bande de fréquence térahertz (0,1 THz ~ 10 THz) est très élevée et n'est actuellement pas réglementée. Elle est considérée comme une bande de spectre ultra-large qui peut réaliser une communication à très haut débit, ce qui peut atténuer la rareté actuelle du spectre et les limitations de capacité. Le faisceau étroit et la courte impulsion du térahertz limitent la possibilité d'écoute clandestine et peuvent réaliser une communication sécurisée et un positionnement de haute précision. La forte pénétrabilité de l'onde térahertz lui confère de

larges perspectives d'application dans les communications sans fil à très haut débit et les communications spatiales. Cependant, la communication térahertz a encore des problèmes techniques à résoudre en termes de composants matériels haute fréquence, de modélisation et d'estimation des canaux et de mise en réseau directionnelle.

4.1.2 Communication par lumière visible (VLC)

Travailler dans la gamme de fréquences de 400 THz à 800 THz est une autre technologie qui devrait atteindre la 6G et qui utilise la lumière visible générée par des matériaux de diodes électroluminescentes (DEL) pour transmettre des données. VLC utilise une bande passante ultra-élevée pour obtenir une transmission et une disponibilité des données à haut débit. Il a de multiples fonctions telles que la communication, l'éclairage et le positionnement qui conviennent aux points chauds intérieurs et à d'autres scénarios. Cependant, VLC est également confronté à des défis tels que la limitation de la bande passante de modulation et la compensation non linéaire.

4.1.3 Technologie de partage dynamique du spectre

Le mode d'attribution du spectre du système existant rend les ressources spectrales complètement occupées là où le taux d'utilisation est faible. La technologie dynamique de partage intelligent du spectre permet aux utilisateurs non autorisés d'utiliser le spectre qui n'est pas entièrement utilisé par l'utilisateur principal dans les dimensions temporelle et géographique, ce qui améliore considérablement l'efficacité du spectre. Pour gérer les connexions à grande échelle dans les applications 6G, des technologies d'atténuation des interférences distribuées et efficaces doivent être utilisées pour améliorer les performances du système. Les technologies de *blockchain* et d'apprentissage en profondeur sont des méthodes efficaces pour le partage dynamique et intelligent du spectre.

4.2 Technologie d'accès sans fil efficace

4.2.1 Amélioration de la technologie traditionnelle de la couche physique

4.2.1.1 Nouvelle modulation de codage

La technologie de codage et de modulation de la 6G nécessite une conception et une optimisation ciblées des caractéristiques de transmission de scénarios d'application de communication complexes tels que le débit en téraoctet, la bande passante de canal ultra-large, la bande de fréquence térahertz élevée, la mobilité ultra-élevée et la stabilité. De plus, la technologie d'IA peut trouver la meilleure méthode de codage et de modulation adaptée à l'environnement de transmission du système actuel grâce à l'apprentissage, à la formation et à l'évaluation, pour fournir une solution qui ne repose pas sur les théories traditionnelles.

4.2.1.2 Multiantenne à grande échelle

La technologie multiantenne à grande échelle de la 6G peut atteindre une efficacité spectrale ultra-élevée grâce au multiplexage par répartition spatiale (SDM). Elle améliore l'efficacité énergétique et réduit le délai en aidant à briser les limites de la cellule, et sa résolution spatiale ultra-élevée soutient un positionnement et une perception environnementale de haute précision. L'application de multiantenne à grande échelle nécessite une percée dans la technologie de l'antenne elle-même. À l'heure actuelle, l'application de surfaces réfléchissantes intelligentes à grande échelle dans les antennes à grande échelle a reçu la plus grande attention.

4.2.2 Moment angulaire orbital (OAM)

La technologie de multiplexage OAM utilise le moment cinétique d'un ensemble d'ondes électromagnétiques orthogonales pour multiplexer plusieurs flux de données sur le même canal et obtenir une efficacité spectrale et une capacité système supérieures. L'OAM dispose d'un grand nombre de modes OAM orthogonaux qui peuvent être multiplexés / démultiplexés ensemble, offrant ainsi une nouvelle façon d'améliorer la capacité du réseau 6G. Cependant, l'application de l'OAM dans les communications sans fil en est encore au stade exploratoire et fait l'objet d'un certain degré de recherche.

4.3 Technologie de base innovante

4.3.1 Chaîne de blocs

La *blockchain* est basée sur la technologie des registres distribués (DLT), ses fonctions inhérentes telles que la protection non centralisée contre les modifications et l'anonymat font de la *blockchain* un choix idéal pour diverses applications. Elle garantit une sécurité renforcée dans l'ensemble du processus de communication des entités du réseau 6G. Le mécanisme de contrôle distribué basé sur la *blockchain* peut établir un lien de communication direct entre les entités du réseau, réduisant ainsi les coûts de gestion. Le système de partage du spectre basé sur la technologie *blockchain* peut améliorer efficacement l'efficacité du spectre. Par conséquent, la technologie *blockchain* est la technologie la plus prometteuse pour assurer la sécurité et la confidentialité du réseau 6G.

4.3.2 Nouveaux matériaux

Le système de communication a connu un grand succès, mais les performances des matériaux semi-conducteurs traditionnels semblent avoir atteint leurs limites et il existe un besoin urgent de matériaux présentant de meilleures caractéristiques haute fréquence et haute température pour une communication ultra-rapide. De nouveaux matériaux tels que le nitrure de gallium, le phosphore d'indium, le silicium-germanium et le graphène sont utilisés pour concevoir des dispositifs de communication de nouvelle génération. De plus, des matériaux fluides sont introduits dans la conception des antennes reconfigurables en fréquence pour offrir une plus grande flexibilité. Les métamatériaux et les métasurfaces peuvent être déployés dans un

environnement sans fil radiocommandé. Les métamatériaux planaires contrôlés par logiciel peuvent réduire les interférences grâce à un contrôle déterministe des caractéristiques environnementales.

4.3.3 Gestion de l'énergie

La demande de calcul cohérent pour le traitement de l'IA et la popularité croissante des appareils IoT posent des défis majeurs à l'efficacité énergétique des appareils de communication. Compte tenu de l'ampleur attendue du réseau 6G, il est nécessaire de tenir compte de la consommation énergétique lors de la conception du système. Une option consiste à utiliser la récupération d'énergie, elle permet aux appareils d'être autoalimentés, ce qui est essentiel pour obtenir un fonctionnement hors réseau, des appareils et capteurs IoT durables, des appareils rarement utilisés et de longs intervalles de veille. De plus, la technologie radio symbiotique et la technologie de gestion intelligente de l'énergie offrent des solutions possibles.

4.4 Nouvelle technologie de communication

4.4.1 Communication et informatique quantique

La 6G doit répondre à des exigences de sécurité plus élevées dans divers scénarios d'application. La communication quantique peut fournir une sécurité renforcée grâce à des clés quantiques. Lorsqu'un indiscret souhaite observer, mesurer et répliquer une communication quantique, l'état quantique est perturbé et le comportement d'écoute indiscret est facilement détectable. En théorie, la communication quantique peut atteindre une sécurité absolue. De plus, la combinaison de la théorie quantique et de l'IA peut développer des algorithmes d'IA plus puissants et efficaces pour répondre aux besoins de la 6G. Cependant, la transmission de données à haut débit et les scénarios d'application à forte demande posent des défis à l'informatique sans fil en 6G.

4.4.2 Communication moléculaire

En biomédecine, le bionano IoT (*Internet of bio-nano things*, IoBNT) peut connecter des nanodispositifs (tels que des nanorobots, des puces implantables et des biocapteurs) et des entités biologiques. La communication moléculaire est une technologie habilitante de l'IoBNT. Elle utilise des molécules biochimiques pour communiquer et transférer des informations entre les nanodispositifs. De plus, la combinaison de l'IoBNT et des jumeaux numériques humains est un réseau sans fil à courte distance, composé d'appareils / capteurs de surveillance portables et d'appareils de détection à l'intérieur ou sur le corps, qui peut fournir une solution complète pour les soins de santé électroniques 6G.

5. Activités de recherche

Avec le succès de la commercialisation à grande échelle des réseaux 5G, l'industrie mondiale, les universités et les instituts de recherche ont officiellement lancé des recherches sur les

exigences potentielles de service 6G, l'architecture de réseau et les technologies habilitantes potentielles en 2019.

5.1 Union européenne

En septembre 2018, le livre blanc « Intelligent Networks in the Next Generation Internet » est publié. Sur cette base, l'Union européenne formule l'agenda stratégique de recherche et d'innovation (SRIA) 6G et la technologie de développement stratégique (SDA) dans le cadre du projet-cadre industrie-université-recherche 2021-2027 au troisième trimestre 2020 et met en œuvre la stratégie dans le premier trimestre 2021.

5.2 Finlande

Le gouvernement finlandais prend l'initiative d'établir le projet phare 6G dirigé et géré par l'université d'Oulu en Finlande en mai 2018. Les membres du projet sont principalement des entreprises, des universités et des instituts de recherche finlandais. L'université d'Oulu prend l'initiative d'organiser deux sommets sans fil 6G en mars de chaque année. Les principaux fabricants et opérateurs ont prononcé des discours au sommet de la technologie 6G. Sur la base des discussions techniques lors de la conférence, ils publient le « Facing 6G Summit » en septembre 2019, et le livre blanc « Drivers of Ubiquitous Wireless Intelligence and Main Research Challenges ». Actuellement, le 6G Wireless Summit rédige des livres blancs sur la technologie 6G et sur 12 sujets techniques dont la 6G drive et les Objectifs de développement durable des Nations unies.

5.3 États-Unis

La Federal Communications Commission (FCC) des États-Unis lance la recherche sur les nouveaux services du spectre térahertz dans la gamme de fréquences de 95 GHz à 3 THz en 2018. En juin 2019, elle délivre une licence de spectre d'essai de 10 ans pour le réseau vendable. Le contenu principal de ses recherches sur le spectre des fréquences comprend :

- la bande de fréquences 95 ~ 275 GHz est partagée par le gouvernement et le non gouvernemental ;
- 275 GHz~3 THz n'interfère pas avec l'utilisation du spectre de fréquences existant ;
- le spectre sans licence a une bande passante totale de 21,2 GHz, y compris 116 ~ 123 GHz, 174,8 ~ 182 GHz, 185 ~ 190 GHz, 244 ~ 246 GHz.

L'Alliance pour les solutions de l'industrie des télécommunications (ATIS) publie une proposition d'action 6G le 19 mai 2020, recommandant que le gouvernement investisse des fonds supplémentaires de R&D dans les percées des technologies de base 6G, et encourage les gouvernements et les entreprises à participer activement à la formulation des politiques nationales du spectre. À l'heure actuelle, les futures technologies de base 6G que les États-Unis

espèrent dominer comprennent l'intégration 5G et le réseau ouvert (ION), des technologies avancées qui prennent en charge l'intelligence artificielle (IA).

Réseaux et services, antennes et systèmes radio avancés (tels que la bande de fréquence térahertz au-dessus de 95 GHz), services de réseau à accès multiple (y compris les réseaux terrestres et non terrestres, auto-induction pour prendre en charge des applications telles que le positionnement ultra-haute définition), services de réseau de santé intelligents (y compris le diagnostic et la chirurgie à distance, l'utilisation de nouvelles fonctionnalités telles que les applications multicapteurs, l'internet tactile et les images 3D à ultra-haute résolution) et les services d'agriculture 4.0 (soutenant l'application unifiée d'eau, d'engrais et de pesticide).

5.4 Japon et Corée du Sud

Le gouvernement japonais lance sa première version d'une stratégie de recherche sur le réseau de communication sans fil 6G à l'été 2020. L'Institut de recherche sur l'électronique et les télécommunications du gouvernement sud-coréen (ETRI) signe un accord de recherche en coopération sur le réseau 6G avec l'université d'Oulu en Finlande en juin 2019. Samsung se concentre sur la 6G, l'intelligence artificielle et la robotique depuis 2019. LG en janvier 2019 et l'Institut coréen des sciences et technologies (KAIST) coopèrent pour créer un centre de recherche 6G. De plus, SKT et les fabricants étudient conjointement les indicateurs de performance clés et les exigences commerciales de la 6G.

5.5 Chine

Le ministère chinois de l'industrie et des technologies de l'information élargit le groupe de promotion IMT-2020 initial au groupe de promotion IMT-2030 pour mener des études de faisabilité sur les exigences 6G, la vision, les technologies clés et les normes mondiales unifiées. Le ministère chinois des Sciences et de la technologie prend l'initiative de lancer en novembre 2019 un groupe de promotion du développement de la technologie 6G impliquant 37 instituts de recherche industriels et universitaires pour mener à bien des projets de coopération industrie-université-recherche sur les exigences, la structure et les technologies habilitantes de la 6G.

La 6G est prévue pour une utilisation commerciale après 2030, sa vitesse élevée, sa faible latence et sa densité de connexion élevée sont attendues à partir de maintenant. Dans la société 6G, il y a un monde où les gens ne sont pas conscients de la communication de fond et c'est comme si la technologie de la communication était intégrée dans la société.

Conclusion

La croissance des données sans fil est presque exponentielle. La commercialisation rapide de divers appareils intelligents, l'émergence de diverses nouvelles applications et la demande d'interconnexion mondiale de toutes choses jettent les bases de la prochaine évolution sans fil vers la 6G. Le réseau sans fil 6G doit améliorer considérablement la qualité des services et assurer un développement durable. Ce chapitre propose des exigences de vision 6G et discute

des scénarios d'application 6G et des technologies clés. Par la suite, deux nouvelles fonctionnalités sont utilisées pour illustrer la nouvelle architecture de réseau de la 6G, notamment l'intégration de réseaux terrestres et non terrestres, et la véritable connexion intelligente prise en charge par l'IA. Enfin, grâce à l'analyse de nombreux articles de machines bien connus au pays et à l'étranger, 12 technologies clés potentielles sont identifiées, y compris l'intégration air-espace-terre-mer, l'intelligence artificielle, le térahertz, la communication par lumière visible, le moment cinétique orbital, la technologie *blockchain*, etc. pour décrire les scénarios d'application, la relation entre les exigences de performance et les technologies clés. Le réseau pré-commercial 6G sera mis en service vers 2030 pour soutenir le développement d'une société mobile intelligente omniprésente.

Chapitre IV ♦ Structure technique d'un réseau 6G exigeant

Introduction

La standardisation des services de cinquième génération crée une concurrence féroce pour développer des services de nouvelle génération à l'échelle mondiale. Les États-Unis, le Japon, la Corée du Sud, d'autres pays développés dont certains pays européens, ont commencé à étudier et à formuler le plan de développement de la technologie de sixième génération^[1].

La technologie 6G intègre des systèmes basés sur les données et l'apprentissage automatique qui concernent les technologies de l'intelligence artificielle (IA). Ces nouvelles architectures de communication sont conçues pour réaliser une auto-organisation complète ainsi qu'un auto-apprentissage, une auto-réparation, une auto-agrégation, une auto-protection et une auto-optimisation du réseau.

L'intelligence artificielle aide également la 6G à fournir des services intelligents hautement personnalisés et connectés aux particuliers, aux entreprises et à d'autres utilisateurs pour répondre à des besoins subtils^[2, 3].

Bien que les spécifications de la 6G, telles que les bandes de fréquences et les exigences en matière de débit de données, n'aient pas encore été définitivement définies, ses applications sont déjà envisagées. La 6G fait désormais l'objet d'un consensus : elle sera un réseau de communication mobile intelligent à très grande échelle qui englobera la 5G^[4, 5]. Quand le réseau 5G, quasi bidimensionnel, ne couvre qu'une partie limitée du globe, le réseau 6G s'étend en trois dimensions, connectant les satellites, les avions, les navires et les infrastructures terrestres avec une couverture au niveau du globe. Les technologies mmWave rendent possible les différentes sortes de connexions sans fil, avec une vitesse et une fiabilité supérieures à la 5G. De plus, l'utilisation du térahertz dans le cadre des bandes de fréquences pour les communications 6G est également proposé^[5], cependant, les dispositifs clés, pendant des puces térahertz, les composants frontaux et les systèmes ne sont pas encore aussi matures et fiables que ceux fonctionnant à des fréquences mmWave pour les communications à longue portée et haute-fidélité. Les communications à ondes millimétriques peuvent exploiter pleinement les ressources spectrales abondantes dans la bande de fréquences au-dessus de 26 GHz

pour réaliser des transmissions de données ultra-rapides et la bande 0,1–10 THz avec l'avènement de l'ère 6G^[6]. Dans le même temps, plus la bande de fréquences de travail des communications à ondes millimétriques est élevée, plus la perte de diffusion est importante. Les chercheurs utilisent souvent de grands réseaux d'antennes pour former des systèmes à entrées et à sorties multiples (MIMO) afin de générer des faisceaux hautement directionnels et de compenser la perte de diffusion.

La 6G est servie par un niveau de débit fulgurant pouvant atteindre 1 To par seconde^[7]. Cela améliore les performances des applications 5G et augmente la capacité à prendre en charge de nouvelles applications innovantes. Pour atteindre ce débit, la 6G doit inclure l'utilisation des très hautes fréquences (ondes millimétriques) du spectre radio. La capacité de bande passante du réseau 5G est due à l'utilisation de hautes fréquences radio ; plus le spectre radio est élevé, plus il est possible de transmettre des données. Le réseau de sixième génération (6G) pourrait à terme se rapprocher de la limite supérieure du spectre radio, atteignant de très hautes fréquences dans les bandes 300 GHz voire térahertz^[2, 8–11].

Plusieurs méthodes et approches pour la sélection de faisceaux d'ondes millimétriques sont proposées. La sélection de faisceau utilise des faisceaux d'émission et de réception simulés pour détecter les canaux MIMO à ondes millimétriques et trouver la combinaison de faisceaux de réception et d'émission avec la plus grande énergie de signal reçu, c'est-à-dire la paire de faisceaux la plus appropriée pour la transmission de canaux, évitant ainsi le recours à des données d'interfaces aériennes de grande dimension.

Inspiré par les grandes percées réalisées par l'apprentissage en profondeur (DL, pour *deep learning*) dans la vision par ordinateur et le traitement du langage naturel, DL est également utilisé pour la communication sans fil^[12]. Par rapport aux méthodes basées sur des modèles mathématiques, DL offre deux avantages clés. En premier lieu, les outils mathématiques sont généralement basés sur des hypothèses idéales, telles que la présence d'un bruit gaussien blanc additif pur, ce qui peut ne pas être compatible avec des conditions réelles. En revanche, DL apprend de manière adaptative les caractéristiques du canal pour prendre en charge un entraînement de faisceau fiable^[13]. En second lieu, les paramètres des modèles DL capturent les caractéristiques de haute dimension du scénario de propagation. Pour ces raisons, ce chapitre étudie les techniques DL pour les faisceaux d'entraînement, car elles sont très bien adaptées pour traiter les propriétés non linéaires et non monotones des fuites de puissance de canal dans les communications à ondes millimétriques.

1. Intelligence artificielle

La première définition de l'IA pourrait remonter à 1950, lorsque Turing^[105] introduit un test, le test de Turing, pour tester l'intelligence d'une machine. Cependant, selon un livre^[34], la terminologie de l'IA est née lors d'un atelier en 1956. Depuis, le domaine est étudié, mais nous n'avons découvert qu'une partie des possibilités. Au début, le développement est lent, mais dans les années 1990 et au début du XXI^e siècle les ordinateurs commencent à avoir suffisamment de puissance de calcul et de données pour que l'IA soit utilisée dans les domaines de la santé et de la pharmacie, de l'exploration de données, de la logistique et d'autres encore. L'IA peut être classée en fonction de sa capacité à émuler la pensée humaine. Ainsi, nous

pouvons différencier comment le système se compare aux humains en termes de polyvalence et de performance. Sur la base de ces critères, nous classons les systèmes d'IA en fonction de leur similitude avec l'esprit humain et de leur capacité à "penser". Cette capacité peut être classée en quatre catégories de base : les machines réactives, les machines à mémoire limitée, la théorie de l'esprit^[101] et l'IA consciente d'elle-même^[21].

- **Réactives** : Les machines réactives n'ont pas de mémoire, elles réagissent seulement à différents stimuli. C'est le type le plus ancien de système d'IA, et c'est aussi le plus "simple". Ces systèmes n'ont aucune mémoire (aucune capacité d'apprentissage) et sont très limités dans ce qu'ils peuvent faire. Ils sont souvent programmés pour faire une seule chose spécifique, réagissant ainsi à un ensemble ou à une combinaison limitée d'entrées. Un exemple populaire d'IA réactive est DeepBlue^[6], une machine qui est devenue maître aux échecs.
- **Machines à mémoire limitée** : Les machines à mémoire limitée ont les mêmes capacités que les systèmes réactifs mais elles utilisent la mémoire pour améliorer leurs réponses. Ces systèmes apprennent à partir de données historiques, qui sont utilisées pour améliorer les décisions. Presque toutes les applications actuelles entrent dans cette catégorie (par exemple, un *chatbot* ou une voiture à conduite autonome).
- **Théorie de l'esprit** : Une IA au niveau de la théorie de l'esprit comprend les entités avec lesquelles elle interagit en comprenant leurs besoins, leurs émotions, leurs croyances et leurs processus de pensée. Il s'agit encore d'un concept ou d'un travail en cours. Pour créer une IA de ce niveau, les chercheurs doivent créer un système qui perçoit les humains comme des individus dont l'esprit peut être défini par de multiples facteurs^[101].
- **IA auto-consciente** : Il s'agit du stade final du développement de l'IA, qui n'existe actuellement que de manière hypothétique. L'IA consciente d'elle-même a évolué au point de ressembler tellement au cerveau humain qu'elle a développé une conscience d'elle-même^[21].

La capacité d'apprendre et d'améliorer à partir des données est si importante qu'elle donne naissance à un autre domaine de l'IA appelé *machine learning* (ML). Le ML déplace le centre d'attention de l'IA vers la résolution de problèmes pratiques. Elle abandonne alors les approches symboliques pour se tourner vers les méthodes et les modèles utilisés en statistique et en théorie des probabilités^[67].

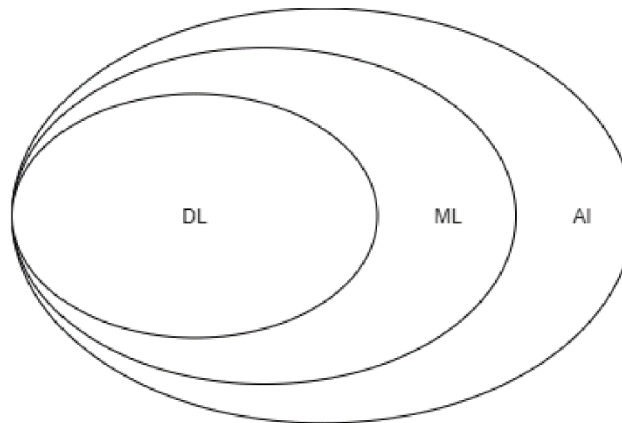
Après un certain temps, Le *et al.*^[68] revisitent l'ancien concept de réseau neuronal profond et créent le premier réseau non supervisé en 2011. Ce réseau est entraîné à reconnaître des concepts de plus haut niveau, comme les chats, uniquement à partir de l'observation d'images non étiquetées^[68]. La théorie des réseaux neuronaux profonds est d'abord abandonnée parce que les chercheurs ne disposent pas d'une puissance de calcul et de données suffisantes pour la prouver, d'où la possibilité de l'appliquer à partir de la dernière décennie. Ces recherches ont déclenché le "boom" de l'apprentissage profond.

L'apprentissage profond fait partie d'une famille plus large d'apprentissage automatique (la relation entre l'IA, le ML et l'apprentissage profond (DL) (qui peut être observée à la figure 26)



basée sur les réseaux neuronaux artificiels (RNA) avec apprentissage par représentation. Les RNA sont des réseaux neuronaux similaires aux réseaux neuronaux biologiques, mais leurs connexions sont statiques et non dynamiques.

Figure 26 : Relation entre IA, ML et DL



1.1 Apprentissage automatique

En 1959, Arthur Samuel définit le ML comme étant la capacité des ordinateurs à apprendre sans être explicitement programmés. En pratique, le programmeur n'a pas besoin d'écrire tous les états possibles. Au lieu de cela, la tâche consiste à trouver un algorithme capable d'extraire des modèles de l'ensemble de données et d'utiliser ces modèles pour construire un modèle prédictif qui se rapproche d'une fonction généralisant les données^[96].

Dans ce contexte, la généralisation signifie obtenir des résultats précis sur de nouvelles données non vues, après avoir traité l'ensemble des données d'apprentissage. L'ensemble des données d'apprentissage est généralement rempli d'exemples provenant d'une distribution de probabilité généralement inconnue^[31].

Comme les ensembles de données d'apprentissage ne sont pas infinis et que l'avenir est incertain, la théorie de l'apprentissage ne permet pas de garantir la performance des algorithmes. Au lieu de cela, les résultats sont fournis avec des limites de certitude. L'exactitude peut être améliorée avec de meilleures données d'apprentissage, ce qui est un sujet en soi, car les données peuvent être sous-ajustées ou sur-ajustées, nous devons donc fournir un équilibre à l'ensemble des données d'apprentissage^[22].

1.2 Types d'apprentissage

Le processus d'apprentissage peut avoir différentes approches. L'apprentissage supervisé et non supervisé sont typiques pour la plupart des applications de ML. Nous n'énumérons ici que quelques-unes des techniques d'apprentissage :

- **Apprentissage supervisé** : L'algorithme d'apprentissage supervisé crée un modèle mathématique qui contient à la fois des entrées et des sorties souhaitées. L'algorithme d'apprentissage supervisé apprend une fonction qui peut être utilisée



pour prévoir la sortie associée à de nouvelles entrées. Les données d'apprentissage contiennent un ensemble de caractéristiques qui incluent une ou plusieurs entrées et leur étiquette assignée^[28].

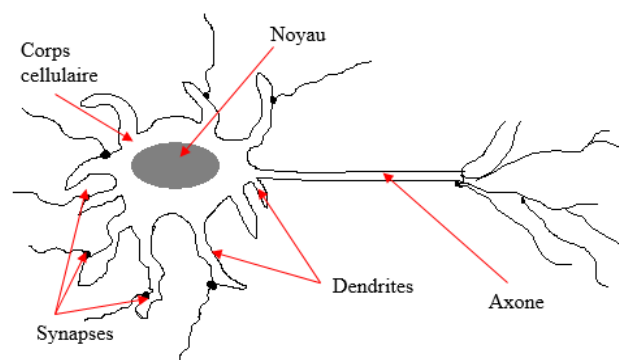
- **Apprentissage non supervisé** : L'apprentissage non supervisé est la capacité de trouver des modèles dans les caractéristiques sans intervention humaine. Cela signifie que l'algorithme ne reçoit que des entrées, et qu'à partir de celles-ci il est supposé classer les données^[28].
- **Apprentissage semi-supervisé** : Il s'agit d'un mélange d'algorithmes d'apprentissage supervisé et non supervisé. Les chercheurs ont découvert que le fait de disposer d'une petite quantité de données étiquetées et de données non étiquetées peut produire de meilleurs résultats (précision d'apprentissage)^[62].
- **Apprentissage par renforcement** : L'apprentissage par renforcement est un domaine de l'apprentissage automatique dans lequel les agents logiciels reçoivent des récompenses pour leur travail. Les agents apprennent progressivement quelles actions produisent la meilleure récompense^[44].
- **Auto-apprentissage** : L'auto-apprentissage est un apprentissage sans récompenses externes et sans conseils d'un enseignant externe^[25].
- **Autres techniques d'apprentissage** : Parmi les autres techniques, citons l'apprentissage par dictionnaire clairsemé^[104], l'apprentissage par transfert^[92], l'apprentissage contradictoire^[81], l'apprentissage de caractéristiques^[30], etc.

2. Réseaux de neurones artificiels

2.1 Neurone biologique

Le neurone biologique (figure 27) est un corps cellulaire constitué de trois parties : les dendrites, qui collectent les informations électriques provenant du système nerveux, ce sont les entrées du neurone ; le soma qui traite ces informations et renvoie un signal électrique de type impulsion ; et l'axone, par lequel le signal sortant est transmis aux neurones voisins.

Figure 27 : Neurone biologique



La faculté d'apprentissage découle des interconnexions complexes entre neurones par l'intermédiaire des synapses. Ceux-ci sont capables de favoriser ou d'empêcher le passage d'un signal électrique et, au fur et à mesure, de faire évoluer leur comportement. C'est ce qu'on attend d'un réseau de neurones artificiels : il doit être capable d'apprendre. À l'instar des synapses, chaque connexion interneuronale du réseau s'adapte et évolue au fil de l'apprentissage. Derrière le formalisme mathématique d'un réseau de neurones artificiels on peut lire en filigrane l'inspiration biologique de plus en plus présente dans les méthodes d'apprentissage.

2.2 Neurone artificiel (ou neurone formel)

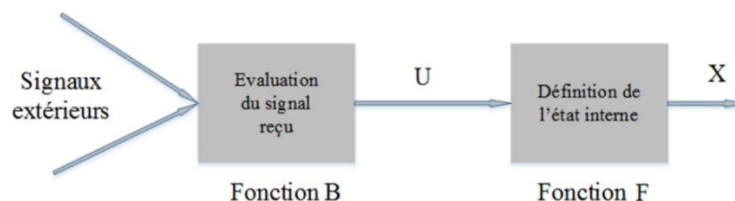
Le neurone formel (figure 28) est l'élément unitaire du réseau de neurones artificiels et renferme deux fonctions essentielles^[153] :

- l'évaluation de la stimulation reçue (fonction E) ;
- l'évaluation de son activation (fonction F).

Il est caractérisé par :

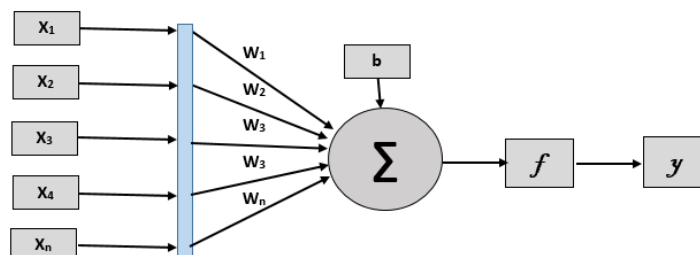
- son état X (binaire, discret ou continu) ;
- le niveau d'activation reçu en entrée U (continu) ;
- le poids des connections en entrée.

Figure 28 : Neurone formel



Le modèle mathématique d'un neurone formel^[154] est le suivant :

Figure 29 : Modèle mathématique du neurone formel



Les entrées X_0 à X_n constituent les vecteurs d'entrée X . La fonction $F(a)$ est appelée fonction d'activation du neurone. Chaque entrée X_i est pondérée par un poids W_i . Ce poids sur les entrées permet au neurone d'apprendre et de modifier ses sorties au fil des itérations. La fonction b présente le biais d'activation. La sortie du neurone est donnée par la relation suivante :

$$y = f(\sum W_i * X_i + b) \quad (20)$$

Il existe plusieurs fonctions d'activation, on peut noter les plus utilisées :

Tableau 7 : Fonctions d'activation usuelles

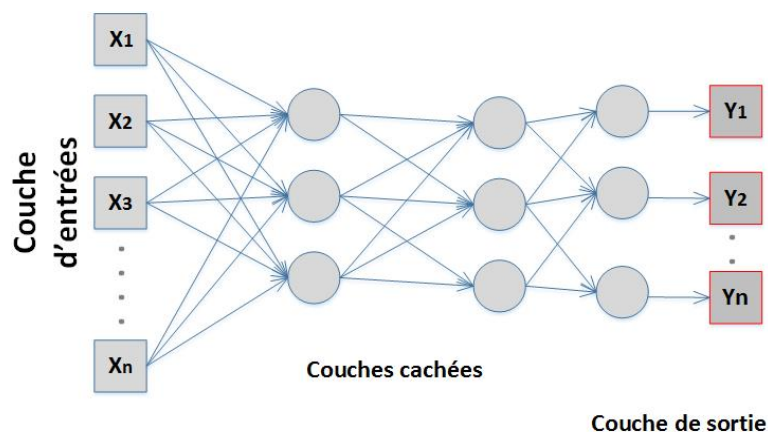
Pas unitaire	$f(x) = \begin{cases} 0 & \text{si } x < 0 \\ 1 & \text{si } x \geq 0 \end{cases}$
Sigmoïde	$f(x) = \frac{1}{1 + e^{-\beta x}}$
Linéaire avec seuil	$f(x) = \begin{cases} 0 & \text{si } x < x_{min} \\ mx + b & \text{si } x_{min} < x < x_{max} \\ 1 & \text{si } x \geq x_{max} \end{cases}$
Gaussien	$f(x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$
Identité	$f(x) = x$

Dans la suite du travail nous utilisons la fonction sigmoïde.

2.3 Réseaux de neurones multicouches

Les réseaux de neurones multicouches comprennent habituellement trois ou quatre couches en tout. La majorité des réseaux de neurones multicouches utilise des perceptrons multicouches (PMC).

Figure 30 : Perceptron multicouche





Un perceptron multicouche est, comme son nom l'indique, un ensemble de neurones en couches, de sorte que les signaux sortants des neurones d'une couche c deviennent des signaux entrants dans les neurones de la couche $c + 1$.

L'architecture générale des PMC consiste en la représentation des neurones en couches (*layers*) consécutives. La première représente la couche d'entrée (*input layer*), la dernière est la couche de sortie (*output layer*). Les couches intermédiaires sont appelées couches cachées (*hidden layers*). Dans cette architecture normalisée^[155], les couches de neurones sont complètement interconnectées, c'est-à-dire que chaque neurone d'une couche est relié à tous les neurones de la couche supérieure.

2.4 Apprentissage

L'apprentissage est une phase du développement d'un réseau de neurones durant laquelle le comportement du réseau est modifié jusqu'à l'obtention du comportement désiré^[156].

Faire qu'un réseau de neurones apprenne, c'est ajuster ses poids de manière à ce que les valeurs données par les neurones de la couche de sortie évoluent dans le sens désiré. Il s'agit dans ce cas d'un apprentissage supervisé : on spécifie la sortie que l'on souhaite voir et le réseau va adapter ses poids et ses biais pour essayer de s'approcher de la valeur de sortie désirée^[157].

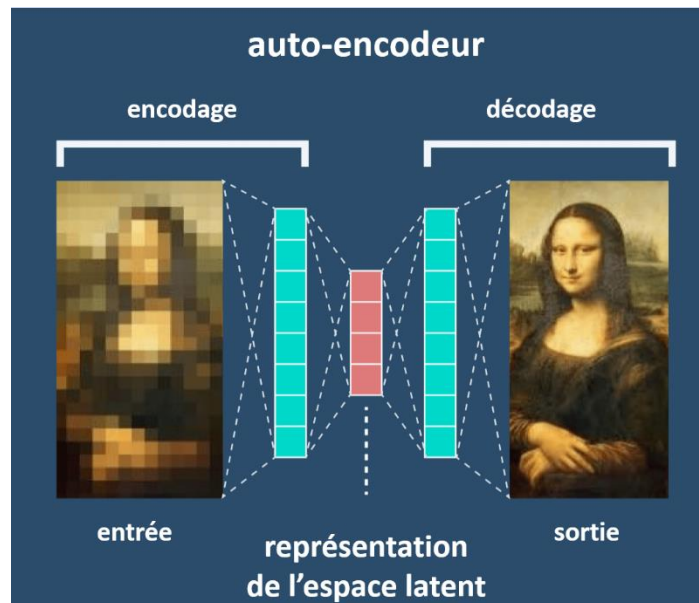
3. Auto-encodeurs

Pour les systèmes les plus courants d'apprentissage automatique, le bruit dans les données reste un problème majeur sur lequel nombre de spécialistes passent des nuits blanches.

Néanmoins, il existe désormais un large éventail de techniques qui permettent aux systèmes informatiques de résoudre plus efficacement les problèmes de compression de données et l'auto-encodeur est l'une d'entre elles, en voici les propriétés en six points :

- Qu'est-ce qu'un auto-encodeur ?
- Architecture d'un auto-encodeur
- Comment entraîner un auto-encodeur ?
- 5 types d'auto-encodeurs
- Applications d'un auto-encodeur
- Synthèse

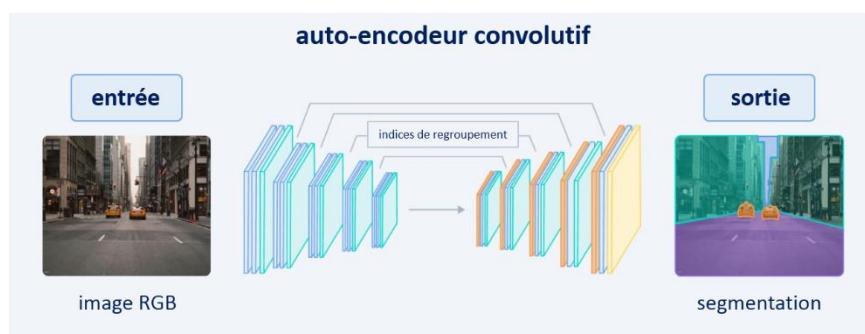
Figure 31 : Auto-encodeur pour le débruitage d'images



3.1 Qu'est-ce qu'un auto-encodeur ?

Un auto-encodeur est un type de réseau neuronal artificiel utilisé pour codifier des données de manière non supervisée. La fonction d'un auto-encodeur est d'interpréter une image (codage) de dimension inférieure et de la traiter en données de dimension supérieure (décodage), généralement pour réduire la dimensionnalité, en entraînant le réseau à capturer les parties les plus significatives de l'image d'entrée.

Figure 32 : Fonction d'un auto-encodeur convolutif



3.2 Architecture d'un auto-encodeur

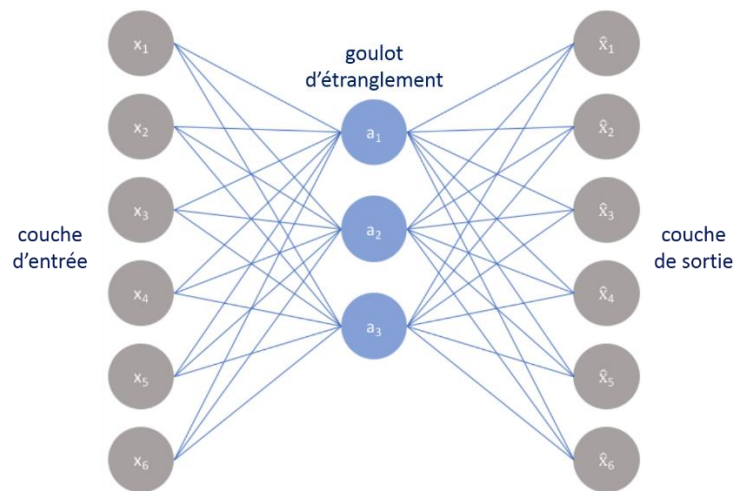
L'architecture d'un auto-encodeur se compose de trois parties :

- L'**encodeur**, module qui compresse les données d'entrée de l'ensemble entraînement-validation-test en une représentation codée qui est généralement inférieure de plusieurs ordres de grandeur aux données d'entrée ;

- Le **goulot d'étranglement**, module qui contient une représentation en données de connaissance compressées qui constitue par conséquent la partie la plus importante du réseau ;
- Le **décodeur**, module qui décompresse la représentation des données de connaissance et reconstruit les données à partir de leur forme codée. Le résultat est ensuite comparé aux valeurs d'entrées.

L'architecture, dans son principe, peut être représentée ainsi :

Figure 33 : Schéma de principe d'un auto-encodeur



L'**encodeur** est un ensemble de blocs convolutifs suivis de modules de « mise en commun » qui compressent les entrées sur le modèle en une section compacte appelée goulot d'étranglement, succédé par le décodeur qui se décompose en une série de modules de suréchantillonnage pour donner à la caractéristique compressée la forme d'une image. Pour un simple auto-encodeur, la sortie est sensiblement la même qu'à l'entrée avec toutefois un bruit réduit ; pour un auto-encodeur variationnel, l'image est totalement nouvelle, formée avec les informations que le modèle fournit en entrée.

Le **goulot d'étranglement**, bien qu'étant le composé le plus important du système, en constitue la plus petite partie. Il limite le flux d'information entre l'encodeur et le décodeur de sorte à ne laisser passer que l'information la plus essentielle. Il est conçu pour capturer le maximum d'information sur l'image d'entrée afin d'en restituer la meilleure représentation en sortie. La structure aide à extraire le maximum d'une image sous forme de données puis à établir des corrélations cohérentes entre les différentes entrées du réseau. En outre, le goulot d'étranglement, en tant que représentation compressée de l'entrée, empêche le réseau neuronal de mémoriser l'entrée et de s'adapter de manière excessive aux données. Il faut retenir qu'en règle générale, plus le goulot d'étranglement est étroit, plus le risque de surajustement est faible. Cependant, un goulot d'étranglement très étroit limite la quantité d'information stockable, ce qui augmente le risque que des informations essentielles échappent aux couches de mise en commun de l'encodeur.



Le **décodeur** est un ensemble de blocs de suréchantillonnage et de convolution qui reconstruit la sortie du goulot d'étranglement. Étant donné que l'entrée du décodeur est une représentation de données instruites compressée, le décodeur sert de décompresseur et reconstruit l'image à partir de ses attributs latents.

3.3 Comment entraîner un auto-encodeur ?

Il s'agit avant tout de définir quatre hyperparamètres :

- la taille du code ou **taille du goulot d'étranglement** est l'hyperparamètre le plus important dans le réglage de l'auto-encodeur. Cette taille détermine le degré de compression des données et peut également servir de terme de régulation ;
- le **nombre de couches**. Comme pour tout réseau neuronal, la profondeur de l'encodeur et du décodeur est également un hyperparamètre important dans ces réglages : si une profondeur élevée augmente la complexité du modèle, une profondeur faible est plus rapide dans le traitement ;
- le **nombre de nœuds par couche** définit les poids utilisés par couche. En général, le nombre de nœuds diminue à chaque couche successive de l'auto-encodeur, car l'entrée de chacune des couches devient plus petite d'une couche à la suivante ;
- la **perte de reconstruction** est une fonction utilisée pour entraîner l'auto-encodeur, elle dépend fortement du type d'entrée et de sortie auquel l'auto-encodeur doit s'adapter. S'il s'agit de données concernant une image, les fonctions de perte les plus populaires pour la reconstruction sont « MSE Loss » et « L1 Loss ». Dans le cas où les entrées et les sorties sont comprises dans la plage $[0,1]$, comme dans MNIST, on peut également utiliser l'entropie croisée binaire comme perte de reconstruction.

3.4 5 types d'auto-encodeur

Le concept d'auto-encodeur pour réseaux neuronaux n'est pas nouveau puisque les premières applications remontent aux années 1980. Initialement utilisé pour la réduction de dimensionnalité et l'apprentissage de caractéristiques, le concept évolue au fil des ans jusqu'à être largement utilisé de nos jours pour l'apprentissage de modèles génératifs de données tels que les cinq plus populaires auto-encodeurs présentés ci-dessous :

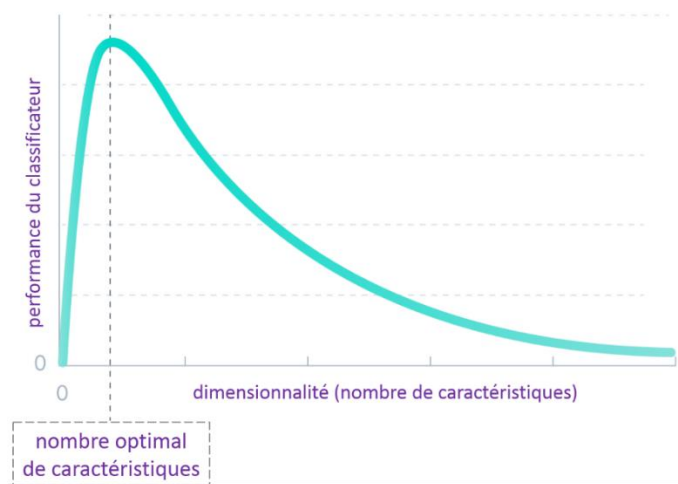
- Auto-encodeur sous-complet ;
- Auto-encodeur épars ;
- Auto-encodeur contractif ;
- Auto-encodeur débruiteur ;
- Auto-encodeur variationnel (pour la modélisation générative).

3.4.1 Auto-encodeur sous-complet

Le fonctionnement d'un auto-encodeur sous-complet est l'un des plus simples. En effet, il capte une image et essaie de prédire la même image en sortie en la reconstruisant à partir de la région compressée au goulot d'étranglement. Les auto-encodeurs sous-complets sont véritablement non supervisés car, la cible étant la même que l'entrée, ils ne prennent aucune forme d'étiquette. Ils sont généralement utilisés pour la génération de l'espace latent (ou goulot d'étranglement) qui forme un substitut compressé des données d'entrée et peuvent être facilement décompressées à l'aide du réseau en cas de besoin.

Cette forme de compression de données peut être modélisée comme une forme de réduction de dimensionnalité :

Figure 34 : Modélisation de la compression de données d'un auto-encodeur sous-complet



La réduction de dimensionnalité évoque des méthodes telles que l'ACP (analyse en composantes principales) qui forment un hyperplan de dimension inférieure pour représenter les données sous une forme de dimension supérieure sans perdre d'information. Mais l'ACP ne peut établir que des relations linéaires, par conséquent elle est désavantagée comparée à l'auto-encodeur sous-complet qui peut apprendre des relations non linéaires et obtenir ainsi de meilleurs résultats en réduction de dimensionnalité.

Cette forme de réduction de dimensionnalité non linéaire où l'auto-encodeur apprend un collecteur non linéaire est également appelée apprentissage du collecteur. En effet, si on supprime toutes les activations non linéaires d'un auto-encodeur sous-complet et qu'on n'utilise que des couches linéaires, on ramène le fonctionnement de l'auto-encodeur sous-complet à pied d'égalité avec l'ACP.

La fonction de perte utilisée pour entraîner un auto-encodeur sous-complet est appelée perte de reconstruction, car elle permet de vérifier la qualité de la reconstruction de l'image à partir de l'entrée. Bien que la perte de reconstruction puisse être n'importe quoi en fonction de l'entrée et de la sortie, on utilise la fonction « L1 Loss » pour décrire le terme (également appelée perte de norme) et représentée par :



$$L = |x - \hat{x}| \quad (21)$$

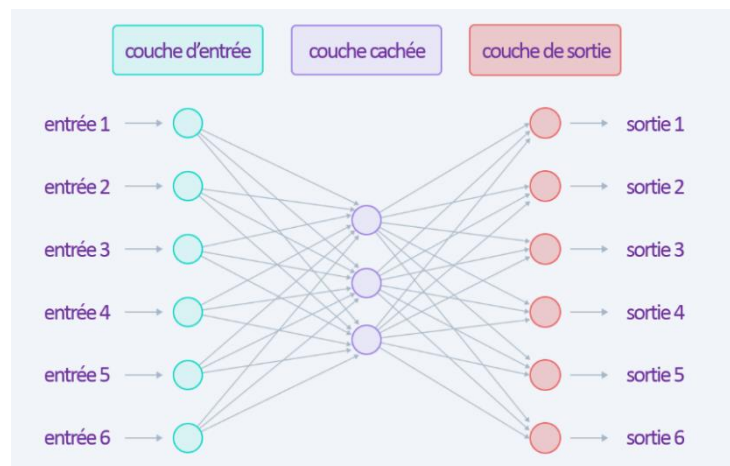
où $\hat{x}$ représente la sortie prédite et x représente les valeurs entrées.

Comme la fonction de perte ne comporte pas de terme de régulation explicite, la seule méthode permettant de s'assurer que le modèle ne mémorise pas les données d'entrée consiste à réguler la taille du goulot d'étranglement et le nombre de couches cachées dans cette partie réseau-architecture.

3.4.2 Auto-encodeur épars

Un auto-encodeur épars est similaire à un auto-encodeur sous-complet dans la mesure où il utilise la même image comme entrée et comme image d'authentification, cependant, le moyen par lequel l'encodage de l'information est réglé est sensiblement différent :

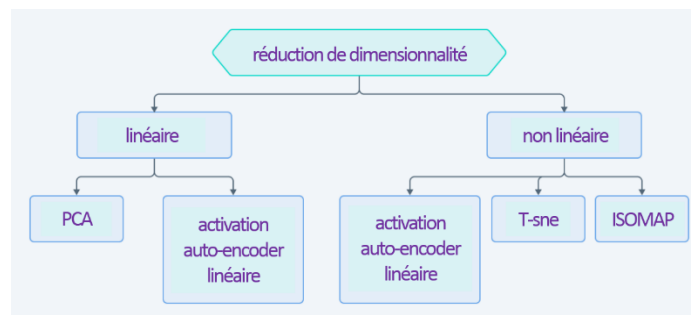
Figure 35 : Schéma de principe d'un auto-encodeur épars



Alors qu'un auto-encodeur sous-complet est finement réglé en ajustant la taille du goulot d'étranglement, l'auto-encodeur épars est régulé en modifiant le nombre de nœuds de chaque couche cachée. Comme il n'est pas possible de concevoir un réseau neuronal dont le nombre de nœuds dans les couches cachées soit flexible, un auto-encodeur épars fonctionne en pénalisant l'activation de certains neurones dans les couches cachées. Autrement dit, la fonction de perte comporte un terme qui calcule le nombre de neurones qui ont été activés et applique une pénalité qui est directement proportionnelle à ce nombre. Cette pénalité, appelée fonction de sparsité, empêche le réseau neuronal d'activer davantage de neurones et sert de régulateur.



Figure 36 : Méthodes de réduction de la dimensionnalité



Sur un auto-encodeur simple la régulation s'applique en créant une pénalité sur la taille des poids aux nœuds, la fonction de sparsité quant à elle, crée une pénalité sur le nombre de nœuds activés. Cette forme de régulation permet au réseau au cours de la reformation d'avoir des nœuds dans les couches cachées dédiés à la recherche de caractéristiques spécifiques dans les images, elle traite le problème de régulation comme un problème distinct de celui de l'espace latent. On peut ainsi fixer la dimensionnalité spatiale latente au niveau du goulot d'étranglement sans se soucier de la régulation.

Il existe deux principales manières d'incorporer le terme de régulation de sparsité dans la fonction de perte :

- 1 - L1 Loss : on ajoute la magnitude du régularisateur de sparsité comme on le fait avec un simple auto-encodeur :

$$L = |x - \hat{x}| + \lambda \sum_i^k |a_i^{(h)}| \quad (22)$$

où h représente la couche cachée, i représente l'image dans le minilot et a représente l'activation.

- 2 - La divergence de Kullback-Leibler (divergence KL) : ici, on considère les activations sur une collection d'échantillons à la fois, plutôt que les additionner comme dans la méthode L1 Loss. On contraint l'activation moyenne de chaque neurone sur cette collection.

En considérant que la distribution idéale est celle de Bernoulli, on inclut la divergence KL dans la perte pour réduire la différence entre la distribution existante des activations et la distribution idéale (Bernoulli) :

$$L = |x - \hat{x}| + \sum_i^k KL(\rho || \hat{\rho}_j) \quad (23)$$

où $\hat{\rho}_j = \frac{1}{m} \sum_i^k [a_i^{(h)}(x)]$ et j désigne le neurone spécifique pour la couche h et m une collection de m échantillons est en cours ici, chacun étant noté x .

3.4.3 Auto-encodeur contractif

Comme les autres auto-encodeurs, l'auto-encodeur contractif a pour tâche d'apprendre une représentation de l'image tout en la faisant passer par un goulot d'étranglement et en la



reconstruisant dans le décodeur. Un auto-encodeur contractif comporte également un terme de régulation pour empêcher le réseau d'apprendre la fonction d'identité et de transformer l'entrée en sortie.

Le fonctionnement d'un auto-encodeur contractif repose sur le postulat que des entrées similaires doivent avoir des codages similaires et une représentation similaire de l'espace latent. Cela signifie que l'espace latent ne doit pas varier considérablement pour des variations mineures à l'entrée. Pour entraîner un modèle qui fonctionne avec cette contrainte, on doit s'assurer que les dérivées des activations de la couche cachée sont petites par rapport à l'entrée.

Mathématiquement :

$$\delta h / \delta x \quad (24)$$

où h représente la couche cachée et x la couche d'entrée.

La grandeur précédente doit être aussi petite que possible.

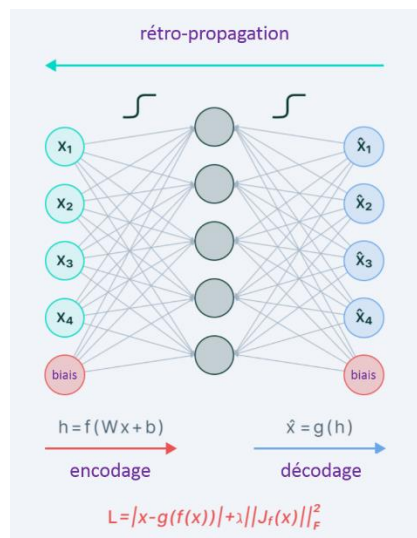
À noter que dans la fonction de perte (formée de la norme des dérivées et de la perte de reconstruction) les deux termes se contredisent. Alors que dans la perte de reconstruction le modèle distingue les différences entre deux entrées et observe les variations dans les données, dans les dérivées de la norme de Frobenius le modèle doit être capable d'ignorer les variations dans les données d'entrée. Le fait de réunir ces deux conditions contradictoires en une seule fonction de perte permet d'entraîner un réseau dont les couches cachées ne capturent plus que les informations les plus essentielles. Ces informations sont nécessaires pour segmenter les images et ignorer les informations qui ne sont pas de nature discriminatoire, et sont donc sans importance.

La fonction de perte totale peut être exprimée mathématiquement comme suit :

$$L = |x - \hat{x}| + \lambda \sum_i^k \|\nabla_x a_i^{(h)}(x)\|^2 \quad (25)$$

où h est la couche cachée pour laquelle le gradient est calculé. Le gradient est additionné sur tous les échantillons d'entraînement, et une norme de Frobenius est prise comme telle $\nabla_x a_i^{(h)}$.

Figure 37 : Modèle de la rétropropagation

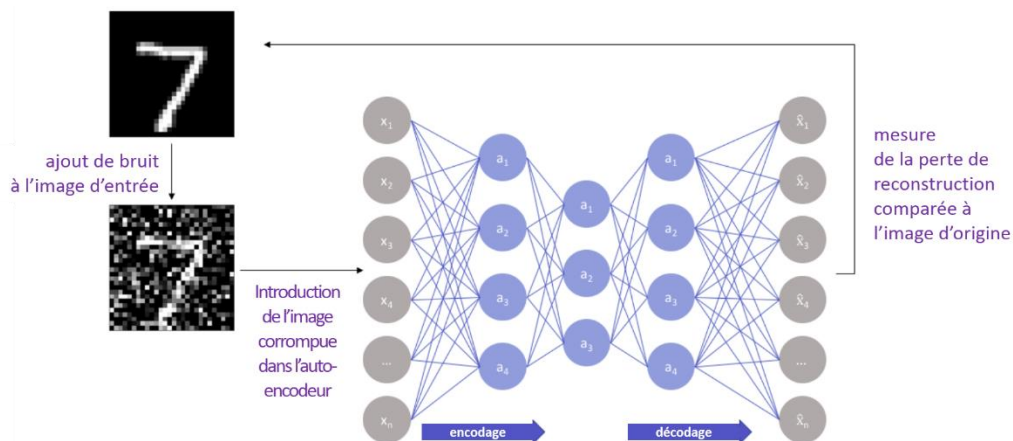


3.4.4 Auto-encodeur débruiteur

Un auto-encodeur débruiteur, comme son nom l'indique, supprime le bruit d'une image. Contrairement aux auto-encodeurs évoqués plus haut, celui-ci est le premier de sa catégorie à ne pas avoir l'image d'entrée comme image d'authentification.

Dans un auto-encodeur débruiteur, une version de l'image dans laquelle du bruit a été ajouté par altérations numériques est fournie. L'image avec bruit est transmise à l'architecture de l'encodeur-décodeur et la sortie est comparée à l'image d'origine.

Figure 38 : Modèle d'un auto-encodeur débruiteur



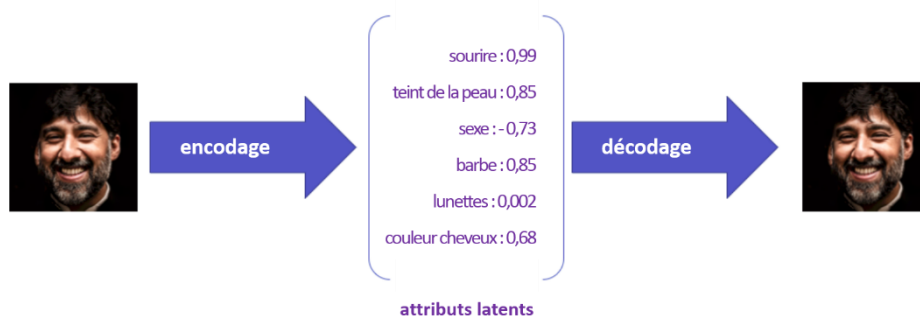
L'auto-encodeur débruiteur se débarrasse du bruit en apprenant une représentation de l'entrée où le bruit peut être filtré facilement. Alors qu'il semble difficile d'éliminer directement le bruit de l'image, l'auto-encodeur s'en charge en « mappant » les données d'entrée dans un collecteur de dimension inférieure (comme pour un auto-encodeur sous-complet) où le filtrage du bruit devient beaucoup plus facile.

Un auto-encodeur débruiteur fonctionne essentiellement à l'aide d'une réduction non linéaire de dimensionnalité. La fonction de perte généralement utilisée dans ces types de réseaux est L2 ou L1 Loss.

3.4.5 Auto-encodeur variationnel

Les auto-encodeurs standard et variationnels apprennent à représenter l'entrée uniquement sous une forme compressée appelée espace latent ou goulot d'étranglement. Par conséquent, l'espace latent formé après l'apprentissage du modèle n'est pas nécessairement continu et, de fait, peut ne pas être facile à interpoler. Par exemple, ci-dessous est représenté ce qu'un auto-encodeur variationnel apprendrait de l'entrée :

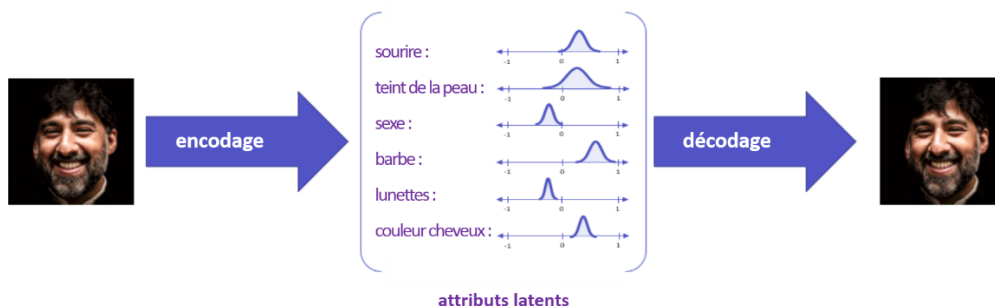
Figure 39 : Modèle de l'auto-encodeur variationnel



Bien que ces attributs expliquent l'image et puissent être utilisés pour reconstruire l'image à partir de l'espace latent compressé, ils ne permettent pas d'exprimer les attributs latents de manière probabiliste. Un auto-encodeur variationnel traite ce sujet spécifique et exprime les attributs latents sous forme de distribution de probabilité, ce qui conduit à la formation d'un espace latent continu qui peut être facilement échantillonné et interpolé.

Lorsqu'il reçoit la même entrée, un auto-encodeur variationnel construit les attributs latents de la manière suivante :

Figure 40 : Construction des attributs latents par l'auto-encodeur variationnel



Les attributs latents sont ensuite échantillonnés à partir de la distribution latente formée et transmis au décodeur qui reconstruit l'entrée.



La motivation derrière l'expression des attributs latents sous forme de distribution de probabilité peut être très facilement comprise par le biais d'expressions statistiques. Ainsi, on cherche à identifier les caractéristiques du vecteur latent z qui reconstruit la sortie en fonction d'une entrée particulière. En d'autres termes, on étudie les caractéristiques d'un vecteur latent sur une certaine sortie : $x [p(z/x)]$.

Si l'estimation de la distribution devient mathématiquement impossible, une option beaucoup plus simple et facile consiste à construire un modèle paramétré qui estime la distribution. Pour cela, il faut minimiser la divergence KL entre la distribution originale et la distribution paramétrée. En exprimant la distribution paramétrée sous la forme q , on peut déduire les attributs latents possibles utilisés dans la reconstruction de l'image.

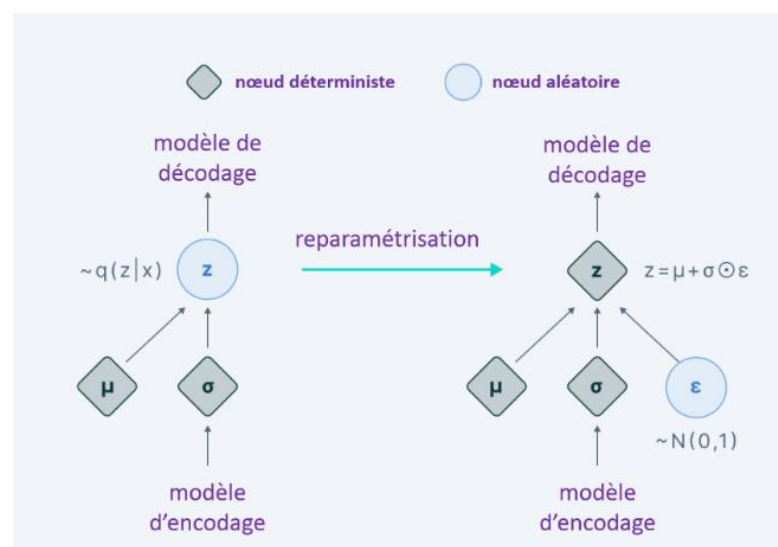
En supposant que l'antériorité z soit un modèle gaussien multivarié, on peut construire une distribution paramétrée contenant deux paramètres : la moyenne et la variance. La distribution correspondante est ensuite échantillonnée et transmise au décodeur qui procède par la suite à la reconstruction de l'entrée à partir des points d'échantillonnage.

Bien que cela semble aisé en théorie, cela devient impossible à mettre en œuvre car la rétro-propagation ne peut pas être définie pour un processus d'échantillonnage aléatoire effectué avant de fournir les données au décodeur. Pour contourner cet obstacle on utilise la reparamétrisation : un moyen astucieux de contourner le processus d'échantillonnage du réseau neuronal.

La reparamétrisation échantillonne aléatoirement une valeur ε à partir d'une gaussienne unitaire, puis elle la met à l'échelle de la variance σ de la distribution latente et la décale de la moyenne μ de cette dernière. Ensuite, on considère le processus d'échantillonnage comme en dehors de ce que le pipeline de rétro-propagation traite, et la valeur échantillonnée ε agit simplement comme une autre entrée du modèle qui est alimentée au niveau du goulot d'étranglement.

Schématiquement, ce qui est atteint peut être exprimé comme suit :

Figure 41 : Modèle de la reparamétrisation





L'auto-encodeur variationnel permet donc d'apprendre des représentations lisses d'états latents des données d'entrée.

Pour entraîner un auto-encodeur variationnel, on utilise deux fonctions de perte : la perte de reconstruction et la divergence KL. Alors que la perte de reconstruction permet à la distribution de décrire correctement l'entrée en se concentrant uniquement sur la minimisation de la perte de reconstruction, le réseau apprend des distributions très étroites, proches des attributs latents discrets. La perte de divergence KL empêche le réseau d'apprendre des distributions étroites et tente de rapprocher la distribution d'une distribution normale unitaire.

La fonction de perte résumée peut être exprimée comme suit :

$$L = |x - \hat{x}| + \beta \sum_i KL(q_i(z|x) || \mathcal{N}(0,1)) \quad (26)$$

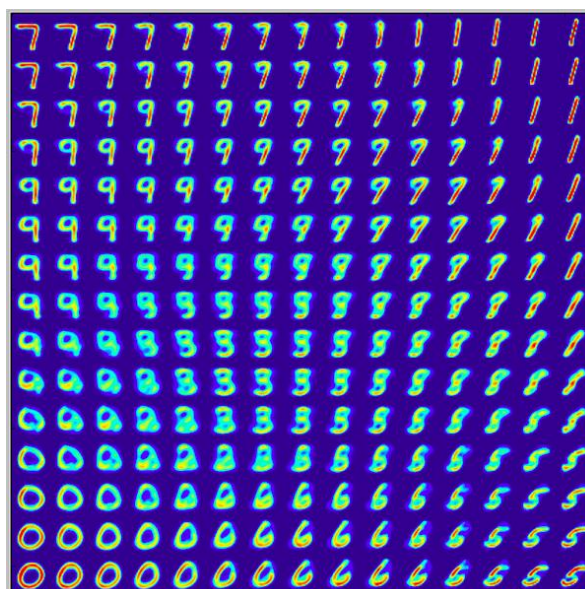
où $\mathcal{N}$ désigne la distribution normale unitaire et β désigne un facteur de pondération.

L'utilisation principale des auto-encodeurs variationnels peut être vue dans la modélisation générative.

L'échantillonnage à partir de la distribution latente entraînée et la transmission du résultat au décodeur peuvent conduire à la génération de données dans l'auto-encodeur.

Un échantillon de chiffres MNIST généré par l'entraînement d'un auto-encodeur variationnel est présenté ci-dessous :

Figure 42 : Échantillon de chiffres MNIST généré par l'entraînement d'un auto-encodeur variationnel



3.5 Applications d'un auto-encodeur

Voici quatre cas les plus courants d'utilisation d'un auto-encodeur.

3.5.1 Réduction de la dimensionnalité

Les auto-encodeurs sous-complets sont ceux utilisés pour la réduction de dimensionnalité. Ils peuvent être utilisés comme étape de prétraitement pour la réduction de dimensionnalité car ils effectuent des réductions de dimensionnalité rapides et précises sans perdre beaucoup d'information. En outre, alors que les procédures de réduction de dimensionnalité telles que l'ACP ne peuvent effectuer que des réductions de dimensionnalité linéaires, les auto-encodeurs sous-complets peuvent effectuer des réductions de dimensionnalité non linéaires à grande échelle.

3.5.2 Débruitage d'image

Les auto-encodeurs tels que l'auto-encodeur débruiteur peuvent être utilisés pour effectuer un débruitage efficace et très précis des images. Contrairement aux méthodes traditionnelles de débruitage, les auto-encodeurs ne cherchent pas le bruit mais ils extraient l'image de données bruitées qui leur sont fournies en apprenant une représentation de celles-ci. Cette représentation est ensuite décompressée pour former une image exempte de bruit. Un auto-encodeur débruiteur peut donc débruiter des images complexes qui ne peuvent être débruitées par les méthodes traditionnelles.

3.5.3 Génération de données d'image et de séries temporelles

Un auto-encodeur variationnel peut être utilisé pour générer des données d'image et de séries temporelles. La distribution paramétrée au niveau du goulot d'étranglement de l'auto-encodeur peut être échantillonnée de manière aléatoire pour générer des valeurs discrètes pour les attributs latents qui sont ensuite transmises au décodeur, ce qui permet de générer des données d'image. Un auto-encodeur variationnel peut également être utilisé pour modéliser des données de séries temporelles comme la musique.

3.5.4 Détection d'anomalies

Un auto-encodeur sous-complet peut également être utilisé pour la détection d'anomalies, par exemple, un auto-encodeur qui est formé sur un ensemble de données spécifique P . Pour toute image échantillonnée dans l'ensemble des données d'entraînement, l'auto-encodeur est tenu de donner une faible perte de reconstruction et est censé reconstruire l'image telle quelle. Cependant, pour toute image qui n'est pas présente dans l'ensemble des données d'apprentissage, l'auto-encodeur ne peut pas effectuer la reconstruction car les attributs latents ne sont pas adaptés à l'image spécifique qui n'a jamais été vue par le réseau. Par conséquent, l'image aberrante génère une perte de reconstruction très élevée et peut facilement être identifiée comme une anomalie à l'aide d'un seuil approprié.



3.6 Synthèse

Un auto-encodeur est une technique d'apprentissage non supervisé pour les réseaux neuronaux qui apprend des représentations de données efficaces (codage) en entraînant le réseau à ignorer le bruit du signal.

Un auto-encodeur peut être utilisé pour le débruitage d'images, la compression d'images et, dans certains cas, la génération de données d'images.

Si le principe d'un auto-encodeur semble abordable à première vue par son cadre théorique simple, il est assez difficile de lui faire apprendre une représentation de l'entrée qui soit significative.

4. Travaux connexes

Les ondes millimétriques sont les composants majeurs des systèmes de communication sans fil, ceci induit des défis dans la formation et la sélection du meilleur faisceau. Plusieurs travaux portent sur la sélection du faisceau optimal d'ondes millimétriques à l'aide de techniques d'apprentissage automatique et d'apprentissage en profondeur. Wang *et al.* proposent une approche d'apprentissage automatique combinée à la connaissance de la situation pour l'entraînement au faisceau. Ils prennent en compte les observations de l'environnement du véhicule, les emplacements des récepteurs et les véhicules environnants^[14]. Les auteurs utilisent des réseaux de neurones multimodaux pour apprendre les faisceaux tout en exploitant un ensemble de données provenant de l'environnement de communication sans fil. De plus, ils évaluent leur approche sur un ensemble de données de véhicules réels^[15]. Ils utilisent la classification des images et les réseaux résiduels pour prédire le faisceau et le blocage directement à partir des caméras RVB et des canaux inférieurs à 6 GHz^[16]. Tarun *et al.* proposent un modèle de réseau neuronal profond pour la prédiction du faisceau d'ondes millimétriques en cartographiant les modèles complexes des intensités du signal reçu au faisceau spatial optimal du récepteur. Les auteurs se concentrent sur un petit nombre de faisceaux pour détecter plus rapidement le meilleur faisceau^[17]. Alrabeiah *et al.* réalisent un ensemble de données avancé avec différentes stations de base, des utilisateurs mobiles et une dynamique riche. Ils proposent également un suivi visuel sans fil du faisceau d'ondes millimétriques avec des réseaux de neurones récurrents^[18]. Bian *et al.* Mettent en place un réseau de neurones à deux entrées, qu'ils nomment fusionNet, pour prédire le meilleur faisceau d'ondes millimétriques. Ils proposent d'utiliser la parcimonie des canaux et l'augmentation des données pour éviter le problème de sur-ajustement^[19]. Ma *et al.* utilisent les informations sur l'état du canal à basse fréquence pour extraire des informations hors bande et prédire le faisceau d'ondes millimétriques à l'aide de l'apprentissage en profondeur^[20]. Catak *et al.* proposent un réseau de neurones artificiels pour prédire les faisceaux d'ondes millimétriques à travers un ensemble de données réaliste, puis appliquent une méthode d'apprentissage automatique contradictoire pour augmenter la sécurité de la prédiction^[21]. Alrabieh *et al.* proposent d'apprendre des fonctions de mappage et de prédire le faisceau optimal et les blocages directement à partir du canal inférieur à 6 GHz^[22]. Ying *et al.* utilisent la reconnaissance d'objets pour localiser les utilisateurs en fonction des images RVB capturées par les caméras. Ensuite, ils utilisent un perceptron multicouche pour prédire les angles entre les utilisateurs et les

caméras. Ils sélectionnent par la suite le faisceau optimal en fonction des informations d'angle extraites et d'un livre de codes de faisceau prédéfini^[23]. Aldalbahi *et al.* étudient la mémoire à long terme pour la prédiction instantanée des directions alternatives du faisceau lorsque le faisceau principal est bloqué^[24].

Dans la même perspective, ce travail porte sur la prédiction du faisceau optimal utilisant une solution simple et en temps réel pour augmenter la fiabilité et la précision des systèmes à ondes millimétriques pour la sixième génération d'applications mobiles sans coût ni temps élevés. Dans notre approche, la station de base des services utilise la séquence de faisceaux qu'elle proposait dans le passé aux utilisateurs de mobiles pour prédire les meilleurs faisceaux. Cela permet d'améliorer la fiabilité du système et de réduire les temps de latence. À cette fin, nous développons un modèle d'apprentissage en profondeur basé sur des perceptrons multicouches et un auto-encodeur empilé pour l'extraction de caractéristiques. Les résultats de la simulation montrent que la solution proposée prédit les faisceaux optimaux avec une erreur minimale, améliorant ainsi la fiabilité du système de grande antenne à ondes millimétriques.

5. Préliminaires et présentation du jeu de données

5.1 Prédiction de la formation de faisceaux à ondes millimétriques

Les réseaux mobiles 6G sont un terrain approprié pour transmettre de grandes quantités de données avec une faible latence. Cependant, avec le déploiement de plus de bande passante, les réseaux 6G sont également confrontés au problème d'une perte élevée de trajet dans la bande MMwave. La technologie de formation de faisceaux est appliquée comme solution pour améliorer le gain d'antenne en établissant une liaison de transmission hautement directionnelle entre l'équipement utilisateur (UE) et les stations de base de nœuds mmWave (gNB).

La formation de faisceaux est une technique de filtrage spatial qui utilise un réseau de radiateurs pour capturer ou émettre de l'énergie dans une direction spécifique d'une ouverture. Une amélioration par rapport à la transmission-réception omnidirectionnelle est le gain de transmission-réception. Dans les systèmes de communication modernes, les systèmes d'antennes intelligentes sont utilisés pour combiner le gain de réseau avec le gain de diversité tout en supprimant les interférences et en augmentant la capacité de la liaison de communication. Ceci est réalisé en utilisant un réseau à commande de phase et un dispositif de rayonnement composé de plusieurs éléments avec une configuration géométrique spécifique pour contrôler le faisceau électronique.

La formation de faisceaux est une exigence pour les réseaux 6G qui fonctionnent dans la bande de fréquence mmWave pour améliorer la couverture du réseau. Selon l'unité qui effectue le traitement du signal, c'est-à-dire dans la bande de base ou la gamme RF, la formation de faisceau est classée soit en formation de faisceau numérique, analogique (ABF) ou hybride (HBF)^[25].

Dans ce travail, nous étudions l'ABF. Dans l'ABF, un signal unique est délivré via des déphaseurs analogiques à chaque élément d'antenne du réseau, au niveau duquel le signal est amplifié et rayonné vers le récepteur souhaité. La variation d'amplitude/phase est appliquée au signal analogique en fin de transmission, en additionnant les signaux des différentes antennes avant



la conversion analogique-numérique. La formation de faisceau analogique est le moyen le plus économique pour construire un réseau de formation de faisceau, mais elle ne peut gérer et générer qu'un seul faisceau de signal. La figure 43 montre l'architecture de la formation de faisceau analogique. Elle est réalisée par un réseau phasé avec uniquement une chaîne de radiofréquence contrôlée par un convertisseur numérique-analogique dans l'émetteur ou un convertisseur analogique-numérique dans le récepteur. Une chaîne RF d'émetteur se compose d'un convertisseur élévateur de fréquence, d'un amplificateur de puissance, etc. ; une chaîne RF réceptrice est constituée d'un amplificateur à faible bruit, d'un abaisseur de fréquence, etc.

La figure 43(a) illustre un émetteur de formation de faisceau analogique : le signal en bande de base transmis est d'abord modulé. Ce signal radio est divisé avec un diviseur de puissance et passe à travers le formateur de faisceau, qui peut modifier l'amplitude et le déphasage (θ_k) des signaux dans chacun des chemins menant à une pile d'antennes. Le diviseur de puissance dépend du nombre d'antennes utilisé dans la pile d'antennes a_k .

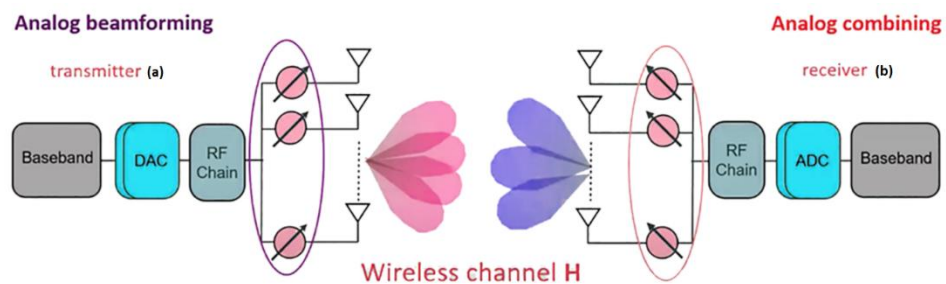
La figure 43(b) illustre un récepteur de formation de faisceau analogique. Comme le montre le schéma fonctionnel du récepteur, une pondération complexe est appliquée au signal de chaque antenne du réseau. La pondération complexe comprend à la fois l'amplitude et le déphasage (formules 1 et 2). Ensuite, les signaux sont combinés pour produire un signal de sortie, qui fournit le modèle directionnel souhaité d'un réseau d'antennes^[25].

$$W_k = a_k * e^{j\sin(\theta)} \quad (27)$$

$$W_k = a_k * \cos(\theta_k) + ja_k \sin(\theta_k) \quad (28)$$

où W_k est les poids complexes de la K-ième antenne de réseau, a_k est l'amplitude relative des poids et θ_k est le déphasage $W_k a_k \theta_k$

Figure 43 : Architecture de formation de faisceau analogique du récepteur et de l'émetteur (RF : radiofréquence, DAC : convertisseurs numérique-analogique, ADC : convertisseur analogique-numérique)



(RF : radiofréquence, DAC : convertisseurs numérique-analogique, ADC : convertisseur analogique-numérique)

Dans les communications à ondes millimétriques, des réseaux d'antennes à grande échelle sont utilisés pour générer des faisceaux hautement directionnels afin de compenser les graves pertes de trajet. La prédiction de faisceau évite d'estimer des matrices de canaux de grande dimension en sélectionnant le meilleur faisceau. L'utilisation d'algorithmes d'apprentissage automatique est une nouvelle solution pour l'apprentissage des ondes millimétriques et le balayage d'un grand nombre de faisceaux étroits. Les faisceaux dépendent des conditions

environnementales telles que l'emplacement de l'utilisateur et l'emplacement de la station de base, les meubles, les arbres, les bâtiments, etc. Il est difficile de définir ces conditions environnementales dans des équations de forme fermée. La solution consiste à utiliser des modèles de faisceaux omnidirectionnels et quasi omnidirectionnels pour prédire le vecteur de formation de faisceau RF optimal. L'avantage d'utiliser ces diagrammes de faisceau est que les réflexions et les diffractions du signal pilote peuvent être prises en compte.

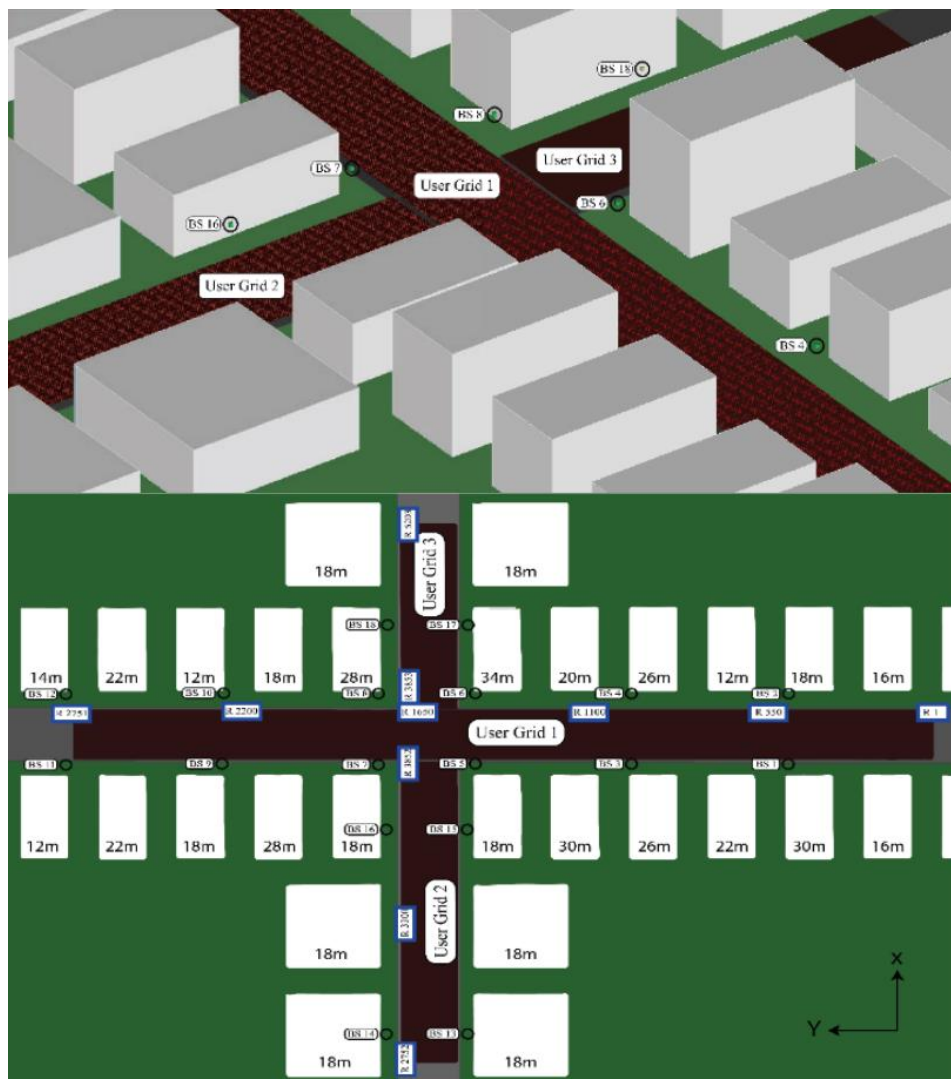
Une solution d'apprentissage en profondeur se compose de deux phases : la formation et la prédiction. Tout d'abord, le modèle d'apprentissage en profondeur apprend le faisceau en fonction des pilotes omnidirectionnels reçus. Ensuite, le modèle utilise les données formées pour prédire l'état actuel du vecteur de formation de faisceau RF. Les sections suivantes présentent les principaux aspects de notre approche de prédiction de formation de faisceau suggérée.

5.2 Présentation du jeu de données

Dans ce travail, nous utilisons le générateur de jeux de données DeepMIMO pour les applications mmWave présentées dans [26], en appliquant le scénario O1. La figure 44 illustre la vue à vol d'oiseau et la vue de dessus du scénario de traçage de rayons O1, qui représente une zone extérieure avec deux rues et une intersection, avec 18 BS et plus d'un million d'utilisateurs et une fréquence de 60 GHz. Pour générer ce jeu de données, nous considérons les paramètres suivants :

- Les stations de base actives : 3, 4, 5 et 6.
- Les utilisateurs actifs : 1000-1300
- Le nombre d'antennes BS : $M_x = 1$, $M_y = 32$, $M_z = 8$
- L'espacement des antennes : 0,5
- Le nombre de sous-porteuses OFDM : 1024
- Le facteur d'échantillonnage OFDM : 1
- La limite OFDM : 64
- Le nombre de chemins : 5

Figure 44 : Vue plongeante et vue de dessus du scénario O1^[26]

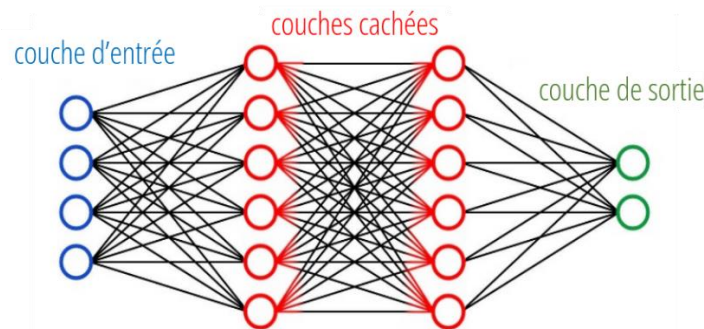


5.3 Modèle étudié de *deep learning*

Les réseaux de neurones artificiels (RNA) sont des réseaux informatiques d'inspiration biologique. Parmi les différents types de réseaux de neurones artificiels, nous nous concentrons sur les perceptrons multicouches (PMC) avec des algorithmes d'apprentissage par rétro-propagation. Les PMC sont couramment utilisés pour résoudre divers problèmes, et sont complémentaires aux réseaux de neurones à anticipation. Ils se composent de trois types de couches : la couche d'entrée, la couche de sortie et la couche masquée, comme illustré dans la figure 45. La couche d'entrée reçoit les entités d'entrée à traiter. Les tâches requises telles que la prédiction et la classification sont effectuées par la couche de sortie. Un nombre arbitraire de couches cachées situées entre les couches d'entrée et de sortie est le véritable moteur de calcul de PMC. Comme pour les réseaux à anticipation, les données dans PMC circulent de la couche d'entrée à la couche de sortie dans le sens direct. Les neurones d'un PMC sont formés à l'aide d'un algorithme d'apprentissage par rétro-propagation. PMC est conçu pour approximer n'importe quelle fonction continue et peut résoudre des problèmes linéairement

inséparables. Les principaux cas d'utilisation de PMC sont la reconnaissance de formes, la classification, la prédiction et l'approximation^[27].

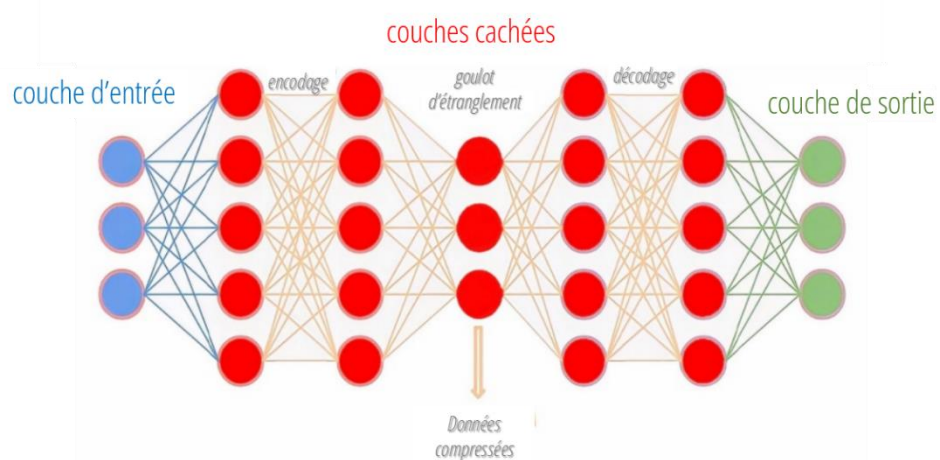
Figure 45 : Architecture des perceptrons multicouches



Un auto-encodeur est un type spécial d'ANN qui condense les entrées dans un code de faible dimension et reconstruit la sortie à partir de cette représentation. Le code est la compression des entrées. Un auto-encodeur est composé de trois composants principaux : l'encodeur, le code et le décodeur. L'encodeur comprime l'entrée et génère un code. Le décodeur utilise alors uniquement ce code pour reconstruire la sortie.

La figure 46 décrit l'architecture de l'auto-encodeur. Tout d'abord, l'entrée passe par l'encodeur, qui est un réseau neuronal artificiel entièrement connecté, pour générer le code (*bottleneck*). Le décodeur a une conception ANN similaire mais utilise uniquement le code compressé pour générer la sortie. Le but est d'obtenir la même sortie que l'entrée. À noter que l'architecture du décodeur est l'image correspondante du codeur. Ce n'est pas obligatoire mais c'est généralement le cas. La seule condition est que les dimensions de l'entrée et de la sortie soient identiques. Tout ce qui se trouve entre les deux peut être joué.

Figure 46 : Architecture de l'auto-encodeur

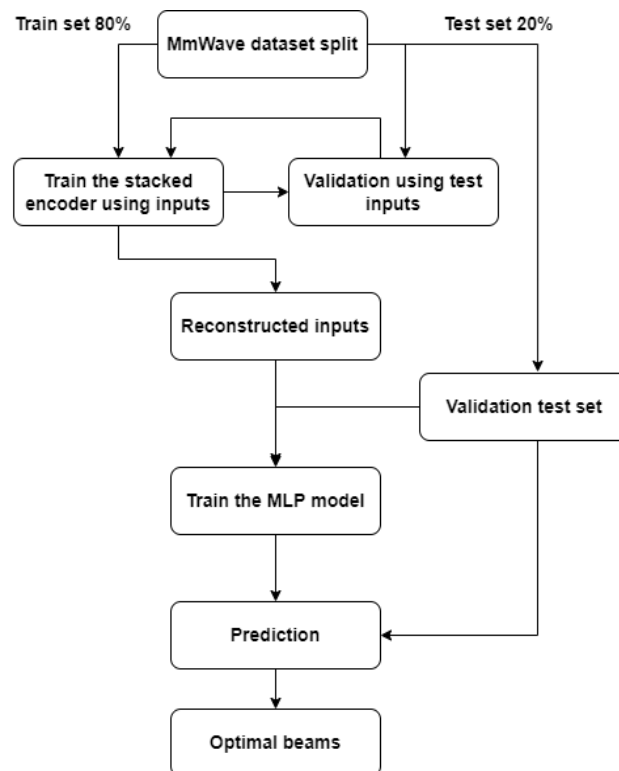




6. Approche proposée

Les données des canaux mmWave sont non linéaires et de nature complexe. Pour la prédiction des canaux mmWave, nous proposons d'utiliser un auto-encodeur empilé pour extraire les caractéristiques du jeu de données, et un modèle PMC pour prédire les faisceaux optimaux. La figure 47 illustre l'architecture de notre modèle construit successivement à travers l'acquisition et le prétraitement des données, l'extraction des caractéristiques et le processus de formation.

Figure 47 : Notre proposition d'entraînement de faisceau utilisant un auto-encodeur empilé et des perceptrons multicouches



6.1 Acquisition et génération de jeux de données

Dans cette étape, nous utilisons l'ensemble de données décrit plus haut. Pour créer le cadre de l'ensemble de données DeepMIMO illustré à la figure 44, nous définissons le scénario et les paramètres de traçage de rayons pour générer l'ensemble de données adopté pour notre application. Plus précisément, les étapes de création de cet ensemble de données peuvent être résumées comme suit :

- Télécharger le fichier de code de génération DeepMIMO à partir du site web de l'ensemble de données DeepMIMO[28] et le décompresser ;
- Télécharger le fichier de sortie du traçage de rayons pour le scénario sélectionné à partir du site web de l'ensemble de données DeepMIMO^[28]. Par exemple, scénario « O1 » et l'étendre ;
- Ajouter le dossier du scénario de traçage de rayons, par exemple le dossier « O1 »,

au chemin "DeepMIMO Dataset Generation/RayTracing Scenarios/" ;

- Ouvrir le fichier "DeepMIMODatasetGeneration.m" et ajuster les paramètres de l'ensemble de données DeepMIMO ;
- Dans la fenêtre de commande MATLAB, appeler la fonction [DeepMIMODataset]=DeepMIMODatasetGenerator(). Cette fonction génère le jeu de données DeepMIMO en fonction du scénario de traçage de rayons défini et des paramètres supposés.

Pour préparer l'ensemble de données pour les étapes suivantes, nous le divisons en deux sous-ensembles : l'ensemble d'apprentissage, qui contient 80 % des données globales, et l'ensemble de test, qui contient les 20 % restants. Il est utilisé pour évaluer la précision de la prédiction.

6.2 Extraction de caractéristiques

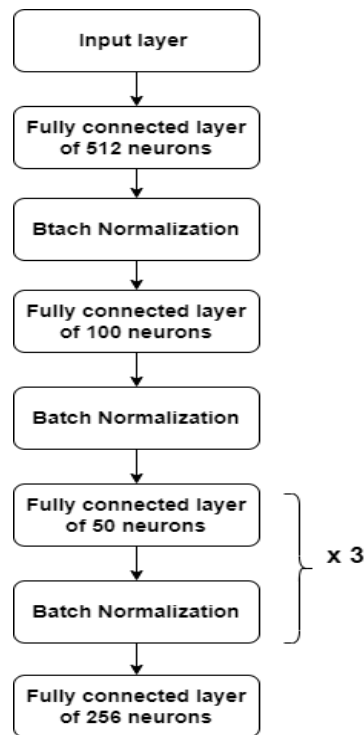
Cette étape consiste à extraire les traits et caractéristiques pertinents des données. La pertinence de l'information sélectionnée est estimée en fonction de sa capacité à influencer la prédiction. Les encodeurs automatiques empilés ont une immense popularité en raison de leur remarquable succès dans divers domaines tels que la classification des images^[29] et le traitement de la parole^[30]. Ce sont des techniques d'extraction de caractéristiques automatisées qui nécessitent peu d'information sur le problème.

Récemment, plusieurs travaux^{[30],[31]} recommandent l'utilisation de réseaux profonds pour apprendre efficacement des fonctionnalités abstraites de manière hiérarchique. Il est démontré que ces architectures profondes fournissent une représentation plus robuste et complète^[31] des données d'entrée.

Dans ce travail, nous utilisons un auto-encodeur empilé pour extraire les informations importantes et reconstruire une représentation d'entrée simple à partir de l'ensemble de données d'origine. L'encodeur empilé commence par entraîner les entrées d'origine, puis nous utilisons le modèle depuis le début jusqu'au goulot d'étranglement (partie code) tout en ignorant sa partie de reconstruction. Cette proportion de code donne une représentation compressée des entrées.

La figure 48 illustre l'architecture de l'auto-encodeur que nous utilisons pour l'extraction de caractéristiques. Il est composé de plusieurs couches neuronales entièrement connectées, séparées par des couches de normalisation par lots. La normalisation par lots stabilise le processus d'apprentissage et diminue considérablement le nombre d'époques de formation nécessaires pour constituer les réseaux de neurones profonds.

Figure 48 : Architecture empilée de l'auto-encodeur proposé



6.3 Processus de formation

Après avoir reconstruit les données d'entrée, nous commençons le processus de formation en appliquant un modèle PMC. Le PMC est souvent recommandé pour les problèmes de classification et de régression. Contrairement aux réseaux de neurones artificiels traditionnels, le PMC peut avoir de nombreuses couches de neurones et est capable d'apprendre des modèles plus complexes. Il est très flexible et peut être utilisé pour apprendre une correspondance entre les entrées et les sorties. Cette flexibilité lui permet d'être appliqué à de nombreux types de données telles que des images, des données de séries chronologiques, des signaux, etc. À cette fin, dans notre travail, nous appliquons un modèle PMC qui utilise le pilote de liaison montante, reçu avec des antennes Omni à plusieurs stations de base pour prédire le meilleur vecteur de formation de faisceau à chacune des stations de base de coordination.

Le PMC est caractérisé par un nombre élevé d'hyperparamètres. Le choix de ces valeurs affecte le comportement des algorithmes en termes de tolérance aux erreurs, de variantes, de nombre d'itérations, etc. Le tableau 8 montre les valeurs optimales d'hyperparamètres que nous adoptons après avoir ajusté plusieurs combinaisons de notre modèle PMC. Nous utilisons cinq couches PMC cachées avec une fonction d'activation d'unité linéaire rectifiée (reLu) et une couche de sortie avec une fonction d'activation linéaire. La fonction reLu est plus puissante que les autres fonctions d'activation car elle évite le problème de la disparition du gradient, ce qui permet au modèle d'apprendre plus rapidement^[32]. Nous utilisons également l'optimiseur Adam car il convient à la plupart des problèmes notamment les problèmes de bruit^[33]. Après avoir entraîné notre modèle avec ces paramètres, nous obtenons les prédictions de faisceau

que nous évaluons à l'aide de l'ensemble de test et de la métrique de performance présentés dans la section suivante.

Tableau 8 : Variables hyperparamètres utilisées pour le modèle PMC

hyperparamètre	valeurs
couches PMC (nombre d'unités) + fonction d'activation	1 couche d'entrée une couche cachée de 200 unités + reLu une couche cachée de 300 unités + reLu une couche cachée de 300 unités + reLu une couche cachée de 300 unités + reLu une couche cachée de 200 unités + reLu 1 couche de sortie de 256 unités + linéaire
taille du lot	128
époques	20
optimiseur	Adam
fonction de perte	erreur quadratique moyenne

7. Résultats expérimentaux

Il est important d'obtenir les réponses les plus précises. Accepter des résultats erronés peut entraîner de lourdes pertes. Dans cette perspective, cette section présente les résultats de notre modèle PMC d'auto-encodeur. Ils sont évalués à l'aide de la mesure de l'erreur quadratique moyenne dans la première sous-section, discutés et comparés dans la deuxième sous-section.

7.1 Indicateur de performances

Pour évaluer les performances d'un modèle, l'écart entre les valeurs prédites et les observations réelles est calculé à l'aide de différentes méthodes statistiques^[34]. Dans notre travail, nous utilisons l'erreur quadratique moyenne (MSE), une mesure de performance très simple et



couramment utilisée. MSE représente la différence ainsi que la distance au carré entre la valeur réelle et la valeur prédite. Cette métrique est utilisée pour éviter l'annulation des termes négatifs, ce qui est l'avantage de MSE. Le MSE est calculé selon la formule 29.

$$MSE = \frac{1}{N} \sum (y - \hat{y})^2 \quad (29)$$

N est le nombre d'échantillons, y est les valeurs réelles et $\hat{y}$ est les valeurs prédites.

7.2 Discussion des résultats

L'un des problèmes ouverts dans la recherche sur la prédiction de la formation de faisceaux est la difficulté à comparer les performances des articles dans la littérature en raison de la diversité des ensembles de données étudiés et des mesures de performance utilisées par les auteurs. Pour cette raison, nous utilisons un PMC simple comme référence et le comparons avec notre modèle proposé.

La figure 49 montre l'historique d'ajustement de la prédiction de formation de faisceau avec PMC uniquement et la figure 50 montre l'historique d'ajustement lors de l'utilisation de PMC et d'auto-encodeurs pour l'extraction de caractéristiques. Nous pouvons noter que la courbe d'ajustement du modèle proposé est très bien ajustée par rapport au modèle PMC. Cela confirme que notre modèle ne connaît ni sur-ajustement ni sous-ajustement.

Figure 49 : Historique du montage du PMC

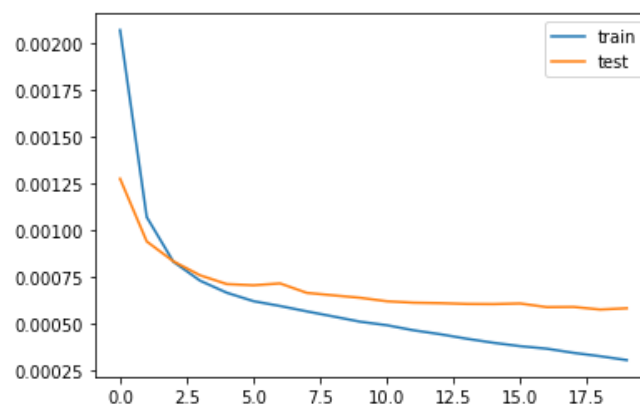
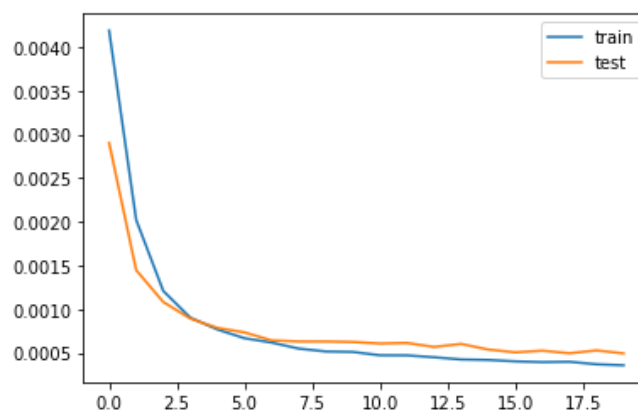
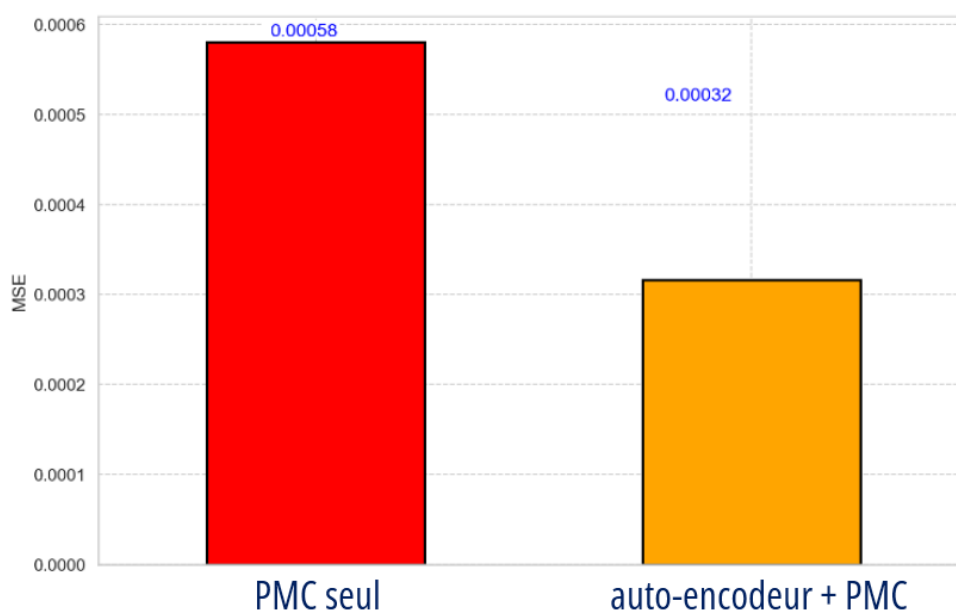


Figure 50 : Historique d'ajustement de l'auto-encodeur empilé (pour l'extraction de caractéristiques + PMC)



La figure 51 montre les résultats quadratiques moyens des faisceaux prédits lors de l'utilisation de PMC uniquement et lors de l'utilisation de l'auto-encodeur pour l'extraction de caractéristiques et de PMC pour la formation. Nous constatons que PMC a une erreur de 0,00058 et notre modèle une erreur de 0,00032. Cela montre l'efficacité de l'auto-encodeur pour l'extraction de caractéristiques car ce nouvel ensemble de caractéristiques résume de nombreuses informations pertinentes contenues dans l'ensemble de données d'origine, ce qui contribue à augmenter la possibilité d'explication du modèle et à en accélérer la formation.

Figure 51 : Erreur quadratique moyenne de notre modèle par rapport au modèle PMC uniquement



Conclusion

Ce chapitre propose une technique d'apprentissage en profondeur pour la prédiction de la formation de faisceaux analogiques à ondes millimétriques. Contrairement aux approches existantes peu pratiques, notre travail est très simple et peu coûteux. Il est basé sur l'application d'un auto-encodeur empilé pour compresser les jeux de données et de perceptrons multicouches afin de former et prédire les meilleurs faisceaux. Notre approche est testée sur un jeu de données disponible et évaluée à l'aide de la métrique d'erreur quadratique moyenne. Les résultats sont prometteurs et surpassent la méthode de référence. Pour de futurs travaux, nous prévoyons d'utiliser différents modèles d'apprentissage en profondeur basés sur l'attention et de les tester dans plusieurs scénarios d'ensembles de données et à différentes fréquences.

Chapitre V $\diamond$ Communication sans fil soutenue par surfaces intelligentes reconfigurables

Introduction

Les ondes électromagnétiques qui transportent les informations dans les communications sans fil interagissent avec les objets et les surfaces sur le chemin qui va de l'émetteur au récepteur. Bien que la superposition de nombreux chemins de propagation donne lieu à des phénomènes aléatoires comme l'évanouissement, chaque chemin de propagation a un comportement constant.

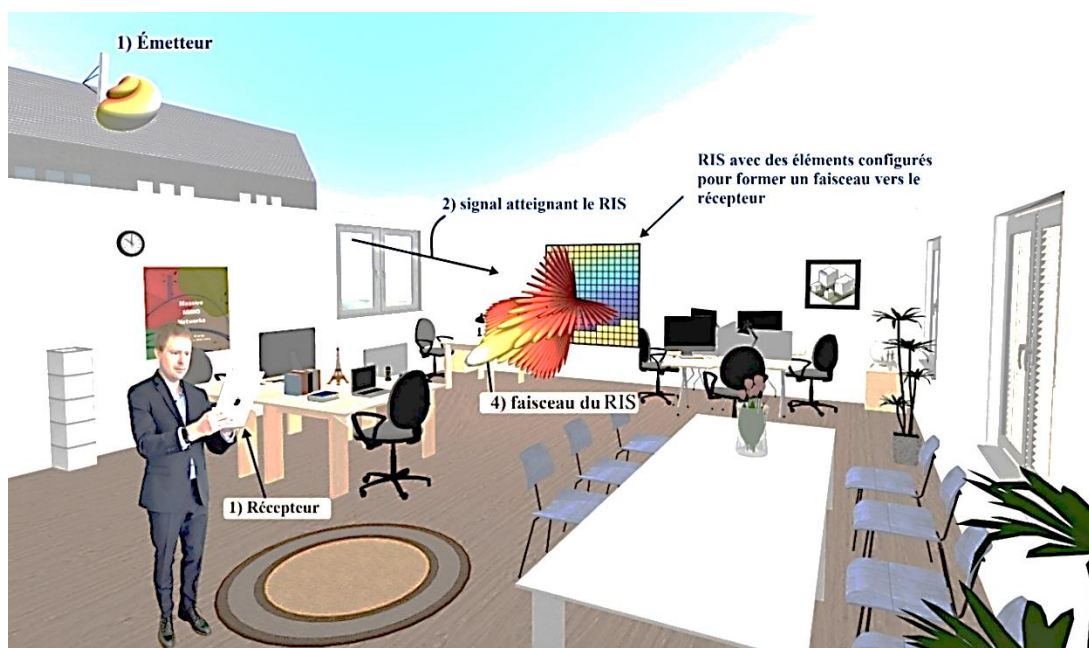
Il existe des matériaux manufacturés destinés à être déployés pour façonner l'environnement de propagation et sur lesquels les ondes sont réfléchies non pas d'une façon standardisée mais reconfigurables. La perspective d'inclure de telles surfaces reconfigurables dites intelligentes (SIR) dans les architectures de réseau au-delà de la 5G suscite beaucoup d'intérêt^{[53], [54]}. Les SIR sont également appelées métasurface contrôlées par logiciel^[55] ou encore surfaces réfléchissantes intelligentes^[56]. Un cas d'utilisation basique de SIR est illustré à la figure 52 où une station de base (BS), érigée sur un toit, transmet à un utilisateur dans un intérieur. Supposons que les matériaux du bâtiment soient tels que le chemin direct à travers le mur subisse des pertes de pénétration massives tandis que le chemin à travers la fenêtre ne subit que des pertes mineures. À l'intérieur de la fenêtre, une SIR est intégrée pour capter l'énergie du signal proportionnellement à sa surface et la réémettre sous la forme d'un faisceau vers le récepteur. Pour que le faisceau se focalise en permanence sur le dispositif de l'utilisateur où qu'il soit dans la pièce, la SIR est reconfigurable. L'utilisation d'une SIR dans cette configuration améliore le rapport signal sur bruit.

Une SIR est une surface mince composée de n éléments, chacun étant un diffuseur reconfigurable : une petite antenne reçoit et réémet sans amplification mais avec un retard configurable^[55]. Pour les signaux à bande étroite, ce retard correspond à un déphasage. Si les déphasages sont correctement ajustés, les n ondes diffusées s'additionnent de manière constructive au niveau du récepteur. Ce principe ressemble à la formation de faisceau traditionnelle : chaque élément a un diagramme de rayonnement fixe, mais la collection de déphasages détermine où se produit l'interférence constructive entre les ondes diffusées. Le



motif en couleur sur la SIR de la figure 52 représente les déphasages nécessaires pour diriger un faisceau vers le récepteur. Chaque élément est sensiblement plus petit que la longueur d'onde (par exemple, un cinquième de la longueur d'onde dans chaque direction^[180]), de sorte qu'il diffuse les signaux presque uniformément, ce qui donne à la surface la capacité de former des faisceaux d'intensité égale dans toutes les directions^[181].

Figure 52 : Cas typique d'utilisation d'une SIR qui reçoit un signal de l'émetteur et le réémet en le focalisant vers le récepteur ; pour focaliser le faisceau dans la bonne direction, la SIR doit être correctement configurée^[182]



L'analyse de la propagation d'une SIR consiste essentiellement à :

- trouver la fonction de Green de la source du signal (une somme d'ondes sphériques si elle est proche ; une onde plane si elle est éloignée) ;
- à calculer le champ incident à chaque élément de la SIR ;
- à intégrer ce champ sur la surface de chaque élément pour trouver la densité de courant ;
- à calculer le champ rayonné par chaque élément ;
- à appliquer le principe de superposition pour trouver le champ au niveau du récepteur.

Comme les éléments sont petits, on peut approcher le champ réémis en prétendant que chaque élément est une source ponctuelle et que le signal reçu est une superposition de signaux de source déphasés et mis à l'échelle en amplitude^[181].

Il existe de nombreux cas d'utilisation potentielle pour les communications sans fil assistées par les SIR. En plus de l'amélioration du rapport signal sur bruit comme dans la figure 52, les SIR peuvent atténuer l'interférence entre les utilisateurs qui sont spatialement multiplexés ou



limiter la fuite du signal en dehors de la zone de couverture prévue et d'atténuer les risques d'écoute clandestine^{[54]–[56]}. La prise en charge du transfert d'énergie sans fil, de la rétrodiffusion et de la modulation spatiale est également concevable ; la plupart des fonctions d'un réseau d'antennes est réalisable par une SIR^[180].

L'idée consiste à prendre en charge la transmission d'une source vers sa destination en adaptant l'environnement de propagation, c'est-à-dire, en configurant la SIR pour qu'elle façonne un faisceau du signal reçu vers la destination. Ce à quoi répondent déjà les relais semi-duplex^[183] mais avec la différence essentielle qu'un relais traite activement le signal reçu pour retransmettre un signal amplifié, alors qu'une SIR répercute passivement le signal, sans amplification, en formant un faisceau. La comparaison entre un relais idéal à amplification et une retransmission (AR) montre d'importants gains d'efficacité énergétique en utilisant une SIR. Cependant, on sait que le relais de décodage et transfert (relais DF) est plus performant que le relais amplification et transfert (AF) en termes de taux réalisables^[184] et constitue donc une meilleure référence.

Dans ce chapitre, nous essayons de faire une comparaison équitable entre la transmission assistée par SIR et le relais DF codé par répétition dans le but de déterminer la taille nécessaire (le nombre d'éléments n) d'une SIR, pour surpasser le relais conventionnel. À cette fin, nous paramétrons les deux technologies en calculant les puissances d'émission optimales et le nombre optimal d'éléments dans une SIR. Plus loin, nous démontrons comment la phase de détection du spectre de la RC peut être améliorée en employant des technologies basées sur les SIR au niveau de l'utilisateur. Nous dérivons des expressions pour les probabilités de fausse alarme et de détection d'un nœud RC, ainsi que pour la probabilité de transmission et le débit. Enfin, nous proposons un modèle pour la gestion dynamique du spectre dans les environnements RC assistés par les SIR.

1. Surfaces intelligentes reconfigurables pour communications sans fil

1.1 Modèle du système étudié

Considérons le système de communication de la figure 53 qui modélise une communication sans visibilité directe (NLOS, pour *non line of sight*) entre une source à antenne unique et une destination à antenne unique (SISO, pour *single input / single output*). Le canal déterministe à évanouissement plat est désigné par $h_{sd} \in \mathbb{C}$. Le signal reçu à la destination est exprimé comme suit :

$$y = h_{sd} \sqrt{P_t} s + n \quad (30)$$

où P_t est la puissance d'émission, s le signal d'information de puissance unitaire et $n \sim \mathcal{NC}(0, \sigma^2)$ représente le bruit reçu. Pour des raisons de simplification de notation, les gains d'antenne sont inclus dans les canaux. L'efficacité spectrale (ES), également connue sous le nom de capacité du canal ou de débit atteignable, est définie par $\log_2(1 + \text{SNR})$, où SNR est le rapport signal / bruit. L'ES de la communication SISO considérée est la suivante :



$$C_{\text{SISO}} = \log_2 \left(1 + \frac{P_t \beta_{\text{sd}}}{\sigma^2} \right) \quad (31)$$

où $\beta_{\text{sd}} = |h_{\text{sd}}|^2$ est le canal déterministe à évanouissement plat de la source à la destination. Le terme $P_t \beta_{\text{sd}}$ représente la puissance reçue à destination P_r , avec

$P_r = P_t \beta_{\text{sd}}$, et le terme $\frac{P_t \beta_{\text{sd}}}{\sigma^2}$ représente le SNR de la communication SISO, avec :

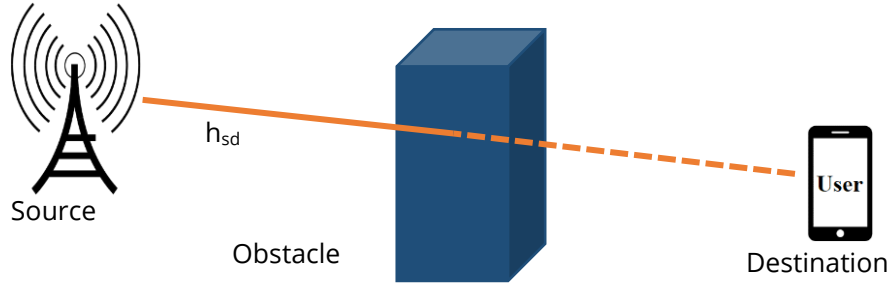
$$\text{SNR}_{\text{SISO}} = \frac{P_t \beta_{\text{sd}}}{\sigma^2} \quad (32)$$

L'affaiblissement de propagation β_{sd} est défini par :

$$\beta_{\text{sd}} = \frac{A}{4\pi d^2} \quad (33)$$

Avec $A = \lambda^2/4\pi$ la surface de l'antenne du récepteur, dans laquelle λ représente la longueur d'onde du signal. Il est évident que $\beta_{\text{sd}} \in [0,1]$ pour la raison que la puissance reçue P_r ne peut pas dépasser la puissance d'émission P_t (loi de conservation de l'énergie).

Figure 53 : Illustration de la communication SISO considérée avec NLOS



À partir de l'expression du débit dans la figure 53, nous pouvons identifier la puissance requise pour atteindre un débit particulier. Pour atteindre un débit spécifique dans le cas SISO, la puissance d'émission requise est :

$$P_{\text{SISO}} = (2^{\bar{R}} - 1) \frac{\sigma^2}{\beta_{\text{sd}}} \quad (34)$$

La puissance totale consommée dans un système de communication se compose de la puissance d'émission de la source et de la puissance dissipée dans le matériel du système. Soient P_s et P_d les puissances dissipées à la source et à la destination, respectivement. Dans le cas SISO, la puissance totale consommée est :

$$P_{\text{total}}^{\text{SISO}} = P_{\text{SISO}}/\nu + P_s + P_d \quad (35)$$



ν représente le rendement de l'amplificateur de puissance.

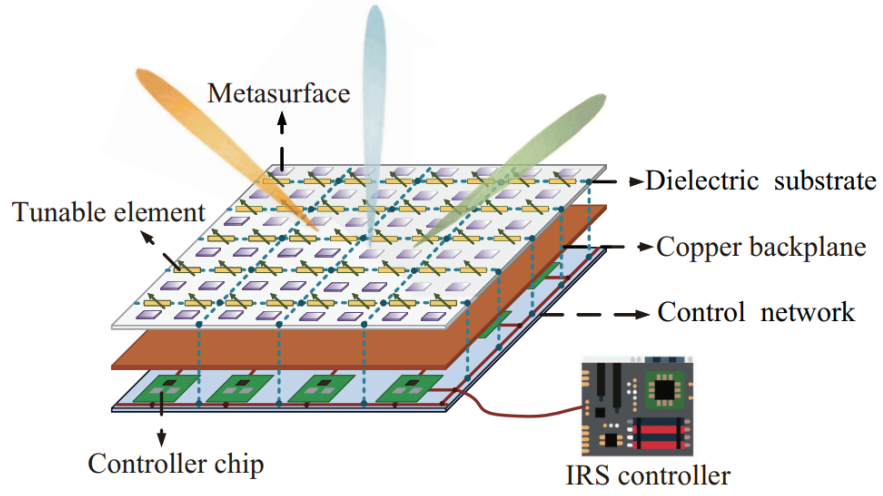
En général, β_{sd} est souvent inférieur à 1. Cependant, l'une des contributions de notre travail est d'obtenir le facteur aussi proche possible de 1 pour augmenter la capacité du canal. Étant donné que nous ne pouvons pas contrôler la distance d dans tous les systèmes de communication, une façon d'augmenter le facteur β_{sd} est d'augmenter la surface A de l'antenne du récepteur. Pour cette raison, nous pouvons employer un équipement avec n antennes avec la même surface A , alors le facteur β_{sd} peut être augmenté n fois à la condition que chaque antenne soit orientée et polarisée pour correspondre au signal reçu. Dans ce cas, la nouvelle capacité du canal s'exprime comme suit :

$$C_N = \log_2 \left(1 + \frac{P_t N \beta_{sd}}{\sigma^2} \right) \quad (36)$$

Par conséquent, en adoptant des équipements supplémentaires dans l'environnement sans fil, la capacité du canal peut être augmentée. La croissance linéaire de la puissance reçue PP est l'idée fondamentale qui est à la base des communications mMIMO dans les récentes générations de transmission sans fil. Dans ce travail, nous étudions comment l'adoption de SIR et relais DF, dans un environnement de radio cognitive, peut augmenter l'efficacité spectrale. Ensuite, nous considérons deux configurations différentes : une SIR, configurée pour orienter passivement le signal incident vers la destination voulue (figure 54) et une configuration de relais mMIMO employée en mode décodage et transfert (DF, pour *decode and forward*) (figure 54). Pour une comparaison équitable SIR-relais DF, les efficacités spectrales sont dérivées et optimisées pour les deux technologies.

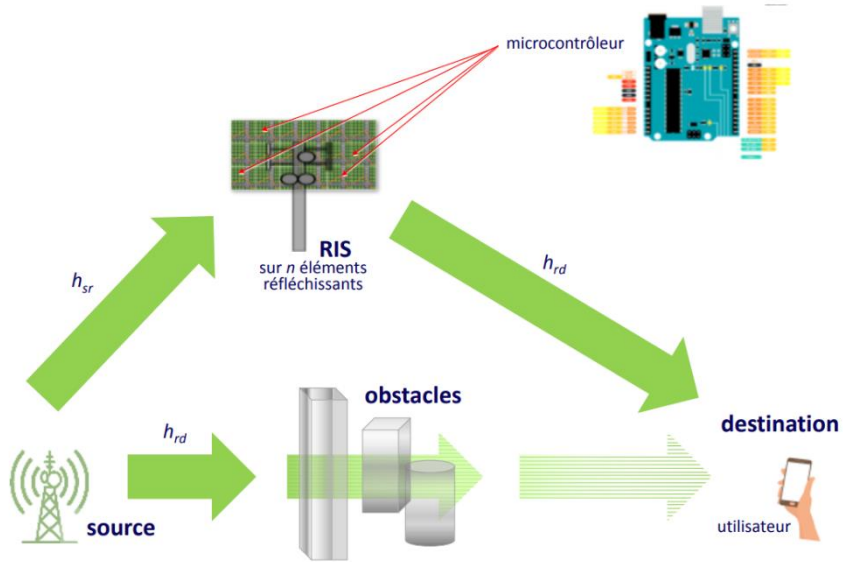
1.2 Transmission supportée par SIR.

Comme le montre la figure 54, la SIR est un réseau d'antennes planaires constitué d'éléments de diffusion imprimés sur trois couches qui ont l'avantage d'être peu coûteux, reconfigurables et quasi passifs. L'élément de diffusion, souvent appelé particule de diffusion ou cellule unitaire, reconfigure intelligemment la phase, l'amplitude, la polarisation et la fréquence des ondes radio incidentes, en temps réel, puis les réfléchit vers le récepteur spécifié avec une amélioration de la qualité du signal et une réduction des interférences.

Figure 54 : Architecture typique de SIR^[185]

Le modèle de système de communication supporté par SIR est similaire à un modèle de relais mMIMO semi-duplex^[186] comme l'illustre la figure 55, avec la différence que les SIR sont conçues pour réfléchir passivement le signal reçu vers sa destination en contrôlant la phase sans augmenter la puissance du signal ni le régénérer.

Figure 55 : Transmission supportée par les SIR



En adoptant les SIR avec n éléments discrets dans le réseau sans fil, l'expression du signal reçu à la destination devient^{[53], [181], [187], [188]} :

$$y = \left(h_{sd} + \mathbf{h}_{sr}^T \Theta \mathbf{h}_{rd} \right) \sqrt{P_t} s + n \quad (37)$$



où h_{sd} , P_t , et n sont les mêmes que dans (31), $[\mathbf{h}_{sr}]_n \in \mathbb{C}^N$ représente le canal déterministe entre l'émetteur et l'élément n^{th} du SIR, et $\mathbf{h}_{rd} \in \mathbb{C}^N$ désigne le canal déterministe du SIR vers la destination. Le facteur Θ est une matrice diagonale qui comprend toutes les propriétés reconfigurables du SIR, il est exprimé par $\Theta = \alpha \text{diag}(e^{j\theta_1}, \dots, e^{j\theta_N})$, dans lequel α représente le coefficient de réflexion d'amplitude fixe et sont les variables de déphasage qui doivent être optimisées par les éléments du SIR pour atteindre la destination. Par conséquent, le rapport signal / bruit de la transmission assistée par SIR est :

$$RSB_{SIR} = \max_{\theta_1, \dots, \theta_N} \frac{P_t |\mathbf{h}_{sd} + \mathbf{h}_{sr}^T \Theta \mathbf{h}_{rd}|^2}{\sigma^2} \quad (38)$$

Sachant que $\mathbf{h}_{sr}^T \Theta \mathbf{h}_{rd} = \alpha \sum_{n=1}^N e^{j\theta_n} [\mathbf{h}_{sr}]_n [\mathbf{h}_{rd}]_n$ alors SNR_{RIS} devient :

$$RSB_{RIS} = \frac{P_t \left(|\mathbf{h}_{sd}| + \alpha \sum_{n=1}^N |[\mathbf{h}_{sr}]_n [\mathbf{h}_{rd}]_n| \right)^2}{\sigma^2} \quad (39)$$

Le SNR maximal est obtenu lorsque les déphasages sont choisis comme $\theta_n = \arg(h_{sd}) - \arg([\mathbf{h}_{sr}]_n [\mathbf{h}_{rd}]_n)$ pour obtenir pour chaque terme $[\mathbf{h}_{sr}]_n$ et $[\mathbf{h}_{rd}]_n$ la même phase que h_{sd} [188]. En intégrant l'affaiblissement de propagation $\beta_{sr} = |h_{sr}|^2$, $\beta_{rd} = |h_{rd}|^2$, et $\beta_{IRS} = \left(\frac{1}{N} \sum_{n=1}^N |[\mathbf{h}_{sr}]_n [\mathbf{h}_{rd}]_n| \right)^2$, nous pouvons exprimer RSB_{RIS} sous la forme suivante :

$$RSB_{RIS} = \frac{P_t \left(\sqrt{\beta_{sd}} + N\alpha\sqrt{\beta_{IRS}} \right)^2}{\sigma^2} \quad (40)$$

Nous pouvons maintenant exprimer la capacité du canal de la transmission assistée par SIR sous la forme suivante :

$$C_{RIS}(N) = \log_2 \left(1 + \frac{P_t \left(\sqrt{\beta_{sd}} + N\alpha\sqrt{\beta_{IRS}} \right)^2}{\sigma^2} \right) \quad (41)$$

Nous pouvons maintenant exprimer la puissance d'émission requise pour atteindre un débit $\bar{R}$ particulier avec SIR comme suit :

$$P_{RIS}(N) = \left(2^{\bar{R}} - 1 \right) \frac{\sigma^2}{\left(\sqrt{\beta_{sd}} + N\alpha\sqrt{\beta_{IRS}} \right)^2} \quad (42)$$



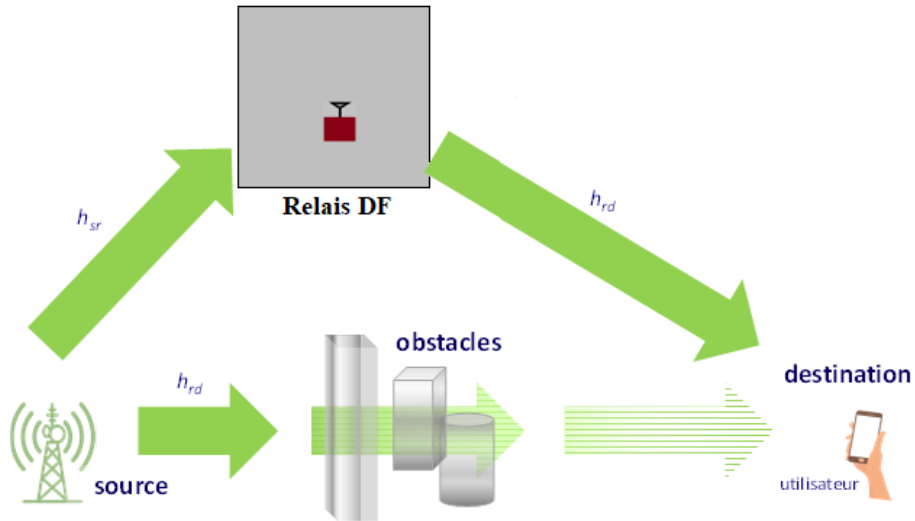
Dans le cas des SIR, la consommation totale d'énergie du système de commutation comprend également la dissipation d'énergie au niveau de chaque élément de SIR, qui est notée P_e [53]. Par conséquent, la consommation totale d'énergie avec SIR est :

$$P_{\text{total}}^{\text{RIS}}(N) = \frac{P_{\text{RIS}}(N)}{N} + P_s + P_d + NP_e \quad (43)$$

1.3 Communication sans fil assistée par relais

Dans la configuration de transmission assistée par relais, nous considérons un relais mMIMO semi-duplex situé à la même position que la SIR, comme le montre la figure 56. La transmission en relais semi-duplex se déroule en deux phases : une transmission de la source vers le relais ; une transmission du relais vers la destination. Aucune liaison directe n'est présente. Parmi les différents protocoles de relais (par exemple [167], [168] et [173] entre autres), nous considérons le protocole de base de décodage et de transmission à répétition codée où un temps égal est alloué aux deux phases.

Figure 56 : Transmission supportée par le relais DF



Dans la première phase, il y a deux transmissions à étudier, la transmission de la source vers la destination, et la transmission de la source vers le relais DF. De ce fait, lorsque la source transmet, le signal capté au niveau du récepteur est le suivant :

$$y_{1d} = h_{sd} \sqrt{P_{t1}} s + n_{1d} \quad (44)$$

h_{sd} , P_{t1} et s sont définis comme dans le cas de communication SISO, et $n_{1d} \sim N_{\mathbb{C}}(0, \sigma^2)$ est le bruit du récepteur. Dans la même phase, le signal reçu au relais est :

$$y_{1r} = h_{sr} \sqrt{P_{t1}} s + n_{1r} \quad (45)$$



h_{sd} désigne le canal entre la source et le relais, tandis que $n_{tr} \sim N_{\square}(0, \sigma^2)$ est le bruit du récepteur. Le relais DF utilise y_{tr} pour décoder l'information, puis la code à nouveau pour la transmettre dans la deuxième phase. Notons que le relais peut être très compact : une antenne, des chaînes émettrices-réceptrices et une unité de bande de base entrent dans les dimensions d'un petit téléphone mobile.

Dans la seconde phase, le relais transmet $\sqrt{P_{t2}}s$ et le signal reçu à destination est :

$$y_{2d} = h_{rd} \sqrt{P_{t2}}s + n_{2d} \quad (46)$$

P_{t2} est la puissance d'émission, h_{rd} désigne le canal entre le relais et la destination, et $n_{2d} \sim N_{\square}(0, \sigma^2)$ est le bruit du récepteur. Selon [183] et en utilisant (44), (45) et (46), pour une combinaison à débit maximal, le débit maximal réalisable à la destination est le suivant :

$$C_{DF} = \frac{1}{2} \min \left\{ \log_2 \left(1 + \frac{P_{t1} |h_{sr}|^2}{\sigma^2} \right), \log_2 \left(1 + \frac{P_{t1} |h_{sd}|^2}{\sigma^2} + \frac{P_{t2} |h_{rd}|^2}{\sigma^2} \right) \right\} \quad (47)$$

L'utilisation du minimum dans (47) est due au fait que le relais commute entre les modes SISO et DF-relais pour transmettre avec une puissance minimale. Pour plus de brièveté, nous utilisons les expressions d'affaiblissement sur le trajet $\beta_{sr} = |h_{sr}|^2$, $\beta_{sd} = |h_{sd}|^2$, $\beta_{rd} = |h_{rd}|^2$. Nous pouvons alors formuler le RSB à la destination pour une transmission supportée par relais DF comme suit :

$$RSB_{relay} = \min \left(\frac{P_{t1} \beta_{sr}}{\sigma^2}, \frac{P_{t1} \beta_{sd}}{\sigma^2} + \frac{P_{t2} \beta_{rd}}{\sigma^2} \right) \quad (48)$$

Par conséquent, nous pouvons reformuler (48) sous la forme plus compacte :

$$C_{DF} = \frac{1}{2} \log_2 \left(1 + \min \left(\frac{P_{t1} \beta_{sr}}{\sigma^2}, \frac{P_{t1} \beta_{sd}}{\sigma^2} + \frac{P_{t2} \beta_{rd}}{\sigma^2} \right) \right) \quad (49)$$

La puissance d'émission nécessaire pour atteindre un débit spécifique $\bar{R}$, dans le cas d'une transmission supportée par un relais DF, dépend des affaiblissements sur le trajet β_{sd} et β_{sr} . Dans le cas où $\beta_{sd} > \beta_{sr}$: l'expression (49) devient $C_{DF} = 1/2 \log_2 (1 + P_{t1} \beta_{sr} / \sigma^2)$. Alors que dans le cas de $\beta_{sd} \leq \beta_{sr}$, elle devient $C_{DF} = 1/2 \log_2 (1 + P_{t1} \beta_{sd} / \sigma^2 + P_{t2} \beta_{rd} / \sigma^2)$.

Pour des raisons de simplicité, l'analyse présentée dans ce travail suppose des canaux déterministes, mais l'extension aux canaux à évanouissement avec une connaissance parfaite du canal est simple : il suffit de prendre les estimations des expressions du débit dans (41) et (49). Par conséquent, toutes les analyses s'appliquent également à ce cas.

Il est évident que $C_{SIR}(N) \geq C_{SISO}$ puisque l'égalité est atteinte pour $N = 0$ et que $C_{SIR}(N)$ est une fonction croissante de N . En fait, le débit augmente comme $O(\log_2(N^2))$ lorsque N est grand, comme noté précédemment dans [188] et expliqué plus en détail dans [190]. La comparaison



entre SIR et relais DF est relativement complexe. Pour qu'elle soit équitable, nous sélectionnons d'abord P_{t1} et P_{t2} de manière optimale, tout en ayant la même puissance moyenne P_t que dans le cas de la SIR.

Supposons que $P_{t1}, P_{t2} \geq 0$ soient sélectionnés sous la contrainte $P_t = \frac{P_{t1} + P_{t2}}{2}$. Si $\beta_{sd} > \beta_{sr}$, il s'avère que $C_{SISO} \geq C_{DR}$ pour toute sélection de P_{t1}, P_{t2} , donc le relais DF est sous-optimal. Si $\beta_{sd} < \beta_{sr}$, le débit avec le relais DF est maximisé par $P_{t1} = \frac{2P_t\beta_{rd}}{\beta_{sr} + \beta_{rd} - \beta_{sd}}$ et $P_{t2} = \frac{2P_t(\beta_{rd} - \beta_{sd})}{\beta_{sr} + \beta_{rd} - \beta_{sd}}$, ce qui conduit à :

$$C_{DF} = \frac{1}{2} \log_2 \left(1 + \frac{2p\beta_{rd}\beta_{sr}}{(\beta_{sr} + \beta_{rd} - \beta_{sd})\sigma^2} \right) \quad (50)$$

Par conséquent, pour atteindre un débit $\bar{D}$, le relais DF a besoin de la puissance d'émission^{[183], [191]}.

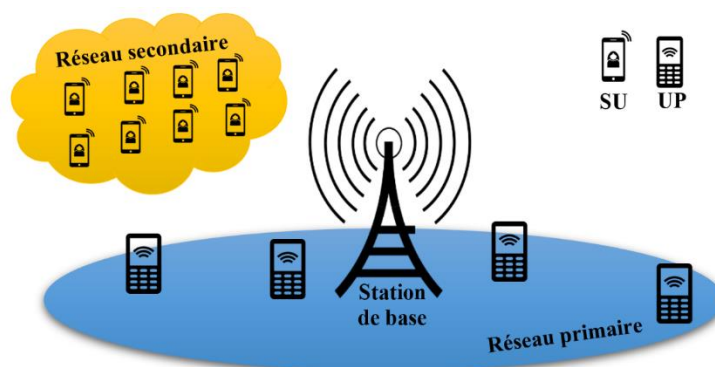
$$P_{DR} = \begin{cases} \left(2^{2\bar{D}} - 1 \right) \frac{\sigma^2}{\beta_{sd}} & \text{if } \beta_{sd} > \beta_{sr} \\ \left(2^{2\bar{D}} - 1 \right) \frac{(\beta_{sr} + \beta_{rd} - \beta_{sd})\sigma^2}{2\beta_{rd}\beta_{sr}} & \text{if } \beta_{sd} \leq \beta_{sr} \end{cases} \quad (51)$$

1.4 SIR pour une gestion dynamique du spectre

1.4.1 Modèle du système proposé

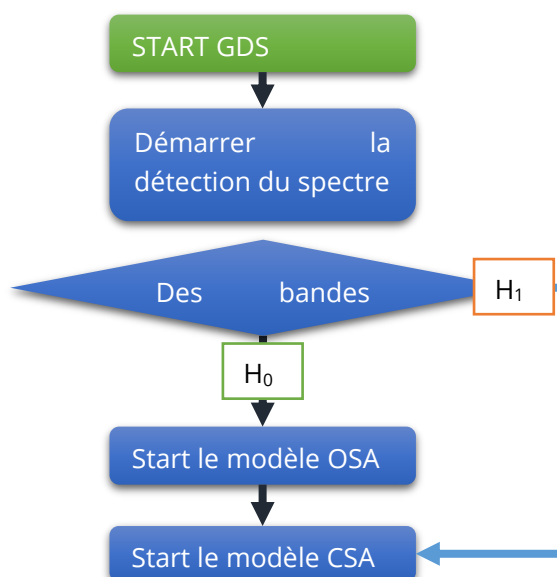
Dans cette section, nous proposons un système hybride unifié OSA-CSA pour une gestion dynamique du spectre souple et très efficace. Le modèle GDS proposé est universel et peut être utilisé sur différentes fréquences et environnements d'exploitation. Le système est basé sur une architecture de réseau RC avec une station de base centralisée (BS) qui peut gérer tous les utilisateurs existant dans sa zone de couverture comme le montre la figure 57. Nous comparons l'allocation du spectre dans trois environnements différents : SISO, SIR, et relais DF. Les UP peuvent coexister avec un ou plusieurs utilisateurs qui peuvent opérer sur la même bande de fréquence en même temps dans une zone déterminée. Les utilisateurs sont administrés par leurs BS correspondantes dans un réseau RC et sont autorisés à accéder de manière opportuniste aux canaux libres ou à utiliser simultanément les canaux occupés sous la contrainte des interférences. En outre, en fonction des interférences générées et des exigences de qualité de service (QoS), les unités de service sont classées en niveaux de priorité faible et élevé.

Figure 57 : Architecture du modèle proposé



Nous utilisons donc un modèle GDS hybride qui intègre les avantages des modèles OSA et CSA. Comme le montre le schéma de la figure 58, en fonction de la disponibilité du spectre, l'une des deux méthodes peut être utilisée par les utilisateurs pour obtenir une meilleure exploitation du spectre et une meilleure efficacité énergétique. Le choix du mode d'accès dépend du résultat de la détection qui déclare l'état des canaux (occupé ou inoccupé). Cependant, lorsque le canal est inoccupé, les utilisateurs peuvent communiquer librement avec une puissance d'émission élevée selon le modèle OSA, tandis que lorsque le canal est occupé ou lorsqu'un utilisateur arrive, ils peuvent transmettre en adaptant leur puissance d'émission selon le modèle CSA. En outre, la sélection du mode d'accès adopté dépend du débit réalisable qu'un utilisateur veut envoyer. Cependant, s'il a besoin de transmettre des informations à faible débit, il choisira directement le modèle CSA sans avoir besoin d'une opération de détection de spectre.

Figure 58 : Schéma d'allocation hybride du modèle proposé



Dans le modèle proposé, on suppose que les utilisateurs existent dans une zone de réseau RC supportée par SIR / relais DF. Les US sont en compétition pour accéder à K canaux de spectre



disponibles dans leur zone, où $K < N$. Chaque utilisateur commence par scanner et analyser son environnement radio et communique avec BS pour déterminer la disponibilité des canaux et le mode d'accès. Pour améliorer la capacité et l'efficacité énergétique des canaux, l'algorithme d'allocation du spectre proposé est appliqué à travers deux technologies : une surface intelligence reconfigurable et un relais DF classique. Dans les deux technologies, nous procédons aux six étapes suivantes :

- Étape 1 : Déterminer le nombre d'utilisateurs qui demandent à communiquer.
- Étape 2 : Lancer l'opération de détection du spectre pour identifier l'état d'occupation ou de non-occupation du spectre où deux scénarios sont envisagés :
 - 1) un utilisateur individuel effectue une détection du spectre en continu pour observer la présence du spectre primaire et identifier les trous dans le spectre ;
 - 2) plusieurs utilisateurs coopèrent pour effectuer la détection du spectre afin d'augmenter la précision de la détection.
- Étape 3 : Estimer l'interférence générée par chaque utilisateur pour chaque canal. Évaluer ensuite l'efficacité spectrale des utilisateurs pour chaque canal.
- Étape 4 : Lorsque l'état du spectre est obtenu, le système commence l'accès dynamique au spectre avec le modèle OSA en attribuant la bande inactive aux utilisateurs qui ont peu de chance d'utiliser le modèle CSA. La priorité d'accès au canal est mise en œuvre sous la forme d'un algorithme de jeu.
- Étape 5 : Le premier utilisateur bénéficiaire met à jour la disponibilité du canal et commence à transmettre sur le canal sélectionné. Ensuite, les autres utilisateurs mettent à jour les informations sur l'état du canal pour une nouvelle allocation.
- Étape 6 : Lorsque tous les canaux libres sont occupés avec le modèle OSA, le système active le modèle CSA. Les autres utilisateurs accèdent simultanément au spectre de fréquences en estimant leur génération de d'interférences et en limitant celle-ci à un niveau acceptable pour préserver la qualité de service. Pour respecter la contrainte de puissance d'interférence, nous utilisons l'algorithme de la théorie des jeux pour contrôler l'interférence générée par les utilisateurs. Dans le modèle CSA, la fonction d'efficacité est utilisée comme fonction objective pour allouer le canal aux autres utilisateurs, toutefois, les utilisateurs ayant des valeurs de l'efficacité spectrale élevées sont sélectionnés comme utilisateurs approuvés pour l'accès simultané au spectre.

1.4.2 Modèle de détection du spectre

La communication sans fil soutenue par SIR présente nombre d'avantages, mais les réseaux de radio cognitive soutenus par SIR font l'objet de peu de recherches. À cet effet et afin de distinguer efficacement le signal utilisateur des signaux de bruit, nous étudions comment l'utilisation des SIR / relais DF dans les environnements RC peut améliorer les performances de détection du spectre. À cette fin, en utilisant les configurations proposées pour les SIR et les



relais DF, nous dérivons des expressions pour P_d et P_{fa} , puis nous analysons la performance du réseau RC avec différentes configurations SIR. Dans la transmission SISO, le signal reçu par un utilisateur RC est formulé comme suit :

$$y = \begin{cases} h_{sd} \sqrt{P_t} s + n & H_1 \\ n & H_0 \end{cases} \quad (52)$$

où y est le signal capté au niveau des canaux utilisateur et examiné à chaque instant, P_t est la puissance d'émission de l'utilisateur et n est le bruit reçu. À l'hypothèse H_1 , le terme $h_{sd} \sqrt{P_t} s + n$ représente le signal que le système veut détecter, où s représente le signal de l'utilisateur. Pour la méthode de détection de l'énergie, P_d et P_{fa} sont généralement définis comme suit :

$$P_d = P_r(y_0 > y_{th} | H_1) \quad (53)$$

$$P_{fa} = P_r(y_0 > y_{th} | H_0) \quad (54)$$

y_0 représente l'énergie du signal reçu et y_{th} le seuil de détection adopté pour protéger les services utilisateur de toute interférence nuisible. Pour un canal BBAG, nous pouvons démontrer que P_d et P_{fa} peuvent être écrits comme suit :

$$P_d = Q \left(\frac{y_{th} - (\sigma^2 + P_r)}{\sqrt{2/N_d} (\sigma^2 + P_r)} \right) \quad (55)$$

$$P_{fa} = Q \left[\frac{y_{th} - \sigma^2}{\sqrt{2/N_d} \sigma^2} \right] \quad (56)$$

N_d est le nombre de données collectées, P_r est la puissance du signal reçu, et σ^2 est la puissance du bruit. Tandis que $Q(\cdot)$ est la fonction de Marcum généralisée, définie par

$Q(u) = \frac{1}{\sqrt{2\pi}} \int_u^\infty e^{-\frac{v^2}{2}} dv$ (existe sur Matlab comme qfunc()). Le seuil de détection y_{th} peut être obtenu à partir de (56) :

$$y_{th} = \sigma^2 \left[Q^{-1}(P_{fa}) \sqrt{2/N_d} + 1 \right] \quad (57)$$

En introduisant (57) dans (58), nous pouvons réécrire la probabilité de détection P_d sous la forme :

$$P_d = Q \left[\frac{\sigma^2 \left[Q^{-1}(P_{fa}) \sqrt{2/N_d} + 1 \right] - (\sigma^2 + P_r)}{\sqrt{2/N_d} \cdot (\sigma^2 + P_r)} \right] \quad (58)$$

Ou sous une forme plus compacte :

$$P_d = Q \left[\frac{Q^{-1}(P_{fa}) \sqrt{2/N_d} - P_r^2 / \sigma^2}{\sqrt{2/N_d} (1 + P_r^2 / \sigma^2)} \right] \quad (59)$$

Sachant que $\gamma = \frac{P_r^2}{\sigma^2}$, avec $SNR(dB) = 10 \log_{10}(\gamma)$, alors P_d devient :

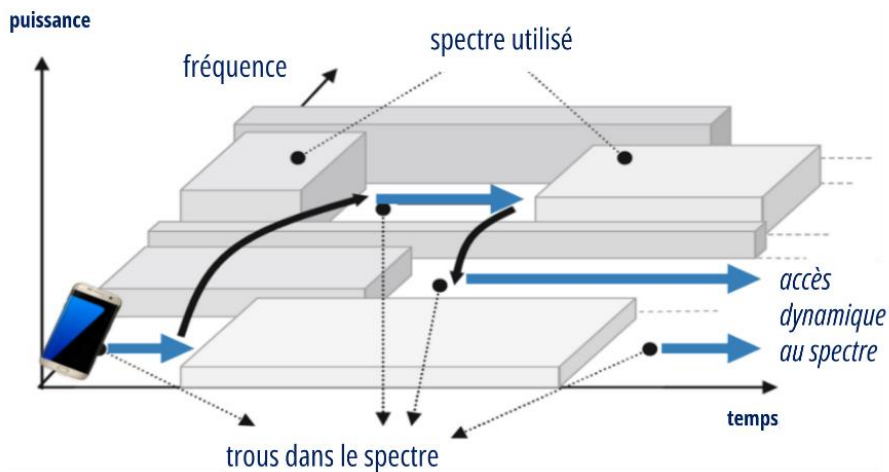
$$P_d = Q \left[\frac{Q^{-1}(P_{fa}) \sqrt{2/N_d} - \gamma}{\sqrt{2/N_d} (1 + \gamma)} \right] \quad (60)$$

La performance du détecteur proposé est étudiée à l'aide de la courbe ROC. Dans les sections suivantes, nous traçons les courbes ROC pour les trois configurations proposées (SISO, SIR et relais DR) et nous comparons leur capacité de détection du signal dans les environnements RC.

2. Modèle hybride pour une gestion dynamique du spectre

Le paradigme de l'accès opportuniste au spectre (OSA) est rationnellement le modèle le plus attrayant pour les utilisateurs afin d'accéder de manière opportuniste au canal sous licence déclaré temporairement inutilisé. Par conséquent, une fois que le modèle de détection du spectre déclare la disponibilité d'un ou de plusieurs trous de spectre, le modèle hybride DSA active l'allocation OSA. Nous proposons un modèle OSA pour permettre une utilisation non licenciée du spectre tout en assurant la protection de la qualité de service des utilisateurs sous licence. La communication des utilisateurs sous licence est généralement intermittente, il y a des bandes de fréquences, des créneaux temporels ou des zones spatiales où l'utilisateur est inactif. Par conséquent, chaque utilisateur est en concurrence avec les autres pour accéder temporairement au spectre inactif sans générer d'interférences nuisibles en contrôlant leur bande passante, leur puissance d'émission, leur schéma de modulation et leur fréquence porteuse. Notons que les utilisateurs bénéficient des capacités de la radio cognitive qui lui permettent d'ajuster intelligemment ses paramètres.

Figure 59 : Utilisation du spectre dans le modèle OSA^[48]





Lorsque la détection du spectre est effectuée, les utilisateurs ajustent leurs paramètres et se démènent pour accéder aux bandes de fréquences disponibles. Cependant, le problème de l'accès au spectre dans l'environnement RC est similaire au processus de sélection des stratégies dans un modèle de jeu. Par conséquent, le problème de la priorité d'accès au canal peut être formé comme un modèle de jeu. Ce modèle de jeu peut organiser l'accès au spectre en atteignant une allocation de canal fiable connue sous le nom de point d'équilibre de Nash. Dans un modèle de jeu, trois objets doivent être définis : les joueurs ou participants, les ensembles de stratégies $\{S_i\}_{i \in N_c}$, et la fonction d'utilité $\{U_i\}_{i \in N_c}$. Dans ce travail, nous utilisons les utilisateurs comme des joueurs, la sélection du canal utilisé pour la transmission comme stratégies de jeu, et les efficacités spectrales développées ci-dessus comme fonction d'utilité. La modélisation mathématique du modèle d'accès au spectre proposé est la suivante :

$$G = \{N_c, \{S_i\}_{i \in N_c}, \{U_i\}_{i \in N_c}\} \quad (61)$$

Ainsi, pour chaque participant RC i dans le jeu G , la fonction d'utilité U_i est une fonction de la stratégie S_i choisie par le participant i et des stratégies adoptées par les autres participants. Dans l'algorithme que nous proposons, la priorité d'accès au canal détecté inactif, dans le modèle OSA, est donnée aux utilisateurs qui possèdent une faible fonction d'utilité. En effet, ces derniers ont peu de chances d'accéder au spectre dans le modèle CSA en raison des interférences nuisibles qu'ils peuvent générer. Ainsi, les utilisateurs avec une fonction d'efficacité élevée peuvent accéder aux canaux disponibles de manière opportuniste avec une priorité élevée.

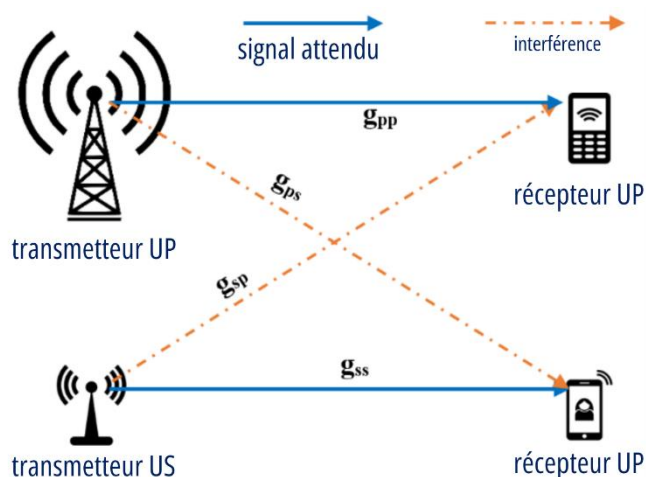
Le modèle OSA permet aux utilisateurs d'utiliser de manière opportuniste tous les canaux inoccupés. Comme les canaux peuvent supporter plus d'un utilisateur en même temps, à condition que les exigences de qualité de service soient assurées, le modèle hybride proposé active le modèle CSA. Le modèle CSA permet aux autres utilisateurs qui ne bénéficient pas du modèle OSA de coexister avec d'autres sur le même canal tant que la contrainte de non-interférence est respectée. Par rapport au modèle OSA, le modèle CSA est devenu récemment le mode d'accès privilégié pour l'utilisation sans licence du spectre. Les principales raisons en trois avantages : le CSA peut permettre à plusieurs utilisateurs de communiquer simultanément sur le même spectre sous licence à condition que l'interférence générée soit contrôlée ; par ailleurs, il n'est pas nécessaire de détecter les espaces blancs et par conséquent, ni la base de données de géolocalisation ni la détection du spectre ne sont nécessaires. Enfin, le modèle CSA atteint une grande efficacité d'allocation spectrale grâce à sa capacité de réutilisation du spectre^[192]. Ainsi, les utilisateurs transmettent librement en continu, que le canal soit utilisé ou non. Le modèle CSA permet d'élaborer des dispositifs cognitifs à faible coût et une faible consommation d'énergie, leur distribution est facile dans les environnement RC, ce qui compte beaucoup pour les dispositifs IoT.

Pour permettre l'attribution du spectre par le modèle CSA, l'émetteur doit contrôler la puissance d'interférence causée au récepteur. Ce contrôle des interférences s'effectue en planifiant les stratégies de communication telles que le débit binaire, la puissance d'émission, la taille des paquets et la largeur de bande, en fonction de l'état du spectre. Afin de réaliser la conception optimale de puissance d'émission, en maximisant la capacité du canal secondaire



sous différents types de contraintes de puissance, nous considérons le système CSA le plus simple et le plus basique composé d'un couple d'utilisateurs comme le montre la figure 60. Nous supposons que chaque dispositif sans fil est équipé d'une seule antenne. Les utilisateurs partagent une seule gamme de fréquences étroite pour la communication. Comme le montre la figure 60, g_{pp} , g_{ps} , g_{sp} et g_{ss} indiquent les gains de puissance instantanés du canal de l'émetteur au récepteur et respectivement.

Figure 60 : Illustration du modèle de partage du spectre de la CSA

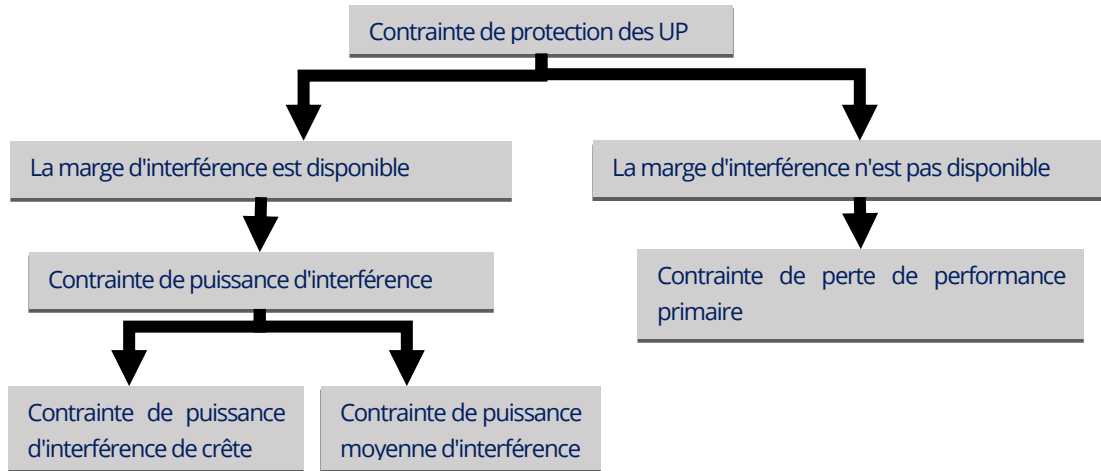


En général, la contrainte de protection des services dépend de la spécification de la marge de brouillage, et elle est donnée selon deux modes :

- lorsque la marge est une valeur prédéterminée, la protection des utilisateurs est modélisée comme une contrainte de puissance de brouillage. Il existe deux types de contraintes de puissance de brouillage, les contraintes de puissance de brouillage de crête et les contraintes de puissance de brouillage moyenne^[193]. Dans la contrainte de puissance d'interférence moyenne, une moyenne de la puissance d'interférence est déterminée sur tous les états du canal. Cette contrainte est la moins rigoureuse car elle peut autoriser le dépassement de la marge de brouillage par la puissance de brouillage pour plusieurs états du canal. En revanche, la contrainte de puissance d'interférence maximale limite la puissance d'interférence à des normes de puissance élevées pour tous les états du canal. Par rapport à la contrainte d'interférence moyenne, la contrainte d'interférence de crête est rigoureuse et peut garantir une protection complète des services, tandis que la contrainte de puissance d'interférence moyenne fournit un débit plus élevé pour les utilisateurs par rapport à la contrainte de puissance d'interférence de crête ;
- lorsque la marge d'interférence exacte est indéterminée, une contrainte de perte de performance primaire est adoptée pour protéger les services UP^{[194], [195]}. La figure 61 présente un schéma qui résume les principales contraintes pour protéger les utilisateurs. Comme indiqué précédemment, la communication simultanée de l'émetteur n'est autorisée que lorsque le brouillage généré par US à UP ne dépasse pas la marge de brouillage. Par conséquent, le contrôle de l'interférence est crucial dans le modèle CSA pour protéger les services.



Figure 61 : Principales contraintes de puissance pour la protection des services UP



Contraintes de puissance d'interférence : dans les environnements RC, les utilisateurs doivent réguler leur puissance d'émission pour protéger les services UP. Les contraintes de protection primaire et la puissance d'émission sont les deux principales classes de contraintes de puissance. Dans la première contrainte, selon le bilan de puissance de l'émetteur, cette contrainte est une restriction de ressource physique qui limite la puissance de chaque utilisateur. Nous supposons que $\nu = (g_{pp}, g_{ps}, g_{sp}, g_{ss})$ et que $P(\nu)$ sont les puissances d'émission US sous ν . Étant donné que P_{av} est la moyenne de la puissance d'émission US et que P_{pk} est la puissance de crête maximale, la contrainte de puissance d'émission peut être formulée comme suit :

$$P(\nu) \geq 0, \quad \forall \nu \quad (62)$$

$$P(\nu) \leq P_{pk}, \quad \forall \nu \quad (63)$$

$$E(P(\nu)) \leq P_{av}, \quad \forall \nu \quad (64)$$

La contrainte de puissance d'émission maximale est représentée par l'équation (63), elle est appliquée pour tenir compte de la non-linéarité des amplificateurs de puissance US. La contrainte de puissance d'émission moyenne est représentée par l'équation (64), elle caractérise le fait que l'utilisation de la puissance US doit être raisonnable sur une longue distance.

Contrainte de protection des utilisateurs : dans cette contrainte, ils ne peuvent transmettre que lorsque les services sont complètement protégés. Par conséquent, la conception RC n'est limitée que par la contrainte de ressources physiques, ce qui la différencie des autres conceptions existantes. En général, il existe deux types de contraintes de protection. La première est appelée contrainte de puissance d'interférence, dans laquelle l'émetteur peut connaître la marge d'interférence de crête Q_{pk} et la marge d'interférence moyenne Q_{av} , la



représentation de la contrainte de protection est similaire à la contrainte de puissance d'interférence :

$$g_{sp}P(v) \leq Q_{pk} \quad , \quad \forall v \quad (65)$$

$$E[g_{sp}P(v)] \leq Q_{av} \quad , \quad \forall v \quad (66)$$

La contrainte de la puissance de brouillage de crête est décrite par la formule (65). Cette contrainte peut permettre la protection complète des services dans n'importe quelle situation d'évanouissement et la protection des services sensibles aux retards. La contrainte de la puissance de brouillage moyenne est décrite par la formule (66). Cette contrainte n'est pas appropriée pour protéger complètement les services, car dans certains environnements d'évanouissement, la puissance de brouillage peut dépasser la marge de brouillage. Par conséquent, elle est appropriée pour protéger les services insensibles au retard.

D'autre part, le deuxième type de contrainte de protection UP est appelé contrainte de perte de performance primaire et dans laquelle les marges Q_{pk} ou Q_{av} ne sont pas accessibles. La contrainte de protection des services peut être exprimée comme suit :

$$\varepsilon_p \leq \varepsilon_0 \quad (67)$$

$$\Delta r_p \leq \delta_0 \quad , \quad \forall v \quad (68)$$

L'équation (67) représente la contrainte d'interruption des utilisateurs, où ε_0 représente la probabilité d'interruption de la transmission qui doit être conservée, tandis que ε_p représente la probabilité d'interruption de la communication en cas de transmission simultanée. En considérant γ_p comme le rapport signal / bruit parasite, nous pouvons formuler ε_p comme suit :

$$\varepsilon_p = \Pr \left\{ \frac{g_{pp}P_p}{g_{sp}P(v) + N_0} < \gamma_p \right\} \quad (69)$$

Alors que l'équation (68) représente la contrainte de la perte de débit primaire^[194], où δ_0 désigne la perte de débit la plus élevée que peut atteindre l'utilisateur. Ensuite, nous supposons que dans les deux équations (67) et (68), l'émetteur peut connaître certaines informations de PU, y compris P_p et g_{pp} . Par conséquent, en combinant la contrainte de puissance d'émission et la contrainte de protection, nous pouvons exprimer les contraintes de puissance CSA comme les combinaisons de contraintes suivantes :



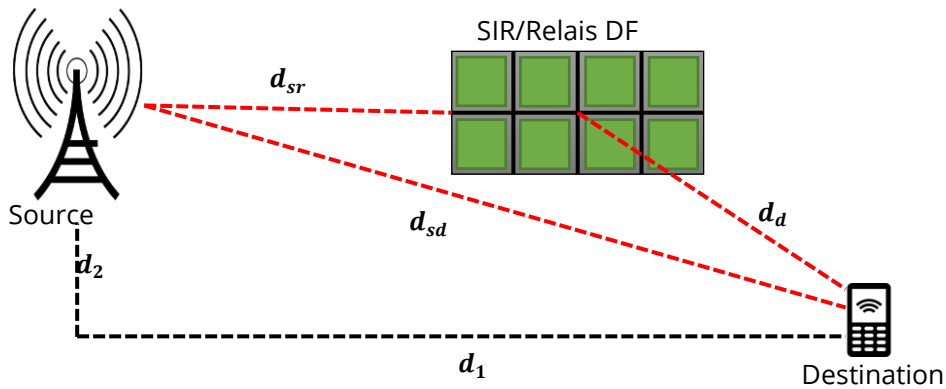
$$\begin{aligned}
 C_1 &= \{P(\nu), (4.33), (4.34), (4.36)\} & C_5 &= \{P(\nu), (4.33), (4.34), (4.38)\}, \\
 C_2 &= \{P(\nu), (4.33), (4.34), (4.37)\} & C_6 &= \{P(\nu), (4.33), (4.35), (4.38)\}, \\
 C_3 &= \{P(\nu), (4.33), (4.35), (4.36)\} & C_7 &= \{P(\nu), (4.33), (4.34), (4.39)\}, \\
 C_4 &= \{P(\nu), (4.33), (4.35), (4.37)\} & C_8 &= \{P(\nu), (4.33), (4.35), (4.39)\} \quad (70)
 \end{aligned}$$

3. Résultats numériques

Dans cette section, les résultats des comparaisons numériques sont présentés pour tous les systèmes proposés. Dans un premier temps, nous comparons la transmission supportée par les SIR et le relais DF, dans le but de déterminer comment le déploiement de la technologie SIR peut améliorer l'efficacité énergétique et la qualité de la communication. Ensuite, nous étudions comment les SIR peuvent soutenir les radios cognitives pour améliorer l'efficacité du spectre en augmentant les performances de détection du spectre et en améliorant l'efficacité énergétique des techniques d'accès dynamique au spectre. À cet effet, nous considérons la configuration de simulation de la figure 62, où la source est positionnée à un endroit fixe, tandis que les emplacements de la destination et du SIR / relais DF peuvent être variables.

Soit d_{sd} , d_{sr} et d_{rd} les distances entre la source et la destination, la source et SIR / relais DF, et SIR / relais DF et la destination, respectivement. Nous supposons que le récepteur peut être un téléphone mobile équipé d'une antenne omnidirectionnelle avec un gain $G_r = 0\text{dBi}$ et une dissipation de puissance matérielle $P_r = 100\text{mW}$, tandis que l'émetteur et la SIR/relais DF sont équipés d'antennes de taille égale avec un gain $G_t = 5\text{dBi}$ et des dissipations de puissance matérielle $P_s = P_d = P_r = 100\text{mW}$ et $P_e = 5\text{mW}$.

Figure 62 : Modèle de simulation proposé



3.1 Comparaison numérique des performances

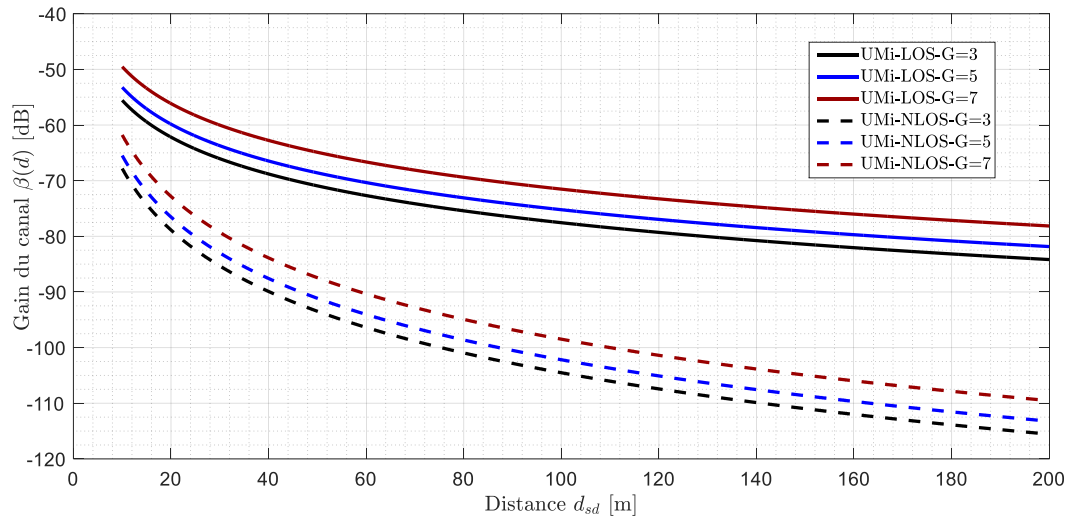
Nous comparons maintenant numériquement les systèmes. Les gains du canal sont modélisés à l'aide du 3GPP Urban Micro (UMi) de [196] avec une fréquence porteuse de 3 GHz. Nous utilisons les versions en visibilité directe (LOS) et non-LoS (NLOS) de l'UMi, qui sont définies pour des distances ≥ 10 m. Nous considérons que G_t et G_r désignent les gains d'antenne (en dBi)



respectivement à l'émetteur et au récepteur. Nous négligeons l'évanouissement d'ombre (*shadow fading*) pour obtenir un modèle déterministe et montrons le gain du canal β en fonction de la distance d_{sd} dans la figure 63 :

$$\beta(d)[\text{dB}] = G_t + G_r + \begin{cases} -37.5 - 22\log_{10}(d/1\text{m}) & \text{if LOS} \\ -35.1 - 36.7\log_{10}(d/1\text{m}) & \text{if NLOS} \end{cases} \quad (71)$$

Figure 63 : Gains typiques du canal en fonction de la distance incluant plusieurs gains d'antenne



La figure 63 montre comment les gains de canal typiques peuvent être affectés par les gains d'antenne à l'émetteur et au récepteur. Cette figure montre qu'un chiffre apparemment faible comme -60 dB, représente en fait un gain de canal très important, alors que les chiffres habituels se situent entre -70 et -110 dB.

Les résultats de la comparaison numérique sont présentés ci-après. Les figures 64 et 65 montrent la puissance d'émission nécessaire pour obtenir respectivement un débit de $\bar{R} = 4 \text{ bits/s/Hz}$ et $\bar{R} = 8 \text{ bits/s/Hz}$. Les deux figures présentent les résultats obtenus pour SISO, relais DF et SIR avec les configurations suivantes : la distance entre la source et la destination est fixée à l'emplacement avec $d_1 = 100\text{m}$ et $d_2 = 15\text{m}$ où $d_{sd} = \sqrt{d_1^2 + d_2^2}$, la largeur de bande est $B = 10\text{MHz}$, le coefficient de réflexion d'amplitude fixe est $\alpha = 1$, la puissance de bruit est -94 dBm et SIR avec trois tailles possibles $N = \{64, 256, 1024\}$.

Figure 64 : Puissance d'émission nécessaire pour atteindre le débit $\bar{R} = 4$ bit/s/Hz en fonction de la distance d_{sr}

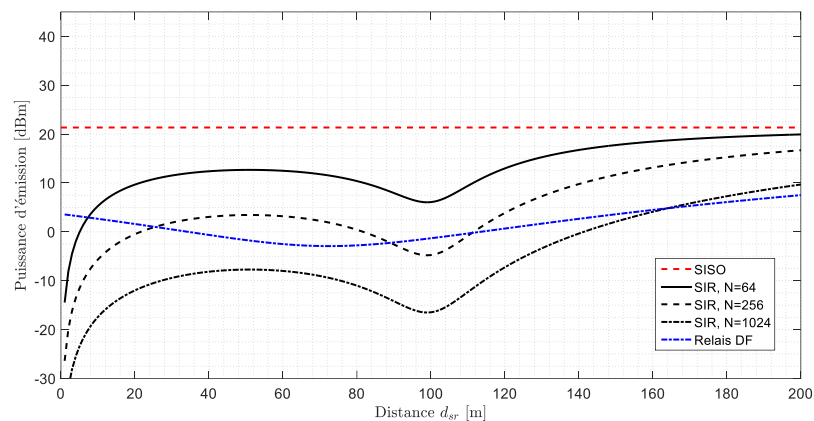


Figure 65 : Puissance d'émission nécessaire pour atteindre le débit $\bar{R} = 8$ bit/s/Hz en fonction de la distance d_{sr}

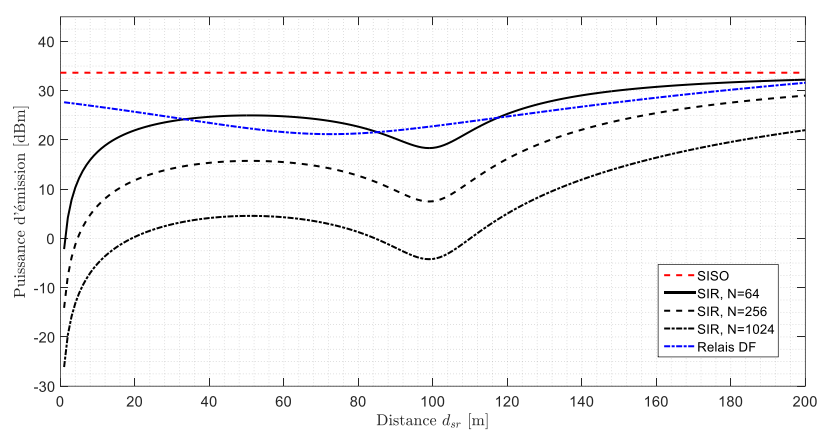
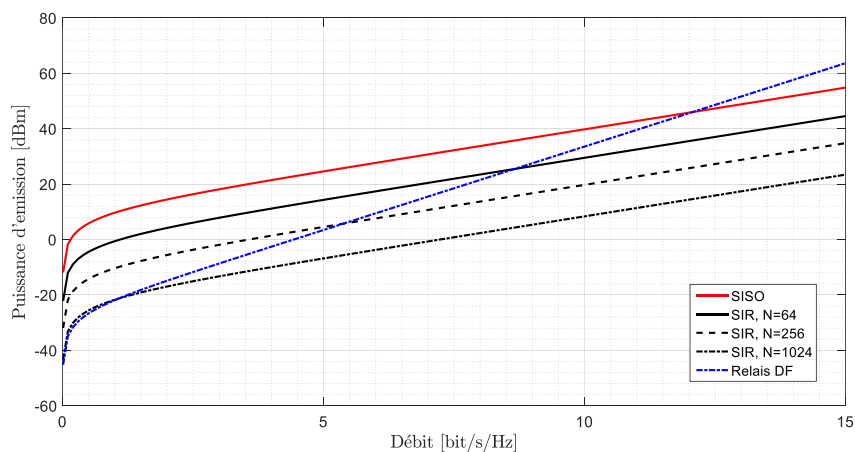


Figure 66 : Puissance d'émission requise en fonction du taux réalisable $\bar{R}$ lorsque $d_1 = 100$ m



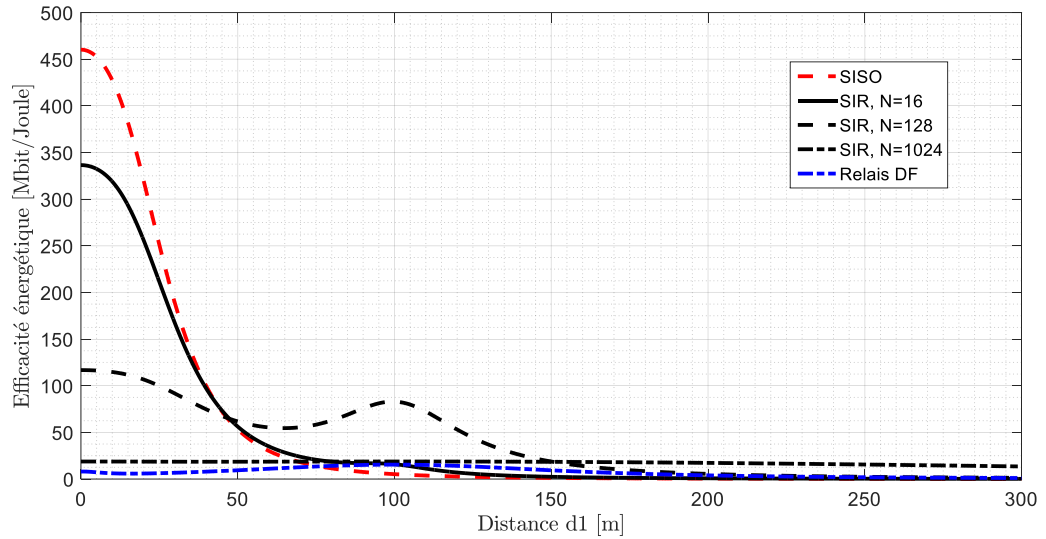


Nous étudions maintenant comment les SIR améliorent l'efficacité énergétique dans les transmissions sans fil, en comparaison avec les systèmes SISO et Relais DF. L'efficacité énergétique (EE) est définie comme suit :

$$EE = B.\bar{R}/P_{total} \quad (72)$$

Nous considérons la configuration de simulation de la figure 66 où nous supposons que l'emplacement de la destination est variable, tandis que la source et la SIR sont maintenant positionnés à des emplacements fixes avec $d_{sr}=100m$ et $\bar{R}=10 \text{ bits/s/Hz}$. La figure 67 montre l'efficacité énergétique en fonction de la distance d_1 pour les trois configurations lorsque N prend différentes valeurs. L'efficacité énergétique est maximale lorsque la distance d_1 est d'environ 100 m, ce qui confirme les résultats précédents. SISO donne une efficacité énergétique élevée pour les courtes distances d_1 , alors que sa performance diminue rapidement lorsque d_1 augmente. Le relais DF fournit une efficacité énergétique plus faible en raison de la valeur élevée du débit supposé. SIR fournit des valeurs d'efficacité énergétique élevées lorsque la destination se rapproche de la source ou de la SIR, mais l'effet du nombre d'éléments N du SIR n'est pas très clair. Par conséquent, nous pouvons optimiser N pour chaque distance d_1 afin d'obtenir la meilleure efficacité énergétique possible avec les SIR.

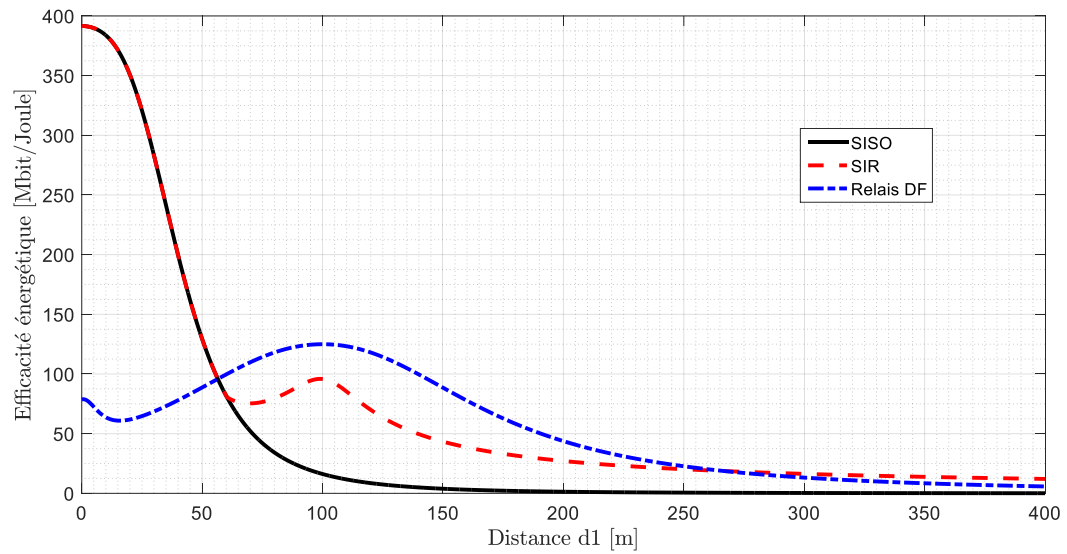
Figure 67 : Efficacité énergétique en fonction de la distance d_1



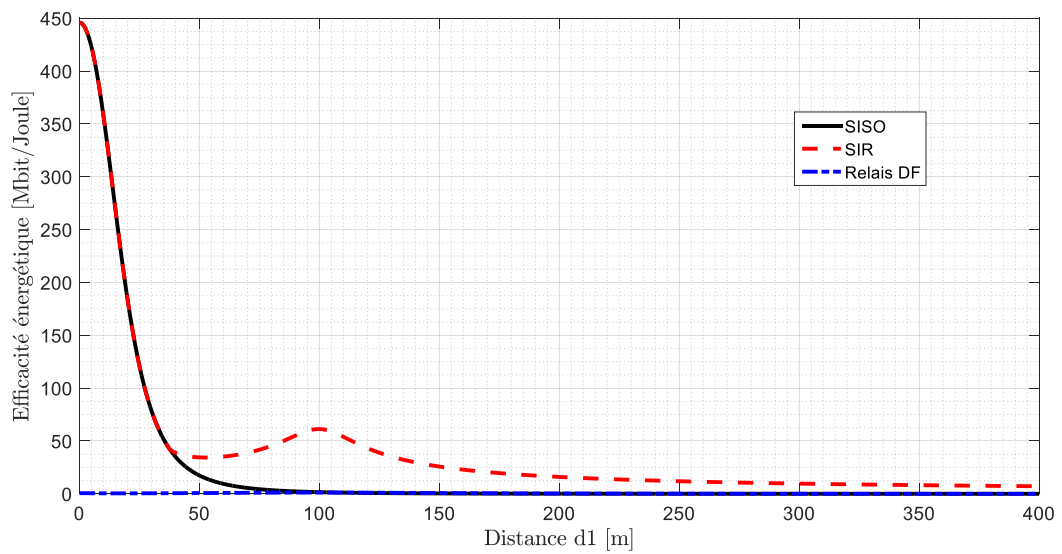
La figure 68 montre l'efficacité énergétique obtenue avec un taux de $\bar{R}=8 \text{ bits/s/Hz}$ et $\bar{R}=12 \text{ bits/s/Hz}$ lorsque le nombre d'élément N est optimisé. Dans le cas du relais DF, nous pouvons obtenir une meilleure efficacité énergétique lorsque le débit $\bar{R}$ diminue, comme le montre la figure 68.a où $\bar{R}=8 \text{ bits/s/Hz}$. Pour ce débit, le relais DF fournit la meilleure efficacité énergétique pour $d_1 \in [56-266]$. Alors que sur la figure 68.b où le taux est augmenté à $\bar{R}=12 \text{ bits/s/Hz}$ l'efficacité énergétique obtenue avec le relais DF est très faible. D'autre part, les SIR avec N

optimisé et SISO fournissent la même efficacité énergétique maximale pour les faibles distances, tandis que les SIR atteignent l'efficacité énergétique maximale lorsque $d_1 > 266m$.

Figure 68 : Efficacité énergétique en fonction de la distance d_1 lorsque N est optimisé



(a) $\bar{R} = 8 \text{ bits} / \text{s} / \text{Hz}$



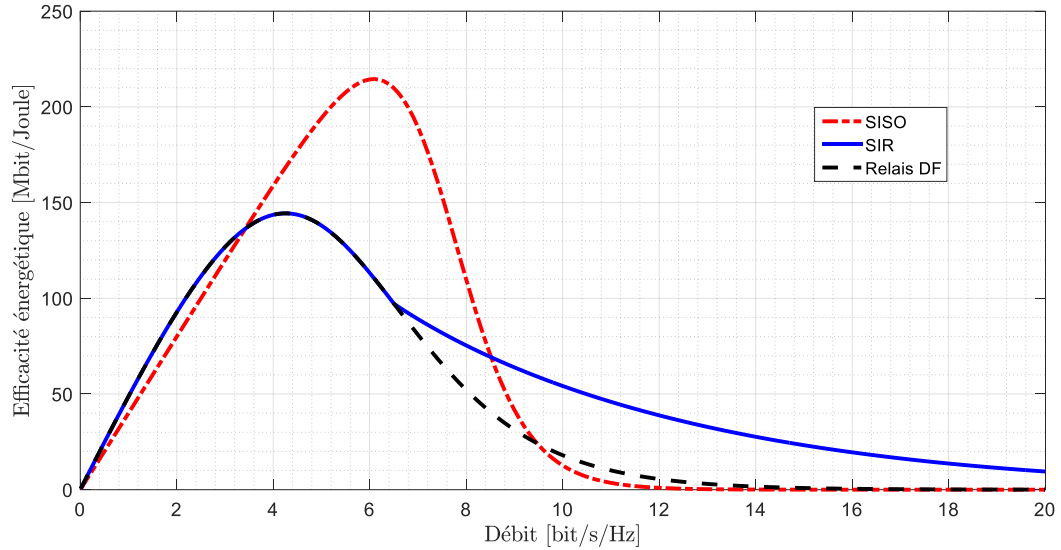
(b) $\bar{R} = 12 \text{ bits} / \text{s} / \text{Hz}$

La figure 69 montre l'efficacité énergétique en fonction du débit réalisable $\bar{R}$. Les SIR et SISO fournissent la même efficacité énergétique maximale pour $\bar{R} \leq 3,42 \text{ bits} / \text{s} / \text{Hz}$, tandis que pour $\bar{R} \in [3,42 - 8,52]$ la technologie du relais DF fournit l'efficacité énergétique la plus élevée. Les SIR performant mieux que SISO et les relais DF lorsque $\bar{R} > 8,52 \text{ bits} / \text{s} / \text{Hz}$ et tentent de maintenir l'efficacité énergétique maximale pour des valeurs de débit élevées. D'après les résultats obtenus, nous pouvons conclure que pour les faibles débits et les courtes distances,



un système avec des modes de commutation entre SISO et SIR est plus efficace pour maximiser l'efficacité énergétique. En outre, la SIR peut prendre en charge la transmission sur de grandes distances à haut débit avec la meilleure efficacité énergétique.

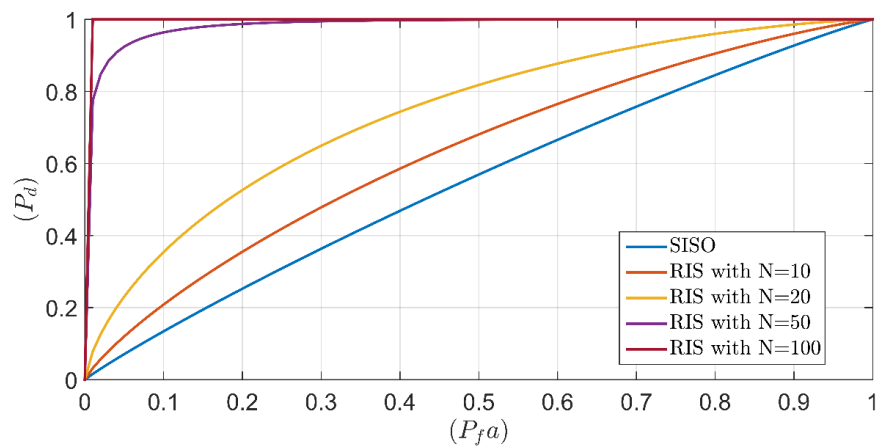
Figure 69 : Efficacité énergétique en fonction du débit réalisable lorsque N est optimisé



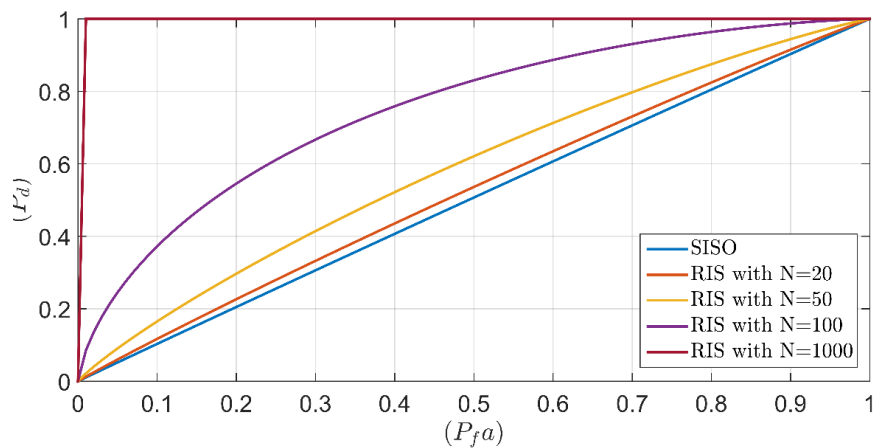
3.2 Réseau radio cognitif supporté par les surfaces intelligentes reconfigurables

Nous examinons maintenant les performances de la technique de détection du spectre pour les trois configurations. Pour étudier la performance d'un détecteur, la méthode la plus couramment utilisée consiste à tracer une courbe ROC qui représente la variation de la probabilité de détection P_d en fonction de la probabilité de fausse alarme P_{fa} . En utilisant (40) et (60), nous obtenons les courbes ROC de la figure 70 où la source utilise une puissance d'émission de $P_t = 0,1W$ ou $P_t = 0,001W$. D'après la figure 70.a, nous observons qu'une meilleure performance est obtenue lorsque les SIR sont équipées d'un grand nombre d'éléments N . La figure 70.b montre que les SIR requièrent plus d'antennes pour obtenir une meilleure performance de détection du spectre lorsque la puissance d'émission diminue à $P_t = 0,001W$. La figure 71 montre la probabilité de détection P_d en fonction du nombre d'éléments SIR et du seuil γ_{th} . La probabilité de détection est largement améliorée lorsqu'un nombre important d'éléments est adopté dans les SIR. Pour $N = 0$ (cas SISO), la probabilité de détection est fortement affectée par le seuil γ_{th} , et nous obtenons de faibles valeurs de P_d même si $\gamma_{th} = 0W$. L'augmentation du nombre d'antennes dans les SIR (à partir de $N = 50$) peut améliorer considérablement les performances de détection, quel que soit le seuil de détection utilisé.

Figure 70 : Courbes ROC pour les modèles SISO et SIR avec différentes configurations

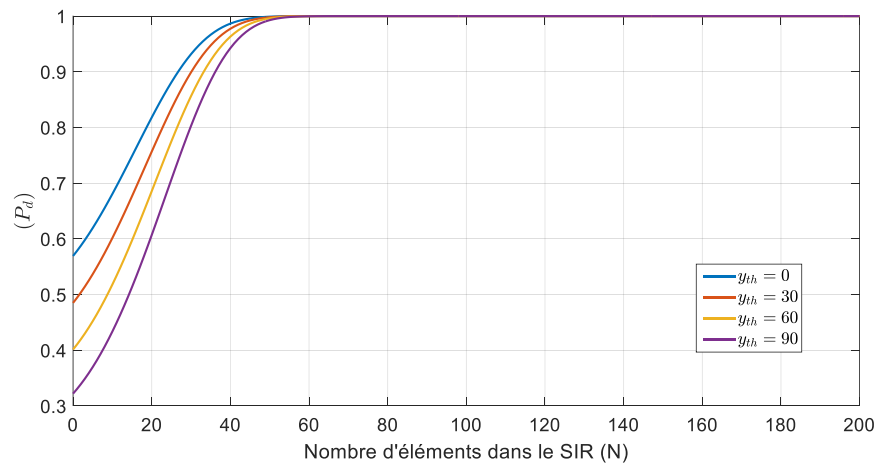


(a) $P_t = 0.1W$



(b) $P_t = 0.001W$

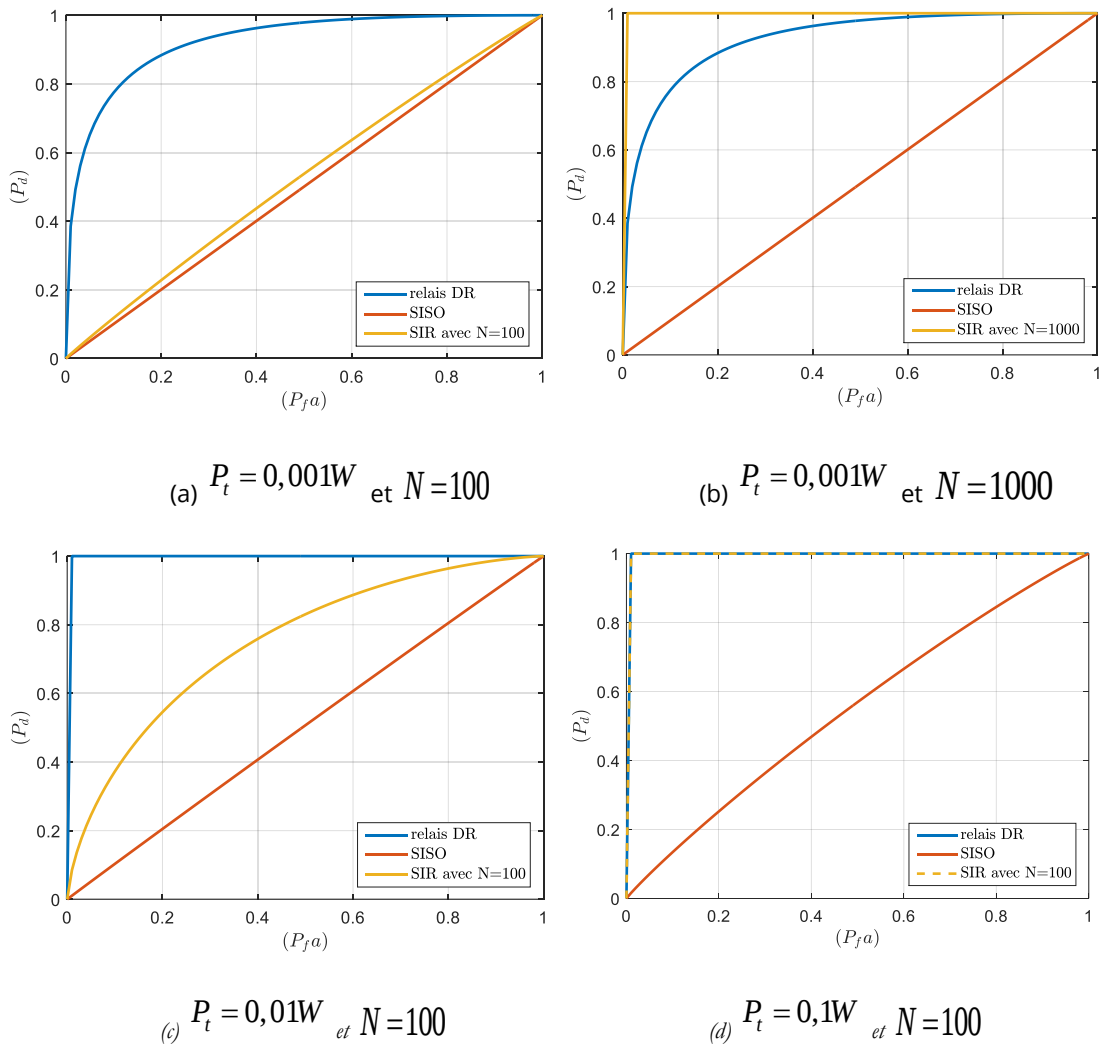
Figure 71 : Courbes ROC pour les modèles SISO et SIR avec différentes configurations SIR lorsque $P_t = 0.001W$





Nous comparons maintenant la performance de détection des trois configurations SISO, SIR et relais DF. La figure 72 illustre les courbes ROC avec les puissances d'émission $P_t \in \{0.001, 0.01, 0.1\}$ et $N \in \{100, 1000\}$. Dans le cas de $P_t = 0.001W$ et $N = 100$, le relais DF fournit la meilleure performance tandis que SIR et SISO fournissent une faible performance, comme le montre la figure 72.a. La figure 72.b indique que SIR peut largement surpasser relais DF lorsque le nombre d'éléments augmente (par exemple lorsque $N = 1000$). Des performances plus élevées sont obtenues à la figure 72.c et à la figure 72.d où la puissance est respectivement augmentée à $P_t = 0.01W$ et $P_t = 0.1W$. Par conséquent, la SIR avec un nombre élevé d'antennes devient plus performante même si la puissance d'émission est faible, elle assure les meilleures performances de détection et une grande efficacité énergétique.

Figure 72 : Courbes ROC des modèles SISO, SIR et relais DF lorsque $N = 100$ et $P_t = 0.001W$





4. Implémentation de l'algorithme hybride GDS et résultats

Nous simulons maintenant le modèle hybride d'accès au spectre. Nous supposons qu'un ensemble $N_c=20$ d'utilisateurs cognitifs uniformément répartis dans une zone carrée de $200m \times 200m$ dans la zone de couverture des utilisateurs autorisés, comme le montre la figure 73. les utilisateurs fonctionnent pour avoir accès à $K=5$ canaux sous licence. Nous supposons que le système secondaire effectue une détection précise avec une probabilité de détection très élevée ($P_d \approx 1$) et une très faible probabilité de fausse alarme $P_{fa} \approx 0$. Le temps de détection est supposé très faible comparé au temps utilisé pour la transmission des données. Nous considérons les dissipations de puissance $P_s, P_d, P_r \in [80, 120]$, $P_e \in [5, 10]$ et le débit $\bar{R}=10 \text{ bits} / \text{s} / \text{Hz}$. La distance minimale entre la destination et SIR / relais est $d_2=10m$, la plage des valeurs d_1 est $d_1 \in [40, 100]$, la plage de distances entre la source et SIR / relais est $d_{sr} \in [70, 150]$. Pour tous les autres paramètres, nous utilisons les mêmes valeurs que précédemment.

Figure 73 : Diagramme du résultat de l'opération de détection du spectre

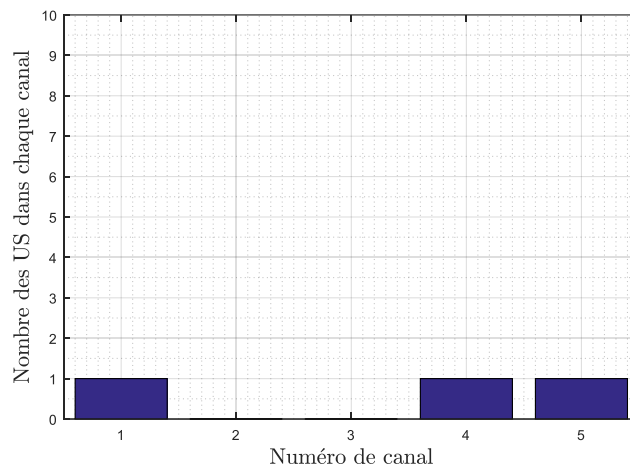


Figure 74 : Diagramme d'allocation initiale pour les utilisateurs secondaires dans chaque canal

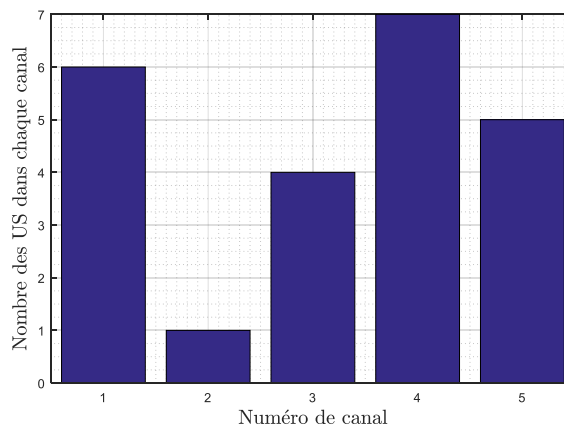


Figure 75 : Diagramme de résultat de l'allocation des canaux par le modèle OSA

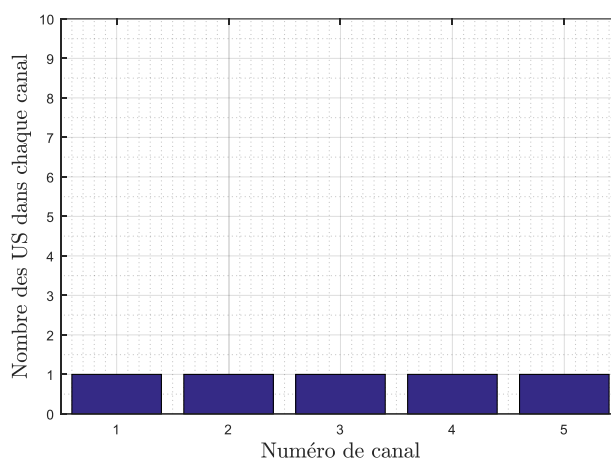
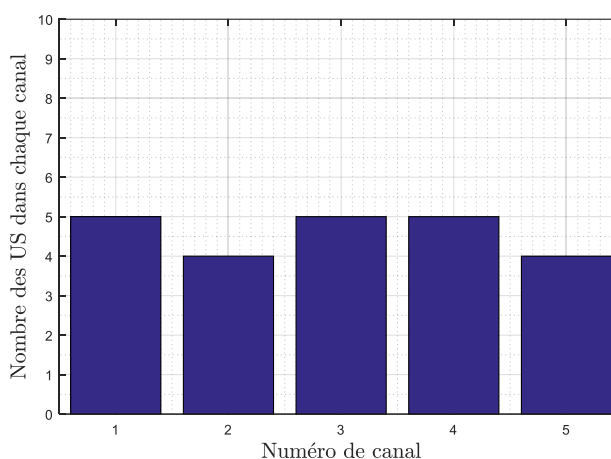


Figure 76 : Diagramme de résultat de l'allocation des canaux basée sur le modèle hybride OSA-CSA

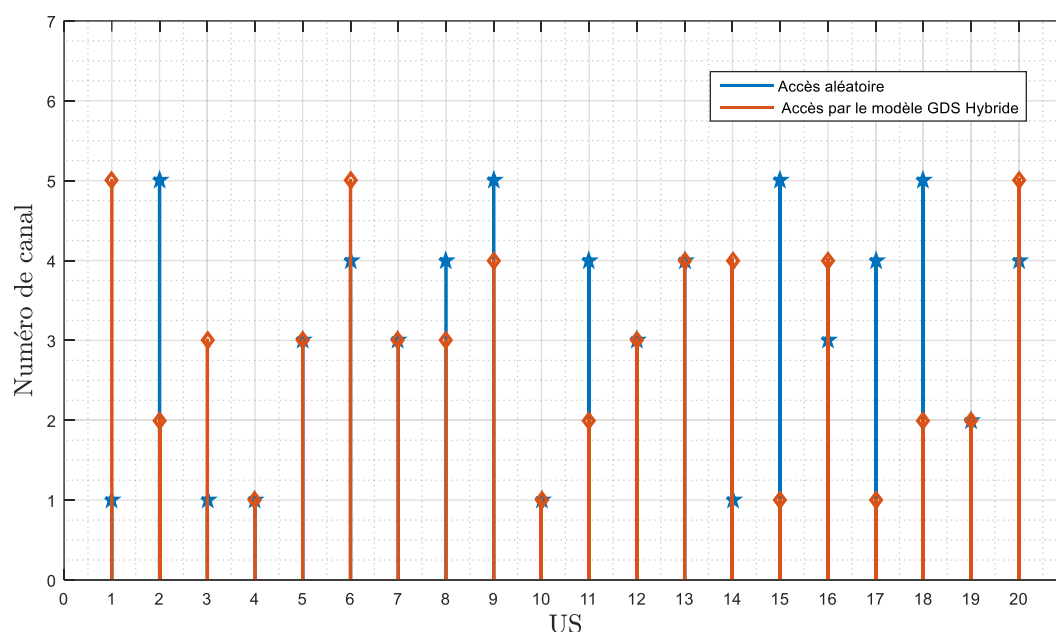


La figure 73 montre le résultat de détection du spectre sur K canaux où le détecteur révèle que les canaux 2 et 3 sont inoccupés tandis que les canaux 1, 4 et 5 sont occupés par leurs utilisateurs licenciés. La figure 74 montre une allocation initiale qui est générée de façon aléatoire. On peut observer que l'occupation des canaux n'est pas uniforme, certains canaux étant très occupés, comme le canal 4 avec 7 utilisateurs, tandis que d'autres sont faiblement exploités comme le canal 2 avec un seul utilisateur. L'allocation aléatoire peut entraîner un gaspillage important des ressources du spectre et générer des niveaux d'interférence élevés. Par ailleurs, la figure 75 montre l'allocation du spectre par l'algorithme OSA. On voit clairement que chaque canal inoccupé est utilisé de manière opportuniste par un utilisateur. La figure 76 illustre l'allocation finale du spectre obtenue à la fin de l'exécution de l'algorithme CSA. L'occupation du spectre est améliorée de manière adaptative par rapport à la première allocation aléatoire. Le nombre d'utilisateurs dans le canal 4 passe de 7 à 5, et le canal 2 devient alloué à 4 utilisateurs au lieu de 1. L'algorithme proposé considère pleinement le problème de

congestion du spectre pour réduire l'interférence des utilisateurs, ce qui protège parfaitement les services.

La figure 77 présente les canaux attribués à chaque utilisateur dans le cas de l'attribution aléatoire initiale et de l'attribution hybride proposée. Avec les capacités de la RC, chaque utilisateur ajuste sa puissance d'émission et change le canal utilisé pour minimiser la consommation d'énergie, le brouillage du système, puis l'efficacité spectrale. Par exemple, l'utilisateur numéro 1 change la transmission du canal 5 au canal 1, tandis que l'utilisateur numéro 10 continue de transmettre sur le canal 1.

Figure 77 : Mobilité des US après la mise en œuvre de l'algorithme hybride OSA-CSA proposé



Nous étudions maintenant l'allocation hybride du spectre avec une transmission supportée par SIR / relais DF. La figure 78 illustre l'efficacité énergétique totale sur chaque canal lorsque les utilisateurs adoptent l'accès dynamique et hybride au spectre. La SIR fournit la meilleure efficacité énergétique pour tous les canaux et surpasse largement les cas relais DF et SISO. La figure 79 présente les résultats lorsque le débit passe à $\bar{R} = 14 \text{ bits} / \text{s} / \text{Hz}$ où l'on constate que les performances diminuent, la SIR conservant sa supériorité en termes de performances. Les résultats confirment que Relais DF offre une performance très faible avec un débit de données élevé. En outre, la figure 80 montre les résultats lorsque le nombre de US passe à $N_c = 50$. Avec un plus grand nombre d'utilisateurs, l'efficacité énergétique de chaque canal augmente car les utilisateurs essaient de réduire leur puissance d'émission pour minimiser les interférences générées. Notons que les figures 78, 79 et 80 sont générées par des simulations de Monte-Carlo avec 200 itérations.

Figure 78 : Efficacité énergétique totale dans chaque canal alloué quand $\bar{R} = 10 \text{ bits/s/Hz}$ et $N_c = 20$

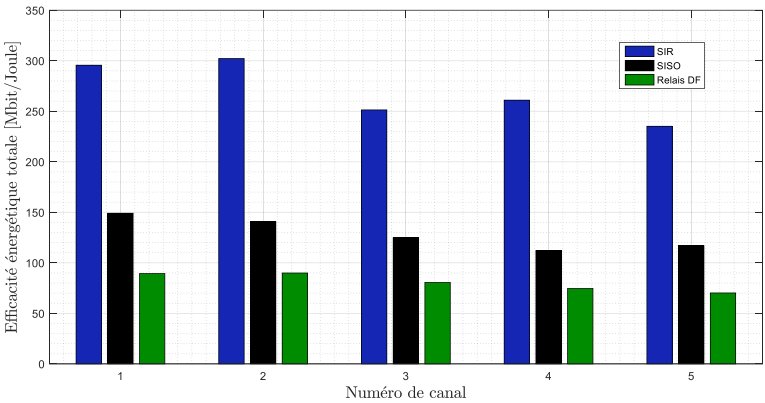


Figure 79 : Efficacité énergétique totale pour chaque canal dans le cas où $\bar{R} = 14 \text{ bits/s/Hz}$ et $N_c = 20$

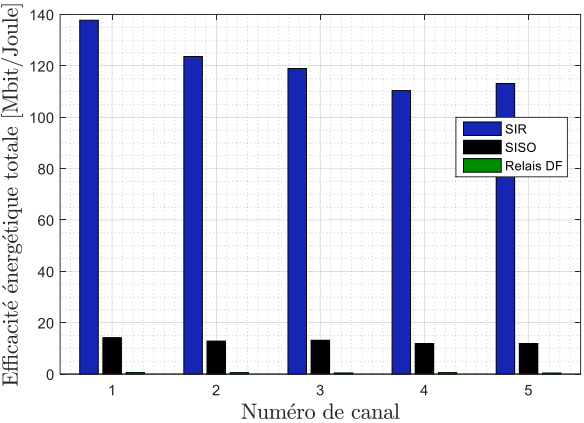
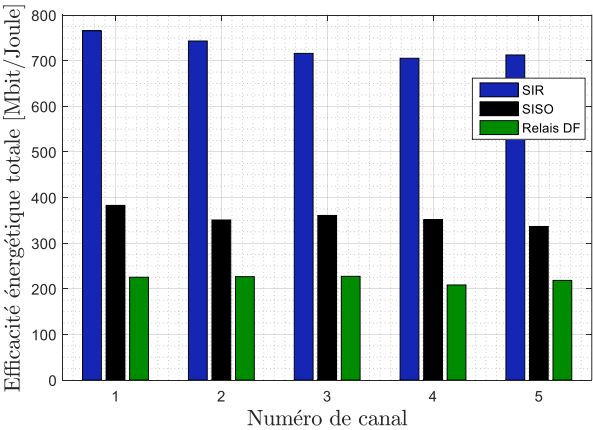


Figure 80 : Efficacité énergétique totale pour chaque canal dans le cas où $\bar{R} = 10 \text{ bits/s/Hz}$ et $N_c = 50$



D'autre part, l'état du spectre peut changer dynamiquement en fonction des activités utilisateurs. Par exemple, la figure 81 illustre un nouvel état du spectre où les canaux 1 et 2 sont occupés. La figure 82 montre la nouvelle occupation du spectre lorsque l'allocation globale du spectre est modifiée. La figure 83 illustre la mobilité des US dans le spectre en conséquence de la nouvelle actualisation de l'état des canaux.

Figure 81 : Nouvelle actualisation de l'état du spectre

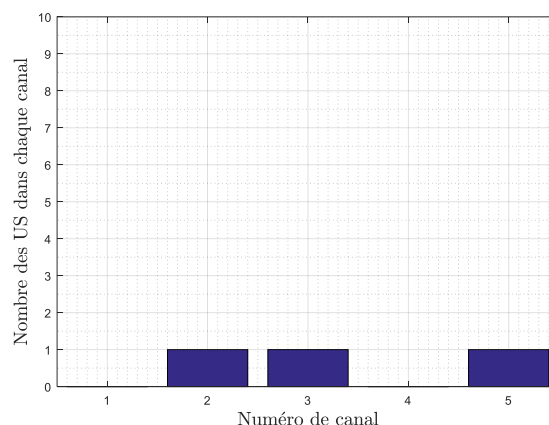
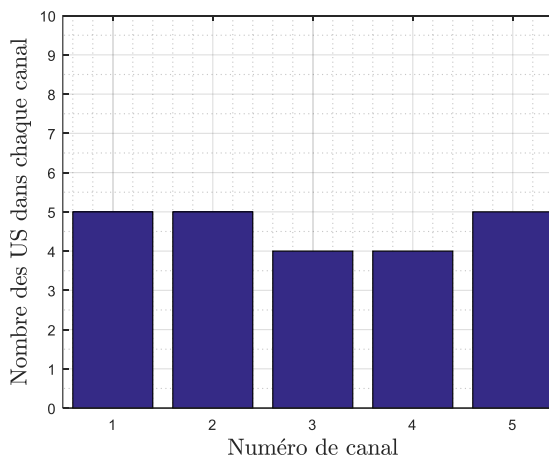
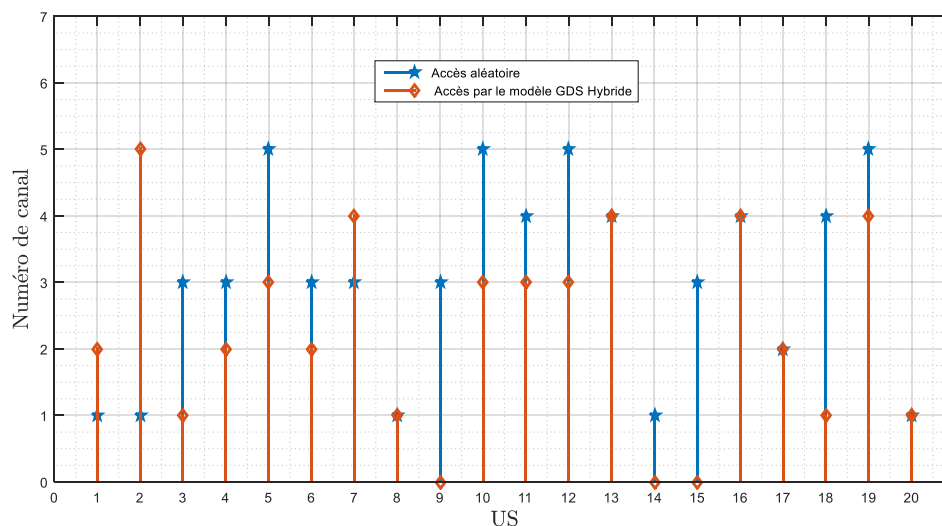


Figure 82 : Allocation du spectre avec le modèle hybride OSA-CSA



Dans la réalité, les dispositifs de communication sans fil sont conçus pour tolérer certains niveaux de brouillage. Cependant, certains utilisateurs ne peuvent pas limiter les interférences qu'ils produisent à un niveau tolérable (marge d'interférence) pour protéger les services. Dans ce cas, d'autres ne peuvent pas bénéficier de l'algorithme hybride d'accès au spectre proposé en raison des contraintes d'interférence. La figure 83 présente les résultats lorsque la contrainte de puissance d'interférence moyenne est adoptée. On peut voir que les utilisateurs 9, 14 et 15 ne peuvent accéder à aucun canal en raison de la contrainte d'interférence.

Figure 83 : Diagramme de la mobilité des US après l'exécution de l'algorithme lorsque la contrainte d'interférence moyenne est adoptée



Conclusion

Dans ce chapitre, nous proposons un système hybride pour la gestion dynamique du spectre (GDS) qui combine les avantages des techniques d'accès opportuniste au spectre et d'accès simultané au spectre pour améliorer l'exploitation des fréquences dans les environnements de radio cognitive soutenus par des surfaces intelligentes reconfigurables (SIR). Nous dérivons des expressions pour les probabilités de détection et de fausse alarme. Nous étudions comment le nouveau concept de surfaces intelligentes reconfigurables améliore les performances des techniques de détection du spectre par rapport au relais classique de type décode et retransmettre (relais DR) en traçant des courbes de caractéristiques opérationnelles du récepteur. Nous comparons ensuite les SIR aux technologies de relais DF en termes d'efficacité énergétique et de capacité de canal réalisable dans des environnements de propagation sans visibilité directe. Nous étudions également l'effet de la configuration SIR et Relais DF sur le débit de données atteint et l'efficacité énergétique. Enfin, nous étudions l'accès dynamique et hybride au spectre lorsque les technologies SIR / Relais DF prennent en charge la transmission. Dans le système GDS hybride, nous prenons en compte la coexistence de plusieurs dispositifs sans fil qui englobe UP et US sur le même canal. Nous commençons par une opération de détection du spectre pour déterminer son état. Puis, un modèle de jeu est développé pour contrôler l'accès des utilisateurs aux bandes libres du spectre. Ensuite, tout en garantissant la protection des services dus, le système permet aux utilisateurs d'accéder aux bandes de spectre allouées.

Les résultats des simulations montrent que les SIR atteignent une efficacité énergétique plus élevée que pour Relais DF et SISO et ce, pour différentes configurations. En outre, les SIR améliorent considérablement la qualité des communications sans fil et les performances de détection du spectre dans un environnement RC. Les résultats numériques montrent également que la gestion dynamique et hybride du spectre soutenue par SIR améliore grandement la capacité du canal et l'efficacité énergétique dans tous les états du canal.

◆ Conclusion générale

Depuis l'avènement de la téléphonie mobile, il y a une quarantaine d'années lors de l'apparition au début des années 1980 de la première génération, les demandes en services, débit et couverture continuent toujours à augmenter à un rythme très accéléré. Un travail continu des opérateurs pour la satisfaction et la fidélisation de leurs abonnés d'une part, et d'autre part des différents équipementiers et instances de normalisation internationales pour la mise au point de technologies et d'infrastructures, est à l'œuvre pour être en mesure de satisfaire les besoins. Cette thèse ayant pour vocation première de tracer quelques pistes pour l'amélioration des services mobiles qui seront susceptibles d'être effectifs avec la 5G et la prochaine génération de communication sans fil 6G, est organisée autour de cinq chapitres.

Dans le premier chapitre nous présentons les différentes générations allant de la 1G à la 5G, et qui se succèdent sur un rythme décennal. Nous mettons l'accent sur les améliorations qui sont apportées par chacune, ainsi que la manière avec laquelle elles répondent aux besoins. Force est de constater que jusqu'à la 4G et ses évolutions, les efforts des acteurs concernés se sont principalement portés sur l'augmentation du débit et que la 5G est considérée comme une révolution par rapport aux générations précédentes dans la mesure où elle introduit trois types de services demandant des débits différents et des qualités de services distinctes, elle permet donc aux opérateurs de proposer différents services de manière flexible.

Nous présentons, dans le deuxième chapitre, la nouvelle interface d'accès radio introduite par la technologie 5G (5G NR) et ses diverses caractéristiques (numérologie, modulation, spectre de fréquences, etc.). À travers ce chapitre, nous montrons que la technologie nécessaire pour garantir une latence aller-retour aussi faible que 1 ms est toujours à l'étude et n'est pas encore prête à relever les défis posés par des services très exigeants en termes de latence, de fiabilité et de qualité de service tel que l'internet tactile. Nous montrons également que la prise en compte de la transmission de petits paquets réduit également la probabilité d'erreur de transmission et que cette probabilité augmente à mesure que la longueur des paquets augmente. D'autre part, nous examinons également la variation de la probabilité d'erreur de transmission en fonction de la variation de la durée de transmission et nous constatons que cette probabilité décroît au-delà d'une certaine valeur de cette durée.

Dans le troisième chapitre nous dressons un vaste panorama de ce que représente aujourd'hui cette sixième génération de réseau mobile. Il a fait l'objet d'une publication dans laquelle nous identifions 12 technologies qui émergent et se confirment dans une vision globale de l'arrivée de la 6G que nous récapitulons à l'aide d'un tableau après avoir étudié leurs caractéristiques. Nous montrons l'étendue des scénarios dans plusieurs domaines, inenvisageables avec la 5G. Les techniques et les thèmes que nous retenons sont primordiaux dans la future architecture technologique de la 6G et l'amélioration des services. Les enjeux en termes de développement durable sont cernés, les changements qui s'annoncent dans notre relation au numérique aussi. Pour conclure ce chapitre, nous illustrons avec quelques exemples significatifs dans plusieurs régions du monde, la grande attention que portent les gouvernements à la 6G,. De ce point de vue le constat est une évidence : la 6G est très attendue.

Nous introduisons le quatrième chapitre par le concept l'intelligence artificielle avec un court historique qui en explique les principes et les grandes directions. Nous développons un certain nombre de concepts, utiles à une optimisation de la ressource spectrale comme les réseaux de neurones artificiels. Puis nous nous intéressons en détail aux différents auto-encodeurs, leur principe de fonctionnement et leurs applications. Alors, pour introduire notre modèle, nous décrivons la prédiction de formation de faisceaux à ondes millimétriques et notre jeu de données, puis nous présentons notre approche avec un auto-encodeur empilé et des perceptrons multicouches que nous évaluons avec la métrique d'erreur quadratique moyenne. Si nous les comparons avec la méthode de référence, les résultats sont très encourageants parce que notre approche de l'apprentissage profond est plus performante, nous aimerions par ailleurs tester encore d'autres scénarios basés sur l'apprentissage en profondeur qui sont prometteurs.

Dans le cinquième et dernier chapitre, nous présentons une nouvelle technologie appelée surfaces intelligentes reconfigurables (SIR) intégrée à des environnements radio cognitive pour une allocation dynamique du spectre. Nous exposons d'abord l'état de l'art sur les SIR en mettant l'accent sur le principe de fonctionnement, l'évaluation des performances, le concept de la formation de faisceaux, la gestion des ressources et l'intégration de cette technologie à la radio cognitive qui améliore le fonctionnement des réseaux sans fil. Puis, nous comparons la transmission assistée par SIR avec la transmission assistée par les relais de décodage et transfert dans le but de déterminer la configuration nécessaire à une SIR pour surpasser la performance d'un relais conventionnel. Nous démontrons ainsi comment la phase délicate de détection du spectre dans la radio cognitive est améliorée en employant cette nouvelle technique que sont les surfaces intelligentes reconfigurables. Nous proposons également un modèle hybride pour la gestion dynamique du spectre dans les environnements RC assistés par SIR et par les relais de décodage et transfert ; ce modèle est basé sur la théorie des jeux qui réalise conjointement l'allocation du spectre et l'atténuation des interférences entre utilisateurs. Dans un premier temps, le problème d'allocation du spectre est transformé en problème d'optimisation conjointe. Ensuite, un algorithme issu de la théorie des jeux est utilisé pour réaliser le schéma d'allocation optimal du spectre. Avec cette technique de gestion dynamique, le spectre est accessible aux deux types d'utilisateurs, avec ou sans licence.

Ce travail nous permet de nous intéresser à d'autres pistes de recherche et à cet égard, nous continuons à travailler sur plusieurs domaines. Nous avons déjà commencé le développement d'un algorithme de colonie d'abeilles artificielles pour maximiser le nombre d'équipements

utilisateurs connectés à une cellule dotée d'antennes intelligentes Massive MIMO et ce, tout en utilisant le minimum d'antennes avec une meilleure efficacité énergétique et une meilleure efficacité spectrale.

Le nouveau protocole internet récemment soumis par la Chine à l'Union internationale des télécommunications, appelé « New IP », retient aussi toute notre attention. En effet, le protocole actuel, TCP/IP, est parvenu aux limites de ses capacités techniques, il ne peut pas répondre aux services futurs (traditionnels et verticaux).

Nous regardons aussi de près la communication entre véhicules autonomes dans le réseau 5G ainsi que la transmission d'hologrammes dans un réseau 6G.



Bibliographie

- [1] A deliverable by the MGMN Alliance, NGMN 5G White paper (2015)
- [2] Recommendation ITU-R M.2083-0, IMT Vision – Framework and overall objectives of the future development of IMT for 2020 and beyond, https://www.itu.int/dms_pubrec/itu-r/rec/m/R-REC-M.2083-0-201509-!!!PDF-E.pdf (2015)
- [3] 5G PPP METIS White Paper on 5G RAN Architecture and Functional Design, March 4th, https://metis-ii.5g-ppp.eu/wp-content/uploads/white_papers/5G-RAN-Architecture-and-Functional-Design.pdf (2016)
- [4] ITU-T Y.3101. Requirements of the IMT-2020 network, , <https://www.itu.int/rec/T-REC-Y.3101-201801-I/en>, (2018)
- [5] Recommendation ITU-R M.2410-0, Minimum requirements related to technical performance for IMT- 2020 radio interfaces, https://www.itu.int/dms_pub/itu-r/opb/rep/R-REP-M.2410-2017-PDF-E.pdf (2017)
- [6] 3GPP TR 38.801, "Technical Specification Group Radio Access Network; Study on new radio access technology: Radio access architecture and interfaces", (2017)
- [7] 3GPP TS 38.300 V16.0.0: Technical Specification Group Radio Access Network, NR, NR and NG-RAN Overall Description (2019).
- [8] 5GPPP: Broadcast and Multicast Communication Enablers for the Fifth Generation of Wireless Systems. D3.2, Air Interface, Version v2.00, (2018).
- [9] T.A.Yahiya and P. Kirci, "Issue and challenges facing low latency in Tactile Internet", in UKH Journal of Science and Engineering, Volume 3, Number 1 (2019)
- [10] Y. Zhao, G. Yu et H. Xu, "Réseau de communication mobile 6G : vision, défis et technologies clés", Sci. Pêche. Informationis, vol. 49, non. 8, pp. 963–987, mai 2019, doi : 10.1360/n112019-00033.
- [11] M. Cave, « À quel point la 5G est-elle perturbatrice ? » Telecomm. Politique, vol. 42, non. 8, pp. 653–658, septembre 2018, doi : 10.1016/J.TELPOL.2018.05.005.
- [12] "Rapport Internet annuel de Cisco - Livre blanc sur le rapport Internet annuel de Cisco (2018-2023) - Cisco." .
- [13] Meer Zafarullah Noohani et Kaleem Ullah Magsi, "Un examen de la technologie 5G : architecture, sécurité et applications étendues", Int. Rés. J.Eng. Technol. , vol. 07, non. 05, p. 3440–3471, 2020, doi : 10.5281/ZENODO.3842353.
- [14] KB Letaief, W. Chen, Y. Shi, J. Zhang et Y.-JA Zhang, « The Roadmap to 6G -- AI Empowered Wireless Networks », IEEE Commun. Mag., vol. 57, non. 8, pp. 84–90, avril 2019, consulté le 11 juillet 2021. [En ligne]. Disponible : <https://arxiv.org/abs/1904.11686v2>.
- [15] S. Chen, YC Liang, S. Sun, S. Kang, W. Cheng et M. Peng, "Vision, exigences et tendance technologique de la 6G : comment relever les défis de la couverture, de la capacité et du débit des données de l'utilisateur et la vitesse de déplacement », IEEE Wirel. Commun., vol. 27, non. 2, pp. 218–228, avril 2020, doi : 10.1109/MWC.001.1900333.



- [18] P. Yang, Y. Xiao, M. Xiao et S. Li, « Communications sans fil 6G : vision et techniques potentielles », *IEEE Netw.*, vol. 33, non. 4, pp. 70–75, juillet 2019, doi : 10.1109/MNET.2019.1800418.
- [19] MW Akhtar, SA Hassan, R. Ghaffar, H. Jung, S. Garg et MS Hossain, « Le passage aux communications 6G : vision et exigences », *Human-centric Comput. Inf. Sci.* 2020 101, vol. 10, non. 1, pp. 1–27, décembre 2020, doi : 10.1186/S13673-020-00258-2.
- [20] H. Tataria, M. Shafi, AF Molisch, M. Dohler, H. Sjöland et F. Tufvesson, « Systèmes sans fil 6G : vision, exigences, défis, idées et opportunités », *Proc. IEEE*, août 2020, consulté le : 09 juillet 2021. [En ligne]. Disponible : <https://arxiv.org/abs/2008.03213v2>.
- [21] MH Alsharif, AH Kelechi, MA Albreem, SA Chaudhry, M. Sultan Zia et S. Kim, « Réseaux sans fil de sixième génération (6G) : vision, activités de recherche, défis et solutions potentielles », *Symmetry (Bâle)*, vol. 12, non. 4, avril 2020, doi : 10.3390/SYM12040676.
- [22] Z. Zhang et al., "Réseaux sans fil 6G : vision, exigences, architecture et technologies clés", *IEEE Veh. Technol. Mag.*, vol. 14, non. 3, p. 28–41, septembre 2019, doi : 10.1109/MVT.2019.2921208.
- [23] K. David et H. Berndt, "Vision et exigences 6G : est-il nécessaire d'aller au-delà de la 5g ?", *IEEE Veh. Technol. Mag.*, vol. 13, non. 3, p. 72–80, septembre 2018, doi : 10.1109/MVT.2018.2848498.
- [24] A. Shahraki, M. Abbasi, MJ Piran et A. Taherkordi, « Une enquête complète sur les réseaux 6G : applications, services de base, technologies habilitantes et défis futurs », *IEEE Trans. Réseau Servir. Gestion*, vol. XX, p. 1, janvier 2021, consulté le : 09 juillet 2021. [En ligne]. Disponible : <https://arxiv.org/abs/2101.12475v2>.
- [25] X. You et al., "Vers les réseaux de communication sans fil 6G : vision, technologies habilitantes et nouveaux changements de paradigme", *Sci. Chine Inf. Sci.* 2020 641, vol. 64, non. 1, pp. 1–74, novembre 2020, doi : 10.1007/S11432-020-2955-6.
- [26] A. Yazar, S. Doğan-Tusha et H. Arslan, « 6G Vision : An Ultra-Flexible Perspective », *ITU J. Futur. Évol. Technol.*, vol. 2020, sept. 2020, consulté le : 09 juillet 2021. [En ligne]. Disponible : <https://arxiv.org/abs/2009.07597v2>.
- [27] F. Nawaz, J. Ibrahim, MA Ali, M. Junaid, S. Kousar et T. Parveen, « Un examen de la vision et des défis de la technologie 6G », *Int. J. Adv. Calcul. Sci. Appl.*, vol. 11, non. 2, p. 643–649, 2020, doi : 10.14569/IJACSA.2020.0110281.
- [28] A. Clemm, MT Vega, HK Ravuri, T. Wauters et F. De Turck, « Vers une communication véritablement immersive de type holographique : défis et solutions », *IEEE Commun. Mag.*, vol. 58, non. 1, pp. 93–99, janvier 2020, doi : 10.1109/MCOM.001.1900272.
- [29] CE Shannon, « Une théorie mathématique de la communication », *Bell Syst. Technologie. J.*, vol. 27, non. 3, pp. 379–423, juillet 1948, doi : 10.1002/J.1538-7305.1948.TB01338.X.
- [30] M. Giordani, M. Polese, M. Mezzavilla, S. Rangan et M. Zorzi, « Vers les réseaux 6G : cas d'utilisation et technologies », *IEEE Commun. Mag.*, vol. 58, non. 3, p. 55–61, mars 2020, doi : 10.1109/MCOM.001.1900411.
- [31] S. Dang, O. Amin, B. Shihada et M.-S. Alouini, "Que devrait être la 6G ?", *Nat. Électron.* 2020 31, vol. 3, non. 1, pp. 20–29, janvier 2020, doi : 10.1038/s41928-019-0355-6.
- [32] L. Zhang, YC Liang et D. Niyato, « 6G Visions : Mobile ultra-large bande, super internet des objets et intelligence artificielle », *China Commun.*, vol. 16, non. 8, pp. 1–14, août 2019, doi : 10.23919/JCC.2019.08.001.
- [33] T. Huang, W. Yang, J. Wu, J. Ma, X. Zhang et D. Zhang, « Une enquête sur le réseau 6G vert : architecture et technologies », *IEEE Access*, vol. 7, pages 175758–175768, 2019, doi : 10.1109/ACCESS.2019.2957648.
- [34] B. Zong, C. Fan, X. Wang, X. Duan, B. Wang et J. Wang, "Technologies 6G : principaux moteurs, exigences de base, architectures système et Technologies », *IEEE Véh. Technol. Mag.*, vol. 14, non. 3, p. 18–27, septembre 2019, doi : 10.1109/MVT.2019.2921398.
- [35] Y. Yuan, Y. Zhao, B. Zong et S. Parolari, "Technologies clés potentielles pour les communications mobiles 6G", *Sci. Chine Inf. Sci.*, vol. 63, non. 8, p. 80300, octobre 2019, doi : 10.1007/s11432-019-2789-y.
- [36] M. Latva-Aho et K. Leppänen, « Facteurs clés et défis de recherche pour l'intelligence sans fil omniprésente 6G - 6G Research Visions 1, septembre 2019 », no. Septembre, p. 1–36, 2019.



- [37] W. Saad, M. Bennis et M. Chen, « Une vision des systèmes sans fil 6G : applications, tendances, technologies et problèmes de recherche ouverts », *IEEE Netw.*, vol. 34, non. 3, p. 134–142, mai 2020, doi : 10.1109/MNET.001.1900287.
- [38] H. Viswanathan et PE Mogensen, « Communications à l'ère 6G », *IEEE Access*, vol. 8, pages 57063–57074, 2020, doi : 10.1109/ACCESS.2020.2981745.
- [39] S. Elmeadawy et RM Shubair, « Enabling Technologies For 6g Future Wireless Communications: Opportunities And Challenges », février 2020, consulté le : 09 juillet 2021. [En ligne]. Disponible : <https://arxiv.org/abs/2002.06068v1>.
- [40] V. Ziegler, T. Wild, M. Uusitalo, H. Flinck, V. Raisanen et K. Hatonen, « Stratification de l'évolution de la 5G et au-delà de la 5G », *IEEE 5G World Forum, 5GWF 2019 - Conf. Proc.*, pp. 329–334, septembre 2019, doi : 10.1109/5GWF.2019.8911739.
- [41] X. You et al., "Vers les réseaux de communication sans fil 6G : vision, technologies habilitantes et nouveaux changements de paradigme", *Sci. Chine Inf. Sci.* 2020 641, vol. 64, non. 1, pp. 1–74, novembre 2020, doi : 10.1007/S11432-020-2955-6.
- [42] S. Chen, Y.-C. Liang, S. Sun, S. Kang, W. Cheng et M. Peng, « Vision, exigences et tendance technologique de la 6G : comment relever les défis de la couverture, de la capacité, du débit de données utilisateur et de la vitesse de déplacement du système » *IEEE Wirel. Commun.*, vol. 27, non. 2, pp. 218–228, février 2020, consulté le : 09 juillet 2021. [En ligne]. Disponible : <https://arxiv.org/abs/2002.04929v1>.
- [43] DANG, Shuping, Oussama AMIN, Basem SHIHADA et Mohamed-Slim ALOUINI. Que devrait être la 6G ? *Nature Électronique* 2020 3:1[en ligne]. 2020, 3(1), 20–29 [vidéo. 2021-07-09]. ISSN 2520-1131. Dostupné z : doi:10.1038/s41928-019-0355-6
- [44] AKHTAR, Muhammad Waseem, Syed Ali HASSAN, Rizwan GHAFAR, Haejoon JUNG, Sahil GARG et M. Shamim HOSSAIN. Le passage aux communications 6G : vision et exigences. *Informatique et sciences de l'information centrées sur l'humain* 2020 10: 1[en ligne]. 2020, 10(1), 1–27 [vid. 2021-07-09]. ISSN 2192-1962. Dostupné z : doi:10.1186/S13673-020-00258-2
- [45] LETAIEF, Khaled B., Wei CHEN, Yuanming SHI, Jun ZHANG et Ying-Jun Angela ZHANG. La feuille de route vers la 6G -- Réseaux sans fil optimisés par l'IA. *Magazine des communications de l'IEEE*[en ligne]. 2019, 57(8), 84–90 [vidéo. 2021-07-11]. Dostupné z : <https://arxiv.org/abs/1904.11686v2>
- [46] SAAD, Walid, Mehdi BENNIS et Mingzhe CHEN. Une vision des systèmes sans fil 6G : applications, tendances, technologies et problèmes de recherche ouverts. *Réseau IEEE*[en ligne]. 2020, 34(3), 134–142 [vidéo. 2021-07-09]. ISSN 0890-8044. Dostupné z : doi:10.1109/MNET.001.1900287
- [47] RAPPAPORT, Theodore S., Yunchou XING, Ojas KANHERE, Shihao JU, Arjuna MADANAYAKE, Soumyajit MANDAL, Ahmed ALKHATEEB et Georgios C. TRICHOPOULOS. Communications et applications sans fil au-dessus de 100 GHz : opportunités et défis pour la 6G et au-delà. *Accès IEEE*[en ligne]. 2019, 7, 78729–78757. ISSN 21693536. Dostupné z : doi:10.1109/ACCESS.2019.2921522
- [48] TRIPATHI, Shuchi, Nithin V. SABU, Abhishek K. GUPTA et Harpreet S. DHILLON. Ondes millimétriques et spectre térahertz pour le sans fil 6G [en ligne]. 2021, 83-121 [vidéo. 2022-02-03]. Dostupné z : doi:10.1007/978-3-030-72777-2_6
- [49] ALSHARIF, Mohammed H., Anabi Hilary KELECHI, Mahmoud A. ALBREEM, Shehzad Ashraf CHAUDHRY, M. SULTAN ZIA a Sunghwan KIM. Réseaux sans fil de sixième génération (6G) : vision, activités de recherche, défis et solutions potentielles. *Symétrie*[en ligne]. 2020, 12(4). Dostupné z : doi:10.3390/SYM12040676
- [50] KHAN, Latif U., Ibrar YAQOOB, Muhammad IMRAN, Zhu HAN et Choong Seon HONG. Systèmes sans fil 6G : une vision, des éléments architecturaux et des orientations futures. *Accès IEEE*[en ligne]. 2020, 8, 147029–147044. ISSN 21693536. Dostupné z : doi:10.1109/ACCESS.2020.3015289
- [51] VOUS, Xiaohu, Cheng-Xiang WANG, Jie HUANG, Xiqi GAO, Zaichen ZHANG, Mao WANG, Yongming HUANG, Chuan ZHANG, Yanxiang JIANG, Jiaheng WANG, Min ZHU, Bin SHENG, Dongming WANG, Zhiwen PAN, Pengcheng ZHU, Yang YANG, Zening LIU, Ping ZHANG, Xiaofeng TAO, Shaoqian LI, Zhi CHEN, Xinying MA, Chih-Lin I, Shuangfeng HAN, Ke LI, Chengkang PAN, Zhimin ZHENG, Lajos HANZO, Xuemin (Sherman) SHEN, Yingjie Jay GUO, Zhiguo DING, Harald HAAS, Wen TONG, Peiying ZHU, Ganghua YANG, Jun WANG, Erik G. LARSSON, ONG Hien Quoc, Wei HONG, Haiming WANG, Debin HOU, Jixin CHEN, Zhe CHEN, Zhangcheng HAO, Geoffrey Ye LI, Rahim TAFAZOLLI, Yue GAO, H. Vincent POOR, Gerhard P. FETTWEIS et Ying-Chang LIANG. Vers les réseaux de communication



- sans fil 6G : vision, technologies habilitantes et nouveaux changements de paradigme. *Science Chine Sciences de l'information* 2020 64: 1[en ligne]. 2020, 64(1), 1–74 [vid. 2021-07-09]. ISSN 1869-1919. Dostupné z: doi:10.1007/S11432-020-2955-6
- [52] TATARIA, Harsh, Mansoor SHAFI, Andreas F. MOLISCH, Mischa DOHLER, Henrik SJÖLAND et Fredrik TUFVESSON. Systèmes sans fil 6G : vision, exigences, défis, perspectives et opportunités. *Actes de l'IEEE*[en ligne]. 2020 [vid. 2021-07-09]. Dostupné z : <https://arxiv.org/abs/2008.03213v2>
- [53] EL METTITI, Abderrahmane et Mohammed OUMSIS. Une enquête sur les réseaux 6G : vision, exigences, architecture, technologies et défis. *Ingénierie des systèmes d'information*[en ligne]. 2022, 27(1), 1–10. ISSN 16331311. Dostupné z: doi:10.18280/ISI.270101
- [54] DAL, Linglong, Ruicheng JIAO, Fumiyuki ADACHI, H. Vincent POOR à Lajos HANZO. Apprentissage en profondeur pour les communications sans fil : un paradigme interdisciplinaire émergent. *Communications sans fil IEEE*[en ligne]. 2020, 27(4), 133–139. ISSN 15580687. Dostupné z: doi:10.1109/MWC.001.1900491
- [55] QI, Chenhao, Yujie WANG et Geoffrey Ye LI. Apprentissage en profondeur pour la formation de faisceaux dans les systèmes MIMO massifs à ondes millimétriques. *Transactions IEEE sur les communications sans fil*[en ligne]. 2020, 1–1. ISSN 1536-1276. Dostupné z: doi:10.1109/TWC.2020.3024279
- [56] WANG, Yuyang, Murali NARASIMHA et Robert W. HEATH. MmWave Beam Prediction with Situational Awareness: A Machine Learning Approach. *Atelier IEEE sur les progrès du traitement du signal dans les communications sans fil, SPAWC*[en ligne]. 2018, 2018-juin. Dostupné z: doi:10.1109/SPAWC.2018.8445969
- [57] CHARAN, Gouranga, Tawfik OSMAN, Andrew HREDZAK, Ngwe THAWDAR et Ahmed ALKHATEEB. Prédiction de faisceau multimodal vision-position à l'aide de jeux de données d'ondes millimétriques réelles [en ligne]. 2021 [vid. 2021-11-30]. Dostupné z : <https://arxiv.org/abs/2111.07574v1>
- [58] ALRABEIAH, Muhammad, Andrew HREDZAK et Ahmed ALKHATEEB. Stations de base à ondes millimétriques avec caméras : prédiction de faisceau et de blocage assistée par vision. *Conférence IEEE sur la technologie des véhicules*[en ligne]. 2020, 2020-Mai. ISSN 15502252. Dostupné z: doi:10.1109/VTC2020-SPRING48590.2020.9129369
- [59] COUSIK, Tarun S, Vijay K SHAH, Jeffrey H Reed WIRELESS@VT , Tugba ERPEK et Yalin E SAGDUYU. Accès initial rapide avec Deep Learning pour la prédiction des faisceaux dans les réseaux 5G mmWave [en ligne]. 2020 [vid. 2021-11-30]. Dostupné z : <https://arxiv.org/abs/2006.12653v1>
- [60] ALRABEIAH, Muhammad, Jayden BOOTH, Andrew HREDZAK et Ahmed ALKHATEEB. ViWi Vision-Aided mmWave Beam Tracking: Dataset, Task, and Baseline Solutions [en ligne]. 2020 [vid. 2021-11-30]. Dostupné z : <https://arxiv.org/abs/2002.02445v3>
- [61] GAO, Feifei, Bo LIN, Chenghong BIAN, Ting ZHOU, Jing QIAN et Hao WANG. FusionNet : prévision de faisceau améliorée pour les communications par ondes millimétriques à l'aide d'un canal sous-6 GHz et de quelques pilotes. *Transactions IEEE sur les communications*[en ligne]. 2020 [vid. 2021-11-30]. ISSN 15580857. Dostupné z: doi:10.1109/TCOMM.2021.3110301
- [62] MA, Ke, Peiyao ZHAO et Zhaocheng WANG. Prédiction de faisceau assistée par apprentissage en profondeur à l'aide d'informations hors bande. *Conférence IEEE sur la technologie des véhicules*[en ligne]. 2020, 2020-Mai. ISSN 15502252. Dostupné z: doi:10.1109/VTC2020-SPRING48590.2020.9128825
- [63] CATAK, Evren, Ferhat Ozgur CATAK et Arild MOLDSVOR. Problèmes de sécurité d'apprentissage automatique contradictoires pour 6G : cas d'utilisation de prédiction de faisceau mmWave. *2021 IEEE International Black Sea Conference on Communications and Networking, BlackSeaCom 2021*[en ligne]. 2021 [vid. 2022-02-03]. Dostupné z: doi:10.1109/BlackSeaCom52164.2021.9527756
- [64] ALRABEIAH, Muhammad et Ahmed ALKHATEEB. Apprentissage en profondeur pour les faisceaux d'ondes millimétriques et la prédiction de blocage à l'aide de canaux inférieurs à 6 GHz. *Transactions IEEE sur les communications*[en ligne]. 2020, 68(9), 5504–5518. ISSN 15580857. Dostupné z: doi:10.1109/TCOMM.2020.3003670
- YING, Ziqiang, Haojun YANG, Jia GAO et Kan ZHENG. Un nouveau schéma de prédiction de faisceau assisté par la vision pour les communications sans fil mmWave. *2020 IEEE 6e Conférence internationale sur l'informatique et les communications, ICCIC 2020*[en ligne]. 2020, 232-237. Dostupné z: doi:10.1109/ICCC51575.2020.9344988



- [65] ALDALBAHI, Adel, Farzad SHAHABI et Mohammed JASIM. Schéma de prédiction de faisceau instantané contre le blocage de liaison dans les communications mmWave. *Sciences Appliquées 2021, Vol. 11, page 5601* [en ligne]. 2021, 11(12), 5601 [vid. 2021-11-11]. ISSN 2076-3417. Dostupné z: doi:10.3390/APP11125601
- [66] SOHRABI, Foad et Wei YU. Conception hybride de formation de faisceaux numérique et analogique pour les réseaux d'antennes à grande échelle. *Journal IEEE sur des sujets sélectionnés dans le traitement du signal* [en ligne]. 2016, 10(3), 501–513. ISSN 19324553. Dostupné z: doi:10.1109/JSTSP.2016.2520912
- [67] ALKHATEEB, Ahmed. DeepMIMO : un ensemble de données d'apprentissage en profondeur générique pour les applications Millimeter Wave et Massive MIMO [en ligne]. 2019 [vid. 2022-02-03]. Dostupné z : <https://arxiv.org/abs/1902.06435v1>
- [68] GUPTA, PRAMOD et NARESH K. SINHA. Réseaux de neurones pour l'identification des systèmes non linéaires : un aperçu. *Soft computing et systèmes intelligents* [en ligne]. 2000, 337–356. Dostupné z: doi:10.1016/B978-012646490-0/50017-2
- [69] *Ensemble de données DeepMIMO* [en ligne]. [vid. 2022-05-26]. Dostupné z : <https://deepmimo.net/>
- [70] MARIA, Joao, Joao AMARO, Gabriel FALCAO et Luís A. ALEXANDRE. Auto-encodeurs empilés utilisant des architectures accélérées à faible consommation d'énergie pour la reconnaissance d'objets dans des systèmes autonomes. *Lettres de traitement neuronal 2015 43: 2* [en ligne]. 2015, 43(2), 445–458 [vidéo. 2022-05-26]. ISSN 1573-773X. Dostupné z: doi:10.1007/S11063-015-9430-9
- [71] MOHAMED, Abdel Rahman, George E. DAHL et Geoffrey HINTON. Modélisation acoustique à l'aide de réseaux de croyances profondes. *Transactions IEEE sur le traitement de l'audio, de la parole et du langage* [en ligne]. 2012, 20(1), 14–22. ISSN 15587916. Dostupné z: doi:10.1109/TASL.2011.2109382
- [72] HINTON, Geoffrey E., Simon OSINDERO et Yee Whye TEH. Un algorithme d'apprentissage rapide pour les réseaux de croyances profondes. *Calcul neuronal* [en ligne]. 2006, 18(7), 1527–1554 [vidéo. 2022-05-26]. ISSN 0899-7667. Dostupné z: doi:10.1162/NECO.2006.18.7.1527
- [73] IAN GOODFELLOW, YOSHUA BENGIO, Aaron Courville. Le livre d'apprentissage en profondeur. *Presse du MIT* [en ligne]. 2017, 521(7553), 785. ISSN 1432122X. Dostupné z: doi:10.1016/B978-0-12-391420-0.09987-X
- [74] KINGMA, Diederik P. et Jimmy Lei BA. Adam : Une méthode d'optimisation stochastique. Dans: *3e Conférence internationale sur les représentations de l'apprentissage, ICLR 2015 - Actes de la piste de conférence*. Bm : Conférence internationale sur les représentations de l'apprentissage, ICLR, 2015.
- [75] BOTCHKAREV, Alexei. Une nouvelle conception de typologie des métriques de performance pour mesurer les erreurs dans les algorithmes de régression d'apprentissage automatique. *Revue interdisciplinaire de l'information, des connaissances et de la gestion* [en ligne]. 2019, 14, 45–76. Dostupné z: doi:10.28945/4184
- [76] C. Huang, A. Zappone, G. C. Alexandropoulos, M. Debbah, and C. Yuen, "Reconfigurable Intelligent Surfaces for Energy Efficiency in Wireless Communication," *IEEE Transactions on Wireless Communications*, vol. 18, no. 8, pp. 4157–4170, Aug. 2019, doi: 10.1109/TWC.2019.2922609.
- [77] . D. Renzo *et al.*, "Smart Radio Environments Empowered by Reconfigurable Intelligent Surfaces: How It Works, State of Research, and The Road Ahead," *IEEE Journal on Selected Areas in Communications*, vol. 38, no. 11, pp. 2450–2525, Nov. 2020, doi: 10.1109/JSAC.2020.3007211.
- [78] C. Liaskos, S. Nie, A. Tsioliaridou, A. Pitsillides, S. Ioannidis, and I. Akyildiz, "A New Wireless Communication Paradigm through Software-Controlled Metasurfaces," *IEEE Communications Magazine*, vol. 56, no. 9, pp. 162–169, Sep. 2018, doi: 10.1109/MCOM.2018.1700659.
- [79] Q. Wu and R. Zhang, "Towards Smart and Reconfigurable Environment: Intelligent Reflecting Surface Aided Wireless Network," *IEEE Communications Magazine*, vol. 58, no. 1, pp. 106–112, Jan. 2020, doi: 10.1109/MCOM.001.1900107.
- [80] N. S. Perović, M. Di Renzo, and M. F. Flanagan, "Channel Capacity Optimization Using Reconfigurable Intelligent Surfaces in Indoor mmWave Environments," *arXiv:1910.14310 [cs, eess, math]*, Feb. 2020. [Online]. Available: <http://arxiv.org/abs/1910.14310>
- [81] C. Huang *et al.*, "Holographic MIMO Surfaces for 6G Wireless Networks: Opportunities, Challenges, and Trends," *IEEE Wireless Communications*, vol. 27, no. 5, pp. 118–125, Oct. 2020, doi: 10.1109/MWC.001.1900534.



- [82] M. Di Renzo *et al.*, "Reconfigurable Intelligent Surfaces vs. Relaying: Differences, Similarities, and Performance Comparison," *IEEE Open Journal of the Communications Society*, vol. 1, pp. 798–807, 2020, doi: 10.1109/OJCOMS.2020.3002955.
- [83] E. Björnson, Ö. Özdogan, and E. G. Larsson, "Intelligent Reflecting Surface Versus Decode-and-Forward: How Large Surfaces are Needed to Beat Relaying?," *IEEE Wireless Communications Letters*, vol. 9, no. 2, pp. 244–248, Feb. 2020, doi: 10.1109/LWC.2019.2950624.
- [84] S. Zhou, W. Xu, K. Wang, M. Di Renzo, and M.-S. Alouini, "Spectral and Energy Efficiency of IRS-Assisted MISO Communication With Hardware Impairments," *IEEE Wireless Communications Letters*, vol. 9, no. 9, pp. 1366–1369, Sep. 2020, doi: 10.1109/LWC.2020.2990431.
- [85] T. Bai, C. Pan, Y. Deng, M. El Kashlan, A. Nallanathan, and L. Hanzo, "Latency Minimization for Intelligent Reflecting Surface Aided Mobile Edge Computing," *IEEE Journal on Selected Areas in Communications*, vol. 38, no. 11, pp. 2666–2682, Nov. 2020, doi: 10.1109/JSAC.2020.3007035.
- [86] Y. Cao and T. Lv, "Sum Rate Maximization for Reconfigurable Intelligent Surface Assisted Device-to-Device Communications," *arXiv:2001.03344 [cs, eess, math]*, Jan. 2020. [Online]. Available: <http://arxiv.org/abs/2001.03344>
- [87] L. Yang, J. Yang, W. Xie, M. O. Hasna, T. Tsiftsis, and M. D. Renzo, "Secrecy Performance Analysis of RIS-Aided Wireless Communication Systems," *IEEE Transactions on Vehicular Technology*, vol. 69, no. 10, pp. 12296–12300, Oct. 2020, doi: 10.1109/TVT.2020.3007521.
- [88] C. Pan *et al.*, "Multicell MIMO Communications Relying on Intelligent Reflecting Surfaces," *IEEE Transactions on Wireless Communications*, vol. 19, no. 8, pp. 5218–5233, Aug. 2020, doi: 10.1109/TWC.2020.2990766.
- [89] Q. Wu and R. Zhang, "Weighted Sum Power Maximization for Intelligent Reflecting Surface Aided SWIPT," *IEEE Wireless Communications Letters*, vol. 9, no. 5, pp. 586–590, May 2020, doi: 10.1109/LWC.2019.2961656.
- [90] Y. Tang, G. Ma, H. Xie, J. Xu, and X. Han, "Joint Transmit and Reflective Beamforming Design for IRS-Assisted Multiuser MISO SWIPT Systems," in *ICC 2020 - 2020 IEEE International Conference on Communications (ICC)*, Jun. 2020, pp. 1–6. doi: 10.1109/ICC40277.2020.9148892.
- [91] C. Pan *et al.*, "Intelligent Reflecting Surface Aided MIMO Broadcasting for Simultaneous Wireless Information and Power Transfer," *IEEE Journal on Selected Areas in Communications*, vol. 38, no. 8, pp. 1719–1734, Aug. 2020, doi: 10.1109/JSAC.2020.3000802.
- [92] Q. Wu and R. Zhang, "Joint Active and Passive Beamforming Optimization for Intelligent Reflecting Surface Assisted SWIPT Under QoS Constraints," *IEEE Journal on Selected Areas in Communications*, vol. 38, no. 8, pp. 1735–1748, Aug. 2020, doi: 10.1109/JSAC.2020.3000807.
- [93] S. Li, B. Duo, X. Yuan, Y.-C. Liang, and M. Di Renzo, "Reconfigurable Intelligent Surface Assisted UAV Communication: Joint Trajectory Design and Passive Beamforming," *IEEE Wireless Communications Letters*, vol. 9, no. 5, pp. 716–720, May 2020, doi: 10.1109/LWC.2020.2966705.
- [94] D. Ma, M. Ding, and M. Hassan, "Enhancing Cellular Communications for UAVs via Intelligent Reflective Surface," in *2020 IEEE Wireless Communications and Networking Conference (WCNC)*, May 2020, pp. 1–6. doi: 10.1109/WCNC45663.2020.9120632.
- [95] A. U. Makarfi, K. M. Rabie, O. Kaiwartya, O. S. Badarneh, X. Li, and R. Kharel, "Reconfigurable Intelligent Surface Enabled IoT Networks in Generalized Fading Channels," in *ICC 2020 - 2020 IEEE International Conference on Communications (ICC)*, Jun. 2020, pp. 1–6. doi: 10.1109/ICC40277.2020.9148610.
- [96] M. Song, C. Xin, Y. Zhao, and X. Cheng, "Dynamic spectrum access: from cognitive radio to network radio," *IEEE Wireless Communications*, vol. 19, no. 1, pp. 23–29, Feb. 2012, doi: 10.1109/MWC.2012.6155873.
- [97] Y. Zhang, J. Zheng, and H.-H. Chen, *Cognitive Radio Networks: Architectures, Protocols, and Standards*. CRC Press, 2016.
- [98] M. S. Gupta and K. Kumar, "Progression on spectrum sensing for cognitive radio networks: A survey, classification, challenges and future research issues," *Journal of Network and Computer Applications*, vol. 143, pp. 47–76, Oct. 2019, doi: 10.1016/j.jnca.2019.06.005.
- [99] M. Nekovee, "Cognitive Radio Access to TV White Spaces: Spectrum Opportunities, Commercial Applications and Remaining Technology Challenges," in *2010 IEEE Symposium on New Frontiers in Dynamic Spectrum (DySPAN)*, Apr. 2010, pp. 1–10. doi: 10.1109/DYSPAN.2010.5457902.



- [100] "Cognitive radio in the context of internet of things using a novel future internet architecture called NovaGenesis - ScienceDirect." <https://www.sciencedirect.com/science/article/abs/pii/S004579061630180X>.
- [101] J. Mitola, "The software radio architecture," *IEEE Communications Magazine*, vol. 33, no. 5, pp. 26–38, May 1995, doi: 10.1109/35.393001.
- [102] J. I. Mitola, "Cognitive radio. An integrated agent architecture for software defined radio.," p. 1, 2002.
- [103] "A survey of clustering algorithms for cognitive radio ad hoc networks | SpringerLink." <https://link.springer.com/article/10.1007/s11276-016-1417-6>.
- [104] A. Ali and W. Hamouda, "Advances on Spectrum Sensing for Cognitive Radio Networks: Theory and Applications," *IEEE Communications Surveys Tutorials*, vol. 19, no. 2, pp. 1277–1304, 2017, doi: 10.1109/COMST.2016.2631080.
- [105] A. Ghasemi and E. S. Sousa, "Spectrum sensing in cognitive radio networks: requirements, challenges and design trade-offs," *IEEE Communications Magazine*, vol. 46, no. 4, pp. 32–39, Apr. 2008, doi: 10.1109/MCOM.2008.4481338.
- [106] F. Hou and J. Huang, "Dynamic Channel Selection in Cognitive Radio Network with Channel Heterogeneity," in *2010 IEEE Global Telecommunications Conference GLOBECOM 2010*, Dec. 2010, pp. 1–6. doi: 10.1109/GLOCOM.2010.5683964.
- [107] "Resource Allocation for Cognitive Small Cell Networks: A Cooperative Bargaining Game Theoretic Approach | IEEE Journals & Magazine | IEEE Xplore." <https://ieeexplore.ieee.org/document/7050288>.
- [108] "Spectrum sharing in Cognitive Radio using game theory | IEEE Conference Publication | IEEE Xplore." <https://ieeexplore.ieee.org/document/6514449>.
- [109] S. S. Oyewobi and G. P. Hancke, "A survey of cognitive radio handoff schemes, challenges and issues for industrial wireless sensor networks (CR-IWSN)," *Journal of Network and Computer Applications*, vol. 97, pp. 140–156, Nov. 2017, doi: 10.1016/j.jnca.2017.08.016.
- [110] G. I. Tsiropoulos, O. A. Dobre, M. H. Ahmed, and K. E. Baddour, "Radio Resource Allocation Techniques for Efficient Spectrum Access in Cognitive Radio Networks," *IEEE Communications Surveys Tutorials*, vol. 18, no. 1, pp. 824–847, 2016, doi: 10.1109/COMST.2014.2362796.
- [111] I. Christian, S. Moh, I. Chung, and J. Lee, "Spectrum mobility in cognitive radio networks," *IEEE Communications Magazine*, vol. 50, no. 6, pp. 114–121, Jun. 2012, doi: 10.1109/MCOM.2012.6211495.
- [112] H. Sun, A. Nallanathan, C.-X. Wang, and Y. Chen, "Wideband spectrum sensing for cognitive radio networks: a survey," *IEEE Wireless Communications*, vol. 20, no. 2, pp. 74–81, Apr. 2013, doi: 10.1109/MWC.2013.6507397.
- [113] A. Naeem, M. H. Rehmani, Y. Saleem, I. Rashid, and N. Crespi, "Network Coding in Cognitive Radio Networks: A Comprehensive Survey," *IEEE Communications Surveys Tutorials*, vol. 19, no. 3, pp. 1945–1973, 2017, doi: 10.1109/COMST.2017.2661861.
- [114] M. López-Benítez, "Overview of Recent Applications of Cognitive Radio in Wireless Communication Systems," in *Handbook of Cognitive Radio*, W. Zhang, Ed. Singapore: Springer, 2018, pp. 1–32. doi: 10.1007/978-981-10-1389-8_59-1.
- [115] H. Al-Hmood, "Performance analysis of energy detector over different generalised wireless channels based spectrum sensing in cognitive radio," *undefined*, 2015. [Online]. Available: <https://www.semanticscholar.org/paper/Performance-analysis-of-energy-detector-over-based-Al-Hmood/4bab068f709b4b2fbe90224c9ece77d1008c3a5a>
- [116] Y. Abdi Mahmoudaliloo, "Cooperative spectrum sensing schemes for future dynamic spectrum access infrastructures," *Jyväskylä studies in computing*, no. 238, 2016. [Online]. Available: <https://jyx.jyu.fi/handle/123456789/49957>
- [117] F. Haider, C.-X. Wang, B. Ai, H. Haas, and E. Hepsaydir, "Spectral/Energy Efficiency Tradeoff of Cellular Systems With Mobile Femtocell Deployment," *IEEE Transactions on Vehicular Technology*, vol. 65, no. 5, pp. 3389–3400, May 2016, doi: 10.1109/TVT.2015.2443046.
- [118] "Spectrum Sensing Methods and Their Performance | SpringerLink." https://link.springer.com/referenceworkentry/10.1007%2F978-981-10-1389-8_6-1.
- [119] M. A. Abdulsattar, "Energy Detection Technique for Spectrum Sensing in Cognitive Radio: A Survey," *International journal of Computer Networks & Communications*, vol. 4, no. 5, pp. 223–242.
- [120] P. Rawat, K. D. Singh, and J. M. Bonnin, "Cognitive radio for M2M and Internet of Things: A survey," *Computer Communications*, vol. 94, pp. 1–29, Nov. 2016, doi: 10.1016/j.comcom.2016.07.012.



- [121] "Cognitive radio networks and spectrum sharing - ScienceDirect." <https://www.sciencedirect.com/science/article/pii/B9780123982810000132>.
- [122] "Cooperative Spectrum Handovers in Cognitive Radio Networks | SpringerLink." https://link.springer.com/chapter/10.1007/978-3-030-15416-5_1.
- [123] "Optimal Multiband Joint Detection for Spectrum Sensing in Cognitive Radio Networks." <https://ieeexplore.ieee.org/abstract/document/4668431/>.
- [124] "Standardization Progress | SpringerLink." https://link.springer.com/chapter/10.1007/978-3-319-15768-9_7.
- [125] M. A. Davenport, J. N. Laska, J. R. Treichler, and R. G. Baraniuk, "The Pros and Cons of Compressive Sensing for Wideband Signal Acquisition: Noise Folding versus Dynamic Range," *IEEE Transactions on Signal Processing*, vol. 60, no. 9, pp. 4628–4642, Sep. 2012, doi: 10.1109/TSP.2012.2201149.
- [126] Z. Tian and G. B. Giannakis, "A Wavelet Approach to Wideband Spectrum Sensing for Cognitive Radios," in *2006 1st International Conference on Cognitive Radio Oriented Wireless Networks and Communications*, Jun. 2006, pp. 1–5. doi: 10.1109/CROWNCOM.2006.363459.
- [127] "A comparative study of spectrum awareness techniques for cognitive radio oriented wireless networks - ScienceDirect." <https://www.sciencedirect.com/science/article/abs/pii/S1874490712000687>.
- [128] S. K. Sharma, T. E. Bogale, S. Chatzinotas, B. Ottersten, L. B. Le, and X. Wang, "Cognitive Radio Techniques Under Practical Imperfections: A Survey," *IEEE Communications Surveys Tutorials*, vol. 17, no. 4, pp. 1858–1884, 2015, doi: 10.1109/COMST.2015.2452414.
- [129] S. M. Mishra, A. Sahai, and R. W. Brodersen, "Cooperative Sensing among Cognitive Radios," in *2006 IEEE International Conference on Communications*, Jun. 2006, vol. 4, pp. 1658–1663. doi: 10.1109/ICC.2006.254957.
- [130] "IS: Intelligent spectrum sensing scheme for cognitive radio networks | EURASIP Journal on Wireless Communications and Networking | Full Text." <https://jwcn-urasipjournals.springeropen.com/articles/10.1186/1687-1499-2013-26> (accessed Dec. 29, 2021).
- [131] D. Cabric, S. M. Mishra, and R. W. Brodersen, "Implementation issues in spectrum sensing for cognitive radios," in *Conference Record of the Thirty-Eighth Asilomar Conference on Signals, Systems and Computers, 2004*, Nov. 2004, vol. 1, pp. 772–776 Vol.1. doi: 10.1109/ACSSC.2004.1399240.
- [132] "Reconfigurable filter implementation of a matched-filter based spectrum sensor for Cognitive Radio systems | IEEE Conference Publication | IEEE Xplore." <https://ieeexplore.ieee.org/document/5938101> (accessed Dec. 29, 2021).
- [133] S. Shobana, R. Saravanan, and R. Muthaiah, "Matched Filter Based Spectrum Sensing on Cognitive Radio for OFDM WLANs," *undefined*, 2013, Accessed: Dec. 29, 2021. [Online]. Available: <https://www.semanticscholar.org/paper/Matched-Filter-Based-Spectrum-Sensing-on-Cognitive-Shobana-Saravanan/711f53647851baefc949c2cad0e71e03b28110fc>
- [134] W. A. Gardner and C. M. Spooner, "Signal interception: performance advantages of cyclic-feature detectors," *IEEE Transactions on Communications*, vol. 40, no. 1, pp. 149–159, Jan. 1992, doi: 10.1109/26.126716.
- [135] J. Proakis and D. Manolakis, "Digital Signal Processing: Principles, Algorithms, and Applications," *undefined*, 1992, Accessed: Dec. 29, 2021. [Online]. Available: <https://www.semanticscholar.org/paper/Digital-Signal-Processing%3A-Principles%2C-Algorithms%2C-Proakis-Manolakis/cb5206cf9ee6a4eb5d1f2deb9736883850fc310>
- [136] S. AparnaP and M. Jayasheela, "Cyclostationary Feature Detection in Cognitive Radio using Different Modulation Schemes," 2012, doi: 10.5120/7472-0517.
- [137] "Higher-Order Cyclostationarity Detection for Spectrum Sensing | EURASIP Journal on Wireless Communications and Networking | Full Text." <https://jwcn-urasipjournals.springeropen.com/articles/10.1155/2010/721695> (accessed Dec. 29, 2021).
- [138] P. D. Sutton, K. E. Nolan, and L. E. Doyle, "Cyclostationary Signatures in Practical Cognitive Radio Applications," *IEEE Journal on Selected Areas in Communications*, vol. 26, no. 1, pp. 13–24, Jan. 2008, doi: 10.1109/JSAC.2008.080103.
- [139] F. K. Jondral, "Cognitive Radio: A Communications Engineering View," *IEEE Wireless Communications*, vol. 14, no. 4, pp. 28–33, Aug. 2007, doi: 10.1109/MWC.2007.4300980.



- [140] A. Tkachenko, D. Cabric, and R. W. Brodersen, "Cyclostationary Feature Detector Experiments Using Reconfigurable BEE2," in *2007 2nd IEEE International Symposium on New Frontiers in Dynamic Spectrum Access Networks*, Apr. 2007, pp. 216–219. doi: 10.1109/DYSPAN.2007.36.
- [141] H.-S. Chen, W. Gao, and D. G. Daut, "Spectrum Sensing Using Cyclostationary Properties and Application to IEEE 802.22 WRAN," in *IEEE GLOBECOM 2007 - IEEE Global Telecommunications Conference*, Nov. 2007, pp. 3133–3138. doi: 10.1109/GLOCOM.2007.593.
- [142] "Optimal and Sub-Optimal Spectrum Sensing of OFDM Signals in Known and Unknown Noise Variance | IEEE Journals & Magazine | IEEE Xplore." <https://ieeexplore.ieee.org/document/5701684> (accessed Dec. 29, 2021).
- [143] "On the Energy Detection of Unknown Signals Over Fading Channels | IEEE Journals & Magazine | IEEE Xplore." <https://ieeexplore.ieee.org/document/4063501> (accessed Dec. 29, 2021).
- [144] K. Tan, K. Kim, Y. Xin, S. Rangarajan, and P. Mohapatra, "RECOG: A Sensing-Based Cognitive Radio System with Real-Time Application Support," *IEEE Journal on Selected Areas in Communications*, vol. 31, no. 11, pp. 2504–2516, Nov. 2013, doi: 10.1109/JSAC.2013.131132.
- [145] H. Urkowitz, "Energy detection of unknown deterministic signals," *Proceedings of the IEEE*, vol. 55, no. 4, pp. 523–531, Apr. 1967, doi: 10.1109/PROC.1967.5573.
- [146] T. S. Shehata and M. El-Tanany, "A novel adaptive structure of the energy detector applied to cognitive radio networks," in *2009 11th Canadian Workshop on Information Theory*, May 2009, pp. 95–98. doi: 10.1109/CWIT.2009.5069530.
- [147] "Performance Analysis of Energy Detection-Based Spectrum Sensing over Fading Channels | IEEE Conference Publication | IEEE Xplore." <https://ieeexplore.ieee.org/document/5600868> (accessed Dec. 29, 2021).
- [148] L. S. CARDOSO, M. DEBBAH, S. LASAULCE, M. KOBAYASHI, and J. PALICOT, "Spectrum Sensing in Cognitive Radio Networks," in *Cognitive Radio Networks*, CRC Press, 2010.
- [149] J. Ma, G. Zhao, and Y. Li, "Soft Combination and Detection for Cooperative Spectrum Sensing in Cognitive Radio Networks," *IEEE Transactions on Wireless Communications*, vol. 7, no. 11, pp. 4502–4507, Nov. 2008, doi: 10.1109/T-WC.2008.070941.
- [150] E. Visotsky, S. Kuffner, and R. Peterson, "On collaborative detection of TV transmissions in support of dynamic spectrum sharing," in *First IEEE International Symposium on New Frontiers in Dynamic Spectrum Access Networks, 2005. DySPAN 2005.*, Nov. 2005, pp. 338–345. doi: 10.1109/DYSPAN.2005.1542650.
- [151] S. Yan, "Collaborative spectrum sensing for cognitive radio networks by using 1-bit compressed sensing," in *Proceedings of the 4th International Conference on Communication and Information Processing*, New York, NY, USA, Nov. 2018, pp. 289–293. doi: 10.1145/3290420.3290479.
- [152] "Collaborative Cyclostationary Spectrum Sensing for Cognitive Radio Systems | IEEE Journals & Magazine | IEEE Xplore." <https://ieeexplore.ieee.org/document/5072368> (accessed Dec. 29, 2021).
- [153] "Distributed Consensus-Based Cooperative Spectrum Sensing in Cognitive Radio Mobile Ad Hoc Networks | SpringerLink." https://link.springer.com/chapter/10.1007/978-1-4419-6172-3_1 (accessed Dec. 29, 2021).
- [154] F. C. Ribeiro, M. L. R. de Campos, and S. Werner, "Distributed cooperative spectrum sensing with adaptive combining," in *2012 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, Mar. 2012, pp. 3557–3560. doi: 10.1109/ICASSP.2012.6288685.
- [155] C. Liu, A. Qi, P. Zhang, L. Bu, and K. Long, "Wideband spectrum detection based on Compressed Sensing in cooperative cognitive radio networks," in *2013 8th International Conference on Communications and Networking in China (CHINACOM)*, Aug. 2013, pp. 638–642. doi: 10.1109/ChinaCom.2013.6694671.
- [156] F. Zeng, Z. Tian, and C. Li, "Distributed Compressive Wideband Spectrum Sensing in Cooperative Multi-Hop Cognitive Networks," in *2010 IEEE International Conference on Communications*, May 2010, pp. 1–5. doi: 10.1109/ICC.2010.5502793.
- [157] F. Zeng, C. Li, and Z. Tian, "Distributed Compressive Spectrum Sensing in Cooperative Multihop Cognitive Networks," *IEEE Journal of Selected Topics in Signal Processing*, vol. 5, no. 1, pp. 37–48, Feb. 2011, doi: 10.1109/JSTSP.2010.2055037.
- [158] "Cooperative Shared Spectrum Sensing for Dynamic Cognitive Radio Networks | IEEE Conference Publication | IEEE Xplore." <https://ieeexplore.ieee.org/document/5198862> (accessed Dec. 29, 2021).



- [159] Q. Zhang, J. Jia, and J. Zhang, "Cooperative relay to improve diversity in cognitive radio networks," *IEEE Communications Magazine*, vol. 47, no. 2, pp. 111–117, Feb. 2009, doi: 10.1109/MCOM.2009.4785388.
- [160] D. Rawat and G. Yan, "Spectrum Sensing Methods and Dynamic Spectrum Sharing in Cognitive Radio Networks: A Survey," *undefined*, 2011, Accessed: Dec. 29, 2021. [Online]. Available: <https://www.semanticscholar.org/paper/Spectrum-Sensing-Methods-and-Dynamic-Spectrum-in-A-Rawat-Yan/de42cde125237a90f3df999293002bbd1a61c556>
- [161] M. Saber, A. E. Rharras, R. Saadane, H. K. Aroussi, and M. Wahbi, "Artificial Neural Networks, Support Vector Machine And Energy Detection For Spectrum Sensing Based On Real Signals," *International Journal of Communication Networks and Information Security (IJCNIS)*, vol. 11, no. 1, Art. no. 1, Apr. 2019, Accessed: May 27, 2021. [Online]. Available: <https://www.ijcnis.org/index.php/ijcnis/article/view/3718>
- [162] C. Cortes and V. Vapnik, "Support-vector networks," *Mach Learn*, vol. 20, no. 3, pp. 273–297, Sep. 1995, doi: 10.1007/BF00994018.
- [163] K. M. Thilina, K. W. Choi, N. Saquib, and E. Hossain, "Machine Learning Techniques for Cooperative Spectrum Sensing in Cognitive Radio Networks," *IEEE Journal on Selected Areas in Communications*, vol. 31, no. 11, pp. 2209–2221, Nov. 2013, doi: 10.1109/JSAC.2013.131120.
- [164] J. G. Proakis and M. Salehi, *Communication Systems Engineering*. Prentice Hall, 1994.
- [165] P. Prandoni and M. Vetterli, *Signal Processing for Communications*. Collection le savoir suisse, 2008.
- [166] "Wireless communications: past events and a future perspective | IEEE Journals & Magazine | IEEE Xplore." <https://ieeexplore.ieee.org/document/1006984> (accessed Dec. 29, 2021).
- [167] "A statistical discrete-time model for the WSSUS multipath channel | IEEE Journals & Magazine | IEEE Xplore." <https://ieeexplore.ieee.org/document/182598> (accessed Dec. 29, 2021).
- [168] J. P. Egan and J. P. Egan, *Signal Detection Theory and ROC-analysis*. Academic Press, 1975.
- [169] J. A. Swets, "Measuring the accuracy of diagnostic systems," *Science*, vol. 240, no. 4857, pp. 1285–1293, Jun. 1988, doi: 10.1126/science.3287615.
- [170] "Receiver-Operating Characteristic Analysis for Evaluating Diagnostic Tests and Predictive Models | Circulation." <https://www.ahajournals.org/doi/full/10.1161/CIRCULATIONAHA.105.594929> (accessed Dec. 29, 2021).
- [171] F. Provost and T. Fawcett, "Robust Classification for Imprecise Environments," *Machine Learning*, vol. 42, no. 3, pp. 203–231, Mar. 2001, doi: 10.1023/A:1007601015854.
- [172] F. J. Provost, T. Fawcett, and R. Kohavi, "The Case against Accuracy Estimation for Comparing Induction Algorithms," in *Proceedings of the Fifteenth International Conference on Machine Learning*, San Francisco, CA, USA, Jul. 1998, pp. 445–453.
- [173] J. Neyman and E. S. Pearson, *Joint Statistical Papers*. University of California Press.
- [174] "RTL2832U - REALTEK." <https://www.realtek.com/en/products/communications-network-ics/item/rtl2832u> (accessed Dec. 29, 2021).
- [175] "Rtl-sdr - rtl-sdr - Open Source Mobile Communications." <https://osmocom.org/projects/rtl-sdr/wiki/Rtl-sdr> (accessed Dec. 29, 2021).
- [176] L. V. Fausett, *Fundamentals of Neural Networks: Architectures, Algorithms And Applications: United States Edition*, US e. édition. Englewood Cliffs, NJ: Pearson, 1993.
- [177] "Neural network design | Guide books." <https://dl.acm.org/doi/abs/10.5555/249049> (accessed Dec. 27, 2021).
- [178] "Multi-Layer Perceptrons | SpringerLink." https://link.springer.com/chapter/10.1007/978-1-4471-5013-8_5 (accessed Dec. 27, 2021).
- [179] R. Reed and R. J. MarksII, *Neural Smithing: Supervised Learning in Feedforward Artificial Neural Networks*. MIT Press, 1999.
- [180] "Pattern recognition using neural networks | Guide books." <https://dl.acm.org/doi/abs/10.5555/262331> (accessed Dec. 27, 2021).
- [181] A. Victor and M. R. Ghalib, "Automatic detection and classification of skin cancer," *International Journal of Intelligent Engineering and Systems*, vol. 10, no. 3, pp. 444–451, 2017, doi: 10.22266/ijies2017.0630.50.



- [182] M. Koklu and I. A. Ozkan, "Skin Lesion Classification using Machine Learning Algorithms," *International Journal of Intelligent Systems and Applications in Engineering*, vol. 4, no. 5, pp. 285–289, 2017, doi: 10.18201/ijisae.2017534420.
- [183] J. R. Quinlan, "Induction of decision trees," *Machine Learning*, vol. 1, no. 1, pp. 81–106, 1986, doi: 10.1007/bf00116251.
- [184] B. HSSINA, A. MERBOUHA, H. EZZIKOURI, and M. ERRITALI, "A comparative study of decision tree ID3 and C4.5," *International Journal of Advanced Computer Science and Applications*, vol. 4, no. 2, 2014, doi: 10.14569/specialissue.2014.040203.
- [185] L. Rokach, "DECISION TREES," *Research gate*, no. November, pp. 0–11, 2014, doi: 10.1007/0-387-25465-X.
- [186] "A comparative assessment of classification methods - ScienceDirect." <https://www.sciencedirect.com/science/article/abs/pii/S0167923602001100> (accessed Oct. 06, 2021).
- [187] "Miscellaneous Clustering Methods," in *Cluster Analysis*, John Wiley & Sons, Ltd, 2011, pp. 215–255. doi: 10.1002/9780470977811.ch8.
- [188] "The Nature of Statistical Learning Theory - Vladimir Vapnik - Google Livres." https://books.google.fr/books?hl=fr&lr=&id=EggACAAAQBAJ&oi=fnd&pg=PR7&dq=The+nature+of+statistical+learning+theory&ots=g4G0ky4V88&sig=Hb96AxRRpVeC05KkQOKim_mB4QQ#v=onepage&q=The%20nature%20of%20statistical%20learning%20theory&f=false (accessed Oct. 07, 2021).
- [189] B. Schölkopf, D. of the M. P. I. for I. in T. G. P. for M. L. B. Schölkopf, rnhard Schölkopf, A. J. Smola, F. Bach, and M. D. of the M. P. I. for B. C. in T. G. P. B. Scholkopf, *Learning with Kernels: Support Vector Machines, Regularization, Optimization, and Beyond*. MIT Press, 2002.
- [190] R. Fletcher, *Practical Methods of Optimization*. John Wiley & Sons, 2013.
- [191] J. Nayak, B. Naik, and H. S. Behera, "A Comprehensive Survey on Support Vector Machine in Data Mining Tasks : Applications & Challenges," *International Journal of Database Theory and Application*, vol. 8, no. 1, pp. 169–186, Feb. 2015.
- [192] I. F. Akyildiz, W. Lee, M. C. Vuran, and S. Mohanty, "A survey on spectrum management in cognitive radio networks," *IEEE Communications Magazine*, vol. 46, no. 4, pp. 40–48, Apr. 2008, doi: 10.1109/MCOM.2008.4481339.
- [193] D. J. Thomson, "Spectrum estimation and harmonic analysis," *Proceedings of the IEEE*, vol. 70, no. 9, pp. 1055–1096, Sep. 1982, doi: 10.1109/PROC.1982.12433.
- [194] F. J. Harris, "On the use of windows for harmonic analysis with the discrete Fourier transform," *Proceedings of the IEEE*, vol. 66, no. 1, pp. 51–83, Jan. 1978, doi: 10.1109/PROC.1978.10837.
- [195] K. Hornik, M. Stinchcombe, and H. White, "Multilayer feedforward networks are universal approximators," *Neural Networks*, vol. 2, no. 5, pp. 359–366, Jan. 1989, doi: 10.1016/0893-6080(89)90020-8.
- [196] I. Goodfellow, Y. Bengio, and A. Courville, *Deep Learning*. MIT Press, 2016.
- [197] C. Szegedy *et al.*, "Going Deeper With Convolutions," 2015, pp. 1–9. Accessed: Dec. 28, 2021. [Online]. Available: https://www.cv-foundation.org/openaccess/content_cvpr_2015/html/Szegedy_Going_Deeper_With_2015_CVPR_paper.html
- [198] A. Krizhevsky, I. Sutskever, and G. E. Hinton, "ImageNet classification with deep convolutional neural networks," *Commun. ACM*, vol. 60, no. 6, pp. 84–90, May 2017, doi: 10.1145/3065386.
- [199] "In the Matter of Unlicensed Operation in the TV Broadcast Bands, Additional Spectrum for Unlicensed Devices Below 900 MHz and in the 3 GHz Band," *Federal Communications Commission*, Dec. 18, 2015. <https://www.fcc.gov/document/matter-unlicensed-operation-tv-broadcast-bands-additional> (accessed Oct. 06, 2021).
- [200] "ECO Documentation." <https://docdb.cept.org/document/267> (accessed Oct. 06, 2021).
- [201] K. J. R. Liu and B. Wang, *Cognitive Radio Networking and Security: A Game-Theoretic View*. Cambridge University Press, 2010.
- [202] "Wireless Communications -- Principles and Practice, Second Edition. (The Book End) - Document - Gale Academic OneFile." <https://go.gale.com/ps/i.do?id=GALE%7CA97115718&sid=googleScholar&v=2.1&it=r&linkaccess=>



- abs&issn=01926225&p=AONE&sw=w&userGroupName=anon%7E42126ea2 (accessed Oct. 06, 2021).
- [203] O. Tsilipakos *et al.*, "Toward Intelligent Metasurfaces: The Progress from Globally Tunable Metasurfaces to Software-Defined Metasurfaces with an Embedded Network of Controllers," *Advanced Optical Materials*, vol. 8, no. 17, p. 2000783, 2020, doi: 10.1002/adom.202000783.
 - [204] Ö. Özdoğan, E. Björnson, and E. G. Larsson, "Intelligent Reflecting Surfaces: Physics, Propagation, and Pathloss Modeling," *IEEE Wireless Communications Letters*, vol. 9, no. 5, pp. 581–585, May 2020, doi: 10.1109/LWC.2019.2960779.
 - [205] E. Björnson, Ö. Özdoğan, and E. G. Larsson, "Reconfigurable Intelligent Surfaces: Three Myths and Two Critical Questions," *IEEE Communications Magazine*, vol. 58, no. 12, pp. 90–96, Dec. 2020, doi: 10.1109/MCOM.001.2000407.
 - [206] J. N. Laneman, D. N. C. Tse, and G. W. Wornell, "Cooperative diversity in wireless networks: Efficient protocols and outage behavior," *IEEE Transactions on Information Theory*, vol. 50, no. 12, pp. 3062–3080, Dec. 2004, doi: 10.1109/TIT.2004.838089.
 - [207] G. Farhadi and N. C. Beaulieu, "On the ergodic capacity of multi-hop wireless relaying systems," *IEEE Transactions on Wireless Communications*, vol. 8, no. 5, pp. 2286–2291, May 2009, doi: 10.1109/TWC.2009.080818.
 - [208] "A Promising Technology for 6G Wireless Networks: Intelligent Reflecting Surface." <http://www.infocomm-journal.com/jcin/EN/abstract/abstract171197.shtml> (accessed Jul. 25, 2021).
 - [209] E. Björnson and L. Sanguinetti, "Power Scaling Laws and Near-Field Behaviors of Massive MIMO and Intelligent Reflecting Surfaces," *IEEE Open Journal of the Communications Society*, vol. 1, pp. 1306–1324, 2020, doi: 10.1109/OJCOMS.2020.3020925.
 - [210] Q.-U.-A. Nadeem, A. Kammoun, A. Chaaban, M. Debbah, and M.-S. Alouini, "Asymptotic Max-Min SINR Analysis of Reconfigurable Intelligent Surface Assisted MISO Systems," *IEEE Transactions on Wireless Communications*, vol. 19, no. 12, pp. 7748–7764, Dec. 2020, doi: 10.1109/TWC.2020.2986438.
 - [211] Q. Wu and R. Zhang, "Intelligent Reflecting Surface Enhanced Wireless Network via Joint Active and Passive Beamforming," *IEEE Transactions on Wireless Communications*, vol. 18, no. 11, pp. 5394–5409, Nov. 2019, doi: 10.1109/TWC.2019.2936025.
 - [212] M. N. Khormuji and E. G. Larsson, "Cooperative transmission based on decode-and-forward relaying with partial repetition coding," *IEEE Transactions on Wireless Communications*, vol. 8, no. 4, pp. 1716–1725, Apr. 2009, doi: 10.1109/TWC.2009.070674.
 - [213] E. Björnson and L. Sanguinetti, "Demystifying the Power Scaling Law of Intelligent Reflecting Surfaces and Metasurfaces," in *2019 IEEE 8th International Workshop on Computational Advances in Multi-Sensor Adaptive Processing (CAMSAP)*, Dec. 2019, pp. 549–553. doi: 10.1109/CAMSAP45676.2019.9022637.
 - [214] "Intelligent Reflecting Surface Versus Decode-and-Forward: How Large Surfaces are Needed to Beat Relaying? | IEEE Journals & Magazine | IEEE Xplore." <https://ieeexplore.ieee.org/abstract/document/8888223> (accessed Aug. 11, 2021).
 - [215] "Modeling and Analysis of K-Tier Downlink Heterogeneous Cellular Networks." <https://ieeexplore.ieee.org/abstract/document/6171996/> (accessed Aug. 23, 2021).
 - [216] "On peak versus average interference power constraints for protecting primary users in cognitive radio networks | IEEE Journals & Magazine | IEEE Xplore." <https://ieeexplore.ieee.org/abstract/document/4907474> (accessed Aug. 23, 2021).
 - [217] R. Zhang, "Optimal Power Control over Fading Cognitive Radio Channel by Exploiting Primary User CSI," in *IEEE GLOBECOM 2008 - 2008 IEEE Global Telecommunications Conference*, Nov. 2008, pp. 1–5. doi: 10.1109/GLOCOM.2008.ECP.184.
 - [218] X. Kang, R. Zhang, Y.-C. Liang, and H. K. Garg, "Optimal Power Allocation Strategies for Fading Cognitive Radio Channels with Primary User Outage Constraint," *IEEE Journal on Selected Areas in Communications*, vol. 29, no. 2, pp. 374–383, Feb. 2011, doi: 10.1109/JSAC.2011.110210.
 - [219] 3GPP, "Further advancements for E-UTRA physical layer aspects (Release 9)," Mar. 2010. <https://portal.3gpp.org/desktopmodules/Specifications/SpecificationDetails.aspx?specificationId=2493>



Université Mohammed V - Faculté des Sciences de Rabat
Laboratoire de recherche en informatique et télécommunications (LRIT)
Discipline : Informatique
Spécialité : Réseaux et télécommunications



Contribution à l'optimisation des ressources spectrales dans la 6G

EL METTITI ABDERRAHME
THESE DE DOCTORAT

■ Rabat, 2022 ■

Notre travail propose des techniques contribuant à l'amélioration des services du futur réseau 6G. Au moyen d'un auto-encodeur empilé pour compresser les jeux de données et de perceptrons multicouches, nous proposons une technique d'apprentissage profond pour la prédiction de formation de faisceaux analogiques à ondes millimétriques. Nous présentons également un système hybride pour la gestion dynamique du spectre dans les environnements radio cognitive à l'aide des surfaces intelligentes reconfigurables pour un accès optimisé aux ressources spectrales. Notre démonstration repose sur un corpus de descriptions allant de la première génération à la 5G, complété par les apports de la 5G nouvelle radio, qui explique les limites technologiques désormais atteintes et que la 6G doit dépasser. Autant de défis identifiés que nous exposons en détail concernant différents domaines dans une vaste vision structurelle de la 6G ; s'il reste de nombreux défis, des avancées significatives sont discutées ici. Nos travaux ont fait l'objet de trois publications qui ont fait part de nos résultats encourageants.

Mots clefs : Machine learning, Deep learning, Auto-encodeur, SIR, GDS, RNA, Perceptron multicouche

Our work proposes techniques that contribute to the improvement of services in the future 6G network. Using a stacked auto-encoder to compress datasets and multi-layer perceptrons, we propose a deep learning technique for the prediction of millimeter wave analog beamforms. We also present a hybrid system for dynamic spectrum management in cognitive radio environments using reconfigurable smart surfaces for optimized access to spectrum resources. Our demonstration is based on a corpus of descriptions from the first generation to 5G, completed by the contributions of the new radio 5G, which explains the technological limits now reached and that 6G must overcome. As many identified challenges that we detail concerning different domains in a broad structural vision of 6G; while many challenges remain, significant advances are discussed here. Our work has been published in three papers, and our results are encouraging.

Keywords: Machine learning, Deep learning, Auto-encoder, SIR, GDS, RNA, Multilayer perceptron

Notre travail propose des techniques contribuant à l'amélioration des services du futur réseau 6G. Au moyen d'un auto-encodeur empilé pour compresser les jeux de données et de perceptrons multicouches, nous proposons une technique d'apprentissage profond pour la prédiction de formation de faisceaux analogiques à ondes millimétriques. Nous présentons également un système hybride pour la gestion dynamique du spectre dans les environnements radio cognitive à l'aide des surfaces intelligentes reconfigurables pour un accès optimisé aux ressources spectrales. Notre démonstration repose sur un corpus de descriptions allant de la première génération à la 5G, complété par les apports de la 5G nouvelle radio, qui explique les limites technologiques désormais atteintes et que la 6G doit dépasser. Autant de défis identifiés que nous exposons en détail concernant différents domaines dans une vaste vision structurelle de la 6G ; s'il reste de nombreux défis, des avancées significatives sont discutées ici. Nos travaux ont fait l'objet de trois publications qui ont fait part de nos résultats encourageants.

Mots clés : *Machine learning, Deep learning, Auto-encodeur, SIR, GDS, RNA, Perceptron multicouche*

Our work proposes techniques that contribute to the improvement of services in the future 6G network. Using a stacked auto-encoder to compress datasets and multi-layer perceptrons, we propose a deep learning technique for the prediction of milimeter wave analog beamforms. We also present a hybrid system for dynamic spectrum management in cognitive radio environments using reconfigurable smart surfaces for optimized access to spectrum resources. Our demonstration is based on a corpus of descriptions from the first generation to 5G, completed by the contributions of the new radio 5G, which explains the technological limits now reached and that 6G must overcome. As many identified challenges that we detail concerning different domains in a broad structural vision of 6G; while many challenges remain, significant advances are discussed here. Our work has been published in three papers, and our results are encouraging.

Keywords : *Machine learning, Deep learning, Auto-encoder, SIR, GDS, RNA, Multilayer perceptron*