

الجمهورية الجزائرية الديمقراطية الشعبية

REPUBLIQUE ALGERIENNE DEMOCRATIQUE ET POPULAIRE

وزارة التعليم العالي والبحث العلمي

Ministère de l'Enseignement Supérieur et de la Recherche Scientifique

جامعة أبي بكر بلقايد - تلمسان

Université Aboubakr Belkaïd – Tlemcen –

Faculté de TECHNOLOGIE



MEMOIRE

Présenté pour l'obtention du **diplôme** de **MASTER PROFETIONNEL**

En : Télécommunications

Spécialité : Télécommunications Optiques

Par : Larkoub Salima

Sujet :

**Simulation d'un système NOMA SCMA basé sur l'intelligence artificielle
pour les réseaux 5G.**

Soutenu publiquement, le 05 /06 / 2024, devant le jury composé de :

Mr.BOUACHA Abdelhafid	Professeur	Univ.Tlemcen	Président
Mme.BELGACEM Nassima	MCB	Univ.Tlemcen	Examinatrice
Mr.BENDIMERAD Yassine	MCA	Univ.Tlemcen	Directeur
Mr.GHOMRI Benmeziane Imad-ddine	Doctorant	Univ.Tlemcen	Co-Directeur

Année universitaire : 2023/2024



Remerciements

En préambule à ce mémoire, je remercie Allah qui m'aide et me donne la patience et le courage durant ces longues années d'étude.

C'est avec un immense plaisir que je réserve ces lignes en signe de gratitude et de remerciements à toutes les personnes qui ont contribué de près ou de loin à l'élaboration de ce projet.

Je remercie messieurs les membres de jury d'avoir accepté d'examiner et d'évaluer mon travail.

Tous mes remerciements s'adressent à mon directeur de mémoire

Mr. Bendimerad Yassine pour la confiance qu'il m'a accordée ainsi que son soutien moral. Les connaissances qu'il a su me transmettre, sa patience et ses judicieux conseils m'ont énormément aidé.

Tous mes remerciements s'adressent également à mon co-directeur de mémoire

Mr. Ghomri Benmeziare Imad-ddine pour son aide, ses remarques pertinentes, sa disponibilité, son écoute, ses conseils avisés et ses appréciations qui m'ont guidé dans mon travail et m'ont aidé à trouver des solutions pour mieux avancer.

Mes sincères remerciements vont également à tous mes professeurs durant ces cinq années à l'université, pour leurs conseils, leur aide et leur soutien.



Je dédie ce projet

À mes chers parents, ma sœur Ikram, mon frère Mustapha pour leur soutien indéfectible,

À mes professeurs ,pour leur enseignement précieux et leurs conseils avisés .

À mes chers collègues et amis de M2 TOP,

Votre soutien inconditionnel et votre amitié ont illuminé mon parcours.

À mes princesses Aridj et Ritedj, et à mon prince Anas,

Votre présence dans ma vie est un cadeau précieux.

Avec gratitude,

Salima

RÉSUMÉ

La technologie NOMA joue un rôle essentiel dans l'évolution des réseaux de communication sans fil, en particulier dans le contexte de la 5G. En autorisant le partage simultané de la même ressource spectrale par plusieurs utilisateurs, NOMA augmente l'efficacité spectrale et propose des performances plus élevées en termes de débit et de capacité par rapport aux techniques d'accès multiple traditionnelles. L'accès multiple par répartition par codes épars (SCMA) apparaît comme une solution prometteuse pour la 5G dans ce contexte. À la différence des schémas de multiplexage classiques, SCMA utilise une combinaison de séquences de codes afin de représenter les données des utilisateurs, ce qui permet d'améliorer l'efficacité spectrale et de résister davantage aux interférences. Cependant, la complexité de détection multi-utilisateur engendrée par le MPA (Message Passing Algorithm) représente un défi majeur à relever. Dans cette étude, une approche innovante utilisant l'intelligence artificielle (IA) est proposée afin de créer un système NOMA SCMA performant dans les réseaux 5G. Le but principal consiste à réduire la complexité de la détection tout en diminuant le taux d'erreur symbolique (SER). Les réseaux de neurones profonds (DNN) sont utilisés pour détecter les signaux dans un environnement NOMA SCMA spécifique à la 5G, selon la méthodologie proposée. La formation des DNNs consiste à détecter et à distinguer les signaux superposés provenant de divers utilisateurs, tout en tenant compte des caractéristiques particulières de la modulation et du canal.

Mots clés : NOMA, SCMA, 5G, Message Passing Algorithm (MPA), SER , Intelligence Artificielle (IA), Réseaux de neurones profonds (DNNs).

ABSTRACT

The NOMA technology plays a key role in the evolution of wireless communication networks, particularly in the context of 5G. By allowing multiple users to share the same spectral resource simultaneously, NOMA increases spectral efficiency and offers higher performance in terms of throughput and capacity than traditional multiple access techniques. Sparse Code Division Multiple Access (SCMA) appears to be a promising solution for 5G in this context. Unlike conventional modulation schemes, SCMA uses a combination of code sequences to represent user data, resulting in improved spectral efficiency and greater resistance to interference. However, the complexity of multi-user detection generated by the MPA (Message Passing Algorithm) represents a major challenge. In this study, an innovative approach using artificial intelligence (AI) is proposed to create a high-performance NOMA SCMA system in 5G networks. The main goal is to simplify MPA detection while decreasing the symbolic error rate (SER). Deep neural networks (DNNs) are used to detect signals in a 5G-specific NOMA SCMA environment, according to the proposed methodology. DNN training consists in detecting and distinguishing superimposed signals from various users, while taking into account specific modulation and channel characteristics.

Keywords: NOMA, SCMA, 5G, Message Passing Algorithm (MPA), SER, Artificial Intelligence (AI), Deep Neural Networks (DNNs).

مُلخَص

تقنية NOMA تلعب دوراً حيوياً في تطوير شبكات الاتصال اللاسلكية، خاصة في سياق الجيل الخامس. من خلال السماح لعدة مستخدمين بمشاركة نفس المورد الطيفي في نفس الوقت، تزيد تقنية NOMA من كفاءة الطيف وتوفر أداءً أفضل من حيث معدل البيانات والسعة مقارنة بالتقنيات التقليدية للوصول المتعدد. يظهر الوصول المتعدد بتقسيم الشفرة المكانية (SCMA) كحلاً واعداً للجيل الخامس في هذا السياق. على عكس التقنيات التقليدية للتشكيل، يستخدم SCMA مجموعة من تسلسلات الشفرة لتمثيل بيانات المستخدم، مما يؤدي إلى تحسين كفاءة الطيف وزيادة مقاومة التداخل. ومع ذلك، تشكل تعقيدات الكشف عن متعدد المستخدمين (MPA) ناتجة عن خوارزمية تمرير الرسائل تحدياً كبيراً. في هذه الدراسة، يتم اقتراح نهج مبتكر باستخدام الذكاء الاصطناعي (AI) لتطوير نظام عالي الأداء يستخدم تقنية NOMA في شبكات الجيل الخامس. الهدف الأساسي هو تبسيط عملية الكشف عن متعدد المستخدمين (MPA) مع إنقاص معدل الخطأ الرمزي (SER). تُستخدم الشبكات العصبية العميقة (DNNs) لكشف الإشارات في بيئة NOMA الخاصة بالجيل الخامس، وفقاً للمنهجية المقترحة. يتم تدريب DNNs على كشف الإشارات المتراكبة من مستخدمين مختلفين وتمييزها، مع مراعاة الخصائص الخاصة للتشكيل والقناة.

الكلمات المفتاحية : NOMA ، SCMA ، 5G، (خوارزمية تمرير الرسائل) MPA، SER، الذكاء الاصطناعي (AI) ، الشبكات العصبية العميقة (DNNs).

TABLE DES MATIÈRES :

LISTE DES ACRONYMES	I
LISTE DES TABLEAUX	V
LISTE DES FIGURES	VI
INTRODUCTION GÉNÉRALE.....	1
CHAPITRE I : GÉNÉRALITÉS SUR LA 5G.	3
I.1.Introduction	3
I.2. Une brève histoire des réseaux de communication mobiles	3
I.2.1. Les Premières Générations : 1G et 2G.....	4
I.2.2. L'Ère de la Donnée : 3G et 4G	4
I.2.3. La Convergence et l'Explosion des Données : 5G	5
I.3.La nouvelle génération des réseaux mobiles 5G	6
I.3.1.Qu'est-ce que la 5G ?.....	6
I.3.2.L'évolution des performances de la 4G vers la 5G	6
I.3.3.Objectifs techniques de la 5G	8
I.4.Cas d'usages de la 5G	8
I.4.1.Massive Machine Type Communications (mMTC)	8
I.4.2. Enhanced Mobile Broadband (eMBB)	9
I.4.3.Ultra-reliable and Low Latency Communications (uRLLC)	9
I.5.Le nouveau spectre pour la 5G	10
I.6.Les avancées technologiques au service de la 5G	11
I.6.1. La 5G New Radio	12
I.6.2.Les Fréquences Millimétriques (mmWave)	12
I.6.3.Le Massive MIMO	13
I.6.4.Le Beamforming	13
I.6.5.Network Slicing	15
I.6.6.Small Cells:.....	16

I.6.7.Le full duplex	18
I.6.8.Non-Orthogonal Multiple Access (NOMA)	18
I.6.9.Approches basées sur l'IoT pour la 5G	18
I.6.10.Virtualisation des fonctions réseau (NFV)	19
I.6.11. Réseau défini par logiciel (SDN) :.....	19
I.6.12. La communication Machine à Machine (M2M)	20
I.7.Les défis de la 5G	20
I.8. Architectures du Réseaux 5G	22
1.8.1. Architectures à deux niveaux	22
1.8.2.Le réseau de radio cognitive (CRN)	23
1.8.3.La communication de dispositif à dispositif (D2D)	23
1.8.4.Architectures basées sur le Cloud	24
I.9-Conclusion :.....	25
CHAPITRE II :LES TECHNIQUES D'ACCÈS MULTIPLE (OMA /NOMA).....	26
II.1.Introduction :	26
II.2.Les techniques d'accès multiple :.....	26
II.3.L'accès multiple orthogonal (OMA) :	27
II.3.1.FDMA:	27
II.3.2.TDMA:	28
II.3.3.CDMA:	28
II.3.4.OFDMA :.....	29
II.4.L'accès multiple non orthogonal (NOMA) :	29
II.4.1.Principales Technologies de NOMA :.....	30
II.4.1.1.Le codage par superposition (SC) :	30
II.4.1.2.L'annulation successive des interférences (SIC) :.....	31
II.4.2. Schémas de NOMA :.....	32

II.4.2.1.PD-NOMA :	32
II .4.2.1.1.Down Link:	32
II .4.2.1.2 .Up Link :	33
II.4.2.2.CD-NOMA :	34
II.4.2.2.1.LDS-CDMA :	34
II .4.2.2.2.MUSA :	35
II.4.2.2.3.PDMA :	36
II.4.2.2.4.GOCA :	37
II .4.2.2.5.SCMA :	38
A) Le système SCMA :	39
B) Les bases du système SCMA :	41
1.Le concept de code book :	41
2.Le concept de codeword :	43
3.Encodage SCMA :	43
4.Encodage SCMA basé sur le principe de rotation de constellation : ...	43
- La constellation	44
- La rotation de constellation	45
C) La technique de détection multi-utilisateurs MPA utilisée dans SCMA:....	47
1. Définition de MPA	47
2. Le graphe de facteurs.....	48
3. L'algorithme MPA	48
4. Les désavantages et les limites du MPA	51
II.4.2.3.Les autres schémas de NOMA	51
II.4.2.3.1.RSMA	51
II.4.2.3.2 .RDMA	52
II.5.Conclusion :	54

CHAPITRE III :LE DEEP LEARNING(L'APPRENTISSAGE EN PROFONDEUR).....	55
III.1.Introduction	55
III.2.L'intelligence artificielle (IA)	55
III.2.1. Bref historique	55
III.2.2. L'évolution de (IA)	56
III.2.3. Les sous domaines de (IA)	58
III.3.L'apprentissage automatique	59
III.3.1.définition	59
III.3.2.Principe	59
III.3.3.Types d'apprentissage automatique	59
III.3.3.1.Apprentissage supervisé	60
III.3.3.2.Apprentissage non supervisé	61
III.3.3.3.Apprentissage semi-supervisé	61
III.3.3.4.Apprentissage par renforcement	62
III.3.4.Limites de L'apprentissage automatique	62
III.4. L'apprentissage en profondeur	63
III.4.1.Définition	63
III.4.2.Les réseaux de neurones	64
III.4.2.1.Définition et présentation	64
III.4.2.2.Le neurone artificiel	64
III.4.2.3.Le fonctionnement d'un neurone	65
III.4.2.4.La fonction d'activation	66
III.4.2.5. Structure des réseaux de neurones	67
III.4.3.Les sous-domaines	68

III.4.3.1.Les réseaux de neurones artificiels (ANN)	68
III.4.3.2.Les réseaux de neurones convolutifs (CNN)	68
III.4.3.3.Les réseaux de neurones récurrents (RNN)	69
III.4.3.4.L'apprentissage par renforcement profond (DRL)	70
III .4.4.Le processus d'apprentissage	70
III .4.4.1.Principe	70
III.4.4.2.La Propagation avant et arrière :	71
III.4.4.2.1.La Propagation avant :	71
III.4.4.2.2.La Propagation arrière :	71
III.4.4.3.Les techniques de prétraitement des données :	72
III.4.4.3.1.Principe et bénéfices :	72
III.4.4.3.2.La standardisation :	73
III.4.4.3.3.La normalisation et la normalisation par lots :	73
III.4.4.3.4. L'encodage des caractéristiques catégorielles :	74
III.4.4.3.5.La discrétisation :	74
III.4.4.4.Les algorithmes d'optimisation:	75
III.4.4.4.1.GD :	75
III.4.4.4.2.SGD :	75
III.4.4.4.3.RMSprop :	75
III.4.4.4.4.Adam:	76
III.4.4.4.5.Adamax :	77
III.4.4.5.Évaluation de modèle et fonctions de perte :	77
III.4.4.6.L'underfitting et l'overfitting:	78
III.4.4.6.1.Définition :	78

III.4.4.6.2. Solutions pour l'underfitting et l'overfitting :	79
A) La validation :	79
B) La régularisation :	80
1. La régularisation L1 et L2 :	80
2. L'augmentation des données :	81
3. Le Dropout :	81
III.4.5. Les avantages de Deep Learning :	82
III.5. Les cas d'utilisation de l'intelligence artificielle dans la 5G :	83
III.6. Conclusion :	85
CHAPITRE IV : SIMULATION ET INTERPRÉTATION DES RÉSULTATS	86
IV.1. Introduction :	86
IV.2. Objectifs :	86
IV.3. Choix des outils de développement :	86
IV.4. Scénario proposé :	88
IV.5. Génération de données d'entraînement :	92
IV.6. Choix des critères :	93
IV.7. Architecture de modèle :	93
1. Modulation BPSK ($M=2$)	94
2. Modulation QPSK ($M=4$) :	94
IV.8. Résultats expérimentale et discussion :	95
IV.8.1. Analyse de l'Évolution de l'exactitude et de la perte au fil des époques, avec considération de l'impact du SNR :	96
IV.8.2. Analyse de l'évolution de SER en fonction de SNR :	100
IV.8. 3. Optimisation des Performances de Décodage par Exploration des Hyperparamètres :	101
IV.8.4. Analyse de l'évolution de SER par rapport à MPA :	104
IV.8.5. Analyse Comparative de SER entre DNN et NN :	105

TABLE DES MATIÈRES :

IV.8.6. Analyse de complexité :.....	108
IV.9. Conclusion :	109
CONCLUSION GÉNÉRALE ET PERSPECTIVES	110
BIBLIOGRAPHIE/WEBOGRAPHIE.....	112

LISTE DES ACRONYMES

1G	Première Génération
2G	Deuxième Génération
3D MIMO	3-Dimensional Multiple Input Multiple Output
3G	Troisième Génération
4G	Quatrième Génération
5G	Cinquième Génération
AdaDelta	Adaptive Delta Algorithm
AdaGrad	Adaptive Gradient Algorithm
Adam	Adaptive Moment Estimation
Adamax	Adaptive Moment Estimation with Maximum
AMSgrad	Adaptive Moment Estimation squared Algorithm
AMPS	Advanced Mobile Phone System
ANN	Artificial Neural Networks
AR	Réalité Augmentée
ARCEP	Autorité de Régulation des Communications Électroniques et des Postes
BBU	Baseband Unit
BPSK	Binary Phase Shift Keying
CDMA	Code Division Multiple Access
CNN	Convolutional Neural Networks
CRN	Coordinated Radio Access Network
C-RAN	Cloud Radio Access Network

LISTE DES ACRONYMES

D2D	Device-to-Device
DL	Deep Learning
DRL	Deep Reinforcement Learning
eMBB	Enhanced Mobile Broadband
EDGE	Enhanced Data rates for GSM Evolution
FEC	Forward Error Correction
FDMA	Frequency Division Multiple Access
FDD	Frequency Division Duplexing
GD	Gradient Descent
GOCA	Group Orthogonal Code Division Multiple Access
GPS	Global Positioning System
GSM	Global System for Mobile Communications
HSPA	High Speed Packet Access
IA	Intelligence Artificielle
IDMA	Interleave Division Multiple Access
IEEE	Institute of Electrical and Electronics Engineers
IGMA	Interference Grouping Multiple Access
IoT	Internet of Things
LDS-CDMA	Low Density Spreading Code Division Multiple Access
LTE	Long Term Evolution
MA	Multiple Access

LISTE DES ACRONYMES

MBS	Macro Base Station
M2M	Machine-to-Machine
MIMO	Multiple Input Multiple Output
ML	Machine Learning
mMIMO	Le MIMO massif
mMTC	Massive Machine Type Communications
MMS	MultiMedia Messaging Service
mmWave	millimeter wave
MPA	Message Passing Algorithm
MUSA	Multi-User Shared Access
NMT	Nordic Mobile Telephone
NFV	Network Functions Virtualization
NOMA	Non-Orthogonal Multiple Access
NAG	Nesterov Accelerated Gradient
NR	New Radio
OMA	Orthogonal Multiple Access
OFDM	Orthogonal Frequency Division Multiplexing
OFDMA	Orthogonal Frequency Division Multiple Access
PDMA	Power Domain Multiple Access
QAM	Quadrature Amplitude Modulation
QPSK	Quadrature Phase Shift Keying
Qos	Quality of Service
RDMA	Remote Direct Memory Access
RMSprop	Root Mean Square Propagation

LISTE DES ACRONYMES

RL	Renforcement Learning
RRH	Remote Radio Head
RNN	Recurrent Neural Networks
RSMA	Resource Sharing Multiple Access
SBS	Small Base Station
SCMA	Sparse Code Multiple Access
SGD	Stochastic Gradient Descent
SINR	Rapport signal sur interférence plus bruit
SDN	Software-Defined Networking
TACS	Total Access Communication System
TDMA	Time Division Multiple Access
TDD	Time Division Duplexing
UMTS	Universal Mobile Telecommunications System
UIT-R	Union Internationale des Télécommunications - Radiocommunications
uRLLC	Ultra-reliable and Low Latency Communications
VR	Réalité Virtuelle
VoIP	Voice over Internet Protocol
WCDMA	Wideband Code Division Multiple Access
WiFi	Wireless Fidelity
WiMax	Worldwide Interoperability for Microwave Access
WSDN	Wireless Software-Defined Networking
WWW	Wireless World Wide Web

LISTE DES TABLEAUX

CHAPITRE I :

Tableau I.1: Une comparaison des vitesses maximales pour différentes générations.	5
Tableau I.2: Tableau comparatif des caractéristiques des différentes générations de téléphonie mobile	7
Tableau I.3: Différents types de cellules.....	23

CHAPITRE III :

Tableau III.1: Les fonctions d'Activation les plus utilisées	69
Tableau III.2: Les fonctions de perte classées en fonction du type de tâche.....	81

CHAPITRE IV :

Tableau IV.1 : Les outils de développement utilisés pour la simulation.	87
Tableau IV.2: Les données de base de système SCMA proposé.	92
Tableau IV.3 : Choix des hyper paramètres dans le cas où $M=2$	94
Tableau IV.4 : Choix des hyper paramètres dans le cas où $M=4$	94
Tableau IV.5 : La répartition des ensembles de données pour les deux modèles	95

LISTE DES FIGURES :

CHAPITRE I :

Figure I.1 : Cycles des générations de téléphonie mobile.	4
Figure I.2 : Comparaison entre 4G et 5G.....	6
Figure I.3 : Les Cas d'usage de la 5G.....	9
Figure I.4 : Spectre pertinent pour l'accès sans fil 5G.	10
Figure I.5 : Les avancées technologiques au service de la 5G.....	11
Figure I.6 : Représentation visuelle des ondes millimétriques.	13
Figure I.7 : Illustration de (a) Beam steering, (b) Beam forming.	14
Figure I.8 : 3D MIMO pour la sectorisation verticale.	15
Figure I.9 : Les réseaux 5G ; a) sans le découpage de réseau, b) avec le découpage de réseau	16
Figure I.10 : Déploiement hétérogène de petites cellules et de cellules macro.	17
Figure I.11 : Illustration du full-duplex, comparé au FDD et TDD.....	18
Figure I.12 : Une architecture multi-niveaux pour les réseaux 5G avec des petites cellules, des petites cellules mobiles, et des communications basées sur D2D et CRN.	22
Figure I.13 : L'architecture de communication D2D.	24
Figure I.14 : Une architecture basée sur le Cloud pour les réseaux 5G.....	24

CHAPITRE II :

Figure II.1 : Le concept d'un système d'accès multiple.	26
Figure II.2 : Les techniques d'accès OMA.	27
Figure II.3 : Le principe de technique d'accès multiple FDMA.....	27
Figure II.4 : Le principe de technique d'accès multiple TDMA.	28
Figure II.5 : Le principe de technique d'accès multiple CDMA.	28

Figure II.6 : Le principe de technique d'accès multiple OFDMA.....	29
Figure II.7 : Classification des techniques d'accès NOMA	30
Figure II.8 : Diagramme de constellation hiérarchique de l'exemple (SIC).	31
Figure II.9 : Représentation d'un système NOMA conventionnel à deux utilisateurs où l'utilisateur-N assiste l'utilisateur-F.	32
Figure II.10 : Une illustration de système NOMA en liaison descendante (PD-Down Link).	33
Figure II.11 : Illustration d'un système NOMA en liaison montante (PD-Up Link)	34
Figure II.12 : Structure de (a) LDS-CDMA en comparaison avec (b) CDMA conventionnel.	35
Figure II.13 : Éléments du code de diffusion complexe MUSA.	36
Figure II.14 : Illustration d'un exemple de motif PDMA.....	37
Figure II.15 : Illustration d'un exemple de génération de séquence dans l'émetteur GOCA.	38
Figure II.16 : L'accès multiple SCMA.	39
Figure II.17 : Illustration d'un modèle émetteur-récepteur SCMA en liaison descendante	41
Figure II.18 : Illustration de la structure du livre de codes SCMA.	42
Figure II.19 : Un exemple d'un codebook SCMA à 8 points.....	42
Figure II.20 : Comptage des étudiants utilisant MPA ; (a) Comptage groupé des étudiants. (b) Comptage par étudiant individuel.	48
Figure II.21 : Graphe de facteurs du système SCMA (4×6).	49

Figure II.22 : Présentation des diagrammes en bloc des deux modes de RSMA.	52
Figure II.23 : Un exemple de motifs de répétition RDMA avec 3 utilisateurs simultanés.....	53
CHAPITRE III :	
Figure III.1 : Chronologie de l'évolution de l'IA.	58
Figure III.2 : Les sous domaines de l'IA.	58
Figure III.3 : Classification des types d'apprentissage automatique.....	60
Figure III.4 : Fonctionnement de a) Supervised Learning, b) Unsupervised Learning	61
Figure III.5: Illustration de l'apprentissage par renforcement.	62
Figure III.6 : Qualité des données de ML et DL.	63
Figure III.7: Le DL, une sous-discipline de ML	63
Figure III.8 : Illustration : de (a) le neurone biologique et (b) le neurone artificiel.....	64
Figure III.9 : Un nœud avec trois entrées.	65
Figure III.10: Représentation de différents structures de RNA ; a) Réseau neuronal à une seule couche, b) Réseau neuronal multicouche , c) Réseau neuronal profond	68
Figure III.11 : L'architecture de base d'un CNN.....	69
Figure III. 12: L'architecture de base d'un RNN.....	70
Figure III.13: Entraîner le réseau neuronal en utilisant l'algorithme de rétropropagation.....	75
Figure III .14 : Exemples de a) sous-ajustement, b) ajustement optimal, c) surajustement ...	79
Figure III.15: Illustration de la division des données d'entraînement pour le processus de validation	80

Figure III.16 : Représentation conceptuelle du a) réseau dense standard en comparaison avec la structure du b) réseau après l'application du dropout.....	82
Figure III.17 : Diverses opportunités d'IA pour les réseaux mobiles	83
CHAPITRE IV :	
Figure IV.1 : Système SCMA proposé basé sur la détection par DL	88
Figure IV.2 : Constellations des 4 REs avec $\Delta=\pi/4$ lors de l'utilisation de la modulation BPSK.....	89
Figure IV.3 : Constellations des 4 REs avec $\Delta=\pi/5$ lors de l'utilisation de la modulation BPSK.....	89
Figure IV.4 : Constellations des 4 REs avec $\Delta=\pi/6$ lors de l'utilisation de la modulation BPSK.....	90
Figure IV.5 : Constellations des 4 REs avec $\Delta=\pi/4$ lors de l'utilisation de la modulation QPSK.....	90
Figure IV.6 : Constellations des 4 REs avec $\Delta=\pi/5$ lors de l'utilisation de la modulation QPSK.....	91
Figure IV.7 : Constellations des 4 REs avec $\Delta=\pi/6$ lors de l'utilisation de la modulation QPSK.....	91
Figure IV.8: L'évolution de : (a) l'accuracy et (b) le loss au fil des époques d'entraînement du modèle dans l'entraînement et la validation pour $M=2$ et $\Delta=\pi/6$	96
Figure IV.9 : L'évolution de : (a) l'accuracy et (b) le loss au fil des époques d'entraînement du modèle dans l'entraînement et la validation pour $M=2$ et $\Delta=\pi/5$	97
Figure IV.10 : L'évolution de : (a) l'accuracy et (b) le loss au fil des époques d'entraînement du modèle dans l'entraînement et la validation pour $M=2$ et $\Delta=\pi/4$	97
Figure IV.11: L'évolution de : (a) l'accuracy et (b) le loss au fil des époques d'entraînement du modèle dans l'entraînement et la validation pour $M=4$ et $\Delta=\pi/6$	98

Figure IV.12 : L'évolution de : (a) l'accuracy et (b) le loss au fil des époques d'entraînement du modèle dans l'entraînement et la validation pour $M=4$ et $\Delta=\pi/5$	98
Figure IV.13 : L'évolution de : (a) l'accuracy et (b) le loss au fil des époques d'entraînement du modèle dans l'entraînement et la validation pour $M=4$ et $\Delta=\pi/4$	99
Figure IV.14 : Évolution de SER en fonction de SNR de modèle ($M=2$) avec différents Δ	100
Figure IV.15 : Évolution de SER en fonction de SNR de modèle proposé ($M=4$) avec différents Δ	100
Figure IV.16 : Évolution des performances de décodage pour un modèle ($M=4$, $\Delta= \pi/6$) avec $a = 10^{-3}$ et trois BS.	102
Figure IV.17: Évolution des performances de décodage pour un modèle ($M=4$, $\Delta= \pi/6$) avec $a = 10^{-4}$ et trois BS différents.....	102
Figure IV.18 : Évolution des performances de décodage pour un modèle ($M=2$, $\Delta= \pi/5$) avec $a = 10^{-4}$ et trois BS différents.	103
Figure IV.19 : Évolution des performances de décodage pour un modèle ($M=2$, $\Delta= \pi/5$) avec $a = 5.10^{-4}$ et trois BS différents.	103
Figure IV.20 : Évolution du SER du modèle proposé ($M=2$) comparé au MPA en fonction du SNR.	104
Figure IV.21: Évolution du SER du modèle proposé ($M=4$) comparé au MPA en fonction du SNR.	105
Figure IV.22 : Comparaison des courbes de SER entre NN et DNN pour $M=2$	106
Figure IV.23 : Comparaison des courbes de SER entre NN et DNN pour $M=4$	106
Figure IV.24 : Comparaison du temps d'exécution entre MPA et DNN en secondes	108

Introduction Générale

Les réseaux de télécommunication mobile sont à l'avant-garde de la révolution numérique, transformant radicalement notre manière de communiquer, de travailler et d'interagir avec le monde qui nous entoure. Ces systèmes sans fil sophistiqués ont évolué à pas de géant depuis leurs modestes débuts, offrant une connectivité universelle et des services innovants à des milliards d'utilisateurs à travers le monde.

Depuis l'avènement des premiers réseaux cellulaires dans les années 1970, chaque génération de communications sans fil cherche à repousser les limites pour créer un monde connecté où la transmission de données signifie une qualité, une fiabilité et une vitesse accrues. Chaque nouvelle génération s'appuie sur des technologies d'accès innovantes et améliorées pour répondre aux divers besoins des utilisateurs.

L'IA est l'une des technologies qui se développent rapidement ces dernières décennies, fusionnant dans plusieurs domaines, y compris les télécommunications. Cette révolution de l'intelligence artificielle ces dernières années offre un potentiel significatif pour optimiser les réseaux sans fil, réduire leur complexité et augmenter leur fiabilité.

Dans l'autre côté, les techniques d'accès non-orthogonales (NOMA) sont prometteuses dans les réseaux 5G en raison de leurs avantages majeurs, tels que la connectivité massive, l'efficacité spectrale et l'amélioration de l'efficacité énergétique. De plus, l'intelligence artificielle joue un rôle important dans l'optimisation des performances de NOMA.

Dans ce contexte, ce projet vise à optimiser les systèmes NOMA- SCMA pour les réseaux 5G en utilisant des techniques d'intelligence artificielle afin d'améliorer le taux d'erreur de symbole (SER) et de réduire la complexité de la détection dans ces systèmes.

Ce mémoire est structuré en quatre chapitres, chacun examinant un aspect clé du sujet abordé.

- Le premier chapitre sera consacré à l'analyse de l'évolution des réseaux de communication mobile, mettant en lumière l'importance croissante de la technologie 5G.
- Dans le deuxième chapitre, on aborde les diverses méthodes d'accès, telles que OMA et NOMA, en mettant l'accent sur SCMA.

- Le troisième chapitre expose en profondeur les concepts de Machine Learning et de Deep Learning.
- Le dernier chapitre décrit une solution basée sur le DL pour réduire le SER et le taux de la détection dans les systèmes SCMA.

Enfin, on terminera bien sûr par une conclusion générale en donnant des perspectives pour les futurs travaux.

CHAPITRE I :

Généralités sur LA 5G.

I.1.Introduction :

Un réseau mobile est un système de réseau téléphonique qui fonctionne grâce à des fréquences formant un spectre hertzien. Ce réseau permet à des millions d'utilisateurs de téléphoner en même temps tout en étant en mouvement, sans aucune contrainte d'immobilité. Les systèmes de communication sans fil sont devenus omniprésents dans notre vie quotidienne, offrant une alternative pratique à l'utilisation de câbles et permettant une connectivité à tout moment et en tout lieu. Avec le développement incessant des réseaux mobiles, plusieurs générations ont vu le jour, chacune apportant des avancées significatives en termes de débit, de bande passante et de capacité.

La cinquième génération (5G) des réseaux de communication mobiles représente une avancée significative dans le domaine des technologies de l'information et de la communication. Nous commencerons par contextualiser l'évolution des réseaux mobiles depuis leurs débuts avec la 1G, puis nous examinerons l'importance croissante de la 5G dans le paysage technologique actuel. Nous examinerons les objectifs techniques de la 5G et les cas d'utilisation qui motivent son développement. Ensuite, nous nous pencherons sur les avancées technologiques clés qui alimentent la 5G. Nous aborderons également les défis auxquels la 5G est confrontée, ainsi que les architectures.

I.2. Une brève histoire des réseaux de communication mobiles:

Au cours des dernières décennies, les réseaux de communication mobiles ont émergé comme un élément crucial du développement des technologies de l'information à l'échelle mondiale. En effet, les réseaux cellulaires ont profondément transformé notre conception de la connectivité et nos méthodes de communication. Cette évolution s'est appuyée sur un rythme soutenu d'innovation scientifique et technologique, organisé en générations successives. Après l'avènement de la 2G en 1991, nous avons vu émerger la 3G dès 2001, puis la 4G en 2010, permettant ainsi de répondre aux besoins croissants en matière de télécommunications mobiles et d'accès à Internet.

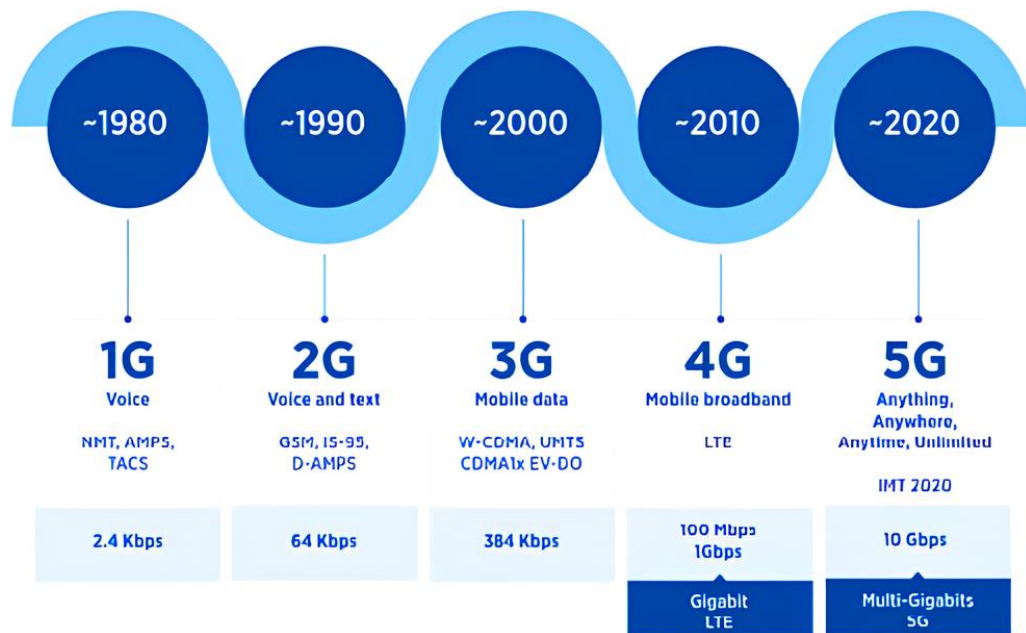


Figure I.1 : Cycles des générations de téléphonie mobile [1].

I.2.1. Les Premières Générations : 1G et 2G

Les premières générations de réseaux mobiles étaient principalement axées sur la voix et offraient des services de communication de base. La 1G introduite dans les années 1980, utilisait la technologie analogique pour les appels vocaux [1]. La 2G arrivée dans les années 1990 a introduit la numérisation des signaux et la capacité d'envoyer des messages textuels. Les normes telles que le GSM et le CDMA ont été des technologies clés de cette époque. Une version (2,5G) a été mise au point, avec un débit plus rapide, d'environ 40 kbit/s, en utilisant le standard GPRS. De la même manière, la version EDGE, qui offre un débit de 3 487Kbit/s, est également disponible sur le marché [2].

I.2.2. L'Ère de la Donnée : 3G et 4G

Le réseau mobile 3G a été le premier à proposer une connexion vocale avec des fonctionnalités multimédias. La 3G reposait sur l'utilisation de la technologie de l'Accès Haut Débit par Paquets (HSPA/HSPA+). De plus, la 3G employait la méthode MIMO afin de renforcer la puissance du réseau sans fil, ainsi que la commutation par paquets pour une transmission rapide des données. 10 années après les débuts de la 2G, l'arrivée de la 3G va bouleverser l'utilisation des téléphones portables, en particulier en augmentant les débits de 384 kbit/s à 2 Mbit/s, puis jusqu'à 42 Mbit/s. Les MMS permettent l'échange de messages multimédias, comme des courriels, des photos et des vidéos, avec des débits d'environ 1,9 Mbit/s. La 3G+ offre un débit de 10Mbit/s, tandis que la 3G++ offre un débit de 42Mbit/s [2].

La 3G a été suivie par la 4G en 2013, qui a considérablement amélioré les vitesses de transmission de données, permettant le streaming vidéo en haute définition et les jeux en ligne sur les Smartphones. Les normes telles que l'UMTS pour la 3G et le LTE pour la 4G ont joué un rôle majeur dans cette évolution. La 4G s'efforce de répondre aux évolutions et à la massification des usages en adoptant des solutions visant à améliorer l'efficacité de l'utilisation des fréquences disponibles. Elle concrétise également la volonté des opérateurs de standardiser les solutions techniques pour véhiculer des informations, dans le but de garantir une cohérence technique et de réduire les coûts. Cela se traduit par l'adoption de protocoles universellement utilisés, permettant ainsi une harmonisation des pratiques à l'échelle mondiale [3].

I.2.3. La Convergence et l'Explosion des Données : 5G

La 5G est un pilier de la transformation numérique ; c'est une véritable amélioration par rapport à tous les réseaux de génération mobile précédents. La 5G offre une vitesse supérieure à la 4G, permettant un fonctionnement télécommandé sur un réseau fiable et réactif, avec des débits descendant pouvant atteindre 20 Gbps. Elle est hautement compatible avec le WWW. Elle assure une connectivité Internet illimitée, rapide, à faible latence, fiable et accessible à tout moment [4][5]. Le **Tableau I.1** présente une comparaison des vitesses maximales autorisées pour différentes générations de réseaux de télécommunications, allant de la 1G à la 5G.

Génération	Vitesse maximale autorisée
1G	40 à 200 calculs/sec
2G - GPRS	0,0125 Mbit/s
2G - Edge	0,0375 Mbit/s
3G - Basic	0,0375 Mbit/s
3G - HSPA	0,9 Mbit/s
3G - HSPA+	2,625 Mbit/s
3G - DC - HSPA+	42 Mbit/s
4G - LTE catégorie 4	150 Mbit/s
4G - LTE Advanced CAT 6	300 Mbit/s
4G - LTE Advanced CAT 9	450 Mbit/s
4G - LTE Advanced CAT 12	600 Mbit/s
4G - LTE Advanced CAT 16	979 Mbit/s
5G	1 à 10 Gbit/s

Tableau I.1: Une comparaison des vitesses maximales pour différentes générations.

I.3. La nouvelle génération des réseaux mobiles 5G :

I.3.1. Qu'est-ce que la 5G ?

La 5G se présente comme une génération de rupture. Selon Arcep « La 5G souhaite se présenter comme la génération de rupture, la génération qui ne s'intéresse plus uniquement au monde des opérateurs de téléphonie mobile et des communications grand public, mais qui ouvre de nouvelles perspectives et permet la cohabitation d'applications et usages extrêmement diversifiés, unifiés au sein d'une même technologie. La 5G se pose en facilitateur de la numérisation de la société et de l'économie ».

L'IEEE définit la 5G comme la cinquième génération de réseaux mobiles, offrant une connectivité sans fil à haut débit et à faible latence. Elle repose sur une combinaison de technologies avancées telles que le Massive MIMO, les ondes millimétriques, la virtualisation du réseau, le network slicing et l'edge computing. La 5G vise à répondre à une demande croissante en connectivité pour prendre en charge un large éventail d'applications, notamment l'Internet des objets (IoT), les véhicules autonomes, la réalité augmentée et virtuelle, ainsi que les applications industrielles et médicales nécessitant une connectivité fiable et à faible latence [4].

I.3.2. L'évolution des performances de la 4G vers la 5G :

Comparée aux réseaux actuels, la 5G offrira une vitesse de transmission considérablement accrue, avec des débits pouvant atteindre 10 Gbit/s, soit 10 à 100 fois plus rapides que la 4G et la 4G-LTE. Cette avancée permettra de dépasser les réseaux ultra haut débit actuels en combinant les technologies existantes telles que l'Internet des objets (IoT), le cloud, le big data, l'intelligence artificielle [6] [7].

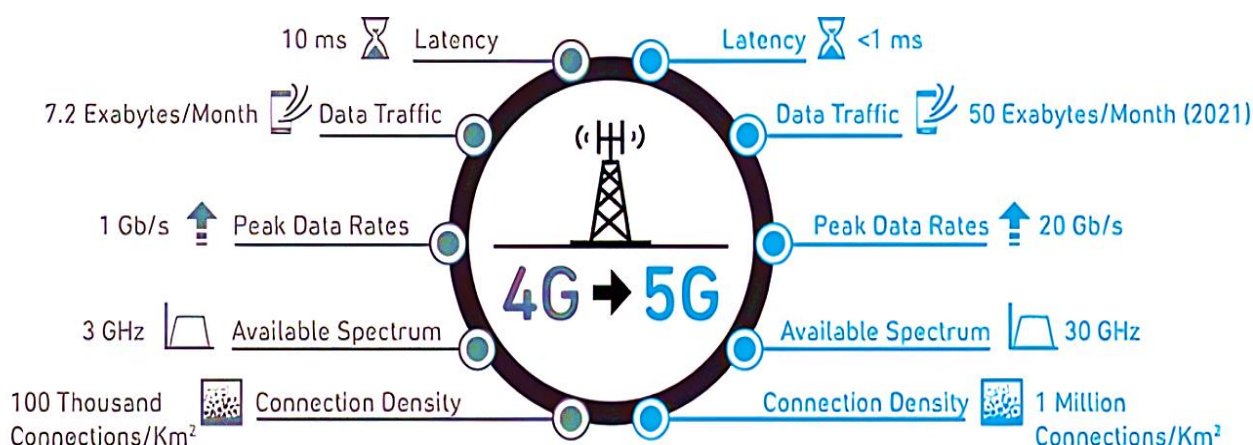


Figure I.2 : Comparaison entre 4G et 5G [7].

La 5G présente également une latence considérablement réduite, avec un temps de réponse inférieur à une milliseconde, permettant des interactions en temps réel. Elle ouvre également la voie à une utilisation optimale de l'IoT grâce à sa largeur de bande élevée et sa connectivité accrue [8].

En conclusion, la 5G ne représente pas seulement une évolution des réseaux mobiles, mais une avancée majeure qui répond aux besoins croissants des utilisateurs, tire parti des avancées technologiques pour faire face à la demande croissante de capacité, maintient la compétitivité mondiale et assure l'interopérabilité grâce à des normes internationales. En somme, la 5G ouvre la voie à une ère de connectivité avancée et de services innovants, propulsant ainsi notre société vers de nouveaux horizons [9][10][11]. Voici un tableau comparatif des caractéristiques des différentes générations de téléphonie mobile :

Caractéristique	1G (Analogique)	2G (Numérique)	3G (Numérique)	4G (LTE)	5G
Années d'introduction	Années 1980	Années 1990	Début des années 2000	Années 2010	Début des années 2020
Technologie principale	AMPS, TACS, NMT	GSM, CDMA	UMTS, CDMA2000	LTE	NR (Nouvelle Radio)
Débit de données	< 10 kbps	9.6 - 144 kbps	0.1 - 2 Mbps	100 Mbps - 1 Gbps	> 1 Gbps
Utilisation principale	Voix	Voix, SMS, données	Voix, SMS, données	Données	Données, IoT, véhicules autonomes
Technologie d'accès multiple	FDMA	TDMA, CDMA	CDMA, WCDMA	OFDMA	OFDMA, SCMA, etc.
Latence	//	100 - 500 ms	50 - 100 ms	10 - 30 ms	< 1 ms
Caractéristiques clés	Analogique, voix seulement	Numérique, SMS	Données mobiles, VoIP	Haut débit, streaming	Haut débit, faible latence, IoT

Tableau I.2 : Tableau comparatif des caractéristiques des différentes générations de téléphonie mobile.

I.3.3.Objectifs techniques de la 5G :

L'objectif de la 5G est de fournir une multitude d'utilités au consommateur à grande vitesse. Les applications développées pour utiliser ces utilités sont très conviviales pour les utilisateurs ; elles réduisent l'intercommunication entre l'application et l'utilisateur. Par exemple, l'intégration de la technologie de reconnaissance vocale dans les interfaces utilisateur faciliterait l'utilisation des applications pour chaque utilisateur. Les exigences techniques pour la 5G comprennent [12]:

- Une bande passante d'au moins 100 MHz.
- Des bandes passantes allant jusqu'à 1 GHz sont nécessaires pour les hautes fréquences (mmWave).
- Une densité de connectivité de 1 million d'appareils par kilomètre carré.
- Un débit théorique de 20 Gbit/s en liaison descendante, et 10 Gbit/s en liaison montante.
- Un débit perçu par l'utilisateur de 100Mbit/s en liaison descendante,et 50Mbit/s en liaison montante.
- Une latence inférieure à 1 ms.
- Une capacité de mobilité jusqu'à 500 km/h.
- Une disponibilité de 99,999 % ainsi qu'une couverture réseau de 100 %.
- Une réduction de la consommation d'énergie jusqu'à 90 %.
- la 5G permettra la diffusion massive de données, avec des débits en gigabits, capable de prendre en charge plus de 60 000 connexions simultanées.

I.4.Cas d'usages de la 5G :

Dans le contexte de la 5G, on distingue trois types d'utilisation différents [11] :

I.4.1.Massive Machine Type Communications (mMTC) :

Cette catégorie vise à répondre à tous les usages liés à l'Internet des objets, nécessitant une couverture étendue, une faible consommation énergétique et des débits relativement modestes. Ces services comprennent le télé-relevé de compteurs d'eau et le suivi d'objets tels que les bouteilles de gaz, nécessitant une connectivité dense sur de vastes territoires [13].

La 5G permet des communications machine à machine à grande échelle, également connues sous le nom d'Internet des Objets (IoT), qui facilitent la connectivité entre de

nombreux dispositifs sans intervention humaine. Ces services enrichissent les applications de la 5G et favorisent la connectivité dans des secteurs tels que l'agriculture, la construction et l'industrie [14][15][16].

I.4.2. Enhanced Mobile Broadband (eMBB):

Il s'agit de fournir un accès à Internet à très haut débit, tant en extérieur qu'en intérieur, avec une qualité de service uniforme même en périphérie des zones de couverture. L'amélioration de la large bande mobile est un aspect clé de la 5G, utilisant des technologies telles que les antennes massives MIMO, les ondes millimétriques et les techniques de formation de faisceaux pour fournir une connectivité à très haut débit sur de vastes zones [17].

I.4.3. Ultra-reliable and Low Latency Communications (uRLLC) :

Désignent des communications extrêmement fiables qui répondent à des besoins critiques avec une latence extrêmement faible, ce qui offre une plus grande réactivité. Il est primordial d'utiliser cette technologie pour les applications qui exigent des temps de réponse extrêmement courts et une fiabilité maximale [11].

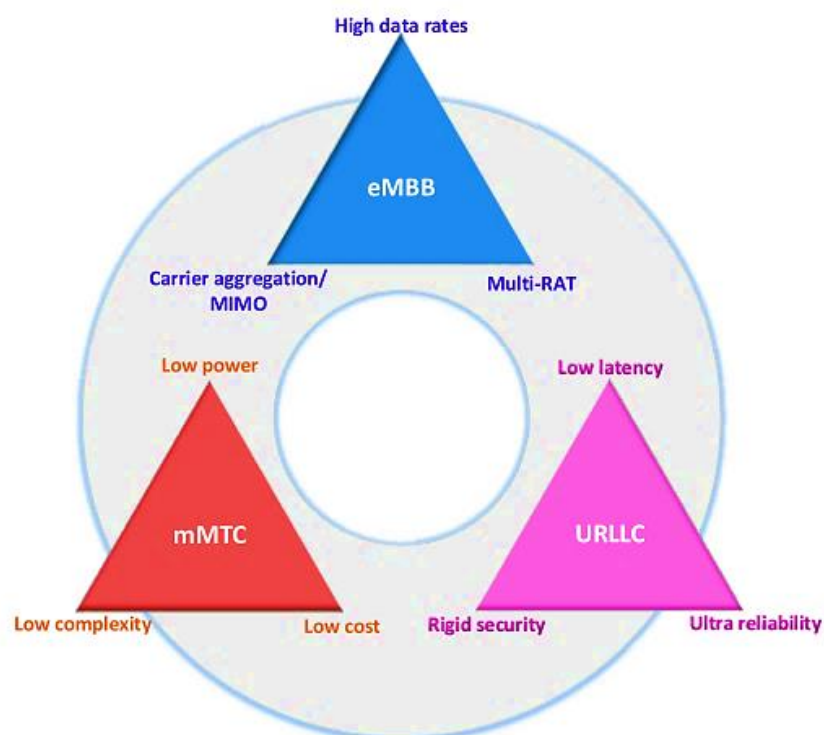


Figure I.3 : Les Cas d'usage de la 5G [16].

I.5. Le nouveau spectre pour la 5G :

Afin de soutenir une capacité de trafic accrue et de permettre les largeurs de bande de transmission nécessaires pour prendre en charge des débits de données très élevés, la 5G étendra la gamme de fréquences utilisées pour la communication mobile. Cela inclut de nouveaux spectres en dessous de 6 GHz, ainsi que des spectres dans des bandes de fréquences plus élevées. Le spectre spécifique candidat pour la communication mobile dans les bandes de fréquences plus élevées reste à être identifié par l'UIT-R ou par les organismes de réglementation individuels. L'industrie mobile reste agnostique quant aux choix particuliers, et l'ensemble de la plage de fréquences jusqu'à environ 100 GHz est en cours de considération à ce stade, bien qu'il y ait un intérêt significatif pour des allocations contiguës importantes qui offrent un spectre dédié et sous licence pouvant être utilisé par plusieurs fournisseurs de réseaux concurrents [19].

Les propriétés de propagation favorisent la partie inférieure de cette plage de fréquences, en dessous de 30 GHz. Au-delà de 30 GHz, il est plus probable que de grandes quantités de spectre et de larges bandes de fréquences de transmission de l'ordre de 1 GHz ou plus soient disponibles. Ainsi, le domaine d'application compatible avec l'accès sans fil à la 5G s'étend de moins de 1 GHz jusqu'à environ 100 GHz, comme illustré dans la **Figure I.4** [19] :

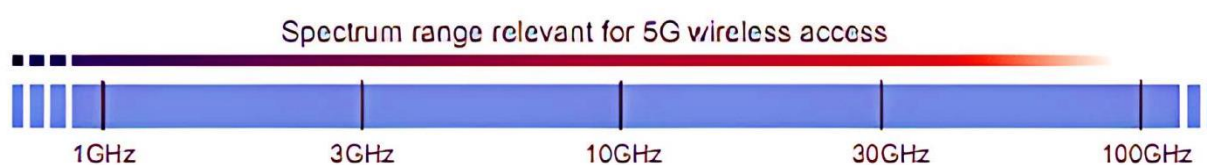


Figure I.4: Spectre pertinent pour l'accès sans fil 5G [19].

Il est important de comprendre que les hautes fréquences, en particulier celles au-dessus de 10 GHz, ne peuvent servir que de complément aux bandes de fréquences plus basses, et fourniront principalement une capacité système supplémentaire et des largeurs de bande de transmission très larges pour des débits de données extrêmes dans des déploiements denses.

I.6. Les avancées technologiques au service de la 5G :

Un ensemble de technologies collaboratives gouverne l'architecture du réseau 5G, comme illustré dans la **Figure I.5**. Le déploiement du réseau défini par logiciel sans fil (WSDN), permet une gestion centralisée et programmable des réseaux sans fil, offrant une flexibilité accrue et une optimisation des performances du réseau. La virtualisation des fonctions réseau (NFV), vise à virtualiser les fonctions réseau traditionnellement exécutées sur des équipements matériels dédiés, permettant ainsi une gestion plus dynamique et efficace des ressources réseau. L'ultra-densification, les communications de périphérie à périphérie (D2D), les ondes millimétriques, le MIMO massif, les nouvelles techniques d'accès, l'Internet des objets (IoT), les communications écologiques, le big data et l'informatique en nuage mobile permettront d'améliorer l'accès radio [16].

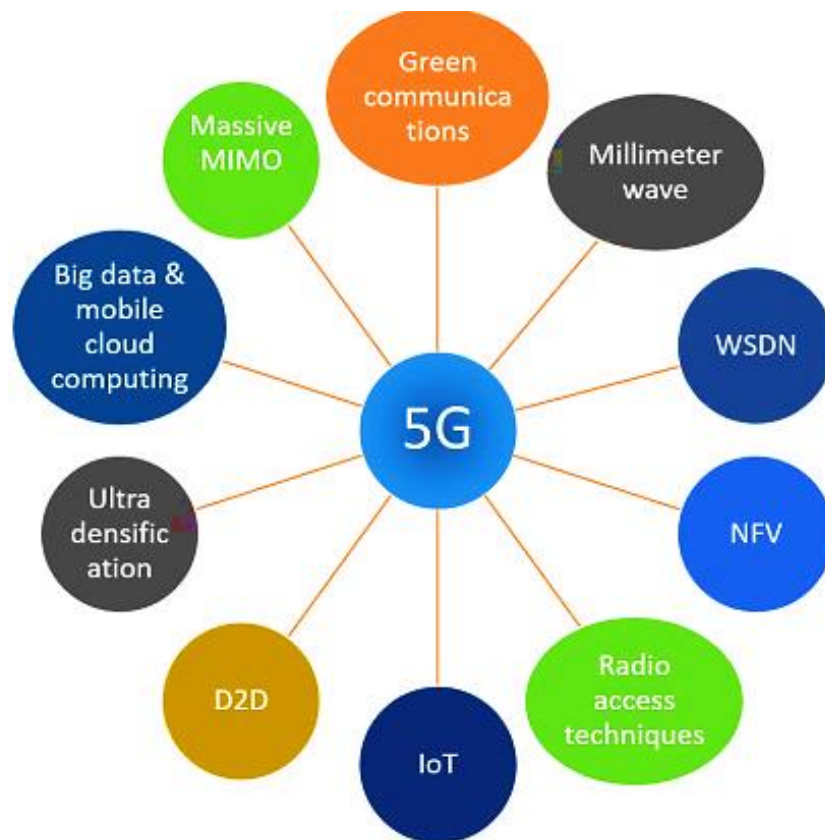


Figure I.5 : Les avancées technologiques au service de la 5G [16].

Cette section explore les avancées récentes dans plusieurs technologies clés de la 5G :

I.6.1. La 5G New Radio :

La technologie mobile 5G New Radio (NR) est la première à exploiter le spectre des ondes millimétriques. Malgré son développement depuis de nombreuses années, les différents acteurs et le public ne comprennent toujours pas pleinement ses caractéristiques et ses performances, ainsi que l'état de l'infrastructure et des appareils. Tandis que le secteur des communications mobiles continue de progresser dans le développement et la mise en place des systèmes 5G NR dans la bande des ondes millimétriques, de nombreuses nouvelles informations sont maintenant disponibles concernant de nouvelles méthodologies, des résultats de performance et des apprentissages tirés de différents essais et déploiements. Ces informations peuvent éclairer le spectre des ondes millimétriques, ses bénéfices et même les défis qui y sont associés.

Tandis que les avancées des systèmes 5G NR dans les fréquences des ondes millimétriques continuent, l'un des résultats les plus attendus est sa mise en œuvre immédiate dans la fourniture de services mobiles haut débit améliorés (eMBB), offrant ainsi aux opérateurs la possibilité de répondre à la demande croissante de données à haute vitesse pour les situations de déploiement courantes en extérieur urbaines denses. La technologie de 5G NR mmWave offre une alternative convaincante aux déploiements Wi-Fi déjà en place, à la fois pour les espaces intérieurs et dans les espaces d'entreprise [20][21] .

I.6.2. Les fréquences millimétriques (mmWave) :

Le spectre des ondes millimétriques, oscillant entre 30 GHz et 300 GHz, est caractérisé par de courtes longueurs d'onde, d'où le terme "millimeter wave". Ces ondes sont utilisées dans le domaine de la 5G avec des fréquences légèrement plus élevées telles que 24 GHz et 28 GHz, partageant des caractéristiques similaires. Cette bande ainsi que les fréquences de 39 GHz, offrent un environnement de diffusion relativement ouvert avec la plupart des canaux de communication situés sous le niveau de bruit. Alors que de nombreuses technologies de communication comme WiMax, GPS, wifi, 4G et 3G opèrent principalement entre 1 et 6 GHz. Le spectre mmWave entre 30 et 300 GHz reste sous-utilisé. Récemment, une bande de fréquence spécifique de 24 à 100 GHz a été réservée pour la 5G [4].

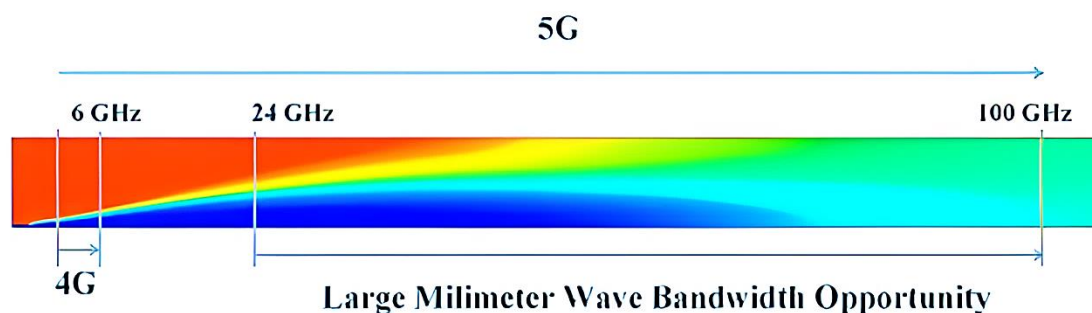


Figure I.6 : Représentation visuelle des ondes millimétriques [4].

Les millimètres d'onde (mmWave) de la 5G présentent trois bénéfices [4]:

- Une bande très récente et très peu fréquemment utilisée.
- Les signaux mmWave sont plus capables de transporter des données que les ondes de fréquence inférieure.
- Il est possible de combiner les mmWave avec des antennes MIMO, ce qui permet d'obtenir une capacité de transmission supérieure par rapport aux systèmes de communication actuels.

I.6.3.Le Massive MIMO :

Le Massive MIMO est une technologie cruciale pour les systèmes sans fil, permettant l'envoi et la réception simultanés de multiples signaux sur un même canal radio. Initialement utilisé pour son efficacité spectrale et énergétique élevée, le MIMO présentait des limitations en termes de débit et de fiabilité de la connexion. C'est là qu'intervient le mMIMO, une évolution de la technologie MIMO adaptée au contexte de la 5G. Cette approche implique l'ajout de centaines voire de milliers d'antennes aux stations de base afin d'augmenter le débit et l'efficacité spectrale. Le mMIMO utilise des antennes supplémentaires pour concentrer l'énergie dans des zones plus restreintes, augmentant ainsi l'efficacité spectrale et le débit. Cette technologie trouve des applications prometteuses dans divers domaines, notamment les voitures autonomes, les réseaux électriques intelligents, les villes intelligentes et les entreprises connectées [22].

I.6.4.Le Beamforming :

La technologie du beamforming, aussi appelée formation de faisceau, a été introduite par la 5G. C'est une méthode de traitement du signal qui permet de focaliser les ondes émises par une antenne vers un téléphone portable particulier, au lieu de diffuser le signal dans toutes les directions de manière uniforme. À la différence des antennes 4G qui arrosent leur

environnement avec le signal. En utilisant le beamforming, le signal est personnalisé en fonction des besoins particuliers de l'utilisateur, ce qui favorise une utilisation plus efficace de l'argent. Les antennes mMIMO 5G n'émettent qu'en cas de besoin, ce qui permet d'économiser une grande quantité d'énergie par rapport aux techniques de diffusion continue. En focaliser le signal [18].

Les antennes utilisées dans la 5G utilisent des fréquences plus élevées, à partir de 3,5 GHz, afin de mettre en place le principe du beamforming. La méthode consiste à regrouper de nombreuses antennes élémentaires pour concentrer l'énergie dans des faisceaux d'émission étroits, qui sont dirigés vers les utilisateurs au moment opportun. Cette méthode permet d'accroître la portée du signal, de réduire au minimum les interférences et d'augmenter le débit de données. À la différence des antennes classiques qui émettent constamment des ondes électromagnétiques dans toutes les directions, le beamforming permet une utilisation plus efficace de l'énergie, réduisant ainsi l'exposition aux champs électromagnétiques environnants des antennes relais. En dirigeant le signal vers des utilisateurs particuliers, cette méthode permet d'améliorer la qualité de connexion et d'améliorer l'expérience utilisateur dans les réseaux 5G [23].

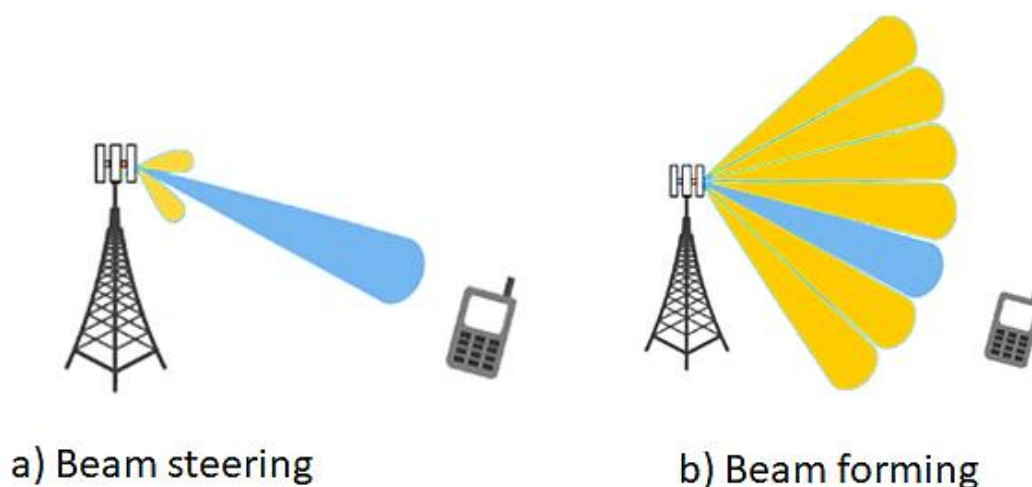


Figure I.7 : Illustration de ; **a)** Beam steering, **b)** Beam forming [22].

Une autre technologie d'application du système MIMO pour le réseau 5G est le mMIMO 3D qui utilise le beamforming 3D pour le transfert de trafic depuis des points d'accès intérieurs connectés à des matrices d'antennes de grande échelle montées sur des bâtiments ainsi que pour la connexion avec les utilisateurs sur différents niveaux de bâtiments de grande hauteur. Cependant, il existe de nombreux défis à surmonter avant que le mMIMO

puisse être appliqué dans le réseau 5G, notamment la fourniture de modèles de canal mMIMO massifs en ondes millimétriques pour divers scénarios [24].

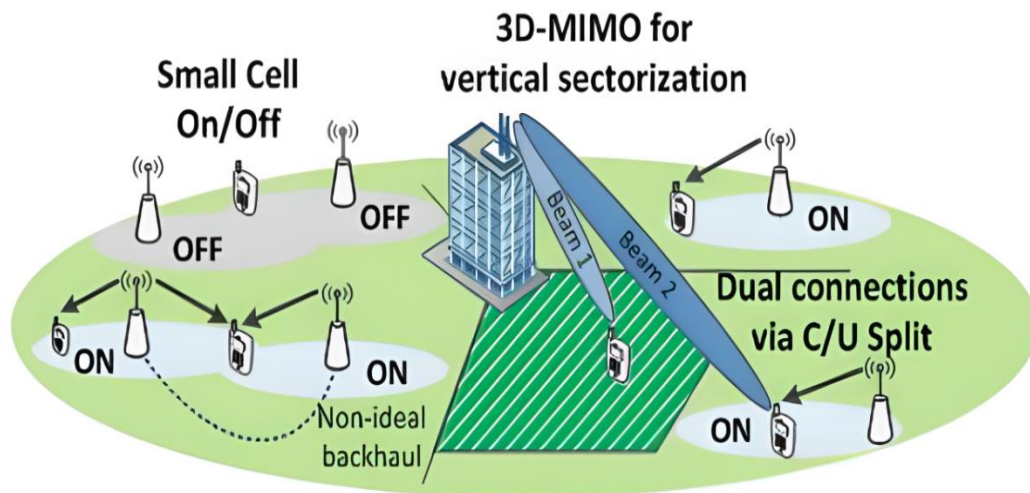


Figure I.8 : 3D MIMO pour la sectorisation verticale [24].

I.6.5.Network Slicing :

Il est essentiel de concevoir le réseau 5G afin de proposer une variété de capacités afin de satisfaire simultanément toutes ces exigences variées. De manière pratique, la méthode la plus logique consiste à élaborer un ensemble de réseaux spécifiques, chacun conçu pour répondre à un type spécifique de client professionnel.

Grâce à ces réseaux spécialisés, il serait possible de mettre en place des fonctionnalités personnalisées et d'exploiter un réseau adapté aux besoins de chaque client professionnel, plutôt que d'adopter une approche taille unique comme celle observée dans les générations de Smartphones actuelles et antérieures, qui ne serait pas économiquement rentable. L'utilisation de plusieurs réseaux dédiés sur une plate-forme commune est une méthode bien plus efficace : c'est précisément ce que permet le network slicing .Avec le découpage de réseau, le réseau 5G est capable de s'adapter à l'environnement externe plutôt que l'inverse [25].

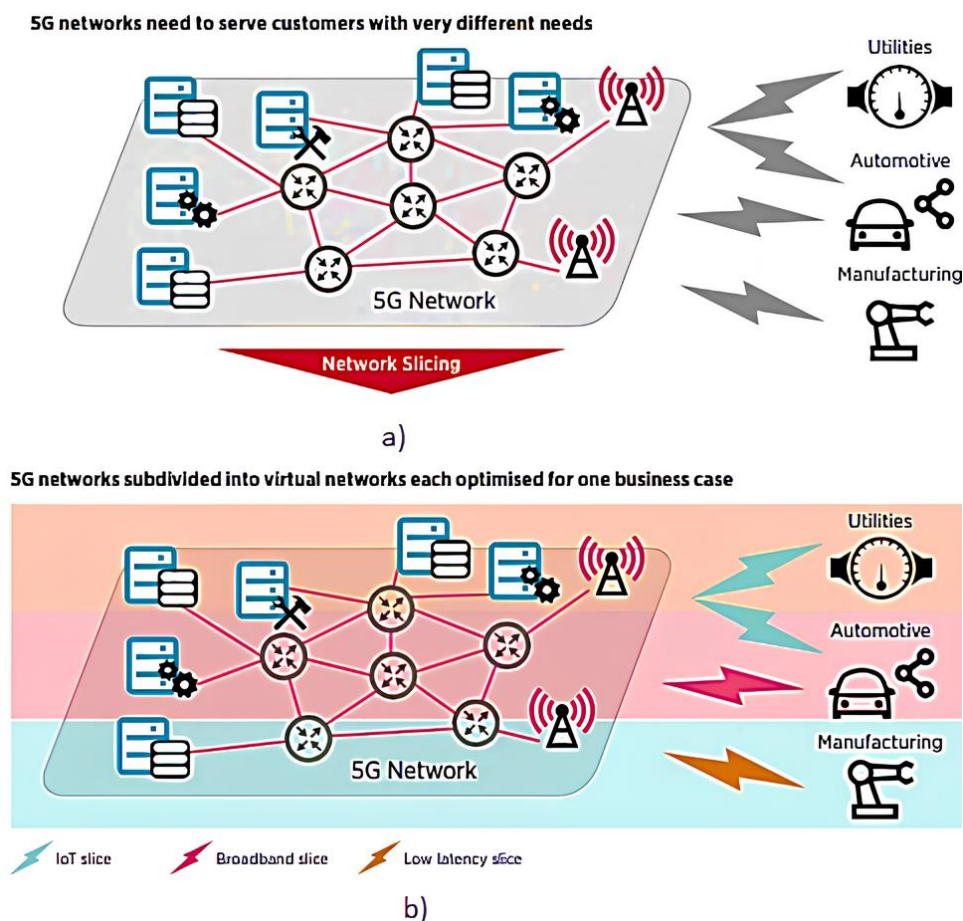


Figure I.9 : Les réseaux 5G ;a) sans le découpage de réseau et b) avec le découpage de réseau [25].

I.6.6.Small Cells:

Généralement, une cellule est une zone couverte par une antenne. Les cellules macro sont la principale composante des réseaux mobiles hérités. On retrouve généralement des cellules de grande taille posées sur un mât ou un toit dans les villes et les villages, le long des autoroutes ou sur des collines rurales. La portée de couverture radio des cellules macro varie de quelques kilomètres à plusieurs dizaines de kilomètres. Toutefois, la croissance de 1000 fois du volume de données mobiles ces dernières années a obligé les opérateurs à renforcer leur réseau. Dans cette optique, l'une des méthodes les plus performantes consiste à optimiser la réutilisation spatiale du spectre restreint en déployant de manière dense de petites cellules pour renforcer les réseaux cellulaires macro existants et avec une station de base cellulaire à haute puissance [26].

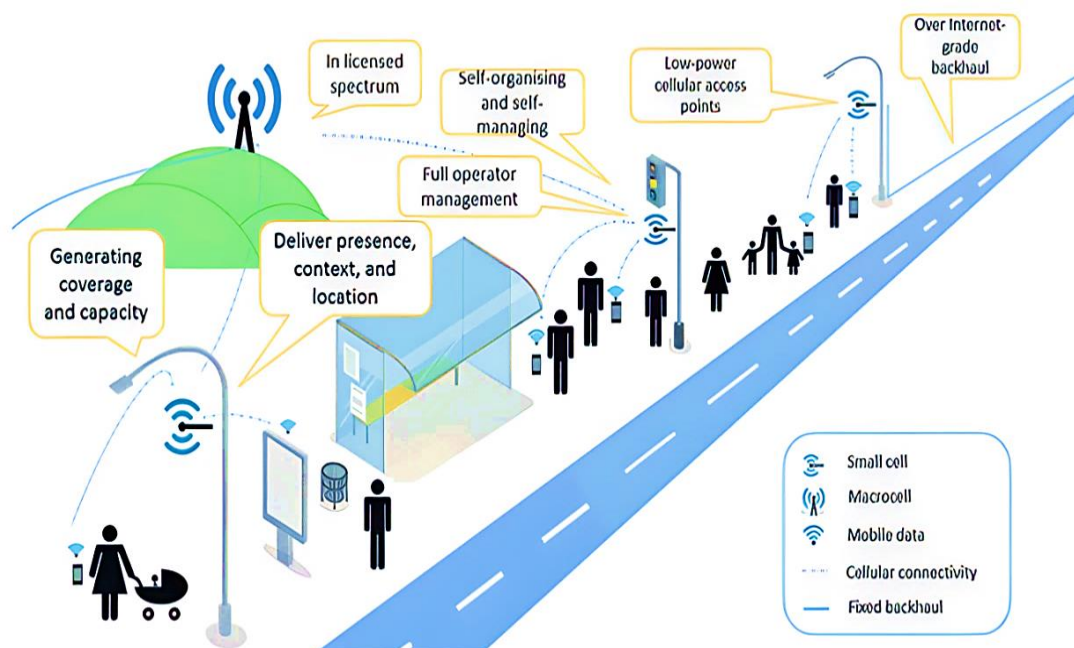


Figure I.10: Déploiement hétérogène de petites cellules et de cellules macro [26].

Les petites cellules ont joué un rôle crucial dans les mises à jour et l'expansion des réseaux 4G, mais elles seront encore plus importantes dans les réseaux 5G en raison de l'introduction de bandes de spectre plus élevées qui exigent des déploiements de réseau plus denses pour faire face à des volumes de trafic plus importants par unité de surface.

Les avantages indéniables des petites cellules pour la couverture réseau résident dans [26] :

- L'amélioration notable de la connectivité à l'intérieur des infrastructures telles que les édifices, ainsi que dans les espaces ouverts ou les régions moins accessibles.
- L'optimisation de l'efficacité spectrale constitue un autre pilier essentiel.
- La réduction des besoins en énergie.

I.6.7. Le full duplex :

Dans les systèmes conventionnels, la transmission et la réception se font soit sur des bandes de fréquences distinctes (FDD), soit à des moments différents (TDD). L'objectif du full duplex est de permettre la transmission et la réception simultanées d'informations sur les mêmes fréquences, au même moment et au même endroit [11].

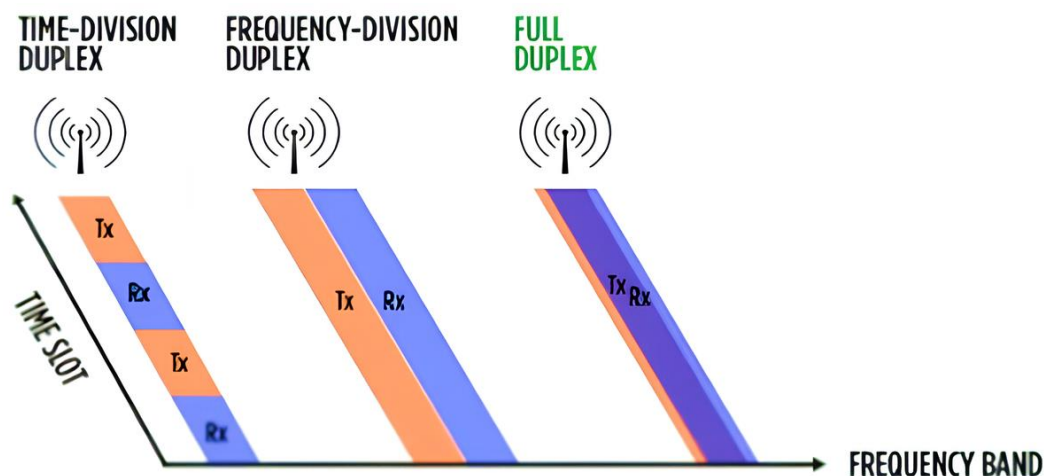


Figure I.11 : Illustration du full-duplex, comparé au FDD et TDD [11].

I.6.8. Non-Orthogonal Multiple Access (NOMA) :

Le NOMA joue un rôle crucial dans la 5G en favorisant une utilisation plus efficace des ressources spectrales. En connectant plusieurs utilisateurs à la même ressource, cela permet d'améliorer le débit global du réseau tout en offrant une gestion souple des ressources. Sa particularité réside dans sa capacité à supporter un grand nombre d'appareils connectés, ce qui le rend particulièrement adapté à la connectivité massive et aux applications IoT. En outre, le NOMA joue un rôle dans la performance énergétique des réseaux sans fil [4] [27].

I.6.9. Approches basées sur l'IoT pour la 5G :

La 5G offre une connexion Internet à un débit très élevé, offrant une grande souplesse et un accès à des bandes de spectre inutilisées, ce qui en fait la technologie la plus performante pour l'IoT. Grâce à l'IoT 5G, il est plus facile de communiquer entre machines et d'échanger des informations entre différents dispositifs sans nécessiter l'intervention humaine. Son utilisation ouvre de nouvelles opportunités dans le domaine des communications mobiles en offrant une connexion Internet à haut débit entre de nombreux appareils. En outre, l'Internet des objets 5G encourage la création de maisons intelligentes, d'appareils connectés, de capteurs, de systèmes de transport intelligents, d'industries automatisées et d'autres

applications, ce qui améliore l'intelligence des environnements pour les agents finaux. L'Internet des objets accepte différents appareils connectés à Internet, offrant ainsi une gamme variée. L'IoT supporte une multitude d'appareils connectés à Internet, offrant ainsi une large sélection d'applications et de services [28].

I.6.10.Virtualisation des fonctions réseau (NFV) :

Dans le contexte de la 5G, la NFV joue un rôle crucial. En permettant aux opérateurs de déployer des fonctions réseau en tant que logiciels plutôt que sur des équipements physiques dédiés, la NFV offre une flexibilité essentielle pour la mise en œuvre et la gestion des réseaux 5G. Cela permet aux opérateurs d'adapter rapidement leurs réseaux aux demandes changeantes des utilisateurs et des applications, tout en réduisant les coûts d'exploitation et en améliorant l'efficacité des ressources. En intégrant la NFV dans l'architecture 5G, les opérateurs peuvent offrir des services plus innovants et répondre plus efficacement aux besoins croissants en connectivité haut débit, faible latence et capacité de traitement élevée exigés par les applications émergentes telles que l'Internet des objets (IoT), la réalité virtuelle et augmentée (VR/AR), et les véhicules connectés [24].

I.6.11. Réseau défini par logiciel (SDN) :

Les architectures SDN divisent les fonctions de contrôle réseau et les fonctions de transfert de données, ce qui permet de programmer les fonctions de contrôle réseau et de gérer les applications et les services réseau par l'infrastructure réseau. En bref, les architectures SDN peuvent être classées en trois parties principales [29] :

- **Plan de données** : Il comprend les équipements réseau physiques qui acheminent les paquets de données à travers le réseau.
- **Plan de contrôle** : Il est centralisé et gère le comportement du réseau, prenant des décisions de routage et de transfert de données en fonction des politiques définies.
- **Plan d'application** : Il héberge les applications et services qui utilisent les fonctionnalités du réseau SDN pour fournir des services aux utilisateurs finaux.

I.6.12. La communication Machine à Machine (M2M) :

La M2M fait référence à la transmission de données entre des appareils réseau, comme des capteurs ou des dispositifs de surveillance, sans l'intervention d'un être humain. Elle est utilisée dans différents secteurs tels que les systèmes de transport intelligents, la santé, la surveillance des bâtiments, les pipelines de gaz et de pétrole, les systèmes de vente au détail intelligents et la sécurité. Toutefois, le développement de la communication M2M comporte divers défis à surmonter, tels que la connectivité des appareils massifs, la gestion des données en rafales, la latence nulle, la capacité à prendre en charge divers appareils, technologies et applications, la livraison rapide et fiable des messages, ainsi que le coût des appareils. En outre, des algorithmes performants doivent être élaborés afin de gérer la localisation, le temps, les groupes, les priorités et la transmission de données multi-sauts pour la communication M2M [4].

I.7.Les défis de la 5G :

La 5G marque une évolution significative dans le domaine des technologies de communication sans fil. Même si cette technologie présente de multiples bénéfices, elle comporte également de nombreux obstacles potentiels. En outre, de la même manière que les générations précédentes, les utilisateurs rencontrent divers problèmes lors de la mise en place de la 5G, tels que [31] :

- **Problèmes liés au spectre :**

Les anciennes générations de communication mobile à haut débit utilisent un spectre de 1340 à 1960 MHz. Cependant, l'arrivée de la technologie 5G entraîne l'anticipation d'une capacité de trafic jusqu'à 1000 fois plus élevée et d'un débit de données utilisateur entre 10 et 100 fois supérieur aux générations traditionnelles. Ainsi, il est essentiel d'augmenter considérablement le spectre et les bandes passantes afin de faciliter le fonctionnement des communications mobiles et sans fil. Autrement dit, la question du coût du spectre est également un enjeu majeur pour la mise en place de la 5G, avec une hausse prévue de plus de trois fois pour les pays en développement par rapport aux pays matures.

- **Coût de déploiement :**

Outre les difficultés liées au spectre, il est crucial de mettre en œuvre des infrastructures accessibles et flexibles pour les systèmes 5G, comme les poteaux, les tours, les antennes et les systèmes de fibres, afin de garantir la capacité et la couverture de la 5G. Autrement dit, cela nécessite des dépenses de développement et de maintenance importantes pour assurer le bon fonctionnement de l'ensemble du système, notamment en ce qui concerne l'infrastructure de fibre optique.

- **La sécurité :**

La 5G est une innovation technologique majeure après la transition à long terme (LTE). Elle se distingue par une nette amélioration de la vitesse de transmission des données sans fil, ainsi que la possibilité de créer des connexions large bande étendues et une mobilité accrue et plus rapide. En outre, la 5G propose une latence réduite et une meilleure qualité de service par rapport aux générations précédentes (comme la 1G à la 4G). De plus, elle supporte un large éventail d'appareils, tels que la communication entre machines (M2M), l'internet des objets (IoT) et les systèmes cyber-physiques. Toutefois, cette avancée technologique engendre aussi de nouveaux enjeux en termes de sécurité.

- **Politique :**

Toutes les technologies de 1G à 4G sont préoccupées par la confidentialité des communications mobiles et de télécommunication. Toutefois, diverses méthodes ont été suggérées afin de préserver la confidentialité des données et éviter les attaques des pirates informatiques. En outre, la mise en place d'un nouveau système structuré 5G pose également des défis de confidentialité plus importants pour les utilisateurs, les parties prenantes et les abonnés mobiles par rapport aux générations précédentes.

I.8. Architectures du Réseaux 5G :

1.8.1. Architectures à deux niveaux :

La 5G, marque une avancée significative dans le domaine des communications sans fil. La 5G est en train de transformer notre manière d'interagir avec la technologie et les services numériques en proposant des débits de données plus rapides, une latence réduite et une connectivité massive. La transformation repose sur une architecture réseau novatrice, élaborée pour satisfaire les besoins grandissants de l'ère numérique. Plusieurs éléments technologiques sont à la base de cette architecture, tels que les petites cellules, les communications entre dispositifs (D2D), les réseaux cognitifs (CRN) et bien d'autres encore. En collaboration, ces éléments constituent un réseau flexible et adaptable, capable de répondre à différentes situations d'utilisation, allant des services haut débit aux services à haut débit aux applications à faible latence en passant par les besoins de connectivité massive [4] .

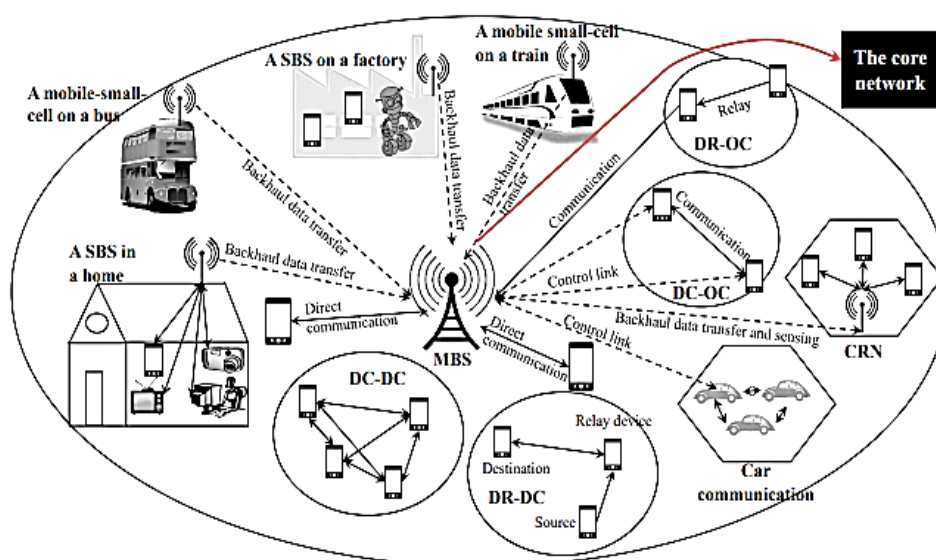


Figure I.12: Une architecture multi-niveaux pour les réseaux 5G avec des petites cellules, des petites cellules mobiles, et des communications basées sur D2D et CRN [4].

Plusieurs architectures à deux niveaux ont été proposées pour les réseaux 5G, où un MBS reste dans le niveau supérieur et les SBS fonctionnent sous la supervision du MBS dans le niveau inférieur. Une macro cellule couvre toutes les petites cellules de différents types, par exemple, la femto cellule, la pico cellule et la microcellule, et les deux niveaux partagent une bande de fréquence identique. La petite cellule améliore la couverture et les services d'une macro cellule, et les avantages des petites cellules sont mentionnés déjà dans la section

précédente. De plus, la communication D2D et la communication basée sur CRN améliorent une architecture à deux niveaux en une architecture à plusieurs niveaux [4].

Type de cellule	Puissance de sortie (W)	Rayon de cellule (km)	Nombre d'utilisateurs	Emplacements
Pico Cellule	0.25 à 1	0.1 à 0.2	30 à 100	Intérieur/Extérieur
Micro Cellule	1 à 10	0.2 à 20	100 à 2000	Intérieur/Extérieur
Femto Cellule	<0.01	Jusqu'à 10mètres	1 à 4	Intérieur
Macro Cellule	10 à >50	8 à 30	>2000	Extérieur

Tableau I.3 : Différents types de cellules.

I.8.2.Le réseau de radio cognitive (CRN) :

Un réseau de radio cognitive (CRN) se compose d'un ensemble de nœuds de radio cognitive (ou processeurs), également connus sous le nom d'utilisateurs secondaires (SUs), qui exploitent de manière opportuniste le spectre déjà existant [29]. Les SUs ont la capacité de s'adapter à différents canaux hétérogènes (ou bandes de fréquences) en l'absence des utilisateurs autorisés (PUs), également connus sous le nom d'utilisateurs primaires. Chaque système de puissance (PU) possède une largeur de bande constante, une puissance d'émission élevée et une fiabilité élevée. Cependant, les SUs sont disponibles dans une variété de largeurs de bande avec une faible puissance d'émission et une haute fiabilité. Dans les réseaux 5G, un CRN est employé afin de développer des architectures multi-niveaux, réduire les interférences entre les cellules et réduire la consommation d'énergie dans le réseau [4].

I.8.3.La communication de dispositif à dispositif (D2D) :

La technologie de communication de dispositif à dispositif (D2D) permet aux appareils d'utilisateurs proches de communiquer directement entre eux sans nécessiter de passer par une station de base centrale. Il est possible de diminuer la charge sur les stations de base, d'améliorer la fiabilité des communications et d'explorer de nouvelles opportunités d'application grâce à cette approche. Néanmoins, sa réalisation demande la résolution de problèmes tels que la gestion des intrusions et la sécurité des données partagées. En dépit de ces difficultés, la D2D possède un énorme potentiel pour optimiser l'efficacité des réseaux sans fil et apporter un soutien à de nouvelles applications [24].

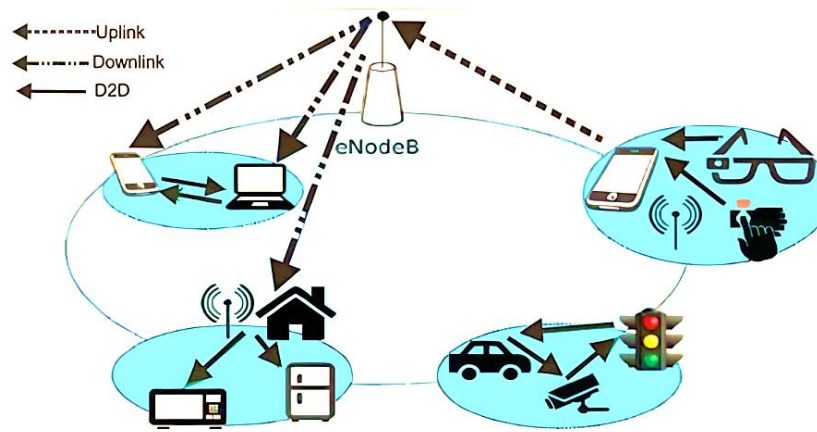


Figure I.13 : L'architecture de communication D2D [24].

I.8.4. Architectures basées sur le Cloud :

Le C-RAN est une structure de réseau qui regroupe les ressources de bande de base pour permettre leur partage entre les stations de base. L'idée principale derrière tout C-RAN est de réaliser la plupart des fonctions d'une MBS dans le cloud, ce qui permet de séparer la fonctionnalité d'une MBS en une couche de contrôle et une couche de données. Les différentes couches de contrôle et de données jouent leurs rôles respectivement dans un cloud et dans une MBS. En particulier, une MBS se compose de deux composants : une unité de bande de base (BBU) pour le traitement de bande de base à l'aide de processeurs de bande de base et une tête de radio (RRH) pour les tâches radio. Dans la plupart des C-RAN, les BBU sont généralement stockées dans le cloud, tandis que les RRH sont conservées dans les MBS. Ainsi, un C-RAN offre une architecture facile à adapter et flexible [4].

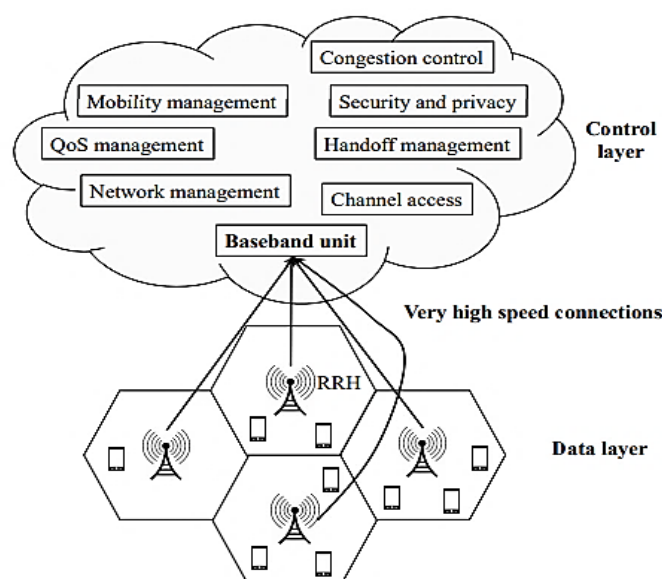


Figure I.14 : Une architecture basée sur le Cloud pour les réseaux 5G [4].

I.9-Conclusion :

Ce premier chapitre constitue une introduction aux réseaux mobiles, en mettant particulièrement l'accent sur la dernière évolution majeure de la 5G. Nous avons étudié les principaux objectifs de la 5G, qui comprennent l'amélioration des performances, une connectivité ultra-rapide et une capacité accrue pour satisfaire les divers besoins des applications émergentes les plus diverses. De plus, nous avons abordé les multiples utilisations possibles de la 5G, soulignant sa polyvalence et sa capacité à soutenir des applications allant de l'IoT aux véhicules connectés. Sur le plan architectural, nous avons étudié les concepts et les éléments clés de la 5G.

Le chapitre suivant sera consacré aux techniques d'accès multiple orthogonal et non orthogonal (OMA /NOMA), avec un accent particulier sur SCMA comme une technique spécifique utilisée dans certains systèmes NOMA.

CHAPITRE II :

*Les techniques d'accès multiple
(OMA /NOMA)*

II.1.Introduction :

Les techniques d'accès multiple sont essentielles pour maximiser l'efficacité et l'utilisation des ressources dans les réseaux de communication sans fil. Deux approches principales dans ce domaine sont OMA et NOMA jouent des rôles clés dans l'évolution des systèmes de communication, notamment dans le cadre de la 5G et au-delà.

Ce chapitre explore les concepts fondamentaux d'OMA et NOMA, mettant en évidence leurs principes de fonctionnement et leurs différences essentielles offrant un aperçu des technologies qui façonnent le futur des réseaux de télécommunication.

II.2.Les techniques d'accès multiple :

L'accès multiple (MA) est un élément essentiel dans le domaine des communications sans fil, car il permet à plusieurs utilisateurs d'accéder aux ressources du réseau sans interruption. Il est possible que ces ressources comprennent des gammes de fréquences, des périodes de temps ou des codes de communication. Le terme « systèmes d'accès multiple » désigne le processus de fusion des signaux de plusieurs utilisateurs. Ensuite, ce signal combiné est transmis via un média commun, comme le montre la **Figure II .1 [16]**.

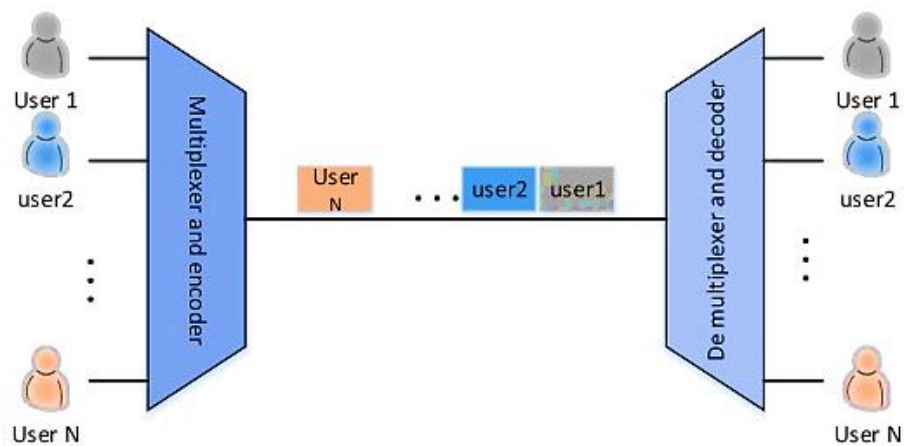


Figure II.1: Le concept d'un système d'accès multiple [16].

On peut distinguer deux catégories principales d'accès multiple, OMA et NOMA. Chaque méthode présente des bénéfices particuliers en ce qui concerne l'efficacité spectrale, la capacité du réseau et les performances globales [30]. Dans cette partie, nous examinerons ces deux types d'accès multiples et les diverses méthodes qui les constituent.

II.3.L'accès multiple orthogonal (OMA) :

L'OMA désigne un ensemble de techniques utilisées pour permettre à plusieurs utilisateurs d'accéder simultanément à un médium de communication partagé sans causer d'interférences. En OMA, chaque utilisateur se voit attribuer une ressource orthogonale unique, telle que la fréquence, le temps ou le code, pour transmettre ses données. Cela garantit que les utilisateurs peuvent communiquer simultanément sans interférer avec les signaux des autres. Des exemples de schémas OMA incluent FDMA dans la première génération 1G, la TDMA avec l'avènement de 2G et en 3G, la CDMA a émergé. Lorsque le nombre d'utilisateurs en 4G a considérablement augmenté, un modèle plus solide et efficace a été proposé, OFDMA [31]. Dans la suite, nous procéderons à une brève présentation de chacun de ces types d'accès.

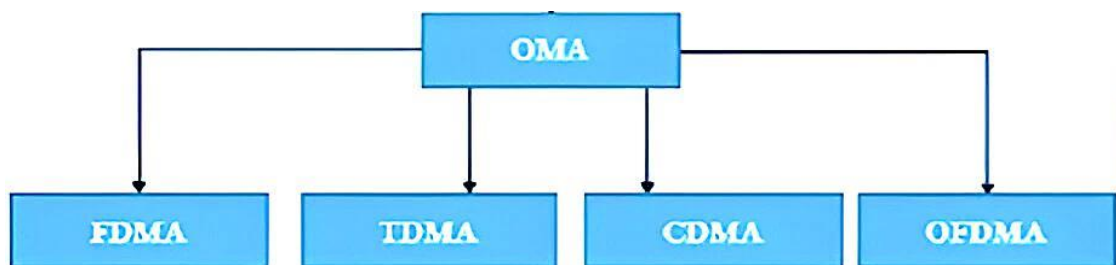


Figure II.2 : Les techniques d'accès OMA [30].

II.3.1.FDMA:

En FDMA, chaque utilisateur se voit attribuer une bande de fréquences unique dans la largeur de bande totale, et il utilise exclusivement cette bande pour sa communication. Les différents utilisateurs sont séparés dans le domaine fréquentiel. Cette séparation basée sur la fréquence garantit que les différents utilisateurs n'interfèrent pas les uns avec les autres, permettant des transmissions simultanées [32].

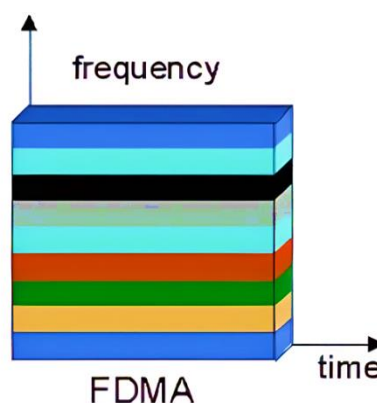


Figure II.3 : Le principe de technique d'accès multiple FDMA [32].

II.3.2.TDMA:

Dans le TDMA, les utilisateurs se voient attribuer des intervalles de temps spécifiques, appelés créneaux temporels, pendant lesquels ils sont autorisés à transmettre des données. Chaque utilisateur utilise toute la largeur de bande de fréquence pendant son créneau temporel attribué, avec une séparation entre les utilisateurs réalisée dans le domaine temporel. Par conséquent, différents utilisateurs transmettent à des moments différents, garantissant qu'un seul utilisateur est actif sur le canal à un moment donné. Cette approche basée sur le temps permet un partage efficace du support de communication entre plusieurs utilisateurs [32].

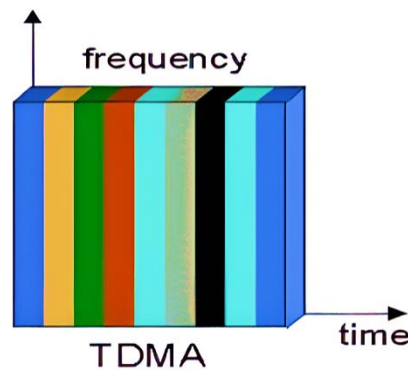


Figure II.4 : Le principe de technique d'accès multiple TDMA [32].

II.3.3.CDMA:

Dans la technique CDMA, chaque utilisateur se voit attribuer une séquence de codes distincte, appelée code de propagation, qu'il utilise pour encoder son signal de données. À la réception, le récepteur, connaissant la séquence de codes de l'utilisateur, décode le signal reçu et récupère ainsi les données originales [32].

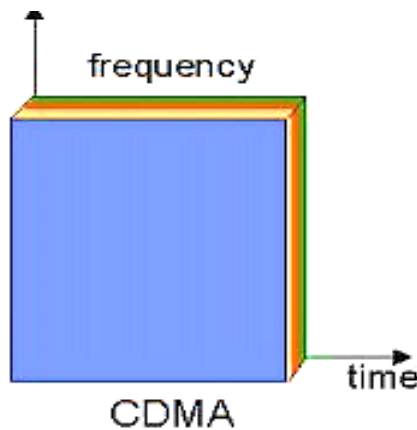


Figure II.5: Le principe de technique d'accès multiple CDMA [32].

II.3.4.OFDMA :

Cette technique permet une utilisation efficace du spectre en permettant à plusieurs utilisateurs de transmettre simultanément sur différentes sous-porteuses sans interférence, car les sous-porteuses sont orthogonales les unes par rapport aux autres. Les ressources radio sont réparties à la fois dans les dimensions temporelles et fréquentielles, ce qui signifie que les ressources radio sont présentées sous la forme d'une grille temps-fréquence [33].

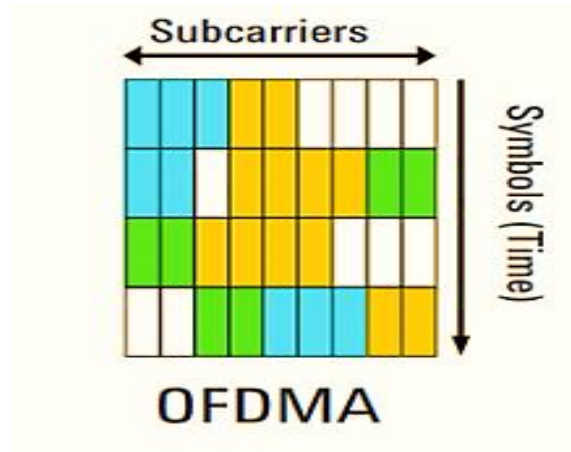


Figure II.6: Le principe de technique d'accès multiple OFDMA [33].

II.4.L'accès multiple non orthogonal (NOMA) :

Au cœur de l'approche radio de la prochaine génération (5G) réside le principe de NOMA. Comme son nom l'indique, la NOMA vise à servir simultanément plusieurs utilisateurs sur le même bloc de ressources. Contrairement à l'approche précédente (OMA), où les ressources sont allouées de manière exclusive à chaque utilisateur [34].

NOMA présente plusieurs subdivisions qui visent à optimiser l'utilisation des ressources spectrales et à améliorer les performances des réseaux. Parmi ces subdivisions, on trouve le NOMA du domaine de puissance (PD-NOMA), où les utilisateurs se voient attribuer différents niveaux de puissance pour partager la même bande passante, et le NOMA du domaine de code (CD-NOMA), où chaque utilisateur utilise un code de modulation unique pour éviter les interférences. Il existe plusieurs techniques d'accès NOMA comme il est représenté dans la **Figure II.7**.

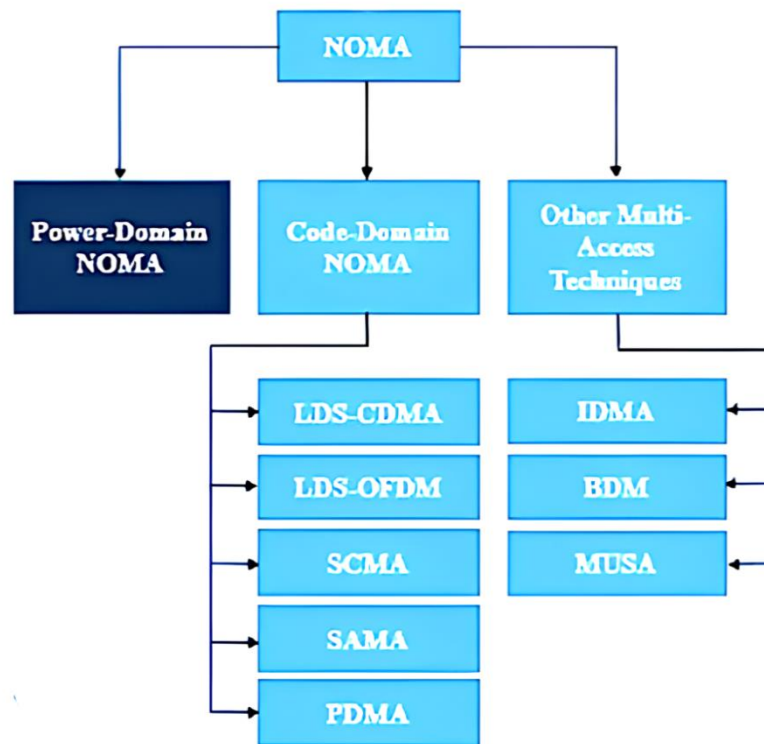


Figure II.7 : Classification des techniques d'accès NOMA [30].

II.4.1.Principales Technologies de NOMA :

La possibilité de NOMA repose sur deux opérations essentielles : le codage par superposition (SC) qui doit être réalisé lors de l'émission, et l'annulation successive des interférences (SIC) à la réception. Ces deux opérations seront brièvement abordées dans la section suivante.

II .4.1.1.Le codage par superposition (SC) :

Le SC consiste à transmettre en même temps des informations à plusieurs récepteurs à partir d'une seule source. Autrement dit, il permet à l'émetteur de communiquer les données de plusieurs utilisateurs simultanément [35].

Prenons un cas où $K = \{\text{utilisateur} - N, \text{utilisateur} - F\}$ est le nombre total d'utilisateurs où utilisateur-N est l'utilisateur proche et utilisateur-F est l'utilisateur distant. L'utilisateur-N et l'utilisateur-F, respectivement, sont représentés par des symboles complexes, x_N et x_F , dans leur constellation de modulation M_{aire} appropriée, avec une puissance totale transmise P_s et les coefficients d'allocation de puissance α_N et α_F par utilisateur-N et utilisateur-F, respectivement. Dans cette situation, les deux utilisateurs seront tous deux utilisant la modulation QPSK.

Dans le SC, on superpose la constellation de l'utilisateur-F (avec une puissance de transmission plus élevée) à celle de l'utilisateur-N (avec une puissance de transmission plus faible) afin de générer une constellation mappée, comme il est indiqué dans la **Figure II.8 [36]** :

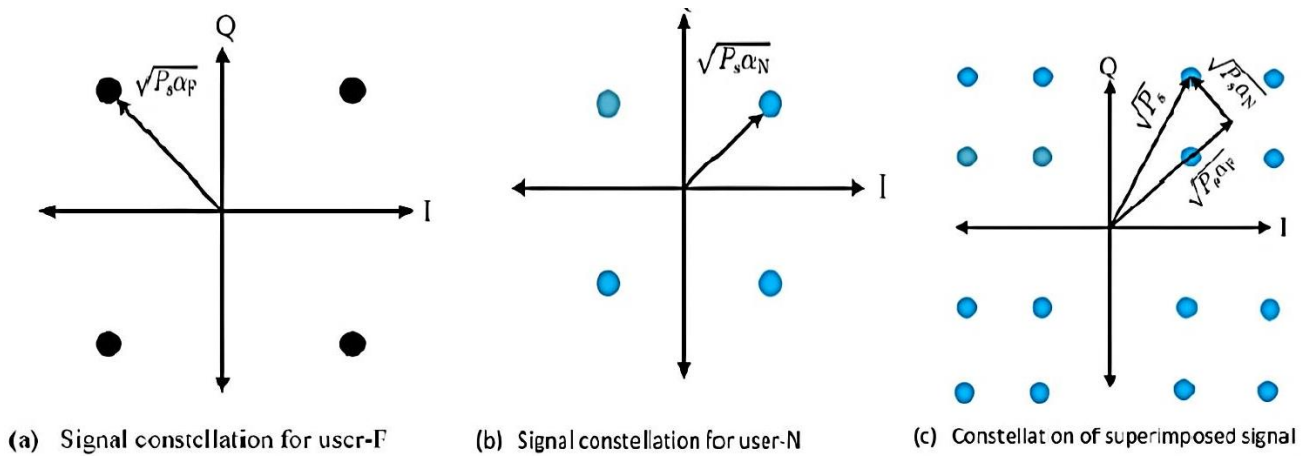


Figure II.8 : Diagramme de constellation hiérarchique de l'exemple de codage par superposition (SC) [36].

II.4.1.2.L'annulation successive des interférences (SIC) :

La SIC repose sur l'idée que les signaux des utilisateurs sont décodés de manière progressive. Une fois que le signal d'un utilisateur est décodé, il est retiré du signal combiné avant de décoder le signal de l'utilisateur suivant. Quand la SIC est mise en place, le signal de l'un des utilisateurs est décodé, en considérant le signal de l'autre utilisateur comme un perturbateur, mais ce dernier est décodé avec l'avantage du signal de l'autre ayant déjà été supprimé précédemment. Toutefois, avant l'utilisation de la SIC, les utilisateurs sont classés en fonction de la puissance de leur signal, afin que le récepteur puisse décoder le signal le plus fort en premier, le supprimer du signal combiné et isoler le signal le plus faible du reste. Il est important de souligner que chaque utilisateur est interprété en considérant les autres utilisateurs interférents comme du bruit lors de la réception du signal [35]. Le processus SIC d'un système NOMA est illustré dans la **Figure II.9**.

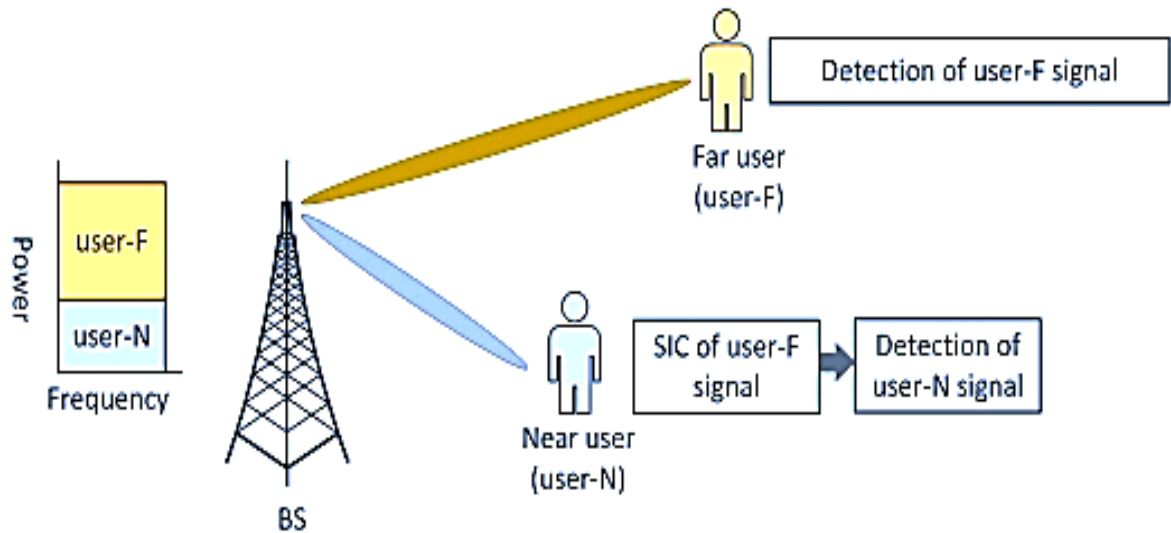


Figure II.9 : Représentation d'un système NOMA à deux utilisateurs, où l'utilisateur-N assiste l'utilisateur-F [36].

II.4.2. Schémas de NOMA :

Les domaines de NOMA peuvent être catégorisés en deux principaux types : le PD-NOMA et le CD-NOMA.

II.4.2.1.PD-NOMA :

Dans le domaine de la puissance, le concept de NOMA permet de fournir un service à plusieurs utilisateurs en utilisant le même créneau temporel ou le même code de diffusion, et l'accès multiple peut être réalisé en allongeant différentes quantités de puissance à différents utilisateurs. Les utilisateurs ayant des gains de canal plus élevés sont désignés comme "utilisateurs forts" dans le PD-NOMA, tandis que ceux ayant des gains de canal plus faibles sont désignés comme "utilisateurs faibles". En utilisant la SIC, l'utilisateur fort supprime le signal de l'utilisateur faible, puis détecte son propre signal. Le signal de l'utilisateur fort est perçu par l'utilisateur faible comme un bruit d'interférence afin de repérer son propre signal [31].

II.4.2.1.1.Down Link:

La **Figure II.10** montre un exemple de transmission NOMA en liaison descendante. Après la signalisation NOMA, une BS transmet un mélange de signaux $(\alpha_1 s_1 + \alpha_2 s_2)$ à l'utilisateur U_1 (condition de canal forte, utilisateur fort) et à l'utilisateur U_2 (condition de canal faible, utilisateur faible), s_1 et s_2 étant des signaux destinés aux deux utilisateurs, α_1 et α_2 étant les coefficients d'allocation de puissance correspondants avec $\alpha_1^2 + \alpha_2^2 = 1$ et $\alpha_1 < \alpha_2$.

Chez U_1 , on procède à la SIC afin de supprimer s_2 et ensuite de récupérer son propre s_1 . Chez U_2 , on interprète directement son signal s_2 en considérant s_1 comme une interférence [37].

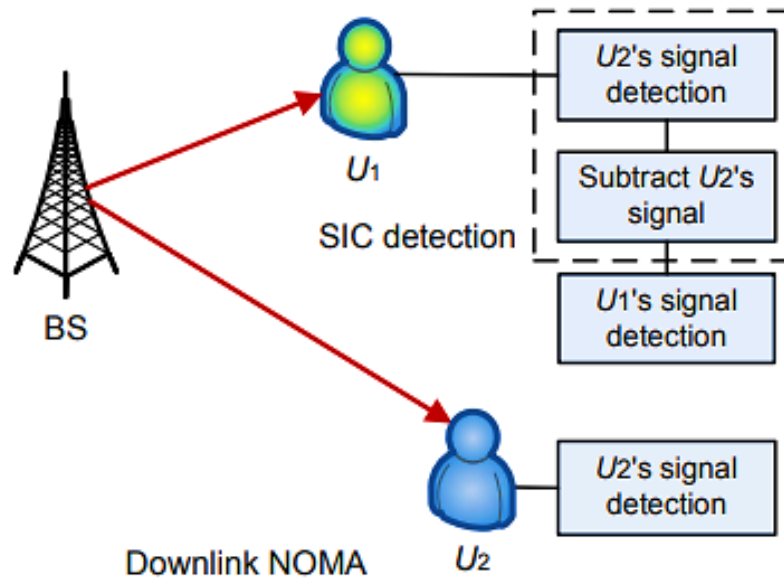


Figure II.10 : illustration d'un système NOMA en liaison descendante (PD-Down Link) [37].

II.4.2.1.2 .Up Link :

En ce qui concerne la transmission NOMA en liaison montante, comme le montre la **Figure II.11**, il est habituel que la BS envoie un message de contrôle contenant des informations sur l'allocation de puissance à l'étape initiale. Par la suite, U_1 (utilisateur fort avec condition de canal forte) et U_2 (utilisateur faible avec condition de canal faible) transmettent des signaux s_1 et s_2 souhaités à la BS en utilisant des niveaux de puissance a_1 et a_2 , respectivement. Après avoir reçu ces signaux, on procède à l'annulation d'interférences successives (SIC) à la BS, et un ordre de détection optimal consisterait à décoder d'abord le signal le plus fort avant de se déplacer vers le signal le plus faible. Pour le traitement SIC, il est recommandé d'avoir une force de puissance reçue différente. Par conséquent, nous optons pour $\alpha_1 > \alpha_2$ dans la transmission NOMA en liaison montante, afin que la BS décrypte d'abord s_1 puis l'annule afin de récupérer s_2 . Ainsi, le faible utilisateur (U_2) profite d'une transmission sans interférence, tandis que le fort utilisateur (U_1) constate une interférence [37].

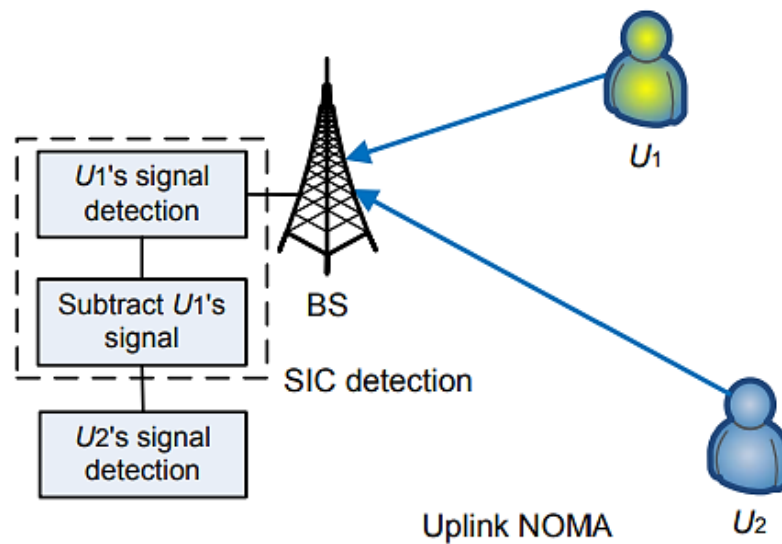


Figure II.11 : Une illustration de système NOMA en liaison montante (PD-Up Link) [37].

II.4.2.2.CD-NOMA :

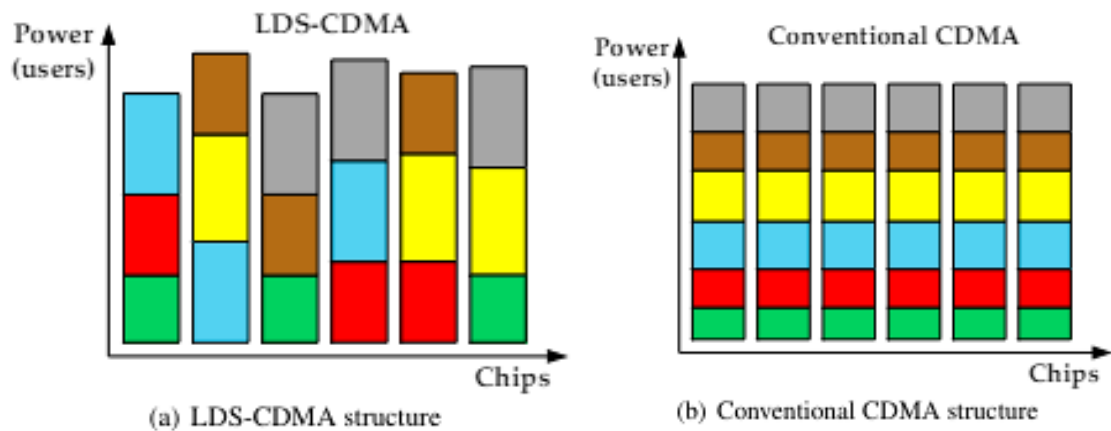
L'autre domaine de multiplexage peut être basé sur le code de signature, qui est appelé CD-NOMA. Ce schéma peut prendre en charge un grand nombre de transmissions d'utilisateurs dans le même bloc de ressources temps/fréquence en attribuant différents codes de signature à différents utilisateurs [31]. Dans le CD-NOMA, plusieurs sous-catégories ont émergé, chacune avec ses propres caractéristiques et applications spécifiques. Nous allons explorer ces sous-catégories dans la section suivante :

II.4.2.2.1.LDS-CDMA :

La technique LDS consiste principalement à désactiver un grand nombre de puces de code de diffusion, ce qui transforme le code en un vecteur clairsemé. Autrement dit, chaque utilisateur transmettra ses informations sur un nombre limité de chips. Par la suite, on organise les chips de manière astucieuse de sorte que le nombre d'interférences par chips soit nettement inférieur au nombre total d'utilisateurs. Le modèle de système LDS-CDMA est représenté dans la **Figure II.12**. Cette structure LDS apporte les avantages suivants [38]:

- Un SINR au niveau de la puce plus élevé peut être atteint, ce qui conduit à un meilleur processus de détection.
- À chaque chips reçue, un utilisateur aura un nombre relativement faible d'interférences, de sorte que l'espace de recherche devrait être plus petit et, par conséquent, des techniques de détection plus abordables peuvent être utilisées.

- Toutes les chips d'un utilisateur seront confrontées à une interférence provenant de différents utilisateurs, ce qui crée une variété d'interférences en évitant que des interférences fortes ne corrompent toutes les chips d'un utilisateur.



II .4.2.2.2.MUSA :

MUSA est un autre schéma NOMA reposant simplement sur le multiplexage dans le domaine des codes, qui peut être considéré comme une amélioration du schéma de style CDMA.

Dans la liaison montante du système MUSA, tous les symboles transmis par un utilisateur spécifique sont multipliés par la même séquence de diffusion. (Notez que différentes séquences de diffusion peuvent également être utilisées pour différents symboles du même utilisateur, ce qui entraîne une moyenne d'interférences bénéfique) [39].

Dans la liaison descendante du système MUSA, les utilisateurs sont séparés en G groupes. Dans chaque groupe, les symboles des différents utilisateurs sont pondérés par des coefficients de mise à l'échelle de puissance différents, puis ils sont superposés. Des séquences orthogonales de longueur G peuvent être utilisées comme séquences de diffusion pour étaler les symboles superposés des G groupes. Plus précisément, les utilisateurs du même groupe utilisent la même séquence de diffusion, tandis que les séquences de diffusion sont orthogonales entre les différents groupes. De cette manière, les interférences intergroupes peuvent être éliminées au récepteur. Ensuite, la SIC peut être utilisée pour effectuer l'annulation des interférences intra-groupe en exploitant la différence de puissance associée [39].

MUSA utilise un code de diffusion complexe comme illustré dans la **Figure II.13**. Le nombre de valeurs disponibles pour chaque élément du code de diffusion complexe illustré dans la **Figure II.13 (a)** est de 4. Pour une longueur de code L , le nombre total de codes disponibles est de 4^L . Par exemple, si la longueur du code est de 4, le nombre total de codes disponibles est de 256, ce qui peut ne pas être suffisant. Ainsi, l'ensemble de valeurs disponibles de la partie réelle et de la partie imaginaire pourrait être considéré pour être étendu afin d'augmenter davantage le nombre de codes disponibles. Un choix préféré est $M = 3$, comme illustré dans la **Figure II.13 (b)**, l'ensemble de valeurs disponibles de la partie réelle et de la partie imaginaire est un ensemble ternaire comprenant 1, 0 et -1, donc chaque élément du code de diffusion complexe est issu de l'ensemble $\{0, 1, 1 + i, i, -1 + i, -1, -1 - i, -i, 1 - i\}$ avant normalisation. Sur cette base, 9^L codes pourraient être générés [40].

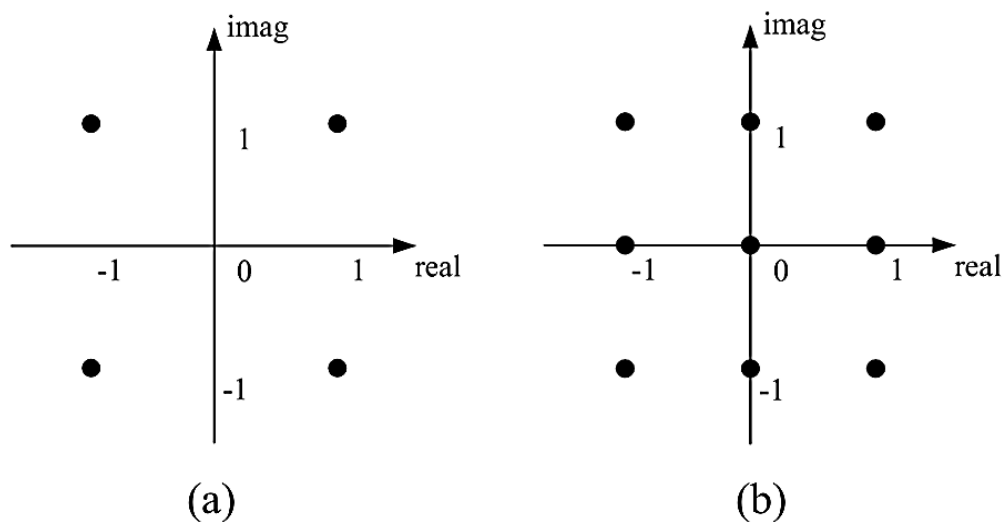


Figure II.13 : Éléments du code de diffusion complexe MUSA [40].

II.4.2.2.3.PDMA :

Le PDMA s'appuie sur la représentation cartographique des informations des utilisateurs qui doivent être transmises dans un ensemble de ressources en fonction d'un motif spécifique. Une ressource peut prendre différentes dimensions telles que la puissance, le code, le temps ou le spectre. Dans cette situation, on définit l'ordre de transmission comme le nombre de ressources rassemblées. De cette façon, on réalise une transmission non orthogonale car plusieurs utilisateurs sont multiplexés sur une même ressource. La **Figure II.14** illustre la situation où 6 utilisateurs ont 4 éléments de ressource communs. Par exemple, l'utilisateur 1 dispose de quatre éléments de ressources et partage les éléments de ressource

(RE) 1 avec l'utilisateur 2, l'utilisateur 3 et l'utilisateur 4. Les 6 utilisateurs ont une diversité de transmission de 4, 3, 2,2, 1,1 respectivement.

Pour détecter les utilisateurs multiples, différents algorithmes peuvent être employés afin de séparer les données des utilisateurs multiplexées sur le même RE, comme l'annulation successive d'interférences (SIC) [41].



Figure II.14 : Illustration d'un exemple de motif PDMA [41].

II.4.2.2.4.GOCA :

GOCA fait partie du domaine de NOMA qui repose sur des séquences. L'utilisation d'une structure à deux niveaux permet aux utilisateurs de créer des séquences orthogonales de groupe, puis de diffuser des symboles modulés dans des ressources temporelles et fréquentielles communes. Un exemple de génération de séquence dans l'émetteur GOCA est illustré dans la **Figure II.15**. La première et la deuxième étape utilisent respectivement une séquence orthogonale et une séquence non orthogonale. En outre, il est possible de subdiviser les séquences GOCA en groupes et chaque groupe possède un ensemble de séquences orthogonales identique et une séquence non orthogonale distincte [42].

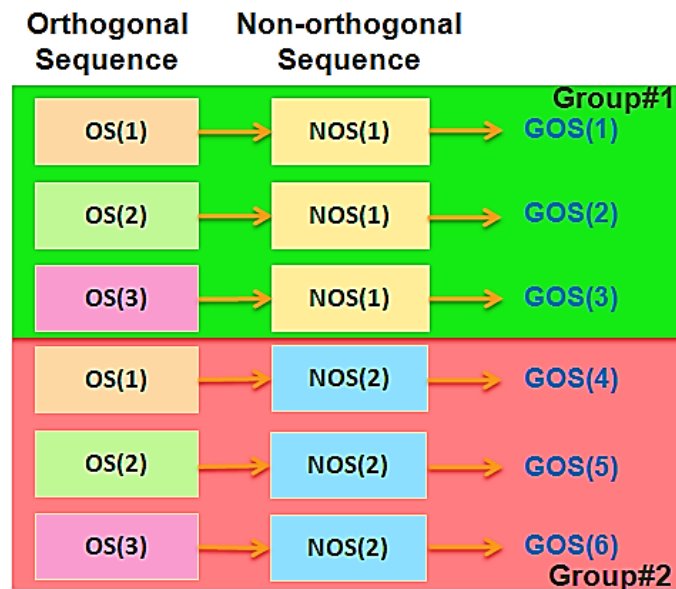


Figure II.15 : Illustration d'un exemple de génération de séquence dans l'émetteur GOCA [42].

II.4.2.2.5.SCMA :

On peut définir SCMA comme un modèle d'accès multiple par répartition de code s'appuyant sur l'idée initiale du LDS-CDMA, qui se distingue par des codebooks clairement définis [43]. Le terme "sparse" ou "épars" dans SCMA fait référence au fait que les codebooks utilisés dans le système sont conçus de manière à ce que chaque utilisateur ne dispose que d'un nombre limité de codewords par rapport à l'ensemble des ressources disponibles. Cette caractéristique de sparsité est importante car elle permet une utilisation efficace des ressources tout en réduisant la complexité du système [44].

La **Figure II.16** présente un cas d'accès multiple de 6 utilisateurs utilisant les codebooks spécifiques à la couche SCMA. Chaque utilisateur dispose d'un codebook SCMA (par exemple, l'utilisateur i utilise le codebook pour la couche i , $i = 1, 2, \dots, 6$). Les bits codés de chaque utilisateur sont ensuite alignés sur le mot de code SCMA après l'utilisation du codeur FEC. Par la suite, les mots de code SCMA sont combinés avec les tons OFDM et les symboles sont transmis en utilisant des blocs SCMA. [44].

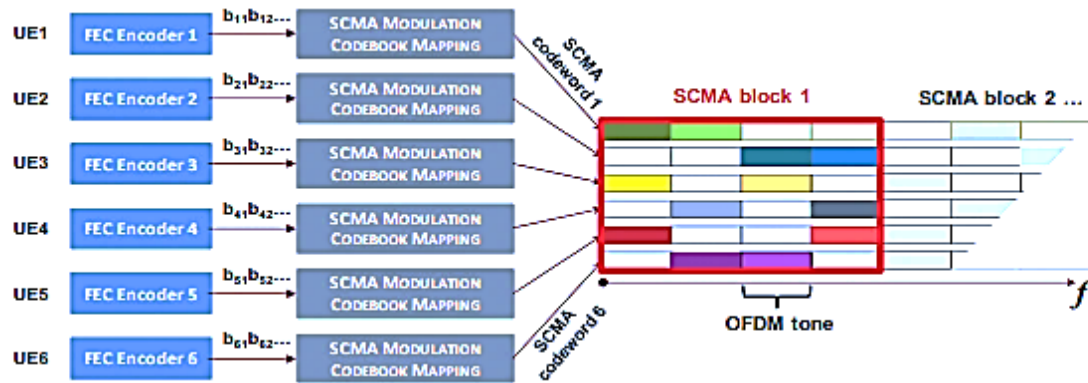


Figure II.16 : L'accès multiple SCMA [44].

L'accès multiple avec SCMA présente les principales caractéristiques suivantes [44] :

- Dans le domaine des codes, la superposition de signal non orthogonale permet de superposer plusieurs symboles provenant de différents utilisateurs sur chaque élément de ressource (RE).
- L'utilisation d'un étalement éparé par SCMA permet de diminuer le nombre de collisions de symboles.
- La modulation multidimensionnelle est utilisée dans SCMA, plutôt que l'étalement linéaire comme dans CDMA.

A) Le système SCMA :

La **Figure II.17** présente le modèle de système SCMA en liaison descendante. On peut définir un système SCMA avec K éléments de ressources et J utilisateurs comme SCMA (K, J) , $J \geq K$. Supposons que chaque utilisateur utilise N ressources spécifiques, $(1 \leq N \leq K)$ pour distinguer les utilisateurs, aucun ne peut utiliser les mêmes N ressources. Afin d'évaluer la capacité de charge d'un système SCMA, nous utilisons la formule $\lambda = J/K$ pour déterminer le taux de charge.

On crée une matrice de mappage $F(K, J)$ afin de représenter la répartition des ressources entre les utilisateurs. L'utilisateur est représenté par les colonnes d'une matrice de mappage, tandis que les éléments de ressources sont représentés par les lignes. $f_k^{(j)}$ peut-être de 0 ou de 1, $1 \leq k \leq K$, $1 \leq j \leq J$. Si $f_k^{(j)} = 1$, cela implique que l'utilisateur U_j utilise la ressource R_k . La relation entre les ressources et les utilisateurs dans SCMA est illustrée par la matrice $F(4,6)$ [45].

$$F_{4 \times 6} = \begin{bmatrix} 1 & 1 & 1 & 0 & 0 & 0 \\ 1 & 0 & 0 & 1 & 1 & 0 \\ 0 & 1 & 0 & 1 & 0 & 1 \\ 0 & 0 & 1 & 0 & 1 & 1 \end{bmatrix} \quad (\text{II.1})$$

Le processus de transmission complet est composé de trois phases : modulation, séquences d'étalement éparses et rotation de constellation, comme illustré dans la **Figure (II.17)**. Dans un premier temps, les informations binaires provenant du j -ième utilisateur sont transformées en un symbole complexe $z^{(j)}$. Par la suite, on effectue simultanément l'étalement de séquence éparses et la rotation de constellation. Ce processus peut être représenté par une multiplication de vecteurs complexes. Ainsi, le signal du j -ième utilisateur est un vecteur complexe de dimension K $x^{(j)} = [x_1^{(j)}, x_1^{(j)}, \dots, x_k^{(j)}]^T$, qui est donné par [45] :

$$x^{(j)} = g^{(j)} \cdot z^{(j)}, \quad (\text{II.2})$$

Où $g^{(j)}$ est égal à $[g_1^{(j)}, g_2^{(j)}, \dots, g_k^{(j)}]^T$ est le vecteur générateur composé de $K - N$ zéros déterminés par F . En cas de liaison descendante, on additionne tous les signaux des J utilisateurs afin de créer un vecteur transmis $s = [s_1, s_2, \dots, s_K]^T$, donné par [45] :

$$s = \sum_{j=1}^J x^{(j)} = \sum_{j=1}^J (g^{(j)} \cdot z^{(j)}) = G \cdot z^T \quad (\text{II.3})$$

$G = [g^{(1)}, g^{(2)}, \dots, g^{(J)}]$ est la matrice génératrice, tandis que $z = [z^{(1)}, z^{(2)}, \dots, z^{(J)}]$ est un vecteur complexe qui représente les symboles modulés de tous les utilisateurs J .

Le signal reçu par le j -ième utilisateur peut être écrit de la manière suivante :

$$r = H \cdot s + n \quad (\text{II.4})$$

$H = \text{diag}(h_1, h_2, \dots, h_k)$ est la matrice du canal de fading plat et $n = [n_1, n_2, \dots, n_K]^T$ il s'agit d'un vecteurs complexe additif blanc Gaussien (AWGN) de dimension K .

L'estimation du symbole transmis nommé $\hat{z}(j)$, est obtenue en effectuant une détection multi-utilisateurs sur r [45].

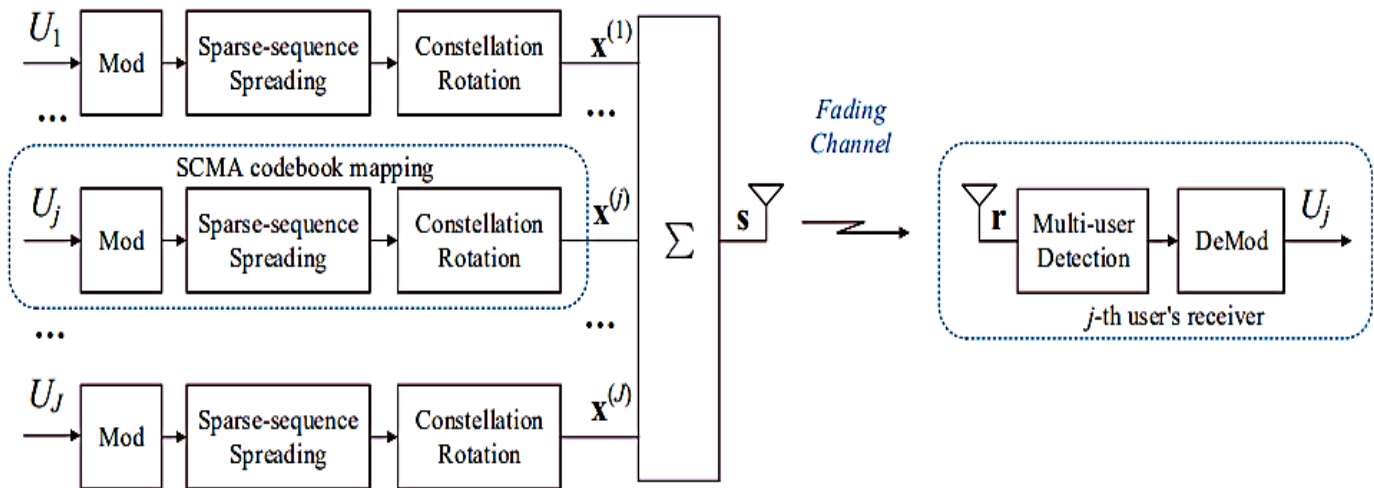


Figure II.17 : Illustration d'un modèle émetteur-récepteur SCMA en liaison descendante [45].

B) Les bases du système SCMA :

Le SCMA apparaît comme une approche novatrice qui se base sur plusieurs concepts essentiels, tels que le concept de livre de codes, le concept de mot de code et le processus d'encodage SCMA.

1. Le concept de code book :

L'objectif de la conception des livres de codes SCMA est d'améliorer les caractéristiques de distance entre les points de la constellation multidimensionnelle, tout en permettant un nombre limité de points de projection sur chaque élément de ressource. Grâce à cette conception, l'efficacité du système est améliorée tout en diminuant la complexité de la réception. La **Figure II.18** montre un exemple d'un ensemble contenant 6 livres de codes pour la transmission de 6 couches de données. Comme on peut le voir, chaque livre de codes contient 8 mots de code complexes multidimensionnels correspondant à 8 points de constellation, respectivement. La longueur de chaque mot de code est de 4, ce qui est identique à la longueur de l'étalement [44].

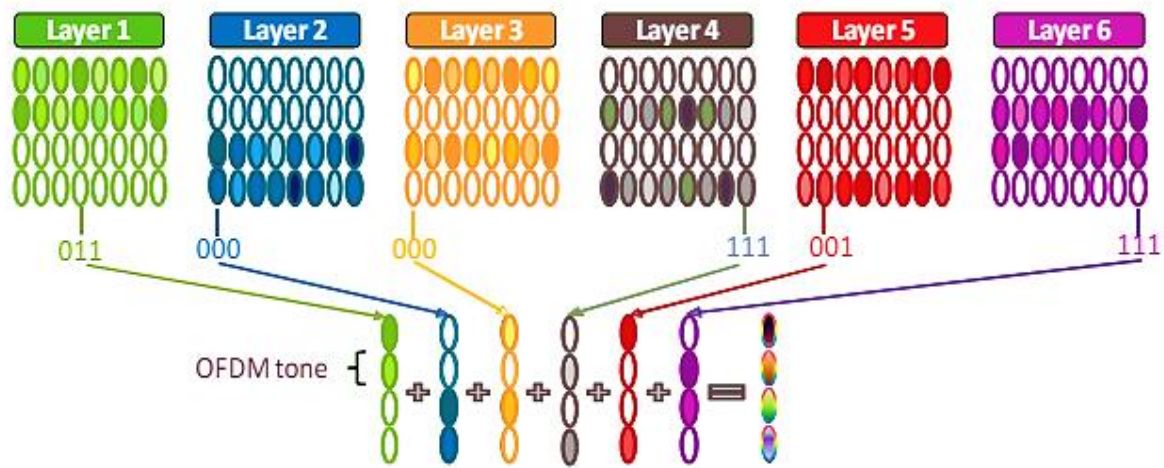


Figure II.18 : Illustration de la structure du livre de codes SCMA [44].

Les livres de codes SCMA présentent également la possibilité d'avoir un nombre réduit de points de projection sur chaque élément de ressource. La raison en est que les livres de codes sont multidimensionnels, ce qui permet à deux points de constellation de se chevaucher sur certains des composants non nuls, car ils peuvent encore être séparés sur les autres composants non nuls. La **Figure II.19** illustre un exemple où les points de constellation correspondant à 100 et 000 se chevauchent sur le premier ton, mais sont séparés sur le deuxième ton, ce qui entraîne un nombre de points de projection de 4 au lieu de 8 [44].

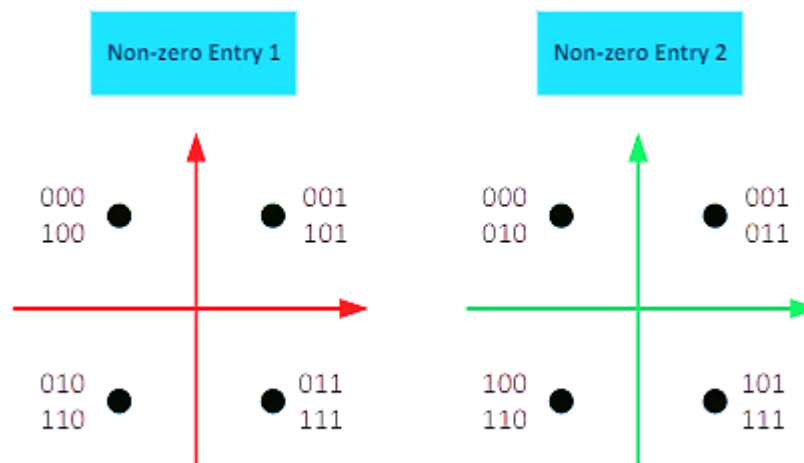


Figure II.19 : Un exemple d'un codebook SCMA à 8 points [44].

2. Le concept de codeword :

Dans un système SCMA, les codewords sont des vecteurs multidimensionnels complexes qui interprètent les symboles modulés transmis par chaque utilisateur. Chaque codeword est choisi dans un codebook dédié à la couche, qui comprend plusieurs points de constellation. Ce sont des codewords épars, c'est-à-dire que la majorité de leurs entrées sont nulles, et seulement quelques entrées sont non nulles. Grâce à cette structure de codeword, il est possible d'utiliser les ressources de manière efficace tout en facilitant la perception des symboles [44].

3. Encodage SCMA :

L'encodage SCMA est une technique de codage utilisée dans les réseaux de communication pour attribuer des ressources de manière efficace à plusieurs utilisateurs. Cette méthode repose sur un mécanisme sophistiqué de mapping des symboles de modulation vers des mots de code. Dans l'encodage SCMA, un mapping est réalisé à partir de m bits vers un codebook complexe de dimension K et de taille M , où $M = 2^m$. Chaque utilisateur J possède son propre codebook parmi l'ensemble des J codebooks. Ci-dessous, une représentation de deux codebooks pour l'utilisateur 1 et 2 avec $K = 4$ [46].

$$CB_1 = \begin{bmatrix} 0 & -0.1815 - 0.1318j & 0 & 0.7851 \\ 0 & -0.6351 - 0.4615j & 0 & -0.2243 \\ 0 & 0.6351 + 0.4615j & 0 & 0.2243 \\ 0 & 0.1815 + 0.1318j & 0 & -0.7851 \end{bmatrix}^T$$

$$CB_2 = \begin{bmatrix} 0.7851 & 0 & -0.1815 - 0.1318j & 0 \\ -0.2243 & 0 & -0.6351 - 0.4615j & 0 \\ 0.2243 & 0 & 0.6351 + 0.4615j & 0 \\ -0.7851 & 0 & 0.1815 + 0.1318j & 0 \end{bmatrix}^T$$

Les colonnes des codebooks sont constituées de mots de code, ce qui signifie que chaque utilisateur peut acheminer $m = 2\text{bits}$ vers l'un des $M = 4$ mots de code en quatre dimensions. Les éléments de ressource partagés (RE) sont utilisés pour transmettre les mots de code SCMA [46].

4. Encodage SCMA basé sur le principe de rotation de constellation :

Les systèmes SCMA utilisent l'encodage SCMA qui repose sur le principe de rotation de constellation afin de lier les bits d'information à des mots de code complexes. Dans cette approche d'encodage, les points de la constellation qui représentent les mots de code sont

déplacés dans le plan complexe afin de générer des éléments de ressources orthogonaux (RE) adaptés à divers utilisateurs. L'encodeur SCMA permet de représenter les mots de code de manière éparsée en tournant les points de la constellation, ce qui permet à plusieurs utilisateurs de partager les mêmes ressources sans interférences [45].

- **La constellation :**

Étant donné que le vecteur transmis s (II.3) est un vecteur complexe de dimension K , il peut être représenté distinctement par un point dans un espace multidimensionnel de $2K$ -dimensions. Ainsi, une constellation multidimensionnelle de $2K$ dimensions, nommée constellation principale, est créée pour produire le vecteur s . L'ordre de modulation, par exemple M , détermine le nombre de points dans cette constellation multidimensionnelle. En prenant en compte le nombre total d'utilisateurs, il y a M^J valeurs potentielles de z . Ainsi, il est prévu qu'il y ait M^J points différents dans cette constellation principale. Dans cette situation, s représente un point de constellation v_m , qui est défini par tous les utilisateurs J , avec une valeur de $1 \leq m \leq M^J$. Deux éléments peuvent être utilisés pour évaluer les caractéristiques d'un codebook SCMA : les distances entre deux points de constellation dans la constellation principale et la puissance moyenne.

La distance entre v_m et $v_{m'}$ peut être représentée par [45] :

$$d_{m \rightarrow m'} = ||v_m - v_{m'}||, \quad m \neq m' \quad (\text{II.5})$$

Et la puissance moyenne est calculée par :

$$\mathcal{E}_{avg} = \frac{1}{M^J} \sum_{m=1}^{M^J} \frac{||v_m||^2}{2} \quad (\text{II.6})$$

La règle générale de conception des codebooks SCMA est donc de maximiser la distance euclidienne entre les points de la constellation lorsque la puissance moyenne est établie.

Il est connu que le vecteur s est constitué de K symboles complexes, tels que $s_1, s_2, \dots, s_K$. La sous-constellation est une constellation 2D qui produit le k -ième composant de s . Puisque chacune des ressources est utilisée exclusivement par d_R utilisateurs particuliers, la k -ième sous-constellation est déterminée par $\{U_j | f_k^{(j)} = 1\}$. Par exemple, trois utilisateurs sont responsables de la décision de la sous-constellation: 1, 4 et 5, d'après (II.1). Ainsi, il est prévu qu'il y ait des points M^{d_R} dans la sous-constellation. Au total, il existe K sous-constellations.

La projection de v_m sur la k -ième sous-constellation est définie par $proj(v_m; k)$. Ainsi, (II.5) peut être déterminé [45] :

$$\|v_m - v_{m'}\|^2 = \sum_{k=1}^K \|proj(v_m; k) - proj(v_{m'}; k)\|^2 \quad (\text{II.7})$$

Ainsi, si au moins un k peut être trouvé pour satisfaire à $\|proj(v_m; k) - proj(v_{m'}; k)\| > 0$, v_m et $v_{m'}$ ne se chevaucheront pas.

- La rotation de constellation :

La rotation de constellation est une méthode qui permet de faire tourner les points de la constellation de modulation dans un espace complexe, habituellement en utilisant des opérations linéaires ou des transformations matricielles. La rotation entraîne une modification de la disposition spatiale des symboles de modulation, ce qui peut avoir un effet important sur la capacité du système à détecter et à démoduler de manière adéquate les données, en particulier dans des environnements où il y a des interférences ou du bruit.

1-Le facteur de rotation de phase :

On optimise la forme de la sub-constellation en définissant les valeurs appropriées des entrées non-zéro dans la matrice génératrice. Soit les entrées non-zéro dans la matrice génératrice purement imaginaires, de sorte que seules les phases des symboles modulés sont modifiées pendant le processus de rotation des constellations. Assigne les entrées non-zéro dans G comme le facteur de rotation de la phase, qui est présenté par [45] :

$$\varphi_a = e^{i.\alpha.\Delta} \quad (\text{II.8})$$

Où α est un entier, $0 \leq \alpha \leq d_R - 1$ et Δ est l'intervalle entre deux phases adjacentes. Selon les équations (II.8) et (II.9) :

$$s_k = \sum_{f_k^{(j)}=1} \varphi_a \cdot z^{(i)} \quad (\text{II.9})$$

α varie en fonction de j .

2-La matrice génératrice latine :

Une matrice génératrice est considérée comme une matrice génératrice latine lorsqu'elle répond aux deux conditions suivantes [45] :

$$g_k^{(j)} \neq g_k^{(j')}, \text{ lorsque } j' \neq j \text{ et } g_k^{(j)} \neq 0; \quad (\text{II.10})$$

$$g_k^{(j)} \neq g_{k'}^{(j)}, \text{ lorsque } k' \neq j \text{ et } g_k^{(j)} \neq 0. \quad (\text{II.11})$$

En premier lieu, créons une matrice de mapping dotée d'une structure particulière. Dans F, la colonne j est représentée par $f^{(j)} = [f_1^{(j)}, f_2^{(j)}, \dots, f_K^{(j)}]^T$. On peut considérer $f^{(j)}$ comme un nombre binaire car ses éléments peuvent être soit 0 soit 1.

On définit une fonction $B(f^{(j)})$ comme suit [45] :

$$B(f^{(j)}) = f_1^{(j)} \times 2^{K-1} + f_2^{(j)} \times 2^{K-2} + \dots + f_K^{(j)} \times 2^0 \quad (\text{II.12})$$

La matrice de mapping $F = [f^{(1)}, f^{(2)}, \dots, f^{(J)}]$ est obtenue en disposant les vecteurs colonnes sous la condition que :

$$B(f^{(1)}) > B(f^{(2)}) > \dots > B(f^{(J)}) \quad (\text{II.13})$$

Ensuite, trouvez l'emplacement du v-ème '1' dans la k-ième ligne de F, et notez son indice de colonne comme (k, v) , $1 \leq k \leq K, 1 \leq v \leq K - 1$. Enfin, l'élément dans cette matrice génératrice correspondant à cet indice est noté comme $g_k^{(ind(k,v))}$, qui est donné par:

$$g_k^{(ind(k,v))} = \varphi_\alpha, \quad \alpha = (k + v - 2) \bmod (K - 1). \quad (\text{II.14})$$

Si la matrice génératrice est construite conformément à (II.14), elle reste latine. L'annexe A présente la preuve. Afin de comparer la matrice génératrice latine avec celle standard, les générateurs (II.15) et (II.16) sont présentés pour un système SCMA (4, 6).

$$G_{4 \times 6}^{La} = \begin{bmatrix} \varphi_0 & \varphi_1 & \varphi_2 & 0 & 0 & 0 \\ \varphi_1 & 0 & 0 & \varphi_2 & \varphi_0 & 0 \\ 0 & \varphi_2 & 0 & \varphi_0 & 0 & \varphi_1 \\ 0 & 0 & \varphi_0 & 0 & \varphi_1 & \varphi_2 \end{bmatrix} \quad (\text{II.15})$$

$$G_{4 \times 6}^{Nla} = \begin{bmatrix} \varphi_0 & \varphi_1 & \varphi_2 & 0 & 0 & 0 \\ \varphi_0 & 0 & 0 & \varphi_2 & \varphi_1 & 0 \\ 0 & \varphi_1 & 0 & \varphi_2 & 0 & \varphi_0 \\ 0 & 0 & \varphi_2 & 0 & \varphi_1 & \varphi_0 \end{bmatrix} \quad (\text{II.16})$$

3-Amélioration supplémentaire :

En augmentant K, nous pouvons créer un nouveau code book pour le système SCMA, à condition qu'il réponde à la contrainte J. Il est indéniable que λ augmente en fonction de l'augmentation de K et J. Toutefois, une augmentation du taux de chargement diminuera la

fiabilité du système. Avec l'augmentation de λ , d_R augmente, car $d_R = N\lambda$. Ainsi, un plus grand nombre de points sont regroupés dans le plan en deux échelles. Lorsque la puissance totale est établie, la distance euclidienne entre deux points sera inévitablement diminuée. Ce problème peut être résolu en maintenant λ petit lorsque K et J augmentent. Cette idée est illustrée par la matrice génératrice latine suivante pour SCMA (6, 9) [45].

$$G_{(9,6)} = \begin{bmatrix} \varphi_0 & \varphi_1 & \varphi_2 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & \varphi_1 & \varphi_2 & \varphi_0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & \varphi_2 & \varphi_0 & \varphi_1 \\ \varphi_1 & 0 & 0 & \varphi_2 & 0 & 0 & \varphi_0 & 0 & 0 \\ 0 & \varphi_2 & 0 & 0 & \varphi_0 & 0 & 0 & \varphi_1 & 0 \\ 0 & 0 & \varphi_0 & 0 & 0 & \varphi_1 & 0 & 0 & \varphi_2 \end{bmatrix} \quad (\text{II.17})$$

C) La technique de détection multi-utilisateurs utilisée dans SCMA:

Pour SCMA, le signal reçu combine les mots de code de chaque utilisateur, entraînant des interférences entre les utilisateurs. Les méthodes de séparation orthogonales classiques ont du mal à distinguer les informations des utilisateurs, ce qui nécessite une détection multi-utilisateurs conjointe (MUD) pour la séparation des données. Parmi les méthodes couramment utilisées, nous détaillons le MPA.

1. Définition de MPA :

Le MPA est appliqué pour la détection multi-utilisateurs, est un outil qui permet d'inférer à partir de modèles graphiques en passant des messages de croyance entre les nœuds. L'exemple illustré dans la **Figure II.20** peut être utile pour comprendre comment l'inférence peut être obtenue à partir d'un graphe. Un groupe d'élèves debout en ligne, chaque élève communique avec ses voisins pour compter le nombre total d'élèves. L'algorithme démarre des deux extrémités de la ligne, où le compteur augmente à mesure que le message passe par chaque élève. Les messages se déplacent dans les deux directions de la ligne, et chaque élève obtient l'information sur le nombre total d'élèves présents. Si un élève reçoit un compte de n d'un côté et de m de l'autre côté, alors le nombre total d'élèves présents est $n + m + 1$ [47].

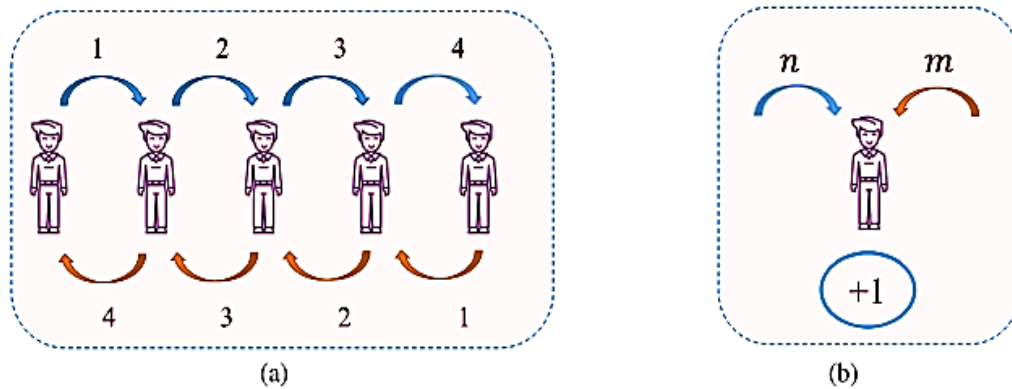


Figure II.20 : Comptage des étudiants utilisant MPA ; **(a)** comptage groupé des étudiants, **(b)** comptage par étudiant individuel [47].

2. Le graphe de facteurs:

La compréhension et l'utilisation de certaines techniques avancées de communication sans fil, comme le SCMA, reposent sur le concept de graphe de facteurs. Un graphe de facteurs représente visuellement les liens entre les variables et les fonctions dans un système complexe. Cette représentation graphique simplifie l'étude et la détection des signaux. La **Figure II.21** illustre un graphe de facteurs du système SCMA avec 6 utilisateurs et 4 RES. Chaque cercle du graphe de facteurs symbolise un utilisateur (nœud variable) tandis que chaque bloc représente un ton (nœud de fonction). Le graphe de facteurs présente des boucles, ce qui signifie qu'aucune des lignes ne peut être calculée en premier [48].

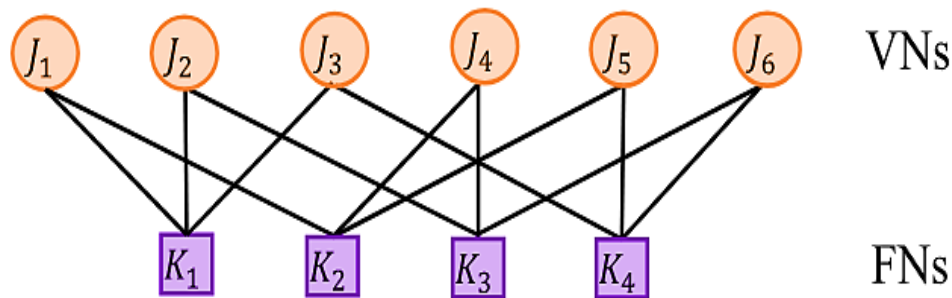


Figure II.21 : Graphe de facteurs du système SCMA (4×6) [47].

3. L'algorithme MPA :

Prenons un diagramme de facteurs avec K FNs et J VNs. Si le graphe de facteurs contient des cycles, le passage de messages dans le graphe de facteurs en utilisant MPA est un processus itératif. Chaque itération comprend deux phases. L'étape 1 consiste à transmettre un message de croyance d'un FN à un VN, tandis que l'étape 2 consiste à la transmission d'un VN à un FN. MPA débute en envoyant des messages à partir des nœuds de feuilles. Un nœud dont le degré (nombre d'arêtes connectées à celui-ci) est d reste inactif jusqu'à ce que des messages

soient reçus sur $d - 1$ arêtes. Soit $\eta_{j \rightarrow k}$ le message envoyé du VN j au FN k et $\eta_{k \rightarrow j}$ le message du FN k au VN j . Soient ζ_j et ξ_k l'ensemble de nœuds directement connectés au VN j et FN k , respectivement. Ensuite, le MPA peut être effectué comme suit [47] :

1) Passage du message du FN k au VN j :

$$n_{k \rightarrow j}(x_j^m) = \sum_{\sim x_j} (\psi_k \prod_{j' \in \xi_k \setminus \{j\}} n_{j' \rightarrow k}(x_{j'}^m)) \quad (\text{II.18})$$

Où ψ_k est la fonction de vraisemblance associée à FN k et $j' \in \xi_k \setminus \{j\}$ désigne tous les VNs de ξ_k sauf le j -ième VN. Aussi $n_{j' \rightarrow k}(x_{j'}^m)$ indique le message du VN j' au FN k correspondant au symbole $x_{j'}^m \in \text{CB}_j$, $m = 1, \dots, M$.

2) Passage du message du VN j au FN k :

$$n_{j \rightarrow k}(x_j^m) = \prod_{k' \in \xi_j \setminus \{k\}} n_{k' \rightarrow j}(x_j^m), \quad (\text{II.19})$$

Où $k' \in \xi_j \setminus \{k\}$ désigne les FN dans ξ_j à l'exception du k -ième FN. Le message est normalisé pour garantir que la somme de toutes les probabilités est égale à un.

$$n_{j \rightarrow k}(x_j^m) = \frac{\prod_{k' \in \xi_j \setminus \{k\}} n_{k' \rightarrow j}(x_j^m)}{\sum_{x_j^l \in \text{CB}_j} \prod_{k' \in \xi_j \setminus \{k\}} n_{k' \rightarrow j}(x_j^l)} \quad (\text{II.20})$$

L'équation (II.19) est plus facile que l'équation (II.20), car il n'existe pas de fonction locale liée à un VN. Après quelques cycles, on calcule le rapport de vraisemblance logarithmique (LLR) de chaque bit afin d'estimer les bits transmis par chaque utilisateur. L'algorithme de décodage du système SCMA ($K \times J$), appelé MPA, est résumé dans l'Algorithme 1 [47].

Algorithme 1 : The Message Passing Algorithm

Inputs:

J : Number of users.

K : Number of REs.

CB_j : CB of j th user, $\forall j = 1, \dots, J$.

y : The received signal vector.

Output:

Estimated value of bits for j th user, $\forall j = 1, \dots, J$.

Initialize:

The prior probability of each codeword

$$\mathbf{n}_{j \rightarrow k}^0 = P(x_j^m) = \frac{1}{M}, \forall j = 1, \dots, J, k \in \zeta_j$$

1: Step 1: Message passing along the nodes

for $t = 1$ to N_t do

a) Update Resource Node:

$$\eta_{k \rightarrow j}^t(x_j^m) = \sum_{\sim x_j} \left(\exp \left\{ \frac{-1}{N_0} \left\| y_k - \sum_{j \in \xi_k} h_{k,j} x_{k,j}^m \right\|^2 \right\} \prod_{j' \in \xi_k \setminus \{j\}} \eta_{j' \rightarrow k}^{t-1}(x_{j'}^m) \right),$$

b) Update Variable Node:

$$\eta_{j \rightarrow k}(x_j^m) = P(x_j) \prod_{k' \in \xi_j \setminus \{k\}} \eta_{k' \rightarrow j}^{t-1}(x_j^m)$$

End
2: Step 2: Bit LLR

The final belief computed at each VN is:

$$I(x_j^m) = \left(\frac{1}{M} \right) \prod_{k \in \zeta_j} \eta_{k \rightarrow j}(x_j^m)$$

The i -th bit LLR ($1 \leq i \leq \log_2(M)$) at j -th VN is :

$$\begin{aligned} LR(b_j^i) &= \log \frac{P(b_j^i = +1)}{P(b_j^i = -1)} \\ &= \log \frac{\sum_{x_j^m \in CB_j | b_j^i = +1} I(x_j^m)}{\sum_{x_j^m \in CB_j | b_j^i = -1} I(x_j^m)} \end{aligned}$$

Where, $b_j^i = +1$ if LLR (b_j^i) is positive otherwise $b_j^i = -1$

4. Les désavantages et les limites du MPA:

Bien que le MPA présente des avantages en termes de performances, il comporte également plusieurs inconvénients et limites qui doivent être pris en considération. Parmi les désavantages les plus fréquents liés au MPA [48] [49] :

- Le nombre d'appareils connectés accroît considérablement la complexité de détection du MPA traditionnel.
- L'algorithme MPA n'est pas applicable dans des conditions de faible rapport signal sur bruit (SNR).
- Le MPA offre une approche efficace pour détecter les signaux multi-utilisateurs en échangeant des messages entre les nœuds d'un graphique dynamique. Mais il peut aussi poser des problèmes concernant la convergence de l'algorithme et la gestion des interférences.
- La sécurité, car chaque utilisateur du système a la capacité de détecter des informations des autres utilisateurs. Cela soulève de nombreux défis en termes de sécurité.

II.4.2.3.Les autres schémas de NOMA :

II.4.2.3.1.RSMA :

Dans RSMA, on superpose un ensemble de signaux provenant de différents utilisateurs, et le signal de chaque utilisateur est réparti sur toutes les ressources fréquentielles/temporelles attribuées au groupe. Les signaux des utilisateurs du groupe ne sont pas forcément orthogonaux les uns aux autres et pourraient éventuellement entraîner des interférences entre eux. Grâce à l'extension des bits sur toutes les ressources, il est possible de décoder le signal en dessous du bruit de fond/interférences. Pour séparer les signaux des différents utilisateurs, RSMA utilise une combinaison de codes de canal à faible débit et de codes de brouillage qui possèdent de bonnes propriétés de corrélation. D'après les cas d'utilisation, cela peut englober [50] :

- Le RSMA à porteuse unique est conçu pour réduire la consommation d'énergie des batteries et étendre la couverture pour les petites transactions de données en utilisant des formes d'onde à porteuse unique et des modulations à très faible PAPR. Il offre la possibilité de transmettre sans allocation de ressources et éventuellement d'accéder de manière asynchrone.

- Le RSMA multi-porteuses est conçu pour offrir un accès à faible latence aux utilisateurs en mode connecté RRC (c'est-à-dire une synchronisation avec eNB déjà acquise) et permet la transmission sans nécessiter d'allocation de ressources. La **Figure II.22** présente les diagrammes en bloc des deux modes.

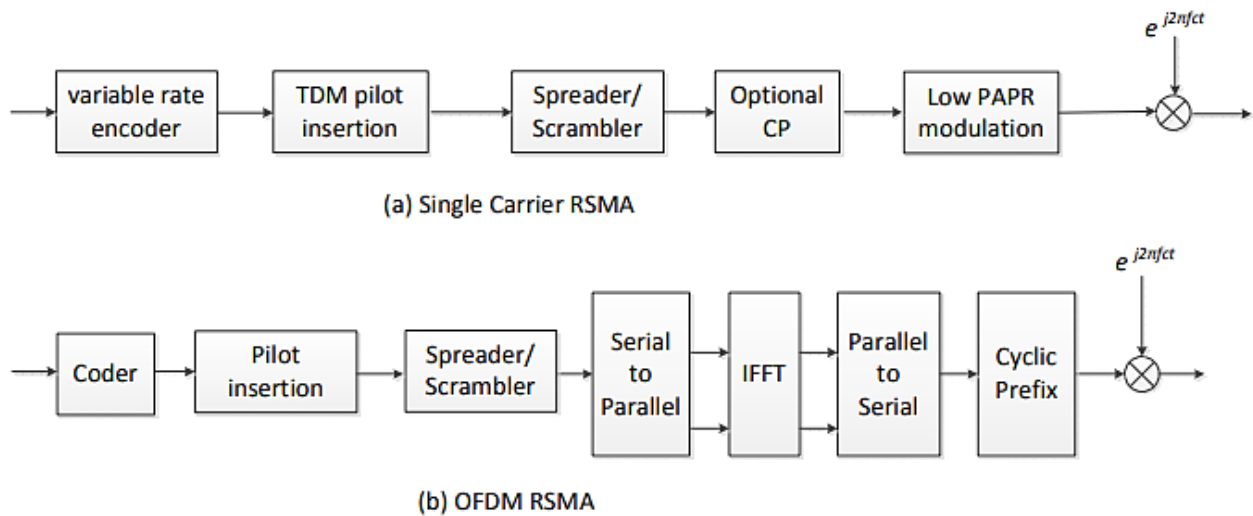


Figure II.22 : Présentation des diagrammes en bloc des deux modes de RSMA [50].

II.4.2.3.2 .RDMA :

Le schéma basé sur l'entrelacement connu sous le nom de RDMA, permet de séparer les signaux des différents utilisateurs et d'utiliser à la fois la diversité temporelle et fréquentielle en répétant simplement un décalage cyclique.

Un exemple de RDMA avec 3 utilisateurs simultanés est illustré dans la **Figure II.23**. Dans cette situation, les utilisateurs partagent un total de 81 (9X9) RE. L'utilisateur- i et l'indice- k , noté $x_i(k)$, répètent chaque symbole modulé M_{rep} fois, soit 3 fois. Les raisons de répétition suivantes semblent entraîner une interférence multi-utilisateurs (MUI) totalement aléatoire et une variété temporelle et fréquentielle pour chaque symbole modulé répété.

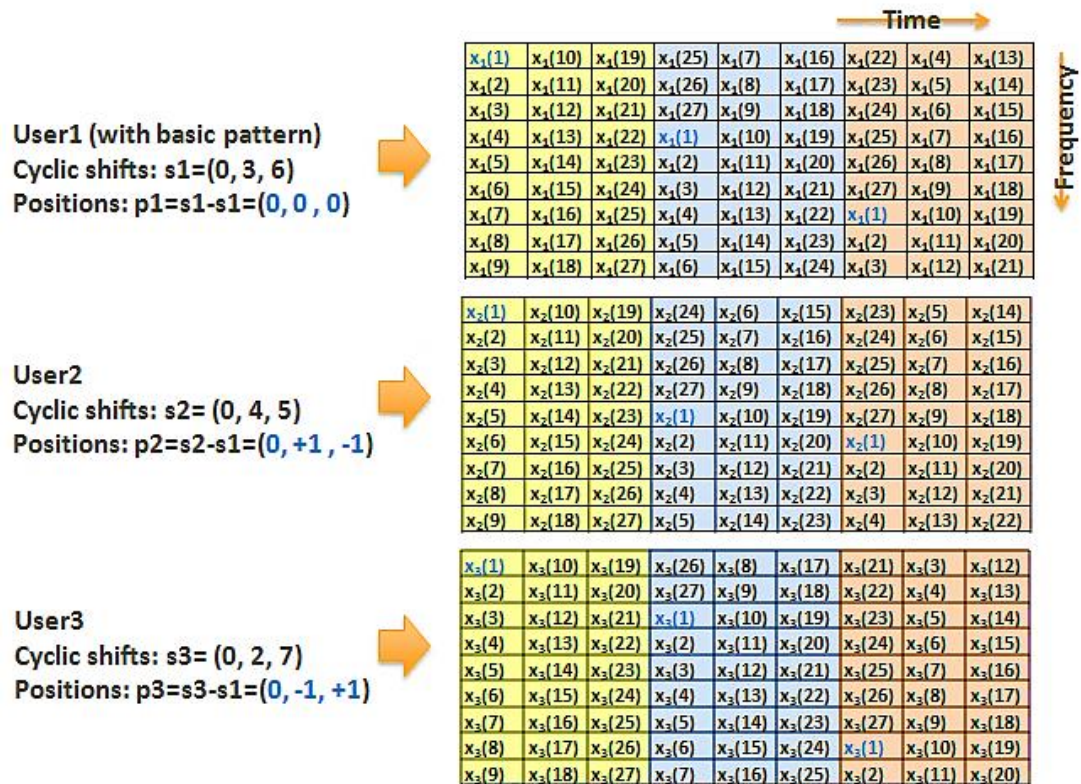


Figure II.23 : Un exemple de motifs de répétition RDMA avec 3 utilisateurs simultanés [42].

On aborde ici les principes de conception du motif de répétition RDMA. Supposons qu'il y ait L positions possible pour une permutation, alors il existe $L^{M_{rep}}$ séquences possible pour M_{rep} permutations. En sélectionnant N_{ue} séquences parmi ces $L^{M_{rep}}$ séquences pour N_{ue} utilisateurs, il est nécessaire que la différence entre deux séquences différentes ait des valeurs différentes pour tous les éléments afin de rendre l'interférence inter-utilisateurs (MUI) totalement randomisée [42].

Effectivement, en plus des schémas NOMA mentionnés précédemment, il existe d'autres approches telles qu'IDMA et IGMA. Ces schémas NOMA offrent des alternatives intéressantes pour répondre aux besoins variés des réseaux de communication modernes. Chacun présente ses propres avantages et inconvénients en fonction du contexte d'application spécifique.

II.5.Conclusion :

En résumé, ce chapitre a jeté les bases de notre étude en présentant un aperçu des méthodes d'accès multiple OMA et NOMA dans les réseaux de communication sans fil, mettant en lumière l'importance de SCMA en tant que technique prometteuse pour améliorer l'efficacité et les performances des réseaux sans fil de prochaine génération.

Dans le prochain chapitre, nous aborderons le domaine du Deep Learning (DL), une discipline de l'intelligence artificielle qui a révolutionné de nombreux domaines, y compris les communications sans fil.

CHAPITRE III :

*Deep Learning (L'apprentissage
en profondeur).*

III.1.Introduction :

La recherche dans le domaine de l'intelligence artificielle (IA) est en perpétuelle évolution, avec pour objectif de concevoir des systèmes informatiques capables de réaliser des tâches qui requièrent habituellement l'intelligence humaine. Deux disciplines interconnectées sont au cœur de l'IA : le Machine Learning (ML) et le Deep Learning (DL). Ces domaines ont transformé notre aptitude à analyser et à utiliser les immenses quantités de données produites quotidiennement, offrant ainsi de nouvelles opportunités et applications dans différents secteurs.

Dans ce chapitre, nous allons examiner les principaux concepts de l'intelligence artificielle, du ML et du DL, les principes fondamentaux qui sous-tendent ces disciplines, ainsi que les algorithmes et techniques couramment employés, seront étudiés.

III.2.L'intelligence artificielle (IA) :

La première utilisation du terme « Intelligence artificielle (IA) » aurait été faite en 1955 par l'informaticien américain John McCarthy. Selon McCarthy à l'aube de ce domaine d'étude, cette définition résume brièvement l'objectif principal de l'intelligence artificielle : "Le domaine de l'informatique qui vise à rendre les ordinateurs semblables aux êtres humains." **John McCarthy, 1956.** L'objectif est de créer des systèmes informatiques qui peuvent reproduire des comportements intelligents humains, tels que la perception, la compréhension du langage, l'acquisition de connaissances et la prise de décision.

Un de ses pionniers définit l'intelligence artificielle comme étant « la création de programmes informatiques capables de réaliser des tâches actuellement réalisées de manière plus efficace par des êtres humains, car elles impliquent des processus mentaux avancés tels que la perception, la mémoire et le raisonnement critique » - **Marvin Lee Minsky [51].**

III.2.1.Bref historique :

L'histoire de l'IA remonte à plusieurs siècles, avec des idées et des concepts qui ont évolué au fil du temps. L'émergence officielle de l'intelligence artificielle (IA) est précédée par une période de gestation prolongée, nourrie par des avancées théoriques et pratiques dans les domaines des automates et de la logique mathématique.

Tout d'abord, du côté des machines automatiques, des jalons importants ont été posés par des figures telles que Charles Babbage avec sa machine "analytique" en 1842, ainsi qu'Alan Turing et sa machine universelle en 1936. Ces travaux ont jeté les bases des premiers calculs mécaniques et universels.

Dans le domaine de la logique mathématique, des contributions significatives ont été apportées par des penseurs tels que Gottfried Wilhelm Leibniz, George Boole et David Hilbert. Les travaux de Kurt Gödel et Alonzo Church, bien que limitant les ambitions, ont également joué un rôle crucial en montrant les limites de la résolubilité algorithmique de certains problèmes.

L'avènement des ordinateurs dans les années 1940 a marqué un tournant décisif dans l'histoire de l'IA. Des pionniers se sont alors posé la question de l'intelligence de ces machines, et le mouvement de la cybernétique, initié par des scientifiques comme Norbert Wiener, a joué un rôle essentiel en introduisant des concepts tels que la régulation et la rétroaction, qui sont fondamentaux pour les réseaux neuronaux et la robotique.

Le perceptron monocouche, inventé par Frank Rosenblatt en 1957, a été le premier outil d'apprentissage basé sur les réseaux neuronaux. Bien que son intérêt ait diminué après 1970 en raison de critiques sur ses limitations, l'émergence du perceptron multicouche et des algorithmes d'apprentissage associés a ravivé l'intérêt pour les réseaux neuronaux.

Aujourd'hui, les réseaux neuronaux profonds, constitués de multiples couches de neurones, sont en tête dans le domaine de l'intelligence artificielle. Leur capacité à fournir des performances exceptionnelles dans une variété d'applications, grâce à l'utilisation de grandes quantités de données et à l'apprentissage profond, en fait des outils incontournables dans le paysage de l'IA moderne [52].

III.2.2.L'évolution de (IA) :

L'intelligence artificielle (IA) est un domaine de recherche et de développement qui vise à créer des systèmes informatiques capables d'exécuter des tâches qui nécessitent normalement l'intelligence humaine. Depuis ses débuts dans les années 1950, l'IA a parcouru un long chemin, passant de concepts théoriques à des applications pratiques dans de nombreux domaines de la vie quotidienne. Cette section propose un aperçu de l'évolution de l'IA [53] :

1. Années 1943 - 1956 : Les premières années de l'IA

- En 1943, un modèle mathématique d'un neurone formel est présenté par Warren McCulloch et Walter Pitts, posant ainsi les fondements théoriques de l'intelligence artificielle.
- En 1950, dans son article intitulé "Computing Machinery and Intelligence", Alan Turing propose le test de Turing pour évaluer l'intelligence des machines .
- En 1956, la conférence de Dartmouth introduit le concept d'intelligence artificielle, ce qui marque le début officiel de la recherche en IA.

2. Années 1960-1970 : La recherche scientifique

- Les premiers logiciels d'IA sont créés, y compris le jeu d'échecs.
- Le langage de programmation Lisp, largement employé dans la recherche en IA, est développé par John McCarthy.
- On assiste à l'apparition de systèmes experts qui exploitent des bases de connaissances afin de simuler le raisonnement humain dans des domaines particuliers.

3. Années 1980 et 1990 : L'émergence de IA

- L'IA est en plein essor, avec des progrès notables dans les domaines de la vision par ordinateur, du traitement du langage naturel et des réseaux neuronaux.
- Les entreprises utilisent de plus en plus les systèmes experts pour prendre des décisions et réaliser des diagnostics.
- Les réseaux de neurones artificiels et les algorithmes génétiques sont des algorithmes d'apprentissage automatique qui commencent à être utilisés avec succès dans différentes applications.

4. Des années 2000 à nos jours : L'émergence du DL et d'IA

- Le DL se présente comme une nouvelle approche en IA, offrant aux machines la possibilité d'acquérir des connaissances sur les représentations de données de manière hiérarchique.
- Les utilisations de l'IA se développent dans divers secteurs, tels que la reconnaissance d'images, la traduction automatique.
- Les avancées dans la compréhension du langage naturel amènent à l'émergence d'assistants virtuels tels que Siri, Alexa et Google Assistant, capables de comprendre et de répondre aux commandes vocales.
- L'IA pose de plus en plus de problèmes éthiques et sociaux, tels que la protection de la vie privée, la sécurité et la disparition des emplois.

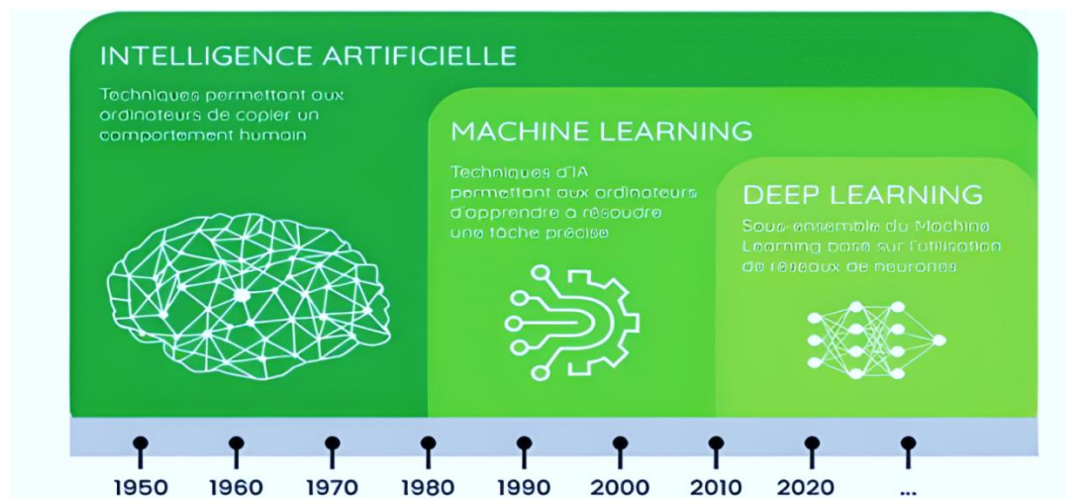


Figure III.1 : Chronologie de l'évolution de l'IA [51].

III.2.3. Les sous domaines de (IA) :

L'IA englobe des domaines fonctionnels tels que les systèmes experts, la planification/optimisation ou la robotique. Toutefois, l'intelligence artificielle inclut également de nouvelles branches telles que l'apprentissage automatique, le traitement du langage naturel, la vision (la capacité d'une machine à percevoir son environnement) et le Speech (texte vers parole ou parole vers texte). L'intelligence artificielle trouve de plus en plus d'applications dans de nombreux domaines [54]. La **Figure III.2** nous expose les principaux domaines fonctionnels traités par l'intelligence artificielle :

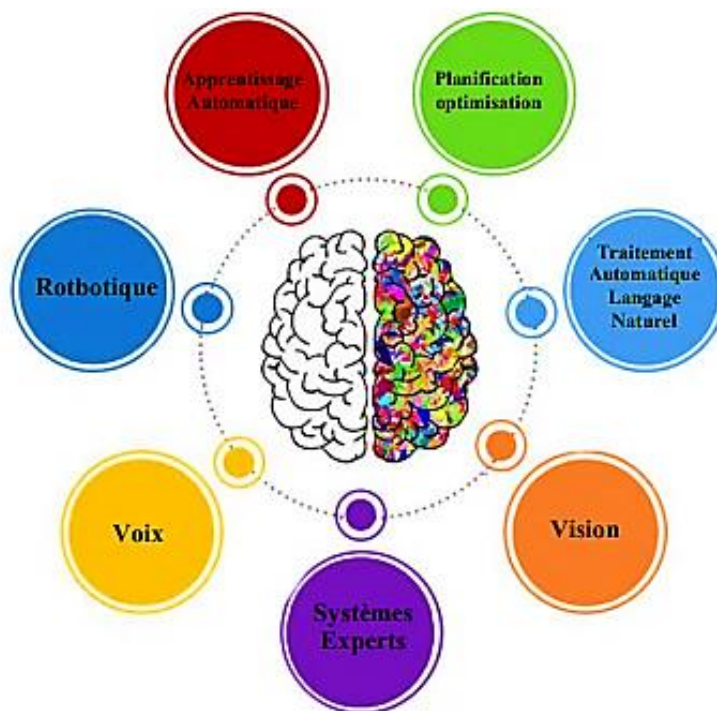


Figure III.2 : Les sous domaines de l'IA [53].

III.3.L'apprentissage automatique :

III.3.1.définition :

L'une des branches les plus importantes de l'IA est l'apprentissage automatique ou le machine Learning en anglais. Une définition proposée par **Fabien Benureau (2015)**: «L'apprentissage est une modification d'un comportement sur la base d'une expérience» [55].L'apprentissage automatique consiste à développer des algorithmes capables d'apprendre à partir des données et d'améliorer leurs performances au fil du temps sans être explicitement programmés pour chaque tâche.

III.3.2.Principe :

Le ML offre la possibilité de résoudre des problèmes pour lesquels il n'y a pas de solutions connues, des problèmes dont les solutions existent mais n'ont pas été formalisées en algorithmes, et des problèmes pour lesquels les solutions existantes sont trop coûteuses en termes de ressources informatiques. Il est employé lorsqu'il y a une grande quantité de données mais que les connaissances requises sont peu disponibles .Le ML repose sur deux piliers fondamentaux [56] :

- Les données, qui sont les exemples à partir duquel l'algorithme va apprendre.
- L'algorithme d'apprentissage, qui est la procédure que l'on fait tourner sur ces données pour produire un modèle. On désigne par "entraînement" l'action de faire fonctionner un algorithme d'apprentissage sur un jeu de données.

III.3.3.Types d'apprentissage automatique :

Les systèmes d'apprentissage automatique peuvent être regroupés en plusieurs catégories, chacune avec ses propres méthodes et applications comme illustré dans la **Figure III.3**. Ces catégories comprennent l'apprentissage supervisé, l'apprentissage non supervisé, semi-supervisé et l'apprentissage par renforcement [57].

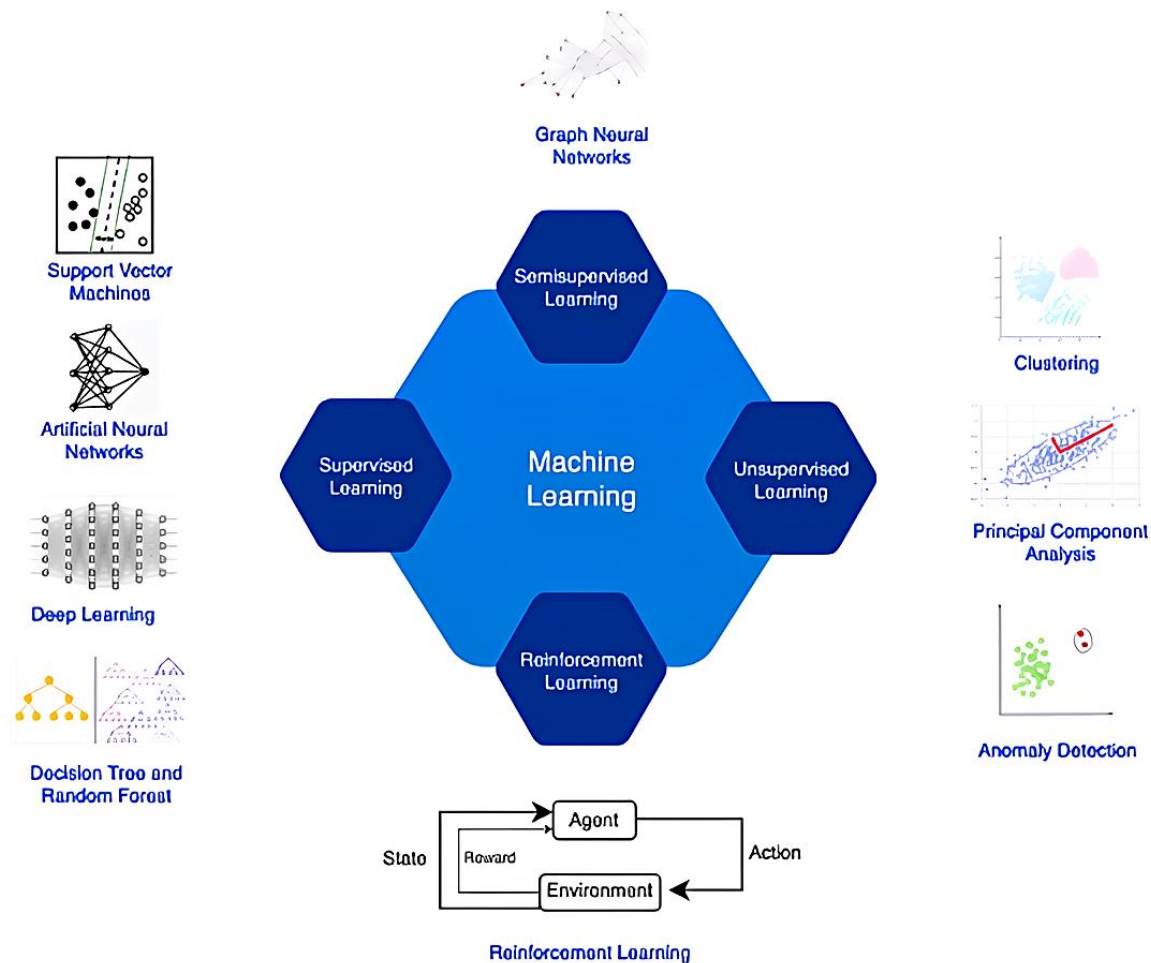


Figure III.3 : Classification des types d'apprentissage automatique [57].

III.3.3.1.Apprentissage supervisé :

Le processus d'apprentissage supervisé est une approche de ML qui se base sur l'entraînement avec des données d'entrée et leurs données de sortie exactes. Ainsi, le modèle utilise l'apprentissage de la relation entre les valeurs données et leurs valeurs d'étiquette (sortie) pour prédire la sortie lorsque des données d'entrée sont fournies au modèle [58].

Parmi les principaux algorithmes d'apprentissage supervisé, on retrouve [57] :

- Les classifieurs de Bayes.
- La régression linéaire et logistique.
- Les arbres de décision et les forêts aléatoires.
- Les machines à vecteurs de support (SVM).
- Les réseaux neuronaux artificiels (en général, les RNA sont entraînés de manière supervisée, mais il existe aussi des RNA non supervisés et semi-supervisés).

III.3.3.2.Apprentissage non supervisé :

L'apprentissage non supervisé peut apprendre sans orientation explicite, comme son nom l'indique. Après avoir fourni des données non étiquetées à un modèle non supervisé, celui-ci peut détecter certains motifs en fonction du type d'algorithme utilisé ou de l'objectif du formateur [58]. La **Figure III.4** décrit le fonctionnement de l'apprentissage supervisé et non supervisé.

Des exemples d'algorithmes de ce type sont [57] :

- Clustering (regroupement)
- Réduction de la dimensionnalité
- Détection d'anomalies

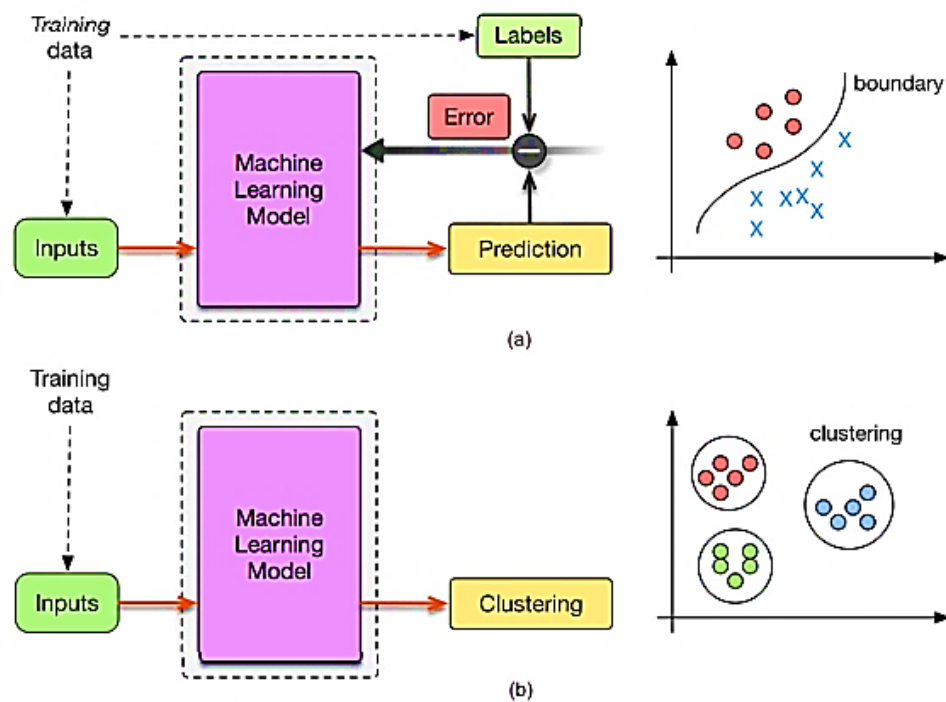


Figure III.4 : Fonctionnement de : **a)** Supervised Learning, **b)** Unsupervised Learning [59].

III.3.3.3.Apprentissage semi-supervisé :

L'apprentissage semi-supervisé est une approche de l'apprentissage automatique où un modèle est entraîné sur un ensemble de données qui contient à la fois des exemples étiquetés et des exemples non étiquetés. Cette méthode est particulièrement bénéfique lorsque l'obtention d'un grand nombre d'exemples étiquetés est coûteuse ou difficile, car elle permet d'utiliser de manière efficace les données non étiquetées afin d'améliorer les performances du modèle [57].

III.3.3.4.Apprentissage par renforcement :

Dans le cadre de l'apprentissage par renforcement, un agent est entraîné en utilisant une récompense pour accomplir certaines tâches dans un cadre. À la différence de l'apprentissage supervisé, le modèle de l'agent n'est pas directement formé à partir de données étiquetées. Au lieu de cela, il est nécessaire que l'agent interagisse avec l'environnement et collecte des données sur les actions qui entraînent la récompense la plus élevée, puis ajuste sa politique en réaction.

On peut observer l'interaction entre l'agent et son environnement dans la **Figure III.5**. Chaque étape de temps donnée offre à l'agent un état et une récompense à travers l'environnement. Par la suite, l'agent doit identifier une action adéquate qui est renvoyée dans l'environnement. L'état suivant de l'environnement est influencé par cela et le cycle se répète [60].

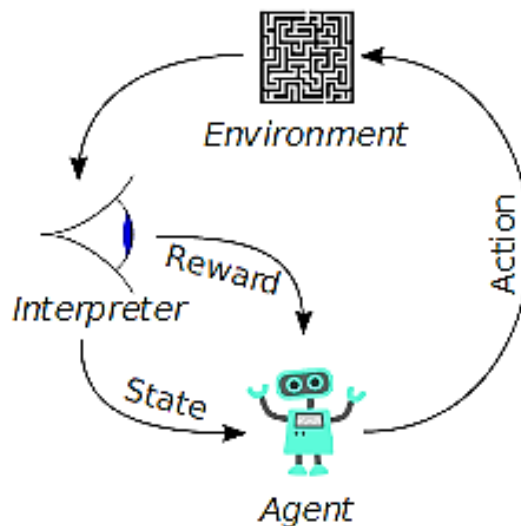


Figure III.5: Illustration de l'apprentissage par renforcement [60].

III.3.4.Limites de L'apprentissage automatique :

Le ML nécessite des données pour détecter des modèles, cependant, des problèmes peuvent survenir en raison de la dimensionnalité et de la stagnation des performances lorsque de nouvelles données sont introduites au-delà d'un seuil critique. Il a été constaté qu'une réduction significative des performances se produit dans de telles situations. La **Figure III.6** suivante met en évidence la réussite du DL en termes de performances, qui s'améliore avec l'augmentation des données, contrairement aux algorithmes de ML classiques, qui atteignent un plateau et montrent leurs limites [55].

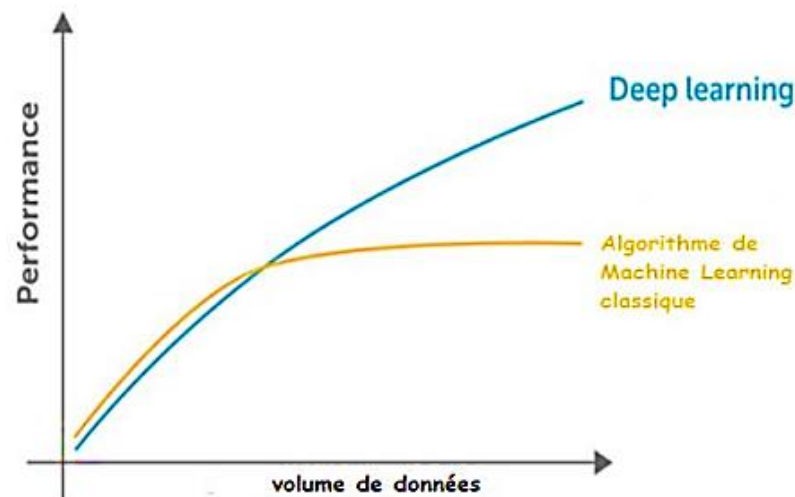


Figure III.6 : Qualité des données de ML et DL [55].

III.4. L'apprentissage en profondeur :

III.4.1.Définition :

Le DL, une sous-discipline de ML basée sur les réseaux de neurones artificiels profonds, a récemment connu un essor considérable, notamment en raison de ses performances remarquables dans des tâches telles que la reconnaissance d'images, la traduction automatique, et la génération de contenu [51]. C'est un algorithme qui tente d'utiliser l'abstraction de haut niveau des données en utilisant plusieurs couches de traitement composées de structures complexes ou de transformations non linéaires multiples [61] .

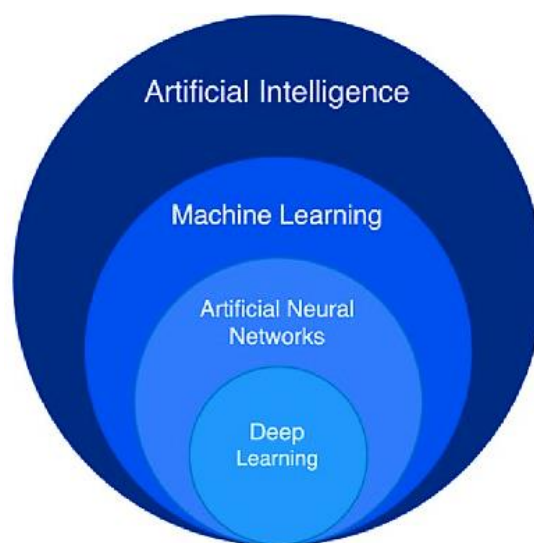


Figure III.7: Le DL, une sous-discipline de ML [57].

III.4.2. Les réseaux de neurones :

III.4.2.1. Définition et présentation :

En tant qu'algorithme d'apprentissage automatique le plus représentatif, le réseau neuronal artificiel est largement utilisé dans divers scénarios d'application et mène continuellement le développement de l'IA. Le premier réseau de neurones artificiels est le perceptron (**Rosenblatt, 1957**) [56]. Ces réseaux, composés d'un grand nombre de neurones, pouvant aller jusqu'à des millions, et de plusieurs couches algorithmiques, sont conçus pour simuler le fonctionnement du cerveau humain. L'une des caractéristiques principales de l'apprentissage profond est son aptitude à traiter de vastes quantités de données par rapport aux méthodes traditionnelles d'apprentissage automatique [51].

III.4.2.2. Le neurone artificiel :

En IA, un neurone est une unité fondamentale de traitement de l'information dans un réseau de neurones artificiels. Inspiré du fonctionnement des neurones biologiques du cerveau humain, un neurone artificiel reçoit des signaux d'entrée pondérés, les combine à l'aide d'une fonction d'activation, et produit une sortie [62].

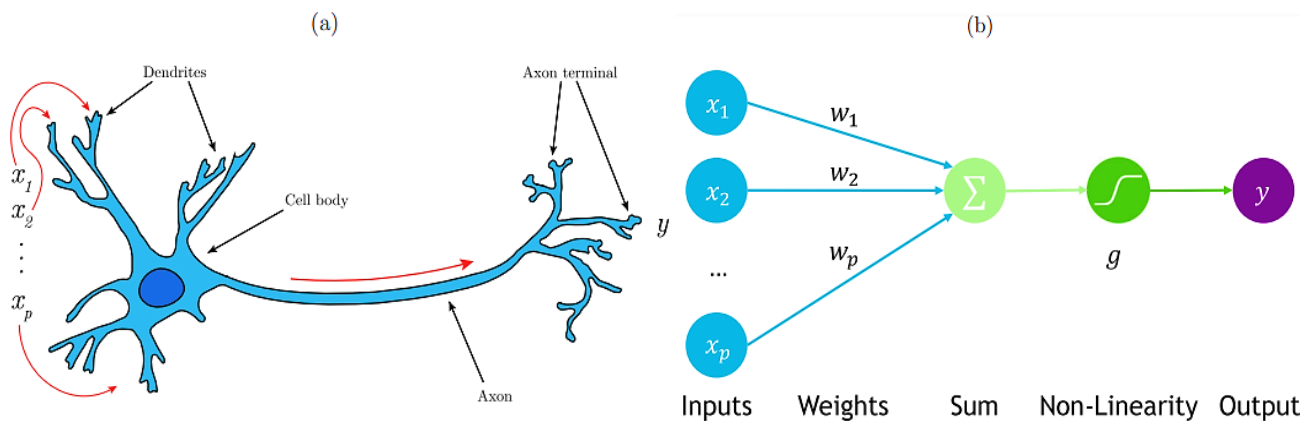


Figure III.8 : Illustration de : (a) le neurone biologique et (b) le neurone artificiel [62].

Dans la **Figure III.8 (a)** illustrée, les synapses forment l'entrée du neurone. Leurs signaux sont combinés, et si le résultat dépasse un seuil donné, le neurone est activé et produit un signal de sortie qui est envoyé à travers l'axone. Dans la **Figure III.8 (b)** le perceptron un neurone artificiel inspiré par la biologie. Il est composé de l'ensemble des entrées (qui correspondent aux informations entrant dans les synapses) x_i , qui sont combinées linéairement avec des poids w_i puis passent à travers une fonction non linéaire g pour produire une sortie y [62].

III.4.2.3. Le fonctionnement d'un neurone :

La **Figure III.9** représente un nœud qui reçoit trois entrées. Le cercle et la flèche dans le schéma représente respectivement le nœud et le flux de signal. x_1 , x_2 et x_3 représentent les signaux d'entrée. Les poids correspondants pour ces signaux sont w_1 , w_2 et w_3 . Enfin, b représente le biais, qui est un autre paramètre lié au stockage de l'information. On peut dire que les informations du réseau neuronal sont encodées sous forme de poids et de biais [63].

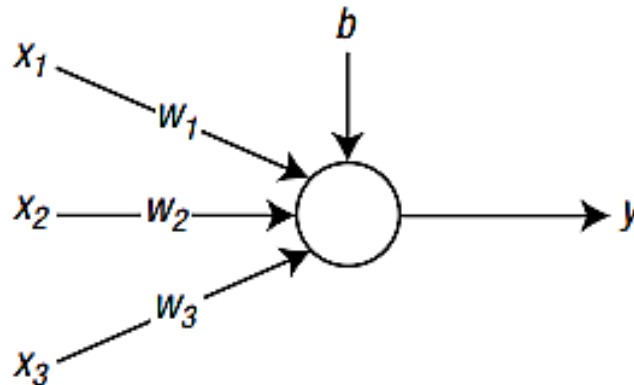


Figure III.9 : Un nœud avec trois entrées [63].

. Le processus suivant se déroule à l'intérieur du nœud du réseau neuronal [63] :

- Le signal d'entrée provenant de l'extérieur est multiplié par le poids avant d'atteindre le nœud. Une fois que les signaux pondérés sont collectés au nœud, ces valeurs sont additionnées pour former la somme pondérée. La somme pondérée de cet exemple est calculée comme suit:

$$v = (x_1 w_1) + (x_2 w_2) + (x_3 w_3) + b \quad (\text{III.1})$$

- Cette équation indique que le signal avec un poids plus important a un effet plus important. L'équation de la somme pondérée peut être écrite avec des matrices comme suit :

$$v = wx + b \quad (\text{III.2})$$

w et x sont définis comme suit :

$$w = [w_1 \ w_2 \ w_3] \quad \text{et} \quad x = \begin{bmatrix} x_1 \\ x_2 \\ x_3 \end{bmatrix}$$

- Enfin, le nœud entre la somme pondérée dans la fonction d'activation et produit sa sortie. La fonction d'activation détermine le comportement du nœud.

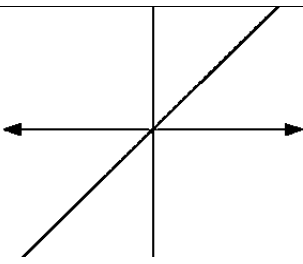
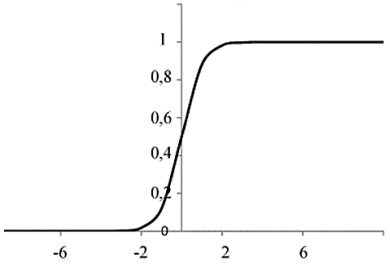
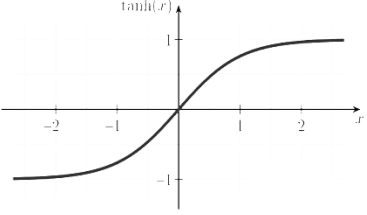
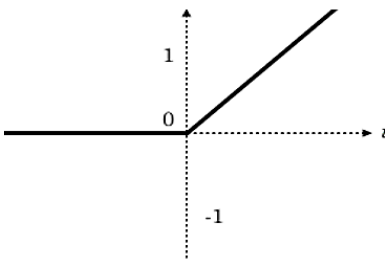
$$y = \mathcal{P}(v) \quad (\text{III.3})$$

$\mathcal{P}$ représente la fonction d'activation.

III.4.2.4. La fonction d'activation :

La fonction d'activation comme son nom l'indique, est une formule mathématique (algorithme) activée dans certaines circonstances. Une fonction de transfert permet de déterminer la valeur de l'état du neurone. Cette valeur sera transmise aux neurones aval.

La fonction de transfert peut prendre différentes formes. Les plus fréquentes sont illustrées dans le **Tableau III.1 [64]**.

Fonction de transfert	Formule	Représentation graphique
Linéaire	$f(x)=x$	
Sigmoïde	$f(x) = \frac{1}{1 + e^{-x}}$	
Tangente hyperbolique	$(f(x) = \tanh(x))$	
ReLU (Rectified Linear Unit)	$f(x) = \max(0, x)$	

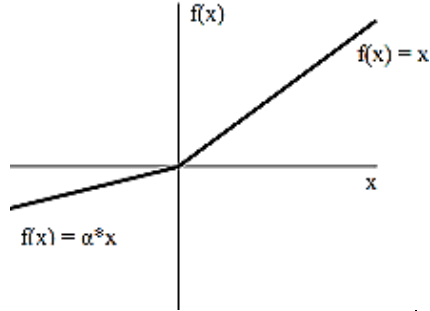
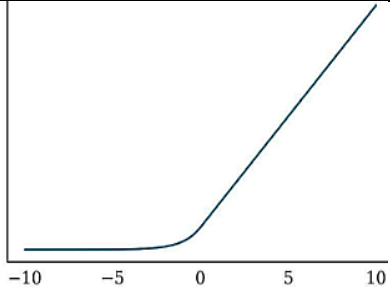
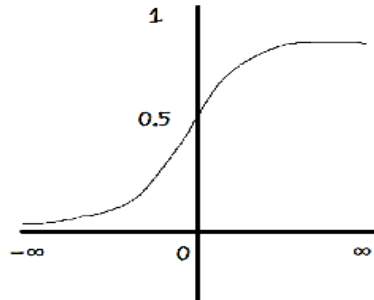
Leaky ReLU	$f(x) = \begin{cases} x & \text{si } x > 0 \\ \text{Pente} \times x & \text{si non} \end{cases}$	
ELU (Exponential Linear Unit)	$f(x) = \begin{cases} x & \text{si } x > 0 \\ \alpha(e^x - 1) & \text{si non} \end{cases}$	
Softmax	$f(x_i) = \frac{e^{x_i}}{\sum_j e^{x_j}}$	

Tableau III.1 : Les fonctions d'activation les plus utilisées.

III.4.2.5. Structure des réseaux de neurones:

Les architectures des réseaux neuronaux ont évolué d'une simplicité initiale vers des structures de plus en plus sophistiquées. À leurs débuts, les pionniers des réseaux neuronaux se contentaient d'une configuration rudimentaire comprenant uniquement des couches d'entrée et de sortie, lesquelles étaient désignées sous le terme de réseaux neuronaux à une seule couche. L'incorporation de couches cachées à ces réseaux neuronaux à une seule couche a marqué l'avènement des réseaux neuronaux multicouches.

Ainsi, un réseau neuronal multicouche est constitué d'une couche d'entrée, d'une couche cachée et d'une couche de sortie. Un réseau neuronal doté d'une unique couche cachée est qualifié de réseau neuronal peu profond ou classique, tandis qu'un réseau neuronal multicouche comportant deux couches cachées ou plus est désigné sous le nom de réseau neuronal profond. La plupart des réseaux neuronaux contemporains employés dans des applications réelles sont des réseaux neuronaux profonds [63].

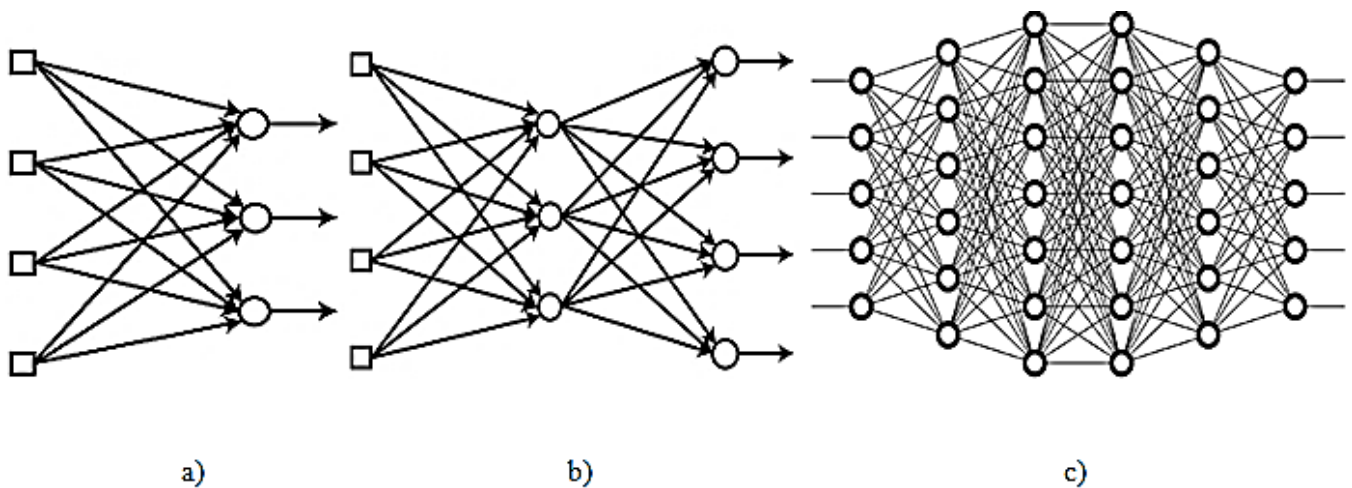


Figure III.10: Représentation de différentes structures de RNA ; **a)** Réseau neuronal à une seule couche, **b)** Réseau neuronal multicouche, **c)** Réseau neuronal profond [63].

III.4.3.Les sous-domaines:

Les réseaux de neurones comprennent de nombreux sous-domaines qui évoluent en permanence. Voici quelques-uns des plus significatifs :

III.4.3.1.Les réseaux de neurones artificiels (ANN) :

Également connu sous le nom de réseau de neurones à propagation avant .Ces réseaux sont constitués de neurones artificiels interconnectés, organisés en couches, et capables d'apprendre à partir des données. Chaque neurone prend des entrées, les combine à l'aide de poids associés, et applique une fonction d'activation pour produire une sortie. Ce type de réseaux neuronaux est l'une des variantes les plus simples des réseaux neuronaux. Sa formulation initiale comportait d'importantes limitations comme l'impossibilité de résoudre des problèmes non-linéaires [65].

III.4.3.2.Les réseaux de neurones convolutifs (CNN) :

Les CNN désignent une sous-catégorie de réseaux de neurones. Cependant, les CNN sont spécialement conçus pour traiter des images en entrée. Leur architecture est alors plus spécifique, contient plusieurs couches différentes avec diverses fonctions. Un CNN est simplement un empilement de plusieurs couches de convolution, pooling, fully-connected. Chaque image reçue en entrée va donc être filtrée, réduite et corrigée plusieurs fois, pour finalement former un vecteur [51].

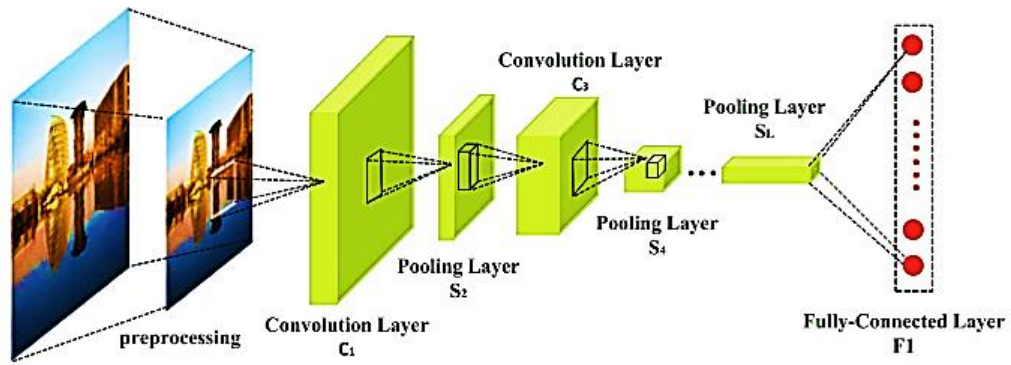


Figure III.11: L'architecture de base d'un CNN [66].

Dans un réseau CNN, la couche de convolution est chargée d'extraire des caractéristiques à partir des données d'entrée, comme des images. Pour parcourir l'image et évaluer des caractéristiques locales, comme les contours, les textures ou les motifs, elle utilise des filtres de convolution. Tous les filtres sont utilisés pour des zones de l'image, connues sous le nom de "fenêtres de convolution", et génèrent une carte de caractéristiques qui met en évidence les zones où la caractéristique spécifique est repérée. Pendant l'apprentissage du réseau, les filtres sont partagés et acquis, ce qui permet au CNN de repérer des motifs locaux, peu importe leur position dans l'image. L'utilisation fréquente de la couche de pooling, qui suit la couche de convolution, vise à diminuer la taille spatiale des cartes de caractéristiques et à gérer le surapprentissage [66].

III.4.3.3. Les réseaux de neurones récurrents (RNN):

Le RNN est un réseau où les connexions entre les unités créent un cycle dirigé. La **Figure III. 12** présente l'architecture la plus simple. Les RNN ont la capacité de traiter des données séquentielles, comme des vidéos et des phrases vocales, en montrant des comportements temporels dynamiques grâce aux connexions récurrentes. Au lieu d'apprendre le posteriori d'un vecteur d'entrée unique x , les RNN apprennent le posteriori d'une séquence $x_1, x_2, \dots, x_t$ et sont plus sensibles aux séquences. La couche cachée h_t à l'instant t est présente dans l'architecture de base des RNN, avec l'entrée à propagation avant x_t et l'entrée récurrente h_{t-1} [66]:

$$h_t = f(W^f x_t + W^r h_{t-1} + b) \quad (\text{III.4})$$

- W^f : la matrice de poids pour l'entrée à propagation avant.
- W^r : pour l'entrée récurrente.
- b : est le vecteur de biais
- h_t : est la sortie des unités cachées.

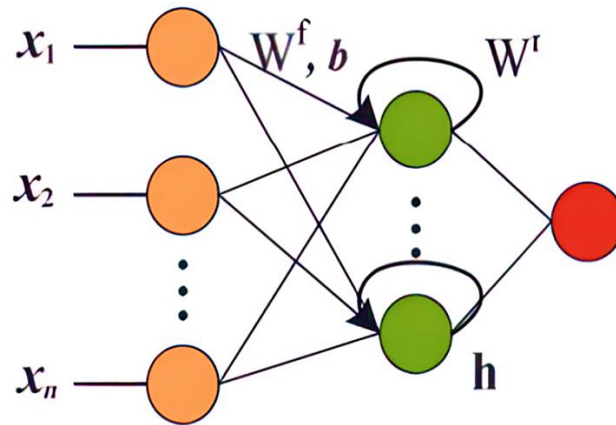


Figure III. 12: L'architecture de base d'un RNN [66].

III.4.3.4.L'apprentissage par renforcement profond (DRL) :

Le DRL est en effet une méthode puissante qui combine les principes de RL avec les techniques du DL. En utilisant des réseaux neuronaux profonds pour approximer la politique, les valeurs d'état-action ou les fonctions de valeur, la DRL peut résoudre des problèmes complexes dans des environnements à grande dimension [67] [68].

Les avantages de la DRL résident dans sa capacité à gérer des espaces d'état à haute dimension et des tâches complexes, ce qui la rend particulièrement adaptée à un large éventail d'applications, notamment la robotique, les véhicules autonomes, les jeux, la finance et les soins de santé [67].

III .4.4.Le processus d'apprentissage:

III .4.4.1.Principe :

Par définition 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. L'apprentissage neuronal fait appel à des exemples de comportement [64].

Apprendre, c'est pouvoir stocker des données pour les retrouver plus tard. Avec les réseaux de neurones, les données sont stockées dans des connexions entre les neurones. Ces connexions sont réduites pendant l'apprentissage pour minimiser l'écart entre les résultats attendus et ceux générés par le réseau. C'est grâce à ces ajustements de poids synaptiques que les réseaux de neurones s'adaptent aux données deviennent et précises. L'entraînement consiste à présenter, de façon répétitive, un ensemble de données pour permettre au réseau de neurones d'apprendre à résoudre ce problème en ajustant correctement ses poids.

III.4.4.2. La Propagation avant et arrière :

La propagation avant et la rétropropagation sont deux concepts clés dans le fonctionnement des réseaux de neurones. La propagation avant consiste à transmettre les données à travers le réseau afin de produire des prédictions, tandis que la rétropropagation est employée pour ajuster les poids du réseau afin d'améliorer ses performances en réduisant l'erreur entre les prédictions et les vraies valeurs. Le principe des propagations avant et arrière est illustré dans la **Figure III.13**.

III.4.4.2.1. La Propagation avant :

Dans la propagation vers l'avant de l'information, le flux de données se déplace strictement des unités d'entrée vers les unités de sortie, sans aucune boucle de rétroaction. Le traitement des données peut impliquer plusieurs unités, mais il n'y a pas de connexion récurrente où les sorties des unités influencent directement leurs entrées [69].

III.4.4.2.2. La Propagation arrière :

Le principe de cette méthode est de déterminer l'effet de chaque poids du réseau sur l'erreur de sortie, puis d'ajuster ces poids en fonction de cette contribution afin de diminuer l'erreur. L'erreur de sortie est propagée à travers les différentes couches du réseau grâce à ce processus itératif, qui calcule les gradients de poids par rapport à celle-ci. Par la suite, les gradients obtenus sont employés dans un algorithme d'optimisation, comme la descente de gradient stochastique, afin de modifier les poids et réduire l'erreur de sortie. On répète ce processus sur toutes les données d'entraînement jusqu'à ce que le réseau se rassemble en un ensemble de poids qui permet de prédire précisément les données d'entraînement [63].

Prenons en compte l'algorithme de backpropagation pour entraîner le réseau neuronal [63]:

1. Initialiser les poids avec des valeurs adéquates.
2. Entrer l'entrée à partir des données d'entraînement {entrée, sortie correcte} et obtenir la sortie du réseau neuronal.
3. Calculer l'erreur de la sortie (e) par rapport à la sortie correcte (y) et le gradient de l'erreur par rapport à la sortie d'un nœud (δ).

$$e = d - y \quad (\text{III.5})$$

$$\delta = \varphi'(v)e \quad (\text{III.6})$$

Propager le delta du nœud de sortie, δ , vers l'arrière, et calculer les deltas des nœuds.

$$e^{(k)} = W^T \delta \quad (\text{III.7})$$

$$\delta^{(k)} = \varphi'(v^{(k)}) e^{(k)} \quad (\text{III.8})$$

4. Répéter l'étape 4 jusqu'à ce qu'elle atteigne la couche cachée qui est immédiatement à droite de la couche d'entrée.
5. Ajuster les poids selon la règle d'apprentissage suivante :

$$\Delta w_{ij} = a \delta_i x_j \quad (\text{III.9})$$

$$w_{ij} \rightarrow w_{ij} + \Delta w_{ij} \quad (\text{III.10})$$

6. Répéter les étapes 2 à 6 pour chaque point de données d'entraînement.
7. Répéter les étapes 2 à 7 jusqu'à ce que le réseau neuronal soit correctement entraîné.

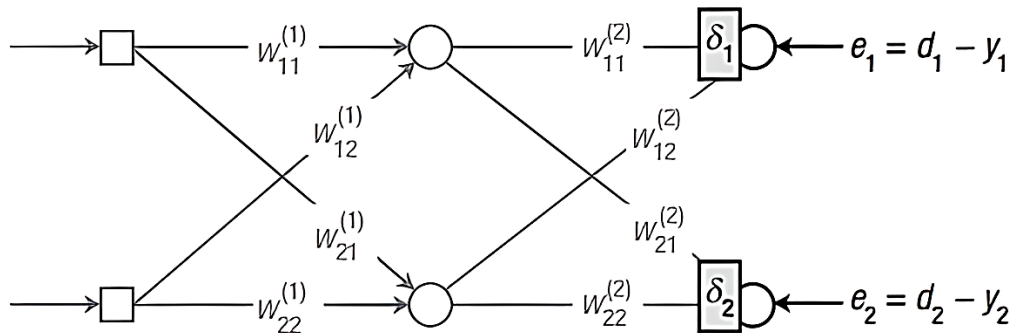


Figure III.13: Entraîner le réseau neuronal en utilisant l'algorithme de rétropropagation [4].

Bien que cet exemple n'ait qu'une seule couche cachée, l'algorithme de rétropropagation est applicable pour entraîner de nombreuses couches cachées.

III.4.4.3. Les techniques de prétraitement des données :

III.4.4.3.1. Principe et bénéfices :

Cette première étape revêt une grande importance car le prétraitement des données est l'une des étapes essentielles lorsqu'on traite de n'importe quel type de modèle. Il n'y a jamais de règle définitive pour cette étape et elle dépend entièrement de l'introduction du problème [70].

Son principale fonction consiste à préparer les données brutes pour une analyse ou une modélisation ultérieure. Cela nécessite diverses mesures, telles que le nettoyage des données afin de repérer et de corriger les erreurs, les valeurs manquantes et les incohérences, la transformation des

données afin de les adapter à l'analyse, la réduction de la dimensionnalité pour gérer la complexité des données, la gestion des valeurs aberrantes afin d'éviter qu'elles ne perturbent les résultats, et l'amélioration de la performance des modèles en diminuant le bruit et en facilitant la convergence des algorithmes. Un traitement préliminaire adéquat des données favorise des résultats plus précis, fiables et compréhensibles, ce qui est crucial pour prendre des décisions éclairées en se basant sur les données.

III.4.4.3.2. La standardisation :

La standardisation est une méthode de prétraitement employée dans le domaine de DL afin de mesurer les caractéristiques de manière à obtenir une moyenne de 0 et un écart type de 1. Cela assure la cohérence des caractéristiques, ce qui facilite l'apprentissage et l'interprétation des relations entre elles.

La formule de la standardisation est exprimée :

$$x_{\text{Standardisé}} = \frac{x - \text{moyenne}(x)}{\text{écart type}(x)} \quad (\text{III.11})$$

Où :

- x : représente la valeur d'une caractéristique,
- $\text{moyenne}(x)$: est la moyenne des valeurs de cette caractéristique dans l'ensemble de données,
- $\text{Ecart type}(x)$: est l'écart type des valeurs de cette caractéristique dans l'ensemble de données.

III.4.4.3.3. La normalisation et la normalisation par lots :

La normalisation est une technique de prétraitement des données qui vise à mettre à l'échelle les valeurs des caractéristiques d'un ensemble de données à un intervalle spécifique ou à une distribution spécifique, généralement entre le 0 et 1. Au cours de l'apprentissage d'un réseau neuronal, la répartition des valeurs d'entrée de chaque couche est influencée par toutes les couches précédentes. Cette variabilité a été résolue grâce à la création de la normalisation par lots, ce qui a accéléré l'apprentissage.

Il est déjà courant de normaliser les valeurs de chaque échantillon avant de les utiliser dans la couche d'entrée du réseau. La normalisation par lots s'étend davantage et établit une normalisation de chaque couche du réseau, et non seulement la couche d'entrée. La normalisation est calculée pour chaque mini-lot [71].

La formule de normalisation est la suivante :

$$x_{\text{Normalise}}' = \frac{\max(x) - \min(x)}{x - \min(x)} \quad (\text{III.12})$$

Ici, x est la valeur d'une caractéristique, $\min(x)$ est la valeur minimale de cette caractéristique et $\max(x)$ est la valeur maximale de cette caractéristique.

La formule complète de la batch normalisation est la suivante :

$$\hat{x} = \frac{x - \mu}{\sqrt{\sigma^2 + \epsilon}} \quad (\text{III.13})$$

Où :

- x : l'activation d'une unité dans la couche,
- μ : la moyenne des activations sur le mini-lot,
- σ^2 : la variance des activations sur le mini-lot,
- ϵ : un terme d'épargne (epsilon) ajouté pour éviter la division par zéro,
- $\hat{x}$: l'activation normalisée.

III.4.4.3.4. L'encodage des caractéristiques catégorielles :

Le prétraitement des variables catégorielles consiste à convertir des données catégorielles, telles que des étiquettes ou des noms de catégories, en une forme numérique appropriée pour être utilisée dans des modèles d'apprentissage automatique. Cela est nécessaire car de nombreux algorithmes d'apprentissage automatique ne peuvent pas traiter directement les données catégorielles. L'objectif est de représenter de manière adéquate les informations catégorielles afin que le modèle puisse les utiliser pour effectuer des prédictions ou des analyses [70].

Il existe plusieurs techniques sont utilisées pour encoder les variables catégorielles comprennent l'encodage one-hot ou Chaque catégorie unique est représentée par un vecteur binaire où une seule dimension est activée (1) et toutes les autres sont désactivées (0), l'encodage ordinal, l'encodage binaire, l'encodage de fréquence, l'encodage cible, et bien d'autres encore [70].

III.4.4.3.5. La discrétisation :

La discrétisation est une technique qui convertit un attribut continu ou numérique en données catégoriques ou nominales. Cette transformation est souvent nécessaire dans diverses tâches lorsque

l'utilisation directe de valeurs numériques n'est pas pertinente d'un point de vue computationnel, ou lorsque des informations utiles ne peuvent pas être extraites à partir de ces valeurs continues. Cela peut être utile pour simplifier l'analyse ou pour rendre les données plus adaptées à certains algorithmes d'apprentissage automatique [72].

III.4.4.4. Les algorithmes d'optimisation:

III.4.4.4.1. GD :

Le GD est une méthode est couramment employée dans le domaine de l'apprentissage automatique et de l'optimisation afin de déterminer les paramètres d'un modèle qui réduisent au minimum une fonction de coût. Il fonctionne en mesurant le gradient de la fonction de coût par rapport aux paramètres du modèle, puis en ajustant ces paramètres dans la direction opposée du gradient, ce qui permet de descendre le long de la pente de la fonction de coût. On répète cette procédure de manière itérative jusqu'à ce qu'une convergence soit obtenue ou qu'un critère d'arrêt prédéfini soit dû. Les paramètres du réseau (θ) à optimiser sont mis à jour selon la formule suivante [73] :

$$\theta_{t+1} = \theta_t - \eta \cdot \nabla_{\theta} J(\theta_t) \quad (\text{III.14})$$

Où η est le taux d'apprentissage, $\nabla_{\theta} J(\theta_t)$ est le gradient de la fonction de perte $J(\theta_t)$ par rapport à θ_t .

III.4.4.4.2. SGD :

La fonction de perte est calculée par l'algorithme SGD pour un seul échantillon d'entraînement à la fois, plutôt que de prendre en compte tous les échantillons de données d'entraînement. Il est possible de prévenir les problèmes de mémoire de cette façon. Le SGD a été développé dans le but de corriger les lacunes de l'algorithme de descente de gradient par segments. Il s'agit de trouver la valeur adéquate du taux d'apprentissage pour éviter les fluctuations et atteindre l'optimum global lors de l'utilisation du SGD. Pour mettre à jour les paramètres, il se sert de l'équation suivante :

$$\theta_{t+1} = \theta_t - \eta \cdot \nabla_{\theta} J(\theta_t; x^{(i)}; y^{(i)}) \quad (\text{III.15})$$

Où $x^{(i)}$ et $y^{(i)}$ représentent les données d'entraînement sous forme de paires entrées-sorties.

III.4.4.4.3. RMSprop :

L'objectif de la technique RMSprop est d'adapter le taux d'apprentissage en fonction de chaque paramètre du modèle. Il évalue la moyenne pondérée des carrés des gradients récents pour

chaque paramètre. Cela permet d'ajuster le taux d'apprentissage de manière à ce qu'il soit plus élevé pour les paramètres avec des gradients plus faibles et inversement, ce qui peut favoriser une convergence plus rapide vers un minimum global dans l'optimisation de la fonction de perte [73].

$$g_t = \nabla_{\theta} J(\theta_t) \quad (\text{III.16})$$

$$v_t = \beta_1 v_{t-1} + (1 - \beta_1) g_t^2 \quad (\text{III.17})$$

$$\theta_{t+1} = \theta_t - \frac{\eta}{\sqrt{v_t + \epsilon}} g_t \quad (\text{III.18})$$

Où g_t est la première dérivée de la fonction de perte et β_2 est le taux de décroissance exponentielle.

III.4.4.4. Adam:

Adaptive Moment Estimation, conserve une moyenne mobile exponentielle des carrés de gradients précédents v_t et garde une moyenne mobile exponentielle des gradients précédents m_t . Cela donne à Adam la possibilité de s'ajuster de manière dynamique aux diverses caractéristiques et structures des données [74]. L'équation de l'algorithme Adam est décrite comme suit:

$$m_t = \beta_1 m_{t-1} + (1 - \beta_1) g_t \quad (\text{III.19})$$

$$v_t = \beta_2 v_{t-1} + (1 - \beta_2) g_t^2 \quad (\text{III.20})$$

$$\hat{m}_t = \left(\frac{m_t}{1 - \beta_1^t} \right) = \hat{v}_t \left(\frac{v_t}{1 - \beta_2^t} \right) \quad (\text{III.21})$$

$$\theta_{t+1} = \theta_t - \frac{\eta}{\sqrt{\hat{v}_t + \epsilon}} \hat{m}_t \quad (\text{III.22})$$

- m_t : est la moyenne mobile des gradients à l'instant t .
- v_t : est le facteur de mise à jour à l'instant t .
- $\hat{m}_t$: est la correction biaisée pour m_t .
- $\hat{v}_t$: est la correction biaisée pour v_t .
- θ_t : représente les paramètres du modèle à l'instant t ,
- α : est le taux d'apprentissage.
- β_1 et β_2 : sont les taux de décroissance exponentielle pour les moyennes mobiles des gradients et des carrés des gradients respectivement.
- g_t : est le gradient à l'instant t .

- ϵ : est un terme d'éviction pour éviter la division par zéro.
- t : est l'itération actuelle de l'algorithme.

III.4.4.4.5.Adamax :

Adamax est une version modifiée de l'algorithme d'optimisation Adam, principalement employée pour l'entraînement de réseaux de neurones. À la différence d'Adam, qui se sert d'un terme de deuxième moment adaptatif pour ajuster le taux d'apprentissage pour chaque paramètre, Adamax utilise une approximation de l'infini-norme des gradients afin de normaliser les mises à jour des paramètres. Cela renforce la résistance d'Adamax aux fluctuations dans la norme des gradients et peut parfois fournir de meilleures performances, notamment lorsque les gradients sont très dispersés [75]. L'équation de l'algorithme Adamax est exprimée comme suit:

$$m_t = \beta_1 m_{t-1} + (1 - \beta_1) g_t \quad (\text{III.23})$$

$$v_t = \max(\beta_2 v_{t-1}, ||g_t||) \quad (\text{III.24})$$

$$\hat{m}_t = \frac{m_t}{1 - \beta_1^t} \quad (\text{III.25})$$

$$\theta_{t+1} = \theta_t - \frac{\eta}{\sqrt{v_t + \epsilon}} \hat{m}_t \quad (\text{III.26})$$

La notation $||g_t||$ représente la norme L2 du gradient g_t .

Il existe plusieurs autres algorithmes d'optimisation populaires en plus de ces algorithmes, notamment, AdaGrad, AdaDelta, AMSgrad, NAG. Chacun de ces algorithmes a ses propres avantages et inconvénients, et leur choix dépend souvent du type de problème à résoudre et des caractéristiques des données.

III.4.4.5.Evaluation de modèle et fonctions de perte :

Trois éléments essentiels sont requis pour développer un modèle d'apprentissage automatique : la représentation, l'évaluation et l'optimisation. L'évaluation implique la définition et le calcul de mesures de performance quantitatives qui reflètent la qualité d'une représentation, c'est-à-dire la capacité d'un modèle spécifique à représenter l'association entre les entrées et les sorties. L'exactitude, la précision, le rappel, l'erreur quadratique moyenne et l'indice de Jaccard sont des mesures de performance fréquemment employées pour évaluer les modèles [76].

On utilise une fonction de perte lors de l'entraînement afin d'optimiser les paramètres du modèle. La fonction de perte évalue la différence entre les résultats prévisibles et attendus du

modèle, et l'objectif l'entraînement est de réduire cette différence. Une fonction de perte permet de quantifier la qualité des sorties du RNA [77]. Les fonctions de perte sont généralement classées en fonction du type de tâche pour lequel elles sont utilisées : classification ou régression. Voici un tableau récapitulatif des fonctions de perte les plus couramment utilisées :

Tâche :	Nom de la fonction :	Equation :
Classification	Entropie croisée binaire (BCE)	$L(y, p) = -(y \log(p) + (1 - y) \log(1 - p))$
	Entropie croisée binaire pondérée (WBCE)	$L = -(w_i \cdot y \log(p) + w_i(1 - y) \log(1 - p))$
	Entropie croisée catégorique (CCE)	$L = -\frac{1}{N} \sum_{i=1}^N \sum_{j=1}^C y_{i,j} \log(p_{i,j})$
	Entropie croisée catégorique (sparse CCE)	$H(Y, \hat{Y}) = -\frac{1}{n} \sum_{i=1}^n \log(\hat{y}_i, y_i),$
	Entropie croisée catégorique (Label Smoothed CCE)	$L(\hat{y}_i, y_i) = -\sum_{c=1}^C [(1 - \epsilon)y_c \log \hat{y}_c + \frac{\epsilon}{C} \log \hat{y}_c]$
	Perte de Hinge	$L(y, f(x)) = \max(0, 1 - y \cdot f(x))$
	Perte Focal	$FL(pt) = -\alpha_t(1 - p_t) \gamma \log(p_t),$
Régression	Erreur quadratique moyenne (MSE)	$MSE = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2,$
	Erreur absolue moyenne (MAE)	$MAE = \frac{1}{n} \sum_{i=1}^n y_i - \hat{y}_i $
	Erreur de Huber	$L(\hat{y}_i, y_i) = \begin{cases} \frac{1}{2}(y - \hat{y})^2 & \text{pour } y - \hat{y} \leq \delta \\ \delta(y - \hat{y} - \frac{1}{2}\delta) & \text{sinon} \end{cases}$
	Erreur de Log-Cosh	$L(\hat{y}_i, y_i) = \frac{1}{n} \sum_{i=1}^n (\cosh(\hat{y}_i, y_i)),$
	Perte de Quantile	$L(y, \hat{y}) = q \cdot \max(y - \hat{y}, 0) + (1 - q) \max(\hat{y} - y, 0)$
	Perte de Poisson	$L(y, \hat{y}) = \frac{1}{n} \sum_{i=1}^n (\hat{y}_i - y_i \log(\hat{y}_i))$

Tableau III.2 : Les fonctions de perte classées en fonction du type de tâche [77].

III.4.4.6.L'underfitting et l'overfitting:

III.4.4.6.1.Définition :

Un modèle a un faible biais s'il prédit bien les étiquettes des données d'entraînement. Si le modèle commet de nombreuses erreurs sur les données d'entraînement, nous disons que le modèle a un biais élevé ou qu'il sous-ajuste. Ainsi, l'underfitting est l'incapacité du modèle à bien prédire les

étiquettes des données sur lesquelles il a été entraîné. Il pourrait y avoir plusieurs raisons pour le sous-ajustement, parmi lesquelles les plus importantes sont [58] :

- Le modèle est trop simple pour les données.
- Les caractéristiques que nous avons conçues ne sont pas suffisamment informatives.

Un autre problème qu'un modèle peut rencontrer est le surajustement ou bien l'overfitting . Les données d'entraînement sont très bien prédites par le modèle qui surajuste, mais les données d'au moins l'un des deux ensembles de validation sont mal prédites. Parmi ses raisons [58] :

- Le modèle est trop complexe pour les données.
- Nombreuses caractéristiques mais peu d'exemples de formation.

La **Figure III.14** illustre la différence entre l'underfitting, good fit, et l'overfitting.

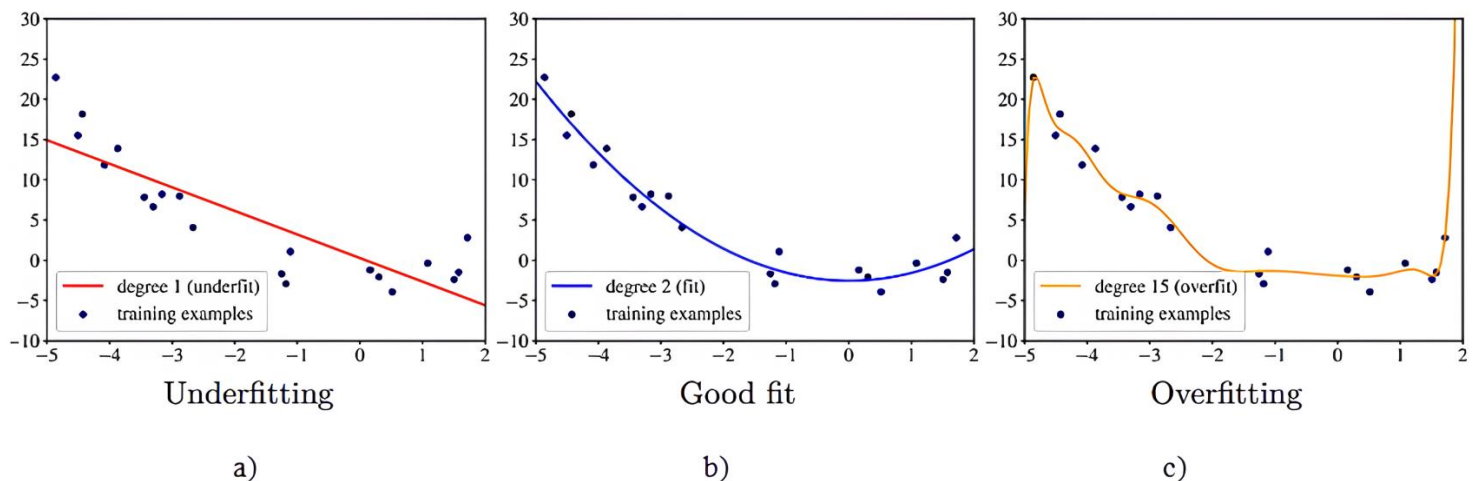


Figure III .14 : Exemples de **a)** sous-ajustement, **b)** ajustement optimal, **c)** surajustement [58].

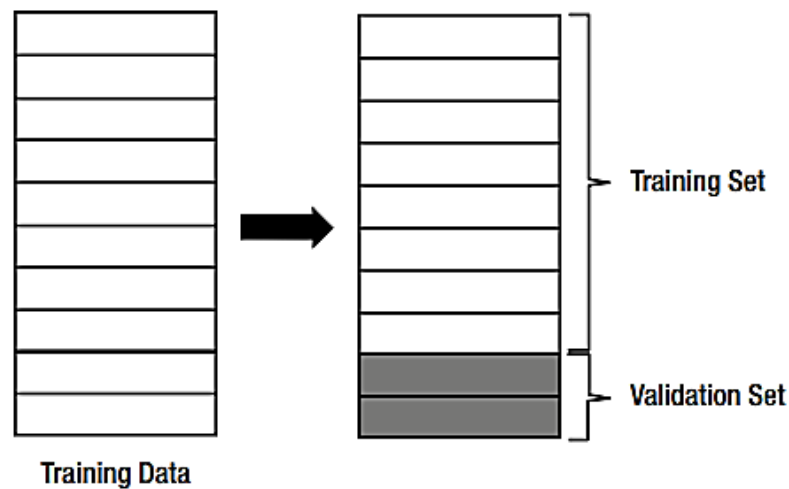
III.4.4.6.2.Des Solutions pour l'underfitting et l'overfitting :

La validation et la régularisation sont deux techniques essentielles en apprentissage automatique pour garantir que nos modèles généralisent bien sur des données qu'ils n'ont pas vues auparavant. Ces deux techniques sont présentées dans la section suivante.

A) La validation :

La validation est un processus qui réserve une partie des données d'entraînement et l'utilise pour surveiller la performance. Les données de validation désignent les informations fournies afin de trouver le meilleur ensemble d'hyperparamètres pour un modèle. Les données non utilisées lors de la

construction du modèle sont appelées ensemble de test (c'est-à-dire, les données non utilisées pour l'entraînement et l'ajustement final du modèle) [63].



La Figure III.15: Illustre la division des données d'entraînement pour le processus de validation [63].

B) La régularisation :

La régularisation est un ensemble de techniques employées dans le domaine de l'apprentissage automatique afin de diminuer l'erreur de généralisation. La majorité des modèles, une fois entraînés, sont extrêmement efficaces sur une minorité de la population globale, mais ne peuvent pas être généralisés de manière efficace. Les mesures de régulation ont pour objectif de diminuer l'utilisation excessive et de maintenir, en même temps, un taux d'erreur de formation le plus faible possible [78].

1. La régularisation L1 et L2 :

L1 et L2 sont les types les plus courants de régularisation. Ils mettent à jour la fonction de coût générale en ajoutant un autre terme appelé terme de régularisation [78]:

Fonction de coût = Perte (disons, entropie croisée binaire) + Terme de régularisation.

En raison de l'ajout de ce terme de régularisation, les valeurs des matrices de poids diminuent car on suppose qu'un réseau neuronal avec des matrices de poids plus petites conduit à des modèles plus simples. Par conséquent, cela réduira également le surapprentissage dans une large mesure. Cependant, ce terme de régularisation diffère entre L1 et L2.

- Dans L2, nous avons:

$$Cost\ function = loss + \frac{\lambda}{2m} * \sum ||W||^2 \quad (III.27)$$

Où : $||w|| = \sum_{i=1}^n w_i^2$ et λ est le paramètre de régularisation. C'est l'hyperparamètre dont la valeur est optimisée pour de meilleurs résultats.

La régularisation L2 est également connue sous le nom de décroissance des poids car elle force les poids à décroître vers zéro (mais pas exactement zéro).

- Dans L1, nous avons :

$$Cost\ function = loss + \frac{\lambda}{2m} * \sum ||W|| \quad (III.28)$$

Dans cette situation, la valeur absolue des poids est mise en péril. À la différence de L2, il est possible de réduire les poids à zéro ici. Ainsi, cela est extrêmement bénéfique lorsque nous cherchons à réduire la taille de notre modèle. Autrement, nous avons tendance à privilégier L2. De la même manière, il est possible d'utiliser la régularisation L1.

2. L'augmentation des données :

L'augmentation des données est une technique de régularisation qui aide un réseau neuronal à mieux généraliser en l'exposant à un ensemble plus diversifié d'exemples de formation. Étant donné que les réseaux de neurones profonds nécessitent un vaste ensemble de données de formation, l'augmentation des données est également utile lorsque nous ne disposons pas de données suffisantes pour former un réseau de neurones [79].

3. Le Dropout :

La méthode de dropout a pour but de réduire le surapprentissage. La clé de son concept est de choisir un modèle qui surapprend et de créer des sous-modèles dérivés de celui-ci en retirant aléatoirement des unités pour chaque lot d'activité. Dans la **Figure III.16**, on peut observer une représentation conceptuelle du réseau dense standard en comparant la structure du réseau après l'utilisation du dropout. En supprimant régulièrement des unités aléatoires, le dropout contraint les unités à être plus robustes, à développer des caractéristiques par elles-mêmes, sans avoir besoin de l'assistance des autres unités. Dans ce cas, on peut l'assimiler à un ensemble de modèles simplex. Le taux de dropout est un nouveau paramètre hyperparamétrique qui définit le nombre d'unités à conserver. Le taux de dropout recommandé est de 0,1 pour la couche d'entrée et de 0,5 à 0,8 pour les couches internes [71].

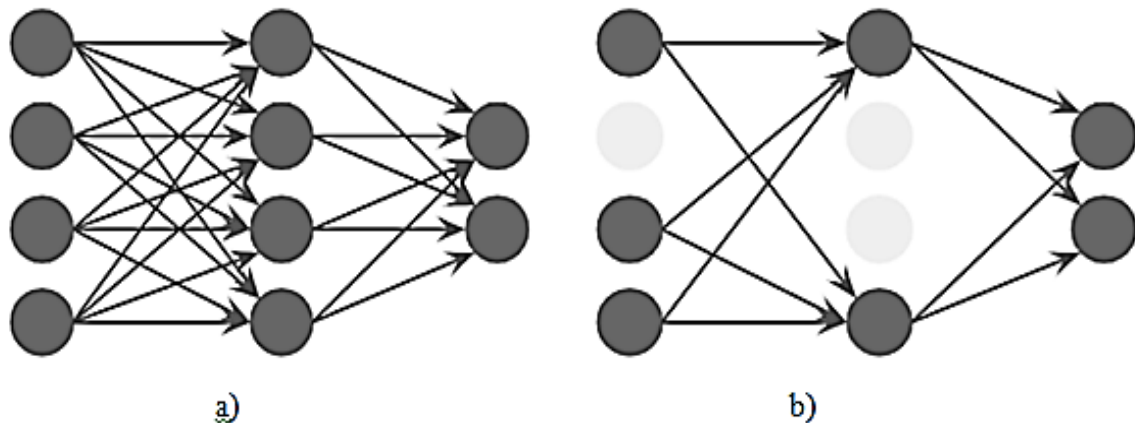


Figure III.16 : Représentation conceptuelle du ; **a)** réseau dense standard en comparaison avec **b)** la structure du réseau après l'application du dropout [71].

Ces méthodes de régularisation sont essentielles pour développer des modèles de ML robustes qui généralisent bien sur de nouvelles données et évitent le surajustement aux données d'entraînement. Il existe aussi une autre méthode qui combine à la fois la régularisation L1 et L2 en ajoutant une pénalité qui est une combinaison linéaire des pénalités L1 et L2. De plus la régularisation par poids maximal, elle limite la taille maximale des coefficients du modèle qui peut aider à éviter l'overfitting.

III.4.5. Les avantages de Deep Learning :

Un des avantages majeurs des modèles de DL par rapport aux méthodes traditionnelles de ML réside dans leur capacité à créer de nouvelles caractéristiques à partir d'un ensemble restreint de caractéristiques dans le système de données entraîné. Ces modèles ont une grande capacité à résoudre non seulement les tâches actuelles, mais ils ont également le potentiel de générer de nouvelles tâches, englobant différents aspects de la vie. Le traitement de grands ensembles de données peut être considérablement accéléré grâce aux algorithmes de DL qui génèrent automatiquement des caractéristiques sans nécessiter l'intervention humaine [80].

III.5. Les cas d'utilisation de l'intelligence artificielle dans la 5G :

L'intelligence artificielle et l'apprentissage automatique connaissent une croissance constante dans différents secteurs, et les communications sans fil et les réseaux ne font pas exceptionnellement exception. La **Figure III.17** décrit diverses opportunités d'IA pour les réseaux mobiles.

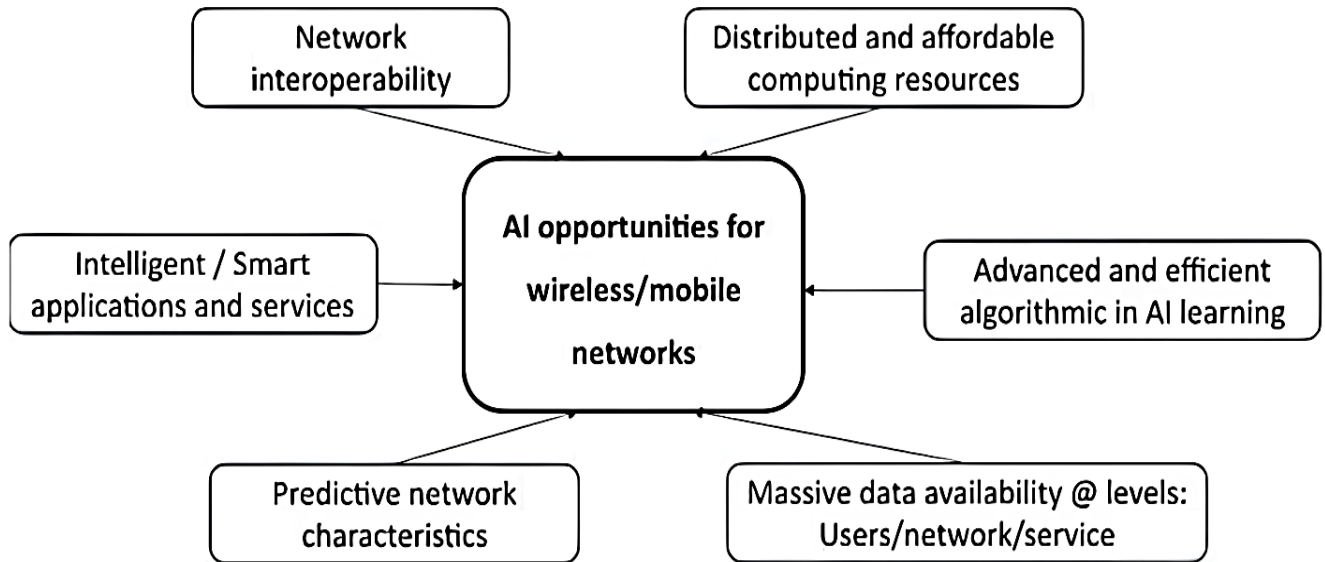


Figure III.17 : Diverses opportunités d'IA pour les réseaux mobiles [81].

La mise en place de l'IA dans la 5G marque une avancée significative dans les réseaux de communication sans fil. Le résultat de cette combinaison présente un potentiel innovant pour optimiser les performances, la sécurité et l'efficacité opérationnelle de la 5G. L'IA peut être employée dans différents secteurs de la 5G, tels que [81] :

- L'intégration croissante de l'intelligence artificielle dans les réseaux 5G permet de proposer des solutions novatrices pour optimiser différents aspects des communications sans fil. L'intelligence artificielle peut être utilisée dans différents domaines tels que la gestion des ressources radio (RRM), la sécurité du réseau, la gestion de la mobilité, l'amélioration de la qualité de service (QoS), la gestion de l'énergie, l'automatisation du réseau et l'analyse prédictive par exemple.
- L'IA est employée dans le domaine de la gestion des ressources radio afin d'optimiser l'utilisation des ressources radio, garantissant ainsi une utilisation optimale du spectre.

Concernant la protection du réseau, l'intelligence artificielle est utilisée afin de repérer et prévenir les éventuelles attaques, assurant ainsi l'intégrité et la confidentialité des données.

- L'IA joue également un rôle positif dans la gestion de la mobilité en prédisant les mouvements des utilisateurs et en optimisant les stratégies de transfert entre les stations de base. En outre, la qualité de service est améliorée en utilisant des algorithmes d'apprentissage automatique qui anticipent et ajustent la QoS pour divers services et applications.
- La gestion efficace de l'énergie dans les réseaux 5G est également favorisée par l'intelligence artificielle, ce qui permet de diminuer les dépenses opérationnelles et l'empreinte écologique. Elle facilite également l'automatisation des opérations réseau, ce qui améliore l'efficacité opérationnelle globale.
- L'IA est utilisée dans l'analyse prédictive afin de prédire les performances du réseau, prévoir les pannes et mettre en place des mesures de maintenance préventive, garantissant ainsi une disponibilité constante des services.

En résumé, l'intégration de l'intelligence artificielle dans la 5G offre une gamme diversifiée de cas d'utilisation qui améliorent la performance, la sécurité et l'efficacité opérationnelle des réseaux de communication sans fil.

III.6.Conclusion :

Dans ce chapitre, on a exploré en détail le domaine de ML, en mettant l'accent sur DL et les réseaux de neurones. On a abordé les bases de ML, les différences entre l'apprentissage supervisé, non supervisé et par renforcement, ainsi que les principaux algorithmes et techniques utilisés. Ensuite, on a examiné DL, une branche de ML utilisant des réseaux de neurones profonds pour construire des modèles à partir de données. On a également examiné les avantages de l'AI et son intégration avec la technologie 5G pour améliorer les performances des réseaux de communication sans fil.

Dans le prochain chapitre, on présentera le système SCMA proposé avec une détection basée sur DL et les résultats des simulations.

CHAPITRE IV :

*Simulation et interprétation
des résultats.*

IV.1.Introduction :

Après avoir posé les bases théoriques dans les chapitres précédents, ce chapitre marque une étape cruciale dans la réalisation de notre recherche. Nous nous concentrons sur la traduction de concepts abstraits en réalité tangible en mettant en œuvre le modèle proposé dans un environnement informatique concret.

La première partie de ce chapitre se concentre sur la mise en œuvre du modèle. Nous examinons les différentes décisions et choix de mise en œuvre, mettant en lumière l'environnement de développement utilisé, les bibliothèques de logiciels sélectionnées. La deuxième partie est consacrée à l'évaluation pratique du modèle mis en œuvre.

IV.2.Objectifs :

Ce chapitre propose une approche innovante pour faire face aux problèmes de complexité et de convergence lente rencontrés par le MPA dans les systèmes SCMA. Au lieu de faire appel à l'algorithme MPA traditionnel, nous allons présenter une alternative qui repose sur les DNNs, ce qui supprime la nécessité d'utiliser cet algorithme. Les objectifs de notre projet sont les suivants :

- Réduire la complexité de la détection multi-utilisateur dans les systèmes SCMA.
- Utiliser des détecteurs en utilisant des DNNs individuels pour chaque utilisateur, contrairement à MPA qui nécessite une connexion partagée, afin de simplifier le système et d'améliorer l'efficacité.
- Renforcer la sécurité du système en entraînant les détecteurs DNN uniquement sur le signal de l'utilisateur concerné, offrant ainsi une meilleure protection par rapport à MPA.

IV.3. Choix des outils de développement :

Le **Tableau IV.1** résume les outils de développement que nous avons utilisés dans notre travail, en précisant les versions lorsque cela est possible. Ces outils ont été essentiels pour le développement, l'entraînement et l'évaluation de modèles de DL, ainsi que pour l'analyse des données et la visualisation des résultats.

Outil de Développement	Logo	Version	Source	Description
Visual Studio Code (VS Code)		1.88	[1]	Un éditeur de code source léger et puissant développé par Microsoft, largement utilisé pour le développement de projets Python et de deep learning grâce à son support étendu pour les langages de programmation et ses nombreuses extensions.
Python		3.9	[2]	Un langage de programmation de haut niveau et polyvalent, largement utilisé dans le domaine du DL pour son écosystème riche en bibliothèques et sa facilité d'apprentissage et d'utilisation.
Tensorflow		2.12.0	[3]	Une bibliothèque open-source développée par Google, largement utilisée pour le développement de modèles de DL. TensorFlow offre une grande flexibilité et des performances élevées
Keras		2.12.0	[4]	Une API de haut niveau construite au-dessus de TensorFlow, qui permet de créer et d'entraîner des modèles de DL de manière simple et intuitive. Keras est souvent utilisé pour sa facilité d'utilisation.
Scikit-learn		1.3.0	[5]	Une bibliothèque open-source pour l'apprentissage automatique en Python, qui propose une variété d'outils pour la préparation des données, le développement de modèles et l'évaluation des performances.
NumPy		1.26.3	[6]	Une bibliothèque fondamentale pour le calcul numérique en Python, qui fournit des structures de données et des fonctions pour travailler efficacement avec des tableaux multidimensionnels.
Matplotlib		3.8.0	[7]	Une bibliothèque pour la visualisation de données en Python, qui permettent de créer des graphiques et des diagrammes pour explorer et présenter les résultats des expériences de DL.

Tableau IV.1 : Les outils de développement utilisés pour la simulation.

IV.4.Scénario proposé :

Le schéma fourni dans la **Figure IV.1** représente un système pour encoder, transmettre et décoder plusieurs flux de bits à travers un canal bruyant en utilisant des réseaux neuronaux profonds (DNN). Dans ce système, nous disposons de données provenant de 6 utilisateurs réparties sur 4 ressources disponibles pour la transmission. Chaque utilisateur est configuré pour envoyer des données simultanément sur deux ressources, de plus, chaque ressource est partagée entre 3 utilisateurs. Le canal AWGN est utilisé pour simuler les conditions réelles de transmission, ce qui nous permet d'évaluer les performances du système en tenant compte des effets du bruit sur la qualité de la transmission. Après le passage à travers le canal AWGN, les REs bruités sont séparés en leurs composantes réelles et imaginaires. Ces composants sont désignés comme RE1 (Réal, Imag), RE2 (Réal, Imag), RE3 (Réal, Imag) et RE4 (Réal, Imag).

Dans la partie réception de notre système SCMA proposé, nous avons remplacé le MPA par le DNN où les signaux reçus sont traités pour estimer les symboles transmis par chaque utilisateur.

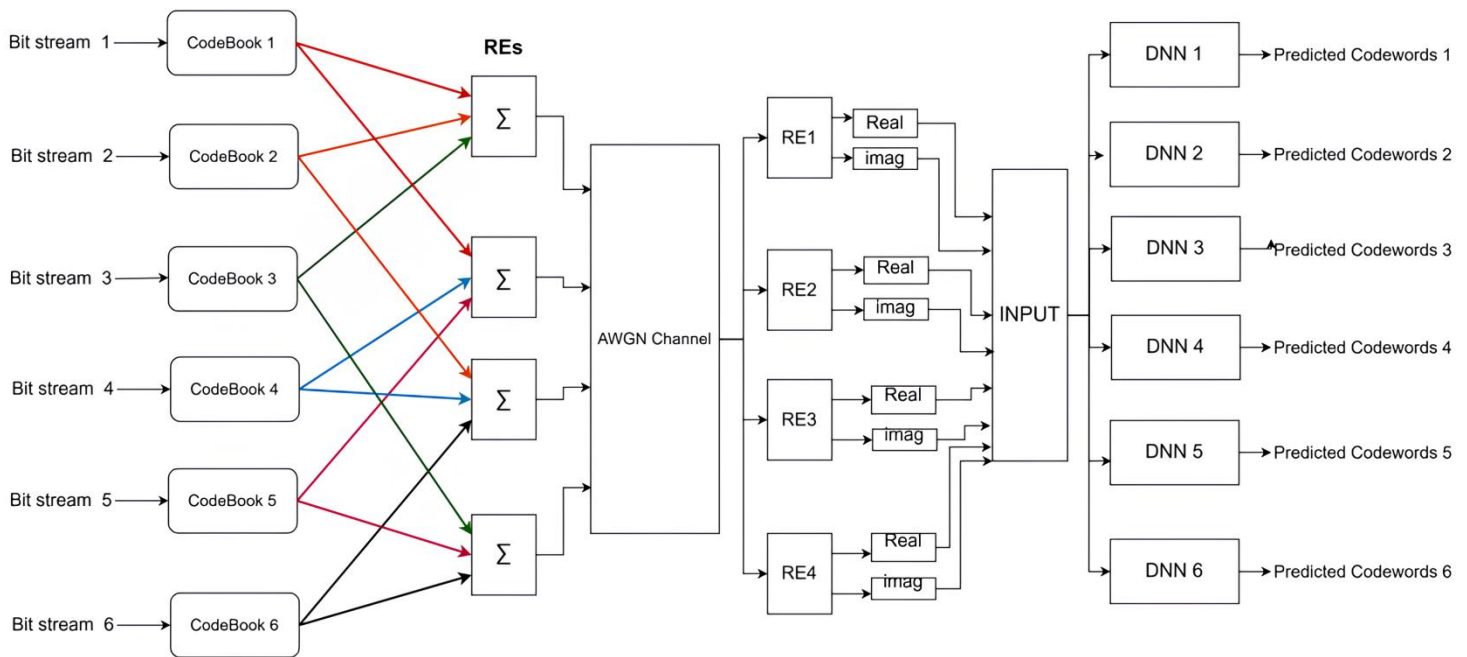


Figure IV.1 : Système SCMA proposé basé sur la détection par DNNs.

Les figures ci-dessus présentent différents diagrammes de constellations en variant les opérateurs de phase Δ en utilisant respectivement la modulation BPSK et QPSK.

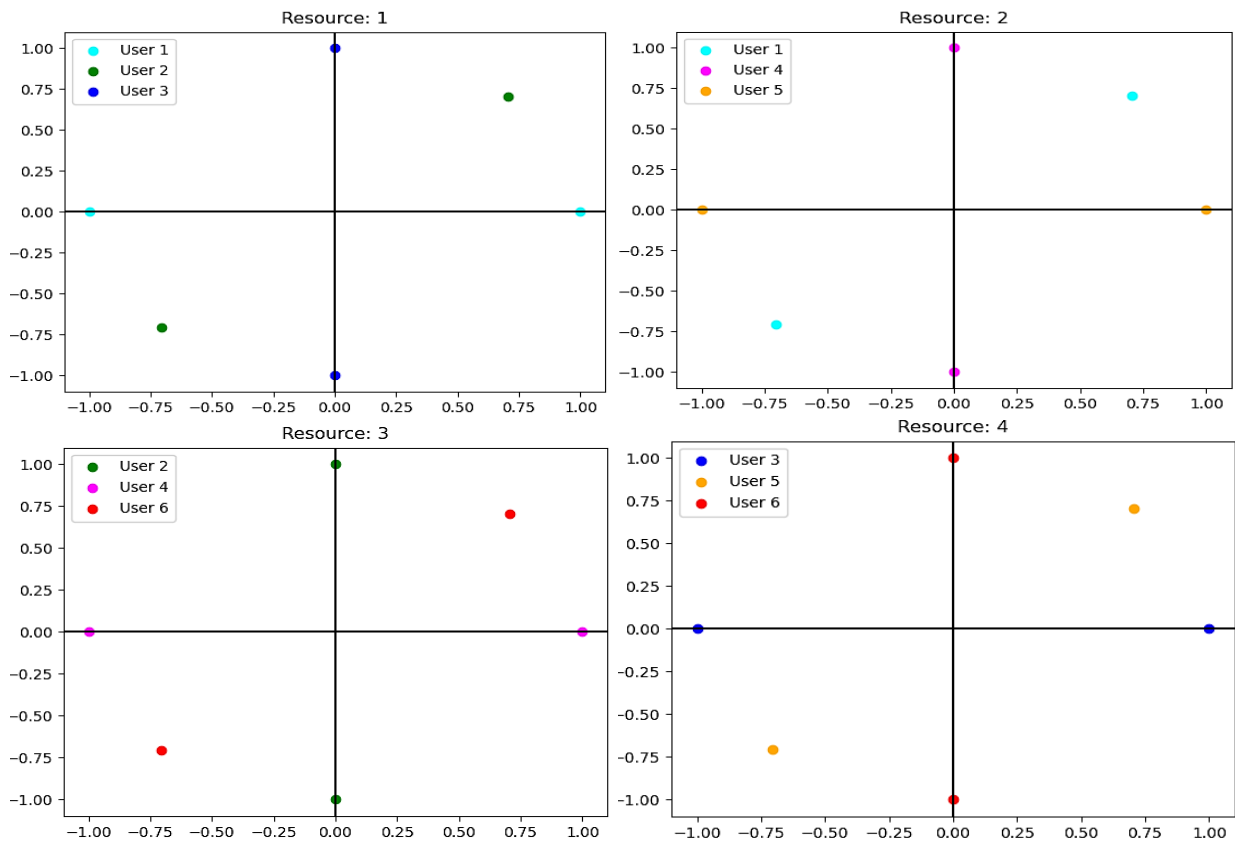


Figure IV.2 : Constellations des 4 REs avec $\Delta = \pi/4$ lors de l'utilisation de la modulation BPSK.

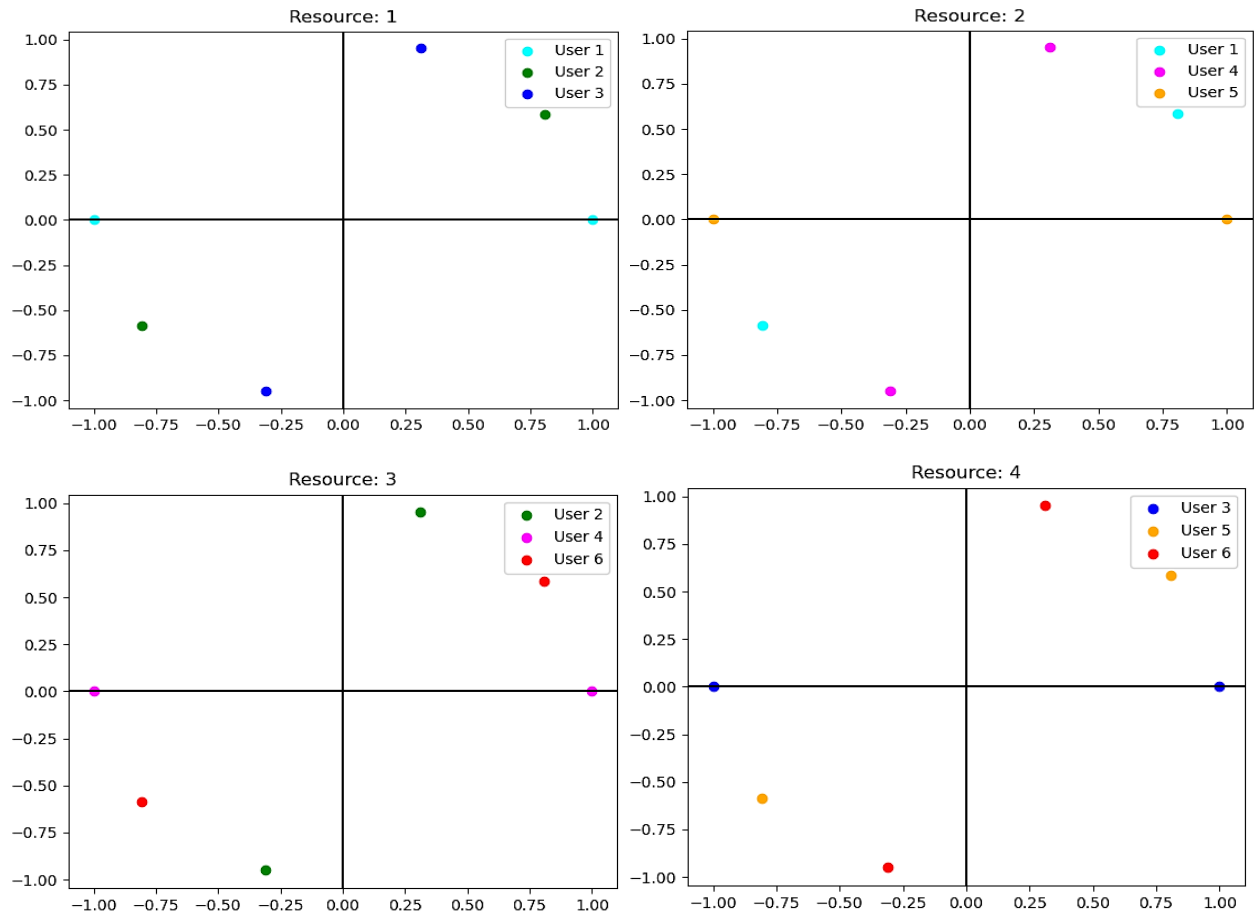


Figure IV.3 : Constellations des 4 REs avec $\Delta = \pi/5$ lors de l'utilisation de la modulation BPSK.

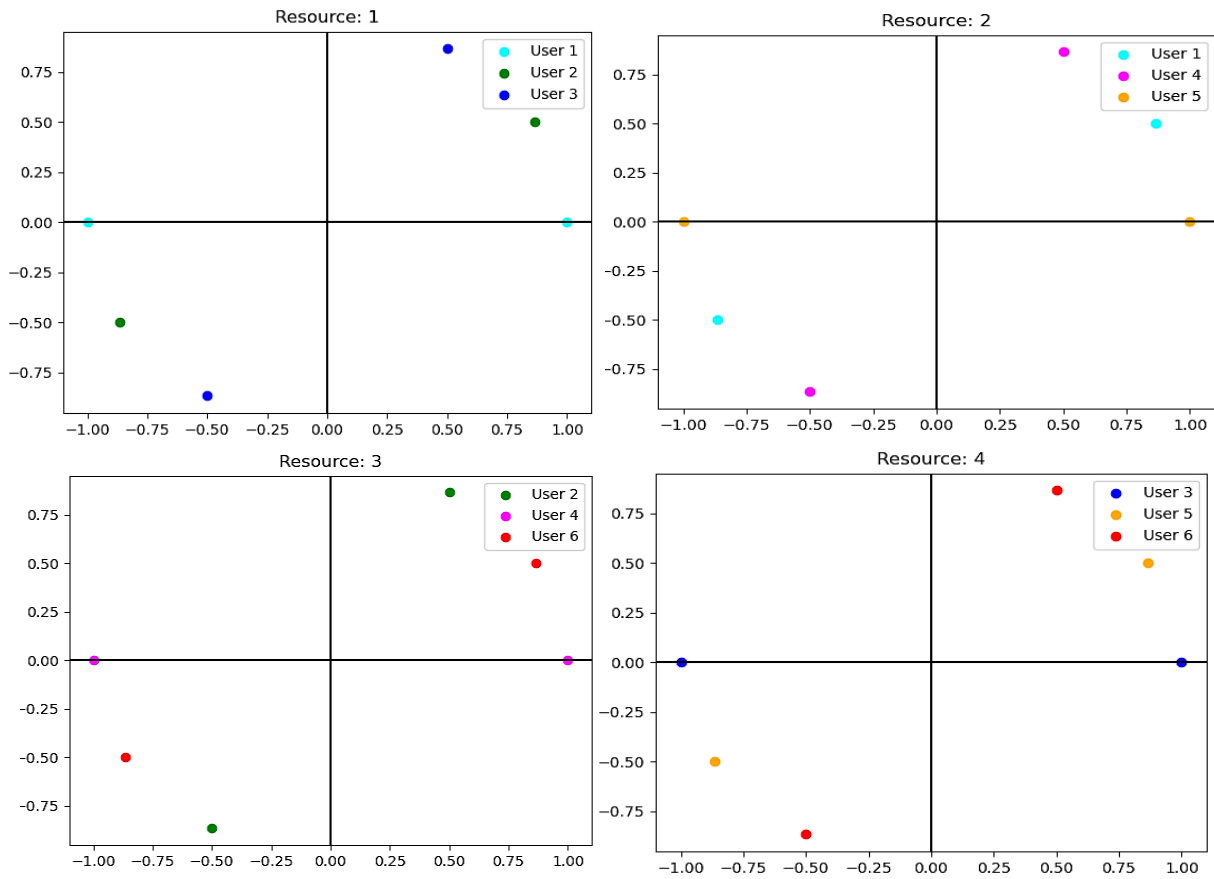


Figure IV.4 : Constellations des 4 REs avec $\Delta=\pi/6$ lors de l'utilisation de la modulation BPSK.

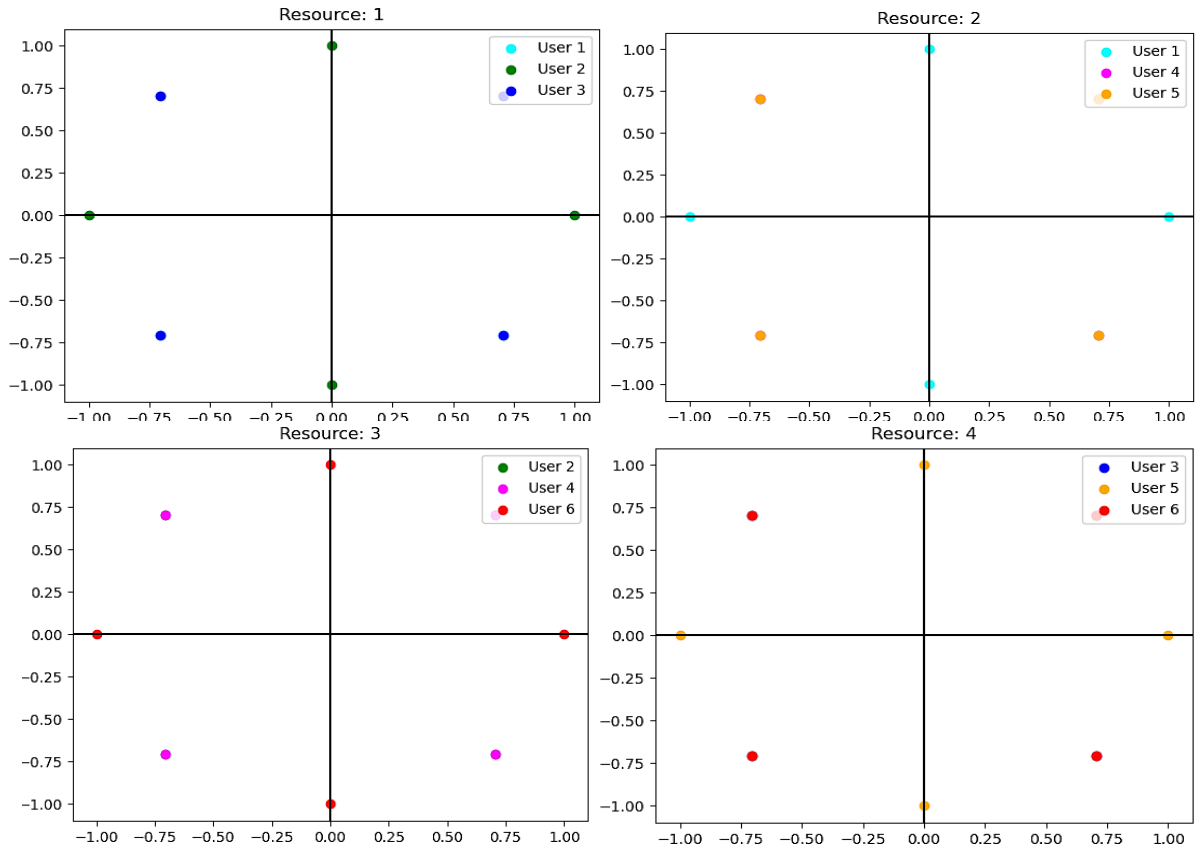


Figure IV.5 : Constellations des 4 REs avec $\Delta=\pi/4$ lors de l'utilisation de la modulation QPSK.

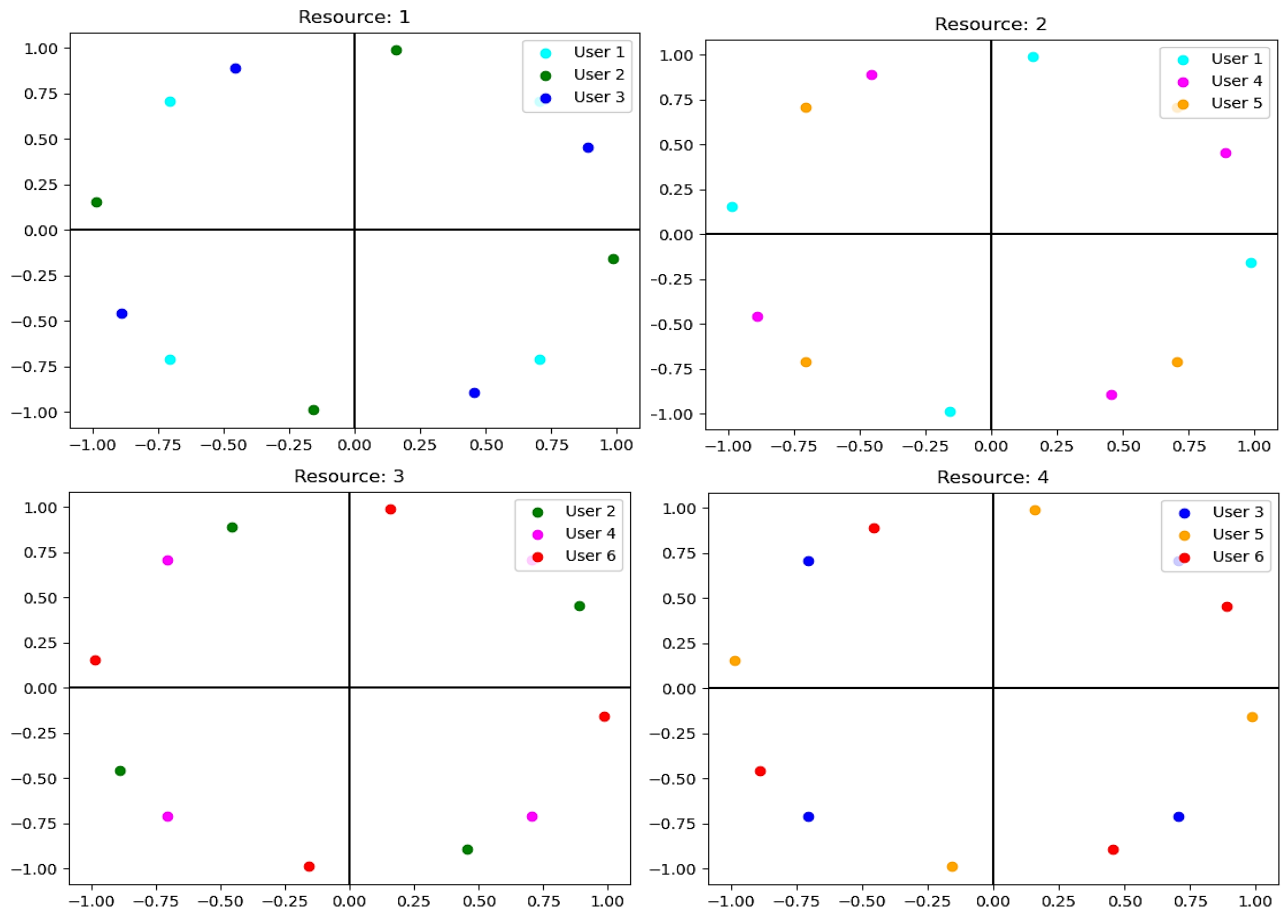


Figure IV.6 : Constellations des 4 REs avec $\Delta=\pi/5$ lors de l'utilisation de la modulation QPSK.

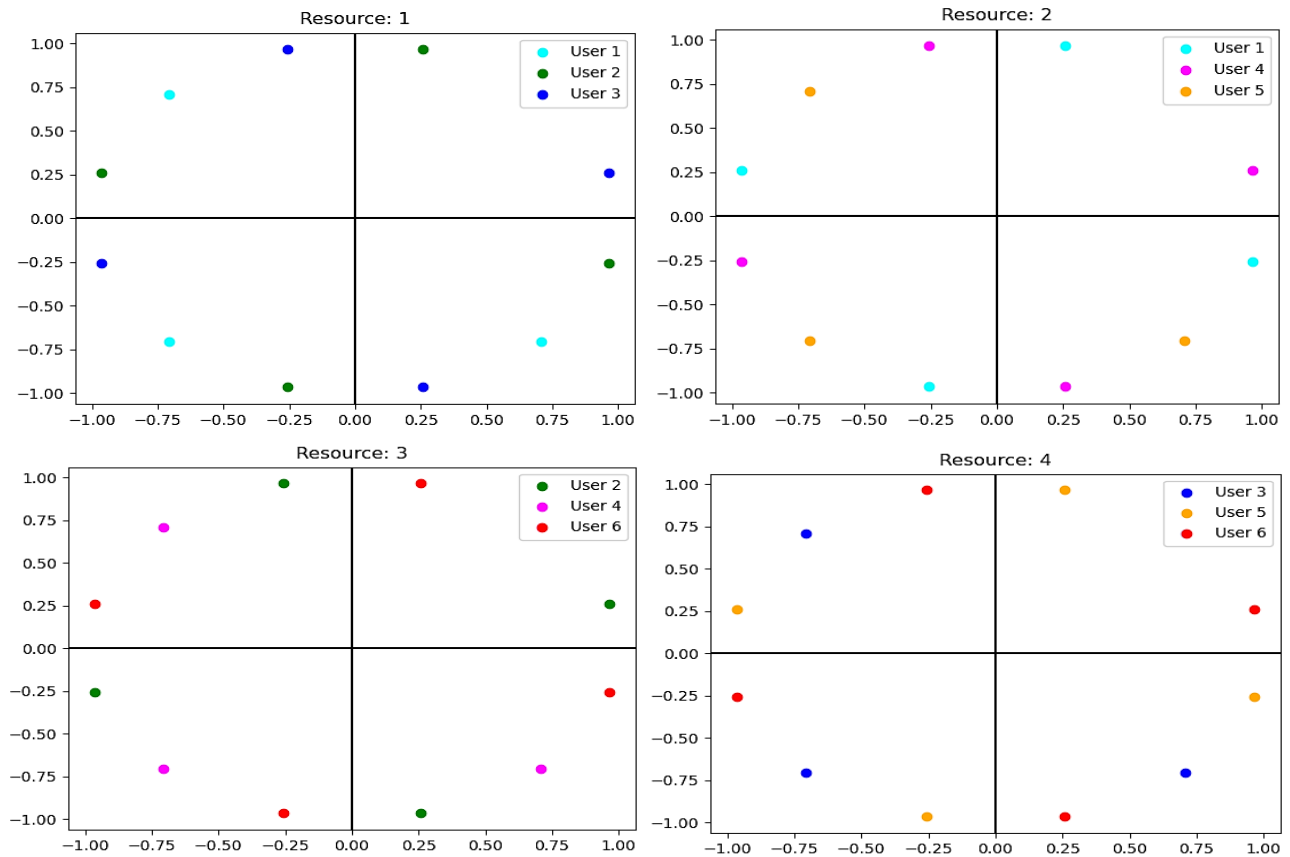


Figure IV.7 : Constellations des 4 REs avec $\Delta=\pi/6$ lors de l'utilisation de la modulation QPSK.

IV.5.Génération de données d'entraînement :

Dans le cadre de système SCMA proposé, nous avons constitué un ensemble de données pour l'entraînement de modèle. Cet ensemble de données est conçu pour refléter fidèlement les conditions réelles du système, en tenant compte de paramètres qui constituent la base de notre modèle de système. Ces paramètres sont représentés dans le **Tableau IV.2**:

Paramètre	Notation	Valeur
Nombre de symboles transmis	N	50000
Nombre d'utilisateurs	J	6
Nombre de ressources	K	4
Nombre de ressources où l'utilisateur transmet	d_v	2
Nombre d'utilisateurs impliqués dans RE	d_f	3
Indice de modulation	M	2 et 4
L'opérateur de rotation	Δ	$\pi/4, \pi/5, \pi/6$
Rapport signal sur bruit	SNR	$\in [-4, 13]$ dB

Tableau IV.2: Les données de base de système SCMA proposé.

L'objectif principal de la création de cet ensemble de données est de permettre au modèle de s'entraîner à détecter et à décoder les signaux transmis dans le système SCMA dans diverses conditions. Chaque exemple d'entraînement dans notre ensemble de données est étiqueté en fonction des signaux réellement transmis, permettant ainsi au modèle de s'entraîner de manière supervisée.

Pour ce faire, on a généré dynamiquement des données d'entraînement qui représentent une variété de scénarios possibles, en prenant en compte la diversité des signaux, des utilisateurs et des ressources. De plus, pour capturer les variations et les défis rencontrés dans un environnement de communication réel, nous avons inclus l'AWGN dans les données d'entraînement. Pour chaque valeur de SNR, nous avons généré un ensemble de données correspondant, garantissant ainsi une couverture exhaustive des conditions de canal possibles. Cela permet au modèle de s'adapter et de généraliser efficacement, quel que soit le niveau de bruit auquel il est confronté lors de son déploiement dans un environnement réel.

En règle générale, les RNA classiques sont développés pour gérer des valeurs réelles plutôt que des valeurs complexes. Chacun des vecteurs d'entrée est donc élaboré de manière à intégrer à la fois la partie réelle et la partie imaginaire des symboles transmis sur chaque ressource du système SCMA pour représenter les signaux dans leur ensemble. De cette manière, chaque vecteur d'entrée possède une dimension en $2 \times K$, où K représente le nombre de ressources disponibles dans le

système. Ainsi, la spécification de l'entrée du modèle est un vecteur de données qui représente les symboles transmis à chaque ressource du système SCMA. La partie réelle et la partie imaginaire des symboles sont intégrées dans ce vecteur pour chaque ressource, ce qui permet de capturer les caractéristiques complètes du signal.

IV.6.Choix des critères :

Lorsqu'on évalue un modèle de DL, il est primordial de considérer divers critères afin de déterminer son efficacité et sa pertinence par rapport aux objectifs prévus dans le projet. Plusieurs critères essentiels ont été identifiés lors de cette évaluation :

- **La précision (Accuracy)** : consiste à évaluer la proportion de prédictions correctes parmi toutes les prédictions réalisées par le modèle.
- **La Perte (Loss)** : analyse de l'erreur du modèle, afin de déterminer à quel point les prédictions du modèle diffèrent des valeurs réelles.
- **Le SER (Symbol Error Rate)** : Pour évaluer l'efficacité d'un modèle de DL pour la détection de symboles dans les communications numériques, le SER peut être un critère essentiel. Un SER faible témoigne de la capacité du modèle à détecter avec précision les symboles transmis, ce qui est essentiel pour garantir une communication fiable et sans erreur.

IV.7.Architecture de modèle :

Dans cet approche, nous avons proposé deux architectures distinctes pour chaque scénario de modulation L'utilisation de différentes modulations, telles que BPSK et QPSK, peut avoir un impact significatif sur la complexité et les caractéristiques des signaux transmis dans un système de communication.

Dans un premier temps, nous avons préparé le modèle en utilisant les bibliothèques mentionnées précédemment. À cette fin, on a effectué une normalisation des données d'entraînement générées, qui sont les entrées de notre modèle, dans un intervalle de $[0,1]$ en utilisant une fonction spécifique. Cette normalisation vise à obtenir la même plage de valeurs pour chaque entrée.

Ensuite, nous avons divisé les données en ensembles d'entraînement, de validation et de test. Et aussi nous avons établi les couches illustrées dans les tableaux **Tableau IV.3** et **Tableau IV.4**, en mesurant le nombre de neurones présents dans chaque couche. L'utilisation de la fonction d'entropie croisée catégorique, aussi appelée « categorical_crossentropy », a été choisie comme fonction de

perte pour nos modèles. Nous avons opté pour l'utilisation de l'algorithme Adamax pour l'optimisation.

Les autres hyper paramètres spécifiques à chaque situation sont présentés ci-dessous, tels que le nombre de couches dissimulées, le taux d'apprentissage, le nombre d'époques d'entraînement, etc., répertoriés dans des tableaux séparés afin de faciliter la compréhension.

1. Modulation BPSK (M=2) :

Le tableau suivant récapitule les choix des hyperparamètres dans le cas où la modulation BPSK :

Couche	Type de Couche	Nombre d'unités	Fonction d'activation	Type de régularisation	Taux d'apprentissage	Taille de lot	Nombre d'époques
1	D'entrée	128	ReLu	/	$a = 10^{-3}$	$BS = 16$	40
2	Dense	64	ReLu	/			
3	Dense	64	ReLu	L2			
4	Dense	32	ReLu	L2			
5	Dense	16	ReLu	/			
6	Dense	8	ReLu	/			
7	Sortie	2	Softmax	/			

Tableau IV.3 : Choix des hyper paramètres dans le cas où M=2.

2. Modulation QPSK (M=4) :

Dans le **Tableau IV.3**, on trouve les hyperparamètres sélectionnés pour le cas où la modulation QPSK :

Couche	Type de Couche	Nombre d'unités	Fonction d'activation	Type de régularisation	Taux d'apprentissage	Taille du lot	Nombre d'époques
1	D'entrée	256	ReLu	/	$a = 5.10^{-4}$	$BS = 16$	40
2	Dense	128	ReLu	/			
3	Dense	128	ReLu	/			
4	Dense	128	ReLu	/			
5	Dense	64	ReLu	/			
6	Sortie	4	Softmax	/			

Tableau IV.4 : Choix des hyper paramètres dans le cas où M=4.

Dans les deux tableaux, le terme "époque" fait référence à une passe complète à travers l'ensemble de données d'entraînement lors de l'entraînement d'un modèle. Ainsi La "taille de lot" (batch size) fait référence au nombre d'échantillons de données d'entraînement utilisés dans une seule itération pour mettre à jour les poids du modèle.

Il est possible que les hyperparamètres d'un modèle de DNN varient en fonction de $M=2$ ou $M=4$ en raison de la complexité du problème, de la dimensionnalité de l'espace de sortie, des caractéristiques spécifiques à chaque modulation et de la capacité de généralisation nécessaire. Toutes les configurations de modulation peuvent avoir besoin de modifications particulières afin d'améliorer les performances du modèle en fonction de ses exigences. Ces hyperparamètres doivent être choisis avec soin pour obtenir un modèle de DNN performant et bien généralisé.

En simulation, les données sont divisées en 2 sous-ensembles (entraînement et validation). Les performances du modèle sont testées directement dans la détection SCMA. Une fois entraîné, le modèle est capable de détecter les symboles transmis dans un environnement de communication avec différents niveaux de bruit et de modulation.

Ensemble de données	M=2	M=4
Entraînement	80%	90%
Validation	20%	10%

Tableau IV.5 : La répartition des ensembles de données pour les deux modèles.

Une fois entraîné et validé sur les données simulées, le modèle a été testé dans un environnement réel de SCMA, démontrant sa capacité à détecter les symboles transmis dans des conditions réelles de communication, avec différents niveaux de bruit et de modulation.

IV.8.Résultats expérimentale et discussion :

Dans la première partie des résultats, nous exposons les métriques d'entraînement du modèle, tels que l'exactitude et la perte, en fonction du nombre d'époques. Grâce à ces mesures, nous pouvons évaluer la convergence et les performances du modèle au fil de l'entraînement. En analysant l'évolution de l'exactitude et du déficit à travers les époques à des niveaux différents de SNR, il est possible de déterminer si le modèle converge vers une solution optimale et si l'entraînement est réalisé de manière stable.

Ainsi, dans la deuxième partie nous étudierons l'évolution du SER pour nos modèles de DNNs et pour l'algorithme MPA à divers niveaux de SNR. En analysant les courbes de SER en fonction du SNR pour chaque méthode, nous serons en mesure de visualiser et de comparer leurs performances respectives dans des conditions de bruit différentes.

Enfin, on a réalisé des simulations avec des modèles de neurones simples, avec seulement les couches d'entrée et de sortie pour les cas $M=2$ et $M=4$, pour comparer leurs performances à celles des modèles de DNN. Dans les systèmes SCMA, cette méthode met en évidence les disparités entre les structures de modèle et évalue l'impact de la complexité du réseau de neurones sur la détection des symboles.

IV.8.1. Analyse de l'Évolution de l'exactitude et de la perte au fil des époques, avec considération de l'impact du SNR :

Dans le contexte d'un système, l'exactitude et la perte sont deux indicateurs essentiels qui déterminent les performances du système. En ce qui concerne le SNR, il joue un rôle essentiel dans la qualité de la transmission des informations. Il est possible d'analyser la relation entre ces trois éléments afin de mieux appréhender le fonctionnement et les limites du système. Ces illustrations présentent l'évolution de l'exactitude et de la perte pour les deux modèles proposés lors de l'entraînement et de la validation, en utilisant différents opérateurs de phase.

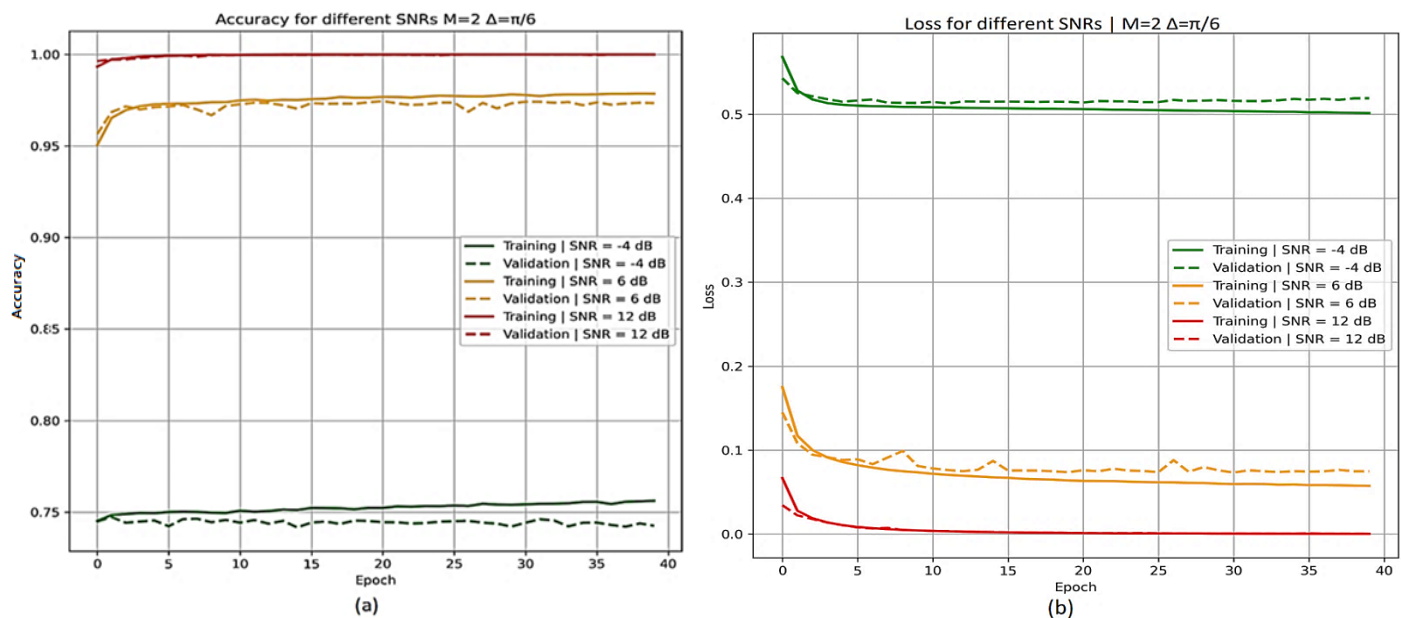


Figure IV.8: L'évolution de : **(a)** l'accuracy et **(b)** le loss au fil des époques du modèle dans l'entraînement et la validation pour $M=2$ et $\Delta=\pi/6$.

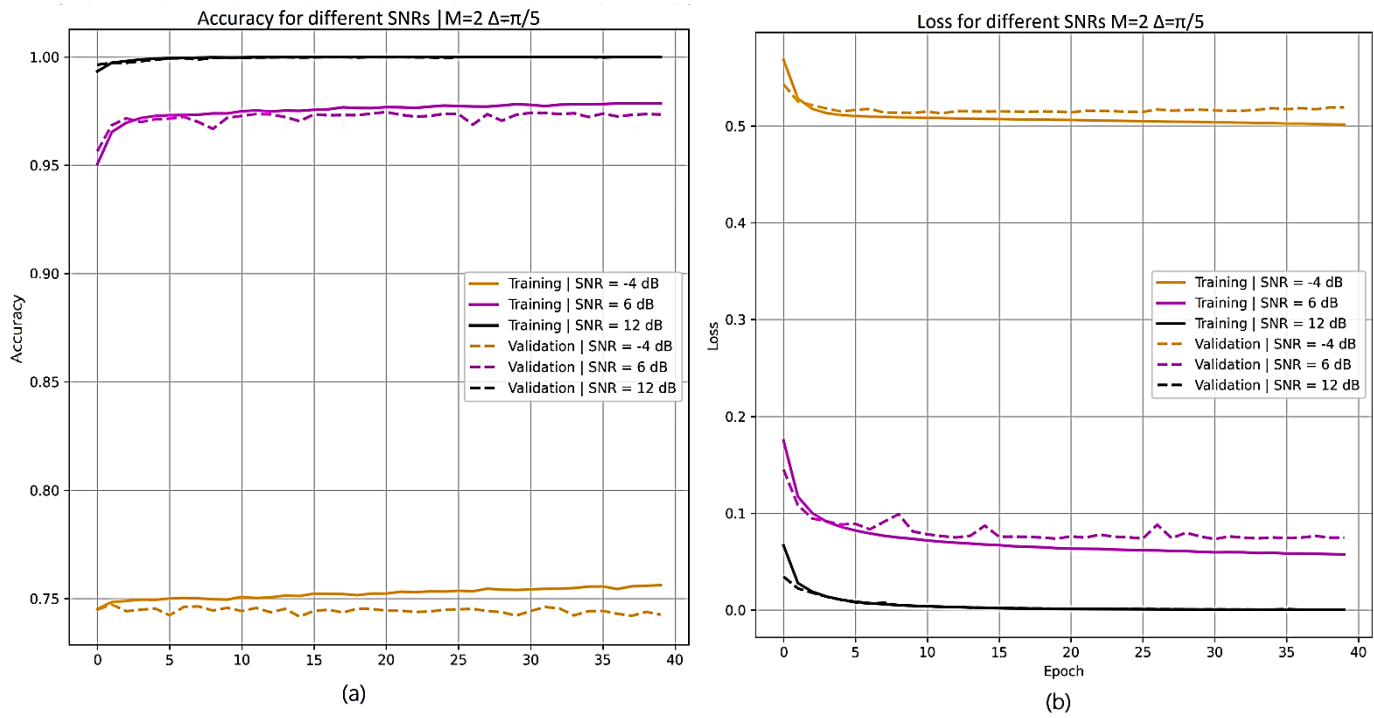


Figure IV.9 : L'évolution de : (a) l'accuracy et (b) le loss au fil des époques du modèle dans l'entraînement et la validation pour $M=2$ et $\Delta=\pi/5$.

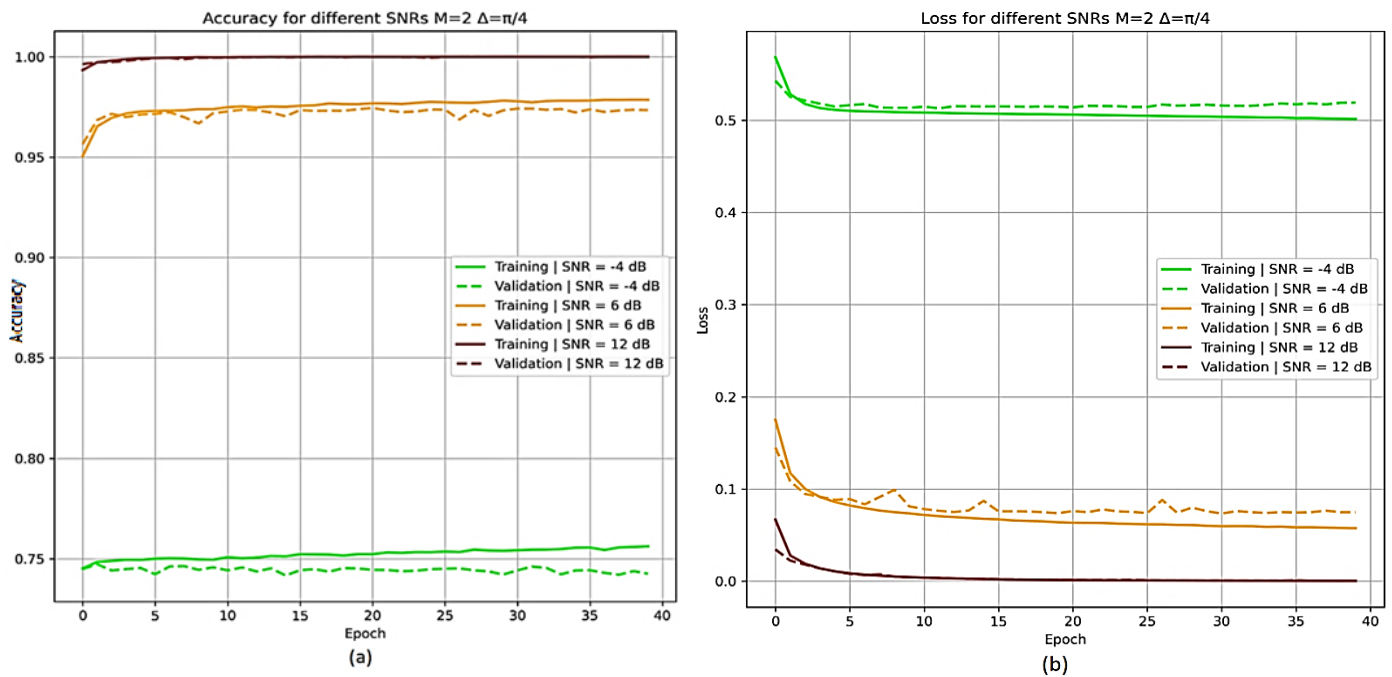


Figure IV.10 : L'évolution de : (a) l'accuracy et (b) le loss au fil des époques du modèle dans l'entraînement et la validation pour $M=2$ et $\Delta=\pi/4$.

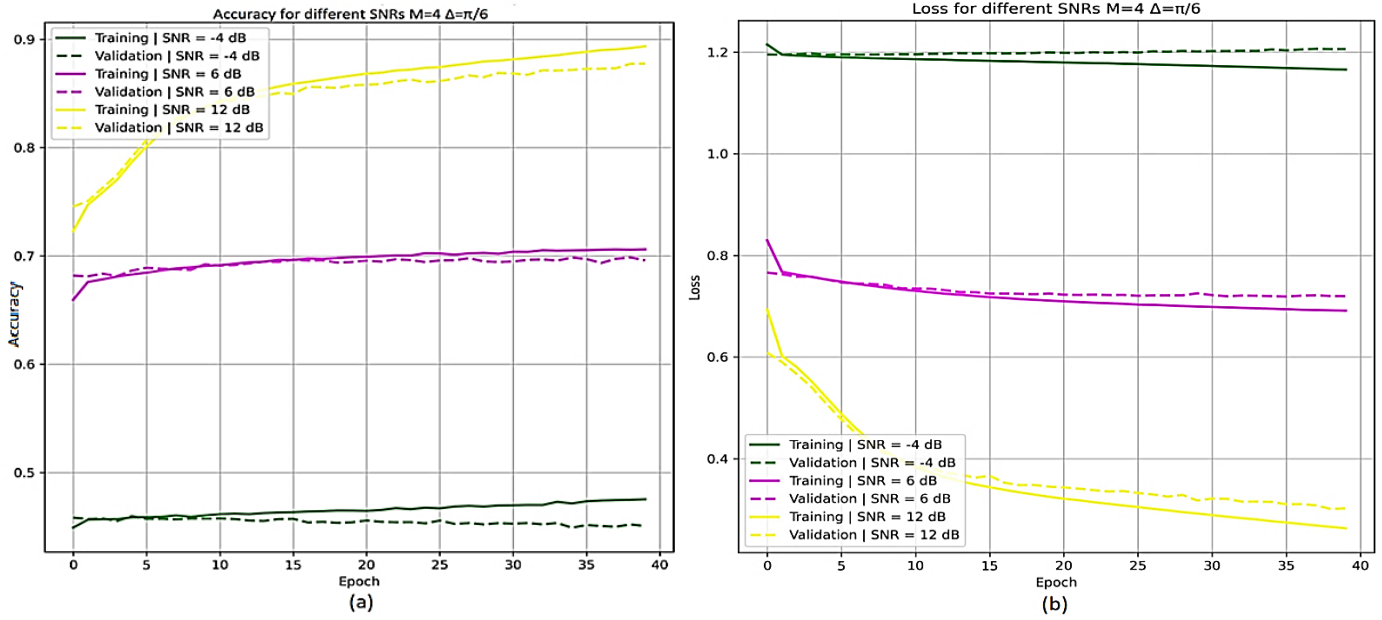


Figure IV.11 : L'évolution de l'accuracy (a) et le loss (b) au fil des époques du modèle dans l'entraînement et la validation pour $M=4$ et $\Delta=\pi/6$.

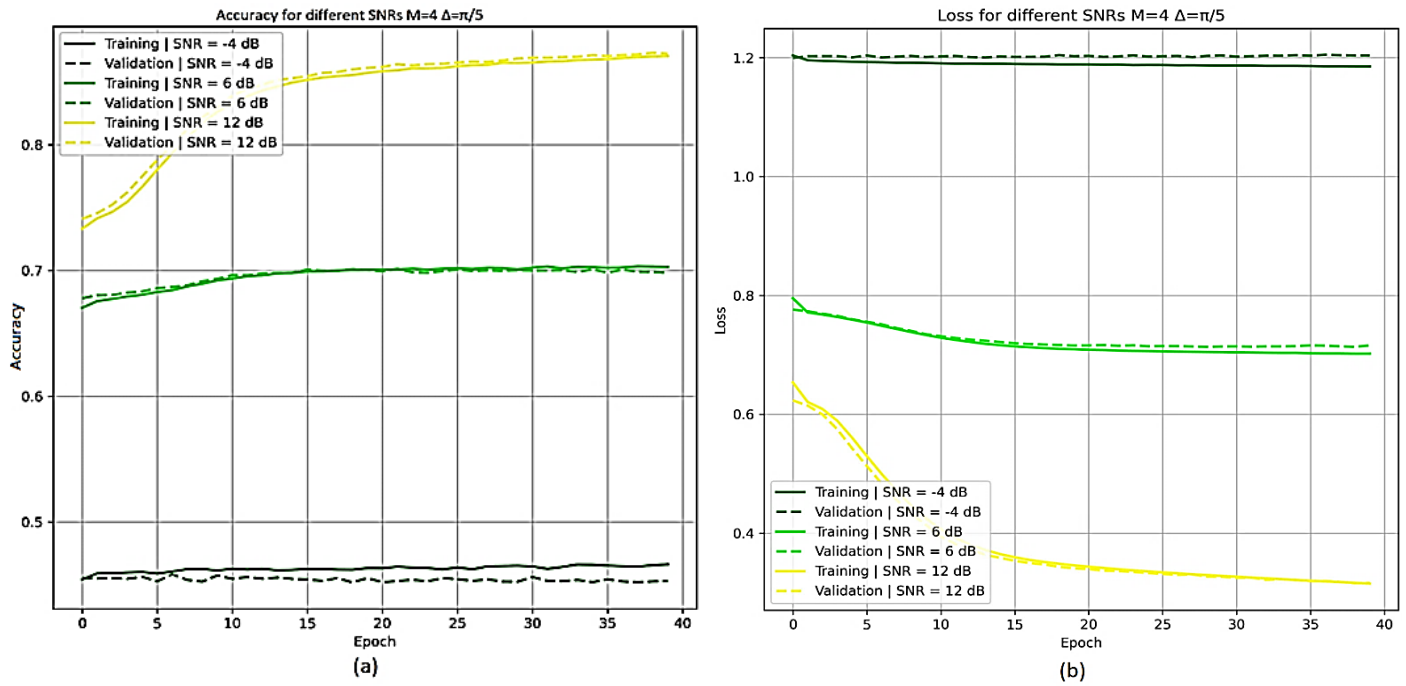


Figure IV.12 : L'évolution de : (a) l'accuracy et (b) le loss au fil des époques du modèle dans l'entraînement et la validation pour $M=4$ et $\Delta=\pi/5$.

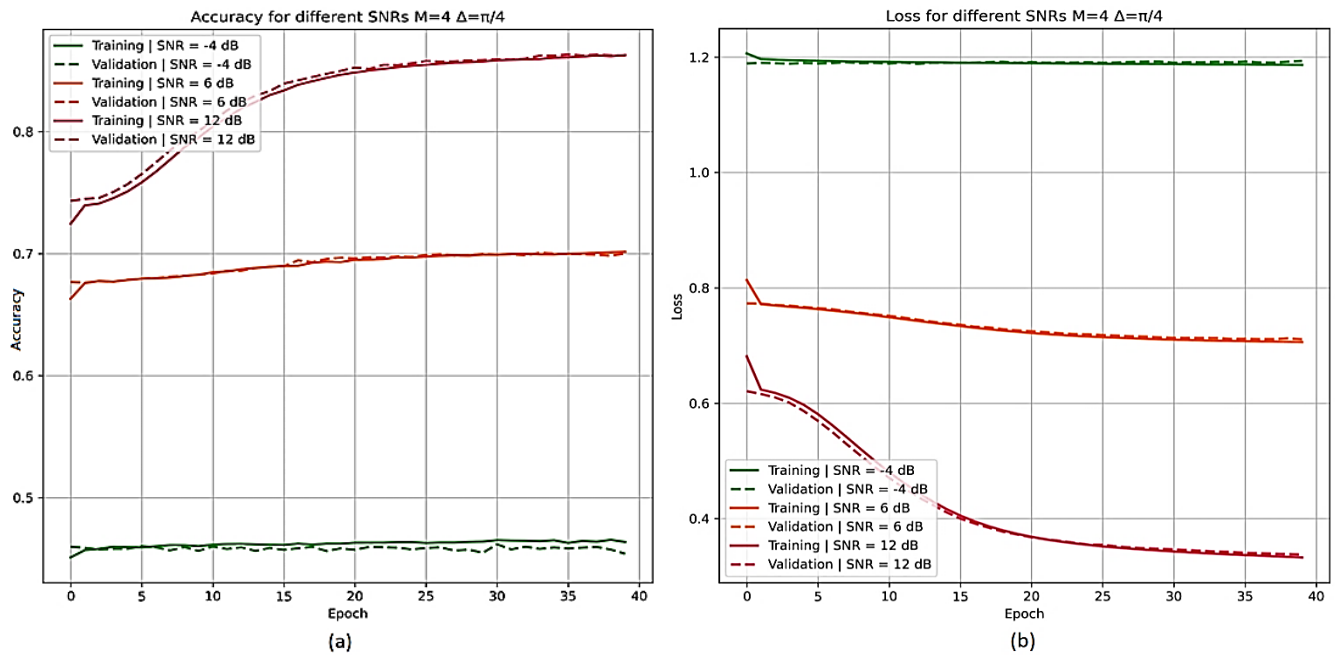


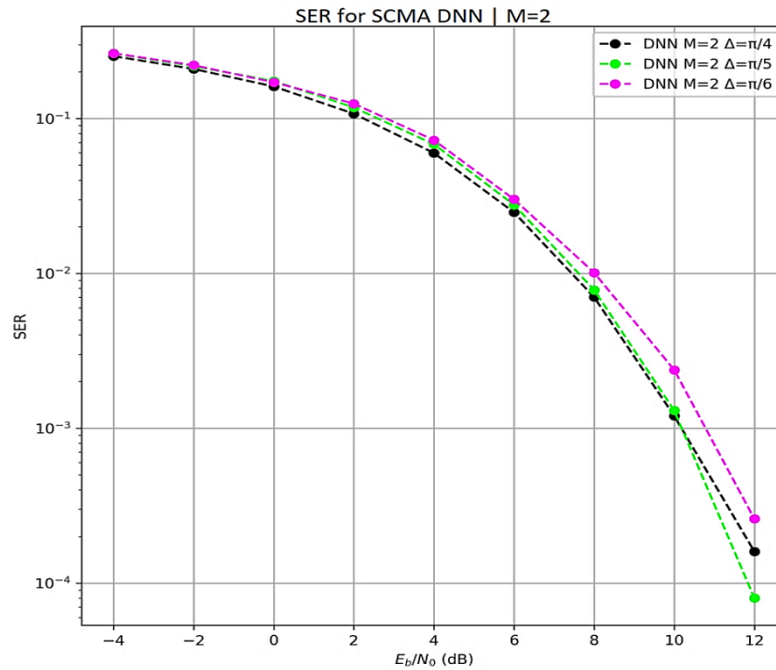
Figure IV.13 : L'évolution de : **(a)** l'accuracy et **(b)** le loss au fil des époques du modèle dans l'entraînement et la validation pour $M=4$ et $\Delta=\pi/4$.

En général, au début de l'entraînement, les pertes sont importantes et l'exactitude est faible car le modèle n'a pas encore acquis une bonne maîtrise de la généralisation des données. Chaque période supplémentaire entraîne une diminution du déficit et une augmentation de l'exactitude, ce qui suggère que le modèle s'améliore progressivement dans sa capacité à prédire de manière précise les données d'entraînement. Toutefois, il est essentiel de surveiller la précision de l'ensemble de validation ainsi que de l'ensemble d'entraînement. Si l'exactitude de l'ensemble de validation commence à baisser tandis que celle de l'ensemble d'entraînement continue d'augmenter, cela peut suggérer un surapprentissage du modèle, où il apprend trop précisément les données d'entraînement et ne généralise pas correctement à de nouvelles données.

Le SNR peut jouer un rôle essentiel dans le contexte de l'évolution de l'exactitude et de la perte au fil des époques dans un système SCMA. Les signaux sont plus évidents et les données sont plus facilement repérables à des niveaux élevés de SNR, ce qui peut entraîner une amélioration de l'exactitude et une diminution de la perte. D'autre part, lorsque le SNR est faible, le bruit peut dissimuler les signaux, ce qui peut entraîner une diminution de l'exactitude et une augmentation de la perte, car le modèle aura plus de difficulté à différencier les données utiles du bruit ambiant.

IV.8.2. Analyse de l'évolution de SER en fonction de SNR :

L'analyse de la variation du SER avec différentes valeurs de SNR permet de saisir l'impact du niveau de bruit dans le canal sur la qualité de la communication. Les **Figures IV.14** et **IV.15** illustrent l'évolution du SER en fonction du SNR pour les deux modèles proposés de décodage, mettant en évidence comment la performance de décodage varie en fonction du niveau de bruit.



Figures IV.14 : Évolution SER en fonction de SNR de modèle proposé (M=2) avec différents Δ .

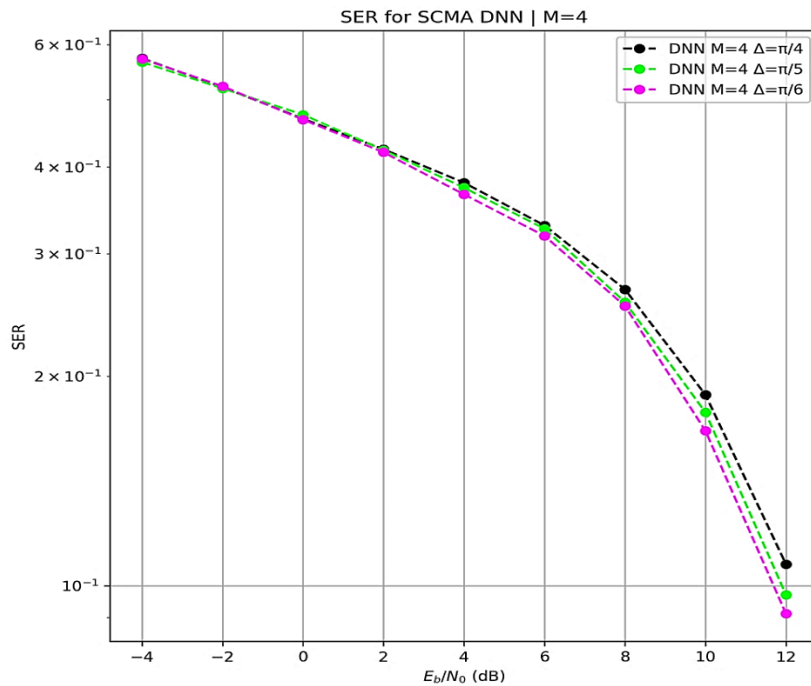


Figure IV.15 : Évolution SER en fonction de SNR de modèle proposé (M=4) avec différents Δ .

En analysant les deux figures, nous pouvons visualiser l'évolution de SER pour les deux modèles et avec différents Δ pour permet d'évaluer la sensibilité de chaque opérateur aux conditions du canal. Le SER diminue progressivement lorsque le SNR est élevé, ce qui témoigne d'une amélioration des performances du modèle dans des conditions de signal fort par rapport au bruit. Néanmoins, même à des niveaux de SNR plus bas, le SER demeure plutôt élevée, ce qui laisse entendre des difficultés persistantes dans le décodage des symboles dans des conditions de bruit élevé.

Selon nos analyses, la performance de détection évaluée par le SER fluctue en fonction de la valeur de Δ . On constate que $\Delta=\pi/5$ ($M=2$) et $\Delta=\pi/6$ ($M=4$) fournissent des performances de détection optimales. Ces valeurs de Δ semblent correspondre à des niveaux où le SER est la plus faible, ce qui suggère une meilleure capacité du système à détecter les symboles transmis avec précision. De ce fait, ces configurations particulières semblent être les plus adaptées pour optimiser les performances du système SCMA dans les conditions spécifiques. Cette observation indique que l'ajustement de Δ en fonction du nombre de bits modulés par symbole peut jouer un rôle significatif dans l'optimisation des performances du système SCMA.

IV.8. 3.Optimisation des Performances de Décodage par Exploration des Hyperparamètres :

Concernant le réglage des hyperparamètres et les performances des modèles de réseaux neuronaux, le taux d'apprentissage et la taille du lot ont été étudiés pour leur impact sur la convergence et l'efficacité des modèles. Des taux d'apprentissage trop élevés peuvent entraîner une convergence instable, tandis que des taux trop bas prolongent le temps d'apprentissage sans amélioration significative. De même, des tailles de lot plus importantes accélèrent l'apprentissage mais augmentent la consommation de mémoire, tandis que des tailles plus petites favorisent une généralisation efficace au détriment d'une convergence plus lente. Les figures correspondantes permettent d'analyser l'impact du Learning Rate et du Batch Size sur les performances de décodage.

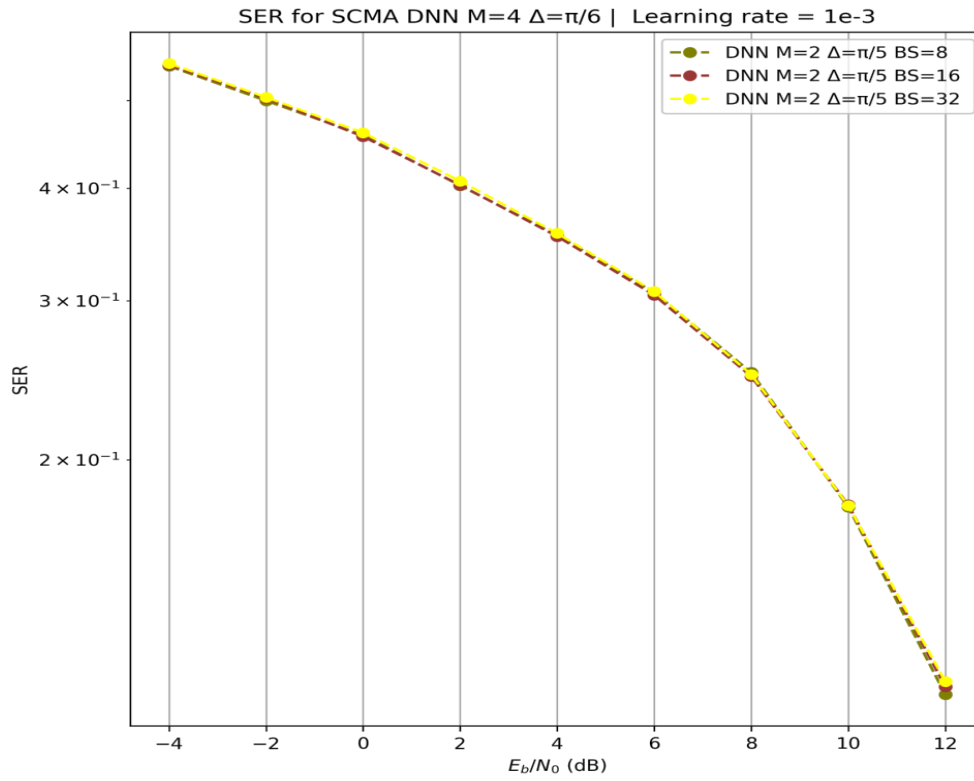


Figure IV.16 : Évolution des performances de décodage pour le modèle ($M=4$, $\Delta= \pi/6$) avec $\alpha = 10^{-3}$ et trois BS différents.

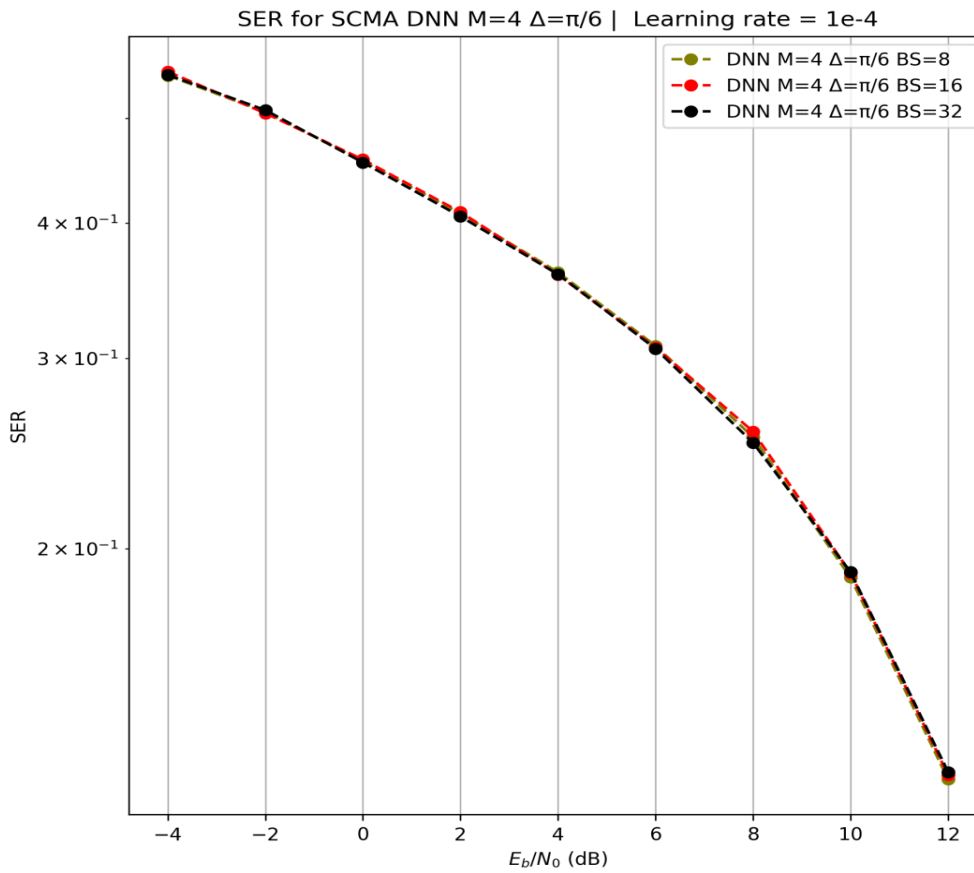


Figure IV.17: Évolution des performances de décodage pour le modèle ($M=4$, $\Delta= \pi/6$) avec $\alpha = 10^{-4}$ et trois BS différents.

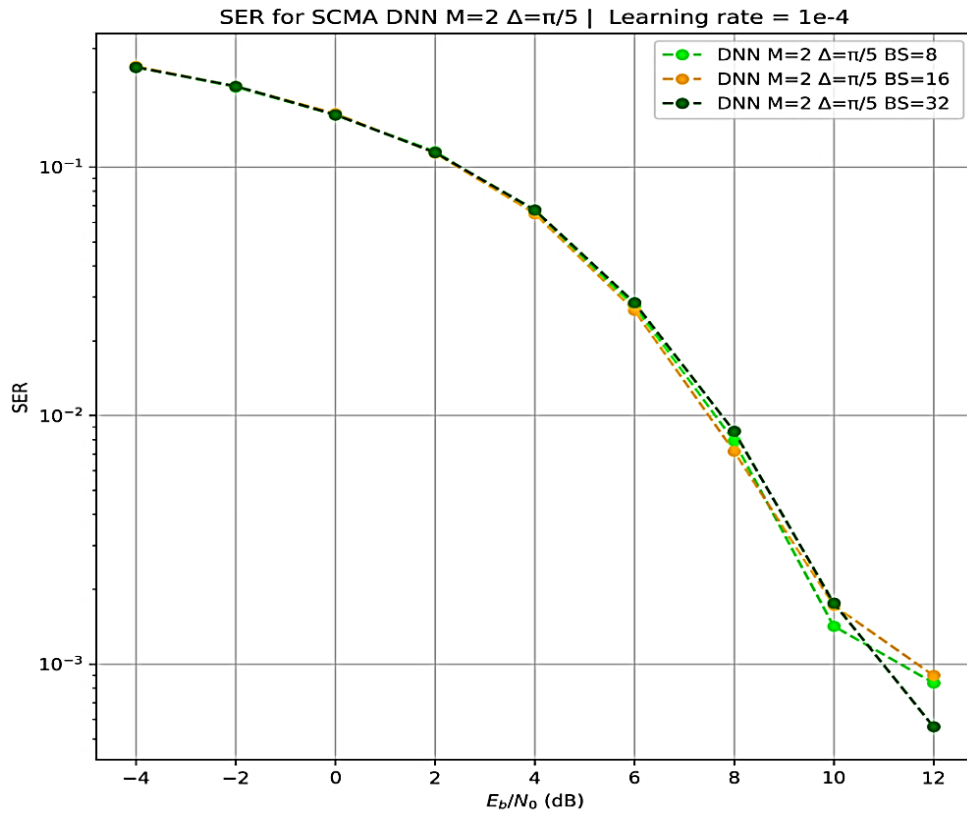


Figure IV.18 : Évolution des performances de décodage pour le modèle ($M=2$, $\Delta=\pi/5$) avec $a = 10^{-4}$ et trois BS différents.

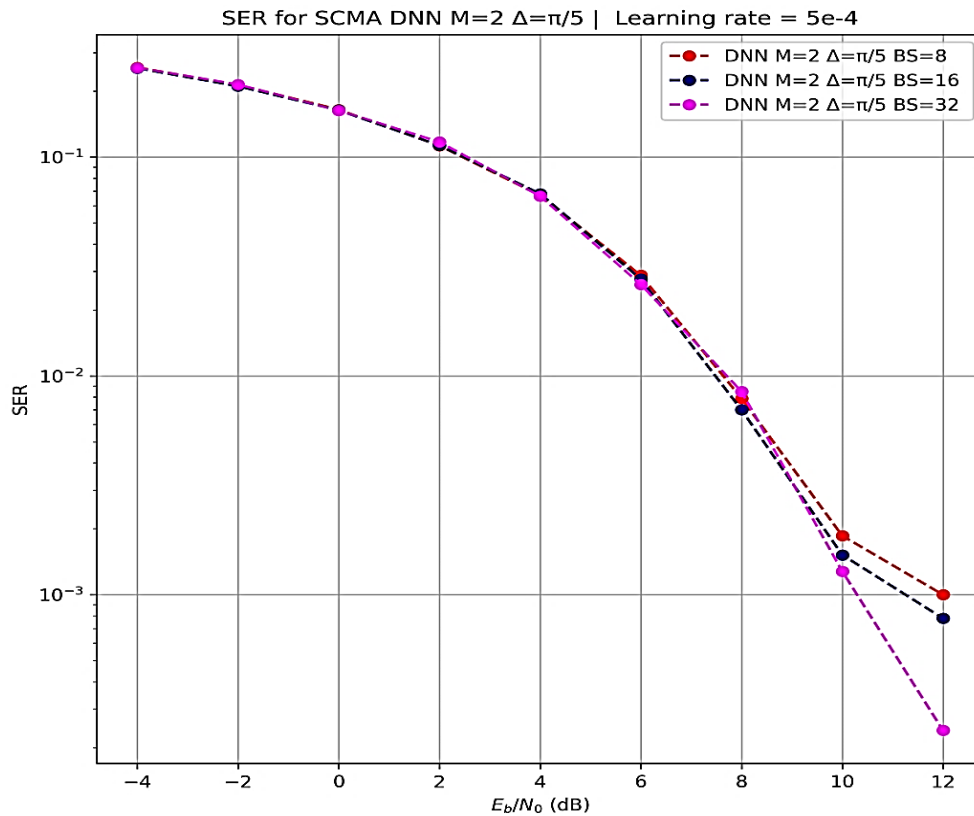


Figure IV.19: Évolution des performances de décodage pour le modèle ($M=2$, $\Delta=\pi/5$) avec $a = 5 \times 10^{-4}$ et trois BS différents.

IV.8.4. Analyse de l'évolution de SER par rapport à MPA :

Maintenant, nous allons analyser l'évolution du SER des deux modèles en le comparant avec MPA. Ces figures ci-dessous nous permettent d'évaluer et de comparer la capacité des modèles DNN et MPA à gérer les variations du SNR. Chaque courbe représentant l'algorithme MPA a une correspondante similaire pour le modèle DNN.

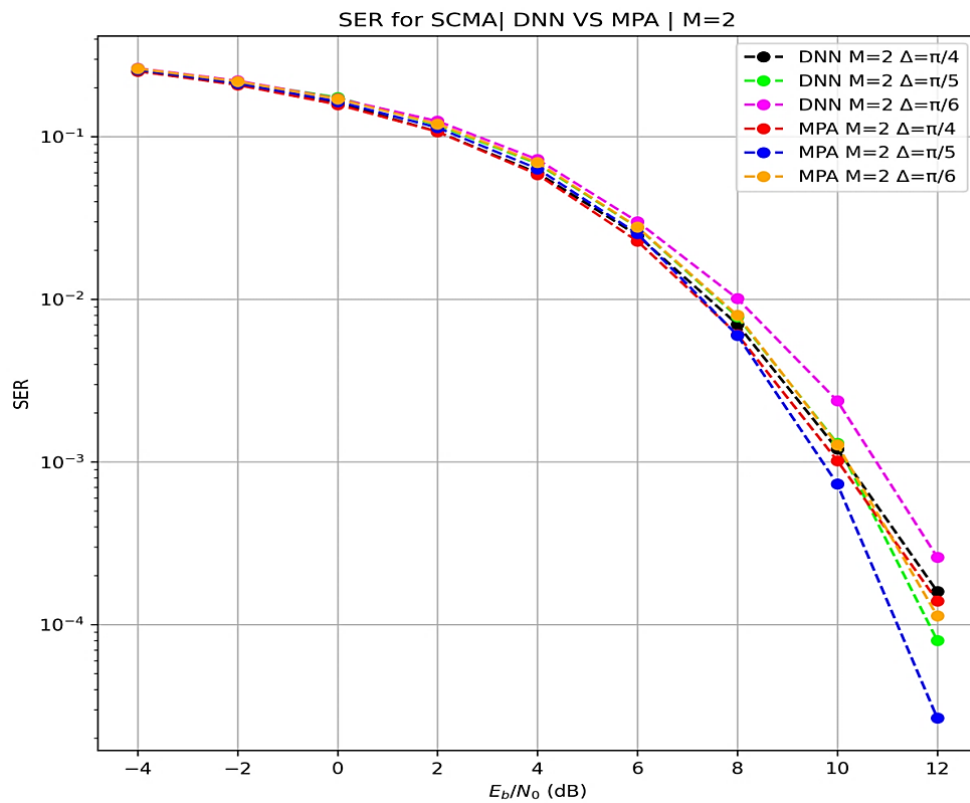


Figure IV.20 : Évolution du SER du modèle proposé (M=2) comparé au MPA en fonction du SNR.

L'analyse attentive des courbes de SER obtenues pour le DNN et le MPA dans la **Figure IV.20** révèle que les algorithmes MPA et DNN présentent en effet des courbes similaires, ce qui démontre que les deux méthodes sont compétitives et offrent des performances similaires dans diverses conditions. Toutefois, l'algorithme MPA a tendance à légèrement dépasser le DNN, en particulier à des SNR plus élevés. Les décisions de Δ ont également un impact similaire sur les performances des deux algorithmes, avec une tendance générale à l'affirmation que $\Delta = \pi / 5$ offre une performance légèrement supérieure.

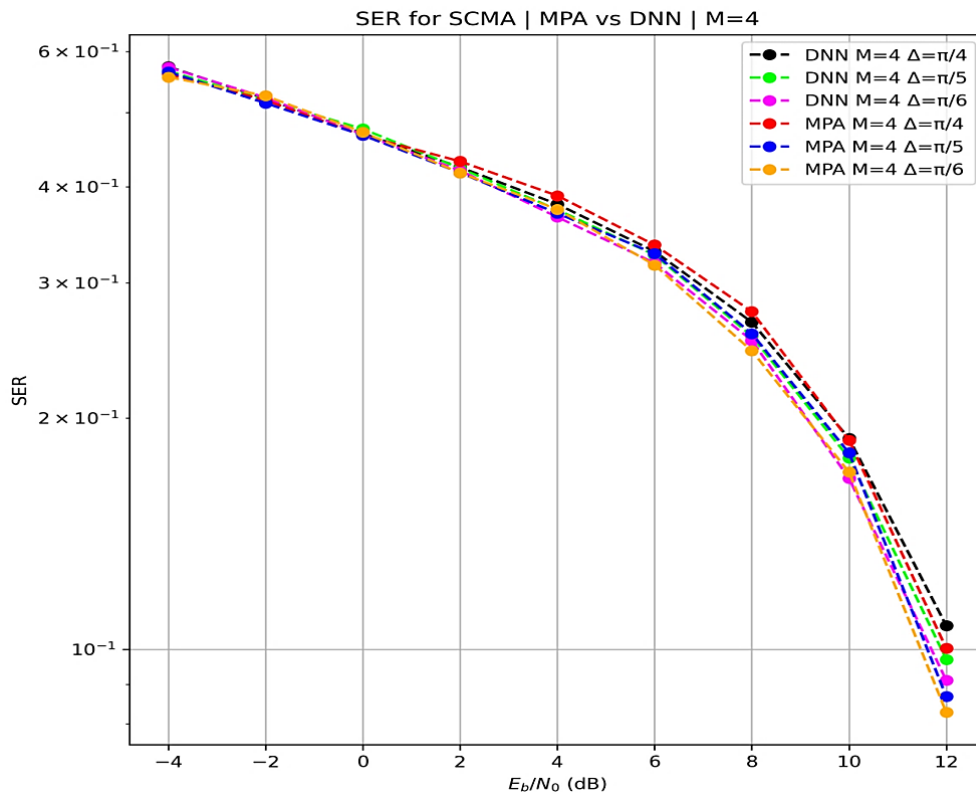


Figure IV.21: Évolution du SER du modèle proposé ($M=4$) comparé au MPA en fonction du SNR.

Dans ce cas, la **Figure IV.21** met en évidence la supériorité du DNN par rapport au MPA dans certains niveaux de SNR (0-10 dB) pour $\Delta = \pi / 4$, (2-6 dB) pour $\Delta = \pi / 6$ et dans l'intervalle de (6-10 dB) pour $\Delta = \pi / 5$. Tandis que le MPA surpasse progressivement le DNN à mesure que le SNR augmente mais les différences ne sont pas majeures. Il convient également de souligner que les courbes de SER des deux algorithmes sont très similaires sur toutes les plages de SNR étudiées.

IV.8.5. Analyse Comparative de SER entre DNN et NN :

Dans ces deux figures, on compare les performances de l'algorithme NN à celles du modèle DNN. Grâce à cette comparaison, il est possible de mesurer l'influence de la profondeur du réseau neuronal sur les capacités de décodage.

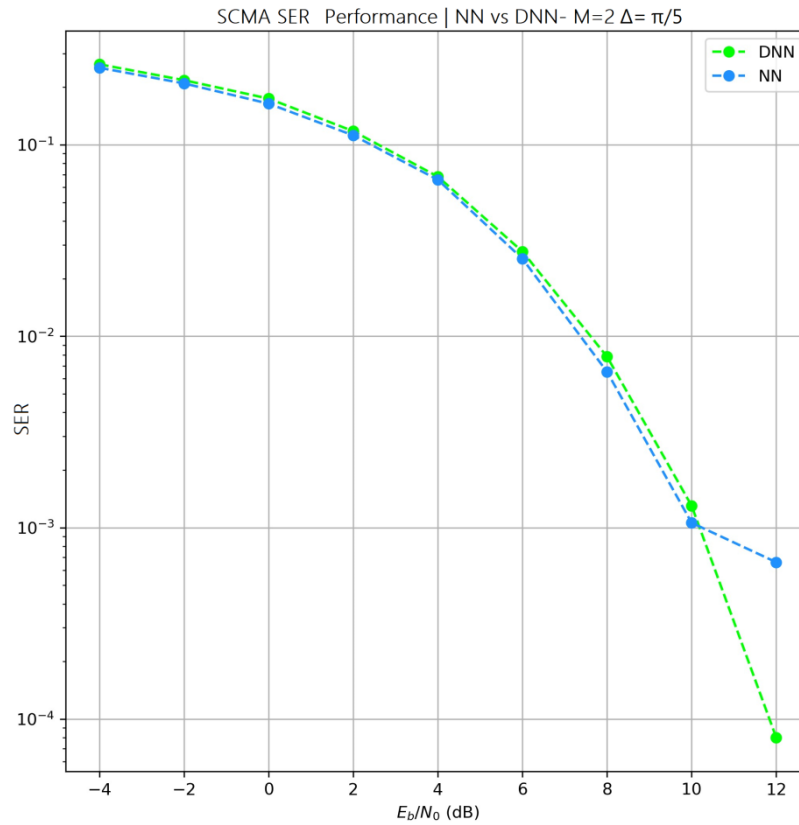


Figure IV.22 : Comparaison des courbes de SER entre NN et DNN pour $M=2$.

En illustrant la **Figure IV.22**, lorsque le SNR varie de -4 à 10 , le NN présente des performances de décodage supérieures à celles du DNN. Toutefois, à partir d'un SNR de 10 et plus, le DNN présente des performances plus élevées, dépassant celles du NN.

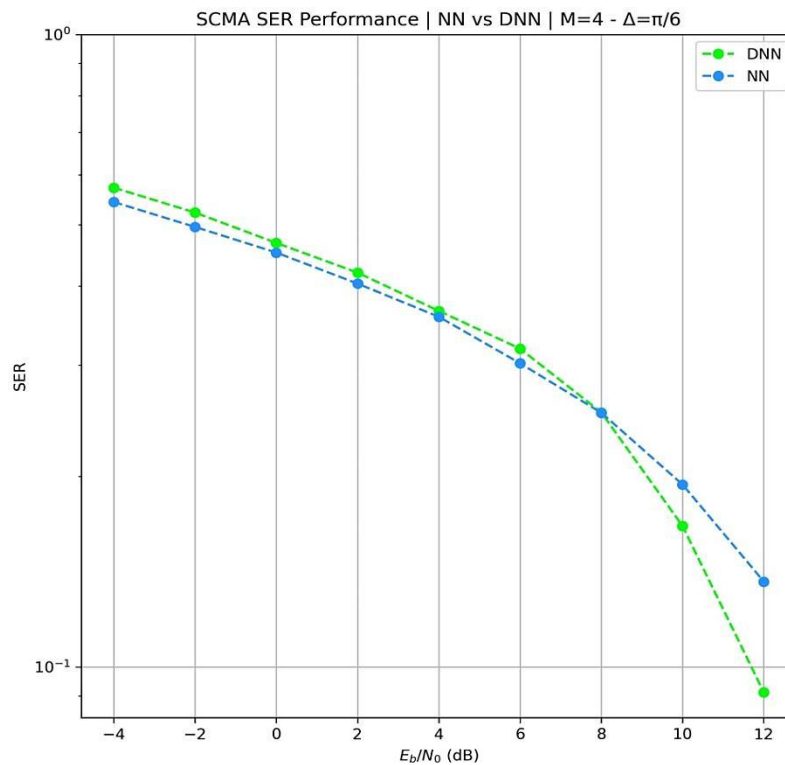


Figure IV.23 : Comparaison des courbes de SER entre NN et DNN pour $M=4$.

Dans la **Figure IV.23**, on constate la même observation : dans la plage de SNR allant de -4 à 8, le NN présente des performances de décodage supérieures à celles du DNN. Toutefois, dès qu'un SNR atteint 8 et plus, le DNN présente des performances plus élevées, dépassant celles du NN..

Les NN peuvent parfois être plus efficaces dans des conditions de faible SNR en raison de leur simplicité, car ils ont moins de chances de surajuster le bruit par rapport aux DNN plus complexes. Ils peuvent être plus résistants et moins sensibles au bruit excessif grâce à leur architecture plus simple. Il arrive parfois que les DNN surajustent le bruit, ce qui entraîne des performances moins optimales dans des environnements à faible SNR. Il peut être difficile de distinguer le signal du bruit en raison de leur complexité lorsque le bruit prédomine.

Les NN fournissent des performances satisfaisantes dans des conditions de SNR élevé, mais peuvent être restreints par leur capacité à détecter des relations complexes dans des données plus précieuses. En utilisant un signal à haut SNR, les DNN se distinguent par leur capacité à capturer des nuances et des détails complexes, ce qui améliore la précision et la fiabilité.

En raison de cette relation non linéaire entre les deux types de réseaux de neurones et le SNR, il est essentiel d'examiner de manière approfondie les performances dans diverses conditions de canal. Il sera donc nécessaire de choisir le modèle approprié en prenant en compte les spécificités du système de communication et les exigences de performance.

IV.8.6. Analyse de complexité :

L'un des principaux objectifs de ce projet est de réduire la complexité de la détection SCMA. Par conséquent, il est crucial de présenter une analyse de complexité comparant la détection DNN proposée avec la MPA. La **Figure IV.24 (a)** montre la comparaison du temps d'exécution entre le MPA et le DNN, tandis que la **Figure IV.24 (b)** offre un aperçu plus détaillé de cette comparaison dans une plage de valeurs plus restreinte. L'analyse de complexité est réalisée sur un CPU Intel(R) Core(TM) i5-12400F de 12e génération, cadencé à 2,50 GHz.

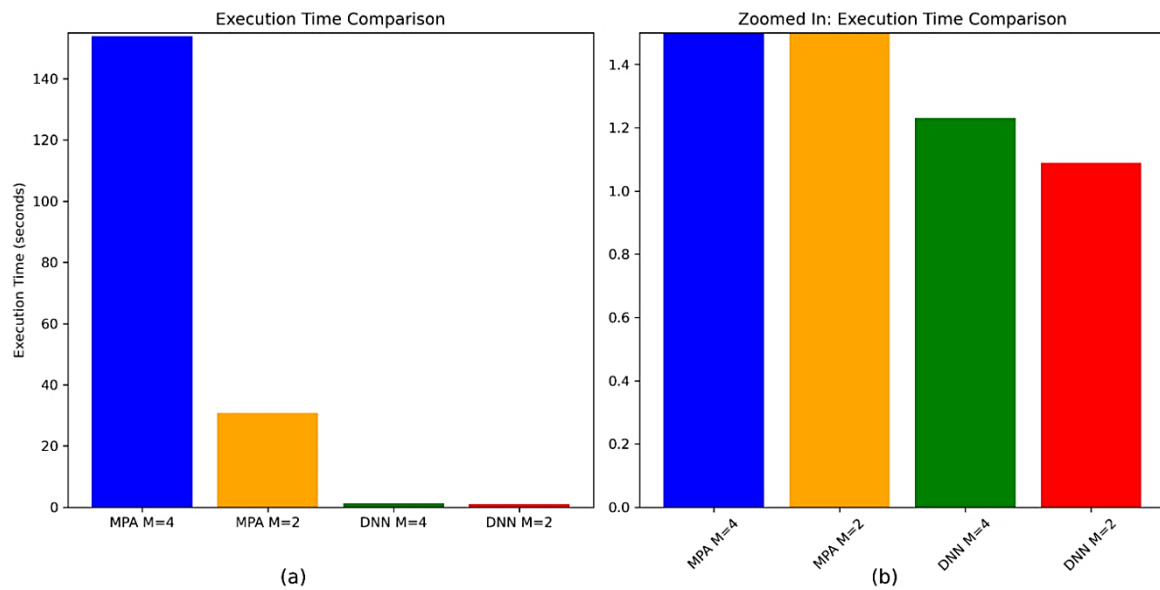


Figure IV.24 : Comparaison du temps d'exécution entre MPA et DNN en secondes.

La comparaison entre les réseaux DNN et le MPA dans un système SCMA révèle des différences significatives en termes de temps d'exécution.

Pour le MPA, le temps d'exécution varie considérablement en fonction du niveau de modulation. Avec $M=4$, le temps d'exécution est de 153.77 secondes, alors que pour $M=2$, il est réduit à 30.74 secondes. Cette augmentation substantielle du temps d'exécution au fur et à mesure que le niveau de modulation augmente suggère que la détection de MPA devient plus exigeante en termes de calcul avec des niveaux de modulation plus élevés.

En revanche, l'approche DNN présente des temps d'exécution cohérents et nettement plus courts pour différents types de modulation. Pour $M=4$ et $M=2$, les temps d'exécution sont remarquablement bas, avec respectivement 1.23 seconde et 1.09 seconde. Ces temps d'exécution

nettement plus courts indiquent que la détection basée sur les DNN reste efficace quel que soit le niveau de modulation.

Cette comparaison souligne le potentiel de la détection basée sur les DNN en tant qu'alternative plus efficace à la MPA, en particulier dans les scénarios où la réduction de la complexité est primordiale. Les DNN offrent non seulement des temps d'exécutions constamment faibles, mais aussi la capacité de gérer des niveaux de modulation plus élevés sans augmentation significative de la charge de calcul, ce qui en fait une solution prometteuse pour la réduction de la complexité dans les systèmes SCMA.

IV.9. Conclusion :

Au cours de ce chapitre, on a réalisé une comparaison entre le MPA et DNN dans le cadre des systèmes SCMA. Dans notre étude, on a analysé comment le SER évolue en fonction du SNR, ainsi que la complexité de calcul liée à chaque technologie.

Dans un premier temps, l'analyse de l'évolution de SER en fonction du SNR révèle des performances similaires entre les deux méthodes, avec des avantages visibles pour le DNN à des niveaux de SNR inférieurs et une meilleure efficacité du MPA à mesure que le SNR augmente. Cette observation souligne l'importance de choisir le bon modèle en fonction des conditions spécifiques du canal de communication et des exigences de performance.

En outre, l'analyse de la complexité de calcul met en évidence que le DNN a généralement des temps d'exécution plus courts que le MPA, mettant en évidence son potentiel pour des applications en temps réel et à faible latence.

Conclusion Générale et Perspectives

Les chercheurs en communications sans fils sont constamment motivés par le besoin d'accroître constamment le nombre d'utilisateurs, d'améliorer le débit, de réduire le temps d'attente et de réduire les erreurs et la perte d'information. Cela les pousse à améliorer constamment les systèmes de communications existants et à envisager de meilleures solutions pour les standards à venir.

Ce mémoire examine les performances de la technique d'accès multiple SCMA, une méthode émergente de multiplexage qui prend de l'importance dans le contexte de la 5G, en mettant l'accent sur son intégration avec l'intelligence artificielle

Dans le premier chapitre , nous avons examiné l'essence même de la 5G. Nous examinons ses objectifs techniques, comme des débits ultra-rapides, une latence minimale et une capacité massive. En outre, nous étudions les bases technologiques de la 5G, ses différentes architectures, ainsi que les difficultés associées à sa mise en place.

Dans le deuxième chapitre, nous avons exploré les techniques d'accès multiple orthogonal (OMA/NOMA), avec une attention particulière portée sur SCMA. Nous avons détaillé ses mécanismes de fonctionnement, ses caractéristiques clés et ses performances.

Au cours de troisième chapitre, nous avons plongé dans le domaine de l'IA en mettant en lumière les concepts de ML et de DL, ainsi que les cas d'utilisation de l'IA dans les réseaux 5G. Nous avons examiné en détail les réseaux de neurones, leurs différentes architectures et leurs applications spécifiques dans les réseaux 5G, notamment dans la prédiction de la qualité de service, la gestion du spectre et l'optimisation des performances du réseau.

Dans le dernier chapitre , En mettant en œuvre des techniques avancées de DL, Nous avons développé deux modèles de réseaux de neurones spécifiquement adaptés à la modulation BPSK et QPSK visant à réduire le taux d'erreur de symbole (SER) et à diminuer la complexité de détection dans les systèmes utilisant la technique SCMA. Nous avons évalué la performance du système en analysant le SER en fonction du SNR. Les résultats obtenus nous ont permis de comparer et d'évaluer les performances des DNNs avec le MPA.

L'approche proposée présente de nombreux avantages par rapport aux méthodes traditionnelles de détection du MPA dans les réseaux 5G, comme le démontrent les résultats de la simulation. D'une part, la réduction de la complexité de détection est significative, ce qui permet d'améliorer l'efficacité opérationnelle du système. D'autre part, bien que l'amélioration du SER ne soit pas observée dans toutes les conditions, il reste dans la même plage que celui du MPA. Cela se traduit par une QoS comparable pour les utilisateurs du système NOMA SCMA dans le contexte de la 5G.

En bref, cette étude joue un rôle important dans l'amélioration des performances des réseaux de communication sans fil dans le cadre de la 5G, en incorporant l'intelligence artificielle dans la conception des systèmes SCMA. Ces résultats encourageants permettent d'explorer de nouvelles applications dans les réseaux 5G, où la gestion des ressources spectrales est efficace

En perspective, l'utilisation de CNN et RNN dans la détection SCMA offre des opportunités passionnantes pour repousser les limites de l'efficacité et des performances des systèmes de communication sans fil, ouvrant ainsi la voie à des réseaux plus intelligents et plus adaptatifs pour répondre aux exigences croissantes des futures normes de communication sans fil.

Bibliographie/WEBOGRAPHIE

- [1] ANFR. (Juillet 2019). "Rapport d'évaluation de l'exposition du public aux ondes électromagnétiques 5G, Volet 1 : présentation générale de la 5G."
- [2] J. P. Damiano, "De la 5G à la 6G: contexte et enjeux!", 2020.
- [3] A. Pouyat, M. Laroche, and J. L. Strauss, "La 5G: des convergences, une ambition, des opportunités et des risques," Académie des technologies & Altran research, Juillet 2022.
- [4] R. Dangi et al, "Study and investigation on 5G technology: A systematic review," Sensors, vol. 22, no. 1, p. 26, 2021.
- [5] T. S. Rappaport et al, "Millimeter wave mobile communications for 5G cellular: It will work!," IEEE access, vol. 1, pp. 335-349, 2013.
- [6] N. Panwar, S. Sharma, and A. K. Singh, "A survey on 5G: The next generation of mobile communication," Physical Communication, vol. 18, pp. 64-84, 2016.
- [7] A. Jaakola, "Unity 3D Visualization Tool for 5G Cell Measurements," 2019.
- [8] D. Li, "5G and intelligence medicine—How the next generation of wireless technology will reconstruct healthcare?," Precision Clinical Medicine, vol. 2, no. 4, pp. 205–208, 2019.
- [9] S. Shakib et al, "mmWave CMOS power amplifiers for 5G cellular communication," IEEE Communications Magazine, vol. 57, no. 1, pp. 98-105, 2019.
- [10] R. Khan et al., "A survey on security and privacy of 5G technologies: Potential solutions, recent advancements, and future directions," IEEE Communications Surveys & Tutorials, vol. 22, no. 1, pp. 196-248, 2019.
- [11] Arcep. (Mars 2017). "Les enjeux de la 5G." Rapport technique.
- [12] M. H. Khan and P. C. Barman, "5G-Future Generation Technologies Of Wireless Communication Revolution 2020," American Journal of Engineering Research, vol. 4, no. 5, pp. 206-215, 2015.
- [13] M. Z. Noohani and K. U. Magsi, "A review of 5G technology: Architecture, security and wide applications," International Research Journal of Engineering and Technology (IRJET), vol. 7, no. 05, pp. 3440-3471, 2020.

- [14] S. Chen et al, "Machine-to-machine communications in ultra-dense networks—A survey," *IEEE Communications Surveys & Tutorials*, vol. 19, no. 3, pp. 1478-1503, 2017.
- [15] S. Li, L. Da Xu, and S. Zhao, "5G Internet of Things: A survey," *Journal of Industrial Information Integration*, vol. 10, pp. 1-9, 2018.
- [16] S. Gamal et al, "Multiple access in cognitive radio networks: From orthogonal and non-orthogonal to rate-splitting," *IEEE Access*, vol. 9, pp. 95569-95584, 2021.
- [17] R. Dangi et al, "Study and Investigation on 5G Technology: A Systematic Review," *Sensors (Basel)*, vol. 22, no. 1, p. 26, Dec. 2021.
- [18] CRIIREM - DV – 0120, "Ce qu'il faut savoir sur la 5G," 2020.
- [19] A. B. Ericsson, "5G radio access—capabilities and technologies," 2016.
- [20] 5G Americas, "Understanding mmWave for 5G Networks," White paper, Dec. 2020.
- [21] F. Rinaldi, A. Raschella, and S. Pizzi, "5G NR system design: A concise survey of key features and capabilities," *Wireless Networks*, vol. 27, no. 8, pp. 5173-5188, 2021.
- [22] S. A. Khwandah et al, "Massive MIMO systems for 5G communications," *Wireless Personal Communications*, vol. 120, no. 3, pp. 2101-2115, 2021.
- [23] I. Lagroye et al, "La 5G et la santé," Nov. 2022.
- [24] O. Idowu-Bismark et al, "5G wireless communication network architecture and its key enabling technologies," vol. 12, pp. 70-82, 2019.
- [25] GSM Association, "An Introduction to Network Slicing," 2017.
- [26] E. Mutafulungwa, J. Favaro, and S. Parker, "White Paper on Small Cells: How Europe Can Accelerate Network Densification for the 5G Era," 2019.
- [27] O. Shental, B. M. Zaidel, and S. S. Shitz, "Low-density code-domain NOMA: Better be regular," in *2017 IEEE International Symposium on Information Theory (ISIT)*, June 2017, pp. 2628-2632

- [28] S. K. Goudos, P. I. Dallas, S. Chatziefthymiou, and S. Kyriazakos, "A survey of IoT key enabling and future technologies: 5G, mobile IoT, semantic web and applications," *Wireless Personal Communications*, vol. 97, pp. 1645-1675, 2017.
- [29] I. F. Akyildiz, W.-Y. Lee, and K. R. Chowdhury, "CRAHNS: Cognitive radio ad hoc networks," *Ad Hoc Networks*, vol. 7, no. 5, pp. 810–836, 2009.
- [30] N. Iswarya and L. S. Jayashree, "A survey on successive interference cancellation schemes in non-orthogonal multiple access for future radio access," *Wireless Personal Communications*, vol. 120, no. 2, pp. 1057-1078, 2021.
- [31] V. K. Deolia, "Code domain non orthogonal multiple access schemes for 5G and beyond communication networks: A review," *Journal of Engineering Research*, vol. 10, no. 4A, 2022.
- [32] Université de Rome La Sapienza, "Telecomunicazioni [Cours de premier cycle en génie électrique]," Rome, Italie, 2007-2008.
- [33] A. A. Driouche and A. Roumane, "Étude des mécanismes de gestion de congestion dans l'EUTRAN pour les applications M2M IoT (cas d'étude PRACH) [Thèse, Institut National des Télécommunications et des Technologies de l'Information et de la Communication]."
- [34] J. Ghosh et al., "On Comparison of Optimal NOMA and OMA in Paradigm Shift of Emerging Technologies," *IEEE*, 2020.
- [35] S. R. Islam, N. Avazov, O. A. Dobre, and K. S. Kwak, "Power-domain non-orthogonal multiple access (NOMA) in 5G systems: Potentials and challenges," *IEEE Communications Surveys & Tutorials*, vol. 19, no. 2, pp. 721-742, 2016.
- [36] M. Ghous et al., "Cooperative power-domain NOMA systems: an overview," *Sensors*, vol. 22, no. 24, p. 9652, 2022.
- [37] L. Lv, J. Chen, Q. Ni, Z. Ding, and H. Jiang, "Cognitive non-orthogonal multiple access with cooperative relaying: A new wireless frontier for 5G spectrum sharing," *IEEE Communications Magazine*, vol. 56, no. 4, pp. 188-195, 2018.
- [38] M. Al-Imari and M. A. Imran, "Low density spreading multiple access," in *Multiple Access Techniques for 5G Wireless Networks and Beyond*, pp. 493-514, 2019.

- [39] L. Dai et al. "A survey of non-orthogonal multiple access for 5G," *IEEE communications surveys & tutorials*, vol. 20, no. 3, pp. 2294-2323, 2018.
- [40] Yuan, Z., Yu, G., Li, W., Yuan, Y., Wang, X., & Xu, J. (2016, May). "Multi-user shared access for internet of things," presented at the 2016 IEEE 83rd Vehicular Technology Conference (VTC Spring), pp. 1-5. IEEE
- [41] Mhedhbi, M., & Boukour, F. E. (2019). "Analysis and evaluation of pattern division multiple access scheme jointed with 5G waveforms," *IEEE Access*, vol. 7, pp. 21826-21833.
- [42] 3GPP TSG RAN WG1 Meeting (2016). "New uplink non-orthogonal multiple access schemes for NR," Gothenburg, Sweden.
- [43] Chaturvedi, S., Liu, Z., Bohara, V. A., Srivastava, A., & Xiao, P. (2022). "A tutorial on decoding techniques of sparse code multiple access," *IEEE Access*, vol. 10, pp. 58503-58524.
- [44] Wang, C., Chen, Y., & Wu, Y. (2018). "Sparse Code Multiple Access for 5G Radio Transmission," Conference Paper.
- [45] Zhou, Y., Yu, Q., Meng, W., & Li, C. (2017, May). "SCMA codebook design based on constellation rotation," presented at the 2017 IEEE International Conference on Communications (ICC), pp. 1-6. IEEE.
- [46] Klimentyev, V. P., & Sergienko, A. B. (2016, September). "A low-complexity SCMA detector for AWGN channel based on solving overdetermined systems of linear equations," presented at the 2016 XV International Symposium Problems of Redundancy in Information and Control Systems (REDUNDANCY), pp. 61-65. IEEE.
- [47] Chaturvedi, S., Liu, Z., Bohara, V. A., Srivastava, A., & Xiao, P. (2021). "A tutorial to sparse code multiple access," arXiv preprint arXiv:2105.06860.
- [48] Han, S., Huang, Y., & Wang, B. (2018). "A method for analysing and improving the multi-user detection algorithm of SCMA," in *Wireless Internet: 10th International Conference, WiCON 2017, Tianjin, China, December 16-17, 2017, Proceedings 10*, pp. 103-115. Springer International Publishing.
- [49] Li, S., Feng, Y., Sun, Y., & Xia, Z. (2022). "A low-complexity detector for uplink SCMA by exploiting dynamical superior user removal algorithm," *Electronics*, vol. 11, no. 7, p. 1020.

- [50] Cao, Y., Sun, H., Soriaga, J., & Ji, T. (2017, September). "Resource spread multiple access-a novel transmission scheme for 5G uplink," presented at the 2017 IEEE 86th Vehicular Technology Conference (VTC-Fall), pp. 1-5. IEEE.
- [51] A. M. H. Bilal, "Polycopié Pédagogique de Cours du module Intelligence Artificielle."
- [52] J.-P. Haton, E. Haton, and M.-C. Haton, *Intelligences artificielles : De la théorie à la pratique*. Malakoff, France, 2023.
- [53] J. Smith and A. Johnson, "Evolution of Artificial Intelligence: From Symbolic Logic to Deep Learning," *Journal of Artificial Intelligence Research*, vol. 45, no. 2, pp. 153-170, 2020.
- [54] M. Djelti and B. Kouninef, "Intelligence artificielle appliquée," Nov. 2023.
- [55] M. Ghalia, "Apprentissage profond," Université Mustapha Benboulaïd de Batna, 2020-2021.
- [56] C.-A. Azencott, *Introduction au Machine Learning*. Paris, France, 2019.
- [57] E. S. Puchi-Cabrera, E. Rossi, G. Sansonetti, M. Sebastiani, and E. Bemporad, "Machine learning aided nanoindentation: A review of the current state and future perspectives," *Current Opinion in Solid State and Materials Science*, vol. 27, no. 4, p. 101091, 2023.
- [58] A. Burkov, *The hundred-page machine learning book*, vol. 1. Quebec City, QC, Canada: Andriy Burkov, 2019, p. 32.
- [59] Y. Fu, S. Wang, C. X. Wang, X. Hong, and S. McLaughlin, "Artificial intelligence to manage network traffic of 5G wireless networks," *IEEE network*, vol. 32, no. 6, pp. 58-64, 2018.
- [60] F. Vahl, "Active Vision for Humanoid Soccer Robots using Reinforcement Learning."
- [61] K. Fukushima and S. Miyake, "Neocognitron: a self-organizing neural network model for a mechanism of visual pattern recognition," 1982.
- [62] G. Varoquaux and O. Colliot, "Machine Learning for Brain Disorders," 2012.
- [63] P. Kim, "Matlab deep learning. With machine learning, neural networks and artificial intelligence," vol. 130, no. 21, 2017.
- [64] C. Touzet, *les réseaux de neurones artificiels, introduction au connexionnisme*. Ec2, 1992.

- [65] Glorot, X. (2015). "Apprentissage des réseaux de neurones profonds et applications en traitement automatique de la langue naturelle."
- [66] Jiang, X., Hadid, A., Pang, Y., Granger, E., & Feng, X. (Eds.). (2019). "Deep learning in object detection and recognition."
- [67] Le, N., Rathour, V. S., Yamazaki, K., Luu, K., & Savvides, M. (2022). "Deep reinforcement learning in computer vision: a comprehensive survey," *Artificial Intelligence Review*, pp. 1-87.
- [68] François-Lavet, V., Henderson, P., Islam, R., Bellemare, M. G., & Pineau, J. (2018). "An introduction to deep reinforcement learning," *Foundations and Trends® in Machine Learning*, vol. 11, no. 3-4, pp. 219-354.
- [69] Labiad, A. (2017). "Sélection des mots clés basée sur la classification et l'extraction des règles d'association" (Doctoral dissertation, Université du Québec à Trois-Rivières).
- [70] Mwamba Kasongo Dahouda & Joe, I. (2021). "A Deep-Learned Embedding Technique for Categorical Features Encoding," *IEEE Access*, vol. 9, pp. 12345-12356.
- [71] Garbin, C., Zhu, X., & Marques, O. (2020). "Dropout vs. batch normalization: an empirical study of their impact to deep learning," *Multimedia tools and applications*, vol. 79, no. 19, pp. 12777-12815.
- [72] Lavangnananda, K., & Chattanachot, S. (2017, February). "Study of discretization methods in classification," presented at the 2017 9th International Conference on Knowledge and Smart Technology (KST), pp. 50-55. IEEE.
- [73] Reyad, M., Sarhan, A. M., & Arafa, M. (2023). "A modified Adam algorithm for deep neural network optimization," *Neural Computing and Applications*, vol. 35, no. 23, pp. 17095-17112.
- [74] Mustapha, A., Mohamed, L., & Ali, K. (2021). "Comparative study of optimization techniques in deep learning: Application in the ophthalmology field," in *Journal of physics: conference series*, vol. 1743, no. 1, p. 012002. IOP Publishing.
- [75] Yi, D., Ahn, J., & Ji, S. (2020). "An effective optimization method for machine learning based on ADAM," *Applied Sciences*, vol. 10, no. 3, p. 1073.
- [76] F. Maleki, N. Muthukrishnan, K. Ovens, C. Reinhold, and R. Forghani, "Machine learning algorithm validation: from essentials to advanced applications and implications for regulatory certification and deployment," *Neuroimaging Clinics*, vol. 30, no. 4, pp. 433-445, 2020

- [77] J. Terven, D. M. Cordova-Esparza, A. Ramirez-Pedraza, and E. A. Chavez-Urbiola, "Loss Functions and Metrics in Deep Learning," arXiv preprint arXiv:2307.02694, 2023.
- [78] S. B. Kamble, H. N. Bhor, R. Patil, M. S. Naik, and V. P. Vasani, "Deep Learning," Université de Mumbai, 2022.
- [79] Liu, T. Y., & Mirzasoleiman, B. (2022). "Data-Efficient Augmentation for Training Neural Networks," *Advances in Neural Information Processing Systems*, vol. 35, pp. 5124-5136.
- [80] Rousseau David & Douarre, C. (2021). "Introduction à l'apprentissage profond (deep learning) de l'intelligence artificielle."
- [81] Bhardwaj, M., Jain, R., Kumar, S., & Kumari, U. (2023). "Use of Machine Learning Approaches in 5G Applications," in *Applications of Artificial Intelligence in Wireless Communication Systems*, pp. 137-149. IGI Global.

WEBGRAPHIE

- [1] <https://colab.google/>
- [2] <https://code.visualstudio.com/>
- [3] <https://www.python.org/>
- [4] <https://www.tensorflow.org/?hl=fr>
- [5] <https://keras.io/>
- [6] <https://scikit-learn.org/>
- [8] <https://matplotlib.org/>
- [7] <https://numpy.org/>