

الجمهورية الجزائرية الديمقراطية الشعبية

République algérienne démocratique et populaire

وزارة التعليم العالي والبحث العلمي

Ministère de l'enseignement supérieur et de la recherche scientifique

جامعة عين تموشنت بلحاج بوشعيب

Université Ain Temouchent Belhadj Bouchaib

Faculté des Sciences et de Technologie

Département d'électronique et des télécommunications



Projet de Fin d'étude

Pour l'obtention du diplôme de master en : Télécommunications

Spécialité : Réseaux et télécommunications

Thème

Classification intelligente des modulations radio par Deep Learning

Présentée Par :

1) M^{elle}. BERRICHI Bouchera

2) M^{elle}. BOUCHIBA kheira samah

Devant le jury composé de :

Pr. DEBBAL Mohammed

Dr. MEKAMI Hayat

Dr. SOUIKI Sihem

Pr UAT.B. B (Ain Temouchent)

MCB UAT.B. B (Ain Temouchent)

MCA UAT.B. B (Ain Temouchent)

Président

Examinatrice

Encadrante

Année Universitaire 2025/2026

Remerciement

Tout d'abord, je remercie Allah de m'avoir accordé la santé, la patience et la volonté nécessaires pour mener à bien ce travail.

Je tiens à exprimer ma profonde gratitude à mon encadrante pour son accompagnement, ses conseils pertinents, sa disponibilité et son soutien tout au long de la réalisation de ce projet.

Je remercie également l'ensemble des enseignants qui ont contribué à ma formation académique durant mon parcours universitaire, pour la qualité de leur enseignement et leurs efforts constants.

Mes remerciements s'adressent aussi à ma famille, qui m'a soutenu moralement et encouragé sans relâche, et qui a été une source de motivation tout au long de mes études.

Enfin, je remercie toutes les personnes qui, de près ou de loin, ont contribué à la réalisation de ce travail.

Dédicace

Nous dédions ce modeste travail à nos chers parents, qui ont toujours cru en nous et nous ont soutenus tout au long de notre parcours académique. Leur amour, leurs sacrifices et leurs encouragements constants ont été notre plus grande source de motivation.

Nous l'adressons également à nos familles et à nos proches pour leur patience, leur soutien moral et leur confiance durant les différentes étapes de ce projet.

Nos sincères remerciements vont aussi à nos enseignants, qui nous ont guidés par leurs conseils, leurs connaissances.

Enfin, nous dédions Ce travail à tous ceux qui nous ont aidés, encouragés et accompagnés, de près ou de loin, dans la réalisation de ce projet de fin d'études.

*Bouchera
Samah*

Résumé : La classification automatique des modulations (AMC) est une technologie clé pour de nombreuses applications sans fil modernes, telles que la radio cognitive, la surveillance du spectre et les réseaux intelligents. Ce mémoire présente l'étude et le développement d'une architecture hybride de Deep Learning destinée à identifier automatiquement le type de modulation de signaux radio reçus sans information préalable. Le modèle proposé combine des couches de réseaux de neurones convolutifs (CNN) pour l'extraction des caractéristiques spatiales, avec un Transformer, un réseau récurrent BiGRU et un mécanisme d'attention afin de capturer efficacement les dépendances séquentielles à long terme des signaux I/Q. L'approche est validée sur la base de données de référence RadioML 2016.10a, comprenant 11 classes de modulations analogiques et numériques. Cette architecture vise à améliorer la robustesse et la précision de la reconnaissance des modulations dans des environnements de transmission complexes, marqués par le bruit, les décalages fréquentiels et les effets de canal.

Mots-clés : Classification automatique des modulations (AMC), Deep Learning, CNN, BiGRU, Mécanisme d'attention, RadioML 2016.10a, Signaux I/Q.

.

Abstract: Automatic Modulation Classification (AMC) is a key technology for many modern wireless applications, such as cognitive radio, spectrum monitoring, and intelligent networks. This thesis presents the study and development of a hybrid Deep Learning architecture designed to automatically identify the modulation type of received radio signals without prior information. The proposed model combines Convolutional Neural Network (CNN) layers for spatial feature extraction with a Transformer, a Bidirectional Gated Recurrent Unit (BiGRU), and an attention mechanism to efficiently capture long-term sequential dependencies in I/Q signals. The approach is validated using the benchmark RadioML 2016.10a dataset, which includes 11 classes of analog and digital modulations. This architecture aims to improve the robustness and accuracy of modulation recognition in complex transmission environments affected by noise, frequency offsets, and channel impairments.

Keywords: Automatic Modulation Classification (AMC), Deep Learning, CNN, BiGRU, Attention Mechanism, RadioML 2016.10a, I/Q Signals

ص خلما:

تقنية أساسية في العديد من التطبيقات (Automatic Modulation Classification - AMC) يُعدّ التصنيف التلقائي للتضمين اللاسلكية الحديثة، مثل الراديو الإدراكي، مراقبة الطيف الترددي، والشبكات الذكية. تقدم هذه المذكرة دراسة وتطوير بنية هجينة للتعلم مصممة للتعرف تلقائيًا على نوع التضمين للإشارات الراديوية المستقبلية دون توفر معلومات مسبقة. يجمع (Deep Learning) العميق لاستخراج الخصائص (Convolutional Neural Network - CNN) النموذج المقترح بين طبقات الشبكات العصبية الالتفافية ثنائية الاتجاه، وآلية الانتباه Bidirectional Gated Recurrent Unit (BiGRU) ووحدة Transformer المكانية، مع نموذج تم التحقق من فعالية هذه I/Q من أجل النقاط التبعيات التسلسلية طويلة المدى بكفاءة داخل إشارات (Attention Mechanism) والتي تضم 11 فئة من التضمينات التماثلية والرقمية. تهدف هذه RadioML 2016.10a المقاربة باستخدام قاعدة البيانات المرجعية البنية إلى تحسين متانة ودقة التعرف على التضمين في بيانات إرسال معقدة تتأثر بالضوضاء، وانزياحات التردد، وتشوهات القناة.

:الكلمات المفتاحية

RadioML، آلية الانتباه، BiGRU، (CNN) التعلم العميق، الشبكات العصبية الالتفافية، (AMC) التصنيف التلقائي للتضمين I/Q، إشارات 2016.10a.

Sommaire

Introduction générale :	1
Chapitre 01 : Généralités sur la modulation numérique.	3
I .1Introduction	4
I.2 Généralités sur la modulation	4
I .2.1Définition	4
I .2.2 La modulation analogique	5
I.2.2.1 Modulation d'amplitude (AM)	6
I .2.2.2Modulation de fréquence (FM)	7
I .2.2.3Modulation de phase (PM)	7
I .3Modulation numérique	9
I .3.1Modulation ASK (Amplitude Shift Keying)	9
I.3.1.1 La modulation OOK (On-Off Keying)	9
I.3.1.2 Modulation M-ASK	10
I.3.1.3 Avantages de la modulation ASK	11
I.3.1.4 Inconvénients de la modulation ASK	11
I .3.2Modulation PSK (Phase Shift Keying)	12
I.3.2.1 Binary Phase Shift Keying (BPSK)	12
I.3.2.2 Avantages de la modulation PSK	13
I.3.2.3 Inconvénients de la modulation PSK	14
I.3.2.4 Diagramme de constellation de la modulation 16-PSK	14
I .3.3 La modulation FSK (Frequency Shift Keying)	15
I.3.3.1 FSK binaire (BFSK)	15
I.3.3.2 Modulation M-FSK	15
I.3.3.3 Avantages de la modulation FSK	16
I.3.3.4 Inconvénients de la modulation FSK	16
I.3.3.5 Diagramme de constellation de la modulation 4-FSK	16
I .3.4 Modulation QAM (Quadrature Amplitude Modulation)	17
I.3.4.1 Avantages de la modulation QAM	18
I.3.4.2 Inconvénients de la modulation QAM	18

I .4 Canal de transmission et perturbations	19
I.4.1 Le bruit blanc gaussien additif (AWGN)	20
I.4.2 L'atténuation du signal	20
I.4.3 Le fading (évanouissement)	21
I.4.4 La distorsion	22
I .5 Performances de la modulation numérique	22
I.5.1 Rapport Signal sur Bruit (SNR)	22
I.5.2 Taux d'Erreur Binaire (BER)	22
I.5.3 Efficacité spectrale	23
I .6 Conclusion	23
Chapitre 02 : les fondements de l'intelligence artificielle.	24
II.1 Introduction	25
II.1.2 Introduction à IA	25
II.1.2.1 Définition d'IA	25
II.1.2.2 Historique d'IA	26
II .1.2.2 Domaines D'applications	27
II .2 Apprentissage automatique (Machine Learning)	28
II.2.1 Définition	28
II.2.2 Types d'apprentissage	29
II.2.2.1 Apprentissage supervisé	29
II.2.2.2 Apprentissage non supervise	30
II .2.2.3Apprentissage par renforcement	31
II .2.2.4 Apprentissage semi-supervisé	32
II.3 Réseaux de neurones artificiels	33
II .3.1Neurone biologique vs neurone artificiel	33
II .3.1.1Neurone biologique	33
II .3.1.2 Neurone artificiel	34
II .4 Perceptron	34
II.4.1 Fonctionnement du perceptron :	35
II.4.1.1 Les entrées (inputs)	35
II.4.1.2 Les poids (Weights)	35
II.4.1.3 Le biais (Bias)	36

II.4.1.4 La fonction sommation	36
II .5 Réseau multi couche (MPL)	36
II .5.2 La couche d'enter	37
II .5.2 Les couches cachées	37
II .5.3 Couche de sortie (Output Layer)	38
II .6 Fonction d'activation	38
II .7 Deep Learning	41
II.7.1 Définition de l'apprentissage profonde (deep learning)	41
II.7.2 La différence entre l'apprentissage automatique et l'apprentissage profond	42
II.7.3 Les avantages et les inconvénients de l'apprentissage profond	44
Défis juridiques et éthiques :	45
II .8 Conclusion	45
Chapitre 03 : État de l'art sur l'AMC basée sur le Deep learning.	47
III.1 Introduction	48
III.2 Taxonomie des méthodes de Deep Learning pour l'AMC	48
III.3 Représentation du signal	49
III.2.1 I/Q temporel	50
III.2.2 Diagramme de constellation	51
III.4 Modèles basés sur CNN	51
III.4.1 CNN typique	52
III.4.2 CDSCNN	54
III.5 RNN	55
III.5.1 LSTM Typique	55
III.5.2 SL-LSTM (LSTM à une seule couche)	56
III.6 Le GNN (Graph Neural Network)	57
III.6.1 AvgNet	58
III.6.2 GNN-FiLM	59
III.7 Transformer	60
III.7.1 Typical model	61
III.7.2 MC former	61
III.8 Modèles hybrides de l'apprentissage profond	62
III.8.1 CNN-LSTM	63

III.8.2 GrrNet	63
III.9 Techniques de régularisation dans l'AMC	65
III.9.1 Régularisation explicite (Explicit Regularization)	65
III.9.2 Régularisation implicite (Implicit Regularization)	66
III.10 Conclusion	67
Chapitre 04 : Résultat et discussion.	68
IV.1 Introduction	69
IV.2 Description de la base de données RadioML 2016.10a	69
IV.2.1 Description de la base de données	69
IV.2.2 Types de modulation et représentation du signal	70
IV.2.3 Effets du canal et distorsions du signal	71
IV.2.4 Importance de base de données pour les applications de Deep Learning	72
IV.3 Préparation et structure des données	73
IV.4 Métriques de performance en classification	74
IV.4.1 Matrice de confusion	74
IV.4.2 Accuracy (Exactitude)	74
IV.4.3 Précision (Precision)	75
IV.4.4 Rappel (Recall ou Sensibilité)	75
IV.4.5 F1-Score	76
IV.5 Architecture conceptuelle du modèle de classification automatique des modulations	76
IV.5.1 présentations de l'architecture de modèle proposé	77
IV.5.1.1 CNN	77
IV.5.1.2 Transformer	78
IV.5.1.3 Le BiGRU (Bidirectional Gated Recurrent Unit)	78
IV.5.1.4 Le mécanisme d'attention	79
IV.5.1.5 Tête de classification	79
IV.6 Analyse de résultats obtenus	80
IV.6.1 Accuracy et Loss	80
IV.6.2 Accuracy vs SNR	81
IV.6.3 La matrice de confusion	82
IV.6.4 Le rapport de classification	84
IV.7 Comparaison avec des modèles des travaux récents	85
IV.8 Conclusion	87

Conclusion Générale	90
Références	90

Liste des figures

Figure I.1: représentation des différentes techniques de modulation	8
Figure I.2:Diagramme de constellation 4-ASk.....	12
Figure I.3:Représentation visuel de 16-PSK	13
Figure I.4:Diagramme de constellation de 16-PSK	15
Figure I.5:Représentation visuel de 4-FSK	16
Figure I.6: Diagramme de constellation 4-FSK.....	17
Figure I.7:Diagramme de constellation 16- QAM.	19
Figure I.8:Canal de transmission et perturbation	20
Figure II.9 : les types d'apprentissage.	29
Figure II.10:Apprentissage supervisé.	30
Figure II.11:Apprentissage non supervisé.....	31
Figure II.12:Apprentissage par renforcement.....	32
Figure II.13:Apprentissage semi –supervisé	32
Figure II.14:ANN vs BNN.....	33
Figure II.15:schéma d'un perceptron simple.	35
Figure II. 16:Architecture d'un perceptron multicouche (MPL).....	37
Figure II.17:Types de fonction d'activation.....	41
Figure III.18:Méthodes d'apprentissage profond (Deep Learning) [51]	49
Figure III.19:Structure générale du RNN.....	55
Figure III.20:Structure générale du LSTM.	57
Figure III.21:Schéma de l'AVgNet	59
Figure III. 22:GNN-FILM.....	60
Figure IV. 23: types des onze modulations	71
Figure IV.24:Architecture conceptuelle du modèle hybride CNN–Transformer–BiGRU–Attention pour la classification automatique des modulations.....	77
Figure IV.25: Évolution des performances du modèle durant l'apprentissage	81
Figure IV.26:Évolution de l'accuracy du modèle en fonction du rapport signal sur bruit (SNR). 82	
Figure IV.27:Matrice de confusion du modèle proposé.....	83
Figure IV.28:Rapport de classification du modèle proposé	85

Liste des tableaux :

Tableau II.1:ML vs DL	43
Tableau II. 2:Avantages et in Inconvénients d'apprentissage profond	45
Tableau IV.3:Comparaison des performances des modèles de classification	86

Liste d'abréviations :

AM: Amplitude Modulation (Modulation d'amplitude)

AM-DSB: Amplitude Modulation Double Sideband (Modulation d'amplitude à double bande latérale)

AM-SSB: Amplitude Modulation Single Sideband (Modulation d'amplitude à bande latérale unique)

AMC: Automatic Modulation Classification (Classification automatique des modulations)

AMC-Net : Automatic Modulation Classification Network

ANN: Artificial Neural Network (Réseau de neurones artificiels)

ASK: Amplitude Shift Keying (Modulation par déplacement d'amplitude)

AWGN: Additive White Gaussian Noise (Bruit blanc gaussien additif)

BER: Bit Error Rate (Taux d'Erreur Binaire - TEB)

BFSK: Binary Frequency Shift Keying (Modulation par déplacement de fréquence binaire)

BiGRU: Bidirectional Gated Recurrent Unit (Unité récurrente intégrée bidirectionnelle)

BNN: Biological Neural Network (Réseau de neurones biologiques)

BPSK: Binary Phase Shift Keying (Modulation par déplacement de phase binaire)

CDSCNN: Coordinate Depthwise Separable Convolutional Neural Network

CNN: Convolutional Neural Network (Réseau de neurones convolutionnels)

CPFSK: Continuous Phase Frequency Shift Keying (Modulation par déplacement de fréquence à phase continue)

CBADNN : Convolutional Block Attention Deep Neural Network

DL: Deep Learning (Apprentissage profond)

FiLM: Feature-wise Linear Modulation (Modulation linéaire des caractéristiques)

FM: Frequency Modulation (Modulation de fréquence)

FSK: Frequency Shift Keying (Modulation par déplacement de fréquence)

FEA-T : Feature Enhancement Attention Transformer

GFSK: Gaussian Frequency Shift Keying (Modulation par déplacement de fréquence gaussienne)

GNN: Graph Neural Network (Réseau de neurones sur graphes)

I/Q: In-phase / Quadrature (Composantes En phase / En quadrature du signal)

IA: Intelligence Artificielle

LSTM: Long Short-Term Memory (Réseau à mémoire court et long terme)

MCFormer: Multi-scale Cross Attention Former

MCLDNN : Multi-Channel Learning Deep Neural Network

ML: Machine Learning (Apprentissage automatique)

MLP: Multi-Layer Perceptron (Perceptron multicouche)

OOK: On-Off Keying (Modulation par tout ou rien)

PAM: Pulse Amplitude Modulation (Modulation d'amplitude pulsée)

PM: Phase Modulation (Modulation de phase)

PSK: Phase Shift Keying (Modulation par déplacement de phase)

PET-CGDNN : Phase-Enhanced Temporal Convolutional Gated Deep Neural Network

QAM: Quadrature Amplitude Modulation (Modulation d'amplitude en quadrature)

QPSK: Quadrature Phase Shift Keying (Modulation par déplacement de phase en quadrature)

RF: Radio Frequency (Radiofréquence)

RMS: Root Mean Square (Valeur efficace)

RNN: Recurrent Neural Network (Réseau de neurones récurrents)

SL-LSTM: Single-Layer Long Short-Term Memory (LSTM à une seule couche)

SNR: Signal-to-Noise Ratio (Rapport Signal sur Bruit - RSB)

SMTrans : Signal Modulation Transformer

WBFM: Wideband Frequency Modulation (Modulation de fréquence à large bande)

Introduction générale :

À l'ère moderne des communications sans fil et avec le développement rapide des technologies numériques, la transmission fiable de l'information est devenue essentielle pour assurer la connectivité mondiale et le bon fonctionnement des systèmes de communication intelligents. La modulation numérique constitue une technique fondamentale dans ces systèmes, car elle permet de transmettre les données en modifiant certains paramètres de l'onde porteuse comme l'amplitude, la fréquence ou la phase. Les différents schémas de modulation numérique jouent ainsi un rôle important dans l'amélioration de l'efficacité spectrale, de la qualité de transmission et de l'utilisation optimale de la bande passante [1].

Avec l'augmentation du nombre d'appareils connectés et la complexité des environnements sans fil, la classification automatique des modulations (AMC) est devenue un domaine de recherche très important. Cette technique permet au récepteur d'identifier automatiquement le type de modulation utilisé par le signal reçu sans information préalable. Elle est utilisée dans plusieurs domaines comme la radio cognitive, la surveillance du spectre, les communications militaires et les réseaux intelligents. Cependant, les méthodes classiques basées sur l'extraction manuelle des caractéristiques présentent souvent des limites lorsque le signal est affecté par le bruit, le fading ou les interférences [2].

Pour résoudre ces problèmes, l'intelligence artificielle (IA), le Machine Learning (ML) et le Deep Learning (DL) ont apporté des solutions très efficaces. Contrairement aux méthodes traditionnelles, ces approches permettent d'apprendre automatiquement les caractéristiques importantes directement à partir des signaux bruts, notamment les signaux I/Q et les diagrammes de constellation [3].

Dans ce travail, nous avons utilisé la base de données RadioML 2016.10a, qui contient plusieurs types de modulations numériques et analogiques sous différents niveaux de SNR. Cette base est largement utilisée pour évaluer les modèles de Deep Learning dans la classification automatique des modulations.

L'objectif principal de ce mémoire est d'étudier les techniques de Deep Learning appliquées à l'AMC, d'analyser leurs performances et leur robustesse dans différentes conditions de transmission, ainsi que d'identifier leurs avantages et leurs limites.

Ce mémoire est organisé en quatre chapitres :

Le premier chapitre, Modulation Numérique, présente les bases des systèmes de communication sans fil, les principales techniques de modulation numérique ainsi que les perturbations du canal de transmission comme le bruit, l'atténuation et le fading.

Le deuxième chapitre traite de l'Intelligence Artificielle, du Machine Learning, du Deep Learning et des architectures des réseaux de neurones, en présentant les concepts fondamentaux de l'IA ainsi que les principales architectures neuronales utilisées dans la classification des signaux radio.

Le troisième chapitre présente l'état de l'art de la classification automatique des modulations basée sur le Deep Learning, avec les différentes représentations du signal, les architectures modernes ainsi que les techniques de régularisation utilisées pour améliorer les performances des modèles.

Le quatrième chapitre décrit la méthodologie adoptée, la base de données utilisée, l'architecture du modèle proposé basée sur les réseaux neuronaux profonds, ainsi que les résultats expérimentaux obtenus et leur analyse.

.

Chapitre 01 : Généralités sur la modulation numérique.

I .1Introduction

Les systèmes de communication sont des modèles fondamentaux qui décrivent comment l'information est échangée entre deux ou plusieurs points, comme un émetteur et un récepteur. Ils retracent le parcours d'un signal, depuis sa source jusqu'à sa destination, en passant par un canal de transmission.

Avant d'être envoyé, un signal doit être préparé pour ce voyage. Il subit plusieurs étapes de traitement, telles que la mise en forme, le codage et la modulation. Ces étapes sont essentielles pour que l'information puisse traverser le canal, un environnement où elle risque d'être dégradée par des perturbations comme le bruit, l'atténuation ou la distorsion, ce qui peut affecter la qualité et l'intégrité du message reçu.

Ces principes fondamentaux sont la base de tous les échanges d'information entre les individus et les appareils qui nous entourent. Ils sont à l'œuvre dans une vaste gamme de technologies, comme la téléphonie, l'Internet ou les réseaux de diffusion. Ils jouent ainsi un rôle crucial dans notre quotidien en assurant une connectivité permanente et des interactions mondiales, que ce soit pour un simple message instantané ou pour une émission de radio internationale.

Dans le cas de ce travail, la compréhension des modulations numériques est une étape préalable indispensable. En effet, notre objectif est de développer un système capable d'identifier automatiquement le type de modulation utilisé dans un signal radio reçu. Pour cela, nous devons d'abord maîtriser les caractéristiques propres à chaque technique de modulation : comment elles encodent l'information, comment elles se comportent face au bruit, et surtout, quelles signatures visibles (notamment dans les diagrammes de constellation) permettent de les distinguer.

I.2 Généralités sur la modulation

I .2.1Définition

La modulation est le processus qui consiste à faire varier les caractéristiques d'un signal porteur, telles que l'amplitude, la fréquence ou la phase, afin de coder l'information à transmettre. Cette technique permet un transfert efficace des données à travers les canaux

de communication, tout en améliorant la puissance du signal et en réduisant les interférences.

Comme illustré dans la figure 1, le signal d'information (signal utile) est combiné avec un signal porteur au niveau du modulateur, produisant ainsi un signal modulé adapté à la transmission.

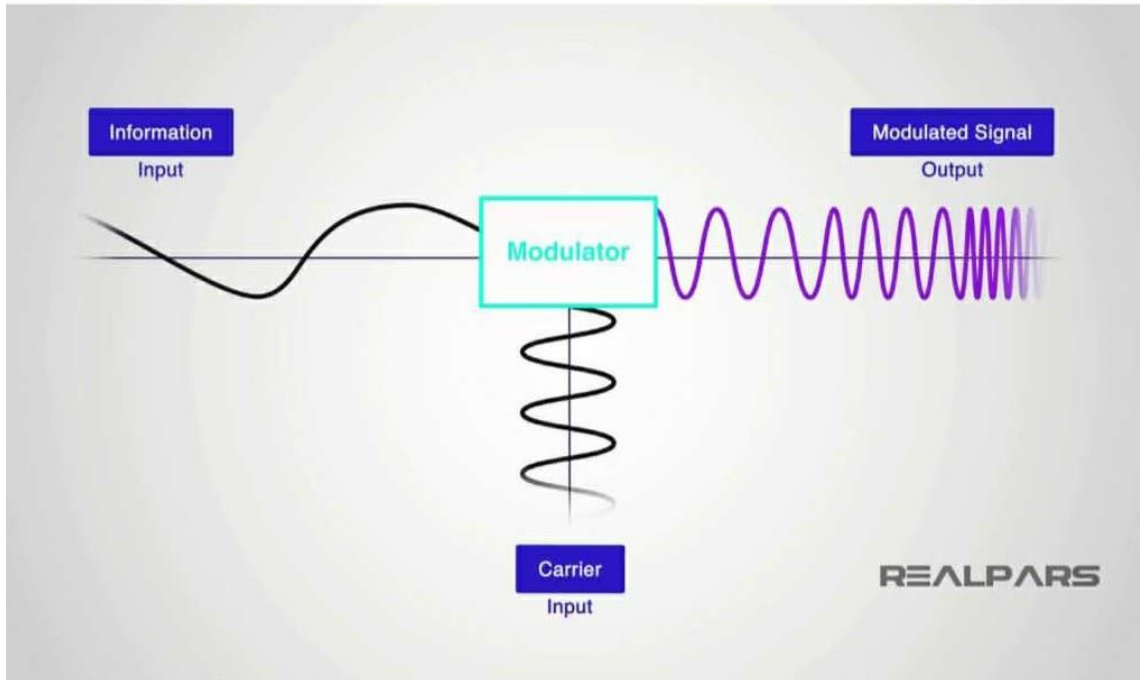


Figure I.1: le principe de la modulation.

Différents types de modulations sont utilisés selon les applications, ce qui fait de la modulation un concept fondamental en traitement du signal et en télécommunications. La modulation est essentielle pour adapter un signal à transmettre au canal de transmission choisi, qu'il soit hertzien, filaire ou optique [4].

I .2.2 La modulation analogique

La modulation analogique est un processus dans lequel le signal est modifié pour correspondre au signal d'information [5]. Les types courants de modulation analogique comprennent :

I .2.2.1 Modulation d'amplitude (AM)

La modulation d'amplitude consiste à transmettre un signal d'information en l'ajoutant à une onde de fréquence élevée, appelée onde porteuse. Le signal d'information peut être de toute sorte en fonction des données qu'il transporte, comme la voix ou les données numériques. La fréquence de ce signal, qui s'appelle aussi signal en bande de base, est très faible.

La fréquence d'un signal est directement reliée à son énergie. Ainsi, quand la fréquence du signal est très faible, il perd de l'intensité après avoir été transmis sur une certaine distance. Pour éviter cette baisse de puissance, le signal d'information est placé sur une onde porteuse à haute fréquence.

Dans le cas de la modulation d'amplitude, c'est l'intensité de l'onde porteuse qui change en fonction de l'intensité du signal d'information. C'est pour cette raison que cette technique porte le nom de modulation d'amplitude. Il faut souligner que la fréquence et la phase de la porteuse restent les mêmes pendant la modulation. Cette relation peut être formalisée mathématiquement par l'expression suivante :

$$S(t) = A_c[1 + K_a m(t)] \cos(2\pi f_c t) \quad (1)$$

Où :

A_c : amplitude de la porteuse.

K_a : indice de modulation.

$m(t)$: signal modulant (signal d'information).

f_c : fréquence de la porteuse.

Le principal problème de la modulation d'amplitude est qu'elle est peu efficace et que le signal transmis est de mauvaise qualité. Le signal modulé ne ressemble pas exactement au signal original car la qualité du signal s'est détériorée. De plus, les systèmes AM ont une immunité au bruit assez limitée [6].

I .2.2.2 Modulation de fréquence (FM)

La modulation de fréquence est un type de modulation dans lequel la fréquence du signal porteur est modifiée en fonction de l'amplitude du signal d'information ou du signal de bande de base, tout en maintenant l'amplitude du signal porteur constante.

Le modulateur de fréquence effectue la modulation en utilisant le signal porteur provenant du générateur de fréquence radio et le signal d'information provenant de la source d'information. Une fois modulé, le signal est envoyé à l'amplificateur RF, qui compense les pertes éventuelles du signal lors de la transmission. Mathématiquement, ce principe de modulation peut être exprimé par l'équation suivante :

$$S(t) = A_c \cos(2\pi f_c t + 2\pi k_f \int_0^1 m(\tau) d\tau + \varphi_0) \quad (2)$$

Où :

$\int_0^1 m(\tau) d\tau$: intégrale du signal modulant.

φ_0 : phase initiale.

Le principal avantage de l'utilisation de la modulation de fréquence est que la qualité du signal transmis reste élevée et ne se détériore pas. Cependant, le système de modulation de fréquence est complexe à concevoir, ce qui rend le coût du système relativement élevé.

De plus, le système de modulation de fréquence est immunisé contre les distorsions de bruit, ce qui fait que l'effet du bruit sur le signal modulé en fréquence est très faible et peut être négligé [7].

I .2.2.3 Modulation de phase (PM)

La modulation de phase est une méthode de modulation du signal dans laquelle la phase du signal porteur haute fréquence change en fonction de l'amplitude du signal d'entrée. En d'autres termes, au lieu de modifier l'amplitude du signal comme dans les méthodes précédentes, cette technique modifie la position de la phase du signal pour transmettre l'information.

La modulation de phase est couramment utilisée dans les transmissions radio et télévisuelles, ainsi que dans la transmission de données numériques via les modems, car elle fonctionne efficacement même en présence de bruit important. L'équation du signal modulé $S(t)$ selon la modulation PM est la suivante :

$$S(t) = A_c \cos(2\pi f_c t + k_p m(t) + \varphi_0) \quad (3)$$

Où :

k_p : sensibilité de phase (indice de modulation PM).

Dans la modulation de phase, la phase de l'onde porteuse varie en fonction du signal d'entrée afin de transmettre l'information à travers différents canaux. La modulation de phase dépend également de la valeur instantanée du signal d'entrée pour ajuster la phase de l'onde porteuse, générant ainsi un signal PM [8].

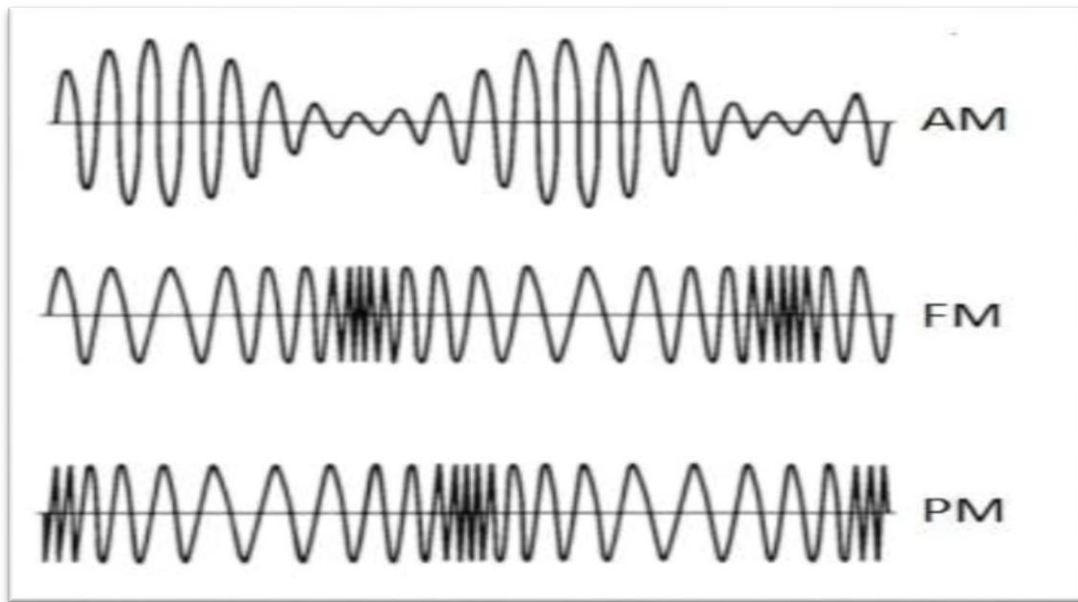


Figure I.1: représentation des différentes techniques de modulation.

I .3Modulation numérique

La modulation numérique est la méthode utilisée pour transformer des informations numériques (bits) en un signal transmis en modifiant ses caractéristiques comme l'amplitude, la phase ou la fréquence. Ce processus d'encodage est important pour définir la bande passante du signal et son aptitude à résister aux perturbations comme le bruit et les interférences dans le canal de communication [9].

I .3.1Modulation ASK (Amplitude Shift Keying)

Dans les communications numériques, diverses méthodes appelées techniques de modulation sont utilisées pour envoyer des données ou des messages d'un émetteur vers un récepteur à travers un canal de communication. L'une de ces méthodes est la modulation par déplacement d'amplitude (Amplitude Shift Keying, ASK).

La modulation ASK consiste à envoyer des signaux numériques en variant l'amplitude de la porteuse. Dans ce schéma, une porteuse de forte amplitude représente un bit binaire '1', tandis qu'une porteuse de faible amplitude représente un bit binaire '0' (ou parfois une absence de signal, appelée OOK - On-Off Keying) [10].

I .3.1.1 La modulation OOK (On-Off Keying)

Le principe de la modulation On-Off Keying (OOK) repose sur la représentation des données binaires par la présence ou l'absence de l'onde porteuse. Lors de la transmission du bit 1, l'onde porteuse est émise avec une amplitude constante et prédéfinie, ce qui indique la présence du signal. En revanche, lors de la transmission du bit 0, l'onde porteuse disparaît complètement, ce qui signifie qu'aucun signal n'est transmis pendant cet intervalle de temps.

La fréquence et la phase de l'onde porteuse restent inchangées pendant le processus de transmission, car seul l'état d'activation ou de désactivation de la porteuse est contrôlé en fonction du flux de bits numériques.

La modulation OOK est considérée comme la forme la plus simple de modulation numérique. Elle est également classée comme un cas particulier de la modulation par déplacement d'amplitude binaire (Binary Amplitude Shift Keying – BASK) [11].

I.3.1.2 Modulation M-ASK

La modulation par déplacement d'amplitude à plusieurs niveaux (M-ary Amplitude Shift Keying, M-ASK) est une technique de communication avancée par rapport à la modulation d'amplitude traditionnelle (ASK), car elle utilise plusieurs niveaux d'amplitude différents au lieu de seulement deux niveaux. Chaque niveau d'amplitude correspond à un ensemble spécifique de bits. Par exemple, dans la modulation à quatre niveaux (4-ary ASK), un symbole peut représenter deux bits d'information, tandis que la modulation à huit niveaux (8-ary ASK) permet de représenter trois bits par symbole. L'un des principaux avantages de cette modulation est l'amélioration du débit de transmission tout en conservant la même largeur de bande. Cela en fait une solution précieuse pour les systèmes nécessitant un transfert de données rapide. Dans ce type de modulation, plusieurs niveaux d'amplitude sont utilisés pour transmettre plusieurs bits par symbole. En associant chaque symbole à une amplitude bien définie, cela permet d'optimiser le codage des données [11]. La figure ci-dessous illustre le principe de la modulation M-ASK en présentant les différents niveaux d'amplitude du signal modulé ainsi que leur correspondance avec le signal porteur et le signal d'entrée :

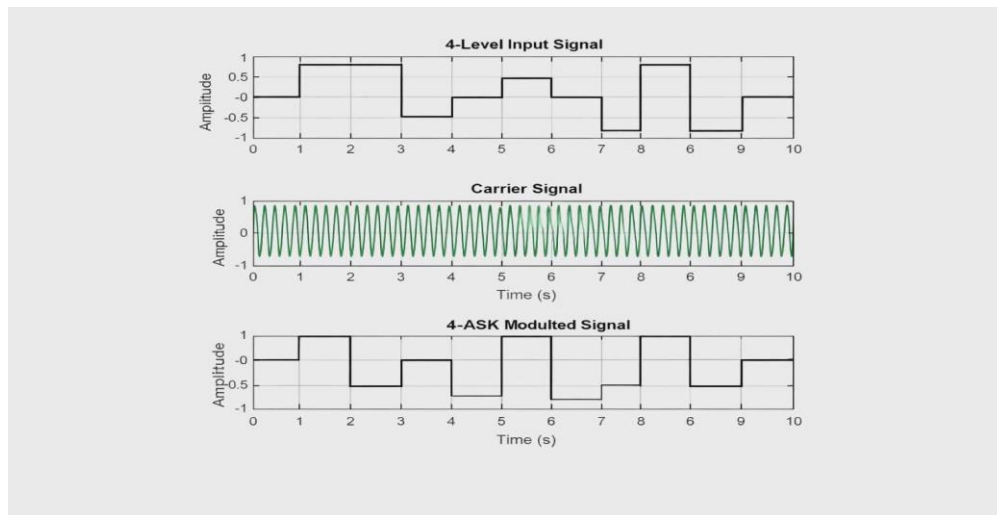


Figure I. 3 : Représentation de modulation 4-ASK

I .3.1.3 Avantages de la modulation ASK

- ❖ **Grande efficacité en bande passante** : Elle utilise moins d'espace de bande passante par rapport à d'autres méthodes de modulation.
- ❖ **Facilité de conception du récepteur** : Décoder les signaux ASK est facile et demande juste un détecteur d'enveloppe.
- ❖ **Peu coûteux et simple à intégrer** : L'ASK est une méthode de modulation relativement simple à mettre en place.
- ❖ **Efficacité énergétique** : Les amplificateurs de puissance fonctionnent dans une zone plus efficace, réduisant la consommation d'énergie.
- ❖ **Compatibilité avec les systèmes numériques** : L'ASK fonctionne bien avec les systèmes numériques, notamment pour la transmission sur fibres optiques [12].

I .3.1.4 Inconvénients de la modulation ASK

- ❖ **Sensibilité au bruit** : L'ASK est très affectée par le bruit, car celui-ci perturbe l'amplitude du signal.
- ❖ **Efficacité spectrale plus faible** : Bien qu'elle soit bonne pour la bande passante, l'efficacité spectrale est souvent moins bonne que d'autres techniques comme la QAM
- ❖ **Limites** pour les applications à haut débit : Difficultés pour atteindre des vitesses très rapides.
- ❖ **Moins robuste face à l'évanouissement multi-trajets** : L'ASK est moins efficace dans des conditions de transmission complexes [12].

I .3.1.5 Diagramme de constellation de la modulation 4-ASK

La figure montre un diagramme de constellation représentant la modulation 4-ASK. Dans ce diagramme, quatre points de signal distincts sont alignés sur l'axe I (In-phase), sans variation sur l'axe Q (Quadrature). Cela indique que le schéma de modulation est unidimensionnel, car seule l'amplitude de la composante en phase est modulée.

Chaque point du diagramme correspond à un symbole unique, et chaque symbole représente 2 bits d'information. Étant donné que les symboles diffèrent uniquement par leur amplitude, la modulation 4-ASK reste sensible au bruit et aux distorsions d'amplitude.

Toutefois, elle est généralement moins sensible que la 8-ASK, car l'espacement entre les niveaux d'amplitude est plus important. La figure ci-dessous présente le diagramme de constellation de la modulation 4-ASK, mettant en évidence les différents niveaux d'amplitude associés à chaque symbole :

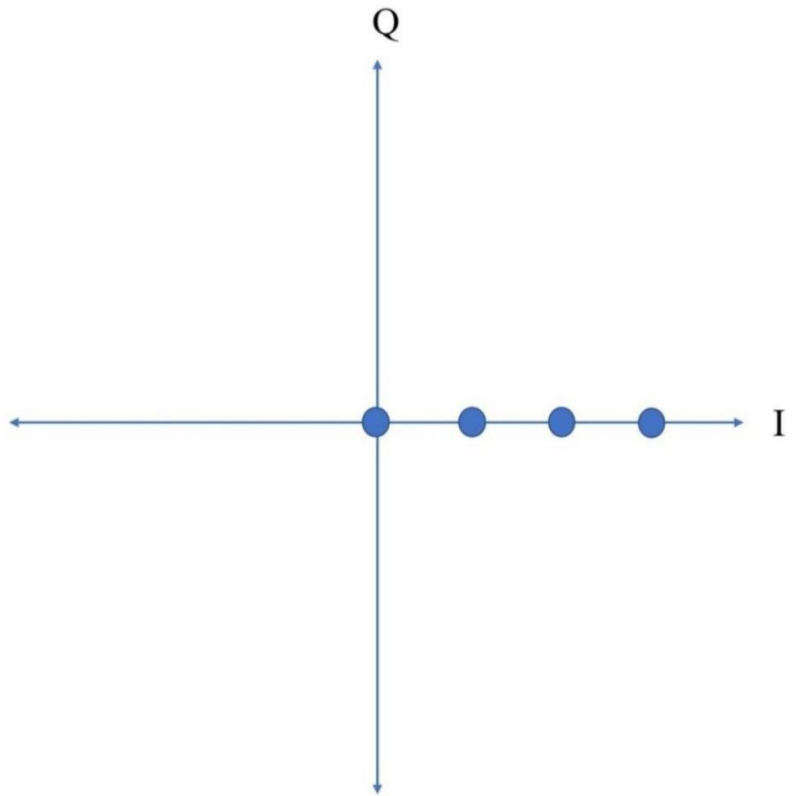


Figure I.2: Diagramme de constellation 4-ASK.

I .3.2 Modulation PSK (Phase Shift Keying)

La modulation par changement de phase, ou PSK, est une méthode numérique qui permet d'envoyer des informations en changeant la phase d'un signal porteur. Dans la modulation numérique, les informations sont représentées par plusieurs variations de phase différentes. En fait, la PSK utilise un nombre défini de décalages de phase, où chaque phase correspond à un symbole binaire spécifique [13].

I .3.2.1 Binary Phase Shift Keying (BPSK)

Le BPSK (Binary Phase Shift Keying) est une technique de modulation numérique dans laquelle la phase de la porteuse sinusoïdale prend deux valeurs possibles : 0° et 180° . Chaque changement de phase représente un bit d'information (0 ou 1). Le BPSK est la

forme la plus simple de modulation par déplacement de phase et peut être interprété comme une modulation à porteuse supprimée, similaire à la DSBSC, mais avec une information codée uniquement dans la phase du signal [13]. La figure ci-dessous illustre le processus de modulation 16-PSK en présentant le signal d'entrée, le signal porteur ainsi que le signal modulé résultant :

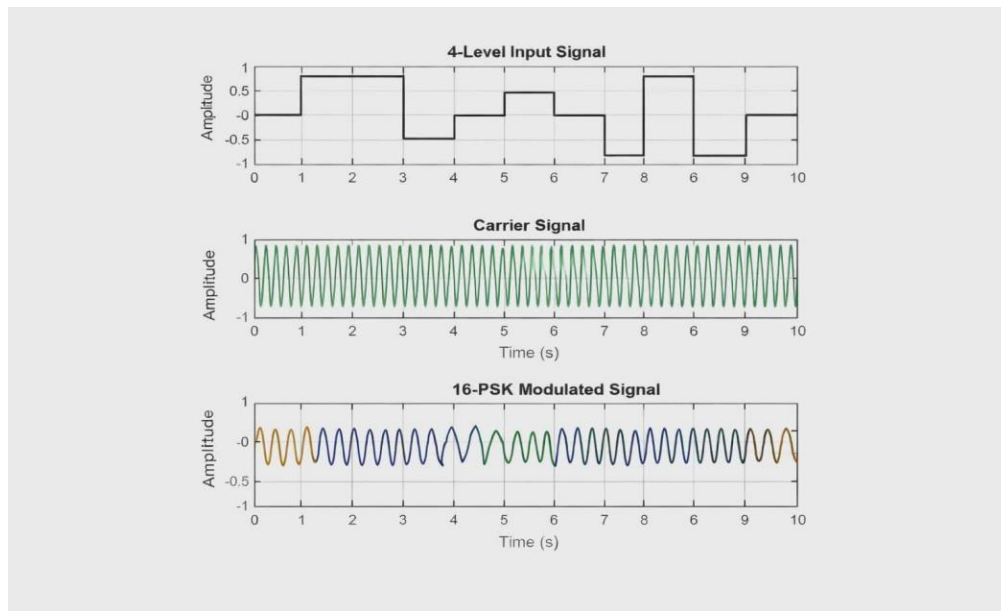


Figure I.3: Représentation visuelle de 16-PSK.

I .3.2.2 Avantages de la modulation PSK

- ❖ **Meilleure transmission des données en RF** : La PSK transmet les informations à l'aide de signaux radiofréquence plus efficacement que d'autres types de modulation.
- ❖ **Moins d'erreurs** : Elle génère moins d'erreurs que la modulation ASK pour une même largeur de bande.
- ❖ **Débits de données plus élevés** : On peut atteindre des vitesses de transmission plus rapides avec des versions avancées comme la QPSK (2 bits par symbole) [14].

I .3.2.3 Inconvénients de la modulation PSK

- ❖ **Efficacité limitée de la bande passante** : La PSK n'utilise pas toujours la bande passante de manière optimale comparée à certaines techniques plus avancées
- ❖ **Décodage complexe** : Le processus de décodage nécessite la détection précise des états de phase
- ❖ **Problèmes de synchronisation de phase** : L'émetteur et le récepteur doivent être parfaitement synchronisés
- ❖ **Sensibilité aux variations de phase** : Les méthodes PSK plus complexes sont plus sensibles aux perturbations de phase [14].

I .3.2.4 Diagramme de constellation de la modulation 16-PSK

Ce diagramme de constellation représente la modulation 16-PSK. En 16-PSK, la modulation est réalisée en faisant varier la phase du signal porteur tout en maintenant une amplitude constante. Les 16 symboles sont répartis uniformément sur un cercle centré à l'origine. Chaque point représente un déphasage unique de la porteuse, espacé de $22,5^\circ$ ($360^\circ/16$), couvrant ainsi la totalité des 360° de phase. La figure ci-dessous présente le diagramme de constellation de la modulation 16-PSK, où les symboles sont répartis uniformément sur un cercle en fonction de leurs différentes phases :

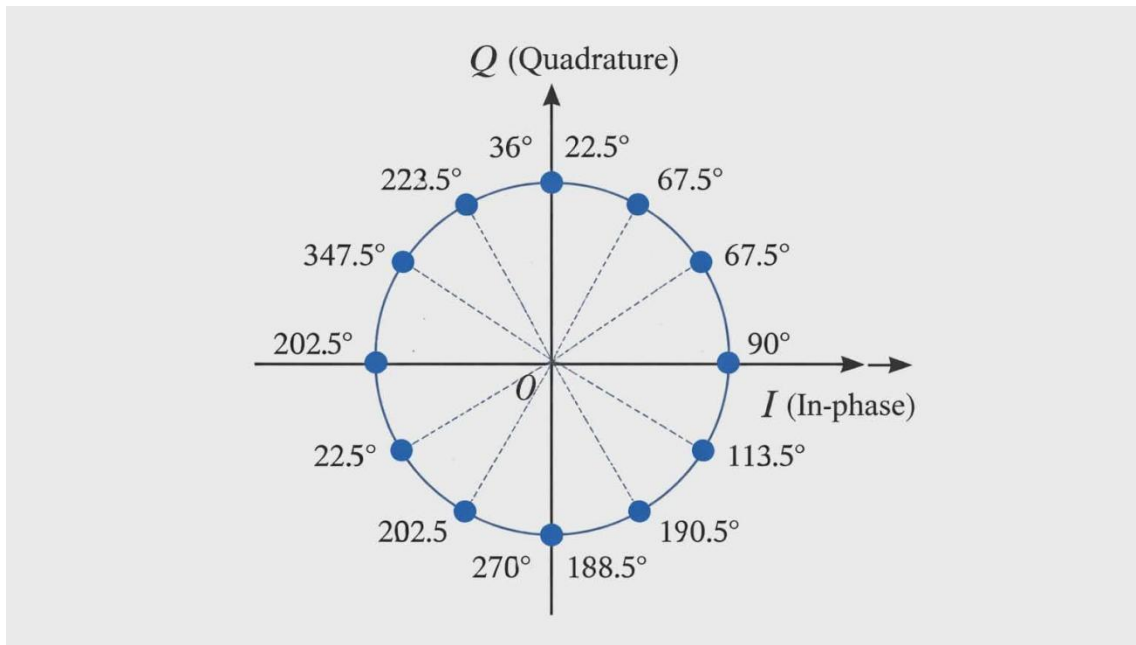


Figure I.4: Diagramme de constellation de 16-PSK.

I .3.3 La modulation FSK (Frequency Shift Keying)

La modulation par déplacement de fréquence (FSK) est une technique de modulation numérique dans laquelle l'information binaire est transmise en faisant varier la fréquence de la porteuse entre plusieurs valeurs discrètes, tandis que son amplitude et sa phase restent constantes. Chaque symbole numérique est ainsi représenté par une fréquence particulière appartenant à un ensemble prédéfini [15].

I .3.3.1 FSK binaire (BFSK)

La FSK binaire correspond au cas particulier où deux fréquences distinctes sont utilisées pour représenter les deux états binaires possibles. Une fréquence f_1 est associée au bit 0 tandis que la fréquence f_2 correspond au bit 1 [16].

I .3.3.2 Modulation M-FSK

La modulation M-FSK constitue une extension de la BFSK dans laquelle le nombre de fréquences utilisées est supérieur à deux. Dans ce cas, chaque symbole est représenté par une fréquence choisie parmi M fréquences possibles, permettant de transmettre plusieurs bits par symbole [17]. La figure suivante présente une représentation visuelle de la modulation 4-FSK, illustrant le signal d'entrée, les fréquences porteuses associées et le signal modulé obtenu :

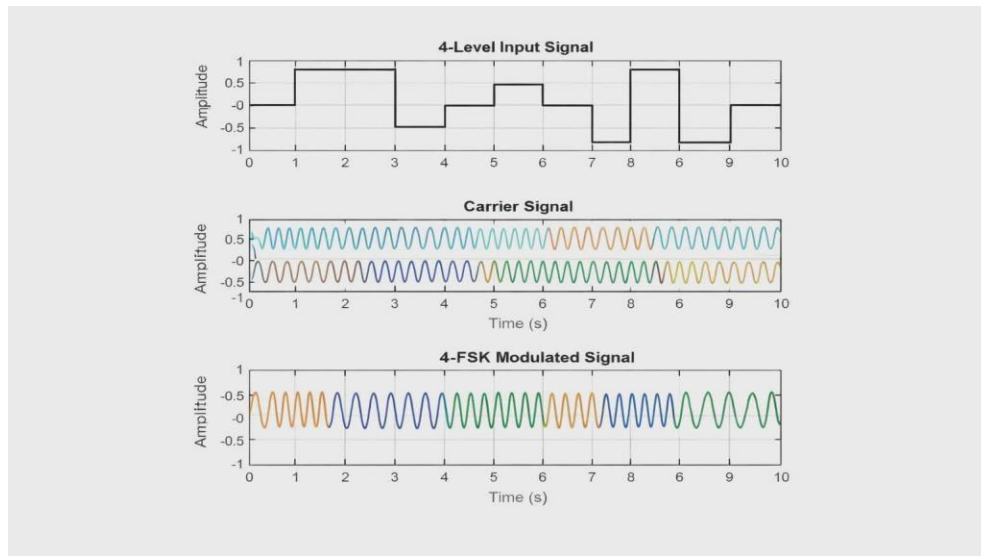


Figure I.5: Représentation visuelle de 4-FSK.

I .3.3.3 Avantages de la modulation FSK

- ❖ **Probabilité d'erreur plus faible** : Meilleure performance en termes de taux d'erreur.
- ❖ Rapport signal sur bruit (SNR) élevé
- ❖ **Bonne résistance au bruit** : Grâce à son amplitude constante, elle est fiable même lorsque le signal s'affaiblit.
- ❖ **Mise en œuvre facile** : Pour les applications utilisant de faibles débits de données [18].

I .3.3.4 Inconvénients de la modulation FSK

- ❖ **Utilisation plus large de bande passante** : Par rapport à ASK et PSK, l'efficacité spectrale est moindre.
- ❖ **Taux d'erreur binaire (BER)** : moins bon dans les canaux AWGN, comparé à la modulation PSK [18].

I .3.3.5 Diagramme de constellation de la modulation 4-FSK

Ce diagramme de constellation représente la modulation 4-FSK. Chaque point f_1 à f_4 correspond à une fréquence porteuse différente. Contrairement aux schémas basés sur

l'amplitude ou la phase, les symboles FSK ne sont pas confinés à une structure circulaire ou en grille. Dans le cas de la 4-FSK, chaque symbole transporte 2 bits. La figure ci-dessous présente le diagramme de constellation de la modulation 4-FSK :

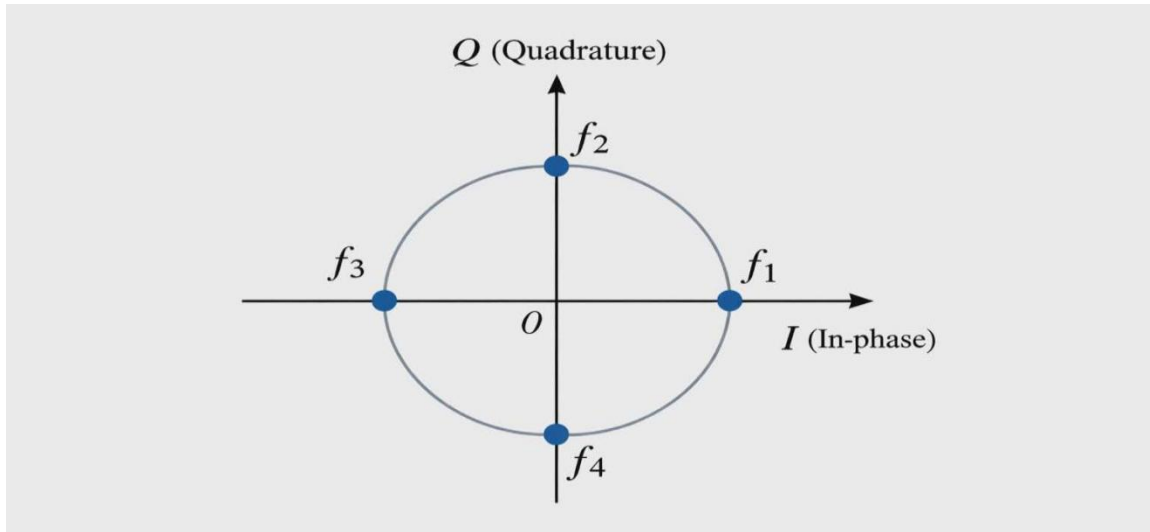


Figure I.6: Diagramme de constellation 4-FSK.

I .3.4 Modulation QAM (Quadrature Amplitude Modulation)

La modulation d'amplitude quadratique (QAM) est une méthode moderne et efficace de transmission de données sur ondes porteuses. Cette technique combine simultanément deux propriétés fondamentales du signal : l'amplitude et la phase. Autrement dit, au lieu de modifier une seule propriété de l'onde porteuse, la QAM module simultanément l'amplitude et la phase pour encoder les données numériques.

Pour comprendre le fonctionnement de la QAM, on peut l'imaginer comme la combinaison de deux signaux différents de même fréquence, mais déphasés de 90 degrés

- ❖ **La composante en phase (I) :** l'amplitude est modulée par une onde cosinus.
- ❖ **La composante en quadrature (Q) :** l'amplitude est modulée par une onde sinusoïdale.

La technologie QAM se divise en plusieurs modèles selon le nombre de symboles : 16-QAM, 64-QAM et 256-QAM. Plus le nombre de symboles est élevé, plus le système est capable de représenter un grand nombre de bits par symbole, ce qui se traduit par un débit de données plus important [19].

I .3.4.1 Avantages de la modulation QAM

- ❖ **Utilisation intelligente de la bande passante** : C'est le plus grand avantage de la QAM
- ❖ **16-QAM** : 4 bits par symbole
- ❖ **64-QAM** : 6 bits par symbole
- ❖ **256-QAM** : 8 bits par symbole
- ❖ **Grande flexibilité** : Adaptable à différents besoins de débit
- ❖ **Largement utilisée** : Wi-Fi, 4G, 5G, communications par satellite, modems [20].

I .3.4.2 Inconvénients de la modulation QAM

- ❖ **Vulnérabilité au bruit** : Les symboles étant plus rapprochés sur le diagramme de constellation, le signal est plus susceptible d'être perturbé.
- ❖ **Récepteurs complexes** : Plus complexes que ceux d'autres types de modulation.
- ❖ **Nécessité de linéarité** : La QAM utilise l'amplitude du signal, nécessitant des amplificateurs linéaires qui consomment plus d'énergie.
- ❖ **Exigence de SNR élevée** : Pour les niveaux élevés de QAM, un bon rapport signal/bruit est indispensable [20].

I .3.4.3 Diagramme de constellation de la modulation 16-QAM

La constellation illustre la modulation 16-QAM. Les 16 points sont disposés sur une grille carrée autour de l'origine. Chaque point représente une combinaison unique de valeurs de composantes I et Q. Contrairement à la PSK, les amplitudes ne sont pas toutes identiques, ce qui permet de transmettre plus de bits par symbole. La figure ci-dessous présente le diagramme de constellation de la modulation 16-QAM :

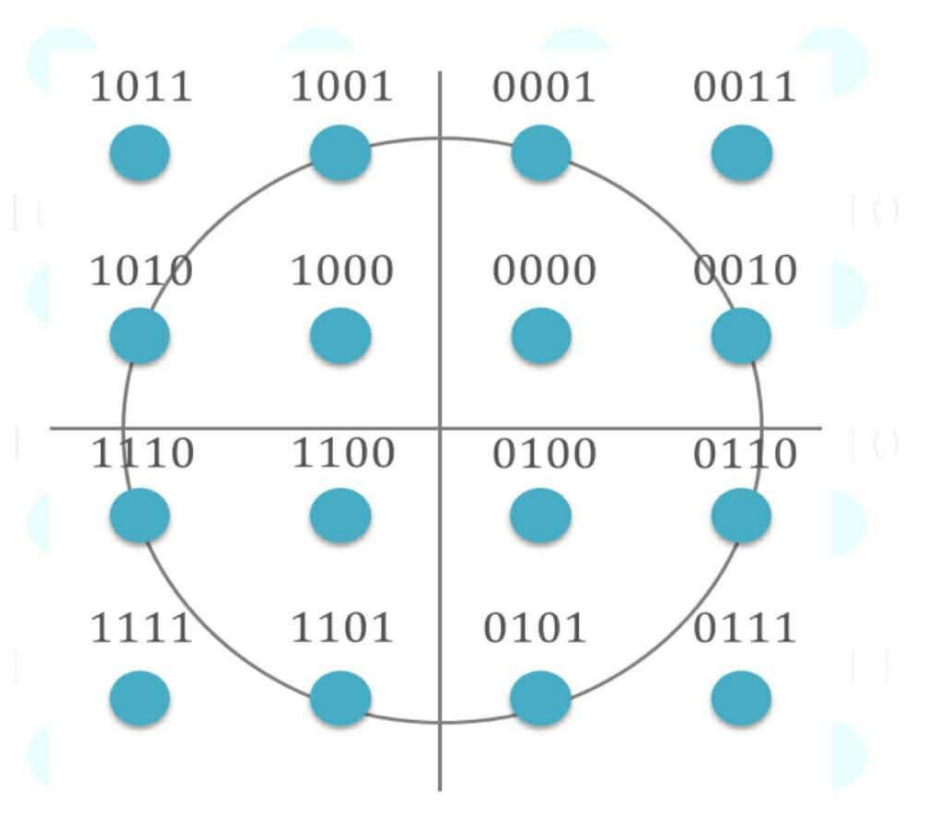


Figure I.7: Diagramme de constellation 16- QAM.

I .4 Canal de transmission et perturbations

Cette image illustre le processus de transmission d'un signal dans un canal de communication. Le signal est d'abord généré par une source d'information, puis envoyé par un émetteur sous forme de signal $s(t)$. Lors de sa propagation dans le canal, il subit plusieurs perturbations telles que l'atténuation, le fading, la distorsion, les interférences et le bruit (notamment le bruit AWGN). Le signal reçu $r(t)$ au niveau du récepteur est donc une version dégradée du signal initial, influencée par la réponse du canal $h(t)$ et le bruit $n(t)$, ce qui peut s'exprimer par la relation

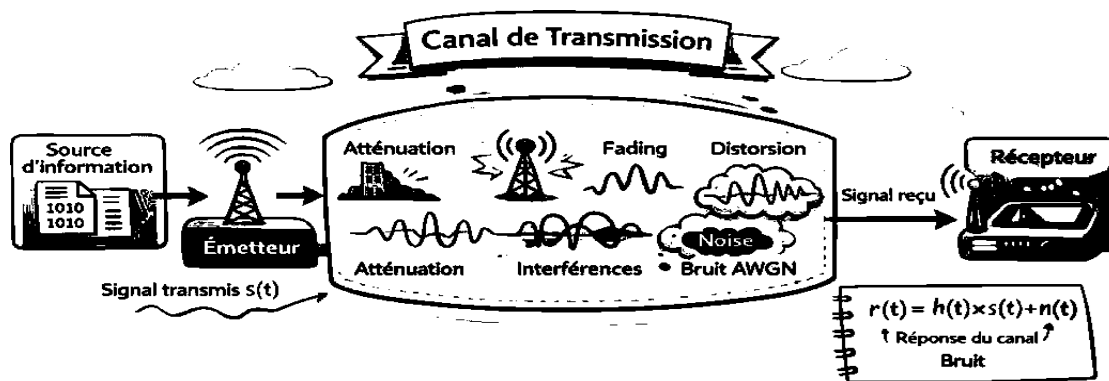


Figure I.8: Canal de transmission et perturbation.

Pendant sa transmission, le signal subit diverses interférences qui affectent directement l'efficacité du système de communication. Les principaux types de dégradations constatées sont le bruit AWGN, l'affaiblissement, la décoloration et la distorsion, exposés ci-dessous.

I .4.1 Le bruit blanc gaussien additif (AWGN)

Le bruit blanc gaussien additif (AWGN) est un modèle statistique fondamental qui joue un rôle crucial dans la compréhension de l'influence du bruit sur les signaux dans les systèmes de communication.

Ce type de bruit est constitué de valeurs aléatoires et non corrélées, et son amplitude suit une loi normale (gaussienne). Grâce à ces propriétés, ce modèle est essentiel pour l'étude et la simulation des systèmes de communication théoriques [21].

I .4.2 L'atténuation du signal

L'atténuation désigne la baisse progressive de l'intensité d'un signal qui se produit quand il se propage à travers un canal de transmission, que ce soit un câble en cuivre, une fibre optique ou l'air.

Ce phénomène existe dans tout système de communication. Du point de vue physique, l'atténuation est en grande partie causée par les pertes dues à la résistance des conducteurs,

aux phénomènes d'absorption et de diffusion dans les matériaux, ainsi qu'aux effets associés à la fréquence du signal envoyé. Dans les systèmes sans fil, cela dépend aussi de la distance, des obstacles et des conditions météorologiques.

L'atténuation est habituellement mesurée en décibels (dB) et se calcule de la manière suivante :

$$A = 10 \log_{10} \left(\frac{P_{\text{émise}}}{P_{\text{reçue}}} \right) \quad (4)$$

P_e : puissance émise.

P_r : puissance reçue.

Pour lutter contre l'atténuation, plusieurs solutions techniques existent :

- ❖ Utilisation d'amplificateurs ou de répéteurs.
- ❖ Choix d'un support de transmission à faible perte.
- ❖ Amélioration de la structure du réseau.
- ❖ Utilisation de la fibre optique pour les longues distances [22].

I .4.3 Le fading (évanouissement)

Le fading désigne une variation non cumulative des canaux de transmission radio, entraînant des fluctuations aléatoires de la puissance du signal dues à la propagation de plusieurs ondes, à la distance et aux obstacles. Le fading peut être classé en trois types principaux :

Fading rapide : Variations rapides du signal causées par la propagation de plusieurs ondes et le mouvement.

Fading lent (ombrage) : Variations causées par des obstacles bloquant la ligne de visée.

Fading sélectif en fréquence : Interférence destructive où certaines fréquences sont atténuées plus que d'autres [23].

I .4.4 La distorsion

La distorsion dans les signaux de communication désigne toute modification involontaire de la forme ou de la structure du signal lors de sa transmission. Contrairement à l'atténuation, qui se limite à réduire l'intensité du signal, la distorsion peut affecter l'amplitude, la fréquence et la phase du signal [24].

I .5 Performances de la modulation numérique

I .5.1 Rapport Signal sur Bruit (SNR)

Le rapport signal sur bruit (SNR) est une mesure importante qui permet de comparer la puissance du signal souhaité à celle du bruit de fond. Il est souvent exprimé en décibels (dB) :

$$SNR = 10 \log 10 \frac{P_s}{P_n} \quad (5)$$

P_s : Puissance de signal utile.

P_n : Puissance de bruit (noise).

Plus la valeur du SNR est élevée, plus le signal est clair et fiable. Le SNR joue un rôle essentiel dans l'évaluation de la clarté, de la robustesse et de l'efficacité des systèmes de communication [25].

I .5.2 Taux d'Erreur Binaire (BER)

Le Taux d'Erreur Binaire (BER) est une mesure utilisée pour évaluer le taux d'erreurs de transmission de bits dans les systèmes de communication numérique. Il représente le rapport entre le nombre de bits erronés reçus et le nombre total de bits transmis :

$$BER = \frac{N_{bits\ erronée}}{N_{bits\ transmis}} \quad (6)$$

Un faible BER indique une transmission fiable, tandis qu'un BER élevé reflète une mauvaise qualité de transmission.

Facteurs influençant le BER :

- Bruit du canal.
- Puissance du signal.
- Distance de transmission.
- Schéma de modulation.
- Distorsion du canal.
- Interférence inter-symboles.
- Débit de données [26].

I .5.3 Efficacité spectrale

L'efficacité spectrale mesure combien de bits peuvent être envoyés en une seconde sur une certaine largeur de bande. Elle s'exprime en bits par seconde par hertz (bits/s/Hz) :

$$\eta = \frac{R_b}{B} \quad (7)$$

η : efficacité spectrale.

R_b : débit binaire.

B : bande passante.

Plus l'efficacité spectrale est grande, plus un système est capable de transmettre de données même lorsque le spectre est limité [27].

I .6 Conclusion

Ce premier chapitre a posé les bases théoriques nécessaires à la compréhension des signaux que nous chercherons à classer dans le cadre de ce PFE. Nous avons exploré les principales techniques de modulation numérique (ASK, PSK, FSK, QAM), leurs caractéristiques, leurs avantages et leurs limites face aux perturbations du canal de transmission.

Nous avons également vu comment les perturbations comme le bruit AWGN, l'atténuation et le fading dégradent la qualité des signaux, et comment des indicateurs comme le SNR, le BER et l'efficacité spectrale permettent d'évaluer les performances des systèmes de communication.

Chapitre 02 : les fondements de l'intelligence artificielle.

II .1 Introduction

L'intelligence artificielle est devenue un domaine très important dans notre vie quotidienne. Elle est utilisée dans plusieurs secteurs comme la santé, l'éducation, l'industrie et les télécommunications. Grâce au développement des technologies et des ordinateurs, les systèmes intelligents sont capables d'apprendre à partir des données et de résoudre différents problèmes automatiquement.

Dans ce travail, nous allons présenter les bases de l'apprentissage automatique (Machine Learning) et de l'apprentissage profond (Deep Learning). Nous allons également expliquer le fonctionnement des réseaux de neurones artificiels, ainsi que leurs principales applications, avantages et limites.

II.1.2 Introduction à IA

II.1.2.1 Définition d'IA

L'intelligence artificielle (IA) peut être décrite comme la capacité d'un ordinateur ou d'un robot à effectuer des tâches qui nécessitent souvent l'intelligence humaine, telles que le raisonnement, la compréhension, la découverte de sens et l'apprentissage à partir de l'expérience. Depuis les années 1940, les chercheurs ont programmé les ordinateurs pour accomplir des tâches complexes, allant de la résolution de problèmes mathématiques à la participation à des jeux stratégiques comme les échecs, avec une grande habileté [28].

Actuellement, l'intelligence artificielle moderne permet aux systèmes et aux machines de reproduire plusieurs compétences de l'intelligence humaine, telles que la résolution de problèmes, la prise de décisions, la créativité et l'exécution de tâches de manière autonome. Les applications actuelles couvrent un large éventail, allant de la reconnaissance d'objets et de la compréhension du langage à l'apprentissage à partir de nouvelles informations et à la prise de décisions de manière indépendante, comme c'est le cas pour les véhicules

autonomes. En outre, les recherches récentes en 2024 se concentrent sur l'intelligence artificielle capable de générer du contenu nouveau, qu'il s'agisse de textes, d'images ou de vidéos [29].

II .1.2.2 Historique d'IA

1950s : En 1950, Alan Turing a publié un article intitulé *Computing Machinery and Intelligence*, dans lequel il a posé une question importante : les machines peuvent-elles penser ? Cette question continue de susciter un grand intérêt jusqu'à aujourd'hui, et ce travail est considéré comme l'une des premières références ayant contribué à l'émergence du domaine de l'intelligence artificielle [30]

1956 – Dartmouth et la naissance de « l'intelligence artificielle » : Le terme « intelligence artificielle » a été utilisé officiellement pour la première fois lors de la conférence de Dartmouth, organisée par John McCarthy. Cette rencontre historique a réuni des figures de proue de l'informatique, telles que Marvin Minsky, Claude Shannon et Nathan Rochester, dans le but d'explorer la possibilité de créer des machines capables de simuler l'intelligence humaine. Depuis, l'intelligence artificielle est devenue un domaine à part entière au sein de l'informatique [31].

1960-1970 : Au cours des années 1960 et 1970, les experts ont développé des programmes pour simuler la pensée humaine. Par exemple, le projet DENDRAL à l'université de Stanford visait à analyser des données chimiques. Des programmes de résolution de problèmes mathématiques et de jeu d'échecs ont également été créés, montrant les capacités des machines [31]. De plus, le système médical MYCIN a été conçu pour aider les médecins à diagnostiquer les maladies et à déterminer les options de traitement, ouvrant la voie à l'intégration de l'intelligence artificielle dans la vie quotidienne [32].

1970-1980 : Au cours de cette période, il y a eu un fort recul du financement de la recherche en intelligence artificielle, et cela pour plusieurs raisons. Au début, les premiers systèmes d'intelligence artificielle n'ont pas pu gérer des tâches plus complexes, ce qui a révélé les limites des principales approches utilisées à l'époque, telles que les systèmes basés sur des

règles. Ces systèmes ont également rencontré des difficultés à faire face à la complexité et à la précision requises dans les environnements réels. Deuxièmement, la puissance de calcul et les données nécessaires pour créer des systèmes d'intelligence artificielle plus avancés n'étaient pas disponibles à cette époque. Cette période est connue sous le nom de l'hiver de l'intelligence artificielle [31].

1990-2010 : Malgré le manque de financement durant la période de l'hiver de l'intelligence artificielle, l'intérêt pour ce domaine est revenu dans les années 1990 grâce au développement des ordinateurs et à l'apparition de grandes quantités de données, ainsi qu'aux progrès réalisés dans les algorithmes d'apprentissage. Parmi les événements les plus marquants de cette période figure celui de 1997, lorsque l'ordinateur Deep Blue, développé par la société IBM, a réussi à battre le champion du monde d'échecs Garry Kasparov, ce qui constitue un accomplissement important ayant démontré la capacité des systèmes d'intelligence artificielle à surpasser l'être humain dans certaines tâches complexes [33].

De 2010 à nos jours : Avec l'émergence de l'apprentissage profond et du Big Data, l'intelligence artificielle est devenue plus développée et répandue, et elle est utilisée dans de nombreux domaines, notamment : la reconnaissance d'images , la traduction automatique , les voitures autonomes, les assistants intelligents , l'analyse de données, en plus de ces domaines, l'intelligence artificielle est utilisée dans de nombreuses autres applications variées, ce qui reflète son importance croissante dans notre vie quotidienne [33].

II .1.2.2 Domaines D'applications

Avec le grand développement que connaît le domaine de l'intelligence artificielle ces dernières années, ses applications sont devenues présentes dans de nombreux domaines différents. Ci-dessous, nous présentons les principales de ces applications :

Santé : L'IA analyse les images médicales (radiographies, IRM) pour aider les médecins à détecter et suivre l'évolution des maladies [34].

Finance : Analyse prédictive et détection des fraudes

Les techniques d'intelligence artificielle aident à analyser les informations financières pour prévoir les tendances du marché et les comportements des consommateurs, et il est également possible de détecter les fraudes en identifiant les comportements anormaux dans les transactions [34].

Commerce : Personnalisation et chatbots

Les systèmes d'intelligence artificielle sont utilisés pour personnaliser les préférences des clients en étudiant leurs habitudes d'achat et en proposant des produits adaptés. De plus, les chatbots fonctionnant avec l'IA offrent une assistance rapide et efficace aux clients [34].

Environnement et Agriculture

- ❖ Optimisation de l'irrigation et l'utilisation des pesticides.
- ❖ Prévisions météorologiques : Étude des changements climatiques et prédiction des phénomènes naturels dangereux [35].

Éducation :

L'intelligence artificielle est utilisée dans des plateformes qui adaptent les leçons au niveau de l'élève, ainsi que dans des assistants virtuels capables de répondre aux questions des étudiants à tout moment [35].

II .2 Apprentissage automatique (Machine Learning)

II .2.1 Definition

L'apprentissage automatique (Machine Learning - ML) est une branche de l'intelligence artificielle (IA) qui se concentre sur l'utilisation des données et des algorithmes afin de simuler la manière dont l'être humain apprend, en améliorant progressivement la précision au fil du temps. Arthur Samuel, informaticien et l'un des pionniers de l'intelligence artificielle, a été le premier à introduire ce concept dans les années 1950, en le définissant comme : « une discipline qui donne aux ordinateurs la capacité d'apprendre sans être explicitement programmés ». L'apprentissage automatique consiste à fournir aux algorithmes informatiques de grandes quantités de données afin qu'ils puissent apprendre à identifier et à découvrir des schémas et des relations au sein de ces données. Lors de leur

exécution, ces algorithmes effectuent leurs propres prédictions ou prennent leurs propres décisions sur la base de leurs analyses. Au fur et à mesure que les algorithmes reçoivent davantage de données, ils continuent d'améliorer leurs choix et de développer leurs performances, de la même manière qu'un être humain s'améliore avec la pratique [36].

II .2.2 Types d'apprentissage de ML:

L'apprentissage automatique comprend différents types, que l'on peut généralement diviser en quatre catégories principales :

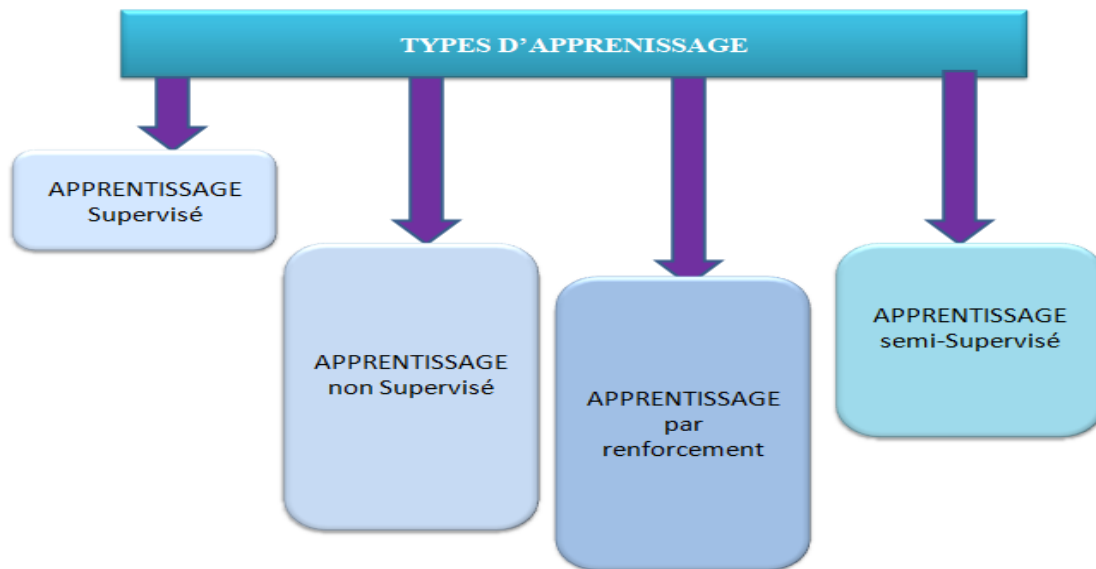


Figure II.9 : les types d'apprentissage.

II.2.2.1 Apprentissage supervisé

L'apprentissage supervisé est une méthode dans laquelle un algorithme est entraîné à l'aide de données étiquetées, où chaque exemple contient des entrées et des sorties correctement définies. L'objectif est d'apprendre la relation entre ces entrées et sorties à partir des exemples fournis. Par exemple, un algorithme peut être entraîné à distinguer entre des chats et des chiens en utilisant des images étiquetées. Il apprend ainsi à reconnaître les caractéristiques de chaque catégorie afin de classer correctement de nouvelles images [37].

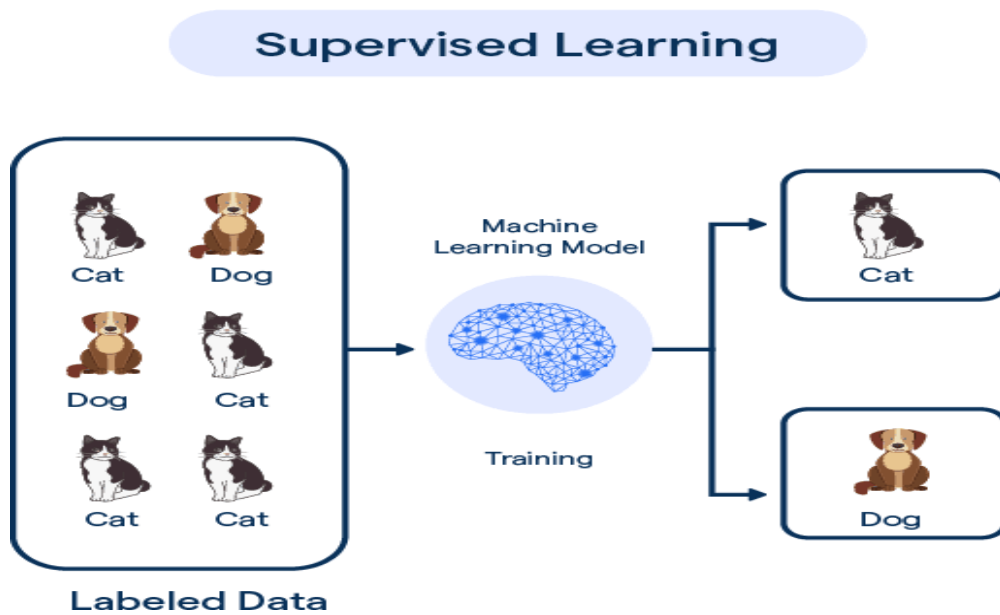
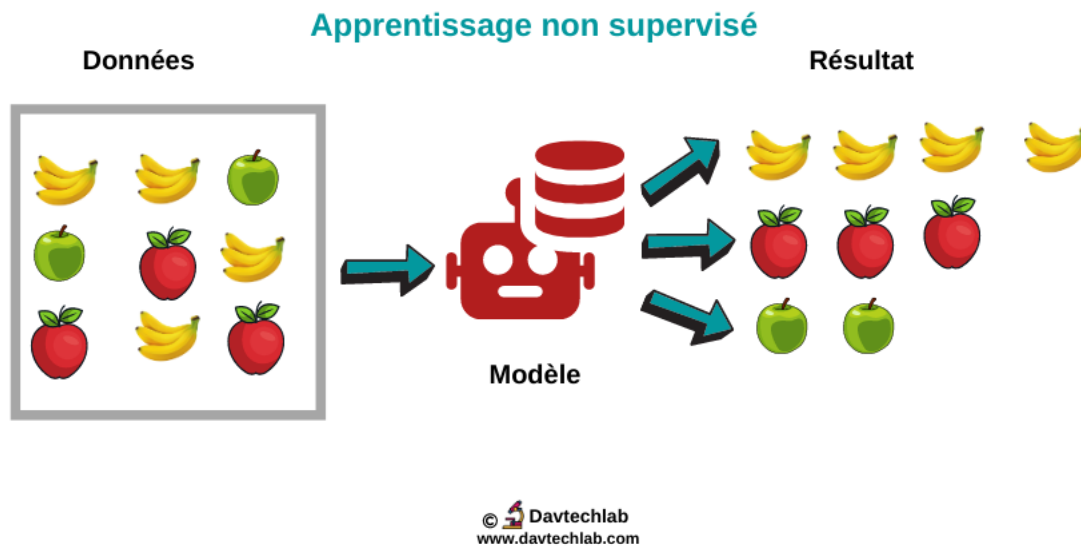


Figure II.10: Apprentissage supervisé.

II.2.2.2 Apprentissage non supervise

L'apprentissage non supervisé est considéré comme l'un des types d'apprentissage automatique, où les algorithmes sont entraînés sur des données ne disposant pas d'étiquettes préalables. Ce type d'apprentissage vise à découvrir les motifs et les relations cachées dans les données. Lors de l'exécution, les algorithmes regroupent les données similaires ou réduisent leur dimension, dans le but d'extraire des informations pertinentes pouvant être exploitées [37].



FigureII.11:Apprentissage non supervisé.

II.2.2.3Apprentissage par renforcement

L'apprentissage par renforcement (Reinforcement Learning - RL) est une branche de l'apprentissage automatique dans laquelle un agent intelligent apprend à prendre des décisions en interagissant avec un environnement spécifique. Contrairement à l'apprentissage supervisé, où les algorithmes sont guidés par des exemples étiquetés, l'agent découvre par lui-même comment agir afin de maximiser la récompense. Pour chaque action qu'il entreprend, l'agent reçoit un signal de rétroaction sous forme de récompense (positive ou négative). L'objectif est de maximiser la somme totale des récompenses sur le long terme, tout en améliorant progressivement ses stratégies de prise de décision, connues sous le nom de politiques (policies) [38].

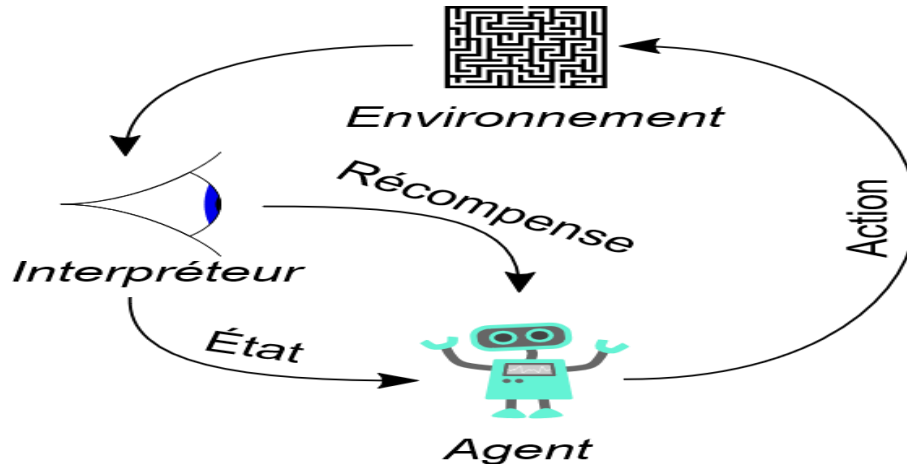


Figure II.12:Apprentissage par renforcement.

II .2.2.4 Apprentissage semi-supervisé

L'apprentissage semi-supervisé est une méthode en apprentissage automatique qui combine des données étiquetées et non étiquetées pour entraîner des algorithmes. Généralement, une petite quantité de données étiquetées est utilisée aux côtés d'une grande quantité de données non étiquetées. L'algorithme effectue des prédictions temporaires pour les données non étiquetées, appelées pseudo-étiquettes (pseudo-labels). Ensuite, ces données étiquetées sont utilisées avec les données étiquetées d'origine pour entraîner le modèle de manière plus efficace, permettant au modèle d'apprendre à partir d'un ensemble de données plus large comparé à l'utilisation uniquement des données étiquetées [39].

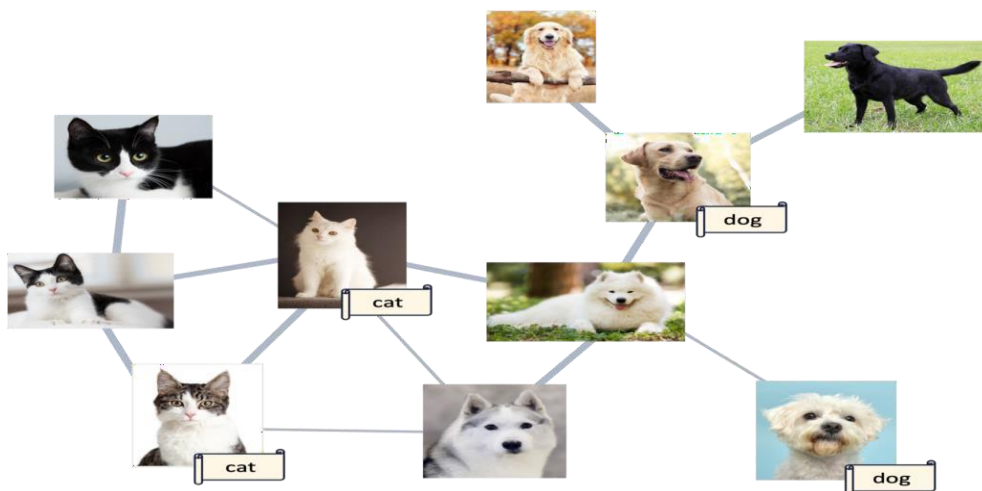


Figure II.13:Apprentissage semi –supervisé.

II.3 Réseaux de neurones artificiels

II .3.1 Neurone biologique vs neurone artificiel

Les réseaux neuronaux jouent un rôle essentiel dans la transmission et le traitement de l'information, que ce soit dans le cerveau humain ou dans les systèmes artificiels. Bien que les réseaux artificiels soient inspirés du cerveau, il existe des différences claires dans la structure, le fonctionnement et l'apprentissage. Alors, quelle est la véritable différence entre les réseaux neuronaux humains et les réseaux neuronaux artificiels ?

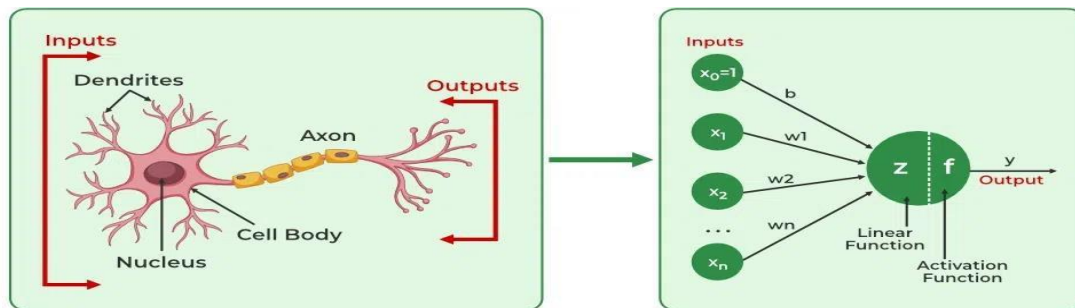


Figure II.14: ANN vs BNN.

II .3.1.1 Neurone biologique

Les réseaux neuronaux représentent la base de la conscience chez les êtres vivants. Le neurone est une cellule spécialisée présente dans le cerveau, la moelle épinière et les nerfs, et elle est responsable de la communication et du traitement de l'information. La cellule neuronale se compose de trois parties principales : les dendrites qui reçoivent les signaux d'autres neurones, le soma qui intègre ces entrées, et l'axone qui transmet le signal résultant aux neurones suivants via les synapses.

Les réseaux neuronaux biologiques sont capables de traiter différents types d'informations rapidement, car ils peuvent gérer plusieurs entrées en même temps. Cependant, leur vitesse reste inférieure par rapport aux systèmes basés sur le silicium en raison de la nature des

signaux électrochimiques, et cela est dû à l'absence d'une unité centrale chargée de l'organisation [40].

II .3.1.2 Neurone artificiel

Les réseaux neuronaux artificiels (ANNs) sont des modèles mathématiques inspirés de la structure des réseaux neuronaux dans le cerveau humain. Les formes les plus simples de ces réseaux suivent le modèle de l'alimentation avant (feedforward), où les données passent de la couche d'entrée à la couche de sortie directement, sans rétroaction. Le réseau est généralement composé de trois types de couches : une couche d'entrée qui reçoit les données, des couches cachées qui les traitent, et une couche de sortie qui produit le résultat final.

Parmi les principaux avantages de ce réseau, on trouve sa capacité à l'apprentissage multitâche, car il peut traiter différents types de données, qu'elles soient textuelles ou non, ce qui le rend utilisable dans divers domaines. Il se distingue également par sa capacité de prédiction, grâce à sa faculté de reconnaître des motifs complexes dans les données, ce qui le rend utile pour les prévisions temporelles telles que les prix des actions ou les tendances économiques. Cependant, les réseaux neuronaux artificiels ne sont pas exempts d'inconvénients, car leur nature interne rend difficile l'interprétation de la façon dont les décisions sont prises à l'intérieur du réseau, certains les qualifiant de « boîte noire ». Leur fonctionnement nécessite également une grande capacité de calcul, ce qui peut constituer un obstacle à leur application dans certains environnements [40].

II .4 Perceptron

Le perceptron est un modèle mathématique inspiré des neurones biologiques, et il est utilisé au sein des réseaux de neurones artificiels. Il a été introduit par Frank Rosenblatt en 1957 et est considéré comme l'un des plus anciens algorithmes dédiés à la classification binaire. Son principe repose sur la capacité à séparer les données en deux classes distinctes à l'aide d'une frontière de décision linéaire, en supposant que les données sont linéairement séparables. Dans la pratique, le perceptron analyse les caractéristiques d'entrée afin d'identifier des relations ou des motifs au sein des données. Ainsi, il joue un rôle important

dans des domaines tels que l'apprentissage automatique et le traitement d'images, où la classification des données constitue une étape essentielle [41].

II.4.1 Fonctionnement du perceptron :

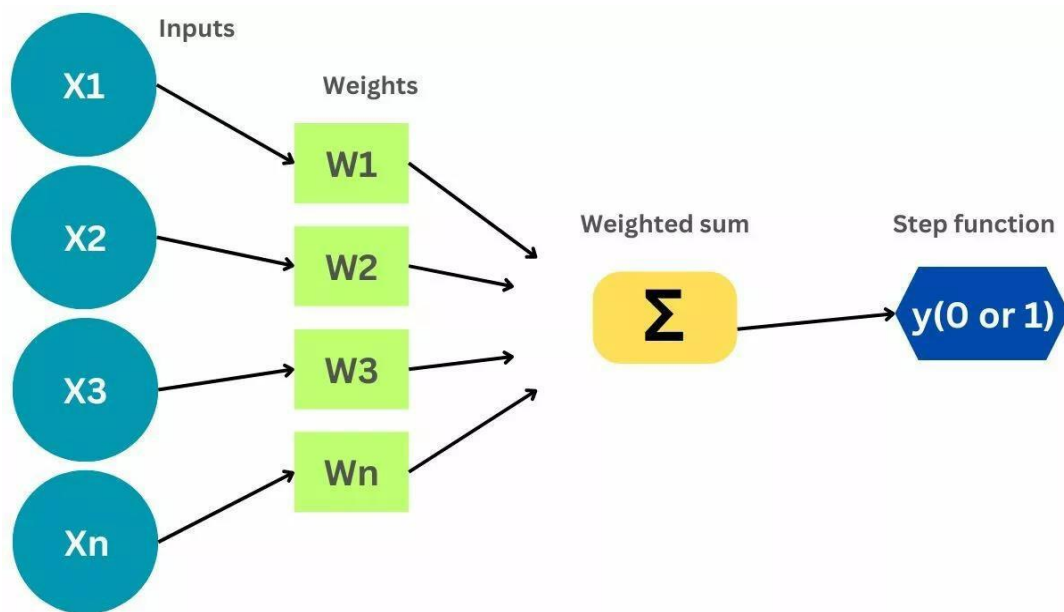


Figure II.15:schéma d'un perceptron simple.

II .4.1.1 Les entrées (inputs)

Les entrées constituent la première étape du fonctionnement du perceptron, où il peut recevoir plusieurs données, que ce soit à partir d'une seule entrée ou de plusieurs entrées, et chaque entrée est représentée par une valeur spécifique qui reflète l'une des caractéristiques des données [41].

II .4.1.2 Les poids (Weights)

À chaque entrée est attribué un nombre représentant son poids, qui indique son importance relative. La valeur de chaque entrée est multipliée par son poids correspondant pour obtenir la valeur pondérée, qui sera utilisée ultérieurement dans le calcul [41].

II .4.1.3 Le biais (Bias)

Le biais est une valeur ajoutée au calcul pour contrôler la sortie du modèle indépendamment des entrées, ce qui lui confère une plus grande flexibilité dans la prise de décisions [41].

II .4.1.4 La fonction sommation

Cette étape est également appelée fonction de somme nette, dans laquelle les valeurs pondérées de toutes les entrées sont additionnées [41].

II .4.1.5 Fonction d'activation

C'est une opération mathématique effectuée sur la somme des entrées et le biais, qui renvoie 1 si la somme pondérée dépasse un certain seuil, et 0 si elle ne le dépasse pas [41].

II .4.1.6 La sortie (Output)

La sortie est la dernière étape du fonctionnement du perceptron, où elle est binaire : 0 ou 1 [41].

II .5 Réseau multi couche (MPL)

Les réseaux de neurones multicouches (MLP) sont l'un des types de réseaux de neurones artificiels les plus courants et sont souvent utilisés pour des tâches telles que la reconnaissance de motifs, la classification et la prédiction. Ce sont des réseaux de neurones à propagation avant (Feedforward), ce qui fait que les données circulent dans une seule direction, et la structure de base de ce réseau se compose de trois couches principales : la couche d'entrée, une ou plusieurs couches cachées, et la couche de sortie [42].

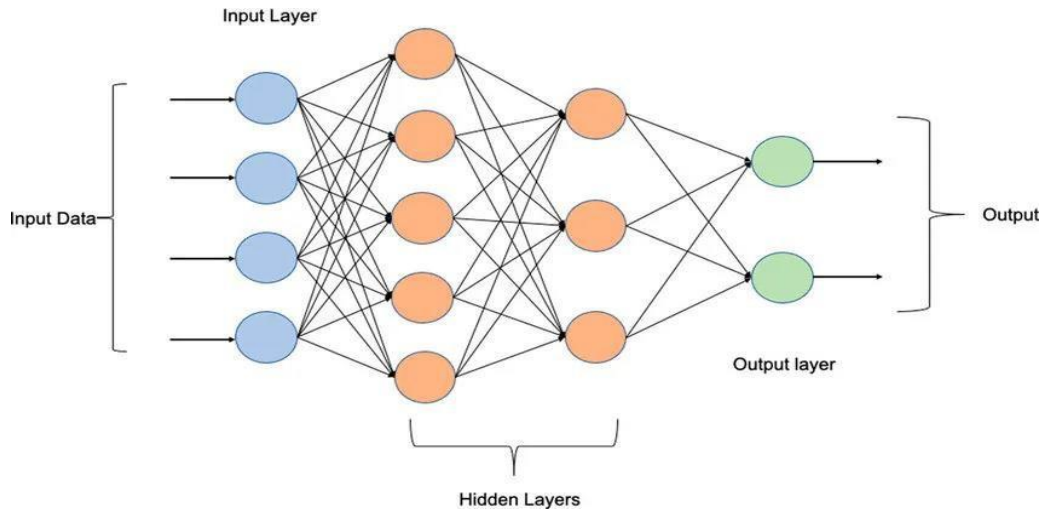


Figure II. 16: Architecture d'un perceptron multicouche (MPL).

II .5.2 La couche d'entrée

La couche d'entrée dans un perceptron multicouche (MLP) reçoit les données d'entrée, qui peuvent être constituées de caractéristiques extraites des échantillons d'entrée dans l'ensemble de données. Chaque neurone de la couche d'entrée représente une seule caractéristique. Les neurones de la couche d'entrée n'effectuent aucun calcul ; ils se contentent de transmettre les valeurs d'entrée aux neurones de la première couche cachée [43].

II .5.2 Les couches cachées

Les couches cachées dans un perceptron multicouche (MLP) sont composées de neurones interconnectés entre eux et qui traitent les données d'entrée. Chaque neurone dans ces couches reçoit des entrées de tous les neurones de la couche précédente. Ces entrées sont multipliées par leurs poids correspondants, notés (w), où ces poids déterminent dans quelle mesure chaque entrée influence la sortie du neurone. En plus des poids, chaque neurone possède un biais (bias), noté (b), qui représente une entrée supplémentaire aidant à ajuster la sortie du neurone. Les poids et les biais sont tous deux appris pendant le processus d'entraînement. La somme pondérée des entrées de chaque neurone est calculée comme suit : $z = \sum w_i x_i + b$

Où (z) représente la somme pondérée des entrées, (w_i) le poids de la i -ème entrée, et (x_i) la valeur de la i -ème entrée. Ensuite, cette somme pondérée est passée à travers une fonction d'activation (f) pour obtenir la sortie du neurone : $a = f(z)$

Où (a) représente la valeur obtenue après application de la fonction d'activation. La fonction d'activation introduit la non-linéarité dans le réseau, ce qui lui permet de représenter et d'apprendre des motifs plus complexes dans les données. Elle détermine également la nature de la sortie du neurone et sa réponse aux différentes entrées [43].

II .5.3 Couche de sortie (Output Layer)

La couche de sortie représente la dernière étape dans l'architecture du réseau de neurones multicouches (MLP), où elle est chargée de produire les résultats finaux ou les prédictions que le modèle cherche à atteindre. Le nombre de neurones dans cette couche est déterminé en fonction de la nature du problème étudié ; dans les cas de classification binaire (Binary Classification), un seul neurone est généralement utilisé, tandis que leur nombre augmente dans les tâches de classification multi-classes (Multi-class Classification) ou dans les problèmes de régression (Regression). Les neurones de cette couche reçoivent les signaux provenant de la dernière couche cachée, puis appliquent une fonction d'activation (Activation Function) qui diffère souvent, dans ses propriétés, de celles utilisées dans les couches précédentes, afin de s'adapter à la nature des sorties et aux valeurs cibles.

Du point de vue de l'entraînement, le réseau cherche à réduire rapidement l'écart (ou l'erreur) entre les sorties prédites et les valeurs réelles présentes dans les données d'entraînement. Ce processus se fait par l'ajustement précis des « poids » associés à chaque connexion d'entrée. Cette mise à jour n'est pas effectuée de manière aléatoire, mais est guidée par des algorithmes d'optimisation avancés, tels que la descente de gradient stochastique (Stochastic Gradient Descent - SGD), qui calcule les gradients de la fonction de perte (Loss Function) par rapport aux poids et les met à jour de manière itérative. Grâce à ce mécanisme, le réseau est capable d'apprendre des modèles et des relations complexes et non linéaires dans les données, ce qui améliore sa capacité à fournir des prédictions précises lorsqu'il est testé sur de nouvelles données qu'il n'a jamais traitées auparavant [43].

II .6 Fonction d'activation

Les fonctions d'activation sont considérées comme une partie importante des réseaux de neurones, car elles représentent une opération mathématique utilisée pour déterminer les

sorties produites par un neurone ou une unité. L'importance de ces fonctions réside dans le fait qu'elles ajoutent une propriété de non-linéarité au modèle, ce qui lui permet de comprendre les relations complexes au sein des données [44]. Elles contribuent également à organiser l'information d'une manière qui facilite la reconnaissance des motifs et des relations, surtout que les données réelles sont souvent non linéaires. En l'absence des fonctions d'activation, le réseau de neurones devient un modèle linéaire similaire à la régression linéaire, ce qui limite sa capacité à traiter des problèmes complexes [45]. Grâce à leur capacité à traiter des données multidimensionnelles, ces fonctions sont utilisées dans divers domaines tels que l'analyse d'images, de sons et de vidéos [44].

Avant d'aborder les types de fonctions d'activation, il est nécessaire de comprendre comment elles fonctionnent au sein du réseau de neurones. En général, une fonction d'activation prend les valeurs issues d'une couche donnée du réseau, qui représentent la somme des poids ajoutée au biais (bias), puis elle transforme ces valeurs en une sortie dans un intervalle déterminé. Pour illustrer cela, on calcule d'abord la somme pondérée des entrées en ajoutant le biais ($w_1x_1 + w_2x_2 + \dots + w_nx_n + b$), puis cette valeur est transmise à la fonction d'activation f , qui fournit finalement la sortie du neurone. Dans ce contexte, chaque valeur (x_i) représente les sorties des neurones de la couche précédente, tandis que (w_i) représente les poids associés à chaque entrée [46].

Et ci-après, nous aborderons les principaux types de fonctions d'activation couramment utilisées dans les réseaux de neurones :

- ❖ **Sigmoid** : La fonction Sigmoid est un type de fonction d'activation qui prend la forme d'une courbe en S, où elle transforme les valeurs dans un intervalle compris entre 0 et 1, ce qui la rend adaptée aux tâches de classification binaire. Elle aide également à représenter les résultats sous forme de probabilités et est utilisée dans les modèles qui reposent sur l'apprentissage par gradient [47].
- ❖ **Tanh** : La fonction Tanh est un type de fonction d'activation, et elle est considérée comme une version modifiée de la fonction Sigmoid. Elle transforme les valeurs dans un intervalle compris entre -1 et 1, ce qui la rend adaptée pour représenter les données de manière plus équilibrée. Elle est également utilisée dans les couches cachées et contribue à améliorer le processus d'apprentissage au sein des réseaux de neurones [47].
- ❖ **La fonction ReLU (Rectified Linear Unit)** : est une fonction d'activation qui renvoie la valeur d'entrée si elle est positive, tandis qu'elle donne zéro si elle est négative. Cette fonction se distingue par sa simplicité d'utilisation et sa rapidité de calcul, et elle est couramment utilisée dans les réseaux de neurones profonds car elle aide à apprendre des motifs complexes [47].
- ❖ **La fonction Softmax** : est un type de fonction d'activation utilisée dans les tâches de classification multi-classes. Elle transforme les valeurs brutes en probabilités dont la somme est égale à 1. Cette fonction aide à déterminer la classe appropriée à laquelle appartient l'entrée en attribuant une probabilité à chaque classe [47].

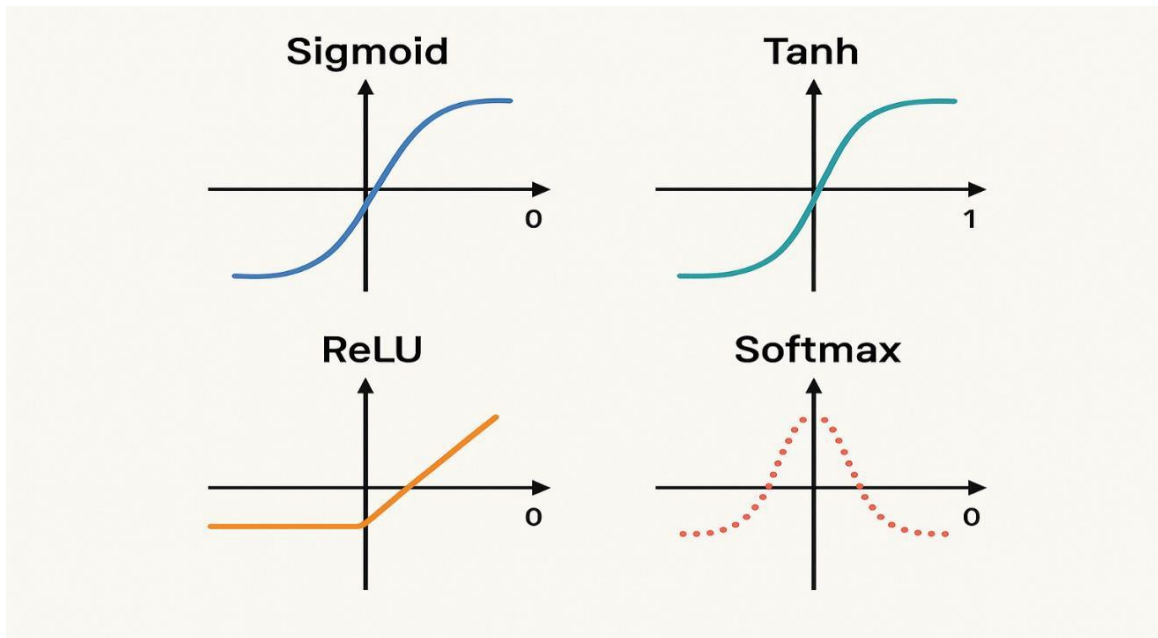


Figure II.17: Types de fonction d'activation.

II .7 Deep Learning

II.7.1 Définition de l'apprentissage profonde (deep learning)

L'apprentissage profond est un type d'apprentissage automatique qui repose sur l'utilisation de réseaux de neurones artificiels imitant le fonctionnement du cerveau humain. Le cerveau est composé de cellules appelées neurones, qui transmettent des signaux pour traiter l'information. De la même manière, les réseaux de neurones profonds sont constitués de plusieurs couches de neurones artificiels qui échangent des informations pour traiter les données. Le terme « profond » fait référence au nombre de couches dans le réseau ; plus le nombre de couches augmente, plus le réseau devient profond. Les réseaux de neurones profonds peuvent apprendre à partir de données étiquetées et non étiquetées, ce qui leur permet de fonctionner dans des environnements d'apprentissage supervisé ou non supervisé. En général, ces réseaux nécessitent de grandes quantités de données pendant l'entraînement afin d'obtenir des performances précises et efficaces.

Ces réseaux améliorent la précision des prédictions pendant l'entraînement en ajustant leurs paramètres internes à l'aide d'algorithmes d'apprentissage, sans intervention humaine directe [48].

II.7.2 La différence entre l'apprentissage automatique et l'apprentissage profond

	ML	DL
Définition	C'est une partie de l'intelligence artificielle qui aide les systèmes à apprendre à partir des informations en utilisant certaines méthodes sans avoir besoin d'écrire du code.	C'est un type particulier d'apprentissage automatique qui utilise des réseaux de neurones complexes pour comprendre les motifs difficiles dans les données.
Exigences en données	Fonctionne bien avec de petites quantités de données structurées.	Nécessite une grande quantité de données, en particulier des données non structurées comme les images, les sons et les textes.
Ingénierie des caractéristiques	Nécessite que les experts sélectionnent et extraient manuellement les informations importantes.	Extrait automatiquement ces informations à partir des données de base lui-même.

Complexité des algorithmes	Utilise des algorithmes plus simples tels que la régression linéaire, les arbres de décision et SVM.	Utilise des architectures complexes telles que CNN et RNN.
Temps d'entraînement	Plus rapide relativement grâce à la simplicité des modèles et à la petite taille des données.	Plus lent en raison de la complexité des modèles et de la grande taille des données.
Puissance de calcul	Peut être entraîné sur des ordinateurs.	Il nécessite des machines puissantes telles que des unités de traitement graphique (GPU) ou des unités Tensor
Interprétabilité	Plus clair et plus facile à comprendre.	Souvent considéré comme une « boîte noire » et difficile à comprendre.
Cas d'utilisation	Bon pour des tâches telles que la détection de spam, la détection de fraude et la maintenance prédictive.	Excellent pour la reconnaissance d'images et de sons, la compréhension du langage naturel et les voitures autonomes.
Intervention humaine	Nécessite davantage d'aide humaine pour sélectionner les informations et améliorer le modèle.	Réduit le besoin d'aide humaine car il apprend tout seul.
Performance avec les données non structurées	Moins efficace et nécessite un prétraitement.	Fonctionne très bien directement avec les données non structurées [49].

Tableau II.1:ML vs DL.

II.7.3 Les avantages et les inconvénients de l'apprentissage profond

Le deep Learning est considéré comme le moteur principal de la révolution de l'intelligence artificielle contemporaine, offrant des perspectives inédites dans la simulation de l'intelligence humaine. Cependant, il soulève également des défis complexes qui nécessitent une analyse précise de ses avantages et de ses inconvénients :

Avantages	Inconvénients
Apprentissage automatique des caractéristiques : L'apprentissage profond peut extraire automatiquement les caractéristiques à partir des données sans intervention humaine, notamment dans des tâches complexes comme la reconnaissance d'images.	Besoin d'une grande quantité de données : L'apprentissage profond nécessite beaucoup de données de qualité pour entraîner correctement les modèles, ce qui demande du temps et des efforts pour les collecter et les préparer.
Traitement des données volumineuses et complexes : Les modèles d'apprentissage profond sont capables de traiter de grandes quantités de données complexes que les méthodes traditionnelles ont du mal à gérer.	Forte exigence en ressources de calcul : L'apprentissage profond nécessite des équipements puissants tels que des processeurs modernes et des cartes graphiques, ainsi que de grandes capacités de stockage et de mémoire, ce qui le rend coûteux.
Amélioration des performances : L'apprentissage profond fournit des résultats précis et des performances élevées dans divers domaines tels que les images, le son et le traitement du langage.	Problème de surapprentissage : Le modèle peut apprendre excessivement les détails des données d'entraînement, ce qui lui permet d'obtenir de très bonnes performances sur ces données, mais il peut échouer lorsqu'il est confronté à de nouvelles données.
Représentation des relations non linéaires : L'apprentissage profond se distingue par sa capacité à comprendre les relations complexes et non linéaires entre les données mieux que les méthodes traditionnelles.	Difficulté d'interprétation : Les modèles d'apprentissage profond sont complexes, et il est souvent difficile de comprendre la manière dont ils prennent leurs décisions par rapport aux méthodes traditionnelles.

Adaptabilité et extensibilité : Les modèles d'apprentissage profond peuvent être adaptés facilement à de nouvelles tâches, même avec une quantité limitée de données.	Défis juridiques et éthiques : Les modèles peuvent refléter les biais présents dans données liées à la confidentialité et aux droits de propriété intellectuelle [50].
---	--

Tableau II. 2:Avantages et in Inconvénients d'apprentissage profond.

II .8 Conclusion

En conclusion, l'intelligence artificielle joue un rôle essentiel dans le développement des technologies modernes. Les techniques de Machine Learning et de Deep Learning permettent aux machines d'apprendre et d'effectuer plusieurs tâches avec efficacité et précision.

À travers cette étude, nous avons découvert les concepts principaux de l'intelligence artificielle, les différents types d'apprentissage et le fonctionnement des réseaux de neurones artificiels. Malgré certaines difficultés comme le besoin de grandes quantités de données et de ressources de calcul, l'intelligence artificielle reste un domaine très prometteur avec de nombreuses applications dans le futur.

Chapitre 03 : État de l'art sur l'AMC basé sur le Deep learning.

III.1 Introduction

Le deep Learning est aujourd'hui une méthode essentielle pour analyser et reconnaître les signaux dans les systèmes de communication. Sa force réside dans sa capacité à apprendre automatiquement les caractéristiques importantes des données, même dans des environnements complexes où le bruit et les interférences sont présents. Mais pour que ces modèles soient efficaces, il est nécessaire de bien choisir la manière de représenter les signaux avant de les introduire dans les réseaux neuronaux.

Différentes techniques sont utilisées à cet effet : les représentations I/Q temporelles, qui conservent les informations de phase et d'amplitude ; les diagrammes de constellation, qui offrent une vision intuitive des symboles modulés ; ou encore des représentations plus avancées basées sur les spectres et les graphes. Ces représentations servent ensuite de base aux architectures de Deep Learning comme les CNN, RNN, GNN, Transformers ou encore les modèles hybrides, chacun apportant des avantages spécifiques pour améliorer la précision et la robustesse de la classification automatique des modulations. Ce schéma présente une classification méthodique des approches récentes basées sur l'apprentissage profond (Deep Learning) pour le traitement des signaux. Il organise ces méthodes en plusieurs catégories principales : d'abord les représentations du signal (I/Q, A/P, spectre, constellation), puis les architectures de réseaux de neurones (CNN classiques ou spécialisés, RNN comme LSTM et BiLSTM, GNN, et Transformers). Le schéma couvre également les modèles hybrides combinant plusieurs types de réseaux (ex. CNN-LSTM), et enfin les techniques de régularisation sous leurs formes explicite et implicite.

III.2 Taxonomie des méthodes de Deep Learning pour l'AMC

Au cours des dernières années, plusieurs méthodes de Deep Learning ont été proposées pour améliorer les performances de la classification automatique des modulations. Ces approches exploitent différentes architectures de réseaux de neurones et différentes représentations des signaux afin d'extraire automatiquement les caractéristiques

pertinentes. La figure suivante présente une taxonomie des principales méthodes utilisées dans la littérature, regroupées selon leur famille de modèles et leurs variantes les plus courantes

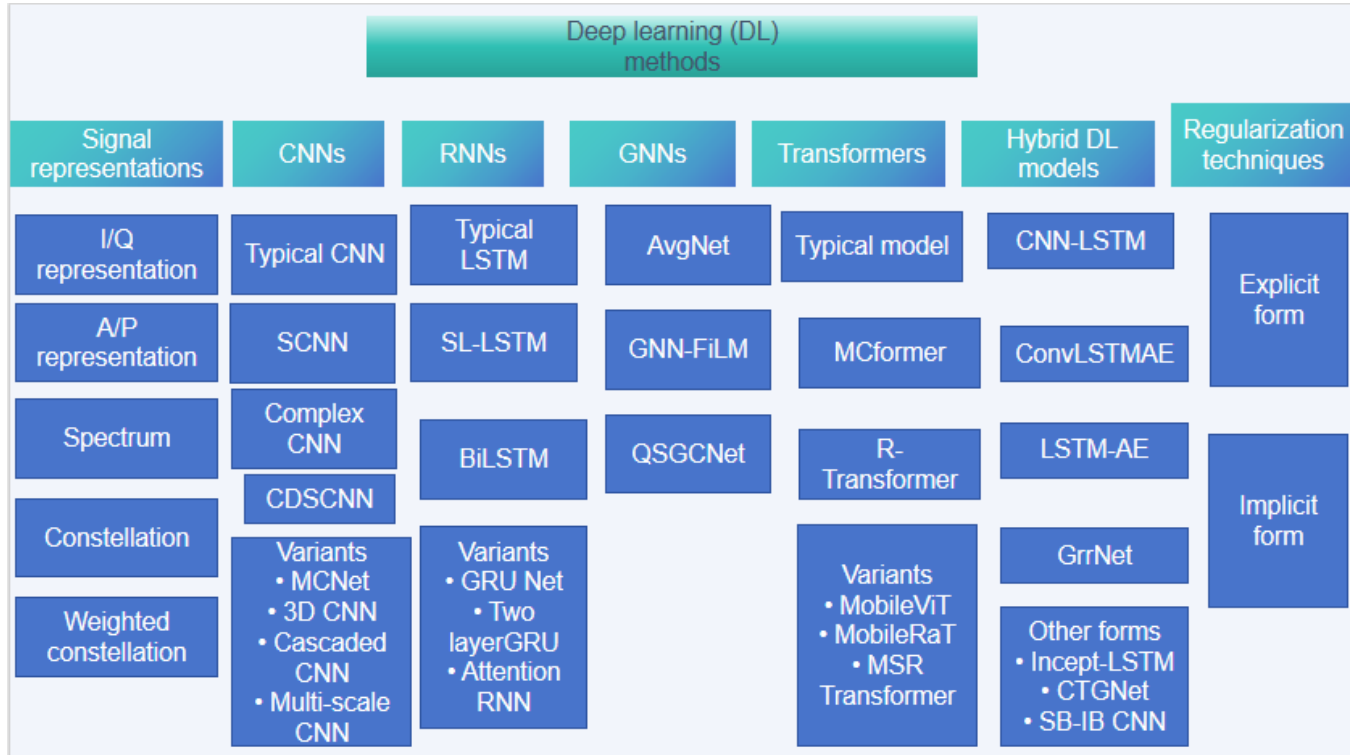


Figure III.18: Méthodes d'apprentissage profond (Deep Learning) [51].

III.3 Représentation du signal

La manière de représenter le signal joue un rôle très important dans la façon d'extraire les caractéristiques principales qui montrent les schémas fondamentaux, et cela est très important pour le développement de l'apprentissage profond. Les différents types de représentation des signaux, que ce soit dans le domaine temporel ou fréquentiel, déterminent l'utilisation de différents types d'architectures d'apprentissage profond. Par conséquent, nous commencerons d'abord par examiner les méthodes courantes de représentation du signal utilisées dans l'apprentissage profond.

III.2.1 I/Q temporel

La technique I/Q est considérée comme une méthode utilisée pour transformer les signaux numériques en signaux analogiques afin de faciliter le processus de transmission. Cela se fait en divisant le signal en deux composantes orthogonales, à savoir la partie réelle (I) et la partie imaginaire (Q). Dans ce cas, le spectre du signal est réparti sur plusieurs bandes de fréquences, ce qui contribue à améliorer l'utilisation du spectre. Grâce à l'orthogonalité de ces deux composantes, le signal I/Q peut supprimer les interférences provenant des canaux adjacents, ce qui le rend adapté à une utilisation dans des environnements complexes tels que la propagation multi-trajets et l'évanouissement sélectif en fréquence. Les représentations I/Q sont également utilisées directement dans les modèles d'apprentissage profond qui traitent les signaux temporels, tels que LSTM et Transformer. Pour un signal I/Q reçu dans les systèmes de communications sans fil, il peut être exprimé par la relation suivante :

$$r(k) = r_I(k) + j, r_Q(k) \quad (8)$$

$r_I(k)$: la composante en phase(In-phase, I);

$r_Q(k)$: la composante en quadrature(Quadrature, Q);

Où représente la composante en phase et la composante en quadrature, et ils peuvent être calculés comme suit :

$$r_I(k) = \text{real}[r(k)] \quad (9)$$

$$r_Q(k) = \text{imag}[r(k)] \quad (10)$$

Dans certaines applications de l'apprentissage profond, les signaux I/Q sont introduits directement dans le modèle afin d'apprendre les caractéristiques qui distinguent les schémas de modulation. Par exemple, une étude a montré que pour une valeur de SNR égale à 0 dB, la précision de classification d'un modèle CNN utilisant des schémas comme 16QAM et 64QAM peut être inférieure à 70 %, tandis qu'elle peut dépasser 90 % lorsque le SNR augmente à 10 dB, ce qui montre l'impact du niveau de bruit sur les performances.

Bien que ce type de représentation soit considéré comme simple, il nécessite des modèles d'apprentissage profond d'extraire automatiquement les informations du domaine fréquentiel, ce qui peut constituer un défi. Par conséquent, des techniques de prétraitement du signal telles que la transformée de Fourier et la transformée en ondelettes sont utilisées afin d'améliorer les performances. [52]

III.2.2 Diagramme de constellation

Dans les schémas de modulation analogique, les paramètres de la porteuse tels que l'amplitude, la fréquence et la phase varient de manière continue. En revanche, dans les schémas de modulation numérique, la variation de ces paramètres est basée sur des signaux numériques discrets. Le diagramme de constellation décrit la distribution des points de symboles et constitue un moyen intuitif de déterminer le taux d'erreur binaire dans les schémas de modulation. Lors de la réception d'un signal, chaque symbole peut être estimé et représenté sur le plan sous forme de diagramme de constellation, qui peut être exprimé comme une fonction de mappage comme suit :

$$C = m(r) \quad (11)$$

En tant que représentation bidimensionnelle, les diagrammes de constellation peuvent être mieux analysés à l'aide de modèles d'apprentissage profond tels que CNN, ce qui permet d'explorer la corrélation entre les symboles. L'utilisation des diagrammes de constellation comme entrées des modèles offre un avantage notable dans la classification des schémas de modulation d'ordre élevé. Des études ont montré qu'à une valeur de SNR de 10 dB, un modèle CNN léger peut atteindre une précision de reconnaissance supérieure à 95 % pour des modulations telles que 16QAM et 64QAM en analysant les caractéristiques de la constellation [53].

III.4 Modèles basés sur CNN

Les modèles basés sur les réseaux de neurones convolutifs (CNN) ont montré des performances excellentes dans le traitement des données présentant des caractéristiques spatiales, telles que la segmentation d'images, la détection d'objets et la classification.

Des études dans le domaine de la reconnaissance automatique des schémas de modulation (AMR) ont introduit les réseaux CNN afin de détecter les schémas du signal en utilisant leur capacité d'extraction de caractéristiques spatiales.

Selon les types de données d'entrée, les approches actuelles basées sur CNN pour la reconnaissance automatique des schémas de modulation peuvent être divisées en deux catégories : les modèles CNN utilisant des entrées I/Q brutes, et les modèles CNN utilisant des entrées prétraitées.

Il existe également une autre catégorie qui concerne la conception d'architectures CNN efficaces, visant à répondre aux exigences de latence et de complexité dans les systèmes de communication avancés. Ces catégories seront analysées ci-dessous.

III.4.1 CNN typique

Kursat et ses collègues ont proposé une architecture de réseau de neurones convolutif (CNN) typique pour la reconnaissance automatique des schémas de modulation. Ce réseau a été construit à l'aide de la bibliothèque Keras, qui est une bibliothèque d'apprentissage automatique open source. L'architecture contient quatre couches de convolution et de pooling, et se termine par deux couches entièrement connectées (couches denses). Pour obtenir de meilleures performances, la première couche du CNN contient 256 noyaux de convolution, tandis que la dernière couche n'en utilise que 64. Avec l'augmentation du nombre de couches, le nombre de noyaux de convolution, ou le nombre de canaux, diminue progressivement. Le nombre total de paramètres entraînables est de 659594.

Après chaque couche de convolution, une fonction appelée ReLU est utilisée pour assurer une transformation non linéaire. Elle est définie comme suit :

$$x_{out} = ReLU(x_{in}) = \max(0, wx_{in} + b) \quad (12)$$

Dans le CNN traditionnel, la couche de max-pooling est utilisée pour réduire la dimensionnalité des caractéristiques produites par les couches cachées, mais son efficacité avec les signaux temporels nécessite encore des recherches supplémentaires. Dans les tâches de classification d'images, la couche Max-Pooling aide à apprendre l'invariance à

la rotation, mais les caractéristiques des signaux temporels sont très différentes de celles des images.

La technique appelée Dropout est utilisée pour supprimer aléatoirement une partie des neurones afin de réduire la complexité du modèle et limiter le sur apprentissage (overfitting). De plus, il est important de noter que deux couches entièrement connectées sont utilisées pour l'intégration des caractéristiques et leur projection vers la couche de sortie. Cela fait que la majorité des paramètres entraînaables se concentrent dans ces couches plutôt que dans les couches convolutionnelles intermédiaires. En d'autres termes, les couches denses sont l'une des principales causes du sur apprentissage en apprentissage profond.

Enfin, le CNN calcule la probabilité de chaque classe à l'aide de la fonction Softmax comme suit :

$$y_i = \text{softmax}(x_{out}) = \frac{e^{x_{iout}}}{\sum_j e^{x_{jout}}} , \quad i, j = 1, 2, \dots, c \quad (13)$$

L'architecture CNN traditionnelle a été utilisée dans des tâches de classification d'images sans prendre en compte les caractéristiques des signaux I/Q. Bien qu'elle ait permis d'obtenir certaines améliorations par rapport aux méthodes classiques, il existe encore un potentiel d'amélioration. Le biais inductif des CNN ne prend pas en considération les propriétés physiques spécifiques des signaux I/Q, ce qui limite fortement les performances du modèle. Dans les CNN, l'accent est mis sur l'invariance spatiale plutôt que sur la dépendance temporelle, alors que les informations essentielles des signaux modulés (telles que les variations continues de phase et les trajectoires de fréquence) sont contenues dans les relations temporelles précises des séquences I/Q. De plus, le CNN traditionnel ne dispose pas d'un mécanisme clair de préservation de la phase, qui est un élément essentiel pour distinguer les modulations d'ordre élevé telles que QAM et PSK, ce qui entraîne une faible robustesse face au bruit de phase à de faibles niveaux de SNR[54].

III.4.2 CDSCNN

Dans les modèles CNN traditionnels, les composantes contenant des signaux de modulation complexes I/Q sont traitées comme des valeurs réelles séparées, ce qui affecte négativement la structure originale du signal et rend la compréhension par le réseau CNN plus difficile. Par conséquent, Xiao et ses collègues ont proposé un réseau de neurones convolutif à valeurs complexes (CDSCNN), qui repose sur des unités d'opérations complexes permettant l'apprentissage automatique de caractéristiques à valeurs complexes, et qui a été conçu spécifiquement pour la tâche de classification automatique de la modulation (AMC).

Les chercheurs ont étudié théoriquement le CDSCNN selon deux aspects : le nombre de paramètres et le nombre d'opérations en virgule flottante. La convolution séparable en profondeur (Depthwise separable convolution) améliore l'efficacité de calcul en décomposant les convolutions classiques simples. Ainsi, la convolution standard est divisée en deux types : la convolution en profondeur (linear depth convolution) et la convolution ponctuelle (pointwise convolution). La convolution en profondeur applique l'opération de convolution séparément à chaque canal des entrées, tandis que la convolution ponctuelle applique une convolution de taille 1×1 sur les sorties.

La convolution ponctuelle ne prend pas en compte les dimensions spatiales des entrées, ce qui permet de réduire davantage le nombre de paramètres du modèle. Cette capacité à réduire la complexité rend cette approche adaptée aux applications industrielles. En outre, la conception d'une fonction d'activation adaptée aux réseaux de neurones à valeurs complexes constitue un défi important. Cela nécessite d'atteindre un équilibre entre la bornitude (boundedness) et l'analyticité (analyticity). Généralement, on utilise des fonctions d'activation à valeurs réelles appliquées séparément aux parties réelle et imaginaire. Par exemple, la fonction ReLU complexe (CReLU) est définie comme suit :

$$CReLU = ReLU_{(xin-I)} + i * ReLU_{(xin-Q)} \quad (14)$$

Où xin_I et xin_Q représentent respectivement la partie réelle et la partie imaginaire de l'entrée [55].

III.5 Modèles basé sur RNN

Par rapport aux réseaux de neurones convolutionnelles (CNN), le réseau RNN est capable de conserver les informations passées, ce qui le rend adapté au traitement des séries temporelles. Chaque neurone reçoit en même temps des informations de l'instant précédent et de l'instant actuel, et cette capacité de mémorisation aide à surmonter les problèmes de dépendance à court terme et à long terme. Les architectures générales de RNN et LSTM sont présentées dans la figure 1, où W , U et V représentent les matrices de poids correspondant respectivement au passage de l'entrée vers la mémoire, de la mémoire vers la mémoire, et de la mémoire vers la sortie. Σ représente la fonction d'activation sigmoïde, tandis que $\oplus$ et $\otimes$ désignent respectivement l'addition et la multiplication élément par élément. I_t , f_t et o_t représentent respectivement les portes d'entrée, d'oubli et de sortie. Enfin, W_i , U_i , W_f , U_f , W_o et U_o sont les matrices de poids responsables de la transformation entre les espaces d'entrée, de mémoire et de sortie.

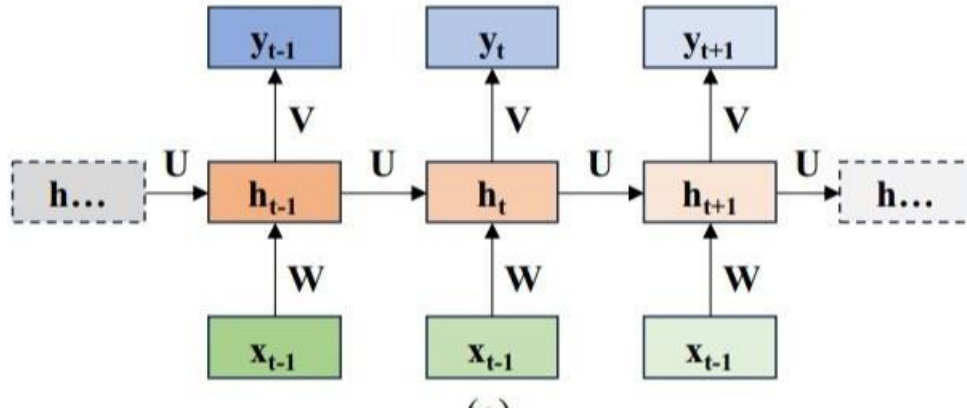


Figure III.19: Structure générale du RNN.

III.5.1 LSTM Typique

Le LSTM traditionnel est un type de réseau de neurones récurrents basé sur un mécanisme de portes, qui contient des unités LSTM. Ces unités aident le modèle à mieux comprendre les relations à dépendance temporelle longue.

L'architecture d'un modèle LSTM classique se compose de plusieurs composants. La porte d'entrée contrôle les informations à mettre à jour à partir des données d'entrée, et détermine la quantité de ces informations à l'aide d'une fonction d'activation spéciale produisant des valeurs comprises entre 0 et 1. De même, la porte d'oubli permet de déterminer la quantité

d'informations à conserver de la mémoire précédente, et elle repose également sur la même fonction d'activation pour produire des valeurs entre 0 et 1, lesquelles sont ensuite utilisées pour multiplier le contenu de l'unité de mémoire.

L'état de la cellule est l'un des éléments les plus importants du LSTM, car il permet de stocker et de transmettre les informations passées au fil du temps, tout en déterminant quelles informations doivent être supprimées ou ajoutées. Quant à la porte de sortie, elle détermine la partie de la mémoire qui sera transmise à la couche suivante ou utilisée comme sortie finale, en utilisant une fonction d'activation produisant des valeurs dans l'intervalle $(-1, 1)$.

À chaque instant, le LSTM reçoit à la fois les données actuelles et l'état caché précédent comme entrées principales, puis contrôle le flux d'informations et la mise à jour de la mémoire grâce au mécanisme de portes. Ainsi, les portes d'entrée et d'oubli déterminent si le contenu de l'état de la cellule doit être mis à jour ou conservé, tandis que la porte de sortie se charge de filtrer et d'extraire les informations de l'unité de mémoire [56].

III.5.2 SL-LSTM (LSTM à une seule couche)

Ont proposé un modèle LSTM à une seule couche (SL-LSTM) basé sur le mécanisme d'attention (Attention Mechanism). Dans la première partie du modèle, un module d'embedding du signal a été conçu afin d'améliorer l'entrée des informations et de permettre l'extraction des caractéristiques discriminantes des signaux de manière plus précise et plus complète.

Ensuite, les sorties de l'état caché (Hidden State) sont transmises au module d'attention, où différents poids sont appliqués à ces états, ce qui aide le modèle à renforcer sa capacité à capturer les caractéristiques temporelles des signaux I/Q.

Dans le contexte de l'utilisation de l'apprentissage profond pour résoudre le problème de la classification automatique de la modulation (AMC), les données à longues séquences peuvent affecter négativement le processus d'entraînement, entraînant une diminution de la précision de classification et de l'efficacité d'inférence. Par conséquent, le SL-LSTM intègre un mécanisme d'attention afin d'accélérer la convergence de la fonction objectif (Objective Function), réduisant ainsi le temps d'entraînement.

De plus, le processus d'embedding contribue à améliorer la représentation des entrées du signal, en contenant de manière plus précise les informations de modulation du signal radio, ce qui réduit la charge sur la capacité d'apprentissage du modèle.

En outre, une technique de suppression aléatoire au niveau des échantillons (Random Erasing) a été adoptée afin d'aider le modèle SL-LSTM à s'adapter à des environnements de communication sans fil de plus en plus complexes [57].

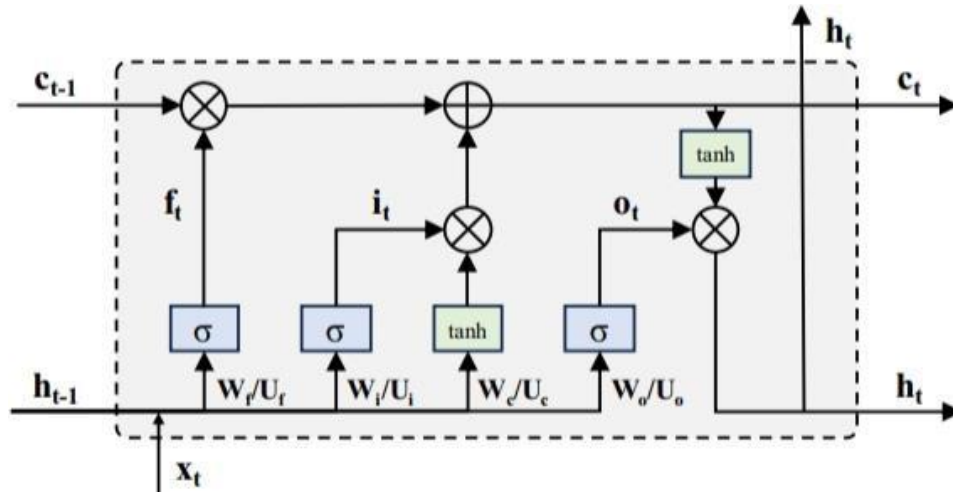


Figure III.20: Structure générale du LSTM.

III.6 Modèles basé sur GNN (Graph Neural Network)

Est un type de réseau de neurones conçu pour traiter des données structurées sous forme de graphes, c'est-à-dire des ensembles de nœuds (nodes) reliés entre eux par des arêtes (edges). Contrairement aux modèles classiques de deep learning qui travaillent sur des vecteurs ou des matrices, le GNN est capable d'exploiter directement les relations entre les éléments de la donnée. L'idée principale du GNN est de représenter chaque nœud par un vecteur de caractéristiques (embedding). Cette représentation n'est pas fixe : elle a amélioré progressivement grâce à un processus appelé message passing. À chaque étape, chaque nœud reçoit les informations de ses voisins, les combine avec ses propres caractéristiques, puis met à jour sa représentation. Ce processus est répété sur plusieurs couches, ce qui permet d'obtenir des représentations de plus en plus riches et globales.

Le GNN utilise donc une fonction d'agrégation et de mise à jour, qui peut être simple (comme une moyenne) ou plus complexe (comme des couches de convolution sur graphe

ou des mécanismes de type gating). Grâce à cela, le modèle peut apprendre à la fois les relations locales et la structure globale du graphe.

Dans le domaine de l'AMC (Automatic Modulation Classification), les GNN sont particulièrement utiles car un signal I/Q peut être représenté comme un graphe où chaque échantillon est un nœud, et les relations entre échantillons (temporelles ou fréquentielles) sont les arêtes. Cela permet de mieux capturer les dépendances complexes des signaux radio.

Cependant, la performance des GNN dépend fortement de la qualité de la structure du graphe et du mécanisme d'agrégation utilisé. Pour améliorer cela, plusieurs variantes ont été proposées, notamment AvgNet et GNN-FiLM.

III.6.1 AvgNet

AvgNet constitue une architecture basée sur les Graph Neural Networks (GNN) dédiée à la classification automatique des modulations (AMC). Cette approche repose sur la modélisation du signal I/Q sous forme de graphe, dans lequel chaque échantillon du signal est représenté par un nœud, tandis que les relations entre les échantillons sont modélisées par des arêtes. Le principe fondamental de cette architecture consiste en une agrégation moyenne des informations des nœuds voisins (Average Aggregation), permettant à chaque nœud de mettre à jour sa représentation en combinant ses propres caractéristiques avec celles de son voisinage. Cette stratégie favorise l'extraction des dépendances temporelles et structurelles du signal tout en améliorant la robustesse face au bruit et aux perturbations du canal. Grâce à sa simplicité architecturale et à son faible coût computationnel, AvgNet représente une solution efficace pour les environnements de communication contraints. Néanmoins, l'utilisation d'une agrégation uniforme peut limiter la capacité du modèle à distinguer l'importance relative des voisins, ce qui peut engendrer une perte d'informations discriminantes essentielles pour certaines modulations complexes [58].

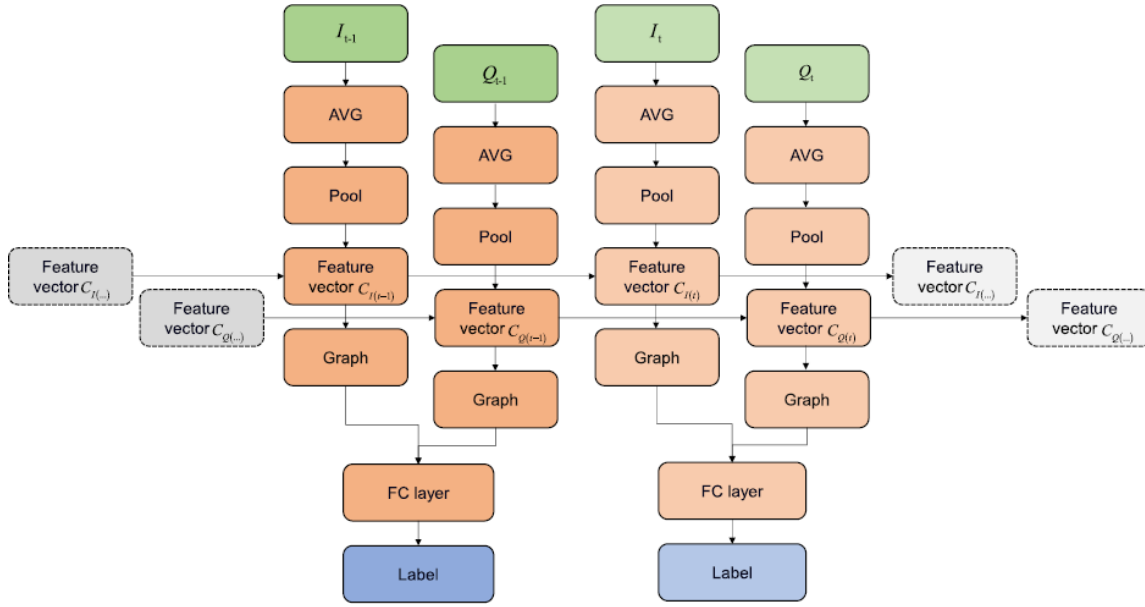


Figure III.21: Schéma de l'AVgNet.

III.6.2 GNN-FiLM

GNN-FiLM (Graph Neural Network with Feature-wise Linear Modulation) représente une extension plus avancée des architectures GNN appliquées à l'AMC, visant à améliorer la qualité de l'apprentissage des relations entre les échantillons du signal. Contrairement à AvgNet, qui adopte une simple agrégation moyenne, GNN-FiLM introduit un mécanisme de modulation linéaire des caractéristiques appelé FiLM (Feature-wise Linear Modulation). Ce mécanisme permet d'adapter dynamiquement l'influence des informations issues des nœuds voisins à travers deux paramètres d'apprentissage, γ (gamma) et β (beta), selon la relation linéaire suivant :

$$y = \gamma x + \beta \quad (15)$$

Cette modulation permet au réseau de pondérer sélectivement les caractéristiques les plus pertinentes et de réduire l'impact des informations moins significatives. Par conséquent, GNN-FiLM améliore la capacité de discrimination entre des schémas de modulation similaires, notamment pour les modulations d'ordre élevé telles que QPSK, 8PSK, 16QAM et 64QAM, en particulier dans des conditions de faible rapport signal sur bruit (SNR). Malgré ses performances supérieures en termes de précision et de robustesse, cette

architecture présente une complexité computationnelle plus importante ainsi qu'un coût d'entraînement plus élevé, ce qui peut limiter son déploiement dans les systèmes embarqués à ressources limitées [59].

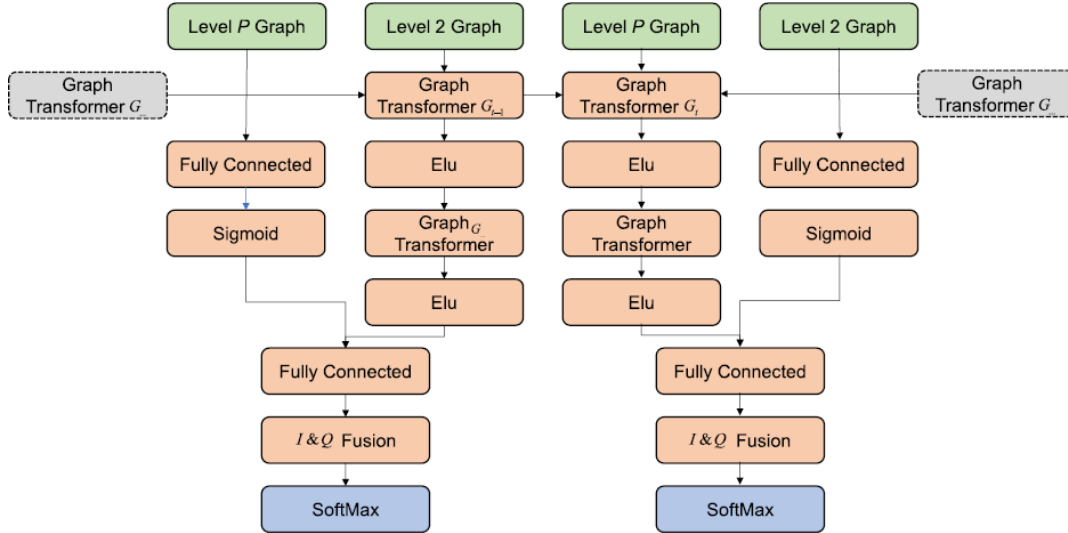


Figure III. 22:GNN-FILM.

III.7 Modèles basé sur Transformer

Le Transformer représente une architecture avancée de l'apprentissage profond basée principalement sur le mécanisme de self-attention, initialement développé pour le traitement du langage naturel, puis adapté avec succès à la classification automatique des modulations (AMC). Contrairement aux réseaux récurrents tels que les RNN et les LSTM, le Transformer ne traite pas les données de manière strictement séquentielle, mais réalise un traitement parallèle de l'ensemble de la séquence, ce qui améliore considérablement la vitesse d'apprentissage et d'inférence. Cette architecture est particulièrement efficace pour les signaux radio complexes, car elle permet de capturer les dépendances globales et les relations à longue distance entre les différentes parties du signal.

III.7.1 Typical model

Le modèle typique du Transformer se compose principalement de deux parties : l'encodeur (Encoder) et le décodeur (Decoder). L'entrée du modèle passe d'abord par une couche d'embedding qui transforme les données en vecteurs de représentation. Comme le Transformer ne prend pas naturellement en compte l'ordre des séquences, un mécanisme de positional encoding est ajouté afin d'introduire l'information de position dans les données d'entrée. Le cœur du modèle repose sur le mécanisme de self-attention, qui permet de calculer la corrélation entre chaque élément de la séquence et tous les autres éléments, afin d'attribuer des poids différents selon leur importance. Pour améliorer davantage la capacité de représentation, le Transformer utilise également le multi-head attention, qui permet d'analyser simultanément plusieurs types de relations à partir de différentes perspectives. Des connexions résiduelles (Residual Connections) ainsi que la normalisation de couche (Layer Normalization) sont également utilisées pour faciliter l'entraînement et améliorer la stabilité du modèle [60].

III.7.2 MC former

Dans le contexte de l'AMC, plusieurs variantes du Transformer ont été développées afin d'adapter cette architecture aux contraintes des signaux radio. Parmi elles, MCFormer (Multi-scale Cross Attention Former) constitue une architecture hybride avancée conçue pour améliorer la reconnaissance des modulations complexes. MCFormer combine les avantages du Transformer avec des mécanismes d'analyse multi-échelle, permettant au modèle de capturer simultanément les informations locales et globales du signal. Grâce au mécanisme de cross-attention, le modèle peut mieux exploiter les relations entre différentes représentations du signal, comme les données I/Q, les spectres et les caractéristiques temporelles. Cette approche améliore fortement la discrimination entre des modulations similaires telles que QPSK, 8PSK, 16QAM et 64QAM, notamment dans des conditions de faible rapport signal sur bruit (SNR). MCFormer offre ainsi une meilleure robustesse et une meilleure généralisation dans les environnements sans fil complexes.

Le Transformer permet globalement une meilleure reconnaissance des modulations complexes, en particulier dans les environnements affectés par le bruit, le fading ou les interférences multi-trajets. Il est particulièrement performant pour l'analyse des signaux à longue séquence, où les modèles classiques comme le CNN ou le LSTM présentent certaines limitations. Des architectures comme TDRA, IQFormer, MAMC et MCFormer illustrent cette évolution vers des modèles plus intelligents et plus performants pour l'AMC.

L'un des principaux avantages du Transformer réside dans sa capacité de calcul parallèle, qui supprime la dépendance séquentielle présente dans les RNN, ainsi que dans son aptitude à capturer efficacement les dépendances globales du signal. Le mécanisme de multi-head attention améliore également la robustesse du modèle en permettant une analyse multi-perspective des données d'entrée. Cependant, malgré ses excellentes performances, le Transformer présente une complexité computationnelle élevée ainsi qu'un grand nombre de paramètres à entraîner, ce qui rend son déploiement difficile sur les plateformes à ressources limitées comme les systèmes embarqués ou les dispositifs IoT. De plus, il nécessite généralement une grande quantité de données d'entraînement pour éviter le surapprentissage (overfitting) et garantir une bonne généralisation. Enfin, l'absence de partage de paramètres réduit son efficacité opérationnelle par rapport à certaines architectures plus légères comme le CNN [61].

III.8 Modèles hybrides de l'apprentissage profond

Les modèles hybrides de l'apprentissage profond (Hybrid Deep Learning Models) représentent une approche avancée dans la classification automatique des modulations (AMC), visant à combiner plusieurs architectures de réseaux de neurones afin d'exploiter les avantages spécifiques de chacune. Contrairement aux modèles individuels tels que CNN, LSTM ou Transformer, les modèles hybrides permettent d'associer plusieurs mécanismes d'apprentissage au sein d'un même système afin d'améliorer la précision de classification, la robustesse face au bruit ainsi que la capacité de généralisation dans des environnements de communication complexes. Cette stratégie est particulièrement utile

dans les systèmes radio sans fil, où les signaux présentent à la fois des caractéristiques spatiales, fréquentielles et temporelles qui nécessitent plusieurs niveaux d'analyse.

III.8.1 CNN-LSTM

Parmi les architectures hybrides les plus utilisées, le modèle CNN-LSTM constitue l'une des solutions les plus performantes pour l'AMC. Cette architecture combine les capacités des réseaux de neurones convolutifs (CNN) et des réseaux Long Short-Term Memory (LSTM). Le CNN intervient en première étape pour extraire automatiquement les caractéristiques locales et discriminantes du signal, telles que les motifs présents dans les représentations I/Q, les spectres ou les diagrammes de constellation. Grâce aux opérations de convolution et de pooling, il réduit la dimension des données tout en conservant les informations essentielles. Ensuite, les caractéristiques extraites sont transmises au réseau LSTM, chargé de modéliser les dépendances temporelles et les relations séquentielles du signal radio. Le LSTM permet ainsi de capturer efficacement les transitions de phase, les variations d'amplitude ainsi que les effets dynamiques du canal comme le fading, le bruit impulsionnel et les interférences. Cette combinaison améliore considérablement la reconnaissance des modulations complexes, notamment dans les environnements à faible rapport signal sur bruit (SNR). Toutefois, cette architecture présente une complexité computationnelle relativement élevée ainsi qu'un temps d'entraînement important [62].

III.8.2 GrrNet

Une autre architecture hybride importante est GrrNet, qui repose sur l'intégration des réseaux récurrents et des mécanismes avancés de représentation des caractéristiques afin d'améliorer la reconnaissance des signaux modulés. GrrNet a été conçu pour renforcer la capacité du modèle à apprendre les dépendances complexes entre les échantillons du signal tout en conservant les informations discriminantes essentielles à la classification. Cette architecture exploite une meilleure interaction entre l'extraction des caractéristiques profondes et la modélisation séquentielle, ce qui permet d'obtenir une meilleure robustesse face aux variations du canal et aux perturbations du bruit. GrrNet vise principalement à surmonter certaines limites des architectures classiques en améliorant la stabilité de

l'apprentissage et la précision de classification pour des schémas de modulation difficiles à distinguer. Son principal avantage réside dans sa capacité à mieux représenter les relations complexes entre les signaux, bien que cela s'accompagne également d'une augmentation de la complexité structurelle et du coût de calcul.

D'une manière générale, les architectures hybrides peuvent être organisées selon deux approches principales : la structure en cascade et la structure de fusion. Dans une structure en cascade, un premier modèle comme le CNN est utilisé pour extraire les caractéristiques importantes du signal, puis ces caractéristiques sont transmises à un second modèle comme le RNN ou le LSTM pour l'analyse séquentielle et la prise de décision finale. À l'inverse, dans une structure de fusion, les sorties de plusieurs modèles sont combinées à l'aide de méthodes telles que la concaténation, la fusion pondérée ou l'attention adaptative. Des architectures comme CNN-LSTM, ConvLSTMAE, LSTM-AE, GrrNet, Incept-LSTM et CTGNet illustrent parfaitement cette approche, qui permet une meilleure exploitation des signaux multimodaux ainsi qu'une amélioration significative des performances en présence de bruit et d'interférences.

Cependant, malgré leurs excellents résultats, les modèles hybrides présentent une complexité structurelle importante ainsi qu'un coût computationnel élevé. La combinaison de plusieurs architectures augmente le nombre de paramètres à apprendre, ce qui rend l'entraînement plus long et l'inférence plus coûteuse, notamment dans les systèmes embarqués ou les dispositifs IoT à ressources limitées. De plus, cette complexité réduit souvent l'interopérabilité du modèle, rendant difficile la compréhension de son processus de décision interne. Ce phénomène est généralement désigné par le terme de « boîte noire » (Black Box), qui constitue l'une des principales limites des approches de deep learning avancées [63].

III.9 Techniques de régularisation dans l'AMC

Bien que les méthodes basées sur le Deep Learning (DL) offrent de bonnes performances pour la reconnaissance des schémas de modulation, leur capacité de généralisation reste un défi important lorsqu'on change de jeux de données ou de conditions de communication.

En effet, les modèles de DL sont souvent sensibles à la distribution des données d'entraînement. Lorsqu'il existe une différence importante entre la distribution des données d'entraînement et celle des données de test (par exemple à cause des variations de canal), les performances du modèle peuvent diminuer de manière significative. Cela limite leur efficacité dans des environnements réels et variés.

De plus, dans les tâches de classification des modulations, les données sont souvent déséquilibrées, certaines classes de modulation étant plus représentées que d'autres. Cela peut entraîner un surapprentissage sur les classes majoritaires, au détriment des classes minoritaires, réduisant ainsi la capacité de généralisation globale du modèle.

Actuellement, l'évaluation de la robustesse des modèles AMC se fait principalement en analysant leurs performances dans différents environnements de propagation et sur plusieurs jeux de données.

Dans ce contexte, les techniques de régularisation jouent un rôle essentiel pour améliorer la généralisation des modèles. Elles peuvent être classées en deux grandes catégories : régularisation explicite et régularisation implicite.

III.9.1 Régularisation explicite (Explicit Regularization)

La régularisation explicite consiste à introduire directement des contraintes supplémentaires dans la fonction d'optimisation du modèle afin de limiter le surapprentissage et d'améliorer la généralisation. Cette approche agit principalement en pénalisant les solutions trop complexes et en favorisant des représentations plus robustes et plus stables. Dans le domaine de l'AMC, plusieurs travaux ont proposé des termes de régularisation spécifiques adaptés aux caractéristiques des signaux radio.

Par exemple, certains chercheurs ont conçu des facteurs de régularisation combinant la séparation inter-classes et la réduction de la variance intra-classe, afin d'augmenter la distance entre différentes modulations tout en rapprochant les échantillons appartenant à la même classe. Cette amélioration de la frontière de décision permet au modèle de mieux distinguer les schémas de modulation, même sous différentes conditions de SNR. D'autres approches utilisent des modules d'optimisation de canaux (Channel Optimization Module), capables d'ajuster dynamiquement les connexions entre les neurones du réseau à partir des informations issues des couches de Batch Normalisation et des cartes de caractéristiques, ce qui renforce la robustesse du modèle face aux variations des conditions de transmission. L'objectif principal de la régularisation explicite est donc de guider directement l'apprentissage vers des solutions mieux généralisables et moins dépendantes des particularités du jeu de données d'entraînement [64].

III.9.2 Régularisation implicite (Implicit Regularization)

Contrairement à la régularisation explicite, la régularisation implicite n'ajoute pas directement de terme supplémentaire dans la fonction de perte, mais améliore la généralisation à travers la stratégie d'apprentissage elle-même, l'architecture du modèle ou les techniques de prétraitement des données. Cette forme de régularisation repose sur des mécanismes indirects qui contribuent à réduire le risque de sur apprentissage.

Parmi les techniques les plus couramment utilisées, on retrouve le Dropout, qui consiste à désactiver aléatoirement certains neurones pendant l'entraînement afin d'éviter une dépendance excessive entre les unités du réseau. La Batch Normalisation constitue également une méthode importante, car elle stabilise la distribution des données entre les couches, accélère la convergence et améliore la robustesse du modèle. L'augmentation de données (Data Augmentation) joue également un rôle essentiel dans l'AMC, notamment par l'ajout de bruit, la variation du SNR, la suppression aléatoire de certains échantillons ou encore la perturbation des nœuds dans les architectures GNN. Ces techniques permettent de simuler des conditions de communication plus réalistes et d'augmenter la diversité des données d'apprentissage [65].

III.10 Conclusion

Ce chapitre a montré que l'apprentissage profond est une solution puissante pour la classification des signaux modulés, mais son efficacité dépend fortement du choix des représentations de signaux. Les méthodes classiques comme les signaux I/Q et les constellations restent utiles, mais elles doivent souvent être complétées par des modèles adaptés. Les architectures CNN, LSTM, GNN et Transformers exploitent chacune des aspects différents du signal spatial, temporel ou structurel et leur combinaison dans des modèles hybrides permet d'obtenir de meilleures performances dans des environnements difficiles.

Cependant, ces approches avancées demandent plus de calculs et de données d'entraînement, ce qui peut limiter leur utilisation dans des systèmes embarqués ou des appareils IoT. Le défi est donc de trouver un équilibre entre précision, rapidité et simplicité, afin de rendre ces modèles utilisables dans la pratique. En résumé, l'évolution des représentations et des architectures de Deep Learning ouvre la voie à des systèmes de communication plus intelligents, robustes et fiables.

Chapitre 04 : Résultats et discussions.

IV.1 Introduction

Ce chapitre expose la conception de notre architecture hybride dédiée à la reconnaissance automatique des modulations (AMC). Pour répondre aux défis de l'identification des signaux dans des environnements bruités, nous proposons un modèle articulé autour de la synergie entre CNN, Transformer, BiGRU et Mécanisme d'Attention. Évaluée sur la base de données RML2016.10a, notre approche se distingue par une exploitation optimale des dimensions spatiales et séquentielles du signal à travers quatre étapes clés :

- Enrichissement de l'entrée : Transformation des signaux bruts I/Q (128×2) en une matrice de caractéristiques (128×5) intégrant l'amplitude, la phase et le SNR.
- Extraction locale et globale : Un bloc CNN extrait d'abord les motifs locaux et réduit la dimension temporelle, tandis qu'un double bloc Transformer capture les interdépendances à longue portée via un mécanisme d'attention multi-têtes.
- Modélisation contextuelle : Un réseau BiGRU récurrent bidirectionnel analyse les dépendances séquentielles, couplé à un mécanisme d'attention qui pondère les instants temporels les plus discriminants.
- Classification : Une tête de classification dense basée sur une fonction Softmax qui discrimine avec précision les 11 classes de modulation cibles.

La suite de ce chapitre détaille l'implémentation, les hyperparamètres et la justification théorique de chacun de ces blocs.

IV.2 La base de données RadioML 2016.10a

IV.2.1 Description de la base de données

Dans ce travail, nous avons utilisé la base de données RadioML 2016.10a, qui constitue une référence publique largement utilisée dans les recherches sur la classification automatique de la modulation (AMC) basée sur les techniques d'apprentissage profond.

Cette base de données a été générée par simulation à l'aide de GNU Radio par l'entreprise DeepSig. Elle a été conçue afin de reproduire des conditions réalistes de propagation des

signaux dans les systèmes de communications sans fil, tout en conservant un environnement expérimental contrôlé. Grâce à sa richesse et à sa pertinence, elle est devenue une référence majeure dans le domaine de la reconnaissance automatique des signaux radio.

La base de données contient environ 220 000 échantillons de signaux, générés par des simulations de type Monte Carlo représentant différents scénarios de transmission dans des environnements radio réalistes.

Ces échantillons sont répartis de manière équilibrée entre plusieurs types de modulations ainsi que différentes valeurs du rapport signal sur bruit (SNR), ce qui permet de garantir un bon équilibre statistique et de réduire les biais lors de la phase d'apprentissage.

Les valeurs du SNR varient de -20 dB à $+18$ dB, avec un pas de 2 dB, ce qui permet d'évaluer les performances des modèles de classification aussi bien dans des conditions fortement bruitées que dans des environnements plus favorables.

Sur le plan structurel, chaque échantillon est représenté sous forme d'une matrice de dimension 128×2 , où chaque canal correspond aux composantes du signal en phase (I) et en quadrature (Q). Cette représentation permet aux modèles d'apprentissage profond d'extraire directement les caractéristiques pertinentes à partir des données brutes, sans nécessiter d'étapes de prétraitement complexes [66].

IV.2.2 Types de modulations et représentation du signal

La base de données RadioML 2016.10a comprend une grande variété de schémas de modulation numériques et analogiques comme il est illustré sur la figure IV.24, permettant une évaluation complète sur différents modèles de communication. Les classes de modulation incluses sont :

- ❖ **Modulation par déplacement de phase (PSK)** : comprend BPSK, QPSK et 8PSK, où l'information numérique est codée à travers les variations de phase du signal.
- ❖ **Modulation d'amplitude en quadrature (QAM)** : comprend 16QAM et 64QAM, dans lesquelles l'amplitude et la phase du signal varient simultanément, offrant une meilleure efficacité spectrale.

- ❖ **Modulation par déplacement de fréquence (FSK) :** comprend BFSK et CPFSK, reposant sur des variations discrètes de fréquence.
- ❖ **Modulation d'amplitude pulsée (PAM) :** comprend PAM4, où l'information est représentée par plusieurs niveaux d'amplitude.
- ❖ **Modulations analogiques :** incluent AM-DSB, AM-SSB et WBFM, représentant les techniques classiques de communication analogique.



Figure IV. 23: types des onze modulations

IV.2.3 Effets du canal et distorsions du signal

L'une des caractéristiques majeures de RadioML 2016.10a est l'intégration de perturbations réalistes du canal de transmission. Contrairement aux jeux de données synthétiques idéalisés, cette base de données introduit plusieurs effets de propagation et de

synchronisation fréquemment observés dans les systèmes de communication sans fil. Parmi ces distorsions :

- ❖ **Bruit blanc gaussien additif (AWGN)** : simule le bruit thermique naturellement présent dans les canaux de communication.
- ❖ **Évanouissement multi-trajets (multipath fading)** : modélise les réflexions du signal et les retards de propagation dus aux obstacles de l'environnement.
- ❖ **Décalage de fréquence (frequency offset)** : représente le désalignement des oscillateurs entre l'émetteur et le récepteur.
- ❖ **Décalage du taux d'échantillonnage** : simule les différences de synchronisation temporelle entre équipements de communication.
- ❖ **Rotation aléatoire de phase** : introduit une incertitude de phase afin de reproduire des déformations réelles du signal.

Ces effets de canal augmentent considérablement la complexité de la tâche de classification et rendent le jeu de données plus représentatif des scénarios réels de communication sans fil.

IV.2.4 Importance de base de données pour les applications de Deep Learning

RadioML 2016.10a est particulièrement adapté aux approches d'apprentissage profond, car il conserve la structure séquentielle et temporelle des signaux. Il peut être directement exploité par des architectures telles que les réseaux de neurones convolutifs (CNN), les réseaux récurrents (RNN), les modèles à mémoire longue durée (LSTM) ainsi que les architectures hybrides.

Son organisation équilibrée et la présence de distorsions réalistes permettent aux chercheurs d'évaluer la robustesse des modèles, leur capacité de généralisation ainsi que leurs performances de classification dans différentes conditions de canal.

En outre, cette base de données est largement utilisée comme benchmark pour comparer différentes architectures et stratégies d'optimisation dans les tâches de reconnaissance automatique de modulation. Malgré ses nombreux avantages, certaines limites ont été

signalées dans la littérature, notamment la confusion entre des schémas de modulation similaires tels que 16QAM et 64QAM à faible SNR. Néanmoins, RadioML 2016.10a demeure l'un des jeux de données les plus utilisés dans les recherches sur les communications sans fil intelligentes et les applications de radio cognitive.

IV.3 Prétraitement de données

Avant de commencer l'entraînement, la base de données RadioML 2016.10a a passé par plusieurs étapes de préparation afin d'améliorer les performances du modèle et d'assurer la cohérence des informations. Les expériences et l'entraînement du modèle ont été réalisés sur la plateforme Google Colab, qui offre un environnement de calcul performant basé sur Python. Afin d'accélérer le processus d'entraînement et d'améliorer les performances de calcul, les ressources GPU fournies par cet environnement ont été exploitées. Afin de préparer les données pour la phase d'apprentissage, les étapes de prétraitement suivantes ont été appliquées :

- **Chargement du dataset IQ** : les données sont importées depuis un fichier pickle contenant les signaux I/Q associés à différentes modulations et valeurs de SNR, puis séparées en signaux, labels et SNR.
- **Réorganisation temporelle** : les données sont restructurées pour représenter chaque signal comme une séquence temporelle de 128 échantillons avec 2 composantes (I et Q), afin de les adapter aux modèles séquentiels.
- **Extraction amplitude + phase** : à partir des composantes I et Q, deux nouvelles caractéristiques sont calculées, l'amplitude (intensité du signal) et la phase (information angulaire), afin d'enrichir la représentation du signal.
- **Ajout du SNR comme feature** : le rapport signal/bruit est normalisé puis répété sur toute la séquence temporelle pour être intégré comme une caractéristique supplémentaire.
- **Normalisation standard** : toutes les caractéristiques sont centrées et réduites (moyenne nulle et variance unitaire) pour améliorer la stabilité et la convergence du modèle.
- **Encodage des labels** : cette technique consiste à transformer les classes de modulation d'un format textuel vers une représentation numérique, afin de permettre aux modèles de traiter et de classifier les données catégorielles.

- **Split train/test** : la base de données a été séparé en 60% pour l'entraînement, 20% pour validation et 20% pour le test en utilisant un échantillonnage équilibré. Cela assure une bonne répartition des classes de modulation dans les deux ensembles et évite les biais lors de l'évaluation.

IV.4 Métriques de performance en classification

L'évaluation des performances des modèles de classification constitue une étape essentielle visant à mesurer la fiabilité du modèle ainsi que sa capacité de généralisation sur des données inconnues. Elle repose sur un ensemble de mesures statistiques permettant d'analyser la qualité des prédictions et de comparer objectivement différents modèles, parmi lesquelles on peut citer les suivantes :

IV.4.1 Matrice de confusion

La matrice de confusion constitue un outil fondamental dans l'évaluation des performances des modèles de classification. Elle permet d'analyser de manière détaillée la correspondance entre les prédictions du modèle et les valeurs réelles observées.

Elle est généralement représentée sous la forme d'un tableau croisé distinguant quatre cas:

- **Vrai Positif (VP)** : instances positives correctement classifiées
- **Vrai Négatif (VN)** : instances négatives correctement classifiées
- **Faux Positif (FP)** : instances négatives incorrectement classifiées comme positives (erreur de type I)
- **Faux Négatif (FN)** : instances positives incorrectement classifiées comme négatives (erreur de type II)

Cette représentation constitue la base de calcul de l'ensemble des métriques d'évaluation utilisées en apprentissage automatique [67].

IV.4.2 Accuracy (Exactitude)

L'accuracy, ou exactitude, mesure la proportion de prédictions correctes parmi l'ensemble des prédictions effectuées par le modèle :

$$Accuracy = \frac{VP + VN}{FP + FNVP + VN} \quad (14)$$

Cette métrique est pertinente lorsque les classes sont équilibrées. Toutefois, elle peut s'avérer trompeuse dans des contextes de déséquilibre de classes, où une classe majoritaire domine largement les données. [67]

IV.4.3 Précision (Precision)

La précision évalue la proportion d'instances correctement identifiées comme positives parmi l'ensemble des instances prédites positives :

$$Precision = \frac{VP}{VP + FP} \quad (15)$$

Elle permet d'estimer la fiabilité des prédictions positives du modèle. Une précision élevée indique que le modèle génère peu de faux positifs.

Cette métrique est particulièrement importante dans les contextes où le coût des faux positifs est élevé [67].

IV.4.4 Rappel (Recall ou Sensibilité)

Le rappel, également appelé sensibilité, mesure la capacité du modèle à identifier correctement les instances positives :

$$Recall = \frac{VP}{VP + FN} \quad (16)$$

Il reflète la proportion des cas positifs effectivement détectés par le modèle. Un rappel élevé signifie que peu de cas positifs sont manqués.

Cette métrique est essentielle dans les applications où les faux négatifs ont des conséquences critiques [67].

IV.4.5 F1-Score

Le F1-score correspond à la moyenne harmonique de la précision et du rappel :

$$F1 = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (17)$$

Contrairement à la moyenne arithmétique, la moyenne harmonique pénalise les écarts importants entre les deux mesures. Ainsi, le F1-score fournit une évaluation plus équilibrée des performances du modèle, en particulier dans le cas de données déséquilibrées [67].

IV .5 Architecture conceptuelle du modèle de classification automatique des modulations

Dans le cadre de l'amélioration des performances de la classification automatique des modulations des signaux radio, nous proposons une architecture hybride combinant plusieurs techniques complémentaires d'apprentissage profond. Ce modèle intègre successivement des couches CNN, des blocs Transformer, des couches BiGRU, ainsi qu'un mécanisme d'attention (Attention). Le choix de cette architecture repose sur la complémentarité de ces approches : les couches CNN permettent d'extraire les caractéristiques locales du signal, tandis que les Transformers capturent les dépendances à long terme. De leur côté, les BiGRU modélisent les relations temporelles dans les deux directions, alors que le mécanisme d'attention met en évidence les informations les plus pertinentes pour la tâche de classification. Cette fusion vise à exploiter efficacement les différentes propriétés des signaux radio afin d'améliorer la robustesse et la précision du système de reconnaissance des types de modulation, comme le montre la figure suivante représentant l'architecture globale du modèle proposé

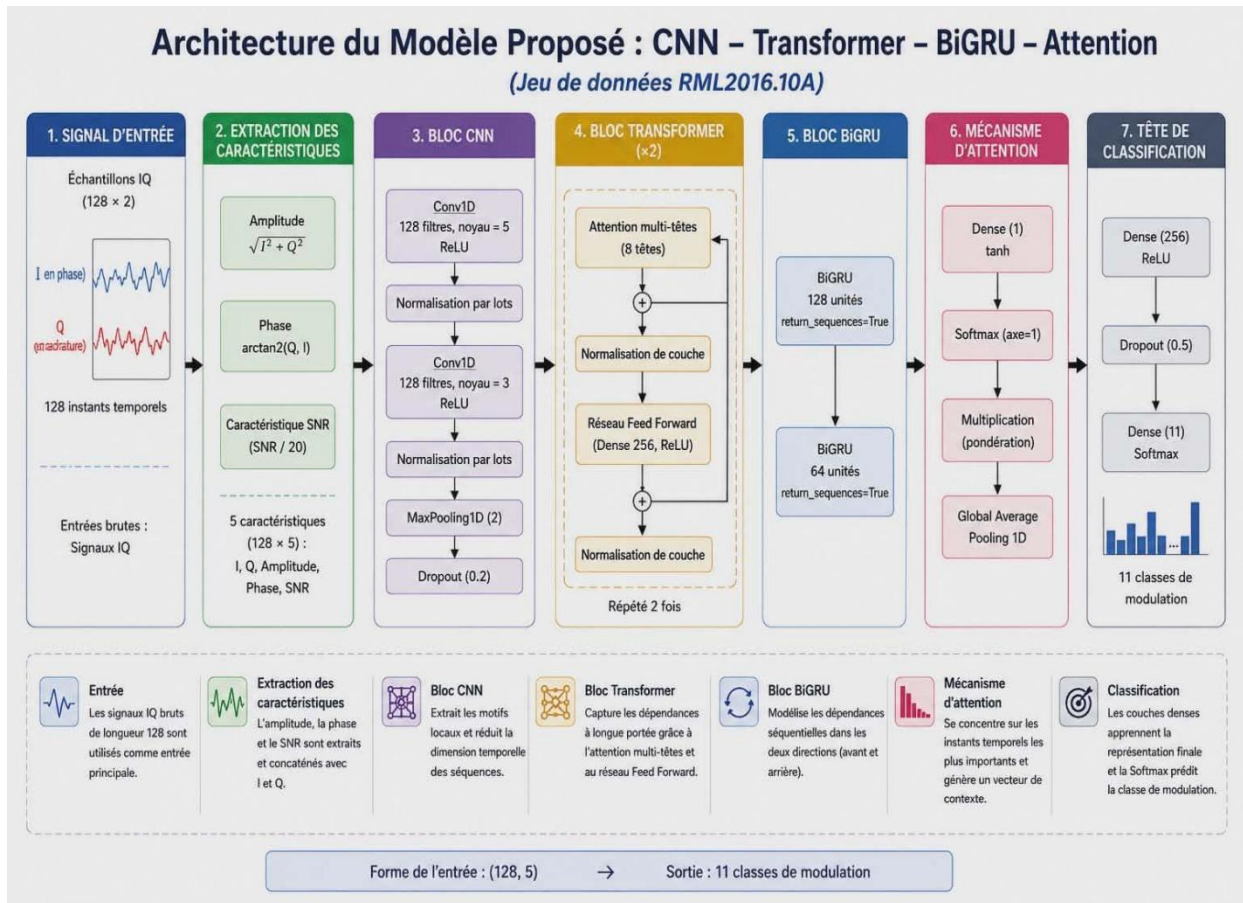


Figure iv.24: Architecture conceptuelle du modèle hybride CNN–Transformer–BiGRU–Attention pour la classification automatique des modulations.

IV.5.1 présentations de l'architecture de modèle proposé

IV.5.1.1 CNN

Les CNN sont des modèles d'apprentissage profond conçus pour traiter des données structurées en grille, comme les images, et peuvent être adaptés en 1D-CNN pour les signaux séquentiels tels que les signaux I/Q [68].

Dans ce travail Le bloc CNN (Convolutional Neural Network) a pour rôle principal d'extraire automatiquement les caractéristiques locales présentes dans le signal. Il débute par une première couche de convolution Conv1D composée de 128 filtres de taille 5, suivie de la fonction d'activation ReLU, qui permet d'introduire de la non-linéarité et d'améliorer la capacité d'apprentissage du modèle. Une étape de Batch Normalization est ensuite

appliquée afin de stabiliser les distributions des données et d'accélérer la convergence durant l'entraînement. Une deuxième couche Conv1D, également constituée de 128 filtres mais avec une taille de noyau de 3, permet d'extraire des caractéristiques plus fines et plus détaillées. Par la suite, l'opération de MaxPooling réduit la dimension temporelle des données tout en conservant les informations les plus pertinentes. Enfin, une couche Dropout de 0,2 est utilisée pour diminuer le risque de surapprentissage en désactivant aléatoirement une partie des neurones pendant l'apprentissage. Grâce à cette architecture, le CNN est capable de détecter efficacement les motifs locaux et les signatures caractéristiques des différentes modulations présentes dans le signal.

IV .5.1.2 Transformer

Le Transformer est une architecture de Deep Learning basée sur le mécanisme de self-attention, utilisée pour le traitement des données séquentielles comme les signaux radio. Il permet d'apprendre les relations entre tous les éléments d'une séquence en même temps, sans traitement séquentiel comme les RNN ou LSTM [69].

Dans ce bloc Le Transformer est utilisé après le CNN pour analyser les relations entre les différentes parties du signal, même lorsqu'elles sont éloignées dans le temps. Contrairement au CNN qui se concentre sur les caractéristiques locales, le Transformer permet de comprendre le contexte global de la séquence. Il s'appuie sur un mécanisme d'attention multi-têtes composé de 8 têtes, qui lui permet d'identifier simultanément plusieurs relations importantes dans le signal. Les informations obtenues sont ensuite traitées par une couche Feed Forward de 256 neurones avec l'activation ReLU, puis normalisées pour améliorer la stabilité de l'apprentissage. Cette opération est répétée deux fois afin d'obtenir une représentation plus riche et plus complète du signal.

IV .5.1.3Le BiGRU (Bidirectional Gated Recurrent Unit)

Le BiGRU (Bidirectional Gated Recurrent Unit) est une architecture de réseaux de neurones récurrents utilisée pour traiter les données séquentielles comme les signaux ou les séries temporelles. Il analyse la séquence dans deux directions (avant et arrière), ce qui permet de capturer à la fois le contexte passé et futur de chaque instant [70].

Dans ce bloc Les caractéristiques extraites sont ensuite envoyées vers un réseau BiGRU, spécialisé dans l'analyse des données séquentielles. Grâce à son fonctionnement bidirectionnel, il traite le signal dans les deux sens (du passé vers le futur et du futur vers le passé), ce qui lui permet de mieux comprendre l'évolution temporelle du signal. Composé de deux couches de 128 puis 64 unités, ce bloc améliore la capture des informations temporelles importantes avant les étapes suivantes du modèle.

IV .5.1.4 Le mécanisme d'attention

Le mécanisme d'attention est une technique de Deep Learning qui permet d'attribuer un poids d'importance différent à chaque élément d'une séquence, afin de se concentrer sur les informations les plus pertinentes pour une tâche donnée [71].

Dans ce bloc Le mécanisme d'attention permet d'identifier les parties les plus importantes du signal pour la classification. Il attribue un poids à chaque instant temporel afin de mettre en valeur les informations les plus pertinentes et de réduire l'influence des moins utiles. Les informations ainsi pondérées sont ensuite résumées en un vecteur compact grâce à une opération de Global Average Pooling, qui sera utilisé pour la classification finale.

IV .5.1.5 Tête de classification

Dans ce bloc Le vecteur produit par le mécanisme d'attention est envoyé vers le bloc final de classification. Une couche dense de 256 neurones avec activation ReLU apprend une représentation de haut niveau des caractéristiques extraites. Une couche Dropout (0,5) est ensuite utilisée pour améliorer la capacité de généralisation du modèle.

La dernière couche est une couche dense de 11 neurones associée à une fonction Softmax. Chaque neurone correspond à une classe de modulation du jeu de données. La sortie Softmax fournit une distribution de probabilités sur les 11 classes, et la classe ayant la probabilité la plus élevée est choisie comme prédiction finale.

IV .6 Analyse de résultats obtenus

IV .6.1 Accuracy et Loss

Les résultats globaux montrent que le modèle atteint une précision de test d'environ 67,80 %, ce qui reflète une capacité relativement satisfaisante à généraliser sur des données inédites, tout en restant cohérent avec les performances observées durant l'apprentissage. L'analyse des courbes de précision révèle que les performances d'entraînement et de validation, initialement faibles (38,6 % et 52,1 % respectivement), augmentent rapidement au cours des premières époques avant de se stabiliser progressivement à partir de l'époque 10, où la précision d'entraînement atteint environ 63,4 % et celle de validation 64,8 %. La proximité entre les courbes d'entraînement et de validation jusqu'à l'époque 30, avec une précision finale d'environ 69,7 % pour l'entraînement et 67,6 % pour la validation, indique un bon équilibre du modèle. Concernant la fonction de perte (Loss), les valeurs diminuent fortement au début du processus d'apprentissage (passant de 1,56 à 0,78 pour l'entraînement), puis convergent progressivement vers environ 0,78 pour l'entraînement et 0,85 pour la validation, ce qui confirme l'efficacité du processus d'optimisation des poids. En outre, l'absence d'un écart significatif entre les courbes d'entraînement et de validation, ainsi que la stabilité de la Validation Loss en fin d'entraînement, témoignent clairement de l'absence de surapprentissage (Overfitting) et d'une bonne capacité de généralisation. Par ailleurs, les courbes atteignent un plateau après environ 25 époques, suggérant que prolonger l'entraînement au-delà de 30 époques n'apporterait probablement pas d'amélioration notable, d'autant plus que le taux d'apprentissage a été réduit à 3×10^{-5} à partir de l'époque 25, ce qui a permis une légère amélioration de la précision. Ainsi, afin d'améliorer la précision globale au-delà de 67,8 %, il serait plus pertinent d'optimiser les hyperparamètres, notamment le taux d'apprentissage (Learning Rate), ainsi que de renforcer l'architecture du modèle en ajoutant des couches ou en augmentant le nombre de neurones afin de mieux représenter la complexité des données, particulièrement pour les classes difficiles comme WBFM (F1-score de 0,39) et AM-DSB (F1-score de 0,57). L'analyse par rapport au SNR montre que le modèle atteint une précision maximale de 93,9 % à 2 dB, tandis que ses performances chutent à 11,7 % à -20 dB, ce qui indique une sensibilité importante au bruit pour les signaux de très faible qualité. La figure ci-dessous présente l'évolution des performances du modèle proposé au cours du

processus d'apprentissage, à travers les courbes de précision (Accuracy) et de fonction de perte (Loss) pour les ensembles d'entraînement (Train) et de validation (Validation) en fonction du nombre d'époques (Epochs).

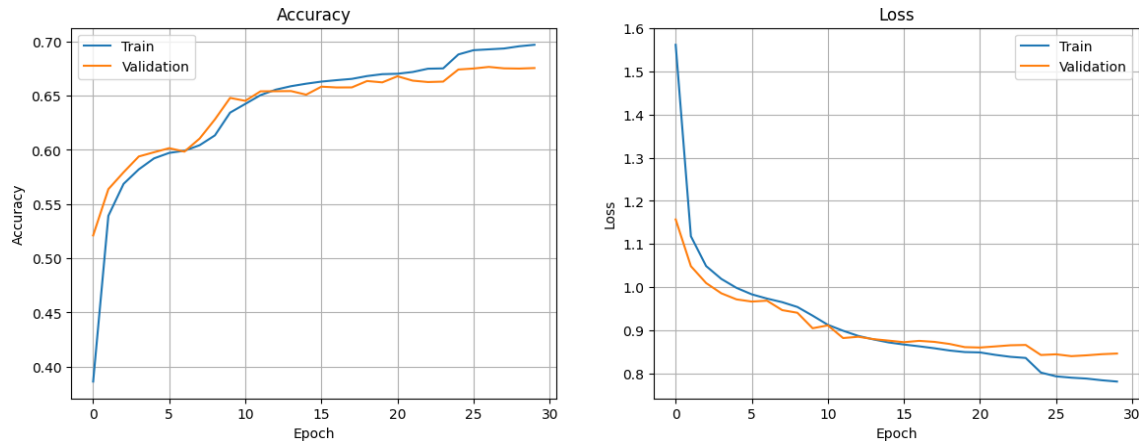


Figure IV.25: Évolution des performances du modèle durant l'apprentissage

IV .6.2 Accuracy vs SNR

La figure IV.27 montre l'influence du rapport signal sur bruit (SNR) sur la précision de classification du modèle proposé. Les résultats indiquent une amélioration progressive de la précision avec l'augmentation du SNR, confirmant l'impact direct de la qualité du signal sur les performances du modèle. Pour les faibles valeurs de SNR (-20 dB à -12 dB), la précision reste limitée (10 % à 30 %) en raison de l'effet important du bruit. Entre -10 dB et -4 dB, une nette amélioration est observée, la précision passant d'environ 43 % à 80 %, ce qui met en évidence la capacité du modèle CNN + Transformer + BiGRU + Attention à extraire efficacement les caractéristiques du signal malgré le bruit. À partir de 0 dB, la précision se stabilise à un niveau élevé (92 % à 94 %) jusqu'à +18 dB, traduisant la robustesse et l'efficacité du modèle dans des conditions de signal plus favorables. Les résultats obtenus sont représentés dans la figure ci-dessous.

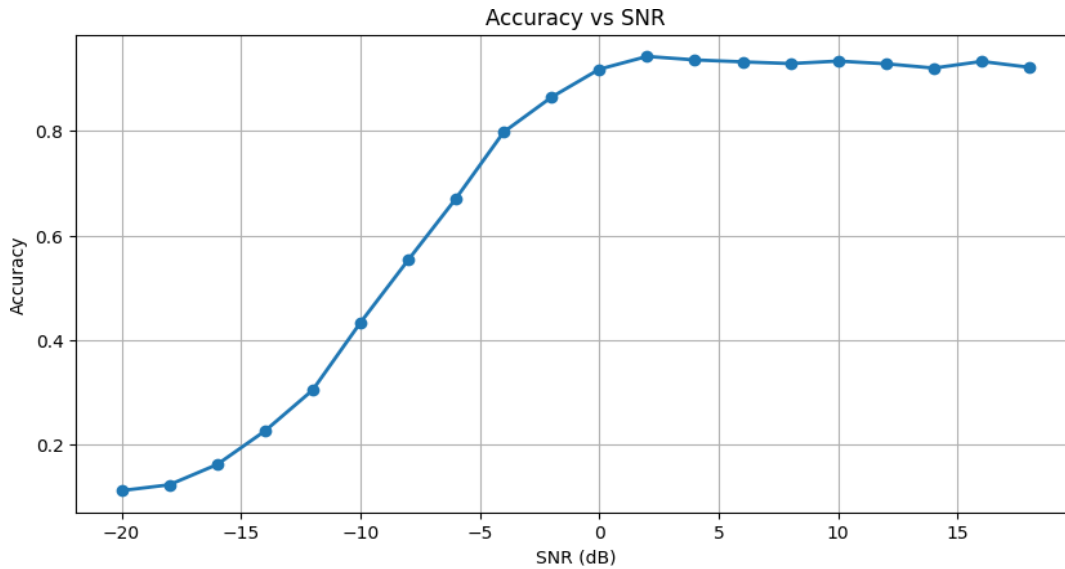


Figure IV.26:Évolution de l'accuracy du modèle en fonction du rapport signal sur bruit (SNR).

IV .6.3 La matrice de confusion

L'analyse de la matrice de confusion sur la figure IV.28 montre que le modèle de classification automatique des modulations présente de bonnes performances globales. La majorité des valeurs élevées sont concentrées sur la diagonale principale, ce qui indique qu'un grand nombre d'échantillons ont été correctement classifiés. Les modulations PAM4 (2 745 classifications correctes sur 4 000, soit 68,6 %), QAM64 (2 722, soit 68,1 %), QAM16 (2579, soit 64,5 %), BPSK (2 375, soit 59,4 %) et 8PSK (2 344, soit 58,6 %) obtiennent les meilleurs résultats en termes de classifications correctes, ce qui démontre la capacité du modèle à extraire efficacement les caractéristiques discriminantes de ces signaux. Toutefois, certaines erreurs de classification subsistent entre des modulations partageant des propriétés similaires. La confusion la plus importante est observée entre WBFM et AM-DSB, où 585 échantillons WBFM sont classés à tort comme AM-DSB, et 442 échantillons AM-DSB sont confondus avec WBFM, ce qui représente une confusion bilatérale significative due à leurs caractéristiques spectrales similaires. Des confusions sont également constatées entre les modulations de la famille PSK, notamment BPSK avec 632 échantillons confondus avec QPSK, QPSK avec 570 échantillons confondus avec BPSK et 8PSK avec 402 échantillons confondus avec QPSK, en raison de la similarité de

leurs caractéristiques de phase. Un chevauchement important apparaît également entre CPFSK et GFSK avec respectivement 398 et 454 échantillons mal classés entre ces deux classes, ce qui est cohérent avec leur appartenance à la même famille de modulations fréquentielles. Les modulations les plus difficiles à classer sont WBFM avec seulement 1 061 classifications correctes sur 4 000 (26,5 % de rappel) et AM-SSB avec 1 876 classifications correctes (46,9 %), indiquant que ces modulations présentent des caractéristiques moins discriminantes pour le modèle. Dans l'ensemble, ces résultats confirment l'efficacité du modèle tout en mettant en évidence certaines difficultés liées à la similarité intrinsèque entre certaines classes, particulièrement pour les modulations analogiques (WBFM, AM-DSB, AM-SSB) qui montrent des performances inférieures aux modulations numériques. La figure ci-dessous présente la matrice de confusion non normalisée du modèle proposé, illustrant le nombre réel d'échantillons correctement et incorrectement classifiés pour chaque type de modulation.

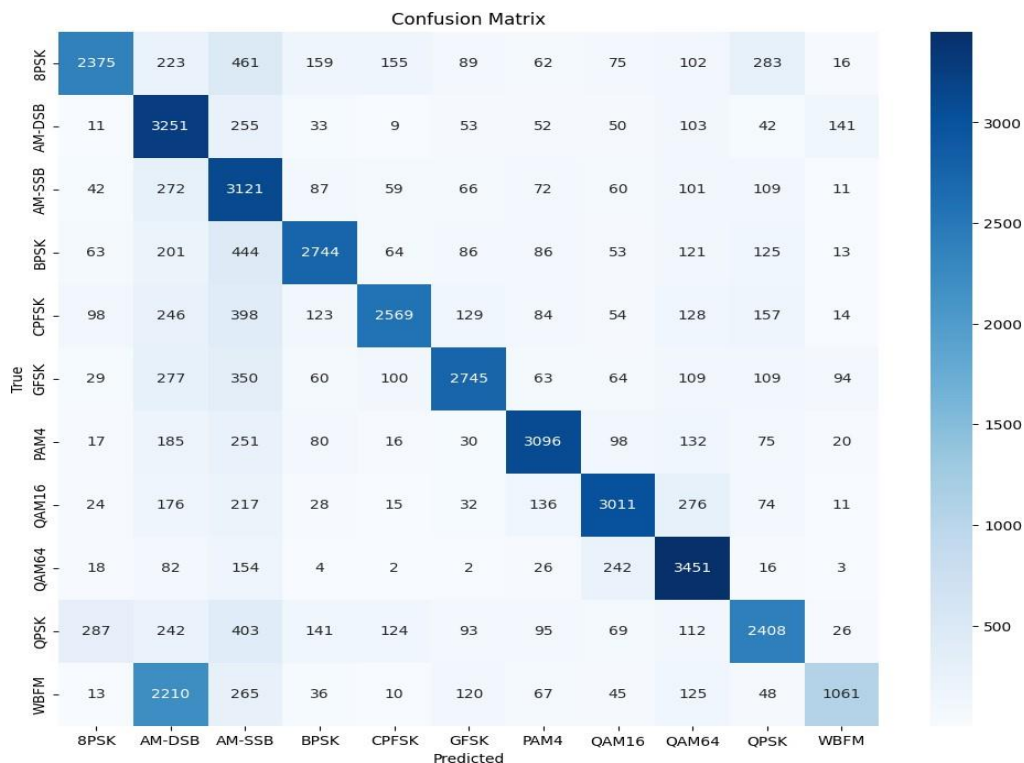


Figure IV.27: Matrice de confusion du modèle proposé

IV .6.4 Le rapport de classification

Le rapport de classification sur la figure (IV.29) montre que le modèle proposé atteint une accuracy globale de 67,80 % sur un ensemble de test équilibré de 44 000 échantillons répartis sur 11 classes de modulation, ce qui traduit une capacité satisfaisante à reconnaître les signaux dans différentes conditions de bruit. Les moyennes macro et pondérée sont identiques grâce à l'équilibre des classes, avec une précision moyenne de 71,89 %, un rappel de 67,80 % et un F1-score de 67,60 %. Les meilleures performances sont obtenues pour les modulations PAM4 (F1-score de 0,7899), QAM64 (0,7879), QAM16 (0,7700) et GFSK (0,7374), tandis que BPSK (0,7322), CPFSK (0,7213) et 8PSK (0,6808) présentent également des résultats satisfaisants.

En revanche, QPSK (F1-score de 0,6468) et AM-DSB (0,5721) enregistrent des performances plus limitées en raison respectivement de la similarité de leurs constellations et de leurs caractéristiques spectrales. Les performances les plus faibles concernent les modulations analogiques, particulièrement WBFM avec un F1-score de 0,3922 et un rappel de seulement 26,52 %, ainsi que AM-DSB (0,5721) et AM-SSB (0,6049), en raison de leur sensibilité au bruit et de leur variabilité intrinsèque. Globalement, les résultats confirment une meilleure reconnaissance des modulations numériques par rapport aux modulations analogiques, avec des performances qui s'améliorent lorsque le SNR augmente, démontrant ainsi l'efficacité et la bonne capacité de généralisation du modèle proposé. La figure ci-dessous présente le rapport de classification du modèle proposé, mettant en évidence les performances de chaque classe de modulation ainsi que les métriques globales obtenues sur l'ensemble de test.

	precision	recall	f1-score	support
8PSK	0.7978	0.5938	0.6808	4000
AM-DSB	0.4414	0.8127	0.5721	4000
AM-SSB	0.4939	0.7802	0.6049	4000
BPSK	0.7851	0.6860	0.7322	4000
CPFSK	0.8226	0.6422	0.7213	4000
GFSK	0.7968	0.6863	0.7374	4000
PAM4	0.8065	0.7740	0.7899	4000
QAM16	0.7880	0.7528	0.7700	4000
QAM64	0.7250	0.8628	0.7879	4000
QPSK	0.6988	0.6020	0.6468	4000
WBFM	0.7525	0.2652	0.3922	4000
accuracy			0.6780	44000
macro avg	0.7189	0.6780	0.6760	44000
weighted avg	0.7189	0.6780	0.6760	44000

Figure IV.28:Rapport de classification du modèle proposé

IV.7 Comparaison avec des modèles des travaux récents

Le tableau IV.3 présente une comparaison des performances du modèle proposé avec celles des principales approches de l'état de l'art pour la classification automatique des modulations radio. Le modèle proposé atteint une précision globale de 67,80 %, surpassant les approches de référence SMTrans (63,27 %), MCLDNN (62,08 %), PET-CGDNN (60,71%), AMC-Net (61,26 %), FEA-T (59,86 %) et CBADNN (64,02 %), soit une amélioration allant jusqu'à +7,94 % par rapport à FEA-T et +3,78 % par rapport à CBADNN, la meilleure référence. Concernant la précision minimale, le modèle proposé enregistre 11,73 %, la plus élevée parmi toutes les approches comparées (contre 9,25 % à 9,94 % pour les références), traduisant une meilleure robustesse face au bruit et une capacité supérieure à reconnaître les modulations les plus complexes ou fortement affectées par un faible SNR. Enfin, la précision maximale atteint 93,87 %, dépassant très légèrement CBADNN (93,66 %) et SMTrans (93,34 %), confirmant la capacité du modèle à exceller également dans des conditions favorables. Ces résultats, supérieurs sur l'ensemble des trois métriques évaluées, témoignent de l'efficacité de l'architecture hybride proposée, qui combine une extraction de caractéristiques spatiales par CNN, une modélisation des dépendances globales par Transformer, une analyse temporelle bidirectionnelle par BiGRU et un mécanisme d'attention focalisant sur les informations les plus pertinentes pour la décision.

	Overall accuracy	Min accuracy	Max accuracy
SMTrans [72]	63.27	9.52	93.34
MCLDNN [73]	62.08	9.25	92.71
PET-CGDNN [74]	60.71	9.40	91.07
AMC-Net [75]	61.26	9.32	90.64
FEA-T [76]	59.86	9.42	89.08
CBADNN [77]	64.02	9.94	93.66
Cnn+transformer+bigru + attention	67.80	11.73	93.87

Tableau IV.3: Comparaison des performances des modèles de classification

IV .8 Conclusion

Dans ce chapitre, nous avons présenté l'implémentation et l'évaluation de notre modèle hybride pour la classification automatique des modulations (AMC). Les expériences ont été réalisées sur la base de données RadioML 2016.10a, qui contient 11 types de modulations analogiques et numériques affectées par différentes perturbations du canal de transmission.

Les résultats obtenus montrent que la combinaison des couches CNN, Transformer, BiGRU et du mécanisme d'attention permet d'atteindre une précision globale de 67,80 % sur l'ensemble de test. Les courbes d'apprentissage indiquent que le modèle converge de manière stable à partir de la 10^e époque, sans signe important de surapprentissage.

L'analyse de la matrice de confusion et du rapport de classification met en évidence de bonnes performances pour les modulations PAM4, QAM16 et QAM64. En revanche, certaines erreurs de classification subsistent entre des modulations ayant des caractéristiques proches, notamment QPSK et 8PSK, surtout lorsque le rapport signal sur bruit (SNR) est faible.

Dans l'ensemble, les résultats confirment l'efficacité du modèle proposé. Le mécanisme d'attention contribue à mieux exploiter les informations temporelles des signaux I/Q, ce qui améliore les capacités de reconnaissance du système et montre son potentiel pour les applications de communications sans fil intelligentes.

Conclusion Générale

L'objectif principal de ce projet de fin d'études était d'implémenter et d'évaluer un système intelligent de classification automatique des modulations (AMC) basé sur le Deep Learning pour répondre aux défis des environnements sans fil modernes. Les méthodes traditionnelles d'extraction manuelle montrant leurs limites face au bruit et à l'évanouissement, l'apprentissage profond s'est imposé comme une solution incontournable pour l'extraction automatique des signatures de signaux.

Au cours de ce travail, nous avons d'abord exposé les concepts fondamentaux des modulations numériques ainsi que les perturbations liées aux canaux de transmission. Nous avons ensuite posé le cadre théorique de l'intelligence artificielle et des réseaux de neurones, avant de mener une étude approfondie de l'état de l'art des architectures appliquées à l'AMC. Enfin, le dernier chapitre a permis de valider notre approche à travers une architecture hybride innovante combinant CNN, Transformer, BiGRU et mécanisme d'attention.

Les tests effectués sur la base de données RadioML 2016.10a confirment la pertinence de cette combinaison, affichant une exactitude globale de 67,80 %. Le modèle fait preuve d'une bonne robustesse pour l'identification de modulations complexes telles que les familles QAM et PAM. Toutefois, certaines difficultés persistent dans la discrimination de modulations présentant des caractéristiques similaires, notamment QPSK et 8PSK, particulièrement dans des conditions de faible rapport signal sur bruit.

Les résultats obtenus démontrent ainsi l'intérêt des architectures hybrides de Deep Learning pour améliorer les performances de la classification automatique des modulations et contribuer au développement de systèmes de communications sans fil plus intelligents et plus adaptatifs.

Perspectives et travaux futurs :

Dans le cadre des travaux futurs, plusieurs pistes d'amélioration peuvent être envisagées. Il serait intéressant d'évaluer le modèle sur des bases de données plus récentes et plus diversifiées, telles que RadioML 2018, afin de tester sa capacité de généralisation

Conclusion générale

L'intégration de mécanismes d'attention plus avancés ou de Transformers plus performants pourrait également permettre d'améliorer la reconnaissance des modulations fortement similaires. Par ailleurs, une optimisation de l'architecture pour réduire sa complexité computationnelle favoriserait son déploiement en temps réel sur des équipements embarqués ou des systèmes radio intelligents. Enfin, l'étude de techniques d'apprentissage auto-supervisé ou de transfert d'apprentissage constitue une perspective prometteuse pour améliorer les performances du modèle dans des environnements réels caractérisés par des conditions de transmission variables et des données limitées.

Références

Références

- [1] Savo G. Glisic, *Advanced Wireless Communications and Internet: Future Evolving*, John Wiley & Sons, 2011.
- [2] F. Zhang, C. Luo, J. Xu, Y. Luo, F. Zheng, «Deep learning based automatic modulation recognition: Models, datasets, and challenges,» *Digital Signal Processing*, 2022.
- [3] F. Rehman, Q. Abbas, M. Karam Shehzad, «DL-AMC: Deep Learning for Automatic Modulation Classification,» 2025.
- [4] TechLTE World, «Modulation in Communication,» 20 Décembre 2023. [En ligne]. Available: <https://techlteworld.com/modulation-in-communication/>.
- [5] Jerry D. Gibson, *Digital Communications: Introduction to Communication Systems*, Éditeur Springer International Publishing, 2023, 2023.
- [6] L. W. Couch, *Digital & Analog Communication Systems*, Upper Saddle River , New Jersey : Pearson, 2012.
- [7] X.-B. Bian et A. D. Bandrauk, «Probing nuclear motion by frequency modulation of molecular high-order harmonic generation,» *Physical Review Letters*, 7 Novembre 2014.
- [8] J. B. Anderson, T. Aulin et C.-E. Sundberg, *Digital Phase Modulation*, New York: Springer Science, 1986.
- [9] J. B. Anderson, T. Aulin et C.-E. Sundberg, *Digital Phase Modulation*, New York: Springer Science, 1986.
- [10] F. Yu et V. Koltun, «Multi-Scale Context Aggregation by Dilated Convolutions,» *arXiv*, 23 November 2015.
- [11] F. A. Rakhmadi, Mitrayana et H. Shalihah, «Design of green laser modulation system based on Arduino Uno microcontroller,» *Conference Proceedings* , 2018.
- [12] R. W. World, «ASK Modulation: Advantages and Disadvantages,» n.d. [En ligne]. Available: <https://www.rfwireless-world.com/terminology/ask-modulation-advantages-disadvantages>.

Références

- [13] S. Beyeler, S. Joly, M. Fries, F.-J. Obermair, F. Burn, R. Mehmood, G. Tabatabai et O. Raineteau, «Targeting the bHLH transcriptional networks by mutated E proteins in experimental glioma,» *Stem Cells*, p. 2583–2595, 15 September 2014.
- [14] R. W. World, «PSK Advantages and Disadvantages,» n.d. [En ligne]. Available: <https://www.rfwireless-world.com/terminology/psk-advantages-disadvantages.html>□
- [15] S. Sarkar, A. Sarrafinazhad et A. Ghoncheh, «Frequency Shift Keying (FSK): Overview,» 31 December 2024. [En ligne]. Available: <https://rahsoft.com/2024/12/31/frequency-shift-keying-fsk-overview/>.
- [16] A. Qaddouri, «Génération et démodulation BFSK,» April 2022. [En ligne]. Available: <https://fr.scribd.com/document/589907784/tp1?query=bfsk>□
- [17] W. D. Group, «MFSK Mode,» n.a. [En ligne]. Available: https://www.w1hkj.org/FldigiHelp/mfsk_page.html□
- [18] R. W. World, «Search results for FSK (Frequency Shift Keying),» n.d. [En ligne]. Available: <https://www.rfwireless-world.com/search?q=fsk>.
- [19] D. International, « What is QAM (Quadrature Amplitude Modulation)?,» 24 March 2025. [En ligne]. Available: <https://info.support.huawei.com/info-finder/encyclopedia/en/QAM.html>□
- [20] R. W. World, «Search results for QAM advantages and disadvantages,» n.d. [En ligne]. Available: <https://www.rfwireless-world.com/search?q=qam+advantage+and+disadvantages>□
- [21] I.-L. A. I.-L. K. Base), «AWGN (Additive White Gaussian Noise) – Definition and Use,» 2024. [En ligne]. Available: <https://ib-lenhardt.com/kb/glossary/awgn>□
- [22] Zaretti, «Atténuation, Distorsion, Bruit et Collisions,» n.d. [En ligne]. Available: <https://cours.zaretti.be/courses/networking/bases-des-reseaux/lessons/attenuation-distorsion-bruit-collision/>□
- [23] GeeksforGeeks, «Fading in Wireless Communication,» 23 July 2025. [En ligne]. Available: <https://www.geeksforgeeks.org/computer-networks/fading-in-wireless-communication/>.
- [24] Wikipédia, « Distorsion (électronique),» n.d. [En ligne]. Available: [https://fr.wikipedia.org/wiki/Distorsion_\(%C3%A9lectronique\)](https://fr.wikipedia.org/wiki/Distorsion_(%C3%A9lectronique))□
- [25] L. Shenzhen Sinoseen Technology Co., «What is Signal-to-Noise Ratio (SNR)? How does it affect embedded vision?,» 13 August 2024. [En ligne]. Available:

Références

- <https://www.sinoseen.com/what-is-signal-to-noise-ratiohow-does-it-effect-embedded-vision>.
- [26] L. GIGALIGHT Technology Co., «What is Bit Error Rate (BER)?», 2021. [En ligne]. Available: <https://www.gigalight.com/resources/what-is-bit-error-rate-ber.html>.
- [27] « Efficacité spectrale », n.d. [En ligne]. Available: <https://dSPACE.univ-guelma.dz/jspui/bitstream/123456789/4024/1/mmoir%20fin%20d%27etude.pdf>.
- [28] B. Copeland, «Artificial intelligence (AI) | Definition, Examples, Types, Applications, Companies, & Facts», Encyclopaedia Britannica, 11 Mars 2026. [En ligne]. Available: <https://www.britannica.com/technology/artificial-intelligence> □
- [29] C. Stryker et E. Kavlakoglu, « What Is Artificial Intelligence (AI)? », 2024. [En ligne]. Available: <https://www.ibm.com/think/topics/artificial-intelligence> □
- [30] S. C. Institute, «History of Artificial Intelligence: A Timeline from Turing to Today», 21 Janvier 2026. [En ligne]. Available: <https://swisscyberinstitute.com/blog/history-artificial-intelligence/> □
- [31] Dr. Hachemi. B, cours «CH1-IA», université Belhadj Bouchaib Ain témouchent, Ain témouchent, s.d.
- [32] Iberdrola, « History of artificial intelligence », 2026. [En ligne]. Available: <https://www.iberdrola.com/about-us/our-innovation-model/history-artificial-intelligence>.
- [33] T. Software, «What is the history of artificial intelligence (AI)?», s.d. [En ligne]. Available: <https://www.tableau.com/data-insights/ai/history> □
- [34] E. Eygay, « Intelligence Artificielle : Définition, Types et Applications », 03 octobre 2025. [En ligne]. Available: <https://www.data-bird.co/blog/intelligence-artificielle-definition> □
- [35] N. indiqué, « Dans qu'elles domaines en peut l'utiliser IA », s.d. [En ligne]. Available: <https://share.google/zWMU2K6sBG08CJlie>.
- [36] Elastic, « Qu'est-ce que le Machine Learning ? », s.d. [En ligne]. Available: <https://www.elastic.co/fr/what-is/machine-learning>.
- [37] KAIZEN, « Types et applications de l'apprentissage automatique », s.d. [En ligne]. Available: <https://kaizen.com/fr/publications/types-applications-apprentissage-automatique/> □

Références

- [38] L. Robots, « L'apprentissage par renforcement : clé de l'IA autonome, » 28 avril 2025. [En ligne]. Available: <https://www.learningrobots.ai/blog/lapprentissage-par-renforcement-cle-de-lia-autonome>
- [39] Oracle, «Ce qu'il faut savoir sur l'apprentissage semi-supervisé, » 2024. [En ligne]. Available: <https://www.oracle.com/fr/artificial-intelligence/machine-learning/semi-supervised-learning/>.
- [40] GeeksforGeeks, «Artificial Neural Network vs Biological Neural Network, » s.d. [En ligne]. Available: <https://www.geeksforgeeks.org/difference-between-ann-and-bnn/> □
- [41] R. Sheldon, « Qu'est-ce qu'un perceptron ?, » 19 janvier 2026. [En ligne]. Available: <https://www.lemagit.fr/definition/Perceptron>.
- [42] K. A. Rashedi, M. T. Ismail, S. A. Wadi, A. Serroukh, T. S. Alshammari et J. J. Jaber, « Multi-Layer Perceptron-Based Classification with Application to Outlier Detection in Saudi Arabia Stock Returns, » *Journal of Risk and Financial Management (JRFM)*, 2024.
- [43] N. Lunge, « A Deep Architecture: Multi-Layer Perceptron, » 24 mars 2024. [En ligne]. Available: <https://medium.com/@nlunge786/a-deep-architecture-multi-layer-perceptron-164bc5ff3842>.
- [44] G. L. E. Team, « What is Activation Function?, » 21 février 2025. [En ligne]. Available: <https://www.mygreatlearning.com/blog/activation-functions/> □
- [45] G. L. E. Team, «What are Activation Functions in Neural Networks?, » 21 février 2025. [En ligne]. Available: <https://www.mygreatlearning.com/blog/activation-functions/> □
- [46] S. Team, «Activation functions in neural networks, » 12 October 2023. [En ligne]. Available: <https://www.superannotate.com/blog/activation-functions-in-neural-networks> □
- [47] GeeksforGeeks, «Activation Functions in Neural Networks, » 2026. [En ligne]. Available: <https://www.geeksforgeeks.org/activation-functions-neural-networks/> □
- [48] S. Tahsildar, « Introduction To Deep Learning And Neural Network, » 31 décembre 2018. [En ligne]. Available: <https://datamites.com/blog/introduction-to-deep-learning-and-neural-networks/> □
- [49] G. L. E. Team, «Machine Learnig vs Deep Learning : Differences,Exampels,and Applications, » 2025. [En ligne]. Available: <https://share.google/nJgrAgqxxNctm9gMg> □

Références

- [50] «Advantages and Disadvantages of Deep Learning,» 31 July 2025. [En ligne]. Available: <https://www.geeksforgeeks.org/deep-learning/advantages-and-disadvantages-of-deep-learning/>
- [51] X. Z. Q. L. B. Q. D. Y. K. W. Z. L. B. J. H. L. X. L. Y. & G. G. Tian, «A survey on deep learning enabled automatic modulation classification methods: Data,» *Signal Processing*, vol. 242, p. 110444, 2026.
- [52] M. A. E.-S. W. A.-S. N. e. a. Abdel-Moneim, «A survey of traditional and advanced automatic modulation classification techniques,» *International Journal of Communication Systems*, vol. 34, 2021.
- [53] S. S. R. Z. L. Chang, «Multitask-Learning-Based Deep Neural Network for Automatic Modulation Classification,» *IEEE Internet of Things Journal*, vol. 9, n° %13, p. 2192–2206, 2022.
- [54] K. A. R. A. a. Tekbiyik, «Robust and fast automatic modulation classification with CNN under multipath fading,» *IEEE 91st Vehicular Technology Conference (VTC)*, p. 1–6, 2020.
- [55] C. S. Z. Xiao, «Complex-valued depth-wise separable convolutional neural network for automatic,» *IEEE Transactions on Instrumentation and Measurement*, vol. 72, p. 2522310, 2023.
- [56] S. Y. P. P. a. Hu, «Robust modulation classification under uncertain noise condition using recurrent,» *IEEE Global Communications Conference (GLOBECOM)*, p. 1–7, 2018.
- [57] Y. W. J. a. Chen, «Automatic modulation classification scheme based on LSTM with random erasing and,» *IEEE Access*, vol. 8, p. 154290–154300, 2020.
- [58] Q. J. K. a. Xuan, «AvgNet: adaptive visibility graph neural network and its application in modulation,» *IEEE Transactions on Network Science and Engineering*, vol. 9, n° %13, p. 1516–1526, 2022.
- [59] N. N. I. a. K. Tonchev, «Automatic Modulation Classification using Graph Convolutional Neural Networks for Time-,» *25th IEEE International Symposium on Wireless Personal Multimedia Communications*, p. 75–79, 2022.
- [60] F. G. C. a. J. Cai, «Signal Modulation Classification Based on the Transformer Network,» *IEEE Transactions on Cognitive Communications and Networking*, vol. 8, n° %13, p. 1348–1357, 2022.

Références

- [61] H. L. W. a. Z. Zhang, «Automatic Modulation Classification Using CNN-LSTM Based Dual-Stream Structure,» *IEEE Transactions on Vehicular Technology*, vol. 69, n° %111, p. 13521–13531, 2020.
- [62] H. L. W. a. Z. Zhang, «Automatic Modulation Classification Using CNN-LSTM Based Dual-Stream Structure,» *IEEE Internet of Things Journal*, vol. 7, n° %18, p. IEEE Internet of Things Journal, 2020.
- [63] R. D. H. a. S. Huang, «Automatic Modulation Classification Using Gated Recurrent Residual Network,» *IEEE Internet of Things Journal*, vol. 7, n° %18, p. 7795–7807, 2020.
- [64] X. T. Y. a. Q. Zheng, «DL-PR: Generalized Automatic Modulation Classification Method Based on Deep Learning,» *Engineering Applications of Artificial Intelligence*, vol. 122, n° %1106082, 2023.
- [65] P. Z. L. a. Q. Zheng, «Spectrum Interference-Based Two-Level Data Augmentation Method in Deep Learning for,» *Neural Computing and Applications*, vol. 33, n° %113, p. 7723–7745, 2021.
- [66] T. J. O'Shea et N. West, «RadioML 2016.10A Dataset,» 2016. [En ligne]. Available: <https://www.deepsig.ai/datasets/>.
- [67] nolasthitnotomorrow, « RadioML2016 DeepSig Dataset,» n.d. [En ligne]. Available: <https://www.kaggle.com/datasets/nolasthitnotomorrow/radioml2016-deepsigcom>.
- [68] N. identifié, n.d. [En ligne].
- [69] G. E. Team, « Architecture and Working of Transformers in Deep Learning,» 18 octobre 2025. [En ligne]. Available: <https://www.geeksforgeeks.org/deep-learning/architecture-and-working-of-transformers-in-deep-learning/> □
- [70] G. E. Team, «Gated Recurrent Unit Networks,» 23 décembre 2025. [En ligne]. Available: <https://www.geeksforgeeks.org/machine-learning/gated-recurrent-unit-networks/> □
- [71] R. Tavernier, «Attention (Deep Learning Book – chapitre),» n.d. [En ligne]. Available: https://rtavenar.github.io/deep_book/fr/content/fr/attention.html.
- [72] Y. W. C. L. J. Y. K. a. L. W. HUO, «SMTrans: An efficient automatic modulation recognition network based on the scale-aware modulation transformer,» vol. 74, p. 102966, 2025.

Références

- [73] J. L. C. P. G. a. L. Y. XU, «A spatiotemporal multi-channel learning framework for automatic modulation recognition,» *IEEE Wireless Communications Letters*, vol. 9, n° %110, p. 1629–1632, 2020.
- [74] F. L. C. X. J. a. L. Y. ZHANG, «An efficient deep learning model for automatic modulation recognition based on parameter estimation and transformation,» *IEEE Communications Letters*, vol. 25, n° %110, p. 3287–3290, 2021.
- [75] J. W. T. F. Z. a. Y. S. ZHANG, «AMC-Net: An effective network for automatic modulation classification,» p. 1–5, 2023.
- [76] Y. D. B. L. C. X. W. a. L. S. CHEN, « Abandon locality: Frame-wise embedding aided transformer for automatic modulation recognition,» *IEEE Communications Letters*, vol. 27, n° %11, p. 327–331, 2023.
- [77] S. W. C. L. J. W. M. Y. K. a. L. W. WU, «A transformer-based framework with complex-valued convolution and enhanced Bi-LSTM for automatic modulation recognition,» *Physical Communication*, vol. 76, p. 102824, 2025.