

République Algérienne Démocratique et Populaire
Ministère de l'enseignement supérieur et de la recherche scientifique
Université 8 Mai 1945 – Guelma
Faculté des sciences et de la Technologie
Département de Electronique et Télécommunication



Mémoire de fin d'étude
Pour l'obtention du diplôme de MASTER

Domaine : **Sciences et Technologie**
Filière : **Électronique**
Spécialité : **Instrumentation**

La reconnaissance palmaire par Réseaux neurones

Présenté par :

Larafa TOUFIK
Gassi ABDELAZIZ

Sous la direction de :

Dr. Boualleg Abdelhalim

Année universitaire 2024-2025

Dédicace

Je tiens, de prime abord, à me prosterner en remerciant Allah le ToutPuissant de m'avoir donné le courage et la patience pour mener à bien ce travail.

Je tiens à exprimer ma gratitude à ceux qui m'ont donné la chance de reprendre mes études à l'université après vingt ans d'absence. Leur confiance en moi a été une source d'inspiration et de motivation inestimable.

*Je remercie mon encadreur, Monsieur **Dr. Boualleg Abdelhalim**, de l'Université 8 Mai 1945 de Guelma, pour m'avoir fait l'honneur de diriger ce travail. Je lui suis reconnaissante pour ses précieux conseils, sa disponibilité, sa confiance scientifique en moi et ses discussions enrichissantes.*

À ceux qui m'ont tant donné sans rien demander, qui m'ont toujours offert leur soutien, n'ont épargné aucun effort pour m'aider, et m'ont inculqué mes principes. Aucun mot ne serait suffisant pour les remercier :

« Mes très chers parents ».

Je tiens également à remercier les membres du jury d'avoir accepté d'y participer et pour l'honneur qu'ils me font en expertisant mon travail.

Résumé

Dans cette étude, nous explorons une approche de reconnaissance palmaire avec contact basée sur l'apprentissage profond et le transfert de connaissances à partir du modèle GoogLeNet. Notre pipeline de traitement est structuré en trois étapes principales : extraction de caractéristiques par transfert d'apprentissage, réduction de dimensionnalité, et classification. Tout d'abord, nous exploitons la capacité de généralisation du réseau GoogLeNet préentraîné sur ImageNet pour extraire des descripteurs profonds à partir des images palmaires des bases MSCASIA et MSPOLY. Ces deux bases comportent des images multispectrales de paumes capturées dans des conditions contrôlées, ce qui les rend adaptées à l'évaluation de méthodes robustes d'identification biométrique. Ensuite, pour maximiser la séparabilité inter-classes tout en réduisant la dimensionnalité des vecteurs de caractéristiques, nous appliquons l'analyse discriminante linéaire (LDA). Cette étape permet de conserver les dimensions les plus discriminantes, en tenant compte des classes cibles, tout en atténuant le bruit et la redondance des données. Enfin, la classification est effectuée à l'aide d'un classifieur k -plus proches voisins (k -NN) utilisant une métrique de similarité cosinus, bien adaptée aux représentations vectorielles normalisées.

Mots clés : *Biométrie, Empreintes palmaires, Multispectral, Identification, Apprentissage profond, CNN, KNN, LDA.*

Abstract

In this study, we explore an approach to contact-based palm recognition based on deep learning and knowledge transfer from the GoogLeNet model. Our processing pipeline is structured in three main steps: feature extraction by learning transfer, dimensionality reduction, and classification. Firstly, we exploit the generalization capability of the GoogLeNet network pre-trained on ImageNet to extract deep descriptors from palm images in the MSCASIA and MSPOLY databases. Both databases feature multispectral images of palms captured under controlled conditions, making them suitable for the evaluation of robust biometric identification methods. Next, to maximize inter-class separability while reducing the dimensionality of feature vectors, we apply linear discriminant analysis (LDA). This step allows us to retain the most discriminating dimensions, taking into account the target classes, while mitigating data noise and redundancy. Finally, classification is performed with a k -nearest neighbor (k -NN) classifier using a cosine similarity metric, well suited to normalized vector representations.

Keywords: *Biometrics, Palm prints, Multispectral, Identification, deep learning, CNN, LDA, KNN.*

ملخص

في هذه الدراسة، نستكشف نهجاً للتعرف على راحة اليد المتلامسة استناداً إلى التعلم العميق ونقل المعرفة من نموذج GoogLeNet. يتم تنظيم خط أنابيب المعالجة الخاص بنا في ثلاث خطوات رئيسية: استخراج السمات عن طريق نقل التعلم، وتقليل الأبعاد، والتصنيف. أولاً، نحن نستغل قدرة التعميم لشبكة GoogLeNet المدربة مسبقاً على ImageNet لاستخراج الوصفات العميقة من صور النخيل في قاعدتي بيانات MSCASIA و MSPOLY. وتحتوي قاعدتا البيانات هاتان على صور متعددة الأطياف لأشجار النخيل التي تم التقاطها في ظروف خاضعة للرقابة، مما يجعلها مناسبة لتقييم طرق التعرف البيومترية القوية. بعد ذلك، ولتعزيز قابلية الفصل بين الفئات مع تقليل أبعاد متجهات السمات إلى أقصى حد، نطبق التحليل التمييزي الخطي (LDA). هذه الخطوة تجعل من الممكن الاحتفاظ بالأبعاد الأكثر تمييزاً، مع مراعاة الفئات المستهدفة، مع تقليل ضوضاء البيانات والتكرار. أخيراً، يتم إجراء التصنيف باستخدام مصنف الجار الأقرب (k -NN) باستخدام مقياس تشابه جيب التمام، وهو مناسب تماماً لتمثيلات المتجهات الطبيعية.

TABLE DES MATIÈRES

I. CHAPITRE I : La biométrie	16
I.1 Introduction	17
I.2 Définitions et propriétés	17
I.3 Architecture d'un système biométrique	18
I.3.1 Principaux Modules	18
I.3.2 Fonctionnement des systèmes biométriques	19
I.4 Modalités biométriques émergentes	20
I.4.1 Modalités morphologiques (physiologiques)	21
I.4.2 Modalités comportementales	24
I.4.3 Modalités biologiques	26
I.5 Représentation comparative entre quelques Techniques Biométriques	28
I.6 Performance des systèmes biométriques	28
I.7 Mesure de la performance d'un système biométrique	30
I.7.1 Taux de faux rejets (False Reject Rate ou FRR)	30
I.7.2 Taux de fausses acceptations (False Accept Rate ou FAR)	30
I.7.3 Taux d'égale erreur (Equal Error Rate ou EER)	31
I.8 Champ d'application des systèmes biométriques	32
I.8.1 Contrôle d'accès physiques aux locaux	33
I.8.2 Contrôle d'accès logiques aux systèmes d'informations	33
I.8.3 Applications légales (juridique)	33
I.9 Les avantages et les inconvénients de la biométrie	34
I.10 Conclusion	34
II. CHAPITRE II Système de reconnaissance des empreintes palmaires	35
II.1 Introduction	36
II.2 Pourquoi la reconnaissance de l'empreinte palmaire ?	36
II.3 Définition de l'empreinte palmaire	37
II.3.1 Avantages de l'empreinte palmaire	37
II.3.2 Caractéristique des empreintes palmaires	37
II.4 Types d'empreintes palmaires	39
II.4.1 Empreinte palmaire latente	40
II.4.2 Empreinte palmaire brevetée	40
II.4.3 Empreinte palmaire en plastique	40
II.5 Les étapes de la reconnaissance de l'empreinte palmaire	40

II.6 Des méthodes de prétraitement	41
II.6.1 Rehaussement des niveaux de gris	41
II.6.2 Égalisation d'histogramme	41
II.6.3 Débruitage par Filtre gaussien	42
II.7 Méthodes d'Extraction des Caractéristiques	43
II.7.1 Extraction basé sur la texture	43
II.7.2 Extraction basé sur le filtrage	45
II.8 Classification	46
II.8.1 Machine à vecteurs de support (SVM)	46
II.8.2 Les k plus proches voisin	46
II.8.3 Les Réseaux de Neurones	46
II.9 Domaines d'application	47
II.10 Conclusion	48
III. CHAPITRE III : Apprentissage En Profondeur (Deep Learning)	49
III.1 Introduction	50
III.2 Qu'est-ce que l'apprentissage en profondeur (Deep Learning) ?	50
III.3 Intelligence artificielle	51
III.4 Machine Learning	52
III.5 L'apprentissage profond	52
III.6 Différent type d'apprentissage	53
III.6.1 L'apprentissage supervisé	54
III.6.2 Apprentissage semi-supervisé	55
III.6.3 Apprentissage non supervisé	56
III.6.4 L'apprentissage par renforcement	56
III.7 Pourquoi d'apprentissage en profondeur ?	57
III.8 L'apprentissage profond et d'apprentissage automatique	58
III.9 Rétropropagation du gradient	59
III.10 Les réseaux de neurones	59
III.10.1 Réseaux de neurones biologiques.....	59
III.10.2 Réseau neuronal artificiel	60
III.11 La fonction d'activation	61
III.12 Les différents types de réseaux de neurones	62
III.12.1 Réseaux de neurones récurrents RNN	62
III.12.2 Réseaux de neurones profonds DNN	62
III.12.3 Le LSTM (LongShort-TermMemory)	63

III.12.4	Réseaux Neurones Convolutifs CNN	64
III.13	La construction des réseaux de neurones	64
III.13.1	Collecte et préparation des données	64
III.13.2	Choix de l'architecture	64
III.13.3	Définition des couches et des paramètres	65
III.13.4	Définition de la fonction de coût et de l'optimiseur	65
III.13.5	Entraînement du réseau	65
III.13.6	Validation et ajustement des hyperparamètres	65
III.13.7	Évaluation des performances	66
III.14	Algorithmes l'apprentissage profond	66
III.14.1	Régression linéaire	66
III.14.2	Régression logistique	66
III.15	Principes fondamentaux de l'apprentissage profond	66
III.16	Exemples de domaines d'application	66
III.17	Quelques architectures connus de réseaux de neurones	67
III.17.1	AlexNet (2012)	67
III.17.2	VGG (2014)	68
III.17.3	ResNet (2015)	68
III.18	Conclusion	69
IV.	CHAPITRE IV : Méthodologie, Expérimentations et résultats	70
IV.1	Introduction	71
IV.2	Présentation des outils de travail (Logiciel) Anaconda {Spyder +Python}	71
IV.2.1	Anaconda	72
IV.2.2	Spyder	73
IV.2.3	Python	74
IV.3	Les bases de données	74
IV.3.1	Base de données de l'empreinte palmaire MS-CASIA	74
IV.3.2	Base de données multi-spectrale (MS-PolyU)	74
IV.4	Méthode proposée	76
IV.5	Vue générale du processus et méthode de reconnaissance palmaire	77
IV.5.1	Acquisition de l'image	77
IV.5.2	Prétraitement	77
IV.5.3	Extraction des caractéristiques	78
IV.5.3.1	GoogLeNet (Inception V1)	78
IV.5.3.2	LDA – Linear Discriminant Analysis.....	86

IV.5.4	Classification	87
IV.5.5	Décision	87
IV.6	Expérimentations et Résultats	88
IV.6.1	Expérimentation avec la base de données MS-PolyU	88
IV.6.2	Discutions des résultats	89
IV.6.3	Expérimentation avec la base de données MS-CASIA (Main Droite)	89
IV.6.4	Expérimentation avec la base de données MSCASIA (Main Gauche)	90
IV.7	Conclusion.....	96
	Conclusion Générale	97
	Bibliographie	98

LISTE DES FIGURES

FIGURE I.1:	STRUCTURE DES SYSTEMES BIOMETRIQUE.....	20
FIGURE I.2:	CATEGORIES DES MODALITES BIOMETRIQUES.....	20
FIGURE I.3:	DISPOSITIF DE RECONNAISSANCE DE LA GEOMETRIE DE MAIN (A) ET GEOMETRIE DE LA MAIN (B) ET(C)	21
FIGURE I.4:	EMPREINTE DIGITALE	22
FIGURE I.5:	L'EMPREINTE PALMAIRE	22
FIGURE I.6:	LE VISAGE	23
FIGURE I.7:	L'IRIS	23
FIGURE I.8:	LA RETINE	24
FIGURE I.9:	L'EMPREINTE DE L'OREILLE	24
FIGURE I.10:	L'ECRITURE (LA SIGNATURE)	25
FIGURE I.11:	LA DYNAMIQUE DE FRAPPE AU CLAVIER	25
FIGURE I.12:	LA VOIX (RECONNAISSANCE VOCALE)	26
FIGURE I.13:	LA DEMARCHE	26
FIGURE I.14:	IMAGE D'ADN	27
FIGURE I.15:	RECONNAISSANCE DES VEINES	28
FIGURE I.16:	ILLUSTRATION DU FFR ET DU FAR	31
FIGURE I.17:	COURBE ROC	32
FIGURE I.18:	COURBE CMC	32
FIGURE I.19:	LES APPLICATIONS DES SYSTEMES BIOMETRIQUES	33
FIGURE II.1:	PAUME DE LA MAIN	38
FIGURE II.2:	LES PLIS DE FLEXIONS DE LA PAUME DE LA MAIN	38
FIGURE II.3:	LES POINTS DE REFERENCE DE L'EMPREINTE PALMAIRE	39
FIGURE II.4:	TROIS CATEGORIES DE TYPE D'EMPREINTES PALMAIRES	39
FIGURE II.5:	LE PRINCIPE DES METHODES D'IDENTIFICATION BIOMETRIQUE	40
FIGURE II.6:	RECHAUSSEMENT DES NIVEAUX DE GRIS	41
FIGURE II.7:	(A) IMAGE DE DEPART, (B) HISTOGRAMME DE L'IMAGE (A)	42
FIGURE II.8:	(A) IMAGE EGALISATION DE L'HISTOGRAMME, (B) HISTOGRAMME APRES EGALISATION	42
FIGURE II.9:	(A) FLOU GAUSSIEN DE RAYON 2, (B) FLOU GAUSSIEN DE RAYON 3	43
FIGURE II.10:	CONSTRUCTION D'UN MOTIF BINAIRE ET CALCUL DE CODE LBP.....	44
FIGURE II.11:	ORGANIGRAMME DE L'ENSEMBLE DES ETAPES NECESSAIRE A LA CONSTRUCTION DU DESCRIPTEUR LPQ	45

FIGURE III.1: LES RELATION ENTRE L'INTELLIGENCE ARTIFICIELLE, L'APPRENTISSAGE AUTOMATIQUE ET L'APPRENTISSAGE PROFOND	50
FIGURE III.2: LES TYPES DE L'INTELLIGENCE ARTIFICIELLE	51
FIGURE III.3: LES TYPE DE L'APPRENTISSAGE AUTOMATIQUE	52
FIGURE III.4: EXEMPLE D'APPRENTISSAGE PROFOND	53
FIGURE III.5: EXEMPLE DE L'APPRENTISSAGE SUPERVISE	54
FIGURE III.6: EXEMPLE D'APPRENTISSAGE SEMI-SUPERVISE	56
FIGURE III.7: EXEMPLE D'APPRENTISSAGE NON SUPERVISE	56
FIGURE III.8: SCHEMA D'APPRENTISSAGE PAR RENFORCEMENT	57
FIGURE III.9: QUALITE DES DONNEES	57
FIGURE III.10: UNE COMPARAISON DES ETAPES DU FONCTIONNEMENT DES ALGORITHMES DL ET ML	58
FIGURE III.11: REPRESENTATION SCHEMATIQUE D'UN NEURONE BIOLOGIQUE.....	60
FIGURE III.12: ARCHITECTURE D'UN NEURONE FORMEL	60
FIGURE III.13: ARCHITECTURE D'UN RESEAU DE NEURONE	61
FIGURE III.14: ARCHITECTURE DE RNN	62
FIGURE III.15: REPRESENTATION SYNOPTIQUE D'UN CNN	64
FIGURE III.16 : ARCHITECTURE ALEXNET	68
FIGURE III.17 : ARCHITECTURE VGG	68
FIGURE III.18: SIX SPECTRES D'UNE SEULE PERSONNE D'EMPREINTES PALMAIRES ROI DANS LA	69
FIGURE IV.1: ENVIRONNEMENT ANACONDA.....	71
FIGURE IV.2: ENVIRONNEMENT SPYDER	72
FIGURE IV.3 : LOGICIEL PYTHON.....	73
FIGURE IV.4: DISPOSITIF D'IMAGERIE MULTI SPECTRALE.....	74
FIGURE IV.5 : SIX IMAGES D'UNE SEULE PERSONNE D'EMPREINTES PALMAIRES DE LA BASE MS-CASIA.....	75
FIGURE IV.6: SIX SPECTRES D'UNE SEULE PERSONNE D'EMPREINTES PALMAIRES ROI DANS LA BASE MS-CASIA	75
FIGURE IV.7 : ÉCHANTILLONS DE ROI D'EMPREINTES PALMAIRES DE LA BASE MS-POLYU (A) : BLEU, (B) : VERT, (C) : NIR, (D) : ROUGE.....	76
FIGURE IV.8: ARCHITECTURE PROPOSEE DU NOTRE SYSTEME DE RECONNAISSANCE D'EMPREINTES PALMAIRES.....	77
FIGURE IV.9 : LES PRINCIPALES TACHES DE PRETRAITEMENT. A. IMAGE D'ENTREE, B. EXTRACTION DE (ROI), C. REDIMENSIONNEMENT DE L'IMAGE.....	78
FIGURE IV.10 : ARCHITECTURE GLOBALE DE GOOGLNET (INCEPTION V1)	79
FIGURE IV.11 : BLOC D'ENTREE.....	79
FIGURE IV.12 : BLOC C1 – PREMIER E COUCHE CONVOLUTIONNELLE.....	80
FIGURE IV.13 : BLOCS C2, C3, C4 – CONVOLUTIONS SUCCESSIVES.....	82
FIGURE IV.14 : BLOC P2 – POOLING 2.....	82
FIGURE IV.15: INCEPTION MODULE.....	83
FIGURE IV.16: BLOC P3 – GLOBAL AVERAGE POOLING.....	83
FIGURE IV.17: FONCTION D'ACTIVATION SOFTMAX.....	84
FIGURE IV.18: TAUX DE RECONNAISSANCE SUR LA BASE MS-POLYU.....	89
FIGURE IV.19: TAUX DE RECONNAISSANCE SUR LA BASE MS-CASIA (MAIN GAUCHE)	91
FIGURE IV.20: TAUX DE RECONNAISSANCE SUR LA BASE MS-CASIA (MAIN DROITE)	92

LISTE DES TABLEAUX

TABLEAU I.1: COMPARAISON ENTRE QUELQUES TECHNIQUES BIOMETRIQUES	28
TABLEAU I.2: LES AVANTAGE ET LES INCONVENIENTS DE LA BIOMETRIE	34
TABLEAU III.1: LES TYPES D'APPRENTISSAGE AUTOMATIQUE	53
TABLEAU III.2: COMPARAISON ENTRE L'APPRENTISSAGE PROFOND ET L'APPRENTISSAGE AUTOMATIQUE	58
TABLEAU III.3: LES FONCTIONS D ACTIVATIONS	61
TABLEAU IV.1 : DIFFERENCE ENTRE LES BASES DE DONNEES.	76
TABLEAU IV.2: DIMENSIONS DANS NOTRE MODELE UTILISANT GOOGLNET	86
TABLEAU IV.3 : RESUME DES RESULTATS PAR SPECTRE POUR MSPOLY.....	89
TABLEAU IV.4 : DES RESULTATS PAR SPECTRE POUR MSCASIA(MAIN GAUCHE)	90
TABLEAU IV.5 : RESULTATS PAR SPECTRE POUR MSCASIA (MAIN DROITE)	92
TABLEAU IV.6: COMPARAISON DE NOTRE METHODE AVEC D'AUTRES APPROCHES SPECIFIANT LE SPECTRE DANS LA BASE DE POLYU.....	94
TABLEAU IV.7: COMPARAISON DE NOTRE METHODE AVEC D'AUTRES APPROCHES SPECIFIANT LE SPECTRE DANS LA BASE DE MS-CASIA.....	95

LISTE DES ABBREVIATIONS

FBI	: Fédéral Bureau of Investigation : Bureau fédéral d'investigation
IAFIS	: Integrated Automated Fingerprint Identification System : Système automatisé intégré d'identification par empreintes digitales
ANSI	: American National Standards Institute : Institut national américain de normalisation
NIST	: National Institute of Standards and Technologie : Institut national des normes et de la technologie
PIN	: Personal Identification Number : Numéro d'identification personnel
ADN	: Acide Désoxyribose Nucléique
FRR	: False Rejection Rate : Taux de faux rejets
FAR	: False Acceptance Rate : Taux de fausse acceptation
ROC	: Receiver Operating Characteristic : Caractéristiques de Fonctionnement du Récepteur
EER	: Equal Error Rate : Taux d'erreur égal
NIR	: Near InfraRed : Proche infrarouge
CCD	: Charge Coupled Device : Dispositif à couplage de charge
CAN	: Convertisseur Analogique-Numérique
ROI	: Region Of Interest : Région d'intérêt
LBP	: Local Binary Pattern : Modèle binaire local
LPQ	: Local phase quantization : Quantification de phase locale.
DFT	: Discrete Fourier Transform : Transformée de fourier discrète
ACP	: Analyse en Composantes Principales
LDA	: Analyse Discriminante Linéaire
CHVD	: Competitive Hand Valley Detection : Détection compétitive de la vallée de la main
PPI	: Pixels Per Inch : Pixels par pouce
NND	: Nearest Neighbour Distance : Distance du voisin le plus proche
BPNN	: BackPropagation Neural Network : Réseau neuronal à rétropropagation
KNN	: K-Nearest Neighbour : Algorithme K-plus proche voisin
DWT	: Discrete Wavelet Transform : Transformée en ondelettes discrète

2DWFT	: Two-dimensional Windowed Discrete Fourier Transform : transformée de Fourier Discrete fenêtrée bidimensionnelle
SVM	: Support Vector Machine : Machine à vecteurs de support
RNA	: Réseaux neuronaux artificiels
MS- CASIA	: MultiSpectral Chinese Academy of Sciences Institute of Automation
MSPOLYU	: MultiSpectral POLYtechnic University: The Hong Kong Polytechnic University
WHT	: WHITE : Lumière Blanche
NIR	: Near InfraRed: Proche infrarouge
KFA	: Kernel Fisher Analysis : Analyse de Fisher du noyau
CNN	: Convolutional Neural Network

INTRODUCTION GENERALE

La nécessité d'un accès sécurisé, automatisé et personnalisé à des environnements physiques ou virtuels est en constante augmentation. Pour répondre à cette exigence, il est crucial de disposer de moyens fiables permettant de vérifier l'identité des individus accédant à ces systèmes. Les méthodes traditionnelles, comme les mots de passe ou les cartes magnétiques accompagnées d'un code personnel, présentent de nombreuses failles : les mots de passe peuvent être oubliés, volés ou partagés, tandis que les cartes d'accès peuvent être perdues ou dérobées.

C'est dans ce contexte que l'utilisation des caractéristiques biométriques, c'est-à-dire des traits physiologiques propres à chaque individu (voix, visage, signature, empreintes digitales, forme de la main, etc.), s'est imposée comme une alternative plus sécurisée. Chaque type de caractéristique est désigné sous le terme de **modalité biométrique**. Les systèmes biométriques jouent ainsi un rôle clé dans la lutte contre la fraude, la sécurisation des transactions financières et commerciales, l'accès aux services publics, ainsi que la prévention du vol d'identité.

Parmi les modalités biométriques émergentes, l'empreinte palmaire (ou *palmprint*) [1] suscite un intérêt croissant. Les recherches menées jusqu'à présent dans ce domaine ont principalement porté sur le prétraitement des images et l'extraction de leurs caractéristiques distinctives afin d'optimiser la phase de classification. Dans notre travail, nous nous concentrons spécifiquement sur cette phase de classification en adoptant une approche basée sur l'apprentissage automatique, en utilisant notamment des méthodes robustes.

Historiquement, l'idée d'utiliser la paume de la main pour identifier les individus remonte à 1858, lorsque William Herschel, en poste en Inde, fit apposer des empreintes de paumes en guise de signature sur des contrats, notamment pour les personnes ne sachant pas écrire.

Plus récemment, en 1994, le premier système automatisé d'identification par empreintes digitales (AFIS) a commencé à intégrer les empreintes palmaires. Entre 2002 et 2004, le FBI a perfectionné cette technologie avec le développement de services nationaux dédiés aux empreintes palmaires, renforçant ainsi les capacités des forces de l'ordre à identifier les criminels avec plus de rapidité et de précision grâce au système automatisé intégré d'identification des empreintes (IAFIS).

Aujourd'hui, l'Australie détient la plus grande base de données d'empreintes palmaires au monde. Son système national automatisé d'identification par empreintes digitales regroupe 4,8 millions d'enregistrements et respecte les normes ANSI/NIST pour l'échange de données, facilitant ainsi la coopération avec des organismes internationaux comme Interpol ou le FBI [2].

Le système de reconnaissance par empreinte palmaire est une technologie avancée, applicable dans de nombreux secteurs et institutions pour l'identification des individus. Ce travail s'appuie sur les recherches antérieures pour explorer les méthodes et techniques de reconnaissance basées sur l'intelligence artificielle. Le processus de reconnaissance se décompose généralement en quatre étapes principales :

1. **Le prétraitement des données,**

2. **L'extraction des caractéristiques,**

3. **Et enfin, la phase de reconnaissance,** qui consiste à faire correspondre les données d'entrée avec celles déjà enregistrées dans le système. [3]

Dans ce mémoire, nous proposons une approche de reconnaissance spectrale des empreintes palmaires, reposant sur l'utilisation de descripteurs locaux, dans le but de concevoir un modèle d'identification hautement sécurisé pour l'utilisateur. Le manuscrit s'articule autour de quatre chapitres, encadrés par une introduction générale et une conclusion générale.

Le premier chapitre nous introduisons tout d'abord la notion de biométrie, en présentant la structure générale d'un système biométrique ainsi que les principaux critères permettant son évaluation. Une attention particulière est portée à la description des différentes modalités biométriques ainsi qu'aux domaines d'application de la biométrie. Nous mettons un accent spécifique sur l'utilisation des empreintes palmaires.

Le deuxième chapitre est dédié à la biométrie palmaire. Nous y introduisons ses principes de reconnaissance ainsi que les différentes techniques d'acquisition des empreintes palmaires. Ce chapitre aborde également les étapes de traitement et de prétraitement des images palmaires, en mettant en lumière les principaux défis rencontrés dans ce domaine ainsi que les solutions proposées pour y faire face.

Dans **le troisième chapitre**, nous présentons les notions d'intelligence artificielle, de machine Learning et d'apprentissage profond. Nous commençons par exposer les différents types d'apprentissage, puis nous abordons les divers types de réseaux de neurones ainsi que leur construction. Avant de conclure ce chapitre, nous présentons quelques algorithmes d'apprentissage profond, suivis d'exemples de domaines d'application.

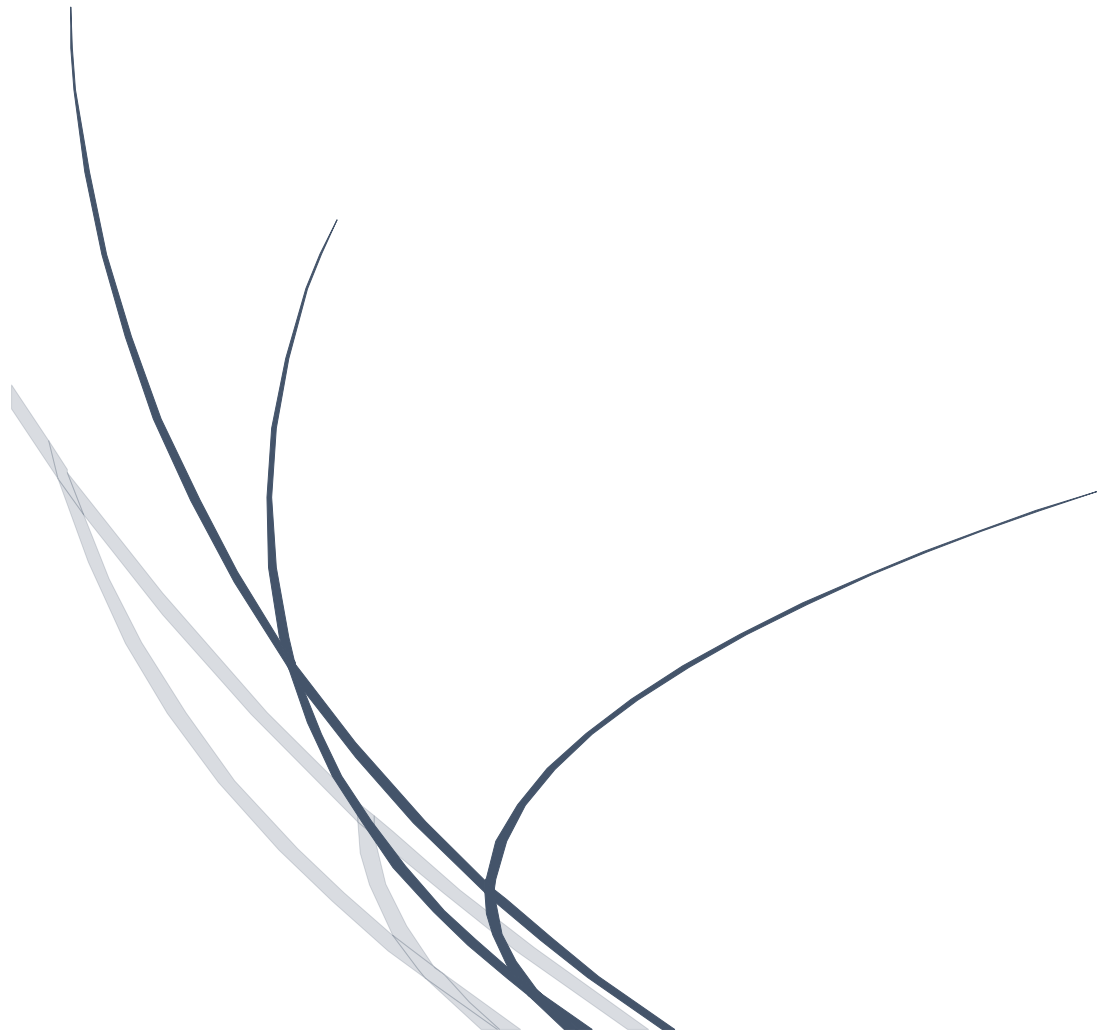
Le **quatrième chapitre** présente les bases de données multispectrales **MSCASIA** et **MSPOLYU**, utilisées pour l'évaluation des performances en reconnaissance palmaire. Il décrit une approche hybride combinant **GoogleNet** pour l'extraction de caractéristiques, **LDA** pour la réduction de dimension, et **k-**

NN pour la classification. Les résultats obtenus montrent des taux de reconnaissance élevés, notamment sur les spectres. Une comparaison avec d'autres méthodes met en évidence la supériorité de cette combinaison. Le chapitre conclut par une discussion sur les limites, la généralisation inter-bases, et propose des pistes de recherche futures.

Enfin, une **conclusion générale** vient clore ce travail en dressant un bilan global des contributions apportées, tout en ouvrant la voie à de futures perspectives de recherche dans le domaine de la reconnaissance biométrique par empreintes palmaires.

CHAPITRE I

LA BIOMETRIE



I.1 Introduction

La biométrie est la science qui détermine l'identité d'un individu, elle se base sur des mesures physiologiques, chimiques ou comportementales d'un ou de plusieurs de ses attributs biologiques. La pertinence de la biométrie dans les sociétés modernes a été augmentée à cause du grand besoin de la sécurité et à la nécessité des systèmes de management (gestion) d'identités à grande échelle, qui s'appuient fonctionnellement sur la détermination précise de l'identité d'un individu, dans un contexte d'applications largement interconnectées. Comme exemples de ces applications : le partage des ressources informatiques dans un réseau public, l'accès de haute sécurité aux zones nucléaires, les transactions bancaires à distance, ou l'embarquement des vols commerciaux. En plus, la prolifération des services web (ex., les banques en ligne) et le déploiement des centres de services clientèles décentralisés (ex., les cartes de crédit) ont souligné la nécessité des systèmes de managements d'identité fiables pouvant accueillir un grand nombre d'individus. Dans ce chapitre, nous introduisons tout d'abord quelques notions et définitions de bases liées à la biométrie, nous décrivons le principe de fonctionnement d'un système biométrique ainsi que les outils d'évaluations utilisés pour mesurer leurs performances, nous donnons un bref aperçu des modalités biométriques les plus répandues, tout en accordant une attention particulière à la reconnaissance par l'empreinte palmaire parmi les autres modalités biométriques, puisqu'elles constituent l'objectif de cette mémoire .[4]

I.2 Définitions et propriétés

Une définition exacte de la biométrie est donnée par Jain et al. la désignant comme la science d'établir automatiquement l'identité d'un individu sur la base de ses caractéristiques physiologiques ou comportementales. Ainsi, la biométrie vérifie l'identité d'un individu par ce qu'il est et non par ce qu'il possède (clé, carte d'accès, ...) ni par ce qu'il sait (mot de passe), ce qui la rend moins vulnérable, que les moyens classiques d'identification, aux tentatives de contrefaçon. En théorie, la plupart des traits physiologiques ou comportementaux humains peuvent être utilisés en tant que modalités biométriques. Toutefois, pour tenir dans un système biométrique, potentiellement précis, pratique et rentable, le trait/caractéristique utilisé doit également satisfaire à une série d'exigences proposées dans:

- L'universalité : toute personne doit posséder le trait biométrique,
- L'unicité : une probabilité quasi nulle que deux personnes soient les mêmes selon le trait caractéristique biométrique,
- La permanence : la stabilité du trait biométrique dans le temps,

- La mesurabilité : la quantification du trait biométrique d'une manière pratique,
- La performance : la précision et la vitesse de la reconnaissance à travers la caractéristique biométrique,
- L'acceptabilité : l'accord du public pour la mesure de la caractéristique,
- La non-circonvension : le degré de facilité/difficulté avec laquelle le système peut être trompé en falsifiant la caractéristique biométrique.

Les modalités adoptées dans les systèmes biométriques possèdent ces propriétés mais à des degrés différents. Un compromis est alors réalisé lors du choix de la modalité en fonction des besoins de l'application biométrique.

De nombreuses modalités biométriques ont été proposées et sont utilisées dans des applications variées. Les modalités physiologiques se basent sur des caractéristiques morphologiques ou biologiques et comprennent le visage, l'oreille, l'iris, la rétine, les empreintes digitale et palmaire, la géométrie de la main, le réseau veineux, l'ADN, l'odeur de la peau, ... Les modalités comportementales utilisent un trait personnel du comportement tel que la voix, la dynamique de la signature, la démarche, la dynamique de frappe au clavier ou encore le mouvement des lèvres. Certains de ces éléments biométriques ont un long historique et peuvent être considérés comme des technologies matures, tandis que d'autres sont encore de jeunes arènes de recherche. [5]

I.3 Architecture d'un système biométrique

I.3.1 Principaux Modules

Le système biométrique est un système pour identifier les tendances et le stockage des données à sauvegarder ou de les identifier dans la forme de matrices. Ensuite, le système est prêt à identifier les intrus. Ce système se compose de quatre unités : l'acquisition, l'extraction des caractéristiques, la comparaison (mesure de similarité) et la décision. L'inscription ou l'enrôlement est utilisé pour une future comparaison tandis que la décision est de reconnaître la personne ou non

- **Acquisition des données** : Cette phase collecte les données biométriques des personnes clients.
Plusieurs processus industriels peuvent être utilisés pour l'acquisition telle qu'un appareil photo, un lecteur d'empreintes digitales, etc.
- **Extraction des caractéristiques** : Les images sont traitées pour en extraire des caractéristiques du procédé. Ce processus sert à éviter les informations inutiles qui existent. Donc, ce module sert à traiter l'image afin d'extraire uniquement les caractéristiques

biométriques, sous forme d'un vecteur ou Template, qui ensuite peuvent être utilisées pour reconnaître les personnes. Ces caractéristiques sont uniques à chaque personne et stable.

- **Comparaison** : Dans ce module, les caractéristiques biométriques extraites sont comparées avec un vecteur précédemment stocké dans la base de données et en marquant le degré de similitude (différence ou distance).
- **Décision** : Vérifie l'identité affirmée par un utilisateur ou détermine l'identité d'une personne basée sur le degré de similitude entre les caractéristiques extraites et le(s) vecteur(s) stocké(s). [6]

I.3.2 Fonctionnement des systèmes biométriques

Un système biométrique peut être un système d'identification (reconnaissance) ou un système d'authentification (vérification), qui sont définis comme suit :

➤ L'identification

L'identification effectue un appariement d'un à plusieurs (1 : N) entre un nouvel échantillon biométrique capturé, et les modèles biométriques stockés dans une base de données biométrique afin de tenter de déterminer l'identité d'une personne inconnue.

➤ L'authentification

L'authentification effectue une correspondance un à un (1 : 1) entre un nouvel échantillon biométrique capturé, et un modèle biométrique spécifique stocké dans une base de données biométrique, pour tenter de vérifier que la personne est bien la personne qu'elle prétend être. er l'identité d'une personne inconnue.

➤ L'enrôlement

Les systèmes identification/authentification comprennent deux phases principales comme le montre la **Figure I.1** : l'enrôlement ou l'apprentissage et la reconnaissance. La phase enrôlement est commune à l'authentification et l'identification. C'est une phase préliminaire de tout système biométrique qui se fait off-line. L'enrôlement consiste à représenter les caractéristiques physiques ou comportementales d'un utilisateur dit "référence" sous forme d'un modèle biométrique appelé

"signature" puis enregistrer ses informations dans une base de données. [7]

Identification

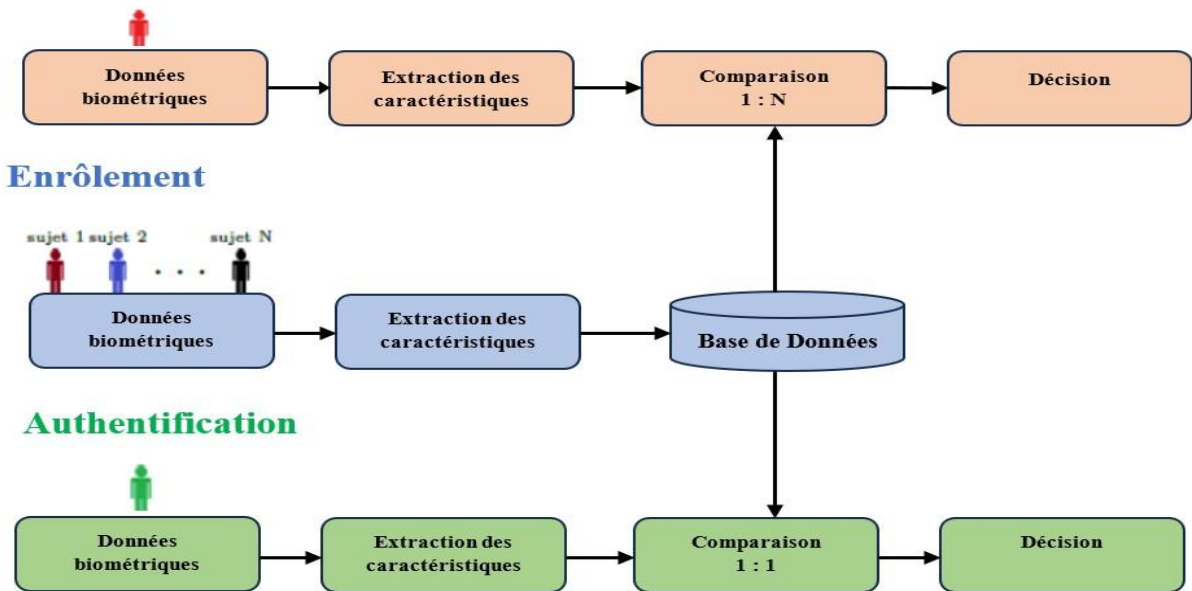


Figure I.1: Structure des systèmes biométrique [7].

I.4 Modalités biométriques émergentes

Dans cette section, nous nous penchons sur les différentes modalités biométriques qui ont été introduites de manière relativement récente. Ces modalités représentent une nouvelle orientation dans les tendances actuelles des recherches en biométrie. On peut les classer en trois catégories.

Figure I.2.

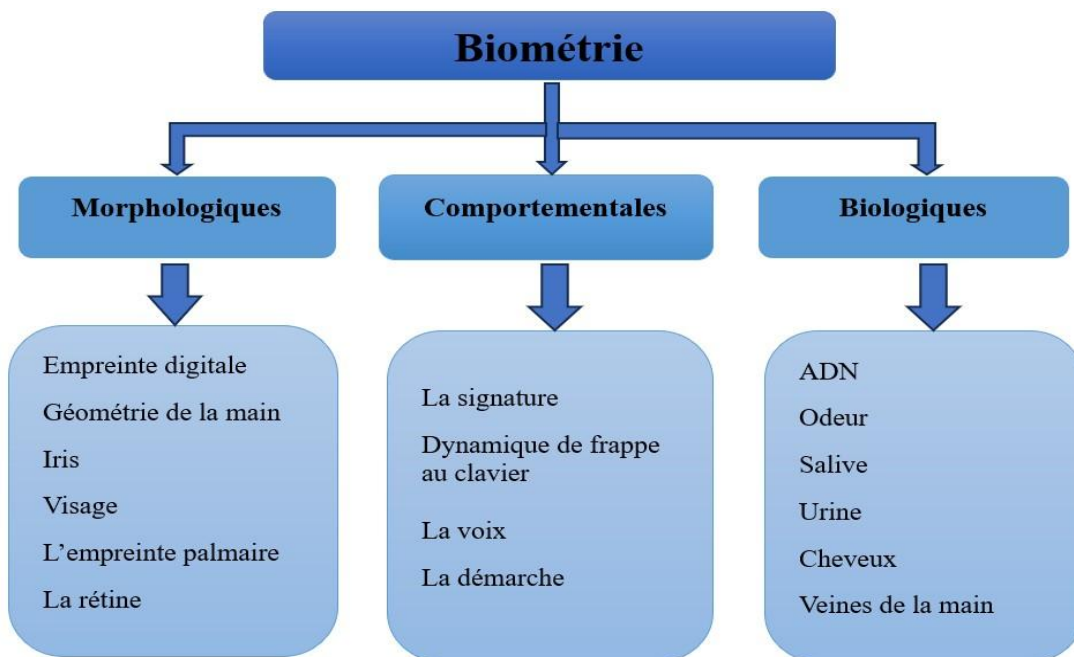


Figure I.2: Catégories des modalités biométriques [15].

I.4.1 Modalités morphologiques (physiologiques)

Elle repose sur l'identification de caractéristiques physiologiques spécifiques, uniques et permanentes à chaque individu. Cette catégorie inclut des éléments tels que la géométrie de la main, l'iris de l'œil, l'empreinte palmaire, les empreintes digitales, les traits du visage, etc.

I.4.1.1 La géométrie de main :

Cette modalité est généralement utilisée pour le contrôle d'accès physique ainsi que pour le pointage horaire, notamment dans certaines administrations. Elle consiste en l'analyse de 90 caractéristiques de la main, telles que la longueur et la largeur des doigts, la forme de la paume, les articulations, ainsi que le dessin des lignes de la main. Lors de la phase de capture, la personne place sa main sur une platine, et les emplacements du pouce, de l'index et du majeur sont marqués, comme illustré sur la photo ci-dessus. Une analyse sous deux angles différents est effectuée pour obtenir un modèle en trois dimensions. **Figure I.3. [8]**

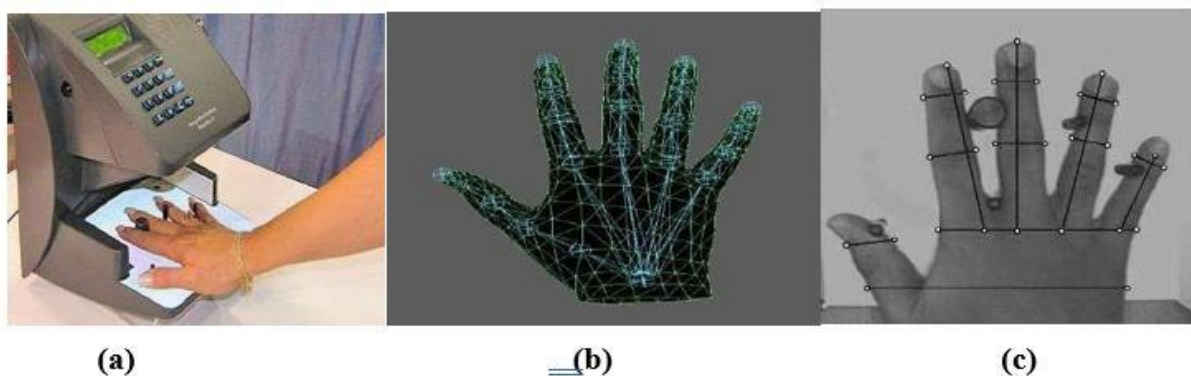


Figure I.3: Dispositif de reconnaissance de la géométrie de main (a) , (b) et (c)

I.4.1.2 Empreinte digitale

Une empreinte digitale est un motif formé par les lignes de la peau des doigts, des paumes des mains, des orteils ou de la plante des pieds. Ce motif se développe pendant la période fœtale. Il existe deux types d'empreintes : l'empreinte directe (qui laisse une marque visible) et l'empreinte latente (résidu de saleté, de sueur ou d'autres substances laissées sur un objet). Les empreintes sont uniques et immuables, ce qui signifie qu'elles ne changent pas au fil du temps, sauf en cas d'accident, comme une brûlure, par exemple. La probabilité de trouver deux empreintes digitales identiques est de 1 sur 10^{24} . Même les jumeaux, bien qu'ils viennent de la même cellule, auront des empreintes très similaires, mais pas identiques. **Figure I.4. [9]**



Figure I.4: Empreinte digitale

I.4.1.3 L’empreinte palmaire

La paume de la main désigne la partie intérieure de la main (celle qui n'est pas visible lorsque la main est fermée), s'étendant du poignet aux racines des doigts. L'empreinte palmaire est donc l'image laissée par la paume de la main lorsqu'elle exerce une pression sur une surface. Autrement dit, elle représente le motif de la paume, illustrant les caractéristiques physiques de la peau, telles que les lignes principales et les rides, les points, les minuties et la texture. [10].

Les principales caractéristiques de l’empreinte palmaire sont les trois lignes principales, appelées : « ligne du cœur », « ligne de la tête » et « ligne de vie », ainsi que les rides et les crêtes Figure I.5. [11]



Figure I.5: L’empreinte palmaire

I.4.1.4 Le visage

La reconnaissance faciale (visage) permet d’adapter la vérification biométrique à toutes les situations. C'est une technologie très efficace qui est utilisée dans de nombreuses applications liées à la sécurité. Elle est par exemple un outil très fiable pour aider les forces de police à identifier des criminels, ou bien pour permettre aux services de douanes de vérifier l’identité des voyageurs. Actuellement, avec la numérisation des échanges, l’usage de cette technologie est en train de s’étendre au monde des entreprises. Utilisée dans des applications commerciales, la reconnaissance faciale permet par exemple de sécuriser des transactions en ligne. La reconnaissance faciale est sans contact et son utilisation ne nécessite aucun outil spécifique, ce qui en fait la solution idéale pour l’identification de personnes dans une foule ou dans des espaces publics. **Figure I.6.** [12]



Figure I.6: Le visage

I.4.1.5 L'iris

Est la membrane colorée située entre le blanc de l'œil et la pupille, l'iris est composé d'une multitude de tubes très fins qui s'entrecroisent, procurant à l'iris une forme particulière et unique qui ne varie que très peu au cours d'une vie. La capture de l'iris se fait à l'aide d'une caméra qui va dans un premier temps positionner l'iris par rapport à l'ensemble de l'œil. Ensuite, la caméra scanne l'image de l'iris pour en analyser les points caractéristiques. Le dispositif analyse notamment la position, la longueur et le relief des tubes qui composent l'iris. Enfin, en ayant retenu au-dessus de 200 points distinctifs, l'ordinateur relié à la caméra procède à la comparaison de l'iris avec la banque de données des identifiants possibles. Le processus d'identification ne prend que quelques secondes. Spécifions que l'image analysée est captée en noir et blanc. La couleur de l'oeil n'est ainsi pas prise en compte dans l'analyse, ce qui annule les biais causés par les changements de couleur de l'iris chez certaines personnes **Figure I.7. [13]**.



Figure I.7: L'iris

I.4.1.6 La rétine

La rétine est la paroi interne de l'œil, opposée à la lentille, sur laquelle se projettent les images que nous percevons. Cette paroi est recouverte d'un réseau de vaisseaux sanguins, qui forme un motif unique à chaque individu. **Figure I.8. [14]**

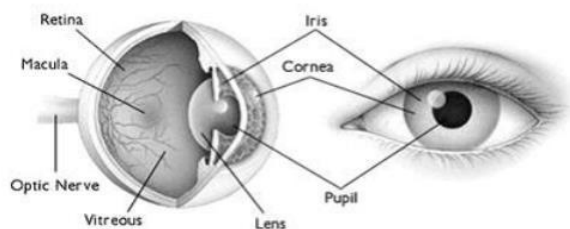


Figure I.8: La rétine

I.4.1.7 L’empreinte de l’oreille

L'utilisation des oreilles pour l'identification des individus est explorée depuis plus d'un siècle. Les recherches continuent de débattre sur le caractère unique des oreilles, ou du moins sur leur suffisance pour être utilisées comme modalité biométrique. Bien que les applications reposant sur la forme de l'oreille ne soient pas encore largement répandues, le sujet demeure pertinent, notamment dans le cadre des enquêtes criminelles. Selon Burge et Burger, la biométrie de l'oreille représente une « approche passive et prometteuse de l'identification humaine automatisée ».

Figure I.9. [16]



Figure I.9: L’empreinte de l’oreille

I.4.2 Modalités comportementale

Elle est basée sur l’analyse de certains comportements d’une personne.

I.4.2.1 L’écriture (la signature)

Les systèmes de reconnaissance de l’écriture analysent les caractéristiques spécifiques d'une signature, telles que la vitesse, la pression exercée sur le stylo, les mouvements, ainsi que les points et les intervalles de temps pendant lesquels le stylo est levé. Ces systèmes reposent généralement sur l'utilisation d'un stylo électronique sur une tablette graphique, tout en étudiant la dynamique complète de la signature, notamment la vitesse, la direction, la pression, la durée pendant laquelle le stylo reste en contact avec le papier, le temps nécessaire pour réaliser la signature, ainsi que les moments où le stylo est relevé ou abaissé. **Figure I.10. [15]**



Figure I.10: L'écriture (la signature)

I.4.2.2 La dynamique de frappe au clavier

Il s'agit d'une technique de reconnaissance des individus basée sur le rythme de frappe propre à chaque utilisateur. Cette méthode est appliquée au mot de passe, rendant ainsi sa reproduction plus difficile. Lors de la mise en place de cette technique, l'utilisateur doit saisir son mot de passe une dizaine de fois de suite. Un algorithme analyse le temps d'appui sur chaque touche ainsi que l'intervalle entre chaque pression. Les différentes saisies sont ensuite moyennées pour créer un « profil de frappe » de l'utilisateur, qui servira de référence. Lors des accès ultérieurs, le mot de passe saisi sera comparé à ce profil de frappe pour vérifier l'identité de l'utilisateur. **Figure I.11 [17].**



Figure I.11: La dynamique de frappe au clavier

I.4.2.3 La voix (Reconnaissance vocale)

La reconnaissance par voix utilise les caractéristiques vocales pour identifier les personnes en utilisant des phrases mot de passe. L'identification de la voix est considérée par les utilisateurs comme une des formes les plus normales de la technologie biométrique, car elle n'est pas intrusive et n'exige aucun contact physique avec le lecteur du système. **Figure I.12 [18].**



Figure I.12: La voix (Reconnaissance vocale)

I.4.2.4 La démarche

Il est également possible de modéliser la démarche d'une personne à l'aide de diverses techniques, mais le problème réside dans le fait que ce système peut être facilement trompé. La biométrie de la démarche se base sur les caractéristiques uniques de la marche d'un individu. Il convient de préciser que la démarche n'est pas influencée par la vitesse de marche de la personne. Certains chercheurs distinguent la démarche de la reconnaissance de la démarche, en soulignant que la démarche représente une combinaison cyclique de mouvements associés à la locomotion humaine, tandis que la reconnaissance de la démarche fait référence à l'identification de styles particuliers de marche, de pathologies, etc. les paramètres communs de l'analyse de la marche sont **Figure I.13 [19]** :

- Paramètres cinématiques tels que le genou, les mouvements de la cheville et les angles.
- Paramètres spatiotemporels tels que la longueur et la largeur des marches, la vitesse de Marche.



Figure I.13: La démarche

I.4.3 Modalités biologiques

Elle est basée sur l'identification de traits biologique particuliers.

I.4.3.1 L'ADN

L'analyse des empreintes génétiques est une méthode d'identification extrêmement précise, directement issue des avancées en biologie moléculaire. L'information génétique d'un individu

est unique, car aucun autre membre de l'espèce n'a la même combinaison de gènes codés dans l'acide désoxyribonucléique (ADN). L'ADN est ainsi l'outil d'identification par excellence. En quelques années, l'analyse des empreintes génétiques est devenue l'un des outils majeurs de la criminalistique, la science qui se spécialise dans l'identification des indices matériels. L'analyse de l'ADN est couramment utilisée en criminologie pour identifier une personne à partir de traces biologiques telles qu'un morceau de peau, un cheveu ou une goutte de sang. **Figure I.14. [20]**

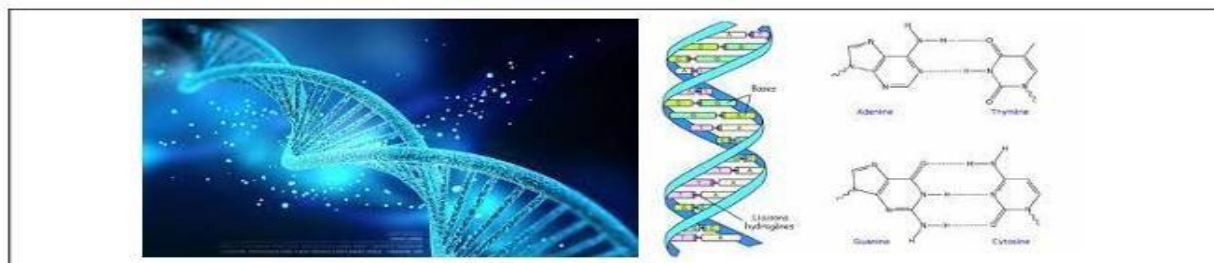


Figure I.14: Image d'ADN

I.4.3.2 Veines de la main

On a longtemps estimé que le modèle des veines dans l'anatomie humaine pourrait être unique à chaque individu. Ainsi, plusieurs types de balayage veineux ont été développés au fil des années, allant du balayage de la main, au balayage du poignet, et plus récemment, au balayage du doigt. Bien que la plupart de ces techniques aient été déployées sur le terrain et aient montré un potentiel pour constituer la base d'un système biométrique fiable de vérification d'identité, elles rencontrent un obstacle majeur : ce n'est pas une question de faisabilité ou d'efficacité technique, mais plutôt une question de réalité du marché. En effet, la domination des systèmes basés sur les empreintes digitales, le visage et l'iris, disponibles à des prix variés, empêche ces technologies distinctes de gagner une part significative du marché, faute d'un avantage clair et irrésistible.

Même les techniques primaires, telles que la géométrie de main, ont une base qui est peu susceptible d'être réalisée par une technique plus récente de performance comparable. En conséquence, pour n'importe quelle nouvelle technique biométrique prenant place dans le marché, elle doit gagner le terrain et offrir des avantages clairs qui ne peuvent pas être réalisés par des méthodes contemporaines. Les diverses réalisations de balayage des veines. Bien qu'assurément intéressantes, ne peuvent lutter que peu dans ce contexte. Cependant, le temps peut s'avérer un niveau intéressant dans ces contextes et les demandes de la technique de balayage de veines peuvent s'accroître. **Figure I.15. [21].**



Figure I.15: Reconnaissance des veines

I.5 Représentation comparative entre quelques Techniques Biométriques

Techniques biométriques	universelles	uniques distinctif	Permanente	Enregistrable Mesurable	Performance Acceptabilité
Empreintes digitales	Moyenne	Haute	Haute	Moyenne	Moyenne
Visage	Haute	Faible	Moyenne	Haute	Haute
Iris	Haute	Haute	Haute	Moyenne	Faible
Rétine	Haute	Haute	Moyenne	Faible	Faible
ADN	Haute	Haute	Haute	Faible	Faible
Signature	Faible	Faible	Faible	Haute	Haute
Voix	Moyenne	Faible	Faible	Moyenne	Haute
Démarche	Moyenne	Faible	Faible	Haute	Haute
Frappe clavier	Faible	Faible	Faible	Moyenne	Moyenne
Veines d main	Moyenne	Moyenne	Moyenne	Moyenne	Moyenne

Tableau I.1: Comparaison entre quelques techniques biométriques

I.6 Performance des systèmes biométriques

L'évaluation des systèmes biométriques peut prendre en compte plusieurs aspects, tels que la précision, la rapidité, la facilité d'utilisation, le coût, la confidentialité, l'acceptation sociale, l'évolutivité, l'interopérabilité et la sécurité. Les méthodes utilisées pour évaluer ces différents critères proviennent de divers domaines, allant de l'ingénierie et de l'informatique, aux sciences sociales et à l'économie. La précision est généralement l'aspect le plus couramment utilisé pour évaluer les systèmes biométriques, et elle s'applique à toutes les modalités biométriques. Toutefois, d'autres critères sont également essentiels pour compléter cette mesure. En particulier, neuf aspects différents doivent être pris en considération.

-
- **La précision** : est une mesure de capacité du système biométrique à distinguer les individus en fonction du trait biométrique utilisé.
 - **La vitesse** : est la quantité de temps nécessaire pour effectuer le processus d'enregistrement et d'authentification. La vitesse est particulièrement importante pour les systèmes biométriques d'identification. De plus, le temps nécessaire pour les différentes étapes du système biométrique doit être considéré séparément. Par exemple, une longue période pour l'étape d'acquisition peut entraîner un système moins utilisable.
 - **L'utilisabilité** : décrit la facilité d'utilisation du système ainsi que la facilité avec laquelle les utilisateurs peuvent apprendre à l'utiliser. L'utilisabilité est mesurée en utilisant le temps d'acquisition et le nombre d'échantillons capturés de manière incorrecte. Cependant, autres facteurs sociaux et personnels peuvent influencer la facilité d'utilisation d'un système.
 - **Le coût** : englobe le coût de la conception, du développement du système matériel et de la mise en œuvre des algorithmes de reconnaissance. Les systèmes biométriques coûteux ont généralement une précision et une vitesse supérieures à celles des systèmes moins coûteux, mais un coût moindre peut favoriser une plus grande utilisation d'un système biométrique.
 - **La confidentialité** : correspond à la possibilité qu'un trait biométrique puisse être volé ou mal utilisé par le système biométrique. Comme les traits biométriques ne peuvent pas être modifiés, un système biométrique doit être en mesure de protéger les données personnelles, et protéger aussi la vie privée des utilisateurs.
 - **L'acceptation sociale** : fait référence à la façon dont le système est perçu par les utilisateurs. Cette mesure peut être liée aux connaissances du public sur la performance du système, les risques perçus sur la vie privée, son caractère invasif et sa facilité d'utilisation. Les facteurs humains peuvent également jouer un rôle important dans l'acceptation sociale d'un système biométrique.
 - **L'évolutivité** : décrit la capacité du système à fonctionner efficacement lorsque la charge augmente, par exemple, un plus grand nombre d'utilisateurs inscrits ou une augmentation du nombre de requêtes à la base de données biométrique centrale. Cet aspect peut être lié à l'architecture matérielle choisie (par exemple, la fréquence du processeur, la vitesse du disque dur, la bande passante du réseau) ou l'efficacité de la mise en œuvre logicielle des algorithmes de reconnaissance biométrique.
 - **L'interopérabilité** : fait référence à la compatibilité de différents systèmes biométriques basés sur le même trait biométrique. Ce facteur peut être influencé par le type et la qualité des échantillons (par exemple, un appareil peu coûteux produira un échantillon différent

d'un appareil haut de gamme), le format de données utilisé pour stocker les modèles, la mesure de similarité calculée lors de phase de comparaison, etc. Pour pallier à ces problèmes, des normes biométriques sont utilisées.

- **La sécurité** : est la robustesse du système contre les attaques. En particulier, la sécurité contre les fausses biométries doit être étudiée pour concevoir un système biométrique efficace. De plus, la robustesse de l'architecture informatique et de l'infrastructure du réseau aux attaques et aux logiciels malveillants doit être prise en compte. [22]

Le tableau suivant présente la Comparaison entre les différentes caractéristiques des modalités biométriques.

- * les caractéristiques (Universalité, Unicité, Permanence, collectable, Performance, précision et sécurité) pour les modalités biométriques les plus utilisées par niveau (E=Elevé, M=Moyen et B= bas).

I.7 Mesure de la performance d'un système biométrique

Les mesures de performance expriment les caractéristiques de fonctionnement du système de reconnaissance et permettent, ainsi de faire des comparaisons entre différents systèmes.

Il existe différents paramètres pour évaluer un système de reconnaissance : les taux de faux rejets (False Reject Rate ou FRR) et taux de fausses acceptations (False Accept Rate ou FAR), taux d'égale erreur (Equal Error Rate ou EER)

I.7.1 Taux de faux rejets (False Reject Rate ou FRR)

Ce taux représente le pourcentage de personnes censées être reconnues, mais qui sont rejetées par le système. Il s'agit du rapport entre le nombre de fausses rejets (*FR*) et le nombre total de tests effectués sur des personnes légitimes.

$$FRR = \frac{\text{le nombre de clients rejetés}}{\text{le nombre total d'accès clients}} \quad (1.1)$$

I.7.2 Taux de fausses acceptations (False Accept Rate ou FAR)

Ce taux représente le pourcentage de personnes censées ne pas être reconnues, mais qui sont tout de même acceptées par le système.

Le FAR, est égal au nombre de fausses acceptations (*FA*) divisé par le nombre total d'accès imposteurs.

$$FAR = \frac{\text{Le nombre d'imposteurs acceptés}}{\text{Le nombre total d'accès imposteurs}} \quad (1.2)$$

I.7.3 Taux d'égale erreur (Equal Error Rate ou EER)

Ce taux est calculé à partir des deux premiers taux préalablement décrits et constitue un point de mesure de performance courant. Ce point correspond à l'endroit où $FRR = FAR$, c'est-à-dire le meilleur compromis entre les faux rejets et les fausses acceptations.

L'**EER** peut être calculé à l'aide de l'équation suivante :

$$EER = \frac{\text{Le nombre de fausses acceptations} + \text{Le nombre de fausses rejets}}{\text{Le nombre total d'accès}} \quad (1.3)$$

La **Figure.I.16** montre le FRR et le FAR à partir de distributions des scores authentiques et imposteurs, tandis que l'EER est représenté sur la **Figure. I.18**.

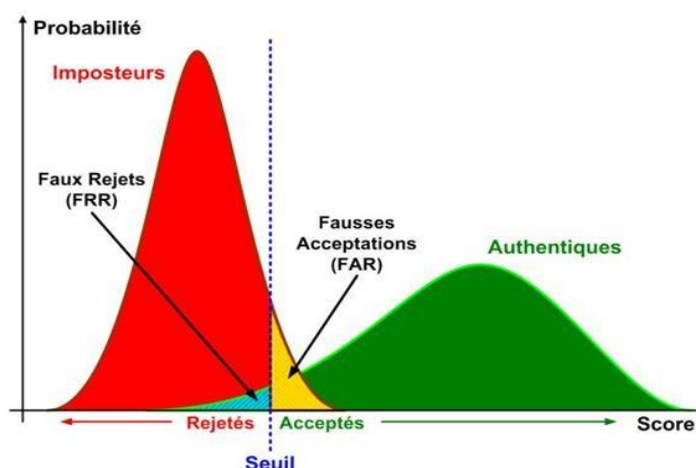


Figure I.16: Illustration du FFR et du FAR

Selon la nature (authentification ou identification) du système biométrique, il existe deux façons d'en mesurer la performance :

- Lorsque le système opère en mode authentification, on utilise ce que l'on appelle **une courbe ROC** (pour Receiver Operating Characteristic en anglais). **La courbe ROC (Figure. I.17)** trace le taux de faux rejets en fonction du taux de fausses acceptations. Plus cette courbe tend à épouser la forme du repère, plus le système est performant, c'est-à-dire possédant un taux de reconnaissance global élevé.

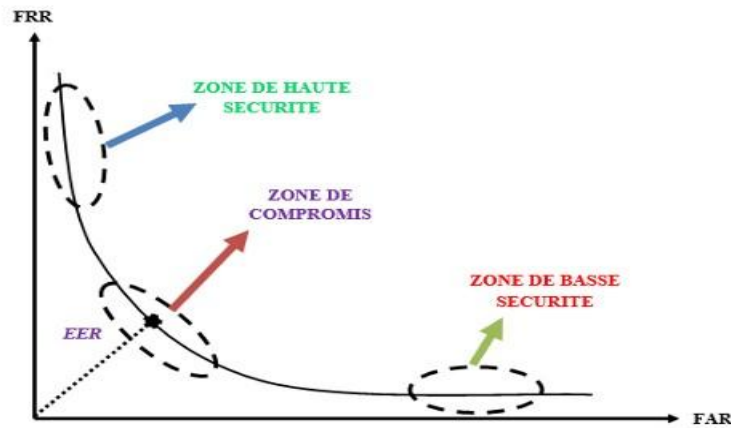


Figure I.17: Courbe ROC

En revanche, dans le cas d'un système utilisé en mode identification, on utilise ce que l'on appelle **une courbe CMC** (pour Cumulative Match Characteristic en anglais). **La courbe CMC** (Figure. 1.18) donne le pourcentage de personnes reconnues en fonction d'une variable que l'on appelle le rang. On dit qu'un système reconnaît au rang 1, lorsqu'il choisit la plus proche image comme résultat de la reconnaissance. On dit qu'un système reconnaît au rang 2, lorsqu'il choisit parmi deux images, celle qui correspond le mieux à l'image d'entrée, etc. On peut donc dire que plus le rang augmente, plus le taux de reconnaissance correspondant est lié à un niveau de sécurité faible. [23]

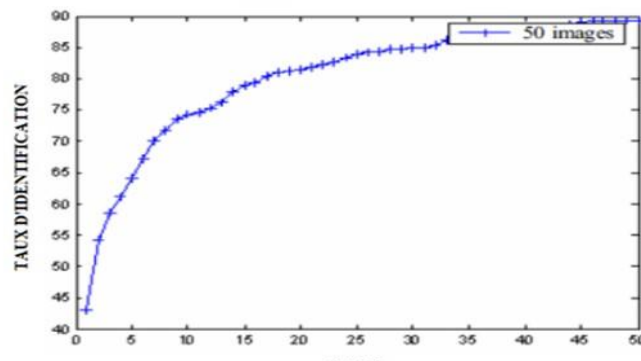


Figure I.18: Courbe CMC

I.8 Champ d'application des systèmes biométriques

En ce qui concerne les applications, l'authentification des utilisateurs via des systèmes biométriques est utilisée dans divers domaines où un accès contrôlé est nécessaire, comme dans les applications bancaires, les lieux hautement sécurisés tels que les sièges du gouvernement, les parlements, l'armée, ou encore les services de sécurité. La reconnaissance biométrique, quant à

elle, est fréquemment employée par la police et les services d'immigration dans les aéroports, ainsi que pour la recherche dans les bases de données criminelles. Elle trouve également des applications dans le secteur civil, notamment pour l'authentification des cartes de crédit, des permis de conduire et des passeports. Parmi les autres domaines pouvant utiliser des systèmes biométriques pour contrôler l'accès, nous citons :

I.8.1 Contrôle d'accès physiques aux locaux

Parmi les autres applications des systèmes biométriques pour le contrôle d'accès, on peut citer : les salles informatiques, les sites sensibles tels que les services de recherche, les installations nucléaires et les bases militaires, les coffres-forts avec serrure électronique, les distributeurs automatiques de billets, le contrôle des adhérents dans un club, les cartes de fidélité, la gestion et le contrôle des temps de présence, les systèmes anti-démarrage pour voitures, les cartes d'identité nationales, les permis de conduire, la sécurité sociale, ainsi que le contrôle des frontières et des passeports.

I.8.2 Contrôle d'accès logiques aux systèmes d'informations

Lancement du système d'exploitation, accès au réseau informatique, commerce électronique, transaction (financière pour les banques, données entre entreprises), tous les logiciels utilisant un mot de passe, terminaux d'accès à internet, téléphones portables.

I.8.3 Applications légales (juridique)

Telles que l'identification de corps et l'identification de cadavre, la recherche criminelle, l'identification de terroriste, etc.



Figure I.19: Les application des systèmes biométriques

I.9 Les avantages et les inconvénients de la biométrie

Avantages de la biométrie	Inconvénients de la biométrie
 Sécurité renforcée : difficile à falsifier	 Risque de piratage des données biométriques
  Identification unique et personnelle	 Problèmes de confidentialité et de protection des données personnelles
 Gain de temps (accès rapide aux systèmes)	 Coût élevé de mise en place des systèmes biométriques
 Facilité d'utilisation (pas besoin de mémoriser de mot de passe)	 Risque d'erreurs (faux positifs ou faux négatifs)
 Impossibilité d'oublier ou de perdre (empreinte, visage, etc.)	 Changements biologiques possibles avec le temps (âge, blessures, etc.)

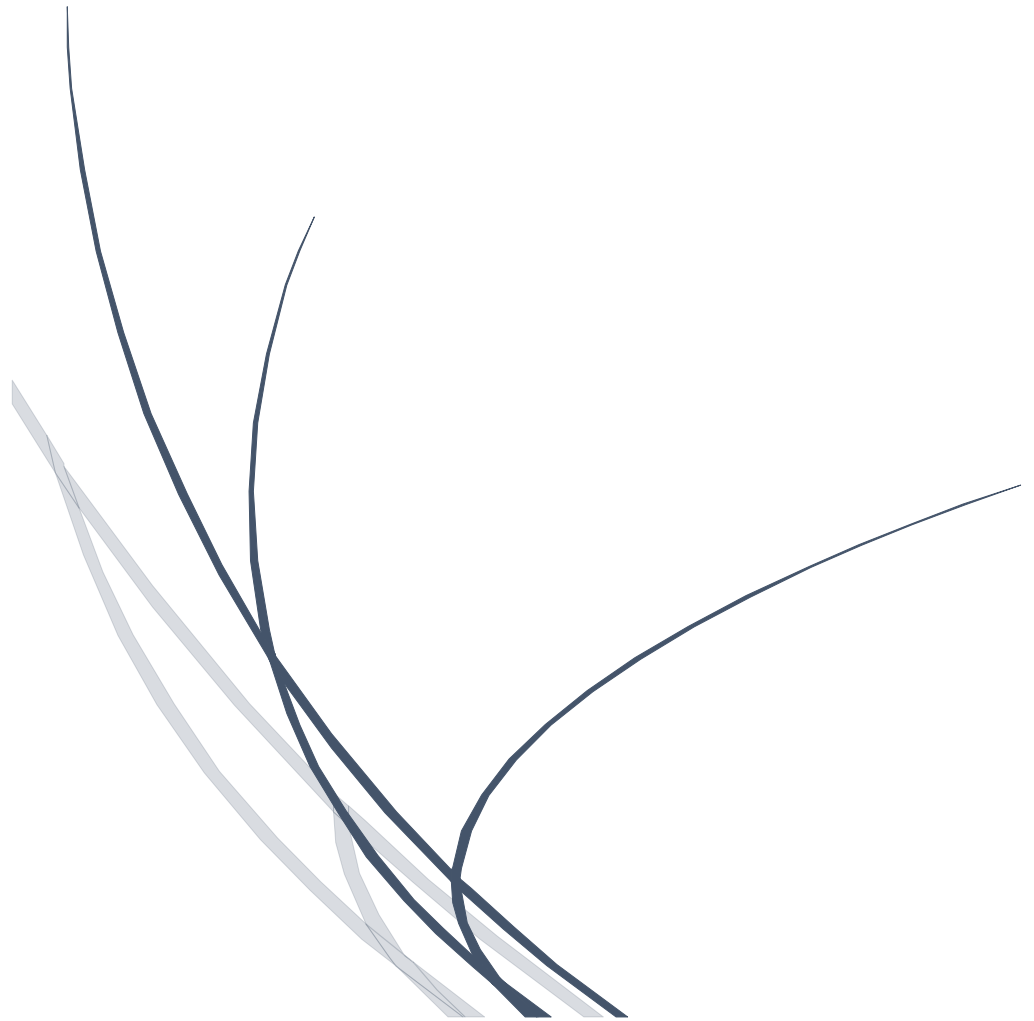
Tableau I.2: Les avantages et les inconvénients de la biométrie

I.10 Conclusion

La biométrie représente une avancée majeure dans le domaine de la sécurité et de l'identification des individus. Grâce à ses caractéristiques uniques et difficilement falsifiables, elle offre un niveau de fiabilité élevé dans de nombreux domaines, notamment l'accès aux systèmes informatiques, le contrôle aux frontières ou encore les services bancaires. Toutefois, cette technologie soulève également des préoccupations liées à la protection de la vie privée et au respect des droits fondamentaux. Il est donc essentiel de l'encadrer par des réglementations strictes et de garantir un usage éthique de ces données sensibles. [24]

CHAPITRE II :

Systeme de reconnaissance des empreintes palmaires



II.1 Introduction

La biométrie est en pleine croissance et tend à rejoindre d'autres technologies de sécurité comme la carte à puce. Dans les systèmes biométriques utilisés aujourd'hui, nous remarquons que la biométrie manuelle est l'un de ceux que les utilisateurs acceptent le plus parce qu'ils ne se sentent pas persécutés dans leur vie privée. Une enquête menée auprès de 129 utilisateurs a montré que l'utilisation du système biométrique de géométrie de la main au centre de loisirs de l'Université Purdue présente de nombreux avantages ; Les participants à l'enquête, 93 % ont aimé l'utilisation de la technologie, 98 % ont aimé sa facilité d'utilisation, et surtout plus personne d'autre ne trouve la technologie intrusive. C'est pourquoi ; De nos jours, la reconnaissance biométrique de la main a été développée avec un grand succès pour l'authentification et l'identification biométriques. Le processus de reconnaissance biométrique permet la reconnaissance d'une personne sur la base de caractéristiques physiques et comportementales. Parce que chaque personne a des caractéristiques qui sont propres pour elle : la voix, les empreintes digitales, les traits de son visage, sa signature... son ADN et d'ailleurs physionomie et physiologie de la main, une vue d'ensemble de ces systèmes peut être trouvée dans. La main est presque appropriée pour certaines situations et scénarios. Pour la modalité biométrique de la main, dans les principales caractéristiques utilisées ; On note : l'analyse de la longueur et de la largeur, la forme des phalanges, les articulations, les lignes de la main... etc La biométrie de la main présente une grande facilité d'utilisation d'un système basé sur. Cependant, le système matériel fait de temps en temps des erreurs en raison de la blessure de la main et de la façon dont la main vieillit. En plus de cela, le système donne une très grande précision avec un niveau de sécurité moyen requis. Cependant, à long terme, la stabilité est en quelque sorte moyenne et doit être améliorée. La plupart des travaux précédents ont élaboré des systèmes basés sur le contact biométrique de la main. [25]

II.2 Pourquoi la reconnaissance de l'empreinte palmaire ?

Avant de parler sur la reconnaissance biométrique des empreintes palmaires, nous devons tout d'abord présenter des généralités concernant cette modalité, son anatomie et leurs spécificités.

II.3 Définition de l'empreinte palmaire

L'empreinte palmaire représente le modèle de la paume de la main humaine illustrant les caractéristiques physiques du motif de sa peau tels que : les lignes (principales et rides), les points, les minuties et sa texture. En d'autres termes, si la partie intérieure de la main qui est non visible lorsque la main est fermée, du poignet aux racines des doigts, comme le montre la Figure 2.1. [26]

II.3.1 Avantages de l'empreinte palmaire

- Les empreintes palmaires contiennent plus d'information que les empreintes digitales ➤ Elles sont plus discriminantes.
- Les sources de capture d'empreintes palmaires sont beaucoup moins chères que celles des empreintes digitales.
- Les empreintes palmaires contiennent des caractéristiques distinctives additionnelles telles que les lignes principales et les ridules.
- En combinant toutes les caractéristiques d'une paume, il est possible d'établir un système robuste de biométrie. [27]

II.3.2 Caractéristique des empreintes palmaires

L'empreinte palmaire représente une surface interne très étendue de la main. Elle contient de nombreux traits caractéristiques qui peuvent être exploités pour la reconnaissance des individus. Grâce à cette large surface et à la richesse des caractéristiques biométriques qu'elle offre, l'empreinte palmaire est considérée comme robuste face aux bruits et unique pour chaque personne. Comparée à d'autres caractéristiques physiques, l'identification basée sur les empreintes palmaires présente plusieurs avantages [28, 29] :

- Traitement possible même à basse résolution ;
- Faible risque d'intrusion ;
- Stabilité des lignes caractéristiques dans le temps ;
- Taux élevé d'acceptation par les utilisateurs



Figure II.1: Paume de la main

II.3.2.1 Caractéristiques géométriques : Comme toute image, l'empreinte palmaire présente des caractéristiques géométriques telles que : la longueur, la largeur, et la surface. Ces caractéristiques ne sont pas distinctives mais peuvent tout de même être utiles pour une première vérification.

II.3.2.2 Les lignes principales :

L'empreinte palmaire est caractérisée par trois lignes principales, dites : plis de flexion (Figure II.2.) :

La ligne de tête.

La ligne de vie.

La ligne du cœur.

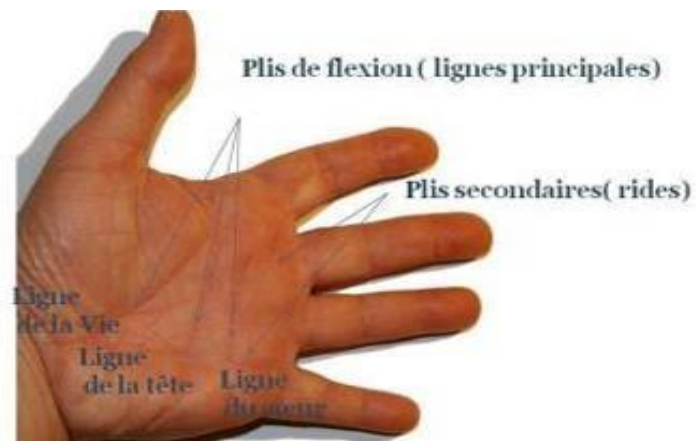


Figure II.2: Les plis de flexions de la paume de la main

II.3.2.3 Les rides (plis secondaires)

L'empreinte palmaire contient de nombreux autres plis qui diffèrent de ceux de flexion du fait qu'ils sont plus minces et plus irréguliers. Certains d'entre eux sont congénitaux, d'autres sont dus aux activités musculaires. Les lignes principales et les rides peuvent être observées facilement sur les images capturées à basse résolution. Comme les lignes principales seules ne fournissent pas une information distinctive sur la santé, les rides jouent un rôle important dans la reconnaissance

palmaire. Combinées aux lignes principales, elles fournissent une information distinctive pour la reconnaissance.

II.3.2.4 Les points de références : Les points de référence représentant les deux extrémités de la paume de la main a et b comme montré dans la Figure2.3

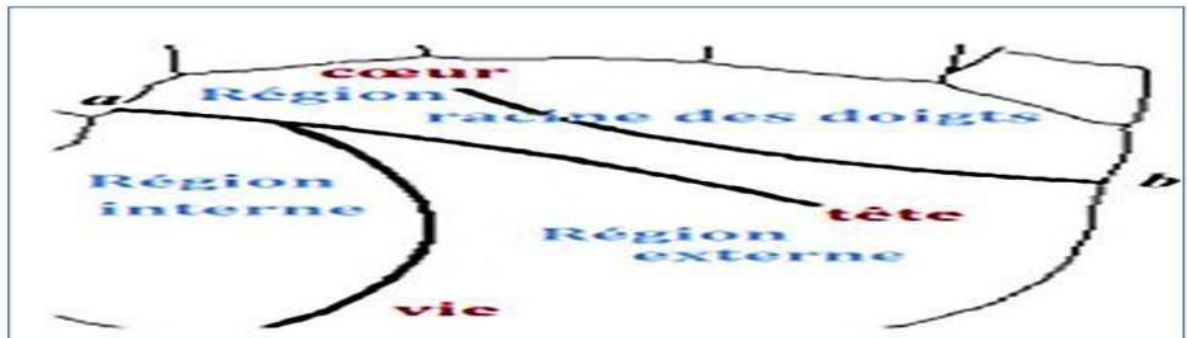


Figure II.3: Les points de référence de l'empreinte palmaire

Ils servent de point de repère lors de l'alignement et l'extraction des caractéristiques de l'empreinte palmaire. La taille de cette dernière peut être aussi estimée grâce à ces deux points et on ajoute que r David Zhang et Shu (chercheurs et professeurs à l'université polytechnique de Hong Kong) en 1996 pour remédier aux problèmes liés à : (i) la non visibilité d'une empreinte digitale, ou bien (ii) le coût élevé des appareils de capture des images de l'iris et de la rétine, ou encore (iii) les faibles taux de reconnaissance des autres modalités biométriques [26].

II.4 Types d'empreintes palmaires

Les empreintes palmaires ont été classées en trois groupes :

- Empreinte palmaire latente
- Empreinte palmaire brevetée
- Empreinte palmaire en plastique

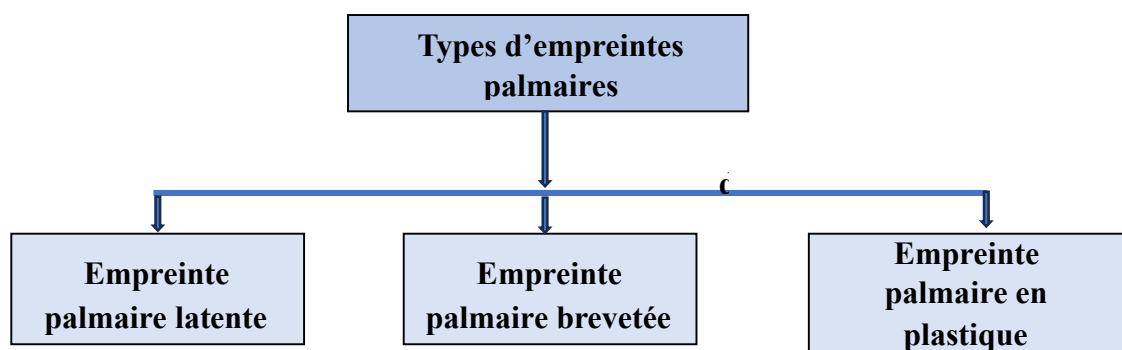


Figure II.4: Trois catégories de type d'empreintes palmaires [30]

II.4.1 Empreinte palmaire latente

La surface de la paume est considérée comme invisible lorsque les empreintes de la paume sont laissées par accident par la peau des crêtes de frottement sur une surface, qu'elles soient visibles ou non au moment de l'épreuve. De nombreuses méthodes peuvent être appliquées pour afficher la fraction ou la totalité de la paume, comme le traitement électronique, physique ou chimique. L'empreinte palmaire latente peut également être créée par l'extraction de la glande eccrine, du sang, de l'huile, de la peinture et de l'encre. Elle peut être partielle, déficiente, déformée, se chevaucher ou se combiner avec d'autres types d'empreintes [31].

II.4.2 Empreinte palmaire brevetée

Elles sont visuelles, évidentes et peuvent se former à la suite du transfert d'un objet étrange sur la surface de la paume. L'empreinte palmaire brevetée est visible et il n'est pas nécessaire de l'améliorer, comme c'est le cas pour le premier type d'empreintes, principalement photographiées [32].

II.4.3 Empreinte palmaire en plastique

L'empreinte palmaire plastique est l'empreinte des crêtes de frottement de la peau de la paume sur un article ou un instrument qui conserve la texture de la paume et la forme des crêtes. Ce type d'empreinte est visible et ne nécessite pas d'amélioration ; elle peut être photographiée et améliorée, comme une empreinte non plastique, et recouverte de la sécrétion naturelle d'un doigt. 46 Comme le type de matière est rarement disponible sur la scène de crime, cette forme de paume est rarement possible [33].

II.5 Les étapes de la reconnaissance de l'empreinte palmaire

Une chaîne de traitement dans un système de reconnaissance comprend plusieurs modules, et plusieurs espaces de travail. L'objectif de la reconnaissance des personnes est de définir une suite d'opérations permettant de passer de l'espace des données ou personnes, à l'espace des classes ou catégories de la personne estimée. Le processus d'un système de reconnaissance des personnes comporte plusieurs étapes qui peuvent être illustrées par le schéma suivant [34] :

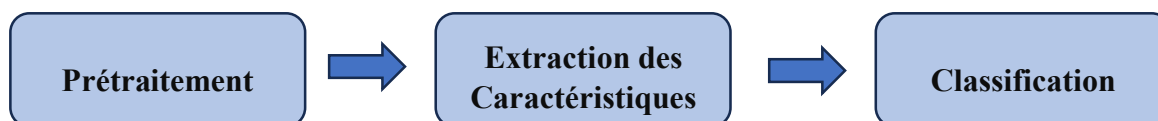


Figure II.5: Le principe des méthodes d'identification biométrique [34]

II.6 Des méthodes de prétraitement

II.6.1 Rehaussement des niveaux de gris

Le prétraitement de rehaussement des niveaux de gris consiste à renforcer certaines plages de niveaux de gris au détriment d'autres plages, pour mettre des objets en valeur.

Soient h_{max} et h_{min} les niveaux de gris maximum et minimum de l'image I. Le rehaussement des niveaux de gris consiste à appliquer aux niveaux de gris de l'image I une fonction croissante f telle que $f(h_{min}) = h_{min}$ et $f(h_{max}) = h_{max}$. Les courbes de la **figure II.6** décrivent ce type de transformation.[35]

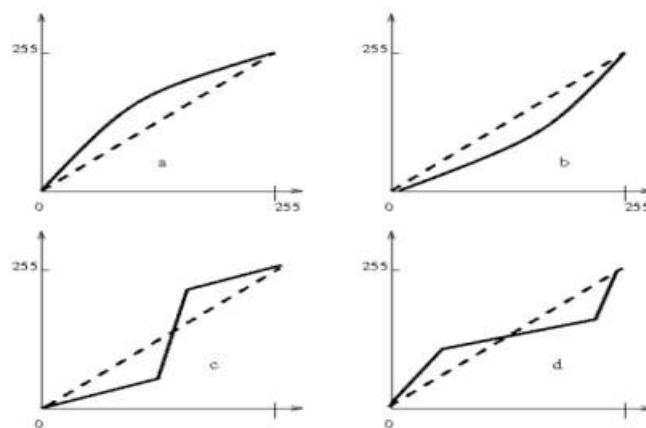


Figure II.6: Rehaussement des niveaux de gris.[35]

- a. Courbe représentant un rehaussement des bas niveaux de gris (zones foncées)
- b. Courbe représentant un rehaussement des hauts niveaux de gris (zones claires)
- c. Courbe représentant un rehaussement des niveaux de gris moyens
- d. Courbe représentant un rehaussement des niveaux de gris extrêmes.

II.6.2 Égalisation d'histogramme

L'égalisation d'histogramme sert à améliorer le contraste. Il faut la faire en s'assurant que les niveaux de gris des pixels de l'image résultante soient uniformément répartis (distribution uniforme des niveaux de gris). Cette transformation consiste à rendre le plus plat possible l'histogramme des niveaux de gris de l'image.

Cette transformation se définit à l'aide de l'histogramme cumulé de l'image. L'histogramme cumulé de l'image I est un vecteur de dimension 256. Chaque élément $h(i)$ représente le nombre de pixels de l'image dont le niveau de gris est inférieur ou égal à i . À un facteur de normalisation près, l'histogramme cumulé représente la fonction de répartition des niveaux de gris

$$hc(i) = \sum_{j=0}^i h(j) \quad (\text{II. 1})$$

On peut définir la transformation d'égalisation d'histogramme de la manière suivante. Soit G le niveau de gris d'un pixel de départ, le niveau de gris de l'image d'arrivée sera :

$$G' = \frac{255}{\text{Nombre depixels}} (G) \quad (\text{II.2})$$

L'outil Image : Calque $\Rightarrow$ Couleurs $\Rightarrow$ Auto $\Rightarrow$ Égaliser permet de réaliser cette opération. Il ne s'applique qu'à une image en niveaux de gris, telle que celle de la figure II.7 (a), dont l'histogramme apparaît figure II.7(b). Le résultat apparaît figure II.8 (a), et l'histogramme transformé figure II.8 (b). [35]

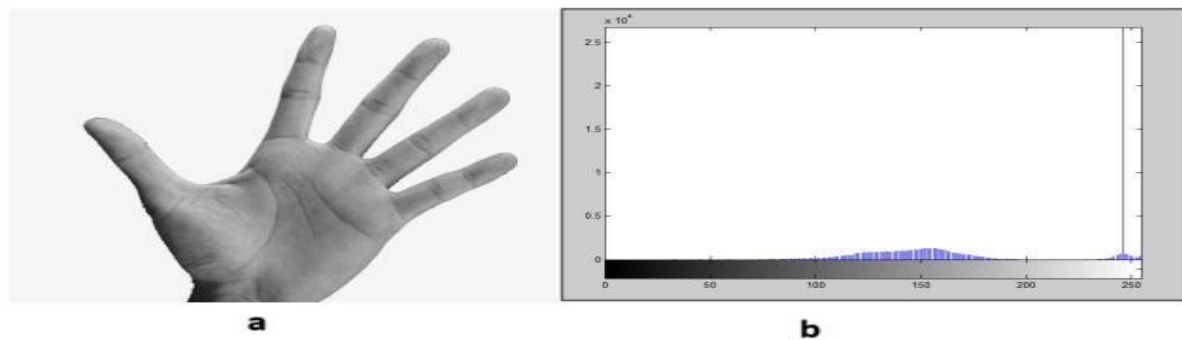


Figure II.7:(a) Image de départ, (b) Histogramme de l'image (a) [35]

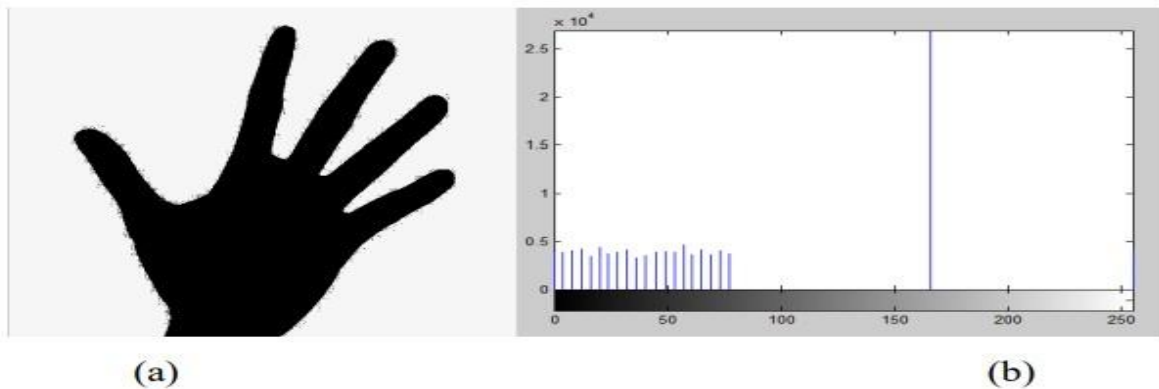


Figure II.8:(a) Image égalisation de l'histogramme, (b) Histogramme après égalisation[35]

II.6.3 Débruitage par Filtre gaussien

C'est également un filtre passe-bas. Une gaussienne à deux dimensions est donnée par l'expression suivante :

$$g(x, y) = \frac{1}{2\pi\sigma} e^{\left(\frac{x^2+y^2}{2\sigma^2}\right)} \quad (\text{II.3})$$

Le paramètre σ permet de régler facilement le degré de filtrage. Comparé au filtre moyenneur, ce filtre accorde une grande importance aux pixels proches du pixel central, et diminue cette importance au fur et à mesure que l'on s'éloigne de lui. Le masque suivant permet une pondération gaussienne : [36]

$$\frac{1}{16} * \begin{pmatrix} 1 & 2 & 1 \\ 2 & 4 & 2 \\ 1 & 2 & 1 \end{pmatrix} \quad (\text{II.4})$$

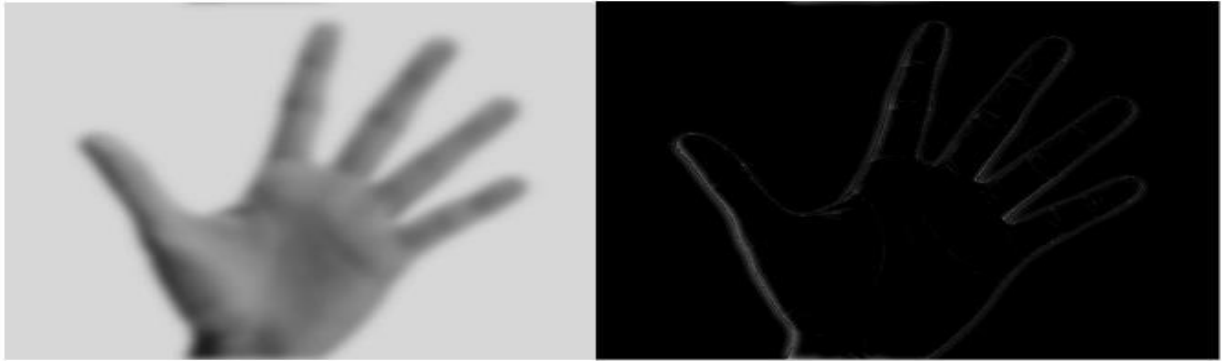


Figure II.9:(a) Flou gaussien de rayon 2, (b) Flou gaussien de rayon 3[36]

II.7 Méthodes d'Extraction des Caractéristiques

Cette étape représente le cœur du système de reconnaissance, on extrait de l'image les informations qui seront sauvegardées en mémoire pour être utilisées plus tard dans la phase de décision.

II.7.1 Extraction basé sur la texture

II.7.1.1 Méthode du Motif Binaire Local (LBP)

Les motifs binaires locaux ont initialement été proposés par Ojala en 1996 afin de caractériser les textures présentes dans des images en niveaux de gris. Ils consistent à attribuer à chaque pixel P de l'image à analyser, une valeur caractérisant le motif local autour de ce pixel. Ces valeurs sont calculées en comparant le niveau de gris du pixel central P aux valeurs des niveaux de gris des pixels voisins. [37]

Le concept du LBP est simple, il propose d'assigner un code binaire à un pixel en fonction de son voisinage. Ce code décrivant la texture locale d'une région est calculé par seuillage d'un voisinage avec le niveau de gris du pixel central. Afin de générer un motif binaire, tous les voisins prendront alors une valeur "1" si leur valeur est supérieure ou égale au pixel courant et "0"

autrement (**Figure II.10**). Les pixels de ce motif binaire sont alors multipliés par des poids et sommés afin d'obtenir un code LBP du pixel courant.

On obtient donc pour toute l'image, des pixels dont l'intensité se situe entre 0 et 255 comme dans une image à 8 bits ordinaire. Plutôt que de décrire l'image par la séquence des motifs LBP, on peut choisir comme descripteur de texture un histogramme de dimension 255.[37]

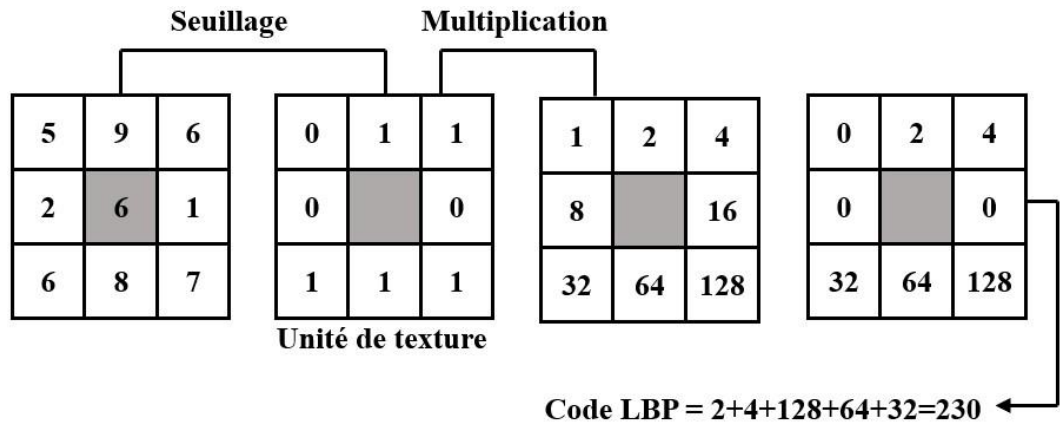


Figure II.10: Construction d'un motif binaire et calcul de code LBP [37]

Pour calculer un code LBP dans un voisinage de P pixels, dans un rayon R, on compte simplement les occurrences de niveaux de gris g_p plus grands ou égaux la valeur centrale

$$LBP(X_c, Y_c) = \sum_{n=0}^{p-1} U(g_i - g_c) * 2^n \quad (II.5)$$

Où x_c et y_c les coordonnées du pixel central, $u()$ est la fonction signe et où g_i et g_c sont respectivement les niveaux de gris d'un pixel voisin et du pixel central. (37)

$$U(x) = \begin{cases} 1 & \text{si } x \geq 0 \\ 0 & \text{autrement} \end{cases} \quad (II.6)$$

II.7.1.2 Descripteur de base LPQ (Local phase quantization) :

L'information de LPQ peut être extraite en utilisant la transformée discrète de Fourier à fenêtre à deux dimensions (2DWFT).

$$F_u(x) = \sum_{m \in N_x} h(m-x) f(m) e^{-j2\pi u^T m} = E_u^T f_x^2 \quad (II.7) \quad \text{Où } E_u, \text{ de taille } = 1$$

$\times M^2$, est un vecteur de base de 2DWFT avec la fréquence u , et f_x , taille = $M^2 \times N$, est un vecteur contenant les valeurs des pixels d'image dans N_x à chaque position x . La fonction fenêtre, $h(x)$ est une fonction rectangulaire.

La transformation est calculée à quatre valeurs de la fréquence, $u = [u_0, u_1, u_2, u_3]$ où $u_0 = [a, 0]$, $u_1 = [0, a]^T$, $u_2 = [a, a]^T$ et $u_3 = [a, -a]^T$. La valeur a est la plus haute fréquence scalaire

pour laquelle $H_{ui} > 0$. Ainsi, seuls quatre fonctions complexes comme un banc de filtres sont nécessaires pour produire huit images résultantes, composées de 4 images de la partie réelle et 4 images de la partie imaginaire de la transformée. Chaque pixel de l'image complexe résultant peut être codé en une valeur binaire représentée dans l'équation (II.4) en appliquant (the quadrant bit coding). [38]

$$B_{u_i}^{Re}(x) = \begin{cases} 1 & \text{si } F_{u_i}^{Re}(x) > 0 \\ 0 & \text{si } F_{u_i}^{Re}(x) \leq 0 \end{cases} \quad B_{u_i}^{Im}(x) = \begin{cases} 1 & \text{si } F_{u_i}^{Im}(x) > 0 \\ 0 & \text{si } F_{u_i}^{Im}(x) \leq 0 \end{cases} \quad (II.8)$$

Ce procédé de codage attribue deux bits pour chaque pixel pour représenter le quadrant dans lequel se trouve l'angle de phase.

En fait, il fournit également la quantification de la fonction de phase de Fourier. En général, LPQ est une chaîne binaire, présentée dans l'expression (II.5), obtenue pour chaque pixel par la concaténation des codes quadrant bits réelles et imaginaires des huit coefficients de Fourier de u_i

$$LPQ(x) = [B_{u_0}^{Re}(x), B_{u_0}^{Im}(x), \dots, B_{u_3}^{Re}(x), B_{u_3}^{Im}(x)] \quad (II.9)$$

La chaîne binaire est convertie en nombre décimal par l'expression (II.6) pour produire une étiquette de LPQ. La **Figure II.11** résume l'ensemble de ces étapes.

$$LPQ(x) = B_{u_0}^{Re}(x) + B_{u_0}^{Im}(x) \times 2^1 + \dots + B_{u_3}^{Re}(x) \times 2^{k-1} + B_{u_3}^{Im}(x) \times 2^k \quad (II.10)$$

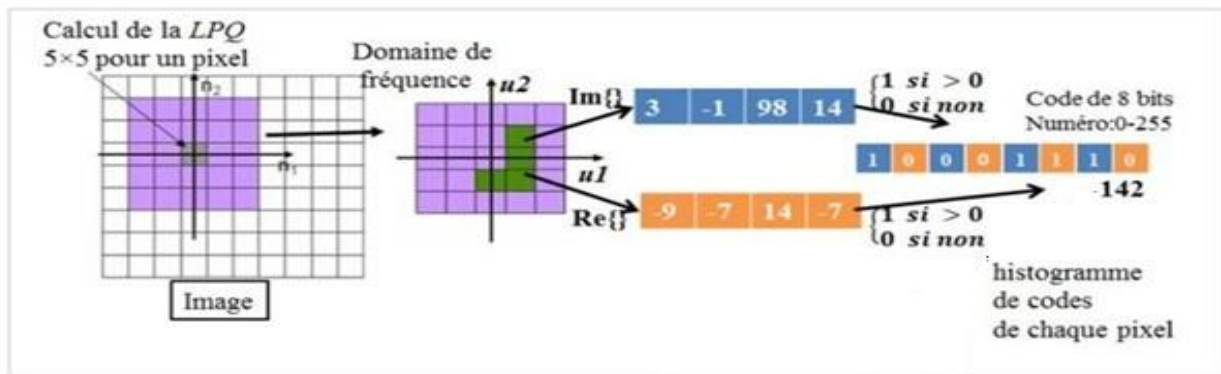


Figure II.11: Organigramme de l'ensemble des étapes nécessaires à la construction du descripteur LPQ

II.7.2 Extraction basé sur le filtrage

Filtre de Gabor

Pour les applications nécessitant une analyse par orientations, les fonctions de Gabor qui produisent une décomposition en ondelettes est très utilisée. De nombreuses applications entraînement d'images font appel à l'utilisation ces types des fonctions, comme par exemple l'analyse de textures ou objets par attributs fréquentiels. En effet, les lignes de l'empreinte sont caractérisées par leur fréquence locale et leur orientation. En utilisant des filtres de Log Gabor,

bien choisis, il est possible d'en extraire les caractéristiques biométriques. Cependant, lorsque ceux-ci sont correctement paramétrés, ils permettent de préserver les lignes et fournissent des informations sur l'orientation locale de la texture. Dans le domaine fréquentiel, la réponse (fGb) de filtre log-Gabor 1D se définit comme [35] :

$$fGb = \exp \left[\frac{-(\log(\frac{f}{f_0}))^2}{2(\log(\frac{\sigma}{f_0}))^2} \right] \quad (II.11)$$

Où f_0 est la fréquence centrale et dénote la variance. Pour l'application du filtre de Log-Gabor, il faut un choix empirique de paramètres de filtre (f_0 et σ). Ces paramètres empiriques sont très difficiles à déterminer et c'est l'un des inconvénients des approches basées sur ce filtre [35].

II.8 Classification

Cette étape consiste à modéliser les paramètres extraits d'une modalité d'un individu en se basant sur leurs caractéristiques communes. Un modèle est un ensemble d'information utiles, discriminantes et non redondantes qui caractérise un ou plusieurs individus ayant dissimilarités. Ces derniers seront regroupés dans la même classe, et ces classes varient selon le type de décision[39].

II.8.1 Machine à vecteurs de support (SVM)

Les machines à vecteurs de support ou séparateurs à vaste marge (en anglais Support Vector Machine, SVM) sont un ensemble de techniques d'apprentissage supervisé destinées à résoudre des problèmes de classification.

Les SVM sont une généralisation des classifiées linéaires. les SVM ont été développés dans les années 1990 à partir des considérations théoriques de Vladimir Vapnik sur le développement d'une théorie statistique de l'apprentissage : la Théorie de Vapnik Chervonenkis. Les SVM ont rapidement été adoptés pour leur capacité à travailler avec des données de grandes dimensions, le faible nombre d'hyper paramètres, leurs garanties théoriques, et leurs bons résultats en pratique.

Les SVM ont été appliqués à de très nombreux domaines (bio-informatique, recherche d'information, vision par ordinateur, finance...). Selon les données, la performance des machines à vecteurs de support est de même ordre, ou même supérieure, à celle d'un réseau de neurones ou d'un modèle de mixture gaussienne. Hyperplan qui est le lieu des points x satisfaisant $\langle w, x \rangle + b = 0$. En orientant l'hyperplan, la règle de décision correspond à observer de quel côté de l'hyperplan se trouve l'exemple x . On voit que le vecteur w définit la pente de l'hyperplan (w est

perpendiculaire à l'hyperplan). Le terme b quant à lui permet de translater l'hyperplan parallèlement à lui-même. ;

Décision $h(s)$. La classe de tous les hyperplans qui en découle sera notée H . [40]

II.8.2 Les k plus proches voisins :

L'algorithme KNN figure parmi les plus simples algorithmes d'apprentissage artificiel. Dans un contexte de classification d'une nouvelle observation x , l'idée fondatrice simple est de faire voter les plus proches voisins de cette observation. La classe de x est déterminée en fonction de la classe majoritaire parmi les k plus proches voisins de l'observation x . Donc la méthode du plus proche voisin est une méthode non paramétrique où une nouvelle observation est classée dans la classe d'appartenance de l'observation de l'échantillon d'apprentissage qui lui est la plus proche, au regard des covariables utilisées. La détermination de leur similarité est basée sur des mesures de distance. [41]

II.8.3 Les Réseaux de Neurones :

Les réseaux de neurones proposent une simulation du fonctionnement de la cellule nerveuse à l'aide d'un automate : le neurone formel. Les réseaux neuronaux sont constitués d'un ensemble de neurones (nœuds) connectés entre eux par des liens qui permettent de propager les signaux de neurone à neurone.

Grâce à leur capacité d'apprentissage, les réseaux neuronaux permettent de découvrir des relations complexes non-linéaires entre un grand nombre de variables, sans intervention externe. De ce fait, ils sont largement utilisés dans de nombreux problèmes de classification (ciblage marketing, reconnaissance de formes, traitement de signal,...) d'estimation (modélisation de phénomènes complexes,...) et prévision (bourse, ventes,...). Il existe un compromis entre clarté du modèle et pouvoir prédictif. Plus un modèle est simple, plus il sera facile à comprendre, mais moins il sera capable de prendre en compte des dépendances trop variées.

II.9 Domaines d'application

Le champ d'application de la biométrie des empreintes palmaires est très vaste. En effet, tous les domaines qui nécessitent de vérifier ou déterminer l'identité d'une personne sont concernés. D'où les applications de la biométrie peuvent être divisées en trois groupes principaux :

- **Applications commerciales** ; telles que l'ouverture d'un réseau informatique, la sécurité des données électroniques, l'e-commerce, l'accès Internet, les cartes de crédit, le contrôle d'accès physique, le téléphone cellulaire, la gestion des registres médicaux, l'étude à distance, etc.

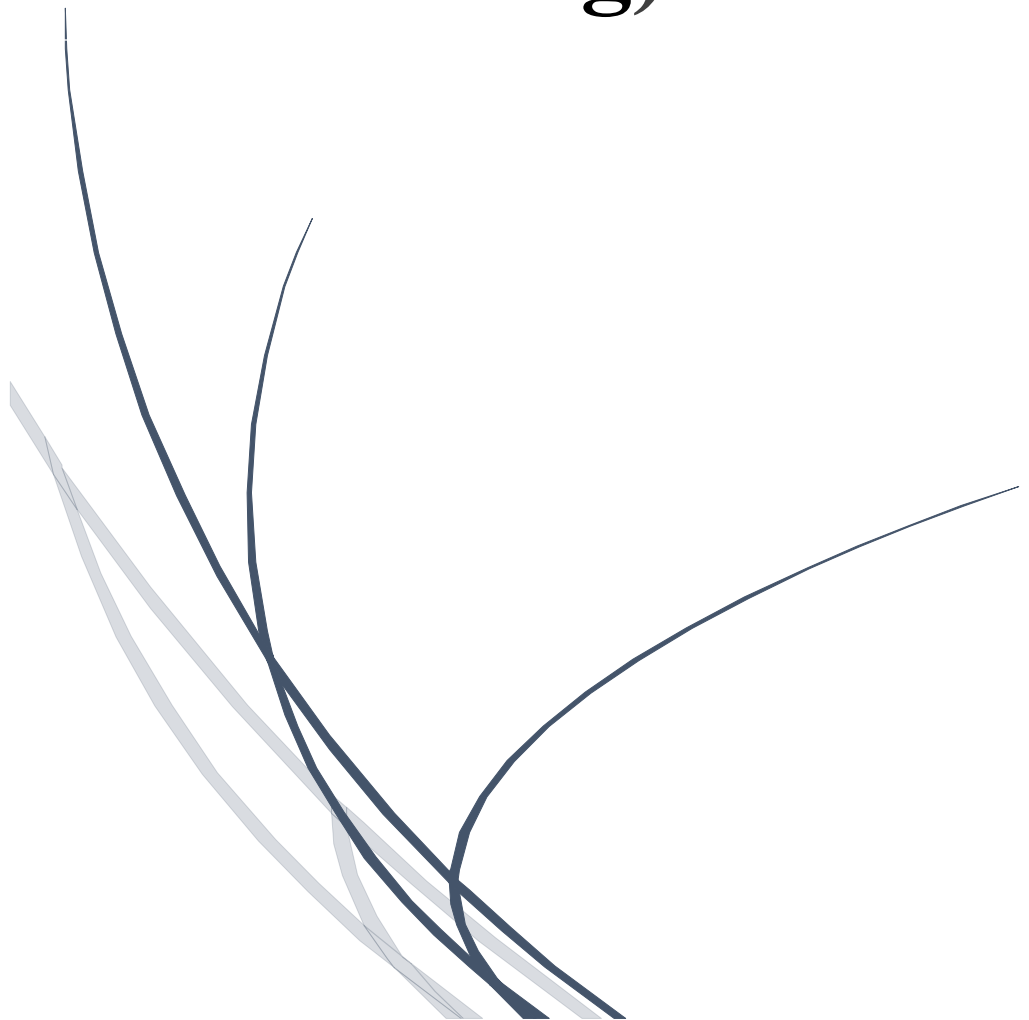
-
- **Applications gouvernementales** ; telles que la carte d'identité nationale, le permis de conduire, la sécurité sociale, le contrôle des frontières, le contrôle des passeports, etc.
 - **Applications légales** ; telles que l'identification de corps, la recherche criminelle, l'identification de terroriste, etc. [42]

II.10 Conclusion

Dans ce chapitre, nous avons présenté les différents recueils d'information des empreintes palmaires qui contiennent une grande quantité d'information, les domaines d'application, les caractéristiques, les composants et les critères d'évaluation de son système d'identification biométrique ont été également décrit. Nous avons détaillé les étapes de son acquisition, et une brève description des algorithmes utilisés pour l'extraction de leurs textures ou leurs primitives. Bien que tous ces algorithmes ont été réussis, ils exigent une étape de prétraitement et de classification. Cependant, les techniques basées sur d'apprentissage profond (deep Learning) regroupent toutes ces étapes de traitement. Dans le chapitre suivant, nous allons étudier la méthode d'apprentissage profond la plus largement répandues qui est basée sur les réseaux de neurones convolutionnels (CNN : Convolutionnal Neural Network). [42]

CHAPITRE III

L'apprentissage En Profondeur (Deep Learning)



III.1 Introduction

L'apprentissage automatique (Machine Learning- ML en anglais) est une technologie de l'intelligence artificielle permettant aux ordinateurs d'apprendre sans avoir été programmés explicitement à cet effet. Pour apprendre et se développer, les ordinateurs ont toutefois besoin de données à analyser et sur lesquelles s'entraîner. De fait, c'est la technologie qui permet d'exploiter pleinement le potentiel du Big Data. Les méthodes d'apprentissage automatiques conventionnelles ont été limitées dans leur capacité de s'occuper des données naturelles dans leur forme brute.

L'apprentissage profond est une fonction de l'IA qui imite le fonctionnement du cerveau humain. C'est l'une des formes de l'apprentissage automatique qui peut être utilisé pour aider à détecter la fraude ou le blanchiment d'argent, le traitement des données pour la détection d'objets, la reconnaissance de la parole, la traduction des langues et la prise de décisions.

Les applications basées sur l'apprentissage profond sont capable d'apprendre sans supervision humaine, en s'appuyant sur des données à la fois non structurées et non étiquetées. Nous présentons dans cette partie les types des deep learning existants ainsi que les architectures connaissent. [43]

III.2 Qu'est-ce que l'apprentissage en profondeur (Deep learning) ?

Avant d'explorer ce qu'est l'apprentissage en profondeur (Deep learning), nous devons réfléchir un instant à la façon dont un ordinateur peut résoudre les problèmes arithmétiques les plus complexes et les résoudre en quelques secondes, alors qu'il lui est très difficile (avant l'émergence d'algorithmes d'apprentissage automatique) de comprendre des choses qui semblent simples et intuitives pour les humains, telles que reconnaître l'image d'une personne ou faire la différence entre deux types de fruits. Afin de comprendre ce qu'est l'apprentissage en profondeur, nous devons d'abord comprendre sa relation avec l'apprentissage automatique et l'intelligence artificielle, et peut-être la meilleure façon de montrer cette relation est montrée dans la Figure III.1 [44]



Figure III.1: Les relation entre l'intelligence artificielle, l'apprentissage automatique et l'apprentissage profond

III.3 Intelligence artificielle

L'intelligence artificielle (IA) est un vaste domaine, où nous essayons d'imiter le comportement humain dans le but de rendre les machines si puissantes pour accomplir de nombreux types de tâche. Plus précisément, c'est la simulation des processus intelligence humaine par des machine, comme les ordinateurs, les systèmes embarqués, industriels, biomédicaux, financiers et bien d'autre encore. Ces processus comprennent l'apprentissage (l'acquisition d'information et de règles d'utilisation de ces informations), le raisonnement (l'utilisation des règles pour parvenir à des conclusion approximative ou définitives) et l'autocorrection. Les applications particulières de l'IA comprennent la reconnaissance vocale, la vision artificielle, le contrôle, la prédiction et divers autres domaines.

Ainsi, l'apprentissage automatique est un sous domaine de l'IA et l'apprentissage profond est un sous domaine de l'apprentissage automatique. [45]

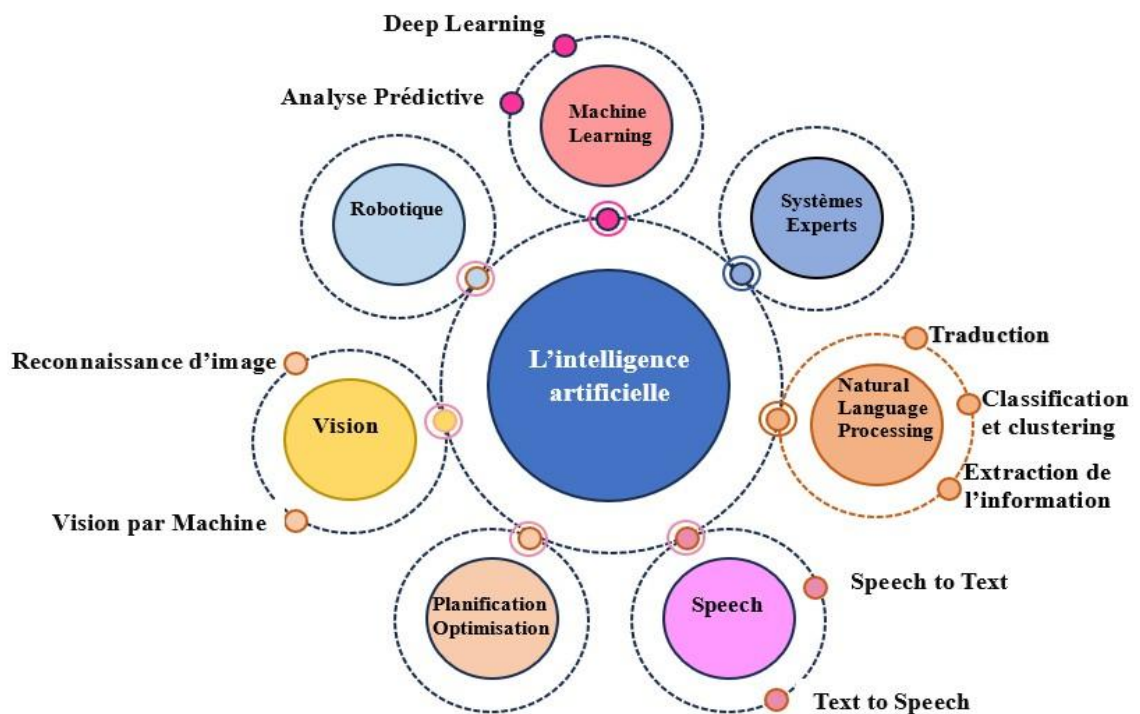


Figure III.2: Les types de l'intelligence artificielle

III.4 Apprentissage Automatique

L'apprentissage Automatique (Machine Learning) est un domaine de recherche en informatique qui traite des méthodes d'identification et de mise en œuvre de systèmes et algorithmes par lesquels un ordinateur peut apprendre, ce champ est généralement lié à l'intelligence artificielle. À l'origine, ce domaine était consacré au développement de l'intelligence artificielle, mais en raison des limites de la théorie de la technologie qui existaient, il est devenu plus logique de concentrer ces Algorithmes sur des tâches spécifiques. La plupart des algorithmes (ML) tels qu'ils existent aujourd'hui se concentrent sur optimisation des fonctions.

Par conséquent, l'utilisation des algorithmes Machine Learning devient fréquemment un processus répétitif d'essais et d'erreurs, dans lequel le choix de l'algorithme parmi les problèmes donne des résultats de performance différents. Cela est acceptable dans certains contextes, mais dans le cas de la modélisation du langage et de la vision par ordinateur, cela devient problématique. [46]

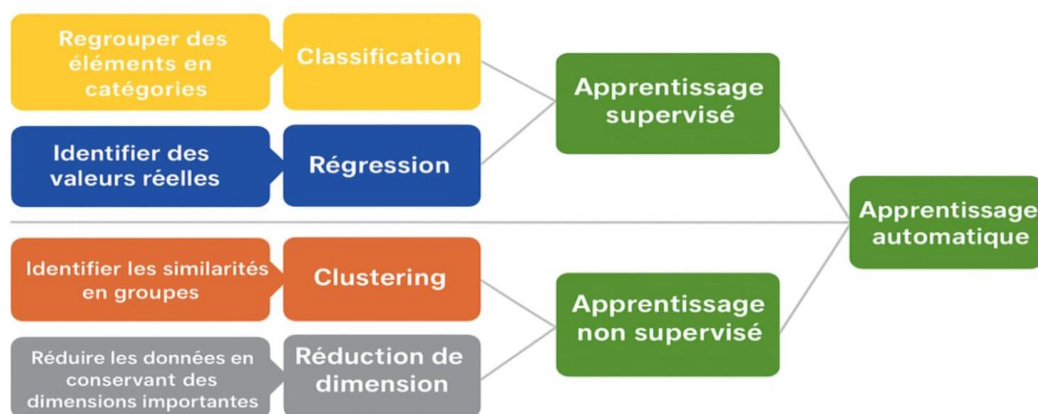


Figure III.3: Les type de l'apprentissage automatique

III.5 L'apprentissage profond

L'apprentissage profond est une forme d'apprentissage automatique qui peut utiliser des algorithmes supervisés ou non supervisés, ou les deux. Bien qu'il ne soit pas nécessairement nouveau, l'apprentissage en profondeur a récemment connu un regain de popularité en tant que moyen d'accélérer la résolution de certains types de problèmes informatiques difficiles, notamment dans les domaines de la vision par ordinateur et du traitement du langage naturel.

Nous trouvons l'idée de base simple : de la même manière qu'un enfant entend les sons en premier, les connecte aux mots puis construit des phrases, les algorithmes d'apprentissage profond (neurones) vont progressivement recueillir et comprendre des informations pour créer de nouvelles connaissances. L'apprentissage profond se caractérise par mieux gérer les concepts abstraits, ce qui distingue de l'apprentissage automatique. [47]

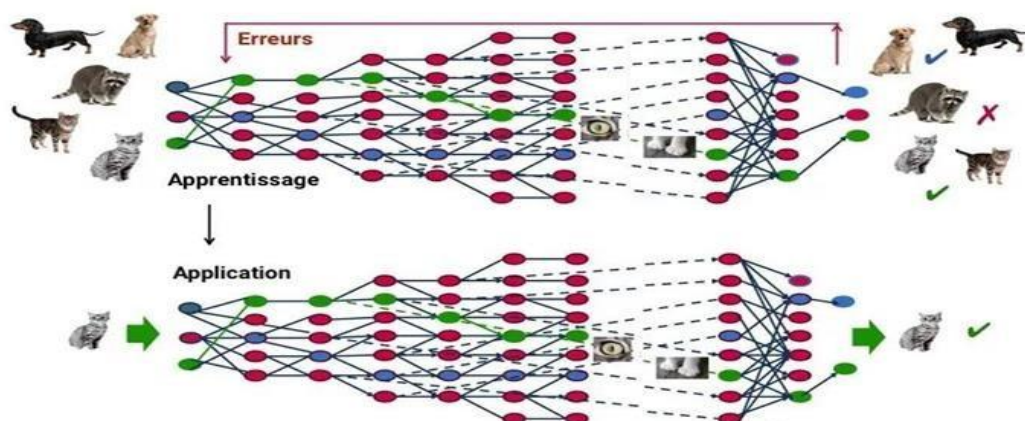


Figure III.4: Exemple d'apprentissage profond [48]

III.6 Différent type d'apprentissage

L'apprentissage est l'étape la plus importante dans le fonctionnement d'un réseau de neurones, c'est un processus dynamique et itératif permettant de modifier les paramètres d'un réseau en réaction avec le stimulus qu'il reçoit de son environnement. Il existe plusieurs types d'apprentissage, nous présentons les deux plus importants et les plus utilisés [40] :

TYPES D'APPRENTISSAGE AUTOMATIQUE			
Apprentissage Supervisé	Apprentissage Non supervisé	Apprentissage Semi-supervisé	Apprentissage Par renforcement
Entrainement avec des données étiquetées (entrée + sortie)	Entrainement avec des données sans étiquettes	Combinaison des données étiquetées et non étiquetées	Apprentissage par interaction avec l'environnement et retour d'information (récompense /pénalité)

Tableau III.5: Les type d'apprentissage automatique

III.6.1 L'apprentissage supervisé

Dans ce type, une information précise sur la sortie désirée est disponible. Le réseau apprend par présentation de pair d'entrée / sortie. Durant l'apprentissage, les valeurs de sorties désirées sont comparées à celles produites par le réseau, l'erreur résultante est alors utilisée pour l'ajustement des poids des connexions. La règle du delta en méthode rétropropagation est utilisée telle qu'elle dans les réseaux multicouches que nous détaillerons par la suite [49]. Avec l'apprentissage supervisé, la machine peut apprendre à faire une certaine tâche en étudiant des exemples de cette tâche. Par exemple, elle peut apprendre à reconnaître une photo de chien après qu'on lui ait montré des millions de photos de chiens. Ou bien, elle peut apprendre à traduire le français en chinois après avoir vu des millions d'exemples de traduction français-chinois.

D'une manière générale, la machine peut apprendre une relation $f x \rightarrow y$: qui relie x à y en ayant analysé des millions d'exemples d'associations.

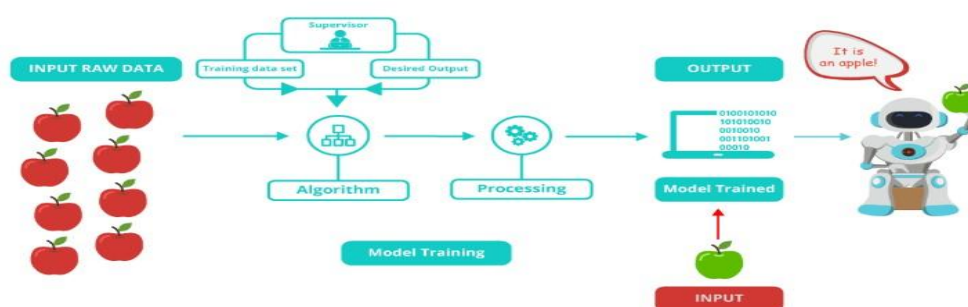


Figure III.5: Exemple de l'apprentissage supervisé

L'apprentissage supervisé fonctionne en 4 étapes :

1. Importer un Dataset (x, y) qui contient nos exemples.
2. Développer un Modèle aux paramètres aléatoires.
3. Développer une Fonction Coût qui mesure les erreurs entre le modèle et le Dataset.
4. Développer un Algorithme d'apprentissage pour trouver les paramètres du modèle qui minimisent la Fonction Coût.

a) Le Dataset

La première étape d'un algorithme de Supervised Learning consiste donc à importer un Dataset qui contient les exemples que la machine doit étudier.[50] Ce Dataset inclut toujours 2 types de variables :

- Une variable objective (target) y.
- Une ou plusieurs variables caractéristiques (features) x.

b) Le modèle : le cœur de programme

Le modèle est en quelque sorte le cœur de programme, c'est lui qui va effectuer la tâche que vous cherchez à accomplir, par exemple reconnaître un animal sur une photo ou prédire le prix d'un appartement.

c) La Fonction coût

Mesure de la performance Pour que la machine trouve le meilleur modèle, il faut déjà qu'elle puisse mesurer la performance d'un modèle donné [50].

d) L'Algorithme d'apprentissage

Parlons peu, parlons bien. Avoir un bon modèle, c'est avoir un modèle qui de petites erreurs.

Logique ?

Ainsi, en Supervised Learning, la machine cherche les paramètres de modèle qui minimisent la Fonction Coût. C'est ça qu'on appelle l'apprentissage. Cette phrase est très importante. C'est l'essentiel de ce qu'il faut comprendre en Machine Learning.

Pour trouver les paramètres qui minimisent la fonction Coût, il existe un paquet de stratégies. On pourrait par exemple développer un algorithme qui tente au hasard plusieurs combinaisons de paramètres, et qui retient la combinaison avec la Fonction Cout la plus faible. C'est un peu comme organiser un concours d'archers pour ne garder que le meilleur. Cette stratégie est cependant assez inefficace la plupart du temps. Une autre stratégie, très populaire en Machine Learning, est de considérer la Fonction Coût comme une fonction convexe, c'est-à-dire une fonction qui n'a qu'un seul minimum, et de chercher ce minimum avec un algorithme de minimisation appelé Gradient Descente [50].

III.6.2 Apprentissage semi-supervisé

L'apprentissage semi-supervisé (SSL) se situe entre l'apprentissage non supervisé et l'apprentissage supervisé et combine une petite quantité de données étiquetées avec une plus grande quantité de données non étiquetées pendant l'entraînement. Plus formellement, l'objectif de SSL est d'exploiter les données non étiquetées (Unlabeled Data) pour produire une fonction de prédiction $f(\theta)$ avec des paramètres entraînables (θ), qui est plus précis que ce qui aurait été obtenu en utilisant uniquement les données étiquetées (Labeled Data). Par exemple, Du peut fournir des informations supplémentaires sur la structure de la distribution des données ($p(x)$) pour mieux

estimer la frontière de décision entre les différentes classes pour améliorer le modèle de connaissance du classifieur [51]

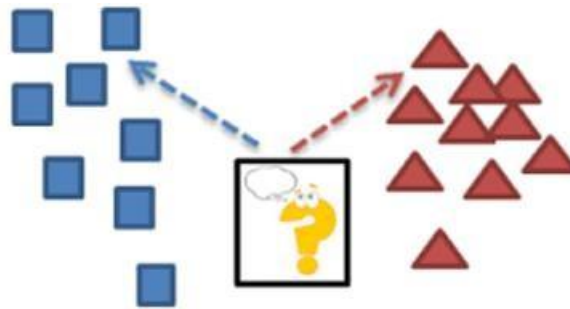


Figure III.6: Exemple d'apprentissage semi-supervisé [51]

III.6.3 Apprentissage non supervisé

L'apprentissage non supervisé est un type d'apprentissage où le réseau ne dispose d'aucune information sur la sortie désirée. Le réseau cherche à extraire des propriétés à partir des données d'entrée, en fonction de la règle d'apprentissage utilisée. Ce type d'apprentissage permet aux systèmes d'IA de trier des informations non classées en fonction de leurs similitudes et différences, sans catégorie spécifique. Il est souvent associé à des modèles d'apprentissage génératifs et peut être utilisé dans divers domaines tels que les chatbots, les véhicules autonomes, la reconnaissance faciale, les systèmes experts et les robots. Dans l'apprentissage non supervisé, les données ne sont ni étiquetées ni classées, et la sortie dépend des algorithmes programmés.

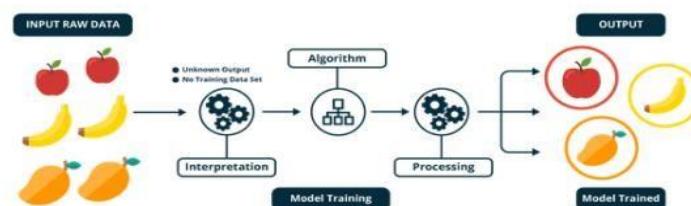


Figure III.7: Exemple d'apprentissage non supervisé [51]

III.6.4 L'apprentissage par renforcement

L'apprentissage par renforcement est un autre type de système d'apprentissage automatique. Un agent dans un "système d'IA" observera l'environnement, effectuera des actions données, puis recevoir des récompenses en retour. Avec ce type, l'agent doit apprendre par lui-même. Liens appelé une politique. Nous pouvons trouver ce type d'apprentissage dans des nombreuses applications robotiques qui apprennent comment faire des réactions dans son environnement. [51]

L'apprentissage par renforcement est le problème auquel est confronté un agent qui doit apprendre son comportement par tâtonnement c'est-à-dire interactions par essais et erreurs (trial-and-error) dans un environnement dynamique. L'agent est connecté à l'environnement par la perception et l'action. A chaque pas de l'interaction, l'agent reçoit en entrée une perception « p » de l'état courant, et génère une action « a » comme sortie qui impacte l'environnement. La valeur de cette transition est communiquée à l'agent par un signal de renforcement ou récompense « r », comme illustré par la Figure I-11. Le comportement de l'agent (ou politique) doit permettre de choisir les actions qui vont maximiser la somme des valeurs des récompenses sur un long terme. La politique fait correspondre les états aux actions. Formellement le modèle est représenté par le triplet : (S, A, R) avec S un ensemble fini d'états, A un ensemble fini d'actions et R une fonction de renforcement ; Où l'environnement est non déterministe : pour une action exécutée deux fois dans un même état, l'état résultant ou la récompense peuvent être différents.

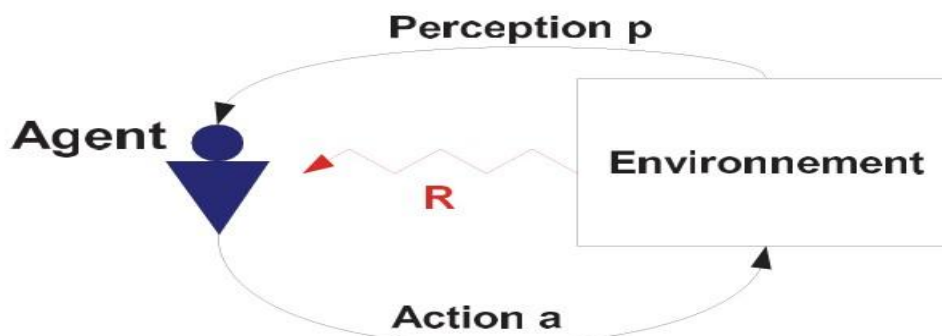


Figure III.8: Schéma d'apprentissage par renforcement [51]

III.7 Pourquoi d'apprentissage en profondeur ?

La réponse à cette question est vraiment essentielle, elle nous permet de comprendre l'engouement autour du Deep Learning. Le graphe ci-dessus l'explique clairement.



Figure III.9: Qualité des données [52]

Sa capacité à améliorer ses performances s'il y a une grande quantité de données (le passage à l'échelle) fait vraiment rêver vu que nous sommes à l'ère du big data et nous voulons exploiter au mieux nos données. [52]

III.8 L'apprentissage profond et d'apprentissage automatique

L'apprentissage profond est une forme spéciale d'apprentissage automatique. La progression de l'apprentissage automatique commence par (les fonctionnalités associées) qui sont extraites manuellement des images, puis les fonctionnalités sont utilisées pour créer un modèle qui classe les objets dans l'image. Avec le deep Learning, les fonctionnalités associées sont automatiquement extraites des images. [53]

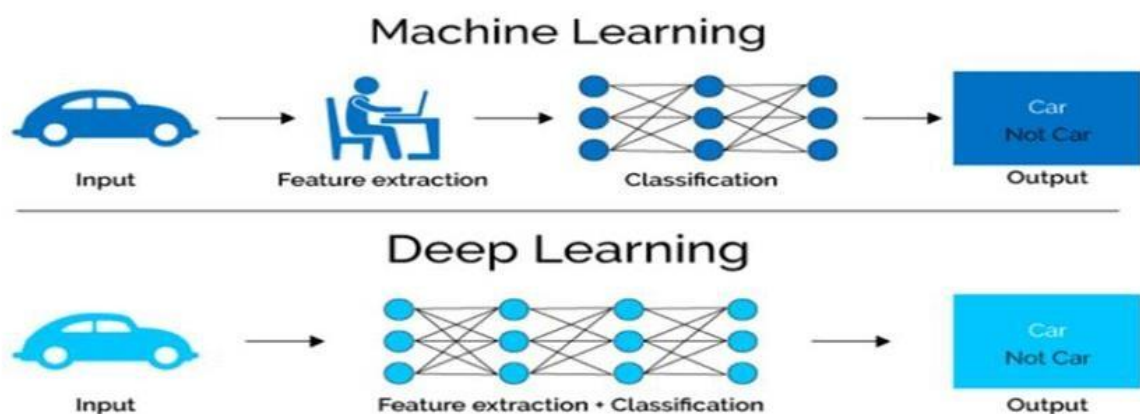


Figure III.10: Une comparaison des étapes du fonctionnement des algorithmes DL et ML

On an résumer la comparaison entre les deux types d'apprentissage dans le tableau suivant :

	Apprentissage profond	Apprentissage automatique
Exigences en matière de données	Nécessite de grandes quantités de données	Peut s'entraîner sur moins de données
Précision	Fournit une grande précision	Donne moins de précision
Temps d'exécution	Prend plus de temps pour s'entraîner	Prend moins de temps pour s'entraîner
Dépendance matérielle	Nécessite un GPU pour s'entraîner correctement	Trains sur CPU
Réglage de hyper-paramètres	Peut être réglé de différentes manières	Capacités de réglage limitées

Tableau III.2: Comparaison entre l'apprentissage profond et l'apprentissage automatique

III.9 Rétropropagation du gradient

La phase d'apprentissage des MLP consiste à adapter les poids des connexions en fonction des erreurs de prédiction constatées à chaque prédiction d'une nouvelle instance. La rétropropagation du gradient (backpropagation) est la méthode la plus utilisée pour l'adaptation de ces poids. Cet algorithme permet de déterminer le gradient de l'erreur pour chaque neurone du réseau en partant de la dernière couche et en arrivant jusqu'à la première couche cachée. L'objectif de la rétropropagation du gradient est d'ajuster les poids des connexions dans le but de minimiser l'erreur quadratique

$$E = \frac{1}{2} \sum_{i=1}^N (ref_i - hyp)^2 \quad (II\ 12)$$

qui représente l'écart entre la sortie attendue (la référence) et la sortie produite par le réseau (hypothèse) correspondant à un vecteur d'entrée donné. N représente la taille des vecteurs en sortie. La rétropropagation du gradient est considérée comme un problème d'optimisation pour lequel nous pouvons utiliser la méthode de descente de gradient. L'exemple le plus simple d'algorithme de descente de gradient est de modifier chacun des poids w dans la direction opposée au gradient du coût par rapport à ce poids :

$$W_n = W_{n-1} - \lambda \frac{\partial C}{\partial w} (W_{n-1}) \quad (11)$$

où λ est le taux d'apprentissage permettant d'ajuster la modification des poids. Étant conçu pour les réseaux non bouclés, l'algorithme de rétropropagation du gradient n'est pas suffisant pour la prise en compte des liaisons temporelles. Une solution à ce problème consiste à considérer une représentation dépliée « hiérarchisée » des RNN. Une des solutions les plus efficaces permettant de pallier ce problème de calcul du gradient se manifeste dans une extension du concept des RNN, à savoir, l'architecture Long Short-Term Memory (LSTM). [54]

III.10 Les réseaux de neurones

Les réseaux de neurones est un système dont la conception est à l'origine schématiquement inspirée du fonctionnement des neurones biologiques, et qui par la suite s'est rapproché des méthodes statistiques.

III.10.1 Réseaux de neurones biologiques

Un neurone (cellule nerveuse) est une cellule qui transporte des impulsions électriques [55], où elles sont connectées les unes aux autres. Ils ne se touchent pas et forment à la place de minuscules espaces appelés synapses. ces lacunes peuvent être des synapses chimiques ou des synapses électriques et transmettre le signal d'un neurone à l'autre. Les composants les plus importants du neurone (Figure III.12) :

- 1) **Dendrite** : Elle reçoit les signaux des autres neurones.
- 2) **Soma (corps cellulaire)**: Il additionne tous les signaux entrants pour générer une entrée.
- 3) **Axone** : Lorsque la somme atteint une valeur seuil, le neurone se déclenche et le signal descend l'axone vers les autres neurones.
- 4) **Synapses** : Le point d'interconnexion d'un neurone avec d'autres neurones. La quantité de signal transmis dépend de la force (poids synaptiques) des connexions.

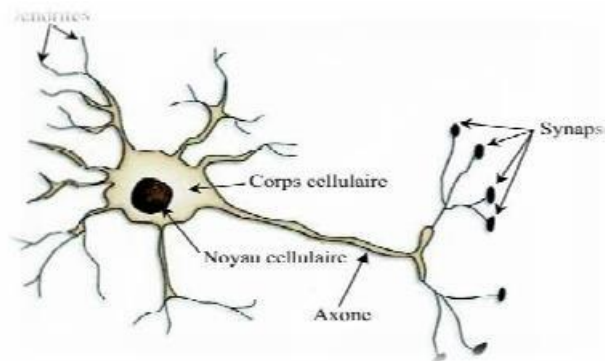


Figure III.11: Représentation schématique d'un neurone biologique

III.10.2 Réseau neuronal artificiel

Un réseau neuronal artificiel (ou réseau neuronal profond) s'inspire de la structure du cerveau humain (imitent la façon dont les neurones biologiques s'envoient mutuellement des signaux). Alors les Réseaux de neurones artificiels (ANN) sont des modèles d'apprentissage automatique composés de deux ou plusieurs couches qui sont constituées d'une couche d'entrée, d'une ou plusieurs couches cachées et d'une couche de sortie neurones, ces derniers interagissent avec un flux de données d'apprentissage pour apprendre à effectuer des tâches. Les réseaux de neurones artificiels sont appliqués à de nombreux domaines : la reconnaissance d'images et la vision par ordinateur, la reconnaissance vocale et, plus généralement, le traitement du langage naturel (NLP).

La figure suivante représente la relation entre des données d'un espace X et un espace de sortie Y dans neurone formel [56].

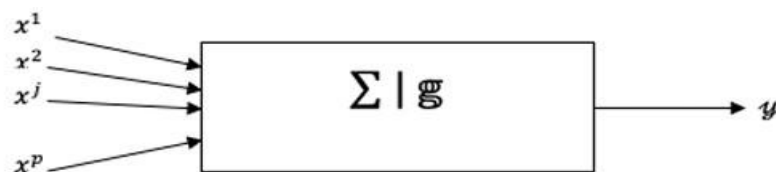


Figure III.12: Architecture d'un neurone formel [57]

Chaque neurone artificiel se connecte à un autre neurone et possède un poids et un seuil associés. Si la sortie d'un neurone dépasse le seuil défini, le neurone est activé et transmet des données à la couche suivante du réseau. En revanche, si la sortie est inférieure au seuil, aucune donnée n'est transmise à la couche suivante du réseau [58]

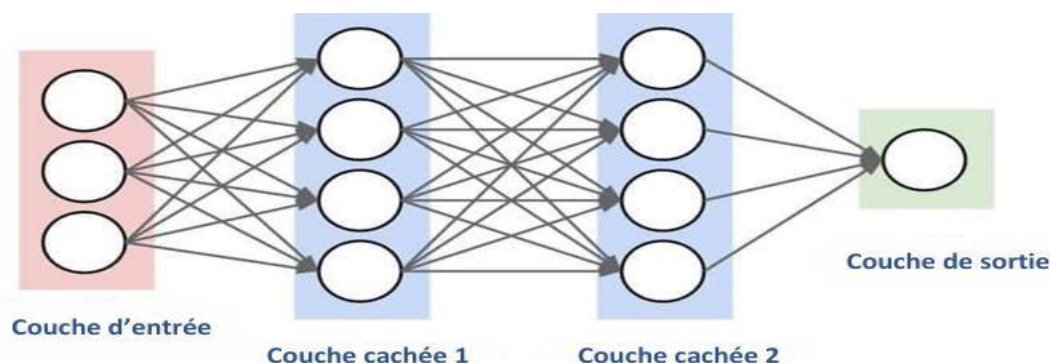


Figure III.13: Architecture d'un réseau de neurone [59]

III.11 La fonction d'activation

La fonction d'activation calcule l'état du neurone, cette valeur sera transmise aux neurones aval. Il existe des nombreuses formes possibles de la fonction d'activation, les plus courantes sont présentées dans le tableau (3.1). [60]

Fonction d'activation	Modèle mathématique	Graphe
Fonction signe	$\varphi(x) = \begin{cases} 1 & x \geq 0 \\ -1 & x \leq 0 \end{cases}$	
Fonction seuil	$\varphi(x) = \begin{cases} 1 & x \geq s \\ 0 & x \leq s \end{cases}$	
Fonction linéaire	$\varphi(x) = x$	
Fonction sigmoïde	$\varphi(x) = \frac{1}{1 + e^{-x}}$	
Fonction gaussienne	$\varphi(x) = \frac{1}{\sqrt{2\pi}\sigma'} e^{-\frac{(x-c)^2}{2\sigma'^2}}$	
Fonction saturation	$\varphi(x) = \begin{cases} 1 & x \geq +1 \\ x \text{ si } -1 \leq x \leq +1 \\ -1 & x < -1 \end{cases}$	

Tableau III.3: Les fonctions d'activations

III.12 Les différents types de réseaux de neurones

III.12.1 Réseaux de neurones récurrents RNN

Recurrent Neural Network (RNN) est considéré comme un réseau avec mémoire. Un réseau neuronal récurrent est un type de réseau neuronal qui contient des boucles utilisées pour stocker des informations dans le réseau. Les RNN utilisent les sorties précédentes comme entrées supplémentaires, ce qui leur permet de prédire très précisément ce qui va se passer ensuite [61].

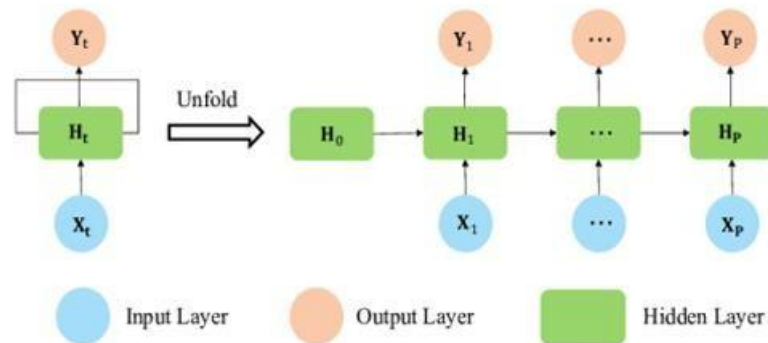


Figure III.14: Architecture de RNN [62]

- X : couche d'entrée
- H : couche cachée
- Y : couche de sortie

À tout moment, d'entrée actuelle est la combinaison des entrées à $x(t)$ et $x(t-1)$.

III.12.2 Réseaux de neurones profonds DNN

Les réseaux de neurones profonds sont également connus sous le nom de perceptrons multi couches. Un réseau de neurones est dit profond s'il contient au moins une couche cachée ; il peut contenir des millions de neurones organisés en couches où l'information ne circule que de la couche d'entrée vers la couche de sortie. Ils sont utilisés en apprentissage profond pour concevoir des mécanismes d'apprentissage supervisés et non supervisés.

III.12.3 Le LSTM (LongShort-TermMemory)

Le LSTM a été proposé en 1997 par Sepp Hochreiter et Jürgen Schmidhuber ils ont introduit des cellules mémoires, une forme de mémoire intermédiaire permettant de stocker des informations importantes sur une période de temps plus longue que les RNN existants [64].

Les LSTM sont des blocs de construction spéciaux des réseaux de neurones récurrents (RNN) doté d'une mémoire à court terme à long terme. C'est une évolution de RNN qui résout le

problème du gradient disparaissant c'est-à-dire que le gradient de poids diminue progressivement pendant l'entraînement, de sorte que le réseau ne stocke plus d'informations utiles. C'est-à-dire que le gradient de poids diminue progressivement pendant l'entraînement, de sorte que le réseau ne stocke plus d'informations utiles.

Les cellules LSTM ont trois types de portes :

- La porte d'entrée régule le flux d'informations.
- La porte de mémorisation et d'oubli veille à ce que les informations non importantes soient oubliées
- La porte de sortie détermine quelles informations sont transmises à l'étape suivante.

III.12.4 Réseaux Neurones Convolutifs CNN

Le réseau de neurones convolutifs est l'un des modèles de classification d'images connus Le plus puissant de l'apprentissage en profondeur. Il permet d'attribuer automatiquement à chaque image donnée en entrée sous forme de matrice de pixels un label correspondant à sa classe [63].

L'architecture du CNN possède deux parties :

— **Partie convolutive** : Son but est d'extraire des informations spécifiques Appelez chaque image une "fonctionnalité" en appliquant une opération de filtrage convolution. Ce processus de filtrage produit de nouvelles images, appelées cartes de contraction, en répétant l'opération, on obtient les caractéristiques de rétrécissement de l'image par rapport à sa taille initiale. Ces va leurs de la dernière caractéristique sont concaténées dans Un vecteur nommé codes *CNN*. Une image se représente en 3 dimensions :

1. Deux dimensions pour une image en niveaux de gris qui correspondent à la largeur et à la hauteur de l'image.
2. Une troisième dimension, de profondeur 3 pour représenter les couleurs fondamentales (Rouge, Vert, Bleu).

— **Partie classification** : Le vecteur obtenu dans la section précédente est L'entrée de ce bloc consiste en un réseau multi-couche. Les valeurs dans le vecteur Transformer l'entrée à l'aide de plusieurs combinaisons et fonctions linéaires Activez pour renvoyer un nouveau vecteur en sortie.

Le but du CNN est de résoudre le problème des réseaux de neurones traditionnels qui souffrent du problème de dimension de ses couches.

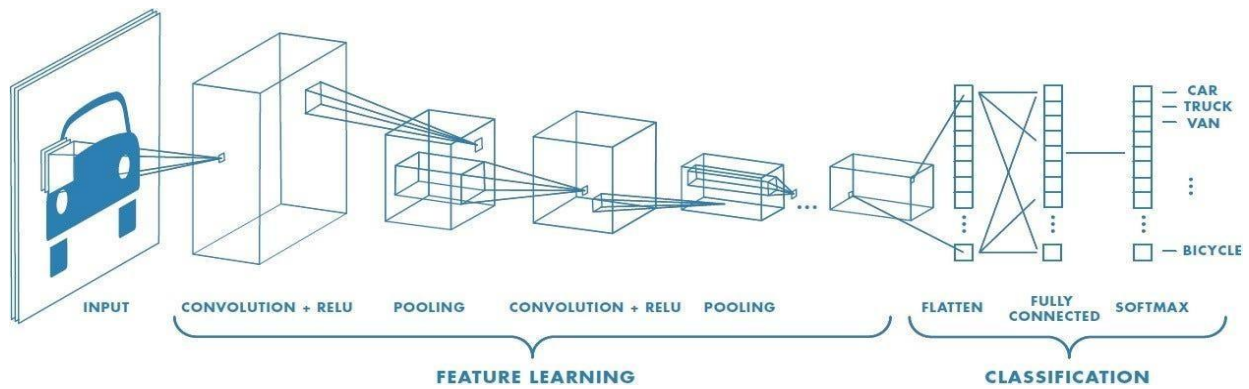


Figure III.15: Représentation synoptique d'un CNN

III.13 La construction des réseaux de neurones

La construction détaillée d'un réseau de neurones artificiels peut varier en fonction de la tâche à accomplir et de l'architecture du réseau choisi. Cependant, voici les étapes générales à suivre :

III.13.1 Collecte et préparation des données

Pour commencer l'apprentissage d'un réseau de neurones, il est nécessaire de collecter les données requises. Ces données doivent être préparées pour accomplir la tâche souhaitée et doivent être divisées en trois ensembles : l'ensemble d'apprentissage, l'ensemble de validation et l'ensemble de test. L'ensemble d'apprentissage est utilisé pour entraîner le réseau, l'ensemble de validation est utilisé pour ajuster les hyperparamètres du réseau, tandis que l'ensemble de test est utilisé pour évaluer les performances du réseau. [65]

III.13.2 Choix de l'architecture

Il existe différents types d'architectures de réseaux de neurones, chacune ayant ses propres avantages et inconvénients. Il est crucial de choisir l'architecture la plus appropriée pour la tâche à réaliser. Les architectures populaires comprennent les réseaux de neurones à couches denses, les réseaux de neurones convolutifs, les réseaux de neurones récurrents et les réseaux de neurones avec attention.

III.13.3 Définition des couches et des paramètres

Une fois l'architecture choisie, il faut définir les différentes couches du réseau, qui sont des unités de traitement de l'information. Les paramètres de chaque couche sont également définis, comme le nombre de neurones et la fonction d'activation.

III.13.4 Définition de la fonction de coût et de l'optimiseur

Pour évaluer la différence entre la sortie prédite par le réseau et la sortie réelle, une fonction de coût est appliquée. Pour réduire cette différence, l'optimiseur est utilisé pour ajuster les paramètres du réseau. L'objectif de l'optimiseur est de minimiser la fonction de coût.

III.13.5 Entraînement du réseau

L'entraînement du réseau de neurones débute avec l'initialisation aléatoire de ses poids et biais. Les données d'apprentissage sont ensuite utilisées pour ajuster les paramètres du réseau en se basant sur la fonction de coût. Le processus d'entraînement continue jusqu'à ce que la fonction de coût soit réduite au minimum, ou lorsque le réseau ne parvient plus à s'améliorer davantage.

III.13.6 Validation et ajustement des hyperparamètres

Une fois l'entraînement terminé, le réseau est évalué sur les données de validation pour déterminer s'il est surajusté (overfitting) ou sous-ajusté (underfitting). Les hyperparamètres du réseau peuvent alors être ajustés pour améliorer les performances.[65]

III.13.7 Évaluation des performances :

Une fois que les hyperparamètres ont été ajustés, le réseau est évalué sur les données de test pour déterminer ses performances sur des données non vues auparavant.

III.14 Algorithmes l'apprentissage profond

Les algorithmes d'apprentissage profond consistent en un ensemble de processus non linéaires avec des structures profondes (d'où leur nom d'apprentissage profond). Ce qui distingue l'architecture de ces algorithmes, c'est qu'ils s'inspirent des réseaux de neurones du cerveau humain, c'est pourquoi les architectures d'apprentissage en profondeur sont appelées réseaux de neurones artificiels.

III.14.1 Régression linéaire

La régression linéaire est un algorithme d'apprentissage supervisé dont le but est d'exprimer la variable y d'un ensemble d'entrée x et de l'exprimer dans une fonction linéaire où :

$$y=h(x)=w_i x+b \quad (2.11)$$

Où w_i est la pente et b est la constante de l'équation.

Afin de trouver les inconnues de l'équation exprimant N paires ou triplets... les entrées qui composent un jeu de données, il existe plusieurs techniques dont les plus connues sont les moindres carrés et la méthode des régressions.

Lors de la recherche des valeurs $\mathbf{b}$ et $\mathbf{w}_i$ par la méthode des moindres carrés, par exemple, ces coefficients permettent de réduire la fonction de coût $c = \sum_{i=1}^n (h(x_i) - y_i)^2$, car cette fonction correspond à la somme des écarts au carré entre les prédictions et les valeurs attendues, et nous appeller ces écarts les valeurs résiduelles [66].

III.14.2 Régression logistique

La régression logistique est essentiellement un algorithme de classification supervisée. Dans les problèmes de classification, la variable cible (ou de sortie) y ne peut prendre que des valeurs discrètes pour un ensemble donné de caractéristiques (ou d'entrée) X .

Le modèle construit un modèle de régression pour prédire la probabilité qu'une entrée de données donnée appartienne à la catégorie numérotée "1". Tout comme la régression linéaire suppose que les données obéissent à une fonction linéaire, la régression utilise la fonction sigmoïde pour modéliser les données.

La régression logistique ne devient une technique de classification que lorsqu'un seuil de décision est introduit dans l'image. La définition du seuil est un aspect très important de la régression logistique et dépend du problème de classification lui-même.

La décision de seuil est principalement influencée par les valeurs de précision et de rappel. Idéalement., nous aimerions que la précision et le rappel soient égaux à 1, mais cela arrive rarement.

III.15 Principes fondamentaux de l'apprentissage profond

L'objectif principal des réseaux de neurones profonds est d'approximer des fonctions complexes grâce à une composition d'opérations simples et prédéfinies d'unités (ou neurones). Les opérations effectuées sont généralement définies par une combinaison pondérée d'un groupe spécifique d'unités cachées avec une fonction d'activation non linéaire, en fonction de la structure du modèle. Ces opérations ainsi que les unités de sortie sont appelées « couches ». L'architecture du réseau neuronal ressemble au processus de perception dans un cerveau, où un ensemble spécifique d'unités est activé compte tenu de l'environnement actuel, influençant la sortie du modèle de réseau neuronal [67]

III.16 Exemples de domaines d'application

Le deep learning peut être appliqué en différent et nombreux domaine tel que :

- ✓ Contrôle vocal d'appareils tels que téléphones, téléviseurs et amplificateurs.
- ✓ Effectuer des tâches de classification directement à partir de la voix, de l'image ou du texte.
- ✓ Conduite conventionnelle, les chercheurs automobiles utilisent l'apprentissage en profondeur pour détecter des éléments tels que les panneaux de stationnement et les feux de circulation. En plus de cela, l'apprentissage en profondeur est utilisé pour détecter les piétons, ce qui contribue à réduire les accidents.
- ✓ L'espace et la défense sont utilisés pour identifier les objets des satellites qui définissent les zones d'importance et pour définir les zones sûres et dangereuses pour les forces militaires.
- ✓ Recherche médicale où les chercheurs sur le cancer utilisent l'apprentissage en profondeur pour dépister les cellules cancéreuses.
- ✓ L'industrie améliore la sécurité des travailleurs autour des machines lourdes en détectant les personnes ou les choses à une distance non sécurisée des machines.
- ✓ L'électronique dans la traduction automatique de l'audition et de la parole.
- ✓ Jeux et compétitions tels que trouver des solutions au jeu avec les mouvements les plus courts.

Donc et en bref , le deep learning est utilisé dans : la reconnaissance d'image, la traduction automatique, la voiture autonome, le diagnostic médical, les recommandations personnalisées, la modération automatique des réseaux sociaux, la prédiction financière et trading automatisé, la détection de malwares ou de fraudes, les chatbots (agents conversationnels), l'exploration spatiale, et les robots intelligents. [47]

III.17 Quelques architectures connus de réseaux de neurones

Il existe différentes architectures de réseaux de neurones convolutifs qui sont utilisées en fonction du contexte. Ces architectures ont souvent fait leurs preuves dans des défis d'apprentissage en profondeur, ce qui les a rendues populaires et largement utilisées. Voici quelques-unes de ces architectures :

III.17.1 AlexNet (2012)

Le pionnier qui a popularisé les CNN. AlexNet était relativement peu profond (8 couches) et utilisait de grands filtres (11x11, 5x5 au début). Il a démontré l'efficacité des GPU pour l'entraînement de grands réseaux et a eu recours au dropout pour régulariser. GoogleNet est beaucoup plus profond et plus efficace

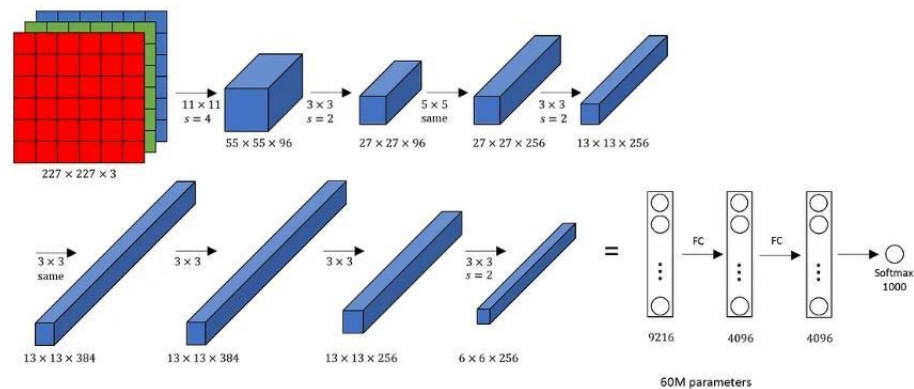


Figure III.16 :architecture AlexNet

III.17.2 VGG (2014)

VGG a montré que la profondeur est cruciale et a simplifié l'architecture en utilisant principalement des filtres 3x3 sur toutes les couches convolutives. Bien que très performant, VGG est notoirement coûteux en termes de calculs et de mémoire en raison de son grand nombre de paramètres (notamment dans ses couches entièrement connectées finales). GoogleNet surpasse VGG en efficacité tout en étant comparable, voire supérieur, en précision.

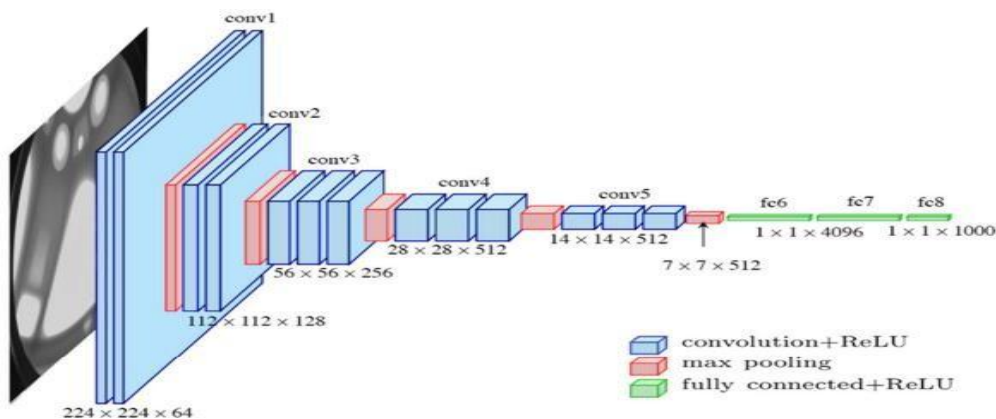


Figure III.17: Architecture VGG

III.17.3 ResNet (2015)

ResNet a introduit les connexions résiduelles pour résoudre le problème du gradient évanescent et permettre la construction de réseaux extrêmement profonds (jusqu'à 1000 couches). Alors que GoogleNet aborde le problème du gradient avec des classificateurs auxiliaires et une conception modulaire, ResNet offre une solution plus directe pour la construction de réseaux très profonds. Les architectures Inception-ResNet ont par la suite combiné les avantages des deux approches.

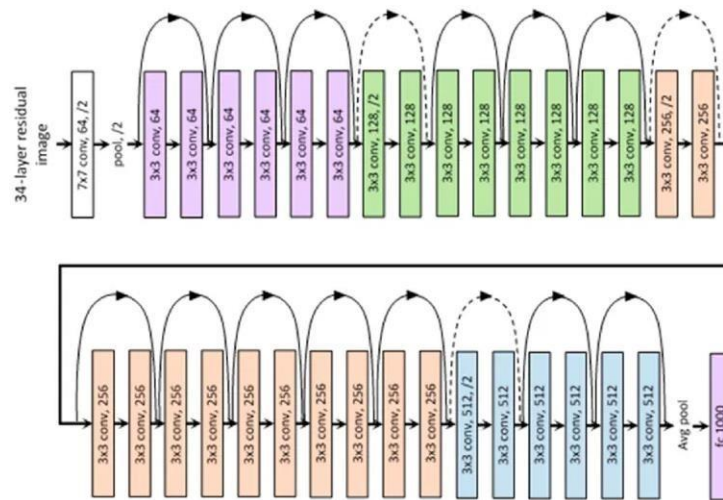


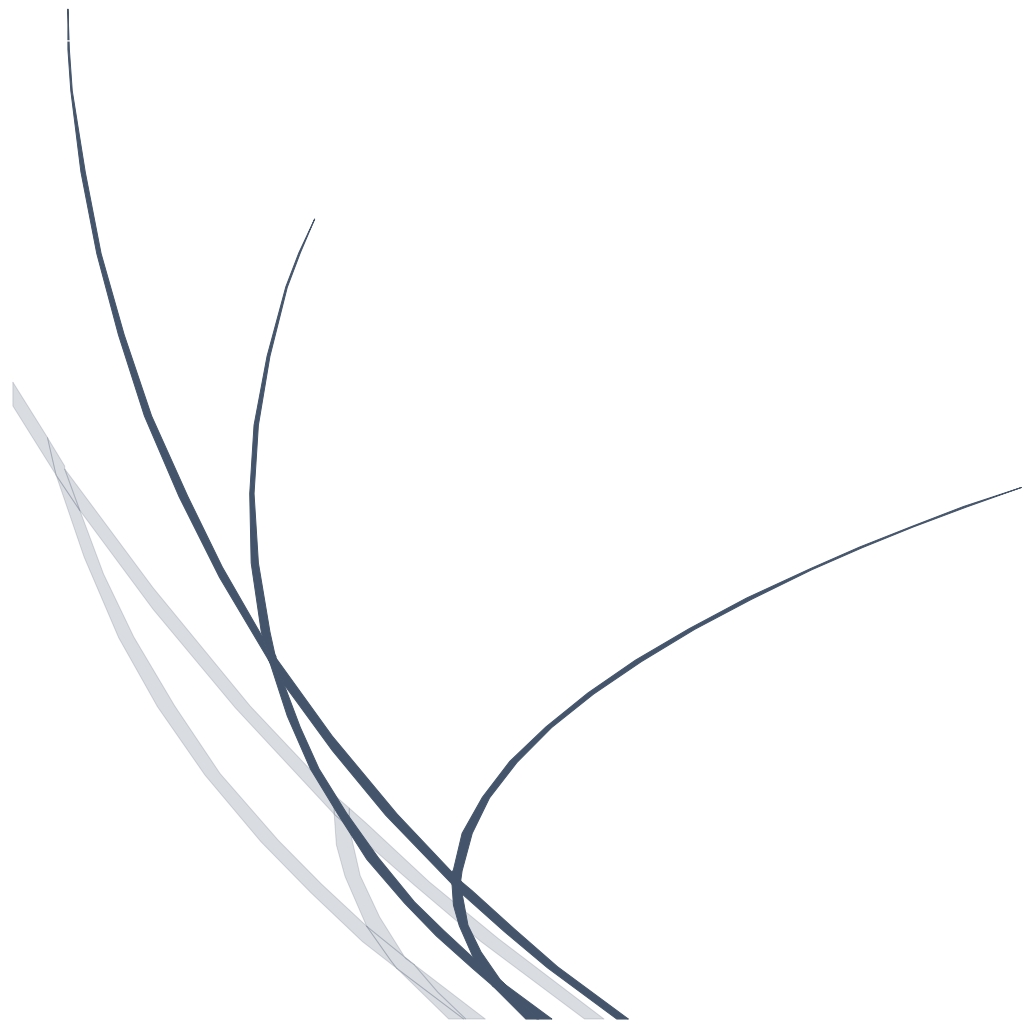
Figure III.18: Architecture *ResNet*

III. 18 Conclusion :

Le deep learning, ou apprentissage profond, a profondément transformé le domaine de l'intelligence artificielle en permettant aux machines d'apprendre automatiquement des modèles complexes et de prendre des décisions avec une intervention humaine minimale. Ses performances remarquables dans des domaines variés tels que la reconnaissance d'images, la compréhension du langage naturel, ou encore les véhicules autonomes, illustrent son potentiel révolutionnaire. Malgré certains défis — comme la nécessité de grandes quantités de données, des ressources de calcul importantes, et un manque de transparence dans les décisions — la recherche continue de progresser pour surmonter ces limites. À mesure qu'il évolue, le deep learning est appelé à jouer un rôle central dans les avancées technologiques de demain. [69]

CHAPITRE IV

METHODOLOGIE EXPREMENTATIONS ET RESULTATS



IV.1 Introduction :

Dans cette étude, nous explorons l'efficacité des descripteurs locaux après application de méthodes de prétraitement visant à atténuer les effets de la lumière. Nous nous concentrons ensuite sur la classification pour la reconnaissance des empreintes palmaires, en exploitant deux bases de données reconnues : MS-CASIA et MS-PolyU. Pour l'extraction des caractéristiques, nous utilisons un réseau de neurones pour apprendre des représentations discriminantes des images. La réduction de dimensionnalité est ensuite étudiée à l'aide de l'Analyse Discriminante Linéaire (LDA) et de l'Analyse en Composantes Principales (PCA), afin de minimiser la redondance et de renforcer la séparabilité des classes. Pour la classification finale, l'algorithme des k plus proches voisins (KNN) est évalué.

Nous avons sélectionné ces bases de données, ces techniques d'apprentissage profond et ces méthodes de réduction de dimensionnalité pour leurs spécificités et leur pertinence dans le domaine de la vision par ordinateur. Notre objectif est de concevoir des systèmes de reconnaissance d'empreintes palmaires robustes et précis, capables de fonctionner efficacement dans des environnements réels et diversifiés. En évaluant ces approches sur des bases de données variées, nous aspirons à surmonter les défis liés à la variabilité des conditions d'acquisition des données et à améliorer significativement la fiabilité et la performance des systèmes de reconnaissance d'empreintes palmaires.

IV.2 Présentation des outils de travail (Logiciel) Anaconda { Spyder +Python }

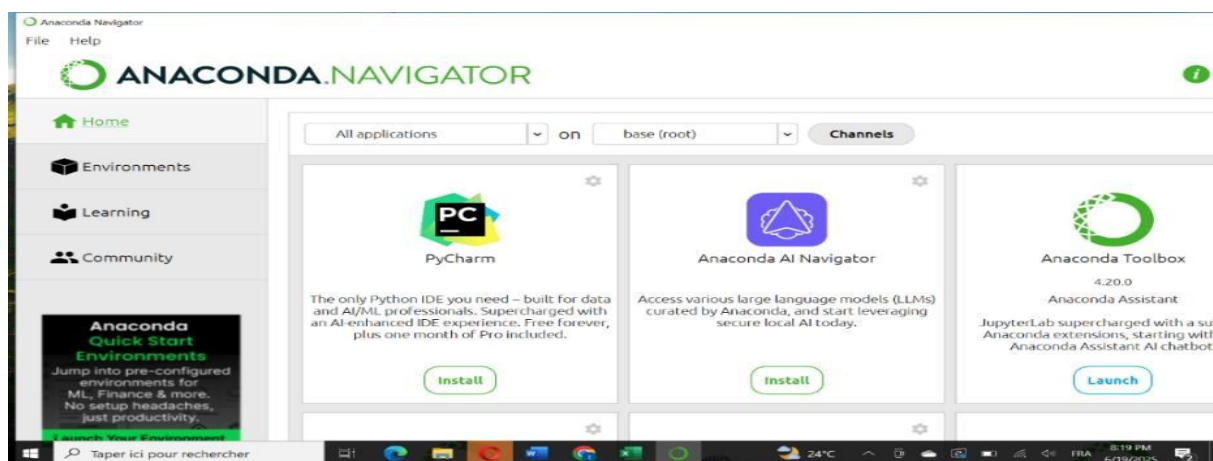


Figure IV.1: Environnement Anaconda

IV.2.1 Anaconda

Anaconda est une distribution libre et open-source de Python (et R) principalement utilisée pour la science des données, le machine learning, et l'analyse scientifique. Elle simplifie l'installation et la gestion des bibliothèques scientifiques grâce à un gestionnaire de paquets performant appelé conda. Anaconda est particulièrement utile dans les projets de deep learning, vision par ordinateur ou traitement de données, car elle intègre de nombreux outils prêts à l'emploi comme Jupyter Notebook, Spyder, NumPy, Pandas, Scikit-learn, TensorFlow, et bien d'autres. Elle permet également de gérer facilement des environnements virtuels, ce qui est essentiel pour isoler les dépendances de différents projets et éviter les conflits entre paquets. Grâce à son interface conviviale (Anaconda Navigator), elle est adaptée aussi bien aux débutants qu'aux utilisateurs avancés.

IV.2.2 Spyder

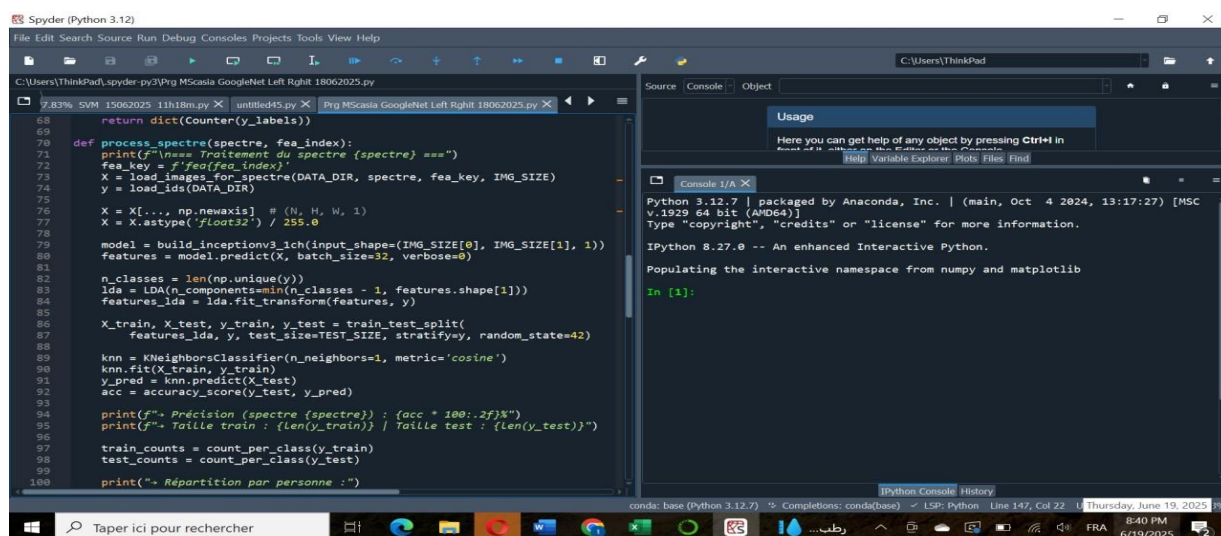


Figure IV.2: Environnement Spyder

Spyder (Scientific Python Development Environment) est un environnement de développement intégré (IDE) conçu spécialement pour la programmation en Python scientifique. Il est inclus par défaut dans la distribution Anaconda, ce qui en fait un choix populaire parmi les chercheurs, ingénieurs et étudiants.

Spyder offre une interface intuitive similaire à celle de MATLAB, ce qui facilite la transition pour les utilisateurs venant du monde de l'ingénierie. Il comprend plusieurs fonctionnalités puissantes :

- Un éditeur de code avancé avec coloration syntaxique, auto-complétion et suggestions intelligentes.
- Une console interactive IPython, idéale pour tester des portions de code en temps réel.
- Un explorateur de variables qui permet de visualiser et modifier les variables en mémoire, comme dans MATLAB.
- Des outils de débogage, profilage et analyse de performance.

Spyder est particulièrement adapté aux projets de data science, de modélisation, de calcul scientifique et de visualisation de données, en rendant l'expérience de développement plus fluide et interactive.

IV.2.3 Python

Python est un langage de programmation interprété, polyvalent, simple à apprendre et très lisible. Il a été créé à la fin des années 1980 par Guido van Rossum et est aujourd'hui l'un des langages les plus utilisés au monde dans des domaines variés comme :

- L'intelligence artificielle et le machine learning
- La science des données et les statistiques
- Le développement web (avec Flask, Django...)
- L'automatisation des tâches et les scripts systèmes
- Le développement de jeux, d'outils ou d'applications Ses points forts incluent :
- Une syntaxe claire et concise, idéale pour les débutants.
- Une immense bibliothèque standard et un grand nombre de modules open-source (NumPy, Pandas, TensorFlow, OpenCV, etc.).
- Une forte communauté mondiale très active, qui contribue à son développement continu.
- Sa compatibilité avec d'autres langages (C/C++, Java, etc.) et plateformes.

Grâce à sa simplicité et sa puissance, Python est devenu un outil incontournable dans le monde académique, industriel et technologique.



Figure IV.3 : Logiciel Python

IV.3 Les bases de données

IV.3.1 Base de données de l'empreinte palmaire MS-CASIA

La base de données que nous avons utilisée dans nos expérimentations est MS-CASIA multi spectral qui contient 7 200 images de paume capturées par 100 personnes différentes à l'aide d'un périphérique d'imagerie spectrale multiple conçu par l'utilisateur, comme illustré à Figure 4.1. Toutes les images palm sont des fichiers JPG de niveau de gris 8 bits. Pour chaque main, nous capturons deux sessions d'images de la paume. L'intervalle de temps entre les deux sessions est supérieur à un mois. Dans chaque session, il y a trois échantillons. Chaque échantillon contient six images de paume qui sont capturées en même temps avec six spectres électromagnétiques différents. Les longueurs d'onde de l'illuminateur correspondant aux six spectres sont respectivement de 460 nm, 630 nm, 700 nm, 850 nm, 940 nm et de la lumière blanche. Entre deux échantillons, nous autorisons un certain degré de variation des postures des mains. Grâce à cela, nous souhaitons augmenter la diversité des échantillons intra-classe et simuler une utilisation pratique.

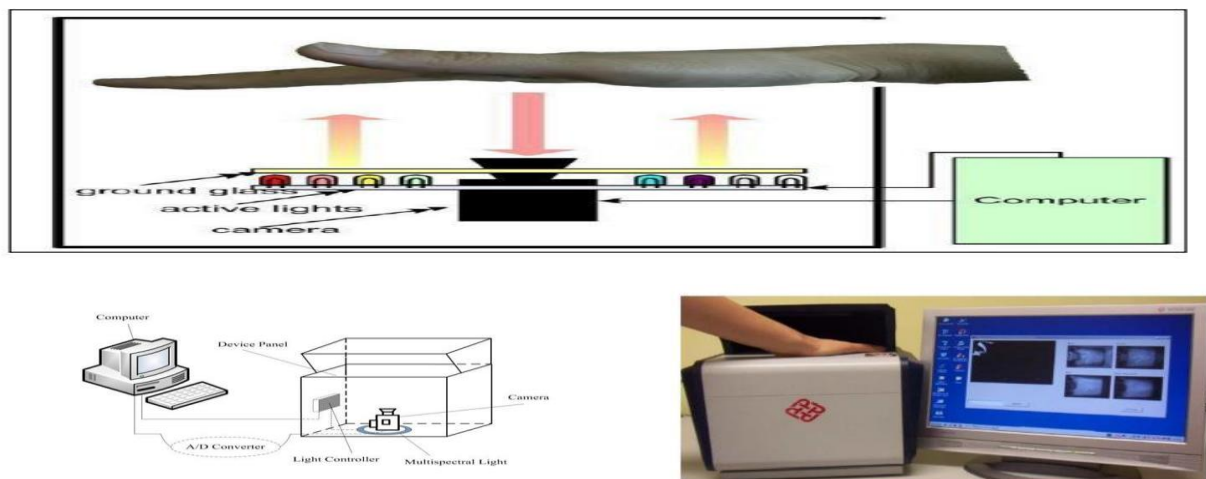


Figure IV.4: Dispositif d'imagerie Multi Spectrale

Dans notre appareil, il n'y a pas de piquets pour limiter les postures et les positions des paumes. Les sujets doivent créer un fond uniforme et coloré. L'appareil fournit un éclairage uniformément réparti et capture des images palm à l'aide d'une caméra CCD située au bas de l'appareil. Nous Résultats expérimentaux et discussions concevons un circuit de contrôle pour ajuster les spectres automatiquement. La Figure IV.5. Montre six images typiques de l'empreinte palmaire dans la base de données.

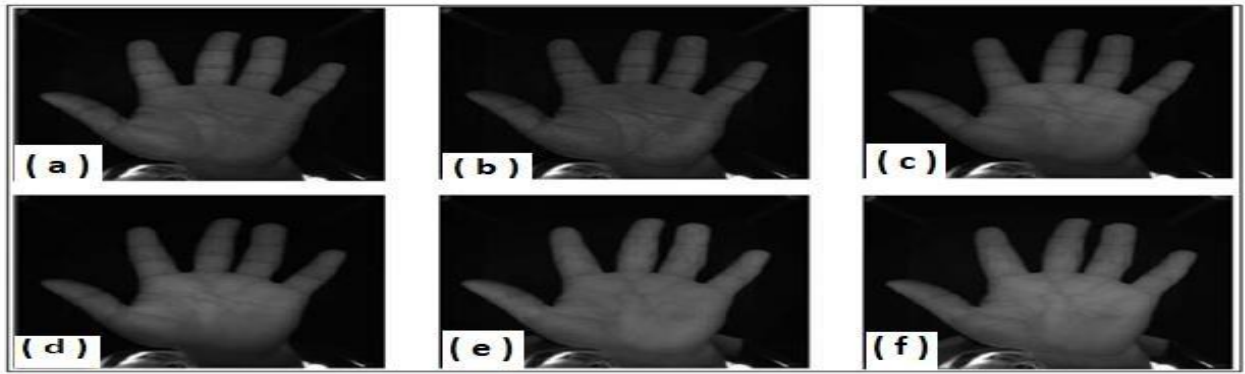


Figure IV.5: Six images d'une seule personne d'empreintes palmaires de la base MS-CASIA
(a) Spectre 460 , (b) spectre 630, (c) spectre 700, (d) spectre 850, (e) spectre 940 et (f) spectre WHT [56]

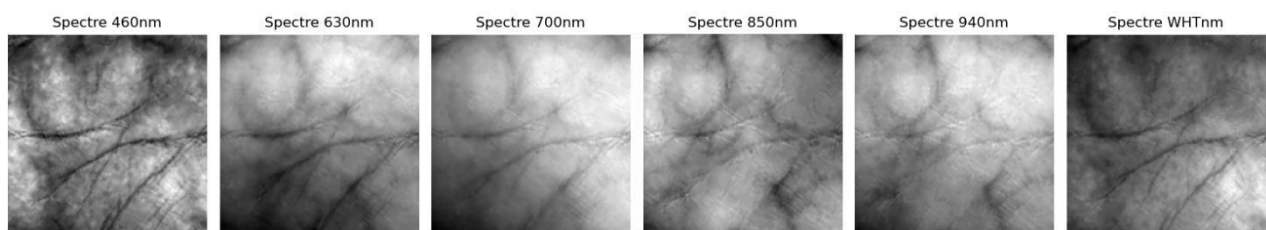


Figure IV.6: Six Spectres d'une seule personne d'empreintes palmaires ROI dans la base MS-CASIA [56]

IV.3.1.1 Formats d'image : Les images de la base de données MSCASIA sont stockées en tant que :

- XXX_ (L / R) _YYY_ZZ .jpg
- XXX : l'identifiant unique de personnes, varie de 000 à 100.
- (L/R): le type de palme, de « L » désigne la paume gauche et «R» désigne paume droite.
- YYY : Spectres électromagnétiques. « WHT » représente la lumière blanche.
- ZZ : l'indice des échantillons allant de 01 à 06. 01 au 03 appartiennent à la première session. 04-06 appartiennent à la deuxième session.

IV.3.2 Base de données multi-spectrale (MS-PolyU) : La base de données MSPolyU a été créée par l'université polytechnique de Hong Kong. Elle se compose de 6000 images d'empreintes digitales recueillies auprès de 500 volontaires. L'âge de chaque volontaire était compris entre 20 et 60 ans. Au cours du processus d'acquisition, chaque volontaire a été échantillonné 12 fois en deux sessions distinctes pour ses paumes gauche et droite [58] [59].

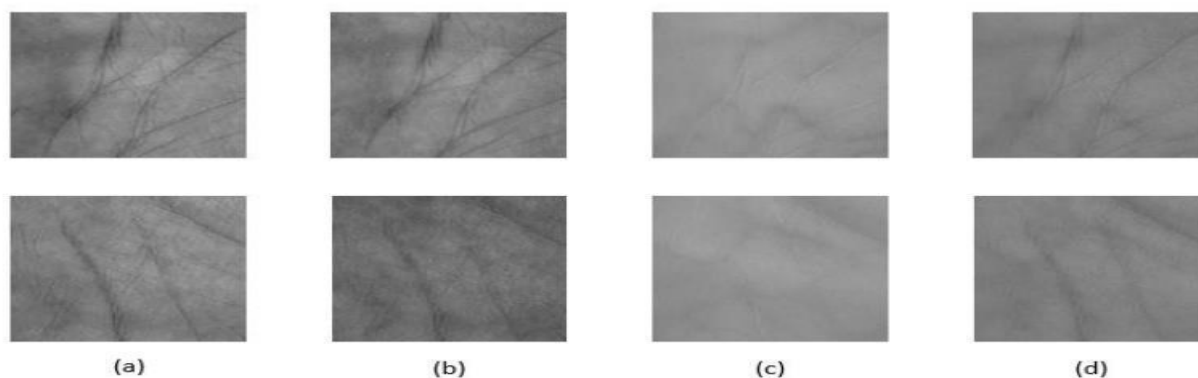


Figure IV.7: Échantillons de ROI d'empreintes palmaires de la base MS-PolyU (a) : bleu, (b) : vert, (c) : NIR, (d) : rouge

Les images d'empreintes palmaires ont été **acquises** dans quatre bandes spectrales, à savoir le rouge (RED), le vert (GREEN), le bleu (BLUE) et le proche infrarouge (NIR). Pour la commodité des chercheurs, l'université polytechnique de Hong Kong fournit les images de la région d'intérêt (ROI) de taille 128×128 . La figure ci-dessus montre quelques échantillons d'empreintes palmaires multispectrales dans la base de données PolyU [58] [59].

Base de données	Base de données MS-CASIA	Base de données MS-POLYU
Nombre de personnes	100	500
Nombre d'images	7200	6000

Tableau IV.1 : Différence entre les bases de données.

IV.4 Méthode proposée:

La base de données CASIA Multispectral Palmprint (CASIA-MS) représente une référence appropriée pour l'évaluation des systèmes biométriques. Notre méthode utilise exclusivement l'analyse discriminante linéaire (LDA) couplée au classifieur des k plus proches voisins (k-NN). La LDA transforme directement les images palmaires multispectrales en un espace de représentation optimisé qui accentue les différences inter-classes tout en réduisant les variations intra-classe. Cette approche permet de préserver et d'amplifier les caractéristiques discriminantes essentielles à la reconnaissance. Le classifieur k-NN opère ensuite dans cet espace pour attribuer chaque échantillon test à la classe correspondant à la majorité de ses k voisins les plus proches. La combinaison LDA/k-NN offre ainsi une solution performante et efficace pour la reconnaissance d'empreintes palmaires, avec l'avantage d'une implémentation simple et d'une interprétation directe des résultats.

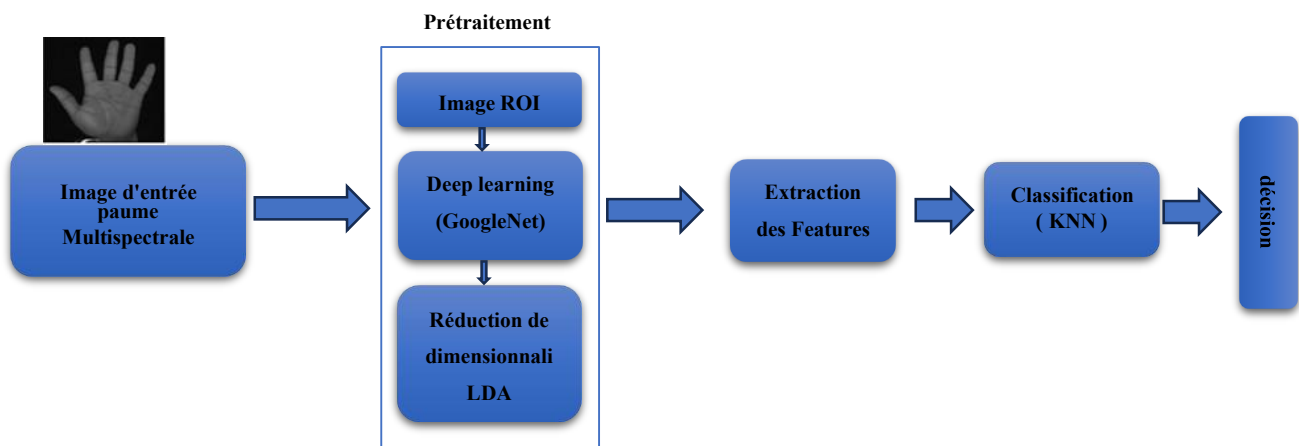


Figure IV.8: Architecture proposée du notre système de reconnaissance d'empreintes palmaires.

IV.5 Vue générale du processus et méthode de reconnaissance palmaire :

IV.5.1 Acquisition de l'image

L'acquisition de l'image constitue la première étape du processus. Elle consiste à capturer la paume de la main à l'aide de capteurs adaptés, tels que des scanners optiques, infrarouges ou capacitifs.

L'image obtenue servira de base pour les étapes de traitement ultérieures, notamment l'extraction des caractéristiques et la reconnaissance.

IV.5.2 Prétraitement

L'étape de prétraitement se compose de trois tâches principales :

- (a) l'utilisation directe de la région d'intérêt (ROI) prétraitée,
- (b) le redimensionnement des images ROI,
- (c) et l'extraction de caractéristiques à l'aide de LDA.

Dans notre cas, les images de la base de données **MSCASIA** ont une taille initiale de **768 × 576 pixels**. À partir de ces images, la région d'intérêt (ROI) est extraite et alignée selon la méthode décrite dans [34], pour isoler la zone contenant les informations discriminantes de la paume. Cette ROI est ensuite centrée et recadrée au format **200 × 200 pixels**.

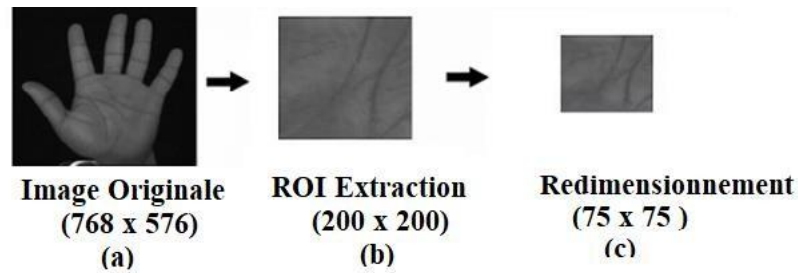


Figure IV.9: Les principales tâches de prétraitement. a. Image d'entrée, b. Extraction de (ROI) , c. Redimensionnement de l'image

IV.5.3 Extraction des caractéristiques

Après l'acquisition de l'image, l'étape suivante consiste à extraire les caractéristiques discriminantes de la paume. Pour cela, nous faisons appel aux réseaux de neurones convolutifs (CNN), qui sont des architectures de deep learning spécialement conçues pour le traitement des images. Un CNN est capable d'apprendre automatiquement des filtres (ou noyaux) qui détectent progressivement des motifs simples (bords, textures) puis des motifs plus complexes (formes, régions caractéristiques) à travers plusieurs couches hiérarchiques.

Dans ce travail, nous adoptons une approche basée sur l'apprentissage profond en utilisant GoogLeNet, un CNN pré-entraîné reconnu pour sa capacité à extraire des descripteurs visuels riches et pertinents. GoogLeNet introduit une architecture innovante basée sur les modules *Inception*, qui permettent de capturer simultanément des informations à différentes échelles (1×1 , 3×3 , 5×5), augmentant ainsi la profondeur tout en maintenant un coût computationnel modéré.

Grâce à cette conception, GoogLeNet permet d'extraire automatiquement des caractéristiques complexes et hiérarchisées à différents niveaux de l'image, offrant ainsi une représentation robuste et discriminative des empreintes palmaires, sans intervention manuelle dans le processus d'extraction.

IV.5.3.1 GoogLeNet (InceptionV1)

GoogLeNet (également connu sous le nom d'Inception V1) est une architecture de réseau neuronal convolutif (CNN) développée par Christian Szegedy et al. chez Google en 2014 (Szegedy et al.

(2014)). Cette architecture a été conçue pour classer les images de l'ensemble de données ImageNet, qui rappelons-le, contient plus de 14 millions d'images annotées, réparties sur plus de 20 000 catégories, dont 1 000 classes sont utilisées pour les compétitions comme le ImageNet Large Scale Visual Recognition Challenge (ILSVRC). GoogLeNet a obtenu des performances

nettement meilleures que les architectures CNN précédentes et a remporté le concours ILSVRC en 2014. Les résultats obtenus ont permis de réduire considérablement la marge d'erreur par rapport aux versions précédentes du concours ainsi qu'à ses concurrents directs. Cela a été rendu possible grâce à l'introduction de sous-réseaux appelés « modules Inception », conçus pour approfondir le réseau tout en optimisant les ressources. GoogLeNet est reconnu pour son efficacité à réduire le nombre de paramètres et le coût computationnel, tout en maintenant une grande précision. Il a également introduit l'idée de classificateurs auxiliaires, qui sont des classificateurs secondaires formés en parallèle avec le classificateur principal afin d'améliorer les performances globales du modèle.

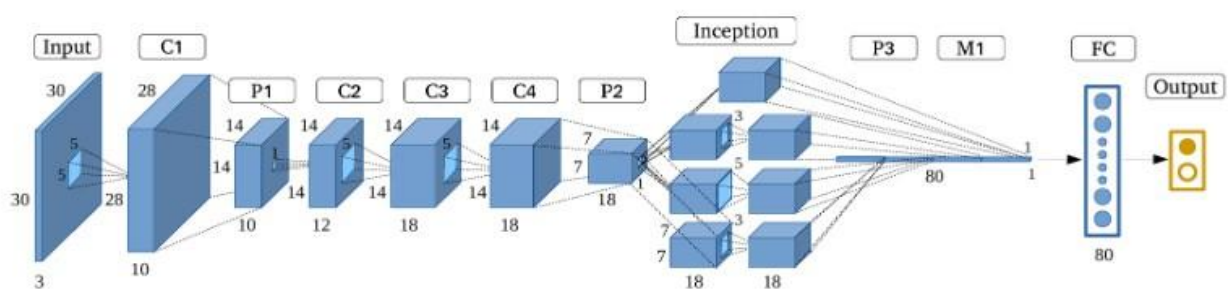


Figure IV.10 : Architecture Globale de GoogLeNet (Inception v1)

1. Bloc d'entrée (Input Layer)

L'image d'entrée, généralement de taille $30 \times 30 \times 3$, représente une image RGB (trois canaux de couleur). Cette couche ne réalise aucun traitement, mais constitue le point de départ du flux d'information. Elle transmet les données brutes du domaine spatial à la première couche convolutive.

L'objectif est de maintenir l'intégrité de l'information spatiale tout en préparant l'image pour les étapes de traitement hiérarchique suivantes.

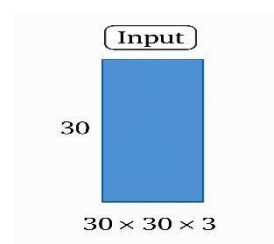


Figure IV.11 : Bloc d'entrée

2. Bloc C1 – Convolution initiale

La première couche convolutive (C1) applique un ensemble de **10 filtres de taille 5×5** avec un stride de 1. Elle permet d'extraire les premières caractéristiques locales telles que les bords, textures simples ou transitions d'intensité. La sortie est une carte de caractéristiques de taille **28×28×10**, réduite en taille par l'absence de padding, tout en augmentant la profondeur. Cette étape pose les fondations de la hiérarchie de représentations nécessaires à la reconnaissance.

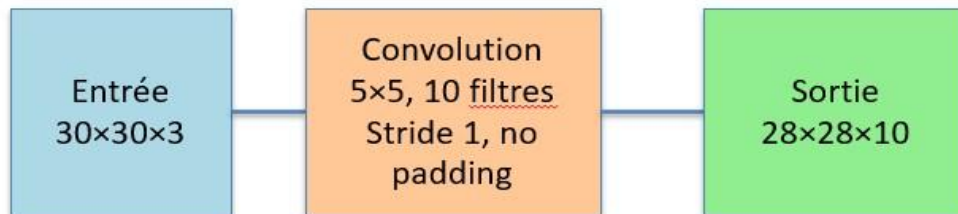


Figure IV.12 :Bloc C1 – Première couche convolutionnelle

3. Bloc P1 – Max Pooling (2×2) :

Le bloc P1 constitue une étape essentielle de réduction spatiale dans l'architecture du réseau. Il applique une opération de max pooling avec une fenêtre de taille 2×2 et un stride (pas) de 2, ce qui signifie que le filtre se déplace de deux pixels à la fois. Ce processus permet de transformer une entrée de taille 28×28×10 (provenant de la couche C1) en une sortie de dimension 14×14×10, en divisant chaque carte de caractéristiques par un facteur de 2 dans les dimensions spatiales (largeur et hauteur), tout en conservant la profondeur (nombre de canaux) intacte.

Le rôle principal du max pooling est de résumer les informations les plus importantes dans chaque sous-région 2×2. Pour chaque région, le réseau conserve uniquement la valeur maximale, ce qui correspond souvent à une activation forte détectée par un filtre convolutif. Ce mécanisme agit comme un filtrage des activations faibles, en mettant l'accent sur les motifs ou contours dominants présents dans l'image.

En somme, le bloc P1 est un composant stratégique qui favorise la robustesse, l'efficacité et la généralisation du réseau convolutionnel. Il équilibre l'objectif de préservation de l'information importante avec celui de simplification de la représentation

4. Blocs C2, C3, C4 – Convolutions intermédiaires

- C2 applique 12 filtres 5×5 , produisant une sortie de $14 \times 14 \times 12$.
- C3 applique 18 filtres $5 \times 5 \rightarrow$ sortie $14 \times 14 \times 18$.
- C4 affine davantage les descripteurs tout en conservant la même taille spatiale.

Ces couches approfondissent la représentation des motifs visuels, permettant au réseau de capturer des structures de plus en plus complexes.

5. Bloc C2 – Convolution 1×1

Le bloc C2 applique une opération de convolution 1×1 , consistant à faire passer un filtre de taille 1 pixel par 1 pixel sur chaque position de l'image d'entrée. Contrairement aux grandes convolutions (3×3 ou 5×5), qui extraient des motifs spatiaux complexes, la convolution 1×1 n'a pas d'effet sur la structure spatiale de l'image (la taille reste 14×14), mais elle opère en profondeur, c'est-à-dire entre les canaux.

Dans le cas du bloc C2, on applique 12 filtres 1×1 , ce qui signifie que l'on transforme une entrée de taille $14 \times 14 \times 10$ (en provenance de P1) en une sortie de taille $14 \times 14 \times 12$. Chaque pixel de la nouvelle carte de caractéristiques est donc une combinaison linéaire des valeurs issues des 10 canaux d'entrée, pondérées par les poids du filtre 1×1 .

L'utilisation de la convolution 1×1 a été popularisée dans l'architecture Inception (GoogLeNet) pour son efficacité à réduire la complexité computationnelle tout en préservant l'information discriminante. Elle permet également d'introduire une forme de projection linéaire sur les canaux, comparable à une couche dense appliquée de façon indépendante à chaque position spatiale de l'image.

Ainsi, le bloc C2, bien qu'il semble simple, joue un rôle stratégique dans le contrôle de la dimensionnalité et l'apprentissage d'interactions inter-canaux à faible coût, contribuant à une architecture plus profonde et plus efficace.

Extraction hiérarchique de caractéristiques

- C2 détecte des **motifs basiques** à moyenne échelle (coins, textures, contours un peu larges). C3 construit à partir des sorties de C2 des **compositions plus complexes** de ces motifs (ex. : motifs texturés, courbes, croisements).
- C4 affine et renforce ces motifs dans un **espace de représentation plus riche**, en mettant en évidence les structures pertinentes pour la tâche (comme la reconnaissance)



Figure IV.13 :Blocs C2, C3, C4 – Convolutions successives

6. Bloc P2 – Pooling 2

Comme pour P1, le bloc P2 applique un pooling pour réduire la taille spatiale à $7 \times 7 \times 18$. Cette réduction permet de concentrer les informations extraites sur des zones plus globales, en préparant la donnée pour l'étape de traitement parallèle dans le module Inception. Cela diminue également la mémoire nécessaire et le risque de surapprentissage.

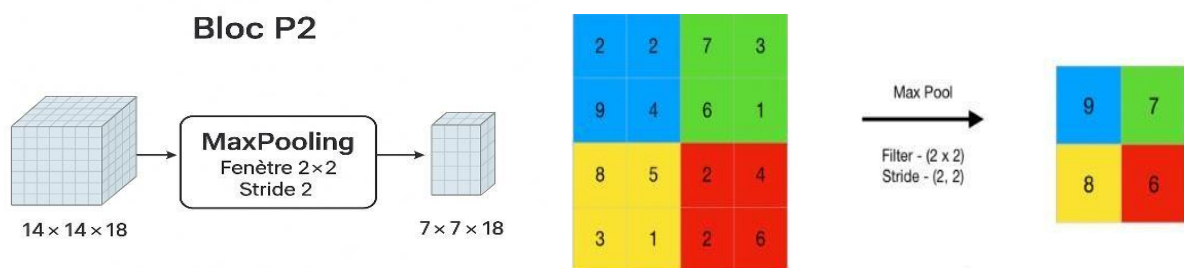


Figure IV.14: Bloc P2 – Pooling 2

7. Bloc Inception (Module Inception)

Le module Inception est la clé de l'architecture. Il applique en parallèle :

- Une convolution 1×1 (réduction de dimension),
- Une convolution 3×3 ,
- Une convolution 5×5 ,
- Une max pooling 3×3 suivie d'une convolution 1×1 .

Chaque branche extrait une forme de caractéristique différente : fine (1×1), locale (3×3), ou plus globale (5×5). La concaténation des sorties permet de fusionner ces diverses perspectives, produisant une sortie riche en information de taille $7 \times 7 \times N$ (N dépend du nombre de filtres dans chaque branche). Cela rend le réseau plus expressif tout en contrôlant le coût computationnel.

Chaque type de filtre extrait des informations différentes sur la même région de l'image :

- Les petits filtres (1×1 , 3×3) détectent des détails précis,
- Les grands filtres (5×5) captent des motifs plus étendus,
- Le pooling renforce la robustesse aux déformations et au bruit.

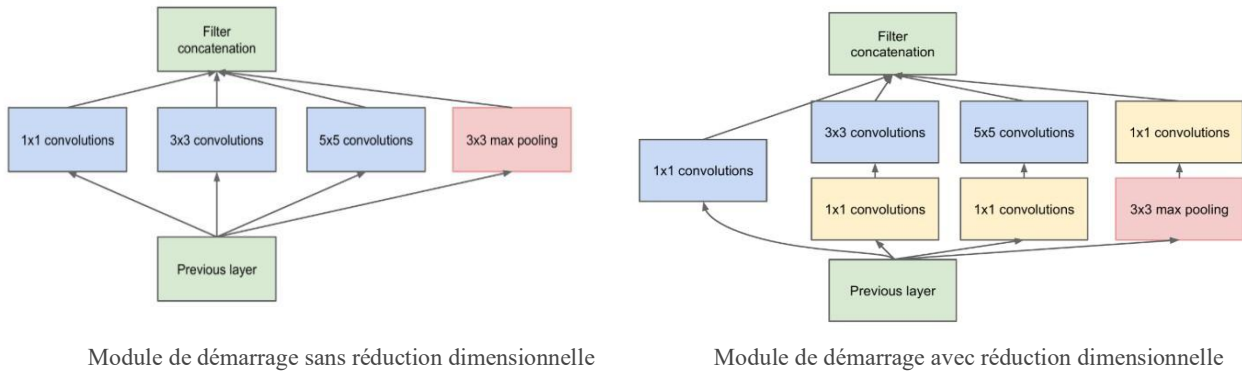


Figure IV.15: Inception module

8. Bloc P3 – Global Average Pooling

Le pooling global (souvent appelé global average pooling) transforme chaque carte de 7×7 en une seule valeur moyenne, produisant une sortie $1 \times 1 \times 80$. Cette méthode réduit drastiquement la dimension sans ajouter de paramètres, et agit comme une couche de régularisation, empêchant le surapprentissage tout en conservant l'essence des activations.

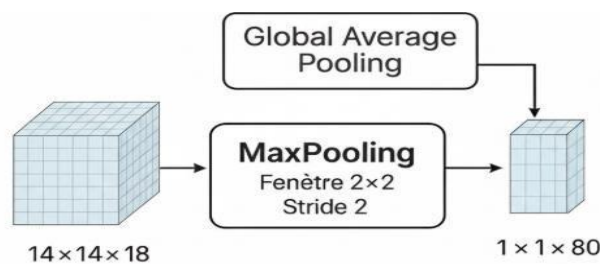


Figure IV.16: Bloc P3 – Global Average Pooling

9. Bloc M1 – Flatten

La sortie $1 \times 1 \times 80$ est aplatie en un vecteur de 80 éléments. Ce vecteur résume les caractéristiques extraites sur l'ensemble de l'image, prêt à être introduit dans une couche entièrement connectée. Ce passage du domaine spatial au domaine vectoriel est essentiel pour la classification finale.

10.Bloc FC – Fully Connected + Sortie

Le vecteur est transmis à une couche entièrement connectée qui calcule la probabilité d'appartenance à chaque classe (via softmax dans les tâches de classification). Le nombre de neurones dans cette couche correspond au nombre de classes cibles. La sortie finale est donc un vecteur de scores ou de probabilités qui permet de prendre une décision de reconnaissance.

11.fonction d'activation Softmax

La fonction Softmax est une fonction d'activation utilisée principalement dans la dernière couche des réseaux de neurones pour la classification multiclasse. Elle transforme un vecteur de scores (logits) en probabilités, facilitant ainsi la prise de décision sur la classe prédite

Pour un vecteur de sortie $\mathbf{z} = [z_1, z_2, \dots, z_K]$ de la couche précédente (par exemple, une couche entièrement connectée), la fonction Softmax est définie par :

$$\text{Softmax}(z_i) = \frac{e^{z_i}}{\sum_{j=1}^K e^{z_j}} \quad \text{pour } i = 1, 2, \dots, K$$

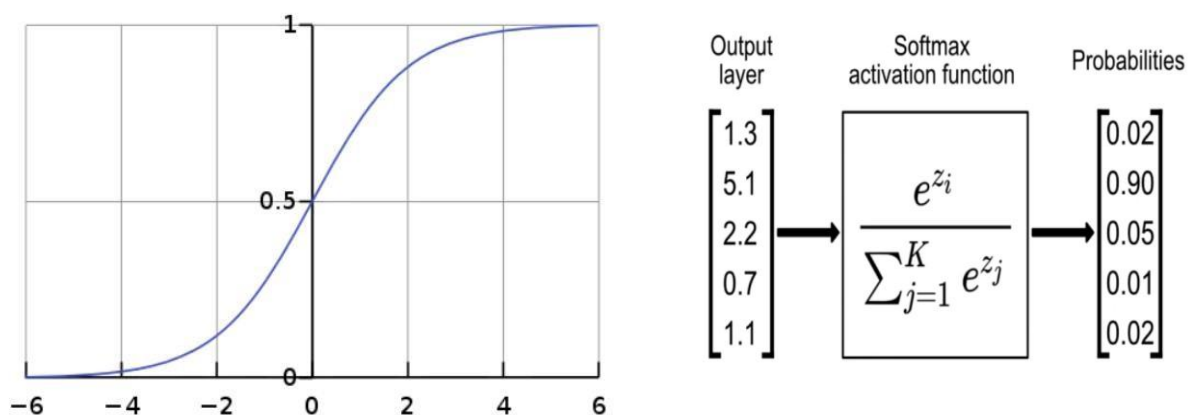


Figure IV.17: Fonction d'activation Softmax

IV.5.3.2 Progression des dimensions spatiales de l'image à travers GoogLeNet de notre système

Le tableau IV.2 ci-dessus décrit l'évolution des dimensions spatiales des cartes de caractéristiques à travers un segment d'un réseau neuronal convolutif (CNN), qui pourrait s'intégrer dans une architecture plus large de type GoogLeNet (Inception). Cette progression

illustre la capacité du réseau à extraire des caractéristiques hiérarchiques tout en gérant efficacement la complexité computationnelle par la réduction de la résolution spatiale.

1. Image Originale : 200×200 pixels L'entrée du modèle est une image de 200 pixels de hauteur sur 200 pixels de largeur. Cette résolution fournit une base suffisante pour capturer les détails nécessaires à l'analyse, tels que les motifs des plis palmaires ou les caractéristiques spectrales dans le cas d'images multispectrales.

2. Couche Convulsive (Conv(5×5) stride=1) : 196×196 pixels La première opération est une convolution avec un filtre de 5×5 pixels et un pas (stride) de 1. L'absence de padding explicite (ou un padding "valid") entraîne une réduction des dimensions de sortie. Pour une image d'entrée H×W et un filtre F×F avec un stride S, la taille de sortie est donnée par $\lfloor (H-F)/S \rfloor + 1$. Dans ce cas : $\lfloor (200-5)/1 \rfloor + 1 = 196$. Cette couche est responsable de l'extraction des caractéristiques de bas niveau, telles que les bords, les coins et les textures rudimentaires, sur de plus grandes régions de l'image d'entrée.

3. Couche de Max Pooling (MaxPool(2×2) stride=2) : 98×98 pixels Suite à la convolution, une couche de max pooling avec une fenêtre de 2×2 et un pas de 2 est appliquée. Le max pooling est une opération de sous-échantillonnage non-linéaire qui réduit la dimensionnalité spatiale et confère une certaine invariance aux petites translations. La taille de sortie est calculée comme $\lfloor H/S \rfloor$. Ici : $\lfloor 196/2 \rfloor = 98$. Cette réduction de taille de 196×196 à 98×98 permet de diminuer la charge computationnelle des couches suivantes et de rendre le modèle plus robuste aux légères variations dans la position des caractéristiques.

4. Couche Convulsive (Conv(3×3) stride=1) : 96×96 pixels Une autre couche de convolution avec un filtre de 3×3 et un pas de 1 est appliquée. Le calcul de la taille de sortie est similaire : $\lfloor (98-3)/1 \rfloor + 1 = 96$. Cette étape permet d'apprendre des caractéristiques de niveau légèrement supérieur à partir des cartes de caractéristiques déjà réduites, souvent des combinaisons de caractéristiques de bas niveau.

5. Couche de Max Pooling (MaxPool(2×2)) : 48×48 pixels Un second max pooling avec une fenêtre de 2×2 (et un pas implicite de 2 car non spécifié, ce qui est la convention par défaut pour le max pooling) réduit la taille à : $\lfloor 96/2 \rfloor = 48$. Cette répétition du max pooling continue la compression des informations spatiales, permettant aux couches profondes de se concentrer sur des caractéristiques plus abstraites et globales.

6. Couche Convulsive (Conv(3×3)) : 46×46 pixels Une troisième couche de convolution avec un filtre de 3×3 (pas de 1 implicite) est appliquée : $\lfloor (48-3)/1 \rfloor + 1 = 46$. Cette couche contribue à affiner davantage les caractéristiques extraites, souvent en détectant des motifs plus complexes ou des relations entre les caractéristiques de niveau inférieur.

7. Couche de Max Pooling (MaxPool(2×2)) : 23×23 pixels La dernière couche de max pooling avec une fenêtre de 2×2 et un pas de 2 réduit la carte de caractéristiques à : $\lfloor 46/2 \rfloor = 23$. À ce stade, la carte de caractéristiques est considérablement compacte. Elle contient une représentation dense et hautement sémantique de l'image originale, où chaque pixel dans la carte de 23×23 correspond à une large région réceptive de l'image d'entrée.

8. Interpolation/Resize : 75×75 pixels Après les phases de réduction de dimensionnalité, une étape d'interpolation (ou de redimensionnement) est appliquée pour amener la carte de caractéristiques de 23×23 à une taille de 75×75 pixels.

Étape	Taille image
Image originale	200 × 200
Conv(5×5) stride=1	196 × 196
MaxPool(2×2) stride=2	98 × 98
Conv(3×3) stride=1	96 × 96
MaxPool(2×2)	48 × 48
Conv(3×3)	46 × 46
MaxPool(2×2)	23 × 23
Interpolation/Resize	75 × 75

Tableau IV.2: Dimensions dans notre modèle utilisant GoogleNet .

IV.5.3.3 LDA :Analyse discriminante linéaire (Linear Discriminant Analysis)

L'analyse discriminante linéaire (LDA – *Linear Discriminant Analysis*) est une méthode supervisée de réduction de dimension utilisée en apprentissage automatique et en reconnaissance de formes.

Contrairement à des techniques non supervisées comme l'ACP (PCA), LDA exploite les étiquettes de classe pour projeter les données dans un espace de dimension réduite tout en maximisant la séparabilité entre les différentes classes. L'objectif fondamental de LDA est de

trouver une transformation linéaire des données qui maximise la distance entre les moyennes des classes (variance inter-classes) tout en minimisant la dispersion des données au sein de chaque classe (variance intra-classe).

Mathématiquement, LDA cherche un vecteur de projection $\mathbf{w}$ qui maximise le rapport suivant :

$$J(\mathbf{w}) = \frac{\mathbf{w}^T S_B \mathbf{w}}{\mathbf{w}^T S_W \mathbf{w}}$$

où S_B est la matrice de dispersion inter-classes, et S_W est la matrice de dispersion intra-classes. La matrice S_B est calculée à partir des différences entre les moyennes de chaque classe et la moyenne globale, tandis que S_W est basée sur les écarts des échantillons individuels par rapport à la moyenne de leur propre classe. Ce problème se résout classiquement par une équation aux valeurs propres :

$$S_W^{-1} S_B \mathbf{w} = \lambda \mathbf{w}$$

Les vecteurs propres associés aux plus grandes valeurs propres λ représentent les directions de projection les plus discriminantes.

En pratique, si l'on dispose de c classes, LDA peut projeter les données dans un espace de dimension maximale $c-1$. Cette propriété est particulièrement utile pour la visualisation et la classification, notamment dans les domaines de la biométrie (empreintes palmaires, reconnaissance faciale), du diagnostic médical, et de l'analyse de texte. LDA est ainsi une méthode puissante pour la classification supervisée, car elle combine efficacement réduction de dimension et discrimination entre classes

IV.5.3.4 K plus proches voisins (KNN) :

L'algorithme des K plus proches voisins (KNN) est une méthode de classification supervisée qui repose sur le principe de similarité : un échantillon inconnu est classé en fonction des classes des K échantillons les plus proches dans l'espace des caractéristiques. Lorsqu'un nouveau point doit être classé, KNN calcule la distance (souvent euclidienne) entre ce point et tous les points du jeu d'entraînement, puis sélectionne les K voisins les plus proches. La classe la plus fréquente parmi ces voisins est alors attribuée au point inconnu. La distance euclidienne est donnée par la formule :

où $\underline{x}$ est le vecteur à classer, x_i un vecteur du jeu d'entraînement, et n le nombre de caractéristiques. KNN est simple, non paramétrique et efficace, mais il peut devenir coûteux pour de grands ensembles de données formule cosine :

$$\text{cosine}(A, B) = \frac{A \cdot B}{\|A\| \cdot \|B\|}$$

- $A \cdot B$ est le produit scalaire des vecteurs A et B ,
- $\|A\|$ est la norme (longueur) du vecteur A ,
- $\|B\|$ est la norme du vecteur B .

IV.5.4 Classification :

Pendant la classification, les caractéristiques extraites de l'image de la paume sont comparées avec celles stockées dans la base de données. Cette étape utilise divers algorithmes pour effectuer la comparaison et déterminer l'identité de l'individu.

IV.5.5 Décision :

La dernière étape du processus est la prise de décision, où le système détermine si l'individu est identifié ou vérifié avec succès. Si la comparaison des caractéristiques extraites avec celles de la base de données dépasse un certain seuil de confiance, l'identification est considérée comme réussie.

IV.6 Expérimentations et Résultats

Cette section présente l'évaluation expérimentale réalisée à l'aide de deux bases de données publiques et largement utilisées dans la reconnaissance palmaire : CASIA Multispectral Palmprint et PolyU Multispectral Palmprint. Dans un premier temps, Ensuite, nous détaillons la configuration expérimentale mise en place pour évaluer notre approche. Enfin, nous présentons et analysons les résultats obtenus dans la section des résultats expérimentaux.

IV.6.1 Expérimentation avec la base de données MS-PolyU

Dans le cadre de nos expérimentations avec la base de données MS-PolyU, nous avons mis en œuvre un protocole de classification **train-test** pour évaluer la performance de nos modèles. Deux configurations distinctes de répartition des données ont été explorées. Dans la première approche, **80% de la base de données a été allouée à l'ensemble d'entraînement**, permettant au modèle d'apprendre sur une proportion majoritaire des données multispectrales d'empreintes palmaires. Les **20% restants ont constitué l'ensemble de test**, servant à évaluer la

capacité de généralisation du modèle sur des données non vues. Une seconde configuration a également été étudiée, où la base de données MS-PolyU a été divisée de manière égale, avec **50% des données dédiées à l'entraînement et 50% au test**. Cette approche vise à analyser l'impact d'une quantité d'entraînement réduite sur la robustesse et la précision de la classification des empreintes palmaires, offrant une perspective comparative sur la performance de nos algorithmes dans différentes conditions de disponibilité des données, les résultats est dans le **tableau IV.3**

50% de l'entraînement et 50% au test			
Longueur d'onde	TEST_SIZE = 0.5	TEST_SIZE = 0.8	TEST_SIZE = 0.2
Blue nm	100.00%	100.00%	100.00%
Green nm	100.00%	100.00%	100.00%
Red nm	100.00%	100.00%	100.00%
NIR nm	100.00%	100.00%	100.00%
Moyenne	100.00%	100.00%	100.00%

Tableau IV.3 Résumé des résultats par spectre pour MSPOLY

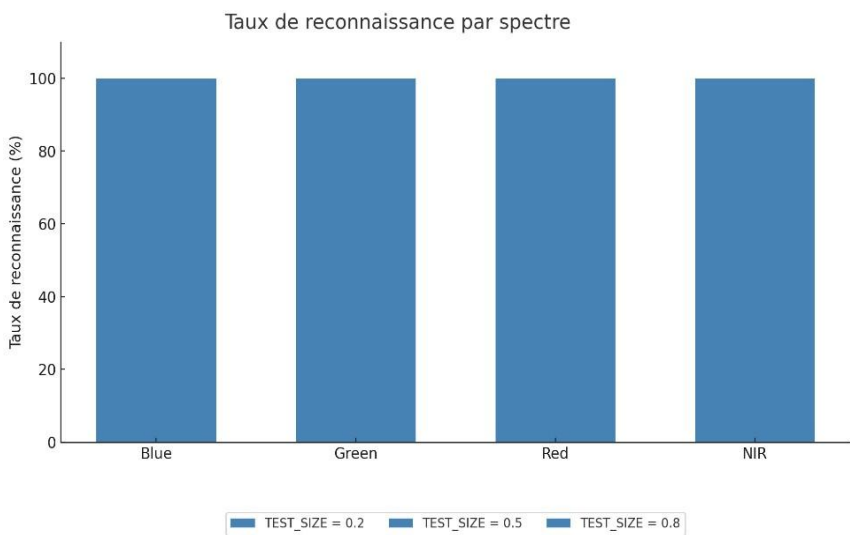


Figure IV.18: Taux de reconnaissance sur la base MS-PolyU

IV.6.2 Discussions des résultats :

Les résultats obtenus avec GoogleNet, affichant une précision de 100% sur toutes les bandes spectrales (Bleu, Vert, Rouge, NIR) et pour l'ensemble des proportions d'entraînement/test (0.5, 0.8, 0.2), sont remarquablement élevés. Cela suggère une performance exceptionnelle du modèle dans l'extraction de caractéristiques discriminantes pour la reconnaissance des empreintes palmaires. La constance de ces résultats, quelle que soit la longueur d'onde utilisée ou la répartition des données entre entraînement et test, indique une grande robustesse et une capacité remarquable de GoogleNet à généraliser sur ces données spécifiques. Une telle perfection pourrait soit témoigner d'un problème de classification particulièrement bien adapté au modèle et à la taille d'image (128x128), soit soulever des questions sur la complexité intrinsèque de la base de données utilisée ou la possibilité d'un surapprentissage si l'ensemble de données est limité. Cependant, sur la base de ces chiffres, le modèle démontre une maîtrise complète de la tâche de classification.

IV.6.3 Expérimentation avec la base de données MS-CASIA (Main Droite)

Main Gauche			
Longueur d'onde	TEST_SIZE = 0.2	TEST_SIZE = 0.5	TEST_SIZE = 0.8
460 nm	98.33%%	99.00%	95.21%
630 nm	100.00%	99.00%	96.88%
700 nm	100.00%	98.67%	95.21%
850 nm	99.17%	98.67%	95.83%
940 nm	100.00%	100.00%	97.92%
WHT nm	96.67%	99.33%	95.62%
Moyenne	99.03%	99.11%	96.11%

Tableau IV.4 des résultats par spectre pour MSCASIA(Main Gauche)

Dans le cadre de nos expérimentations avec la base de données **MS-CASIA**, nous avons rigoureusement appliqué un protocole de classification **train-test** afin d'évaluer la performance de nos modèles spécifiquement sur les **empreintes palmaires de main droite**. Deux configurations distinctes de répartition des données ont été mises en œuvre. Dans la première

approche, **80% de la base de données multispectrale a été consacrée à l'ensemble d'entraînement**, fournissant une quantité substantielle de données pour l'apprentissage des modèles, tandis que les **20% restants ont été réservés à l'ensemble de test**, permettant d'évaluer leur capacité de généralisation sur des données de main droite non vues. Une seconde configuration a également été explorée, où la base de données MS-CASIA a été divisée équitablement, avec **50% des données dédiées à l'entraînement et 50% au test**. Cette diversité dans les proportions de division vise à analyser l'impact de la taille de l'ensemble d'entraînement sur la robustesse et la précision de la reconnaissance des **empreintes palmaires de main droite**, et à fournir une compréhension approfondie de la performance de nos algorithmes sous différentes contraintes de disponibilité des données, les résultats dans le **tableau IV.4**.

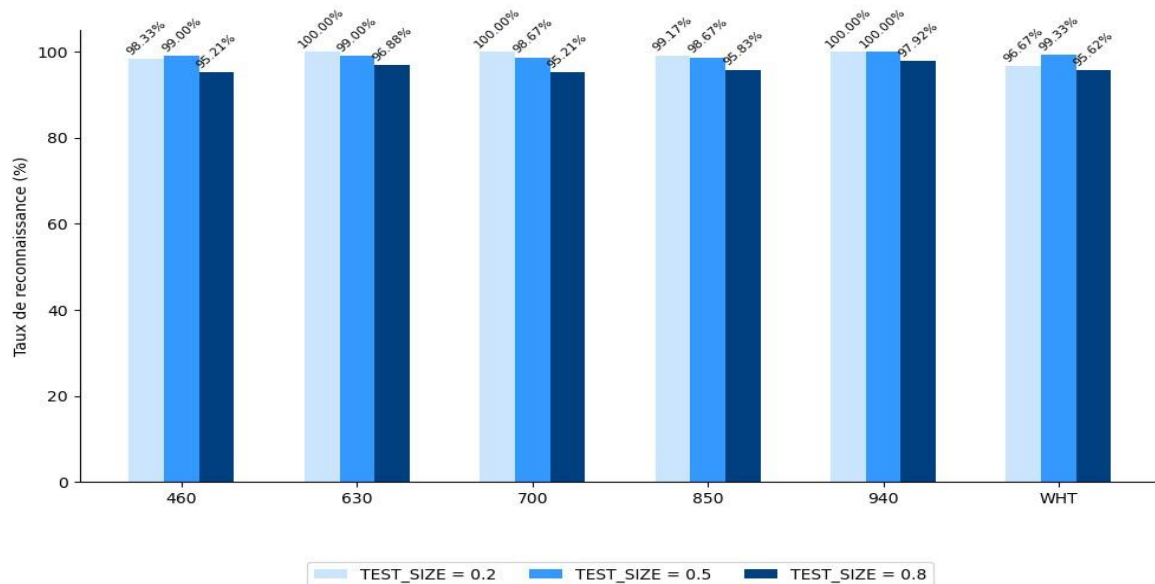


Figure IV.19: Taux de reconnaissance sur la base MS-CASIA (Main Gauche)

IV.6.4 Résultats de MSCASIA Main Gauche

Les résultats obtenus témoignent d'une performance globale élevée du système de reconnaissance biométrique palmaire basé sur des images multispectrales. La précision varie en fonction de la longueur d'onde utilisée ainsi que de la taille du jeu de test.

Lorsque 80 % des données sont consacrées à l'apprentissage et 20% de test, les performances sont déjà excellentes, avec une moyenne de 99.03 %, et plusieurs spectres, comme 630 nm, 700 nm et 940 nm, atteignent 100 % de moyenne. A taille de teste de 50%, c'est-à-dire avec un équilibre entre apprentissage et test, le système atteint sa meilleure performance moyenne avec 99.11 %, illustrant la forte capacité de généralisation du modèle GoogleNet combiné à LDA et k-NN.

En revanche, lorsque 80 % des données sont réservées au test, les performances diminuent légèrement mais demeurent solides, avec une moyenne encore élevée de 96.11 %. Ce résultat démontre la résilience du système, même avec un faible volume d'apprentissage (20%).

Sur le plan spectral :

- Le spectre 940 nm se démarque avec un taux constant de 100 % pour taille de test 20% et 50%, et 97.92 % à 80%, confirmant la supériorité des longueurs d'onde infrarouges dans la captation des détails biologiques profonds.
- Le spectre 630 nm offre également des performances exceptionnelles (100.00 %, 99.00 %, 96.88 %).
- Les longueurs d'onde visibles, telles que 460 nm et 700 nm, affichent une légère variabilité (460 nm allant de 98.33 % à 95.21 %) pouvant être attribuée à une plus grande sensibilité au bruit et à un contraste spectral moins accentué sur les motifs palmaires.

Globalement, le modèle présente une robustesse remarquable face à la variabilité des données, et confirme la pertinence des spectres infrarouges (850 et 940 nm) pour une reconnaissance fiable et précise de la main. Ce comportement est cohérent avec les résultats présentés dans le tableau récapitulatif.

Main droite			
Longueur d'onde	TEST_SIZE = 0.2	TEST_SIZE = 0.5	TEST_SIZE = 0.8
460 nm	97.50%	96.67%	94.17%
630 nm	100.00%	99.33%	97.08%
700 nm	98.33%	99.33%	96.04%
850 nm	100.00%	99.33%	98.33%
940 nm	98.33%	99.67%	98.12%
WHT nm	96.67%	99.00%	96.04%
Moyenne	98.47%	98.89%	96.63%

Tableau IV.5 résultats par spectre pour MSCASIA (Main droite)

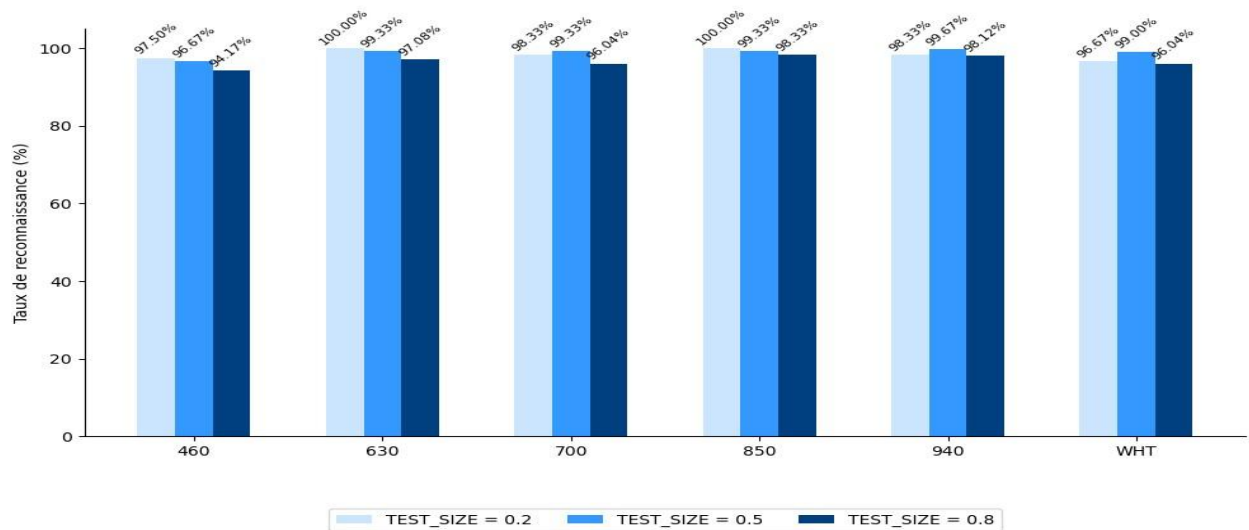


Figure IV.20: Taux de reconnaissance sur la base MS-CASIA (Main Droite)

IV.6.5 Résultats de MSCASIA Main Droite

Les résultats expérimentaux obtenus pour la main droite confirment également la très haute efficacité du système de reconnaissance palmaire multispectrale basé sur GoogleNet, LDA et k-NN, avec des taux de reconnaissance supérieurs à 94 % dans tous les cas. Lorsque le modèle est entraîné avec 80 % des données apprentissage 20%, il atteint une moyenne de 98.47 %, illustrant sa stabilité dès les premiers niveaux d'apprentissage.

À apprentissage 50%, c'est-à-dire avec un équilibre entre les données d'apprentissage et de test, la précision augmente légèrement pour atteindre une moyenne de 98.89 %, montrant une capacité de généralisation optimale sur des échantillons inconnus. Lorsque la quantité de données d'entraînement est la plus réduite apprentissage 80%, la performance reste très solide, avec une moyenne de 96.63 %, ce qui souligne une robustesse notable du système même dans des scénarios de données limitées.

Du point de vue des spectres :

Les meilleures performances sont obtenues avec les spectres 630 nm, 850 nm et 940 nm, qui atteignent tous plus de 97 % pour les trois valeurs de taille de test, et même 100 % dans certains cas, notamment pour le spectre 630 nm à apprentissage 20%, et 850 nm à apprentissage 20%.

Les spectres visibles comme 460 nm et WHT affichent des résultats légèrement inférieurs, avec respectivement 94.17 % et 96.04 % pour apprentissage 80%. ce qui pourrait s’expliquer par une moindre pénétration optique dans les tissus ou une variabilité accrue des motifs de surface.

Ces résultats démontrent que l’approche est hautement performante, et que la main droite peut être reconnue avec une très grande fiabilité à partir d’images multispectrales. Les longueurs d’onde proches de l’infrarouge (notamment 850 nm et 940 nm) apparaissent comme les plus robustes, consolidant leur intérêt dans les systèmes biométriques modernes.

Référence	Années	Base de données	Descripteur / Classifieur	Taux de reconnaissance%			
				BLEU	RED	NIR	GREEN
[64]	2019	MS-Polyu	CNN	/	99.99	99.99	99.99
[69]	2019	MS-Polyu	transform (UDTCWT)	99.49	99.67	99.98	99.93
[65]	2019	MS-Polyu	CNN-Alexnet	99.99	99.99	/	/
[66]	2021	MS-Polyu	MAHCO	99.33	99.96	99.50	99.46
		MS-Polyu	HP	99.93	100	99.90	99.73
[67]	2023	MS-Polyu	BSIF_26/KNN	100	100	100	100
[70]	2024	MS-Polyu	BSIF26_/LDA KFA	100	100	100	100
Proposée	2025	MS-Polyu	GoogleNet (InceptionV3) LDA+KNN	100	100	100	100

Tableau IV.6: Comparaison de notre méthode avec d'autres approches spécifiant le spectre dans la base de POLYU

L’analyse comparative des performances sur la base de données multispectrale MS-Polyu met en évidence une progression significative des techniques de reconnaissance palmaire au fil des années, tant en termes de descripteurs que de classifieurs. Dès 2019, les réseaux de neurones convolutifs (CNN classiques ou basés sur AlexNet) ont permis d’atteindre des taux de reconnaissance remarquables, supérieurs à 99.99 % sur certains spectres (notamment RED et NIR), bien que certaines longueurs d’onde comme BLEU et GREEN aient été partiellement

omises. Les méthodes basées sur des transformations, comme UDTCWT ([69], 2019), ont également affiché des performances très élevées, confirmant la pertinence des représentations fréquentielles dans le contexte multispectral. Par la suite, des approches plus spécialisées comme MAHCO et HP ([66], 2021) ont contribué à renforcer la robustesse du système sur l’ensemble des spectres, atteignant quasiment les 100 % de précision. Néanmoins, les résultats les plus remarquables proviennent des approches combinant des descripteurs texturaux tels que BSIF26 avec des classifieurs puissants comme k-NN ou LDA/KFA ([67], [70], 2023–2024), atteignant une précision parfaite (100 %) sur tous les spectres, confirmant ainsi l’efficacité des descripteurs binaires appliqués à des images multispectrales. Enfin, la méthode proposée en 2025, intégrant le réseau GoogleNet (InceptionV3) pour l’extraction automatique de caractéristiques couplé à LDA + k-NN, égale ces performances optimales, tout en offrant la robustesse et la généralisation propres à l’apprentissage profond. Ces résultats traduisent une convergence des performances vers l’excellence grâce à l’hybridation entre descripteurs apprenants et classifieurs supervisés.

Référence	Années	Base de données	Descripteur / Classifieur	Taux de reconnaissance 50 %					
				460L	630L	940	850	700	WHT
[69]	2019	MS-CASIA	transform (UDTCWT)	94.45	95.84	97.52	97.80	87.48	/
[68]	2022	MS-CASIA	Filter de gabor +KNN SVM	99.67	99.34	98.16	98.35	99.09	99.5
[67]	2023	MS-CASIA	BSIF26+KNN	100	97	96.33	97	95.33	99
[69]	2023	MS-CASIA	Siamese network(SNN)	95.6%					
[70]	2024	MS-CASIA	BSIF26+KNN (Tantrrings)	100	99.33%	96.33	96	96.67	100
Proposée	2025	MS-CASIA	GoogleNet (InceptionV3) LDA+KNN	99.00%	99.00%	98.67%	98.67%	100.00%	99.33%

Tableau IV.7: Comparaison de notre méthode avec d'autres approches spécifiant le spectre dans la base de MS-CASIA

L’évolution des performances des systèmes de reconnaissance biométrique palmaire sur la base de données multispectrale MS-CASIA met en évidence les avancées significatives en matière de descripteurs, de classifieurs et de techniques d’apprentissage. Les premières approches, telles que celle de [69] en 2019 utilisant la transformée UDTCWT, affichaient des taux de reconnaissance plus modestes, notamment sur certaines longueurs d’onde comme 700 nm

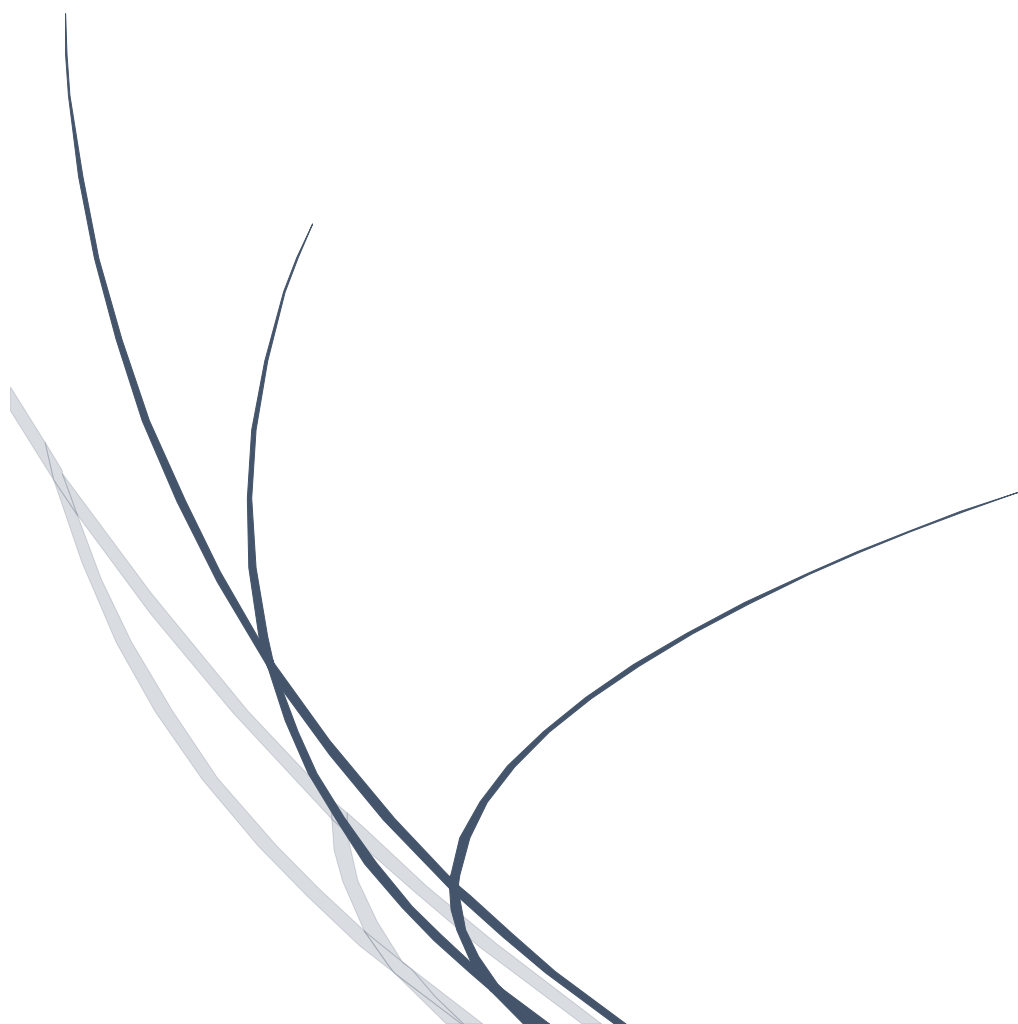
(87.48 %), soulignant les limites des méthodes traditionnelles face à la variabilité spectrale. En revanche, les travaux ultérieurs, à l'instar de [68] en 2022 et [70] en 2024, ont introduit des descripteurs texturaux robustes comme les filtres de Gabor ou les BSIF26, couplés à des classifieurs performants (k-NN, SVM), atteignant des précisions proches ou égales à 100 % sur plusieurs spectres, notamment à 460 nm et WHT. L'introduction de l'apprentissage profond avec les réseaux siamois ([69], 2023) a permis d'améliorer la représentation des caractéristiques, bien que les performances restent légèrement en retrait par rapport aux méthodes hybrides. Enfin, la méthode proposée en 2025, combinant GoogleNet (InceptionV3) pour l'extraction automatique

de caractéristiques avec LDA pour la réduction supervisée de dimension et k-NN pour la classification, offre une solution à la fois efficace et stable, atteignant jusqu'à 100 % de précision sur certains spectres (notamment à 700 nm), tout en maintenant une robustesse élevée sur l'ensemble des longueurs d'onde. Ces résultats illustrent la pertinence des approches hybrides combinant apprentissage profond et méthodes statistiques pour les tâches de reconnaissance multispectrale.

IV.7 Conclusion

En conclusion, l'analyse comparative des différentes approches appliquées à la base de données MS-CASIA démontre une nette amélioration des performances au fil du temps, grâce à l'intégration progressive de descripteurs plus discriminants et de techniques d'apprentissage avancées. Les méthodes classiques ont été progressivement supplantées par des approches hybrides alliant apprentissage profond et classifieurs statistiques supervisés, permettant d'atteindre des taux de reconnaissance très élevés, voire parfaits, sur plusieurs longueurs d'onde. Ces résultats confirment que la combinaison de réseaux convolutifs profonds comme GoogleNet, associés à des techniques de réduction de dimension comme LDA et à des classifieurs robustes comme k-NN, constitue une stratégie efficace et fiable pour la reconnaissance biométrique multispectrale.

CONCLUSION GENERALE ET PERSPECTIVE



Conclusion Générale

Cette étude présente une approche hybride innovante pour la reconnaissance biométrique palmaire à partir d'images multispectrales, appliquée à la base de données MS-CASIA. L'approche proposée combine trois éléments clés : l'extraction automatique de caractéristiques via un réseau préentraîné (GoogleNet / InceptionV3), la réduction de dimension par LDA (Analyse Discriminante Linéaire), et une classification simple mais efficace par 1-NN avec distance cosinus. Cette combinaison a permis d'atteindre des taux de reconnaissance très élevés, dépassant parfois les 99 % voire 100 %, mettant en évidence la pertinence de l'architecture choisie.

Le recours à l'apprentissage par transfert avec GoogleNet, même sur des images en niveaux de gris, a permis d'extraire des descripteurs riches et discriminants. La spécificité multispectrale (longueurs d'onde de 460 nm à WHT) a renforcé la robustesse du système, chaque spectre révélant des aspects physiologiques différents de la paume (épiderme, vaisseaux, pigmentation, etc.). Cette diversité spectrale a amélioré la résistance du système aux variations environnementales et aux tentatives de fraude.

La réduction de dimensionnalité par LDA a joué un rôle crucial en concentrant l'information discriminante dans un espace réduit, tout en améliorant la séparation entre classes. Elle a également permis de limiter les effets du « fléau de la dimensionnalité ». Le classifieur 1-NN avec la distance cosinus, appliqué sur les données réduites, a démontré une efficacité remarquable, adaptée à la structure des descripteurs extraits.

Les résultats expérimentaux obtenus confirment la supériorité des approches multispectrales et hybrides face aux systèmes biométriques classiques basés sur un seul spectre. Ce travail ouvre ainsi la voie à des systèmes plus fiables, généralisables et résistants aux perturbations. Les perspectives incluent l'utilisation d'autres architectures CNN comme ResNet50 ou EfficientNet et AlexNet..... etc., la fusion des spectres en amont de LDA, et l'évaluation sur d'autres bases biométriques ou en conditions réelles pour renforcer la validité et la robustesse de la méthode proposée.

Bibliographie

- [1] Boukhari, W., « Identification Biométrique des Individus par leurs Empreintes Palmaires » Mémoire de Magister, UST Oran, Octobre 2007.
- [2] H. Shao and D. Zhong, «Few shot palmprint recognition via graph neural networks» Electron Lett, vol. 55, no. 16, pp. 890–892, 2019.
- [3].<http://journal.uad.ac.id/index.php/Jiteki> Email: jiteki@ee.uad.ac.id Palm Print Recognition Using Intelligent Techniques.
- [4] Amir BENZAOUI « Identification Biométrique par Descripteurs de Texture Locaux : Application au Visage & Oreill », Thèse de Doctorat ; L'université de Guelma ; 2015
- [5] Mme Nefissa KHIARI-HILI « Biométrie multimodale basée sur l'iris et le visage », thèse de doctorat ; L'UNIVERSITE DE TUNIS EL MANAR et de L'UNIVERSITE PARIS-SACLAY ; 2016
- [6] Abdelatif GHACHOUA, Ibrahim KAHLAOUI, "Reconnaissance de personnes en utilisant L'empreintes Palmaires multispectral basés sur L'apprentissage approfondi" MEMOIRE MASTER ACADEMIQUE, UNIVERSITE KASDI MERBAH OUARGLA, 2016
- [7] Mme. Fatima Zohra AMARA « IDENTIFICATION BIOMÉTRIQUE PAR FUSION MULTIMODALE DE L'EMPREINTE DIGITALE » Mémoire master, CENTRE UNIVERSITAIRE BELHADJ BOUCHAIB D'AÏNTÉMOUCHENT, 2018
- [8].<http://biometrics.over-blog.com/pages/Liris-2019780.html>(2014)
- [9]. <https://fr.wikipedia.org/wiki/Biométrie>.
- [10]. S. Prabhakar, S. Pankanti, and A. K. Jain, "Biometric recognition: security and privacy concerns," IEEE Securite. Priv. Mag., vol. 1, no. 2, 2003.
- [11]. François Lamar doctor de Telecom Sud Paris thèse OCT en phase pour la Reconnaissance Biométrique par empreintes digitales et sa sécurisation. 21 mars 2016.
- [12] ADJAINE Elmechri, BENSLIMAN Abdelkarim « Authentification et Identification biométrique des personnes par les empreintes palmaires » Mémoire MASTER ACADEMIQUE, UNIVERSITE KASDI MERBAH OUARGLA, 2019
- [13].<http://biometrics.over-blog.com/pages/Liris-2019780.html>.
- [14]. [http://www.linternaute.com/science/biologie/dossiers/06/0607-](http://www.linternaute.com/science/biologie/dossiers/06/0607-biometrie/retine.shtml)
- [15]. Fatiha Saidat djemaa-Gueziz. Identification des personnes par l'empreinte de l'articulation des Doigts. Ouargla, Université Kasdi Merbah, 2016.
- [17] [biometrie/retine.shtml](http://biometrie-online.net/technologies/frappe-du-clavier). <https://www.biometrie-online.net/technologies/frappe-du-clavier>.
- [18] <https://www.biometrie-online.net/technologies/voix>
- [19]. Ibtissam, Benchenane. Etude et mise au point d'un procédé biométrique multimodale pour La reconnaissance des individus. Université des Sciences et de la Technologie d'Oran Mohamed Boudiaf, 2016.
- [20] <https://www.biometrie-online.net/technologies/modalites-comparatif> Adn.
- [21] Abes, A. Ben Khalif. Identification d'individus par reconnaissance d'empreintes palmaires. Ouargla, Université Kasdi Merbah, 2008
- [22]. Dehache Ismahèn Thèse, Doctorat en Sciences Approches immunologique pour la reconnaissance des formes Option Intelligence Artificielle Année 2017-2018.

- [23] Mlle. DRIS Chima, Mlle. BOUGHERARA « Khedidja Proposition et Evaluation d'un système biométrique de reconnaissance à base d'empreintes des articulations des doigts. » Mémoire Master, Université Kasdi Merbah -Ouargla,2023
- [24] Commission nationale de l'informatique et des libertés (CNIL), "La biométrie", www.cnil.fr
- [25] International Journal of Advances in Engineering & Technology, Jan 2012.
- [26] Rima Khelif, Asma Saidani « Identification et reconnaissance biométrique par l'utilisation des empreintes palmaires par une approche hiérarchique » mémoire Master, Université de BBA,2021
- [27] Soumia, BENOUAER Aichouche-TAHRINE « Système biométrique basé sur les motifs locaux binaires orientés (LBP⁰) ». OUARGLA : UNIVERSITE KASDI MERBAH OUARGLA, 2016.
- [28] W. BOUKHARRI et M. BENYETOU « Identification Biométrique des Individus par leurs Empreintes Palmaires : Classification par la Méthode des Séparateurs à Vaste Marge (SVM) » Mémoire de magister, Université USTOran, 2007.
- [29] D. SANTOS MARTINE. « biometric recognition based on the texture along palmprint principal lines », thèse de masters, university de porto. July 2011
- [30]. P.Tamrakar & D.Khanna. Occlusion invariant palmprint recognition with ULBP histograms : Procedia Comput.Sci, 2015.
- [31]. AK.Jain. Latent Palmprint Matching : IEEE, 2009.
- [32] I.Hanson. Forensic Archeology-Approaches to Crime Scene Investigation. : NSW Branch Newsletter, 2011.
- [33]. M.A. Al-Garadi & A. Mohamed. A survey of machine and deep learning methods for internet of things (IoT) security: IEEE, 2020.
- [34] M. TAYEB LASKRI et D. CHEFROUR. `Système d'identification de visage humains. 2002
- [35] Abdelatif GHACHOUA et Ibrahim KAHLAOUI « reconnaissance de personnes en utilisant L'empreintes Palmaires multispectral basés sur L'apprentissage profondi », MEMOIRE MASTER ACADEMIQUE ; Université Kasdi Merbah ; 2016
- [36]. biometrie-online.net/technologies/modalites-comparatif#Adn. biometrie online.net. [En ligne] 10.05. 2018. <https://www.biometrie-online.net/technologies/modalites-comparatif#Adn>.
- [37] Bouzidi adel « Système de reconnaissance des empreintes palmaires » mémoire master ; Université Mohamed Khider de Biskra ;2018
- [38]. OUAMANE, Abdelmalik « Reconnaissance Biométrique par Fusion Multimodale du Visage 2D et 3D. biskra » Université Mohamed Khider, 2015. Doctorat en sciences en Electronique
- [39] Ababsa, Guerfi « Authentification d'individus par reconnaissance de caractéristiques biométriques liées aux visages 2D/3D » . Evry : Université d'Evry Val d'Essonne, 2008
- [40] BOUCHER, ALAIN. OUTIL D'AIDE A L'ANNOTATION. 22 janvier 2007.
- [41] Boughaba Mohammed et Boukhris Brahim « L'apprentissage profond (Deep Learning) pour la classification et la recherche d'images par le contenu » Mémoire Master Professionnel, UNIVERSITE KASDI MERBAH OUARGLA,2017
- [42] Kertous Meroua , Rezai Afaf « Classification des empreintes palmaires multispectrales par réseaux de neurones convolutifs » mémoire de master académique; Université Mohammed Seddik Ben Yahia -JIJEL,2021
- [43] Mohamed Abd Elmoumen DJABALLAH « Système de prédiction de la consommation d'énergie basé Deep Learning » mémoire Master, Université de 8 Mai 1945 – Guelma ;2021
- [44] BOUMAZA Amel et MOGDAD Yamina « Emploi de la Technique Deep Learning dans Les Systèmes Mobiles 5G à Onde Millimétrique » MASTER ACADEMIQUE, Université Kasdi Merbah Ouargla,2021

- [45] BENTOUATI SARA, et BENTIBA FELLA, « L'application du Deep Learning dans la classification d'image », Mémoire Master professionnel, Université SAAD DAHLAB de BLIDA, 2021.
- [46] M.R.K.M.A ZACCONE giancarlo, deep learning with tensorflow, 2017
- [47] BOUMAZA Amel et MOGDAD Yamina « Emploi de la Technique Deep Learning dans Les Systèmes Mobiles 5G à Onde Millimétrique » MASTER ACADEMIQUE, Université Kasdi Merbah Ouargla, 2021
- [48] Deluzarche C. Définition | Deep Learning - Apprentissage profond | Futura Tech. Futura n.d. <https://www.futura-sciences.com/tech/definitions/intelligence-artificielle-deep-learning-17262/> (accessed May 11, 2023).
- [49] "Convolutional Neural Networks for Visual Recognition" de Fei-Fei Li, Justin Johnson et Serena Yeung
- [50] Aissougui Iheb et Dahemechi Tayeb Nizar « Modélisation des micromachines/capteurs en utilisant les réseaux de neurones artificiels » mémoire master, Université 8 Mai 1945 – Guelma, 2023
- [51] MAKHLOUF Lazhar « L'apprentissage profond appliqué à la reconnaissance des anomalies mammaires. » mémoire master, Université de 8 Mai 1945 – Guelma, 2022
- [52] BOUMAZA Amel et MOGDAD Yamina « Emploi de la Technique Deep Learning dans Les Systèmes Mobiles 5G à Onde Millimétrique » MASTER ACADEMIQUE, Université Kasdi Merbah Ouargla, 2021
- [53] Serardi Souraya et Benhamouda Mohamed El houcin « Vers un système de détection d'intrusion dans l'Internet des Objets » mémoire master, UNIVERSITE IBN KHALDOUN – TIARET ; 2023
- [54] ZIGHED Narimane « Contrôle de la Qualité et Prédiction de la Maintainabilité Logicielle » thèse de doctorat, Université Badji Mokhtar – Annaba, 2022
- [55] A. FULLICK, " Guide de révision de la biologie du GCSE d'Edexcel International ", 24 février 2015
- [56] Dreyfus G. LES RÉSEAUX DE NEURONES n.d.
- [57] st-m-app-rn.pdf n.d.
- [58] ARTIFICIAL NEURAL NETWORKS. n.d.
- [59] Réseau de neurones artificiels : qu'est-ce que c'est et à quoi ça sert ? n.d. <https://www.lebigdata.fr/reseau-de-neurones-artificiels-definition> (accessed May 12, 2023)
- [60] I. CHETTIBI, et B. ZENIFECHE, « Commande adaptative par Réseaux de Neurones Artificiels d'un Système non Linéaire incertain », Mémoire Master Académique, Université de Jijel, Algérie, 2017/2018.
- [61] Tomas MIKOLOV et al. « Recurrent neural network based language model. " In : Interspeech. T. 2. 3. Makuhari. 2010 .
- [62] <https://www.researchgate.net/figure/The-folded-and-unfolded-structure-of-recurrent-neural-networks-1-RNN-Similar> (visité au 21/05/2023)
- [63] Rafael C GONZALEZ. " Deep convolutional neural networks [Lecture Notes] ". In : IEEE Signal Processing Magazine (2018).
- [64] Ikram Chraïbi Kaadoud, "Apprentissage de séquences et extraction de règles de réseaux récurrents : application au traçage des schémas techniques", thèse de doctorat en informatique, sous la direction de Frédéric Alexandre, Université de Bordeaux, 2018.
- [65] Goodfellow, I., Bengio, Y., & Courville, A. (2016). Deep learning. MIT press.

-
- [66] Vikramaditya JAKKULA. " Tutorial on support vector machine (svm) " In : School of EECS, Washington State University (2006).
 - [67] C. Zhang, P. Patras, and H. Haddadi, "Deep Learning in Mobile and Wireless Networking: A Survey," article IEEE IEEE Communications surveys & tutorials, 2019.
 - [68] Mitchell et al. (2019). *Model Cards for Model Reporting*
 - [69] LeCun, Y., Bengio, Y., & Hinton, G. (2015). *Deep learning*. Nature, 521(7553), 436–444. <https://doi.org/10.1038/nature1453>
 - [70] Larafa Adra, " Reconnaissance multi spectrales des empreintes palmaires basée sur Le descripteur local BSIF Mémoire Master, Université 8 mai 1945 Guelma ;2024.