



**MINISTERE DE L'ENSEIGNEMENT SUPERIEUR
ET DE LA RECHERCHE SCIENTIFIQUE
UNIVERSITE ABDELHAMID IBN BADIS DE MOSTAGANEM**

**Faculté des Sciences Exactes et d'Informatique
Département de Mathématiques et d'Informatique
Filière Informatique**

**MEMOIRE DE FIN D'ETUDES
Pour l'Obtention du Diplôme de Master en Informatique
Option : Ingénierie des Systèmes d'Information**

THEME :

**Utilisation d'une méthode bio-inspirée pour l'analyse
des données médicales**

Etudiante :

✓ **BOURAS Salima**

Encadrant :

✓ **DJAHAFI Fatiha**

*Deuxième Année Master Ingénierie de Systèmes d'Information
Année Universitaire 2016/2017*

Utilisation d'une méthode bio-inspirée pour l'analyse des données Médicales

Résumé :

Un diagnostic médical est un processus de classification et la nécessité d'automatiser ce domaine de recherche est devenu indispensable. L'utilisation des techniques à partir des systèmes naturels apportent une aide non négligeable au diagnostic médical puisqu'ils réduisent les erreurs dues à la fatigue. Le sujet proposé vise à utiliser une méthode bio-inspirée efficace pour bien classer les data set médicale et optimiser les performances de ces méthodes. Nous pensons plus particulièrement aux méthodes qui s'inspirent de la nature déjà connues comme les réseaux de neurones ou les systèmes immunitaires artificiels pour leur efficacité dans plusieurs domaines d'application comme le diagnostic médical.

Mots-clés : classification, apprentissage, méthodes bio-inspirées

..Table Des Matières ..

• Liste des acronymes	i
• Liste des figures	ii
• Liste des tables	iii
• Introduction générale	1
• Chapitre I : Initialisation à la reconnaissance de formes	
1. Introduction	3
2. Le diagnostic médical	3
2.1. La classification médicale	3
3. La reconnaissance de formes	4
3.1. Définition	5
4. Les étapes de la reconnaissance des formes	5
4.1. Prétraitement des données	6
4.2. L'apprentissage	7
4.2.1. Les type d'apprentissage	8
4.3. La classification	11
4.3.1. Les types de classification	11
5. Méthodes de reconnaissance de formes	12
6. Domaines d'application de la reconnaissance de formes	14
4. Conclusion	15
• Chapitre II : les approches bio-inspirées	
1. Introduction	17
2. État de l'art des approches bio-inspirées	17
2.1. Les méta-heuristiques	17
2.1.1. Définition	17
2.1.2. Les caractéristiques des méta-heuristiques	17
2.2. La bio inspiration	18
2.3. Informatique bio-inspirée	18
2.4. le processus d'un modèle inspiration	18
3. Classification de Méthodes bio-inspirées	19
3.1. Algorithmes évolutionnaires	19
3.2. Intelligence en essaim	20
4. Les méthodes bios inspirées	21
4.1. Les réseaux de neurones	21
4.1.1. Propriétés des Réseaux de Neurones	22
4.1.2. Structures d'interconnexion	22
4.1.3. Architecture des Réseaux de Neurones	24

4.2. L'algorithme génétique	25
4.3. La colonie de fourmis	27
4.4. La colonie d'abeille artificielle	28
4.5. L'optimisation par les essaims de particules	29
4.6. Système immunitaire artificiel	30
5. Domaine d'application des méthodes bio-inspirées	30
6. Quelques travaux réalisés des méthodes bio-inspirées	32
7. Conclusion	33
• Chapitre IV : Méthode appliqué	
1. Introduction	35
2. Neurone et le réseau de neurones.	35
3. Le modèle MLP (perceptron multicouche).	36
3.1. Définition.	36
3.2. Architecture du MLP.	36
3.3. Apprentissage du MLP.	37
3.4. Calcul de fonction de l'activation.	39
3.5. Algorithme de la rétro propagation des erreurs.	40
4. Algorithme de classification de MLP.	42
5. Caractéristique du MLP.	42
6. Avantages et Inconvénients du MLP.	42
7. Domaine d'application de MLP.	43
8. Conclusion.	44
• Chapitre IV : Implémentation	
1. Introduction	45
2. la description de la base de données	45
3. Les taux de classification	4
4. Environnement de développement	46
4.1. Les ressources physiques	47
4.2. La ressource logicielle	47
5. Description de travail réalisé	47
6. Résultats	52
6.1. Etude de comparaison entre les méthodes	53
7. Conclusion	54
• Conclusion générale	55
• Références	56

Introduction générale

La Nature est naturellement une grande et immense source d'inspiration pour résoudre des problèmes difficiles et complexes en informatique, car elle présente un phénomène extrêmement diversifié, dynamique, robuste, complexe et fascinant. Les algorithmes inspirés de la nature sont méta heuristiques qui imite la nature pour résoudre les problèmes d'optimisation d'ouvrir une nouvelle ère dans le calcul. Pour les dernières décennies, de nombreux efforts de recherche a été concentré dans ce domaine particulier. Toujours jeune et les résultats étant très étonnants, élargit la portée et la viabilité des algorithmes bio-inspirées explorant de nouveaux domaines d'application et plus d'opportunités en informatique.

L'informatique bio inspirée est la discipline qui consiste à étudier les systèmes naturels afin d'y trouver des solutions pour des problèmes informatiques qui ne peuvent être résolus par les méthodes classiques. En effet, les chercheurs ont noté beaucoup de similitudes entre les problèmes que rencontrent les systèmes naturels et les problèmes informatiques.

La vraie beauté des algorithmes inspirés de la nature réside dans le fait qu'elle reçoit son inspiration unique de la nature. Ils ont la capacité de décrire et de résoudre des relations complexes à partir de conditions initiales intrinsèquement très simples et des règles avec peu ou pas de connaissance de l'espace de recherche nature est l'exemple parfait pour l'optimisation, car si nous examinons de près chaque fonctionnalité ou phénomène dans la nature, il faut trouver la stratégie optimale, tout en abordant les interactions complexes entre les organismes allant des microorganismes aux êtres humains à part entière, en équilibrant l'écosystème, en maintenant la diversité, l'adaptation, Même si la stratégie derrière, la solution est simple et les résultats sont étonnants.

Nature est le meilleur professeur et ses conceptions et ses capacités sont extrêmement énormes et mystérieuses que les chercheurs essaient d'imiter la nature dans la technologie. Aussi les deux domaines ont une connexion beaucoup plus forte depuis, il semble tout à fait raisonnable que les problèmes nouveaux ou persistants en informatique pourrait avoir beaucoup en commun avec les problèmes nature a rencontré et résolu il ya longtemps. Ainsi, une cartographie facile est possible entre la nature et la technologie.

L'informatique inspirée de la biologie est devenue une nouvelle ère de l'informatique qui englobe un large éventail d'applications, couvrant la plupart des domaines tels que les réseaux informatiques, la sécurité, la robotique, le génie biomédical, les systèmes de contrôle, le traitement parallèle, l'exploration de données, les systèmes d'alimentation, l'ingénierie de production et beaucoup plus.

D'après la nature on a plusieurs méthodes bio-inspirés ou inspirées de la nature ont été alors utilisées pour résoudre le problème posé. Depuis quelques années, les chercheurs informaticiens ont trouvé dans le monde naturel, parmi ces méthodes il ya les réseaux de neurones.

Durant ces dernières années, les réseaux de neurones artificiels ont été utilisés pour résoudre un bon nombre de problèmes. La majorité des problèmes sont liés à la classification. Cette utilisation est souvent justifiée par leur capacité à résoudre des problèmes difficiles sans avoir une connaissance détaillée de la méthode de résolution.

Introduction générale

Dans ce contexte, les réseaux de neurones multicouches (MLP) sont utilisés intensivement comme des boîtes noires pour répondre à des questions liés à l'intelligence artificielle et à la reconnaissance des formes. Même si ces réseaux donnent des résultats satisfaisants pour un nombre conséquent d'applications, des exceptions existent toujours vu le caractère non déterministe des réseaux MLP ainsi que le problème des minima locaux souvent rencontré.

Par voie de conséquence, ce rapport comporte deux chapitres :

Dans le premier chapitre nous présentons une initiation de la reconnaissance des formes, l'apprentissage et la classification.

Le deuxième chapitre sera consacré à l'exposé de l'état de l'art des différentes approches bio-inspirées en faisant appel aux techniques d'optimisation ainsi que leurs domaines d'application.

Le troisième chapitre explique le concept de méthode proposée de l'approche hybride entre les réseaux de neurone MLP.

Le dernier chapitre présente la méthode qui il appliqué sur la base de données de cancer WDBC. Ainsi que la discussion des différents résultats obtenus d'implémentation.

Enfin nous terminerons notre travail par une conclusion générale.

LISTE DES ACRONYMES

SVM : Support Vector Machine.

RNA : Les réseaux de neurones artificiels.

AG : Algorithmes génétiques.

ACO : Ant Colony Optimization.

ABC : Artificial Bee Colony

AE : Les abeilles employées.

AS : Des abeilles spectateur.

PSO : L'optimisation par les essaims de particules.

SIA : Un système immunitaire artificiel

LISTE DES TABLES

Table II.1 : Modélisation des Réseaux de Neurones.	22
Table II.2 : Quelques domaines d'application.	32
Table II.3 : Un historique de différents travaux.	34
Table III.1 : Avantages et Inconvénients du MLP.	43
Table IV.1 : Les attributs de la base Wisconsin.	45
Table IV.2 : Les résultats.	53

LISTE DES FIGURES

Figure I.1 : Les étapes principales d'une approche de reconnaissance des formes	6
Figure I.2 : schéma exilique la classification supervisée	11
Figure I.3: schéma exilique la classification non supervisée	12
Figure I.4 : Reconnaissance d'empreintes	15
Figure I.5: reconnaissance de la parole	15
Figure I.6 : reconnaissance caractères manuscrits	15
Figure I.7 : reconnaissance de visage	15
Figure II.1: Processus d'inspiration d'un phénomène naturel	19
Figure II.2: les différentes méthodes bio-inspirée.	20
Figure II.3: L'architecture de neurone biologique et artificiel.	21
Figure II.4 : Définition des couches d'un réseau multicouche.	23
Figure II.5 : Réseau à connexion locales.	23
Figure II.6 : Réseau à connexion récurrents.	23
Figure II.7 : Réseau à connexion complète	24
Figure II.8 : Réseau de neurones non bouclés.	24
Figure. II.9 : Réseau de neurones bouclés .	25
Figure II.10 : l'organigramme de fonctionnement d'un algorithme génétique.	26
Figure II.11: Expérience de sélection des branches les plus courtes par une colonie de fourmis.	27
Figure II.12 : l'organigramme de fonctionnement de l'optimisation par les essaims de particules	29
Figure III.1 : Neurone artificiel.	35
Figure III. 2: Différents types de fonctions de transfert pour le neurone artificiel.	36
Figure III.3 : Perceptrons multicouches.	37
Figure III.4 : i le neurone amont, j le neurone aval et w_{ij} le poids de la connexion.	37
Figure III.5: Algorithme de l'apprentissage du MLP.	38
Figure III.6 : Organigramme de la rétro propagation des erreurs.	41

Figure IV.1 : la page d'accueil	48
Figure IV.2 : Fenêtre principale	48
Figure IV.3 : La fenêtre pour le chargement de la base de données	49
Figure IV.4: Base de données chargée	49
Figure IV.5 : Paramètres de MLP	50
Figure IV.6 : Les deux phases d'apprentissage et test.	51
Figure IV.7: La fenêtre e lancement d'apprentissage.	51
Figure IV.8 : Affichage de graphe d'erreur global.	52
Figure IV.9 : L'affichage de résultat.	52
Figure .IV.10 : Diagramme de comparaison entre MLP et Clonclass.	54

I.1 Introduction :

L'un des principaux objectifs de l'ordinateur est d'automatiser les tâches habituellement effectuées par l'homme. La plupart des activités humaines reposent sur une faculté importante de notre cerveau : pouvoir distinguer entre les formes. Pour pouvoir automatiser ce genre de tâches, l'ordinateur doit obligatoirement acquérir la faculté de reconnaître les formes.

La reconnaissance des formes est la discipline qui tente d'automatiser les mécanismes de reconnaissance utilisés par les êtres vivants en général, et par les êtres humains en particulier. Malgré plusieurs décennies de recherche, aucune solution générale n'a pu être développée. Il est devenu évident que les performances optimales ne peuvent être atteintes qu'en construisant des solutions spécifiques. Dans ce chapitre nous allons présenter des concepts fondamentaux de la reconnaissance de formes ainsi que son domaine d'application.

Pour appréhender et comprendre son environnement l'Homme a toujours cherché à classer, et les médecins n'ont pas échappé à cette démarche.

Mais aujourd'hui, au-delà de l'intérêt scientifique pour les choses classées, les enjeux liés à la santé publique, l'économie de la santé et à la gestion du système sanitaire donnent un relief particulier aux classifications médicales.

I.2 Le diagnostic médical :

L'histoire de la Classification Internationale des Maladies commence en 1853 avec William Farr, et se poursuit avec Jacques Bertillon en 1891, chargé par l'International Statistical Institute de préparer une "Nomenclature internationale des causes de Décès". A partir de 1898, le principe de la Révision décennale est adopté... et les révisions se succèdent sur un siècle. Lors de sa 6ème révision en 1948, cette Nomenclature devient "Classification Internationale des Maladies", sous l'égide de l'OMS. La Neuvième révision rentre en vigueur en 1975 et la Dixième en 1995[1].

I.2.1. La classification médicale :

Le diagnostic médical est un élément très important dans le domaine de reconnaissance et traitement des maladies. La bio puce est l'une des techniques modernes qui nous aide à faire le diagnostic.

L'informatique médicale est l'application des techniques issues de l'informatique au domaine médical. L'informatique médicale est une science à part entière ; aux confluent

des sciences de l'information et de la médecine, c'est aussi l'une des technologies nécessaire au développement de l'E-médecine.

Elle permet d'affiner et d'accélérer ou automatiser certains moyens d'investigation médicale et de diagnostic. Elle apporte de nouveaux mécanismes et moyens d'interprétation et de raisonnement médical.

Mais aujourd'hui, au-delà de l'intérêt scientifique pour les choses classées, les enjeux liés à la santé publique, l'économie de la santé et à la gestion du système sanitaire donnent un relief particulier aux classifications médicales.

A côté de ces grands systèmes de classifications médicales, il existe un certain nombre de classifications pour des domaines limités des sciences médicales, que nous ne mentionnons que brièvement ici :

Cancer : - Classification de Thomson, surtout destinée aux tirés à part (Pickett-Thomson Research Laboratory Fulmer, Stough, Bucks (Angleterre), 1946).

Cardiologie : - Classification très détaillée de Lepeschkin (Université du Vermont) fondée sur la Standard nomenclature of diseases, 1956.

Chirurgie : - Classification de E. Wanderley à l'Université de Recife (Brésil), 1946.

Dentisterie : - Révision par Black de la Classification décimale de Dewey, 1955.

Épilepsie : - Classification internationale des crises d'épilepsie, proposée par la Ligue internationale contre l'épilepsie, in : Epilepsia, 1964.

Maladie de Chagas : - Classification très spécialisée de Medeiros-Valeri-Palavra du Service de prophylaxie du paludisme, à São Paulo (Brésil).

Médecine aéronautique : - Classification de Gallohin (Centre d'étude de biologie aéronautique du Service de santé de l'Air), 1950.

Paralysie cérébrale : - Schéma préparé par S. J. Teague (England's Center for Spastic Children), 1956.

Pharmacie : - Classification pharmaceutique élaborée par les Laboratoires Eli Lilly, Indianapolis, 1915.

Psychiatrie : - Classification préparée par J.S.A. Miller pour les hôpitaux psychiatriques à Rockland State Hospital Library, Orangeburg (New York).

Psychologie : - Classification de Loutitt, 1941, à l'Université d'Indiana.

Santé publique : - Classification préparée par la « National Health Library » (National Health Council) [2].

I.3 La reconnaissance de formes :

I.3.1 Définition :

La reconnaissance de formes (RDF) est de classer les nouvelles formes en utilisant un classifieur qui génère une fonction d'appartenance pour chaque classe. Ainsi la classification d'un nouveau point peut se faire en fonction de la valeur d'appartenance qu'il obtient par rapport à chaque classe [3].

On désigne par reconnaissance de formes (ou parfois reconnaissance de motifs) un ensemble de techniques et méthodes visant à identifier des motifs à partir de données brutes afin de prendre une décision dépendant de la catégorie attribuée à ce motif. On considère que c'est une branche de l'intelligence artificielle qui fait largement appel aux techniques d'apprentissage automatique et aux statistiques.

La reconnaissance de formes est réalisée en deux phases : l'apprentissage à partir des données connues et la classification des nouvelles données. En amont de ces deux phases, une étape de prétraitement est utilisée pour trouver l'ensemble minimal de paramètres informatifs nécessaire à l'établissement de l'espace de représentation [3].

Plusieurs méthodes et algorithmes ont été développés pour chacun des modules de reconnaissance. Chaque méthode étant plus adaptée à un certain type de problèmes. Au fur et à mesure que la puissance des machines croît, de nouvelles applications pour la reconnaissance voient le jour, pour lesquels il faut souvent développer de nouvelles solutions. La complexité de certains problèmes est telle que, les méthodes algorithmiques standards ne sont pas applicables.

I.4. Les étapes de la reconnaissance des formes :

Durant les dernières décennies, système de reconnaissance de formes peut être divisé en trois modules qui sont la préparation des données, l'apprentissage, la classification et traitement. Pour résoudre un problème de reconnaissance revient à construire un système qui donne les meilleures performances pour l'application ciblée. Cela implique l'optimisation de chacun Les modules du système pour donner les meilleures solutions possibles.

Les principales phases de RDF sont présentées dans la figure ci-dessous :

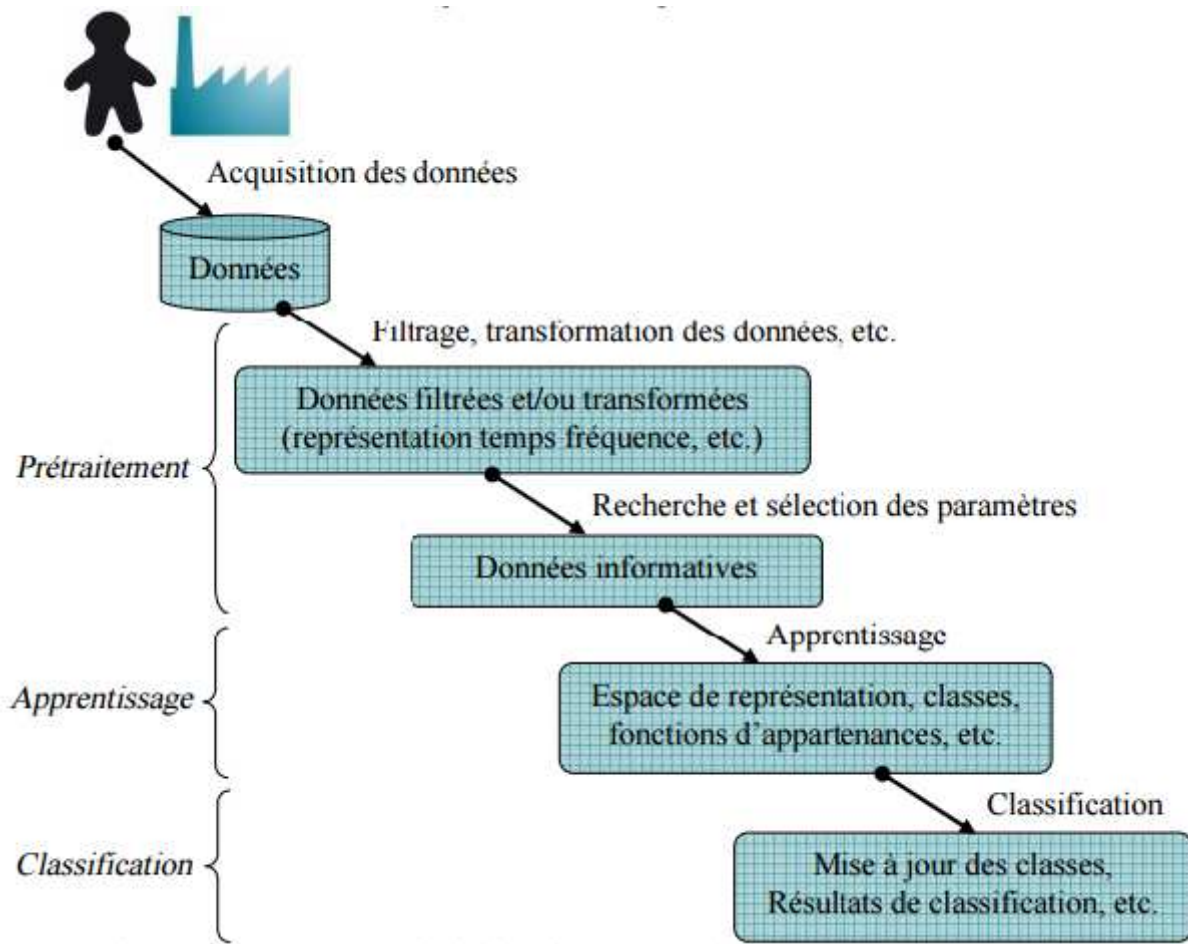


Figure I.1. Les étapes principales d'une approche de reconnaissance des formes [4].

I.4.1. Prétraitement des données :

La première phase de prétraitement est une étape essentielle du processus d'extraction des connaissances à partir des données. Elle permet d'extraire l'information utile des données.

La phase du prétraitement concerne le filtrage et la transformation des données qui permettent d'obtenir des données plus adaptées à la recherche de caractéristiques informatives.

Consiste à sélectionner dans l'espace de représentation l'information nécessaire au domaine d'application. Cette sélection passe souvent par l'élimination du bruit, la normalisation des données, ainsi que par la suppression de la redondance.

L'objectif des prétraitements est de faciliter la caractérisation de la forme (caractère, chiffre, mot) ou de l'entité à reconnaître soit en nettoyant la forme (élimination du bruit) ou en réduisant la quantité d'information à traiter pour ne garder que les informations les plus significatives.

La phase de prétraitement dépend fortement du domaine d'application, par exemple pour les systèmes de reconnaissance de texte les opérations de prétraitement les plus usuelles sont la correction de l'inclinaison du texte, l'amélioration de la qualité (élimination du bruit), ainsi que la correction de l'orientation du texte. Vient ensuite la phase de calcul des représentations qui va se charger de représenter les données sous la forme d'ensembles d'attributs. Souvent, d'autres traitements sont nécessaires avant la phase de calcul des représentations. Par exemple, pour les systèmes de reconnaissance de texte, la segmentation est une étape cruciale de la phase de préparation des données [4].

I.4.2. L'apprentissage :

L'apprentissage (machine learning en anglais), fait référence au développement, à l'analyse et à l'implémentation de méthodes qui permettent à une machine d'évoluer, dans la résolution d'une catégorie de problèmes, grâce à un processus d'apprentissage, et ainsi de remplir des tâches qu'il est difficile ou impossible de remplir par des moyens algorithmiques plus classiques. L'apprentissage est une problématique ancienne qui a été abordée dans de nombreux domaines dont la Psychologie, la Statistique, la Didactique, l'Intelligence Artificielle,...

Le terme apprentissage correspond à une caractéristique des machines. Il s'agit plus précisément de leur capacité à organiser, construire et généraliser des connaissances pour une utilisation ultérieure, de leur capacité de tirer profit de l'expérience pour améliorer la résolution d'un problème.

L'apprentissage est utilisé pour remplacer l'acquisition des règles de classification auprès d'un expert du domaine, il permet ainsi de contourner la difficile tâche de l'acquisition des connaissances. Une règle de classification permet de prédire la classe d'un nouvel exemple, la classe étant une variable discrète d'une importance particulière. A titre d'illustration, un exemple peut correspondre à un patient décrit par un ensemble

d'attributs comme l'âge, le sexe, la pression sanguine, ... et la classe pourrait être un attribut binaire concluant ou non à l'affectation du patient par une maladie particulière. L'apprentissage automatique consisterait ici à « apprendre » des règles de classification à partir d'un ensemble de descriptions de patients.

Nous nous intéressons à un type d'apprentissage particulier qui est l'apprentissage inductif encore appelé apprentissage à partir d'exemples ou d'observations. Le processus de l'apprentissage inductif peut être considéré comme une recherche de descriptions générales plausibles qui expliquent les données d'entrée et qui sont utiles pour en prédire de nouvelles (résultats, données de sortie, etc.) et autres termes, l'inférence inductive essaie de dériver une description complète et correcte d'un phénomène donné à partir d'observations spécifiques de ce phénomène ou de parties de lui.

L'apprentissage va permettre d'optimiser les paramètres du classificateur pour le problème à résoudre, en utilisant des données d'entraînement. Lorsque les données d'entraînement sont préalablement classées.

L'apprentissage peut être utilisé pour doter un ordinateur de la perception de son environnement, comme par exemple la reconnaissance d'objets (visages, schémas, formes, etc.

La capacité d'apprentissage est une caractéristique des êtres vivants. De la naissance à l'âge adulte, les êtres vivants acquièrent de nombreuses capacités qui leur permettent de survivre dans leur environnement. L'apprentissage d'un langage, de l'écriture et de la lecture sont de bons exemples des capacités humaines, et des phénomènes mis en jeu : apprentissage par cœur, apprentissage supervisé par d'autres êtres humains [5].

Donc, l'apprentissage automatique pour la machine est qu'avec l'ensemble de tâches que la machine doit réaliser, elle utilise l'ensemble d'expériences telles que sa performance sur est améliorée, et les principaux domaines d'applications de l'apprentissage automatique sont les fouilles de données et l'intelligence artificielle. Extraire et exploiter automatiquement l'information présente dans un jeu de données, et de résoudre automatiquement des problèmes complexes par la prise de décisions sur la base des échantillons de ces problème.

I.4.2.1 Les type d'apprentissage :

Il existe trois grands types d'apprentissage : l'apprentissage supervisé, l'apprentissage non supervisé et l'apprentissage semi-supervisé.

L'apprentissage supervisé :

L'apprentissage supervisé nécessite une base d'apprentissage où chaque instance est préalablement étiquetée avec sa classe respective. Il nécessite que les observations soient étiquetées ; on connaît a priori le mode de fonctionnement associé à chacune de ces observations. Supposons que nous disposions d'une grande quantité de données non étiquetées. Un opérateur étiquette manuellement autant de données que possible pour les utiliser dans un algorithme d'apprentissage supervisé, dans le but d'apprendre un classifieur. Notez ici que les données d'apprentissage à étiqueter manuellement sont choisies préalablement et au hasard parmi les données disponibles. On dit alors que l'apprentissage se fait de façon passive. Et comme des méthodes d'apprentissage on a Support Vector Machine (SVM), régression logistique et réseaux de neurones, etc [5].

par exemple : en fonction de points communs détectés avec les symptômes d'autres patients connus, le système peut catégoriser de nouveaux patients au vu de leurs analyses médicales en risque estimé (probabilité) de développer telle ou telle maladie [5].

▪ L'apprentissage non-supervisé :

L'apprentissage non supervisé ne nécessite pas l'étiquetage des données de la base d'apprentissage, il cherche à former des classes dans un ensemble d'observations non étiquetées pour lesquelles aucune connaissance n'est disponible sur leur appartenance aux différentes classes. Ce type d'apprentissage peut être utilisé pour découvrir des clusters formés par l'ensemble des données. Il doit découvrir par lui-même la structure des données, Le système doit ici dans l'espace de description (la somme des données) cibler les données selon leurs attributs disponibles, pour les classer en groupe homogènes d'exemples. Cette forme d'apprentissage travaille avec un ensemble de données non annotées, l'objectif étant de permettre une extraction de connaissance organisée à partir de ces données.

par exemple : Un épidémiologiste pourrait par exemple dans un ensemble assez large de victimes de cancers du foie tenter de faire émerger des hypothèses explicatives, l'ordinateur pourrait différencier différents groupes, qu'on pourrait ensuite associer par exemple à leur provenance géographique, génétique, à

l'alcoolisme ou à l'exposition à un métal lourd ou à une toxine telle que l'aflatoxine [5].

▪ **L'apprentissage semi supervisé :**

Il s'agit d'une d'apprentissage a mi-chemin entre supervisé, qui utilise un ensemble de donnes dont seulement une partie (typiquement faible) est annotée. Dans l'apprentissage semi-supervisé, l'idée de base est de réduire le coût d'étiquetage en utilisant peu de données étiquetées et en exploitant les données non étiquetées. À titre d'exemple, une méthode d'apprentissage semi-supervisée très basique consiste à apprendre un modèle de classification à partir des données étiquetées et prédire les classes des données non étiquetées. Les données non étiquetées les plus sûres (pour lesquelles le classifieur est le moins incertain), avec leurs étiquettes prédites, sont alors ajoutées à la base d'apprentissage pour apprendre un nouveau modèle de classification [7].

Il utilise l'ensemble d'apprentissage étiqueté pour estimer les fonctions d'appartenance de chaque classe connue et l'apprentissage non supervisé pour améliorer la précision de cette estimation, détecter les nouvelles classes et apprendre leurs fonctions d'appartenance. Cette combinaison est particulièrement efficace quand l'ensemble d'apprentissage étiqueté contient peu de points de chaque classe. Cela est souvent le cas pour les classes représentant les modes de fonctionnement défaillants [7].

De plus, les modes de fonctionnement, surtout défaillants, sont rarement tous représentés dans l'ensemble d'apprentissage. En effet pour des raisons de coût ou de sécurité, certains modes de fonctionnement défaillants ou dangereux ne peuvent pas être provoqués et sont donc absents de l'ensemble d'apprentissage. Un apprentissage non supervisé permet donc de détecter et d'apprendre les nouveaux modes de fonctionnement au fur et à mesure qu'ils apparaissent afin d'enrichir la connaissance initiale. L'apprentissage s'applique à un grand nombre d'activités humaines et convient en particulier au problème de la prise de décision automatisée, il s'agira comme ces exemples :

- D'établir un diagnostic médical à partir de la description clinique d'un patient.
- De donner une réponse à la demande de prêt bancaire de la part d'un client sur la base de sa situation personnelle.

- De déclencher un processus d'alerte en fonction de signaux reçus par des capteurs.
- De la reconnaissance des formes.
- De la reconnaissance de la parole et du texte écrit.
- De contrôler un processus et de diagnostiquer des pannes [5].

I.4.3. La classification :

En analyse de données, la classification consiste à regrouper des ensembles d'exemples (souvent, de manière non-supervisée) en classes. Ces classes sont généralement organisées en une structure (groupes ou grappes).

La classification automatique est la catégorisation algorithmique d'objets. Elle consiste à attribuer une classe ou catégorie à chaque objet (ou individu) à classer, en se basant sur des données statistiques. Elle fait couramment appel à l'apprentissage automatique et est largement utilisée en reconnaissance de formes.

I.4.3.1 Les type de classification :

- **Classification supervisée :**
 - On dispose d'éléments déjà classés Exemple : articles en rubrique économie, politique, sport, culture...
 - On veut classer un nouvel élément Exemple : lui attribuer une étiquette parmi économie, politique, sport, culture... [8].

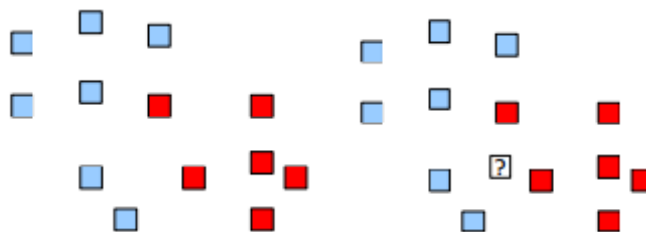


Figure I.2 : schéma exilique la classification supervisée [8].

- **Classification non supervisée :**

- On dispose d'éléments non classés Exemple : mots d'un texte
- On veut les regrouper en classes Exemple : si deux mots ont la même étiquette, ils sont en rapport avec une même thématique... [8].



Figure I.3: schéma exilique la classification non supervisée [8].

La classification non-supervisée est utilisée lorsque que l'on possède des documents qui ne sont pas classés et dont on ne connaît pas de classification. A la fin du processus de classification non-supervisée, les documents doivent appartenir à l'une des classes générées par la classification. On distingue deux catégories de classifications non-supervisées : hiérarchiques et non-hiérarchiques.

- Classification hiérarchique : Les groupes que l'on crée sont emboîtés les uns dans les autres de façon hiérarchique.
- Classification non hiérarchique : En classification non hiérarchique, on ne cherche pas à structurer de manière hiérarchique les sous-ensembles entre-eux [8].

❖ **Les problèmes de classification :**

Il y a deux grandes catégories de problèmes de classification. Si, à chaque observation est associée une classe a priori et que l'objectif de la classification est de respecter, au mieux, ces classes a priori alors nous sommes dans un problème de discrimination ou de classification supervisée ou de l'apprentissage avec professeur. S'il n'y a pas de classification a priori et que l'objectif de ce classement est de regrouper ces individus dans des classes homogènes en fonction de l'ensemble de variables sélectionnées. Ce type de problème est un problème de classification

automatique ou de classification non supervisée ou bien d'apprentissage sans professeur [9].

❖ Les critères classification :

L'objectif principal des techniques de classification est de trouver une partition où les objets d'une classe devraient être semblables (entre eux), les objets de différentes classes devraient être différents, une bonne classification devrait accomplir différents critères :

✓ **Validité** : Elle peut se définir par :

Chaque classe d'une partition doit être homogène : Les objets qui appartiennent à la même classe doivent être semblables.

Les classes doivent être isolées entre elles : Les objets de différentes classes doivent être différents.

La classification doit s'adapter aux données : La classification doit pouvoir expliquer la variation des données.

✓ **Interprétabilité** : Les classes doivent avoir une interprétation substantive c'est-à-dire qu'il est possible de donner des noms aux classes, dans le meilleur des cas les noms doivent correspondre aux types déduits d'une certaine théorie.

✓ **Stabilité** : Les classes doivent être stables ça veut dire la petite modification dans les données et dans les méthodes ne doivent pas changer les résultats.

Parfois la taille et le nombre de classes sont employés en tant que critères additionnels : le nombre de classes doit être aussi petit que possible, et la taille des classes ne doit pas être trop petite [10].

I.5. Méthodes de reconnaissance de formes :

Parmi les méthodes de la reconnaissance de formes nous pouvons citer :

✓ Mise en correspondance de graphes : application aux réseaux routiers extraits d'un couple carte/image

.

- ✓ Méthode bayésienne : est une méthode d'inférence permettant de déduire la probabilité d'un événement à partir de celles d'autres événements déjà évalués. Elle s'appuie principalement sur le théorème de Bayes.
- ✓ Classifieur linéaire : le terme de classifieur linéaire représente une famille d'algorithmes de classement statistique. Le rôle d'un classifieur est de classer dans des groupes (des classes) les échantillons qui ont des propriétés similaires, mesurées sur des observations
- ✓ Réseau de neurones : est un ensemble d'algorithmes dont la conception est à l'origine très schématiquement inspirée du fonctionnement des neurones biologiques, et qui par la suite s'est rapproché des méthodes statistiques
- ✓ Support vector machine (SVM) : est un ensemble de techniques d'apprentissage supervisé destinées à résoudre des problèmes de discrimination et de régression. Les SVM sont une généralisation des classifieurs linéaires.
- ✓ Un algorithme bien connu pour la détection de formes, la transformée de Hough, est une méthode d'estimation paramétrique [10].

I.6. Domaines d'application de la reconnaissance de formes :

La reconnaissance de formes est utilisée dans plusieurs domaines d'activités. Parmi ces domaines nous pouvons citer :

- Recherche d'images par le contenu : Cette technologie est actuellement intéressante pour la recherche de données sur l'imagerie médicale, ou cartographiques.
- Classification de documents : Son activité est essentielle dans de nombreux domaines économiques : elle permet d'organiser des corpus documentaires, de les trier, et d'aider à les exploiter dans des secteurs tels que l'administration, l'aéronautique, la recherche sur internet, les sciences.
- Reconnaissance de l'écriture manuscrite : Cette technologie fait appel à la reconnaissance de forme, mais également au traitement automatique du langage naturel. Cela veut dire que le système, tout comme le cerveau humain, reconnaît des mots et des phrases existant dans un langage connu plutôt qu'une succession de caractères. Ceci améliore grandement la robustesse [11].

La reconnaissance des formes est un domaine très vaste dans l'intelligence artificiel, dont on peut trouver la reconnaissance des visages, les empreintes, la parole, l'écriture et autres domaines qui ne sont pas moins important de ce qu'on a citer .

Comme dans les figures suivantes :



Figure I.4: Reconnaissance d'empreintes [12].

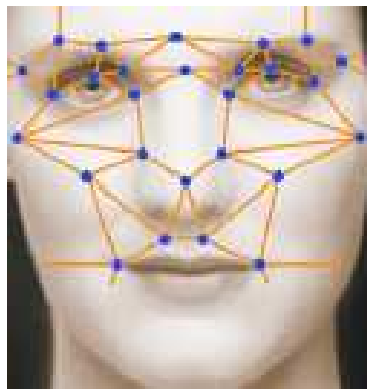


Figure I.5: reconnaissance de visage [12].



Figure I.6 : reconnaissance de la parole [12].



Figure I.7 : reconnaissance caractères manuscrits [12].

I.7. Conclusion :

De manière générale, dire que les systèmes de reconnaissance des formes sont des fonctions de classification qui associent à toute forme inconnue sa classe la plus probable. Malgré la difficulté de la tâche, les systèmes de reconnaissance des formes ont connus un grand succès dans des domaines variés.

Un système de reconnaissance des formes contient principalement quatre modules. Un module de préparation de données qui se charge de faciliter la tâche de reconnaissance en représentant les données en entrée sous une forme adaptée à la classification. Un module d'apprentissage qui va optimiser la fonction de classification pour le domaine d'application ciblé. Un module de classification qui va classer les formes inconnues et un module de post traitement qui va corriger les résultats du classificateur en utilisant des mécanismes spécifiques au domaine d'application.

II.1. Introduction :

Il est de bon ton, de nos jours, de chercher dans la nature la solution à tous nos problèmes. C'est faire fi de siècles de luttes où la nature était un ennemi implacable contre lequel il fallait se défendre, qu'il fallait dominer et, in fine, exploiter. À quelques remarquables exceptions près, comme les travaux de Léonard de Vinci dont les réussites sont plus intellectuelles que pratiques.

Dans ce chapitre nous essayons de présenter les différentes approches bio-inspirées ainsi que ses domaines d'application.

II.2. État de l'art des approches bio-inspirées :

Si les méthodes de résolution exactes permettent d'obtenir une solution dont l'optimalité est garantie, dans certaines situations, on peut cependant chercher des solutions de bonne qualité, sans garantie d'optimalité, mais au profit d'un temps de calcul plus réduit. Pour cela, On applique des méthodes appelées méta-heuristiques, adaptées à chaque problème traité, avec cependant l'inconvénient de ne disposer en retour d'aucune information sur la qualité des solutions obtenues.

II.2.1. Les méta-heuristiques :

II.2.1.1. Définition :

Le terme méta-heuristique vient des mots grecs meta (au delà) et heuriskein (trouver). Il n'y a pas clairement de consensus sur la définition exacte des heuristiques et des métas-heuristiques. Nous allons adopter celles-ci :

- Une heuristique est une technique de résolution spécialisée à un problème. Elle ne garantit pas la qualité de la solution obtenue.
- Une méta-heuristique est une heuristique générique qu'il faut adapter à chaque problème [13].

Les méta-heuristiques sont des algorithmes généraux d'optimisation applicables à une grande variété des problèmes. Elles sont apparues dans le but de résoudre au mieux des problèmes d'optimisation, elles sont généralement inspirées de phénomènes naturels : de la biologie (Algorithme génétiques, système immunitaire, etc.), l'éthologie (colonie de fourmis, colonie d'abeilles, etc.).

Concernant les algorithmes des fourmis, ils s'inspirent respectivement de la théorie de l'évolution et du comportement de fourmis à la recherche de nourriture.

II.2.1.2. les caractéristiques des méta-heuristiques :

Les propriétés fondamentales des méta-heuristiques sont comme suit :

- Les méta-heuristiques sont des stratégies qui permettent de guider la recherche d'une solution optimale.
- Le but visé par les méta-heuristiques est d'explorer l'espace de recherche efficacement afin de déterminer des solutions (presque) optimales.
- Les techniques qui constituent des algorithmes de type méta-heuristique vont de la simple procédure de recherche locale à des processus d'apprentissage complexes.

- Les méta-heuristiques sont en général non-déterministes et ne donnent aucune garantie d'optimalité.
- Les méta-heuristiques peuvent contenir des mécanismes qui permettent d'éviter d'être bloqué dans des régions de l'espace de recherche.
- Les concepts de base des méta-heuristiques peuvent être écrits de manière abstraite, sans faire appel à un problème spécifique.
- Les méta-heuristiques peuvent faire appel à des heuristiques qui tiennent compte de la spécificité du problème traité, mais ces heuristiques sont contrôlées par une stratégie de niveau supérieur.
- Les méta-heuristiques peuvent faire usage de l'expérience accumulée durant la recherche de l'optimum, pour mieux guider la suite du processus de recherche [14].

II.2.2. La bio inspiration :

La bio-inspiration illustre dans l'examen de la façon dont les structures et les fonctions des organismes vivants ont inspiré de nouvelles formes de technologie. Les structures et les processus biologiques ont profité de milliards d'années de raffinement durant le processus d'évolution de la vie sur Terre. Cette rencontre a rassemblé des scientifiques de diverses disciplines pour discuter de la façon dont la nature a inspiré leur travail et comment cela peut contribuer à l'évolution technologique [14].

II.2.3. Informatique bio-inspirée :

L'informatique bio-inspirée ou nature-inspirée s'intéresse à l'usage des systèmes naturels comme source d'inspiration pour le développement des systèmes informatiques.

Une des sources de cette inspiration est apparue via l'observation des systèmes biologiques tels que : le cerveau, la génétique des populations, le système immunitaire, les écosystèmes, les colonies d'insectes sociaux...

Les activités de recherche sur cet axe, comprennent, tant l'intégration de méthodes bio-inspirées dans des systèmes existants, que la conception de solutions originales.

Les chercheurs vont d'abord identifier le problème ou le besoin de l'Homme. Puis ils vont chercher des exemples biologiques qui répondront au besoin. Ils vont créer des hypothèses, mettre en place des schémas et tester leurs idées. Puis enfin, l'évaluer afin de savoir si c'est une solution durable ou non. C'est un processus long et il est difficile à évaluer car on ne peut jamais réellement savoir si le processus créé sera stable et de ce fait durable.

II.2.4. le processus d'un modèle inspiration :

Le processus de modèle inspiration est montré dans la figure ci-dessous (Figure II.1).

Ce processus se déroule en trois phases principales:

- **Observation** : Avant toute chose, il faut observer et étudier la nature.

Cette observation s'effectue en collaboration avec des biologistes. Lorsqu'un phénomène, une 'astuce' de la nature est observée, vient ensuite l'étape de la compréhension.

- **Modélisation** : L'observation ayant eu lieu, l'objectif est maintenant de mathématiser ce qui a été observé. Il ne suffit pas de copier, il faut comprendre. Il s'agit là du point essentiel pour permettre d'exploiter le phénomène observé.

- **Implémentation** : L'étape de la compréhension étant parfaitement réalisée, il devient alors possible d'implémenter l'observation première sur des dispositifs artificiels [15].

La bio-inspiration permet de travailler sur les différents domaines qui gravitent autour de la robotique : l'intelligence artificielle, la locomotion, la vision, etc.

La nature donne l'inspiration aux chercheurs de développer une observation d'un phénomène naturel particulier. Premièrement, Les développeurs créent et testent un modèle, utilisent des simulations mathématiques qui aident à raffiner le modèle original [15].

Le modèle raffiné sera utilisé pour extraire une méta-heuristique qui peut être utilisée comme une base pour finalement concevoir un algorithme inspiré de la nature. Pour passer d'un phénomène naturel à un algorithme inspiré de la nature, la figure suivante nous montre ce processus d'inspiration :

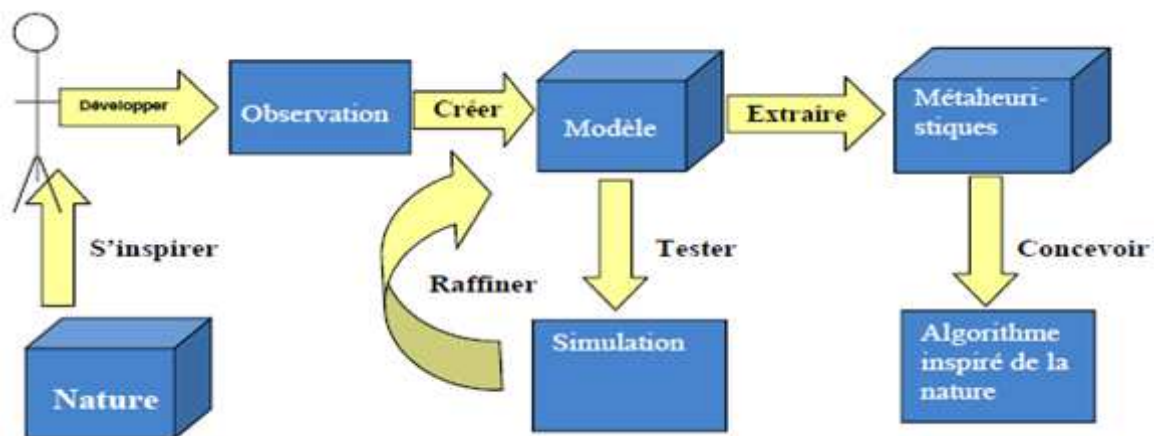


Figure II.1 : Processus d'inspiration d'un phénomène naturel [15].

II.3. Classification de Méthodes bio-inspirées :

Les plus prédominantes classes des algorithmes bio-inspirés sont les algorithmes évolutionnaires (inspirés de la sélection naturelle) et les algorithmes basés essaim (inspirés du comportement collectif chez les animaux). Ces deux classes du domaine bio-inspiré [1].

II.3.1 Algorithmes évolutionnaires :

Les Algorithmes Evolutionnaires (appelés aussi les Algorithmes Evolutionnistes) (Evolutionary Algorithms en anglais), sont une famille d'algorithmes s'inspirant de la théorie de l'évolution pour résoudre plusieurs problèmes connus dans le domaine d'informatique. Ils font ainsi évoluer un ensemble de solutions à un problème donné, dans le but de trouver de meilleurs résultats. Ce sont des algorithmes stochastiques, car ils utilisent itérativement des processus aléatoires. La

grande majorité de ces méthodes sont utilisées pour résoudre des problèmes d'optimisation [15].

II.3.2. Intelligence en essaim :

L'intelligence en essaim (Swarm Intelligence en Anglais) est le comportement collectif décentralisé, auto-organisé des systèmes naturels ou artificiels. Le concept est employé dans les travaux sur l'intelligence artificielle. L'expression a été introduite par Gerardo Beni et Jing Wang en 1989, dans le cadre de systèmes robotiques cellulaires. La source d'inspiration de cette approche vient de l'observation du comportement des insectes sociaux: une population d'agents extrêmement simples, interagissant et communiquant indirectement à travers leur environnement, constitue un algorithme massivement parallèle pour résoudre une tâche donnée (par exemple le fourragement, l'attroupement ou flockage, la division de tâches, la capture de proies...) [15].

Les méthodes de recherche à population, comme leur nom l'indique, travaillent sur une population de solutions et non pas sur une solution unique. On peut trouver d'autres noms génériques pour ces méthodes, le plus en vogue étant sans doute évolutionnaires algorithmes. Le principe général de toutes ces méthodes consiste à combiner des solutions entre elles pour en former de nouvelles en essayant d'hériter des bonnes caractéristiques des solutions parents.

Cette figure représente les différentes méthodes bio-inspirées :

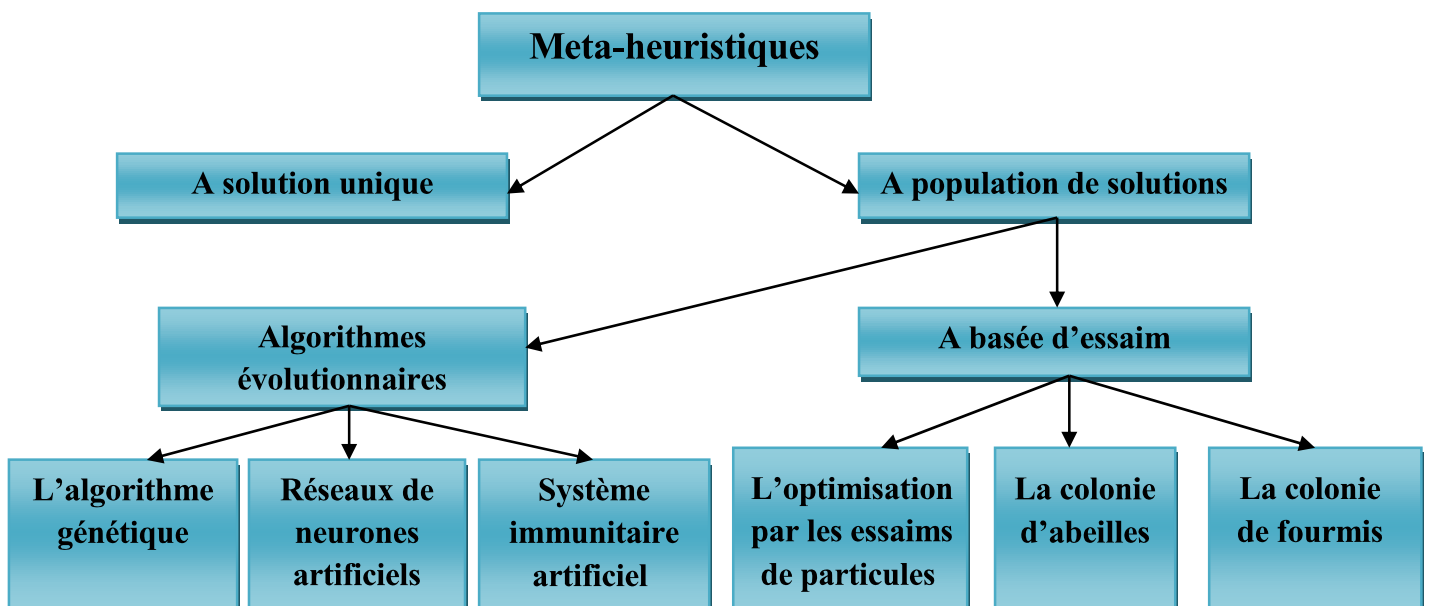


Figure II.2 : les différentes méthodes bio-inspirées.

II.4. Les méthodes bio-inspirées :

Nous présentons dans cette partie quelques méthodes d'optimisation bio-inspirés :

II.4.1. Les Réseaux de neurones :

- **Définition :**

Un réseau de neurones est un graphe orienté, chaque sommet du graphe est un neurone qui possède des entrées, des sorties et des paramètres. Les connexions du graphe transportent les sorties de certains neurones vers les entrées de certains autres [16].

- **Réseaux de neurones artificiels :**

Les réseaux de neurones artificiels (RNA) sont des réseaux fortement connectés de processeurs élémentaires fonctionnant en parallèle. Chaque processeur élémentaire calcule une sortie unique sur la base des informations qu'il reçoit. Toute structure hiérarchique de réseaux est évidemment un réseau.

- **Le neurone biologique:**

Le neurone est une cellule composée d'un corps cellulaire et d'un noyau. Le corps cellulaire se ramifie pour former ce que l'on nomme les dendrites. C'est par les dendrites que l'information est acheminée de l'extérieur vers le corps du neurone.

L'information traitée par le neurone chemine ensuite le long de l'axone (unique) pour être transmise aux autres neurones.

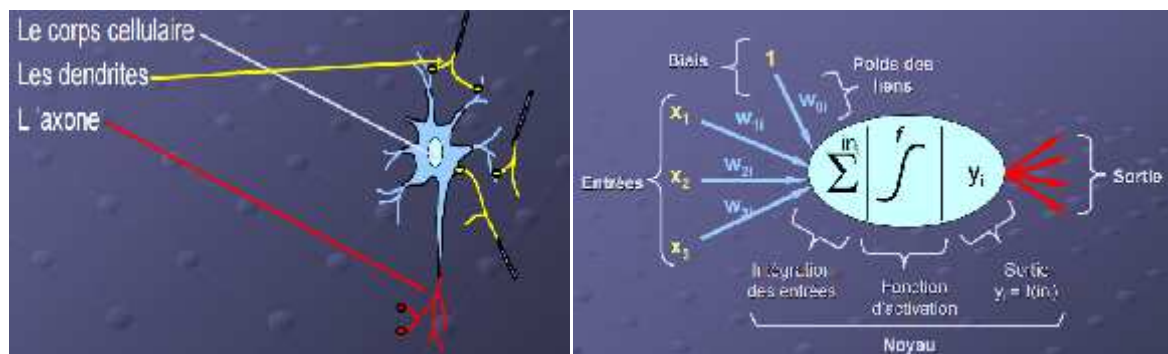


Figure II.3 : L'architecture de neurone biologique et artificiel.

- Le tableau suivant représente la différence entre les réseaux de neurones biologiques et les réseaux de neurones artificiels :

Neurone biologique	Neurone artificiel
Corps cellulaire	Fonction de transfert
Axone	Elément de sortie
Synapse	Poids

Table II.1:Modélisation des Réseaux de Neurones [16].

II.4.1.1. Propriétés des Réseaux de Neurones :

Les réseaux de neurones sont pourvus trois grandes propriétés :

- Approximation de fonction : le réseau de neurones est une méthode qui présente les qualités d'approximation universelle comme les méthodes de polynômes, de séries de Fourier. Ces méthodes permettent par régression d'approcher une fonction à un degré de précision fixé.
- Classification: Les réseaux de neurones sont capables de classifier. En fait, cette propriété découle de la capacité d'approximation.
- Apprentissage : apporte des propriétés d'adaptabilité des réseaux de neurones. Ainsi une même structure peut approcher une multitude de fonctions grâce à l'apprentissage [16].

II.4.1.2. Structures d'interconnexion :

Les connexions entre les neurones qui composent le réseau décrivent la topologie du modèle. Elle peut être quelconque, mais le plus souvent il est possible de distinguer une certaine régularité. Ils existent plusieurs types de topologie.

- **Réseau multicouche (au singulier) :** les neurones sont arrangés par couche. Il n'y a pas de connexion entre neurones d'une même couche et les connexions ne se font qu'avec les neurones des couches avalent. Habituellement, chaque neurone d'une couche est connecté à tous les neurones de la couche suivante et celle-ci seulement. Ceci nous permet d'introduire la notion de sens de parcours de l'information (de l'activation) au sein d'un réseau et donc définir les concepts de neurone d'entrée, neurone de sortie. Par extension, on appelle couche d'entrée l'ensemble des neurones d'entrée, couche de sortie l'ensemble des neurones de sortie. Les couches intermédiaires n'ayant aucun contact avec l'extérieur sont appelés couches cachées [16].

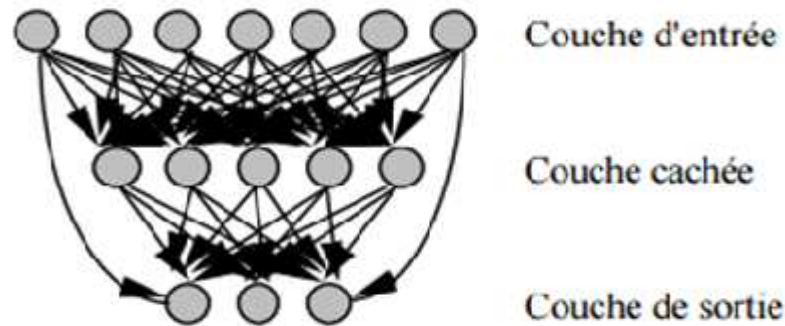


Figure II.4 : Définition des couches d'un réseau multicouche [17].

- **Réseau à connexions locales** : Il s'agit d'une structure multicouche, mais qui à l'image de la rétine, conserve une certaine topologie. Chaque neurone entretient des relations avec un nombre réduit et localisé de neurones de la couche avale. Les connexions sont donc moins nombreuses que dans le cas d'un réseau multicouche classique [10].

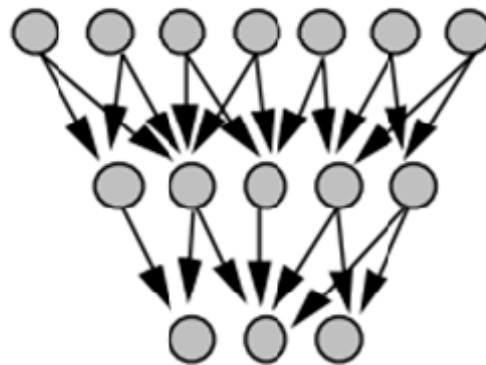


Figure II.5 : Réseau à connexion locales [17].

- **Réseau à connexions récurrentes** : les connexions récurrentes ramènent l'information en arrière par rapport au sens de propagation défini dans un réseau multicouche. Ces connexions sont le plus souvent locales [17].

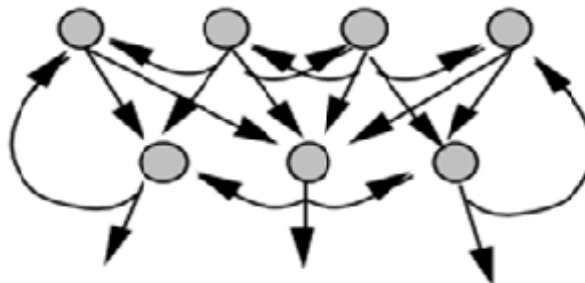


Figure II.6 : Réseau à connexion récurrents [17].

- **Réseau à connexion complète** : c'est la structure d'interconnexion la plus générale. Chaque neurone est connecté à tous les neurones du réseau (et à lui-même) [17].

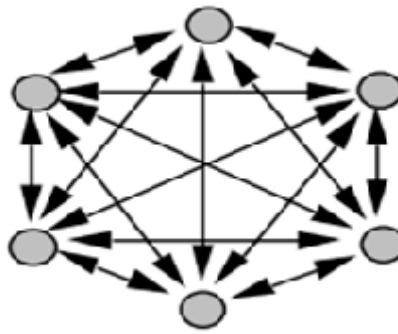


Figure II.7 : Réseau à connexion complète [16].

II.4.1.3. Architecture des Réseaux de Neurones :

On distingue deux grandes familles des réseaux de neurones :

➤ **Réseaux de neurones non bouclés :**

Un réseau de neurones non bouclé est représenté graphiquement par un ensemble de neurones "connectés" entre eux, l'information circulant des entrées vers les sorties sans "retour en arrière" ; si l'on représente le réseau comme un graphe dont les nœuds sont les neurones et les arêtes les "connexions" entre ceux-ci, le graphe d'un réseau non bouclé est acyclique montré dans la figure ci-dessus. Le terme de "connexions" est une métaphore : dans la très grande majorité des applications, les réseaux de neurones sont des formules [16].

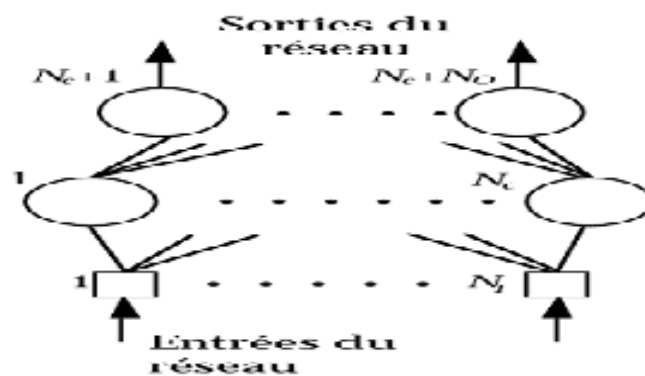


Figure II.8 : Réseau de neurones non bouclés [16].

➤ **Réseaux de neurones bouclés :**

Un réseau est bouclé, ou dynamique, si son graphe possède au moins un cycle comme dans la figure suivante, lorsqu'on se déplace dans le réseau en suivant le sens des connexions, il est possible de trouver au moins un chemin qui revient à son point de départ.

En pratique, les réseaux de neurones bouclés sont utilisés pour effectuer des tâches de modélisation de systèmes dynamiques [16].

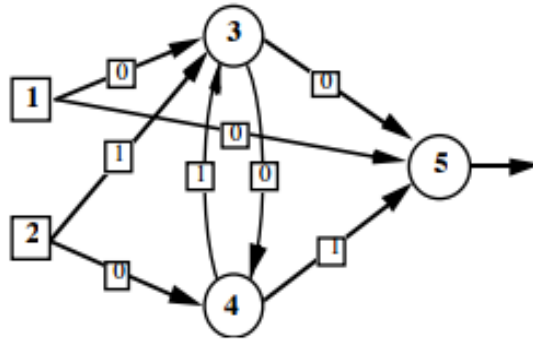


Figure. II.9 : Réseau de neurones bouclés [16].

Les avantages :

- Capacité de représenter n'importe quelle fonction, linéaire ou pas, simple ou complexe.
- Faculté d'apprentissage à partir d'exemples représentatifs, par " rétro propagation des Erreurs ". L'apprentissage (ou construction du modèle) est automatique.
- Résistance au bruit ou au manque de fiabilité des données.
- Simple à manier, beaucoup moins de travail personnel à fournir que dans l'analyse statistique classique.
- Comportement moins mauvais en cas de faible quantité de données.
- Temps de réponse, C'est l'un des avantages principaux du réseau de neurones : en effet une fois que le réseau a appris, il peut sortir quasi-instantanément la réponse. En fait, les opérations que fait un réseau de neurones sont très simples du point de vue informatique, et peu gourmandes en CPU.

Les inconvénients :

- L'absence de méthode systématique permettant de définir la meilleure topologie du réseau et le nombre de neurones à placer dans la (ou les) couche(s) cachée(s).
- La connaissance acquise par un réseau de neurone est codée par les valeurs des poids sont inintelligibles pour l'utilisateur.

II.4.2. L'algorithme génétique :

Les algorithmes génétiques sont des algorithmes d'optimisations stochastiques fondées sur les mécanismes de la sélection naturelle et de la génétique. Ils ont été initialement développés par John Holland.

- Voici un algorithme synthétique pour la Programmation génétique :
 1. Génération aléatoire de la population (1 individu = 1 programme).

2. Évaluation du fitness de chacun des individus de la population (Évaluation de l'adéquation des programmes au problème à résoudre).
3. Sélection des individus les mieux adaptés.
4. Application des opérateurs de croisement, mutation, reproduction sur la Population afin de créer une nouvelle population.
5. Répéter l'étape 2 un certain nombre de fois [17].

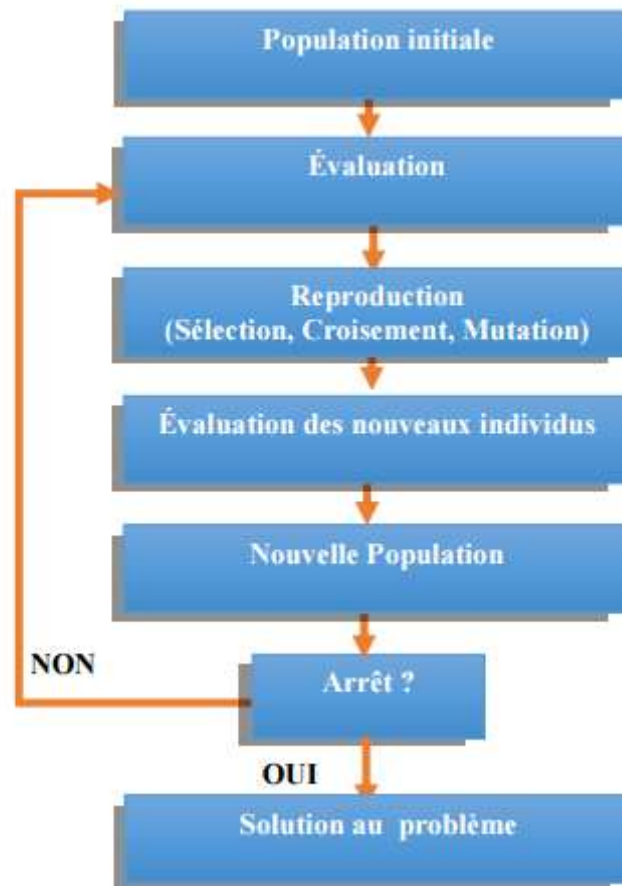


Figure II.10: l'organigramme de fonctionnement d'un algorithme génétique

- **L'avantage :**

L'avantage des algorithmes génétiques est leur simplicité. Néanmoins, les algorithmes génétiques seuls ne sont pas très efficaces dans la résolution d'un problème. Cependant, ils apportent une solution acceptable assez rapidement. Sa combinaison avec un algorithme déterministe peut améliorer la solution d'une manière assez efficace [17].

- **L'inconvénient :**

Malheureusement, il y a de nombreux inconvénients. En effet, il n'y a pas de garantie quant à l'obtention de la solution optimale au problème posé en un temps fini. L'utilisation de ces algorithmes est souvent coûteuse en temps de calculs (les algorithmes sont lents).

II.4.3. La colonie de fourmis :

Optimisation par colonie de fourmis (AntColonyOptimization ACO en Anglais) est un algorithme d'optimisation inspiré de la nature qui est motivé par le comportement naturel de recherche de nourriture des espèces de fourmis. Son principe repose sur le comportement particulier des fourmis qui, lorsqu'elles quittent leur fourmilière pour explorer leur environnement à la recherche de nourriture, finissent par élaborer des chemins qui s'avèrent fréquemment être les plus courts pour aller de la fourmilière à une source de nourriture intéressante [17].

Les fourmis utilisent les pistes de phéromones pour marquer leur trajet, par exemple entre le nid et une source de nourriture. Une colonie est ainsi capable de choisir (sous certaines conditions) le plus court chemin vers une source exploitée, sans que les individus aient une vision globale du trajet [17].

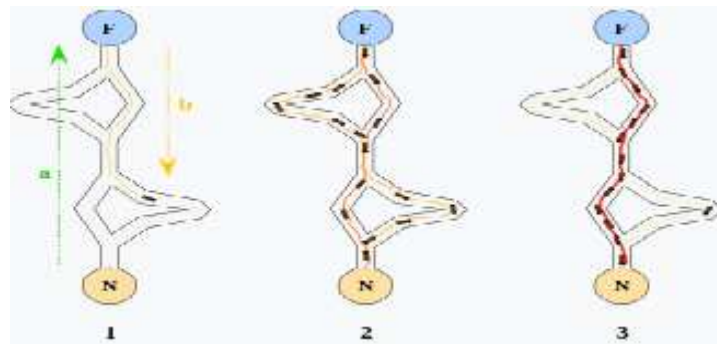


Figure II.11: Expérience de sélection des branches les plus courtes par une colonie de fourmis.

- 1) la première fourmi trouve la source de nourriture (F), via un chemin quelconque (a), puis revient au nid (N) en laissant derrière elle une piste de phéromone (b).
- 2) les fourmis empruntent indifféremment les quatre chemins possibles, mais le renforcement de la piste rend plus attractif le chemin le plus court.
- 3) les fourmis empruntent le chemin le plus court, les portions longues des autres chemins perdent leur piste de phéromones.

Un modèle expliquant ce comportement est le suivant :

1. une fourmi (appelée « éclaireuse ») parcourt plus ou moins au hasard l'environnement autour de la colonie.
2. si celle-ci découvre une source de nourriture, elle rentre plus ou moins directement au nid, en laissant sur son chemin une piste de phéromones.
3. ces phéromones étant attractives, les fourmis passant à proximité vont avoir tendance à suivre, de façon plus ou moins directe, cette piste.
4. en revenant au nid, ces mêmes fourmis vont renforcer la piste.
5. si deux pistes sont possibles pour atteindre la même source de nourriture, celle étant la plus courte sera, dans le même temps, parcourue par plus de fourmis que la longue piste.

6. la piste courte sera donc de plus en plus renforcée, et donc de plus en plus attractive.
7. la longue piste, elle, finira par disparaître, les phéromones étant volatiles.
8. à terme, l'ensemble des fourmis a donc déterminé et « choisi » la piste la plus courte.

- **Avantage :** Très grande adaptabilité.

Parfait pour les problèmes basés sur des graphes [15].

- **Inconvénients :** Un état bloquant peut arriver.

Temps d'exécution parfois long.

Ne s'applique pas à tous type de problèmes [15].

II.4.4. La colonie d'abeille artificielle :

L'algorithme de colonie d'abeille artificielle (Artificial Bee Colony ABC en Anglais) est l'un des plus récemment introduit dans la base des Algorithmes basés essaim. La colonie d'abeille artificielle simule le comportement intelligent de recherche de nourriture d'un essaim d'abeilles. L'algorithme de colonie d'abeille artificielle a été introduit par Karaboga [9].

- **Principe de l'algorithme :**

Chaque solution représente une position de nourriture potentielle dans l'espace de recherche et la qualité de la solution correspond à la qualité de la position alimentaire. Agents (abeilles artificielles) recherche à exploiter les sources de nourriture dans l'espace de recherche. La colonie d'abeille artificielle utilise trois types d'agents : les abeilles employés (Employed Bee EB en Anglais), les abeilles spectateur (Onlookers Bee EO en Anglais), et les scouts (Scouts Bee SB en Anglais). Les abeilles employées (AE) sont associés à des solutions actuelles de l'algorithme. À chaque étape de l'algorithme un AE tente d'améliorer la solution, il la représente en utilisant une étape de recherche locale, après il va essayer de recruter des abeilles spectateur (AS) pour sa position actuelle. Les AS choisissent parmi les postes promus en fonction de leur qualité, ce qui signifie que de meilleures solutions attireront plus d'AS. Une fois un AS a choisi un AE et donc une solution, il cherche à optimiser la position de l'AE par une étape de recherche locale. AE mise à jour sa position si un AS recruté était en mesure de repérer une meilleure position, sinon il reste sur sa position actuelle. En outre, un AE va abandonner sa position si elle n'était pas en mesure d'améliorer sa position pour certain nombre d'étapes. Quand une AE abandonne sa position, il devient une Abeille éclaireuse, ce qui signifie qu'elle sélectionne une position aléatoire dans l'espace de recherche et devienne une Abeille employée à cette position.

La majorité des problèmes qui ont étaient résolus par cette méthode, ont donné de très bons résultats concernant la valeur de la fonction objective, et le temps d'exécution qui a été acceptable [15].

- **Les Avantages :**

L'algorithme de la colonie d'abeille artificiel présente les avantages d'une grande robustesse, une convergence rapide, une grande flexibilité et moins de paramètres de contrôle [15].

II.4.5. L'optimisation par les essaims de particules :

L'optimisation par les essaims de particules (Particle Swarm Optimization ou PSO) a une inspiration biologique. Elle a été introduite en 1995 par Kennedy et Eberhart, psychologue social et ingénieur en électricité.

L'optimisation par les essaims de particules est une méthode heuristique basée sur la population découverte grâce à la simulation des modèles sociaux de volées d'oiseaux, bancs de poissons ou essaims d'insectes. Supposons le scénario suivant: un groupe d'oiseaux recherche aléatoirement de la nourriture dans une région. Il n'y a qu'un seul morceau de nourriture dans la zone fouillée. Les oiseaux ne savent pas où est la nourriture, mais ils savent à quelle distance elle se trouve et les positions de leurs pairs. Alors, quelle est la meilleure stratégie pour trouver la nourriture? Une stratégie efficace est de suivre l'oiseau qui est le plus proche de la nourriture. L'optimisation par les essaims de particules se base sur ce scénario et l'utilise pour résoudre les problèmes d'optimisation. Dans L'optimisation par les essaims de particules, chaque solution unique est comme un "oiseau" dans l'espace de recherche, qui est appelé "particule" [18].

L'algorithme commence avec une initialisation aléatoire de l'essaim de particules dans l'espace de recherche. Chaque particule est modélisée par sa position dans l'espace de recherche et par sa vitesse. A chaque instant, toutes les particules ajustent leurs positions et vitesses, donc leurs trajectoires, par rapport :

1. à leurs meilleures positions.
2. à la particule ayant la meilleure position dans l'essaim.
3. à leur position actuelle.

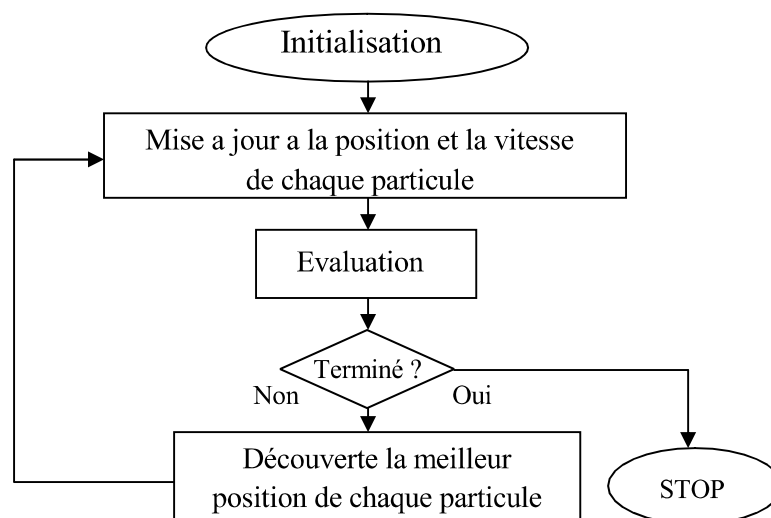


Figure II.12: l'organigramme de fonctionnement de l'optimisation par les essaims de particules [19].

II.4.6. Système immunitaire artificiel :

Plus récemment, on assiste à un intérêt croissant accordé à l'utilisation d'un autre système biologique qui est le système immunitaire comme source d'inspiration qui est doté de capacité de traitement de l'information y compris l'identification du modèle l'apprentissage, la mémorisation. ce domaine de recherche est appelé le Système Immunitaire Artificiel [16].

Un système immunitaire artificiel (SIA) est une catégorie d'algorithme d'optimisation, inspirée par les principes et le fonctionnement du système immunitaire naturel des vertébrés. Ces algorithmes exploitent typiquement les caractéristiques du système immunitaire naturel pour ce qui est de l'apprentissage et de la mémorisation comme moyen de résolution de problèmes. Ils sont liés à l'intelligence artificielle et très proche des travaux sur les algorithmes génétiques [16].

Les algorithmes des AIS peuvent être divisés en trois grandes parties, chacune d'elles s'inspirant d'un comportement ou d'une théorie du système immunitaire biologique : la sélection négative, la sélection clonale et les réseaux immunitaires.

➤ **Sélection négative :**

La sélection négative, permet d'assurer que tout lymphocyte T reconnaissant trop les cellules du soi comme antigène ne sorte pas du thymus afin d'éviter les maladies auto-immunes.

➤ **Sélection clonale :**

La sélection clonale est la théorie expliquant comment le système immunitaire interagit avec les antigènes.

Les principales différences entre tous les algorithmes de sélection clonale sont les méthodes utilisées dans la génération aléatoire des détecteurs, la mutation et l'affinité entre les détecteurs et les antigènes.

➤ **Les Réseaux immunitaires :**

Le réseau immunitaire décrit la manière dont les cellules répondent entre elles dans le système immunitaire. Alors est un système immunitaire autorégulé de molécules et de cellules qui se reconnaissent entre elles, même en l'absence d'antigène. Parmi les algorithmes de ce modèle L'algorithme AiNet qui utilise la notion d'image interne pour représenter les regroupements de données dans un réseau.

II.5. Domaine d'application des méthodes bio-inspirées :

Les méthodes bio-inspirées sont utilisées dans plusieurs domaines d'activités. Parmi ces domaines nous pouvons citer dans cette table :

Les méthodes bio-inspirées	Domaine d'application
les réseaux de neurones	• Traitement d'images : comparaison d'images, reconnaissance de caractères et de signatures, reconnaissance des formes et de motifs. • Traitement du signal : traitement de la parole, filtrage, classification. • Contrôle : diagnostic de pannes, commande de processus, contrôle qualité, robotique.
L'algorithme génétique	• Optimisation : optimisation de fonctions, planification, etc ... • Apprentissage : classification, prédiction, robotique, etc ... • Programmation automatique : programmes LISP, automates cellulaires, etc ... • Etude du vivant, du monde réel : marchés économiques, comportements sociaux, systèmes immunitaires, etc ...
La Colonie de fourmis	Les algorithmes de colonies de fourmis ont été appliqués à un grand nombre de problèmes d'optimisation combinatoire, allant de l'affectation quadratique au repli de protéine ou au routage de véhicules. Comme beaucoup de méta-heuristiques, l'algorithme de base a été adapté aux problèmes dynamiques, en variables réelles, aux problèmes stochastiques, multi-objectifs ou aux implémentations parallèles, etc.
La colonie d'abeille artificielle	• La résolution du problème du voyageur de commerce (TSP) qui a été faite par Lucic et Teodorovic et qui a donné de très bons résultats. • L'apprentissage des réseaux de neurones tels que MLP, RBF, SNN, LVQ. • Conception électroniques et mécanique. • Clustering de données. • Contrôle de robot. • L'ordonnancement de tâches.
	• Télécommunications. • le contrôle.

L'optimisation par les essaims de particules	• l'exploration de données. • l'optimisation combinatoire. • les systèmes d'alimentation. • le traitement du signal.
Système immunitaire artificiel	• Système informatique hybrides d'analyse de données bios informatique • Algorithme de génération d'emploi de temps

Table II.2 : Quelques domaines d'application.

II.6. Quelques travaux réalisés des méthodes bio-inspirées :

Le tableau suivant récapitule un historique de différents travaux qui ont déjà réalisé par des méthodes bio-inspiration :

les méthodes bio_inspirées	Année	Les chercheurs	Les travaux	Référence
Réseau de neurones artificiels	1943	Mac Culloch et Pitts	Présentation d'un neurone formel	[16]
	1949	Donald hebb	Loi de hebb	[16]
	1958	Widrow et Hoff	Modèle avec processus d'apprentissage, perceptron	[16]
	1969	Minsky et Papert	Limite des perceptrons	[16]
	1972	Kohonen	Mémoires associatives	[16]
	1980	Rumelhart – Mc Clelland	Perceptron multicouches, mécanismes d'apprentissage performants	[16]
L'algorithme génétique	1960	Charles Darwin	Livre intitulé L'origine des espèces au moyen de la sélection naturelle	[20]
	1966	L. J. Fogel	Programmation évolutionnaire	[20]
	1973	I. Rechenberg	Stratégie d'évolution	[20]
	1975	John Holland	Le premier modèle formel des algorithmes génétiques	[20]

La colonie de fourmis	1991	Dorigo & al	Optimisation distribuée par Ant Colonies Proceedings	[21]
	1997	Wodrich & Bilchev	La métaphore de la colonie de fourmis pour la recherche d'espaces de conception continue Notes de cours en informatique	[21]
	2000	Monmarché et al	Sur la façon <i>Pachycondyla apicalis</i> fourmis suggèrent un nouvel algorithme de recherche Future Generation Computer Systems	[21]
La colonie d'abeilles	2004	Craig A. Tovey à Georgia Tech	Collaboration avec SUNIL NAKRANIC	[22]
	2005	Xin-She Yang à l'Université de Cambridge	A développé le Virtual Bee Algorithm (VBA) pour résoudre des problèmes d'optimisation numérique	[22]
	2006	B. Basturk et D. Jarabogo en Turquie	Ils ont développé un algorithme appelé ARTIFICIAL BEE COLONY (ABC) pour l'optimisation de fonction numérique	[22]
L'optimisation par les essaims de particules	1995	Kennedy et Eberhart	L'optimisation par essaim de particules s'inspire du comportement social des oiseaux évoluant en groupe et des bancs de poissons.	[23]
	1980	Craig W. Reynolds	Il développe ce modèle permettant de simuler d'un groupe d'oiseaux	[23]
	2004	Wachowiak et al	L'enregistrement multimodal d'images biomédicales	[23]
Système immunitaire artificiel	1994	Forrest S et al	Sécurité informatique et réseau	[16]
	1997	Somayaji A. et al.		[16]
	1999	Hofmeyr, S.A., Forrest, S		[16]

	1997	Mori, M. et al	Logarithme	[16]
	1998	Hart, E. et al		[16]
	1999	Russ, S.H. et al		[16]
	2002	Costa, A.M. et al		[16]
	2000	De Castro, L.N, Von Zuben F.J	Modèle d'identification multi-usages (combinatoire)	[16]
	2002	De Castro, L.N. , Timmis, J	Modèle d'identification multi-usages (combinatoire)	[16]

Table II.3: Un historique de différents travaux.

Nous pouvons dire que les méthodes bio-inspirées sont utiles lorsque:

- on traite des problèmes qui n'ont pas de solution algorithmique
- on traite des données imprécises, bruitées, incomplètes
- on cherche des solutions suffisamment bonnes à des problèmes pour lesquels les solutions exactes ou optimales sont trop chères ou trop difficiles à trouver.

II.7. Conclusion :

Dans ce chapitre, nous avons dressé un état de l'art sur les méthodes d'optimisation bio-inspirées qui ils sont basé sur les méta-heuristiques les plus connues. Ces dernières sont divisées en deux classes : les algorithmes évolutionnaires qui s'inspirent de la théorie de l'évolution pour résoudre des problèmes divers utilisent itérativement des processus aléatoires (stochastiques) et les algorithmes basé essaim inspirant du comportement collectif décentralisées, auto-organisés des systèmes naturels où une population d'agents extrêmement simples, interagissant et communiquant indirectement à travers leur environnement, constitue un algorithme massivement parallèle pour résoudre une tâche donnée.

III.1.1. Introduction :

Les réseaux de neurones artificiels constituent l'une des approches d'intelligence artificielle dont le développement se fait à travers les méthodes par lesquelles l'homme essaye toujours d'imiter la nature et de reproduire des modes de raisonnement et de comportement qui lui sont propres.

Ce chapitre est consacré à une étude détaillée des réseaux de neurones multicouches (MLP Multi Layer Perceptron) qui sont l'unité de base de l'approche implémentée dans ce travail pour la classification des données médicales.

III.2. Neurone et le réseau de neurones :

Il existe un grand nombre de modèles de neurones. Les plus utilisés sont basés sur le modèle développé par McCULLOCH & PITTS [25].

La figure suivante montre un schéma comportant la structure générale d'un neurone artificiel.

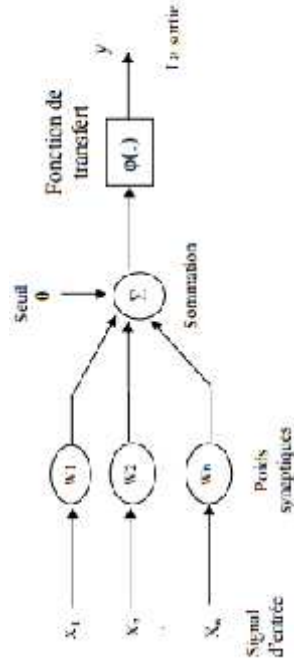


Figure III.1 : Neurone artificiel.

Un neurone artificiel est considéré comme un élément élémentaire de traitement de l'information. Il reçoit les entrées et produit un résultat à la sortie.

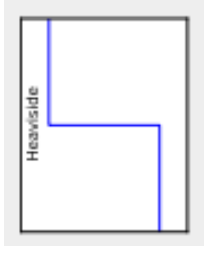
Chaque neurone artificiel est un processeur élémentaire. Il reçoit un nombre variable d'entrées en provenance de neurones amont. À chacune de ces entrées est associé un poids w , abréviation de weight (poids en anglais), représentatif de la force de la connexion.

Chaque processeur élémentaire est doté d'une sortie unique, qui se ramifie ensuite pour alimenter un nombre variable de neurones aval.

neurones biologiques dont l'état est binaire, la plupart des fonctions de transfert sont continues, offrant une infinité de valeurs possibles comprises dans l'intervalle $[0, +1]$ ou $[-1, +1]$ [25].

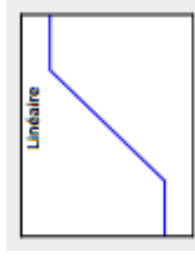
Heaviside (seuil θ)

- Si $x < \theta$ alors $f(x) = 0$
- Si $x \geq \theta$ alors $f(x) = 1$



Linéaire (seuil θ_1, θ_2)

- Si $x < \theta_1$ alors $f(x) = 0$
- Si $x > \theta_2$ alors $f(x) = 1$
- Si $\theta_1 \leq x \leq \theta_2$ alors $f(x) = x$



Sigmoïde / Tangente hyperbolique

- $f(x) = 1 / 1 + \exp(-x)$



Figure III. 2: Différents types de fonctions de transfert pour le neurone artificiel [25].

III.3. Le modèle MLP (perceptron multicouche) :

III.3.1. Définition : Le réseau MLP (Multi Layer Perceptron) est un modèle de réseau de neurones comprenant une ou plusieurs couches cachées, dont les connexions sont directes et totales entre les couches. Werbos, pour sa thèse de doctorat en 1974, développa la théorie de la rétro propagation du gradient de l'erreur. Cet algorithme fut popularisé en 1986 par Rumelhart qui l'utilisa pour la modification des poids et des biais dans les réseaux de neurones multicouches. Le Perceptron multicouche est un Classifieur linéaire de type réseau neuronal formel organisé en plusieurs couches au sein desquelles une information circule de la couche d'entrée vers la couche de sortie uniquement : Chaque couche est constituée d'un

La première couche est reliée aux entrées, ensuite chaque couche est reliée à la couche précédente. C'est la dernière couche qui produit les sorties du MLP. Les sorties des autres couches ne sont pas visibles à l'extérieur du réseau, et elles sont appelées pour cette raison couches cachées [27].

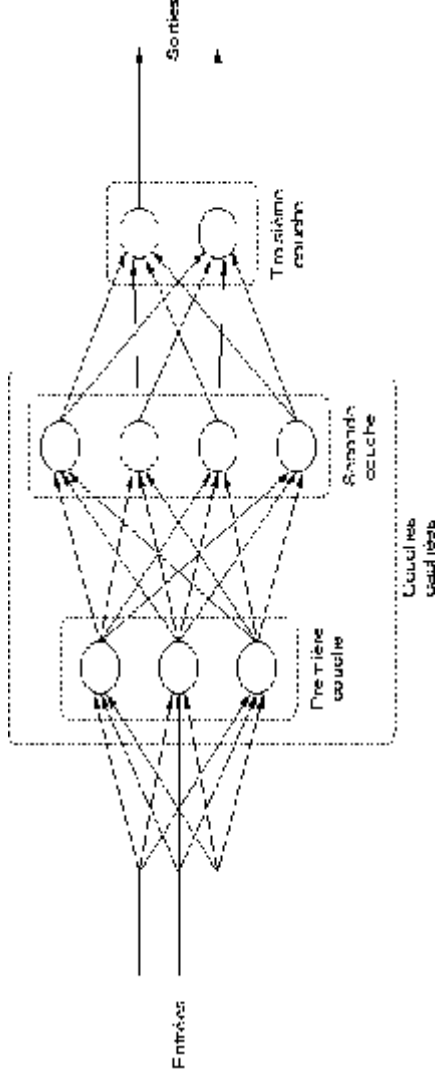


Figure III.3 : Perceptrons multicouches[27].

III.3.3. Apprentissage du MLP :

L'apprentissage est une phase du développement d'un réseau de neurones durant laquelle le comportement du réseau est modifié jusqu'à l'obtention du comportement désiré. L'apprentissage neuronal fait appel à des exemples de comportement.

La loi de Hebb, un exemple d'apprentissage non supervisé La loi de Hebb (1949) s'applique aux connexions entre neurones, comme le représente la figure



Figure III.4 : i le neurone amont, j le neurone aval et w_{ij} le poids de la connexion.

MLP utilise la règle de rétro propagation du gradient (Back-propagation), Il a fallu attendre le début des années 1980 pour qu'une règle efficace soit mise au point pour l'apprentissage des réseaux multicouches. Cette règle, découverte simultanément par des

1-Début : Initialisation des poids du MLP**2-Faire** :

- Propagation Directe : Pour chaque vecteur de la base d'apprentissage :

2-1 Affecter $X_k = X_n$ ($n=1,...,N$) ; envoyer aux neurones cachés.

2-2 Chaque neurone caché calcule son entrée nette :

$$Z_in_j = V_{j0} + \sum_{1 \leq n \leq N} T_n V_{jn}$$

2-3 Appliquer sa fonction de sortie $Z_j = F(Z_in_j)$

2-4 Chaque neurone de sortie calcule son entrée

$$Y_in_m = W_{j0} + \sum_{1 \leq j \leq J} Z_j W_{mj}$$

2-5 Appliquer sa fonction de sortie

$$Y_m = F(Y_in_m)$$

- Rétropropagation : Chaque neurone de sortie reçoit son étiquette t_m

2-6 Calculer les signaux d'erreur de sortie

$$\delta_m = (t_m - Y_m) F'(Y_in_m)$$

2-7 Calculer la modification des poids-couche de sortie

$$\Delta W_{mj} = \eta Z_j \delta_m$$

2-8 Rétropropager les signaux d'erreurs δ_m vers la couche cachée précédente.

2-9 Chaque neurone caché calcule sa contribution à partir de la rétro propagation des signaux d'erreurs de sortie.

$$\delta_in_j = \sum_{1 \leq m \leq M} \delta_m W_{mj}$$

2-10 Chaque neurone caché calcule son signal d'erreur

$$\delta_j = \delta_in_j F'(Z_in_j)$$

2-11 calculer les modifications de poids-couche caché

L'apprentissage d'un MLP est défini comme un problème d'optimisation qui consiste à trouver les coefficients du réseau en minimisant une fonction d'erreur globale, la définition de cette fonction est primordiale, car celle-ci sert à mesurer l'écart entre les sorties désirées du modèle et les sorties du réseau obtenus. Dont la définition est :

Pour chaque exemple n ($n \in N$) calculer une fonction d'erreur quadratique selon l'équation suivante :

$$e(n) = \frac{1}{2} \sum_{1 \leq j \leq J} [d_j(n) - Y_j(n)]^2 \quad \dots\dots(1)$$

Pour tout l'ensemble d'apprentissage N , la fonction d'erreur quadratique moyenne (EQM) :

$$E(n) = \frac{1}{N} \sum_{n=1}^N e(n) \quad \dots\dots (2)$$

III.3.4. Calcul de fonction de l'activation :

La fonction d'activation est de type sigmoïde exponentiel :

$$F(x) = 1 / [1 + \exp(-x)] \quad \dots\dots(3)$$

Une fonction $F(x)$ est dite sigmoïde si c'est une fonction continue, non-linéaire, bornée et monotone croissante telle qu'il existe une valeur unique pour laquelle la dérivée $F'(x)$ atteint un maximum global d'où :

$$F'(x) = F(x) * (1 - F(x)) \quad \dots\dots(4)$$

III.3.5. Algorithme de la rétropropagation des erreurs :

On a vu que les poids de la structure MLP ont fait l'objet d'une optimisation par l'algorithme de la rétro propagation des erreurs [27].

Soit le réseau MLP composé de $L+1$ couches (couche cachées plus une couche de sortie).

q : un indice qui représente le numéro de la couche.

W_{ji} : le poids de connexion entre le neurone ' j ' de la couche ' q ' et le neurone ' i ' de la couche précédente ' $q-1$ '.

H_0 : le nombre d'itération maximum.

h : un indice qui représente le numéro d'itération.

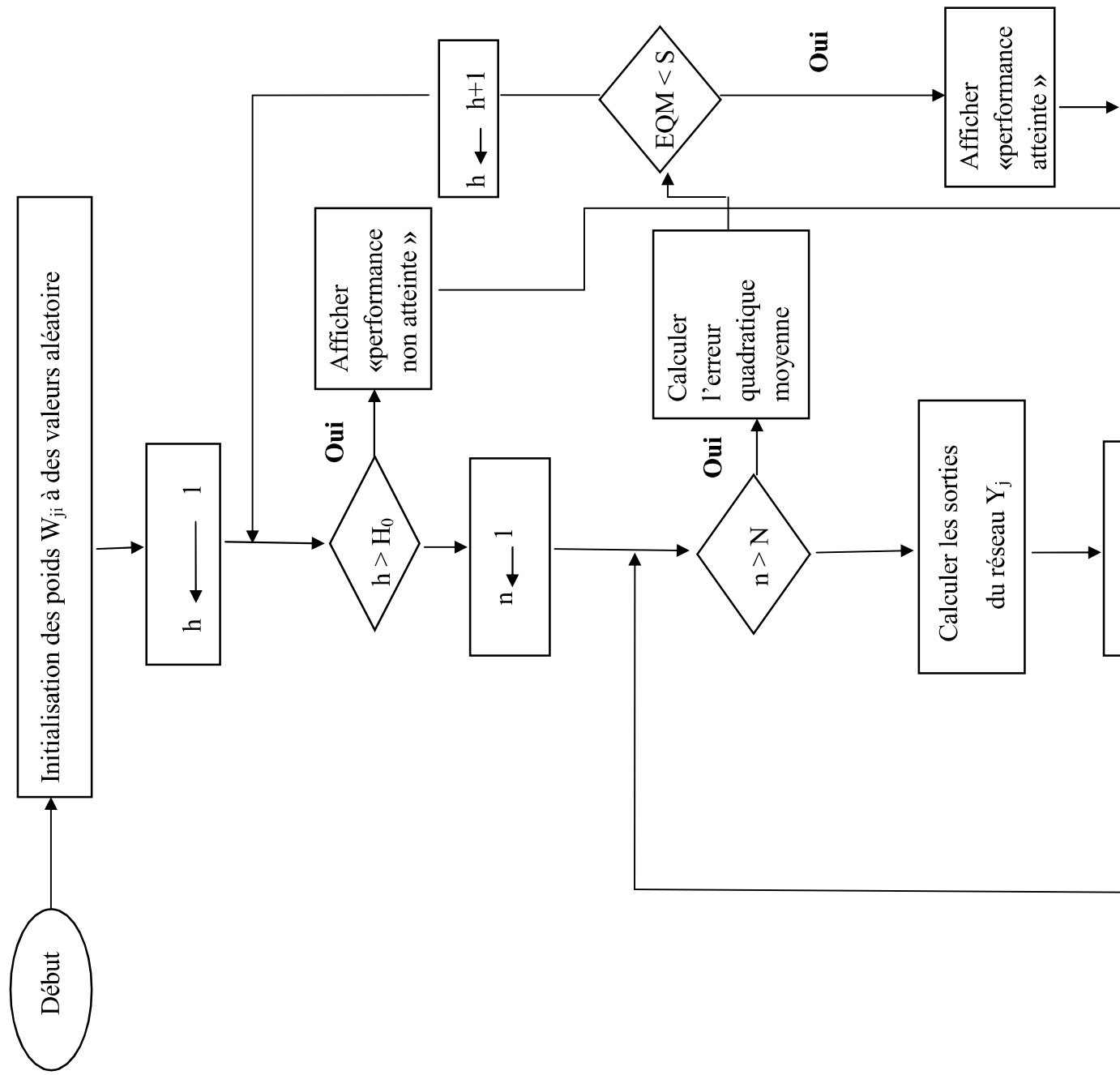
EQM : l'erreur quadratique moyenne sur l'ensemble d'apprentissage N .

Y_j : la sortie ' j ' du réseau.

S : un seuil qui représente la valeur minimale d'EQM qu'on désire obtenir.

N : le nombre d'éléments de la base d'apprentissage.

n : un indice qui représente le numéro d'élément de la base d'apprentissage



La Figure III.5, au dessus représente l'organigramme de l'algorithme de la rétropropagation des erreurs, qui est composé essentiellement de trois boucles imbriquées l'une dans l'autre.

La première boucle sert au contrôle du nombre d'itération 'h', la deuxième boucle contrôle le nombre d'exemples d'apprentissage, quand à la dernière boucle elle contrôle la propagation de l'erreur sur les différents couches du réseau [27].

On remarque que la fin de l'algorithme peut être causée par la condition : EQM est inférieure au seuil S ou bien on attend le nombre maximum d'itérations.

III.4. Algorithme de classification de MLP :

La phase de classification de MLP est définie comme suit [27] :

1. Sélectionner un exemple de la phase de test.
2. Faire propagation et puis calculer la sortie de chaque classe.
3. Si la sortie obtenue est plus proche que la sortie désirée alors l'exemple est bien classé.
4. Boucler jusqu'à ce que tous les exemples soient présentés.

III.5 Caractéristique du MLP :

Les perceptrons multicouches(MLP) ont d'abord été introduits pour résoudre des problèmes complexes de classification.

Un perceptron multicouche est capable d'approximer des fonctions de forme très différente et l'approximation faite n'est pas global mais local, Le choix du nombre de couches et du nombre de neurones est primordial dans un perceptron. En ajoutant des neurones ou des couches on améliore les capacités du réseau et donc la finesse de l'approximation, mais l'apprentissage devient plus long (particulièrement en augmentant le nombre de couches) et le risque de sur-apprentissage augmente[27].

Généralement PMC est caractérisé par :

- Une seule couche d'entrée.
- Une seule couche de sortie.

moment permet au réseau de répondre non seulement au gradient local mais aussi aux tendances récentes de la surface d'erreur (hyper surface de la fonction de coût vis-à-vis des paramètres libres du réseau, c'est-à-dire les poids synaptiques). Agissant comme un filtre passe-bas, le moment permet au réseau d'ignorer des petits détails de la surface réduisant dès lors la probabilité que l'algorithme d'apprentissage ne reste coincé dans un minimum local non globalement optimal. Le moment peut être ajouté à l'apprentissage par rétro propagation en gardant une mémoire du changement de poids le plus récent [28].

Le perceptron multicouche est le plus connu et le type de réseau neuronal le plus fréquemment utilisé. Sur la plupart du temps, les signaux sont transmis dans les réseaux dans une direction: de l'entrée à la sortie. Là N'est pas une boucle, la sortie de chaque neurone n'affecte pas Le neurone lui-même. Cette architecture s'appelle feedforward [29].

III.6 Avantages et Inconvénients du MLP :

Certains avantages et inconvénients du MLP sont récapitulés dans la table suivante :

Avantages	Inconvénients
Gestion des relations non linéaires entre variable	Boite noire
Non paramétrique	Convergence non assurée
Versatilité	Essentiellement applicable aux variables continues
Robustesse au bruit	Difficulté de mise en œuvre (expertise)

Table III.1 : Avantages et Inconvénients du MLP.

Lorsque l'on augmente le nombre de couches, on constate que l'algorithme d'entraînement nécessite de plus en plus d'itérations pour converger vers un résultat. C'est

- De très nombreuses applications
 - machine à prononcer l'anglais
 - analyse de courbes, classification de formes (EEG, diagrammes,
 - analyse financière
 - aide à la prise de décision
 - reconnaissance de lettres, de paroles, etc...
- Exemple de la modélisation du transport, mais en étudiant simultanément le transport par rail et par route
- Identification de séries temporelles et prévision
- Consommation d'eau, d'électricité

III.8. Conclusion :

A travers ce chapitre nous avons exploré la capacité des réseaux de neurones artificiels dans le domaine de la reconnaissance des formes. Ainsi ses différents modèles et son algorithme d'apprentissage.

Nous nous sommes concentrés principalement dans ce chapitre sur une architecture neuronale qui est concernée par notre travail de recherche et qui est généralement la plus utilisée pour la classification des anomalies cardiaques: le perceptron multicouche.

Les perceptrons multicouches sont les plus couramment utilisés types de réseaux neuronaux.

Le prochain chapitre sera la mise en œuvre des algorithmes détaillés ainsi que les résultats d'application.

IV.1. Introduction :

Le cancer reste le danger qui menace la sécurité car au problème de son diagnostic.

Dans le présent chapitre nous abordons le problème de classification automatique en utilisant le principe de l'algorithme de MLP en l'appliquant sur la base de **Wisconsin**.

IV.2. La base de Wisconsin :

La base de données « Wisconsin BreastCancer Data base » (WDBC) a fait l'objet de nombreuses études dans le domaine du diagnostic du cancer du sein. La thématique de recherche de créateurs de cette base de données Wolberg, Street Heisy et Managasian , était l'automatisation du diagnostic par le traitement d'images, l'apprentissage automatique ; cette base a été créer par le travail clinique de Dr WilliamH.Wolberg à l'université de Wisconsin à Madison en 1992.

Notre choix de la base de Wisconsin a été motivé par sa large utilisation ce qui nous permette de se situé par rapport aux autres travaux et de comparer nos résultats avec d'autres.

Cette base est constituée de 699 instances distribuées sur deux classes : 458 cas bénins et 241 cas malins.

Chaque instance comporte la valeur de chacun des 11 attributs **Tab IV.1:**

Numéro	Attribut	Domaine
01	Sample code number	Id number
02	ClumpThickness	1-10
03	Uniformity of Cell Size	1-10
04	Uniformity of Cell Shape	1-10
05	Marginal Adhesion	1-10
06	Single EpithelialCell Size	1-10
07	BareNuclei	1-10
08	BlandChromatin	1-10
09	Normal Nucleoli	1-10
10	Mitoses	1-10

IV.3. Les taux de classification :

Les taux de classification et d'erreurs permettent d'évaluer la qualité du classifieur par rapport au problème pour lequel il a été conçu. Ces taux sont évalués grâce à une base de test qui contient des formes étiquetées par leur classe réelle d'appartenance comme celles utilisées pour l'apprentissage afin de pouvoir vérifier les réponses du classifieur.

Pour que l'estimation du taux de reconnaissance soit la plus fiable possible, il est important que le classifieur n'ait jamais utilisé les échantillons de cette base pour faire son apprentissage, de plus cette base de test doit être suffisamment représentative du problème de classification.

Les performances en termes de taux de classification sont alors déterminées en présentant au classifieur chacun des exemples de la base de test et en comparant la classe donnée en résultat à la vraie classe.

Le taux de classification correcte est défini par :

❖ Mesures de la qualité (classification supervisée à 2 classes)

- Matrice de confusion :

	Classé +	Classé -
Exemple +	V_p	F_p
Exemple -	F_n	V_n

- Erreur : $\text{accuracy} = \frac{F_p + F_n}{F_p + F_n + V_p + V_n}$

V_p : Vrai Positif : nombre de positifs classés positifs.

V_n : Vrai Négatif : nombre de négatifs classés négatifs.

F_p : Faux Positif : nombre de négatifs classés positifs.

F_n : Faux Négatif : nombre de positifs classés négatifs.

L'erreur (ou taux d'erreur) ne fait pas de distinction entre les erreurs : pas forcément une bonne mesure de qualité d'un apprentissage.

IV.4.1. Les ressources physiques :

- ❖ Processeur Intel ® Core (TM) i3 CPU M380 @ 2.53 GHz
- ❖ Une Mémoire vive d'une capacité 4.00 Go

IV.4.2. Les ressources logicielles :

- ❖ Système d'exploitation : Windows 7
- ❖ Editeur utilisé est le Netbeans Version 8.0

Notre logiciel est écrit en NetBeans IDE est un environnement de développement intégré (EDI) de premier ordre pour Windows. Le projet NetBeans consiste en un EDI Open Source et en une plate-forme d'application, il est un outil pour les programmeurs pour écrire, compiler, déboguer et déployer des programmes et de créer rapidement des applications Web, d'entreprise, de bureau et mobiles à l'aide de la plate-forme Java, Il est écrit en Java - mais peut supporter n'importe quel langage de programmation. Il y a également un grand nombre de modules pour étendre l'EDI NetBeans.

Le langage **Java** est un langage de programmation informatique orienté objet créé par James Gosling et Patrick Naughton, il permet de réaliser de façon très simple l'interface des applications et de relier aisément le code utilisateur aux événements Windows, quelle que soit leur origine.

Il repose sur un ensemble très complet de composants visuels prêts à l'emploi. La quasi-totalité des contrôles de Windows (boutons, boîtes de saisies, listes déroulantes, menus et autres barres d'outils) y sont représentés, regroupés par famille.

Java possède un langage de programmation à la fois puissant et simple d'utilisation, il permet d'exprimer les problèmes et les solutions d'une façon aisée. Il s'impose dans le monde universitaire et industriel comme un outil puissant de réalisation des applications interactives.

- Java est aujourd'hui un langage aussi rapide que le c++ pourvu qu'on ne l'utilise pas pour une application très lourde (jeux en ligne, logiciel de traitement d'image, encodage vidéo etc...) Java est organisé, il contient des classes bien conçues et bien réparties.
- Java est connu et donc plus de chance de trouver des développeurs Java; pour

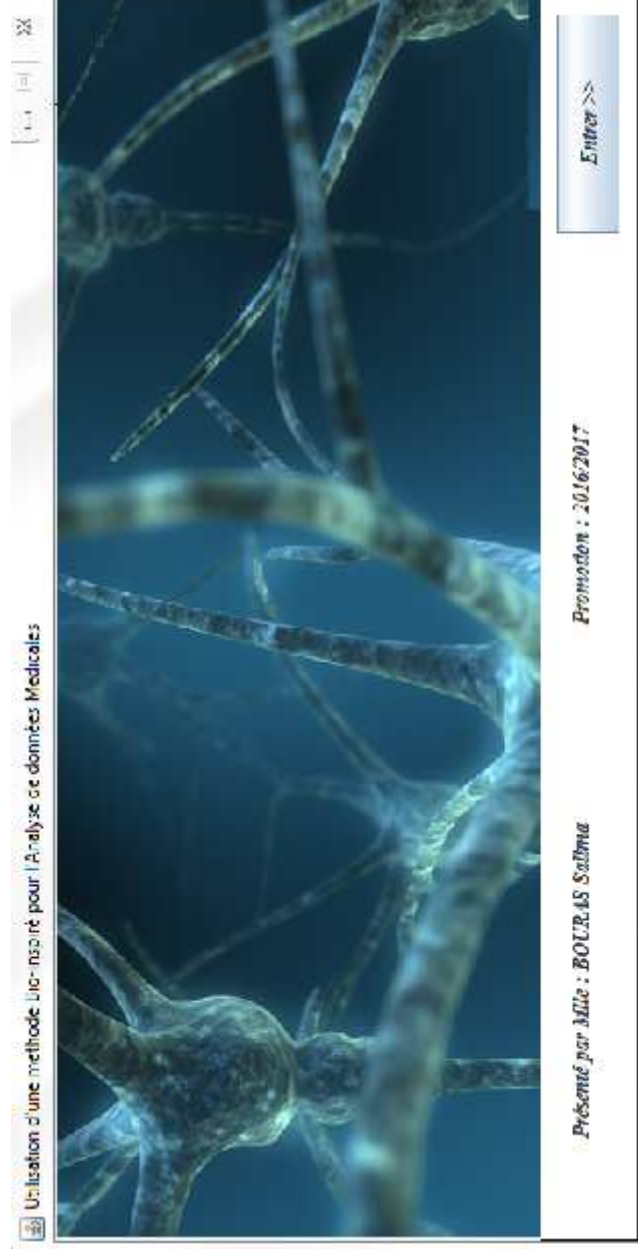
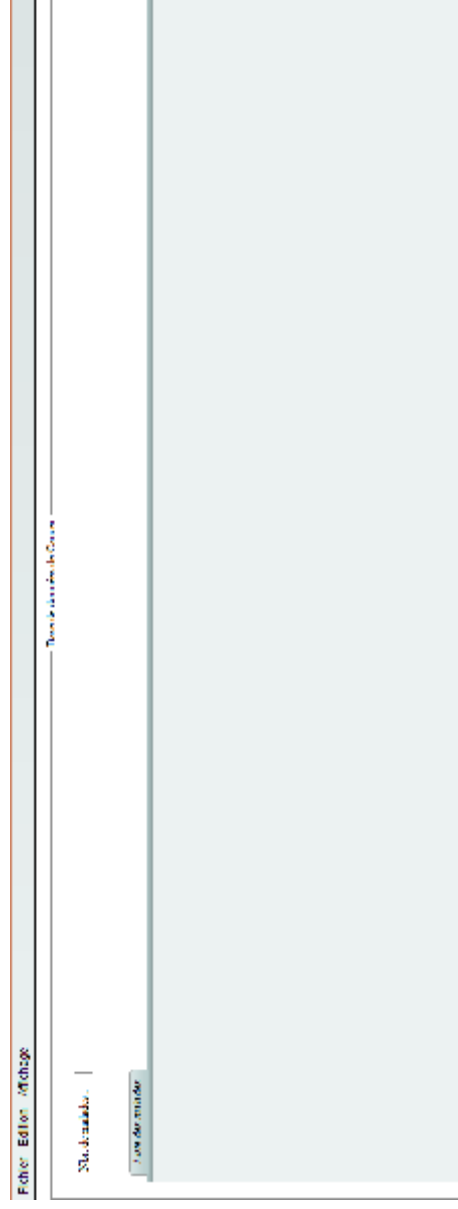


Figure IV.1. 1a page d'accueil

- Puis en cliquant sur « Entrer » il va être affiché cette fenêtre, qui elle est la fenêtre principale.



- En appuyant sur le bouton « fichier » puis « charger » on va charger la base de données.

- En appuyant sur le bouton « Edition » puis « Paramètres MLP », la fenêtre ci-dessous va être affichée, c'est la fenêtre qui contient les paramètres de le réseau MLP et permet de faire l'application leur algorithme sur la base WBCD.

Paramètre MLP

Paramètres

Nb. de couche d'entrée :

Nb. de couche cachée :

Nb. de couche de sortie :

Taux d'apprentissage : %

ϵ :

Nbr. Iteration :

Générer Liste d'apprentissage

MLP

Apprentissage **Tests**

Propagation MLP

Lancer Apprentissage **Grapher réseau**

Temps d'apprentissage :

Nbr. Iteration effectuée :

Figure IV.5. Paramètres de MLP

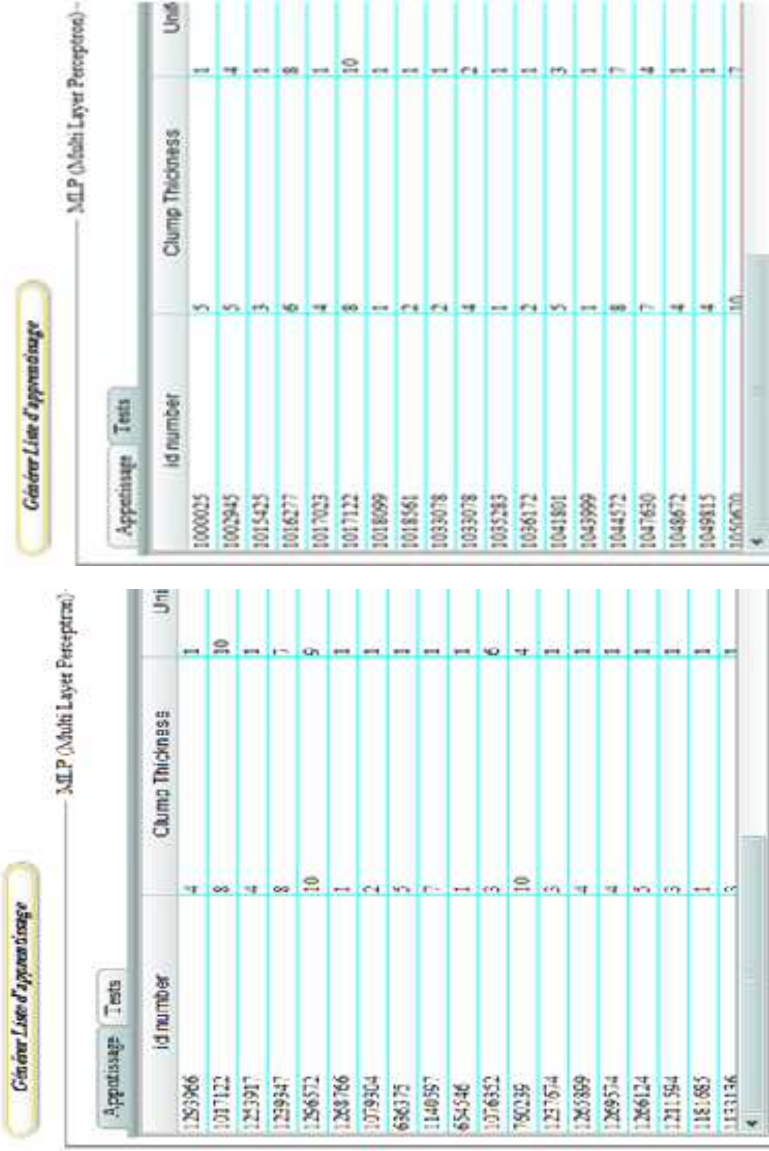
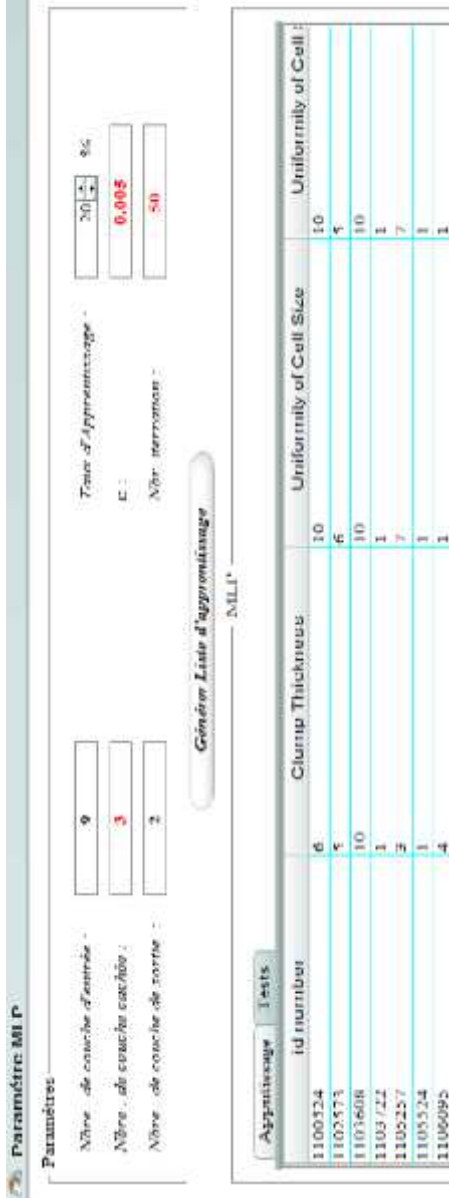


Figure IV.6. Les deux phases d'apprentissage et test.

➤ En cliquant sur « Lancer Apprentissage »



- Pour afficher le graphe, en appuyant sur « Graphe erreur »

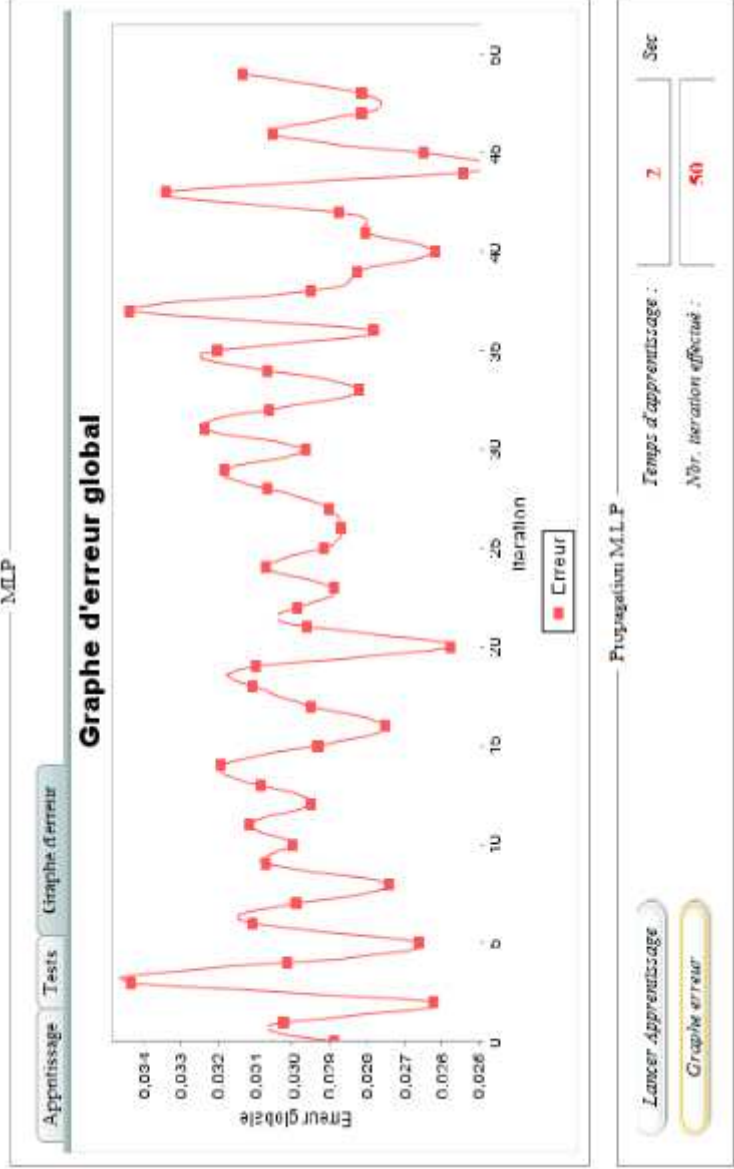


Figure IV.8. Affichage de graphe d'erreur global.

- Cette fenêtre est pour le test, en cliquant sur le bouton « Lancer le test de reconnaissance » pour obtenir le résultat.

- ✓ Couche de sortie qui contient 2 neurones éventuellement selon le nombre de classes.
- Après le déroulement du programme pour **50** itérations, les résultats sont présentés dans le tableau suivant :

Classification Test			
Pourcentage(%) des exemples	% d'exemples Reconnus	% d'exemples Non Reconnu	Taux d'erreur
80%	67,33%	7.15%	25.54%

Table IV.2. Les résultats.

Discussion :

En remarquant que les biens classés est supérieur que les males classés, on peut lier ce résultat a la propriété d'adaptabilité des réseaux de neurones.

IV.6.1 Etude de comparaison entre les méthodes :

On va faire une comparaison entre deux méthodes le MLP et Clonclass (l'algorithme de system immunitaire artificiel), Les deux approches sont favorables et fiables et aussi simplifier les opérations de classification. Pour les deux méthodes la

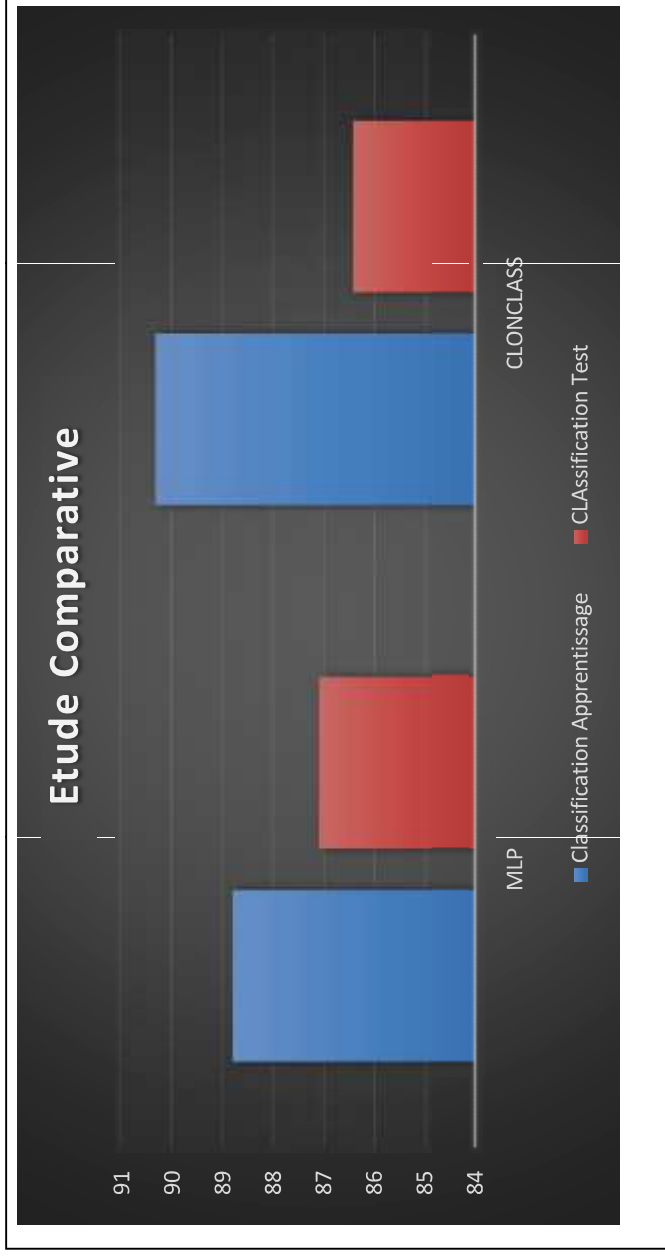


Figure .IV.10 Diagramme de comparaison entre MLP et Clonclass.

IV.7.Conclusion :

Dans ce chapitre nous avons présenté l'approche des Réseaux de neurones qui sont appliqué sur la base de données de cancer WDBC. Selon les résultats obtenus nous pouvons confirmer l'efficacité et la capacité de cette approche pour donner des solutions aux différents problèmes en changeant évidemment quelques paramètres donnés.

Conclusion générale

Conclusion générale :

L'évolution rapide de l'informatique fut associée à une montée éclatante de problèmes informatiques de tous genres. La technologie des ordinateurs s'est développée à un rythme effréné, rendant possible le traitement de problèmes de plus en plus compliqués. Mais les problèmes sont devenus tellement complexes qu'il n'est plus possible à l'être humain de compter uniquement sur la puissance de son ordinateur ni même sur celle de son super calculateur.

La complexité des problèmes actuels est due à l'hétérogénéité et la diversité des acteurs en jeu (humains, systèmes, appareils électroniques, etc.), à la masse importante des données à manipuler et l'augmentation de leur volume, à la distribution des informations manipulées ainsi qu'à la dynamique des environnements dans lesquels ils sont plongés. En effet, dans un environnement dynamique, des événements imprévisibles tels que l'introduction de nouvelles entités, la modification des contraintes, des pannes ou même la détection de situations non prévues à la conception du système peuvent apparaître. Généralement, un problème complexe peut être exprimé sous la forme d'un problème d'optimisation, dans lequel on définit une fonction objectif que l'on recherche à minimiser ou à maximiser selon le contexte.

Différentes techniques de résolution ont été développées pour résoudre ses problèmes complexes. La nature a été toujours une source d'inspiration de nouvelles méthodes de calcul. Ainsi, la théorie de l'évolution a inspiré les algorithmes évolutionnaires, l'étude de l'organisation de groupes d'animaux a donné naissance aux méthodes d'optimisation par essais de particules.

Les méthodes bio-inspirées ont prouvées leur efficacité pour résoudre des problèmes complexes divers, dans plusieurs domaines de recherche et d'ingénierie. Mais, malgré leur grand succès connu cette dernière décennie Le travail de ce projet a mis l'accent sur l'utilité de l'hybridation des méthodes bio-inspirés devant la complexité grandissantes des problèmes actuels.

Dans ce mémoire nous avons présenté la méthode qui est utilisée afin de traiter le problème du diagnostic du cancer. Au départ nous avons étudié le principe de réseaux de neurones qui est un domaine de recherche comble les domaines de l'immunologie, de l'informatique et l'ingénierie. Ensuite, nous avons décrit globalement leurs différents modèles.

Références

- [1] http://www.sfm.org/data/generateur/generateur_fiche/ Date de consultation : 10/03/2016.
- [2] <http://bbf.ensib.fr/> Date de consultation : 25/04/2017.
- [3] Laurent HARTERT ; « Reconnaissance des formes dans un environnement dynamique appliqué au diagnostic et au suivi des systèmes évolutifs » ; thèse de doctorat ; Novembre 2010.
- [4] <https://tel.archives-ouvertes.fr/> Date de consultation : 16/12/2016.
- [5] Marref Nadia ; « Apprentissage Incrémental & Machines à Vecteurs Supports » ; Mémoire de Magister, Batna ; 2015.
- [6] Mohamed-Rafik bouguelia ; « Classification et apprentissage actif a partir d'un flux de données évolutif en présence d'étiquetage incertain » ; Thèse de Doctorat ; 25 mars 2015
- [7] Cours « Classification supervisée et non supervisée ; 14/02/2014, Date de consultation : 10/12/2016
- [8] Cours « Informatique orientation logiciels classification automatique » ; Date de consultation : 11/12/2016
- [9] Cours « Méthodes en classification automatique » ; Date de consultation : 13/12/2016
- [10] <http://dspace.univ-tlemcen.dz/bitstream>; Date de consultation : 14/12/2016.
- [11] <http://damz-fish.fr/download/M%C3%A9moire.pdf>; Date de consultation : 13/01/2016
- [12] « Systèmes Bio-inspirés » ; Andrés Pérez-Urbe ; 2004 ;Date de consultation : 13/12/2016

- [16] **TAOUI Nassima**, « Etude comparative entre méthode Neuronale et Immunitaire appliqué sur une base de cancer du sein » ; Diplôme de Master ; 2016.
- [17] <http://www.touzet.org/Claude/Web-Fac> Date de consultation : 15/01/2017
- [18] **Melle MENGHOUR Kamilia** ; « Cycle Approches Bio-inspirées pour la Sélection d'Attributs » ; thèse de Doctorat ; 2013.
- [19] **Maxime BOMBRUN Abdoulaye SENE** ; « L'optimisation par essai particulière pour des problèmes d'ordonnement » ; le 24/03/2011.
- [20] <http://sis.univ-tln.fr/~tollari/TER/AlgoGen1/node5.html>; Date de consultation : 25/12/2016
- [21] <https://tel.archives-ouvertes.fr> ; Date de consultation: 29/12/2016.
- [22] <http://www.baladesentomologiques.com> ; Date de consultation: 01/01/2016.
- [23] <http://www.isima.fr> ; Date de consultation : 29/12/2016.
- [24] <https://www.labri.fr/perso/nrougier/downloads/Perceptron.pdf> Date de consultation : 15/12/2016.
- [25] **Nicolas P. Rougier** ; Perceptron simple Perceptron multi-couches; Master 2 ;- Sciences Cognitives ; Université de Bordeaux.
- [26] <https://msirtadgrp2.files.wordpress.com/2015/02/perceptron-multicouche.pdf>
- [27] **B.Gosselin** : « Application De Réseaux De Neurones Artificiels A La Reconnaissance Automatique De Caractères Manuscrits », Thèse de Doctorat, Faculté Polytechnique de Mons, 1995-1996.
- [28] <http://constellation.uqac.ca>; Date de consultation: 02/05/2016
- [29] **F. BADRAN et S. THIRIA** « Les Perceptrons Multicouches de la régression non linéaire aux problèmes inverses », Laboratoire d'Océanographie Dynamique et de Climatologie, PARIS, France.