



L'utilisation de l'intelligence de artificielle dans la découverte de nouveaux candidats-médicaments et dans l'imagerie médicale

Mazari Zemihi

► To cite this version:

Mazari Zemihi. L'utilisation de l'intelligence de artificielle dans la découverte de nouveaux candidats-médicaments et dans l'imagerie médicale. Sciences pharmaceutiques. 2020. hal-03298153

HAL Id: hal-03298153

<https://hal.univ-lorraine.fr/hal-03298153>

Submitted on 3 Nov 2022

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



AVERTISSEMENT

Ce document est le fruit d'un long travail approuvé par le jury de soutenance et mis à disposition de l'ensemble de la communauté universitaire élargie.

Il est soumis à la propriété intellectuelle de l'auteur. Ceci implique une obligation de citation et de référencement lors de l'utilisation de ce document.

D'autre part, toute contrefaçon, plagiat, reproduction illicite encourt une poursuite pénale.

Contact : ddoc-thesesexercice-contact@univ-lorraine.fr

LIENS

Code de la Propriété Intellectuelle. articles L 122. 4

Code de la Propriété Intellectuelle. articles L 335.2- L 335.10

http://www.cfcopies.com/V2/leg/leg_droi.php

<http://www.culture.gouv.fr/culture/infos-pratiques/droits/protection.htm>



UNIVERSITÉ
DE LORRAINE



FACULTÉ
DE PHARMACIE

UNIVERSITE DE LORRAINE

2020

FACULTE DE PHARMACIE

THESE

Présentée et soutenue publiquement

Le 16/09/2020, sur un sujet dédié à :

L'utilisation de l'intelligence artificielle dans la découverte de nouveaux candidats-médicaments et dans l'imagerie médicale.

pour obtenir

le Diplôme d'Etat de Docteur en Pharmacie

par Mazari ZEMIHI, né le 03/08/1977

Membres du Jury

Président : M. Michel BOISBRUN, Maître de conférences titulaire HDR

Directeur : M. Gérard MONARD, Professeur

Co-directeur : Mme Alexandrine LAMBERT, Maître de conférences

Juges : M. Nicolas VERAN, Radiopharmacien

M. Antoine CAROF, Maître de conférences

UNIVERSITÉ DE LORRAINE
FACULTÉ DE PHARMACIE
Année universitaire 2019-2020

DOYEN

Raphaël DUVAL

Vice-Doyen

Julien PERRIN

Directrice des études

Marie SOCHA

Conseil de la Pédagogie

Présidente, Brigitte LEININGER-MULLER

Vice-Présidente, Alexandrine LAMBERT

Collège d'Enseignement Pharmaceutique Hospitalier

Présidente, Béatrice DEMORE

Commission Prospective Facultaire

Président, Christophe GANTZER

Vice-Président, Jean-Louis MERLIN

Commission de la Recherche

Présidente, Caroline GAUCHER

Chargés de Mission

Communication

Innovation pédagogique

Référente ADE

Référente dotation sur projet (DSP)

Référent vie associative

Aline BONTEMPS

Alexandrine LAMBERT

Virginie PICHON

Marie-Paule SAUDER

Arnaud PALLOTTA

Responsabilités

Filière Officine

Filière Industrie

Filière Hôpital

Pharma Plus ENSIC

Pharma Plus ENSAIA

Pharma Plus ENSGSI

Cellule de Formation Continue et Individuelle

Commission d'agrément des maîtres de stage

ERASMUS

Caroline PERRIN-SARRADO

Julien GRAVOULET

Isabelle LARTAUD,

Jean-Bernard REGNOUF de VAINS

Béatrice DEMORE

Marie SOCHA

Jean-Bernard REGNOUF de VAINS

Xavier BELLANGER

Igor CLAROT

Luc FERRARI

François DUPUIS

Mihayl VARBANOV

DOYENS HONORAIRES

Chantal FINANCE

Francine PAULUS

Claude VIGNERON

PROFESSEURS EMERITES

Jeffrey ATKINSON

Max HENRY

Pierre LEROY

Philippe MAINCENT

Claude VIGNERON

PROFESSEURS HONORAIRES

Jean-Claude BLOCK
 Pierre DIXNEUF
 Chantal FINANCE
 Marie-Madeleine GALTEAU
 Thérèse GIRARD
 Pierre LABRUDE
 Vincent LOPPINET
 Alain NICOLAS
 Janine SCHWARTZBROD
 Louis SCHWARTZBROD

ASSISTANTS HONORAIRES

Marie-Catherine BERTHE
 Annie PAVIS

MAITRES DE CONFERENCES HONORAIRES

Monique ALBERT
 Mariette BEAUD
 François BONNEAUX
 Gérard CATAU
 Jean-Claude CHEVIN
 Jocelyne COLLOMB
 Bernard DANGIEN
 Marie-Claude FUZELLIER
 Françoise HINZELIN
 Marie-Hélène LIVERTOUX
 Bernard MIGNOT
 Blandine MOREAU
 Dominique NOTTER
 Francine PAULUS
 Christine PERDIAKIS
 Marie-France POCHON
 Anne ROVEL
 Gabriel TROCKLE
 Maria WELLMAN-ROUSSEAU
 Colette ZINUTTI

ENSEIGNANTS

Section CNU

*

Discipline d'enseignement

PROFESSEURS DES UNIVERSITES - PRATICIENS HOSPITALIERS

Danièle BENSOUSSAN-LEJZEROWICZ	82	Thérapie cellulaire
Béatrice DEMORE	81	Pharmacie clinique
Jean-Louis MERLIN	82	Biologie cellulaire
Jean-Michel SIMON	81	Economie de la santé, Législation pharmaceutique
Nathalie THILLY	81	Santé publique et Epidémiologie

PROFESSEURS DES UNIVERSITES

Ariane BOUDIER	85	Chimie Physique
Christine CAPDEVILLE-ATKINSON	86	Pharmacologie
Igor CLAROT	85	Chimie analytique
Joël DUCOURNEAU	85	Biophysique, Acoustique, Audioprothèse
Raphaël DUVAL	87	Microbiologie clinique
Béatrice FAIVRE	87	Hématologie, Biologie cellulaire
Luc FERRARI	86	Toxicologie
Pascale FRIANT-MICHEL	85	Mathématiques, Physique
Christophe GANTZER	87	Microbiologie
Frédéric JORAND	87	Eau, Santé, Environnement
Isabelle LARTAUD	86	Pharmacologie
Dominique LAURAIN-MATTAR	86	Pharmacognosie
Brigitte LEININGER-MULLER	87	Biochimie
Patrick MENU	86	Physiologie
Jean-Bernard REGNOUF de VAINS	86	Chimie thérapeutique
Bertrand RIHN	87	Biochimie, Biologie moléculaire

MAITRES DE CONFÉRENCES DES UNIVERSITÉS - PRATICIENS HOSPITALIERS

Alexandre HARLE	82	<i>Biologie cellulaire oncologique</i>
Julien PERRIN	82	<i>Hématologie biologique</i>
Loïc REPPÉL	82	<i>Biothérapie</i>
Marie SOCHA	81	<i>Pharmacie clinique, thérapeutique et biotechnique</i>

MAITRES DE CONFÉRENCES

Xavier BELLANGER ^H	87	<i>Parasitologie, Mycologie médicale</i>
Emmanuelle BENOIT ^H	86	<i>Communication et Santé</i>
Isabelle BERTRAND ^H	87	<i>Microbiologie</i>
Michel BOISBRUN ^H	86	<i>Chimie thérapeutique</i>
Cédric BOURA ^H	86	<i>Physiologie</i>
Sandrine CAPIZZI	87	<i>Parasitologie</i>
Antoine CAROF	85	<i>Informatique</i>
Sébastien DADE	85	<i>Bio-informatique</i>
Dominique DECOLIN	85	<i>Chimie analytique</i>
Natacha DREUMONT ^H	87	<i>Biochimie générale, Biochimie clinique</i>
Florence DUMARCAY ^H	86	<i>Chimie thérapeutique</i>
François DUPUIS ^H	86	<i>Pharmacologie</i>
Reine EL OMAR	86	<i>Physiologie</i>
Adil FAIZ	85	<i>Biophysique, Acoustique</i>
Anthony GANDIN	87	<i>Mycologie, Botanique</i>
Caroline GAUCHER ^H	86	<i>Chimie physique, Pharmacologie</i>
Stéphane GIBAUD ^H	86	<i>Pharmacie clinique</i>
Thierry HUMBERT	86	<i>Chimie organique</i>
Olivier JOUBERT ^H	86	<i>Toxicologie, Sécurité sanitaire</i>

ENSEIGNANTS (suite)

Section CNU

*

Discipline d'enseignement

Alexandrine LAMBERT	85	Informatique, Biostatistiques
Julie LEONHARD	86/01	Droit en Santé
Christophe MERLIN ^H	87	Microbiologie environnementale
Maxime MOURER ^H	86	Chimie organique
Coumba NDIAYE	86	Epidémiologie et Santé publique
Arnaud PALLOTTA	85	Bioanalyse du médicament
Marianne PARENT	85	Pharmacie galénique
Caroline PERRIN-SARRADO	86	Pharmacologie
Virginie PICHON	85	Biophysique
Sophie PINEL ^H	85	Informatique en Santé (e-santé)
Anne SAPIN-MINET ^H	85	Pharmacie galénique
Marie-Paule SAUDER	87	Mycologie, Botanique
Guillaume SAUTREY	85	Chimie analytique
Rosella SPINA	86	Pharmacognosie
Sabrina TOUCHET	86	Pharmacochimie
Mihayl VARBANOV	87	Immuno-Virologie
Marie-Noëlle VAULTIER	87	Mycologie, Botanique
Emilie VELOT ^H	86	Physiologie-Physiopathologie humaines
Mohamed ZAIOU ^H	87	Biochimie et Biologie moléculaire

PROFESSEUR ASSOCIE

Julien GRAVOULET	86	Pharmacie clinique
------------------	----	--------------------

PROFESSEUR AGREGÉ

Christophe COCHAUD	11	Anglais
--------------------	----	---------

^H Maître de conférences titulaire HDR* Disciplines du Conseil National des Universités :

80 : Personnels enseignants et hospitaliers de pharmacie en sciences physico-chimiques et ingénierie appliquée à la santé

81 : Personnels enseignants et hospitaliers de pharmacie en sciences du médicament et des autres produits de santé

82 : Personnels enseignants et hospitaliers de pharmacie en sciences biologiques, fondamentales et cliniques

85 ; Personnels enseignants-chercheurs de pharmacie en sciences physico-chimiques et ingénierie appliquée à la santé

86 : Personnels enseignants-chercheurs de pharmacie en sciences du médicament et des autres produits de santé

87 : Personnels enseignants-chercheurs de pharmacie en sciences biologiques, fondamentales et cliniques

11 : Professeur agrégé de lettres et sciences humaines en langues et littératures anglaises et anglo-saxonnes

SERMENT DE GALIEN

En présence des Maitres de la Faculté, je fais le serment :

D'honorer ceux qui m'ont instruit(e) dans les préceptes de mon art et de leur témoigner ma reconnaissance en restant fidèle aux principes qui m'ont été enseignés et d'actualiser mes connaissances

D'exercer, dans l'intérêt de la santé publique, ma profession avec conscience et de respecter non seulement la législation en vigueur, mais aussi les règles de Déontologie, de l'honneur, de la probité et du désintéressement ;

De ne jamais oublier ma responsabilité et mes devoirs envers la personne humaine et sa dignité

En aucun cas, je ne consentirai à utiliser mes connaissances et mon état pour corrompre les mœurs et favoriser des actes criminels.

De ne dévoiler à personne les secrets qui m'auraient été confiés ou dont j'aurais eu connaissance dans l'exercice de ma profession

De faire preuve de loyauté et de solidarité envers mes collègues pharmaciens

De coopérer avec les autres professionnels de santé

Que les Hommes m'accordent leur estime si je suis fidèle à mes promesses.

Que je sois couvert(e) d'opprobre et méprisé(e) de mes confrères si j'y manque.

« LA FACULTE N'ENTEND DONNER AUCUNE APPROBATION,
NI IMPROBATION AUX OPINIONS EMISES DANS LES
THESES, CES OPINIONS DOIVENT ETRE CONSIDEREES
COMME PROPRES A LEUR AUTEUR ».

Remerciements

Je remercie tout d'abord les membres de mon jury pour avoir pris le temps de lire cette thèse.

Je tiens ensuite à remercier tout particulièrement le Professeur Monard d'avoir accepté d'encadrer ce sujet de thèse.

Je remercie également Mme Lambert, Maître de conférences à la faculté de Pharmacie, d'avoir bien voulu être co-directeur de thèse.

Je remercie les professeurs de la faculté de Pharmacie de Nancy pour la transmission de leur savoir.

Je remercie les collègues de la promotion avec qui j'ai pu discuter et travailler.

Mes remerciements affectueux à mes parents, à mes frères et à ma sœur, et à mes amis pour leurs encouragements.

Table des matières

Introduction.....	6
A Intelligence artificielle	8
1 Histoire de l'intelligence artificielle	8
2 Définition de l'intelligence artificielle	9
3 Principes de base de l'apprentissage automatique.....	11
3.1 Algorithmes d'apprentissage.....	11
3.2 Apprentissage supervisé.....	12
3.3 Apprentissage non supervisé.....	13
3.4 Apprentissage semi supervisé	14
3.5 Apprentissage multi-instancés	14
3.6 Apprentissage par renforcement.....	14
3.7 Apprentissage par transfert.....	14
4 Algorithmes de régression et de classification	15
5 Mesure de la performance (P) des algorithmes	16
6 Généralisation	19
7 Jeu d'entraînement, de test et de validation	19
7.1 Les données	20
7.2 MoleculeNet.....	21
7.3 Logiciels utilisés en pratique	23
8 Méthodologie de mise en place d'un algorithme	23
B Applications de l'IA	25
1 Drug Discovery.....	25
1.1 Contexte économique	25
1.2 Description du processus de R&D	26
1.3 Les différentes approches dans la découverte d'un médicament	28
1.4 La représentation des molécules	30
1.5 <i>Deep learning</i> ou apprentissage profond	32
1.5.1 Application de l'apprentissage profond dans la détection de toxicophores	34

1.5.2 AlphaFold : prédicteur de la structure protéique	36
1.5.3 La recherche d'un traitement contre le virus à Ebola	37
1.5.4 Covid-19 et IA ?	38
1.5.5 Découverte d'un nouvel antibiotique potentiel	43
1.5.6 ClinicalTrial.gov et IA	46
2 Application à l'imagerie médicale	47
2.1 Définition d'une image	47
2.2 Les réseaux de neurones convolutifs	48
2.3 Vers une médecine personnalisée	52
2.4 Limites actuelles de l'IA	52
3 Résumé des applications de l'IA.....	54
C Ethique et réglementation.....	56
1 L'éthique en IA	56
2 Données personnelles.....	57
2.1 Le consentement	58
2.2 La confidentialité des données.....	59
2.3 La sécurisation des données	60
3 Responsabilité juridique en cas de mauvais diagnostic	61
4 Explicabilité des algorithmes	62
5 Acceptabilité des algorithmes	63
6 Impact sur l'emploi	63
Conclusion et perspectives	67
Bibliographie	70
Annexes.....	78

Table des figures

Figure 1 Diagramme de Venn de l'IA	10
Figure 2 Schéma d'un algorithme de régression et de classification	15
Figure 3 Représentation d'une courbe de ROC	18
Figure 4 Validation croisée à 5 parties, chaque point appartient à un des cinq jeux de test (en blanc) et aux quatre autres jeux d'entraînements (en orange), d'après la source [16].....	20
Figure 5 Schéma simplifié du processus de développement d'un projet d'IA	24
Figure 6 Contenu et durée des différentes phases de recherche et développement d'après la référence [23]	27
Figure 7 Exemple de graphe avec 6 sommets reliés par des arêtes	30
Figure 8 Molécule de benzène transformée en graphe moléculaire (à droite) d'après [32]...	31
Figure 9 Représentation d'un neurone formel d'après la référence [39]	33
Figure 10 Topologie basique d'un réseau de neurones artificiels	33
Figure 11 Détection de pharmacophores toxiques par un réseau de neurones profond d'après la source [40].....	35
Figure 12 Processus de prédiction de la structure protéique extraite de la publication [45]..	36
Figure 13 Modélisation en 3D du coronavirus d'après CDC (Centers for Disease Control and Prevention)	40
Figure 14 Schéma résumant l'article de cette référence [55].....	45
Figure 15 Structure en 2D de la molécule d'halicine	46
Figure 16 Image de CT-scanner représentée en pixels.....	48
Figure 17 Exemple d'architecture d'un CNN d'après la référence [67]	49
Figure 18 Filtres appliqués à une image en entrée d'après la référence [68].....	50
Figure 19 Exemples de sortie du logiciel Transpara™	51
Figure 20 Méthodes d'apprentissage automatique et leurs applications dans la découverte de médicaments et en imagerie d'après [80].....	54

Liste des tableaux

Tableau I Taxonomie des algorithmes d'apprentissage automatique [6]	16
Tableau II Matrice de confusion	17
Tableau III Matrice de confusion d'un cas d'usage fictif.....	18
Tableau IV Matrice de combinaison de molécules	21
Tableau V Propriétés prédéfinies des molécules à générer	42
Tableau VI Tableau de synthèse récapitulatif des applications mentionnées dans ce document.....	55

Abréviations et acronymes

AAK1	<i>AP2-Associated Protein Kinase 1</i>
ACE2	<i>Angiotensin-Converting Enzyme II</i>
AT2	<i>Angiotensin Type II</i>
AMM	Autorisation de Mise sur le Marché
ASCII	<i>American Standard Code for Information Interchange</i>
BATX	Baidu, Alibaba, Tencent, Xiaomi
BI-RADS	<i>Breast Imaging Reporting And Data System</i>
CCNE	Comité Consultatif National d'Éthique pour les sciences de la vie et de la santé
CEPD	Contrôleur Européen de la Protection des Données
CERNA	Commission de réflexion sur l'Éthique de la Recherche en sciences et technologies du Numérique d'Allistene
CNIL	Commission Nationale Informatique et Libertés
CNOM	Conseil National de l'Ordre des Médecins
Covid-19	<i>CoronaVirus Disease 2019</i>
CT	<i>Computerized Tomography</i>
DFT	<i>Density Functional Theory</i>
DNN	<i>Deep Neural Network</i>
ECFPs	<i>Extended Connectivity Fingerprints</i>
EMA	<i>European Medicines Agency</i>
FDA	<i>Food and Drug Administration</i>
GAFAMI	Google, Amazon, Facebook, Apple, Microsoft et IBM
GPU	<i>Graphics Processing Units</i>
HAS	Haute Autorité de Santé

IA	Intelligence artificielle
INDS	Institut National des Données de Santé
IRM	Imagerie par Résonance Magnétique
IUPAC	<i>International Union of Pure and Applied Chemistry</i>
JAK	Janus Kinase
LBVS	<i>Ligand-based virtual screening</i>
MERS-Cov	<i>Middle East Respiratory Syndrome Coronavirus</i>
NIH	<i>National Institutes of Health</i>
PCR	Polymerase Chain Reaction
PMC	Perceptron MultiCouche
PMDA	<i>Pharmaceuticals and Medical Devices Agency</i>
QSAR	<i>Quantitative Structure-Activity Relationships</i>
R&D	Recherche & développement
RNA	Réseau de Neurones Artificiels
RNP	Réseau de Neurones Profonds
ROC	<i>Receiver Operating Characteristic</i>
SBVS	<i>Structure-based virtual screening</i>
SMILES	<i>Simplified Molecular Input Line Entry System</i>
SNDS	Système National des Données de Santé
SOSA	<i>Selective Optimization of Side Activities</i>
SRAS	Syndrome Respiratoire Aigüe Sévère

Introduction

- Ok Google, l'adresse du campus Brabois de santé ?
- Et voilà, 7 avenue de la Forêt de Haye, 54500 Vandœuvre-lès-Nancy.

Cet exemple de dialogue avec l'assistance vocale de notre smartphone illustre le fait que nous utilisons l'intelligence artificielle (IA) dans notre vie quotidienne. Notre requête a été formulée en langage naturel et a été comprise par un algorithme de reconnaissance vocale. L'intelligence artificielle est déjà présente dans les objets de grande consommation tels que les smartphones, les enceintes connectées, les voitures autonomes. Cet exemple montre également l'importance des entreprises comme Google dans notre vie personnelle et professionnelle. Ces entreprises du numérique s'intéressent également au domaine de la santé et aux données qu'ils peuvent récupérer des utilisateurs à travers leur moteur de recherche, leurs sites de réseaux sociaux, leurs objets connectés. A l'instar des autres domaines (transport, internet), la technologie permet de « disrupter » le monde de la santé, en automatisant des tâches répétitives et fastidieuses effectuées manuellement avec une promesse de réduction du temps et de coûts. Le changement de paradigme dans le processus de découverte et développement de nouvelles molécules dans l'industrie pharmaceutique est également en cours. Certains dispositifs médicaux embarquent déjà cette technologie. En médecine, le diagnostic de certaines pathologies est effectué par une machine, en particulier dans la lecture des radiographies. Ce changement s'inscrit dans une continuité de la recherche scientifique. L'intelligence artificielle est née presque au même moment que l'informatique vers les années 1950. En tant que domaine scientifique, elle a connu des périodes de recherches intenses et d'autres périodes de désintérêt ou de désengagement de la part des pouvoirs publics ou des sociétés privées. Cependant, nous assistons depuis les années 2000 à un regain d'intérêts. L'augmentation de la puissance de calcul des ordinateurs et la présence d'une quantité massive de données (*big data*) ont permis cet essor.

Cette thèse présente l'essentiel de l'état de l'art des solutions d'IA appliquées aux domaines de la santé actuellement. Le plan de cette thèse s'articule en trois grandes parties. La première partie présente la genèse de l'intelligence artificielle. Nous donnerons une définition des différents concepts de cette technologie. Nous présenterons les principes de base de l'apprentissage automatique et les différents types d'apprentissage. La mesure de la performance de ces algorithmes sera également définie. Nous parlerons de l'importance des données dans les systèmes informatiques d'IA. Le dernier point traité dans cette première partie sera la méthodologie à mettre en place dans l'utilisation d'un algorithme d'apprentissage automatique.

La deuxième partie aborde la mise en application de l'intelligence artificielle dans la découverte et le développement de candidats-médicaments. Les enjeux économiques pour

l'industrie pharmaceutique justifient le recours à des technologies qui permettent d'améliorer la compétitivité. Nous verrons le processus de recherche et découverte de nouvelles molécules. Nous présenterons quelques approches adoptées pour découvrir de nouveaux médicaments.

Dans le domaine du diagnostic médical en particulier dans l'imagerie, nous verrons que l'IA a dépassé le stade des preuves de concept et actuellement en service dans certains hôpitaux. La troisième partie traite de l'éthique et de la réglementation encadrant cette technologie. Nous nous intéresserons en particulier aux données sensibles que sont les données de santé. Nous questionnerons la responsabilité en cas de mauvais diagnostic. Ensuite, nous parlerons de l'explicabilité et de l'acceptabilité des algorithmes d'apprentissage automatique. Enfin, nous évaluerons l'impact que pourrait avoir l'utilisation de cette technologie sur l'emploi.

A Intelligence artificielle

1 Histoire de l'intelligence artificielle

Le terme intelligence artificielle est apparue en 1956 lors d'un séminaire d'été à Dartmouth aux Etats-Unis d'Amérique et c'est à John McCarthy qu'on attribue la paternité du nom de cette nouvelle discipline [1]. Cependant Warren McCulloch et Walter Pitts ont été les précurseurs des travaux appartenant à l'IA [1]. Ils se sont inspirés des connaissances sur la physiologie et le fonctionnement des neurones dans le cerveau, de l'analyse formelle de la logique propositionnelle de Russell et Whitehead, et de la théorie du calcul de Turing [2]. Un modèle de neurones artificiels a été donc défini en 1943 par McCulloch et Pitts [2]. En 1957, Frank Rosenblatt inventa le perceptron, un algorithme d'apprentissage qui simulait le fonctionnement d'un neurone. L'objectif premier du perceptron était de classer des images. Une période de désintérêt et de désengagement financier intervient au début des années 1970, à la suite des échecs de l'IA à l'époque de tenir ses promesses dans le domaine du langage, de la traduction automatique. Les chercheurs ont nommé cette période « l'hiver de l'IA », en anglais « *AI Winter* ». Après une lueur d'espoir avec l'essor des premiers ordinateurs personnels et une puissance de calcul plus importante disponible, l'IA connaît son deuxième « hiver » à la fin des années 1980. C'est l'apparition d'Internet et la quantité massive de données créées dans son sillage au milieu des années 1990 qui a remis l'IA sur le devant de la scène. En 2001, l'apparition des premières cartes graphiques dédiées (*GPU* pour *Graphics Processing Units*) pour la gestion de l'affichage des écrans a parallèlement permis d'augmenter la puissance de calcul des ordinateurs. En effet, les scientifiques ont par exemple pu effectuer des calculs matriciels à l'aide de ces processeurs graphiques programmables. L'accès à ces ressources de calcul a été facilité par la mise à disposition de la bibliothèque CUDA par le fabricant Nvidia [3]. D'autres dates avec une portée symbolique importante sont résumées ci-dessous :

- 1951 : SNARC (*Stochastic Neural Analog Reinforcement Calculator*), Marvin Minsky construit avec un étudiant de Princeton la première machine neuronale avec 40 « synapses ».
- 1966 : premier robot mobile *Shakey the robot*.
- 1972 : reconnaissance vocale, premier appareil commercialisé après 20 ans de recherche.
- 1975 : MYCIN ; système expert d'aide à l'identification des infections aiguës et de recommandation des traitements antibiotiques [4].
- 1989 : reconnaissance d'image à partir des travaux de Yann Le Cun.

- 1997 : Deep Blue d'IBM gagne aux échecs contre le champion du monde Garry Kasparov. La méthode utilisée n'est pas de l'IA, mais cet exemple illustre les capacités de calcul immense des ordinateurs modernes.
- 2005 : DARPA Challenge démontre le caractère opérationnel de la voiture autonome.
- 2012 : La compétition ImageNet de reconnaissance d'objets dans une image est remportée par Geoffrey Hinton et son équipe en utilisant du *deep learning*.
- 2016 : AlphaGO, système conçu par DeepMind bat le champion de GO Lee Sedol.

Après ce bref rappel historique qui permet de resituer le contexte actuel du développement de l'IA et de chasser cette idée que l'IA serait récente, nous allons tenter d'en donner une définition.

2 Définition de l'intelligence artificielle

Dans son article datant de 1950 « *Computing Machinery and Intelligence* », Alan Turing définit les fondements de l'intelligence artificielle et développe un test qui porte son nom dans lequel une personne doit identifier si son interlocuteur est un humain ou une machine. On peut donc définir l'intelligence artificielle comme un ensemble de techniques permettant aux ordinateurs de simuler et d'imiter l'intelligence humaine. Ce diagramme de Venn simplifié (Figure 1), inspiré de cette référence [1] schématise le fait que l'apprentissage profond est un sous-ensemble de l'apprentissage automatique qui est lui-même un élément de l'IA.

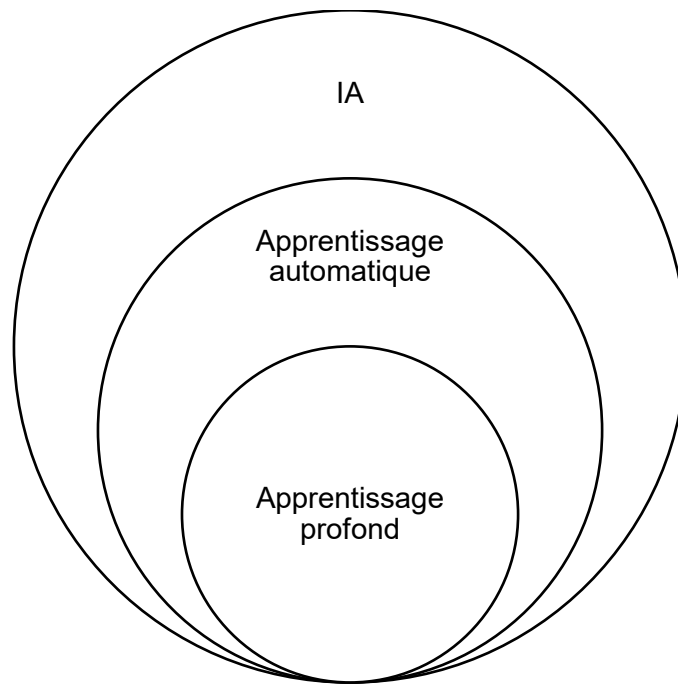


Figure 1 Diagramme de Venn de l'IA

L'IA donne la capacité à une machine de reconnaître une image, de traduire un texte, d'automatiser la conduite d'une voiture ou le pilotage d'un procédé industriel. Il y a différentes approches de l'IA. L'approche cognitive cherche à comprendre les processus des raisonnements humains et à les représenter en algorithmes [5]. Dans l'approche pragmatiste, l'IA est considérée comme une boîte noire avec en entrée les données du problème et en sortie les résultats sans que l'on sache expliquer le fonctionnement de cette boîte [5]. L'approche connexionniste consiste en une modélisation du fonctionnement des neurones d'un cerveau humain [5]. L'IA est constituée de plusieurs briques scientifiques : l'informatique, les mathématiques, les statistiques, les sciences cognitives, la neurobiologie, la psychologie, la philosophie. Cette transdisciplinarité de l'IA fait sa richesse en termes d'innovation mais participe également à sa complexité. Dans la littérature, nous trouvons la distinction entre l'IA forte et l'IA faible. L'IA forte ou IA générale aurait l'ambition d'imiter le cerveau humain avec une conscience et une sensibilité. Ce serait une version améliorée de l'approche connexionniste. Cela reste du domaine de la fiction et sert surtout à alimenter la machine à fantasmes de certains auteurs. L'IA faible qui est utilisée aujourd'hui est spécifique à une tâche et réalise parfois de meilleures performances qu'un humain.

Cette définition succincte de l'IA nous permet d'introduire les principes de base de l'apprentissage automatique.

3 Principes de base de l'apprentissage automatique

3.1 Algorithmes d'apprentissage

Il y a plusieurs expressions similaires pour l'apprentissage machine (*machine learning*) : l'apprentissage automatique, l'apprentissage statistique. Un algorithme d'apprentissage machine est un algorithme qui peut apprendre des données [4]. Arthur Samuel, l'un des pionniers de l'intelligence artificielle, définit en 1959, le *machine learning* comme le champ d'étude visant à donner la capacité à une machine d'apprendre sans être explicitement programmée. En 1997, Tom Mitchell (professeur à l'université de Carnegie Mellon), donne une définition plus précise :

« *A computer program is said to learn from experience E with respect to some class of tasks T and performance measure P, if its performance at tasks T, as measured by P, improves with experience E* ». « On dit qu'un programme informatique apprend de l'expérience E, étant données certaines classes de tâches T et une mesure de performance P, si sa performance de réalisation des tâches T, telle que mesurée par P, s'améliore avec l'expérience E » [6]. Le *machine learning* est une informatique probabiliste basée sur des informations réelles, à l'opposé d'une informatique impérative basée sur des hypothèses [6]. Une tâche T est un élément de la décomposition d'un problème, susceptible d'être représenté par un programme informatique [7]. Les tâches d'apprentissage machine sont décrites en fonction des variables (ou caractéristiques, *features* en anglais) traitées par le système d'apprentissage machine. Une collection de caractéristiques est nommée un exemple qui est représenté mathématiquement comme un vecteur $x \in \mathbb{R}^n$ où chaque entrée x_i du vecteur est une caractéristique. Par exemple, une image est caractérisée par les valeurs des pixels qui la composent. L'apprentissage automatique a la capacité d'accomplir de nombreux types de tâches, comme la classification, la régression, la transcription, la traduction automatique. La mesure de la performance P, qui est spécifique à la tâche T effectuée par le système permet d'évaluer la capacité d'un algorithme d'apprentissage machine. Par exemple, les métriques peuvent être la précision, le taux d'erreur, un score. L'expérience E donne une information sur le mode d'apprentissage qui peut être supervisé, non supervisé, semi-supervisé, multi-instances, par renforcement, par transfert. Nous allons détailler ci-après chacun de ces apprentissages.

3.2 Apprentissage supervisé

L'apprentissage supervisé présuppose que l'on dispose d'un ensemble d'exemples (étiquetés ou labélisés), caractérisés par des variables prédictives pour lesquels on connaît les valeurs de la variable cible. L'objectif durant la phase d'apprentissage est de généraliser l'association observée entre variable explicative et variable cible pour construire une fonction de prédiction f [8].

Une représentation mathématique de cette approche est :

Soit les n couples $(x_1, y_1), (x_2, y_2), (x_3, y_3), \dots, (x_i, y_i)$ qui sont les données d'exemple annotées, la prédiction de la sortie sur de nouvelles observations est : $x^ \rightarrow y^*$.*

Par exemple, nous avons un corpus de documents texte scientifiques annotés en fonction du sujet traité (cancérologie, pneumologie, cardiologie, etc). La classification d'un nouveau document dans une catégorie en fonction du domaine pathologique est une tâche supervisée. Nous pouvons utiliser un algorithme de traitement du langage naturel qui regroupe automatiquement ces documents par sujets prédéfinis [8]. Le traitement automatique des langues naturelles est le domaine de l'informatique qui s'intéresse à l'interprétation et à la production par des machines de phrases ou de textes dans des langues telles que le français ou l'anglais [9]. Il existe de nombreuses applications pratiques à l'utilisation des techniques du traitement du langage naturel comme la traduction automatique, le système de questions-réponses, la recherche d'information. Elles sont basées sur des modèles de langage qui définissent une distribution de probabilités sur des séquences de mots, de caractères, ou d'octets dans un langage naturel [10]. Ce traitement informatique de données en langage naturel utilise des sources qui peuvent provenir de forums de sites internet, de livres, de documents techniques, de publications scientifiques, méls. L'objectif de ce traitement est donc de représenter et de communiquer l'information contenue dans les données textuelles. Cette science du traitement de texte chevauche plusieurs disciplines : la linguistique, l'intelligence artificielle, les neurosciences, l'informatique. Les reconnaissances d'entités nommées sont appliquées à un corpus de textes. Cela consiste à repérer des éléments textuels dans un corpus de textes et à les catégoriser dans des classes définies au préalable. Par exemple, les noms d'un médicament, d'une maladie ou d'une classe thérapeutique. Cela se matérialise par un balisage encadrant les entités lors de l'annotation. La reconnaissance d'entités nommées est un sous-ensemble de l'extraction d'information. Historiquement, cette technique a été appliqué à des corpus tirés de la presse traitant des sujets géopolitiques [11]. Actuellement, la reconnaissance d'entités nommées est appliquée à des corpus spécialisés, et en ce qui nous concerne dans le domaine biomédical. La reconnaissance d'entités nommées comme composante interne du traitement du langage naturel est utile indirectement dans :

- L'analyse syntaxique.

- La coréférence ou le traitement des chaînes anaphoriques.
- La désambiguïsation lexicale : opération consistant à déterminer le sens d'un mot en contexte.
- La traduction automatique.

Dans une application directe, la reconnaissance d'entités nommées est mise en œuvre dans :

- L'extraction d'information et veille documentaire
- La tâche de questions-réponses
- L'anonymisation [12]

Le traitement du langage naturel peut être appliqué au volume considérable de publications scientifiques (PubMed) dans le cadre d'une recherche d'informations dans un projet de développement d'un candidat médicament.

Un autre exemple, nous avons un jeu de données de composés chimiques qui ont été préalablement annotés par un chimiste en molécules actives ou inactives, il s'agit de classer par inférence une nouvelle molécule dans ces deux catégories. La performance de ces algorithmes supervisés est meilleure que celle des autres algorithmes, mais il n'est pas toujours possible d'obtenir des données parfaitement labellisées par un expert du domaine. Cette limite est surmontée par l'utilisation d'algorithmes non supervisés.

3.3 Apprentissage non supervisé

L'apprentissage non supervisé ne présuppose aucun étiquetage préalable des données d'apprentissage. L'objectif est que le système parvienne, par lui-même, à regrouper en catégories les exemples fournis en entrée. Cela présuppose qu'il existe une notion de distance ou de similarité entre observations qui peut être utilisée à cet effet. L'interprétation des catégories identifiées reste à la charge d'un expert humain. Si nous modélisons cette définition par :

Soit $x_1, x_2, x_3, \dots, x_n$, les variables aléatoires représentant les observations brutes, on souhaite découvrir la relation avec des variables latentes structurelles suivante : $x_i \rightarrow y_i$.

En anglais, le terme *clustering* (ou partitionnement en français) est employé pour identifier les groupes dans les données [8]. En continuant sur l'exemple précédent, nous avons des documents textes non libellés cette fois-ci en entrée du système. Une méthode possible de regrouper les textes par catégories est la cooccurrence de mots qui se trouvent dans cet ensemble de textes. Ainsi, si nous avons les deux termes « virus » et « infection », tous les documents contenant ces deux occurrences peuvent être regroupés automatiquement dans la même catégorie. En épidémiologie, l'identification des groupes de patients présentant des

symptômes identiques aide à cibler les variants d'une maladie, est un autre exemple de partitionnement.

3.4 Apprentissage semi supervisé

L'apprentissage semi supervisé est un mélange des deux approches précédentes d'apprentissage. Il est appliqué lorsque nous avons beaucoup de données en entrée (X) mais très peu d'exemples étiquetés en sortie (Y). En reprenant l'exemple précédent d'annotation de documents scientifiques, une fois l'étape de regroupement effectuée, un expert humain affectera un sujet à chaque groupe en prélevant aléatoirement un ou deux des documents qu'il contient. L'un des avantages de cet apprentissage est qu'il décharge l'expert du domaine de la tâche fastidieuse d'annotation de l'ensemble des exemples d'apprentissage.

3.5 Apprentissage multi-instances

L'apprentissage multi-instances est un apprentissage supervisé qui permet la résolution de problèmes dans lesquels les données étiquetées appartiennent à un ensemble d'instances, au lieu d'une instance individuelle [13]. Dans le domaine de la classification de cancers en radiologie, l'ensemble des instances correspond à l'intégralité des coupes d'un organe et une instance représente une coupe [13].

3.6 Apprentissage par renforcement

Ces algorithmes d'apprentissage par renforcement intègrent une boucle de rétroaction entre le système d'apprentissage et ses expériences en optimisant une fonction de récompense. Cette boucle de rétroaction est un retour cyclique d'observations entre un agent et son environnement. Ce fonctionnement est inspiré du circuit de la récompense dopaminergique des neurones des êtres vivants [14]. AlphaGo Zero est une application de ces algorithmes dans le jeu de Go. La machine a joué des parties contre elle-même pour progresser.

3.7 Apprentissage par transfert

L'apprentissage par transfert permet de réutiliser des connaissances apprises à partir de tâches antérieures sur un nouveau problème ayant des similitudes avec les exemples précédents. Cette approche a l'avantage de réduire le besoin et l'effort de collecte des

données d'entraînement [15]. Ainsi, la réutilisation des résultats d'un corpus de texte généraliste annoté est une pratique usuelle dans le domaine de l'apprentissage automatique du langage. En effet, la sortie d'un algorithme qui classe les étiquettes identifiant un sujet, un verbe ou un complément dans une phrase extraite d'un corpus non scientifique peut être appliquée à un ensemble de texte scientifique.

4 Algorithmes de régression et de classification

La distinction entre régression et classification concerne les algorithmes supervisés. Un modèle de régression (voir figure 2 à gauche) est un modèle de *machine learning* dont la variable cible est quantitative, c'est-à-dire pouvant être classées en variables continues (par exemple la taille, le poids). La régression linéaire en est un exemple (voir tableau I). Un modèle de classification (voir figure 2 à droite) est un modèle d'apprentissage machine dont la variable cible est qualitative, pouvant être classées en variables catégorielles (les couleurs des yeux) ou ordinales (l'intensité d'une douleur classée en nulle, faible, moyenne, importante). La régression logistique est un exemple où la variable cible est binaire (voir tableau I).

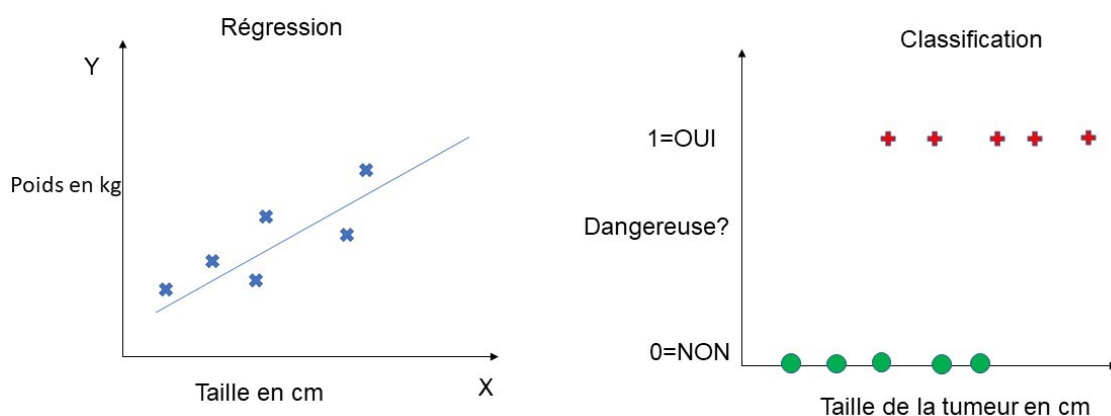


Figure 2 Schéma d'un algorithme de régression et de classification

Le graphique de gauche illustre la problématique suivante : quel est le poids d'une personne en fonction de sa taille ? Ce poids peut prendre une infinité de valeurs dans l'ensemble des réels $\mathbb{R}$, c'est un problème de régression. Dans le graphique de droite, nous voulons savoir si une tumeur est bénigne ou maligne en fonction de sa taille. C'est un problème de classification, ici binaire [6]. Si les étiquettes sont discrètes et correspondent donc à plusieurs classes, il s'agit de classification multi-classe [16].

5 Mesure de la performance (P) des algorithmes

Une grande variété d'algorithmes de *machine learning* est disponible pour développer un modèle (voir tableau I).

Tableau I Taxonomie des algorithmes d'apprentissage automatique [6]

Algorithme	Mode d'apprentissage	Type de problème à traiter
Régression linéaire univariée	Supervisé	Régression
Régression linéaire multivariée	Supervisé	Régression
Régression polynomiale	Supervisé	Régression
Régression régularisée	Supervisé	Régression
<i>Naive Bayes</i>	Supervisé	Classification
Régression logistique	Supervisé	Classification
<i>Clustering</i> hiérarchique	Non supervisé	-
<i>Clustering</i> non hiérarchique	Non supervisé	-
Arbres de décision	Supervisé	Régression ou classification
Les forêts aléatoires (<i>Random forest</i>)	Supervisé	Régression ou classification
<i>Gradient boosting</i>	Supervisé	Régression ou classification
Les machines à vecteurs de support (<i>Support Vector Machine</i>)	Supervisé	Régression ou classification
Analyse en composantes principales	Non supervisé	-

Pour évaluer la qualité d'un problème de régression, différentes mesures de performance du modèle peuvent être utilisées :

- L'erreur moyenne absolue (MAE pour *Mean Absolute Error*).
- La racine carrée de la moyenne du carré des erreurs (RMSE, *Root Mean Squared Error*).
- Le coefficient de détermination (R^2). Il mesure si un modèle est adapté aux données observées. Ce coefficient a une valeur comprise entre 0 et 1. 1 correspond à une adaptation parfaite [6].
- Le critère d'information d'Akaike (AIC) prend en compte la vraisemblance L du modèle. Un AIC petit signifie que le modèle est correct [6].

Pour les problèmes de classification, l'évaluation du modèle s'effectue à l'aide d'une matrice de confusion qui prend en compte les données prédites et les données observées.

Tableau II Matrice de confusion

		Observations		Total
		+	-	
Prédictions	+	Vrais positifs (VP)	Faux positifs (FP)	Total des positifs prédits (VP+FP)
	-	Faux négatifs (FN)	Vrais négatifs (VN)	Total des négatifs prédits (FN+VN)
	Total	Total des vrais positifs observés (VP+FN)	Total des vrais négatifs observés (FP+VN)	Taille totale de l'échantillon N

Les vrais positifs sont les exemples positifs correctement classifiés. Les faux positifs sont les exemples négatifs étiquetés positifs par le modèle. Les vrais négatifs sont les exemples négatifs correctement classifiés. Les faux négatifs sont les exemples positifs étiquetés négatifs par le modèle [8].

A partir de cette matrice, nous pouvons calculer :

- Le taux d'erreur défini par le taux de mauvaise classification : $(FN+FP) / N$.
- Le taux de vrais positifs, appelé également rappel (*recall* en anglais) ou sensibilité qui correspond à la proportion d'exemples positifs correctement identifiés comme tels : $VP / (VP+FN)$.
- La précision ou valeur prédictive positive (VPP) est la proportion de prédictions correctes parmi les prédictions positives : $VP / (VP+FP)$.
- La spécificité est le taux de vrais négatifs qui correspond à la proportion d'exemples négatifs correctement identifiés comme tels : $VN / (VN+FP)$.
- La F-mesure (F-score ou F_1 -score en anglais) est la moyenne harmonique de la précision et du rappel :

$$F_1\text{-score} = \frac{2 (\text{Précision} \times \text{Rappel})}{\text{Précision} + \text{Rappel}} = \frac{2VP}{2VP + FP + FN}$$

Le F_1 -score permet de comparer plusieurs modèles.

Une autre méthode d'évaluation des problèmes de classification est la courbe ROC (*Receiver Operating Characteristic*). Deux indicateurs de performance sont nécessaires pour tracer une courbe ROC : la sensibilité et la spécificité (voir la figure 3). Pour comparer des modèles, nous pouvons utiliser l'aire sous la courbe de ROC (AUC = *Area Under the Curve*). Plus grande est l'AUC, meilleur est le modèle [6].

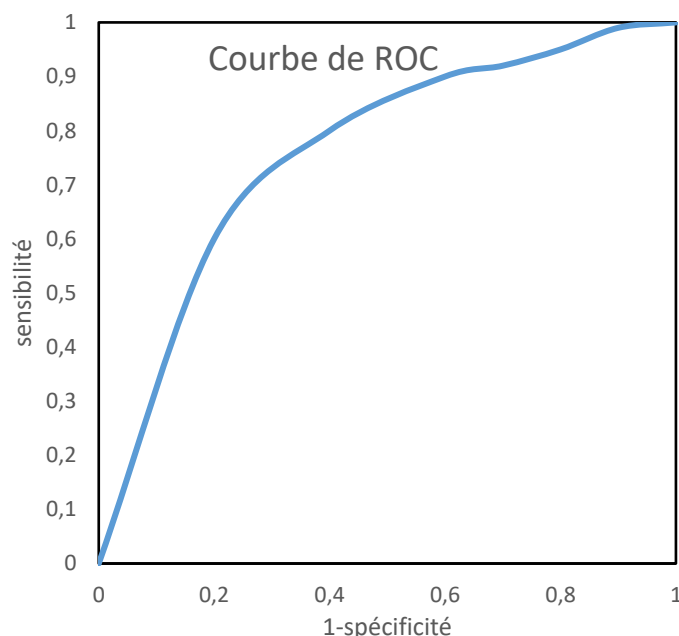


Figure 3 Représentation d'une courbe de ROC

Prenons un exemple illustrant une matrice de confusion :

Nous avons à notre disposition une base de données de mammographies annotées par des radiologues. Ils ont entouré les tumeurs sur les images. Un modèle de classification est entraîné sur ce jeu de données annotées. L'objectif est de prédire sur un nouveau cliché si une tumeur est présente ou pas. Les résultats pour 100 nouvelles mammographies sont les suivants :

Tableau III Matrice de confusion d'un cas d'usage fictif

		Observations		Total
		Cancer	Pas de cancer	
Prédictions	+	19	21	40
	-	1	59	60
	Total	20	80	100

Le rappel est : $19 / (19+1) = 0.950$.

La spécificité est : $59 / (21+59) = 0.740$

La précision est : $19 / (19+21) = 0.475$.

Le F₁ score est égal à : $2 \cdot (0.475 \cdot 0.950) / (0.475 + 0.950) = 0.633$

Lorsque l'algorithme du modèle détecte un cancer, c'est vrai à 47.5%. De plus, il détecte 95% des cancers. Une alerte fausse est moins dangereuse dans ce cas d'usage, que de passer à côté de la tumeur et de ne pas traiter la personne.

6 Généralisation

La généralisation est la capacité à donner la bonne réponse pour des exemples que la machine apprenante n'a pas vus durant l'apprentissage [4].

Le surapprentissage (ou *overfitting* en anglais) caractérise un modèle qui s'ajuste parfaitement aux données présentes, mais qui ne se généralise pas à d'autres données non utilisées. Une des causes du surapprentissage est la complexité du modèle. En application du principe du rasoir d'Ockham (ou Occam), plus un modèle d'apprentissage machine est simple, moins il surapprend les particularités des données traitées. Ce principe n'est pas spécifique à l'apprentissage automatique.

Le sous-apprentissage (ou *underfitting* en anglais) est le risque opposé au surapprentissage. C'est un modèle dont les performances sur les données d'apprentissage sont faibles.

Le compromis biais-variance permet de savoir si le modèle est dans un cas de surapprentissage. L'erreur de généralisation peut être définie comme la somme de trois erreurs : le biais, la variance et l'erreur irréductible. Le biais est dû à de mauvaises hypothèses sur les données. Un modèle avec un biais élevé sera plus enclin au sous-apprentissage. La variance est la conséquence d'une sensibilité excessive du modèle à des variations minimales dans le jeu d'entraînement. Une variance élevée est un signe de surapprentissage. L'erreur irréductible provient du bruit présent dans les données. Les hyperparamètres sont des paramètres de l'algorithme d'apprentissage qui permettent de déterminer le modèle, par exemple le nombre de couches de neurones d'un réseau de neurones profond.

7 Jeu d'entraînement, de test et de validation

Pour faire fonctionner ces algorithmes d'apprentissage automatique, un jeu de données est nécessaire. Il sera divisé en deux ou trois parties : un jeu d'entraînement, un jeu de validation, et un jeu de test. Ce dernier n'est pas utilisé pour le choix du modèle et son entraînement. Le calcul de la performance de l'algorithme se fait sur ce jeu de test. Cependant cette méthode a un inconvénient qui est de ne pas utiliser la totalité du jeu de données pour entraîner ou tester le modèle. De plus, si le jeu de test est complexe ou à l'inverse trop simple à prédire, nous estimerions la performance de l'algorithme avec un biais. Pour surmonter ces écueils, la validation croisée (*cross validation* en anglais) est utilisée (Figure 4). Cette méthode permet l'utilisation de l'intégralité du jeu de données pour la phase d'entraînement et la phase de validation. Le principe est de diviser le jeu de données en k parties (*folds* en anglais) plus ou moins égales entre elles. Chacune des k

parties servira de jeu de test alternativement. Le restant des données (l'union des k-1 autres parties) sera utilisé pour le jeu d'entraînement.

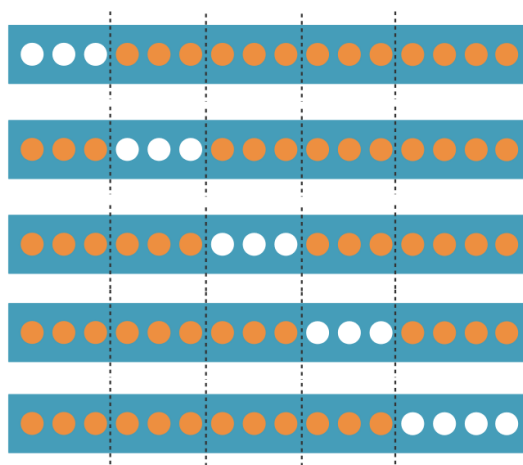


Figure 4 Validation croisée à 5 parties, chaque point appartient à un des cinq jeux de test (en blanc) et aux quatre autres jeux d'entraînements (en orange), d'après la source [16].

D'autres méthodes existent dans l'utilisation des données pour l'entraînement et le test.

7.1 Les données

L'IA a réalisé des progrès impressionnants depuis cette dernière décennie grâce à la qualité et à la quantité massive de données (*big data*) utilisée pour alimenter les algorithmes d'apprentissage. Un ordre de grandeur de la taille des bases de données est d'une centaine de milliers à plusieurs millions de points. L'IA ne sert à rien s'il n'y a pas des données de qualité en entrée. Une expression anglaise résume cette idée : « *garbage in garbage out* ». Si les données sont de mauvaises qualités, peu nombreuses ou incomplètes, la sortie ou la prédiction sera fausse. Dans le contexte de la chimie, les données peuvent être une représentation d'une molécule avec son énergie libre de solvation ou toutes autres propriétés [17]. Pour être utilisé par un algorithme d'apprentissage statistique, ces descripteurs moléculaires sont modélisés mathématiquement. C'est ce que l'on appelle les caractérisations moléculaires ou en anglais *molecular featurizations*. Ces représentations comprennent des modélisations en graphe 2D des molécules, les interactions électrostatiques en 3D, les fonctions des orbitales atomiques. Les descripteurs chimiques ou moléculaires peuvent-être :

- Le nombre d'atomes
- La masse moléculaire
- Le volume de Van Der Waals ou volume atomique
- Le nombre de charges partielles
- L'énergie totale

- L'énergie de l'orbitale la plus haute occupée (HOMO)
- L'énergie de l'orbitale la plus basse vacante (LUMO)
- Etc...

Il existe des bibliothèques chimiques qui peuvent être :

- Une base de données interne de la société pharmaceutique
- Une base de données commerciale de grande taille (1 à 2 millions de molécules)
- Des bases de données en libre accès (*open source*)
- De l'espace chimique virtuel (ensemble des composés chimiques qu'il serait possible de synthétiser, de l'ordre de 10^{60} molécules).

Ces données sont produites par criblage virtuel à haut débit, ou par la chimie combinatoire.

Le criblage à haut débit permet de tester des milliers à plusieurs centaines de milliers de molécules en une journée. Ensuite, il faut synthétiser ces molécules qui présentent une activité potentielle pour les purifier et les caractériser. La chimie combinatoire a la capacité de concevoir des chimiothèques en utilisant le résultat des réactions des ensembles de réactifs. Le principe est qu'un ensemble de composés d'une réactivité A (A_1 à A_n) est opposé à un ensemble de composés de réactivité B (B_1 à B_n). Ce principe est illustré dans le tableau III suivant.

Tableau IV Matrice de combinaison de molécules

	B_1	B_j	B_n
A_1	A_1-B_1	A_1-B_j	A_1-B_n
A_i	A_i-B_1	A_i-B_j	A_i-B_n
A_n	A_n-B_1	A_n-B_j	A_n-B_n

Regardons un exemple de ces données chimiques dans une bibliothèque appelé MoleculeNet.

7.2 MoleculeNet

Parmi les bibliothèques de composés chimiques accessibles librement, nous pouvons citer *MoleculeNet* qui est un ensemble de jeux de données de référence pour évaluer la qualité des algorithmes d'apprentissage automatique [18], à l'instar d'autres domaines comme la reconnaissance d'objets (*ImageNet*) ou le traitement du langage naturel (*WordNet*). *MoleculeNet* dispose de données sur les propriétés de 700 000 composés. Les propriétés sont divisées en quatre catégories : la mécanique quantique, la physique-chimie, la biophysique et la physiologie. Les données proviennent de plusieurs bases de données publiques (par exemple GDB voir annexe 2). Les différentes catégories de jeux de données dans MoleculeNet [18] sont :

- **Catégorie mécanique quantique :**

- *QM7/QM7b* : ces deux jeux de données sont extraits de la base de donnée GDB-13 [19] (voir annexe 2). Cette dernière base contient des données de un million de petites molécules organiques jusqu'à treize atomes de C, N, O, S et Cl en respectant les règles simples de stabilité chimique et de faisabilité synthétique. QM7 et QM7b encodent les propriétés électroniques des petites molécules déterminées par la théorie de la densité fonctionnelle (DFT pour *density functional theory*).
- *QM8* : ce jeu de donnée intègre les spectres électroniques et les états d'énergie excités des molécules calculés à partir de plusieurs méthodes de mécanique quantique.
- *QM9* : les informations contenues dans ce jeu de donnée sont les propriétés géométriques, énergétiques, électroniques, thermodynamiques produites par DFT.

- **Catégorie physique-chimie :**

- *ESOL* contient les données sur la solubilité des molécules organiques dans l'eau.
- *Freesolv* présente les valeurs représentant l'énergie libre d'hydratation des molécules dans l'eau issues d'expériences et de calcul.
- *Lipophilicity* correspond aux résultats expérimentaux du coefficient de partage octanol/eau (log P au pH 7.4)

- **Catégorie biophysique :**

- *PCBA* est un jeu de données extraites de la base PubChem BioAssay qui renseignent sur l'activité biologique des molécules identifiées par le criblage à haut débit (HTS).
- *MUV* (*maximum unbiased validation*) est un sous-ensemble de PubChem BioAssay conçu pour valider les techniques de criblage virtuel.
- *HIV* est une base de mesures expérimentales de la capacité des molécules à inhiber la réplication du virus de l'immunodéficience acquise.
- *PDBbind* contient les données concernant les liaisons affines des molécules biologiques complexes.
- *BACE* illustre les données quantitatives (concentration inhibitrice IC50) et qualitatives d'un ensemble d'inhibiteurs de la bêta-sécrétase1 humaine (BACE-1).

- **Catégorie physiologique :**

- *BBBP* (*blood-brain barrier penetration*) informe sur la perméabilité de la barrière hémato-encéphalique à l'égard des substances chimiques.
- *Tox21* décrit les mesures de la toxicité sur des cibles biologiques, tels que des récepteurs nucléaires et des voies de signalisation du stress oxydant.
- *ToxCast* renseigne sur les données toxicologiques d'une chimiothèque de composés mis en évidence par un criblage à haut débit *in vitro*.

- *SIDER* est une base de données sur les effets indésirables des médicaments commercialisés classés par système d'organes.
- *ClinTox* possède des données qualitatives sur les médicaments approuvés par la FDA et sur ceux dont les essais cliniques ont échoué pour des raisons de toxicité [18].

Ces différents jeux de données sont utilisés par des méthodes de régression ou de classification. MoleculeNet est inclus dans la bibliothèque DeepChem qui est construite à partir de la plateforme TensorFlow de Google [20]. Elle est utilisée pour implémenter des algorithmes d'apprentissage profond. DeepChem propose des modèles, des algorithmes, des jeux de données adaptés à la résolution de problèmes en chimie. La société Schrödinger (voir annexe 1) développe des applications à partir de DeepChem.

7.3 Logiciels utilisés en pratique

L'implémentation de ces différents algorithmes est généralement basée sur des *framework* ou logiciels en *open source* (voir annexe 3). Le plus connu et le plus utilisé actuellement est Tensorflow développé par une équipe de la société Google [20]. Une autre bibliothèque de logiciels libre de droit est Pytorch créée par l'entreprise Facebook [21]. Ces deux exemples montrent que l'écosystème des logiciels d'apprentissage automatique est dominé par les grandes entreprises du numérique. Même si le code est libre de droit, une certaine dépendance s'installe pour les utilisateurs. En effet, ces entreprises fournissent des services pour faciliter l'installation de ces logiciels sur leur infrastructure. Mais les avantages à réutiliser ces logiciels sont la maintenabilité, la robustesse de leur code et la présence d'une communauté d'utilisateurs très active.

Une fois que nous savons où nos données se trouvent pour entraîner un modèle algorithmique, nous devons choisir une méthode à mettre en place.

8 Méthodologie de mise en place d'un algorithme

Dans le cadre d'un projet intégrant une solution d'IA (voir figure 5), une méthodologie de mise en place d'un algorithme d'apprentissage automatique peut être la suivante :

- Définir les mesures de la performance : le taux d'erreur, la précision, le rappel (*recall*), le F-score, la courbe ROC.
- Le choix du modèle s'effectue en fonction de la tâche à accomplir et du type de données disponible en entrée de l'algorithme.
- Estimer la quantité et la qualité de données à recueillir.

Les données sont généralement stockées dans des bases de données. Si le format de donnée ne correspond pas à celui utilisé par l'algorithme, une phase de pré-traitement sera nécessaire. Il faudra s'assurer de l'intégrité des données.

- Choisir les hyperparamètres par un réglage manuel, automatique. Les hyperparamètres diffèrent en fonction de l'algorithme d'apprentissage mis en œuvre. Dans un réseau de neurones, les hyperparamètres correspondent au choix du nombre de neurones et de couches, fixer le taux d'apprentissage.

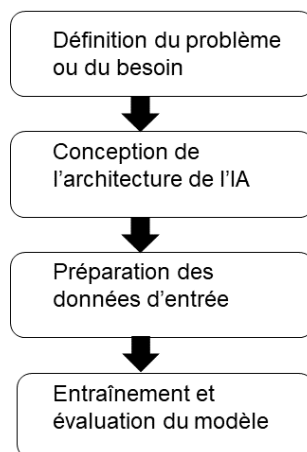


Figure 5 Schéma simplifié du processus de développement d'un projet d'IA

L'application de cette méthode de travail est du ressort d'un *data scientist* (scientifique des données) en collaboration avec un expert métier. Ce dernier définit le besoin et pose la problématique à résoudre. Il intervient également dans la phase d'entraînement du modèle pour qualifier les données d'apprentissage. Cette méthode est effectuée selon un processus itératif par essais et erreurs.

Dans cette première partie, nous avons mis en place les différents concepts et le vocabulaire utilisé dans le domaine de l'intelligence artificielle. Maintenant, nous allons aborder les cas d'usage dans la découverte de nouvelles molécules et l'aide au diagnostic en radiologie.

B Applications de l'IA

Les applications de l'IA dans la découverte de nouvelles entités chimiques font l'objet de publications nombreuses dans la littérature scientifique. Elles concernent les différentes approches utilisées par les chimio-informaticiens ou les bio-informaticiens dans l'identification de nouveaux candidats-médicaments. Dans cette partie, nous allons resituer le contexte économique de la découverte de nouveaux traitements. Dans un premier point, nous décrirons succinctement le processus de R&D. Dans un second point, nous aborderons les différentes approches employées pour trouver des médicaments. Après, nous parlerons de la représentation des molécules nécessaire à un algorithme d'apprentissage automatique pour être entraîné. Nous présenterons l'apprentissage profond qui est l'une des techniques d'IA les plus performantes à l'heure actuelle de l'état de l'art. Enfin dans un dernier point, nous illustrerons cette application de l'apprentissage profond par des cas d'utilisation dans les domaines pharmaceutique et médical, en particulier l'imagerie.

1 Drug Discovery

1.1 Contexte économique

L'un des objectifs de la chimie médicinale ou thérapeutique est la découverte de nouveau médicament. Partant du constat que le nombre de nouveaux médicaments mis sur le marché a considérablement diminué, alors que le coût du développement n'a cessé de croître, l'industrie pharmaceutique cherche à optimiser la découverte de nouvelles molécules. De plus, les besoins thérapeutiques pour certaines pathologies sont insatisfaits. C'est le cas notamment des maladies liées au vieillissement de la population, par exemple la maladie d'Alzheimer. D'autre part, la concurrence des médicaments génériques et la perte d'exploitation des brevets pour les molécules princeps sont responsables de l'érosion des parts de marché de certains laboratoires pharmaceutiques. Pour répondre à ces enjeux de santé publique et économique, l'industrie pharmaceutique est à la recherche de solutions innovantes pour développer de nouveaux médicaments avec un service médical rendu important. L'intelligence artificielle est une piste prometteuse pour relever ce défi. Les promesses d'entreprises utilisant des systèmes d'IA sont de proposer des candidats-médicaments à un coût réduit, dans un délai court et avec un taux d'attrition très inférieur à celui que les entreprises pharmaceutiques ont connu jusqu'à présent. Le nombre de molécules actives commercialisées est de 1200. Ces molécules sont présentes dans des milliers de médicaments vendus (12 718 en France d'après la base de données publique des

médicaments en 2019). Elles sont d'origine végétale (une quarantaine) ou fabriquées par synthèse chimique et biotechnologie [22]. De plus, il existe 330 cibles moléculaires dont 270 codées par l'ADN humain et le reste est exprimé par le génome des organismes pathogènes (virus, bactéries, parasites). Ces cibles ne sont qu'une infime partie des cibles potentielles d'origine humaine ou d'organismes pathogènes. Le réservoir de cibles est donc énorme pour la recherche pharmaceutique [22].

1.2 Description du processus de R&D

Trois étapes générales constituent ce processus : la recherche en amont ou *drug discovery*, le développement pré-clinique, le développement clinique [23].

La première étape de la R&D dure habituellement de deux à cinq ans et se déroule selon les points suivants :

- Détermination de la cible biologique.
- *Hit generation* ou génération des touches.
- *Lead identification* ou identification d'un chef de file ou tête de série.
- *Lead optimization*, optimisation d'un chef de file.

Les cibles peuvent être des récepteurs, des protéines, des gènes, des enzymes. Les méthodologies sont nombreuses, nous pouvons citer la recherche *in vitro*, la fouille de données dans des bases existantes, le criblage phénotypique, etc. L'IA peut accélérer et augmenter l'efficacité de ces méthodes dans la découverte de la cible. Ainsi, l'IA permet de trouver des corrélations là où l'Homme peut avoir des difficultés, par exemple dans la modélisation des relations structures-activités des molécules (QSAR).

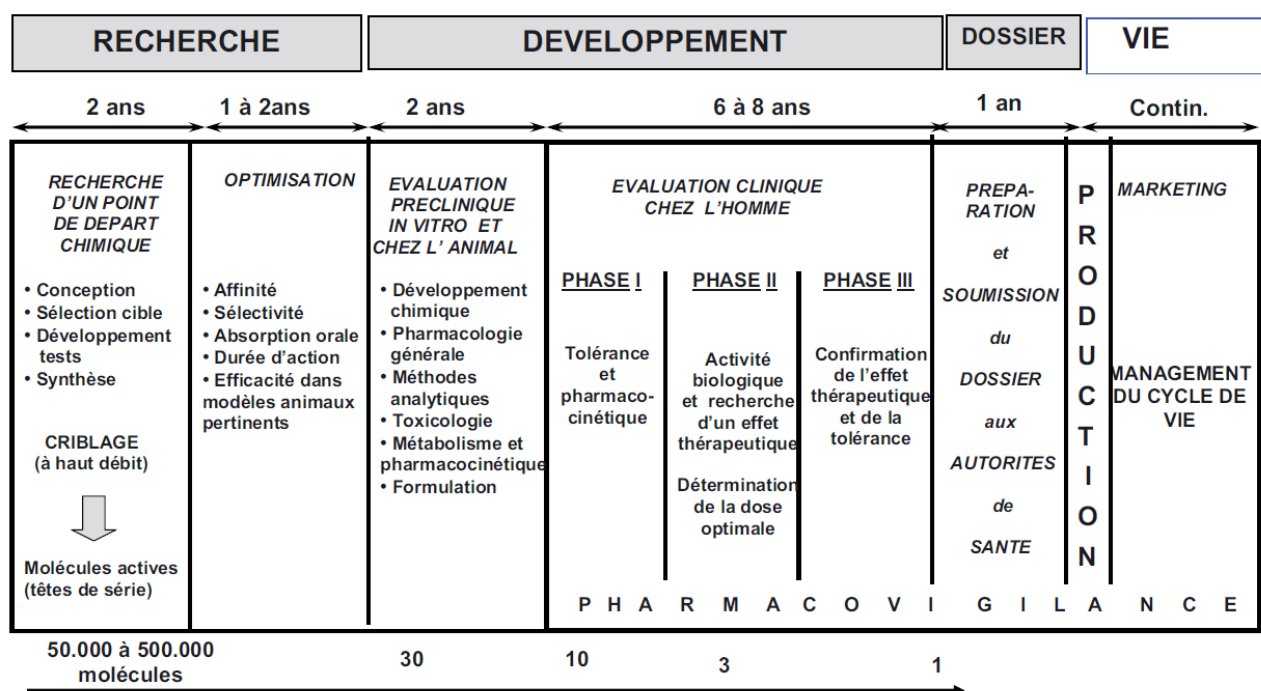


Figure 6 Contenu et durée des différentes phases de recherche et développement d'après la référence [23]

Dans la phase d'optimisation (voir à gauche de la figure 6), les molécules synthétisées sont caractérisées en fonction de leur affinité, leur sélectivité, leur absorption orale, leur durée d'action et leur efficacité. Les tests *in vitro* sont effectués sur des cellules recombinantes ou des cellules natives exprimant la cible moléculaire ou cellulaire. Les tests *in vivo* sont réalisés dans des modèles d'animaux présentant la pathologie. A ce stade de la sélection du candidat médicament, des études pharmacocinétiques et toxicologiques sont menées avant de passer à l'étape suivante le développement préclinique [23]. Il ne reste que 30 molécules sur les 50 000 à 500 000 molécules de départ de la recherche. Des études analytiques du principe actif vont être réalisées. Il s'agit de préparer, identifier, la substance active, d'en vérifier sa pureté, sa stabilité. La synthèse à grande échelle du produit est nécessaire pour obtenir quelques kilos. Un travail de recherche de la forme galénique est également effectué. Ensuite, l'étude clinique chez l'Homme est engagée selon trois temps :

- Phase I : la dose maximale tolérée de substance active chez des volontaires sains est déterminée. Les résultats de l'étude préclinique chez l'animal renseignent le choix de la dose initiale.
- Phase II : les objectifs sont d'étudier les propriétés pharmacodynamiques étudiées chez l'animal et d'améliorer la connaissance pharmacocinétique de l'entité chimique. La détermination de la posologie optimale est recherchée.
- Phase III : l'efficacité thérapeutique du composé doit être prouvée en tenant compte du rapport bénéfice/risque potentiel. Un essai comparatif permet de comparer le nouveau traitement à un traitement de référence ou à un placebo. A l'issue de cette phase, une autorisation de mise sur le marché (AMM) est

demandée. Les critères d'obtention de l'AMM sont la qualité, la sécurité et l'efficacité du médicament.

La phase IV débute après la mise sur le marché du médicament, avec une surveillance en vie réelle des effets indésirables du produit sur les patients [23].

D'autre part, les règles de Lipinski [24] limitent la sélection des composés pour le futur médicament. On la dénomme « la règle des 5 », et les critères énoncés sont :

- La masse moléculaire ≤ 500 Da (Dalton).
- Le $\log P \leq 5$, où $\log P$ représente le coefficient de partage octanol/eau.
- Les sites accepteurs de liaisons H ≤ 10 , H pour hydrogène.
- Les sites donneurs de liaisons H ≤ 5 .

Cette règle permet de prédire l'absorption et la perméabilité de la molécule.

Il existe d'autres règles dénommées les règles de Veber qui reprennent les précédentes en les complétant :

- La molécule ne doit pas avoir plus de 5 sites donneurs de liaisons hydrogène.
- Elle ne doit pas avoir plus de 10 sites accepteurs de liaisons H.
- La masse moléculaire doit être inférieure à 500 Dalton.
- Le nombre d'atomes doit être compris entre 20 et 70.
- L'aire de la surface polaire de la molécule doit être plus petite que $140 \text{ \AA}^2$.

Un autre critère pris en compte dans la poursuite du développement du candidat-médicament est la capacité à le synthétiser avec peu d'étapes dans le procédé industriel. En effet, si la réaction de synthèse est trop complexe, le coût de production sera trop élevé.

1.3 Les différentes approches dans la découverte d'un médicament

En remontant le fil de l'Histoire jusqu'à l'Antiquité, nous observons que l'Humanité a utilisé les plantes, certains animaux ou minéraux pour se soigner sans en comprendre les mécanismes d'action. Cette connaissance empirique prit fin avec l'amélioration du savoir scientifique, en particulier en médecine, en chimie. L'isolement des premiers principes actifs d'origine végétale (morphine par exemple) et la synthèse de molécules médicamenteuses comme l'aspirine commencèrent vers la fin du XIX^e siècle.

Le développement de la chimie combinatoire (synthèses automatisées à haut débit) à la fin des années 1970 rendit possible la préparation de dizaines de milliers de produits par les chimistes. L'apparition des techniques de *screening* à haut débit avec des robots augmenta le nombre de composés testés (50 000 par jour).

1.3.1 L'approche empirique

Le rôle du hasard dans la découverte de médicaments est nommé la sérendipité. Dans leur livre les auteurs Bohuon et Monneret recensent plusieurs médicaments qui ont

été mis au point de manière fortuite [25]. Ainsi les découvertes de la pénicilline, du valium, du modafinil, du cisplatine, et d'autres molécules, sont dues à la sérendipité. Ce mot issu de l'anglais « *serendipity* » peut-être défini par l'exploitation heureuse et opportune du hasard. Cependant comme le disait si bien Pasteur, « le hasard ne favorise que les esprits préparés ». C'est la faculté d'observation et l'esprit de déduction du chercheur qui ont permis à ces découvertes de devenir des médicaments.

1.3.2 L'approche fonctionnelle

Dans le passé, la recherche de nouvelles molécules au laboratoire était basée sur un criblage fonctionnel (*screening* phénotypique). Les essais s'effectuaient sur des cellules isolées ou *in vivo* chez un modèle animal de la pathologie humaine. Le mécanisme d'action de la substance sur les cibles moléculaires était souvent inconnu. Cette approche est chronophage et ne permet de tester qu'un nombre limité de molécules. Cependant, elle a été à l'origine de la découverte de grands médicaments qui sont encore utilisés de nos jours [26]. Nous pouvons citer la cyclosporine, le taxol, la vinblastine, le clopidogrel. Les progrès dans le domaine de la biologie moléculaire et en biochimie à la fin des années 80 ont permis le développement d'un criblage biochimique *in vitro* plus rapide et à haut débit sur des cibles moléculaires précises [23].

1.3.3 L'approche récente ou rationnelle

a) Identification et validation de la cible

La connaissance du mécanisme de la maladie permet d'avoir un point de départ pour identifier une cible thérapeutique potentielle. Cela peut être une protéine impliquée dans une voie de signalisation nécessaire au développement de cette maladie.

Le séquençage du génome humain couplé aux techniques de bio-informatique, de génomique et de protéomique permet également d'identifier des milliers de cibles potentielles [27]. Par exemple, dans le cas de la maladie d'Alzheimer, les études génétiques des formes familiales précoces ont identifié des mutations sur les gènes codant pour le précurseur de la protéine bêta-amyloïde [28]. Trois gènes responsables ont été identifiés. Cependant, aujourd'hui aucun traitement efficace pour guérir cette maladie n'est disponible. Un anticorps monoclonal l'aducanumab produit par le laboratoire américain Biogen en partenariat avec le japonais Eisai donne de l'espoir et en cours d'étude clinique [29]. Depuis octobre 2019, la société a demandé une AMM à la FDA. Un article du journal Lancet critique les résultats statistiques des études et émet des réserves quant à leur validité [30]. Cet exemple illustre le long parcours du cycle de développement d'un médicament.

b) La connaissance acquise sur les récepteurs

Les médicaments se fixent sur leurs cibles cellulaires qui peuvent être des récepteurs, des enzymes, des canaux ioniques, ADN, ... La synthèse des ligands de ces

cibles permet de trouver des substances possédant des propriétés biologiques et thérapeutiques à fort potentiel. Certains récepteurs peuvent être membranaires.

L'optimisation des activités secondaires de principes actifs qui sont déjà connus est une autre méthode de découverte de médicaments. Cette approche SOSA (*Selective Optimization of Side Activities*) part du postulat qu'un médicament agit sur différents récepteurs ou cibles biologiques. Cette optimisation améliore l'activité secondaire du principe actif pour une autre cible que celle prévue initialement [31].

1.4 La représentation des molécules

Les molécules organiques de la matière vivante sont constituées d'atomes de carbone qui peuvent s'associer les uns avec les autres et également fixer d'autres atomes tels que l'hydrogène, l'oxygène, l'azote, le phosphore et le soufre. Ces atomes forment des liaisons chimiques entre eux. Ces liaisons peuvent être fortes ou faibles. Plus la liaison est forte, plus l'énergie nécessaire pour la rompre est importante. Dans les liaisons faibles (10 à 100 kJ.mol^{-1}), on trouve la liaison hydrogène et l'interaction de van der Waals. Dans les liaisons fortes (100 à 1000 kJ.mol^{-1}), on a la liaison ionique, la liaison covalente et la liaison métallique.

Les liaisons faibles jouent un rôle crucial dans les processus biologiques et dans le *drug design*. Ainsi, la reconnaissance entre les acides aminés des sites catalytiques des protéines et les groupements fonctionnels de leurs substrats fait intervenir ce type de liaisons. De même, la plupart des médicaments interagissent avec les biomolécules du corps humain par l'intermédiaire de liaisons faibles.

Pour modéliser ce concept de liaison et d'atome, les chimio-informaticiens utilisent la théorie des graphes. Un graphe est un modèle mathématique représentant des sommets connectés ensemble par des arêtes (Figure 7). Dans cette description, si on considère une molécule comme un graphe, alors les atomes qui la composent sont les sommets et les liaisons sont les arêtes (Figure 8).

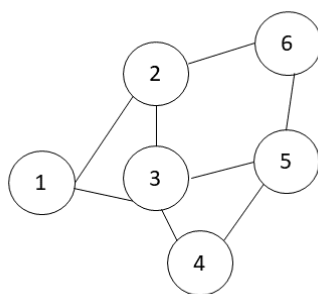


Figure 7 Exemple de graphe avec 6 sommets reliés par des arêtes

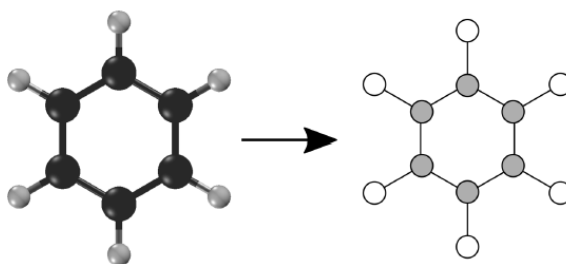


Figure 8 Molécule de benzène transformée en graphe moléculaire (à droite) d'après [32]

Les graphes moléculaires peuvent être représentés par des chaînes de caractères en utilisant la notation SMILES. *Simplified Molecular Input Line Entry System* (SMILES) est un langage symbolique de description dans lequel les molécules et les réactions chimiques sont formalisées avec la norme de codage de caractères informatiques ASCII (*American Standard Code for Information Interchange*) [33] [34]. Ainsi, le benzène est représenté par c1ccccc1.

Les molécules peuvent aussi être représentées par leur formule brute (1D), par leur projection en deux dimensions (Fischer par exemple), par leur structure tridimensionnelle (représentation de Cram). Les molécules sont décrites par une empreinte moléculaire appelé ECFPs, *Extended Connectivity Fingerprints* [35]. Ce descripteur traduit la représentation de la molécule en 2D sous la forme d'un vecteur numérique. Il est utilisé dans la recherche de groupes fonctionnels ou sous-structures moléculaires dans des bases de données chimiques. Les ECFPs servent également aux tâches liées à la prédiction et à la compréhension de l'activité des substances actives médicamenteuses. Les applications utilisant les ECFPs sont le criblage virtuel à haut débit, la modélisation de la relation structure-activité, et l'analyse des bibliothèques chimiques [35]. Les méthodes de similarité, de classification, de regroupement (*clustering*) de molécules peuvent intégrer les empreintes moléculaires.

Preuer K, Klambauer G, Rippmann F, et al ont appliqué un algorithme d'apprentissage profond en utilisant les ECFPs pour déterminer la présence ou l'absence de pharmacophores [36]. Une définition de la notion de pharmacophore est donnée par l'IUPAC (*International Union of Pure and Applied Chemistry*) : « Un pharmacophore est l'ensemble des caractéristiques stériques et électroniques nécessaires à l'établissement d'interactions supramoléculaires avec une cible spécifique, afin de déclencher (ou d'inhiber) une réponse biologique » [37].

Pour résumer, la représentation des molécules se fait de quatre manières différentes : les descripteurs moléculaires, les empreintes moléculaires, le format SMILES, et les grilles (matrices). Ces représentations doivent préserver les propriétés d'invariance des molécules. Ces propriétés sont en nombre de trois :

- L'invariance de la permutation : la représentation ne doit pas changer l'ordre des atomes spécifiés de la molécule.

- L'invariance de la translation : la représentation ne doit pas être modifiée par une translation dans l'espace.
- L'invariance de la rotation : la représentation ne doit pas être altérée par une opération de rotation [17].

L'implémentation informatique de ces caractérisations moléculaires est fournie par des bibliothèques comme le RDKit [38]. Ce dernier est un ensemble d'outils libre de droit à destination des chimio-informaticiens. Il propose des fonctions de lecture, écriture, de schématisation de molécules. Il offre la possibilité de manipuler les molécules en 2D ou 3D. La fonctionnalité de recherche de sous-structures dans une molécule est présente également.

1.5 *Deep learning* ou apprentissage profond

Les neurologues Warren McCulloch et Walter Pitts ont développé le premier modèle mathématique du neurone biologique, le neurone formel (voir figure 9). Par exemple, la fonction d'activation de ce dernier représente l'atteinte d'un seuil par la somme pondérée des influx arrivant au neurone. Les réseaux de neurones profonds sont des modèles paramétriques flexibles. Un des premiers réseaux de neurones artificiels est le perceptron qui date de 1957, inventé par le psychologue Rosenblatt. Ce perceptron est composé d'une couche d'entrée de p neurones, ou unités, qui correspondent chacune à une variable d'entrée. Ces unités sont connectées au seul neurone de l'unique couche du perceptron (après la couche d'entrée).

Cependant, en 1969, M. Minsky et S. Papert ont démontré les limites des perceptrons, notamment leur incapacité à résoudre le problème de classification du OU exclusif (XOR) [1]. Pour dépasser ces contraintes, l'idée d'empiler plusieurs perceptrons a germé. Ce réseau de neurone artificiel a été dénommé perceptron multicouche (PMC). Un PMC a une couche d'entrée, une ou plusieurs couches cachées et une couche de sortie. On parle de réseaux de neurones profonds (RNP) lorsque le réseau de neurones artificiels (RNA) est constitué de deux couches cachées ou plus (voir figure 10). En anglais, le terme de *deep neural network* (DNN) est utilisé dans la littérature scientifique. Le calcul des réseaux de neurones (ou connexionnistes) est basé sur la propagation des informations entre des unités élémentaires de calcul sur lesquelles des poids (w) de connexion sont appliqués (voir figure 9) [39]. L'équation mathématique d'un neurone formel est :

$$\hat{y} = f(\langle w, x \rangle + b), \text{ où } b \text{ représente le biais.}$$

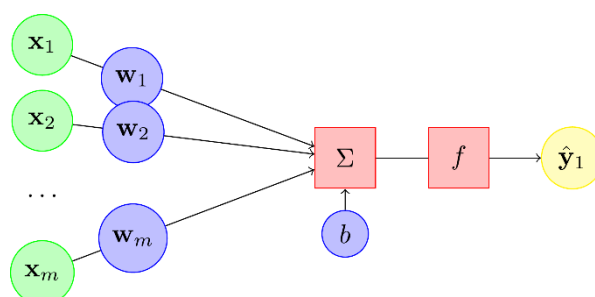


Figure 9 Représentation d'un neurone formel d'après la référence [39]

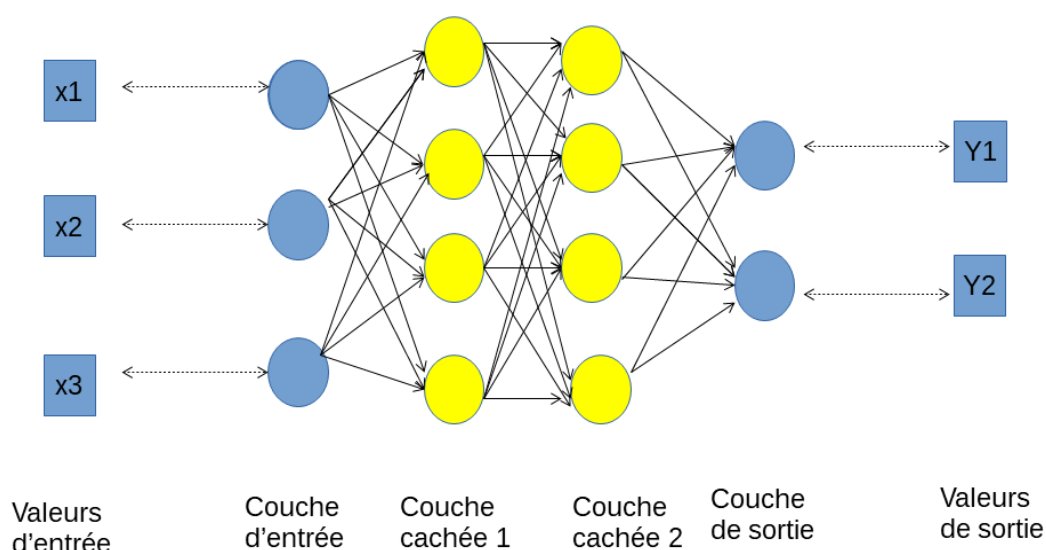


Figure 10 Topologie basique d'un réseau de neurones artificiels

La construction d'un réseau de neurones s'effectue en empilant les couches de neurones où la sortie d'une couche correspond à l'entrée de la suivante. La dernière couche du réseau de neurones représente la ou les valeurs prédites. Une fonction de perte est associée à cette dernière pour mesurer l'erreur de classification [1]. L'objectif lors de la phase d'apprentissage est de calculer les poids en minimisant la fonction de perte. L'apprentissage est réalisé par la méthode de rétropropagation qui consiste à partir de la dernière couche à calculer les poids en appliquant un algorithme de descente de gradient stochastique.

Nous avons posé les définitions et les concepts sous-jacents de l'apprentissage profond avec une volonté délibérée d'occulter l'aspect mathématique de la modélisation. Dans le paragraphe suivant, nous allons passer de la théorie à la pratique en illustrant nos propos par des cas d'utilisations concrets. Le premier exemple montre la mise en œuvre d'un algorithme d'apprentissage profond pour prédire la toxicité des molécules chimiques.

1.5.1 Application de l'apprentissage profond dans la détection de toxicophores

Le défi nommé Tox21 Data Challenge proposé par l'agence américaine de la santé NIH (*National Institutes of Health*) en 2014 a été remporté par une équipe qui a mis en place un algorithme d'apprentissage profond avec une méthode multi-tâches (plusieurs effets toxiques différents pour un composé chimique). Ce défi avait pour objectif de prédire la toxicité des composés chimiques sur la santé humaine et d'évaluer la capacité de ces outils dans la réduction des tests toxicologiques *in vivo* et *in vitro* [40]. Les gagnants de ce concours ont développé DeepTox, une architecture de prédiction basée sur un algorithme d'apprentissage profond qui détermine les caractéristiques similaires des pharmacophores toxiques (appelés « toxicophores » dans la publication [40]). Les données sont représentatives de 12 effets toxiques différents, par exemple l'effet de la réponse au stress ou l'effet sur les récepteurs nucléaires. 12000 produits chimiques environnementaux et des médicaments constituent cette base de données. Le jeu d'entraînement est de 11764 composés, le jeu de validation de 296 et le jeu de test de 647. Le critère de performance est l'AUC, qui indique la capacité du modèle à classer un composé chimique toxique d'un autre composé non-toxique. Ces molécules chimiques sont encodées au format SDF (*Structure-Data File*), dans lequel la structure chimique est représentée en graphe labellisé et non-directif. Chaque nœud et arête de ce graphe représentent un atome et une liaison. Deeptox normalise la représentation chimique des molécules. Il calcule de nombreux descripteurs chimiques utilisés en entrée du modèle. Un entraînement de celui-ci est effectué. L'apprentissage profond permet l'abstraction des caractéristiques des molécules, et par conséquent convient très bien à la prédiction de la toxicité. Cette représentation moléculaire dans ce réseau de neurones profonds est hiérarchique.

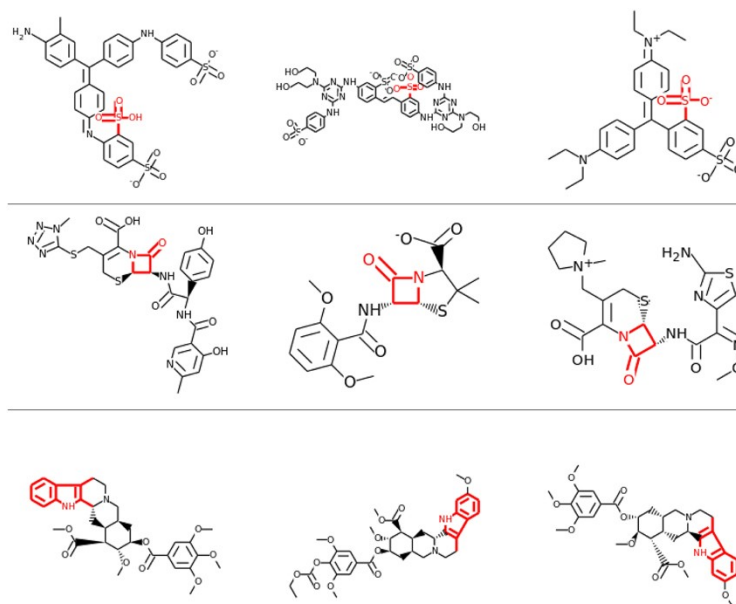


Figure 11 Détection de pharmacophores toxiques par un réseau de neurones profond d'après la source [40]

Dans la figure 11, chaque ligne représente une couche cachée du réseau de neurones corrélée à un pharmacophore possédant le caractère toxique de la molécule, dessiné en rouge. Par exemple, la première couche cachée représente le groupe acide sulfonique détecté dans les trois molécules.

La limite des réseaux de neurones profonds est la quantité de données d'entraînement nécessaire pour les faire tourner. Des travaux de recherche ont mis en évidence un autre type de modèle appelé « *one-shot learning* » qui utilise peu de données. Cette architecture est basée sur une mémoire à long et court terme itérative (*long short-term memory* en anglais), combinée à un graphe de réseaux de neurones convolutifs. Cet algorithme est pertinent dans l'étape d'optimisation du chef de file dans la découverte du candidat-médicament, car il y a très peu de données à ce stade [41].

Chen H, Engkvist O, Wang Y, et al concluent sur les avantages des algorithmes d'apprentissage profond par rapport aux autres méthodes d'apprentissage automatique [42]. Ils estiment que ces premiers ont des architectures plus flexibles et permettent de répondre à une problématique spécifique. Cependant, l'un des désavantages est que l'apprentissage profond nécessite plus de données d'entraînement. D'autres études ont mis en évidence la meilleure performance des modèles d'apprentissage profond multi-tâches par rapport à une approche tâche unique et à un modèle de forêts aléatoires (*Random forests*) [43]. Ce modèle multi-tâche est implémenté dans le logiciel *open source* DeepChem. Un des problèmes que soulignent des chercheurs japonais de l'université de Tokyo, est le manque de diversité dans la génération de nouvelles molécules en utilisant des méthodes d'apprentissage profond. Ils proposent un algorithme génétique, appelé ChemGE, pour *Chemistry Grammatical Evolution* [44]. Cet outil utilise le format SMILES pour la représentation des entités chimiques.

1.5.2 AlphaFold : prédicteur de la structure protéique

La prédiction de la structure d'une protéine permet de déterminer sa fonction. Le géant du web, Google à travers sa filiale dédié à la santé et l'intelligence artificielle Deepmind, a participé à une compétition intitulée l'expérience Casp (pour *Critical Assessment of Protein Structure Prediction*) dont l'objectif est de prédire le repliement d'une protéine en trois dimensions à partir de sa séquence d'acides aminés donnée de manière linéaire [45]. Son algorithme d'intelligence artificielle a surpassé ses concurrents avec une précision de prédiction de 15% à 20% supérieure. Son programme nommé AlphaFold, qui est une déclinaison de son célèbre AlphaGo (programme qui a battu le champion du monde de jeu de Go en 2016), est basé sur des réseaux de neurones et sur des techniques d'apprentissage profond. Ceux-ci incorporent les données des angles de torsion (ϕ, ψ) du squelette peptidique et les distances deux à deux des acides aminés. Les prédictions de distance sont de meilleurs indicateurs de la structure protéique que les prédictions des liaisons de contact. Les données d'entraînement du modèle proviennent de la base de données PDB (*Protein Data Base*). Les scientifiques de Deepmind ont extrait les domaines non-redondant grâce à une base de données intitulée CATH (*Class Architecture, Topology/fold, Homologous superfamily*). Cette base est une classification des structures protéiques provenant de PDB. Le jeu de données généré est formé de 31247 domaines, celui-ci est divisé en un jeu d'entraînement et d'évaluation (29427 et 1820 protéines respectivement).

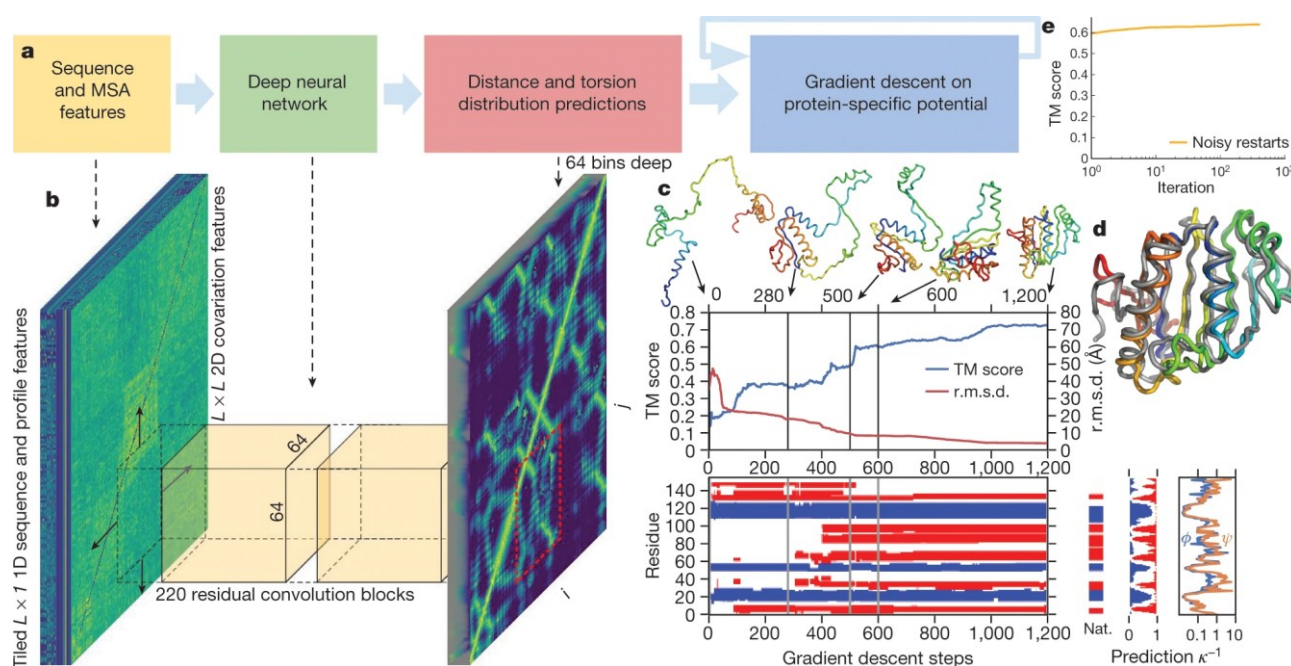


Figure 12 Processus de prédiction de la structure protéique extraite de la publication [45]

La figure 12 synthétise la méthode employée par les chercheurs de Deepmind. La partie (a) représente les différentes étapes de la prédiction de la structure protéique. La partie (b) montre la prédiction de la distribution des distances et des angles de torsion des structures des protéines, à partir des caractéristiques d'alignement des séquences multiples (MSA pour *Multiple Sequence Alignment*). Le schéma en (c) affiche les métriques de performance du modèle, le TM (*Template Modelling*) score associé avec le r.m.s.d (*root mean square deviation*) de la distance en Å en ordonnées sur le graphique. En abscisse, nous trouvons le nombre de pas d'itération de la fonction descente de gradient. 5 images de structures protéiques accompagnent ce graphique. En dessous de ce graphique, un autre représente le nombre d'acides aminés en fonction des pas de la descente de gradient. Le modèle est constitué d'un réseau de neurones convolutifs identique à celui utilisé dans la reconnaissance d'images. Il prédit la probabilité de distribution de la distance entre chaque atome de carbone C_{α} des paires d'acides aminés. Celles-ci sont vectorisées dans une matrice de dimension 64 x 64 représentant les régions de la L x L matrice de distance. 64 x 64 symbolise les distances par paire entre 64 résidus consécutifs et un autre groupe de 64 résidus.

De nouvelles perspectives s'ouvrent à l'industrie pharmaceutique dans le « design » de protéines-médicaments capables de cibler avec précision les sites actifs d'autres protéines, situées par exemple sur le VIH.

Dans cet exemple d'application de prédiction de la structure protéique, nous avons vu l'utilisation d'un réseau de neurones profonds. Les réseaux de neurones convolutifs peuvent être utiles pour la découverte d'un traitement contre un virus, et nous allons en prendre connaissance par la suite.

1.5.3 La recherche d'un traitement contre le virus à Ebola

La société Atomwise (voir annexe 1) a développé une plateforme de prédiction de l'activité biologique des molécules basée sur leur structure. Cette technologie appelé AtomNet repose sur des réseaux de neurones convolutifs profonds. Elle extrait des informations issues de millions de mesures expérimentales d'affinité et de milliers de structures protéiques. L'objectif est de prédire la liaison des molécules aux protéines [46]. Les auteurs ont appliqué le même concept de reconnaissance d'image par les réseaux de neurones convolutifs aux molécules. Ils ont décomposé l'interaction entre un candidat-médicament et un système biologique en groupes chimiques d'interaction plus petits, à l'échelle de la liaison hydrogène par exemple. La force d'attraction ou de répulsion peut varier en fonction de la distance, de l'angle et du type de liaisons. Ce phénomène est

localisé au sein du complexe ligand-protéine. Le logiciel AtomNet a été utilisé avec succès pour développer des nouveaux médicaments pour deux maladies : Ebola et la sclérose en plaques [47]. La maladie à virus Ebola est aiguë et grave, souvent mortelle chez l'homme. La société Atomwise a noué un partenariat avec l'Université de Toronto pour son expertise médicale et scientifique et avec IBM pour ses compétences dans l'infrastructure matérielle. Les chercheurs académiques canadiens ont identifié la cible biologique, la glycoprotéine 2 du virus Ebola. Cette dernière permet l'entrée et la fusion du virus dans la cellule hôte humaine. Cette protéine est constituée de trois hélices externes se repliant autour d'un noyau trimérique central. L'hypothèse des chercheurs était que bloquer le changement de conformation de cette glycoprotéine pourrait empêcher l'entrée du virus dans la cellule. La stratégie d'Atomwise a été d'utiliser un jeu de données de 7000 molécules évaluées précédemment au cours d'essais cliniques de phase II ou supérieurs, pour cribler la région d'intérêt de la protéine. Ces composés ayant des données de sécurité clinique validées chez les patients, pouvaient être proposés à des essais cliniques pour le traitement ou la prévention des infections à virus Ebola [46]. Cette stratégie a potentiellement permis de proposer rapidement un traitement dont l'innocuité chez l'Homme a déjà fait ses preuves.

L'actualité récente concernant l'épidémie de coronavirus nous rappelle que des virus potentiellement dangereux pour la santé humaine et animale sont encore présents. Pourrait-on appliquer la même stratégie d'Atomwise pour la découverte d'un traitement contre le virus à Ebola pour ce nouveau virus ? Nous tenterons de répondre à cette question dans le point suivant.

1.5.4 Covid-19 et IA ?

La recherche d'un traitement efficace ou d'un vaccin contre le coronavirus 2019 (ou SRAS-CoV-2 pour Syndrome Respiratoire Aiguë Sévère-CoronaVirus-2) à l'aide de l'IA sera un cas d'école pour les entreprises spécialisées dans le domaine de la découverte de candidats-médicaments. Le but est d'accélérer la guérison des malades et de diminuer la durée du portage viral de façon à limiter la transmission en proposant des traitements médicamenteux. Avant d'aborder les pistes de recherche, rappelons le contexte de cette pandémie.

a) Contexte

Un patient qui a des symptômes d'une grippe se présente aux urgences d'un hôpital en Chine dans la région du Wuhan. Cependant, son état se dégrade et il a du mal à respirer, les médecins décident de le diriger vers un service de réanimation. Nous sommes le 8 décembre 2019. Ce cas officiel de pneumopathie atypique est signalé par un jeune médecin

le Dr Li Wenliang (mort le 7 février 2020) sur un site de réseau social chinois, où il informe des amis qu'une nouvelle épidémie risque de survenir. Une analyse biologique des prélèvements révèle qu'il s'agit d'une infection virale. Une analyse plus poussée par PCR (réaction en chaîne par polymérase) indique qu'il s'agit d'un virus appartenant à la famille des *Coronaviridae*, et qu'il est proche d'un virus connu le SRAS (Syndrome Respiratoire Aiguë Sévère). La microscopie électronique montre que le virus a la forme d'une couronne et confirme qu'il appartient à la famille des coronavirus. Après une enquête épidémiologique, les autorités sanitaires du pays apprennent que le patient du 8 décembre est un vendeur qui travaille dans un marché aux poissons de la ville où des animaux sauvages destinés à la consommation alimentaire humaine sont commercialisés. Les médecins suspectent que ce nouveau virus utilise ces animaux comme vecteur pour infecter l'être Humain. Nous sommes donc en présence d'une zoonose. Elle se propage rapidement et en quelques jours, de nombreux malades affluent vers les hôpitaux. Les autorités politiques et sanitaires chinoises prennent conscience de l'étendue de la propagation de la maladie. Ce triste décor planté, intéressons-nous maintenant aux coronavirus.

Les coronavirus (CoVs) appartiennent à la famille des *Coronaviridae*, de l'ordre des *Nidovirales*. Leur diamètre peut aller de 80 nm à 160 nm. Ces virus enveloppés à ARN simple brin ont une capsidie à symétrie hélicoïdale, d'où leur forme particulière en couronne dont dérive leur nom (voir figure 13). Ils sont répartis en quatre groupes actuellement : alphacoronavirus, bêtacoronavirus, gammacoronavirus, deltacoronavirus [48]. Ils sont caractérisés par des protéines de structure intervenant dans la morphologie du virus et elles sont au nombre de cinq :

- Les protéines de nucléocapside en se liant à l'ARN forment la capsidie pendant l'assemblage du virion.
- Les protéines de membrane (M, confère la figure 13) entourant la capsidie du virion se lient aux précédentes
- Les protéines d'enveloppe (E) en interaction avec les protéines membranaires constituent l'enveloppe virale.
- Les protéines de spicule (S), glycoprotéines membranaires qui donnent l'aspect hérissé de la couronne observée en microscope électronique. Ces protéines comportent deux sous-unités S1 et S2, formant la partie globulaire et la tige du spicule respectivement.

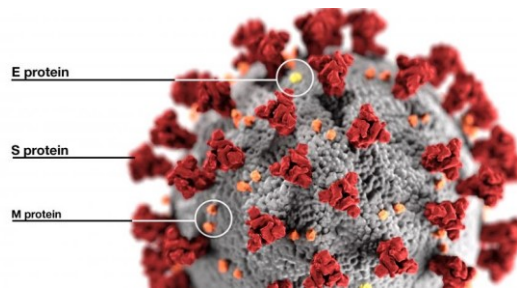


Figure 13 Modélisation en 3D du coronavirus d'après CDC (Centers for Disease Control and Prevention)

L'intelligence artificielle pourrait contribuer à endiguer l'épidémie actuelle dans plusieurs domaines : la détection de signaux faibles dans l'émergence de maladies, la proposition de médicaments pour un traitement efficace, et la conception d'un vaccin.

b) Applications de l'IA à l'épidémie de Covid-19

Tout d'abord, des technologies utilisant du traitement du langage automatique peuvent détecter les signaux faibles en surveillant les sites de réseaux sociaux ou la presse internationale. Par exemple, l'entreprise Blue Dot affiche sur son site Internet qu'elle a été la première au monde à identifier le risque émergent de Covid-19 dans la province chinoise du Hubei. Elle a informé ses clients à partir de sa plateforme Insights. Celle-ci analyse 100 000 sources officielles dans 65 langues par jour pour détecter des épidémies de plus de 150 agents pathogènes différents en temps réel. La dispersion des maladies est prédite en utilisant les données agrégées et anonymisées des itinéraires de vol des voyageurs des aéroports, des données de téléphones mobiles, des données météorologiques, de la capacité du système de santé des pays à gérer l'épidémie, de la population d'animaux et d'insectes. Ainsi, le 31 décembre 2019, le logiciel de détection aurait identifié un article rédigé en chinois où il était relaté un cas de pneumonie atypique de cause inconnue dans un marché à Wuhan. Une alerte aurait été émise à ses clients. Une semaine après, le 9 janvier 2020, l'OMS a notifié au public des cas de grippe en Chine. Les prédictions de l'IA sont analysées et validées par des épidémiologistes pour limiter le risque de fausse alerte [48].

Ensuite, l'Allen Institute for AI et plusieurs groupes de recherche ont mis à disposition un ensemble de données de recherche ouvert Covid-19 (CORD-19 pour Covid-19 Open Research Dataset). Ce corpus de textes est constitué de plusieurs milliers d'articles scientifiques sur la famille des coronavirus et sur la maladie Covid-19. En effet, les coronavirus ont par le passé provoqué des infections respiratoires chez l'Homme et une abondante littérature scientifique a été produite. Ainsi en 2003, le syndrome respiratoire aigüe sévère (SRAS) a touché 8096 personnes et fait 774 morts. Plus récemment en 2012, le MERS-Cov (coronavirus du syndrome respiratoire du Moyen-Orient en français) apparu au Moyen-Orient, a une létalité de 37% avec 1026 cas confirmés par un laboratoire dont 376

patients décédés [49]. Pour ces coronavirus, le réservoir est localisé chez les chiroptères (chauves-souris) qui transmettent le virus à un hôte intermédiaire. Une personne infectée a la capacité de transmettre le virus à une autre personne par contact direct, par la projection de gouttelettes de salives (éternuement, postillons).

Un défi CORD-19 a été publié sur la plateforme spécialisée Kaggle [50] pour encourager la communauté de chercheurs à appliquer des algorithmes de traitement du langage naturel. L'objectif est de générer de nouvelles connaissances pour soutenir la lutte contre cette maladie infectieuse.

L'entreprise BenevolentAI a appliqué un graphe de connaissance issu d'un apprentissage automatique pour proposer un traitement potentiel contre le Covid-19 [51]. Cette information médicale structurée a été extraite de la littérature scientifique en utilisant des techniques du traitement du langage naturel. Les chercheurs ont identifié dans les médicaments approuvés ceux qui inhibent le processus d'infection virale et le baricitinib devrait réduire la capacité du virus à infecter les cellules pulmonaires. La majorité des virus entrent dans les cellules par endocytose médiée par des récepteurs. L'endocytose peut être définie par l'internalisation dans la cellule de substances provenant des espaces extracellulaires ou de la surface membranaire. Le virus du Covid-19 utiliserait le récepteur *ACE2* (*Angiotensin-Converting Enzyme II*), l'enzyme de conversion de l'angiotensine 2, une protéine de surface cellulaire présente dans différents organes (reins, vaisseaux sanguins, cœur et poumons). En particulier, nous la retrouvons dans les cellules épithéliales alvéolaires de type II (AT2) du poumon. La protéine kinase 1 liée à l'AP2 (*AAK1 = AP2-Associated Protein Kinase 1*) est un régulateur de l'endocytose. D'après le résumé des caractéristiques du produit (RCP) du médicament Olumiant (baricitinib), il a une indication dans le traitement de la polyarthrite rhumatoïde active modérée à sévère chez les patients adultes. Le médicament baricitinib est un inhibiteur sélectif et réversible des Janus kinases (JAK)1 et JAK2.

L'entreprise Insilico Medicine a appliqué un modèle génératif profond pour construire informatiquement des structures moléculaires *de novo* avec des propriétés prédéfinies (voir tableau V) [52]. L'architecture du modèle est basée sur un auto-encodeur adverse qui génère une empreinte moléculaire. Les scientifiques ont ciblé la protéase 3C-like du virus Covid-2019. Ils ont produit des nouvelles structures avec trois approches simultanées : l'approche basée sur la poche, sur le ligand et sur la génération d'un modèle homologue.

Tableau V Propriétés prédéfinies des molécules à générer

Propriétés moléculaires	Valeur cible
Log P	1.49-6.00
Masse Moléculaire	400-800
Nombre de liaisons donneurs d'hydrogène	1-10
Nombre de liaisons accepteurs d'hydrogène	2-10
Topologie de l'aire de surface polaire	80-210
MCE-18	40-180
Nombre de centres stéréochimiques	0-3

DeepMind apporte sa contribution à l'effort de recherche sur le Covid-19 en modélisant la structure protéique du virus à l'aide de sa technologie d'IA AlphaFold system. Les chercheurs émettent des mises en garde car pour l'instant aucune publication, ni aucune revue par des pairs n'a été effectué sur leurs travaux [53].

En conclusion, il faut rappeler qu'il n'y a pas de vaccins actuellement, ni de traitement efficace contre ce nouveau virus. De plus, la prudence doit être de mise car les publications citées ici sont en pré-impression et n'ont pas été revues par des pairs. Les auteurs le disent clairement dans leurs articles. D'autre part, plusieurs essais cliniques dans le monde et en France sont en cours avec une approche de repositionnement de molécules sans l'aide de l'IA. L'essai appelé Discovery, conçu dans le cadre du consortium européen *REACTing* (Research and Action targeting emerging infectious diseases), inclus 3200 personnes dont 800 en France. L'objectif est de tester plusieurs médicaments dont des antiviraux seuls ou combinés. Les cinq modalités de traitement sont définies :

- Soins standards (groupe témoin)
- Soins standards plus lopinavir et ritonavir
- Soins standards plus lopinavir, ritonavir et interféron bêta
- Soins standards plus remdesivir (testé contre Ebola).
- Soins standards plus hydroxy-chloroquine

Cet essai clinique randomisé et ouvert a débuté le 22 mars 2020. En France, l'Inserm coordonne l'étude. L'analyse des résultats pour évaluer l'efficacité et la sécurité des traitements sera disponible dans deux semaines après inclusion de chaque patient [54].

Si les virus font l'actualité en ce début d'année 2020, d'autres pathogènes en particulier les bactéries posent également des problèmes de santé publique importants. La recherche de nouveaux antibiotiques pour parer à la résistance des bactéries aux anciennes molécules devient crucial. Un article publié récemment (20 février 2020), relate la découverte d'un antibiotique potentiel à l'aide de l'IA.

1.5.5 Découverte d'un nouvel antibiotique potentiel

Les antibiotiques sont des médicaments essentiels dans l'arsenal thérapeutique de la médecine pour soigner les infections bactériennes. Cependant, avec le temps et le recours abusif à ces antibiotiques, une antibiorésistance s'est développée. Stokes J M, Yang K, Swanson K, et al ont appliqué un modèle de réseau de neurones profonds qui représente les molécules avec une propriété spécifique, celle d'inhiber la croissance de la bactérie *E. coli* [55]. Ils ont utilisé une approche par passage de messages directs (traduction littérale de l'expression anglaise « *directed message passing approach* »). D'autre part, la stratégie de recherche appliquée est le repositionnement d'anciennes molécules qui sont au stade préclinique, au stade clinique et dont l'approbation par les instances réglementaires n'a pas abouti ou des médicaments approuvés pour d'autres aires thérapeutiques que l'infectiologie. Les chercheurs ont posé comme prérequis pour la construction de leur jeu d'entraînement que les données soient peu coûteuses, chimiquement diversifiées et ne nécessitant pas de ressources de laboratoires sophistiquées pour effectuer les expériences. Le premier jeu d'entraînement est constitué par 1760 molécules approuvés par la FDA, plus 800 produits naturels issus de plantes, d'animaux, et de sources microbiologiques. Ce qui fait un total de 2560 molécules. Ils ont ensuite supprimé les doublons présents dans la base, et finalement 2335 composés uniques ont été utilisés pour le jeu d'entraînement. Ainsi, les chercheurs du MIT (*Massachusetts Institute of Technology*) ont entraîné leur modèle avec ce jeu de données pour classer les molécules en actives ou inactives. Par la suite, ils ont utilisé ces données pour entraîner un modèle de classification dont la sortie est la probabilité qu'une nouvelle molécule en fonction de sa structure inhibe la croissance d'*E. coli*. Ce modèle d'apprentissage profond est une approche par passage de messages directs qui consiste à traduire les informations présentes dans la représentation géométrique d'une molécule en un vecteur continu en partant des liaisons des atomes voisins (voir figure 14 extraite de [55]). Cette représentation moléculaire est construite en additionnant de façon itérative les caractéristiques de chaque atome et liaison. Cette itération est appliquée plusieurs fois pour avoir des informations étendues des groupes fonctionnels chimiques. Ils ont amélioré le modèle, par la suite, avec un jeu de caractéristiques moléculaires, et avec l'optimisation d'hyperparamètres. En fonction du score de prédiction du modèle établi à 80% d'inhibition de croissance d'*E. coli*, une première liste de 120 candidats-médicaments a été sélectionnée. Une autre bibliothèque chimique appelé *Drug Repurposing Hub* [56] (plateforme de repositionnement de médicaments) de la Broad Institute a été utilisée. Cette bibliothèque contient 6111 composés chimiques à la date de janvier 2020 à différents stades de développement. 99 molécules ont été prédites avec un pouvoir antibactérien potentiel par le modèle. Après une vérification expérimentale par un test d'inhibition de croissance d'*E. coli* et en prenant comme seuil de décision une DO_{600} (densité optique à 600) inférieure à 0.2

pour quantifier le nombre de bactéries, il en résulte que 51 molécules sont effectivement inhibitrices de la croissance bactérienne. En prenant en compte d'autres critères de sélection, les chercheurs ont identifié qu'une molécule sur les 51 vérifiées ces exigences : halicine (nouveau nom donné au composé SU 3327 qui existe depuis 2009 [57]). D'autre part, ce modèle a été utilisé avec deux autres bibliothèques chimiques, la WuXi anti-tuberculose (9997 molécules en janvier 2020) de la Broad Institute et un extrait de la base de données Zinc15 (107349233 molécules en janvier 2020) [58] pour identifier d'autres potentiels composés chefs de file avec une activité anti- *E. coli*. L'approche innovante des réseaux de neurones est dans leur capacité à apprendre cette représentation moléculaire automatiquement, en cartographiant les molécules en vecteurs continus, utilisés pour prédire leurs propriétés. Les auteurs de cet article [55] ont réussi à mettre en évidence qu'un inhibiteur de la protéine kinase au niveau du N-terminal de l'enzyme *c-Jun N-terminal kinase* (*JNK*) a des propriétés antibactériennes en effectuant des tests à la « paillasse » sur des boîtes de Petri avec un protocole expérimental classique en microbiologie. Cet inhibiteur qui était en phase préclinique de développement pour une autre aire thérapeutique, le SU3327 [57] a été renommé en halicine en référence à HAL 9000, un supercalculateur doté d'intelligence artificielle de la saga des livres de science-fiction Odyssées de l'espace du romancier britannique Clarke A C. Initialement, ce composé était pressenti pour le traitement du diabète [57], mais des tests non-concluants ont arrêté son développement. La structure moléculaire de l'halicine (voir figure 15) diffère des antibiotiques connus actuellement. Nous pouvons observer la présence d'un groupe thiadiazole avec une amine en position 2, un groupe thiazole et un groupe nitro. Sa formule brute est $C_5H_3N_5O_2S_3$ et son nom IUPAC est le : 5-[(5-Nitro-1,3-thiazol-2-yl) sulfanyl] -1,3,4-thiadiazol-2-amine.

Ce candidat médicament est un antibiotique à large spectre et il a une activité bactéricide sur d'autres pathogènes tels que : *Mycobacterium tuberculosis*, *Clostridioides difficile*, *Acinetobacter baumannii*, *Enterobacteriaceae* résistantes aux carbapénèmes. Pour l'instant, les essais ont été effectués sur des souches bactériennes issues de patients et des modèles murins. En revanche, il est inefficace contre *Pseudomonas aeruginosa*, car sa membrane cellulaire serait probablement imperméable à l'halicine. Le mécanisme d'action d'halicine est spécial car ce composé induirait une baisse de la régulation des gènes impliqués dans la mobilité de la cellule bactérienne et une augmentation de la régulation des gènes nécessaires dans l'homéostasie du fer. Le fer est un élément essentiel pour la croissance et le développement des microorganismes. La chélation de halicine avec le fer perturberait le gradient électro-chimique de la bactérie qui voit sa production d'adénosine triphosphate (ATP) bloquée [59]. Et comme toute cellule vivante, sans énergie, la bactérie meurt. Ce principe de fonctionnement fait que la bactérie ne peut pas acquérir une mutation qui lui permettrait de résister à l'antibiotique. Cette étude montre qu'*E. coli* n'a développé aucune résistance à l'halicine pendant une période de 30 jours de traitement [55]. Les scientifiques

veulent effectuer des essais cliniques et sont à la recherche d'un partenariat avec un laboratoire pharmaceutique ou une organisation à but non lucratif pour prendre en charge les coûts de développement de ces essais. Halicine reste donc actuellement un candidat médicament prometteur qui est au stade préclinique. La recherche scientifique est ici confrontée à la réalité économique du marché des médicaments. Est-ce qu'un investissement financier dans une molécule qui a déjà été en phase préclinique et qui a peut-être un dépôt de brevet en cours est rentable pour un laboratoire pharmaceutique ? La réponse à cette question reste en suspens et l'avenir nous dira si halicine obtiendra une autorisation de mise sur le marché prochainement.

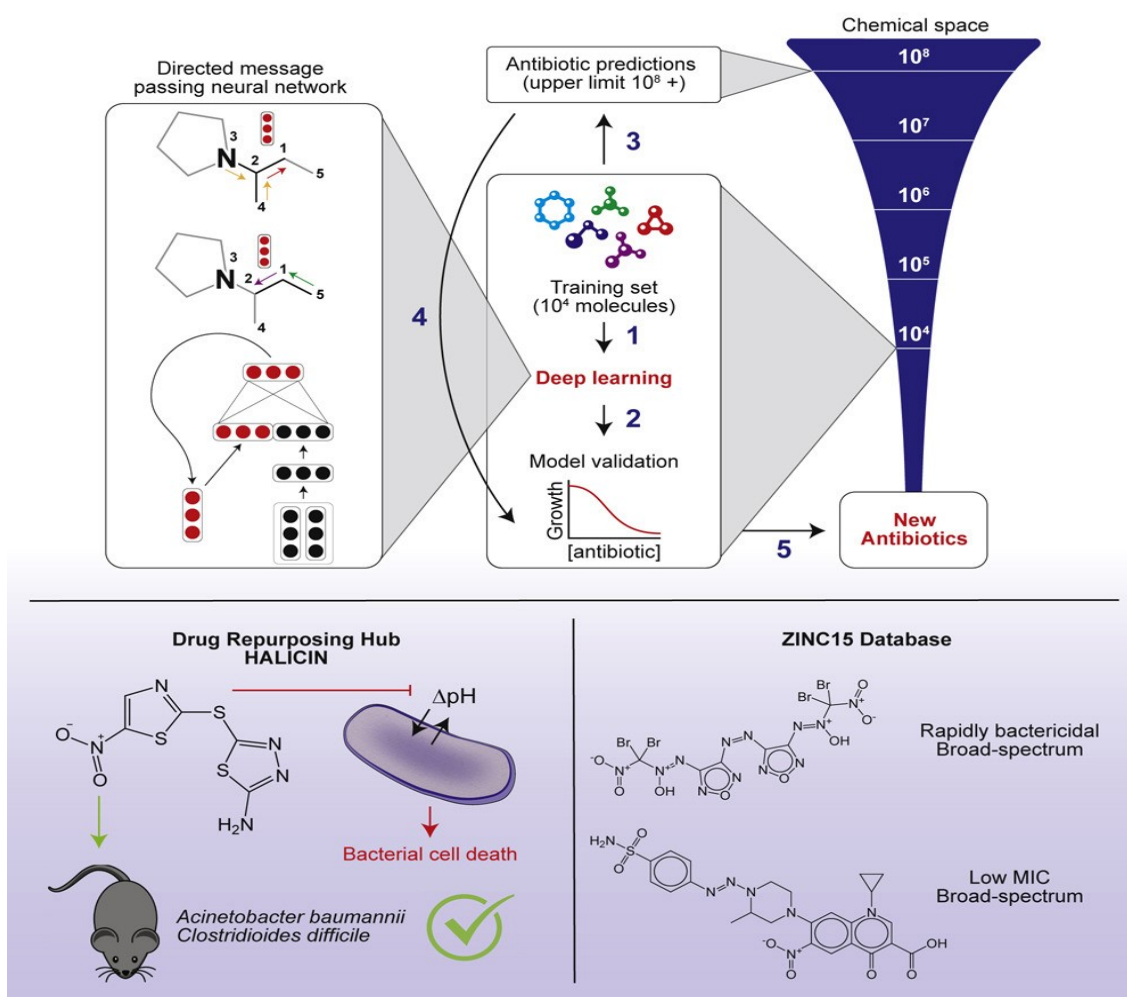


Figure 14 Schéma résumant l'article de cette référence [55]

Le code informatique du modèle est disponible à cette adresse :

<https://github.com/chemprop/chemprop>.

Une démonstration du modèle entraîné est accessible sur ce site internet :

<http://chemprop.csail.mit.edu/>. Vous pouvez cliquer sur l'onglet « predict » et entrer le nom SMILES de la molécule d'halicine : C1=C(SC(=N1)SC2=NN=C(S2)N)[N+](=O)[O-].

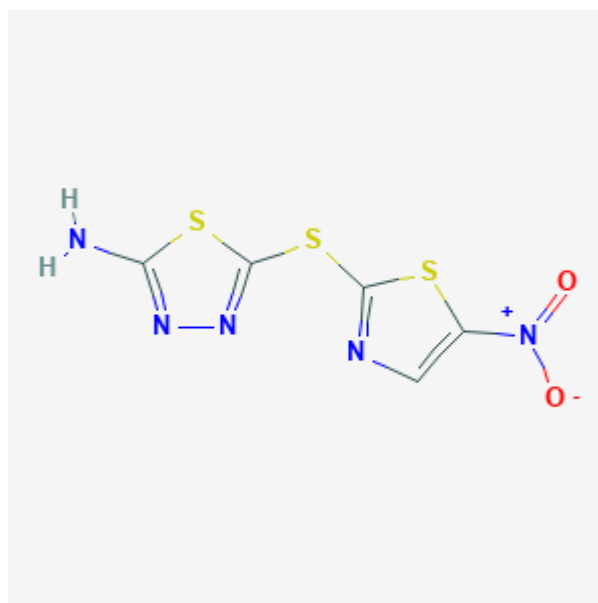


Figure 15 Structure en 2D de la molécule d'halicine

1.5.6 ClinicalTrial.gov et IA

Une recherche sur la base de données américaine des essais cliniques à l'adresse : <https://clinicaltrials.gov/> avec les termes « *artificial intelligence* » comme mots clés entrés dans le champ intitulé « Other terms » nous renvoie un résultat de 249 essais cliniques. 86 études ont lieu en Chine, 50 aux Etats-Unis et 50 en Europe dont 16 en France. Cette requête a été soumise le 08 avril 2020. Nous observons la prédominance de la Chine qui a enregistré environ cinq fois plus d'essais cliniques que la France. La majorité de ces essais portent sur des dispositifs médicaux. L'utilisation de l'IA dans le développement de médicaments ou dans des outils d'aide au diagnostic est effective, cependant ce n'est que le début, et il est difficile de savoir combien de médicaments ont obtenu une autorisation de mise sur le marché à la suite d'une proposition d'une technologie d'IA. De plus, l'IA intervient dans certaines étapes du processus de découverte, mais elle n'est pas présente partout. Ainsi, par exemple, l'entreprise britannique Exscientia et son partenaire pharmaceutique, le japonais Sumitomo Dainippon Pharma ont commencé une étude clinique en phase I du DSP-1181. Ce candidat-médicament créé à l'aide de la plateforme d'IA Centaur Chemist™ aura comme indication le traitement du trouble obsessionnel compulsif (TOC) au Japon. Cette annonce a été publiée le 30 janvier 2020 [60]. DSP-1181 aurait une durée d'action longue et serait un agoniste potentiel du récepteur à la sérotonine 5-HT_{1A}. Les avancées scientifiques de l'apprentissage machine dans la découverte de nouveaux candidats médicaments sont une réalité difficile à quantifier (nombres de médicaments mis sur le marché), comme nous venons de le voir avec l'halicine. Cependant, c'est dans le domaine de l'analyse d'images que les algorithmes d'apprentissage ont fait leur preuve de concept et les applications informatiques utilisant de l'IA sont en production.

2 Application à l'imagerie médicale

L'objectif de l'utilisation des techniques d'IA dans l'imagerie médicale est de permettre le développement d'outils d'aide au diagnostic et au suivi de pathologies, principalement de cancers. Ces nouveaux logiciels amélioreront la rapidité et la précision de la détection de la lésion. Actuellement, l'état de l'art est l'utilisation des réseaux de neurones convolutifs pour concevoir ces logiciels. Les disciplines médicales bénéficiant de l'application de l'IA à l'imagerie médicale sont par exemple la cancérologie, la dermatologie, l'ophtalmologie, la cardiologie, la neurologie, l'anatomopathologie, la radiothérapie et la chirurgie guidées par l'image [61] [62] [63] [64]. Toutes les disciplines médicales sont ou seront concernées par l'IA appliquée à l'imagerie.

L'IA peut capitaliser sur l'utilisation de l'informatique dans le domaine de l'imagerie médicale depuis de nombreuses années comme le serveur d'images *Picture Archiving and Communication System* (PACS) contenant des images digitalisées sous le format de fichier DICOM (*Digital Imaging and Communications in Medicine*).

Les logiciels actuels tels que les CAD (*Computer Aided Diagnosis*) sont encore largement utilisés pour l'interprétation. Cependant, les nouveaux outils intégrant de l'IA tendent à les remplacer.

Avant d'aborder la présentation de quelques cas d'usage de ces systèmes informatiques d'IA, nous donnerons une définition d'une image numérique, et présenterons les réseaux de neurones convolutifs.

2.1 Définition d'une image

D'un point de vue mathématique, une image est une fonction qui quantifie l'intensité lumineuse de chaque point de cette image. Ainsi une image en noir et blanc est quantifiée par l'intensité de son niveau de gris (voir figure 16). Un point sombre correspond un niveau de gris faible. Pour une image en couleurs, l'intensité d'un point désigne sa couleur qui est vue comme un mélange de trois couleurs primaires (rouge, vert et bleu). L'image numérique se caractérise par sa définition et sa résolution. La définition est le résultat de sa hauteur multipliée par sa largeur, avec comme unité le pixel. La résolution est le nombre de pixels par unité de longueur de l'image analogique. Cette mesure est un indicateur de la qualité de l'image numérique. Les images médicales peuvent être en 2D, 3D, ou en 4D.

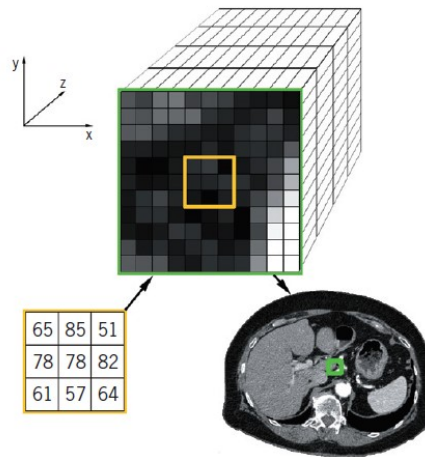


Figure 16 Image de CT-scanner représentée en pixels

Les modalités d'acquisition des images sont le scanner, l'IRM (imagerie par résonance magnétique), l'échographie, la radiographie.

L'échographie est la réflexion d'ondes ultra-sonores sur les interfaces des tissus.

La radiographie est la projection sur un film ou un écran d'un faisceau de rayons X atténué par les différents tissus. Le contraste dans les images est dû au degré d'absorption des faisceaux de rayons X par les tissus. Plus le coefficient d'absorption est élevé, plus le faisceau de rayons X sera atténué. Ainsi, une structure de forte intensité sera visible sur les images. Le contraste est défini sur une échelle de gris à quatre niveaux de densité : la densité calcique (l'os et les calcifications), hydrique (le tissu mou et les liquides), grasseuse, et la densité aérique [65].

Le scanner ou *Computerized Tomography* (voir figure 16) est la reconstruction d'images en coupes à partir des profils d'atténuation des rayons X dans différentes directions.

L'IRM permet d'obtenir des images en coupes basées sur les propriétés magnétiques des tissus sous l'effet d'un champ magnétique externe.

2.2 Les réseaux de neurones convolutifs

Les CNN ou ConvNet pour *Convolutional Neural Network* sont actuellement les meilleurs modèles pour classifier des images [66]. Ils ont été développés par Yann Le Cun en 1988-1989 pour les premières versions. Ils sont un sous-ensemble des réseaux de neurones, possédant une architecture spécifique spécialisée dans le traitement des images en entrée. Cette architecture est constituée de deux composantes principales :

- Un extracteur de *features* ou caractéristiques
- Un vecteur de sortie contenant autant d'éléments qu'il y a de classes

Un CNN comporte habituellement des couches de convolution entrecoupées de couches d'agrégation (*pooling layers*) ou de sous-échantillonnage suivies de couches entièrement connectées (*fully connected layers*) (Figure 17) [67].

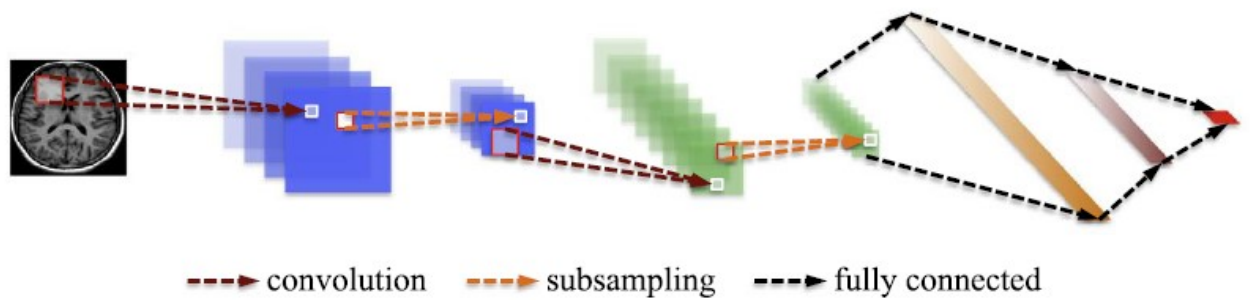


Figure 17 Exemple d'architecture d'un CNN d'après la référence [67]

La méthode de *pooling* permet de réduire une image de grande taille sans perte d'informations importantes. Le principe est de faire glisser une fenêtre de petite taille (par exemple 3×3 pixels) sur l'image initiale et de garder la valeur maximum de cette fenêtre à chaque déplacement. Les termes de filtres ou noyaux sont également utilisés à la place de fenêtres dans la littérature des algorithmes. Ceux sont des matrices de dimensions $n_k \times n_k$, avec n_k entier. Par exemple, si nous avons une image médicale de taille 512×512 en entrée d'un réseau de neurone convolutif, et que chaque neurone doit prendre en compte chaque pixel de cette image, alors une énorme quantité de mémoire informatique serait nécessaire. Pour économiser les ressources matérielles, seuls les paramètres des opérateurs des filtres spécialisés, appelés convolutions, sont appris. Cette opération mathématique consiste en la multiplication des pixels voisins d'un pixel donné par un petit tableau de paramètres appris appelés noyau. La figure 18 décrit la convolution d'une image par un noyau de taille 3×3 . Les pixels de l'image sont multipliés par les neuf valeurs d'un noyau 3×3 (carré rouge) et additionnés pour produire la valeur du pixel bleu (carré bleu dans la figure 18). Cette opération est répétée pour couvrir toute l'image. En apprenant des noyaux significatifs, cette opération reproduit l'extraction d'éléments visuels tels que les contours et les formes, comme le fait le cortex visuel des mammifères [68].

Les réseaux convolutifs sont utilisés pour des tâches de classification, de détection ou de segmentation. L'algorithme classe les images d'entrée en deux ou plusieurs catégories en sortie. Par exemple, dans une mammographie, les kystes sont bénins ou malins.

Dans la tâche de détection, l'algorithme doit localiser les structures d'intérêts (métastases par exemple sur une image CT (*Computerized Tomography*) dans un espace en deux ou trois dimensions.

La segmentation consiste à délimiter un organe ou une pathologie.

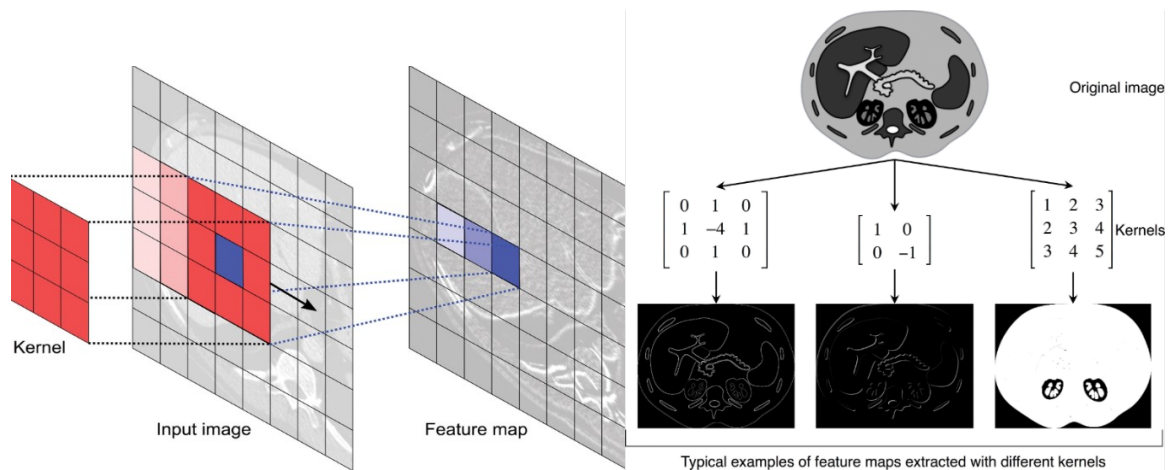


Figure 18 Filtres appliqués à une image en entrée d'après la référence [68]

2.2.1 Application clinique concrète

Nous présentons ci-dessous quelques exemples de logiciels utilisés en pratique dans la détection de différentes pathologies en imagerie médicale.

a) En ophtalmologie

Un logiciel, nommé IDx-DR a été certifié par l'agence nationale du médicament américaine (FDA) dans l'aide au diagnostic de la rétinopathie diabétique en avril 2018 [69]. Ce système d'IA est capable d'établir un diagnostic sans être supervisé par un médecin (voir annexe 4).

La rétinopathie diabétique est une atteinte évolutive de la rétine touchant la moitié des patients diabétiques. Le risque à long terme pour ces personnes est la cécité si cette pathologie n'est pas détectée précocement. Une surveillance régulière est instaurée pour évaluer son apparition.

L'algorithme est précis à 90% (d'après un essai clinique réalisé sur 819 patients), un taux de précision supérieur à celui des ophtalmologistes en moyenne. C'est surtout la rapidité de détection qui est remarquable : une ou deux minutes [70] [71].

b) En cancérologie

Transpara™ est un logiciel d'IA pour l'interprétation des mammographies développé par l'entreprise ScreenPoint Medical [72] [73]. Il a obtenu la certification 510(k) de la part de la FDA en novembre 2018. Il avait déjà le marquage CE en Europe. Cette application d'intelligence artificielle est utilisée pour détecter les cancers du sein en examinant les anomalies suspectes des mammographies. Elle aide également les radiologues à décider si un suivi des anomalies est nécessaire. Ce logiciel détecte automatiquement les lésions des tissus mous et les microcalcifications [74]. Il établit un score de suspicion de cancer en combinant ces résultats. Ce logiciel implémente un réseau de neurones convolutifs pour effectuer cette tâche [75].

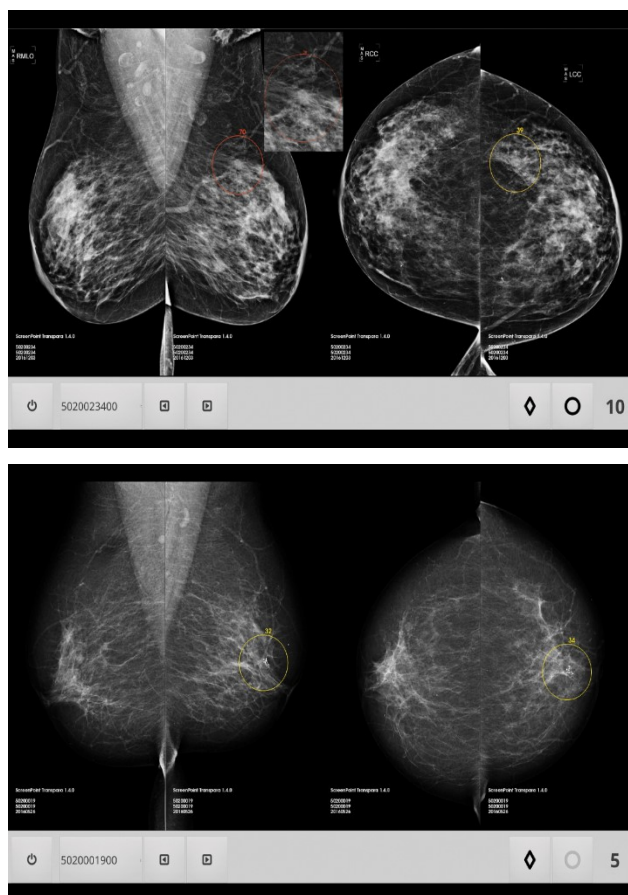


Figure 19 Exemples de sortie du logiciel Transpara™

La figure 19 à gauche (image extraite de la référence [75]) représente des mammographies chez une femme de 71 ans présentant un carcinome canalaire invasif (zone entourée en rouge sur la figure 19 avec un score attribué par le système informatique). La patiente a été rappelée (score BI-RADS ≥ 3) par quatre des quatorze radiologistes lors d'une lecture sans aide du logiciel et par onze des quatorze radiologues utilisant un système d'intelligence artificielle (voir annexe 5 pour une image agrandie).

La figure 19 à droite illustre des mammographies chez une femme de 62 ans sans cancer, qui a été rappelée (score BI-RADS, ≥ 3) par douze radiologues sur quatorze en lecture non assistée par le logiciel et par sept lecteurs sur quatorze lors de l'utilisation du système d'IA (confère annexe 6 pour une image agrandie) [75]. BI-RADS est l'acronyme de *Breast Imaging Reporting And Data System*.

La conclusion de cette étude [75] est que les radiologistes ont amélioré leur détection du cancer lors de la lecture d'une mammographie lorsqu'ils utilisaient un système d'intelligence artificielle comme aide au diagnostic, sans nécessiter de temps de lecture supplémentaire.

c) En neurologie

La détection des accidents vasculaires cérébraux (AVC) [76] par le logiciel d'aide à la décision clinique Viz.AI Contact [77] permet une prise en charge plus rapide du patient par les neurologues. Cette application est conçue pour analyser les images

tomodensitométries du cerveau et prévenir par une notification (un message textuel) un neurochirurgien si une obstruction a été identifiée. Cette plateforme logicielle est construite à partir d'un réseau de neurones convolutifs.

d) En imagerie ostéo-articulaire

Le CHRU de Nancy vient de s'équiper d'un scanner de dernière génération embarquant une technologie d'IA. Il est nommé Aquilion Precision par le constructeur Canon. Ce scanner permet d'obtenir une résolution d'image améliorée et une granularité fine des détails. Le logiciel, appelé AiCE (*Advanced Intelligent Clear-IQ Engine*), est basé sur des méthodes de réseaux de neurones convolutifs profonds. L'équipement est certifié 510(K) par la FDA. L'objectif de cette acquisition est d'optimiser la prise en charge des patients [78]. Il sera utilisé pour établir un diagnostic et en radiologie interventionnelle.

2.3 Vers une médecine personnalisée

La mise en place des dossiers médicaux électroniques permet de rassembler en théorie les données hétérogènes du patient : les résultats des bilans sanguins, la liste des traitements médicamenteux, les images médicales. Les techniques d'IA permettent d'exploiter ces bases de données comportant plusieurs paramètres pour personnaliser les traitements en fonction de chaque phénotype, génotype du patient. D'après cette revue générale intitulée « Intelligence artificielle appliquée à la radiothérapie » [62], l'apprentissage machine appliqué au domaine de l'oncologie peut prédire la radiothérapie. Nous allons vers une médecine des « 4 p » : personnalisée, précise, préventive et prédictive. L'IA permet de modéliser des représentations numériques du patient, c'est le concept de jumeau numérique.

2.4 Limites actuelles de l'IA

Dans la découverte de médicaments, les solutions d'IA ont proposé des molécules qui sont actuellement dans les *pipelines* des entreprises à différents stades de développement. Cependant, à ce jour aucun médicament n'a été approuvé par les instances réglementaires américaine, européenne ou japonaise (FDA, EMA, PMDA). Il est donc prématuré de quantifier le retour sur investissement des laboratoires pharmaceutiques ou des entreprises spécialisés en IA. Les annonces faites par les différents acteurs de la recherche en IA appliquée à la santé peuvent aussi servir à leur communication d'entreprise. Cela leur permet de réaliser une publicité peu coûteuse et d'attirer des investisseurs pour leur levée de fonds nécessaires à leur développement.

D'autre part, la découverte de candidat-médicaments par ces entreprises ou par les chercheurs académiques doit continuer son cycle de développement. Ces entreprises ou laboratoires publics n'ont pas les moyens financiers ou les compétences pour effectuer les autres phases du développement. Ainsi, si aucun partenariat avec un laboratoire pharmaceutique n'a été engagé, des molécules prometteuses pourraient ne pas obtenir d'autorisation de mise sur le marché. Par exemple, le futur antibiotique halicine vu précédemment (partie B) serait potentiellement utile pour la santé publique, mais si une étude de marché d'une entreprise pharmaceutique estime que ce n'est pas rentable, ce nouveau traitement contre plusieurs bactéries ne verra pas le jour. Cette remarque peut s'appliquer aux maladies émergentes comme celle que nous vivons actuellement (Covid-19). La recherche fondamentale sur les coronavirus n'est pas récente et depuis l'apparition du syndrome respiratoire aigüe sévère en 2003, soit 17 ans, aucun candidat-vaccin n'a été validé et mis sur le marché. L'épidémie a certes « disparu », mais le MERS-Cov en 2012 a été un rappel de la dangerosité de certains coronavirus. Alors que l'industrie pharmaceutique développe et met sur le marché un vaccin contre la grippe saisonnière chaque année. La recherche est une chose et la réalité économique en est une autre. Nous pouvons constater (voir annexe 1) que le domaine thérapeutique de prédilection des entreprises d'IA est l'oncologie. Cet axe de recherche est fondé sur des besoins de santé non satisfaits, mais peut-être aussi sur le fait que l'oncologie représente un secteur lucratif pour les laboratoires pharmaceutiques.

Pour rappel, nous avons vu que la qualité d'un algorithme est liée à la qualité de la base d'images qui a été utilisée pour la conception de son modèle et sur laquelle il a été testé [79]. La représentativité de la base entre dans ce critère de qualité. Les images médicales et les comptes rendus sont-ils issus d'un échantillon de patients avec des caractéristiques ethniques spécifiques. Par exemple, l'utilisation d'un modèle entraîné sur une base de données chinoise appliquée à la détection du cancer du sein ne sera pas transposable sur une population européenne. Cela constitue un biais statistique de l'échantillon de la population étudiée. Une autre limite est celle de la méthode d'apprentissage supervisé utilisée qui nécessite des données labellisées par des experts du domaine en quantité importante. Cela coûte cher et est chronophage pour les radiologues.

Un autre point irritant est le phénomène de boîte noire qui traduit le manque d'explicabilité des décisions prises par les algorithmes des réseaux de neurones et peut s'avérer être un frein à son usage, notamment dans le diagnostic des pathologies. Cela pose également des questions d'ordre éthique.

3 Résumé des applications de l'IA

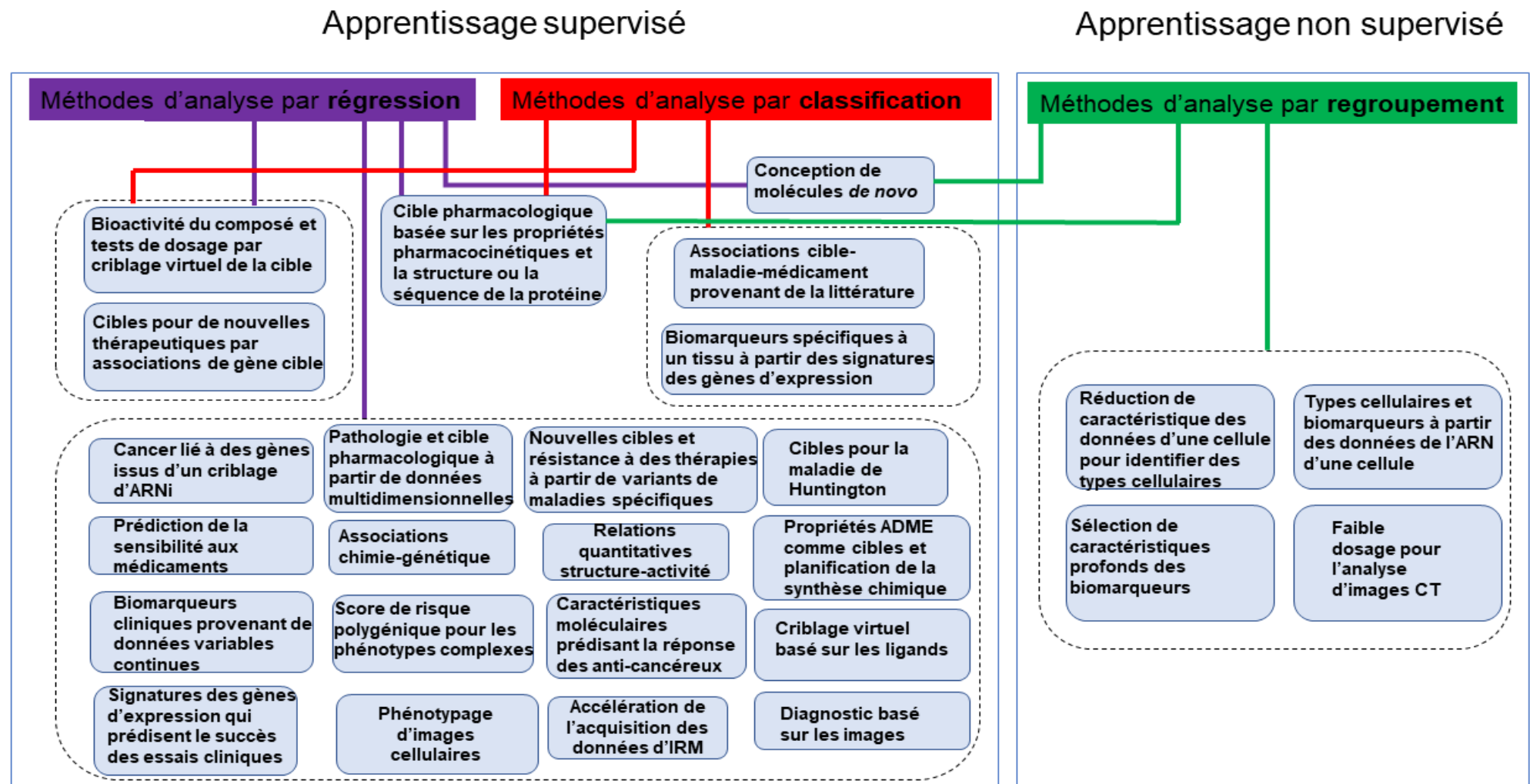


Figure 20 Méthodes d'apprentissage automatique et leurs applications dans la découverte de médicaments et en imagerie d'après [80]

La figure 20 synthétise les différentes applications de l'IA dans la découverte de nouvelles molécules et en imagerie. Pour compléter cet instantané de l'état de l'art en IA, nous allons récapituler les différents exemples vus précédemment dans le tableau VI suivant.

Tableau VI Tableau de synthèse récapitulatif des applications mentionnées dans ce document

Applications	Nom du logiciel	Références
Détection de toxicophores	Deeptox	[40]
Prédiction de la structure protéique	AlphaFold	[45]
Prédiction de l'activité biologique	AtomNet	[46][47]
Repositionnement de molécules	Chemprop	[55]
Repositionnement de molécules	Centaur Chemist™	[60]
Diagnostic de la rétinopathie diabétique	IDx-DR	[70][71]
Diagnostic du cancer du sein	Transpara™	[72][73][74][75]
Détection des AVC	Viz.AI	[76][77]
Amélioration de la précision des scanners	AiCE	[78]

La partie suivante traite de la réglementation et de l'éthique dans le domaine de l'IA. Nous nous intéresserons aux données personnelles qu'utilisent les algorithmes d'apprentissage machine. Nous poserons la question de la responsabilité juridique en cas de mauvais diagnostic fait à l'aide d'un logiciel à base d'IA. Nous aborderons l'explicabilité et l'acceptabilité des algorithmes. Enfin, nous verrons quelles pourraient être les conséquences de l'utilisation de l'IA dans le monde du travail.

C Ethique et réglementation

De nombreuses questions juridiques et éthiques sont soulevées par l'utilisation des données de santé dans les applications d'intelligence artificielle [81]. Le 2 décembre 2019, un communiqué de presse du Comité Consultatif National d'Éthique pour les sciences de la vie et de la santé (CCNE) annonce la création d'un comité pilote d'éthique du numérique. Il a été chargé par le Premier ministre de constituer ce comité qui traitera des enjeux éthiques de l'IA [82].

1 L'éthique en IA

Le rapport de synthèse du débat public animé par la CNIL (Commission Nationale Informatique et Libertés) sur les enjeux éthiques des algorithmes et de l'intelligence artificielle de décembre 2017 énonce deux principes importants :

- La loyauté : l'outil algorithmique ne doit pas trahir la confiance de l'utilisateur (en sa qualité de consommateur ou de citoyen) dans le traitement des données personnelles ou non.
- La vigilance/ réflexivité : mise en place de procédures qui interrogent régulièrement et méthodiquement les différents acteurs de la chaîne algorithmique (concepteur, utilisateur) dans le cadre d'un comité d'éthique par exemple [83].

Ce rapport préconise également six recommandations opérationnelles : «

- i. Former à l'éthique tous les acteurs-maillons de la « chaîne algorithmique »
- ii. Rendre les systèmes algorithmiques compréhensibles
- iii. Travailler le design des systèmes algorithmiques au service de la liberté humaine
- iv. Constituer une plateforme nationale d'audit des algorithmes
- v. Encourager la recherche sur l'IA éthique et lancer une grande cause nationale participative autour d'un projet de recherche d'intérêt général
- vi. Renforcer la fonction éthique au sein des entreprises [83] ».

Un autre rapport sur l'éthique de la recherche en apprentissage machine réalisé par la commission CERNA , la Commission de réflexion sur l'Ethique de la Recherche en sciences et technologies du Numérique d'Allistene (l'Alliance des Sciences et Technologies du Numérique), nous indique que la question éthique des algorithmes est un sujet de réflexion pour les chercheurs [84]. Le travail de la commission énonce les propriétés générales des systèmes numériques.

Au niveau international, l'association professionnelle internationale du numérique *l'Institute of Electrical and Electronics Engineers*, a initié *l'IEEE Global Initiative for Ethical Considerations in Artificial Intelligence and Autonomous Systems* qui a rédigé fin 2016 un rapport d'étape

Ethically Aligned Design [85].

Aux Etats-Unis, la réflexion est également engagée sur l'impact de l'IA sur les individus et la société, notamment par l'Université de Stanford qui a publié un rapport en septembre 2016 [86].

En Europe, une démarche identique a été initiée par L'*European Data Protection Supervisor* (EDPS) ou en français le Contrôleur Européen de la Protection des Données (CEPD) [86] [87].

Les institutions nationales ou internationales, les chercheurs, les utilisateurs, les citoyens ont pris la mesure de l'importance des interrogations éthiques que soulèvent le numérique et l'IA en particulier [84].

Les entreprises privées tels que les GAFA (Google Amazon Facebook Apple) y vont également de leurs comités d'éthiques sur l'IA. Est-ce un effet d'annonce à la suite de scandales sur la protection des données personnelles, comme celui du scandale Cambridge Analytica ? Cette société britannique a utilisé des données provenant du site de réseau social Facebook de façon illégale. L'objectif était de manipuler l'opinion publique dans le cadre d'élections politiques aux Etats-Unis et d'autres pays. Elle avait également des activités de *lobbying* auprès d'entreprises. La société a fermé en 2018. La réponse à la question précédente pencherait vers l'affirmative car aussitôt créé, Google a décidé de fermer son comité d'éthique suite à une plainte de ses employés pour discrimination [89].

2 Données personnelles

Le développement de l'IA repose sur des données qui doivent être accessibles, interopérables, et de préférence en libre accès (*open data*), tout en préservant leur protection.

Le règlement général sur la protection des données (RGPD) du 27/04/2016 est entré en vigueur le 25/05/2018. Il donne une définition des données de santé qui est : « les données à caractère personnel relatives à la santé physique ou mentale d'une personne physique, y compris la prestation de services de soins de santé, qui révèlent des informations sur l'état de santé de cette personne ». Ces données peuvent être les informations présentes dans les dossiers électroniques des médecins, des pharmaciens (dossier pharmaceutique DP). Les hôpitaux collectent également des données sur les patients. Les cabinets d'imagerie médicale fournissent également des images numérisées sur la santé des patients. La liste des professionnels de santé collectant des données de santé n'est pas exhaustive. Actuellement, d'autres acteurs proposent du matériel (dispositifs connectés) ou des services (tests génétiques) aux clients-patients qui consciemment ou non enrichissent les bases de données de ces sociétés.

En France, les bases de données médico-administratives utiles à la facturation, au remboursement de soins sont une autre source d'information.

Les citoyens ont des droits sur l'ensemble de ces données collectées. Ils ont un droit d'accès, un droit de rectification, un droit d'effacement, un droit d'opposition [90].

Cependant, le citoyen n'est pas propriétaire légalement de ces données personnelles, il peut en disposer mais non les vendre. De même, l'organisme qui collecte les données personnelles ne peut revendiquer un droit de propriété que lorsque la base de données a été anonymisée et agrégée.

Le *Health Data Hub* a été mis en place pour faciliter l'accès aux données de santé pour les chercheurs, les *startups*. Cette plateforme nationale sécurisée dont l'objectif est d'élargir les données médico-administratives du SNDS (Système National des Données de Santé) aux données cliniques, permettra aux opérateurs publics et privés d'exploiter ses informations. Cet accès sera autorisé selon un cadre réglementaire précis garantissant l'anonymat et le consentement des patients [91] [92].

2.1 Le consentement

Le consentement, dans le contexte médical, est la capacité d'une personne de recevoir des informations et de donner un consentement libre et éclairé à un acte médical qu'on lui propose.

Il y a une obligation de recueillir le consentement de la personne lorsqu'un traitement de données est effectué par un professionnel de santé. Le consentement concerne également la prise en charge du patient.

La délégation de consentement éclairé à un algorithme pourrait affaiblir la réalité du consentement du patient [93]. Deux enjeux éthiques majeurs sont soulignés par la commission CERNA et doivent être régulés : «

- Le risque de priver, en pratique par « délégation de consentement », le patient d'une large partie de sa capacité de participation à la construction de son processus de prise en charge face aux propositions de décisions fournies par un algorithme ;
- Le danger d'une minoration de la prise en compte des situations individuelles dans le cadre d'une systématisation de raisonnements fondés sur des modèles dont des limites peuvent être liées à leur capacité à prendre en compte l'ensemble des caractéristiques et des préférences de chaque patient. » [93].

Lorsqu'une utilisation secondaire des images radiologiques est réalisée, par exemple pour entraîner des modèles d'apprentissage automatique, il est difficile voire impossible de redemander le consentement explicite de chaque patient [94]. Dans l'Union Européenne, le règlement général sur la protection des données permet aux patients de donner un « consentement en ce qui concerne certains domaines de la recherche scientifique, dans le

respect des normes éthiques reconnues en matière de recherche scientifique. Les personnes concernées devraient pouvoir donner leur consentement uniquement pour ce qui est de certains domaines de la recherche ou de certaines parties de projets de recherche, dans la mesure où la finalité visée le permet » [93].

De plus, le RGPD précise que « le traitement des catégories particulières de données à caractère personnel peut être nécessaire pour des motifs d'intérêt public dans les domaines de la santé publique, sans le consentement de la personne concernée. » [93].

2.2 La confidentialité des données

L'anonymisation des données de santé est l'impossibilité d'identifier une personne dans un jeu de données extrait d'une base de données. La désidentification se réfère à la suppression ou à l'obscurcissement d'informations d'identification personnelle. La pseudonymisation consiste à remplacer les identifiants personnels par des identifiants artificiels.

Dans cet article publié dans le journal de l'association canadienne des radiologistes [95], Jaremko JL, Azar M, Bromwich R, et al donnent un exemple d'une reconstruction en 3D du visage à partir d'une image d'IRM. Cette image reconstruite du visage pourrait être utilisée par un algorithme d'IA de reconnaissance faciale et ainsi réidentifier le patient (voir annexe 7). La solution pour anonymiser cette image consiste à lui appliquer un pré-traitement à l'aide d'un algorithme dit de « démaquillage du crâne » ou « *skull stripping* » en anglais. L'actualité récente nous informe que cette anonymisation pourrait n'être que partielle [96]. En effet, d'après ce journal en ligne Usine digitale, le Royaume-Uni revendrait les données de santé de ces concitoyens à des laboratoires pharmaceutiques américains sans les anonymiser. Les données de millions de patients du *National Health System* (NHS), l'équivalent anglais de l'Assurance maladie française, auraient été vendues à des fins de recherche. Certains experts en lien avec le NHS estiment que ces données permettraient d'identifier les patients. Ce mésusage des données serait une pratique courante des laboratoires pharmaceutiques pour sélectionner les malades dont le dossier médical aurait un intérêt élevé. Les autorités de santé anglaises réfutent catégoriquement ces allégations et assurent que les données ne sont commercialisées qu'après un traitement de confidentialité et d'anonymisation ne leurs soient appliquées. L'accès à ces données s'effectue en contractant une licence commerciale et académique. L'organisme en charge de vendre ces licences se nomme le *Clinical Practice Research Datalink* (CPRD). Une estimation de la valeur totale des 55 millions de dossiers médicaux anglais serait de l'ordre de 10 milliards de livres par an, soit environ 11 milliards d'euros. D'autres acteurs, notamment Amazon s'intéressent aux données des patients anglais. Nous assistons à une marchandisation de la santé sans la garantie de préservation de l'anonymat des patients. Est-ce que cette

problématique du non-respect de la confidentialité peut-elle se poser dans le cadre du HDH (Health Data Hub) ? Il est encore tôt pour répondre à cette question, mais un autre problème a surgi concernant cette nouvelle plateforme. En effet, une demande d'ouverture d'une enquête pour « favoritisme » après la décision d'attribuer le marché de l'hébergement des données du HDH à Microsoft, a été formulée par plusieurs sociétés françaises du numérique. Les doléances adressées à Monsieur le Ministre de la santé Olivier Véran concernent les règles de publicité des marchés publics. Les entreprises françaises s'estiment lésées dans cette affaire en argumentant par un texte du code de la commande publique, l'article 3 précisément. Celui-ci définit un « principe d'égalité » garanti par « les principes de liberté d'accès et de transparence des procédures ». Le choix du géant de Redmond s'est fait dans une certaine opacité, car aucun appel d'offre n'a été publié sur le site du Bulletin officiel des annonces des marchés publics (BOAMP), ainsi que sur celui de l'Union des groupements d'achats publics (UGAP). Ce manquement aux règles de la commande s'il est avéré serait pénalement condamnable. La cheffe de mission du HDH, Madame Stéphanie Combes, pour sa défense, argue qu'aucune irrégularité n'a été commise dans l'attribution du marché à la société américaine. Cet article de Médiapart met en lumière les contradictions des politiques qui clament que les données de santé des français sont protégées [97]. Mais en hébergeant ces données par une société américaine, le gouvernement français prend le risque d'une divulgation de celles-ci. En effet, les Etats-Unis ont adopté, le 23 mars 2018, un texte nommé Cloud Act, qui oblige tout fournisseur de service informatique américain à transmettre aux autorités les données qu'il héberge, indifféremment du lieu de stockage. La sécurisation des données est également un enjeu important permettant de préserver la confiance des patients.

2.3 La sécurisation des données

Les systèmes d'information de santé peuvent être la cible d'attaque informatique malveillante. Le 15 novembre 2019, le CHU de Rouen a été victime d'une cyberattaque qui a paralysé son fonctionnement. Ce virus informatique de type rançongiciel a bloqué l'accès aux données de santé des patients. Une somme d'argent est demandée pour restaurer l'accès à ces données [98]. Nous pourrions envisager une autre menace plus sournoise, moins mercantile, qui aurait pour objectif de modifier l'intégrité des données. La conséquence serait d'introduire des biais dans les algorithmes à l'insu des utilisateurs. Prenons l'exemple des images en radiologie, la modification des pixels pourrait modifier le diagnostic. La santé des patients serait en danger. Un autre risque est la perte ou le vol des données. Un autre cas récent de la vulnérabilité des serveurs des systèmes d'archivage d'images PACS a été mise à jour par la société Greenbone [99]. Elle a observé que l'accès par adresse IP (*Internet Protocol*) à ses serveurs n'était pas sécurisé et que ses archives

étaient facilement téléchargeables à distance et sans avoir besoin de connaissances informatiques. Ces fichiers DICOM contiennent des informations médicales et personnelles sensibles de millions de personnes dans le monde. D'après l'analyse de cette société, ces données ne sont pas protégées par un mot de passe ou chiffrées.

Les techniques de sécurisation des données sont la pseudonymisation, l'authentification, la traçabilité, le contrôle, la sensibilisation des administrateurs et également des systèmes d'IA. En effet, dans le domaine de la sécurité des réseaux, l'IA peut permettre de détecter des signaux faibles d'intrusion par sa capacité à analyser des quantités importantes de données.

La Haute Autorité de Santé (HAS) dans le guide sur les spécificités d'évaluation clinique d'un dispositif médical connecté (DMC) en vue de son accès à son remboursement par la CNEDiMTS (La Commission Nationale d'Evaluation des Dispositifs Médicaux et des Technologies de Santé) rappelle que le respect de la sécurité des données est un prérequis à son évaluation [100]. Elle précise que les dispositifs médicaux utilisés exclusivement par un professionnel de santé ou entre professionnels de santé (par exemple, les dispositifs d'aide à la décision professionnelle, les logiciels d'aide à la prescription ou à la dispensation, ceux de télé-expertise, les dispositifs d'imagerie médicale d'aide au diagnostic ou à la décision thérapeutique, etc.) ne sont pas concernés par ce guide. La HAS a également rédigé un rapport d'analyse prospective en 2019 [101] intitulé : « Numérique quelle (R)évolution ? ». Elle y aborde la sécurité des données personnelles en précisant qu'un corpus de textes réglementaires existe à ce sujet et qu'il n'est pas judicieux de le modifier du fait de l'intelligence artificielle ou du *big data*. En effet, la HAS rappelle que le régime de protection des données personnelles date de plus de vingt ans et qu'il est efficace, en particulier le secret professionnel. Elle préconise de mieux informer le citoyen concernant ses droits dans la proposition 22 de ce rapport : « Rendre plus visible le régime de protection des données personnelles, dont les droits des utilisateurs en matière d'effacement ou de portabilité, et faciliter les voies de recours ».

3 Responsabilité juridique en cas de mauvais diagnostic

Dans le cas de la prescription, la délivrance d'un traitement à un patient, les professionnels de santé (médecins, pharmaciens) s'interrogent sur le bénéfice et le risque des médicaments prescrits et délivrés. Il en est de même des recommandations émises par une IA dans le cas d'un diagnostic. D'un point de vue juridique, une machine avec de l'IA ou pas est une chose. La responsabilité incombe toujours à une personne. Cette personne responsable peut être le concepteur du logiciel, le médecin radiologue.

La loi Informatiques et Liberté de 1978 énonce déjà les principes des algorithmes dans l'article 1 :

« L'informatique doit être au service de chaque citoyen. Son développement doit s'opérer dans le cadre de la coopération internationale. Elle ne doit porter atteinte ni à l'identité humaine, ni aux droits de l'homme, ni à la vie privée, ni aux libertés individuelles ou publiques. » [102].

Les différents acteurs travaillant dans le domaine de l'IA peuvent avoir un rôle de concepteur, d'entraîneur, ou d'utilisateur. Chacun de ces rôles, a une responsabilité dans la chaîne algorithmique d'un outil d'aide au diagnostic.

Le rapport du Conseil national de l'ordre des médecins [103] de janvier 2018, « Médecins et patients dans le monde des data, des algorithmes et de l'intelligence artificielle », contient des recommandations. Citons la recommandation 22 : « Il (Le CNOM) engage à examiner le régime juridique des responsabilités : celle du médecin dans ses usages de l'aide à la décision et celle des concepteurs des algorithmes quant à la fiabilité des data utilisées et les modalités de leur traitement informatisé. »

Si d'après les performances des logiciels d'analyse d'imageries médicales, l'IA dépiste mieux certains cancers que le médecin, il serait tentant de ne se fier qu'à la décision de la machine. Nous pouvons nous interroger sur la perte de compétences du radiologue que cette situation de délégation pourrait induire. A contrario, si le médecin ne tient pas compte de l'avis de l'IA et qu'il se trompe dans le diagnostic, est-ce que sa responsabilité est majorée par cette attitude ?

Les discussions sur la responsabilité de la décision prise par une IA restent ouvertes et sont largement débattues par la communauté de chercheurs et des professionnels de santé. Certaines personnes se demandent s'il ne faut pas accorder un statut juridique autonome aux algorithmes.

Il est utile de rappeler qu'en France, le médecin reste le dernier décisionnaire.

L'objectif de la réglementation est d'éviter de bloquer l'innovation par un cadre légal trop rigide en France ou en Europe. Le risque ce serait de devoir importer des solutions développées dans un autre pays où les considérations éthiques ne sont pas prises en compte. L'enjeu est de trouver un équilibre entre la réglementation et la liberté pour les acteurs industriels d'entreprendre. Le CNOM propose la création d'une instance de régulation [103].

4 Explicabilité des algorithmes

L'explicabilité peut être définie par la capacité à expliquer le fonctionnement technique d'une solution d'IA et les décisions humaines qui ont permis son développement [104]. Le cheminement et la traçabilité des décisions prises par un système d'IA devraient être possibles dans une documentation.

Le comité consultatif national d'éthique pour les sciences de la vie et de la santé a émis un avis rendu public le 29 mai 2019, intitulé « Données massives et santé : une nouvelle approche des enjeux éthiques » [105]. Le CCNE émet le souhait qu'une recherche scientifique de haut niveau en informatique et mathématiques fondamentales puisse expliquer les différentes étapes d'un algorithme d'apprentissage automatique et permettre ainsi une généralisation des applications dans le domaine de la santé. En effet, bien que la méthode d'apprentissage, le choix des données d'entrée ont été déterminés par le concepteur du logiciel d'IA, les raisons qui expliquent comment et pourquoi la matrice de plusieurs millions de paramètres arrive à effectuer la tâche de reconnaissance d'image attendue demeurent obscures. On parle de phénomène de « boîte noire » [106]. Ce peut être un obstacle lorsqu'il s'agit de déterminer les responsabilités en cas d'erreur, par exemple dans le diagnostic médical. Une exigence de transparence est demandée au système d'IA. De nombreux chercheurs travaillent sur l'interprétabilité des réseaux de neurones et notamment l'agence de recherche de l'armée américaine DARPA (*Defense Advanced Research Project Agency*) [107]. Ils conçoivent des outils qui permettront de mieux comprendre la manière dont les réseaux de neurones prennent une décision. De la Torre J, Valls A, Puig D, et al ont développé un modèle d'apprentissage profond pertinent de propagation par couche qui évalue l'importance de chaque pixel d'entrée de l'image dans la décision finale de la classification dans le diagnostic de la rétinopathie diabétique. Ils ont utilisé une méthode statistique, l'analyse en composantes indépendantes en combinaison avec un score de visualisation technique [108].

5 Acceptabilité des algorithmes

Le fait divers lié à la voiture autonome de l'entreprise Uber qui a percuté un cycliste ayant traversé la chaussée en dehors du passage piéton, pose la question de l'acceptabilité de l'IA. La confiance des professionnels de santé donnée à ces applications d'IA dépendra de la preuve de leur efficacité. Les nouveaux outils seront adoptés si le risque d'erreur est réduit, le processus de prise en charge est rapide, les pratiques des médecins améliorées [109]. Ces outils devront être intuitifs et simples à l'utilisation.

6 Impact sur l'emploi

Un rapport destiné à la ministre du Travail et au secrétaire d'État auprès du Premier ministre, chargé du Numérique présente l'impact que pourrait avoir l'IA sur le travail [109]. Trois secteurs d'activité sont analysés : les secteurs des transports, bancaire, et de la santé.

Les professionnels de santé sont concernés par la généralisation prévisible de la lecture d'images automatisée. Sur des tâches bien précises, l'IA peut aider au diagnostic effectué par des médecins non-radiologues, comme pour la rétinopathie diabétique par le logiciel IDx-DR. Une adaptation du cadre réglementaire actuel sera nécessaire pour mettre en place ce genre d'outil en France [110].

La dermatologie, la radiologie sont donc directement concernées par le changement de l'organisation de leurs métiers induite par l'IA. Le rôle du radiologue dans la prise en charge du patient ne se limite pas à l'acquisition d'une donnée d'imagerie et à l'interprétation de celle-ci. Le médecin-radiologue communique les résultats au médecin généraliste. Il est le garant de la qualité, de l'amélioration de la qualité dans l'exercice de son métier. Il effectue beaucoup d'autres tâches qui ne peuvent être automatisées. L'IA ne se substituera pas au professionnel de santé dans la prise de décision. Cependant, des mutations importantes vont modifier les métiers, les rôles, les fonctions, les responsabilités des différents acteurs de la santé.

L'IA questionne également la relation médecin-patient. Le colloque singulier entre un soignant et un soigné, entre un « sachant » et un « ignorant » informé par internet sur sa pathologie, fait intervenir une tierce partie, le logiciel. Celui-ci détient un savoir médical, et il est meilleur expert dans certaines tâches qu'un professionnel de santé. Si l'IA a la capacité d'augmenter la connaissance du médecin, elle pourra aussi le faire pour le patient. La valeur ajoutée du médecin est son sens critique, par exemple, dans l'interprétation des images radiologiques réalisées par l'IA.

Dans l'industrie pharmaceutique, la R&D connaît une vague de suppression d'emploi due à une externalisation, au recours à la sous-traitance, aux transferts d'activités. Ainsi, le laboratoire français Sanofi a supprimé 299 postes en France et 167 en Allemagne sur la base de départs volontaires et de retraites anticipées [111]. Dans ce contexte, l'IA accentuerait cette tendance à la restructuration, mais ne serait pas le facteur principal [112].

La formation initiale des étudiants devrait anticiper ces bouleversements dans la pratique de leur future profession. Ainsi, après l'université Paris Descartes, la faculté de médecine de Lille propose un diplôme universitaire (DU) d'IA dont les cours ont commencé le 14 novembre 2019 [113]. La formation continue des professionnels en activité intégrera également ces changements. Les compétences nécessaires seront transversales et pluridisciplinaires.

La création de nouveaux métiers (*data manager*, *data scientist*) avec le développement de l'IA est une réalité actuellement. L'enjeu est de former et de retenir les talents.

Le management de l'IA, à l'instar du management des organisations et des personnes, devrait être pris en compte par les décideurs des entreprises. La gestion des relations entre l'Humain avec des systèmes informatiques dont l'autonomie décisionnaire s'améliore, sera une problématique pour les sociétés [112]. Les risques psychosociaux induits par le

déploiement de l'IA devraient être anticipés, notamment la perte d'autonomie du salarié, l'augmentation de la charge de travail. En effet, si les cas simples et routiniers sont traités par l'IA, les cas complexes dévolus au radiologue seraient majoritaires et nécessiteront plus de concentration. A l'inverse, si des tâches compliquées effectuées par l'humain sont automatisables par un système d'IA, une perte de compétences pourrait survenir et entraîner une déqualification de la fonction [109]. Le transfert d'activités (lecture d'imagerie médicale) vers d'autres professionnels de santé (infirmier, manipulateur-radio) augmentera leur qualification avec l'aide au diagnostic des dispositifs médicaux intégrant de l'IA. Ce rôle de superviseur de dispositif demandera des compétences numériques. D'autre part, les compétences relationnelles, l'empathie, seront approfondies chez les professionnels de santé. Le temps libéré par l'utilisation de l'IA permettra une meilleure prise en charge des patients.

Quels sont les répercussions de l'IA sur l'emploi ? Actuellement, il est difficile d'être catégorique sur l'une ou l'autre des hypothèses émises par les différents rapports d'analyse : création ou remplacement de postes. Une seule certitude c'est que le changement dans les métiers est en cours et qu'il faut l'accompagner.

Un autre point d'ordre stratégique est à prendre en compte dans le développement de l'IA : le risque de concentration par les entreprises du numérique américain, les GAFAMI (Google, Amazon, Facebook, Apple, Microsoft et IBM) et leurs équivalents chinois, les BATX (Baidu, Alibaba, Tencent, Xiaomi) de la valeur économique qui les placeraient en situation de monopole dans leurs marchés respectifs. Cette bipolarité entre les Etats-Unis d'Amérique et la Chine est déjà présente dans l'économie mondiale. Ces entreprises ont une capacité de recruter les meilleurs chercheurs pour développer leurs outils et elles ont une telle avance technologique qu'il sera difficile de les concurrencer. La perte de souveraineté des pays pourrait avoir des répercussions sociales et économiques [112]. En effet, les pays qui investissent le plus sont ces deux locomotives de l'économie mondiale, les Etats-Unis et la Chine. Nous pouvons nous interroger sur la place de l'Europe et de la France en particulier dans cette course à l'innovation technologique. La France a mis en place des programmes d'investissement, mais à une échelle plus modeste par rapport aux sommes que les américains et les chinois ont décidé de mettre sur la table. Cependant, les atouts français ne manquent pas et l'hexagone peut s'appuyer sur la qualité de ses chercheurs, sur la disponibilité de données de santé de qualité. Il faut rester prudent et ne pas recommencer les mêmes schémas de développement économique qui ont permis une délocalisation partielle de la production chimique vers des pays où la main d'œuvre est à bas coût et les normes environnementales sont moins strictes. En effet, l'industrie pharmaceutique européenne importe majoritairement les matières premières ou les médicaments de base (exemples de l'ibuprofène, l'aspirine, l'héparine) de Chine ou d'Inde. Cela pose régulièrement des problèmes d'approvisionnement et engendre des ruptures au niveau de la

chaîne du médicament. Cette dépendance forte peut-elle se reproduire dans la technologie et dans le domaine de l'IA en particulier ? Quelles seront les conséquences de cette dépendance ? Ces questions méritent d'être posées aux décideurs des grands groupes industriels pharmaceutiques et aux gouvernants du pays.

Conclusion et perspectives

L'intelligence artificielle est une science qui n'est pas récente. Actuellement, les technologies d'apprentissage machine, en particulier les réseaux de neurones profonds à couches multiples sont les plus déployées dans les cas d'usage notamment dans l'analyse d'images. Cependant, la quantité d'exemples nécessaires à l'entraînement des algorithmes d'apprentissage est un facteur limitant. Un être humain n'a besoin de quelques images de chat pour être capable de généraliser ce concept, alors qu'il faut des centaines de milliers à un algorithme de reconnaissance d'images pour effectuer la même tâche. L'avenir de l'IA passera par l'utilisation des algorithmes non supervisés pour réduire le nombre de données étiquetées qui ont un coût non négligeable. Dans la découverte de nouvelles molécules, les solutions d'IA se développent et les entreprises pharmaceutiques les ont intégrées à leurs outils de R&D. L'IA peut être une aide à chaque étape du cycle de développement du médicament, comme d'autres techniques ou méthodes l'ont été ou le sont encore. L'intérêt de l'IA est de détecter les molécules prometteuses et d'écarter le plus tôt possible celles qui vont échouer aux stades avancés du développement. Finalement, une diminution des coûts de développement est réalisée, car les essais cliniques représentent une part importante des dépenses de R&D.

Actuellement, le recul nécessaire n'est pas suffisant pour affirmer que l'IA a révolutionné le champ de la découverte de nouvelles entités chimiques. En effet, il y a très peu d'informations qui nous renseignent sur le nombre de candidats-médicaments proposés par une solution d'IA ayant obtenu une autorisation de mise sur le marché. A l'heure actuelle, aucun candidat-médicament n'est entré en phase III (essais chez l'Homme), qui est l'indicateur de référence pour les médicaments expérimentaux.

Par ailleurs, l'encadrement juridique de l'utilisation des données et de l'IA doit tenir compte de l'équilibre entre le soutien à l'innovation et la protection des droits des patients. Les instances réglementaires veilleront à ce que ce progrès technique bénéficie à tout le monde, et sans nuire à personne. Les débats que suscite l'utilisation de l'IA dans les activités professionnelles interrogent sur la capacité à créer ou à détruire des emplois grâce à cette technologie. Nous pouvons faire un parallèle avec la robotisation dans l'industrie et ses conséquences sur les effectifs d'ouvriers. Les pays les plus fortement robotisés comme l'Allemagne ont un taux de chômage bas. D'autres facteurs interviennent dans le marché de l'emploi.

Enfin, la question éthique soulevée par l'utilisation de l'IA n'a pas qu'une seule réponse, et dépend de la culture du pays ou de la région du monde. En tant que professionnels de santé, nous devons garantir la qualité, l'efficacité, et la sécurité du médicament, ainsi que des dispositifs médicaux.

Les perspectives de l'IA dans le domaine de la santé seront par exemple la mise en place d'un jumeau numérique du patient, où toutes les caractéristiques d'une personne comme son génome, ses antécédents médicaux, seront digitalisées. Ce clone numérique permettra de personnaliser un traitement en évaluant s'il est répondant ou non à cette molécule en fonction de ses données génétiques. Les essais cliniques seront plus efficaces avec l'aide de l'IA qui permettra d'identifier des groupes de volontaires avec des caractéristiques particulières. De plus, la pharmacovigilance pourrait également bénéficier de l'apport de l'IA dans la détection des signaux faibles. Actuellement, les solutions logicielles à base d'IA sont spécifiques à une tâche, nous pouvons imaginer une solution généraliste qui intégrerait tout le cycle de développement d'un médicament. Par ailleurs, l'industrie pharmaceutique numérise de plus en plus ses outils de production et collecte des quantités importantes de données à l'aide de capteurs connectés. L'IA a ou aura un rôle important dans cette industrie 4.0 en interaction avec d'autres technologies telles que la robotique, la 5G, l'internet des objets, la *blockchain*. Quel sera le rôle du pharmacien dans cette industrie du futur ?

N.B. : Je n'ai aucuns conflits d'intérêts, les entreprises citées dans ce document ont valeur d'exemples.

Bibliographie

1. Goodfellow I, Bengio Y, Courville A. Deep learning book. [En ligne]. Site disponible sur : <https://www.deeplearningbook.org> (page consultée le 13 septembre 2019).
2. Russel S, Norvig P. Artificial Intelligence A modern Approach. Third Edition. Pearson Education France. 2010.
3. Julia L. L'intelligence artificielle n'existe pas. FIRST Editions. Edi8. 2019.
4. Le Cun Y. Quand la machine apprend. Odile Jacob. 2019.
5. Alliot J-M, Schiex T, Brisset P, Garcia F. Intelligence artificielle & informatique théorique. Cépadues. 2002.
6. Biernat E, Lutz M. Data science : fondamentaux et études de cas Machine learning avec Python et R. 2015.
7. Larousse É. Définitions : tâche - Dictionnaire de français Larousse. [En ligne]. Site disponible sur : <https://www.larousse.fr/dictionnaires/francais/t%C3%A2che/76339> (page consultée le 19 septembre 2019).
8. Azencott C-A. Introduction au Machine Learning. Dunod. 2018.
9. Gilloux M. Traitement automatique des langues naturelles. Ann Télécommunications. 1989 ; **44**(5) : 301-16.
10. Mikolov T, Chen K, Corrado G, Dean J. Efficient Estimation of Word Representations in Vector Space", 2013. arXiv : 1301.3781.
11. Lee K, Salant S, Kwiatkowski T, Parikh A, Das D, Berant J. Learning Recurrent Span Representations for Extractive Question Answering, arXiv : 1611.01436.
12. Ehrmann M. Les Entités Nommées, de la linguistique au TAL : Statut théorique et méthodes de désambiguïsation. [En ligne]. Site disponible sur : <https://hal.archives-ouvertes.fr/tel-01639190/document> (page consultée le 02 mars 2020).
13. Remedios SW, Wu Z, Bermudez C, Kerley CI, Roy S, Patel MB, et al. Extracting 2D weak labels from volume labels using multiple instance learning in CT hemorrhage detection.
14. Mnih V, Kavukcuoglu K, Silver D, Rusu AA, Veness J, Bellemare MG, et al. Human-level control through deep reinforcement learning. *Nature*. 2015 ;**518**(7540) : 529-533.
15. Pan SJ, Yang Q. A Survey on Transfer Learning. *IEEE Trans Knowl Data Eng*. 2010 ; **22**(10) : 1345-1359.
16. OpenClassrooms. Mettez en place un cadre de validation croisée. [En ligne]. Site disponible sur : <https://openclassrooms.com/fr/courses/4297211-evaluez-et-ameliorez-les-performances-dun-modele-de-machine-learning/4308241-mettez-en-place-un-cadre-de-validation-croisee> (page consultée le 20 septembre 2019).
17. Mater AC, Coote ML. Deep Learning in Chemistry. *J Chem Inf Model*. 2019 ; **59**(6) : 2545-2559.

18. Wu Z, Ramsundar B, Feinberg EN, Gomes J, Geniesse C, Pappu AS, et al. MoleculeNet : A Benchmark for Molecular Machine Learning. ArXiv170300564 Phys Stat. 2018.
19. Blum LC, Reymond J-L. 970 Million Druglike Small Molecules for Virtual Screening in the Chemical Universe Database GDB-13. *J Am Chem Soc.* 2009 ; **131**(25) : 8732-8733.
20. Google. TensorFlow. [En ligne]. Site disponible sur : <https://www.tensorflow.org/> (page consultée le 16 décembre 2019).
21. Facebook. PyTorch. [En ligne]. Site disponible sur : <https://www.pytorch.org> (page consultée le 16 décembre 2019).
22. Landry Y, Gies J-P, Sick E, Niederhoffer N. Pharmacologie des cibles à la thérapeutique. Dunod. 2019.
23. Scatton B. Le processus de découverte du médicament dans l'industrie pharmaceutique. *J Société Biol.* 2009 ; **203**(3) : 249-269.
24. Lipinski CA, Lombardo F, Dominy BW, Feeney PJ. Experimental and computational approaches to estimate solubility and permeability in drug discovery and development settings. *Adv Drug Deliv Rev.* 2001 ; **46** : 3-26.
25. Bohuon C, Monneret C. Fabuleux hasards Histoire de la découverte de médicaments. EDP Sciences. 2009.
26. Bach J-F, Blanchard-Desce M, Couvreur P, Dardel F, Giovannangeli C, Maffrand J-P, et al. La chimie et la santé au service de l'homme. EDP Sciences. 2010.
27. Donald J-A. Burger's Medicinal Chemistry and Drug Discovery ; Volume 2 Drug Discovery and Drug Development (6th Ed.). Wiley-Interscience. 2003.
28. Morrisette DA, Parachikova A, Green KN, LaFerla FM. Relevance of Transgenic Mouse Models to Human Alzheimer Disease. *J Biol Chem.* 2009 ; **284**(10) : 6033-6037.
29. Biogen. Biogen Plans Regulatory Filing for Aducanumab in Alzheimer's Disease Based on New Analysis of Larger Dataset from Phase 3 Studies. [En ligne]. Site disponible sur : <http://investors.biogen.com/news-releases/news-release-details/biogen-plans-regulatory-filing-aducanumab-alzheimers-disease> (page consultée le 10 décembre 2019).
30. Schneider L. A resurrection of aducanumab for Alzheimer's disease. *Lancet Neurol.* 2020 ; **19**(2) : 111-112.
31. Wermuth CG. Selective Optimization of Side Activities : Another Way for Drug Discovery. *J Med Chem.* 2004 ; **47**(6) : 1303-1314.
32. Ramsundar B, Eastman P, Walters P, Pande V. Deep learning for the Life Sciences. Oreilly. 2019.
33. Daylight. [En ligne]. Site disponible sur : <https://www.daylight.com/smiles/index.html> (page consultée le 4 octobre 2019).
34. Wikipédia. Simplified Molecular Input Line Entry Specification. [En ligne]. Site disponible sur :

https://fr.wikipedia.org/w/index.php?title=Simplified_Molecular_Input_Line_Entry_Specification&oldid=161919412 (page consultée le 4 octobre 2019).

35. Rogers D, Hahn M. Extended-Connectivity Fingerprints. *J Chem Inf Model*. 2010 ; **50**(5) : 742-754
36. Preuer K, Klambauer G, Rippmann F, Hochreiter S, Unterthiner T. Interpretable Deep Learning in Drug Discovery. ArXiv190302788 Cs Q-Bio Stat. 2019.
37. Wermuth CG, Ganellin CR, Lindberg P, Mitscher LA. Glossary of terms used in medicinal chemistry (IUPAC Recommendations 1998). *Pure Appl Chem*. 1998 ; **70**(5) : 1129-1143.
38. RDKit. [En ligne]. Site disponible sur : <https://www.rdkit.org/> (page consultée le 9 décembre 2019).
39. OpenClassrooms. Initiez-vous au *deep learning*. [En ligne]. Site disponible sur : <https://openclassrooms.com/fr/courses/5801891-initiez-vous-au-deep-learning/5801898-decouvrez-le-neurone-formel> (page consultée le 20 septembre 2019).
40. ResearchGate. DeepTox : Toxicity prediction using deep learning. [En ligne]. Site disponible sur : https://www.researchgate.net/publication/320828461_DeepTox_Toxicity_prediction_using_deep_learning (page consultée le 7 octobre 2019).
41. Altae-Tran H, Ramsundar B, Pappu AS, Pande V. Low Data Drug Discovery with One-Shot Learning. *ACS Cent Sci*. 2017 ; **3**(4) : 283-293.
42. Chen H, Engkvist O, Wang Y, Olivecrona M, Blaschke T. The rise of deep learning in drug discovery. *Drug Discov Today*. 2018 ; **23**(6) : 1241-1250.
43. Ramsundar B, Liu B, Wu Z, Verras A, Tudor M, Sheridan RP, et al. Is Multitask Deep Learning Practical for Pharma ? *J Chem Inf Model*. 2017 ; **57**(8) : 2068-2076.
44. Yoshikawa N, Terayama K, Honma T, Oono K, Tsuda K. Population-based de novo molecule generation, using grammatical evolution. *Chem Lett*. 2018 ; **47**(11) : 1431-1434.
45. Senior A W, Evans R, Jumper J et al. Improved protein structure prediction using potentials from deep learning. *Nature*. 2020 ; **577** : 706–710.
46. Wallach I, Dzamba M, Heifets A. AtomNet : A Deep Convolutional Neural Network for Bioactivity Prediction in Structure-based Drug Discovery. ArXiv151002855 Cs Q-Bio Stat. 2015.
47. Atomwise. New Ebola Treatment Using Artificial Intelligence. [En ligne]. Site disponible sur : <https://www.atomwise.com/2015/03/24/new-ebola-treatment-using-artificial-intelligence/> (page consultée le 22 novembre 2019).
48. BlueDot. [En ligne]. Site disponible sur : <https://bluedot.global/> (page consultée le 07 mars 2020).
49. Wang L-F, Cowled C, Bats and viruses A new Frontier of Emerging Infectious Diseases, Wiley Blackwell, 2015.

50. Kaggle. [En ligne]. Site disponible sur : <https://www.kaggle.com/allen-institute-for-ai/CORD-19-research-challenge> (page consultée le 06 avril 2020).
51. Richardson P, Griffin I, Tucker C, et al. Baricitinib as potential treatment for 2019-nCoV acute respiratory disease, *Lancet*. 2020 ; **395** : 497-506.
52. Zhavoronkov A, Aladinskiy V, Zhebrak A, et al. Potential 2019-nCoV 3C-like protease inhibitors designed using generative deep learning approaches. [En ligne]. Site disponible sur : <https://insilico.com/ncov-sprint/> (page consultée le 10 mars 2020).
53. Jumper J, Tunyasuvunakool K, Kohli P, et al. Computational predictions of protein structures associated with COVID-19. [En ligne]. Site disponible sur : <https://deepmind.com/research/open-source/computational-predictions-of-protein-structures-associated-with-COVID-19> (page consultée le 20 mars 2020).
54. Inserm. Lancement d'un essai clinique européen contre le Covid-19. [En ligne]. Site disponible sur : <https://presse.inserm.fr/lancement-dun-essai-clinique-europeen-contre-le-covid-19/38737/> (page consultée le 25 mars 2020).
55. Stokes J M, Yang K, Swanson K, et al, A Deep Learning approach to antibiotic Discovery, *Cell*. **180** : 688-702.
56. Corsello S M, Bittker J A, Liu Z, Gould J, McCarren P, Hirschman J E, Johnston S E, Vrcic A, Wong B, Khan M, et al, The Drug Repurposing Hub : a next-generation drug library and information resource, *Nat. Med.* 2017 ; **23** : 405-408.
57. De S K, Stebbins J L, Chen L H, Riel-Mehan M, Machleidt T, Dahl R, Yuan H, Emdadi A, Barile E, Chen V, et al. Design, synthesis, and structure-activity relationship of substrate competitive, selective, and in vivo active triazole and thiadiazole inhibitors of the c-Jun N-terminal kinase, *J. Med. Chem.* 2009 ; **52** : 1943–1952.
58. Irwin JJ, Shoichet BK. ZINC—a free database of commercially available compounds for virtual screening. *J Chem Inf Model*. 2005 ; **45**(1) : 177–182
59. Jang S, Yu LR, Abdelmegeed MA, Gao Y, Banerjee A, et Song B J. Critical role of c-jun N-terminal protein kinase in promoting mitochondrial dysfunction and acute liver injury. *Redox Biol*. 2015 ; **6** : 552–564.
60. Exscientia. Sumitomo Dainippon Pharma and Exscientia Joint Development New Drug Candidate Created Using Artificial Intelligence (AI) Begins Clinical Trial. [En ligne]. Site disponible sur : <https://www.exscientia.ai/news-insights/sumitomo-dainippon-pharma-and-exscientia-joint-development> (page consultée le 07/03/2020).
61. Zemouri R, Devalland C, Valmary-Degano S, Zerhouni N. Intelligence artificielle : quel avenir en anatomie pathologique ? *Ann Pathol*. 2019 ; **39**(2) : 119-129.
62. Bibault J-E, Burgun A, Giraud P. Intelligence artificielle appliquée à la radiothérapie. *Cancer/Radiothérapie*. 2017 ; **21**(3) : 239-243.
63. Fernandez-Maloigne C, Guillevin R. L'intelligence artificielle au service de l'imagerie et de la santé des femmes. *Imag Femme*. 2019 ; **29**(4) : 179-186.
64. Heeke S, Delingette H, Fanjat Y, Long-Mira E, Lassalle S, Hofman V, et al. La pathologie cancéreuse pulmonaire à l'heure de l'intelligence artificielle : entre espoir, désespoir et perspectives. *Ann Pathol*. 2019 ; **39**(2) : 130-136.

65. Collège des enseignants de radiologie de France, Collège national des enseignants de biophysique et de médecine nucléaire. *Imagerie Médicale Radiologie et Médecine Nucléaire*. Elsevier Masson. 2015.
66. Litjens G, Kooi T, Bejnordi BE, Setio AAA, Ciompi F, Ghafoorian M, et al. A survey on deep learning in medical image analysis. *Med Image Anal*. 2017 ; **42** : 60-88.
67. Zhou SK, Greenspan H, Shen D. *Deep Learning for Medical Image Analysis*. Elsevier ; 2017.
68. Chartrand G, Cheng PM, Vorontsov E, Drozdal M, Turcotte S, Pal CJ, et al. Deep Learning : A Primer for Radiologists. *RadioGraphics*. 2017 ; **37**(7) : 2113-2131.
69. FDA (Food and Drug Administration). FDA permits marketing of artificial intelligence-based device to detect certain diabetes-related eye problems. [En ligne]. Site disponible sur : <http://www.fda.gov/news-events/press-announcements/fda-permits-marketing-artificial-intelligence-based-device-detect-certain-diabetes-related-eye> (page consultée le 28 octobre 2019).
70. Abràmoff MD, Lou Y, Erginay A, Clarida W, Amelon R, Folk JC, et al. Improved Automated Detection of Diabetic Retinopathy on a Publicly Available Dataset Through Integration of Deep Learning. *Investig Ophthalmology Vis Sci*. 2016 ; **57**(13) : 5200-5206.
71. IDx Technologies Inc. [En ligne]. Site disponible sur : <https://www.eyediagnosis.co/> (page consultée le 28 octobre 2019).
72. FDA (Food and Drug Administration). FDA Clears ScreenPoint Medical's AI System for Reading Mammograms. [En ligne]. Site disponible sur : <https://www.fdanews.com/articles/189314-fda-clears-screenpoint-medicals-ai-system-for-reading-mammograms> (page consultée le 6 novembre 2019).
73. ScreenPoint Medical BV. Transpara. [En ligne]. Site disponible sur : <https://www.screenpoint-medical.com/transpara> (page consultée le 6 novembre 2019).
74. Mordang J-J, Janssen T, Bria A, Kooi T, Gubern-Mérida A, Karssemeijer N. Automatic Microcalcification Detection in Multi-vendor Mammography Using Convolutional Neural Networks. *Springer International Publishing Switzerland*. 2016 ; **9699** : 35-42.
75. Rodríguez-Ruiz A, Krupinski E, Mordang J-J, Schilling K, Heywang-Köbrunner SH, Sechopoulos I, et al. Detection of Breast Cancer with Mammography : Effect of an Artificial Intelligence Support System. *Radiology*. 2019 ; **290**(2) : 305-314.
76. FDA (Food and Drug Administration). FDA permits marketing of clinical decision support software for alerting providers of a potential stroke in patients. [En ligne]. Site disponible sur : <http://www.fda.gov/news-events/press-announcements/fda-permits-marketing-clinical-decision-support-software-alerting-providers-potential-stroke> (page consultée le 12 novembre 2019).
77. Viz.ai, Inc. [En ligne]. Site disponible sur : <https://www.viz.ai> (page consultée le 12 novembre 2019).
78. Canon Medical Systems USA. Aquilion Precision CT Scanner. [En ligne]. Site disponible sur : <https://us.medical.canon/products/computed-tomography/aquilion-precision> (page consultée le 5 décembre 2019).
79. Deslandes M, Chave L, Pommier M, Detraz J, Nord B, Panassie L, et al. État de l'art en imagerie médicale. *IRBM News*. 2019 ; **40**(2) : 45-61.

80. Vamathevan J, Clark D, Czodrowski P, et al. Applications of machine learning in drug discovery and development. *Nature Reviews*. 2019 ; **18** : 463-477.
81. Pesapane F, Volonté C, Codari M, Sardanelli F. Artificial intelligence as a medical device in radiology : ethical and regulatory issues in Europe and the United States. *Insights Imaging*. 2018 ; **9**(5) : 745-53.
82. CCNE (Comité Consultatif National d'Éthique pour les sciences de la vie et de la santé). Création du comité pilote d'éthique du numérique. [En ligne]. Site disponible sur : https://www.ccne-ethique.fr/sites/default/files/communiquelancement_comite_numerique.pdf (page consultée le 3 décembre 2019).
83. CNIL (Commission Nationale Informatique et Libertés). COMMENT PERMETTRE À L'HOMME DE GARDER LA MAIN ? Les enjeux éthiques des algorithmes et de l'intelligence artificielle. [En ligne]. Site disponible sur : https://www.cnil.fr/sites/default/files/atoms/files/cnil_rapport_garder_la_main_web.pdf (page consultée le 3 novembre 2019).
84. CERNA (Commission de réflexion sur l'Éthique de la Recherche en sciences et technologies du Numérique d'Allistene). Éthique de la recherche en apprentissage machine. [En ligne]. Site disponible sur : http://cerna-ethics-allistene.org/digitalAssets/53/53991_cerna___thique_apprentissage.pdf (page consultée le 3 novembre 2019).
85. Chatila R, Havens J C. The IEEE Global Initiative on Ethics of Autonomous. [En ligne]. Site disponible sur : <https://standards.ieee.org/industry-connections/ec/autonomous-systems.html> (page consultée le 4 novembre 2019).
86. Stanford University. Artificial intelligence and life in 2030. [En ligne]. Site disponible sur : https://ai100.stanford.edu/sites/g/files/sbiybj9861/f/ai_100_report_0901fnlc_single.pdf (page consultée le 4 novembre 2019).
87. Smith J. Le Contrôleur Européen de la Protection des Données - European Data Protection Supervisor. Éthique. [En ligne]. Site disponible sur : https://edps.europa.eu/data-protection/our-work/ethics_fr (page consultée le 4 novembre 2019).
88. CEPD (Contrôleur européen de la protection des données). Rapport annuel 2018. [En ligne]. Site disponible sur : https://edps.europa.eu/sites/edp/files/publication/ar2018_executive_summary_fr.pdf (page consultée le 4 novembre 2019).
89. Le Monde.fr. Le comité d'éthique de Google sur l'intelligence artificielle n'aura existé qu'une semaine. [En ligne]. Site disponible sur : https://www.lemonde.fr/pixels/article/2019/04/05/intelligence-artificielle-google-renonce-a-son-comite-d-ethique-une-semaine-apres-son-lancement_5446456_4408996.html (page consultée le 4 novembre 2019).
90. LOI n° 2018-493 du 20 juin 2018 relative à la protection des données personnelles. JO du 21 juin 2018.
91. Health Data Hub. [En ligne]. Site disponible sur : <https://www.health-data-hub.fr> (page consultée le 7 novembre 2019).
92. INDS (Institut National des Données de Santé). Impact de la loi relative à l'organisation et à la transformation du système de santé sur les données de santé. [En ligne]. Site

disponible sur : <https://www.indisante.fr/node/1161> (page consultée le 7 novembre 2019).

93. CCNE (Comité Consultatif National d'Éthique pour les sciences de la vie et de la santé). NUMÉRIQUE & SANTÉ QUELS ENJEUX ÉTHIQUES POUR QUELLES RÉGULATIONS ?. [En ligne]. Site disponible sur : https://www.ccne-ethique.fr/sites/default/files/publications/rapport_numerique_et_sante_19112018.pdf (page consultée le 3 novembre 2019).
94. CNIL (Commission Nationale Informatique et Libertés). Le règlement général sur la protection des données - RGPD. [En ligne]. Site disponible sur : <https://www.cnil.fr/fr/reglement-europeen-protection-donnees> (page consultée le 14 novembre 2019).
95. Jaremko JL, Azar M, Bromwich R, Lum A, Alicia Cheong LH, Gibert M, et al. Canadian Association of Radiologists White Paper on Ethical and Legal Issues Related to Artificial Intelligence in Radiology. *Can Assoc Radiol J*. 2019 ; **70**(2) : 107-18.
96. Usine digitale. Le Royaume-Uni revend les données de patients à des laboratoires américains sans les anonymiser. [En ligne]. Site disponible sur : <https://www.usine-digitale.fr/article/le-royaume-uni-revend-les-donnees-de-patients-a-des-laboratoires-americains-sans-les-anonymiser> (page consultée le 12 mars 2020).
97. Mediapart. Données de santé : l'Etat accusé de favoritisme au profit de Microsoft. [En ligne]. Site disponible sur : <https://www.mediapart.fr/journal/france/110320/donnees-de-sante-l-etat-accuse-de-favoritisme-au-profit-de-microsoft> (page consultée le 12 mars 2020).
98. Le Monde.fr. Le CHU de Rouen perturbé par une attaque informatique. [En ligne]. Site disponible sur : https://www.lemonde.fr/sante/article/2019/11/17/le-chu-de-rouen-perturbe-par-une-attaque-informatique_6019491_1651302.html (page consultée le 4 décembre 2019).
99. Greenbone. *Confidential patient data freely accessible on the internet*. [En ligne]. Site disponible sur : https://www.greenbone.net/wp-content/uploads/Confidential-patient-data-freely-accessible-on-the-internet_20190918.pdf (page consultée le 16 janvier 2020).
100. HAS (Haute Autorité de Santé). Guide sur les spécificités d'évaluation clinique d'un dispositif médical connecté (DMC) en vue de son accès au remboursement. [En ligne]. Site disponible sur : https://www.has-sante.fr/upload/docs/application/pdf/2019-02/guide_sur_les_specificites_devaluation_clinique_dun_dmc_en_vue_de_son_acces_au_remboursement.pdf (page consultée le 3 novembre 2019).
101. HAS (Haute Autorité de Santé). Rapport d'analyse prospective 2019 Numérique : quelle (R)évolution?. [En ligne]. Site disponible sur : https://www.has-sante.fr/upload/docs/application/pdf/2019-07/rapport_analyse_prospective_20191.pdf (page consultée le 7 novembre 2019).
102. Loi n° 78-17 du 6 janvier 1978 relative à l'informatique, aux fichiers et aux libertés. JO du 7 janvier 1978. p.227.
103. CNOM (Conseil National de l'Ordre des Médecins) Médecins et patients dans le monde des data, des algorithmes et de l'intelligence artificielle. [En ligne]. Site disponible sur : https://www.conseil-national.medecin.fr/sites/default/files/external-package/edition/od6gnt/cnomdata_algorithmes_ia_0.pdf (page consultée le 26 septembre 2019).

104. European Commission. Ethics guidelines for trustworthy AI. 8 avril 2019. [En ligne]. Site disponible sur : <https://ec.europa.eu/digital-single-market/en/news/ethics-guidelines-trustworthy-ai> (page consultée le 3 décembre 2019).
105. CCNE (Comité Consultatif National d'Éthique pour les sciences de la vie et de la santé). Données massives et santé : Une nouvelle approche des enjeux éthiques. [En ligne]. Site disponible sur : https://www.ccne-ethique.fr/sites/default/files/publications/avis_130.pdf (page consultée le 5 novembre 2019).
106. Castelveccchi D. Can we open the black box of AI ?, *Nature*. 2016 ; **538**(7623) : 20-23.
107. Mueller ST, Hoffman RR, Clancey W, Emrey A, Klein G. Explanation in Human-AI Systems : A Literature Meta-Review, Synopsis of Key Ideas and Publications, and Bibliography for Explainable AI. ArXiv190201876 Cs.
108. De La Torre J, Valls A, Puig D, Romero-Aroca P. Identification and Visualization of the Underlying Independent Causes of the Diagnostic of Diabetic Retinopathy made by a Deep Learning Classifier. ArXiv180908567 Cs Stat.
109. France Stratégie. Intelligence artificielle et travail. [En ligne]. Site disponible sur : https://www.strategie.gouv.fr/sites/strategie.gouv.fr/files/atoms/files/fs-rapport-intelligence-artificielle-28-mars-2018_0.pdf (page consultée le 3 novembre 2019).
110. Décret n° 2016-1672 du 5 décembre 2016 relatif aux actes et activités réalisés par les manipulateurs d'électroradiologie médicale. 2016-1672 5 décembre 2016.
111. Le Monde.fr. Sanofi supprime 466 postes, dont près de 300 en France. [En ligne]. Site disponible sur : https://www.lemonde.fr/economie/article/2019/06/19/sanofi-poursuit-ses-restructurations-en-france_5478347_3234.html (page consultée le 16 novembre 2019).
112. Veyssière M, Robeveille R. Manager l'Intelligence artificielle. Gereso. 2019.
113. Université de Lille. Ouverture du DU Intelligence artificielle en santé. [En ligne]. Site disponible sur : <http://medecine.univ-lille.fr/index.php> (page consultée le 10 décembre 2019).
114. Deep Knowledge Analytics. AI in Drug Discovery. [En ligne]. Site disponible sur : <https://ai-pharma.dka.global/> (page consultée le 22 novembre 2019).

Annexes

Nom entreprise	Lien site internet	Domaine thérapeutique	Partenariat avec laboratoires pharmaceutiques
Ardigen	https://ardigen.com/	Immunologie, oncologie, Gastro-entérologie	
Atomwise	https://www.atomwise.com/	Neurologie, oncologie	Hansoh Pharmaceutical Group Company Limited, Eli Lilly's, Charles River Laboratories, Pfizer
Benevolent.AI	https://benevolent.ai/	Neurologie, oncologie, Néphrologie, Pneumologie	Novartis, Neuropore Therapies, AstraZeneca
Biovista	https://www.biovista.com/	Maladies rares, neurologie, oncologie	Astellas Pharma, Pfizer, Novartis
C4X discovery	https://www.c4xdiscovery.com/	Maladies inflammatoires, neurologie, oncologie	Indivior, AstraZeneca
Cyclica	https://cyclicarx.com/	Oncologie, neurologie	Yuhan, Enamine,
CytoReason	https://www.cytoreason.com/	Oncologie	GSK, Pfizer
Deep Genomics	https://www.deepgenomics.com/	Ophtalmologie, maladies métaboliques	Wave Life Sciences
DeepMind Health	https://deepmind.com/	Oncologie, neurologie	
e-therapeutics	https://www.etherapeutics.co.uk/	Neurologie, oncologie	
Exscientia	https://www.exscientia.ai/	Neurologie, immuno-oncologie, maladies métaboliques, maladies rares, cardiologie	Sumitomo Dainippon Pharma, Sunovion Pharmaceuticals, Evotec, Sanofi, GSK, Roche, Celgene, GT Apeiron Therapeutics, Rallybio, Bayer
GNS Healthcare	https://www.gnshealthcare.com/	Oncologie	Novartis, Janssen, Celgene, Bristol-Myers Squibb
iCarbonX	https://www.icarbonx.com/en/		
Insilico Medicine	https://insilico.com/		
Insitro	http://insitro.com/	Maladies métaboliques	Gilead
Lantern Pharma	https://www.lanternpharma.com/	Oncologie	EISAI and MGI

			Pharma
Numerate	http://www.numerate.com/	Cardiologie, neurologie, oncologie	Lundbeck Pharmaceutical
Nuritas	https://www.nuritas.com/	Dermatologie, maladies métaboliques	
PathAI	https://www.pathai.com/	Maladies métaboliques	Bristol- Myers Squibb, Gilead
Recursion Pharmaceutical s	https://www.recursionpharma.com/	Cardiologie, neurologie, oncologie, dermatologie, ophtalmologie, immunologie	
Schrödinger	https://www.schrodinger.com/	Oncologie	Bayer, Sanofi, Takeda, Gilead, SPARC (Sun Pharma Advanced Research Company), Agios
TwoXAR	http://www.twoxar.com/	Oncologie, immunologie, ophtalmologie	SK Biopharmaceuticals , Ono Pharmaceutical, 1ST Biotherapeutics, Adynxx
Vyasa Analytics	https://vyasa.com/	Maladies rares	Pfizer
WuXi NextCODE	https://www.wuxinextcode.com/	Oncologie, cardiologie, maladies inflammatoires	
XtalPi	http://www.xtalpi.com/		Pfizer

Annexe 1 Les 25 entreprises d'IA, leaders mondiaux pour la découverte de nouveau médicaments d'après [114]

Nom	Description	URL du site Internet
Structures et propriétés calculées		
AFLOWLIB	Référentiel de la structure et des propriétés issus de calculs ab initio à haut débit des matériaux inorganiques	http://aflowlib.org
Computational Materials Repository	Infrastructure permettant la collecte, le stockage, la récupération et l'analyse de données à partir de codes de structure électronique	https://cmr.fysik.dtu.dk
GDB	Bases de données de petites molécules organiques hypothétiques	http://gdb.unibe.ch/downloads
Harvard Clean Energy Project		https://cepdb.molecularspace.org
Materials Project	Propriétés calculées de matériaux connus et hypothétiques réalisées à l'aide d'un schéma de calcul standard	https://materialsproject.org
NOMAD	Fichiers d'entrée et de sortie issus de calculs utilisant une grande variété de structures électroniques	https://nomad-repository.eu
Open Quantum Materials Database	Propriétés calculées de structures principalement hypothétiques réalisées à l'aide d'un schéma de calcul standard	http://oqmd.org
NREL Materials Database	Propriétés calculées de matériaux pour des applications d'énergie renouvelable	https://materials.nrel.gov
TEDesignLab	Propriétés expérimentales et calculées pour aider à la conception de nouveaux matériaux thermoélectriques	http://tedesignlab.org
ZINC	Molécules organiques disponibles dans le commerce en formats 2D et 3D	https://zinc15.docking.org
Structure et propriétés issues de l'expérimentation		
ChEMBL	Des molécules bioactives aux propriétés médicamenteuses	https://www.ebi.ac.uk/chembl
ChemSpider	Base de données sur la structure de la Royal Society of Chemistry, contenant des Propriétés issues de calculs et d'expériences provenant de	https://chemspider.com

		diverses sources	
Citrination		Propriétés des matériaux issues de calculs et d'expérimentation	https://citrination.com
Crystallography Database	Open	Structures de composés organiques, inorganiques, métallo-organiques et minéraux	http://crystallography.net
CSD		Référentiel des structures cristallines des petites molécules organiques et métallo-organiques	https://www.ccdc.cam.ac.uk
ICSD		Base de données sur la structure des cristaux inorganiques	https://icsd.fiz-karlsruhe.de
MatNavi		Plusieurs bases de données ciblant des propriétés telles que la supraconductivité et la conductance thermique	http://mits.nims.go.jp
MatWeb		Fiches techniques pour divers matériaux d'ingénierie, y compris les thermoplastiques, les semi-conducteurs et des fibres	http://matweb.com
NIST Chemistry WebBook		Données thermochimiques et spectroscopiques en phase gazeuse de haute précision	https://webbook.nist.gov/chemistry
NIST Materials Data Repository	Data	Référentiel pour télécharger les données de matériaux associées à des publications spécifiques	https://materialsdata.nist.gov
PubChem		Activités biologiques de petites molécules	https://pubchem.ncbi.nlm.nih.gov

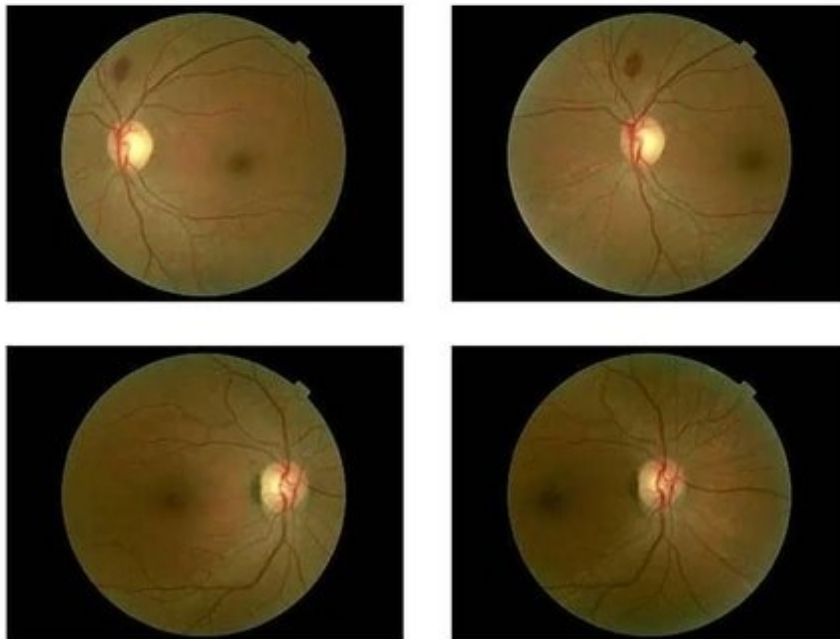
Annexe 2 Bases de données chimiques et biologiques

Nom du logiciel (<i>framework</i>)	Description
Généraliste	
<u>Chainer</u>	Logiciel d'apprentissage profond écrit en langage python
<u>Caffe</u>	Logiciel d'apprentissage profond écrit en langage python développé par la Berkeley AI Research (BAIR) et par une communauté de contributeurs
<u>Keras</u>	Logiciel d'apprentissage profond écrit en langage python
<u>Pytorch</u>	Logiciel d'apprentissage machine écrit en langage python développé par l'entreprise Facebook
<u>Tensorflow</u>	Logiciel d'apprentissage machine écrit en langage python développé par la société Google
<u>Theano</u>	Logiciel d'apprentissage profond écrit en langage python développé par Mila (Institut québécois d'intelligence artificielle)
Spécialisé en biologie et chimie	
<u>Chainer Chemistry</u>	Bibliothèque d'extension du logiciel Chainer pour la biologie et la chimie
<u>Deepchem</u>	Bibliothèques en langage python pour les algorithmes d'apprentissage profond développées par l'université de Stanford et par la société Schrödinger

Annexe 3 Logiciels d'apprentissage automatique en open source

IDx-DR Analysis Report

Patient ID: 2016-09-206:09:44PM
IDx Submission ID: 22-1
Exam Analysis Date: 2016-09-20
Exam Analysis Time: 6:05:11 PM
Exam Result: Moderate diabetic retinopathy detected



WARNING The above images are reduced resolution, compressed versions of the original images used by IDx-DR Client.
Do NOT use these images for diagnostic purposes.

DR-BJ 2.0.2

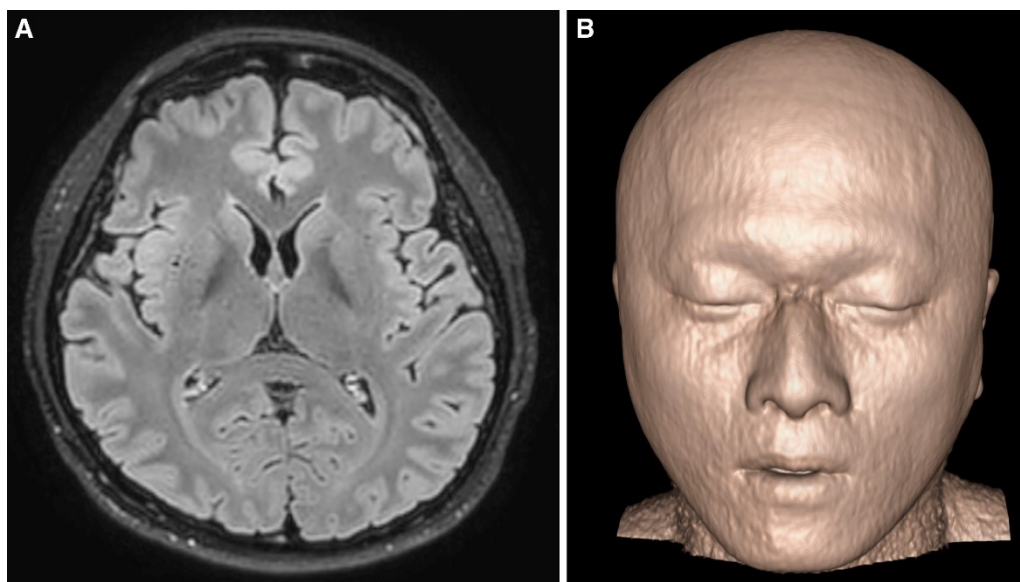
Annexe 4 Exemple de résultat de l'analyse d'images par le logiciel IDx-DR [71]



Annexe 5 Exemple de sortie du logiciel Transpara™ [75]



Annexe 6 Exemple de sortie du logiciel Transpara™ [75]



Annexe 7 Reconstruction d'un visage à partir d'une image IRM d'après la référence [95]

DEMANDE D'IMPRIMATUR

Date de soutenance : 16/09/2020

DIPLOME D'ETAT DE DOCTEUR EN PHARMACIE présenté par : Mazari Zemihi <u>Sujet</u> : L'utilisation de l'intelligence artificielle dans la découverte de nouveaux candidats-médicaments et dans l'imagerie médicale. <u>Jury</u> : Président : M. Michel BOISBRUN, Maître de Conférences, Pharmacien Directeur : M. Gérard MONARD, Professeur Co-directeur : Mme Alexandrine LAMBERT, Maître de conférences Juges : M. Nicolas VERAN, Radiopharmacien M. Antoine CAROF, Maître de Conférences	Vu, Nancy, le 23/06/2020 <table border="0"><tr><td>Le Président du Jury</td><td>Directeur de Thèse</td><td>Co-directeur</td></tr><tr><td>M. BOISBRUN</td><td>M. MONARD</td><td>Mme LAMBERT</td></tr><tr><td></td><td></td><td></td></tr></table>	Le Président du Jury	Directeur de Thèse	Co-directeur	M. BOISBRUN	M. MONARD	Mme LAMBERT			
Le Président du Jury	Directeur de Thèse	Co-directeur								
M. BOISBRUN	M. MONARD	Mme LAMBERT								
										
Vu et approuvé, Nancy, le 18.08.2020 Doyen de la Faculté de Pharmacie de l'Université de Lorraine,  Pr. Raphaël DUVAL	Vu, Nancy, le Le Président de l'Université de Lorraine,  Pierre MUTZENHARDT N° d'enregistrement : 11279C									

N° d'identification :

TITRE

L'utilisation de l'intelligence artificielle dans la découverte de nouveaux candidats-médicaments et dans l'imagerie médicale.

Thèse soutenue le 16/09/2020

Par Mazari ZEMHI

RESUME

L'IA est un sujet d'actualité. Mais qu'est-ce que c'est exactement ? Une science, une technologie, un domaine informatique ? Dans cette thèse, nous tenterons d'en donner une définition. Nous verrons que l'industrie pharmaceutique porte un intérêt à l'IA et l'utilise dans des cas d'usage bien spécifiques, notamment dans la découverte de nouvelles entités chimiques. Dans le domaine médical, c'est la radiologie qui a tiré profit des capacités des solutions d'IA. Cependant, l'IA suscite des interrogations éthiques et des inquiétudes légitimes à propos de son impact sur l'emploi.

MOTS CLES : intelligence artificielle, apprentissage automatique, *drug discovery*, imagerie médicale, éthique.

Directeur de thèse	Intitulé du laboratoire	Nature
Mr Monard Gérard	Physique et Chimie Théoriques UMR 7019, Université de Lorraine, CNRS Boulevard des Aiguillettes B.P. 70239 F-54506 Vandoeuvre-les-Nancy, FRANCE	Expérimentale ☐ Bibliographique ☒ Thème ☐

Thèmes

1 – Sciences fondamentales

2 – Hygiène/Environnement

③ – Médicament

4 – Alimentation – Nutrition

5 – Biologie

6 – Pratique professionnelle