



Version 01 : Octobre 2025

**FMD2101 : Les fondamentaux du
Machine Learning et du Deep Learning**



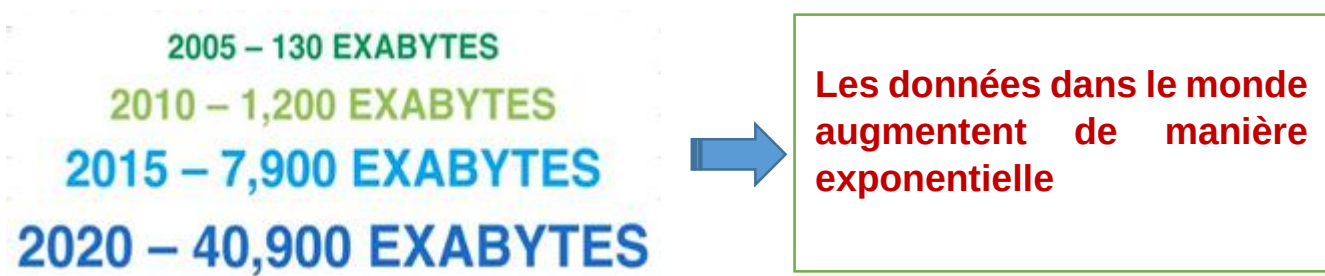
Cours
Semestre 1
Spécialité : CIO-BC

FMD2101 : Les fondamentaux du Machine Learning et du Deep Learning

Leçon 1: INTRODUCTION A L'IA, MACHINE LEARNING ET DEEP LEARNING

Introduction :

Avec l'avènement de l'informatique puis d'Internet, le volume des données a littéralement explosé, au cours des dernières décennies. De ce fait, la croissance des données s'est multipliée par quarante (40) si on s'en tient aux périodes couvrant 2010 et 2020. Cette complexité et cette abondance du Big Data "**Données massives**" rendent les méthodes traditionnelles d'analyse inefficaces voire obsolètes.



C'est dans ce contexte que l'apprentissage automatique et l'apprentissage profond apparaissent comme une nécessité. Il offre des outils avancés permettant d'extraire automatiquement des connaissances, de détecter des modèles complexes, et de prédire des tendances à partir de ces grands ensembles de données de diverses formes. Comprendre les enjeux de ces méthodes soulève des questions aux mathématiciens, statisticiens et informaticiens.

Dans cette leçon se veut d'établir plus clairement les concepts clés de l'intelligence Artificielle, ainsi que ce qui relèvent ou non du machine Learning.

Objectifs :

- **Appréhender** les enjeux du Big Data et son rôle dans l'émergence du Machine Learning ;
- **Distinguer** les approches classiques de programmation des approches par apprentissage automatique ;
- **Identifier** les principaux types d'apprentissage automatique (supervisé, non supervisé, etc.) ;
- **Comprendre** les concepts fondamentaux de l'Intelligence Artificielle (IA), du Machine Learning (ML) et du Deep Learning (DL) ;
- **Saisir** les liens hiérarchiques et fonctionnels entre l'IA, le ML et le DL.

I. La data science au cœur des enjeux contemporains

1. Un peu d'histoire

En 2020 Netflix affichait un chiffre d'affaire de 25 milliards de dollars pour 200 millions d'abonnés à travers le monde. Mais Netflix est sans doute la première entreprise dans les médias à prendre le traitement des données au sérieux, à chaque fois que vous en tant que client faites un événement sur la plateforme, elle l'enregistre et le traite.

Par exemple pour chaque abonné vous avez :

- ~ Le type de film ou séries que vous regardez ;
- ~ Lorsque vous mettez pause, revenez-vous en arrière ou passez rapidement ?
- ~ La façon dont vous regardez une série (premier visionnage sans lendemain, visionnage dans son intégralité, boulimique ou étalé, etc...) ;
- ~ Le jour et l'heure de votre visionnage ;
- ~ L'appareil utilisé ;
- ~ Votre évaluation sur l'application ;
- ~ Le type de recherche que vous faites sur leur moteur ;
- ~ Etc...

Plusieurs informations pour des millions d'abonnés, on parle donc ici de ****Big Data****. **L'Intelligence Artificielle, ce n'est pas du nouveau !**

2. Le Big data qu'est-ce que c'est ?

Depuis les années 2000, le terme ***Big Data*** a connu un certain succès, pour finalement devenir un concept phare dans le domaine du traitement de données. Ainsi le Big Data se définit donc comme des masses de données à la fois structurées et non structurées, générées souvent en flux continu, stockées dans des Data Lakes et qui pourraient être interprétées en quasi-temps réel.

Afin de caractériser ce type de données, les **<< 3V du Big Data >>** ont été popularisés par le cabinet d'étude Gartner :

- ~ **Volume** : Le volume de données produites est considérable et sa croissance est exponentielle. Les sources de données se multiplient, avec Internet et les IOT.
- ~ **Vélocité** : C'est la rapidité de production des données, les bases de données traditionnelles étaient relativement stables, s'incrémentant progressivement au fur et à mesure de l'enregistrement de données. Ici les données sont produites de façon automatique ou semi-automatique et de manière continue en fonction du trafic sur un site, de l'utilisation d'un objet, de l'enregistrement des activités commerciales, etc. L'enjeu est de traiter ces données à la volée.
- ~ **Variété** : Les données traditionnelles étaient structurées alors que la plupart de celles dont on dispose aujourd'hui ne le sont plus (texte, image, sons, vidéos, etc...). Les outils de traitement peuvent donc faire appel aux techniques de reconnaissance visuelle par exemple, ou de traitement de la voix. La difficulté est d'exploiter des masses d'information aussi variées.

Remarque : A ces V l'on peut ajouter la **Valeur** et la **Véracité** si l'on veut parler des 5 V's du Big Data.

3. A quoi sert le Big Data ?

Le Big data va aider les entreprises dans diverses activités telles que :

Le marketing : avec la grande quantité de données analysée, il est possible de mieux cerner les clients et prospects et ainsi d'adapter et personnaliser le contenu et les communications que leur seront dédiés.

Le développement d'un produit : en effet, les données récoltées vont permettre d'anticiper la demande des clients puisqu'elles modélisent de nouveaux produits selon les nouvelles innovations et les succès commerciaux.

L'expérience client : avec la Big Data, les entreprises disposent d'une meilleure visibilité sur l'expérience de leurs clients et elles peuvent donc proposer des services plus adaptés et traiter plus efficacement les problèmes qu'ils rencontrent.

Améliorer l'utilisation des machines et autres appareils : grâce à la Big Data, les entreprises peuvent désormais analyser et évaluer leur production de sorte à être capable d'anticiper d'éventuelles pannes et tenter de les réduire.

Anticiper la demande client : avec l'analyse des données de la Big Data, il est possible de prédire la demande future au vu de l'analyse du marché, de la demande passée et de bien d'autres facteurs etc.

4. Les métiers de la Data Science le Big Data ?

La science des données plus connu sous le nom data science, est la discipline qui s'appuyant sur les statistiques, les mathématiques et l'informatique, consiste à collecter, agréger, nettoyer et classer les données pour ensuite les analyser, les interpréter, mettre en place des algorithmes traitant ces données afin d'en sortir des insights.

Et donc la data science est la science qui permet d'analyser les données du big data.

Plusieurs métiers en découlent et dont les plus connu sont entre autres:

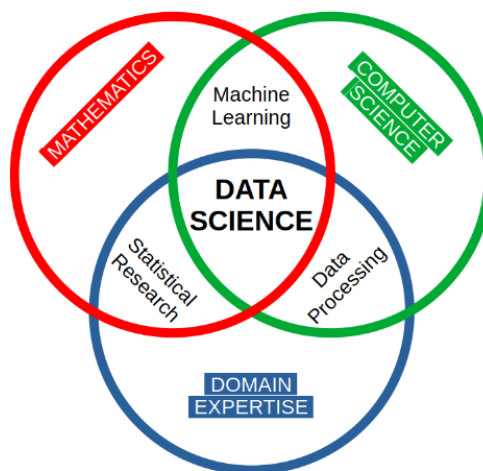
- ~ **Data Scientist** : Il a pour rôle d'analyser de grands ensembles de données pour extraire des insights et des tendances significatives. Concevoir, développer et mettre en œuvre des modèles prédictifs et des algorithmes statistiques pour résoudre des problèmes métier ;
- ~ **Data Engineer** : Il a pour mission de concevoir, construire et entretenir les infrastructures de données et les pipelines pour recueillir, stocker, et traiter les données. Assurer la qualité et la fiabilité des données ;

~ **Business Analyst :**

Le Business Analyst fait le lien entre les besoins métier et les solutions techniques. Il analyse les processus, identifie les besoins, rédige les spécifications et accompagne le projet jusqu'à sa mise en œuvre.

- ~ **Data Analyst** : Il a pour rôle d'explorer et analyser les données pour identifier des tendances, des modèles et des insights qui peuvent être utilisés pour prendre des décisions commerciales. Créer des rapports et des visualisations de données pour communiquer les résultats.
- ~ **Data Architect** : Il a pour rôle de concevoir et gérer l'architecture des systèmes et des bases de données pour assurer la qualité, la sécurité et la scalabilité des données. Définir les normes et les bonnes pratiques en matière de gestion des données.
- ~ **Data Product Manager** : Il est responsable de la valeur des produits data (dataset, Dashboard, APIs, algorithmes, plateformes...) qu'il conçoit pour répondre aux besoins métiers tout en assurant leur viabilité technique et leur impact business. Il est le chef d'orchestre entre les équipes data (data engineers, data scientists, data analysts), les métiers et parfois les équipes tech classiques.

Et donc faire de la data science ne consiste pas seulement à faire de la collecte, du stockage, de la description et de l'analyse de données, il y'a aussi un volet qui consiste à déterminer les schémas sous-jacents existant dans les données pour ensuite les utiliser afin d'effectuer des prédictions concernant de nouveaux exemples. On parle donc d'apprentissage automatique ou de **Machine Learning**, qui lui-même appartient à la **GRANDE FAMILLE DE L'INTELLIGENCE ARTIFICIELLE**.



II. De l'Intelligence Artificielle au Machine Learning :

1. L'Intelligence Artificielle.

1.1. L'Intelligence Artificielle, ce n'est pas du nouveau !

Lorsque l'on pose la question aux citoyens lambda quant à la date de naissance de l'IA, beaucoup indique les années 2000. Mais, il n'en est rien !

A la genèse, c'est lors de l'été 1956 qu'a vu officiellement le jour l'Intelligence Artificielle au Dartmouth Collège (New Hampshire, Etats-Unis) lors d'une université d'été (**18 juin au 17 août**) organisée par **John McCarthy**. Académique qu'est l'Intelligence Artificielle suppose que toutes les fonctions cognitives humaines peuvent être décrites de façon très précise pouvant alors donner lieu à une reproduction de celle-ci sur ordinateur. Il serait alors possible de créer des systèmes capables d'apprendre, de calculer, de mémoriser, et pourquoi pas de réaliser des découvertes scientifiques ou encore de la créativité artistiques.

1.2. Mais qu'est-ce que L'IA ?

Comme nous venons de le voir, l'IA n'est pas une science nouvelle. Cependant, qui peut affirmer avoir vu une IA ? Personne n'ose s'aventurer sur ce chemin sans issue car l'IA est inodore et invisible. Les robots, les voitures autonomes ne sont pas ce que l'on peut appeler des IA, ce sont plutôt des machines utilisant cette intelligence.

L'intelligence artificielle, définie comme l'ensemble des techniques mises en œuvre afin de construire des machines capables de faire preuve d'un comportement que l'on peut qualifier d'intelligent, fait aussi appel aux sciences cognitives, à la neurobiologie, à la logique, à l'électronique, à l'ingénierie et bien plus encore.

A retenir

Au demeurant, l'Intelligence Artificielle (IA) est la science dont le but est de faire exécuter par une machine, des tâches que l'homme accomplit en utilisant son intelligence (**Debois, 2021**).

Une grande variété d'activités entre dans le champ disciplinaire de l'IA notamment les systèmes experts, la reconnaissance faciale et vocale, le traitement des langages naturels et de l'imagerie médicale. Par ailleurs, L'IA continue de prendre son envol en proposant d'une part des algorithmes capables d'apprendre et de s'adapter à de nouveaux problèmes et d'autre part, en construisant des systèmes de plus en plus indépendants (**Honneurs, 2008**).

1.3. Les intérêts de l'IA pour l'Homme

L'IA a complètement bouleversé et transformé sous l'impulsion des algorithmes, le quotidien de l'Homme dans plusieurs domaines. Elle accomplit avec une grande simplicité et une fiabilité extrême, des tâches informatisées colossales qui, jusqu'à alors, étaient exécutées de manière répétitive et manuellement par l'Humain (**Hervé, 2023**). Partant de ce fait, l'IA propose des solutions novatrices dans de nombreux secteurs tout en améliorant la performance des processus de prise de décision endossés aux fonctionnements des domaines concernés. Nous répertorions un avant-goût des avantages de l'IA pour des spécialités (**Trey, 2022**). Ainsi, calquée pour fonctionner comme l'intelligence Humaine, l'IA :

- Améliore la productivité ;
- Autonomise des tâches répétitives (pour libérer le temps aux humains) ;
- Aide à la prise de décision basée sur des données à temps réel ;
- perfectionne l'utilisation des données ;
- réalise une analyse poussée des données grâce aux méthodes du Big Data ;
- augmente la précision de détection des certains phénomènes ou maladies.

1.4. Les catégories d'IA.

Les domaines d'application de l'IA sont transversaux et peuvent être regroupés en diverses catégories, à savoir : les systèmes experts, la reconnaissance des formes, des visages et des voix, le traitement des langages naturels et de l'imagerie médicale et le Machine Learning.

- **Les systèmes experts** : Un système expert est un outil capable de reproduire les mécanismes cognitifs d'un expert, dans un domaine particulier. Il s'agit de l'une des voies tentant d'aboutir à l'intelligence artificielle. Plus précisément, un système expert est un logiciel capable de répondre à des questions, en effectuant un raisonnement à partir de faits et de règles connues. Il peut servir notamment comme outil d'aide à la décision. Un système expert se compose de trois parties : une base de faits, une base de règles et un moteur d'inférence. Le moteur d'inférence est capable d'utiliser les faits et les règles pour produire de nouveaux faits, jusqu'à parvenir à la réponse de la question experte posée (Maria, 2008). Nous avons pour exemple, le premier système expert Dendral. Il permettait d'identifier les constituants chimiques (Trey, 2022). La figure ci-dessous proposée par (Mougnutou, 2023) illustre son fonctionnement.

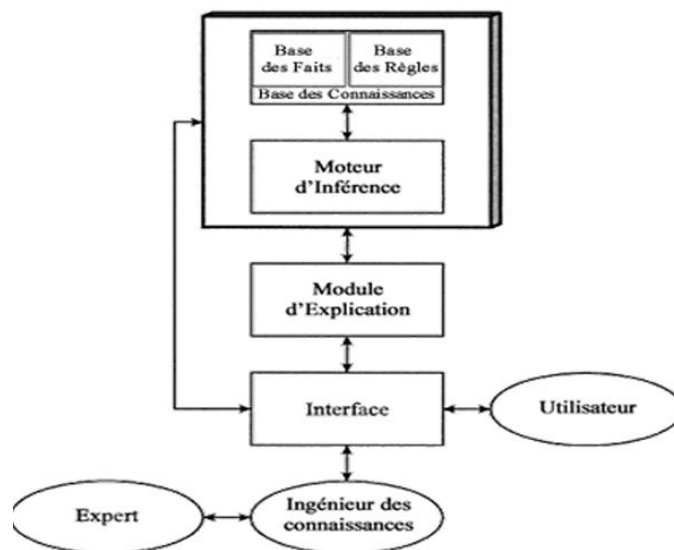


Figure 1 : Structure fondamentale d'un système expert

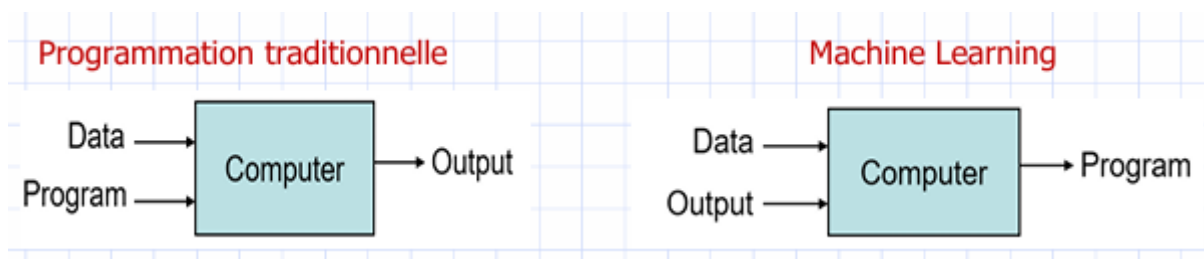
- **La reconnaissance des formes, des visages et des voix** : L'IA au fil des temps a engendré plusieurs sous disciplines parmi lesquelles la reconnaissance forme (Rdf). Celle-ci admet pour dérivée, la vision par ordinateur, la reconnaissance vocale (Rasa, Cortona, Alexa, Siri...) et faciale. La reconnaissance faciale est une technologie ou un algorithme automatique qui consiste à analyser une image et à reconnaître ses caractéristiques (le nez, la texture de la peau, les yeux et la bouches). Ce processus se base sur la détermination par calcul de l'emplacement exact et la taille de ces traits humains. Toutefois, la précision de cet outil est fonction de la qualité de l'image, de la luminosité de la pièce et l'angle de prise de vue du visage (Sharma, 2023). Nous l'avons compris, avec le développement exponentiel des domaines d'application de l'IA, il est difficile, voire impossible de faire une liste exhaustive de ses différentes utilisations puisqu'elle est devenue un abri pour tous les chercheurs.

2. Le Machine Learning

2.1. Qu'est-ce que le Machine Learning ?

L'IA est à notre sens un mot défini par **Marvin Minsky** comme étant une science dont **le but principal est de faire réaliser par la machine des tâches que l'homme accomplit en utilisant son intelligence.**

Un pan de cette activité cognitive se réalise grâce à l'apprentissage machine. Celle-ci se définit comme est un domaine d'étude de l'intelligence artificielle. Sa particularité face aux programmes classiques est que, plutôt que de programmer explicitement une solution, c'est dire qui utilise une procédure et les données qu'il reçoit en entrée pour produire en sortie des réponses, à un programme d'apprentissage automatique, utilise les données et les réponses afin de produire la procédure qui permet d'obtenir les secondes à partir des premières.



Exemple : Problème traditionnelle & Machine Learning

1. Supposons que la scolarité d'une université veut connaître le total des notes obtenus dans une ECUÉ pendant par un étudiant en Licence 1 de l'UVCI à la fin du premier semestre. Pour y arriver, il suffit d'appliquer un programme traditionnel (*algorithme classique*) à savoir une simple addition : un algorithme d'apprentissage n'est pas nécessaire.
2. Supposons maintenant que l'on veut utiliser ces notes pour déterminer l'ECUE dans laquelle il est susceptible d'avoir la bonne note au premier semestre de Master 1. Bien que cela soit vrai semblablement lié, nous n'avons manifestement pas toutes les informations nécessaires pour ce faire. Cependant, si nous disposons de l'historique notes d'un grand nombre d'étudiants, il devient possible d'utiliser un algorithme de machine Learning pour qu'il en tire un modèle prédictif nous permettant d'apporter une réponse à notre question

Le Machine Learning sert à résoudre les problèmes :

- Que l'on ne sait pas résoudre avec les algorithmes classiques (*comme la prédiction des notes d'un étudiant*) ;
- Que l'on sait résoudre, mais dont on ne sait formaliser la solution (*c'est le cas de la reconnaissance des images ou de la compréhension du langage naturel*) ;
- Que l'on sait résoudre, mais avec des procédures fastidieuses, longue et gourmandes en ressources informatiques (*c'est le cas par exemple de la prédiction d'interactions entre molécules de grande taille, pour lesquelles les simulations sont très lourdes*)

En somme, **Le machine Learning est donc utilisé quand les données sont abondantes (relativement), mais les connaissances peu accessibles ou peu développées.**

A retenir

De façon simple le Machine Learning c'est :

- Un vaste ensemble d'outils pour modéliser et comprendre les données complexes
- Le fait de comprendre et apprendre un comportement à partir d'exemples dans le but de faire des prédictions, des classifications ou prendre des décisions automatiques.

Se référer à [Machine Learning, Laurent Rouviere, Septembre 2020
(https://lrouviere.github.io/machine_learning/cours.pdf)

2.2. Les piliers du Machine Learning ?

Le Machine Learning repose sur deux piliers fondamentaux :

- D'une part, les données, qui sont les exemples à partir duquel l'algorithme va apprendre ;
- D'autre part, l'algorithme d'apprentissage, qui est la procédure que l'on fait tourner sur ces données pour produire un modèle.

Ces deux piliers sont aussi importants l'un que l'autre car aucun algorithme d'apprentissage ne pourra créer un bon modèle à partir de données qui ne sont pas de bonne. On parle des concepts **garbage in, garbage out**.

Toutefois, il ne faut pas négliger qu'une part importante du travail de **machine learner** ou de **data scientist** est un travail d'ingénierie consistant à préparer les données (data processing) afin d'éliminer les données aberrantes, gérer les données manquantes, choisir une représentation pertinente, etc.

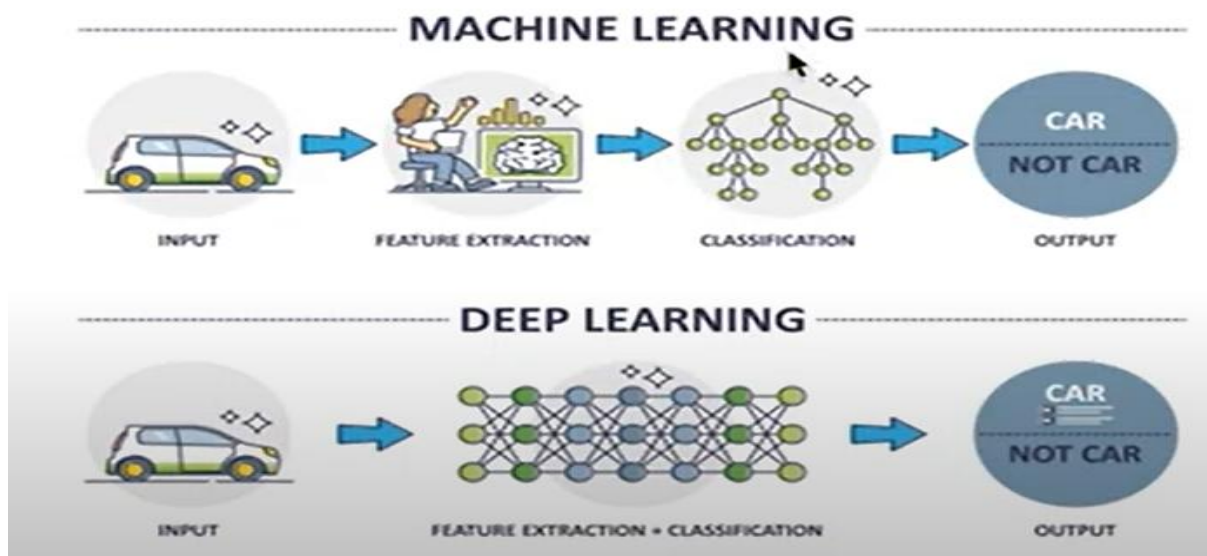
Attention

1. Bien que l'usage soit souvent d'appeler les deux du même nom, il faut distinguer l'algorithme d'apprentissage automatique du modèle appris : le premier utilise les données pour produire le second, qui peut ensuite être appliqué comme un programme classique ;
2. En somme, un algorithme d'apprentissage permet donc de modéliser un phénomène à partir d'exemples ;
3. Le Machine Learning repose d'une part sur les mathématiques, Ainsi, le machine Learning repose d'une part sur les mathématiques, et en particulier les statistiques, pour ce qui est de la construction de modèles et de leur inférence à partir de données, et d'autre part sur l'informatique, pour ce qui est de la représentation des données et de l'implémentation efficace d'algorithmes d'optimisation.

3. Le Deep Learning

3.1. Définition du Deep Learning

Deep Learning ou apprentissage profond est un sous domaine du Machine Learning qui s'appuie sur l'usage de neurones artificiels s'inspirant de la structure du cerveau humain. Ces neurones sont organisés en plusieurs couches donnant alors une notion de profondeur (deep) au réseau de neurones. En gros, avec le Deep Learning, plus on va en profondeur dans la couches, plus la représentation des données est de bonne qualité. La figure ci-dessous montre la différence entre les techniques de ML et du Deep Learning.



3.2. Les caractéristiques du Deep Learning

- **Architecture en Réseaux de Neurones Profonds**

Elle constitue la principale caractéristique. Le « Deep » (Profond) fait référence de **nombreuses couches cachées** entre la couche d'entrée et la couche de sortie.

- **Machine Learning Classique** : Souvent 1 ou 2 couches (on parle de modèles "shallow" ou superficiels).
- **Deep Learning** : Peut avoir des dizaines, des centaines, voire des milliers de couches (ex: ResNet). Ces couches successives permettent d'extraire des caractéristiques de plus en plus complexes et abstraites.
- **Analogie** : Pour reconnaître un visage.
 - La première couche détecte des bords et des contours ;
 - La couche suivante combine ces contours pour détecter des formes (yeux, nez) ;
 - Les couches profondes assemblent ces formes pour reconnaître des parties du visage, puis le visage dans son ensemble.

- **Traitement de Données non Structurée**

Le Deep Learning excelle là où les méthodes traditionnelles étaient limitées : **Les données non structurées**.

- **Images** (reconnaissance d'objets, segmentation) ;
- **Son** (reconnaissance vocale, synthèse) ;
- **Texte** (traduction, analyse de sentiment) ;
- **Vidéo** (détection de mouvements)

Ces données, riches en information, étaient très difficiles à modéliser avec des features manuelles.

- **Forte Dépendance aux Données (Data Hunger)**

C'est à la fois force et une faiblesse des modèles du Deep Learning. Ils nécessitent **des volumes de données considérables** (parfois des millions) pour bien généraliser et éviter le surapprentissage (*overfitting*). En effet, le nombre de paramètres (*poids et biais*) à prendre est énorme. Il en est pour données dont un très grand volume est nécessaires pour ajuster correctement ces paramètres.

- **Besoin en Puissance de Calcul (GPU)**

L'entraînement de réseaux profonds sur de grands jeux de données est extrêmement gourmand en calculs. Cette tâche est rendue possible par l'utilisation de **cartes graphiques (GPU) ou Unité de Traitement Graphique**. Les GPU sont optimisés pour effectuer des millions d'opérations matricielles simples en parallèle, ce qui correspond parfaitement aux calculs nécessaires à l'apprentissage des réseaux de neurones.

- **Scalabilité croissante avec plus de données**

Contrairement à de nombreux algorithmes classiques qui plafonnent en performance après un certain volume de données, les modèles de Deep Learning **continuent généralement à s'améliorer** à mesure que la quantité de données d'entraînement augmente. Cette propriété d'extensibilité (scalability) est cruciale à l'ère du Big Data.

A retenir :

FORCES

1. **Performance de pointe** sur des tâches complexes (vision, langage).
2. **Automatisation** de l'ingénierie des caractéristiques.
3. **Extensibilité** avec la quantité de données.
4. **Polyvalence** grâce à des architectures spécialisées.

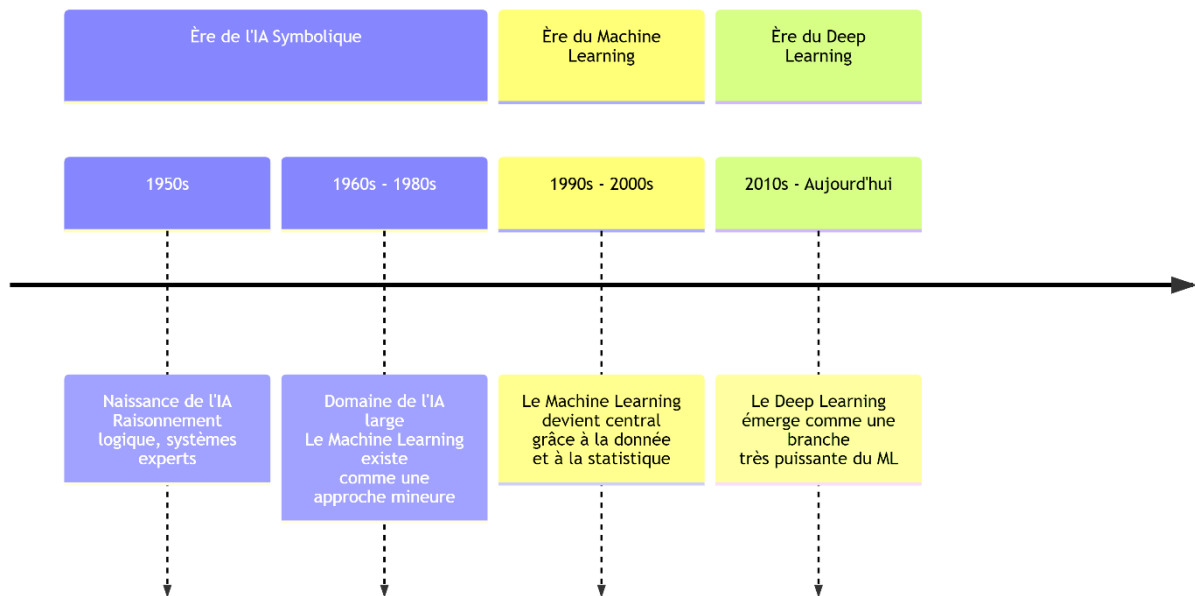
DEFIS

1. **Besoins massifs en données** (étiquetées de préférence).
2. **Besoins élevés en puissance de calcul** (coût et énergie).
3. **Complexité et "boîte noire"** : difficile d'interpréter les décisions du modèle.
4. Temps d'entraînement **parfois très longs**.

3.3. Lien entre IA, ML et DL : Une histoire Emboîtée dans le Temps

On peut visualiser leur relation comme une série de poupées russes ou de cercles concentriques qui apparaissent successivement dans l'histoire. **L'IA est le domaine le plus large et le plus ancien. Le Machine Learning en est une sous-partie cruciale, et le Deep Learning est une sous-partie spécialisée et récente du Machine Learning.** L'échelle temporelle ci-dessous s'efforce de mieux illustrer le rapport entre ces trois notions.

Évolution de l'IA, du ML et du Deep Learning



L'essentiel : IA & ML & DL

- **IA** : Machines capables d'imiter certains comportements de l'intelligence humaine ;
- **ML** : Algorithmes et modèles capables d'apprendre à partir des données (expériences) sans intervention humaine (*Extraction manuelle de caractéristiques ou Handcrafted features extraction en anglais*)
- **DL** : Basé essentiellement sur les réseaux de neurones profonds pour imiter le comportement du cerveau humain (*Extraction automatique des caractéristiques*)
- **ML et DL** sont des outils d'aide à la décision que les Data Scientist mettent à la disposition des experts de tous les domaines selon la problématique initiale.

Conclusion

Cette leçon a permis de situer le **Machine Learning (ML)** comme une composante essentielle de l'intelligence artificielle, elle-même née dans les années 1950. Face à l'explosion du volume des données (Big Data), le ML s'impose comme une solution pour extraire des connaissances à partir de jeux de données complexes, là où les méthodes traditionnelles échouent. De ce fait, le paysage du Machine Learning est marqué par une évolution **continue** : des approches classiques, interprétables mais limitées face aux données non structurées, vers des modèles profonds, performants mais gourmands en données et en calcul à l'instar du **Deep Learning**. Cette complémentarité entre familles de méthodes qui a permis d'adapter au fil du temps la solution au problème et aux données disponibles, nous s'impose à mieux cerner les spécificités de chacune de ces approches connexionniste de l'IA.

RESSOURCES WEBOGRAPHIQUES ET BIBLIOGRAPHIQUES INDICATIVES

- Bakir, G., Hofmann, T., Schölkopf, B., Smola, A. J., Taskar, B., et Vishwanathan, S. V. N. (2007). Predicting Structured Data. MIT Press, Cambridge, MA. [predicting-structured-data. https://mitpress.mit.edu/books/](https://mitpress.mit.edu/books/predicting-structured-data)
- Barto, R. S. et Sutton A. G. (1998). Reinforcement Learning: An Introduction. MIT Press, Cambridge, MA. <http://incompleteideas.net/book/the-book-2nd.html>
- Benureau, F. (2015). Self-Exploration of Sensorimotor Spaces in Robots. Thèse de doctorat, Université de Bordeaux.
- Samuel, A. L. (1959). Some studies in machine learning using the game of checkers. IBM Journal of Research and Development, 44(1.2) : 206–226.
- Scott, D. W. (1992). Multivariate density estimation. Wiley, New York.