

# Race Against the Machine: Deep Learning vs. Theory – Illustration with Deep Hedging

dans le cadre du Petit Déjeuner de l'AFGAP

Le 19 décembre 2019

**Le programme FaiR** (Finance and Insurance Reloaded), dirigé par Charles-Albert Lehalle, renforce la dynamique de recherche interdisciplinaire entre industriels et académiques de L'Institut Louis Bachelier - l'ILB. Il propose une nouvelle collection de publications réalisées à partir d'une série de tables rondes autour de l'impact des nouvelles technologies sur le monde de la finance et de l'assurance.

Fédérant plus de 60 chaires et initiatives de recherche, issues des plus grandes institutions académiques, l'Institut Louis Bachelier s'appuie sur les moyens de deux fondations (L'Institut Europlace de Finance - IEF et la Fondation du Risque - FdR) pour soutenir les programmes existants et permettre l'émergence de nouveaux projets. De par son activité, l'Institut Louis Bachelier favorise les passerelles entre la recherche et l'entreprise, éclaire les pouvoirs publics sur les enjeux d'aujourd'hui et de demain et valorise l'excellence de la recherche française. L'Institut Louis Bachelier réunit aujourd'hui plus de 400 chercheurs et près d'une cinquantaine d'entreprises et organisations mécènes.

Créée en 1991 à l'initiative de banquiers du CCF et de la Compagnie Bancaire pour échanger sur des questions techniques de gestion de bilan susceptibles d'intéresser toutes les banques, l'**AFGAP** est une association à but non lucratif qui regroupe des professionnels de la gestion de bilan d'institutions financières et de grandes entreprises françaises. Elle réunit près de 200 membres de l'ensemble des institutions françaises. Elle vise à faciliter les échanges sur ce sujet sensible, au coeur de la gestion financière des banques, à apporter des réponses aux questions techniques et conjoncturelles de la gestion du bilan, et à éclairer les autorités de tutelle sur les conséquences des évolutions prudentielles et de leurs mises en oeuvre.

L'AFGAP a également un rôle dans la formation de professionnels aptes à contribuer à la fonction de gestion de bilan. L'association organise un petit déjeuner par mois avec des intervenants extérieurs sur de multiples thèmes allant des taux bas à la gestion de la liquidité ou encore sur les évolutions prudentielles et comptables; et plusieurs événements dans l'année avec des associations sœurs en Europe (notamment l'Asset and Liability Management Association, ALMA, contrepartie en Grande Bretagne, françaises ou internationales. L'association compte environ 200 membres.



## TABLE DES MATIÈRES

|                                                                              |      |
|------------------------------------------------------------------------------|------|
| <b>Présentation</b>                                                          | p. 4 |
| <b>Le Machine learning dans l'industrie financière</b>                       | p. 5 |
| <b>Statisticiens vs. Data scientists</b>                                     | p.7  |
| <b>Une introduction rapide au deep learning</b>                              | p.9  |
| <b>Deep learning pour la valorisation et l'analyse des risques</b>           | p.12 |
| <b>Les challenges du deep learning</b>                                       | p.14 |
| <b>Des technologies de la désintermédiation</b>                              | p.16 |
| <b>L'Intelligence Artificielle : une "Technologie Générique"</b>             | p.18 |
| <b>L'impact des technologies de l'IA en finance de marché</b>                | p.19 |
| <b>Le rôle des données, et la complexité de l'automatisation</b>             | p.21 |
| <b>Faut-il ne réguler que les technologies ?</b>                             | p.23 |
| <b>Le cas des données de marché</b>                                          | p.24 |
| <b>Le cas spécifique du deep hedging</b>                                     | p.25 |
| <b>De l'extension des produits dérivés à l'apprentissage statistique</b>     | p.28 |
| <b>Le rôle de la modélisation dans la validation de résultats numériques</b> | p.30 |
| <b>Les technologies de rupture transforment notre vision du monde</b>        | p.32 |
| <b>Le raisonnement au secours de l'exploitation aveugle des données</b>      | p.33 |

# PRÉSENTATION

## Stéphane Denise

Aujourd'hui, nous avons le plaisir d'accueillir trois orateurs.

- Sandrine Ungari, Head of Cross Asset Research à la Société Générale à Londres, donc elle vient de loin pour vous.
- Nicole El Karoui, professeure émérite de mathématiques appliquées à l'Université Paris Sorbonne et
- Charles-Albert Lehalle qui est Head of Data Analytics chez Capital Fund Management.

L'Association Française des Gestionnaires Actif Passif (AFGAP) remercie vivement ces 3 orateurs, d'autant plus que la question du big data et de ses applications aux activités bancaires et d'assurance est une problématique importante qui suscite l'intérêt des directions financières de nos institutions.

Ils vont vous offrir trois perspectives différentes sur ce sujet et notamment une réflexion sur l'articulation entre modélisation, exploitation des gros volumes de données et nécessité de fournir à nos directions générales la logique et la traçabilité des propositions de décision que nous formulons sur la base de ce type de modèle. Nous avons bien conscience que ces modèles font l'objet de controverses, justement parce qu'ils peuvent être perçus par les directions générales ou par les superviseurs comme des « boîtes noires » à la robustesse incertaine. C'est pourquoi ces trois témoignages ce matin sont précieux pour l'industrie bancaire et assurantielle.

Encore merci d'avoir accepté notre invitation.

## Sandrine Ungari

Merci Stéphane, merci beaucoup pour l'invitation. Je suis ravie de présenter ce matin les travaux de mon équipe sur le Machine Learning en général et le deep learning en particulier.

A la Société Générale, je m'occupe d'une équipe d'une dizaine de personnes qui fait de la recherche quantitative à destination des clients investisseurs de la banque. Dans les années récentes, nous nous sommes intéressés à des problématiques d'investissement quantitatifs et systématiques, maintenant connus sous le nom de « Quantitative Investment Strategy » (QIS). Dans ce cadre, nous regardons les différentes techniques qui s'offrent au fil du temps aux investisseurs.

Le but de la présentation aujourd'hui est multiple. D'abord, j'aimerais revenir sur les notions qui se cachent derrière le terme « machine learning »: d'où vient-on, pourquoi en parle-t-on. J'aimerais ensuite introduire rapidement les notions de deep learning et deep hedging, pour conclure sur les enjeux, tels que perçus par les investisseurs, de ce type d'approche.

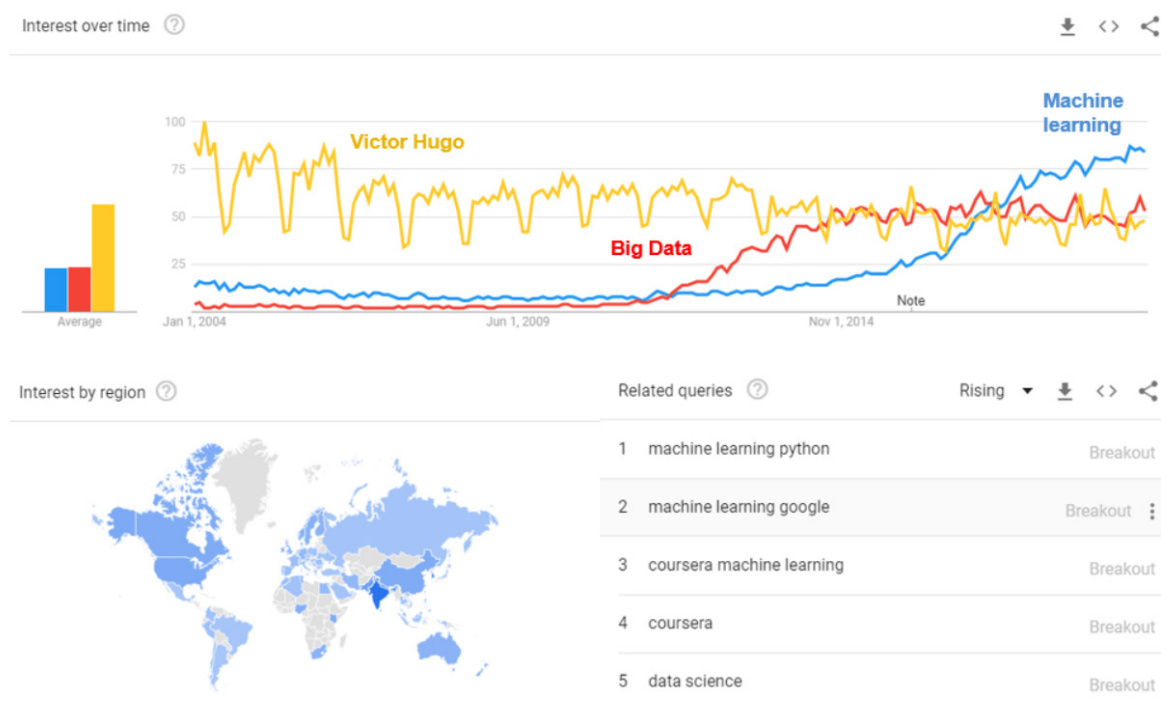
# 1 LE MACHINE LEARNING DANS L'INDUSTRIE FINANCIÈRE

## Sandrine Ungari

Une observation simple pour commencer. D'après « Google trend », au début des années 2010 la mode est au big data, avec un nombre croissant de recherches utilisant ce terme. Cette tendance s'est essouffée au milieu de la décennie, pour laisser place à la mode du « machine learning ». Quand on regarde la liste des recherches associées au terme « machine learning », on voit apparaître des mots comme « Python » ou « Coursera ». Les internautes cherchent non seulement de l'information sur le thème, mais ils cherchent également à se former sur les techniques associées.

A titre de comparaison, le machine learning a récemment dépassé Victor Hugo en nombre de recherches sur Google.

**Graph 1. Machine learning vs. Victor Hugo**



Source : <https://trends.google.co.uk/trends/explore?date=all&q=machine%20learning>

Intéressons-nous maintenant à ce qui se passe dans notre industrie. Je vais utiliser tout au long de ma présentation une étude qui a été publiée en octobre 2019 par la Bank of England (BOE) et la Financial Conduct Authority (FCA).

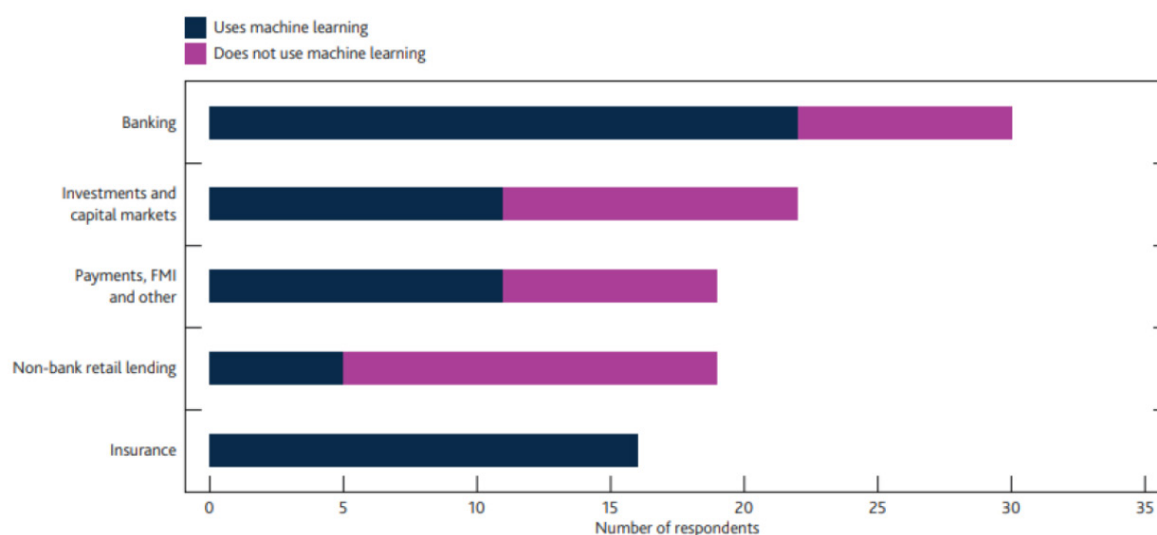
En 2019, la FCA et la BOE ont réalisé un sondage auprès des institutions financières au Royaume-Uni qui leur a permis de collecter une centaine de réponses. Dans leur rapport publié en Octobre 2019, on trouve un résumé détaillé de ce que font et anticipent les différents acteurs du secteur.

**100% DES PERSONNES INTERROGÉES ANTICIPENT UNE AUGMENTATION  
DE L'UTILISATION DU MACHINE LEARNING EN FINANCE**

Selon cette étude, 100% des participants dans le monde de l'assurance affirment utiliser le machine learning. A peu près 75% des participants dans la banque affirment utiliser le machine learning dans leur processus de production. De manière unanime, tous les participants anticipent une augmentation du nombre d'utilisation du machine learning dans l'industrie.

Le "machine learning" est donc un sujet important pour l'industrie. C'est probablement la raison pour laquelle nous nous retrouvons ce matin pour débattre du thème.

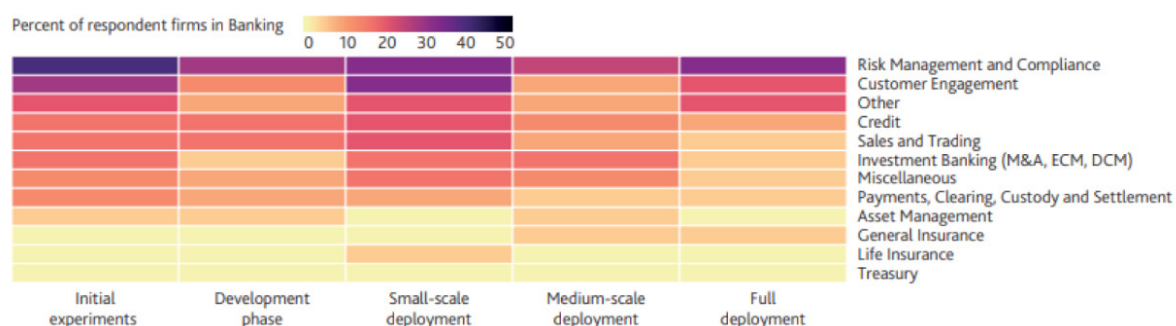
**Graphe 2. Deux tiers des personnes interrogées utilisent le machine learning en production**



Source: Machine Learning in UK Financial Services – Bank of England, FCA

Un autre élément développé dans le rapport de la BOE et de la FCA : la maturité du machine learning par type de business dans le secteur bancaire. Le stade de développement du « machine learning » est le plus abouti dans les fonctions de « risk management » et de « compliance » : détection de fraude ou de blanchiment d'argent font partie des applications les plus courantes. Les banques sont en phase de développement et de recherche sur des thématiques comme le risque de crédit, la vente, le trading, l'exécution, le pricing. Et elles n'ont pas encore commencé à regarder le thème dans les fonctions Asset and Liability Management (ALM) et Assurance-vie.

**Graphe 3. Maturité du Machine Learning, par business, dans le secteur bancaire**



Source: Machine Learning in UK Financial Services – Bank of England, FCA

## 2 STATISTICIENS VS. DATA SCIENTISTS

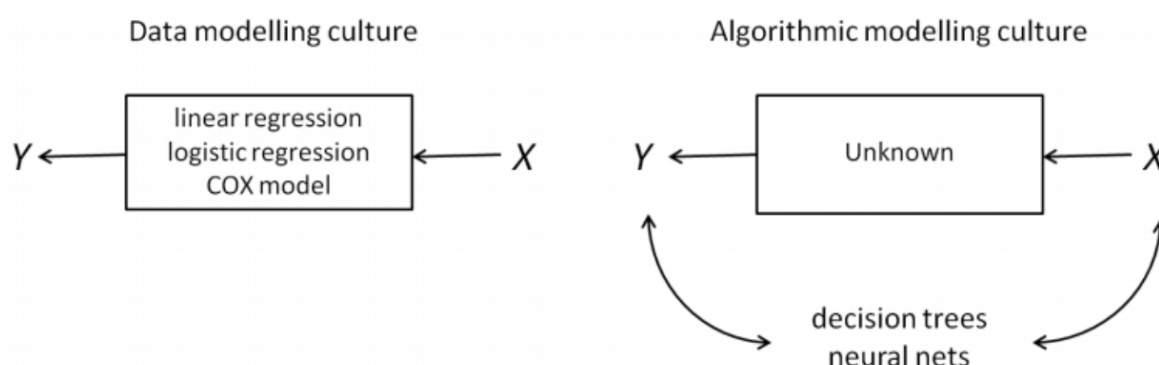
Alors qu'est-ce que le machine learning? Dans quel cadre s'inscrit-on lorsque l'on commence à parler de Machine Learning?

Revenons sur l'analyse de Leo Breiman dans « Statistical Modeling: Two Cultures », publié en 2012 dans Statistical Science (Vol 16, No. 3 199-231).

Dans une approche traditionnelle, de type économétrique, l'accent est mis sur le modèle et ses hypothèses. La connaissance du modélisateur, son intuition de marché et sa connaissance de la macro-économie est mis au centre de la démarche. Le modèle est validé s'il est capable de comprendre la donnée sur laquelle il est calibré. Il est en ligne avec l'intuition du modélisateur, et l'interprétation est facile.

Dans une approche de type machine learning, l'accent est mis sur la performance et le modèle devient une inconnue. Il est choisi en fonction de sa capacité à comprendre la donnée sur laquelle il est calibré, mais aussi en fonction de sa performance sur un jeu de données qui n'appartiennent pas au jeu de calibration. Le choix du modèle est le résultat d'un processus clé dans le Machine Learning, qui s'appelle la cross validation.

**Graphe 4. Le choc des cultures**



**Source : Statistical Modeling: Two Cultures, Leo Breimann**

Entre les deux approches, existe un compromis classique : interprétabilité vs. flexibilité. D'un côté du spectre se trouvent les approches classiques de type macro-économique, avec des modèles bien connus de tous, régression linéaire, le modèle de Black & Scholes pour les modèles de pricing.

De l'autre se trouvent les data scientists, qui développent des modèles très flexibles mais peu interprétables. Au milieu du spectre, en finance, se situent les analystes quantitatifs et ingénieurs financiers, qui essaient d'apporter la connaissance du domaine aux uns, et de faire progresser les techniques utilisées par les autres.

### LA PLUPART DES ALGORITHMES SONT ACCESSIBLES GRATUITEMENT SUR INTERNET

Voilà pour les grands principes. En pratique, le machine-learning, c'est tout un univers d'algorithmes possibles : deep learning, ensemble learning, neural networks, regularization methods, rule-based systems, regression, Decision tree, clustering... La liste est longue.

Graphe 5. La boîte à outil du data scientist



Source : [MachineLearningMastery.com](https://machinelearningmastery.com)

Au cours des années récentes, non seulement la boîte à outils disponible s'est développée, mais elle est également devenue plus facilement accessible. Aujourd'hui n'importe qui peut aller sur internet et se former à l'utilisation de ces différents algorithmes. L'accessibilité à l'algorithme est devenue extrêmement facile, ouvrant de-facto tout un champ de possibilités pour explorer de nouvelles voies.

### 3 UNE INTRODUCTION RAPIDE AU DEEP LEARNING

Le sujet qui nous réunit aujourd'hui est plus spécifiquement celui du deep learning. Le deep learning est une petite branche de l'univers des algorithmes possibles. Cependant, le deep learning devient de plus en plus populaire, grâce notamment au succès rencontré dans les applications de la vie de tous les jours, comme la reconnaissance d'image.

Le deep-learning est une technique déjà ancienne, qui se base sur des réseaux de neurones. Les réseaux de neurones étaient à la mode dans les années 90, mais ont été abandonnés la décennie suivante. Au début des années 2010, il y a eu un point de bascule. En 2009, aux États-Unis, les universités de Princeton et de Stanford lancent le projet ImageNet. Le but de ce projet est de collecter une base de données d'images publiques. Tous les ans, un concours est organisé. Ce concours consiste à essayer de labéliser de manière la plus précise possible une collection d'images.

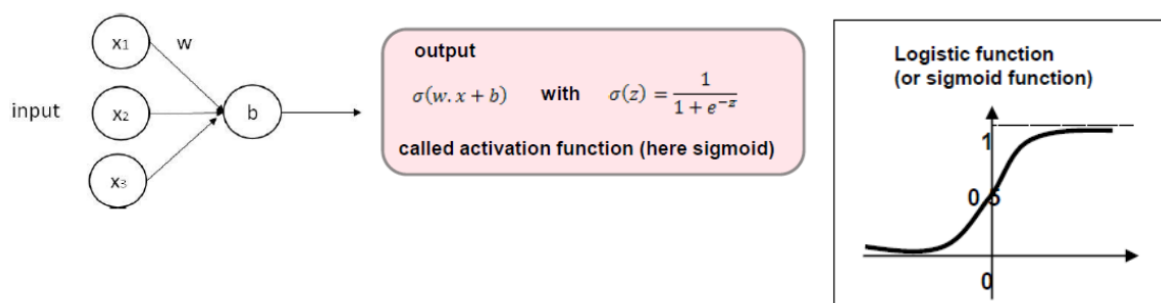
#### LE PROJET IMAGENET A DÉMOCRATISÉ L'UTILISATION DU DEEP LERNING

Entre 2009 et 2011, le taux d'erreur de classification est d'environ 25%. Un taux d'erreur de classification encore très élevé. En 2012, une équipe décide de tester et d'utiliser des réseaux de neurones. Leurs travaux ont permis de réduire énormément l'erreur de classification qui passe de 25% à 15%. Les années suivantes, les chercheurs se concentrent sur des algorithmes à base de réseaux de neurones et continuent à améliorer le taux d'erreur de classification. Depuis 2016, le taux d'erreur est inférieur à 5%, ce qui correspond au taux d'erreur moyen d'un opérateur humain. C'est l'un des premiers cas, où on a vu la machine faire mieux que l'homme dans une tâche essentiellement intellectuelle. À partir de ce moment, les équipes de recherche du monde entier et dans tous les domaines se sont mises à déployer ces techniques pour d'autres applications.

Qu'est-ce que le deep learning? Le deep-learning ou apprentissage profond est une technique à base de réseaux de neurones. Les réseaux de neurones existent depuis longtemps. L'idée du deep-learning est d'empiler des couches de neurones. C'est pour cette raison qu'on appelle ça apprentissage profond, car on a plusieurs couches de réseaux de neurones. Les réseaux de neurones les plus compliqués peuvent avoir des dizaines de couches. Mais un réseau à quatre couches est déjà une architecture très profonde.

Un neurone est un objet mathématique relativement simple similaire à une porte logique. Dans sa forme la plus simple, un neurone prend deux valeurs, zéro ou un. Le neurone sigmoïde est le plus utilisé en pratique. Une sigmoïde est une fonction qui varie entre zéro et un. Cette fonction prend un certain nombre de variables en entrée. Chaque variable d'entrée est pondérée par un poids,  $w$ . On ajoute parfois une constante appelée biais. La valeur de sortie du neurone varie entre 0 et 1, comme illustré sur le graphique ci-dessous.

Graphe 6. Qu'est-ce qu'un neurone ?

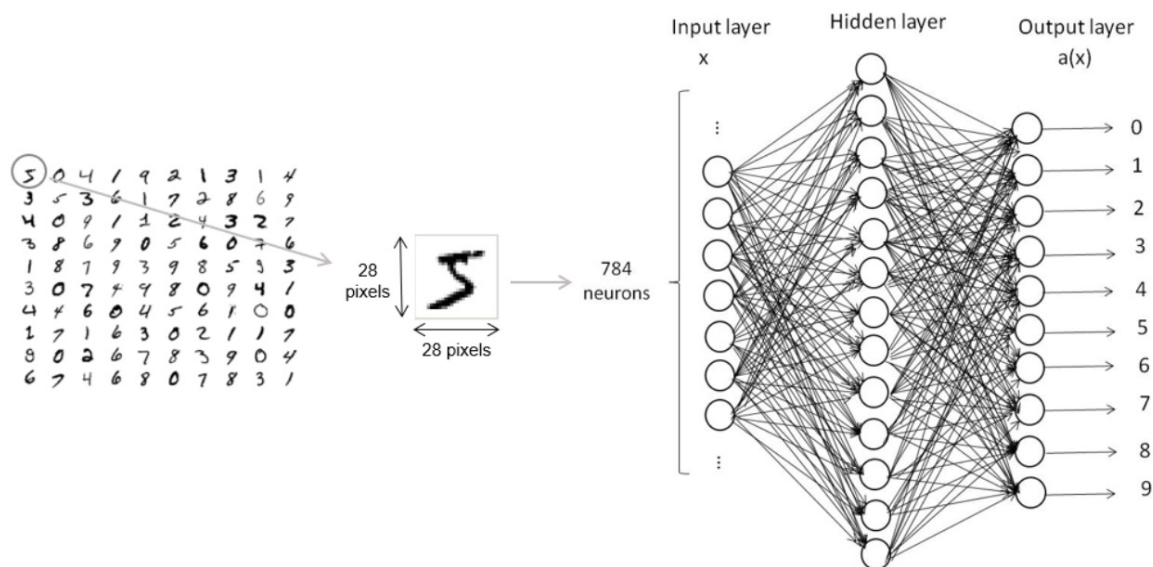


Source : SG Cross Asset Quant Research

Un réseau de neurones est une combinaison de neurones. Les neurones sont reliés entre eux, la sortie d'un neurone, devient l'entrée d'un autre neurone.

Un exemple classique et relativement simple est la reconnaissance de chiffres. Un chiffre manuscrit peut se modéliser par un carré de 28 pixels par 28 pixels par exemple. Un chiffre est donc décrit par un vecteur de 284 éléments contenant des zéros ou des uns : 0 si le pixel est activé, 1 sinon. On met les 284 éléments de notre vecteur en entrée du réseau, chaque élément devient un neurone. En sortie du réseau, on positionne dix éléments qui valent zéro ou un, en fonction de la donnée en entrée. A partir d'une banque d'image labellisée, on peut construire une base de données de vecteurs associés avec un label, qui deviennent les variables d'entrée sur lesquelles le réseau va être calibré.

Graphe 7. Un réseau simple de reconnaissance d'image



Source : SG Cross Asset Quant Research

L'enjeu du réseau de neurones est de trouver quels sont les poids et les biais de la fonction d'activation des neurones du réseau. Dans le cas présent, si on spécifie un réseau avec une seule couche interne à 10 neurones, le nombre de poids à trouver est large :  $784 \times 10$  poids pour transiter de la couche d'entrée à la couche interne, puis  $10 \times 10$  poids pour transiter de la couche interne à la couche de sortie. Ce qui fait un total de 7940 poids à trouver. Les poids sont optimisés pour minimiser une certaine fonction de coût qui est l'erreur de classification du réseau. C'est un problème d'optimisation complexe de grande dimension.

## LE DEEP LEARNING NÉCESSITE DE RÉSOUDRE DES PROBLÈMES D'OPTIMISATION COMPLEXES

Des algorithmes d'optimisation ont été développés pour résoudre ce problème, comme les algorithmes de back-propagation, combinés avec des gradients de descente. Mais il faut également énormément de données. Un tel modèle doit être calibré sur les dizaines de milliers d'images pour atteindre une convergence satisfaisante.

Cet exemple est un exemple simple. En réalité, à chaque type de problème correspond une architecture différente. Par exemple, les réseaux de neurones à convolution sont efficaces pour résoudre les problèmes de reconnaissance d'image. Les réseaux récurrents sont utilisés pour résoudre la modélisation des séries temporelles. Dans un réseau récurrent, l'état d'un neurone dépend non seulement de l'état des neurones qui l'entourent mais aussi de son état précédent. Les auto-encodeurs sont des réseaux relativement peu profonds et avec peu de neurones sur les couches internes. Ils servent à faire de la compression de données.

Le travail d'un « data scientist » aujourd'hui consiste à comprendre un problème et de trouver quelle est l'architecture neuronale adaptée à ce problème-là. Tout ceci reste une démarche très empirique. Il n'y a pas de recettes magiques. Il faut faire pas mal d'essais et de tests pour comprendre quel est l'architecture optimale, et quels sont les méta-paramètres associés.

L'utilisation du deep learning se répand de plus en plus en finance, avec une multiplication des cas d'application : détection des fraudes, détection de blanchiment d'argent, prédiction des risques, analyse plus fine du risque de crédit à l'aide de données alternative, comme les données de texte. Le deep learning se prête plutôt bien au text mining, qui est analyse automatique de texte. Un autre exemple est l'analyse des communications des banques centrales pour quantifier la probabilité d'un changement de politique monétaire.

## 4 DEEP LEARNING POUR LA VALORISATION ET L'ANALYSE DES RISQUES

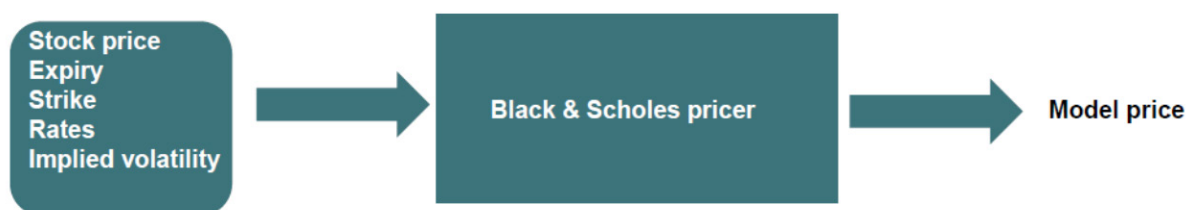
Sur les années très récentes, de plus en plus de chercheurs essayent d'appliquer ces approches pour la valorisation et l'analyse de risque des produits dérivés.

Utiliser le deep-learning pour la valorisation des dérivés n'est pas une idée nouvelle. Le premier papier sur le sujet a été publié en 1994 dans The Journal of Finance par J.M. Hutchinson, A.W. Lo and T.A. Poggio (A Nonparametric Approach to Pricing and Hedging Derivative Securities Via Learning Networks). Les auteurs y développent l'idée d'utiliser une approche neuronale pour valoriser des contrats dérivés.

Cette approche n'a pas été poursuivie à l'époque. Elle revient sur le devant de la scène aujourd'hui, grâce aux progrès en optimisation et grâce à la puissance de calcul des ordinateurs modernes.

Dans une approche classique, le modèle prend une série de paramètres en entrée : le prix du sous-jacent, le strike, la maturité, les taux, la volatilité implicite. On applique une formule, la formule de Black Scholes par exemple, et on obtient un prix.

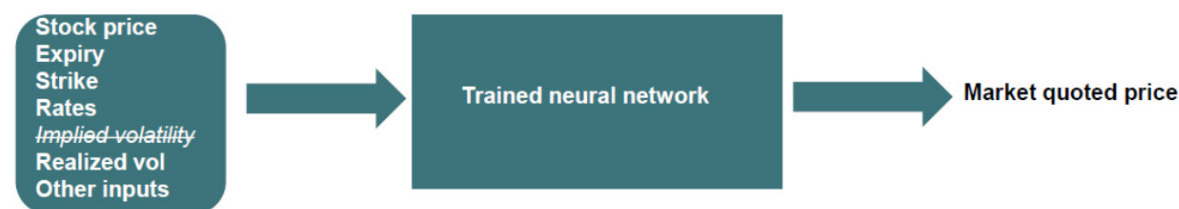
**Grappe 8. Valorisation des dérivés dans un model classique**



Source : SG Cross Asset Quant Research

Avec les réseaux de neurones, on prend un maximum de données en entrée, historiques ou simulées. Et on calibre le réseau de sorte à minimiser l'écart entre les prix qui sortent du réseau et les prix de marché. Le modèle est non-paramétrique et devient le résultat d'une calibration.

**Grappe 9. Valorisation des dérivés avec un réseau de neurones**

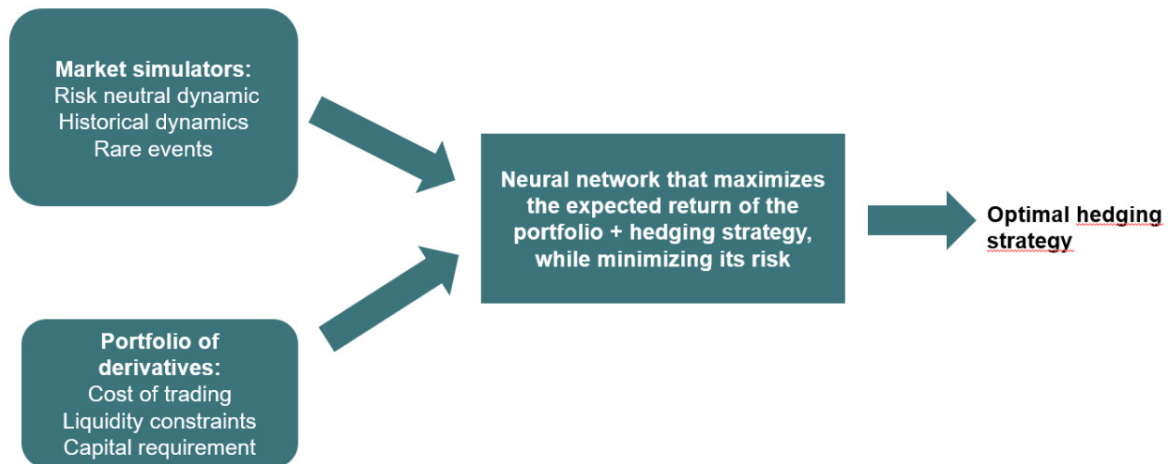


Source : SG Cross Asset Quant Research

Le deep-hedging reprend la même idée. C'est une approche qui a été développée et proposée par un groupe de chercheurs à ETH Zurich et JP.Morgan (Deep Hedging, the Journal of Quantitative Finance, Jan 2019, Hans Bühler, Lukas Gonon, Josef Teichmann, Ben Wood). Dans cette approche, on simule plein des données de marchés et on optimise la stratégie de hedging.

Plutôt que de calibrer le réseau sur le prix d'un instrument, on calibre le réseau de telle sorte que l'on minimise la variance du portefeuille qui contient les instruments de hedge et on laisse ce réseau décider de la stratégie de hedging optimale.

**Graph 10. Gestion du risque de couverture avec le deep pricing**



La théorie et la démarche derrière le concept est beaucoup plus compliquée. Pour l'instant ce genre de technique se base essentiellement sur des simulations. Les données de marchés disponibles ne sont pas suffisantes pour calibrer des réseaux de neurones. Il faut donc générer des données artificielles. Il y a toute une branche de recherche, populaire en ce moment, qui consiste à étudier comment simuler de la donnée réaliste, et qui colle à la réalité des marchés financiers.

## 5 LES CHALLENGES DU DEEP LEARNING

Le dernier point que je voulais aborder ce sont les contraintes et les challenges liés à ce type d'approches. Dans l'étude publiée par la FCA et la Bank of England, on voit émerger plusieurs questions sur la mise en place de ces techniques.

Donc la première question est la gestion des systèmes historiquement utilisés. Ce n'est pas un problème propre au deep learning mais un problème qui se pose dès que l'on cherche à faire quelque chose de nouveau dans une industrie comme la banque ou l'assurance. Le passage d'un ancien système à une technologie plus moderne est toujours compliqué à gérer.

Plus spécifique au Machine Learning (ML), il y a l'intégration de ces techniques dans les processus utilisés pour faire tourner le business. Cette question est liée à la formation nécessaire pour amener les utilisateurs à utiliser ces nouvelles techniques. Il faut que les utilisateurs soient familiers avec les concepts abordés au début de cette présentation. C'est tout un changement culturel à faire au sein des organisations.

### LE MANQUE D'EXPLICABILITÉ EST UN ÉCUEIL MAJEUR A LA MISE EN PLACE DU ML

Un autre point qui souvent mentionné est le manque d'explicabilité. C'est un problème majeur de l'approche deep learning. Dans une approche non-paramétrique, on ne peut pas changer un paramètre et voir comment le résultat du modèle est impacté. On manque d'intuition sur comment le modèle réagit et sur l'origine des résultats. Par exemple, Deep Mind a créé un algorithme capable de jouer au Go, AlphaGo. Lors de la partie où AlphaGo a battu pour la première fois un joueur de Go professionnel, l'ordinateur a sorti un coup que personne n'a compris ou n'a pu expliquer sur le moment. A posteriori, c'était un coup de génie.

**Graph 11. Les contraintes et les enjeux du Machine Learning en pratique**



Source: Machine Learning in UK Financial Services – Bank of England, FCA

Avec le deep learning, on peut avoir des résultats complètement inattendus. En tant qu'investisseur, en tant que trader, en tant que risk manager, on ne sait pas si c'est un coup de génie, ou juste un problème numérique. En pratique, ce genre de comportement peut devenir problématique.

L'insuffisance des données est également un thème qui est régulièrement mentionné. Le Machine Learning est une approche vorace en données. Or nous ne disposons que d'une dizaine d'années de données sur la plupart des marchés financiers. Ça peut être un problème dans la mise en place de ce genre d'algorithmes.

Lié à cela, d'autres questions se posent sur l'utilisation de données alternatives : quelle doit être la régulation autour de l'utilisation des données personnelles, quels sont les processus de gouvernance sur la donnée, comment la gère-t-on, comment s'assure-t-on de sa fiabilité.

### COMMENT S'ASSURER QU'UNE DONNÉE EST FIABLE ?

C'est une bonne question. On ne peut pas avoir une garantie à 100% que la donnée soit fiable. Par contre, on peut mettre en place de vérifications pour avoir la meilleure représentation possible de la réalité. Par exemple si on prend des données trading, on peut observer les prix traités, et il faut s'assurer qu'il n'y a pas une erreur dans les bases de données.

Voici donc une brève introduction au machine learning, au deep learning, et au deep hedging, qui soulève déjà bon nombre de questions.

## Charles-Albert Lehalle

Bonjour, je m'appelle Charles-Albert Lehalle et je suis responsable, à Capital Fund Management (CFM), d'une équipe qui s'appelle "Data Analytics". CFM est un fond d'investissement de type "hedge fund", qui gère environ 10 milliards de dollars. Cette équipe est en charge de l'intégration de "données alternatives" : ce sont par exemple les images satellites, les données de carte de crédit, les transcriptions de nombreux textes concernant les entreprises et l'économie. La plupart du temps, les outils nécessaires à l'analyse de ces données sont évidemment les statistiques, mais surtout le "machine learning" car ce sont des données en grande dimension. Pour en extraire de l'information, il faut passer par des outils non linéaires. Je suis donc amené à m'interroger sur le sens de ces outils, et la confiance que l'on peut avoir dans les résultats qu'on obtient en les utilisant. Je cherche aussi à comprendre comment et pourquoi ils fonctionnent, sans doute parce que j'ai fait ma thèse sur ce sujet il y a une vingtaine d'années.

## 6 DES TECHNOLOGIES DE LA DÉSINTERMÉDIATION

Mais avant d'entrer dans le détail, je pense que cela vaut la peine de s'interroger de façon plus globale et macroscopique à propos de l'impact de ces technologies sur le monde financier, ou plutôt sur l'écosystème financier.

### L'INFLUENCE DES TECHNOLOGIES DE LA DÉSINTERMÉDIATION EST LOIN D'ÊTRE ANECDOTIQUE.

Est-ce que ça va fondamentalement changer la façon dont on travaille dans l'univers financier ou est-ce que c'est juste un saut technologique de plus ? Il est assez facile de se convaincre qualitativement que leur influence est loin d'être anecdotique. Il suffit de regarder la façon dont ces technologies ont déjà transformé certains secteurs, comme celui des médias ou des transports. Prenons par exemple les médias : quand j'étais jeune et que j'achetais le journal Le Monde, je payais pour la totalité des articles du Monde, alors que je n'en lisais que 1/10ème. Rapidement, les moteurs de recherche ont permis de faire du "cherry picking" : ne consommer que ce qui nous intéresse. A travers votre question et votre historique de lecture sur internet, ils vous proposent de ne lire qu'un sous ensemble des articles qui sont produits par un journal. Cela modifie drastiquement la façon de produire de l'information : les responsables du journal ne peuvent plus utiliser la marge qu'ils font sur des nouvelles faciles à produire pour financer des reportage au long cours ; ils ne peuvent plus répartir eux même les coûts de production sur plusieurs produits en imposant un menu. Car le consommateur ne va payer que pour ce qu'il veut. C'est le genre de désintermédiation que produisent ces technologies : un producteur de services a beaucoup moins de leviers pour imposer la consommation d'une gamme de service, pour être un "one stop shop".

### DANS SES INTERACTIONS AVEC GOOGLEBRAIN, LE MOTEUR DE RECHERCHE GOOGLE, L'HUMAIN S'EST ADAPTÉ À LA MACHINE.

Certains d'entre vous s'interrogent peut-être sur le rapport avec le deep learning, l'apprentissage profond ? Il faut que vous vous rendiez compte que le meilleur système de machine-learning, que nous utilisons tous depuis longtemps, est le moteur de Google search. C'est un système de NLP, de Natural Language Processing, qui utilise essentiellement de l'apprentissage profond, et qu'il faut considérer comme une instance très particulière de moteur conversationnel. Le moteur de recherche de google, c'est un énorme réseau de neurones qui passe son temps à essayer d'analyser toutes les données qui ont été capturées par Google pour pouvoir répondre à des questions. Ce qui est intéressant est que les interactions avec Google Search sont "transformative" au sens où l'humain s'est adapté au système : la façon dont la plupart d'entre nous posons une question à Google Search, n'est pas la façon dont nous la poserions à un être humain. Néanmoins vous pouvez aussi vraiment poser une question avec un point d'interrogation à Google Search, ça marchera aussi. Cette interaction est passionnante : en ce qui concerne la façon de poser les questions, l'humain s'est adapté à la machine. Nous avons, de façon collaborative, créé un nouveau dialecte qui n'est pas le langage humain naturel, afin d'obtenir des réponses plus précises à nos questions. Google Search et chacun d'entre nous parlons ce dialecte. Néanmoins Google Search comprend le langage humain, afin de digérer les pages internet, les livres, les vidéos et autres podcasts qu'il stocke et auxquels il adapte les paramètres de ses réseaux de neurones profonds, afin de nous en restituer l'essentiel.

Donc tout ça pour dire que les industries qui ont été révolutionnées par ces nouvelles technologies, comme les médias et les transports, l'ont été en se faisant désintermédier. Les transports par Uber, les chaînes de télévision pour Youtube puis Netflix, les journaux par Google News. Tout ceci en s'appuyant sur une bonne connaissance des habitudes des consommateurs, souvent acquises grâce aux données récoltées par réseaux sociaux.

Comme cela peut-il modifier le fonctionnement des participants de marché, aujourd'hui nous nous intéressons plutôt aux assurances et aux banques d'investissement ? Pour le comprendre il faut commencer par reconnaître que ces entreprises proposent encore pour beaucoup un menu complet à leurs clients. Les banques d'investissement, que je connais mieux que les assurances, proposent souvent des services d'exécution, de l'analyse financière, des produits dérivés et structurés, des prêts, etc. En bref des services qui tournent autour de l'intermédiation de bilan et de l'intermédiation de marché. Si possible elles factorisent leurs marges en les organisant autour d'un service global par client : elles négocient lentement un "menu de services" par gamme ou par type de client, et ajustent ensuite trimestre après trimestre une marge globale pour toute cette collection de services. Il me paraît clair que les technologies de la désintermédiation vont bousculer ce genre de "business model", car les clients vont faire de plus en plus de "cherry picking", en ne consommant qu'un service marginalement plus rentable pour eux chez une banque et un autre chez une autre.

## 7 L'INTELLIGENCE ARTIFICIELLE : UNE "TECHNOLOGIE GÉNÉRIQUE"

Une fois que l'on a posé le débat en terme de désintermédiation des services financiers, on peut s'interroger sur la nature de ce qu'on a envie d'appeler "les technologies de l'Intelligence Artificielle", même si on peut débattre des heures sur s'il faut vraiment les appeler ainsi. Force est de constater que tout le monde utilise cette dénomination aujourd'hui, je vais donc le faire aussi. Quelle est la particularité de l'IA lorsqu'il s'agit de la mettre en oeuvre ? Je pense que c'est parce qu'elle appartient à ce que certains économistes appellent des Technologies Génériques.

Je vais vous donner d'autres exemples de technologies génériques : la machine à vapeur et l'électricité. Ce n'est pas parce que vous avez utilisé une machine à vapeur pour faire avancer une locomotive que c'est immédiat d'utiliser une machine à vapeur pour faire fonctionner un métier à tisser. Il faut re-développer ce qu'on appelle des innovations secondaires. L'intelligence artificielle, c'est pareil : ça a beau marcher très bien sur en traitement du signal pour reconnaître des images ou du son, si vous voulez l'utiliser dans une banque il faut refaire plein de choses choses, ré-adapter l'IA à chaque gamme de problème. Cela exige des investissements.

Par conséquent il faut que l'écosystème s'organise, afin de mettre en contact les experts génériques de cette technologie, et des experts spécifiques, pendant des petits projets à court-terme, pour arriver à fabriquer des innovations secondaires qui vont répondre à une industrie ou entreprise en particulier.

### **LE DÉPLOIEMENT DE L'IA DANS LA BANQUE ET L'ASSURANCE NE PEUT PAS SE FAIRE SANS INVESTISSEMENT ET NÉCESSITERA BON NOMBRE D'INNOVATIONS SECONDAIRES.**

La distinction entre les innovations secondaires et la technologie générique, et surtout le fait qu'il faille investir pour aboutir à des innovations secondaires crée une distorsion compliquée. Puisque d'un côté je peux lire chaque jour dans le journal que l'IA c'est génial, que ça marche pour des tas d'applications. Mais d'un autre côté quand je veux l'appliquer sans investir du temps et de l'argent dans un projet bien ficelé ça ne fonctionne pas sur les problèmes qui m'intéressent. Aujourd'hui le rapport Villani a bien exposé cela et fait en sorte qu'il y aura bientôt en France des centres d'excellence qui vont mettre en contact des experts génériques, en IA, en big data, en data sciences et des besoins spécifiques, avec leurs experts métier : les instituts interdisciplinaires d'intelligence artificielle (3IA). Typiquement en région parisienne un de ces 3IA s'appelle PRAIRIE. Les entreprises vont pouvoir faire des partenariats, avec des durées de un à trois ans, pour essayer de fabriquer des innovations secondaires. Les grandes universités anglaises l'ont fait de leur côté avec l'Institut Turing créé en 2015, et ça semble bien fonctionner. Aux Etats Unis chaque grande université américaine a créé un "petit" centre de cette nature, comme le Center for Data Sciences de NYU, créé en 2013 par Yann LeCun.

### **LES CENTRES D'EXCELLENCE UNIVERSITAIRE SONT MIS EN PLACE EN FRANCE POUR METTRE EN CONTACT EXPERTS GÉNÉRIQUES EN IA ET EXPERTS MÉTIERS.**

Je voulais souligner, avant d'entrer dans le vif du sujet, cette nécessité d'investir, sans doute en collaboration avec des universitaires, afin d'obtenir les innovations secondaires nécessaires à bien faire fonctionner l'IA en finance de marché.

## 8 L'IMPACT DES TECHNOLOGIES DE L'IA EN FINANCE DE MARCHÉ

Maintenant nous pouvons nous demander : comment tout cela peut-il toucher l'industrie financière ? On peut dire que les innovations secondaires vont avoir, ou ont déjà lieu dans trois grandes directions : la distribution et la personnalisation des produits financiers, l'extraction d'information de données non financières, et la gestion et l'intermédiation du risque. Je vais rapidement les passer en revue.

### **LES INNOVATIONS SECONDAIRES VONT DANS TROIS GRANDES DIRECTIONS : LA DISTRIBUTION ET LA PERSONNALISATION DES PRODUITS FINANCIERS, L'EXTRACTION D'INFORMATION DE DONNÉES NON FINANCIÈRES, ET LA GESTION ET L'INTERMÉDIATION DU RISQUE.**

La partie "distribution et personnalisation des produits financiers" est le domaine qui est plus proche de ce que l'on voit se produire dans d'autres secteurs. Le mot à la mode c'est personnalisation de l'expérience client. Fabriquer des produits financiers à la demande, comme des ETF, des swaps ou de produits structurés, pour chaque client, c'est quelque chose qui est en train de se produire et ces technologies vont permettre de le faire à partir des "habitudes de consommation" de chaque client : quand est-ce qu'il achète une exposition à ce genre de risque ? C'est une sorte de segmentation très fine des habitudes des clients, dans le monde de la consommation de produits financiers, c'est à dire de l'exposition à des facteurs de risques, et de l'échange de flux financiers.

Ensuite, il y a un axe sur lequel je travaille plus particulièrement en ce moment, et que j'ai envie d'appeler "une meilleure connection à l'économie réelle". Il s'agit d'injecter des données non financières dans le monde financier ; et je vous en donnerai un exemple tout à l'heure. Typiquement l'utilisation des images satellite, puisqu'on sait que l'IA obtient de bons résultats sur ce genre d'images. Il est possible de faire beaucoup d'autres choses, avec beaucoup de données hétérogènes, et finalement aboutir à des indicateurs économiques avancés, un peu comme les chiffres fournis par les agences gouvernementales chaque trimestre, sauf qu'il est possible d'en construire des estimateurs en temps réel, en y intégrant toutes les informations dont on dispose.

Dans cette catégorie il y a l'utilisation des textes, qui nécessite des réseaux de neurones profonds, je pense aux modèles de langage de la famille de BERT (qui veut dire "Bidirectional Encoder Representations from Transformers"). En ce qui concerne les applications en finance on peut les utiliser sur les "Earning Calls" des dirigeants d'entreprise listées aux Etats-Unis. Il s'agit de la présentation aux investisseurs des résultats et de la stratégie de l'entreprise, chaque trimestre. On peut analyser la syntaxe, le vocabulaire utilisé pour essayer de comprendre l'influence qu'ont ces commentaires de résultats sur la valeur de l'entreprise avant que les analystes financiers nous l'expliquent. Avec une IA, on peut traiter des centaines de textes très rapidement, pas forcément avec une précision très fine mais beaucoup plus rapidement qu'un humain, ou qu'une équipe d'humain pourrait le faire, et ainsi se faire une idée de la modification de la valeur relative d'un très grand nombre d'entreprises assez rapidement.

Voici un autre exemple : vous pouvez récupérer toutes les paiements par carte de crédit de 10% des américains, presque en temps réel. Vous disposez de la géolocalisation du paiement. Si c'est dans un magasin, vous connaissez le magasin, il appartient à une société cotée, et vous agrégez ça et vous avez une vision d'une partie de la comptabilité de ces entreprises. Vous voyez une partie des flux sortant : ce qui est vendu, mais pas l'état des stocks ni les flux entrants ; vous ne voyez pas les approvisionnements. Néanmoins pour certains secteurs d'activité, c'est très informatif.

Il est facile de se convaincre que toutes ces approches permettent de construire des indicateurs, qui correspondent à ceux qui sont fabriqués en France par exemple par l'INSEE, de façon périodique, mais qu'on peut avoir jour après jour. Par la suite on peut prendre des positions, que l'on espère gagnantes, sur les marchés financiers et par ce biais accélérer la diffusion de l'information dans les prix. Ce qui est intéressant est que l'on injecte ainsi des informations qui proviennent de monde réel, du monde économique, dans la valorisation des actifs sur les marchés financiers. On peut espérer que ces

approches vont réduire les éventuelles écarts entre des valorisations purement “financières” et la réalité de la valeur des entreprises.

Et le troisième point c’est d’améliorer l’intermédiation du risque, et c’est là que se situe typiquement le deep-hedging. Il s’agit de façon plus générale de mieux comprendre et de mieux couvrir les risques portés par un intermédiaire financier. Les technologies de l’IA devraient pouvoir améliorer ces aspects non seulement grâce au deep hedging, mais aussi en fabriquant des scénarios de stress test sur mesure et adaptatifs, et aussi en permettant aux techniques d’évaluation des risques de fonctionner en plus grande dimension.

## 9 LE RÔLE DES DONNÉES, ET LA COMPLEXITÉ DE L'AUTOMATISATION

La nouveauté des approches dont nous parlons aujourd'hui, les méthodes d'apprentissage profond, réside non seulement dans leur nature profondément non linéaire et qui semble pouvoir s'adapter à de nombreuses relations entre des variables que l'on ne connaît pas parfaitement, mais elle provient aussi du fait que les données qu'elles ingèrent influencent énormément les résultats obtenus. La mise en oeuvre d'un algorithme d'apprentissage nécessite

- Des bibliothèques de codes "externes", souvent fournies par Google, par exemple avec TensorFlow, ou Facebook, avec pyTorch,
- Du code développé pour l'application, qui précise par exemple le critère à minimiser, le choix de la technique d'optimisation et ses paramètres
- Une base de données qui va servir de référence à l'algorithme.

### COMMENT PROTÉGER LE RÉSULTAT D'UN ÉVENTUEL BIAIS DANS LES DONNÉES OU MÊME D'UNE MANIPULATION DES DONNÉES QUI VA CHANGER LE RÉSULTAT DE L'ALGORITHME ?

Une famille des questions que l'on peut se poser concerne le niveau de qualité qu'il faut exiger des données. Comment protéger le résultat d'un éventuel biais dans les données ou même d'une manipulation des données qui va changer le résultat de l'algorithme ? Il y a le sujet de la cyber-sécurité bien entendu, qui est d'ailleurs plutôt une problématique industrielle : comment sécuriser les accès aux multiples interfaces qui résultent de l'utilisation de composants distribués ? On peut penser à la cyber attaque des systèmes de paiement de la banque centrale mexicaine de mai 2018 : ce genre d'attaques sera-t-il facilité lorsque les participants de marché utiliseront des algorithmes qui se connectent beaucoup plus de fois à beaucoup plus d'intermédiaires ?

Mais ce n'est pas le sujet pour aujourd'hui, focalisons-nous sur les problématiques spécifiques à la mise en oeuvre des technologies de l'IA. Les difficultés peuvent provenir des plusieurs éléments :

- Les bibliothèques externes : comment mesurer de façon fiable l'impact d'un changement de version d'une des grosses bibliothèques d'apprentissage, comme TensorFlow ou pyTorch ? Il faut être capable de relancer des apprentissages avec la nouvelle version, sur les mêmes données que celles qui ont été utilisées avec l'ancienne version, et de comparer les "résultats". D'ailleurs que sont exactement les résultats : les paramètres issus de l'apprentissage (i.e. les "poids" des réseaux de neurones) ? Ou les prédictions obtenues par ces réseaux de neurones ? Sans doute les deux. Mais la nature aléatoire des apprentissages, car n'oublions pas que la plupart de ces méthodes reposent sur la descente de gradient stochastique, rend cette comparaison difficile.
- Les données elles-mêmes peuvent changer, être mises à jour, lorsqu'on y trouve une inexactitude, ou bien tout simplement on peut augmenter le volume de ces données. Il faut alors relancer des apprentissages et de nouveau comparer ces résultats, en recourant là encore à des critères statistiques. Il faut vérifier que les résultats ne sont pas "statistiquement différents".
- Avant même de parler d'une modification des données, il faut se pencher sur l'influence possible de "biais" dans les données. L'exemple classique est celui d'un biais homme - femme : si la base d'apprentissage n'est pas équilibrée, alors les résultats ne le seront pas non plus. Si on pense à l'attribution de prêts, il ne faut pas qu'un algorithme accepte plus facilement de prêter aux hommes qu'aux femmes, par exemple. En finance de marché il y a bien d'autres biais, dont celui très important des "données contributives" sur lequel je reviendrai par la suite.
- Bien entendu le code qui met en oeuvre la fonctionnalité spécifique à une application, très souvent développé en interne par l'institution financière, et qui comporte un critère à optimiser, appelé en apprentissage statistique la "fonction de perte", a une place centrale. Il joue un très grand rôle dans le résultat de l'apprentissage. Quelle fonction choisir ? Et en particulier vaut-il mieux choisir un critère de précision statistique, ou bien un "critère opérationnel", c'est-à-dire par exemple un gain en euros ? Faut-il attribuer une pénalité à la prise de risque, et comment la combiner avec le critère de gain ?

Les industries de l'assurance et de la finance de marché sont particulièrement régulées, à juste titre puisqu'elles peuvent avoir de très fortes externalités négatives. Il est donc nécessaire de trouver un cadre clair permettant d'encadrer l'utilisation de ces technologies.

# 10 FAUT-IL RÉGULER UNIQUEMENT LES TECHNOLOGIES ?

Une fois que l'on entame une réflexion vers la régulation des technologies de l'IA, et pour éviter de se lancer dans le sujet, un peu "tarte à la crème" de "l'éthique des algorithmes", il convient de prendre un peu de recul.

## **POURQUOI EXIGER DE CONTRÔLER LE BIAIS D'UN ALGORITHME ALORS QUE L'ON N'EXIGE PAS DE CONTRÔLE SUR LE BIAIS DES DÉCISIONS DE MÊME NATURE, MAIS PRISES PAR DES HUMAINS**

A titre personnel, je me rappelle des débats autour du trading haute fréquence, auxquels j'ai beaucoup contribué. Le débat n'était pas très différent à l'époque de celui d'aujourd'hui, puisque le trading haute fréquence recourt à des algorithmes de négociation automatiques, calibrés à l'aide de données, et qui ont une relative autonomie.

Une fois de plus, j'ai le sentiment que le fait de poser la problématique dans un cadre automatisé et technologique a la vertu de faire émerger les vraies questions. Lorsqu'on étudie un processus qui implique beaucoup de décisions humaines, on peut toujours penser : "comme il y a des humains dans la boucle, la description que l'on fait là est un peu approximative et jamais exactement respectée donc on ne va pas trop réfléchir à ses implications". Lorsqu'on se trouve face à des processus actionnés par des machines, on se pose les vraies questions de façon beaucoup plus claire et beaucoup plus concises. Il faut donc en profiter pour mieux réguler. Bien réguler les machines bien sûr, mais aussi mieux réguler les humains. Moi ce qui m'a beaucoup frappé dans le débat sur le trading haute fréquence, c'est qu'il y a eu à un moment des décisions du genre : "régulons les machines et essentiellement les machines". Mais non, la façon dont les humains faisaient les prix, en tant qu'intermédiaire, sur certains produits peu liquides étaient évidemment à réguler autant que la façon dont les machines font des prix de façon systématique, même si c'est souvent sur des produits plus liquides.

## **IL FAUT AUSSI SUPERVISER LES BIAIS QUI PROVIENNENT DES DÉCISIONS HUMAINES**

Pour en revenir aux biais, il faut aussi superviser les biais qui proviennent des décisions humaines. Il faut organiser des "revues de prises de décision", qui à partir de critères statistiques vont permettre d'établir si les décisions prises le mois dernier par l'institution sont biaisées, et si oui il lui faudra justifier pourquoi. Et si ce n'est pas justifié elle devra démontrer qu'elle prend des mesures pour éviter que cela se reproduise. Que ces décisions soient prises par des machines ou par des humains.

Je prend donc comme une chance que l'automatisation et l'utilisation des technologies de l'IA permettent de poser ces questions. Et je souhaite que les conclusions qui seront issues des débats permettent d'accoucher de régulations qui s'appliqueront aussi bien aux processus gérés par des humains que par des machines. Et ça c'est vraiment une richesse, même si cela soulève plusieurs problèmes, mais je pense qu'ils seront mieux posés.

# 11 LE CAS DES DONNÉES DE MARCHÉ

Les deux cas importants qu'il faut garder à l'esprit si on pense à l'influence du biais sur les algorithmes obtenus par apprentissage en finance de marché sont : la non stationnarité du processus de formation des prix et les prix contribués.

## LA NON STATIONNARITÉ DU PROCESSUS DE FORMATION DES PRIX ET LES PRIX CONTRIBUÉS INFLUENCENT LE BIAIS SUR DES ALGORITHMES OBTENUS PAR APPRENTISSAGE

Le processus de formation des prix est un terme générique utilisé pour décrire la dynamique jointe des prix (ou des rendements) ; il souligne que les prix se "forment", ils ne sont pas le fruit d'un modèle abstrait, mais résulte du mélange, en temps réel, d'informations exogènes qui s'incorporent aux prix (par exemple des dépêches de presse, ou des communication des entreprises ou des gouvernements) et d'un processus endogène (les actions des participants de marché dépendent elles-mêmes des prix). Ce processus est influencé par l'état de l'économie, et a une certaine inertie ne serait-ce qu'à cause des inventaires au sens large des entreprises, mais aussi à cause des cycles économiques ou des changements géopolitiques. Ceci engendre a priori une très forte non stationnarité de la formation des prix. Dans ce cas comment penser que des systèmes appris sur des prix ou des rendements puissent rester valides quelques mois ou années plus tard. Un exemple simple connu depuis longtemps est que durant les crises, les corrélations entre actifs se modifient : leur valeur absolue augmente.

Pour zoomer sur le deep hedging : si les processus utilisés pour "apprendre" le hedging optimal, c'est-à-dire la stratégie de réplcation du risque optimale, sachant les conditions de marché récentes, qui sait si leur comportement sera valide une fois que les conditions de marché auront changées ? Cela veut dire qu'on ne peut se passer, lorsqu'on a "appris" sur des données, de faire une analyse critique des données, de leurs biais, et ainsi développer un point de vue comportemental sur ces données. En fait : choisir des données est équivalent à faire un choix de modèle, qu'il est nécessaire de rendre explicite, à travers des analyses statistiques poussées.

Un autre cas intéressant, typique pour les banques d'investissements, est l'utilisation de données contribuées. Il s'agit de données qui sont fournies par d'autres banques d'investissement sur une base déclarative, ou en tout cas sans que personne ne puisse affirmer qu'il s'agit de prix auxquels des transactions ont véritablement eu lieu. Apprendre sur ces données incorporent nécessairement un biais, qu'il est pour le coup extrêmement difficile de qualifier.

## 12 LE CAS SPÉCIFIQUE DU DEEP HEDGING

Il est temps de discuter du cas du deep hedging. Le deep hedging est un moyen de mettre au point des stratégies de réplcation du risque d'une position incertaines qui s'inspirent de l'apprentissage par renforcement, en anglais le "reinforcement learning". Rappelons le principe de la réplcation du risque : si j'ai une position qui a des expositions non linéaires à des processus stochastiques, c'est-à-dire à des dynamiques incertaines de variables économiques comme la volatilité, les prix, des taux ou des corrélations, combien dois-je acheter ou vendre chaque jour de produits négociables afin de réduire cette exposition ? En fin de course, lorsque le portefeuille à couvrir arrive à expiration, si ma stratégie est bien faite je dispose exactement de la somme à livrer à mon client, ou je suis à découvert de la somme qu'il va me payer. Ces stratégies de réplcation sont très importantes pour les banques d'investissement car non seulement elles permettent de diminuer l'exposition aux risques économiques, mais en plus elles fournissent une estimation au coût moyen de la couverture tout le long de la vie du produit, et donc permettent de lui attribuer un prix.

### **DANS UNE TECHNIQUE DE HEDGING "TRADITIONNEL": LA DYNAMIQUE DES PRIX EST SIMULÉE À PARTIR DE GRANDEURS IMPLICITÉES**

Si on regarde le principe de fonctionnement du hedging "traditionnel": la dynamique des prix est simulée à partir de grandeurs implicites. De nombreuses trajectoires sont générées et la combinatoire d'une très grande partie des stratégies de couverture possible est explorée et confrontées aux simulations, seules les stratégies les plus efficaces sont retenues et forment les composantes trajectorielles de la stratégie dite optimale. Il faut préciser que le procédé d'implication revient à ne pas utiliser une base de données historiques de prix de marchés des instruments négociables, mais de n'utiliser que les prix actuels des produits dérivés, c'est-à-dire de produits de même nature que celui couvert, qui donnent eux aussi lieu à la mise en oeuvre de stratégies de couverture. On dit que l'on n'utilise pas la mesure historique, qui est celle de la base de données qu'utiliserait un algorithme d'apprentissage, mais la mesure risque neutre. Cette mesure a beaucoup de propriétés sympathiques, celle que je préfère est qu'elle assure que si, pendant la vie du produit réplqué, il existe d'autres produits de même nature disponibles à l'achat ou à la vente, alors la banque pourra "revendre son risque" et n'aura plus aucune exposition d'aucune sorte. Le prix qui aura été ainsi calculé en début de vie du produit sera plutôt juste. Ceci serait moins le cas sous probabilité historique. En utilisant cette mesure risque neutre pour générer des trajectoires permettant d'évaluer la performance de stratégies de couverture, la banque remplit ainsi parfaitement son rôle d'intermédiaire financier. Elle se met temporairement entre plusieurs acheteurs et vendeurs de différents risques et ne prend pas de "risque historique", comme un fond d'investissement pourrait le faire.

### **AVEC UN RÉSEAU DE NEURONES PROFOND, ON PEUT GÉNÉRER MOINS DE STRATÉGIES DE RÉPLICATION, POUR SE CONCENTRER SUR CELLES QUI ONT L'AIR DE MIEUX FONCTIONNER.**

Une fois que l'on comprend bien ce principe de générer des trajectoires afin d'évaluer la performance de stratégies potentielles, on peut utiliser des techniques d'apprentissage pour l'améliorer d'une première façon : lorsqu'on parcourt de façon systématique presque toutes les stratégies de réplcation possibles, c'est très coûteux informatiquement, surtout quand on augmente la dimension du problème, c'est-à-dire quand on veut couvrir des portefeuilles avec beaucoup d'actifs et de nombreuses non linéarités. Avec un réseau de neurones profond, un "deep neural network", on peut générer moins de stratégies de réplcation, en faisant confiance aux paramètres du réseau de neurones pour se concentrer rapidement sur celles qui ont l'air de mieux fonctionner, et ainsi converger plus rapidement vers la stratégie optimale, ou en tout cas qui semble optimale. Suffisamment de comparaisons ont déjà été faites par des universitaires pour que l'on soit convaincu que c'est effectivement plus économique en temps de calcul et en mémoire.

## **LES MÉTHODES TRADITIONNELLES GARANTISSENT UNE ERREUR MAXIMUM “LOCALE”, ALORS QUE LES RÉSEAUX DE NEURONES NE GARANTISSENT QU’UNE ERREUR “EN MOYENNE”.**

Evidemment, une fois qu’on se rend compte de cela, on n’a qu’une envie c’est d’augmenter la dimension du problème pour traiter des cas qui étaient tout bonnement impossible à traiter avec les méthodes traditionnelles. Et en effet c’est possible. Mais cela soulève la question, puisqu’on n’a plus d’éléments de comparaison, et qu’on est en grande dimension, de valider que ces stratégies sont effectivement valides. En effet il faut comprendre que les méthodes traditionnelles garantissent aujourd’hui une erreur maximum “locale”, au sens où l’erreur est contrôlée de la même façon partout, alors que les réseaux de neurones ne garantissent qu’une erreur “en moyenne”. Pour un niveau d’erreur mesurée on ne sait pas si le réseau de neurones fait une erreur de cet ordre de grandeur partout ou bien s’il en fait de très grandes à certains endroits et de très petites à d’autres. Et en grande dimension, on sait que ce qu’on appelle la “malédiction des grandes dimensions” trompe beaucoup les calculs “en moyenne”. Cela veut dire qu’il reste à mettre au point une méthode qui traduise cette garantie “en moyenne” en une garantie “presque partout”. Il y a des pistes mais rien d’établi à ma connaissance.

L’étape d’après, si on est à l’aise avec l’utilisation de méthodes d’apprentissage pour mettre au point des stratégies de couverture, c’est de recourir à des méthodes du même genre pour générer des trajectoires à partir des données plutôt que de faire des simulations à partir de processus modélisés dont les paramètres sont implicites sous mesure risque neutre. Et c’est encore un très gros changement. D’un côté on peut se dire que c’est l’occasion de se rapprocher d’une réalité, notamment en prenant en compte des auto-corrélations de façon non paramétriques, mais d’un autre côté il faut éviter de quitter les propriétés sympathiques du monde risque neutre. Certains chercheurs travaillent sur ce sujet mais c’est compliqué. En effet, personne ne veut que des stratégies de couvertures soient influencées par six mois récent de marchés haussiers ; ce n’est pas aux intermédiaires que sont les banques d’investissement de prendre ce genre de paris, ce genre de risques, qui sont typiquement des risques de fonds d’investissement.

## **LE SUCCÈS DE CE GENRE DE PROJET REPOSE SUR UNE PETITE ÉQUIPE D’EXPERT QUI TRAVAILLE DE CONCERT AVEC DES EXPERTS MÉTIERS SUR DES PROJETS BIEN DÉFINIS.**

Je vais conclure en récapitulant quelques points essentiels à mes yeux : en premier mettre en place une innovation secondaire provenant de l’IA n’est pas chose facile, ça demande des efforts et des investissements. Et la bonne façon de favoriser le succès de ce genre de projet est d’avoir une petite équipe d’expert qui travaille de concert, pendant un temps court, avec des experts métiers sur des projets bien définis. Il faut toujours garder à l’esprit qu’un algorithme apprenant est composé de trois choses qu’il faut monitorer toutes les trois : des librairies informatiques externes, des codes internes et des bases de données. Cela engendre une certaine complexité de versionning, de déploiement et de monitoring. Un autre aspect très important provient du biais dans les données : il ne faut pas croire que parce qu’il n’est pas nécessaire d’écrire des équation pour mettre en oeuvre des algorithmes apprenant, on ne fait pas de choix de modèle. Le choix de modèle est caché dans les données, et dans la façon dont l’algorithme va faire de son mieux pour s’adapter aux données. Il est nécessaire de garder un oeil critique, de modélisateur, sur les données. Et lorsqu’on en vient au deep hedging, il faut se remémorer tous ces points, et surtout celui du biais : on ne veut pas qu’une stratégie de réplcation soit influencée par des spécificités récentes épisodiques de la dynamique des prix des instruments négociables. Cela soulève de beaux sujets pour les chercheurs en mathématiques financières. Certains s’y attaquent déjà, et je suis convaincu qu’ils vont y apporter des réponses, et elles ne seront certainement pas qu’il faut fermer les yeux et “faire confiance aux données”.

## Nicole El Karoui

Bonjour à tous et merci à Stéphane pour l'invitation et surtout à Sandrine et Charles-Albert pour ces présentations très pédagogiques et complémentaires sur les problèmes des Big Data et de l'IA. Quant à moi, je ne suis pas du tout une spécialiste du deep-learning, même si le sujet me passionne, à la fois pour ma recherche académique, pour la formation au sein du Master Probabilités et Finance, et plus généralement l'évolution de la société.

# 13 DE L'EXTENSION DES PRODUITS DÉRIVÉS À L'APPRENTISSAGE STATISTIQUE

On peut dire que j'ai été une sorte de « compagnon de route » des marchés « du risque financier », notamment sur les questions de « hedging » des produits dérivés, et de l'émergence de leur régulation (sans oublier les crises ) pendant ces trente dernières années. Curieusement, il me semble qu'il y a un certain intérêt à confronter les évolutions actuelles issues du Big Data et de ses outils technologiques extraordinaires à celles de la finance des années 1980.

Quand les marchés à terme (MATIF et MONEP) ouvrent en France en 1985-86, dix ans après les Etats Unis, nous sommes en pleine révolution informatique, qui correspond à l'abandon des gros ordinateurs IBM au profit de PC. Les banques française en profitent pour monter leur activité autour des produits dérivés de manière très quantitative : notamment à la Société Générale. Ce sont des ingénieurs qui développent le business, misant prioritairement sur les possibilités offertes par l'informatique.

Parallèlement, en Faculté des sciences et à l'Ecole des Ponts et Chaussées, puis plus tard dans d'autres universités et grandes Ecoles, se montent des formations en mathématiques financières très quantitatives, la plupart dirigées par des femmes (une exception française dans ce monde très technologique). La croissance rapide de la puissance des ordinateurs (loi de Moore) permet le développement de méthodes probabilistes numériques (de type Monte Carlo) efficaces (pour l'époque). De la période expérimentale (de type start-up) des années 90, où la question du sens et de l'efficacité du dynamic hedging des produits dérivés était sans cesse questionnée, on est passé à la phase « industrielle » des années 2000-2008, où l'innovation financière, l'expansion des marchés (parallèle à celle des ordinateurs) explosent, malgré une régulation accrue. Tous les supports deviennent autorisés, les dérivés de crédit, commodities, sans autres réflexions sur les nouveaux risques créés, puisque cela était très rentable... Les aspects théoriques été très délaissés.

## LA MONTÉE EN PUISSANCE DES PRODUITS DÉRIVÉE A ÉTÉ CAUSÉE, COMME LE MACHINE LEARNING AUJOURD'HUI, PAR LA TECHNOLOGIE

L'analogie avec la période actuelle concerne la technologie. Comme pour les dérivés de crédit en finance, tout semble possible, faisable : cela devient d'abord une question d'accélération de calculs et d'algorithmes, à propos desquels on pense avant tout technique et problème informatique, ensuite seulement analyse et interprétation. Comme le souligne Sandrine, la connaissance du modélisateur, et son intuition de marché n'est plus au centre de la démarche de validation.

Cela fait écho aux années 60, lorsque les statisticiens français autour de Benzécri ont été parmi les pionniers à développer l'analyse des données, connue aussi comme analyse en composantes principales etc. L'idée était de représenter géométriquement les données, et de les projeter sur les axes les plus significatifs. L'une des innovations majeures avait été d'afficher la représentation graphique de ces données dans les nouveaux axes. Plus de modèle statistique, les données parlaient d'elles-mêmes, et de tels graphiques se trouvaient par exemple dans des journaux non spécialisés comme le « Nouvel Observateur ». Mais cette parole pouvait être très difficile à interpréter, et à valider, en médecine en particulier lorsqu'on était confronté à des « mélanges » de facteurs explicatifs. Sa simplicité en faisait sa beauté...mais a fini par limiter son utilisation.

Cette question rejoint les préoccupations actuelles d'interprétabilité, même si aujourd'hui les dimensions des données sont infiniment plus nombreuses.

## **DISPOSER D'UN GRAND NOMBRE DE DONNÉES N'ÉLIMINE PAS L'INCERTITUDE**

Un autre biais est lié à la mythologie qu'un grand nombre de données limite l'incertitude, en donnant une information "probabilisable". Mais, cela est franchement erroné si les données ne sont pas indépendantes ou stationnaires, d'où l'importance de la connaissance empirique du panel de données sur lequel on travaille, en rapport avec le problème l'on cherche à résoudre.

Par exemple, l'important dans un taux de change est son drift, sa tendance, si on fait de l'investissement et sa volatilité si on fait du « hedge », cela n'a pas grand chose à voir et influe significativement sur le choix de données et de modèles à prendre en considération.

Ce biais est moins important lorsqu'on se réfère à des phénomènes physiques, qui se reproduisent de manière relativement similaires au cours du temps. Dans le monde de la banque et de l'assurance, et pour beaucoup de données sociétales en général, nous sommes confrontés au fait que les marchés et les sociétés apprennent vite ; les données passées sont en conséquent rarement informatives pour les projections dans le futur, d'où la nécessité d'une vigilance accrue.

# 14 LE RÔLE DE LA MODÉLISATION DANS LA VALIDATION DE RÉSULTATS NUMÉRIQUES

En fait, tous les trois, nous avons soulevées sous une forme ou sur une autre la question de la validation des résultats. La première phase est empirique, et concerne la cohérence des données produites, les « output » : données aberrantes, comparaison avec les résultats de la veille... La deuxième phase est indissociable d'une bonne représentation des problèmes, et cette question concerne aussi bien les données que les modèles.

La première des choses c'est d'avoir une bonne représentation des problèmes. Par exemple, en finance de marché la question de l'échelle de temps à prendre en considération n'est pas aussi triviale qu'il y paraît : l'échéance du produit peut être très éloignée, mais l'horizon du « hedging » est à la journée, réglementairement ou presque, sans compter toutes les échéances intermédiaires dues à la régulation. Ce n'est pas évident de savoir comment prendre ces éléments en considération en même temps dans la vision des données et des algorithmes qu'on va mettre en oeuvre. C'est aussi vrai dans les choix de modèles mathématiques, la seule différence est qu'on peut en mesurer théoriquement plus facilement les conséquences, et remettre en cause les hypothèses de départ si les résultats ne conviennent pas.

## LES MATHÉMATIQUES EXIGENT L'UTILISATION DE TERMES PRÉCIS, CE QUI PERMET DE STRUCTURER L'EXPLOITATION DE DONNÉES

Un autre apport fondamental des mathématiques aux sciences appliquées, préalable à toute modélisation porte sur les contraintes qu'elles font porter sur le vocabulaire. Il ne peut y avoir d'échanges ou de développements dans l'ambiguïté.

Quand j'ai commencé en finance, il y avait à peu près une dizaine de notions sous le nom de volatilité : la volatilité de marché, implicite, locale, moyenne, à la monnaie, et bien d'autres qualificatifs qui n'étaient pas spécifiés. Imaginez ce que cela a comme implications sur la récolte des données.

D'ailleurs, dans les sujets qui nous préoccupent, les questions sont les mêmes : par exemple les données, que sont elles pour vous, pour moi ? des données brutes, re-traitées, de grandes dimensions (laquelle ?), reproductibles, originelles, générées par des algorithmes de la banque, achetées, vendues, personnelles, récupérées sur le web, stockées sur le Cloud, dans votre informatique, obsolètes, historiques, etc. On est vite proche d'un inventaire à la Prévert. Avec ma culture bancaire, je suis tentée d'y associer des notions de risques.... et alors cet inventaire ne se termine plus.

Depuis trente ans que je regarde (moi académique) fonctionner le monde de la banque et de l'assurance, dont le risque est par définition inscrit dans le temps, je n'ai jamais vu se mettre en place une culture de signaux (faibles parfois) alertant sur des ruptures potentielles. Cela devrait être une culture nécessaire dans le monde incertain, que la vigilance « bottom up » soit de mise. Cela suppose d'augmenter la diversité des regards, points de vue mais aussi la formation.

Il faut dire que bien que mon métier de mathématicienne appliquée soit de décrypter autrement les univers que j'étudie, fondamentalement dynamiques et incertains, (marchés financiers, régulation, assurance vie et longévité), je trouve que les pistes sont souvent brouillées, et que j'ai parfois du mal à m'y retrouver, même si le challenge est souvent excitant. Je vais m'expliquer un peu mieux, en prenant un exemple un peu académique:

La première fois que j'ai vu que pour tirer des nombres de manière parfaitement aléatoire entre 0 et 1, on utilisait un algorithme déterministe, cela m'a plus que surprise, estomaquée on peut dire. C'est l'algorithme qui sous-tend la touche random de votre ordinateur. Il faut savoir que c'est la base de toutes

les simulations dont nous vous avons parlé. Donc il n'y a pas d'aléa...et moi qui avait un « vrai faible » pour le hasard...

J'ai été « rassurée » quand j'ai découvert que l'algorithme de recherche de la meilleure page (contenant un mot donné) de Google consiste dans sa recherche du meilleur classement à perturber aléatoirement le système pour l'éviter de rester piéger dans un « puits » non optimal. « ouf », je pouvais continuer à considérer « mes » probabilités comme géniales...

Mais, le monde devient vraiment très imbriqué ; il est très difficile de chiffrer les importances respectives des aspects déterministes et aléatoires.

Cela rejoint ce qu'on évoquait. On s'habitue à interpréter des données dans un monde où souvent on ne sait pas comment et qui les ont créées.

# 15 LES TECHNOLOGIES DE RUPTURE TRANSFORMENT NOTRE VISION DU MONDE

**Nicole El Karoui**

Bien sûr, ce ne sont que des questions, mais qui me paraissent importantes dans le champ qui nous concerne. D'ailleurs, cela ne concerne pas seulement l'interprétation mais aussi une forme de conditionnement. J'ai remarqué que je ne réfléchis plus de la même façon en maths depuis que je tape à l'écran, comparativement à l'époque où j'écrivais sur du papier. Tous ces effets là créent une culture Google comme tu dis, Charles-Albert, sur la recherche, sur la manière de voir, sur le vocabulaire, qui est très conditionnante ; il faut y réfléchir sérieusement, pour qu'on ne perde pas trop de crédibilité dans la manière de penser et d'évoluer dans l'univers de l'IA. Par exemple en ce qui concerne le « formatage » des algorithmes : quand c'est toi qui a programmé l'algorithme, tu sais ce qu'il y a dedans, tu t'es fatigué à le faire... et tu connais en partie ses limites ; mais quand tu vas utiliser un logiciel « open source » ou l'un des algorithmes en Python ou R, il y a toute une partie dont tu n'es plus maître, que tu t'es appropriée de l'extérieur, une boîte noire. A un certain niveau, qu'est-ce que ça laisse comme degré de liberté de réfléchir par exemple. Où sont les espaces singuliers pour se poser les bonnes questions?

## UN NOUVEAU RAPPORT AUX LIBRAIRIES INFORMATIQUES EST EN TRAIN D'ÉMERGER

**Charles-Albert Lehalle**

Je suis tout à fait d'accord. Un changement significatif s'est produit lorsqu'il y a dix ans le MIT a arrêté d'utiliser Scheme comme langage pour apprendre l'informatique aux étudiants et est passé à Python. La justification était complètement incroyable, et conforme à ce que tu dis. Les responsables des études d'informatiques au MIT de l'époque, par ailleurs créateur du langage Scheme, a expliqué que la façon d'utiliser l'informatique aujourd'hui avait considérablement changée, ce n'est plus celui de quelqu'un qui va maîtriser de bout en bout la façon dont c'est fabriqué ; c'est un point de vue d'un scientifique extérieur qui observe des librairies et qui veut comprendre à l'aide d'un protocole d'essais - erreurs afin de comprendre comment elles fonctionnent. C'est un changement de paradigme considérable sur la nature de l'outil informatique, que le MIT a complètement accepté et justifié il y a dix ans déjà. C'est exactement le même genre de changement que tu évoques. Si on veut conserver ton analogie : «la chaîne de la logique mathématique est rompue », le programmeur n'a plus un point de vue de bout en bout qui va jusqu'à l'axiomatique. Les informaticiens étaient comme les mathématiciens qui utilisent des théorèmes qu'ils savent démontrer de bout en bout ; désormais il y a un sorte d'étape d'expérimentation. Aujourd'hui ils expérimentent, font des allers et retours entre essais et erreurs sur des données, et finissent par se dire : ça marche donc comme cela. C'est deux fois plus piégeant, parce que on se fait une intuition de comment ça fonctionne à partir d'essais erreurs sur des données qui ont des biais, et ensuite, alors qu'on est à l'aise psychologiquement, on va les réutiliser sur des données qu'on ne connaît pas. Il y a deux niveaux de manque de compréhension : au moment de l'appropriation des outils puis au moment de leur utilisation.

Merci beaucoup, tu as exprimé beaucoup mieux que moi ce que je voulais faire sentir.

Finalement il y a énormément de choses qu'on finit par apprendre et par réutiliser sans s'interroger sur le pourquoi et le comment, sans questionner l'origine des données et du problème. D'un certain point de vue c'est heureux car ce sont des questions sans fin, mais dans tous les systèmes de risques, de surveillance, de communication, d'information ce sont des points sur lesquels il est très important d'être vigilant. Il ne s'agit pas de les remettre en cause fondamentalement, mais de réfléchir à la manière de se les approprier, les apprivoiser, les critiquer surtout.

# 16 LE RAISONNEMENT AU SECOURS DE L'EXPLOITATION AVEUGLE DES DONNÉES

J'ai évoqué le domaine des probabilités, où on continue malgré tout à faire des simulations, des maths et des probabilités théoriques, même si on peut faire maintenant beaucoup d'expérimentation avec Monte Carlo. D'abord, parce que l'on sait depuis longtemps que cela ne marche pas toujours si bien que cela, et ce sont les mathématiques seules qui ont la capacité de démontrer pourquoi. Mais il y a un autre problème de taille : l'échelle de temps pour résoudre un problème un peu compliqué en mathématiques est plutôt d'un an ou deux, parfois plus ; cette temporalité ne résiste pas aux 15 jours que l'on va mettre pour améliorer l'algorithme de manière acceptable pour l'utilisateur, sans pour autant résoudre le problème.

Au département de Statistique de Berkeley, il y a 4-5 ans les « statisticiens théoriques » qui étudient par exemple la (vitesse de) convergence des algorithmes ; sont devenus l'exception. La notion de preuve « académique » a presque disparu au profit de celle des informaticiens, c'est à dire par validation par l'expérience, sur un jeu de données test, en montrant qu'on améliore les résultats précédents. Par suite, Il ne sert plus à rien de justifier mathématiquement pour les grandes données, les matrices aléatoires ou autre chose, pourquoi la pratique peut conduire à des erreurs, mêmes avec des outils très classiques comme le bootstrap. Il faut tout de même dire que souvent les seuils critiques ne sont pas encore atteints.

## LES NOUVELLES TECHNOLOGIES SONT PORTEUSES DE RÉVOLUTIONS CONCEPTUELLES: QUELLES SONT CELLES QUE NOUS RÉSERVE L'IA?

La question dépasse en fait l'idée de validation déjà évoquée ; elle relève plutôt d'une réflexion sur La Réforme du Vrai, titre de l'ouvrage passionnant d'un physiologiste végétal et philosophe, Nissim Amzallag, paru en 2010, dont je vais relater l'exemple de l'horloge mécanique que j'ai trouvé assez extraordinaire. Vers 1350, apparaissent les premières horloges mécaniques. C'est une horloge qui est complètement autonome, grâce à un régulateur articulé à un balancier. Avant il y avait des horloges à eau pas toujours très fiables car dépendant des variations de la pression, et des cadrans solaires. L'horloge mécanique est un assemblage très ingénieux et novateur d'une roue crantée et du balancier, qui découpe de temps en segments égaux. Il y a donc eu à cette période des artisans particulièrement innovants. Les conséquences en ont été spectaculaires. D'abord contrairement aux horloges à eau, ces horloges étaient parfaitement reproductibles, donnant à chacun la possibilité d'avoir la même échelle de temps, plutôt que de référer à un temps subjectif, de l'enfance, des saisons, différent la nuit... Cette machine change radicalement la perception du temps, qui devient sécable en petites unités. Ce progrès conceptuel a été l'amorce de la grande révolution scientifique de la Renaissance. Un progrès technologique, spectaculaire induit une révolution scientifique, d'une ampleur exceptionnelle.

La révolution technologique de l'IA, a au sein de notre société beaucoup de points communs avec celle née de l'horloge mécanique, amplifiée par la découverte peu après de l'imprimerie.

Quelles sont les révolutions conceptuelles que toutes ces nouvelles technologies vont nous obliger à faire ? Comment vont changer nos manières de penser ? Cela a manifestement déjà commencé, avec l'addiction au portable par exemple, mais je pense aussi que de vraies surprises « philosophiques » nous attendent.

C'est donc un vrai challenge qui nous est posé, à nous femmes et hommes, afin que ce levier technologie n'oriente pas notre société démocratique dans des directions que nous ne pourrions plus maîtrisées...

## BIOGRAPHIES DES PARTICIPANTS



### Charles-Albert LEHALLE

Responsable de l'analyse de données chez Capital Fund Management (CFM, Paris) et Chercheur invité à l'Imperial College (Londres), Charles-Albert a étudié le Machine-Learning pour le contrôle stochastique lors de son doctorat il y a 20 ans. Il a commencé sa carrière en tant que Responsable des projets IA au centre de recherche de Renault et a rejoint le milieu de la finance en 2005 avec l'apparition du trading automatisé. En 2016, Charles-Albert reçoit le prix du meilleur article en finance de l'Institut Europlace de Finance (IEF) et il publie plus de cinquante articles académiques et chapitres de livres. Il est co-auteur du livre « Market Microstructure in Practice » (World Scientific Publisher, 2e édition 2018), analysant la principale caractéristique des marchés actuels. Il est le Directeur scientifique du Programme de recherche Interdisciplinaire «Finance and Insurance Reloaded» de l'Institut Louis Bachelier; ce programme explore l'influence des nouvelles technologies (de la blockchain à l'intelligence artificielle) sur l'industrie du secteur banque/finance/assurance.



### Sandrine Ungari

Sandrine Ungari est actuellement responsable de l'équipe de recherche quantitative cross-asset à la Société Générale. Au sein du groupe Cross-Asset Research, l'équipe de recherche quantitative est active dans les stratégies de primes de risque, les produits dérivés et structurés, la modélisation des risques de portefeuille, et fournit des recherches aux investisseurs du monde entier. Le groupe a été reconnu comme un leader du marché de la recherche quantitative et a été classé n° 1 dans l'enquête Extel dans la catégorie Stratégies quantitatives. Sandrine a rejoint la Société Générale en 2006. Auparavant, elle a travaillé en tant qu'analyste quantitative chez HBOS Treasury et chez Reech Sungard à Londres. Elle est diplômée de l'ENSTA (Paris) et titulaire d'un Master en Finance Quantitative de l'Université Paris VI. Elle est maître de conférences à l'Université Paris Diderot.



## Nicole El Karoui

Nicole El Karoui est professeur Emérite à Sorbonne Université (Paris VI), après avoir été pendant dix ans professeur à l'École Polytechnique. Spécialiste du contrôle stochastique, elle a utilisé ses idées dans les marchés financiers, contribuant à l'émergence du domaine international des «mathématiques financières». L'École Polytechnique lui a donné l'opportunité de créer une équipe de recherche dans ce domaine, au rayonnement international. Ses recherches actuelles concernent le risque de longévité vu sous l'angle de la dynamique de population et l'utilité dynamique.

Très soucieuse de transmettre aux jeunes, elle crée en 1990 à Paris VI, le DEA (Master de Probabilités et finance), le premier dans une université scientifique. Par cette formation, elle s'est fait connaître dans les marchés financiers du monde entier, allant jusqu'à faire la «une» du Wall Street Journal en 2006. Les «quants» français sont très recherchés dans les banques d'investissement.

À la suite de la création de la chaire «Risques financiers» par la Société Générale, elle œuvre activement à la création de la Fondation du Risque en 2007 et à l'Institut Louis Bachelier, persuadée que cette structuration des liens entre académiques et professionnels était un atout très important de la France.