

IA, qui décide ? Elle ou nous ?

2e partie – Ça dépend...

Dans la première partie de cet article, inspiré par le livre « Faut-il encore décider ? » d'Éric Azan et Olivier Sibony¹, nous avons compris que, techniquement, il y a deux types d'IA aux principes de fonctionnement très différents. Dans cette deuxième partie, je vais tenter de résumer la pensée des auteurs sur la question quasi philosophique de savoir comment adapter son processus de décision à ces nouveaux outils.

« L'IA a déjà tellement transformé le monde que cela impose que nous nous transformions nous-mêmes pour continuer à exister dans ce nouvel environnement ». Je retiens cette phrase clé des auteurs, non pas que ce besoin d'adaptation soit une surprise, elle est surtout un encouragement, ou un rappel salvateur, dans cette période qui va si vite et qui remet tellement de modèles en cause. Avançons !

L'IA a déjà conquis plusieurs environnements où sa capacité de « décision » est plus rapide et plus fiable que celle de l'humain. On peut citer quelques exemples :

- La finance, tant dans la décision automatique d'acheter et de vendre des actions à la vitesse de la lumière que dans le *credit scoring* ou la détection de fraude.
- La gestion des stocks, où là aussi, les algorithmes sont plus fiables et plus rapides que les humains pour décider de quand et à qui, acheter tel ou tel produit en intégrant de multiples critères comme la prévision météo ou les derniers sondages « d'ambiance ».
- La sécurité, où les caméras « intelligentes » détectent la menace, voire, en associant d'autres systèmes, participent à la prédiction de risque.
- L'imagerie médicale, où la capacité des IA à détecter les cancers de façon plus fiable, et surtout mieux anticipée, que les médecins spécialistes est aujourd'hui bien connue.

Trois conditions caractérisent ces domaines où l'IA a déjà conquis la décision :

- L'IA doit être concentrée sur un environnement bien délimité. Ces IA sont en quelque sorte les expertes de chacun de ces métiers : un oncologue n'est pas un chirurgien cardiaque. Il est rassurant de voir que la notion d'expertise reste un élément clé.
- Les objectifs fixés à l'IA doivent être clairs. Comme dit le proverbe, si l'on ne sait pas où l'on va, on ne risque pas d'y arriver. La vérité reste, c'est l'outil qui change.
- L'IA a un besoin impérieux d'une masse importante de données pour son apprentissage : c'est l'exemple connu depuis le début des IA, elle doit avoir vu des dizaines de milliers de photos de chats pour savoir reconnaître un chat.

¹ Flammarion 2026

Quand ces trois conditions sont réunies, l'acceptabilité de la décision de l'IA se fait vite par les personnes concernées. Cette acceptabilité entraîne ensuite la confiance dans la décision proposée par l'IA. Ainsi, le médecin peut mieux s'occuper de son patient, le logisticien de son client et le trader, devenu agrégé de mathématiques, de ses algorithmes.

Observons ici que sur ces terrains déjà conquis par une IA meilleure que l'homme, il est assez facile d'étudier le niveau de fiabilité des décisions proposées par l'IA et de le comparer à celui de l'humain. Quand l'étude confirme qu'elle est meilleure, on peut la laisser décider. Ceci ouvre alors d'autres questions comme celle de la responsabilité de la décision ou du risque de perte de compétences des humains.

Au-delà de ces exemples où les éléments sont facilement mesurables et comparables, le territoire des IA génératives, les LLM (*Large Language Model*), ouvre un champ de questions plus large quant à la décision. Basés sur des techniques statistiques et probabilistes, elles génèrent mécaniquement des erreurs qu'il est d'autant plus difficile de détecter que leur système de « raisonnement » est une boîte noire dans laquelle il est quasi impossible de pénétrer. Ces IA sont un peu comme le *Système 1* de notre cerveau (voir première partie de l'article). Elles ont le niveau de fiabilité de l'intuition humaine ; nous savons aujourd'hui qu'il est très variable suivant les individus ou les domaines où elle s'exprime². C'est là qu'il est indispensable de savoir que l'IA utilisée est basée sur un LLM. À ce stade de la maturité de ces systèmes, il semble déconseillé de laisser la décision de ces IA sans contrôle. Quelques exemples célèbres viennent illustrer les limites des IA génératives :

- Aux États-Unis, des procès ont été annulés, ou des décisions revues, quand, après étude, la justice s'est aperçue que l'IA avait inventé de toutes pièces des jurisprudences. Le contrôle de l'humain, s'il avait été fait, aurait probablement provoqué une autre décision de la part des avocats mis en cause.
- Un grand cabinet de conseil mondial, Deloitte en Australie, a été pris en défaut dans le rapport conclusif d'une étude à quelques centaines de milliers de dollars. Probablement que de jeunes cadres brillants et très bien formés s'étaient laissé berné par leur passion pour les nouvelles technologies. Le rapport contenait des références inventées et des citations fausses. C'était beau, c'était bien présenté, c'était cher, mais le travail n'avait pas été contrôlé par un humain.

Plus on s'éloigne des domaines scientifiques et techniques, les sciences dures, aux IA majoritairement algorithmiques, plus il est difficile de faire confiance aux décisions de l'IA. Il est facile de comprendre que si le modèle utilisé pour l'apprentissage de l'IA est biaisé — et il est facile de le biaiser —, la décision sera biaisée. Le modèle d'apprentissage peut être biaisé volontairement ou non. Deux exemples illustrent cela :

- Dans le domaine de la décision de justice, si au cours de son apprentissage, l'IA est alimentée de façon disproportionnée par des conclusions où le coupable est majoritairement de couleur noire, elle va avoir tendance à surpondérer l'importance de la couleur de peau noire du justiciable dans sa décision.
- Une IA utilisée par Amazon pour son recrutement a été abandonnée car jugée biaisée. Elle avait été entraînée sur les décisions de recrutement... d'Amazon. Involontairement, elle reproduisait les biais des recruteurs, mais en toute « bonne foi ». Biaisée peut-être, mais prenait-elle majoritairement de bonnes décisions de recrutement ? Meilleures ou moins bonnes qu'avant ?

² Voir mes travaux sur le management de la complexité. Thèse soutenue en 2021 et livre « Le management de la complexité », 2022, disponible sur Amazon (chapitre 6 : l'intuition).

Les auteurs de l'ouvrage proposent cinq conditions qu'il faut respecter avant de pouvoir penser que la décision proposée par une IA est probablement la bonne ou la meilleure. Dans cette tentative de résumé, j'en retiens trois qui concernent le processus de décision des cadres et des dirigeants dans leur contexte professionnel.

L'IA doit être validée de façon indépendante.

Dans les domaines techniques ou scientifiques, le taux de fiabilité des décisions de l'IA est facilement comparable à celui des humains du domaine concerné. La validation de l'IA est alors facile par ses « pairs ». Une fois validée, on peut leur faire confiance.

Dans les domaines qui appellent des aspects d'éthique, de valeurs ou de jugements subjectifs, le domaine des IA génératives, on comprend bien que l'indépendance est impossible. Éthique, valeur, équité n'ont pas le même sens à Moscou, au Japon ou au Mali. Il faudrait ici faire quelques retours sur le concept du libre arbitre en convoquant Saint Augustin, Descartes, Spinoza et quelques-uns de leurs amis. Posons simplement que le libre arbitre d'une IA générative est aujourd'hui une illusion. Elle a le libre arbitre de quelqu'un d'autre.

Le processus de décision de l'IA doit être explicable

L'humain qui s'appuie sur une IA pour prendre une décision doit avoir une compréhension suffisante du fonctionnement de l'IA. Cela lui permettra de savoir à quel moment et dans quelles conditions il doit se montrer plus vigilant quant à la décision proposée. Comme un collaborateur défendant une position face à son patron, l'IA doit pouvoir expliquer pourquoi elle arrive à telle ou telle conclusion. C'est là que l'origine des références et des informations qu'elle a exploitées donnent des pistes utiles pour évaluer la crédibilité du résultat proposé. Si ce n'est pas le cas, prudence.

Les biais de l'IA doivent avoir été identifiés et corrigés

Le biais de la machine sera le biais de l'humain qui l'a construite. Dans ce domaine, de nombreux travaux sont engagés car si le biais (on parle ici des biais involontaires) est identifié, il est assez facile de corriger le processus de « réflexion » de l'IA sur ce point... à condition de ne pas réintroduire involontairement un nouveau biais.

Conclusion

Le livre « *Faut-il encore décider ?* » aborde de nombreuses autres questions que je vous encourage à découvrir. Comme lors de toutes les révolutions précédentes de l'informatique, car l'IA n'est qu'une nouvelle étape dans cette révolution industrielle, cela ne sert à rien d'être inquiet, il faut comprendre et apprendre. L'IA n'est qu'un nouvel outil de la boîte à outils. Ne nous laissons pas séduire par sa pseudo-empathie ; elle ne pense pas, elle « *raisonne comme un tambour* » pour reprendre une expression chère à ma grand-mère.

Comme un marteau dans la boîte à outils, l'IA n'a pas (aujourd'hui) de créativité. Cependant, elle vous évitera de commettre des erreurs qui ont déjà été faites. Elle vous évitera d'oublier un point qui aurait

pu vous être fatal. Et c'est avec ces avantages-là, sous le contrôle de votre libre arbitre toujours vigilant (votre Système 2...), que vous aurez un peu plus de temps pour vous perdre dans une belle errance intellectuelle qui vous fera trouver cette idée que personne n'a jamais trouvée.

25 juin 2026

Michel MATHIEU

Doctor of business administration.