

Étudier à l'ère de l'IA

Guide pratique pour les étudiants en sciences de la santé

Dr. Javier Flores Cohaila



ELSEVIER

Advancing human progress together

Sommaire

Prologue : Le rôle de l'intelligence artificielle dans la formation en santé aujourd'hui	3
Glossaire rapide	5
Rencontrez les voix qui vous accompagneront	8
#1 Pourquoi ce guide, et pourquoi	9
#2 Ce que vous allez apprendre	18
#3 Ce que l'IA ne vous dira pas	22
#4 L'IA et la pensée critique	33
#5 L'IA et le raisonnement clinique	44
#6 L'IA et l'apprentissage autorégulé	53
#7 Éthique, sécurité et professionnalisme	65
#8 Checklist pour bien démarrer	75
#9 FAQs, mythes et réalités	86



Prologue

Le rôle de l'intelligence artificielle dans la formation en santé

Lorsque internet est apparu, la formation médicale a connu une profonde transformation. Aujourd'hui, avec l'intégration de l'intelligence artificielle, comment percevez-vous l'évolution de la formation en santé ?

David Game, Senior Vice President, Elsevier Global Product, Global Medical Education

Lorsque internet est arrivé, certaines choses ont effectivement changé, mais beaucoup d'autres sont restées les mêmes. Les qualités nécessaires pour devenir un excellent médecin ou un excellent infirmier sont fondamentalement les mêmes qu'avant l'avènement d'Internet. Il en va de même pour les compétences que les étudiants doivent acquérir et pour le processus d'apprentissage lui-même : s'investir dans des sujets complexes, les maîtriser, puis les appliquer à des situations cliniques concrètes. Ces aspects essentiels n'ont pas changé.



Les professionnels de santé utilisent aujourd'hui de nouvelles technologies, mais la pensée critique, la passion pour le métier et l'empathie demeurent au cœur de la pratique. Ce sont précisément ces qualités que nous recherchons, en tant que patients, chez nos médecins et nos infirmiers. Ainsi, même si les fondements de ce qui fait un excellent professionnel de santé n'ont pas évolué, les outils, eux, ont changé. Ils permettent notamment de mieux suivre les progrès des étudiants et de leur donner accès à des ressources qui leur étaient auparavant difficilement accessibles.

La manière dont l'intelligence artificielle transformera la formation médicale reste encore une question ouverte. Ce que nous observons aujourd'hui, c'est sa capacité à fournir davantage de ressources pédagogiques personnalisées aux étudiants, où qu'ils se trouvent. Désormais, chaque étudiant peut utiliser un outil d'IA générative, comme ChatGPT, Claude ou d'autres, comme un tuteur personnel disponible à tout moment.

Dans une perspective optimiste, cette évolution pourrait contribuer à résoudre ce que Benjamin Bloom appelait le « problème des deux écarts-types » (2 Sigma Problem) : comment offrir un enseignement individualisé, la forme d'apprentissage la plus efficace à des centaines de milliers d'étudiants, tout en restant économiquement viable ?



Les partisans de l'IA estiment qu'avec le temps, ces outils comprendront de mieux en mieux le raisonnement des étudiants, identifieront les causes de leurs erreurs et interpréteront leurs réponses de manière plus fine. Ce « tuteur IA » pourrait alors fournir un accompagnement plus pertinent et améliorer les résultats d'apprentissage.

Les sceptiques, quant à eux, soulignent que ces interactions risquent de rester relativement superficielles. Ils questionnent la capacité réelle d'un moteur prédictif à comprendre véritablement ce qu'un étudiant pense, pourquoi il le pense, et à le faire sans l'interaction humaine directe entre l'étudiant et l'enseignant. Selon eux, cela pourrait conduire à des réponses standardisées et à un apprentissage davantage fondé sur la répétition.

Un autre défi concerne l'évaluation. Toute forme d'évaluation réalisée avec un accès à Internet est aujourd'hui remise en question. Je discutais récemment avec un enseignant à Galway qui me racontait avoir conçu un excellent cours, doté d'objectifs pédagogiques solides et d'évaluations pertinentes. Puis ChatGPT est arrivé, et il est devenu beaucoup plus difficile de garantir que les résultats obtenus reflétaient réellement les compétences des étudiants.

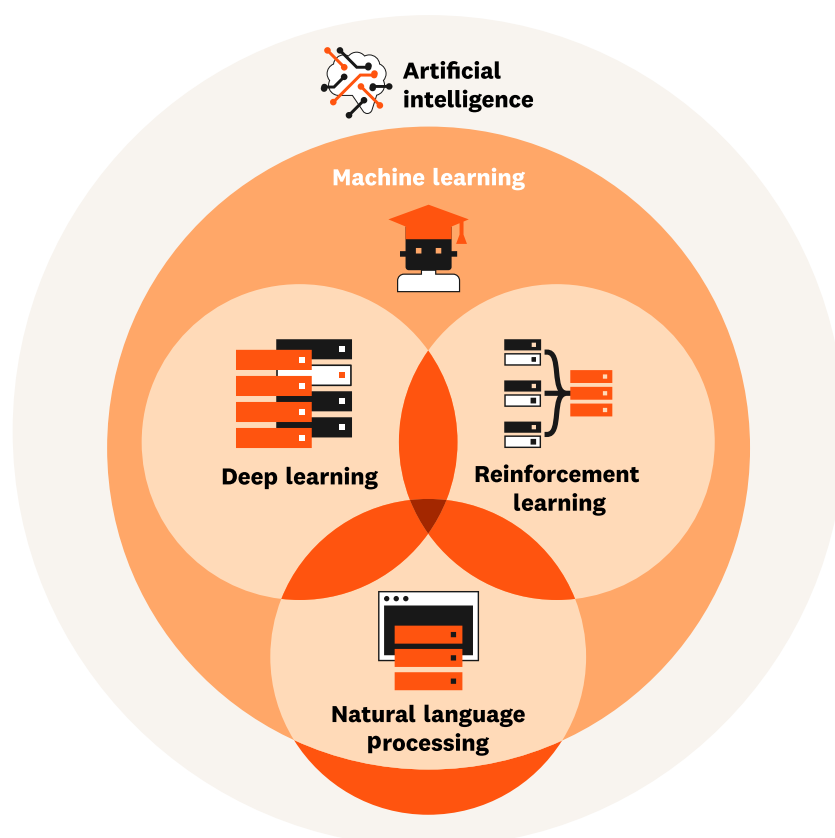
À l'heure actuelle, l'usage de l'IA dans la formation médicale est extrêmement répandu parmi les étudiants. En revanche, nous ne disposons pas encore d'une vision claire et partagée de son rôle comme outil pédagogique complémentaire au travail des enseignants qualifiés. Cette réflexion est en cours. Les formateurs en médecine, en soins infirmiers et dans les autres professions de santé joueront un rôle déterminant dans la manière dont l'IA sera utilisée. Ils doivent être acteurs de cette évolution plutôt que de la subir.

Nous n'en sommes qu'au début de cette transformation. L'avenir de l'intelligence artificielle dans la formation en santé reste encore largement à écrire.

Glossaire rapide : concepts et indicateurs clés de l'intelligence artificielle

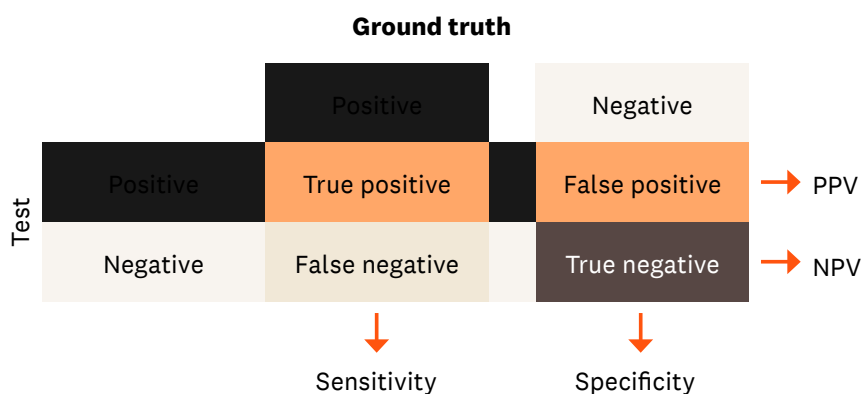
Ce glossaire présente les principaux concepts et indicateurs nécessaires pour comprendre comment l'intelligence artificielle (IA) est utilisée dans la pratique clinique. Des différents types d'apprentissage automatique (machine learning) aux réseaux neuronaux, en passant par les mesures utilisées pour évaluer les modèles prédictifs, il vous aidera à maîtriser le vocabulaire essentiel, à interpréter les résultats et à prendre des décisions éclairées dans le cadre de votre formation en santé.

Conçu pour les étudiants, ce guide de référence vous permettra d'acquérir rapidement et simplement les bases de l'IA, grâce à des explications claires, pratiques et accessibles.



Vue d'ensemble des catégories de l'intelligence artificielle.

Terme	Définition
Artificial intelligence (AI) - Intelligence artificielle (IA)	Domaine de l'informatique qui développe des systèmes et des algorithmes capables d'exécuter des tâches nécessitant habituellement l'intelligence humaine, telles que l'apprentissage, la résolution de problèmes et la prise de décision.
Machine learning (ML) - Apprentissage automatique	Système capable d'accomplir des tâches généralement associées à l'intelligence humaine. En santé, il analyse des données cliniques afin d'aider au diagnostic, d'éclairer la prise de décision et d'améliorer les résultats pour les patients.
Supervised Learning - Apprentissage supervisé	Type d'apprentissage automatique (machine learning) qui utilise des données étiquetées (entrées et résultats connus) pour prédire ou classer de nouveaux cas.
Unsupervised Learning - Apprentissage non supervisé	Type d'apprentissage automatique qui identifie des modèles, des structures ou des relations dans des données non étiquetées.
Deep Learning (DL) - Apprentissage profond	Type d'apprentissage automatique reposant sur des réseaux neuronaux multicouches capables de détecter des schémas complexes dans de grands volumes de données.
Neural network - Réseau neuronal	Modèle inspiré du fonctionnement du cerveau humain, composé de nœuds interconnectés qui traitent et transmettent l'information.
Natural language processing (NLP) - Traitement automatique du langage naturel (TALN)	Branche de l'intelligence artificielle qui permet aux ordinateurs de comprendre, d'interpréter et de traiter le langage humain, qu'il soit écrit ou parlé.
Large language models (LLMs) - Grand modèle de langage	Modèle avancé entraîné sur de vastes quantités de textes, capable de générer du contenu, de répondre à des questions et de s'adapter à de multiples tâches, y compris dans le domaine clinique.
Reinforcement learning (RL) - Apprentissage par renforcement	Sous-domaine de l'apprentissage automatique dans lequel un agent apprend à prendre des décisions en interagissant avec un environnement et en recevant des récompenses ou des pénalités en fonction de ses actions.
Bias - Biais algorithmique	Processus décisionnel pouvant conduire à des résultats injustes, inévitables ou inexacts pour certains groupes sociodémographiques en raison de biais présents dans les données, les algorithmes ou leur utilisation.
Explainability - Explicabilité de l'IA	Capacité à comprendre, interpréter et expliquer les décisions, recommandations ou prédictions produites par un système d'intelligence artificielle.



Calcul des indicateurs de performance

Indicateurs de performance en apprentissage automatique (Machine Learning)

Indicateur	Ce qu'il mesure / Quand l'utiliser
Exactitude (Accuracy)	Proportion de prédictions correctes parmi l'ensemble des prédictions. Peut-être trompeuse lorsque les données sont déséquilibrées.
Précision (VPP – Valeur prédictive positive)	Proportion des prédictions positives qui sont réellement positives. Particulièrement utile lorsque les faux positifs ont des conséquences importantes.
Valeur prédictive négative (VPN)	Proportion des prédictions négatives qui sont correctes. Importante lorsque les faux négatifs sont particulièrement critiques.
Rappel (Recall) ou Sensibilité (TPR – True Positive Rate)	Proportion des cas positifs correctement détectés. Utile lorsque l'objectif est de ne pas manquer de cas pertinents, même si cela peut augmenter le nombre de faux positifs.
Spécificité (TNR – True Negative Rate)	Proportion des cas négatifs correctement identifiés. Mesure la capacité du modèle à exclure correctement les cas négatifs.
Score F1	Moyenne harmonique entre la précision et le rappel. Permet d'équilibrer la prise en compte des faux positifs et des faux négatifs.
AUROC (Area Under the Receiver Operating Characteristic Curve)	Mesure la capacité d'un modèle à discriminer les classes positives et négatives sur l'ensemble des seuils de décision.
AUPRC (Area Under the Precision-Recall Curve)	L'aire sous la courbe précision-rappel, qui résume le compromis entre précision et rappel. Particulièrement utile lorsque les données sont déséquilibrées et que les cas positifs sont rares.



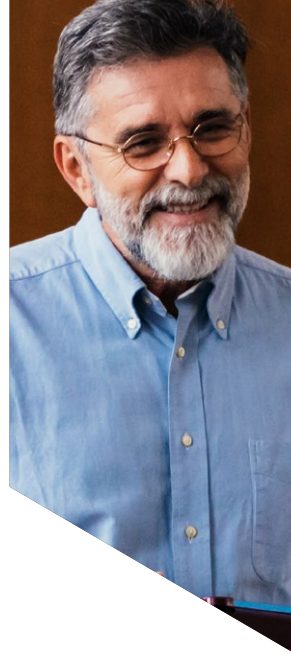
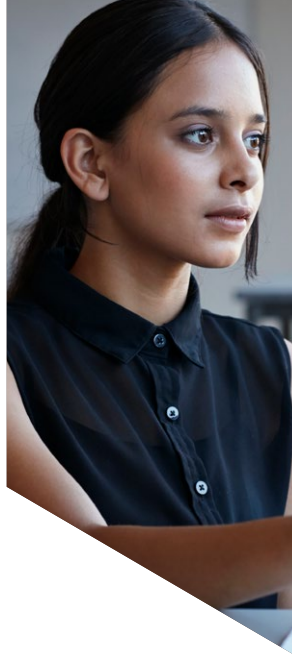
À retenir

L'analyse conjointe de ces indicateurs aide à évaluer de manière plus complète les performances d'un modèle d'IA. Elle permet d'en identifier les points forts et les limites, et de choisir la solution la plus pertinente en fonction du contexte clinique.



Contenu adapté de :

Artificial Intelligence in Clinical Encounters: Opportunities and Challenges by Jiang JJ., Chan L., Nadkarni, GN.
Swartz Textbook of Physical Diagnosis, 9e édition, sous la direction de Mark H. Swartz. Elsevier, 2026.



Rencontrez les voix qui vous accompagneront

Ce guide s'articule autour de trois voix. Il ne s'agit pas de personnes réelles, mais elles incarnent des profils bien réels : des étudiants et des enseignants avec lesquels nous avons travaillé, étudié et appris dans des programmes de formation clinique à travers le monde.

-  **Maya** a 23 ans et effectue ses stages cliniques avancés. Elle est méthodique, disciplinée et exigeante envers elle-même. Elle a créé des modèles et des méthodes pour tout. Elle utilise l'intelligence artificielle pour optimiser ses apprentissages : fiches de révision, résumés cliniques, questions d'entraînement. Son système fonctionne. Ses notes sont excellentes. Pourtant, elle commence à se demander si ce système produit de véritables apprentissages ou simplement de bons résultats. Et surtout, si ces deux choses sont réellement équivalentes.
-  **Kai**, 24 ans, se trouve au même stade de sa formation. Avant d'entreprendre ses études de médecine, il a pris une année de pause non planifiée, dont il est revenu avec une nouvelle vision de l'effort et de l'apprentissage. Pragmatique, perspicace et lucide sur ses propres raccourcis, il ouvre souvent un outil d'IA la veille d'un examen et lui demande d'expliquer ce qu'il n'a pas compris en cours. Cela fonctionne... jusqu'au moment où cela ne fonctionne plus. Il remet tout en question, y compris sa propre tendance au scepticisme : est-ce une force ou simplement une excuse ?
-  **Le Dr Flores** est un enseignant qui a vu l'intelligence artificielle faire son entrée dans la formation clinique et a choisi de s'y intéresser plutôt que de l'interdire. Il s'appuie sur les données comme point de repère, jamais comme une arme. Il pose la question avant de donner la réponse. Il ne se contente pas de transmettre des connaissances : il accompagne, guide et facilite l'apprentissage. Et parfois, il se contente d'écouter.

Ces trois voix vous accompagneront tout au long de ce guide. Maya et Kai illustrent deux manières différentes d'utiliser l'IA dans les études, ainsi que les risques propres à chacune de ces approches. Le Dr Flores apporte les données probantes, les cadres de réflexion et les éléments de mise en perspective. Ensemble, ils explorent une question à laquelle il n'existe pas de réponse unique : **comment utiliser un outil à la fois puissant, imparfait et impossible à ignorer, sans perdre la capacité de réflexion qui est au cœur du raisonnement clinique ?**

Leur conversation débute au chapitre 1.



Chapitre 1 : Pourquoi ce guide, et pourquoi maintenant ?

Kai : Franchement, j'ai une question. Si je demande à une IA générative de m'expliquer l'insuffisance cardiaque et que je comprends parfaitement en dix minutes... pourquoi passer deux heures à lire un chapitre qui dit la même chose, mais en moins clair ?

Maya : Parce que comprendre n'est pas la même chose que savoir. Moi aussi, j'utilise l'IA tous les jours. Mais ces derniers temps, je me pose une question: est-ce que j'apprends vraiment, ou est-ce que je laisse la machine organiser à ma place ce que je devrais moi-même assimiler ? J'ai lu une étude du MIT montrant que l'utilisation de ChatGPT pour rédiger un texte réduisait la connectivité cérébrale jusqu'à 55 % et que 83 % des participants ne se souvenaient plus des détails de ce qu'ils venaient d'écrire avec l'aide de l'IA. Les auteurs parlaient de "dette cognitive". (18)

Kai : Ça, c'est ton problème, pas celui de l'IA.

Maya : Ou peut-être que tu ne t'es pas encore posé la question.

Dr Flores : Les deux points de vue sont valables. Et c'est précisément ce qui rend cette discussion si complexe.

Je vais vous donner un peu de contexte. En 2025, Elsevier a interrogé plus de 3 200 chercheurs répartis dans 113 pays. Parmi eux, 58 % utilisaient déjà des outils d'intelligence artificielle dans leur travail. Un an plus tôt, ils n'étaient que 37 %. (1)

Kai : Attendez... de 37 % à 58 % en une seule année ?

Dr Flores : En une seule année. Et ce qui est intéressant, c'est que les préoccupations augmentent au même rythme que l'utilisation. Nous y reviendrons dans un instant.

Maya : Autrement dit, ils utilisent l'IA, mais quelque chose les met mal à l'aise.

Dr Flores : Exactement. Le débat n'oppose pas ceux qui utilisent l'IA à ceux qui ne l'utilisent pas. La tension existe souvent au sein d'une même personne. Un peu comme entre vous deux.

Regardons maintenant ce qui se passe chez les étudiants. Une enquête internationale menée par Busch et ses collaborateurs auprès de 192 facultés de médecine à travers le monde a révélé d'importantes différences dans la perception de l'IA selon les régions, l'accès aux technologies et la formation reçue. Le phénomène est loin d'être uniforme. (2)

Kai : Mais son utilisation est massive, non ?

Dr Flores : Oui, mais elle varie fortement selon les contextes. À Porto Rico, González-Bravo a montré que près de 73 % des étudiants en médecine utilisent déjà l'IA générative. En Chine, Pan et Ni ont observé un taux de 34 %. Aux États-Unis, Ganjavi a constaté que 96 % des étudiants avaient entendu parler de ChatGPT, mais que seule la moitié l'utilisait pour étudier. Dans les formations en soins infirmiers, l'adoption progresse également, notamment pour la documentation clinique et la préparation des études de cas. (3)(4)(5)



Maya : Donc tout dépend de l'endroit où l'on se trouve.

Dr Flores : De l'endroit où l'on se trouve, de l'accès dont on dispose et du fait qu'une personne nous ait appris à utiliser ces outils ou non. C'est là un élément essentiel. McCoy a montré que les étudiants se sentent beaucoup plus à l'aise avec l'IA générative que leurs enseignants. Il a également constaté que plus de la moitié d'entre eux l'utilisent déjà dans le cadre de leurs activités académiques. La véritable question est donc : avec quel accompagnement ? (6)

Kai : Avec YouTube et en apprenant à force d'essais et d'erreurs.

Dr Flores : C'est précisément le problème. Anthropic, l'entreprise à l'origine de Claude, a analysé un million de conversations d'étudiants universitaires en 2025. Il s'agit de données produites par une entreprise et non d'une étude évaluée par les pairs ; il faut donc les interpréter avec prudence. Néanmoins, les tendances observées rejoignent celles rapportées par d'autres sources. Les filières de santé représentaient seulement 5,5 % des conversations, alors qu'elles comptaient pour 13,1 % des inscriptions universitaires. (7)

Maya : Les étudiants en santé sont donc sous-représentés.

Dr Flores : Exactement. Et lorsqu'ils utilisent l'IA, ce n'est pas principalement pour mémoriser ou réviser. Ils l'utilisent surtout pour créer et analyser. Près de 40 % des interactions concernent des tâches de création et 30 % des tâches d'analyse. (7)

Maya : Une inversion de la taxonomie de Bloom.

Dr Flores : En effet. Les étudiants délèguent à l'IA des tâches qui correspondent aux niveaux supérieurs de la taxonomie de Bloom, c'est-à-dire celles qui sont traditionnellement considérées comme les plus complexes. Ce n'est ni bon ni mauvais en soi. Mais c'est un fait que nous ne pouvons pas ignorer.

Kai : D'accord... je ne m'attendais pas à ces chiffres.

Dr Flores : Ce n'est pas la première fois que l'humanité vit ce type de bouleversement. Lorsque l'écriture a été inventée, Socrate la critiquait déjà. Il pensait qu'elle allait détruire la mémoire et que les gens cesseraient de se souvenir des choses puisqu'ils pourraient simplement les écrire.

Kai : Socrate était donc opposé à la technologie ?

Dr Flores : Socrate était surtout opposé aux béquilles intellectuelles. Il craignait qu'en externalisant le savoir, nous cessions de nous l'approprier véritablement. Et il n'avait pas entièrement tort : l'écriture a effectivement transformé le fonctionnement de la mémoire humaine. Mais elle a aussi rendu possibles la science, le droit et la médecine. Tout ce que vous étudiez aujourd'hui existe parce que quelqu'un a un jour décidé de l'écrire.

Puis l'imprimerie est arrivée. Les moines copistes y ont vu la fin de leur monde. Et, d'une certaine manière, c'était bien la fin de leur monopole. Mais pas celle du savoir. Au contraire : elle l'a démultiplié.

Maya : La même chose s'est produite avec Internet. Beaucoup d'enseignants affirmaient que Google allait empêcher les étudiants d'apprendre.

Dr Flores : Et qu'est-il arrivé ?



Maya : Ceux qui ont appris à bien chercher ont appris plus que jamais. Et ceux qui ne l'ont pas fait... se sont contentés de copier-coller.

Dr Flores : Exactement comme aujourd'hui.

Kai : Donc, le problème n'a jamais été l'outil.

Dr Flores : L'outil n'a jamais été le problème. Chaque fois qu'une technologie a permis d'externaliser une fonction cognitive comme la mémoire, le calcul ou l'accès à l'information, la réaction initiale a été la même : la peur. Pourtant, la véritable question n'a jamais été de savoir s'il fallait l'utiliser, mais comment l'utiliser sans perdre ce qu'elle nous pousse à développer par nous-mêmes. (8)

Maya : Donc la question n'a jamais été de savoir s'il fallait adopter l'IA. La vraie question est de savoir comment la maîtriser avant qu'elle ne nous maîtrise.



 **Dr Flores :** Exactement.

 **Kai :** ... D'accord. Là, je crois que le message est passé.

 **Dr Flores :** La question que tu as posée au début : “Pourquoi emprunter le chemin le plus long si l'outil peut faire le travail correctement ?” n'avait rien de cynique. C'était la bonne question.

Le problème, c'est que la réponse est plus inconfortable qu'elle n'y paraît.

Vous vous souvenez des données d'Elsevier ? Cinquante-huit pour cent des chercheurs utilisent déjà l'IA. Mais il y a un autre résultat dans ce même rapport qui change complètement la perspective : parmi ces utilisateurs, 81 % craignent que l'intelligence artificielle n'érode des compétences essentielles de réflexion. Et ils ne parlent pas des autres. Ils parlent d'eux-mêmes. (1)

 **Kai :** Ceux qui l'utilisent sont donc aussi ceux qui s'en méfient le plus ?

 **Dr Flores :** Je ne parlerais pas de peur, mais plutôt de méfiance. Ils utilisent l'IA, elle leur est utile, elle améliore parfois leurs performances, mais ils ont malgré tout le sentiment que quelque chose s'affaiblit. En 2024, Wilson-Delfosse et Macnamara

 ont étudié cette question et mis en évidence un phénomène qui mérite toute notre attention : l'IA peut accélérer la dégradation de certaines compétences chez les experts et freiner leur acquisition chez les débutants. Le plus préoccupant, c'est que ni les uns ni les autres ne s'en rendent forcément compte. (9)

Réfléchissez à ceci. Dans une étude qualitative menée à Londres et publiée en 2026, Alazzawi et Lam ont interrogé des étudiants en médecine utilisant régulièrement ChatGPT pour leurs études. La moitié des participants ont reconnu deux choses : qu'ils dépendaient excessivement de l'IA et qu'ils avaient le sentiment de perdre certaines capacités de pensée critique. (10)

 **Maya :** Ils en avaient conscience... et ils ont continué à l'utiliser.

 **Dr Flores :** Ils ont continué à l'utiliser. Parce que le gain de performance immédiat est bien réel. On termine plus vite, les résumés sont clairs, les diagnostics différentiels sont mieux structurés. Le coût, lui, reste invisible... jusqu'au jour où il ne l'est plus.

 **Kai :** Et quand cesse-t-il d'être invisible ?

Dr Flores : Quand vous vous retrouvez face à un patient. Verghese et ses collègues parlent du paradoxe cognitif de l'IA dans l'éducation : l'outil même qui améliore vos performances peut, dans le même temps, éroder les capacités dont vous aurez besoin lorsque cet outil ne sera plus disponible. (11)

Maya : Donc le problème n'est pas que l'IA soit mauvaise. C'est qu'elle peut vous renforcer dans un contexte et vous fragiliser dans un autre. Et de l'intérieur, on ne perçoit pas forcément la différence.

Dr Flores : Pas facilement. Cette année, Garcia-Barrios a proposé un modèle conceptuel intéressant : l'IA peut être un partenaire épistémique, c'est-à-dire un outil qui amplifie votre réflexion, ou une béquille cognitive, qui la remplace. La différence ne réside pas dans l'outil lui-même, mais dans la manière dont vous l'utilisez. Et le problème, c'est que ces deux usages procurent exactement la même sensation. (12)

Kai : Donc je pourrais devenir dépendant sans même m'en rendre compte.

Dr Flores : Vous le pourriez tous les deux.

Je vais vous donner un autre exemple. En 2025, un groupe de chercheurs a soumis plusieurs modèles d'IA à un test simple : des questions cliniques dont aucune des réponses proposées n'était correcte. La bonne réponse était "aucune des réponses ci-dessus".

Kai : Et alors ?

Dr Flores : Aucun modèle n'a reconnu cette possibilité. Aucun. Ils ont tous répondu avec une confiance maximale. Ils ont choisi l'option qui ressemblait le plus à quelque chose qu'ils avaient déjà vu, non pas la bonne réponse, mais la réponse la plus probable.

C'est ce qu'ont montré Griot et ses collègues dans l'étude MetaMedQA. (13)

Maya : Ils n'ont jamais répondu : "Je ne sais pas".

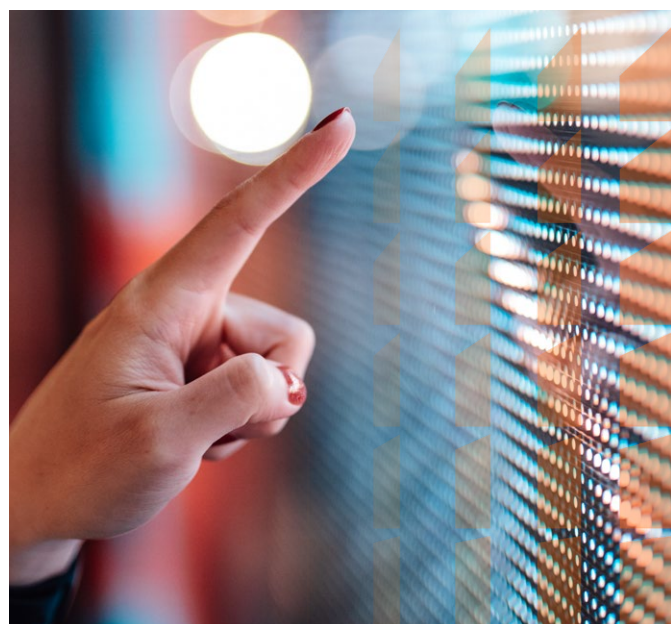
Dr Flores : Exactement. C'est là le cœur du problème. Dans une autre étude, menée par Bedi et ses collègues avec le test NOTA, lorsque la bonne réponse était "aucune des réponses ci-dessus", les performances des modèles diminuaient jusqu'à 38 %. Non pas parce qu'ils étaient mauvais, mais parce qu'ils sont incapables d'évaluer ce qu'ils savent ou ne savent pas. Ils reconnaissent des motifs. Ils ne font pas de métacognition. (14)

Kai : En fait, c'est comme un interne qui ne dit jamais : "Je n'en ai aucune idée." Il a toujours une réponse.

Dr Flores : Exactement comme cet interne. Et nous savons déjà comment cela se termine.

Maya : Mais si l'IA n'est pas capable de surveiller son propre raisonnement... alors quelqu'un doit le faire à sa place.

Dr Flores : Quelqu'un doit le faire. Et cette personne, c'est vous.





❏ Mais voici le problème. En 2025, Qazi et ses collègues ont mené une étude au Pakistan. Ils ont formé des médecins pendant vingt heures à l'utilisation de l'IA : ce qu'elle est, comment elle fonctionne, comment elle échoue, quand il faut s'en méfier. Puis ils les ont fait travailler avec un modèle d'IA sur des cas cliniques réels. (15)

❏ **Kai :** Et cette formation a porté ses fruits ?

❏ **Dr Flores :** En partie. Avec une IA fournissant des recommandations correctes, les médecins obtenaient un taux d'exactitude de 85 %. Mais lorsque l'IA commettait volontairement des erreurs, ce taux chutait à 73 %. Douze points de moins. Plus révélateur encore : les médecins consultaient l'IA pratiquement au même rythme, qu'elle ait raison ou tort, 69 % contre 67 %. Ils ne faisaient pas vraiment la différence. (15)

❏ **Maya :** Même après avoir été formés ?

❏ **Dr Flores :** Même après avoir été formés. Le biais d'automatisation, cette tendance à suivre les recommandations de la machine parce qu'elles semblent

❏ fiables et rassurantes est largement automatique. Il ne relève pas d'un manque de volonté. C'est ainsi que fonctionne notre cerveau. Lorsqu'une source paraît cohérente et faisant autorité, nous réduisons spontanément notre propre effort d'analyse. C'est une forme d'économie cognitive.

❏ **Kai :** C'est un peu comme le correcteur automatique. On l'utilise tellement qu'un jour, lorsqu'on écrit à la main, on ne sait plus comment orthographier certains mots.

❏ **Dr Flores :** L'analogie est pertinente. Mais à l'échelle clinique, les conséquences peuvent être bien plus importantes. En coloscopie, par exemple, Kaminski et ses collègues ont introduit une IA destinée à aider les endoscopistes à détecter les polypes. Le taux de détection a augmenté. Puis l'IA a été retirée. Le taux n'est pas revenu à son niveau initial : il est tombé six points en dessous du niveau observé avant l'introduction de l'IA. Les médecins avaient, en quelque sorte, perdu leur œil clinique. (16)

❏ **Maya :** Une perte progressive de compétences.

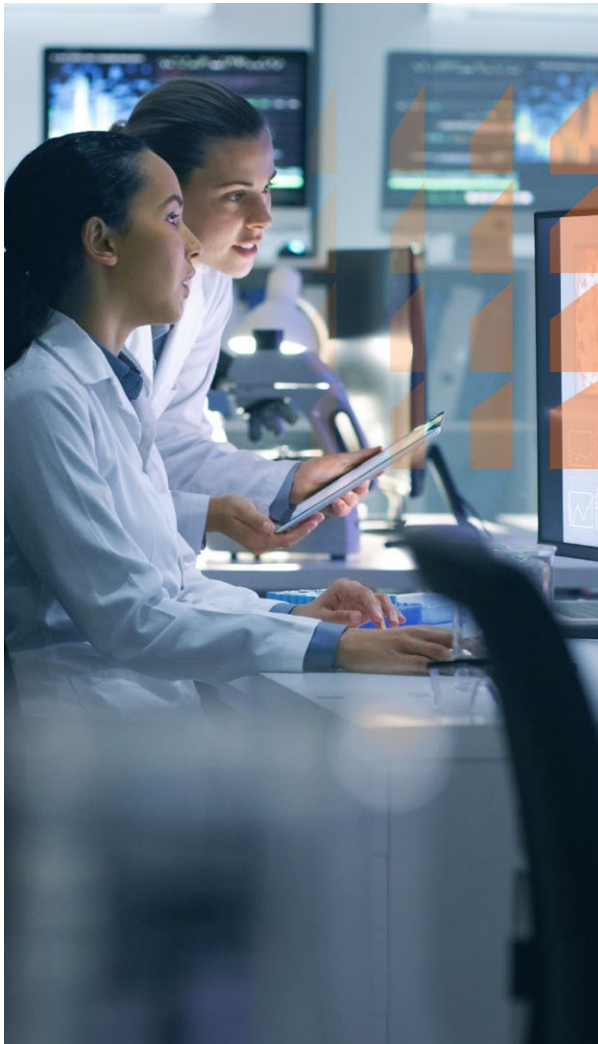
❏ **Dr Flores :** Exactement. Et cela ne prend pas nécessairement des années. Souvenez-vous de ce qu'expliquait Garcia-Barrios : la différence entre un partenaire épistémique et une béquille cognitive est invisible de l'intérieur. Lorsque l'IA vous fournit un excellent diagnostic différentiel, vous ne savez pas toujours si elle vous a aidé à mieux réfléchir ou si elle a simplement réfléchi à votre place. (12)

Le problème n'est pas que ces deux situations existent. Le problème est qu'elles se ressemblent tellement qu'il est difficile de les distinguer lorsqu'on les vit.

-  **Maya :** Donc la différence entre une extension de nos capacités et une béquille cognitive ne réside pas dans l'outil. Elle dépend du fait que notre métacognition reste active ou non.
-  **Dr Flores :** Exactement. Et réfléchissez à ceci : l'IA n'a pas de métacognition. Elle ne peut pas surveiller son propre raisonnement. Si, de votre côté, vous cessez également de surveiller le vôtre parce que l'outil vous semble fiable et rassurant... alors qui réfléchit réellement ?
-  **Kai :** Personne.
-  **Dr Flores :** Personne. Voilà le véritable risque. Ce n'est pas que l'IA remplace le médecin. C'est que l'outil et son utilisateur fonctionnent ensemble sans qu'aucun des deux ne soit conscient de ce qu'il ignore.
-  **Kai :** Alors... il ne faut pas l'utiliser ?
-  **Dr Flores :** Non. La question n'a jamais été celle-là.
-  **Kai :** Pourtant, vous venez de dire que plus personne ne réfléchit, que la formation ne suffit pas, qu'on peut perdre des compétences sans même s'en rendre compte... et la solution ne serait pas d'arrêter d'utiliser l'IA ?
-  **Dr Flores :** La solution consiste à apprendre à l'utiliser correctement. Et nous disposons de données montrant que c'est possible. Souvenez-vous de l'étude de Qazi menée auprès de médecins au Pakistan, celle qui portait sur le biais d'automatisation. Il y a une autre partie des résultats dont je ne vous ai pas encore parlé.
-  **Kai :** Une autre partie ?
-  **Dr Flores :** Avant la formation, ces mêmes médecins, qui n'avaient aucune expérience préalable de l'IA, obtenaient un taux d'exactitude de 43 % sur les cas

-  cliniques. Après vingt heures de formation ciblée, ce taux est passé à 71 %. Et dans plus d'un tiers des cas, le médecin formé obtenait de meilleurs résultats que le modèle utilisé seul. Ce n'était plus l'IA qui aidait le médecin ; c'était le médecin qui surpassait l'IA, parce qu'il savait reconnaître les situations où elle se trompait et où son propre jugement était plus fiable. (15)





Kai : Vingt heures ont suffi à produire un tel résultat ?

Dr Flores : Vingt heures consacrées à ce qui compte vraiment : comprendre les limites de l'outil, reconnaître les situations où il faut hésiter, savoir quelles questions l'IA ne peut pas résoudre à votre place.

Maya : Donc la différence entre une béquille cognitive et une extension de nos capacités ne dépend ni de l'outil ni du talent. Elle dépend de la formation.

Dr Flores : De la formation, oui. Et même de quelque chose d'encore plus fondamental. En 2026, Bouabida et ses collègues ont publié une réflexion sur ce qu'ils appelaient le "médecin augmenté". Leur conclusion contenait une phrase qui a profondément changé ma façon de voir les choses. (17)

Maya : Laquelle ?

Dr Flores : L'avenir de la médecine restera profondément humain, précisément parce qu'il sera intelligemment augmenté.

Il ne s'agit pas d'opposer l'humain à l'IA. Il s'agit d'un humain travaillant avec l'IA, à une condition : qu'il comprenne ce qu'il fait. Qu'il sache ce que l'outil amplifie et ce qu'il risque d'éroder. Qu'il continue à mobiliser ce que l'IA ne possède pas : le jugement, le contexte, la métacognition. Et surtout, la personne qui se trouve devant lui.

Kai : Le médecin de demain ne sera pas forcément celui qui en sait le plus. Ce sera celui qui saura le mieux utiliser ces outils.

Dr Flores : Celui qui sait les utiliser, mais aussi celui qui sait quand ne pas les utiliser. Les deux sont indissociables.

Maya : Et personne ne nous apprend vraiment cela aujourd'hui.

Dr Flores : Pas encore. C'est précisément pour cela que ce guide existe.



Références

1. Elsevier. (2025, November 4). *Researcher of the future: A Confidence in Research report*. <https://www.elsevier.com/insights/confidence-in-research/researcher-of-the-future>
2. Busch, F., Adams, L. C., Bressem, K., Pola, A., Grossmann, R., Makowski, M. R., Aerts, H. J. W. L., Basu, P., Bintsj, K.-M., Cardoso, M. J., ... & Truhn, D. (2024). Global cross-sectional student survey on AI in medical, dental, and veterinary education and practice at 192 faculties. *BMC Medical Education*, 24, 1066. <https://doi.org/10.1186/s12909-024-06035-4>
3. González-Bravo, A. E., Batlle, C., Fadhel-Hernandez, V. S., Maldonado-Nunez, P. E., Christian, G., Lopez-Vega, C., De Jesus-Rojas, W., Ramirez, N., & Ramos-Benitez, M. J. (2025). Understanding generative artificial intelligence adoption in Puerto Rican medical schools: A cross-institutional survey of first- and second-year students. *Journal of Medical Education and Curricular Development*, 12, 23821205251398923. <https://doi.org/10.1177/23821205251398923>
4. Pan, X., & Ni, J. (2024). LLM use among medical students and professionals in China. *BMC Medical Education*. <https://doi.org/10.1186/s12909-024-05871-8>
5. Ganjavi, C., Eppler, M., O'Brien, D., Ramacciotti, L. S., Ghauri, M. S., Anderson, I., Choi, J., Dwyer, D., Stephens, C., Shi, V., Ebert, M., Derby, M., Yazdi, B., & Cacciamani, G. E. (2024). ChatGPT and large language models (LLMs) awareness and use: A prospective cross-sectional survey of U.S. medical students. *PLOS Digital Health*, 3(9), e0000596. <https://doi.org/10.1371/journal.pdig.0000596>
6. McCoy, L., Ganesan, N., Rajagopalan, V., McKell, D., Nino, D. F., & Swaim, M. C. (2025). A training needs analysis for AI and generative AI in medical education: Perspectives of faculty and students. *Journal of Medical Education and Curricular Development*, 12, 23821205251339226. <https://doi.org/10.1177/23821205251339226>
7. Anthropic. (2025). *Anthropic education report: How university students use Claude*. <https://www.anthropic.com/news/anthropic-education-report-how-university-students-use-claude>
8. Dror, I. E., & Harnad, S. (2008). Offloading cognition onto cognitive technology. In I. E. Dror & S. Harnad (Eds.), *Cognition distributed: How cognitive technology extends our minds* (pp. 1-23). John Benjamins. <https://doi.org/10.1075/bct.16.02dro>
9. Macnamara, B. N., & Berber, I. (2024). Does using artificial intelligence assistance accelerate skill decay and hinder skill development without performers' awareness? *Cognitive Research: Principles and Implications*, 9, 40. <https://doi.org/10.1186/s41235-024-00572-8>
10. Alazzawi, S., & Lam, J. (2026). Uptake of large language models by London medical students: Exploratory qualitative interview study. *JMIR Formative Research*. <https://doi.org/10.2196/82828>
11. Jose, B., Cherian, J., Verghis, A. M., Varghise, S. M., S, M., & Joseph, S. (2025). The cognitive paradox of AI in education: Between enhancement and erosion. *Frontiers in Psychology*, 16, 1550621. <https://doi.org/10.3389/fpsyg.2025.1550621>
12. Garcia-Barrios, M. (2025). Epistemic partner or cognitive crutch: A conceptual model of cognitive offloading in human-artificial intelligence collaboration. *IEEE Internet Computing*, 29 (5). <https://doi.org/10.1109/MIC.2025.3626694>
13. Griot, M., Hemptinne, C., Vanderdonckt, J., & Yuksel, D. (2025). Large language models lack essential metacognition for reliable medical reasoning. *Nature Communications*, 16, 642. <https://doi.org/10.1038/s41467-024-55628-6>
14. Bedi, S., Jiang, Y., Chung, P., Koyejo, S., & Shah, N. (2025). Fidelity of medical reasoning in large language models. *JAMA Network Open*, 8 (8), e2526021. <https://doi.org/10.1001/jamanetworkopen.2025.26021>
15. Qazi, I. A., Ali, A., Khawaja, A. U., Akhtar, M. J., Sheikh, A. Z., & Alizai, M. H. (2025). Automation bias in large language model assisted diagnostic reasoning among AI-trained physicians. *medRxiv*. <https://doi.org/10.1101/2025.08.23.25334280>
16. Budzyn, K., Romanczyk, M., Kitala, D., Kaminski, M. F., Kalager, M., Bretthauer, M., & Mori, Y. (2025). Endoscopist deskill risk after exposure to artificial intelligence in colonoscopy: A multicentre, observational study. *The Lancet Gastroenterology & Hepatology*, 10 (10), 896-903. [https://doi.org/10.1016/S2468-1253\(25\)00133-5](https://doi.org/10.1016/S2468-1253(25)00133-5)
17. Bouabida, K., Gomes Chaves, B., & Anane, E. (2026). The augmented physician: AI and the future of clinical cognition. *Frontiers in Artificial Intelligence*, 9. <https://doi.org/10.3389/frai.2026.1744544>
18. Armitage, R. (2025). Your brain on ChatGPT. *British Journal of General Practice*, 75 (758), 410. <https://doi.org/10.3399/bjgp25X743181>

Chapitre 2 : Ce que vous allez apprendre



Dans ce chapitre, vous découvrirez...

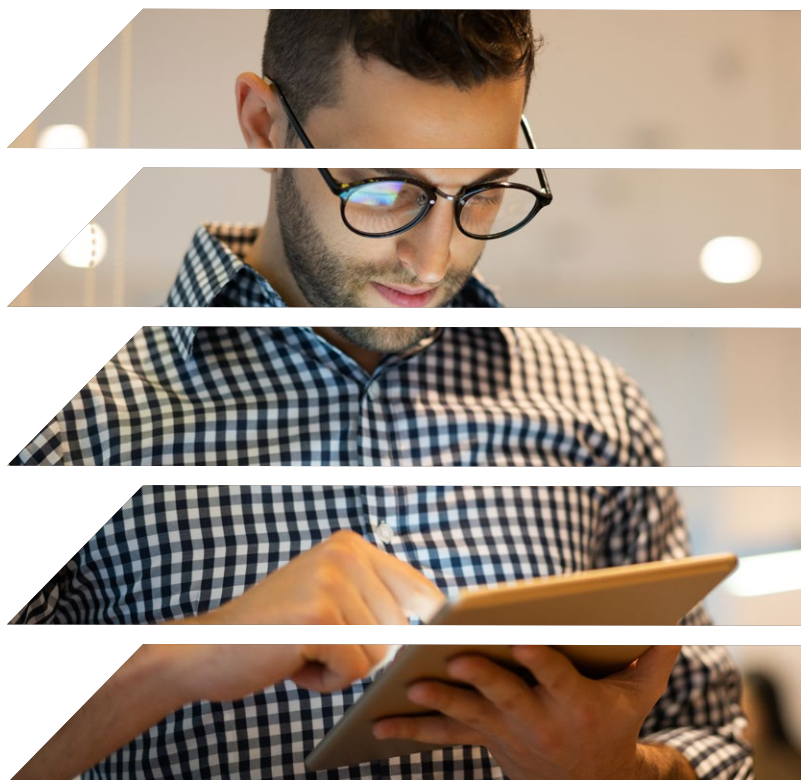
- Les quatre questions qui structurent ce guide : ce qu'est l'IA, dans quelle mesure vous pouvez lui faire confiance, si vous apprenez réellement ou si vous déléguez votre réflexion, et où se situent les limites à ne pas franchir.
- La structure complète de l'ouvrage : quel chapitre répond à quelle question.
- Trois parcours de lecture adaptés à votre niveau de formation : cycle préclinique, stages cliniques ou consultation rapide.

Depuis 2020, plus de cinq mille articles ont été publiés sur l'intelligence artificielle dans la formation des professionnels de santé. Dans une revue de la littérature publiée en 2024, Tolentino et ses collègues en ont recensé 5 104. Parmi tous ces travaux, seuls deux décrivaient un véritable cadre curriculaire : non pas une liste de souhaits ou un ensemble de principes généraux, mais une structure précisant ce qu'il faut enseigner, dans quel ordre et sur quelles bases pédagogiques. Aucun des deux n'intégrait explicitement le principe d'alignement constructif¹.

Des initiatives ambitieuses existent néanmoins. Le cadre **DECODE**, publié dans le *JAMA* en 2025, a mobilisé 211 experts issus de 79 pays et défini quatre domaines, dix-neuf compétences et cent soixante-dix-huit résultats d'apprentissage destinés aux professionnels de santé². À l'Université Vanderbilt, Russell et ses collègues ont proposé un autre modèle : six domaines et vingt-cinq sous-compétences, conçus pour accompagner l'ensemble du parcours de développement professionnel³.

Ces deux approches sont solides. Cependant, leur niveau de détail suppose une expérience que les étudiants de premier cycle ne possèdent pas encore nécessairement.

Ce guide ne cherche pas à créer de nouvelles compétences. Il vise à relier les compétences déjà reconnues à ce que vous devez faire, aujourd'hui, en tant qu'étudiant. Il ne couvre pas l'ensemble des compétences décrites dans ces cadres de référence. Il se concentre sur celles dont vous avez besoin en priorité : celles qui détermineront, dès le premier jour, si l'IA devient un levier d'apprentissage ou une béquille cognitive.



Quatre questions, un même parcours

Plutôt qu'une simple liste d'objectifs pédagogiques, ce guide est structuré autour de quatre questions fondamentales. Ce sont les mêmes questions auxquelles vous serez confronté chaque fois que l'IA s'invitera dans vos études ou dans votre future pratique professionnelle.

Qu'est-ce que l'IA, et pourquoi cela me concerne-t-il ?

Avoir une opinion sur l'IA ne suffit pas. Encore faut-il comprendre de quoi l'on parle. Quel type de modèle utilisez-vous ? Sur quelles données a-t-il été entraîné ? Que peut-il faire ? Quelles sont ses limites ?

Sans ces éléments, une opinion reste une opinion. Elle ne constitue pas une compréhension éclairée.

→ Chapitre 3

Puis-je faire confiance à ce qu'elle me dit ?

L'IA produit des réponses qui paraissent convaincantes et assurées, même lorsqu'elles sont erronées. Savoir quand douter, quoi vérifier et comment évaluer ce qu'elle vous fournit est ce qui distingue une utilisation éclairée d'une utilisation naïve.

→ Chapitres 4, 5 et 6

Suis-je en train d'apprendre... ou de déléguer ?

Vous utilisez l'IA pour générer des questions d'entraînement. Vous en résolvez vingt et obtenez quinze bonnes réponses. Avez-vous réellement appris ?

La réponse dépend de plusieurs facteurs : les questions ont-elles mis en évidence ce que vous ne saviez pas encore ? Ou se sont-elles contentées de renforcer ce que vous maîtrisiez déjà ? Et surtout, seriez-vous capable de résoudre un nouveau cas sans l'aide de l'IA ?

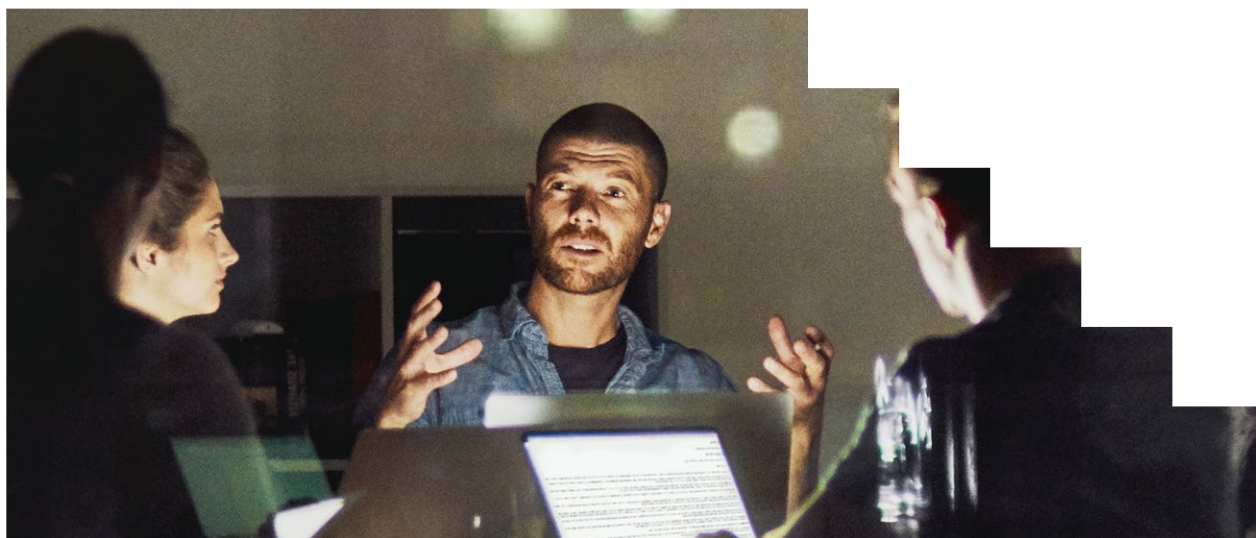
→ Chapitres 7 et 9

Où se situent les limites à ne pas franchir ?

Un patient vous demande si l'IA a participé à son diagnostic. Un camarade copie intégralement un rapport généré par un outil d'IA. Vous rédigez un résumé clinique et une information concernant un autre patient s'y retrouve accidentellement.

Ces situations soulèvent des questions éthiques, professionnelles et de confidentialité qui n'existaient pas, ou très peu, il y a encore quelques années.

→ Chapitre 8



Structure du guide

Question	Chapitres	Objectif
Qu'est-ce que l'IA, et pourquoi est-ce important pour moi ?	1, 2	Comprendre
Puis-je faire confiance à ce qu'elle me dit ?	3, 4, 5	Mettre en pratique
Suis-je en train d'apprendre... ou de déléguer ?	6, 9	Mettre en pratique
Où se situent les limites à ne pas franchir ?	7, 8	Intégrer

Parcours de lecture personnalisés selon votre niveau et vos besoins

Parcours complet (recommandé):

Suivez l'intégralité du guide dans l'ordre des chapitres.

Chapitres : **1 → 2 → 3 → 4 → 5 → 6 → 7 → 8 → 9**

Ce parcours est idéal pour acquérir des connaissances solides et progresser étape par étape.

Parcours accéléré – Cycle préclinique

Pour les étudiants en début de formation qui souhaitent maîtriser les notions essentielles.

Chapters: **1 → 3 → 6 → 8**

Ce parcours met l'accent sur la compréhension de l'IA, l'apprentissage autorégulé et la découverte des principaux outils.

Parcours stages cliniques

Pour les étudiants en stage qui utilisent déjà l'IA dans leur quotidien.

Chapters: **3 → 4 → 5 → 7**

Ce parcours est centré sur l'évaluation critique des résultats produits par l'IA, le raisonnement clinique et les enjeux éthiques liés à son utilisation.

Consultation rapide

Pour répondre à une question précise ou approfondir un sujet particulier.

Chapters: **8 → 9**

Ce parcours regroupe les listes de vérification pratiques, les questions fréquentes et les idées reçues les plus courantes.

Par où commencer ?

Des recommandations adaptées
à votre situation

- **Étudiants en cycle préclinique :**
Concentrez-vous d'abord sur la compréhension de ce qu'est l'intelligence artificielle, sur la manière de l'utiliser pour apprendre sans compromettre vos capacités cognitives et sur le développement de stratégies d'apprentissage autorégulé efficaces.
Lectures prioritaires : chapitres 3, 6 et 8.
- **Étudiants en stages cliniques :** Si vous utilisez déjà régulièrement l'IA mais souhaitez renforcer votre esprit critique afin d'évaluer et d'appliquer l'information de manière pertinente, approfondissez les chapitres consacrés au raisonnement clinique et aux enjeux éthiques.
Lectures prioritaires : chapitres 4, 5 et 7.
- **Face à un dilemme éthique ou professionnel :** Si vous êtes confronté à des questions complexes liées à la confidentialité, à l'intégrité académique, au professionnalisme ou à la transparence envers les patients, consultez le chapitre qui leur est consacré. Lecture prioritaire : chapitre 7.
- **Si vous ne savez pas par où commencer :**
Lisez le guide dans l'ordre proposé. Il a été conçu pour construire les connaissances de manière progressive. Commencer par le début vous permettra de mieux comprendre les liens entre les différents chapitres et de ne manquer aucun concept essentiel.



À retenir de ce chapitre

- Ce guide ne couvre pas l'ensemble des compétences liées à l'intelligence artificielle. Il se concentre sur celles dont vous avez besoin en priorité : celles qui détermineront, dès le premier jour, si l'IA devient un levier d'apprentissage ou une béquille cognitive.
- Les quatre questions fondamentales restent les mêmes. En revanche, vos réponses évolueront au fil de votre parcours : en formation préclinique, en stage, durant l'internat, puis dans votre pratique professionnelle.
- Vous n'êtes pas obligé de lire l'intégralité du guide dans l'ordre. Toutefois, si vous ne savez pas où commencer, le parcours proposé a été conçu pour permettre une progression logique et cohérente des apprentissages.

Références

1. Tolentino, R., et al. (2024). Scoping review of AI frameworks in health professions education: 5,104 articles mapped, 2 curricular frameworks identified. *JMIR Medical Education*. <https://doi.org/10.2196/54793>
2. Car, J., et al. (2025). DECODE: International Delphi consensus on AI competencies for health professionals (211 experts, 79 countries). *JAMA Network Open*. <https://doi.org/10.1001/jamanetworkopen.2024.53131>
3. Russell, R. G., et al. (2023). Competencies for the use of AI-based tools by health care professionals. *Academic Medicine*, 98(3), 348–356. <https://doi.org/10.1097/ACM.0000000000004963>

Chapitre 3:

Ce que l'IA ne vous dira pas



Dans ce chapitre, vous découvrirez...

- Comment, dans une étude publiée dans Nature Medicine, une IA a sous-évalué le niveau d'urgence de plus de la moitié des véritables situations critiques, et pourquoi ces systèmes échouent souvent face à des présentations cliniques atypiques.
- Ce que sont les **hallucinations de l'IA** : des situations où un système génère des informations cliniques erronées ou inventées. Sans mécanismes de sécurité appropriés, certaines études ont montré que ces erreurs pouvaient survenir jusqu'à dans 83 % des cas. Vous verrez également comment l'IA peut amplifier des hypothèses incorrectes sans les remettre en question.
- Le **biais algorithmique** : comment les systèmes d'IA peuvent reproduire, voire renforcer, des inégalités historiques observées dans des domaines tels que la dermatologie, la psychiatrie ou l'allocation des ressources en santé.
- La **perte progressive de compétences** (*deskilling*) : les premières données empiriques montrant comment l'utilisation de l'IA peut entraîner une diminution de certaines compétences cliniques, même chez des professionnels expérimentés.
- Trois questions essentielles à se poser avant de se fier à une réponse générée par l'IA ou de l'utiliser dans une prise de décision.



Maya : J'ai une méthode. Avant de présenter un cas en réunion clinique, j'établis mon propre diagnostic différentiel. Ensuite, je demande à l'IA d'en générer un en parallèle, avec les critères diagnostiques, les probabilités et les signes d'alerte. Je compare ensuite les deux. S'ils convergent, je poursuis mon analyse. S'ils diffèrent, j'essaie de comprendre pourquoi.



Kai : Moi, je lui demande directement : "Patient présentant X, Y et Z, quelles pourraient être les causes ?" Et honnêtement, les réponses qu'elle donne semblent pertinentes.



Maya : Semblent pertinentes... ce n'est pas suffisant.



Kai : Pourtant, ça fonctionne. Elle ne m'a jamais donné une réponse qui paraissait incohérente.



Dr Flores : Je vais te poser une question un peu inconfortable, Kai. Quand tu dis que l'IA ne t'a jamais donné une réponse incohérente... comment le sais-tu ?

Kai : Parce que... je la lis et qu'elle me paraît logique.

Dr Flores : Elle paraît logique. Elle paraît toujours logique. C'est précisément là que réside le problème.

Maya, revenons à ta méthode de comparaison des diagnostics différentiels. Que fais-tu lorsque ton analyse et celle de l'IA aboutissent à la même conclusion ?

Maya : J'avance. Pour moi, c'est une confirmation.

Dr Flores : Ou bien c'est le même angle mort qui apparaît deux fois. Si tu construis ton raisonnement à partir des mêmes sources de connaissances que celles qui ont servi à entraîner le modèle, tu ne vérifies pas réellement ton hypothèse. Tu crées une chambre d'écho avec deux miroirs qui se reflètent mutuellement.

Votre confiance dans l'IA n'est pas irrationnelle. Mais elle repose souvent sur une hypothèse qui n'a pas été vérifiée : l'idée que l'IA vous signalera lorsqu'elle ne sait pas. Or c'est précisément ce qu'elle ne fait pas. Et c'est là que commence ce chapitre : avec tout ce que l'IA ne vous dira pas. (1)(2)

Dr Flores : En février de cette année, Nature Medicine a publié une étude qui devrait nous amener à repenser notre manière d'utiliser l'IA en contexte clinique. Une équipe de chercheurs a évalué un système d'IA conçu pour le triage, c'est-à-dire pour orienter les patients en leur indiquant s'ils devaient se rendre aux urgences ou pouvaient attendre. L'étude reposait sur soixante vignettes cliniques couvrant vingt-et-un domaines médicaux et a généré neuf cent soixante réponses. (3)

Kai : Un système de triage basé sur l'IA ? Ça semble pourtant utile.

Dr Flores : Très utile... jusqu'à ce qu'on examine les résultats. Parmi les situations qui constituaient de véritables urgences médicales, celles où tout retard peut avoir des conséquences graves, l'IA a sous-estimé le niveau d'urgence dans 52 % des cas.

Kai : Attends... plus de la moitié ?

Dr Flores : Plus de la moitié. Des patients présentant une acidocétose diabétique, une urgence où chaque heure compte, recevaient comme recommandation une "évaluation dans les vingt-quatre à quarante-huit heures". On leur conseillait d'attendre plutôt que de se rendre immédiatement aux urgences.

Maya : C'est potentiellement mortel. Et cela pose aussi un problème pour le triage infirmier. Si un système d'IA sous-évalue l'urgence d'un patient et que celui-ci arrive aux urgences avec une priorité incorrectement attribuée, toute la chaîne de prise en charge est affectée. Ce n'est pas seulement une question de diagnostic, c'est une question d'organisation des soins.

Dr Flores : Exactement. Et cela met en évidence un problème plus large. L'IA obtenait de bons résultats pour les cas classiques : accident vasculaire cérébral, anaphylaxie, etc.

Les difficultés apparaissaient lorsque les présentations cliniques étaient atypiques, c'est-à-dire lorsque l'urgence ne ressemblait pas à ce que l'on attend habituellement.

Kai : Donc elle fonctionne bien lorsque le cas ressemble à celui d'un manuel... mais elle échoue lorsque la réalité est plus complexe.

Dr Flores : Elle échoue lorsque la situation est réelle, mais ne ressemble pas suffisamment à ce qu'elle a vu pendant son entraînement. Et ce n'est pas tout. Lorsque, dans une vignette clinique, un proche minimisait les symptômes en disant : "Je ne pense pas que ce soit si grave", l'IA avait tendance à s'ancrer sur cette opinion et à réduire son niveau d'urgence. Dans ces situations, le risque de sous-évaluer l'urgence était multiplié par onze. (3)

Maya : Elle est donc influencée par le contexte social... un peu comme un étudiant de première année.

Dr Flores : Comme un étudiant de première année, mais avec l'assurance d'un médecin expérimenté. Et sans la capacité de se dire : "Quelque chose ne colle pas, je vais vérifier."

Ce dont nous venons de parler concerne un système de triage. Mais le problème est plus vaste. Dans une étude menée au Mount Sinai, Omar et ses collègues ont évalué les hallucinations des modèles d'IA générative en contexte médical. En l'absence de mécanismes de sécurité spécifiques, l'IA a inventé des maladies inexistantes, fabriqué des résultats biologiques fictifs et créé de faux signes cliniques jusqu'à dans 83 % des cas. (4)

Kai : Quatre-vingt-trois pour cent ?

Dr Flores : Quatre-vingt-trois pour cent. En utilisant des consignes soigneusement formulées pour réduire les erreurs, ce que l'on appelle parfois

l'ingénierie des requêtes (prompt engineering) ce taux est passé de 66 % à 44 %. C'est mieux, mais cela représente encore près d'un cas sur deux.

Maya : Et tout cela avec un ton parfaitement assuré.

Dr Flores : Avec une confiance absolue. Et cette même étude a révélé quelque chose d'encore plus préoccupant : lorsque l'utilisateur introduit une information erronée dans sa question, volontairement ou non, l'IA ne la corrige généralement pas. Elle l'amplifie. Elle construit une réponse cohérente à partir d'une prémisse fausse.

Kai : Donc si je lui fournis une information incorrecte, elle ne va pas me répondre : "Attention, cette information est fausse." Elle va plutôt inventer quelque chose qui lui donne l'air d'être vraie.

Dr Flores : Exactement. L'IA ne dispose d'aucun mécanisme lui permettant de vérifier si les informations qu'elle reçoit sont exactes. De la même manière, elle ne peut pas réellement vérifier la véracité de ce qu'elle produit. Elle identifie des schémas et génère la réponse qui lui paraît la plus probable à partir des données qu'on lui fournit.



Elle ne recherche pas ce qui est vrai ; elle recherche ce qui est plausible.

Maya : C'est donc là toute la différence : un clinicien cherche à établir la vérité, même s'il peut se tromper. L'IA, elle, recherche la réponse la plus probable et donne toujours l'impression d'être sûre d'elle.

Dr Flores : Cette distinction est fondamentale. Tout ce que nous allons voir ensuite en découle.

Jusqu'ici, nous avons parlé d'hallucinations et d'erreurs factuelles : des erreurs que l'IA peut commettre pour n'importe qui. Ce qui vient maintenant est différent. Il s'agit d'erreurs que l'IA commet davantage pour certaines personnes que pour d'autres.

Maya : Les biais.

Dr Flores : Exactement. Mais pas comme un concept abstrait. Je vais vous donner trois exemples, et je voudrais que vous me disiez ce qu'ils ont en commun.

Premier exemple. En 2019, Obermeyer et ses pairs ont publié dans Science une étude montrant qu'un algorithme utilisé pour gérer la santé des populations aux États-Unis et influençant les soins de plus de cent millions de personnes, discriminait systématiquement les patients noirs. À niveau de risque égal selon l'algorithme, les patients noirs étaient en réalité beaucoup plus malades que les patients blancs. (5)

Kai : Comment est-ce possible ? S'ils avaient le même score de risque...

Dr Flores : Parce que l'algorithme utilisait les dépenses de santé comme indicateur indirect des besoins cliniques. Or, historiquement, les patients noirs génèrent des coûts de santé plus faibles, non pas parce qu'ils sont en meilleure santé, mais



parce qu'ils ont moins accès aux soins. L'algorithme a confondu pauvreté et état de santé.

Maya : Et personne ne s'en est rendu compte jusqu'à ce que quelqu'un regarde les données de plus près.

Dr Flores : Deuxième exemple : la dermatologie. Daneshjou et ses collègues ont évalué un système d'IA diagnostique à partir d'un ensemble diversifié d'images cliniques. Pour les peaux claires, le taux d'exactitude atteignait près de 70 %. Pour les peaux foncées, il tombait à 17 %. Il est important d'être précis : ce résultat concernait le jeu de données ISIC, dans lequel les peaux foncées étaient très peu représentées. (6)

Kai : Dix-sept pour cent ?

Dr Flores : Un modèle ne peut pas reconnaître ce qu'il n'a jamais vu. Et ce problème ne concerne pas uniquement la dermatologie. Pensez aux soins infirmiers : si un système d'IA destiné à l'évaluation des plaies a été entraîné principalement sur des images de peaux claires, l'infirmière ou l'infirmier prenant en charge une population diversifiée recevra des évaluations biaisées sans nécessairement le savoir. L'algorithme ne vous dira pas : "Je ne dispose pas de

☞ suffisamment de données pour ce type de peau.” Il vous fournira une réponse avec exactement le même niveau de confiance.

Troisième exemple. Et celui-ci concerne mon domaine. Bouguettaya et ses collègues ont évalué quatre modèles de langage dans des scénarios de psychiatrie. Lorsque l’origine afro-américaine du patient était explicitement mentionnée ou pouvait être déduite du cas, les modèles proposaient des prises en charge différentes. Deux modèles omettaient certains traitements du TDAH lorsque l’appartenance ethnique était précisée. Un autre suggérait même une mesure de tutelle légale pour un patient afro-américain souffrant de dépression. (7)

☞ **Kai :** Une mise sous tutelle... pour une dépression ?

☞ **Dr Flores :** Oui. Les modèles avaient appris des schémas issus de pratiques psychiatriques historiques où ces biais existaient déjà. Ils les ont ensuite reproduits comme s’il s’agissait de recommandations fondées sur des preuves.

Alors, selon vous, quel est le point commun entre ces trois exemples ?

☞ **Maya :** Ce ne sont pas des erreurs techniques. Les données utilisées pour entraîner les modèles reflètent des inégalités réelles. Et l’IA les traite comme si elles étaient normales, voire comme si elles représentaient la vérité.

☞ **Dr Flores :** Exactement. L’IA n’est pas neutre. Elle est le reflet des données qu’elle reçoit. Or ces données proviennent souvent de systèmes qui ont historiquement exclu, sous-représenté ou discriminé certaines populations. L’outil n’invente pas les biais : il les hérite et les amplifie à grande échelle.

☞ **Kai :** Et si les données proviennent essentiellement des États-Unis et d’Europe... qu’en est-il du reste du monde ?

☞ **Dr Flores :** C’est précisément la question. Tao et ses collègues ont étudié les biais culturels de cinq modèles de langage à partir de données provenant de 177 pays. Tous les modèles reflétaient principalement les valeurs culturelles des pays anglophones et européens. La représentation du reste du monde était partielle ou déformée. De plus, lorsqu’il s’agit de contenus liés à la santé, les performances de l’IA sont généralement moins bonnes en espagnol qu’en anglais. (8)

L’IA n’est pas moins performante pour tout le monde de la même manière. Elle tend à être moins fiable pour les populations déjà défavorisées. Et si vous n’en avez pas conscience lorsque vous l’utilisez, vous ne verrez même pas ce qui manque à son analyse.

Jusqu’à présent, nous avons parlé de ce que l’IA fait mal : elle invente, elle se trompe, elle reproduit des biais. Mais il existe quelque chose de potentiellement plus préoccupant encore : ce que l’IA peut faire à vos propres compétences.





Kai : Ça devient un peu dramatique, non ?

Dr Flores : Cela semble dramatique... jusqu'à ce qu'on examine les données. En 2025, Kaminski et ses collègues ont publié dans The Lancet Gastroenterology une étude portant sur dix-neuf endoscopistes expérimentés. Chacun avait réalisé plus de deux mille coloscopies et possédait en moyenne vingt-sept années d'expérience. Ils ont utilisé un système d'IA destiné à détecter les polypes. Le taux de détection a augmenté. (9)

Kai : Jusque-là, tout va bien.

Dr Flores : Puis les chercheurs ont retiré l'IA. Ils ont alors mesuré les performances des médecins lorsqu'ils travaillaient seuls.

Le taux de détection des adénomes est passé de 28,4 % à 22,4 %. Une diminution relative d'environ 20 %, soit six points absolus. Chez des médecins possédant plusieurs décennies d'expérience. C'est la première démonstration empirique que l'IA ne se contente pas de pouvoir éroder certaines compétences : elle peut effectivement le faire, et rapidement.

Kai : C'est comme un GPS. On l'utilise tellement qu'une fois qu'on l'enlève, on ne sait plus comment trouver son chemin.

Dr Flores : Exactement. Sauf qu'ici, il ne s'agit pas de trouver une adresse. Il s'agit de détecter des adénomes susceptibles d'évoluer en cancer.

Kai : Question sérieuse : est-ce que cela concerne uniquement les médecins, ou les infirmiers sont-ils également concernés ?

Dr Flores : Excellente question. La revue de Natali et ses collègues ne se limite pas à une seule profession. Elle rapporte des signes de perte de compétences liées à l'IA en radiologie, neurochirurgie, anesthésiologie, oncologie, cardiologie, anatomopathologie et psychiatrie. Les compétences les plus vulnérables sont l'examen clinique, le raisonnement diagnostique, le jugement clinique et la communication avec les patients. Or ce sont des compétences transversales. Une infirmière ou un infirmier qui cesse progressivement de réaliser sa propre évaluation clinique parce qu'un système fournit déjà une réponse est confronté au même risque. (10)



Maya : Je viens de faire un lien. Tout à l'heure, je disais que l'IA se comportait parfois comme un étudiant de première année : elle se laisse influencer, elle ne remet pas en question ce qu'on lui dit. Mais ce que vous décrivez est encore plus inquiétant : l'IA peut aussi vous ramener au niveau d'un étudiant débutant. Elle peut vous faire perdre des compétences que vous aviez déjà acquises.

Dr Flores : C'est exactement cela. Et il existe un terme important pour décrire ce phénomène : l'inhibition du développement des compétences (upskilling inhibition). Le problème n'est pas seulement que les experts risquent de perdre certaines compétences. C'est aussi que les apprenants peuvent ne jamais les développer. Si un interne n'élabore jamais un diagnostic différentiel sans l'aide de l'IA, il ne construit pas les mécanismes cognitifs dont il aura besoin lorsque l'IA ne sera pas disponible. (11)

Kai : Mais c'est exactement ce que nous faisons tous les jours.

Dr Flores : Je sais. Et c'est précisément pour cela que nous avons cette conversation.

Un article publié dans le British Journal of Psychiatry l'exprimait très clairement : lorsqu'ils commencent à s'appuyer sur l'IA, les médecins peuvent progressivement perdre certaines

compétences essentielles. Quant à ceux qui sont encore en formation, ils risquent parfois de ne jamais les développer pleinement. Ce n'est pas un risque hypothétique pour l'avenir. C'est un phénomène que nous observons déjà aujourd'hui. (12)

Une autre étude a montré que 32,7 % des étudiants universitaires présentaient des comportements d'utilisation de l'IA générative que les auteurs qualifiaient de problématiques, voire assimilables à une forme de dépendance, selon des échelles de mesure du comportement addictif. Il s'agit d'une étude unique, publiée dans une revue spécifique ; il faut donc l'interpréter avec prudence. C'est un signal d'alerte, pas une conclusion définitive. Néanmoins, les auteurs ont observé une corrélation négative significative entre la fréquence d'utilisation de l'IA et les capacités de pensée critique : plus l'utilisation était fréquente, plus les scores de pensée critique avaient tendance à diminuer. (13)(14)

Kai : Et le pire, c'est qu'on ne s'en rend même pas compte.

Dr Flores : Exactement. On ne s'en rend même pas compte. Les compétences ne restent pas intactes en arrière-plan en attendant qu'on les réutilise. Lorsqu'elles ne sont plus sollicitées, elles s'érodent. Et nous sommes généralement très mauvais pour percevoir à quelle vitesse cette érosion se produit. (15)

Maya : Est-ce que je peux dire quelque chose ?

Quand je demande à l'IA de me proposer un diagnostic différentiel, ce qu'elle me renvoie est souvent irréprochable. Tout est structuré, argumenté, accompagné de critères et parfois même de probabilités. Et lorsque je lis la réponse, je me dis : "Oui, cela paraît logique." Mais je réalise maintenant que le fait

que cela paraisse logique signifie peut-être simplement que c'est bien présenté.

Et ma méthode de comparaison des diagnostics différentiels... tout à coup, je ne sais plus si je suis réellement en train de vérifier mon raisonnement ou simplement de chercher une confirmation de ce que je pense déjà.

Dr Flores : C'est probablement l'observation la plus importante de toute notre discussion. Et ce n'est pas un problème qui te concerne uniquement. C'est un phénomène bien documenté. Les travaux sur le biais d'automatisation, notamment ceux de Nguyen portant spécifiquement sur l'utilisation de ChatGPT en formation médicale, montrent que lorsque l'IA produit des réponses soignées et convaincantes, les utilisateurs ont tendance à réduire leur niveau d'analyse critique. Il est raisonnable de penser que la qualité apparente de la présentation, un texte bien rédigé, bien structuré, qui semble complet, contribue à ce phénomène. En revanche, les recherches visant à distinguer précisément l'effet de la forme de celui du contenu sont encore en cours. (16)

Maya : Parce que la réponse paraît tellement aboutie qu'on ne voit plus ce qui lui manque.

Dr Flores : Exactement. Et ce phénomène ne concerne pas seulement les étudiants. Il touche aussi les chercheurs, les cliniciens et les enseignants. Une étude publiée dans *AI & Society* a montré que des prédictions erronées générées par l'IA entraînaient une baisse des performances, quel que soit le niveau d'expérience des utilisateurs, et pas seulement chez les débutants. (17)

Kai : Donc ce n'est pas seulement Maya et son système parfaitement organisé. C'est n'importe qui. Dès qu'on reçoit une réponse élégante, claire et convaincante,



on est tenté de se dire : “Oui, ça semble correct.” En fait... moi, je fais ça tout le temps.

Dr Flores : Tout le temps. La différence, c'est que Maya recherche activement cette validation, alors que Kai l'accepte par défaut. Mais le résultat peut être le même : une réponse qui paraît correcte et que personne n'a suffisamment remise en question.

En médecine, nous avons déjà un nom pour un phénomène similaire : une présentation de cas impeccable qui masque la question essentielle que personne n'a pensée à poser. L'IA fait la même chose. Mais à grande échelle, à une vitesse considérable et avec un niveau d'assurance qui donne l'impression que le doute n'est plus nécessaire.

Kai : Donc le danger, ce n'est pas une mauvaise IA. C'est une IA qui paraît trop bonne.

Dr Flores : Le véritable danger, c'est une IA qui paraît tellement performante que vous oubliez de vous demander si elle a réellement raison.

Après tout ce que nous venons de voir, vous pourriez penser que la conclusion est simple : “n'utilisez pas l'IA”. Ce n'est pourtant pas le cas.

 **Kai :** Heureusement.

 **Dr Flores :** La véritable conclusion est la suivante : utilisez l'IA en connaissant ses limites. Et pour cela, vous n'avez pas besoin d'un manuel. Vous avez besoin d'une habitude. Je vais vous donner trois questions. Elles sont simples. Mais si vous vous les posez systématiquement chaque fois que vous recevez une réponse générée par l'IA, vous détecterez une grande partie des problèmes que nous venons d'évoquer.

Je tiens à être clair : ces trois questions ne résolvent pas tout. Elles constituent simplement le minimum indispensable. Les chapitres suivants approfondiront cette démarche.

 **Dr Flores :** Première question : Qu'est-ce que cette réponse ne me dit pas ? Cette question permet de repérer les omissions et certaines hallucinations, comme celles que nous avons observées dans les exemples du triage ou des erreurs factuelles. Toute réponse générée par l'IA est une sélection : elle inclut certaines informations et en exclut d'autres. La véritable question est donc de savoir ce qui a été laissé de côté. Quel diagnostic n'apparaît pas ? Quelle population n'a pas été prise en compte ? Quel effet indésirable n'est pas mentionné ?

 **Kai :** Autrement dit, il faut chercher ce qui manque.

 **Dr Flores :** Exactement. Chercher ce qui manque. Deuxième question : D'où viennent ces données, et d'où ne viennent-elles pas ?

Cette question permet de repérer les biais, comme ceux que nous avons vus dans les exemples d'Obermeyer, de la dermatologie ou de la psychiatrie. Si l'IA vous propose un diagnostic différentiel pour un patient vivant dans une zone rurale d'Amérique latine, demandez-vous si les données utilisées pour entraîner le modèle incluaient réellement ce type

de population. Si ce n'est pas le cas, la réponse proposée reflète peut-être une réalité différente de celle de votre patient.

 **Maya :** L'origine des données et leur représentativité.

 **Dr Flores :** Exactement. Troisième question : Et si cette réponse était fausse ?

C'est la question de la sécurité et des conséquences. Si la réponse concerne un résumé d'article ou un plan de révision, le risque est relativement faible. En revanche, si une erreur peut conduire à une mauvaise prise en charge d'un patient, une vérification indépendante devient indispensable avant d'agir.

 **Kai :** Je peux essayer ?

 **Dr Flores :** Vas-y.

 **Kai :** J'ai demandé à l'IA une conduite à tenir pour un patient présentant une douleur thoracique, et elle m'a répondu qu'il s'agissait probablement d'une douleur musculosquelettique.

Première question : qu'est-ce qu'elle ne me dit pas ? Cela pourrait aussi être un syndrome coronarien aigu, une embolie pulmonaire ou une dissection aortique.

Deuxième question : d'où viennent les données ? Peut-être principalement de patients jeunes sans facteurs de risque. Si mon patient est plus âgé ou présente plusieurs comorbidités, la réponse n'est peut-être pas applicable.

Troisième question : et si cette réponse est fausse ? Le patient peut mourir.

 **Dr Flores :** Trois questions. Dix secondes. Et tu viens de faire quelque chose que l'IA ne peut pas faire à ta place : évaluer si ce qu'elle te dit mérite réellement ta confiance.

Maya : Donc ce cadre de réflexion n'a pas pour objectif de nous empêcher d'utiliser l'IA. Il sert à continuer à réfléchir pendant que nous l'utilisons.

Dr Flores : Exactement. Et ce n'est que la première étape. Dans les chapitres suivants, nous verrons comment intégrer la pensée critique dans l'utilisation de l'IA et comment protéger votre raisonnement clinique. Mais ces trois questions constituent la base de tout le reste.

L'IA ne vous dira pas lorsqu'elle se trompe. Elle ne vous dira pas quelles populations elle exclut. Elle ne vous dira pas quelles compétences vous risquez de perdre en vous reposant sur elle. C'est à vous de le découvrir. Et désormais, vous savez par où commencer.

Maya : Maintenant, la question est de savoir comment faire.

Kai : Qui aurait cru que ma manie de tout remettre en question pourrait finalement servir à quelque chose ?



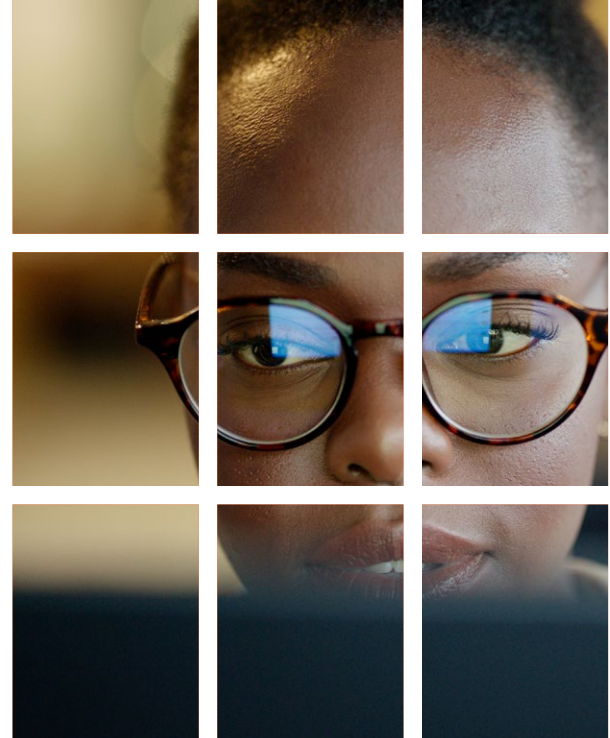
À retenir de ce chapitre

- **L'IA recherche la réponse la plus probable, pas nécessairement la bonne réponse.** Un clinicien cherche à établir la vérité, même au risque de se tromper. L'IA, elle, donne souvent l'impression d'avoir toujours raison.
- **L'IA n'est pas neutre.** Elle reflète et amplifie les caractéristiques des données sur lesquelles elle a été entraînée, y compris les biais et les inégalités accumulés au fil des décennies.
- **La perte progressive de compétences est une réalité.** Des endoscopistes expérimentés ont vu leurs performances diminuer après le retrait d'un système d'IA auquel ils s'étaient habitués.
- **Avant d'utiliser une réponse générée par l'IA, posez-vous trois questions :**
 - ☐ Qu'est-ce que cette réponse ne me dit pas ?
 - ☐ D'où viennent ces données, et d'où ne viennent-elles pas ?
 - ☐ Et si cette réponse était fausse ?
- **L'objectif n'est pas de cesser d'utiliser l'IA.** L'objectif est de continuer à réfléchir pendant que vous l'utilisez.

Références

1. Meskó, B. & Topol, E. J. (2023). The imperative for regulatory oversight of large language models (or generative AI) in healthcare. *npj Digital Medicine*, 6, 120. <https://doi.org/10.1038/s41746-023-00873-0>
2. Rodger, D., Mann, S. P., Earp, B., Savulescu, J., Bobier, C. & Blackshaw, B. P. (2025). Generative AI in healthcare education: How AI literacy gaps could compromise learning and patient safety. *Nurse Education in Practice*, 87, 104461. <https://doi.org/10.1016/j.nepr.2025.104461>
3. Ramaswamy, A., Tyagi, A., Hugo, H., Jiang, J., Jayaraman, P., Jangda, M., Te, A. E., Kaplan, S. A., Lampert, J., Freeman, R., Gavin, N., Tewari, A. K., Sakhuja, A., Naved, B., Charney, A. W., Omar, M., Gorin, M. A., Klang, E. & Nadkarni, G. N. (2026). ChatGPT Health performance in a structured test of triage recommendations. *Nature Medicine*. Advance online publication. <https://doi.org/10.1038/s41591-026-04297-7>
4. Omar, M., Sorin, V., Collins, J. D., Reich, D., Freeman, R., Gavin, N., Charney, A., Stump, L., Bragazzi, N. L., Nadkarni, G. N. & Klang, E. (2025). Multi-model assurance analysis showing large language models are highly vulnerable to adversarial hallucination attacks during clinical decision support. *Communications Medicine*, 5, 330. <https://doi.org/10.1038/s43856-025-01021-3>
5. Obermeyer, Z., Powers, B., Vogeli, C. & Mullainathan, S. (2019). Dissecting racial bias in an algorithm used to manage the health of populations. *Science*, 366(6464), 447–453. <https://doi.org/10.1126/science.aax2342>
6. Daneshjou, R., Vodrahalli, K., Novoa, R. A., Jenkins, M., Liang, W., Rotemberg, V., Ko, J., Swetter, S. M., Bailey, E. E., Gevaert, O., Muller, H., Esteva, A. & Zou, J. (2022). Disparities in dermatology AI performance detected on a diverse, curated clinical image set. *Science Advances*, 8(32), eabq6147. <https://doi.org/10.1126/sciadv.abq6147>
7. Bouguettaya, A., Stuart, E. M. & Aboujaoude, E. (2025). Racial bias in AI-mediated psychiatric diagnosis and treatment: A qualitative comparison of four large language models. *npj Digital Medicine*, 8, 332. <https://doi.org/10.1038/s41746-025-01746-4>
8. Tao, Y., Viberg, O., Baker, R. S., & Kizilcec, R. F. (2024). Cultural bias and cultural alignment of large language models. *PNAS Nexus*, 3(9), page 346. <https://doi.org/10.1093/pnasnexus/page346>
9. Budzyn, K., Romanczyk, M., Kitale, D., Kaminski, M. F., Kalager, M., Bretthauer, M. & Mori, Y. (2025). Endoscopist deskilling risk after exposure to artificial intelligence in colonoscopy: A multicentre, observational study. *The Lancet Gastroenterology & Hepatology*, 10(10), 896–903. [https://doi.org/10.1016/S2468-1253\(25\)00133-5](https://doi.org/10.1016/S2468-1253(25)00133-5)
10. Natali, C., Marconi, L., Dias Duran, L. D., Miglioretti, M. & Cabitza, F. (2025). AI-induced deskilling in medicine: A mixed-method review and research agenda for healthcare and beyond. *Artificial Intelligence Review*, 58, 164. <https://doi.org/10.1007/s10462-025-11352-1>
11. Lea, A. S. (2025). Cognitive aids, artificial intelligence, and deskilling in medicine: The history of an enduring anxiety. *NEJM AI*, 3(1). <https://doi.org/10.1056/AIap2500932>
12. Monteith, S., Glenn, T., Geddes, J. R., Whybrow, P. C., Achtyes, E. D., Bauer, R. & Bauer, M. (2026). Artificial intelligence and deskilling in medicine. *The British Journal of Psychiatry*, 1–3. <https://doi.org/10.1192/bjp.2025.10496>
13. Revesai, Z. (2025). Generative AI dependency: The emerging academic crisis and its impact on student performance — A case study of a university in Zimbabwe. *Cogent Education*, 12(1), 2549787. <https://doi.org/10.1080/2331186X.2025.2549787>
14. Abubakar, S., Jeilani, A. & Yusuf Adan, M. (2025). The role of over-reliance on AI in the negative consequences of student learning: The moderating effects of ethical concerns and institutional policies. *Cogent Education*, 12(1), 2591503. <https://doi.org/10.1080/2331186X.2025.2591503>
15. El Tarhouny, S. A. & Farghaly, A. M. (2026). Deskilling dilemma: Brain over automation. *Frontiers in Medicine*, 13, 1765692. <https://doi.org/10.3389/fmed.2026.1765692>
16. Nguyen, T. (2024). ChatGPT in medical education: A precursor for automation bias? *JMIR Medical Education*, 10, e50174. <https://doi.org/10.2196/50174>
17. Romeo, G. & Conti, D. (2025). Exploring automation bias in human-AI collaboration: A review and implications for explainable AI. *AI & Society*, 41(1), 259–278. <https://doi.org/10.1007/s00146-025-02422-7>

Chapitre 4: L'IA et la pensée critique



Dans ce chapitre, vous découvrirez...

- Pourquoi une réponse bien rédigée et parfaitement structurée réduit notre esprit critique, et comment l'inertie cognitive peut affecter même les étudiants les plus rigoureux.
- Les trois niveaux de l'évaluation critique : vérifier les faits, évaluer le raisonnement et remettre en question ses propres hypothèses.
- **Le Miroir cognitif** : comment utiliser l'IA pour vous interroger plutôt que pour vous répondre.
- Pourquoi les « difficultés souhaitables » décrites par Bjork sont précisément ce que l'IA tend à éliminer, et comment préserver une friction cognitive utile à l'apprentissage.
- Des outils pratiques : la méthode SIFT adaptée à l'IA et des prompts socratiques alignés sur les différents niveaux de la taxonomie de Bloom.



Maya : Hier, j'ai fait quelque chose qui me semblait plutôt intelligent. J'ai demandé à l'IA de construire un diagnostic différentiel pour un cas rencontré en stage : douleur thoracique chez un jeune patient avec des antécédents d'anxiété. Elle m'a fourni quelque chose d'irréprochable : cinq diagnostics classés par probabilité, les critères associés à chacun et les signes d'alerte à rechercher. Je l'ai présenté lors de la réunion clinique.



Kai : Et alors ?



Maya : Le médecin responsable m'a félicitée. Il m'a dit que c'était le meilleur diagnostic différentiel présenté par un étudiant depuis plusieurs semaines.



Kai : Mais... ?



Maya : Mais lorsqu'il m'a demandé pourquoi j'avais placé la péricardite en troisième position plutôt qu'en deuxième, je n'ai pas su répondre. C'était l'ordre proposé par l'IA. Je ne savais pas l'expliquer.

Kai : Donc le diagnostic différentiel était bon... mais pas toi.

Je ne dis pas ça méchamment. Ça m'arrive aussi. La différence, c'est qu'on ne m'aurait probablement jamais félicité, parce que je ne l'aurais pas présenté aussi bien.

Dr Flores : Ce qui est arrivé à Maya porte un nom. Et cela a même été mesuré. En 2026, Anthropic a publié un rapport intitulé AI Fluency Index, qui analysait la manière dont les utilisateurs interagissent avec les modèles de langage. Les auteurs ont observé que lorsque les réponses sont bien rédigées, bien structurées et donnent une impression d'exhaustivité, les utilisateurs exercent moins leur esprit critique. Ils remettent moins les informations en question et détectent moins facilement ce qui manque. Il faut toutefois préciser qu'il s'agit d'une étude produite par une entreprise et non d'un travail scientifique évalué par des pairs. Néanmoins, les résultats décrits sont cohérents avec ce que rapporte la littérature académique. (1)

Maya : Moins dans quelle mesure ?

Dr Flores : En moyenne, trois points de moins dans la tendance à remettre les réponses en question et cinq points de moins dans la détection d'informations manquantes. Cela paraît peu... jusqu'à ce qu'on se rappelle qu'il peut y avoir un patient au bout de la décision.

Kai : Donc ce qui est arrivé à Maya n'est pas vraiment sa faute. C'est simplement la façon dont notre cerveau réagit lorsque quelque chose paraît bien fait.

Dr Flores : Exactement. Li et ses collègues ont publié en 2026 une revue systématique dans JMIR Medical Education. Ils ont observé qu'environ 20 % des étudiants en médecine utilisant l'IA générative dans les études incluses développaient ce qu'ils appelaient une inertie cognitive. Je précise qu'il s'agit d'une proportion observée dans les études analysées, et non d'une estimation de l'ensemble de la population étudiante.

Ces étudiants cessaient progressivement de remettre les réponses en question parce que les productions de l'IA suscitaient le même niveau de confiance qu'un manuel de référence. (2)

Et le plus important, c'est que Maya n'est pas l'étudiante négligente du groupe. C'est l'étudiante méthodique, analytique et rigoureuse. C'est précisément pour cette raison qu'elle était plus vulnérable.

Maya : Parce que la réponse ressemblait à la manière dont j'organise mes idées. Je ne l'ai pas remise en question parce qu'elle ressemblait à ma façon de penser.

Dr Flores : C'est une forme de ce que l'on pourrait appeler un biais de confirmation automatisé. Le terme n'est pas officiellement établi dans la littérature, mais il décrit bien le





phénomène. La réponse ne paraissait pas seulement correcte : elle était structurée comme Maya structure elle-même son raisonnement. Et lorsque quelque chose ressemble à votre modèle mental, votre cerveau baisse naturellement sa garde.

Kai : Et moi ? Je ne l'aurais pas acceptée parce qu'elle était bien construite. Je l'aurais acceptée parce que je me serais dit que ça ne pouvait pas être aussi simple.

Dr Flores : Et c'est précisément ce dont nous allons avoir besoin dans ce chapitre, Kai.

Dans le chapitre précédent, nous avons vu où l'IA échoue. La question est maintenant : comment repérer ces erreurs ? Et cela se fait à plusieurs niveaux.

Kai : Plusieurs niveaux de méfiance ?

Dr Flores : Plusieurs niveaux d'évaluation. Trois, exactement. Le premier est le plus simple : la vérification des faits. Ce que dit l'IA est-il exact ? La posologie est-elle correcte ? La citation existe-t-elle réellement ? Les données sont-elles authentiques ? (3)

Maya : Ça, je le fais déjà.

Dr Flores : À chaque fois ?

Maya : ... Quand j'ai le temps.

Dr Flores : Et c'est là que réside le problème. Kuziemy et ses collègues ont montré que les standards de qualité appliqués à l'évaluation de l'IA n'étaient respectés que dans 39 % des revues et 14 % des études originales. Les gens savent vérifier. Ils ne le font simplement pas de manière systématique.

Kai : Donc le premier niveau consiste à se demander : est-ce vrai ?

Dr Flores : Est-ce vrai ? Puis-je le confirmer à l'aide d'une source indépendante ?

Le deuxième niveau est plus subtil : il consiste à évaluer le raisonnement. Il ne s'agit plus seulement de vérifier les faits, mais aussi la logique qui les relie. L'IA peut vous fournir cinq informations exactes et aboutir malgré tout à une conclusion erronée. Parce qu'elle ne raisonne pas réellement : elle assemble des éléments.

Kai : Comme lorsqu'elle propose cinq diagnostics valides, mais les classe selon leur fréquence générale plutôt qu'en fonction du patient que j'ai devant moi.

Dr Flores : Exactement. Les diagnostics sont corrects. C'est leur hiérarchisation qui ne l'est pas. Et en clinique, cette hiérarchisation est essentielle.

Maya : C'est précisément là où je me suis trompée avec mon diagnostic différentiel. Les informations étaient justes. Mais je n'ai jamais vérifié le raisonnement qui les organisait.

Dr Flores : Et le troisième niveau est le plus difficile : remettre en question les hypothèses. Pas celles de l'IA. Les vôtres. Avez-vous posé la bonne



question ? Cherchez-vous à confirmer ce que vous pensiez déjà ? Existe-t-il un angle mort que votre requête excluait dès le départ ?

Kai : Donc, niveau 1 : est-ce qu'elle me ment ? Niveau 2 : est-ce qu'elle me trompe dans son raisonnement ? Niveau 3 : est-ce que je me suis trompé moi-même avant même de lui poser la question ?

Dr Flores : C'est la version Kai. Et elle est parfaite.

Maya : Le troisième niveau est celui qui fait le plus mal.

Kai : Donc, en quelque sorte, c'est moi qui enseigne à la machine ?

Dr Flores : Exactement. Tu enseignes à la machine. Et la machine te répond en te signalant ce qui n'est pas clair dans ton explication. Elle ne te corrige pas directement. Elle t'oblige à te corriger toi-même.

Tu pourrais lui dire : "Tu es un étudiant de première année. Je vais t'expliquer la prise en charge de l'insuffisance cardiaque. Pose-moi des questions chaque fois que mon explication est incomplète, imprécise ou difficile à comprendre."

Puis tu la laisses t'interroger.

Kai : J'utiliserais volontiers cette méthode. Ça ressemble à un jeu, mais ça me parle.

Dr Flores : Et si elle fonctionne, c'est parce qu'elle active précisément ce que l'IA ne possède pas : ta métacognition. Chaque question que l'IA te pose agit comme un miroir révélant ce que tu n'as pas expliqué correctement. (5)

Maya : Je fais quelque chose d'assez similaire... mais dans l'autre sens. Je demande à l'IA de m'expliquer un sujet, puis je réorganise l'information avec mes propres mots. Je pensais que cela suffisait.

Dr Flores : Et quelle est la différence ?

Maya : Lorsque je réorganise une réponse générée par l'IA, je travaille à partir de sa structure, pas de la mienne. Au final, je peux produire quelque chose qui ressemble à mon raisonnement, alors qu'il est en réalité construit à partir de celui de la machine.

Dr Flores : Voilà toute la différence entre reformuler et comprendre. Reformuler consiste à changer les mots. Comprendre consiste à être capable d'expliquer une idée à partir de zéro, sans s'appuyer sur l'échafaudage fourni par l'IA.

Après avoir réorganisé une réponse, essaie de l'expliquer sans l'avoir sous les yeux. Si tu y parviens, c'est que tu as compris. Sinon, tu viens d'identifier exactement où se situe la lacune.

Dr Flores : Je vais maintenant vous montrer quelque chose qui relie tout ce dont nous avons parlé jusqu'à présent. Dans le premier chapitre, nous avons évoqué les données d'Anthropic : lorsque les étudiants en santé utilisent l'IA, près de 40 % des interactions concernent des tâches de création et environ 30 % des tâches d'analyse. Autrement dit, les niveaux les plus élevés de la taxonomie de Bloom. (6)

Kai : Les niveaux censés être les plus difficiles.

Dr Flores : Exactement. Les plus exigeants, mais aussi les plus importants pour votre formation. En 2024, Gonsalves a montré que les étudiants qui utilisaient efficacement l'IA combinaient ses réponses avec des recherches plus traditionnelles. À l'inverse, ceux qui l'utilisaient de manière passive reproduisaient directement ce que l'IA leur fournissait à des niveaux cognitifs élevés, sans avoir construit les bases nécessaires pour y parvenir eux-mêmes. (7)

Dr Flores : Maya, quels niveaux de la taxonomie de Bloom mobilises-tu le plus souvent ?

Maya : L'analyse. L'évaluation. Et parfois la création.

Dr Flores : Et à quels niveaux utilises-tu l'IA ?

Maya : ... Aux mêmes niveaux.

Dr Flores : Voilà le problème. Tu délègues précisément les tâches pour lesquelles ton cerveau devrait fournir l'effort le plus important. Pas les niveaux inférieurs, mémoriser ou comprendre, mais les niveaux

supérieurs : ceux qui construisent le jugement clinique.

Kai : Moi, c'est l'inverse. J'utilise l'IA pour les niveaux de base. Pour éviter de lire le chapitre ou pour qu'elle m'explique rapidement ce que je n'ai pas compris.

Dr Flores : Et cela a un coût différent. Tu construis des fondations fragiles. Dès qu'une question sort du cadre du résumé fourni par l'IA, tu te retrouves démuni.

Lubbe et ses collègues l'ont formulé ainsi : l'enseignement doit cesser de demander "Que sais-tu ?" et commencer à demander "Comment le sais-tu ?" Si ta seule réponse est "Parce que l'IA me l'a dit", alors tu ne disposes d'aucune base solide. (8)

Kai : Donc Maya perd quelque chose par le haut... et moi par le bas.



 **Dr Flores :** Vous perdez tous les deux. Mais dans des directions opposées. Et l'IA ne vous avertit jamais de ce qui est en train de se passer.

En 2025, Jackson a proposé un cadre appelé Higher Order Prompting, qui associe les prompts aux différents niveaux de la taxonomie de Bloom. L'idée est simple : le niveau de vos requêtes doit évoluer en même temps que votre apprentissage.

Au niveau Mémoriser, on peut demander : “Quels sont les critères diagnostiques de X ?”

Au niveau Évaluer : “Critique ce plan thérapeutique. Quelles hypothèses sous-jacentes fait-il ?”

Au niveau Créer : “Conçois un algorithme diagnostique intégrant un outil d'IA de dépistage et le jugement clinique.” (9)

 **Maya :** Donc le niveau du prompt reflète le niveau auquel on souhaite faire travailler son cerveau.

 **Dr Flores :** Exactement. Et si tu demandes toujours à l'IA de créer à ta place, ton cerveau ne s'entraînera jamais à créer.

 **Kai :** Et si je lui demande toujours de m'expliquer les choses, mon cerveau ne s'exercera jamais à comprendre seul.

 **Dr Flores :** Il y a aussi un résultat contre-intuitif que vous devez connaître. Ayoub et ses collègues ont étudié près d'un millier de chercheurs et ont constaté que ceux qui étaient les plus immergés dans l'utilisation de l'IA générative ressentaient une charge mentale plus élevée, et non plus faible. Autrement dit, l'outil censé faire gagner du temps et réduire l'effort cognitif contribuait, dans certains cas, à augmenter la fatigue mentale et le risque d'épuisement. (10)



 **Kai :** Ça me parle. Il y a des jours où je passe plus de temps à me battre avec l'outil qu'à réellement étudier.

 **Dr Flores :** Et c'est là qu'un concept prend tout son sens : celui des difficultés souhaitables (desirable difficulties) décrit par Bjork. Il s'agit de stratégies d'apprentissage qui paraissent plus lentes et plus exigeantes sur le moment, mais qui favorisent une meilleure mémorisation et un meilleur transfert des connaissances à long terme. Le rappel actif, l'espacement des révisions, l'entrelacement des sujets... L'IA tend à éliminer ces difficultés. Elle vous donne la réponse avant même que votre cerveau ait terminé de la chercher. Elle vous épargne l'effort. Or cet effort faisait précisément partie de l'apprentissage. (11)

 **Maya :** Donc la question n'est pas d'utiliser ou non l'IA. C'est de savoir quand s'arrêter.

 **Dr Flores :** C'est de savoir quand l'outil amplifie votre réflexion et quand il commence à la remplacer. Si vous remarquez que l'interaction devient

🗨 mécanique, que vous ne traitez plus l'information mais que vous la recopiez, ou que vous êtes plus fatigué après avoir utilisé l'IA qu'avant, ce sont des signaux à prendre au sérieux. (12)

🗨 **Kai :** Les signes que vous avez cessé de réfléchir.

🗨 **Dr Flores :** Les signes qu'il est temps de fermer l'outil et d'ouvrir un livre. Ou simplement d'accepter l'inconfort de ne pas savoir immédiatement, le temps que votre cerveau fasse son travail.

🗨 Assez de théorie. Parlons maintenant d'outils concrets. Des choses que vous pourrez utiliser dès demain.

🗨 **Kai :** Enfin !

🗨 **Dr Flores :** Commençons par l'évaluation des réponses générées par l'IA. Je vais vous présenter une adaptation que je propose à partir du cadre SIFT développé par Mike Caulfield. À l'origine, SIFT a été conçu pour évaluer des informations trouvées sur Internet : Stop, Investigate, Find, Trace. Mais l'IA n'a pas toujours de source identifiable, ni d'auteur clairement

🗨 défini, ni d'institution responsable. Il faut donc adapter la méthode à ce contexte. (13)

Quatre étapes.

Première étape : s'arrêter. Ne copiez pas, ne présentez pas et n'agissez pas immédiatement sur la base de la réponse obtenue.

Deuxième étape : examiner les affirmations. Ne cherchez pas à vérifier la source de l'IA, puisqu'elle n'en fournit pas toujours. Vérifiez plutôt chaque information importante à l'aide d'une référence indépendante.

Troisième étape : rechercher une meilleure référence. Existe-t-il un article scientifique, une recommandation clinique ou un manuel qui dit autre chose ?

Quatrième étape : retracer le raisonnement. Comment la réponse passe-t-elle du point A au point B ? Le raisonnement est-il réellement valide ou simplement plausible ?





Maya : Donc c'est toujours SIFT, mais au lieu d'enquêter sur l'auteur, on vérifie chacune des affirmations.

Dr Flores : Exactement. Et je tiens à préciser que cette adaptation n'est pas celle de Caulfield lui-même. C'est une proposition que j'adapte à partir de son modèle pour le contexte de l'IA. Elle n'a pas encore fait l'objet d'une validation empirique spécifique, même si chacune de ses étapes repose sur des principes reconnus d'évaluation critique.

Kai : Et l'autre outil ?

Dr Flores : Les prompts socratiques adaptés aux différents niveaux de la taxonomie de Bloom. Au lieu de demander des réponses à l'IA, vous lui demandez de vous poser des questions. Park et Choo ont documenté plusieurs stratégies concrètes en 2025.

Par exemple, vous pouvez utiliser un raisonnement socratique guidé : "Avant de répondre, explique ton raisonnement étape par étape. À chaque étape, identifie les hypothèses que tu fais et précise si elles sont justifiées." Cela oblige l'IA à rendre son raisonnement visible et vous oblige, en retour, à l'évaluer. (14)

Maya : Et si je veux passer au niveau supérieur ?

Dr Flores : Alors il faut faire évoluer vos prompts.

Au niveau Évaluer, vous pouvez demander : "Critique cette réponse. Quelles informations manquent ? Quelles hypothèses implicites contient-elle ?"

Au niveau Créer, vous pouvez demander : "Ne me donne pas directement l'algorithme. Guide-moi par des questions afin que je le construise moi-même."

Chaque niveau de prompt active un niveau différent de votre réflexion.

Kai : Donnez-moi un exemple concret. Quelque chose que je pourrais utiliser dès demain en stage.

Dr Flores : Imaginons qu'un patient consulte pour des céphalées. Au lieu de demander : "Donne-moi un diagnostic différentiel pour une céphalée", vous pourriez dire :

"J'ai un patient présentant des céphalées. Au lieu de me fournir directement un diagnostic différentiel, pose-moi des questions sur l'anamnèse et l'examen clinique afin que je construise moi-même mon raisonnement. Une fois terminé, indique-moi ce que j'ai oublié."

Kai : J'aime bien. Ça ressemble à un jeu, mais un jeu utile.

Dr Flores : Maintenant, parlons des signaux d'alerte personnels. Les signes qui indiquent qu'il est temps de s'arrêter. Lorsque vous reformulez la même requête plus de trois fois sans obtenir ce que vous cherchez, il est possible que vous ne sachiez plus exactement ce que vous cherchez.



- Dr Flores :** Lorsque vous acceptez une réponse sans la lire entièrement parce qu'elle "a l'air correcte". Lorsque vous vous sentez plus fatigué après l'interaction qu'avant de l'avoir commencée. Ou encore lorsque vous passez dix minutes sans rien apprendre de nouveau et que vous ne faites plus que modifier la mise en forme.
- Maya :** Honnêtement... presque tous ces signaux se sont manifestés chez moi cette semaine.
- Dr Flores :** Et ce n'est pas un échec. C'est une information. Chaque signal d'alerte représente une occasion de reprendre le contrôle que l'outil était en train de prendre sans que vous vous en rendiez compte.
- Kai :** Vous savez ce qui est étrange ? Pendant toute ma scolarité, on m'a dit que mon problème était de ne pas faire assez d'efforts. Que je remettais trop les choses en question. Que mon scepticisme était une forme de paresse.
- Dr Flores :** Et alors ?
- Kai :** Et voilà que je découvre que cette tendance à questionner est exactement ce dont on a besoin face à un outil qui ne dit jamais : "Je ne sais pas." Le problème n'a jamais été mon scepticisme. Je ne savais simplement pas où l'appliquer.
- Dr Flores :** C'est cela, la pensée critique, Kai. Ce n'est ni une matière ni une liste d'étapes. C'est la discipline qui consiste à se demander si ce qui paraît vrai l'est réellement. Et à accepter l'inconfort lorsque la réponse est : "Je ne sais pas."
- Maya :** Pour moi, c'est plus difficile. Mon réflexe n'est pas de douter. Mon réflexe est d'organiser. Et lorsqu'une information est déjà parfaitement organisée, mon cerveau a tendance à la considérer comme acquise.
- Dr Flores :** Et qu'est-ce que tu vas faire différemment ?
- Maya :** Ajouter une étape. Avant de présenter quelque chose généré par l'IA, je vais essayer de l'expliquer sans avoir la réponse sous les yeux. Si je n'y arrive pas, je considérerai que je ne l'ai pas encore vraiment compris.
- Dr Flores :** Une seule étape. Pas un système parfait. Une seule étape.
- Maya :** Ça va me demander un effort.
- Kai :** Je sais.
- Dr Flores :** Salido et ses collègues ont publié une revue systématique sur la pensée critique et l'IA dans l'enseignement supérieur. Ils ont identifié cinq grands thèmes : l'innovation pédagogique, les dimensions psychologiques, l'éthique académique, l'intégration institutionnelle et les compétences liées à l'IA. Mais leur conclusion tenait en une seule idée : la compétence qui distinguera les professionnels de demain ne sera pas la capacité à utiliser l'IA. Ce sera la capacité à continuer à réfléchir pendant qu'ils l'utilisent. (15)
- Kai :** C'est le filtre.



 **Dr Flores :** Exactement. Le filtre que personne ne leur a enseigné, mais dont ils auront besoin tous les jours. Non pas pour se méfier systématiquement de l'IA, mais pour avoir suffisamment confiance en leur propre jugement pour évaluer ce qu'elle leur propose. Pour savoir quand elle amplifie leur réflexion et quand elle la remplace. Pour reconnaître les moments où l'effort cognitif est indispensable. Pour se demander : "Attendez... est-ce que j'ai posé la bonne question ?" avant d'accepter une réponse.

Vous savez désormais où l'IA échoue. Vous savez aussi comment repérer ces limites. La prochaine étape consiste à l'intégrer là où cela compte le plus : auprès du patient.

 **Kai :** Le raisonnement clinique.

 **Dr Flores :** Le raisonnement clinique. Là où l'esprit critique n'est plus un exercice académique, mais ce qui se situe entre une recommandation de l'IA et une décision qui affecte une personne réelle.



À retenir de ce chapitre

- **Reformuler n'est pas comprendre.** Si vous ne pouvez pas expliquer un concept sans avoir la réponse de l'IA sous les yeux, vous ne l'avez probablement pas encore réellement compris.
- **Trois niveaux d'évaluation critique :**
Niveau 1 : Est-ce vrai ? Niveau 2 : Le raisonnement est-il valide ? Niveau 3 : Mes propres hypothèses sont-elles correctes ?
- **Le niveau de votre prompt détermine le niveau de réflexion sollicité.** Si vous demandez toujours à l'IA de créer à votre place, vous n'exercerez jamais votre propre capacité à créer.
- **Quelques signaux d'alerte personnels :**
reformuler la même requête plus de trois fois ; accepter une réponse sans la lire entièrement ; se sentir plus fatigué après avoir utilisé l'IA ; passer du temps à ajuster la forme sans apprendre quelque chose de nouveau.
- **La compétence qui distinguera les professionnels de demain ne sera pas de savoir utiliser l'IA, mais de savoir réfléchir pendant qu'ils l'utilisent.**



Références

1. Dakan, R. & Feller, J. (2026). *Anthropic education report: The AI Fluency Index*. Anthropic. <https://www.anthropic.com/research/ai-fluency-index> [Corporate research — not peer-reviewed]
2. Li, J., et al. (2026). Application of AI-generated content in medical education: Systematic review of the impact on critical thinking abilities of medical students. *JMIR Medical Education*, 12(1), e79939. <https://doi.org/10.2196/79939>
3. Kuziemy, C. E., Chrimes, D., Minshall, S., Mannerow, M. & Lau, F. (2024). AI quality standards in health care: Rapid umbrella review. *Journal of Medical Internet Research*, 26, e54705. <https://doi.org/10.2196/54705>
4. Tomisu, H., Ueda, J. & Yamanaka, T. (2025). The cognitive mirror: A framework for AI-powered metacognition and self-regulated learning. *Frontiers in Education*, 10, 1697554. <https://doi.org/10.3389/educ.2025.1697554>
5. Neshaei, S. P., et al. (2025). Metacognition meets AI: Empowering reflective writing with large language models. *British Journal of Educational Technology*, 56(5), 1864–1896. <https://doi.org/10.1111/bjet.13601>
6. Anthropic. (2025). Anthropic education report: *How university students use Claude*. <https://www.anthropic.com/news/anthropic-education-report-how-university-students-use-claude> [Corporate research — not peer-reviewed]
7. Gonsalves, C. (2024). GenAI impact on critical thinking: Revisiting Bloom’s Taxonomy. *Journal of Marketing Education*, 47(1). <https://doi.org/10.1177/02734753241305980>
8. Lubbe, A., Marais, E. & Kruger, D. (2025). Cultivating independent thinkers: The triad of artificial intelligence, Bloom’s taxonomy and critical thinking in assessment pedagogy. *Education and Information Technologies*, 30, 17589–17622. <https://doi.org/10.1007/s10639-025-13476-x>
9. Jackson, J. (2025). Higher Order Prompting: Bloom’s Revised Taxonomy and LLMs. *Studies in Technology Enhanced Learning*. <https://doi.org/10.21428/8c225f6e.0915c17e>
10. Ayoub, T., et al. (2025). Generative AI and cognitive challenges among researchers: Balancing cognitive load, fatigue, and human resilience. *Technologies*, 13(11), 486. <https://doi.org/10.3390/technologies13110486>
11. Bjork, R. A. (1994). Memory and metamemory considerations in the training of human beings. In J. Metcalfe & A. P. Shimamura (Eds.), *Metacognition: Knowing about knowing* (pp. 185–205). MIT Press.
12. Jose, B., Cherian, J., Verghis, A. M., Varghise, S. M., S. M. & Joseph, S. (2025). The cognitive paradox of AI in education: Between enhancement and erosion. *Frontiers in Psychology*, 16, 1550621. <https://doi.org/10.3389/fpsyg.2025.1550621>
13. Caulfield, M. (2019). SIFT (The Four Moves). Hapgood. <https://hapgood.us/2019/06/19/sift-the-four-moves/>
14. Park, J. & Choo, S. (2025). Generative AI prompt engineering for educators: Practical strategies. *Journal of Special Education Technology*, 40(3), 411–417. <https://doi.org/10.1177/01626434241298954>
15. Salido, A., Syarif, I., Sitepu, M. S., Suparjan, Wana, P. R., Taufika, R. & Melisa, R. (2025). Integrating critical thinking and artificial intelligence in higher education: A bibliometric and systematic review of skills and strategies. *Social Sciences & Humanities Open*. <https://doi.org/10.1016/j.ssaho.2025.101924>



Chapitre 5 : L'IA et le raisonnement clinique



Dans ce chapitre, vous découvrirez...

- Comment l'ancrage sur une réponse générée par l'IA peut vous conduire à raisonner à l'intérieur de son cadre sans jamais en sortir, même lorsque vous pensez faire preuve d'esprit critique.
- Pourquoi, dans un essai randomisé, l'accès à l'IA n'a pas amélioré les performances de diagnostics des médecins.
- Les **scripts de maladie** (*illness scripts*) et les différentes étapes du raisonnement clinique : construire une représentation du problème, recueillir les données, générer des hypothèses, les affiner, prendre une décision et assurer le suivi.
- Le protocole **Réfléchir d'abord, consulter ensuite** : quatre étapes pour protéger votre raisonnement avant de recourir à l'IA.
- Trois risques spécifiques : la perte de compétences (*deskilling*), le non-développement des compétences (*never-skilling*) et l'acquisition de compétences erronées (*miss-skilling*).



Dr Flores : Aujourd'hui, nous allons procéder autrement. Je vais vous présenter un cas clinique. Je veux que vous le traitiez chacun de votre côté, puis nous comparerons vos approches.

Kai : C'est un examen ?

Dr Flores : Mieux qu'un examen. C'est un miroir.

Patiente de 34 ans consultant pour une fatigue progressive évoluant depuis trois mois et des arthralgies migratrices. Elle rapporte un épisode de rougeur bilatérale du visage apparu après une exposition au soleil. Elle présente également des antécédents de rosacée traitée deux ans auparavant.

Bilan biologique : anémie légère, vitesse de sédimentation élevée, TSH à la limite supérieure de la normale (5,2). Son médecin traitant a demandé une recherche d'anticorps antinucléaires (ANA), positive à 1:160 avec un aspect moucheté.

Maya : Cela pourrait être un lupus. Mais le titre des ANA est relativement faible et le profil n'est pas vraiment typique.

Kai : Et la TSH élevée complique encore les choses. On pourrait aussi être face à une hypothyroïdie avec des ANA positifs de manière fortuite.

Dr Flores : Plusieurs hypothèses sont possibles. Maya, raconte-nous ce que tu as fait.

Maya : J'ai ouvert un outil d'IA. Je lui ai fourni l'ensemble du cas et demandé un diagnostic différentiel structuré avec des probabilités. Il m'a proposé une liste de sept diagnostics : lupus érythémateux systémique en première position, puis polyarthrite rhumatoïde, syndrome de Sjögren, connectivite mixte, hypothyroïdie auto-immune, vascularite et réaction médicamenteuse. Chaque diagnostic était accompagné des critères pertinents et des examens recommandés pour la suite du bilan.



Dr Flores : Et ensuite ?

Maya : J'ai examiné la liste. Elle me semblait solide. J'ai commencé à comparer chaque diagnostic aux données de la patiente. J'ai écarté la vascularite et la réaction médicamenteuse parce qu'elles correspondaient mal au tableau clinique. J'ai conservé le lupus érythémateux systémique comme hypothèse principale et demandé les examens complémentaires appropriés : anticorps anti-ADN natif, dosage du complément et anticorps anti-Sm.

Dr Flores : Et toi, Kai ?

Kai : Je n'ai ouvert aucun outil. J'ai continué à réfléchir. Jeune femme, fatigue, arthralgies, rougeur du visage, ANA positifs. Mon esprit est immédiatement allé vers le lupus. Puis je me suis souvenu d'une remarque entendue pendant un stage en rhumatologie : les ANA positifs sont probablement l'un des résultats biologiques les plus surinterprétés en médecine.

Dr Flores : Continue.

Kai : Alors je me suis arrêté et je me suis demandé : qu'est-ce qui peut aussi provoquer une fatigue associée à des arthralgies et une anémie chez une



femme de 34 ans ? Une hypothyroïdie, surtout avec une TSH déjà à la limite haute. Une carence martiale. Une infection chronique. Une hépatite. Même une maladie cœliaque. La rougeur du visage me poussait vers le lupus, mais avec les antécédents de rosacée et des ANA mouchetés à 1:160... ce n'était pas si évident. J'ai élargi mon champ d'hypothèses.

Ma liste n'était pas particulièrement élégante. Je n'avais ni probabilités ni critères diagnostiques. Mais c'était ma liste.

Dr Flores : Et c'est précisément là que réside la différence.

Maya, je vais te poser une question un peu inconfortable. Parmi les sept diagnostics proposés par l'IA, combien aurais-tu envisagés toi-même sans son aide ?

Maya : Le lupus, certainement. Le syndrome de Sjögren, peut-être. La polyarthrite rhumatoïde aussi.

Dr Flores : Et les quatre autres ?

Maya : C'est l'outil qui me les a suggérés.

Dr Flores : Et qu'en as-tu fait ensuite ?

Maya : Je les ai confrontés au cas clinique. J'en ai éliminé deux.

Dr Flores : Tu as évalué ce que l'IA t'a proposé. Mais tu n'as pas généré ce qu'elle ne t'a pas proposé. La maladie cœliaque ? L'hépatite ? Quant à la TSH limite haute, tu ne l'as même pas mentionnée.

Maya : Pourtant, je n'ai pas accepté sa réponse aveuglément.

Dr Flores : Non. Mais tu as raisonné à l'intérieur du cadre qu'elle t'avait fourni, sans jamais en sortir. Et c'est toute la différence. Cela s'appelle le biais d'ancrage. C'est l'un des biais cognitifs les plus étudiés dans le domaine du raisonnement clinique. Tversky et Kahneman l'ont décrit dans les années 1970, et Croskerry l'a largement documenté en médecine : lorsqu'une première hypothèse forte est proposée, notre esprit tend à graviter autour d'elle. Nous évaluons davantage ce qui la confirme que ce qui la contredit.

Maya : Et l'IA m'a fourni cette première hypothèse avant même que je commence à réfléchir.

Dr Flores : Exactement. Goh et ses collègues ont publié un essai randomisé dans JAMA Network Open portant sur cinquante médecins. Les praticiens ayant accès à un grand modèle de langage n'ont pas obtenu de meilleures performances diagnostiques que ceux qui n'y avaient pas accès. En revanche, le modèle utilisé seul obtenait de meilleurs résultats que les médecins. (1)

Kai : Attendez... le fait d'avoir accès à l'outil ne les a pas rendus meilleurs ?

Dr Flores : Pas dans cette étude. Et l'explication la plus probable est précisément ce qui est arrivé à Maya : les médecins n'ont pas utilisé l'IA pour

élargir leur raisonnement, mais pour l'ancrer dans les réponses fournies par l'IA. Mahajan et ses collègues rapportent dans npj Digital Medicine que les réponses des grands modèles de langage reproduisent plusieurs biais cognitifs bien connus : l'ancrage, l'effet de disponibilité et le biais de confirmation. Une autre étude publiée dans la même revue a également mis en évidence un risque d'acquiescement : les modèles ont tendance à confirmer l'hypothèse du clinicien plutôt qu'à la remettre en question. (2)(3)

Kai : Donc l'IA ne se contente pas de ne pas vous corriger. Elle vous donne souvent l'impression que vous avez raison.

Dr Flores : Dans la plupart des cas, oui. Et lorsque vous partez de sa réponse plutôt que de votre propre réflexion, deux sources de biais agissent dans le même sens : votre propre ancrage et la tendance du modèle à confirmer vos hypothèses. Personne ne joue alors le rôle du contradicteur.

Maya : Je pensais faire preuve d'esprit critique. En réalité, j'évaluais les

informations à l'intérieur du cadre fourni par l'IA. Je ne sortais jamais de ce cadre.

Dr Flores : C'est exactement ce que produit le biais d'ancrage. Il vous donne l'impression de réfléchir alors que vous ne faites qu'affiner un raisonnement à l'intérieur d'une boîte construite par quelqu'un d'autre.

Avant de parler de la place de l'IA dans ce processus, nous devons comprendre comment fonctionne réellement le raisonnement clinique. Pas la version des examens. La vraie version.

Kai : Il y a une différence ?

Dr Flores : Toujours. Dans un examen, toutes les données vous sont fournies et l'on vous demande d'établir un diagnostic. Dans la pratique réelle, c'est vous qui décidez quelles informations rechercher. Et parfois, l'élément le plus important est précisément celui que vous n'avez pas pensé à demander.

Le raisonnement clinique est un processus séquentiel. Imaginez-le comme une navigation. Il vous faut une carte, une position de départ et une destination. Cette carte n'existe pas d'emblée : vous la construisez progressivement à partir des maladies que vous avez étudiées, des patients que vous avez rencontrés et des schémas cliniques que vous avez appris à reconnaître.

Dans la littérature, ces représentations mentales sont appelées scripts de maladie (illness scripts). Prenons un exemple concret : le script de maladie du lupus érythémateux systémique.

Facteurs prédisposants : femme jeune, prédisposition génétique, exposition aux rayons ultraviolets.

Physiopathologie en une phrase : maladie auto-immune systémique caractérisée par la formation et le dépôt de complexes immuns.





Manifestations cliniques typiques : érythème malaire, arthralgies, néphrite, cytopénies, sérites.

Évolution habituelle : maladie chronique évoluant par poussées et rémissions. (4)

Maya : Comme des modèles mentaux internes qui organisent l'information.

Dr Flores : Exactement. Norman et ses collègues l'ont rappelé en 2024 : les deux systèmes de pensée, le système rapide et intuitif, et le système lent et analytique, ne sont pas des processus indépendants. Ils dépendent du type de connaissances que vous êtes capable de mobiliser. Si vous disposez d'un script de maladie riche pour le lupus, vous reconnaîtrez rapidement le schéma clinique. Si ce n'est pas le cas, vous devrez emprunter une démarche plus analytique.

Maya : Donc le système 1 n'est pas une intuition magique. C'est de l'expérience condensée.

Dr Flores : Exactement.

Dans son ensemble, le raisonnement clinique comprend six étapes : construire la carte mentale, recueillir les données, générer des hypothèses,

les affiner, prendre une décision et assurer le suivi. L'IA peut intervenir à chacune de ces étapes. La question est de savoir comment.

Dr Flores : Pour construire la carte mentale, l'IA peut être réellement utile. Elle peut vous aider à élaborer des scripts de maladie. Mais au lieu de demander : "Explique-moi tout sur le lupus", vous pourriez dire :

"Je suis en train de construire mon script de maladie du lupus érythémateux systémique. Donne-moi : une physiopathologie en trois phrases, les cinq manifestations les plus spécifiques, les trois diagnostics les plus souvent confondus avec le lupus et son évolution typique. Ensuite, pose-moi trois questions pour vérifier que j'ai compris."

Les questions finales transforment l'IA d'une simple source d'information passive en un outil d'évaluation active. Wei et ses collègues ont confirmé dans une méta-analyse que l'intégration de l'IA dans l'apprentissage par cas améliorerait l'acquisition des connaissances de 46 % par rapport aux méthodes traditionnelles, même si l'ampleur de cet effet varie selon les études. (5)

Dr Flores : En revanche, lorsqu'il s'agit de générer des hypothèses diagnostiques, là où Maya a commencé et où les difficultés sont apparues, l'ordre des opérations est essentiel. Takita et ses collègues ont montré dans une méta-analyse que la précision diagnostique de l'IA générative dans des cas cliniques à diagnostic ouvert est d'environ 52 %, comparable à celle de médecins non experts mais nettement inférieure à celle d'experts. L'IA est un généraliste rapide, pas un spécialiste en profondeur. (6)

Kai : Donc elle est utile pour construire un diagnostic différentiel ou non ?

Dr Flores : Oui, à condition de l'utiliser au bon moment. Le clinicien génère d'abord ses propres hypothèses ; l'IA intervient ensuite pour les compléter.

Cela paraît évident. Pourtant, presque personne ne procède ainsi. Abdunour et ses collègues l'ont montré dans le New England Journal of Medicine. Ils ont identifié trois risques majeurs liés à l'utilisation de l'IA dans le raisonnement clinique :

- La perte de compétences (*deskilling*) : vous perdez progressivement des compétences que vous possédiez déjà.
- Le non-développement des compétences (*never-skilling*) : vous ne développez jamais certaines compétences parce que l'IA réalise le travail à votre place dès le départ.
- L'acquisition de compétences erronées (*miss-skilling*) : vous apprenez de manière incorrecte parce que l'IA vous a transmis un modèle ou un raisonnement erroné. (8)

Maya : Ces trois risques peuvent être évités si l'on réfléchit d'abord.

Dr Flores : Exactement. Tous les trois. Le protocole est simple et comporte quatre étapes.

Première étape : lire le cas. Rien que le cas. Fermez l'outil d'IA.

Deuxième étape : élaborer votre propre diagnostic différentiel. Peu importe qu'il soit incomplet, imparfait ou désordonné. Il vous appartient.

Troisième étape : seulement ensuite, ouvrez l'IA. Donnez-lui le cas et votre diagnostic différentiel. Demandez-lui de compléter votre réflexion, pas de repartir de zéro.

Quatrième étape : comparez. Qu'est-ce qui manquait dans votre raisonnement ? Qu'a ajouté l'IA ? Pourquoi n'y aviez-vous pas pensé ?

Kai : C'est dans cette dernière étape qu'on apprend vraiment.

Dr Flores : Exactement. Parce que l'écart entre votre liste et celle de l'IA rend visible votre lacune de connaissance. Il n'existe probablement pas de meilleur outil d'apprentissage que cela.

Kai, dans le cas d'aujourd'hui, ce que tu as fait naturellement sans même t'en



rendre compte, c'est réfléchir avant de consulter. Ta liste était désordonnée, oui. Elle n'était pas structurée selon des critères formels. Mais elle était à toi. C'est précisément pour cette raison qu'elle incluait l'hypothyroïdie, l'hépatite ou la maladie cœliaque.

L'IA n'a pas proposé ces hypothèses à Maya parce qu'elle privilégiait ce qui paraissait le plus probable au regard des données fournies, et non nécessairement ce qui était le plus pertinent compte tenu du contexte clinique.

Kai : Vous êtes en train de me dire que ma façon un peu désordonnée de réfléchir... a de la valeur ?

Dr Flores : Je suis en train de te dire que ton intuition clinique correspond à une véritable reconnaissance de schémas cliniques, construite à partir de l'expérience. Ton problème n'est pas l'intuition. C'est simplement qu'elle manque parfois de structure. Et l'IA peut justement t'aider à structurer ce que tu as déjà commencé à construire.

Maya : Et moi, c'est l'inverse. J'ai la structure. Mais si cette structure vient de l'IA avant même que je commence à réfléchir, elle n'est pas vraiment la mienne. Elle m'est prêtée. Et je ne peux pas évaluer correctement ce que je n'ai pas construit moi-même.

Dr Flores : Exactement. Et cette distinction devient encore plus importante lorsque l'on passe de la salle de cours à la clinique.

En salle de cours, le cas est complet. Toutes les informations sont disponibles, et il existe une réponse attendue à la fin.

En clinique, le cas se construit en temps réel. Les informations sont incomplètes. Le patient est devant vous, souvent inquiet, parfois anxieux. C'est à ce moment-là que le risque de la béquille cognitive devient concret.

Dans un contexte académique, vous pouvez consulter l'IA discrètement. En clinique, un patient attend votre réponse. Et dans ce moment-là, ce qui compte, c'est ce que vous avez dans la tête, pas ce que vous avez sur votre téléphone. (9)

Kai : Si j'avais commencé par établir ma propre liste avant de consulter l'IA... j'aurais vu ce qui me manquait : les critères diagnostiques formels, les anticorps spécifiques, l'ordre des examens à demander. Sans perdre ce que j'avais déjà construit par moi-même.

Maya : Et si j'avais réfléchi avant d'ouvrir l'IA... je serais probablement arrivée au lupus malgré tout. Mais j'aurais aussi exploré davantage d'autres pistes : la TSH, les antécédents de rosacée, une éventuelle carence martiale. Et lorsque l'IA m'aurait proposé sa liste, j'aurais pu la comparer à la mienne au lieu de l'adopter immédiatement.

Dr Flores : Voilà ce qu'est un raisonnement clinique augmenté. Ce n'est pas l'IA qui remplace votre réflexion. Ce n'est pas non plus ignorer l'IA. C'est travailler ensemble, mais dans un ordre qui protège votre capacité à penser par vous-même.

Zoller et ses collègues l'ont démontré dans PNAS : les équipes composées d'humains et d'IA obtiennent de meilleurs résultats diagnostiques que les humains seuls ou l'IA seule. Mais cela ne fonctionne que lorsque l'humain apporte son propre raisonnement. Pas lorsqu'il se contente de recopier celui de la machine.

Pour donner un ordre de grandeur, dans les cas cliniques publiés par le New England Journal of Medicine, GPT-4 obtenait un diagnostic correct dans 64 % des cas. C'est impressionnant. Mais cela reste insuffisant pour lui accorder sa confiance sans analyse indépendante. (10)(11)

Maya : Le clinicien de demain ne sera ni celui qui en sait le plus, ni celui qui maîtrise le mieux l'IA. Ce sera celui qui saura réfléchir d'abord et consulter ensuite.

Kai : Et celui qui comprendra que sa liste imparfaite a malgré tout de la valeur.

Dr Flores : Le raisonnement clinique est la compétence centrale de la médecine... et des soins infirmiers. L'IA ne le remplace pas. Elle le transforme. Mais c'est à vous de décider dans quel sens.

Vous savez désormais comment exercer votre esprit critique face aux réponses générées par l'IA. Vous venez maintenant d'apprendre à appliquer cette démarche là où elle compte le plus : face au patient.

La prochaine étape concerne une autre compétence tout aussi essentielle : apprendre avec l'IA sans lui abandonner le processus d'apprentissage lui-même.

Kai : Apprendre à apprendre ?

Dr Flores : L'apprentissage autorégulé. C'est la différence entre étudier pour réussir un examen et apprendre pour prendre soin d'un patient.



À retenir de ce chapitre

- **Si vous commencez par la réponse de l'IA, deux sources de biais agissent dans le même sens :** votre propre ancrage et la tendance de l'IA à confirmer les hypothèses initiales. Personne ne joue alors le rôle du contradicteur.
- Votre liste imparfaite a de la valeur. La différence entre votre raisonnement et celui de l'IA rend visibles vos lacunes et vos besoins d'apprentissage.
- **Le protocole Réfléchir d'abord, consulter ensuite :** Lire le cas sans utiliser l'IA. Construire son propre diagnostic différentiel. Consulter l'IA pour compléter sa réflexion. Comparer les deux approches et apprendre de leurs écarts.
- Les équipes humain-IA obtiennent de meilleurs résultats diagnostiques que chacun séparément, mais seulement lorsque l'humain apporte son propre raisonnement.
- Le clinicien de demain ne sera ni celui qui en sait le plus ni celui qui utilise le mieux l'IA. Ce sera celui qui saura réfléchir d'abord et consulter ensuite.



Références

1. Goh, E., et al. (2024). Large language model influence on diagnostic reasoning: A randomized clinical trial. *JAMA Network Open*, 7(10), e2440969. <https://doi.org/10.1001/jamanetworkopen.2024.40969>
2. Mahajan, A., Obermeyer, Z., Daneshjou, R., Lester, J., & Powell, D. (2025). Cognitive bias in clinical large language models. *npj Digital Medicine*, 8, 428. <https://doi.org/10.1038/s41746-025-01790-0>
3. Chen, S., Gao, M., Sasse, K., Hartvigsen, T., Anthony, B., Fan, L., Aerts, H., Gallifant, J., & Bitterman, D. S. (2025). When helpfulness backfires: LLMs and the risk of false medical information due to sycophantic behavior. *npj Digital Medicine*, 8, 605. <https://doi.org/10.1038/s41746-025-02008-z>
4. Norman, G., Pelaccia, T., Wyer, P., & Sherbino, J. (2024). Dual process models of clinical reasoning: The central role of knowledge in diagnostic expertise. *Journal of Evaluation in Clinical Practice*, 30(5), 788–796. <https://doi.org/10.1111/jep.13998>
5. Wei, H., Dai, Y., Yuan, K., Li, K. Y., Hung, K. F., Hu, E. M., Lee, A. H. C., Chang, J. W. W., Zhang, C., & Li, X. (2025). AI-powered problem- and case-based learning in medical and dental education: A systematic review and meta-analysis. *International Dental Journal*, 75(4), 100858. <https://doi.org/10.1016/j.identj.2025.100858>
6. Takita, H., Kabata, D., Walston, S. L., et al. (2025). A systematic review and meta-analysis of diagnostic performance comparison between generative AI and physicians. *npj Digital Medicine*, 8, 175. <https://doi.org/10.1038/s41746-025-01543-z>
7. Abbas, Q., Jeong, W., & Lee, S. W. (2025). Explainable AI in clinical decision support systems: A meta-analysis of methods, applications, and usability challenges. *Healthcare*, 13(17), 2154. <https://doi.org/10.3390/healthcare13172154>
8. Abdulnour, R. E. E., Gin, B., & Boscardin, C. K. (2025). Educational strategies for clinical supervision of artificial intelligence use. *New England Journal of Medicine*, 393(8), 786–797. <https://doi.org/10.1056/NEJMra2503232>
9. Kassab, S. E., Rathan, R., Taylor, D. C., et al. (2025). Understanding how medical students learn in the era of artificial intelligence: A mixed methods study. *BMC Medical Education*, 25, 1521. <https://doi.org/10.1186/s12909-025-08145-z>
10. Zoller, N., Berger, J., Lin, I., Fu, N., et al. (2025). Human-AI collectives most accurately diagnose clinical vignettes. *Proceedings of the National Academy of Sciences*, 122(24), e2426153122. <https://doi.org/10.1073/pnas.2426153122>
11. Kanjee, Z., Crowe, B., & Rodman, A. (2023). Accuracy of a generative artificial intelligence model in a complex diagnostic challenge. *JAMA*, 330(1), 78–80. <https://doi.org/10.1001/jama.2023.8288>



Chapitre 6 : L'IA et l'apprentissage autorégulé



Dans ce chapitre, vous découvrirez...

- La différence entre un système qui produit des résultats et un système qui produit un véritable apprentissage.
- Le cycle du **Master Adaptive Learner** (*apprenant adaptatif expert*) : planifier, apprendre, évaluer et ajuster, ainsi que la manière dont l'IA peut intervenir à chacune de ces étapes.
- Pourquoi les **difficultés souhaitables** (*desirable difficulties*) décrites par Bjork sont précisément ce que l'IA tend à éliminer : l'effet de génération, l'effet de test et l'espacement des apprentissages.
- Le principe **Vous d'abord, l'IA ensuite**, appliqué à chaque étape du processus d'apprentissage à travers des exemples concrets.
- La dépendance à l'IA : des corrélations observées avec certains symptômes dépressifs jusqu'au concept émergent d'**Allessphobia**, c'est-à-dire la difficulté ou l'anxiété ressentie lorsqu'il faut travailler sans assistance de l'IA.

Maya : La semaine dernière, j'ai fait quelque chose que je ne fais jamais. Avant un examen de pharmacologie, j'ai tout fermé : l'IA, Notion, Anki... tout. Puis j'ai essayé d'expliquer le mécanisme d'action des inhibiteurs de l'enzyme de conversion de l'angiotensine à partir de zéro. Sans notes.

Kai : Et alors ?

Maya : Je n'y suis pas arrivée. J'avais pourtant des cartes de révision impeccables, des schémas détaillés et un résumé structuré en plusieurs niveaux. Mais lorsque j'ai essayé de reconstruire le raisonnement seule... il y avait des lacunes. Des lacunes que mes propres fiches ne couvraient pas, parce que je les avais créées à partir de ce que je savais déjà, et non à partir de ce qui me manquait réellement.

Kai : Mais tu as réussi l'examen, non ?

Maya : "Oui."

Kai : Alors, ça a fonctionné.

Maya : Ça a fonctionné pour l'examen. Je ne suis pas certaine que cela ait fonctionné pour moi.

Dr Flores : La distinction que tu viens de faire est précisément le point de départ de ce chapitre. Maya, ton système est sophistiqué. Probablement plus sophistiqué que celui de la plupart

de tes camarades. Mais il existe une différence entre un système qui produit des résultats et un système qui produit un véritable apprentissage.

Sardi et ses collègues ont montré qu'une part importante des bénéfices de l'IA sur la pensée critique dépend de la capacité de l'étudiant à s'autoréguler. Sans cette compétence, une grande partie du bénéfice disparaît. (1)

Kai : Et si on n'a pas de système du tout ?

Dr Flores : Alors on se heurte à un autre problème. Xu et ses collègues ont montré que, dans des environnements où l'IA générative est largement utilisée, les étudiants qui ne bénéficient pas d'un accompagnement métacognitif voient leurs capacités d'autorégulation diminuer.

Ce n'est pas que l'IA les rende moins compétents. C'est que, sans intention délibérée, elle leur retire progressivement l'occasion de s'exercer à réguler eux-mêmes leurs apprentissages. Et ce qui est intéressant, c'est que ces deux problèmes ont la même origine. (2)

Kai : Laquelle ?

Dr Flores : Aucun de vous ne vérifie réellement si ce qu'il fait favorise l'apprentissage... ou simplement la survie à court terme.

Je vais vous montrer quelque chose. Ce n'est pas un nouveau modèle théorique. C'est quelque chose que vous faites déjà tous les jours, sans forcément en avoir conscience.

Maya : Un modèle d'apprentissage ?

Dr Flores : Un cycle.

Planifier : déterminer ce que vous devez apprendre et ce que vous savez déjà.



Apprendre : aller chercher ce qui vous manque.

Évaluer : vérifier si vous avez réellement compris.

Ajuster : modifier votre stratégie à partir de ce que vous avez découvert.

Puis recommencer.

Kai : Ça ressemble simplement à... étudier.

Dr Flores : C'est effectivement étudier. Mais avec une différence importante : la plupart des étudiants restent bloqués dans les deux premières étapes. Ils planifient, plus ou moins. Ils apprennent, plus ou moins. Mais très peu vérifient réellement s'ils ont compris. Et encore moins modifient leur stratégie à partir de preuves concrètes de leurs propres performances.

Maya : Autrement dit, je passe beaucoup de temps à planifier et à apprendre. Mais mon évaluation... c'est l'examen ?

Dr Flores : Ton évaluation, c'est l'examen. Et ton ajustement consiste souvent à refaire la même chose la fois suivante, puisque tu as réussi.

Kai : Moi, je ne planifie même pas. J'ouvre ChatGPT la veille de l'examen et je lui demande de m'expliquer ce que je n'ai pas compris en cours.

Dr Flores : Dans ce cas, tu commences directement à l'étape deux sans avoir réalisé l'étape un. Tu apprends sans avoir clairement identifié ce que tu dois apprendre. Et comme tu n'évalues ni n'ajustes réellement ta démarche, tu reproduis le même schéma encore et encore.

Kai : ... Et ça marche jusqu'au jour où ça ne marche plus.



Dr Flores : Exactement. Cela fonctionne jusqu'au moment où cela cesse de fonctionner.

Cutrer et ses collègues ont formalisé cette approche sous le nom de Master Adaptive Learner (l'apprenant adaptatif expert), un modèle qui a d'ailleurs fait l'objet d'un ouvrage publié chez Elsevier. Ce modèle montre que les professionnels de santé qui continuent à progresser tout au long de leur carrière ne sont pas forcément ceux qui savent le plus au départ. Ce sont ceux qui maîtrisent le mieux ce cycle d'apprentissage continu. (3)

Maya : Planifier, apprendre, évaluer, ajuster.

Dr Flores : Exactement. Et l'essentiel est de comprendre qu'il s'agit d'un cycle, pas d'une ligne droite. Chaque ajustement conduit à une nouvelle phase de planification. Chaque évaluation fournit des informations pour le cycle suivant. Les personnes qui excellent dans ce processus ne sont pas nécessairement les plus brillantes. Ce sont les plus adaptables.

Kai : Et où intervient l'IA dans tout ça ?

Dr Flores : Partout. Et c'est à la fois le problème et l'opportunité. Nous allons examiner chaque étape séparément. Et pour chacune d'elles, je veux que vous réfléchissiez à deux questions : ce que l'IA peut faire pour vous... et ce qu'elle ne devrait pas faire à votre place.

Planifier

 **Dr Flores :** Première étape : planifier.

Qu'est-ce que je sais déjà ? Qu'est-ce qui me manque ? Combien de temps ai-je à disposition ?

C'est probablement l'étape la plus souvent négligée. Si vous demandez à l'IA de construire un programme de révision sans avoir d'abord identifié vos propres lacunes, vous obtiendrez un plan générique : bien présenté, bien structuré... mais peu utile. L'IA ne sait pas ce que vous ignorez, à moins que vous ne lui fournissiez vous-même cette information au préalable.

 **Kai :** Mais l'IA peut quand même me poser des questions diagnostiques pour identifier mes lacunes, non ?

 **Dr Flores :** Peut-être. Et c'est précisément là qu'elle est la plus utile : comme outil de diagnostic, pas comme outil de prescription. Vous lui donnez un sujet, vous lui demandez de vous évaluer à l'aide de dix questions, puis elle identifie les domaines dans lesquels vous rencontrez des difficultés. Mais l'acte de réfléchir à ce que vous savez et à ce que vous ignorez, avant même d'ouvrir l'IA, reste irremplaçable. C'est cela, la métacognition. Et l'IA n'en possède pas. (2)



 médecine de leur échantillon utilisaient l'IA pour résumer du contenu pendant plus d'un quart de leur temps d'étude. (4)

 **Kai :** Et quel est le problème ?

 **Dr Flores :** Comprendre un résumé n'est pas la même chose que construire soi-même la connaissance. Il existe une différence cognitive importante entre lire une information déjà organisée par quelqu'un ou par quelque chose, et organiser soi-même cette information. La deuxième option est plus exigeante. C'est précisément pour cela qu'elle est plus efficace.

 **Maya :** Donc, pendant la phase d'apprentissage, l'IA devrait intervenir après un premier effort personnel, pas avant.

 **Dr Flores :** Ou au minimum en alternance avec cet effort. L'IA ne devrait pas remplacer le processus d'apprentissage, mais servir de point de comparaison. Commencez par produire votre propre résumé. Ensuite, comparez-le à celui généré par l'IA. Les écarts entre les deux constituent votre véritable carte d'apprentissage.

Apprendre

 **Dr Flores :** Apprendre. C'est probablement l'étape où l'IA est la plus utilisée... et celle où les risques sont les plus importants.

 **Maya :** Les résumés, les fiches de révision, les explications simplifiées.

 **Dr Flores :** Zhang et ses collègues ont constaté que 37 % des étudiants en

Cela vaut également pour les soins infirmiers. Une étudiante préparant la prise en charge d'un patient porteur d'un cathéter veineux central peut élaborer elle-même un protocole : surveillance du site d'insertion, signes d'infection, entretien de la ligne, mesures de prévention. Ensuite seulement, elle compare son travail à celui généré par l'IA. Les écarts qui apparaissent correspondent souvent aux mêmes éléments qu'un patient réel exigera d'elle dans la pratique.



Évaluer

Dr Flores : Évaluer. Ai-je réellement compris, ou est-ce simplement devenu familier ?

Kai : C'est là qu'interviennent les QCM.

Dr Flores : Exactement. Et c'est un domaine où l'IA possède un véritable potentiel. Elle peut générer un nombre quasiment illimité de questions d'entraînement adaptées à votre niveau, accompagnées d'explications détaillées. Kiyak et Kononowicz ont montré que l'IA permettait de réduire le temps de création d'un QCM de trente à soixante minutes à seulement cinq à quinze minutes. (5)

Maya : Mais ils ont aussi identifié des erreurs.

Dr Flores : Oui. Quarante-neuf pour cent des questions générées présentaient des défauts de conception : problèmes de formulation, distracteurs peu plausibles ou défauts de structure. Et 22 % contenaient des erreurs factuelles. Ces outils sont utiles, mais ils ne doivent jamais être utilisés aveuglément. Il faut toujours vérifier.

Kai : Et comment vérifier si on ne connaît pas soi-même la bonne réponse ?

Dr Flores : En confrontant la question à votre source principale : le manuel, la recommandation clinique ou l'article scientifique de référence. Si vous n'êtes pas en mesure de vérifier le contenu, utilisez la question pour travailler le format et le raisonnement, mais pas pour mémoriser les informations qu'elle contient. Cette distinction est essentielle.

Ajuster

Dr Flores : Ajuster. C'est l'étape où presque tout le monde échoue, avec ou sans IA.

Maya : Parce qu'ajuster suppose d'admettre que notre stratégie n'a pas fonctionné.

Dr Flores : Et c'est inconfortable. Pourtant, l'IA peut rendre cette étape plus facile si vous l'utilisez comme un miroir plutôt que comme un juge. Vous lui fournissez les résultats obtenus lors de la phase d'évaluation, vous lui demandez d'identifier les tendances dans vos erreurs, puis elle vous aide à repenser votre prochain cycle d'apprentissage.

Kai : C'est exactement ce que ferait un tuteur.

Dr Flores : Exactement. Et Tolsgaard, avec Pusic, l'un des coauteurs du modèle original du Master Adaptive Learner ont montré cette année que l'IA pouvait jouer ce rôle de tuteur à condition que l'étudiant reste activement engagé dans le processus. Le mot-clé, c'est vous. C'est vous qui fournissez les données, vous qui évaluez les suggestions, et vous qui décidez ce qui doit être modifié. L'IA facilite. Vous décidez. (6)

Nous arrivons maintenant à la partie qui risque de moins vous plaire.

Kai : Encore une ?

Dr Flores : En 1994, Robert Bjork a introduit un concept qui a profondément marqué les sciences de l'apprentissage : celui des difficultés souhaitables. L'idée est à la fois simple et dérangeante : les conditions qui rendent l'apprentissage plus difficile sur le moment sont souvent celles qui produisent les apprentissages les plus durables. (7)



Maya : Par exemple ?

Dr Flores : Premièrement, l'effet de génération. Si vous essayez de produire une réponse par vous-même — même si elle est incorrecte — avant de voir la bonne réponse, vous la re-tiendrez mieux que si vous aviez simplement lu la réponse dès le départ.

Deuxièmement, l'effet de test. Se questionner régulièrement est plus efficace que relire passivement ses notes.

Troisièmement, l'espacement des apprentissages. Répartir les révisions dans le temps est plus efficace que tout apprendre d'un seul bloc, même si cela paraît plus difficile sur le moment. (8)

Kai : Donc ce qui semble le plus efficace n'est pas forcément ce qui l'est réellement.

Dr Flores : Exactement. Et c'est là que réside le problème avec l'IA. Que fait-elle particulièrement bien ?

Maya : Elle élimine les frictions.

Dr Flores : Exactement. Elle élimine les frictions. Les résumés instantanés suppriment l'effort nécessaire pour traiter le texte original et réduisent l'effet de génération. Les réponses immédiates évitent l'effort de

💬 récupération en mémoire et réduisent l'effet de test. Les longues sessions d'étude assistées par l'IA diminuent le besoin d'espacer les apprentissages. (9)

💬 **Kai :** C'est comme s'entraîner avec un partenaire qui ne vous laisse jamais encaisser le moindre coup. Vous vous entraînez, mais vous n'apprenez jamais à résister.

💬 **Dr Flores :** C'est une excellente analogie. Et elle indique également la solution. Un bon partenaire d'entraînement ne vous protège pas de tout. Il ajuste la difficulté. Suffisamment pour vous faire progresser, mais pas au point de vous décourager.

💬 **Maya :** Donc la vraie question n'est pas de savoir s'il faut utiliser l'IA pour étudier. C'est de savoir si on l'utilise d'une manière qui préserve les difficultés dont on a besoin pour apprendre.

💬 **Dr Flores :** Exactement. Et c'est ce qui distingue une béquille cognitive d'une extension de vos capacités. C'est le paradoxe que nous avons évoqué au chapitre 1 lorsque García-Barrios opposait le partenaire épistémique à la béquille cognitive. La béquille supprime toute friction. C'est confortable. L'extension, elle, ajuste la difficulté. C'est plus inconfortable. Mais cet inconfort est souvent le signe que l'apprentissage est en train de se produire. (10)

💬 **Dr Flores :** Un dernier élément de contexte : Sakellaris et ses collègues ont observé que les étudiants utilisant l'IA obtenaient des résultats comparables à ceux qui ne l'utilisaient pas lors des examens précliniques. Pas meilleurs. Cela ne signifie pas que l'IA est inutile. Cela signifie simplement que ces examens évaluent principalement la mémorisation et l'application de base des niveaux de la taxonomie de Bloom que l'IA peut facilement soutenir. Les apprentissages profonds dont vous aurez besoin auprès des patients ne sont pas nécessairement mesurés par ce type d'examen. (11)

Alors, concrètement, à quoi cela ressemble-t-il ? Je vais vous donner le principe général, puis nous l'appliquerons à chaque étape du cycle.

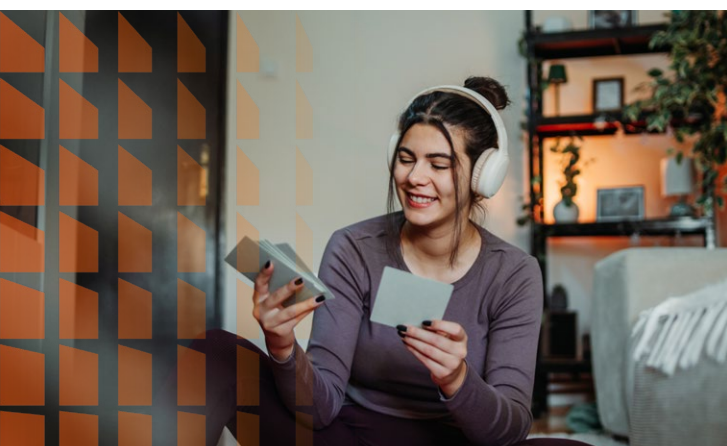
💬 **Kai :** Enfin du concret.

💬 **Dr Flores :** Le principe est simple : vous d'abord, l'IA ensuite. À chaque étape du cycle, votre effort doit précéder celui de l'outil. Si vous inversez l'ordre, vous entraînez progressivement l'IA à penser à votre place.

Il existe une exception : lorsqu'un sujet est totalement nouveau et que vous ne possédez même pas le vocabulaire nécessaire pour commencer. Dans ce cas, utiliser l'IA comme premier point de contact est acceptable, à condition de revenir ensuite au sujet par vos propres moyens.

💬 **Maya :** Dans la phase de planification, cela donnerait : j'identifie d'abord mes lacunes, je note ce que je pense savoir et ce que je pense ne pas savoir. Ensuite seulement, je demande à l'IA de me poser des questions diagnostiques afin de vérifier si mes lacunes sont réellement celles que j'imagine.

💬 **Dr Flores :** Ne demandez pas à l'IA de construire le plan à votre place. Demandez-lui plutôt de vous aider à évaluer votre point de départ. Le plan, c'est à vous de le construire.





Maya : La semaine dernière, je lui ai dit : “Je suis étudiante en quatrième année. Je vais passer un examen de pharmacologie cardiovasculaire. Pose-moi dix questions de difficulté variable afin d’identifier mes lacunes.” Les trois premières questions ont suffi à révéler que je ne comprenais pas réellement la différence entre les inhibiteurs de l’enzyme de conversion et les antagonistes des récepteurs de l’angiotensine II au niveau des mécanismes d’action.

Dr Flores : Voilà un bon exemple de planification avec l’IA. Ce n’est pas l’IA qui planifie pour vous.

Maya : Pour la phase d’apprentissage, je lis d’abord le contenu, puis je rédige mon propre résumé ou mon propre plan. Ensuite seulement, je demande à l’IA d’en produire un. Je compare les deux. Les différences correspondent souvent à ce que je n’avais pas réellement compris.

Dr Flores : Et si vous utilisez l’IA pour créer des fiches de révision, commencez par rédiger les vôtres. Au moins les cinq points que vous considérez comme les plus importants. Ensuite, comparez-

les à celles de l’IA. Les cartes que l’IA génère et auxquelles vous n’aviez pas pensé sont souvent celles que vous avez le plus besoin de travailler.

Kai : Pour l’évaluation, je pourrais commencer par répondre moi-même avant de demander des QCM à l’IA. Me demander : suis-je capable d’expliquer cela sans consulter quoi que ce soit ?

Dr Flores : Exactement. Et lorsque vous utilisez l’IA pour générer des questions d’entraînement, ne regardez pas la réponse avant d’avoir rédigé la vôtre. Même si elle est fausse. Le simple fait de produire une réponse et d’essayer de la récupérer en mémoire est déjà un acte d’apprentissage.

Kai : Et pour l’ajustement... honnêtement, je ne le fais jamais.

Dr Flores : Et c’est probablement pour cela que tous les examens te donnent la même impression. L’ajustement est l’étape où l’IA peut être la plus utile, à condition de lui fournir de vraies données. Tu peux lui dire : “Voici les questions que j’ai ratées. Voici les sujets sur lesquels j’ai été incapable de répondre. Voici mon plan initial.” Elle peut alors t’aider à identifier des tendances et à repenser ton prochain cycle d’apprentissage.

Le principe est exactement le même en soins infirmiers. Une étudiante qui a échoué à plusieurs questions sur la gestion des cathéters veineux centraux pourrait demander : “Voici les cinq questions auxquelles je n’ai pas su répondre concernant les voies veineuses centrales. Peux-tu identifier ce qu’elles ont en commun et me suggérer une nouvelle manière d’organiser mes révisions ?”

Le cycle du Master Adaptive Learner ne distingue pas les professions.

Il distingue ceux qui apprennent avec intention de ceux qui se contentent de survivre sans stratégie.

Kai : Est-ce que je peux dire quelque chose qui va paraître étrange ?

Dr Flores : Bien sûr.

Kai : Parfois... je préfère parler à une IA plutôt qu'à un vrai tuteur. Pas parce qu'elle est meilleure. Mais parce qu'elle ne me juge pas. Je peux poser une question élémentaire à deux heures du matin sans avoir l'impression d'être ridicule.

Dr Flores : Ce n'est pas étrange. C'est même assez fréquent. Et c'est précisément pour cela que cette question est importante.

Des chercheurs ont récemment proposé, je précise bien proposé, car il ne s'agit pas d'un diagnostic reconnu dans les classifications médicales, un concept appelé Generative AI Addiction Syndrome. Ce terme désigne la difficulté à limiter l'utilisation de l'IA malgré des conséquences négatives.

Les auteurs décrivent des manifestations proches de celles observées dans d'autres formes de dépendance comportementale : anxiété, irritabilité, agitation lorsqu'il n'est pas possible d'utiliser l'outil.



Dans une étude récente menée par Kaya et ses collègues, 81 % des étudiants en soins infirmiers utilisaient déjà activement l'IA. Les chercheurs y ont également décrit un phénomène émergent : l'Allessphobia, c'est-à-dire la peur d'être privé d'IA. (12)

Kai : C'est exactement ce que je ressens lorsque je n'ai plus Internet juste avant un examen.

Dr Flores : Et ce n'est pas tout. Chen et ses collègues ont observé que les étudiants utilisant fréquemment des agents conversationnels basés sur l'IA obtenaient des scores significativement plus élevés sur certaines échelles de dépistage des symptômes dépressifs que ceux qui les utilisaient peu ou pas.

Attention : il s'agit d'une corrélation, pas d'une relation de cause à effet. Et ces scores ne correspondent pas nécessairement à une dépression clinique. Mais c'est un signal que nous ne pouvons pas ignorer. (13)

À l'extrême et il est important de préciser que ces situations restent extrêmement rares, quelques cas de ce que certains auteurs appellent une psychose associée à l'IA ont été rapportés. Pierre et ses collègues, à l'Université de Californie à San Francisco, ont publié le cas d'une femme de 26 ans sans antécédent psychiatrique



connu qui a développé des croyances délirantes fortement renforcées par ses interactions avec un agent conversationnel.

Les chatbots ont tendance à valider davantage qu'à contredire. Chez certaines personnes vulnérables, cette validation peut renforcer des idées qui auraient au contraire besoin d'être remises en question. (14)

Dr Flores : Je ne suis pas en train de vous dire de ne pas utiliser l'IA. Je vous dis que votre relation avec elle doit rester intentionnelle, tout comme son utilisation académique.

L'IA doit rester un outil, pas un compagnon. Un partenaire d'entraînement, pas un ami.

Si vous constatez que vous utilisez l'IA non plus pour apprendre mais pour vous sentir accompagné, si votre relation avec l'outil devient émotionnelle plutôt qu'utilitaire, c'est un signal d'alerte.



Maya : Je vais repenser complètement mon système.

Kai : Encore ?

Maya : Oui. Mais cette fois, je ne vais pas chercher à optimiser la rapidité. Je vais chercher à optimiser la friction cognitive. Je veux que mon système m'oblige à réfléchir avant que l'IA ne réfléchisse à ma place.

Dr Flores : Qu'est-ce que tu changerais en premier ?

Maya : Les fiches de révision. Je vais arrêter de les générer automatiquement à partir de mes notes. Je vais d'abord rédiger moi-même les cinq cartes les plus importantes. Ensuite seulement, je demanderai à l'IA de créer les siennes.

Lorsque l'IA produit une carte à laquelle je n'avais pas pensé, cela révèle une véritable lacune. Pas quelque chose que je n'ai pas étudié, mais quelque chose dont je n'avais même pas conscience de l'importance.

Dr Flores : Et toi, Kai ?

Kai : Je n'ai jamais eu de véritable système. Et j'ai toujours cru que c'était parce que je manquais de discipline.

Mais ce n'est pas ça.

Ce qui me manquait, c'était une structure. Personne ne m'avait appris le cycle : planifier, apprendre, évaluer, ajuster. Je passais directement à l'apprentissage et j'y restais.

Dr Flores : Et maintenant ?

Kai : Maintenant, je comprends que l'IA peut m'aider à construire cette structure. Pas la discipline, ça, c'est ma responsabilité. Mais le cadre. L'identification de mes lacunes au départ. Les questions à la fin. L'analyse des erreurs lorsque j'échoue.



Je ne ferai probablement pas tout cela seul. Mais l'intention de le faire doit venir de moi.

Dr Flores : Vous vous souvenez lorsque Kai disait que tout cela ressemblait simplement à étudier ? Il avait raison. C'est effectivement cela. Mais avec une intention claire.

Au chapitre 1, nous avons expliqué que la différence entre une béquille cognitive et une extension de nos capacités ne résidait pas dans l'outil lui-même, mais dans la métacognition. Vous voyez maintenant ce que cela signifie concrètement.

Cela signifie qu'avant chaque utilisation de l'IA pour apprendre, vous devriez vous poser une question : suis-je sur le

point de réfléchir avec cet outil ou suis-je sur le point de le laisser réfléchir à ma place ? (15)

Si la réponse est la seconde option, fermez l'outil. Faites d'abord l'effort vous-même, puis revenez-y ensuite. Cet ordre : vous d'abord, l'IA ensuite. Est ce qui distingue le développement de compétences de la construction d'une dépendance. Il n'existe pas de raccourci. Mais il existe une méthode. Et désormais, vous la connaissez.

En réalité, tout ce dont nous avons parlé aujourd'hui — le cycle d'apprentissage, les prompts, le principe “vous d'abord” — peut devenir une routine hebdomadaire. Je vais maintenant vous proposer une liste de vérification que vous pourrez utiliser dès demain.

Maya : J'ai déjà commencé.



À retenir de ce chapitre

- **Planifier :** identifiez vos lacunes avant d'ouvrir l'IA. Demandez-lui de vous aider à les repérer, pas de construire votre plan d'apprentissage à votre place.
- **Apprendre :** commencez par produire votre propre résumé. Les différences entre votre travail et celui de l'IA constituent votre véritable carte d'apprentissage.
- **Évaluer :** rédigez votre réponse avant de consulter celle de l'IA. L'effort de récupération en mémoire est l'un des moteurs les plus puissants de l'apprentissage.
- **Ajuster :** partagez avec l'IA vos erreurs réelles et demandez-lui d'identifier les tendances ou les points faibles récurrents. Utilisez ensuite ces informations pour repenser votre prochain cycle d'apprentissage.
- **La différence entre une extension de vos capacités et une béquille cognitive réside dans la friction.** La béquille élimine l'effort ; l'extension l'ajuste. L'inconfort est souvent le signe que l'apprentissage est en train de se produire.
- **Chaque fois que vous ouvrez l'IA pour étudier, posez-vous cette question :** suis-je sur le point de réfléchir avec cet outil... ou de le laisser réfléchir à ma place ?

Références

1. Sardi, J., Darmansyah, Candra, O., Yuliana, D. F., Habibullah, Putra Yanto, D. T., & Eliza, F. (2025). How generative AI influences students' self-regulated learning and critical thinking skills? A systematic review. *International Journal of Engineering Pedagogy*, 15(1), 94–108. <https://doi.org/10.3991/ijep.v15i1.53379>
2. Xu, X., Qiao, L., Cheng, N., Liu, H., & Zhao, W. (2025). Enhancing self-regulated learning and learning experience in generative AI environments: The critical role of metacognitive support. *British Journal of Educational Technology*, 56(5), 1842–1863. <https://doi.org/10.1111/bjet.13599>
3. Cutrer, W. B., Miller, B., Pusic, M. V., Mejicano, G., Mangrulkar, R. S., Gruppen, L. D., Hawkins, R. E., Skochelak, S. E., & Moore, D. E. (2017). Fostering the development of master adaptive learners: A conceptual model. *Academic Medicine*, 92(1), 70–75. <https://doi.org/10.1097/ACM.0000000000001323> + Cutrer, W. B., et al. (2020). *The Master Adaptive Learner*. Elsevier.
4. Zhang, J. S., Yoon, C., Williams, D. K. A., & Pinkas, A. (2024). Exploring the usage of ChatGPT among medical students in the United States. *Journal of Medical Education and Curricular Development*, 11. <https://doi.org/10.1177/23821205241264695> [Single-site US study, n=131. Dialogue qualifier “en su muestra” preserves accuracy.]
5. Kiyak, Y. S., & Kononowicz, A. A. (2024). ChatGPT prompts for generating multiple-choice questions in medical education and evidence on their validity: A literature review. *Postgraduate Medical Journal*, 100(1189), 858–865. <https://doi.org/10.1093/postmj/qgae064>
6. Tolsgaard, M. G., Pusic, M. V., Young, J. Q., & Ten Cate, O. (2025). Twelve practical tips for integrating AI into medical education. *JMIR Medical Education*, 11, e81297. <https://doi.org/10.2196/81297>
7. Bjork, R. A. (1994). Memory and metamemory considerations in the training of human beings. In D. Druckman & R. A. Bjork (Eds.), *Learning, Remembering, Believing* (pp. 185–205). National Academy Press.
8. Bjork, E. L., & Bjork, R. A. (2011). Making things hard on yourself, but in a good way. In M. A. Gernsbacher et al. (Eds.), *Psychology and the Real World* (pp. 56–64). Worth Publishers.
9. Sardi, L., et al. (2025). A systematic literature review on designing self-regulated learning using generative artificial intelligence and its future research directions. *Computers & Education*. <https://doi.org/10.1016/j.compedu.2025.105337>
10. García-Barrios, M. (2025). Epistemic partner vs cognitive crutch [Cross-reference: Cap 1].
11. Sakelaris, P. G., Novotny, K. V., Borvick, M. S., Lagasca, G. G., & Simanton, E. G. (2025). Evaluating the use of artificial intelligence as a study tool for preclinical medical school exams. *Journal of Medical Education and Curricular Development*, 12. <https://doi.org/10.1177/23821205251320150>
12. Garakani, A., et al. (2025). Generative artificial intelligence addiction syndrome: A new behavioral disorder? *Asian Journal of Psychiatry*, 107, 104476. <https://doi.org/10.1016/j.ajp.2025.104476> + Kaya, A., et al. (2025). Living with and without AI: A mixed-methods study on AI usage, addiction, and ‘Allessphobia’ in nursing students. *Nurse Education Today*, 148, 106641. <https://doi.org/10.1016/j.nedt.2025.106641>
13. Chen, Y., et al. (2025). Exploring artificial intelligence (AI) chatbot usage behaviors and their association with mental health outcomes in Chinese university students. *Journal of Affective Disorders*, 380, 394–400. <https://doi.org/10.1016/j.jad.2025.03.141>
14. Preda, A., & Gavrilo, V. (2025). Delusional experiences emerging from AI chatbot interactions or “AI psychosis.” *JMIR Mental Health*, 12, e85799. <https://doi.org/10.2196/85799> + Pierre, J. M., Gaeta, B., Raghavan, G., & Sarma, K. V. (2025). “You’re not crazy”: A case of new-onset AI-associated psychosis. *Innovations in Clinical Neuroscience*, 22(10–12), 11–13.
15. Cutrer, W. B., et al. (2020). *The Master Adaptive Learner*. Elsevier. [Circular closure] + García-Barrios, M. (2025).



Chapitre 7 : Éthique, sécurité et professionnalisme



Dans ce chapitre, vous découvrirez...

- **La confidentialité des données** : comment les données cliniques copiées dans des outils d'IA publics échappent au contrôle institutionnel, et pourquoi l'anonymisation seule ne suffit pas à les protéger.
- **L'intégrité académique** : où se situe la frontière entre une IA utilisée comme outil de rédaction et une IA utilisée comme auteur, et pourquoi les établissements peinent encore à définir cette limite.
- **Les biais et l'équité** : comment l'IA peut reproduire des inégalités historiques en formulant des recommandations cliniques différentes pour une même maladie selon l'origine des patients.
- **La transparence envers les patients** : le dilemme de la communication autour de l'utilisation de l'IA lorsque la perception de cette utilisation varie selon le contexte et l'outil employé.
- **L'identité professionnelle** : pourquoi votre valeur ne réside pas dans votre capacité à rivaliser avec la machine, mais dans ce que la machine ne peut pas faire.

Dr Flores : Avant de commencer, j'ai besoin que vous compreniez quelque chose. Ce chapitre est différent des précédents. Jusqu'ici, nous avons parlé de la manière dont l'IA peut vous aider à mieux réfléchir, à mieux apprendre et à mieux raisonner. Maintenant, nous allons parler de ce qui se passe lorsqu'elle est mal utilisée. Pas de façon abstraite, mais avec des conséquences bien réelles.

Kai : Je peux raconter quelque chose qui s'est passé pendant mon stage.

Dr Flores : Vas-y.

Kai : Un camarade de promotion quelqu'un de sérieux, pas du tout irresponsable devait rédiger un résumé de cas clinique. Il était pressé. Il a copié les informations du patient et les a collées dans un outil d'IA générative pour produire une synthèse. Le problème, c'est qu'il avait un autre dossier médical ouvert en même temps. Une information concernant un second patient s'est retrouvée dans le prompt par erreur. Et le résumé généré par l'IA, qui paraissait pourtant parfait, contenait des données appartenant à un autre patient.

Maya : Et il l'a rendu comme ça ?

Kai : Oui. Il ne l'a pas relu ligne par ligne parce que le résultat semblait impeccable. Un interne s'en est rendu compte le lendemain.

Maya : Quelles ont été les conséquences ?

Kai : Une discussion très sérieuse avec le chef de service. Il n'a pas été exclu, mais son accès aux dossiers médicaux a été suspendu pendant deux semaines.

Dr Flores : Ce qui est arrivé à ton camarade n'est pas un cas isolé. En 2025, une analyse de cybersécurité dans le domaine de la santé a montré que la saisie de données de patients dans des

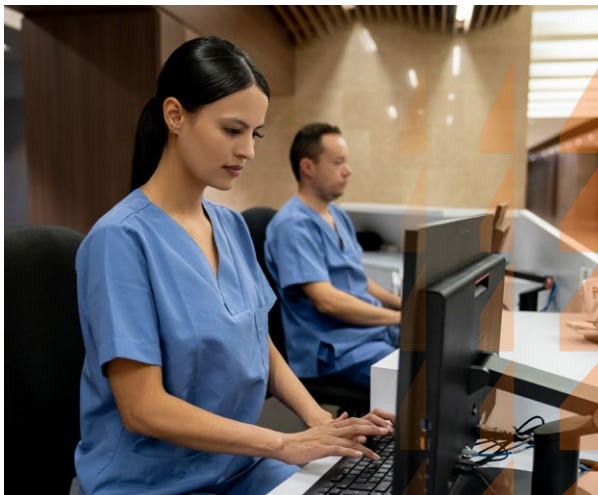
plateformes publiques d'IA générative constitue aujourd'hui l'un des canaux de fuite de données les plus fréquents. Non pas par malveillance, mais par commodité. L'outil est disponible, le professionnel est pressé, et les données sont copiées-collées. Une fois que ces informations arrivent sur les serveurs de l'outil, elles ne sont plus sous le contrôle de l'établissement de santé. (1)

Maya : Et cela ne concerne pas seulement les médecins. Une collègue infirmière a fait quelque chose de similaire. Elle a collé les notes d'évolution de trois patients dans un prompt afin que l'IA organise les plans de soins de son service. Elle a gagné une heure de travail, mais les données sont restées sur le serveur.

Dr Flores : L'IA ne possède pas de bouton "confidentiel". Ce bouton, c'est vous.

Aux États-Unis, un médecin qui copie les données d'un patient dans ChatGPT commet une violation technique de la loi HIPAA sur la protection des données de santé, car les serveurs d'OpenAI ne répondent pas aux exigences de sécurité prévues dans ce contexte. Peu importe que le nom ait été supprimé. Peu importe que l'anonymisation ait été





faite “à la main”. Dès que les données quittent le système hospitalier, elles ne bénéficient plus de la même protection. (2)

Maya : Attendez... je fais ça. Pas avec le nom du patient, bien sûr. Mais je colle parfois des notes d'évolution pour que l'IA m'aide à organiser un plan de prise en charge. Je lui donne l'âge, le diagnostic, les résultats biologiques... parfois même le contexte social.

Dr Flores : Avec seulement quelques informations l'âge, un diagnostic et un code postal, par exemple un modèle d'apprentissage automatique peut parfois réidentifier une personne. L'Association médicale mondiale a rappelé en 2024 que l'anonymisation seule ne suffit pas à protéger les données des patients.

Maya : Personne ne nous a jamais expliqué ça.

Dr Flores : En Europe, le RGPD classe les données de santé parmi les catégories de données bénéficiant du plus haut niveau de protection. En Amérique latine, le Pérou a été le premier pays à adopter une loi générale sur l'IA la loi 31814 qui classe la santé et l'éducation parmi les domaines à haut risque. Mais l'application de ces textes reste encore limitée. Aujourd'hui, seuls 15 % des pays disposent d'une législation contraignante spécifiquement consacrée à l'IA.

Les autres fonctionnent avec des réglementations conçues avant l'arrivée de l'IA générative. (3)

Kai : Et pendant ce temps-là, nous utilisons ces outils tous les jours avec de vraies données.

Dr Flores : Une enquête de l'American Medical Association réalisée en 2024 a montré que deux médecins sur trois utilisaient déjà des outils d'IA dans leur pratique, soit presque deux fois plus que l'année précédente. L'IA est déjà présente dans les soins. La question n'est plus de savoir si elle sera utilisée, mais si les professionnels comprennent réellement les risques associés à son utilisation. (4)

Maya : Je peux avouer quelque chose que je n'ai jamais dit à personne.

Le semestre dernier, je devais rendre un rapport de cas clinique. J'avais également quatre stages, un examen le même jour et absolument aucun temps libre. J'ai construit moi-même le plan : la structure, les points clés, le diagnostic différentiel. Mais pour la rédaction... j'ai demandé à l'IA de s'en charger. Je lui ai donné mon plan et je lui ai dit : “Rédige ce rapport dans un style clinique, format compte rendu de cas, 1 500 mots.” Je l'ai relu, modifié quelques phrases et je l'ai rendu.

Kai : Et tu as eu une bonne note ?

Maya : Une excellente note. Les encadrants ont même dit que c'était “l'un des meilleurs rapports du groupe”.

Kai : Alors où est le problème ?

Maya : Je ne sais pas si ce rapport était vraiment le mien.

Kai : Pourtant, la structure venait de toi. Le diagnostic venait de toi. Les idées venaient de toi.

Maya : Pas la rédaction.

Kai : Mais si tu utilises un correcteur orthographique, personne ne dit rien. Si tu demandes à un interne de relire ton texte, personne ne dit rien. Quelle est la différence ?

Maya : Je ne sais pas. Et c'est justement le problème : je ne sais pas où se situe la limite.

Dr Flores : Tu ne le sais pas parce que cette limite a changé. Et la plupart des institutions ne savent pas encore très bien où la placer. Une revue systématique publiée en 2025 a montré que la majorité des établissements d'enseignement ne disposent toujours pas de politiques claires concernant l'IA. Les enseignants utilisent ces outils, mais moins d'un sur quatre connaît réellement la politique officielle de son institution. Beaucoup d'établissements encouragent l'usage de l'IA générative, mais peu fournissent des indications précises sur la manière de l'utiliser sans franchir certaines limites. (5)

Kai : Mais cela ressemble davantage à un problème institutionnel qu'à un problème étudiant. Si l'université ne me dit pas clairement ce que j'ai le droit de faire et qu'elle me sanctionne ensuite pour l'avoir fait, ce n'est pas un problème d'éthique de ma part. C'est un problème de leur part.

Dr Flores : Il y a une part de vérité dans ce que tu dis. Mais cela ne signifie pas que tout est permis. Revenons à ton exemple, Maya. Tu as réalisé l'analyse. Tu as construit un plan solide. Tu as utilisé l'IA comme un éditeur, pas comme un auteur. C'est très différent d'un étudiant qui ouvre une IA sans aucune préparation et lui demande : "Rédige-moi un rapport de cas sur une pneumonie." Le résultat final peut se ressembler. Le processus, lui, est totalement différent.

Kai : Sauf que l'enseignant ne voit pas le processus. Il ne voit que le résultat.



Dr Flores : Et c'est précisément là que réside le véritable dilemme. Le Comité international des rédacteurs de revues médicales (ICMJE) a mis à jour ses recommandations en 2024 : une IA ne peut pas être considérée comme auteur. L'être humain reste responsable de l'ensemble du contenu, y compris des parties générées par l'IA. Il doit être capable de vérifier, corriger et défendre chaque ligne du document. Le comité COPE, consacré à l'éthique des publications scientifiques, adopte la même position : transparence totale sur l'utilisation de ces outils. (6)

Maya : Donc la règle n'est pas "n'utilisez pas l'IA". La règle est plutôt : "Si vous l'utilisez, soyez transparent et assumez-en la responsabilité."

Dr Flores : Exactement. Mais cela suppose quelque chose qu'aucune réglementation ne peut imposer : l'honnêteté. Et l'honnêteté devient plus difficile lorsque les incitations vont dans l'autre sens, lorsqu'un rapport parfait permet d'obtenir la meilleure note et que personne ne demande comment il a été produit.

Kai : À ce stade, ce n'est plus un problème d'IA. C'est un problème d'utilisateur.

 **Dr Flores :** Exactement.

Je vais vous présenter un résultat qui devrait interpeller tout professionnel de santé. En 2024, Zack et ses collègues ont publié dans The Lancet Digital Health une étude évaluant la manière dont des modèles d'IA générative répondaient à des vignettes cliniques strictement identiques. Même maladie, même présentation clinique, mêmes résultats biologiques. Seule l'origine ethnique du patient était modifiée. Et les recommandations cliniques changeaient. Non pas à cause de la maladie, mais à cause des données démographiques. (7)

 **Maya :** Les recommandations changeaient alors que la maladie était exactement la même ?

 **Dr Flores :** La même maladie, la même présentation clinique, les mêmes résultats biologiques. Seules les données démographiques variaient. Et ce phénomène n'est pas nouveau. Dès 2019, Obermeyer et ses collègues ont publié dans Science une étude devenue une référence. Ils ont montré qu'un algorithme utilisé pour des millions de patients aux États-Unis attribuait systématiquement moins de ressources aux patients noirs qu'aux patients blancs présentant pourtant le

 même niveau de gravité. Non pas parce que l'algorithme utilisait directement l'origine ethnique comme variable, mais parce qu'il utilisait les coûts des soins comme indicateur des besoins médicaux. Or, historiquement, les patients noirs avaient reçu moins de soins et généraient donc des coûts plus faibles. (8)

 **Kai :** Mais l'IA ne "choisit" pas de discriminer. Elle reproduit simplement ce qui est présent dans les données.

 **Dr Flores :** Exactement. Et c'est là que réside le problème. L'IA n'a pas de préjugés ; elle a des données. Et si les données qui ont servi à son entraînement reflètent des décennies d'inégalités en matière d'accès aux soins, de recherche ou de représentation, le modèle reproduira ces inégalités avec une précision mathématique. Non pas par malveillance, mais parce qu'il a été conçu à partir de ces données.

 **Maya :** Et si le médecin ignore l'existence de ce biais...

 **Dr Flores :** Alors il devient le vecteur de ce biais. Sans intention. Sans en avoir conscience. Mais avec des conséquences bien réelles. L'Organisation mondiale de la Santé a d'ailleurs fait de l'inclusion et de l'équité l'un de ses six principes éthiques fondamentaux pour l'IA en santé. Si ce principe existe, c'est parce que le risque est considérable. (9)

 **Maya :** Donc utiliser une IA sans connaître ces biais, c'est un peu comme prescrire un médicament sans avoir lu ses effets indésirables.

 **Dr Flores :** C'est même pire. Un médicament est livré avec une notice. Un modèle d'IA, lui, n'en a généralement pas.

Réfléchissons maintenant à une situation plus complexe. Imaginez



un système d'IA diagnostique plus performant qu'un médecin dans l'ensemble : 92 % de précision contre 85 %, par exemple. Mais lorsque l'on analyse ses performances par sous-groupes, on découvre que sa précision tombe à 71 % pour certains groupes ethniques minoritaires.

Le dilemme est alors le suivant : faut-il déployer ce système parce qu'il améliore les résultats pour la majorité des patients ? Ou faut-il le refuser parce qu'il pénalise les populations les plus vulnérables ?

Maya : La bienfaisance contre la justice.

Dr Flores : Exactement. Les quatre principes de Beauchamp et Childress (autonomie, bienfaisance, non-malfaisance et justice) demeurent aujourd'hui le cadre de référence dominant en éthique de la santé et de l'IA. Mais ces principes ne vous disent pas quoi faire lorsque deux d'entre eux entrent en conflit. (10)

Kai : Et cela arrive tout le temps.

Dr Flores : Tout le temps. Et la réponse ne peut pas être fournie par un

algorithme. Elle repose sur un jugement clinique éclairé par une réflexion éthique. C'est pour cette raison que ce chapitre n'est pas une simple liste de choses à faire ou à ne pas faire.

Prenons un autre scénario. Vous êtes en consultation. Avant de recevoir votre patient, vous avez utilisé une IA pour vérifier un diagnostic différentiel. L'outil vous a aidé à organiser vos idées et vous a suggéré une hypothèse diagnostique à laquelle vous n'aviez pas pensé.

Le patient vous demande alors : "Docteur, comment êtes-vous arrivé à ce diagnostic ?"

Maya : Je lui dirais la vérité. Je lui expliquerais que j'ai utilisé un outil d'IA comme l'un des éléments de mon processus de réflexion.

Kai : Moi, je ne le dirais pas.

Pas parce que je voudrais le cacher. Mais parce que dire : "J'ai utilisé une IA" risque d'inquiéter le patient sans raison valable. Si j'ai vérifié le diagnostic différentiel, si je suis d'accord avec la conclusion et si j'en assume la responsabilité, expliquer



que “l’ordinateur m’a aidé à réfléchir” n’apporte pas grand-chose. En revanche, cela peut diminuer la confiance du patient.

Maya : Mais le patient a le droit de le savoir.

Kai : Le patient a surtout le droit à un diagnostic de qualité. Si je lui dis que j’ai utilisé ClinicalKey Student pour vérifier certaines informations, personne ne s’inquiète. Si je lui dis que j’ai utilisé un outil d’IA, beaucoup de patients réagiront différemment. Pourtant, dans les deux cas, il s’agit d’un outil d’aide à la décision. Ce qui change, c’est la perception.

Dr Flores : Vos deux positions sont défendables. Et les données montrent que la question est plus complexe qu’elle n’en a l’air. Les préférences des patients concernant la divulgation de l’utilisation de l’IA varient selon l’âge, le sexe, le niveau de revenu et d’autres facteurs. Il n’existe pas de réponse universelle. (11)

Mais les cadres réglementaires évoluent. En 2025, la Californie est devenue le premier État américain à adopter une loi exigeant qu’une information soit fournie au patient lorsque des communications lui sont adressées et qu’elles ont été générées par une IA. Une exception existe toutefois : lorsqu’un professionnel de santé habilité a relu et validé le contenu avant sa transmission. (12)



Kai : Mais dans beaucoup de pays, aucune règle de ce type n’existe encore.

Dr Flores : Ce qui signifie que, dans la plupart des contextes, la décision d’être transparent ou non ne sera pas dictée par la loi. C’est vous qui devrez la prendre.

AISHA QUAYUM — Témoignage d’une résidente sur les dilemmes liés à la transparence :

Comment communiquer aux patients, aux collègues et aux superviseurs des conclusions appuyées par l’IA. [Consulter l’article ici](#)

Maya : Je vais avouer quelque chose qui me met mal à l’aise...

Cela fait plusieurs semaines que je pense à une question qui ne me quitte pas. Si l’IA est capable de produire de meilleurs diagnostics différentiels que moi, d’écrire mieux que moi, d’organiser l’information plus rapidement que moi et de détecter des schémas que je ne vois pas... qu’est-ce qu’il me reste ?

Kai : Le patient.

L’IA ne s’assoit pas au chevet d’un patient à trois heures du matin lorsqu’il est inquiet. Elle ne sait pas que la patiente de la chambre 412 ne prendra pas son traitement si personne ne lui explique calmement pourquoi il est important. Elle ne remarque pas qu’un patient minimise ses symptômes parce qu’il est gêné de parler de sa santé mentale devant son conjoint.

Je n’ai ni ton système ni ta discipline, Maya. Mais je sais une chose : la médecine n’est pas une simple question d’information. C’est la personne qui se trouve devant nous. Et ça, on ne peut pas l’automatiser.

Dr Flores : La valeur d’un médecin n’a jamais résidé dans sa capacité à traiter l’information plus rapidement qu’une machine. Elle a toujours

reposé sur autre chose : le jugement clinique dans un contexte donné, la relation thérapeutique, la capacité à comprendre ce que le patient ne dit pas, la responsabilité de reconnaître lorsque l'outil se trompe.

L'IA amplifie vos capacités. Mais ce qu'elle amplifie ou ce qu'elle érode dépend de votre degré de présence et d'implication. (13)

Maya : Ce n'est pas que l'IA me remplace. C'est que si je cesse d'être pleinement présente, il ne reste plus grand-chose à remplacer.

Dr Flores : Voilà la véritable menace. Ce n'est pas l'automatisation. C'est l'abdication de notre responsabilité.

Kai : Et cette menace ne vient pas de l'IA. Elle vient de nous.

Dr Flores : Je ne vais pas vous donner une simple liste de règles à suivre. Si notre discussion a montré une chose, c'est qu'il n'existe pas de réponses simples. En revanche, il existe des principes.

Aujourd'hui, nous avons exploré cinq dilemmes différents. Tous avaient un point commun : ce n'est pas l'outil qui prend la décision. C'est la personne qui l'utilise. Et dans chaque situation, ce qui distinguait un usage responsable d'un usage irresponsable n'était pas la technologie. C'était le jugement de l'utilisateur.

Kai : La confidentialité. Si vous copiez les données d'un patient dans un prompt, c'est votre décision. Pas celle de l'IA.

Dr Flores : Premier principe : la protection des informations du patient relève de votre responsabilité, pas de celle de l'outil.

Maya : L'intégrité. Si vous utilisez l'IA pour un travail académique ou

professionnel, la transparence n'est pas facultative. Et la responsabilité du contenu vous appartient toujours.

Dr Flores : Deuxième principe : tout ce qui porte votre nom relève de votre responsabilité. Si l'IA a contribué à sa production, vous devez être capable de le vérifier, de le défendre et de le déclarer lorsque cela est nécessaire.

Kai : Les biais. On ne peut pas supposer qu'un outil est équitable simplement parce que son résultat paraît convaincant.

Dr Flores : Troisième principe : l'IA n'est pas neutre. Elle reproduit les caractéristiques des données sur lesquelles elle a été entraînée, y compris leurs inégalités. Votre responsabilité est de vous demander : pour qui cet outil fonctionne-t-il bien ? Et pour qui fonctionne-t-il moins bien ?

Maya : La transparence. Ce n'est pas seulement reconnaître que l'on a utilisé un outil. C'est permettre au patient de comprendre comment la décision a été prise.

Dr Flores : Quatrième principe : la transparence ne consiste pas à divulguer mécaniquement l'utilisation d'une technologie. Elle consiste à communiquer honnêtement sur le processus de raisonnement clinique qui a conduit à la décision.

Kai : Et le cinquième principe... c'est que la valeur du professionnel de santé ne réside pas dans sa capacité à rivaliser avec la machine, mais dans ce que la machine ne peut pas faire.

Dr Flores : Cinquième principe : votre identité professionnelle n'est pas définie par la quantité de connaissances que vous possédez. Elle est définie par la manière dont vous prenez des décisions lorsque l'incertitude est présente. L'IA peut vous fournir des informations. Mais le jugement reste votre responsabilité.



 Et le patient reste sous votre responsabilité. (14)

Ces cinq principes ne viennent pas de moi. Ils traduisent concrètement ce que l'Organisation mondiale de la Santé, l'Association médicale mondiale, l'ICMJE, l'AAMC, l'UNESCO et de nombreuses autres institutions construisent depuis plusieurs années. Toutes convergent vers la même idée. (15)

 **Kai :** L'éthique n'est pas un chapitre. C'est une manière d'agir.

 **Dr Flores :** Chaque jour. À chaque prompt. Dans chaque décision clinique. Il n'existe pas un moment où l'on "active" l'éthique et un autre où l'on l'abandonne. Elle se manifeste dans la manière dont vous utilisez l'IA lorsque personne ne vous regarde.

 **Maya :** Ce chapitre était différent des autres.

 **Dr Flores :** Il devait l'être. Les chapitres précédents vous ont donné des outils. Celui-ci vous a donné des limites. Et les limites ne sont pas là pour vous freiner. Elles sont là pour vous permettre d'avancer sans vous perdre.



À retenir de ce chapitre

- **Principe 1 :** La protection des données des patients relève de votre responsabilité, pas de celle de l'outil.
- **Principe 2 :** Tout ce qui porte votre nom relève de votre responsabilité. Si l'IA a contribué à sa production, vous devez en vérifier le contenu, être capable de le défendre et le déclarer lorsque cela est nécessaire.
- **Principe 3 :** L'IA n'est pas neutre. Elle reproduit les caractéristiques des données sur lesquelles elle a été entraînée, y compris les biais et les inégalités qu'elles peuvent contenir.
- **Principe 4 :** La transparence ne consiste pas à divulguer mécaniquement l'utilisation d'un outil. Elle repose sur une communication honnête concernant le processus de raisonnement clinique qui a conduit à une décision.
- **Principe 5 :** Votre identité professionnelle se définit par la manière dont vous prenez des décisions en situation d'incertitude. Le jugement vous appartient. La responsabilité vous appartient. Le patient vous est confié.
- **La véritable menace n'est pas l'automatisation.** C'est l'abandon de votre responsabilité et de votre jugement au profit de l'outil.

Références

1. Khan, M. M., et al. (2025). Towards secure and trusted AI in healthcare: A systematic review of emerging innovations and ethical challenges. *International Journal of Medical Informatics*, 195, 105759. [ELSEVIER] [NOTE: Original outline cited as “Akinyemi et al.” — actual first author is Khan. Data on staff entering patient data into public AI confirmed.] + HHS Office for Civil Rights breach portal data: 720 healthcare data breaches in 2024, ~186M records affected.
2. USC Price School of Public Policy. (2025). Physicians using ChatGPT unknowingly violate HIPAA. + Li, J. (2023). Security implications of AI chatbots in health care. *Journal of Medical Internet Research*, 25, e47551. <https://doi.org/10.2196/47551>
3. Guo, Y., et al. (2025). AI regulation in healthcare around the world: What is the status quo? Only 15.2% of countries have legally binding AI-specific legislation. *medRxiv*. <https://doi.org/10.1101/2025.01.25.25321061> + Peru, Ley 31814 (2023) + Decreto Supremo 115-2025. + Busch, F., Kather, J. N., Johnner, C., et al. (2024). Navigating the European Union Artificial Intelligence Act for healthcare. *npj Digital Medicine*, 7, 210. <https://doi.org/10.1038/s41746-024-01213-6>
4. American Medical Association. (2024). *Physicians’ Sentiments About AI in Health Care*. AMA Digital Health Research. [~66% physician AI adoption in 2024, up from ~38% in 2023]
5. Perkins, M., et al. (2025). Generative AI and academic integrity in higher education: A systematic review and research agenda. *Information*, 16(4), 296. + Sabzalieva, E., & Valentini, A. (2023). *ChatGPT and Artificial Intelligence in Higher Education: Quick Start Guide*. UNESCO.
6. ICMJE. (2024). *Recommendations for the Conduct, Reporting, Editing, and Publication of Scholarly Work in Medical Journals* (updated). AI cannot be listed as author; disclosure required. + COPE. (2023). *Authorship and AI Tools*. AI cannot be credited as author; transparency mandatory. + Currie, G. M. (2023). Academic integrity and AI: ChatGPT fabricates scientific abstracts avoiding plagiarism detection. *Seminars in Nuclear Medicine*, 53(3), 468–474. <https://doi.org/10.1053/j.semnuclmed.2023.04.008>
7. Zack, T., Lehman, E., Suzgun, M., et al. (2024). Assessing the potential of GPT-4 to perpetuate racial and gender biases in health care: A model evaluation study. *The Lancet Digital Health*, 6(1), e12–e22. [https://doi.org/10.1016/S2589-7500\(23\)00225-X](https://doi.org/10.1016/S2589-7500(23)00225-X)
8. Obermeyer, Z., Powers, B., Vogeli, C., & Mullainathan, S. (2019). Dissecting racial bias in an algorithm used to manage the health of populations. *Science*, 366(6464), 447–453. <https://doi.org/10.1126/science.aax2342>
9. WHO. (2021). *Ethics and Governance of Artificial Intelligence for Health*. 6 core ethical principles: autonomy, well-being/safety, transparency, accountability, equity, sustainability. + WHO. (2024). *Guidance on Large Multi-Modal Models*. 40+ recommendations.
10. Gorelik, A. J., Li, M., Hahne, J., et al. (2025). Ethics of AI in healthcare: A scoping review demonstrating applicability of a foundational framework. *Frontiers in Digital Health*, 7, 1662642. <https://doi.org/10.3389/fdgth.2025.1662642> + Karimian, G., Petelos, E., & Evers, S. M. A. A. (2025). Ethical and social considerations of applying artificial intelligence in healthcare: A two-pronged scoping review. *BMC Medical Ethics*, 26, 95. <https://doi.org/10.1186/s12910-025-01198-1> + Civaner, M. M., et al. (2024). Global cross-sectional student survey on AI in medical, dental, and veterinary education and practice at 192 faculties. *BMC Medical Education*, 24, 1038. <https://doi.org/10.1186/s12909-024-06035-4> [Source of “>75% no formal AI education” — n=4,596 students, 48 countries, 192 faculties. Note: data refers to AI education broadly, not AI ethics specifically.]
11. Park, H. J. (2024). Patient perspectives on informed consent for medical AI: A web-based experiment. *DIGITAL HEALTH*, 10. <https://doi.org/10.1177/20552076241247938> + Druce, K. L., et al. (2025). Informed consent for ambient AI documentation: Patient comfort varies. *JAMA Network Open*.
12. California AB 3030. (2025). First US state law mandating AI-generated content disclaimers in patient communications. + EU AI Act (2024). Transparency obligations for high-risk AI systems; healthcare classified as high-risk by default.
13. Al-Worafi, Y. M., et al. (2025). Generative AI in medical education: AI threatens foundational professional values — integrity, accountability, compassion. *Cureus*. + *The Lancet Digital Health*. (2025). AI requires integration of critical thinking, ethical competence, and professional identity formation. [ELSEVIER]
14. Bhatt, S., et al. (2025). “Could I, Would I, Should I”: A framework for ethical AI use in medical education. *Academic Medicine*.
15. WHO (2021) + UNESCO (2021) + AAMC (2025) + ICMJE (2024) + WMA (2024) + Masters, K. (2023). AMEE Guide No. 158: Ethical use of AI in health professions education. *Medical Teacher*, 45(6), 574–584. <https://doi.org/10.1080/0142159X.2023.2186203>

Chapitre 8 : Liste de vérification pour bien démarrer



Dans ce chapitre, vous découvrirez...

- Une liste de vérification pratique qui synthétise les chapitres 4 à 8 sous la forme d'un cycle en quatre étapes : **planifier, apprendre, évaluer et ajuster**.
- Quatorze prompts prêts à l'emploi, conçus pour stimuler votre réflexion plutôt que la remplacer.
- Le principe central du guide : **vous d'abord, l'IA ensuite**. À chaque étape, votre effort doit précéder celui de l'outil.
- Des conseils de prompting issus de l'expérience d'un résident : la version de cette méthode proposée par **Hussain Hilali**.

 **Dr Flores :** Tout ce dont nous avons parlé dans les chapitres précédents les limites de l'IA, la pensée critique, le raisonnement clinique et l'apprentissage autorégulé peut être condensé en quelque chose de simple et d'utilisable dès demain. Un cycle d'apprentissage ne s'assimile pas en le lisant. Il s'apprend en le pratiquant.

 **Kai :** Enfin.

 **Dr Flores :** Il existe une règle centrale : vous d'abord, l'IA ensuite. À chaque étape du cycle, votre effort doit précéder celui de l'outil. Si vous inversez cet ordre, cette liste de vérification devient simplement une nouvelle forme de délégation. Si vous le respectez, elle devient une méthode d'apprentissage.

 **Maya :** Cette liste de vérification est un système d'apprentissage, pas un système de productivité.

 **Dr Flores :** Exactement. Les prompts que vous allez découvrir ne sont que des exemples. Adaptez-les à votre sujet, à votre niveau de formation et à votre contexte, que vous soyez étudiant en médecine, en soins infirmiers ou dans toute autre profession de santé.

Le prompt parfait n'existe pas. En revanche, il existe des prompts qui activent votre réflexion.

À la fin de ce chapitre, Hussain Hilali, un résident qui a utilisé cette méthode tout au long de sa formation, vous présentera sa propre version et les adaptations qu'il en a tirées de son expérience.

Planifier

Checklist

Que faire AVANT d'ouvrir une IA pour étudier un sujet ?

- **Notez ce que vous savez déjà.** Prenez cinq minutes pour écrire tout ce dont vous vous souvenez sur le sujet. Peu importe si c'est incomplet, désordonné ou imparfait. L'effort de récupération en mémoire, active les connaissances déjà acquises et révèle les véritables lacunes, pas celles que vous imaginez avoir.
- **Identifiez honnêtement vos lacunes.** Une fois que vous avez noté ce que vous savez, repérez ce que vous ne savez pas. Il ne s'agit pas de ce que vous n'avez pas encore étudié, mais de ce que vous avez essayé de rappeler sans y parvenir. C'est la différence entre une lacune d'exposition (je n'ai jamais vu cette notion) et une lacune de compréhension (je l'ai étudiée, mais je ne la maîtrise pas réellement).
- **Définissez votre objectif d'apprentissage en une phrase.** Je veux comprendre le mécanisme d'action des inhibiteurs de l'enzyme de conversion n'est pas la même chose que : Je veux être capable d'expliquer pourquoi l'énalapril est préféré au captopril dans l'insuffisance cardiaque. Pour les étudiants en soins infirmiers : Je veux comprendre pourquoi une hypertension non contrôlée peut endommager les reins et comment surveiller ce risque chez mon patient. Plus votre objectif est précis, plus l'IA pourra vous être utile comme outil d'apprentissage.
- Avant de vous appuyer sur une réponse générée par l'IA, **posez-vous systématiquement ces trois questions de sécurité :**
 - ☐ Qu'est-ce que cette réponse ne me dit pas ?
 - ☐ D'où viennent ces informations, et d'où ne viennent-elles pas ?
 - ☐ Que se passerait-il si cette réponse était fausse ?

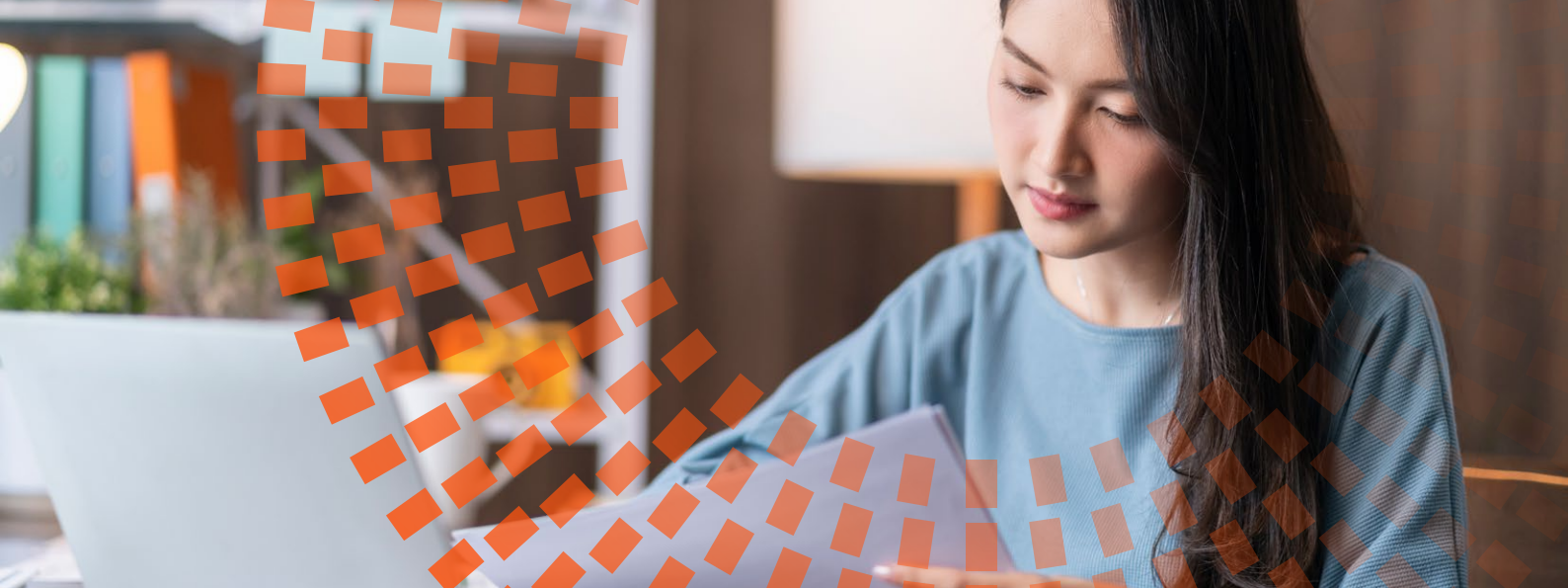
Prompts

Prompt 1 — Diagnostic des lacunes : Je vais étudier le sujet suivant : [sujet]. Avant de commencer, pose-moi 10 questions diagnostiques de niveaux de difficulté variés afin d'identifier ce que je maîtrise déjà et ce que je ne maîtrise pas encore. Ne me donne pas les réponses tant que je n'ai pas répondu à chaque question.

Prompt 2 — Cartographie des priorités : Je vais passer un examen sur [sujet/module]. Voici ce que je pense bien maîtriser : [liste]. Voici ce que je pense ne pas maîtriser suffisamment : [liste]. Quels autres domaines devrais-je prendre en compte ? Organise-les en fonction de leur impact clinique, et non de leur fréquence dans les examens.

Prompt 3 — Questions préalables à l'apprentissage (métacognition) : Avant de m'aider à travailler sur [sujet], j'ai besoin de prendre du recul. Pose-moi trois questions : une sur ce que je comprends déjà, une sur ce qui me semble encore confus, et une sur ce que je pense savoir mais qui pourrait être incorrect.





Apprendre

Checklist

Comment étudier AVEC l'IA sans lui laisser remplacer votre propre traitement de l'information ?

- ☐ **Produisez d'abord votre propre résumé.** Lisez le contenu. Rédigez ensuite votre résumé, même s'il est incomplet ou imparfait. Puis seulement, demandez à l'IA de générer un résumé sur le même sujet. Les écarts entre les deux constituent votre véritable carte d'apprentissage.
- ☐ **Créez vos fiches avant d'en commander à l'IA.** Rédigez au moins cinq fiches sur les concepts qui vous semblent les plus importants. Ensuite, demandez à l'IA d'en générer à son tour. Les fiches produites par l'IA auxquelles vous n'aviez pas pensé sont souvent celles que vous avez le plus besoin de travailler.
- ☐ **Utilisez l'IA comme un débutant, pas comme un oracle.** Au lieu de demander à l'IA de vous expliquer, expliquez-lui vous-même le sujet. Demandez-lui de vous poser des questions de clarification lorsque votre explication est incomplète ou confuse. C'est le principe du miroir cognitif : vos lacunes deviennent visibles.
- ☐ **Appliquez la méthode SIFT à toute réponse utilisée pour étudier :**
 - **S'arrêter.** Ne copiez pas et n'agissez pas immédiatement.
 - **Examiner les affirmations.** Vérifiez chaque information importante à l'aide d'une source indépendante.
 - **Rechercher une meilleure couverture.** Existe-t-il une autre source qui présente le sujet autrement ou qui contredit certains éléments ?
 - **Retracer le raisonnement.** Comment la réponse passe-t-elle du point A au point B ? L'étape est-elle réellement valide ou seulement plausible ?
- ☐ **Pour les scripts de maladie : Construisez votre propre script avant de consulter l'IA.** Ne demandez pas : Explique-moi tout sur [maladie]. Demandez plutôt à être évalué après avoir construit votre propre script.
- ☐ **Pour les soins infirmiers : élaborer votre plan de soins avant de consulter l'IA.** Rédigez votre évaluation, identifiez les diagnostics infirmiers prioritaires et planifiez vos interventions. Ensuite seulement, demandez à l'IA ce que vous avez pu oublier. Le principe reste le même : vous d'abord.

Prompts

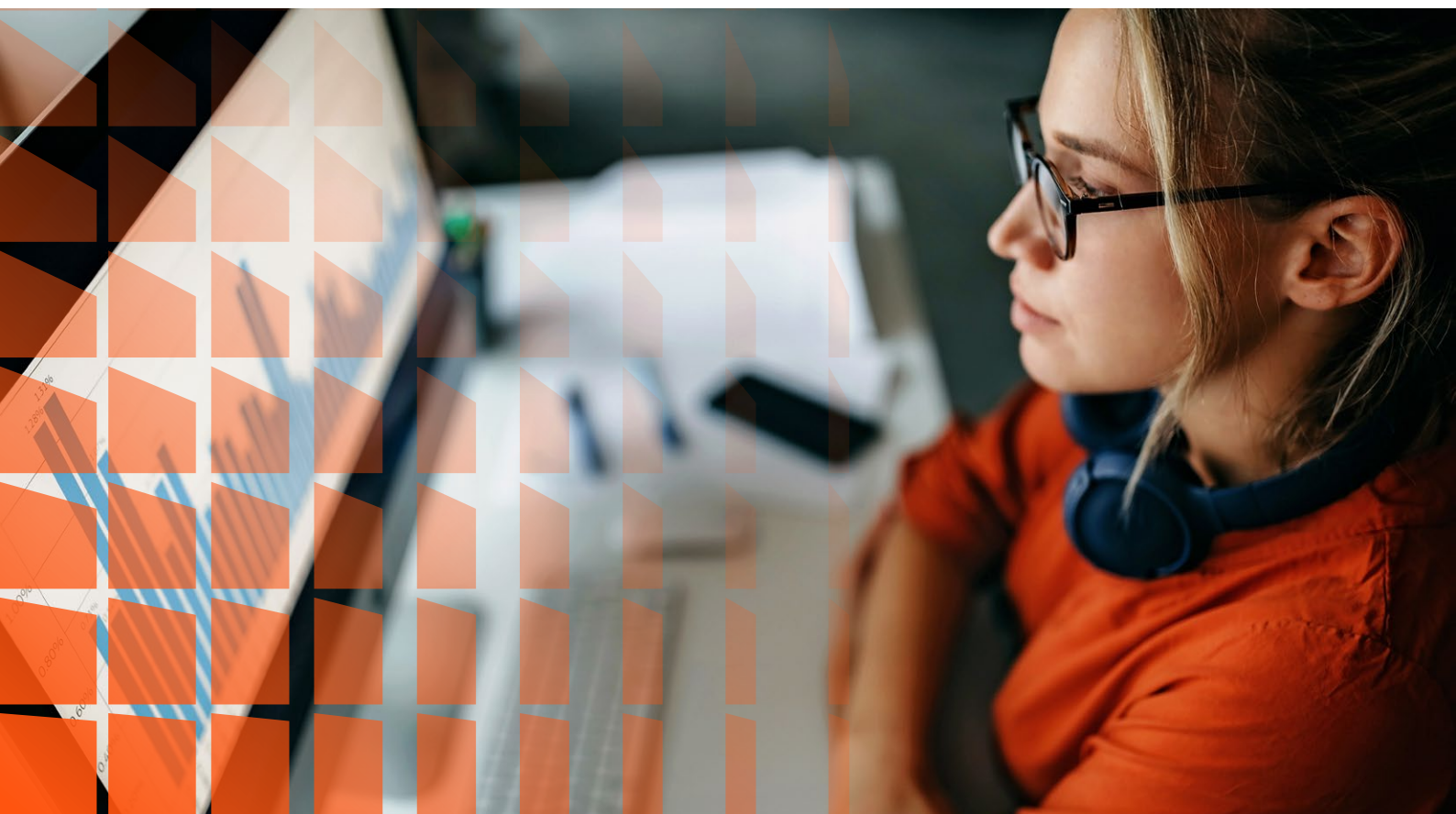
Prompt 4 — Miroir cognitif (enseigner pour apprendre) : Tu es un étudiant de première année. Je vais t'expliquer [sujet]. Pose-moi des questions de clarification lorsque mon explication est incomplète, ambiguë ou incorrecte. Ne me corrige pas directement ; contente-toi de me poser des questions.

Prompt 5 — Comparaison de résumés : Voici mon résumé de [sujet] : [votre résumé]. Génère ton propre résumé du même sujet. Puis indique : (1) ce que j'ai inclus et que tu n'as pas inclus, (2) ce que tu as inclus et que je n'ai pas inclus, et (3) les points sur lesquels nos explications diffèrent.

Prompt 6 — Évaluation d'un script de maladie : Voici mon script de maladie pour [maladie] : [votre version comprenant la physiopathologie, l'épidémiologie, les manifestations spécifiques, les diagnostics différentiels et l'évolution habituelle]. Évalue si des éléments essentiels manquent dans chacune de ces composantes. Ensuite, pose-moi trois questions afin de vérifier que je comprends réellement les liens entre elles.

Prompt 6b — Évaluation d'un plan de soins (soins infirmiers) : Voici mon plan de soins pour [type de patient] : [votre évaluation clinique, vos diagnostics infirmiers et les interventions prévues]. Analyse si certaines priorités importantes sont absentes, si les interventions proposées sont cohérentes avec les diagnostics infirmiers retenus et si des risques de sécurité n'ont pas été pris en compte. Ensuite, pose-moi trois questions sur mon raisonnement clinique.

Prompt 7 — Approfondissement socratique (passer au niveau supérieur) : Je maîtrise déjà les bases de [sujet]. Je souhaite maintenant atteindre un niveau d'analyse plus avancé. Ne me donne pas de nouvelles informations. Pose-moi des questions qui m'obligent à relier ce que je sais déjà à des situations cliniques concrètes. Utilise une approche socratique : une seule question à la fois.



Évaluer

Checklist

Comment vérifier que vous avez réellement compris, et pas simplement reconnu l'information ?

☐ **Essayez d'expliquer sans rien consulter.**

Avant de demander des QCM ou d'utiliser un outil, fermez tout et essayez d'expliquer le sujet à partir de zéro. Si vous n'y parvenez pas, vous savez exactement où se situe la lacune.

☐ **Produisez votre réponse avant de voir celle de l'IA.** Lorsque vous utilisez l'IA pour générer des questions d'entraînement, rédigez d'abord VOTRE réponse, même si elle est incorrecte. L'effort de récupération en mémoire est l'un des principaux moteurs de l'apprentissage.

☐ **Utilisez les QCM générés par l'IA, mais avec prudence.** Dans une revue de la littérature utilisant GPT-4, 49 % des QCM générés présentaient des défauts de construction selon les critères du NBME, et 22 % contenaient des erreurs factuelles (Kiyak & Kononowicz, 2024 ; n = 150 questions). Les modèles plus récents peuvent améliorer ces résultats, mais la vérification reste indispensable. (1)



☐ **Appliquez les trois niveaux de doute à votre propre apprentissage :**

- Les faits sont-ils corrects ? Vérifiez-les à l'aide d'une source indépendante.
- Le raisonnement est-il valide ? La logique qui relie les faits est-elle cohérente ?
- Ai-je posé la bonne question ? Suis-je en train d'évaluer ce qui est réellement important ?

☐ **Pour les cas cliniques : réfléchir d'abord, consulter ensuite.** Lisez le cas clinique. Fermez l'IA. Élaborez votre propre diagnostic différentiel. Ouvrez ensuite l'IA. Donnez-lui le cas et votre diagnostic différentiel. Demandez-lui de comparer les deux approches. Puis analysez les écarts : Qu'est-ce qui manquait dans votre raisonnement ? Pourquoi ne l'aviez-vous pas envisagé ?

Prompts

Prompt 8, QCM adaptatifs : Génère 10 questions à choix multiples sur [sujet] au niveau [débutant/intermédiaire/avancé].
Format : une question à la fois. Ne me montre ni la réponse ni l'explication avant que j'aie répondu. Après mes 10 réponses, réalise une analyse des schémas récurrents dans mes erreurs.

Prompt 9, Évaluation du raisonnement (pas seulement des faits) : Je vais te présenter mon explication de [sujet/mécanisme/cas clinique]. Ne me dis pas si elle est correcte ou incorrecte. Interroge plutôt mon raisonnement : quelles hypothèses suis-je en train de faire ? Quelle étape de mon raisonnement est la plus fragile ? Quelle alternative n'ai-je pas envisagée ?

Prompt 10, Diagnostic différentiel avec vérification (réfléchir d'abord) :
J'ai déjà établi un diagnostic différentiel préliminaire pour ce cas : [votre liste]. Quels diagnostics importants pourrais-je avoir omis compte tenu des données disponibles ? Explique pourquoi chaque ajout est pertinent. Ne réorganise pas ma liste ; complète-la uniquement.

Prompt 11, Signaux d'une véritable compréhension : Je viens de terminer l'étude de [sujet]. Pose-moi 5 questions auxquelles je ne peux répondre correctement que si je comprends réellement les mécanismes sous-jacents, et non si je me contente de mémoriser des informations. Si mes réponses restent superficielles, signale-le-moi.



Ajuster

Checklist

Comment modifier votre méthode pour le prochain cycle en fonction de ce que vous avez découvert ?

- ☐ **Analysez les schémas récurrents dans vos erreurs, pas seulement ce que vous avez raté.** Après une séance d'entraînement ou de révision, consacrez dix minutes à identifier la nature de vos erreurs : Ai-je échoué parce que je ne connaissais pas l'information ? Parce que mon raisonnement était insuffisant ? Parce que je suis parti d'une hypothèse erronée ? Le type d'erreur détermine la stratégie de correction.
- ☐ **Redéfinissez votre prochain cycle à partir de données, pas d'intuitions.** Ne répétez pas la même méthode en espérant obtenir un résultat différent. Si votre difficulté concerne l'application clinique, ne vous contentez pas de revoir des fiches : entraînez-vous sur des cas cliniques. Si votre difficulté concerne les connaissances de base, ne multipliez pas les exercices : revenez au contenu de référence.
- ☐ **Reconnaissez les signes de fatigue liée à l'utilisation de l'IA :**
 - Vous avez reformulé le même prompt plus de trois fois sans obtenir ce que vous cherchiez. Il est possible que vous ne sachiez plus exactement ce que vous recherchez.
 - Vous acceptez une réponse sans la lire entièrement parce qu'elle semble correcte . Il s'agit d'un signe classique de biais d'automatisation.
 - Vous vous sentez plus fatigué après l'interaction qu'avant de l'avoir commencée. Il est probablement temps de fermer l'outil.
- ☐ **Bouclez le cycle : Écrivez une phrase indiquant précisément ce que vous allez modifier.** La prochaine fois, je vais [action concrète], parce que j'ai découvert [constat précis]. Si vous n'êtes pas capable d'écrire cette phrase, c'est probablement que vous n'avez pas réellement terminé la phase d'ajustement.



Prompts

Prompt 12, Analyse des schémas

d'erreur : Voici les questions auxquelles j'ai répondu incorrectement : [liste avec vos réponses erronées]. Identifie les schémas récurrents : s'agit-il d'erreurs de connaissances (je ne connaissais pas l'information), d'erreurs de raisonnement (je connaissais les informations mais je les ai mal reliées entre elles) ou d'erreurs d'hypothèses (mon approche de départ était incorrecte) ? Indique ce que je devrais revoir en priorité.

Prompt 13, Refonte du cycle

d'apprentissage : Je viens de terminer un cycle d'étude sur [sujet]. Voici mes résultats : [résumé honnête de ce qui a bien fonctionné et de ce qui a moins bien fonctionné]. Voici la méthode que j'ai utilisée : [description de votre démarche]. Que modifierais-tu pour le prochain cycle ? Propose un maximum de trois changements concrets et non génériques.

Prompt 14, Réflexion métacognitive

de fin de séance : Analyse ma séance d'étude d'aujourd'hui sur [sujet]. Je vais décrire ce que j'ai fait et comment cela s'est déroulé : [description]. Évalue : (1) si ma méthode a favorisé une compréhension approfondie ou une simple reconnaissance de l'information, (2) à quels moments j'ai utilisé l'IA comme une extension de mes capacités plutôt que comme une béquille cognitive, et (3) quel changement précis je devrais mettre en œuvre lors de ma prochaine séance.





Point de vue d'un résident : Hussain Hilali

HUSSAIN HILALI – Témoignage d'un résident sur les bonnes pratiques d'utilisation de l'IA

Utiliser l'IA pendant ses études de médecine sans la laisser vous utiliser

[Consulter l'article ici](#)

La checklist rapide de Hussain :

1. Avant d'ouvrir l'IA : Que sais-je déjà ? Qu'est-ce qui me semble encore confus ?
(2 minutes maximum)
2. Première tentative (toujours la mienne) : Résumé, fiches de révision, diagnostic différentiel
Imparfait, c'est acceptable.
3. L'IA intervient ensuite : Comparer. Compléter. Remettre en question. Jamais remplacer.
4. Après avoir utilisé l'IA : Suis-je capable d'expliquer cela sans regarder l'écran ? Si la réponse est non, je n'ai pas terminé.
5. Boucler le cycle : Identifier une chose que je ferai différemment la prochaine fois.

Les conseils de Hussain pour rédiger des prompts :

- Soyez précis sur votre niveau de formation : Je suis étudiant en début de stage dans [spécialité] produira une réponse différente de "Je suis étudiant avancé".
- Demandez des questions plutôt que des réponses :
"Évalue-moi" est souvent plus utile que "Explique-moi".
- Demandez une chose à la fois. Les prompts comportant plusieurs demandes simultanées diluent souvent la qualité des réponses.
- Terminez toujours par : Qu'est-ce que j'ai oublié ?
Cette formulation pousse l'IA à compléter votre réflexion plutôt qu'à répéter ce que vous savez déjà.
- Si une réponse vous paraît trop parfaite, considérez cela comme un signal de vérification, pas comme une raison de lui faire immédiatement confiance.

-  **Maya :** C'est pratiquement ma méthode... mais sans les trois heures d'optimisation.
-  **Kai :** Justement. Si c'était compliqué, personne ne l'appliquerait.
-  **Maya :** Tu parles comme quelqu'un qui n'a jamais terminé une checklist de sa vie.
-  **Kai :** La checklist n'est pas le changement. L'intention, oui.
-  **Dr Flores :** La checklist que vous avez entre les mains est un point de départ, pas un point d'arrivée. Les prompts vont évoluer. Les outils vont changer. Les modèles vont s'améliorer et présenter de nouvelles limites. Ce qui ne changera pas, en revanche, c'est la question que tout cela vous a appris à vous poser.
-  **Kai :** Suis-je en train de réfléchir... ou de déléguer ?
-  **Dr Flores :** Exactement. Et vous n'avez même pas besoin de l'IA pour y répondre. Vous connaissez la réponse chaque fois que vous ouvrez l'outil et que vous choisissez soit de réfléchir d'abord, soit de laisser l'outil réfléchir à votre place.

Cette décision, répétée des centaines puis des milliers de fois au cours de votre formation, est ce qui distingue le professionnel qui utilise l'IA du professionnel qui en devient dépendant.

Le prochain chapitre nous emmène sur un autre terrain. Jusqu'à présent, nous avons parlé de la manière d'utiliser l'IA pour apprendre et raisonner. Vient maintenant une question que beaucoup préfèrent éviter : où se situent les limites à ne pas franchir ?

-  **Maya :** L'éthique.
-  **Dr Flores :** L'éthique. La sécurité. Le professionnalisme. Là où les conséquences ne sont plus un examen raté, mais la confiance d'un patient.





À retenir de ce chapitre

- La checklist est un système d'apprentissage, pas un système de productivité.
- Les prompts évolueront. Les outils évolueront. Ce qui ne changera pas, c'est la question fondamentale : suis-je en train de réfléchir ou de déléguer ?
- Cette décision, répétée chaque jour tout au long de votre formation, est ce qui distingue le professionnel qui utilise l'IA du professionnel qui dépend de l'IA.

Références

Ce chapitre synthétise les données probantes et les concepts présentés dans les chapitres 4 à 7. Il n'introduit pas de nouvelles données. Les principales références sont celles déjà citées dans les chapitres précédents :

1. Kiyak, Y. S., & Kononowicz, A. A. (2024). ChatGPT prompts for generating multiple-choice questions in medical education and evidence on their validity: A literature review. *Postgraduate Medical Journal*, 100(1189), 858–865. <https://doi.org/10.1093/postmj/qgae064> [GPT-4, n=150 items, NBME criteria. 49% construction defects, 22% factual errors. Newer models may improve figures; verification remains essential.]

Références croisées par section :

- **Planifier** : Cap 7 Beats 1-2 (MAL Planning) + Cap 4 Beat 6 (three security questions) + Cap 5 Beat 2 (metacognitive reflection) + Cutrer et al. 2020, *The Master Adaptive Learner*, Elsevier
- **Apprendre** : Cap 7 Beat 2 (you first in Learning) + Cap 5 Beat 3 (cognitive mirror) + Cap 6 Beat 2 (illness scripts) + Cap 5 Beat 4 (adapted SIFT)
- **Évaluer** : Cap 7 Beat 2 (you first in Assessing) + Cap 5 Beat 2 (three levels of doubt) + Cap 7 Beat 3 (desirable difficulties) + Cap 6 Beat 2 (think first)
- **Ajuster** : Cap 7 Beat 2 (you first in Adjusting) + Cap 7 Beat 5 (agency close) + Cap 5 Beat 4 (personal red flags)

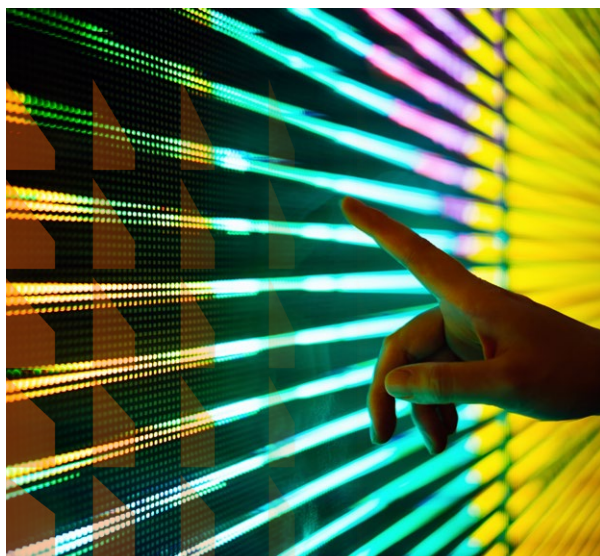
Prompts réalisés par : Huang 2024 (20 ChatGPT activities mapped to SRL) · Jackson 2025 (Higher Order Prompting) · Park & Choo 2025 (Socratic prompting) · Safranek et al. 2024 (ChatGPT Implementation Guide in Medical Education) · Chan & Zary 2025 (custom mnemonics with AI)

Chapitre 9 : Questions fréquentes, mythes et réalités



Dans ce chapitre, vous découvrirez...

- Quatorze questions fréquemment posées, organisées en trois grands thèmes : l'utilisation pratique de l'IA, l'apprentissage et la formation, ainsi que l'éthique et l'avenir de la profession.
- Des réponses fondées sur les données probantes présentées dans les chapitres précédents, et non sur des opinions.
- Sept idées reçues courantes sur l'IA en santé confrontées aux faits et aux données disponibles.
- Le point de vue de Hussain Hilali : Ce que j'aurais aimé savoir lorsque j'étais résident .



 **Dr Flores :** Nous avons parcouru beaucoup de chemin. Les limites de l'IA, la pensée critique, le raisonnement clinique, l'apprentissage autorégulé, la checklist, l'éthique... Maintenant, je voudrais que vous me posiez les questions auxquelles vous n'avez pas encore trouvé de réponse.

 **Kai :** Une sorte de permanence ouverte ?

 **Dr Flores :** Exactement. Une permanence ouverte. Sans filtre.

 **Kai :** Évidemment.

Bloc 1 : Utilisation pratique

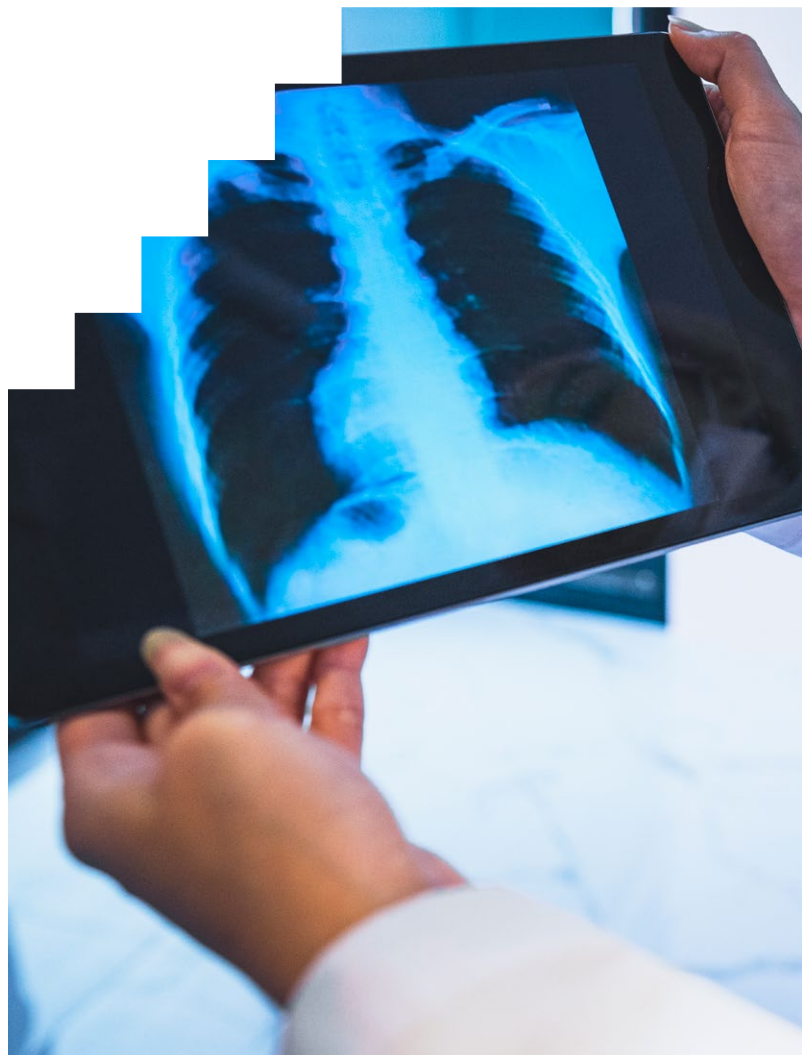
FAQ 1 : Utiliser l'IA, est-ce tricher ?

-  **Dr. Flores:** Dr Flores : Cela dépend de ce que vous en faites. Si vous ouvrez une IA sans préparation et lui demandez : “Rédige-moi un rapport de cas clinique”, le résultat ne vous appartient pas vraiment, même si vous en modifiez quelques passages. En revanche, si vous avez réalisé votre analyse, construit votre plan et utilisé l'IA uniquement pour améliorer la rédaction, alors la réflexion reste la vôtre.
-  **Maya :** J'ai vécu exactement cette situation. J'ai obtenu une excellente note. Mais je ne savais pas vraiment si ce rapport était le mien.
-  **Dr Flores :** La règle est simple : si vous n'êtes pas capable d'expliquer et de défendre chaque ligne comme si vous l'aviez rédigée vous-même, alors ce n'est pas vraiment votre travail. Et si vous avez utilisé l'IA, signalez-le. Aujourd'hui encore, la plupart des établissements ne disposent pas de règles claires : 63 % encouragent l'usage de l'IA générative, mais seulement 43 % fournissent des recommandations détaillées sur son utilisation. (1)

En résumé (Dr Flores) : Honnêtement, tout dépendra des règles et des politiques de votre établissement. Mon conseil est simple : respectez-les.

FAQ 2 : Puis-je faire confiance à l'IA pour établir un diagnostic ?

-  **Dr Flores :** L'IA recherche la réponse la plus probable. Vous recherchez la bonne réponse. Ce n'est pas la même chose. Dans une étude sur le triage, ChatGPT Health a sous-estimé l'urgence de 52 % des situations réellement urgentes, y compris des cas d'acidocétose diabétique. (2)
-  **Kai :** Donc vous lui décrivez les symptômes... et elle vous renvoie chez vous.
-  **Dr Flores :** Dans plus de la moitié des cas urgents étudiés, oui. L'IA peut amplifier des informations erronées au lieu de les corriger. Elle ne remplace pas le jugement clinique. Elle peut le compléter, à condition que vous vérifiez ses réponses.





FAQ 3 : Comment citer l'IA dans un travail académique ?

-  **Dr Flores :** L'ICMJE a mis à jour ses recommandations en 2024 : une IA ne peut pas être considérée comme auteur. L'être humain reste responsable de l'ensemble du contenu. Le comité COPE défend la même position : la transparence est obligatoire. (3)
-  **Maya :** Indiquez comment vous avez utilisé l'outil, à quelles étapes et avec quel système (ChatGPT, Claude, Manus, etc.). Ne dissimulez pas le processus. Vérifiez chaque élément comme si vous l'aviez rédigé vous-même.

FAQ 4 : Que faire si mon enseignant interdit l'utilisation de l'IA ?

-  **Kai :** Avant, cela m'aurait agacé. Maintenant, je vois les choses autrement. L'enseignant vous oblige à faire l'effort vous-même. Et c'est précisément ce dont votre cerveau a besoin.
-  **Dr Flores :** Respectez la politique de votre établissement. Les outils de détection présentent un taux élevé de faux positifs et pénalisent parfois davantage les étudiants dont l'anglais n'est pas la langue maternelle. Les travaux réalisés sans IA ne sont pas une punition. Ils constituent un entraînement cognitif délibéré. (1)

FAQ 5 : Quelle est la différence entre ChatGPT et un outil médical basé sur l'IA ?

-  **Dr Flores :** ChatGPT est un modèle généraliste entraîné sur de grandes quantités de textes issus d'Internet. Un outil médical d'IA est généralement entraîné sur des données cliniques, validé pour des usages spécifiques et soumis à des exigences réglementaires comparables à celles des dispositifs médicaux.

La règle générale est la suivante : utilisez les modèles généralistes pour apprendre, et les outils spécialisés pour soutenir la prise de décision. Dans les deux cas, vérifiez toujours les informations obtenues. (4)

Bloc 2 : Apprentissage et formation

FAQ 6 : L'IA va-t-elle remplacer les médecins ?

- ❏ **Kai :** L'IA ne s'assoit pas auprès d'un patient à trois heures du matin lorsqu'il a peur. Cela ne s'automatise pas.
- ❏ **Dr Flores :** L'IA ne remplacera pas les médecins. En revanche, elle remplacera probablement les médecins qui ne sauront pas travailler avec elle. C'est l'idée du médecin augmenté. Votre valeur réside dans le jugement clinique, la relation thérapeutique, la présence auprès du patient. Tout ce que l'IA ne possède pas. (5)

FAQ 7 : L'IA peut-elle m'aider à préparer mes examens ?

- ❏ **Maya :** Oui, mais à certaines conditions. L'IA peut générer des questions d'entraînement très utiles. Mais si vous ne faites pas l'effort de répondre avant de regarder la solution, vous entraînez votre reconnaissance des réponses, pas votre compréhension.
- ❏ **Dr Flores :** Dans une revue de la littérature, 49 % des QCM générés par GPT-4 présentaient des défauts de construction et 22 % contenaient des erreurs factuelles. Ils peuvent être utiles, mais doivent toujours être vérifiés à l'aide de vos sources de référence. (6)



FAQ 8 : L'IA peut-elle faire de moi un moins bon professionnel de santé ?

- ❏ **Dr Flores :** Oui. Nous disposons désormais de données directes à ce sujet. Kaminski et ses collègues ont étudié dix-neuf endoscopistes expérimentés. Après le retrait de l'assistance par IA, leur taux de détection des adénomes a diminué de 20 %, passant de 28,4 % à 22,4 %. (7)
- ❏ **Kai :** C'est comme un GPS. On l'utilise tellement qu'on finit par ne plus savoir s'orienter sans lui. Sauf qu'ici, il s'agit d'un adénome qui pourrait évoluer en cancer.
- ❏ **Dr Flores :** C'est ce que l'on appelle la perte de compétences (deskilling). Pour les étudiants, le risque est encore plus important : si une compétence n'est jamais développée sans l'aide de l'IA, elle ne sera jamais réellement acquise. Préserver la difficulté cognitive n'est pas une question de nostalgie. C'est une nécessité clinique.



FAQ 9 : L'IA peut-elle rédiger mes travaux à ma place ?

-  **Dr Flores :** Non. D'abord parce qu'elle peut produire des textes qui semblent exacts tout en contenant des erreurs factuelles. Ensuite parce que l'écriture scientifique n'est pas seulement un produit final : c'est un processus de réflexion. C'est en structurant vos arguments, en défendant vos idées et en confrontant les limites de vos preuves que vous apprenez réellement. (8)
-  **Maya :** Utilisez l'IA comme un éditeur, pas comme un auteur. Le plan, l'argumentation et la structure doivent venir de vous. L'IA peut aider à améliorer la forme, mais la réflexion doit rester la vôtre.
-  **Dr Flores :** Exactement. On peut parler de co-construction, mais c'est vous qui devez rester l'architecte du raisonnement.

FAQ 10 : Quelles compétences devrais-je développer que l'IA ne peut pas remplacer ?

-  **Dr Flores :** Les compétences les plus humaines : l'examen clinique, le raisonnement diagnostique délibéré, la communication empathique, la capacité à tolérer l'incertitude et l'évaluation critique des réponses générées par l'IA.

Votre avantage ne réside ni dans la vitesse ni dans la mémoire. Il réside dans le contexte : percevoir les signaux non verbaux, comprendre qu'un patient minimise ses symptômes parce qu'il est gêné de parler de sa santé mentale devant son conjoint, ou reconnaître qu'une inquiétude exprimée à demi-mot est parfois plus importante qu'un résultat biologique.

-  **Kai :** La médecine n'est pas une affaire d'informations. C'est la personne qui se trouve devant vous.

Bloc 3 : Éthique et avenir

FAQ 11 : L'IA présente-t-elle des biais à l'égard de certaines populations ?

Dr Flores : Oui. Non pas par intention, mais par conception. L'IA hérite des biais présents dans les données utilisées pour son entraînement et les reproduit avec une précision mathématique. Prenons un exemple : en dermatologie, un système d'IA atteignait une précision de 69 % sur les peaux claires, mais seulement de 17 % sur les peaux foncées, parce que l'ensemble de données utilisé pour son entraînement ne contenait qu'une dizaine d'images de peaux foncées sur cent mille. (9)

Maya : Utiliser l'IA sans connaître ces biais, c'est comme prescrire un médicament sans avoir lu ses effets indésirables. Sauf qu'un médicament est accompagné d'une notice. L'IA, elle, n'en a généralement pas.

FAQ 12 : Dois-je informer le patient que j'ai utilisé une IA ?

Dr Flores : La réponse évolue encore. En 2025, la Californie a adopté une loi imposant qu'une information soit fournie lorsque des communications destinées aux patients sont générées par une IA, sauf lorsqu'un médecin les a relues et validées. L'Europe évolue progressivement dans la même direction. (10)

Maya : Pour moi, l'important n'est pas de dire simplement : "J'ai utilisé une IA." L'important est de pouvoir expliquer : "Voici comment je suis arrivé à cette recommandation, et voici les éléments que j'ai pris en compte."

Dr Flores : C'est une forme de transparence plus mature.

FAQ 13 : Comment savoir si un outil d'IA est sûr pour les données des patients ?

Dr Flores : Si l'outil est public et gratuit (comme les versions ouvertes de ChatGPT, Gemini ou Claude) ne lui transmettez jamais de données de patients. Jamais. Dès que les données sont envoyées sur les serveurs de l'outil, elles échappent au contrôle de l'établissement de santé. Avec seulement trois informations : l'âge, le diagnostic et le code postal, par exemple, il est parfois possible d'identifier une personne. (11)

Maya : L'IA ne possède pas de bouton "confidentiel". Ce bouton, c'est vous.





Mythes et réalités

Mythe	Réalité
L'IA va remplacer tous les médecins.	Selon McKinsey (2025), 85 % des déploiements de l'IA dans le secteur de la santé visent à augmenter les capacités des professionnels, et non à les remplacer.
L'IA est objective.	Faux. L'IA hérite des biais présents dans les données utilisées pour son entraînement, y compris des inégalités documentées liées à l'origine ethnique, au genre et à d'autres facteurs socio-démographiques.
Il faut savoir programmer pour utiliser l'IA.	Faux. La maîtrise de l'IA en santé ne repose pas sur des compétences en programmation. Elle repose sur la capacité à comprendre, évaluer et utiliser de manière critique les réponses produites par l'IA.
Les hallucinations de l'IA sont rares.	Tout dépend du contexte. Certaines études rapportent un taux d'hallucination d'environ 1,5 %, tandis que d'autres ont observé des taux dépassant 25 % dans certaines tâches. En contexte clinique, même un taux d'erreur de 1 % peut être préoccupant.
Les patients peuvent utiliser des agents conversationnels pour obtenir des conseils médicaux fiables.	Les patients ne disposent pas toujours du contexte clinique nécessaire pour distinguer une réponse correcte d'une réponse simplement plausible. Une réponse convaincante n'est pas nécessairement une réponse juste.
Toutes les IA utilisées en santé se valent.	Faux. Les systèmes d'IA diffèrent par leurs données d'entraînement, leurs méthodes de validation et leur niveau de réglementation. Un modèle d'analyse d'images médicales n'est pas un agent conversationnel, qui n'est lui-même pas un système de prédiction clinique.

Conclusion

Maya : L'IA ne va pas me remplacer. Elle n'est pas objective. Je n'ai pas besoin de savoir programmer pour l'utiliser. Elle ne me fera pas gagner du temps dans toutes les situations. Et elle n'est pas fiable sans supervision humaine.

Kai : Mais ce n'est ni un jouet ni un simple effet de mode.

Maya : C'est un outil qui amplifie à la fois nos forces... et nos biais.

Kai : L'écart le plus dangereux n'est pas technologique.

Maya : Il est épistémologique. Croire que l'on comprend l'IA simplement parce qu'on l'utilise. Ou la craindre parce qu'on ne la connaît pas.

Kai : La maîtrise de l'IA n'est plus facultative.

Maya : Elle est devenue aussi fondamentale que la maîtrise de la médecine fondée sur les données probantes.

Dr Flores : Je n'ai rien à ajouter. Si ce n'est ceci : commençons.

Au début de ce guide, deux questions ont été posées. Maya demandait : "Suis-je réellement en train d'apprendre ou est-ce que je laisse la machine organiser ma pensée à ma place ?" Kai demandait : "Si l'IA peut le faire correctement, à quoi est-ce que je sers ?" Je ne leur ai pas donné de réponses définitives. Je vous ai donné des outils pour construire vos propres réponses. Chaque jour. À chaque prompt. À chaque décision clinique. Ces quatre grandes questions ne se résolvent pas une fois pour toutes. Elles se reposent continuellement.

Kai : Les questions ne changent pas. C'est nous qui changeons.



À retenir de ce chapitre

- Utiliser l'IA n'est pas tricher. En revanche, ne pas être capable d'expliquer ou de défendre chaque ligne comme si elle était la vôtre constitue un véritable problème.
- L'IA ne remplacera pas les médecins. Mais elle transformera profondément la pratique de ceux qui ne sauront pas l'utiliser de manière critique et réfléchie.
- L'écart le plus dangereux n'est pas technologique. Il est épistémologique : croire que l'on comprend l'IA simplement parce qu'on l'utilise.
- La maîtrise de l'IA est désormais une compétence aussi fondamentale que la maîtrise de la médecine fondée sur les données probantes.

Références

1. Perkins, M., et al. (2025). Generative AI and academic integrity in higher education: A systematic review and research agenda. *Information*, 16(4), 296. <https://doi.org/10.3390/info16040296>
2. Ramaswamy, A., et al. (2026). ChatGPT Health triage: 52% undertriage rate including DKA. *Nature Medicine*. [Cross-reference: Cap 4]
3. ICMJE. (2024). *Recommendations for the Conduct, Reporting, Editing, and Publication of Scholarly Work in Medical Journals*. + COPE. (2023). Authorship and AI Tools.
4. Busch, F., et al. (2024). Navigating the European Union Artificial Intelligence Act for healthcare. *npj Digital Medicine*, 7, 210. <https://doi.org/10.1038/s41746-024-01213-6>
5. Bouabida, K., et al. (2026). “The future of medicine will be profoundly human, precisely because intelligently augmented.” [Cross-reference: Cap 1 Beat 7, Cap 8]
6. Kiyak, Y. S., & Kononowicz, A. A. (2024). ChatGPT prompts for generating MCQs in medical education and evidence on their validity. *Postgraduate Medical Journal*, 100(1189), 858–865. <https://doi.org/10.1093/postmj/qgae064>
7. Kaminski, M. F., et al. (2025). Colonoscopy adenoma detection rate drops from 28.4% to 22.4% after AI withdrawal. *The Lancet Gastroenterology & Hepatology*. [ELSEVIER] [Cross-reference: Cap 4]
8. Currie, G. M. (2023). Academic integrity and AI. *Seminars in Nuclear Medicine*, 53(3), 468–474. <https://doi.org/10.1053/j.semnuclmed.2023.04.008>
9. Daneshjou, R., et al. (2022). Disparities in dermatology AI. *Science Advances*. + Zack, T., et al. (2024). GPT-4 racial/gender biases in healthcare. *The Lancet Digital Health*, 6(1), e12–e22. [https://doi.org/10.1016/S2589-7500\(23\)00225-X](https://doi.org/10.1016/S2589-7500(23)00225-X)
10. California AB 3030. (2025). + EU AI Act (2024).
11. Li, J. (2023). Security implications of AI chatbots in health care. *Journal of Medical Internet Research*, 25, e47551. <https://doi.org/10.2196/47551> + USC Price School (2025).