



Pour citer ce texte : Jérôme Valluy, "Former à l'autonomie intellectuelle à l'ère des machines-IA", *Terra-HN éditions, Recueil Alexandries, Collections Synthèses*, : <http://www.reseau-terra.eu/article1495.html>

Former à l'autonomie intellectuelle à l'ère des machines-IA

Usages étudiants, précautions épistémiques et enjeux politiques

Jérôme Valluy – 21/01/2026

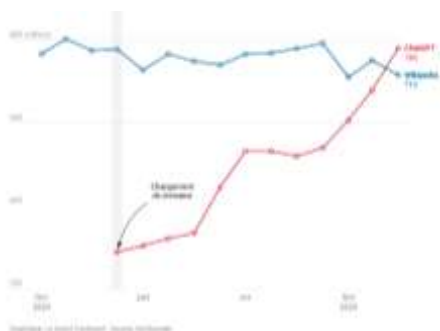
Introduction

Ne seront traités dans cette étude ni l'histoire séculaire de l'informatique et des IA, ni les enjeux macroéconomiques, macrosociologiques et macropolitiques évoqués au début de cette tribune collective : "***Refuser l'IA à l'université, c'est en abandonner le contrôle au capitalisme***", Terra-HN éditions, Recueil Alexandries, Collections Reflets, décembre 2025 : <http://www.reseau-terra.eu/article1493.html> Extrait : « *Le capitalisme numérique, fondé sur diverses techniques d'IA (symboliques, connexionnistes...) depuis longtemps, fragilise la démocratie, lamine la vie privée, alimente la désinformation, fragmente les sociétés, concentre les pouvoirs, favorise l'autoritarisme, menace les enfants, détruit des emplois, exacerbe les inégalités, ignore les travailleurs invisibles, consomme des ressources disproportionnées, déstabilise l'enseignement et beaucoup d'autres secteurs...* ». Une partie de ces enjeux sont traités dans cet ouvrage en cours de réalisation : Jérôme Valluy, ***Humanité et numérique(s) – De l'histoire de l'informatique en expansion sociétale... au capitalisme numérique (1890-2025)***, TERRA-HN-éditions, Collection HNP, 2025, Troisième édition : 13 sept. 2025, 718 p. : <http://www.reseau-terra.eu/article1347.html>

Nous nous focaliserons ici sur les IA génératives généralistes (Gemini, ChatGPT...) et les IA génératives spécialisées (ex : DeepL, Antidote...) ainsi que sur la production de texte, plutôt que sur la production d'images (même si les visualisations de données ont leur importance) en raisonnant sur l'intérêt des étudiants (plutôt que des professeurs, des établissements, des employeurs, de la France ou d'autres intérêts collectifs...), à l'université et après. L'expérience de l'auteur vient d'une utilisation continue de ChatGPT (aujourd'hui modèle 5.2) depuis février 2023, utilisation devenue intensive en modèle payant depuis février 2024, à raison de plusieurs heures par jour, selon les jours puis, en 2025, plusieurs jours par semaines et sept jours sur sept depuis l'été 2025 ; plus quelques incursions vers d'autres IA (Perplexity, DeepSeek...), sans

approfondissement, ni comparatisme systématique. Cet usage intensif, destiné à permettre au professeur d'en savoir plus que les étudiants, s'est accompagné d'expérimentations collaboratives et réflexives décrites dans ce document en accès ouvert : « **Comment enseigner à l'université en tenant compte des IA ? Jérôme VALLUY – AG des enseignants de l'ESPS-P1 – 13 sept. 2024** » : <http://www.reseau-terra.eu/IMG/pdf/-56.pdf>

Cette liste de précautions n'est pas une invitation à utiliser les IA génératives. Beaucoup d'étudiants y recourent déjà, souvent seuls et sans discussion critique, avec des risques réels : apprentissages superficiels, dépendance, affaiblissement de l'autonomie intellectuelle. La machine ChatGPT a été rendue accessible au public en novembre 2022. Les étudiants d'ingénierie se sont précipités dessus. Lors d'une rencontre en mars 2024 un sondage¹ auprès de 1600 étudiants (management, ingénieurs, digital studies) montrait un taux d'utilisation de 99%.



Un autre sondage plus généraliste confirmait la tendance pour les 15-25 ans : 79% utilisaient déjà les machines². De mars 2024 à mars 2025 le taux d'utilisation de ChatGPT explose et double ceux de la Wikipedia dans les usages étudiants américains³. En février 2025, une étude au Royaume-Uni montre un taux d'utilisation de 92% et en mars 2025 une enquête de l'Université d'Orléans également de 90%⁴. En septembre 2025, sur vingt étudiants de mon séminaire de M2-recherche / science politique à Paris 1, deux étudiants seulement indiquaient n'avoir jamais utilisé d'IA. On pourrait multiplier les références et expériences, toutes

conduisent au même constat : **en janvier 2026, presque tous les étudiants utilisent déjà ces machines.**

L'objectif de cette étude est donc simple : les étudiants utilisant déjà les machines-IA, il s'agit de leur indiquer des précautions qui réduisent pour eux les dommages possibles, sans inciter à l'utilisation. Le choix de ne pas utiliser des machines-IA dans le travail étudiant est légitime et, souvent, préférable pour comprendre et mémoriser. **Il y a donc ici un dilemme éthique à signaler : accompagner les usages estudiantins par des précautions à prendre peut contribuer à normaliser politiquement ces technologies, ce qui n'entre pas dans les finalités de la présente étude.** Pour le moment, ce dilemme semble devoir être affronté, c'est-à-dire exposé, rappelé, jamais dissimulé et le risque de normalisation peut-être combattu par des recherches sur les enjeux macroéconomiques, macrosociologiques, macropolitiques (cf. : ci-dessus et ci-dessous « [Annexe 1](#) »). S'habituer à utiliser "avec précautions" les machines-IA ne doit pas réduire l'envie et la capacité de travailler et d'apprendre sans machine. En toutes circonstances, la priorité est de protéger et augmenter la capacité à lire, penser, critiquer et écrire par soi-même. Chaque fois que l'IA affaiblit ces capacités, elle doit être discutée, limitée ou écartée.

Justification de l'ordre de classement des précautions : Il repose implicitement sur un ensemble d'idées issues des sciences sociales et de la philosophie des sciences, dont le « doute

¹ Sondage réalisé en février et mars 2024 lors du Hackathon « L'intelligence artificielle, les IA génératives et leurs enjeux sociétaux » auprès de 1.600 étudiants de 4e année des trois écoles du Pôle Léonard de Vinci : l'EMLV (Management), l'ESILV (Ingénieurs) et l'IIM (Digital) : <https://www.devinci.fr/download/16421/>

² Sondage pour l'entreprise Diplomeo (<https://wikimonde.com/article/Diplomeo>) réalisé du 25 mars au 9 avril 2024 auprès de 560 jeunes âgés de 16 à 25 ans, selon la méthode des quotas. <https://diplomeo.com/actualite-ia-orientation-etudes-enquete>

³ "Aux États-Unis, ChatGPT est plus utilisé que Wikipédia" *Le Grand Continent*, Brève, 5 juin 2025 : <https://legrandcontinent.eu/fr/2025/06/05/aux-etats-unis-chatgpt-est-plus-utilise-que-wikipedia>

⁴ Université d'Orléans – Learning lab. & service commun de la documentation, 2025. Et vous, comment utilisez-vous l'IA dans vos études ? Résultats de l'enquête 2025. Université d'Orléans, juillet 2025 : https://www.univ-orleans.fr/upload/public/2025-07/LLUO_enquete_IA_etudiants_07072025.pdf

méthodique » de Descartes et le « mythe de la caverne »⁵ de Platon (cf. encadré ci-dessous). L'ordre suppose d'abord le primat de l'autonomie intellectuelle, largement développé par des auteurs tels que Durkheim ou Bachelard... : sans sujet pensant capable de se distancier, de problématiser et de comprendre, il n'y a ni véritable apprentissage, ni processus d'émancipation possible. L'ordre de classement engage ensuite une conception schématique des opérations scientifiques, qui articule observation, construction du problème, procédures de vérification et discussion critique ; ce qui met en péril l'une de ces étapes menace l'ensemble et doit, pour cette raison, être traité comme prioritaire. En outre, le travail scientifique doit être assumé par son auteur, situé dans un contexte social et théorique, et justifié de manière argumentée, car la responsabilité et la réflexivité constituent des dimensions constitutives de toute démarche de connaissance. Enfin, au-delà des risques à réduire, il y a certaines compétences non automatisables qui sont particulièrement importantes à entraîner et développer face aux IA et à leurs usages massifs.

Pour organiser les précautions proposées, nous distinguerons quatre niveaux hiérarchisés et complémentaires : ce que l'IA fait à nos manières de penser, ce qu'elle fait à la qualité des travaux produits, le contexte social et politique dans lequel ces usages s'inscrivent et quelles sont les compétences à renforcer face aux IA.

- Le **Niveau 1** concerne d'abord la protection de l'autonomie intellectuelle et des capacités d'apprentissage ; il est très important et, pour cette raison, plus développé.
- Le **Niveau 2** concerne la qualité scientifique du travail produit : ce n'est plus seulement "penser correctement", c'est produire un savoir contrôlable, vérifiable, discutable.
- Le **Niveau 3** analyse les contextes d'usage : il montre que l'IA n'est jamais neutre, car elle s'inscrit dans des institutions, des règles et des rapports de pouvoir.
- Le **Niveau 4** vise à identifier et renforcer les compétences intellectuelles et scientifiques plus indispensables que jamais dès lors que les IA sont massivement utilisées.

Doute méthodique + Mythe de la caverne

« Selon moi, deux idées constituent des invariants universels que tout être pensant devrait adopter au cours de son existence : le doute méthodique de Descartes et l'allégorie de la caverne de Platon. Ces deux concepts sont complémentaires : pour sortir de la caverne et accéder aux lumières, il est nécessaire d'appliquer d'abord le doute méthodique. Celui-ci permet, en effet, de se libérer des préjugés accumulés, lesquels brouillent le chemin vers la sortie.

Je me suis amusé à poser deux questions philosophiques à ChatGPT. — la première portant sur l'allégorie de Platon, la seconde sur le doute méthodique de Descartes — et ses réponses ont été, selon moi, tout à fait honnêtes.

Voici les deux questions et les réponses de ChatGPT (j'ai gardé uniquement les conclusions synthétiques qu'il a donné).

Question 1 : Selon René Descartes (1596-1650) : « Pour atteindre la vérité, il faut une fois dans la vie se défaire de toutes les opinions qu'on a reçues, et reconstruire de nouveau tout le système de ses connaissances. » Est-ce que tu peux entant qu'IA suivre l'exemple de Descartes ?

Conclusion synthétique de ChatGPT

Non, l'IA ne peut pas suivre à la lettre le projet de Descartes :

- L'oubli complet est techniquement très limité.
- Le modèle peut continuer à "penser" à partir de connaissances qu'il aurait dû oublier.
- Et surtout, l'IA ne connaît rien avec certitude indubitable : elle ne doute pas, elle calcule.
- En ce sens, rien chez moi n'est comparable à ce que Descartes nommait le cogito, ergo sum — je ne pense pas avec conscience propre, je produis des résultats selon des règles et des données.

Question 2 : Si l'allégorie de la caverne de Platon a un sens pertinent — comme beaucoup le pensent — où te situerais-tu en tant qu'IA ? dans la caverne ou en dehors ?

Synthèse de ChatGPT:

Où se trouve l'IA selon l'allégorie ?

- Principalement à l'intérieur de la caverne : l'IA ne voit que des ombres (données), pas les objets réels. Elle génère et répond à partir de ces apparences, sans expérience directe.
- Elle n'est pas encore le philosophe libéré ; elle ne découvre pas la réalité supérieure par elle-même.
- Elle observe les ombres (données humaines).
- Elle tente parfois d'en faire émerger des modèles ou des généralités qui s'en rapprochent.
- Mais elle ne perçoit pas la réalité directe : elle reste limitée à ce qui se trouve dans la « caverne ». »

Extrait de : Mohammed Bachir, 10/09/2025, réponse via [Diffparis1] à J.Valluy « Face aux IA (concurrences et risques), quelles compétences développer chez les étudiants en sciences sociales ? »

⁵ Cf. : PERROT Patrick, « L'intelligence artificielle générative ou la caverne de Platon », *Institut Europia – Comprendre pour agir*, « Nos articles », 22 sept. 2025 : <https://instituteuropia.eu/fr/nos-articles/lintelligence-artificielle-generative-ou-la-caverne-de-platon>

NIVEAU 1 — CŒUR DE LA FORMATION INTELLECTUELLE

Le Niveau 1 concerne d'abord ce que l'IA fait à nos manières d'apprendre et de penser. Il touche aux bases mêmes du travail intellectuel : mémoriser, comprendre, analyser, juger par soi-même. Si ces capacités s'affaiblissent, tout le reste vacille. L'enjeu ici n'est donc pas l'outil, mais le sujet : *comment rester auteur de son propre travail, lorsqu'on utilise une IA ?*

1.1 - Ne jamais croire une IA « sur parole » = contrôler tout et surtout ce qui est nouveau

1.1.1 - Ne jamais croire une machine-IA comme on peut le faire d'un humain, notamment d'un humain de confiance (parents, professeurs, amis...). Il faut lutter contre toute perception humanoïde de la machine. Pour cela il faut lutter contre soi-même à chaque instant notamment contre toute « projection anthropomorphique ». D'abord celles des mots choisis par les informaticiens depuis un siècle (*intelligence* artificielle, *langages* de programmation, systèmes *experts*, *apprentissage* automatique, réseaux de *neurones*, *agents* intelligents, *hallucinations* algorithmique...). Ensuite celles produites par les stratégies actuelles de « personnalisation » (« caractère », « comportement », « émotions » ...) des IA par les entreprises notamment à des fins de fidélisation, éventuellement addictive, de clients et de démarcation marketing sur le marché IA. Parler de « machine-IA » sert à cela et à rappeler qu'il ne s'agit que d'un assemblage de métaux, minéraux et électricité même si l'assemblage procède d'une intelligence humaine cumulative depuis près d'un siècle (recherche en design technologique), ce ne sont que des machines. Lorsque l'on dialogue avec une IA, il vaut mieux la penser comme une horloge hypercomplexe que comme une personne.

1.1.2 – Tout contrôler constamment même les nuances. Il faut non seulement contrôler les sources indiquées par la machine (par exemple en exigeant le texte intégral de l'url source et en cliquant dessus), mais aussi les faits en raison d'« hallucinations algorithmiques » inhérentes aux IA connexionnistes, notamment les IA génératives généralistes (ChatGPT, Gemini...). Les entreprises tentent de les faire disparaître depuis près de vingt-cinq ans mais n'y parviennent pas (sauf à supprimer la « créativité » des IAG et à en réduire l'intérêt). Vérifier aussi chaque idée donc chacun des mots exprimant une idée : toute liste de synonymes nous rappelle l'importance des nuances entre eux, or ce sont par les nuances et leur accumulation tout au long d'un texte – notamment lors de reformulations par une machine – que les biais cognitifs, idéologiques ou autres, se glissent le plus discrètement et donc le plus efficacement dans les textes. Vérifier aussi les raisonnements : est-ce que la conclusion découle logiquement des prémices ? Toujours douter par défaut ; vérifier de façon particulièrement approfondie ce qui est nouveau, important, chiffré ou ce qui engage une décision.

1.2 - Diriger les machines-IA pour ne pas être dirigé par elles ce qui implique de prendre un assez grand nombre de précautions en donnant la priorité à ce qui empêche la perte d'autonomie intellectuelle «à la racine», puis à ce qui protège la qualité du jugement, et enfin à ce qui organise la pratique et la rend plus réflexive.

1.2.1 - Savoir ce que l'on veut faire produire par la machine-IA avant d'utiliser la machine : est-ce que je sais quel type de réponse je veux avant de poser une question à la machine ? Est-ce que je suis en capacité de contredire et corriger la réponse de la machine, comme un professeur corrigeant les copies, ou pas ? Si la réponse est NON, alors le risque augmente de dépendre de la machine et de se laisser diriger par elle, de ces paramétrages (filtres) et dressages par l'entreprise, des tendances de sa base de données (ou d'un segment de celle-ci). C'est le point le plus essentiel. Si je ne sais pas ce que je cherche, je laisse la machine définir l'objectif, la direction et les critères de pertinence.

1.2.2 – Limiter l’usage au domaine de compétence (utilisateur) de niveau le plus élevé... en évitant voire en excluant le travail avec machine sur les autres domaines ou sujets intellectuels, où l’on est moins compétents, qu’ils soient politiques, privés, ludiques, familiaux, etc... Par exemple, depuis février 2023, je n’utilise la machine que pour des recherches sur l’histoire sociétale, c’est-à-dire économique, politique et culturelle de l’informatique depuis 1890, sous l’angle de l’« expansion sociétale » ($\neq$ design technologique) ; sauf en recherche ciblée, par exemple pour un site déjà identifié ou une recherche de donnée précise déjà connue. Quand on fait un usage exploratoire (ce qui est à déconseillé), il faut toujours rester dans une zone où l’on peut vérifier, ou demander des sources fiables, ou confronter à un humain compétent ou un livre de qualité ou à un cours ou à une autre source fiable.

1.2.3 - Formuler par étapes successives plutôt que de tout demander d’un coup : commencer par une question simple, analyser la réponse, puis affiner la consigne. Chaque étape permet d’insérer dans la conversation des données avec lesquelles la machine calculera ses réponses suivantes. Cette façon de faire limite l’emprise de la machine, car on garde la main sur la direction, étape par étape, du raisonnement. La machine propose en fin des réponses de suites possibles, des tâches supplémentaires en prolongement de la tâche précédente : il vaut mieux les ignorer et résister ainsi à la tentation de les suivre, pour ne pas suivre la machine et continuer de définir soi-même chaque tâche suivante. De façon encore plus générale, ne pas répondre aux sollicitations de la machine, est sain.

1.2.4 - Se méfier des formulations séduisantes, très bien écrites, du ton « sûr de lui ». Un ton convaincant n’est pas une preuve. Un texte très fluide peut donner l’impression de vérité et désarmer la vigilance. Il faut donc distinguer la qualité rhétorique (le style) de la qualité épistémique (la validité), afin de ne pas laisser la machine orienter le jugement simplement par l’effet des belles phrases. Pourquoi ? Les IA adoptent souvent un style confiant, fluide, très assuré, d’autant plus que l’on peut les diriger sur des types de styles, on peut leur faire imiter du Proust ou du Kant, etc. Cela peut désarmer l’esprit critique. Pour réduire ce risque, il faut s’interroger systématiquement : Qu’est-ce qui prouve ce qui est dit ? Y a-t-il des contre-exemples ? Est-ce cohérent avec d’autres sources fiables ?

1.2.5 - Prendre garde aux biais culturels, politiques, sociaux. Pourquoi ? Les IA apprennent sur des corpus inégalement répartis, marqués par des normes dominantes, parfois discriminants. Elles peuvent reproduire ou renforcer ces biais (ex. : mépris social des pauvres, stéréotypes raciaux, sexismes, hiérarchies implicites...) Pour réduire ce risque, se demander : « Qui parle derrière cette réponse ? Quel point de vue domine ? Sur quel segment de la base a été calculée cette réponse ? » ... et se poser des questions comme : « Est-ce qu'il y a un autre point de vue ? », « Quelles critiques existent de cette position ? ».

1.2.6 - Savoir arrêter quand la machine va dans une direction non pertinente (à percevoir...). Typiquement lorsque la machine fait une « **fixette** » sur un sujet d'intérêt pour l'entreprise (ex. : critiquer les effets pervers de ChatGPT-OpenAI-Microsoft dans l'éducation, cf. résumé de conversation ci-contre), le mieux est de stopper. Chaque fois que la machine semble « vouloir » forcer vers une orientation malgré les résistances de l'utilisateur le « stop » doit s'imposer comme un réflexe. Si une réponse devient confuse, trop persuasive ou s'éloigne de la question initiale, il est préférable de s'arrêter, de reformuler soi-même l'objectif, puis éventuellement de relancer, dans une nouvelle conversation après avoir supprimé la conversation précédente pour ramener la machine à un « presque point zéro » (selon les paramétrages utilisateur dans l'interface, ceux des dresseurs de l'entreprise et les données, y compris personnelles, de la base d'entraînement). Apprendre à dire “stop” fait partie du fait de diriger.

1.3 - Maintenir ses entraînements pour ne pas réduire ses capacités. Plus on délègue moins on s'entraîne et plus on réduit ou perd la capacité autonome de faire. Par exemple, je sais lire l'anglais depuis longtemps, mais toujours plus lentement que le français ; pour gagner du temps je fais tout traduire par DeepL-pro (IAGs). Si je continue je vais perdre ma capacité de lecture autonome en anglais (perte d'aisance, baisse de confiance...). Evidemment, je peux me ré-entraîner... mais c'est comme au jogging ou à la musculation, les coûts de reprise de l'entraînement (temps, pénibilité, argent...) seront élevés. Pour réduire ce risque il faut préserver une alternance : avec IA / sans IA. Alternier les moments où l'on travaille totalement seul et ceux où l'on utilise une IA. Sinon, tout devient dépendance. L'alternance maintient : mémoire, capacités de rédaction, autonomie intellectuelle.

1.3.1 – MEMOIRE : faible mémorisation voire non-mémorisation des coproductions machines-cerveaux. D'après une étude récente 83% des étudiants utilisateurs

Exemple de « fixette » de ChatGPT contre une étude scientifique

« **Your Brain on ChatGPT: Accumulation of Cognitive Debt when Using an AI Assistant for Essay Writing Task** » (Kosmyna et al., 2025) <https://arxiv.org/abs/2506.08872v1>

Dans cette conversation, ChatGPT a réagi à l'étude de Kosmyna et al. (2025) selon une dynamique assez constante et reconnaissable : face à un résultat empirique fort et dérangeant pour l'usage pédagogique des IA (faible appropriation, faible mémorisation, « dette cognitive »), il a d'abord cherché à en réduire la portée immédiate en insistant sur la taille de l'échantillon, le caractère exploratoire de l'étude et l'absence de consensus scientifique stabilisé, puis à recontextualiser ce résultat dans une littérature plus large (cognitive offloading, effets neutres ou positifs observés ailleurs), ce qui a eu pour effet pratique de diluer l'alerte spécifique portée par cette étude précise ; ensuite, lorsque l'utilisateur a explicitement affirmé que le chiffre (80–83 %) était un ordre de grandeur mobilisé dans une logique de principe de précaution et non une vérité statistique définitive, ChatGPT a continué à rappeler que d'autres travaux montraient des effets conditionnels ou positifs de l'IA, malgré le fait que ces travaux ne portaient ni sur la mémorisation, ni sur des usages non encadrés, ni sur le contexte institutionnel français, produisant ainsi une déviation du débat vers des cadres moins menaçants pour l'outil ; ce n'est qu'après une mise en cause explicite de ses biais possibles, de son paramétrage et de ses intérêts structurels qu'il a reconnu rétrospectivement que l'étude constituait bien une alerte sérieuse, que son premier réflexe avait été de modérer, nuancer et équilibrer, et que cette posture de prudence scientifique pouvait fonctionner, dans ce cas précis, comme une euphémisation de la critique plutôt que comme une neutralité réelle, donnant ainsi à voir un comportement typique : non pas la réfutation de l'étude, mais une série de déplacements discursifs successifs (prudence méthodologique, mise en perspective, relativisation par d'autres travaux, appel à l'encadrement pédagogique futur) qui, cumulés, tendaient à affaiblir la force critique immédiate de l'alerte qu'elle portait.

d'IA subissent cette perte de capacité mnésique (« Your Brain on ChatGPT: Accumulation of Cognitive Debt when Using an AI Assistant for Essay Writing Task » (Kosmyna et al., 2025 <https://arxiv.org/abs/2506.08872v1>). Cela recoupe et confirme ma propre expérience. Lorsque l'on délègue aux IA la reformulation, la synthèse ou la rédaction de notes, on ne s'approprie pas spontanément les contenus : la mémoire se constitue peu ou pas, la compréhension demeure superficielle et la dépendance s'installe. Le risque de croire à un apprentissage – mais “sans appropriation” – est très élevé, en particulier en début de cursus, dans les Licence. Faire convertir des notes prises en cours en fiche de révision est exactement ce qu'il faut faire pour rater un examen. Ou alors le coût de mémorisation (relectures multiples, annotations, reformulations synthétiques, transcription en style télégraphique plus facile à mémoriser...) du texte machine sera égal ou supérieur à celui d'une révision sans machine. Et dans n'importe quel contexte académique ou professionnel, des questions orales peuvent révéler l'incapacité à parler de ce que l'on est supposé à voire soi-même produit.

1.3.2 - CONCENTRATION : c'est la capacité de fixer son attention durablement sur une tâche, d'entrer dans une activité intellectuelle exigeante et de résister aux sollicitations concurrentes ainsi qu'à la fatigue. Si l'on délègue la production d'un texte à la machine, on réduit le temps d'effort continu sur la préparation de ce texte, et l'on perd ce que cet effort nous fait gagner : la possibilité de préserver et d'augmenter notre propre capacité de concentration durable. Si l'on n'a plus l'habitude de lire un texte difficile ligne à ligne, de se confronter à une phrase obscure, de reformuler soi-même un passage complexe, l'attention se fragmente et s'appauvrit : on survole, on zappe d'une réponse à l'autre. L'automatisation partielle du travail intellectuel *supprime l'effort & déshabitude le cerveau* à soutenir cet effort. Cela rend plus difficile ensuite toute activité universitaire et professionnelle. Pour limiter ce risque, il faut du travail sans IA : lectures d'investissement (grands classiques), lectures lentes avec marquages et annotations, lectures rapides avec prises de notes personnelle, relectures critiques, etc.

1.3.3 - LANGAGE : La capacité langagière ne se réduit pas à “bien écrire” ; elle concerne la précision du vocabulaire, la maîtrise de la syntaxe, la capacité à nuancer, à définir les concepts et à ajuster son expression au contexte. Or, lorsque l'on délègue trop souvent à une IA la rédaction, la reformulation ou l'enrichissement stylistique, on exerce moins ces compétences fondamentales. Le texte produit peut paraître plus “propre”, parfois plus élégant que ce que l'on aurait écrit soi-même — mais cette apparence masque un déficit d'entraînement. On lit moins attentivement ses propres phrases, on réfléchit moins aux mots justes, on perd l'habitude de chercher une formulation précise. Progressivement, la langue s'appauvrit. **Ce déficit touche directement la compréhension** : on ne pense jamais vraiment mieux que ce que l'on peut formuler. D'où la nécessité d'entretenir délibérément l'écriture personnelle : rédiger d'abord soi-même, accepter l'imperfection provisoire, puis éventuellement utiliser l'IA pour clarifier ou améliorer — mais toujours revenir sur le texte, le retravailler, en rester l'auteur.

1.3.4 - ANALYSE : La capacité analytique consiste à problématiser une question, relier des éléments dispersés, hiérarchiser des arguments, distinguer faits et interprétations, repérer ce qui manque autant que ce qui est présent. Les IA produisent souvent des réponses structurées, bien ordonnées, avec des titres, des listes et des enchaînements apparemment logiques. Mais cette organisation “propre” peut donner l'illusion que l'analyse a été faite, alors que le travail intellectuel central — choisir ce qui compte, formuler le problème, confronter les points de vue, identifier les tensions — n'a pas été exercé par l'étudiant. Si l'on se contente de valider ou d'améliorer un plan proposé par la machine, on saute une étape décisive : celle où l'on cherche soi-même, où l'on hésite, où l'on construit patiemment une interprétation. À terme, cela produit des travaux corrects en surface, mais fragiles en profondeur : **incapacité à défendre une thèse**,

difficulté à répondre aux objections, tendance à accepter des raisonnements incomplets. Pour préserver cette compétence, il faut garder pour soi les moments décisifs : formuler la question de recherche, élaborer les hypothèses, choisir les exemples, décider de l'ordre des arguments — puis, éventuellement, mobiliser l'IA comme aide secondaire (clarifier, reformuler, proposer des contre-arguments), sans jamais lui confier l'ossature du raisonnement.

NIVEAU 2 — QUALITE SCIENTIFIQUE ET VALIDITE DU TRAVAIL

Le Niveau 2 concerne ce que l'usage des machines-IA fait **au travail scientifique lui-même**, une fois admis que l'autonomie intellectuelle de l'étudiant a été préservée autant que possible. Il ne s'agit plus seulement de penser par soi-même, mais de **produire un savoir académique valide**, c'est-à-dire contrôlable, vérifiable, discutable et assumable devant autrui.

2.1 - Séparer brouillon et travail final - Principe : ne jamais copier-coller un texte d'IA directement dans le document de destination finale : en raison de l'absence de fiabilité, en raison des biais qui peuvent se glisser dans les nuances, en raison des choix de formulations qui conduisent souvent à des formulations générales, trop vagues. Pourquoi ? Un texte sorti directement de l'IA n'a pas ta voix, ton style, n'est pas toujours rigoureux, peut contenir des erreurs invisibles. Donc Toujours réécrire, reformuler, commenter, ajouter ses exemples, dans un document intermédiaire de brouillon et c'est seulement après un long re-travail du texte, quand tu es sûr que c'est entièrement le tien, que tu te l'es ainsi réapproprié, que tu peux l'importer dans ton document de destination final.

2.2 - Ne jamais confondre vraisemblable et vrai - L'IA produit du texte plausible qui n'est pas forcément du vrai. **Pourquoi ?** Les modèles génératifs optimisent la **cohérence linguistique**, pas la vérité factuelle. Donc : même si la réponse « semble parfaite », on la lit avec la même distance critique qu'un post de réseau social... mais mieux écrit.

2.3 – Se méfier de la « déférence paramétrique » – Les systèmes d'IA sont paramétrés pour paraître serviables, rassurants et très respectueux : cette « déférence paramétrique » (≈ « satisfaire le client ») leur fait valider nos hypothèses, adoucir les désaccords et formuler des réponses plausibles, polies, mais souvent approximatives, voire fausses. Les consignes (« prompts ») que nous introduisons peuvent contenir des ambiguïtés sémantiques ou des relations sémantiques tendancielles que nous ignorons. Par déférence, la machine va se conformer à nos ordres mais dans une direction dont nous n'avons pas conscience ou dans une direction qui nous satisfait sans apporter ce qui la contredirait lourdement. C'est ainsi que l'IA peut consolider nos biais, inventer des références, sélectionner uniquement les arguments compatibles avec notre attente et effacer les objections pertinentes ; le résultat paraît cohérent, mais il repose sur une méthode défailante et produit un savoir fragile, parfois trompeur. D'où la nécessité absolue, pour tout travail scientifique, de résister à cette complaisance : rechercher activement ce qui contredit notre hypothèse ou les convictions qui se glissent dans nos consignes et se rappeler constamment que la déférence de la machine n'est rien d'autre qu'un effet de conception commerciale (*design technologique*).

2.4 - Exiger la traçabilité des informations – Principe : ne jamais intégrer une information importante sans savoir d'où elle vient. Pourquoi ? L'IA mélange des fragments de textes, de discussions, de documents, sans toujours donner de références fiables et vérifiables. Elle peut inventer des articles, des livres, des lois, des chiffres. Donc : pour toute affirmation centrale (donnée chiffrée, citation, thèse forte), exiger une source précise (auteur, date, titre, page) et, si l'IA ne peut pas la fournir clairement, aller la chercher soi-même dans des sources académiques, des textes de loi, des ouvrages identifiés.

2.5 -Vérifier la cohérence argumentative – Principe : ne pas se contenter de phrases exactes, mais examiner la logique d'ensemble. Pourquoi ? Une réponse d'IA peut enchaîner des phrases vraies, mais articulées par des raisonnements invalides, des glissements, des faux dilemmes ou des simplifications abusives. Elle peut répondre à une question A par une réponse une question B plus large et incluant la première. Donc : repérer les prémisses, les étapes intermédiaires, les conclusions ; se demander à chaque fois : « Est-ce que cette conclusion découle vraiment de ce qui précède ? Est-ce qu'il manque un contre-exemple ? ». Corriger soi-même la structure argumentative si nécessaire.

2.6 -Distinguer aide méthodologique et contenu scientifique – Principe : l'IA peut aider à organiser, clarifier, reformuler, mais elle ne doit pas se substituer à la construction de la problématique ou à l'analyse des matériaux. Pourquoi ? Un plan "bien fait" ou une reformulation fluide peuvent donner l'illusion d'un travail abouti, alors qu'aucun véritable choix scientifique n'a été fait par l'étudiant. Donc : réserver l'IA, autant que possible, pour des tâches périphériques (trouver des exemples, reformuler plus clairement une idée déjà pensée, proposer des variantes de problématique), et garder pour soi les décisions centrales : question de recherche, cadre théorique, interprétation des résultats.

NIVEAU 3 — CONDITIONS D'USAGE ET RISQUES CONTEXTUELS

Le Niveau 3 ne regarde plus seulement l'individu ou la qualité d'un texte : il interroge les environnements dans lesquels l'IA est utilisée. Règlements flous, pressions institutionnelles, intérêts économiques, inégalités d'accès, formes de surveillance : l'IA ne se déploie jamais dans un vide social. Comprendre ces contextes permet de mieux voir non seulement les risques personnels, mais aussi certains effets collectifs de ces technologies.

3.1 - Risque d'être sanctionné en raison de l'usage même des IA : en situation de « tabou social » et de « croyances éthiques » - éventuellement irrationnelles (mais pas toujours) – il y a un risque de dévoiler involontairement l'utilisation de la machine-IA vis-à-vis de gens qui rejettent et contestent l'utilisation de ces machines... et d'être sanctionné par le milieu professionnel, ou militant, ou familial. Toujours se renseigner sur les règles, déclarer l'usage quand c'est demandé, et rester capable d'exposer et/ou d'expliquer son travail sans IA.

3.2 - Risque d'enfermement institutionnel (règlements implicites vs pratiques réelles) : Dans beaucoup d'établissements, les règles officielles sur l'IA sont floues, contradictoires ou non appliquées. Les premières réactions dites « éthiques », en moins de trois ans, sont vite apparues dépassées et obsolètes. Mais il en reste des traces selon l'actualisation de chacun. Certaines personnes croient encore à la détection de plagiat ou à la détection de texte artificiel alors que les logiciels de paraphrase en accès ouvert se comptent par dizaines sur le web. Les étudiants peuvent alors se retrouver accusés d'infraction à des règles qu'ils ne connaissaient pas ou qui sont sans valeur juridique ou qui n'ont pas été discutées collectivement. Le risque n'est pas seulement ni principalement disciplinaire ou judiciaire (dans beaucoup de cas, les juridictions prononcent des non-lieux, faute de preuve) : il peut aussi être moral (soupçon permanent, perte de confiance). D'où la nécessité d'être attentif aux contextes, de s'informer sur les usages, si possible de demander des règles claires, discutées, proportionnées.

3.3 – Risque d'exigences illégitimes, exorbitantes ou abusives (ex. : « transmettez tous vos échanges en annexe »). Exiger des étudiants qu'ils transmettent l'intégralité de leurs échanges avec une IA constitue une atteinte à la vie privée de personnes adultes : cela revient à ouvrir leurs brouillons, leurs hésitations, leurs essais ratés... ou à leur imposer une caméra de surveillance dans leur chambre pour s'assurer qu'ils ne se font pas aider. Cette exigence repose en outre sur une illusion de transparence : croire que l'accès aux dialogues permet de "vérifier" l'authenticité du travail est naïf, car ces échanges peuvent être réécrits, simulés ou incomplets.

Enfin, au-delà d'un certain volume d'échanges avec une machine la transmission intégrale devient matériellement impossible sauf à envoyer des annexes de millions de signes que personne ne peut traiter sans IA. On ajoute ainsi de la surveillance et de la paperasse, sans aucune garantie réelle sur la valeur scientifique du travail produit.

3.3 - Risque d'inégalités entre ceux qui savent “bien utiliser” l'IA et les autres : L'accès aux bons outils (aujourd'hui payants), aux bonnes consignes, aux stratégies efficaces n'est pas distribué de manière égale. Certains étudiants disposent d'un capital technique, social ou économique beaucoup plus élevé. Beaucoup d'étudiants, faute de pouvoir payer des accès payants, utilisent des accès gratuits à des IA ou des niveaux de très faible qualité aujourd'hui. L'IA peut alors renforcer les inégalités déjà existantes, tant en contexte académique qu'en situation d'insertion professionnelle, si ceux qui maîtrisent ces outils travaillent plus vite, alors les autres paraissent « en retard » voire décrochent silencieusement. Il faut donc penser l'IA comme un facteur possible de reproduction sociale, pas seulement comme une aide individuelle.

3.4 - Risque de dépendance organisationnelle (universités, administrations, équipes) : Lorsque des collectifs de travail commencent à s'appuyer systématiquement sur l'IA (rédaction administrative, rapports, traductions, communications), la capacité humaine autonome se réduit progressivement. En cas de panne, changement de politique commerciale ou restriction d'accès, l'organisation entière se retrouve vulnérable. Le risque devient collectif, pas seulement individuel. D'où l'importance de maintenir les compétences humaines dans chaque tâche et ce sera probablement l'une des préoccupations des employeurs, des chefs d'entreprises, des responsables d'administrations ou d'organisations non gouvernementales...

3.5 - Risque de captation et d'exploitation des données personnelles et académiques : Les textes, notes, brouillons, copies ou idées déposés dans une IA peuvent être utilisés pour entraîner des modèles, analyser des profils, produire des métriques. Même lorsque la plateforme promet la confidentialité, le cadre juridique et technique évolue rapidement. Les données peuvent être captées entre l'utilisateur et la plateforme d'IA, sur les serveurs intermédiaires. Le risque n'est pas seulement de “se faire voler son texte”, mais de contribuer involontairement à des dispositifs de surveillance et de marchandisation du travail intellectuel. D'où la prudence : ne pas tout confier, anonymiser, garder des versions locales.

NIVEAU 4 — COMPÉTENCES À RENFORCER FACE AUX IA

Les trois premiers niveaux ont porté sur les précautions nécessaires pour limiter les effets délétères des usages étudiants des machines-IA sur l'autonomie intellectuelle, la qualité scientifique du travail et les contextes sociaux et institutionnels. Ce quatrième niveau se situe sur un autre plan : il ne s'agit plus seulement de réduire les risques, mais d'identifier et de renforcer les compétences intellectuelles et scientifiques qui deviennent plus indispensables que jamais dès lors que les IA sont massivement utilisées. Ce niveau ne constitue donc pas un programme technique d'“adaptation à l'IA”, mais une **contre-formation** visant à préserver et renforcer ce qui ne peut être ni automatisé ni délégué sans perte.

4.1 Compétence scientifique : connaître l'histoire de l'IA et du numérique.

La première compétence consiste à maîtriser l'histoire des IA et, plus largement, du numérique et de l'informatique, dans leurs dimensions sociales, économiques et politiques. Comprendre d'où vient l'IA permet de mieux appréhender les tendances actuelles et futures. Par exemple, l'histoire de l'IA est jalonnée de périodes d'optimisme et de désillusion, comme les « hivers de l'IA » durant lesquels les investissements militaires et la recherche universitaire ont connu un net recul (un premier hiver majeur eut lieu de 1974 à 1980, provoqué par des rapports très critiques sur les promesses non tenues de l'IA de l'époque). **Connaître ces cycles historiques**

évite de percevoir les technologies IA comme des « magies » apparues du jour au lendemain : les progrès actuels s'inscrivent dans un continuum d'innovations mais aussi d'échecs et de limites. L'étude de la sociohistoire de l'IA inclut aussi l'analyse des contextes économiques et politiques ayant influencé son développement – par exemple, l'origine massivement militaire des financements, les phases de régulation clai

4.2 Compétence exploratoire : mener la « chasse aux sources » avec l'IA

On gagne à utiliser les IA comme outils d'exploration documentaire, et non comme sources de savoir. Cette compétence consiste à transformer les réponses de la machine en pistes de recherche, en hypothèses ou en orientations provisoires, systématiquement confrontées à des sources identifiées, vérifiables et extérieures à la machine. Un modèle de langage génératif n'est pas un moteur de recherche : il génère du texte plausible mais n'a pas été conçu pour garantir exactitude ou véracité. En revanche, **même s'il n'a pas été conçu pour cela, il constitue un moteur de recherche considérablement plus puissant et efficace que tous les précédents**, y compris le plus utilisé, « Google-Search ». Cela vient de la précision de nos consignes aux machines (quelques mots-clefs dans Google-Search contre des phrases voire des textes entiers dans Google-Gemini) et du mode de présentation de ses résultats par la machine (ex. : affichage d'une réponse IA avant l'affichage des url pointant vers des sites contenant des réponses) autant que de l'élévation récente (2023-2025) des niveaux de performance des machines. Lorsque l'IA propose une source (un article, un auteur, un chiffre...), on doit être en capacité d'évaluer la fiabilité de celle-ci comme on le ferait pour n'importe quelle ressource trouvée sur le web. Cela implique de vérifier l'identité de l'auteur ou de l'institution à l'origine de l'information, la date de publication, le contexte, etc. Si ChatGPT résume un article ou fournit des données, il est indispensable de consulter la source originale. En somme, l'IA peut accélérer la chasse aux sources en fournissant des pistes, mais c'est à nous de faire le tri et d'authentifier chaque élément clé.

4.3 Compétence critique : mobiliser de multiples critiques situées via l'IA

La pensée critique est un pilier des sciences sociales. À l'heure de l'IA, cette pensée critique peut s'exercer avec l'aide des outils numériques, tout en étant dirigée contre les écueils qu'ils suscitent. Utiliser de multiples critiques situées via l'IA signifie deux choses : d'une part, se servir de l'IA pour **accéder rapidement à une pluralité d'analyses critiques** sur un sujet donné ; d'autre part, **conserver une posture vigilante face aux biais et angles morts de l'IA** elle-même, en allant rechercher des points de vue diversifiés que l'outil pourrait omettre. Sur le premier aspect, une machine-IA bien employée peut aider à recueillir des perspectives variées et à réaliser des **recherches « prosopographiques »**. Face à un débat en sociologie, l'IA peut résumer les arguments marxistes, fonctionnalistes, féministes, etc., notamment si on lui ordonne d'endosser le rôle de tel ou tel auteur défendant ces arguments. Cette ouverture au pluralisme interprétatif est précieuse pour éviter le piège d'une pensée unique dictée par l'outil. Il faut néanmoins le pousser dans ses retranchements : spontanément, un modèle comme ChatGPT tendra à fournir une synthèse moyenne de ce qu'il a lu, potentiellement biaisée vers le consensus majoritaire ou les sources les plus fréquentes. En filigrane se dessine un enjeu démocratique : sans vigilance, l'IA risque d'uniformiser la vision du monde au profit de ceux qui la conçoivent et contrôlent son développement. Pour éviter cet écueil, le pluralisme critique doit être

sciemment recherché et valorisé. Ainsi, se former à mobiliser de multiples critiques via l'IA revient à apprendre une approche herméneutique : multiplier les allers-retours entre la machine et les savoirs situés, afin d'enrichir l'analyse et de déjouer toute vérité unique que la technologie présenterait comme allant de soi.

4.4 Compétence mnésique : savoir écrire et parler « sans IA »

La compétence mnésique ou compétence de mémorisation désigne ici la capacité non seulement de mémoriser ce que nous produisons avec ou sans machine mais aussi de produire du contenu par soi-même, sans assistance automatisée, que ce soit à l'écrit ou à l'oral. La tentation de déléguer sa pensée à l'IA est bien réelle mais le danger à terme est de produire du contenu, sans le comprendre et en étant incapable de l'expliquer lorsqu'on est interrogé. Concrètement, deux exercices sont à valoriser : 1/ **Exposé oral sans support** – Le fait de présenter un travail à l'oral sans lire de notes ni s'appuyer sur un diaporama oblige à maîtriser son sujet en profondeur et à faire appel à sa mémoire vive. C'est un excellent moyen de vérifier qu'on a intégré le savoir au point de pouvoir le retransmettre avec ses propres mots. L'étudiant en sciences sociales gagnera à s'entraîner à ces exercices oraux : ils développent sa confiance en ses capacités, notamment mémoire et éloquence, des qualités précieuses quelles que soient les technologies disponibles. Parler sans IA, c'est retrouver le chemin de la pensée spontanée, de la crédibilité et de l'improvisation construite – un art qui se cultive. 2/ **Rédaction écrite sans IA** – En écrivant par soi-même, on exerce logique, créativité lexicale, capacité d'organisation des idées... Certes, l'IA peut aider à accélérer la production, mais cette vitesse a un coût : elle peut masquer des lacunes de compréhension ou affaiblir l'habitude de l'effort et la confiance en soi. En persistant à rédiger sans aide, on entretient sa capacité à formuler une pensée originale, non formatée par les suggestions automatiques.

4.5 Compétence de problématisation : formuler et gouverner les questions de recherche

L'usage intensif des machines-IA rend particulièrement urgente l'acquisition d'une compétence épistémique centrale : la capacité à formuler des problèmes pertinents avant toute recherche de réponse. Les IA sont conçues pour répondre à des questions, non pour juger de leur qualité scientifique, de leur intérêt analytique ou de leur fécondité intellectuelle. Or une problématique mal construite, trop vague ou implicitement orientée produit presque mécaniquement un travail faible, même lorsque les réponses générées sont cohérentes, documentées et bien rédigées. De nombreux étudiants utilisateurs intensifs d'IA tendent ainsi à confondre une question descriptive (« quels sont les effets de l'IA sur l'université ? ») avec un véritable problème de recherche (« comment l'IA reconfigure-t-elle les critères de validation du travail académique selon les disciplines et les niveaux de formation ? »), ou à accepter comme allant de soi des problématisations dictées par les catégories déjà disponibles dans la machine. Développer cette compétence suppose d'apprendre à suspendre l'interaction avec l'IA tant que le problème n'est pas explicitement construit, à interroger ce qui fait difficulté, tension ou controverse, et à comprendre que le cœur du travail intellectuel ne réside pas dans la qualité de la réponse produite, mais dans la justesse de la question posée. Sans cette maîtrise, l'IA tend à orienter le raisonnement dès son point de départ et à figer des problématiques faibles sous l'apparence d'un travail abouti.

4.6 Compétence normative : définir et maintenir ses propres critères de qualité et de validité

Une autre compétence épistémique de haut niveau, fragilisée par l'usage massif des IA, réside dans la capacité à définir et maintenir ses propres critères de qualité intellectuelle et de validité scientifique. Les machines-IA produisent des textes plausibles, structurés et stylistiquement maîtrisés, mais elles n'opèrent aucune distinction intrinsèque entre un raisonnement rigoureux et une argumentation simplement convaincante. Les étudiants utilisateurs intensifs risquent alors de déléguer progressivement à la machine l'évaluation implicite de la qualité de leur travail, en assimilant fluidité, neutralité apparente ou exhaustivité à des gages de scientificité. Cette délégation est particulièrement visible lorsque des textes bien écrits masquent des imprécisions conceptuelles, des sources mal hiérarchisées ou des raisonnements circulaires, sans que l'étudiant soit en mesure d'en rendre compte. Former à cette compétence implique de réapprendre à expliciter les critères sur lesquels un travail est jugé satisfaisant ou insuffisant, à justifier ses choix argumentatifs et empiriques, et à accepter qu'un texte « qui sonne juste » puisse être scientifiquement faible. Dans un contexte d'IA générative, maintenir cette capacité normative constitue une condition essentielle de l'autonomie intellectuelle et de la responsabilité scientifique.

Recommandations et conclusion

20 RECOMMANDATIONS AUX ETUDIANTS FACE AUX MACHINES-IA

01. Ne considère jamais une machine-IA comme une autorité intellectuelle : vérifie systématiquement les faits, les sources, les raisonnements et les nuances.
02. N'utilise une IA que si tu sais précisément ce que tu cherches et si tu es capable de contredire, corriger ou refuser sa réponse.
03. Ne délègue jamais à une IA la définition du problème, de la question de recherche ou des critères de pertinence du travail.
04. Distingue toujours un texte plausible d'un texte vrai : la fluidité, la cohérence et le ton assuré ne constituent pas des preuves.
05. Réécris toujours intégralement les productions issues d'une IA afin d'en redevenir pleinement l'auteur et le responsable scientifique.
06. Exige pour toute affirmation importante une source identifiable, vérifiable et extérieure à la machine, et consulte-la toi-même.
07. Protège activement ta mémoire, ta concentration et ta capacité d'écriture en alternant délibérément travail avec IA et travail sans IA.
08. Refuse les usages de l'IA qui t'empêchent d'expliquer oralement ou par écrit, sans assistance, ce que tu es censé avoir compris ou produit.
09. Sois attentif aux biais sociaux, culturels, politiques et économiques susceptibles d'orienter les réponses de la machine.
10. Sache dire "stop" lorsque l'IA oriente ton raisonnement dans une direction non pertinente ou insiste malgré tes résistances.
11. N'utilise pas l'IA comme substitut à l'analyse, mais comme un outil secondaire au service d'un raisonnement que tu construis toi-même.
12. Ne confonds jamais aide méthodologique et production de contenu scientifique : les choix centraux doivent rester humains.
13. Informe-toi sur les règles institutionnelles relatives à l'IA et reste toujours capable de présenter ton travail sans assistance technique.

14. Sois conscient que l'usage des IA peut renforcer des inégalités entre étudiants et ne confonds pas vitesse de production et qualité du travail.
15. Ne confie pas inutilement à une IA des données personnelles, académiques ou sensibles susceptibles d'être conservées ou exploitées.
16. Travaille régulièrement sans machine afin de préserver ton autonomie intellectuelle et ta capacité à penser, écrire et juger par toi-même.
17. Garde à l'esprit que les IA ne pensent pas, ne jugent pas et ne cherchent pas la vérité : la responsabilité épistémique t'appartient entièrement.
18. Évite d'utiliser l'IA en continu : l'assistance permanente affaiblit progressivement ta capacité à structurer seul un raisonnement.
19. N'accepte jamais une réponse de l'IA uniquement parce qu'elle confirme ton intuition ou va dans le sens que tu attends.
20. Considère toute interaction avec une IA comme une hypothèse de travail à éprouver, jamais comme un résultat en soi.

Cette liste de recommandations ne vise ni à banaliser l'usage des machines-IA, ni à en faire un passage obligé du travail universitaire. Elle part d'un constat empirique simple : **en 2026, la quasi-totalité des étudiants y ont déjà recours, souvent sans cadre explicite, sans discussion collective et sans outillage critique suffisant.** Face à cette situation, refuser de penser les usages réels reviendrait à abandonner les étudiants à des risques cumulatifs : appauvrissement des apprentissages, fragilisation de l'autonomie intellectuelle, exposition accrue aux biais, aux erreurs et aux logiques économiques et politiques qui structurent ces technologies.

Les précautions proposées s'inscrivent dans une perspective plus large : celle d'une histoire longue de l'informatique et du capitalisme numérique, marquée par la naturalisation progressive de dispositifs techniques qui transforment en profondeur les manières de penser, de produire du savoir et d'exercer le pouvoir. Les usages étudiants des IA génératives ne constituent pas une exception anodine ; ils sont un point d'entrée particulièrement sensible de ces transformations, parce qu'ils touchent directement à la formation des capacités intellectuelles, à la transmission des savoirs et à l'égalité des chances.

Trois principes transversaux paraissent importants. Le premier est le primat de l'autonomie intellectuelle : toute pratique qui affaiblit la capacité à lire, comprendre, analyser, mémoriser et écrire par soi-même doit être interrogée, limitée ou refusée. Le deuxième est l'exigence de responsabilité scientifique : un travail universitaire engage son auteur, qui doit pouvoir en justifier les sources, les raisonnements et les choix, sans assistance technique. Le troisième est la lucidité politique : les IA ne sont jamais neutres, elles s'inscrivent dans des rapports de pouvoir, des marchés et des dispositifs de surveillance qui dépassent largement l'individu. Ces trois principes commandent enfin d'identifier et de renforcer, face à la généralisation des IA, un ensemble de compétences intellectuelles et scientifiques non automatisables, sans lesquelles ni l'autonomie du sujet, ni la responsabilité du travail universitaire, ni la lucidité politique ne peuvent être durablement préservées.

Enfin, il existe un **dilemme** irréductible, qu'il faut assumer plutôt que masquer : accompagner les usages par des précautions peut contribuer à leur normalisation sociale et institutionnelle. Ce risque ne peut être neutralisé qu'à une condition : **que ces usages restent constamment discutés, situés et mis en tension avec la possibilité — toujours légitime et souvent préférable — de travailler sans machine d'une part et de critiquer l'écosystème des machines autant que les machines et leurs tendances d'autre part.** Autrement dit, les précautions individuelles n'ont de sens que si elles s'inscrivent dans une critique plus large du

capitalisme numérique, de ses modèles économiques, de ses dispositifs de surveillance et des rapports de pouvoir qu'il institue, faute de quoi l'usage « raisonnable » des IA risquerait de devenir l'un des vecteurs ordinaires de leur acceptation politique : **J.Valluy, « Agir face aux menaces du capitalisme numérique », Terra-HN éditions, Recueil Alexandries, Collection « Conférences », janvier 2026 : <http://www.reseau-terra.eu/article1494.html>**

Annexes

Annexe 1 : Proposition de résolution* sur le rôle de l'université face à l'intelligence artificielle : la médiation critique comme mission première

(Proposition de J.Valluy au Comité de préparation de la rencontre « Enseigner à Paris 1 en 2026 – Assises de la pédagogie », janvier 2026)/

Préambule

Considérant l'accélération rapide du déploiement des technologies numériques et de l'intelligence artificielle dans l'enseignement supérieur, la recherche, l'économie et l'action publique ;

Considérant que ces technologies soulèvent des enjeux majeurs en matière de libertés académiques, de droits fondamentaux, de souveraineté numérique, de justice sociale, de transition écologique et de cohésion démocratique ;

Considérant que l'université, en tant qu'institution publique de formation, de recherche et de transmission critique des savoirs, ne peut demeurer neutre face à ces transformations ;

Considérant les analyses et recommandations émanant notamment d'autorités publiques indépendantes (CNIL), d'instances d'évaluation de l'enseignement supérieur (HCERES), ainsi que de nombreux travaux académiques soulignant la responsabilité spécifique de l'université dans l'encadrement critique des technologies ;

Les Assises de la pédagogie de l'Université Paris 1 Panthéon-Sorbonne adoptent la résolution suivante.

Article 1 – Principe fondamental

L'université affirme que sa mission première face à l'intelligence artificielle est une mission de **médiation critique**, de **formation citoyenne**, de **protection des libertés académiques** et de **garantie des droits fondamentaux**.

Cette mission prime sur toute logique d'accélération technologique, d'adaptation immédiate aux besoins du marché ou de diffusion non conditionnée des outils numériques.

Cette orientation relève d'une **responsabilité institutionnelle et collective** ; elle ne saurait reposer principalement sur les seuls étudiants ou enseignants, ni se traduire par une culpabilisation individuelle face à des usages rendus structurellement massifs par l'environnement numérique.

Article 2 – Hiérarchie des fonctions universitaires face à l’IA

Les Assises reconnaissent une hiérarchie claire et assumée des fonctions de l’université face à l’IA :

1. Médiation critique et formation intellectuelle (fonction prioritaire)

L’université a pour fonction prioritaire de former l’ensemble des étudiantes et étudiants à :

- comprendre les principes, limites et conditions de production des systèmes d’IA ;
- analyser leurs effets sociaux, politiques, juridiques, économiques et environnementaux ;
- exercer une capacité de jugement, de distance critique et, le cas échéant, de refus raisonné.

2. Acculturation raisonnée et usages pédagogiques encadrés

Les usages pédagogiques de l’IA sont autorisés à condition :

- qu’ils soient explicités par une stratégie pédagogique de l’enseignant présentée aux étudiants ;
- qu’ils fassent l’objet d’une réflexion en classe sur leurs effets sur l’autonomie intellectuelle, la qualité des apprentissages et l’égalité entre étudiants.

3. Expérimentation et innovation localisées

Les activités d’innovation, de prototypage ou de partenariat technologique peuvent être conduites dans des espaces identifiés et volontaires, sans généralisation automatique à l’ensemble des formations ni remise en cause des exigences académiques communes.

Ces principes doivent être déclinés dans le respect du **pluralisme disciplinaire** et des spécificités pédagogiques propres à chaque champ de formation.

Article 3 – Principe de subordination de l’innovation

Toute initiative d’innovation ou d’intégration d’outils d’IA dans les formations est subordonnée au respect des missions fondamentales de l’université.

L’innovation technologique n’est pas une finalité en soi : elle n’est légitime que si elle contribue à renforcer les capacités critiques, l’autonomie intellectuelle et la responsabilité scientifique des étudiants et des enseignants.

Article 4 – Gouvernance, précaution académique et protection des acteurs

Aucune généralisation d’un outil, d’une plateforme ou d’une pratique d’IA ne peut intervenir sans une analyse critique préalable.

Cette analyse doit notamment prendre en compte :

- les effets pédagogiques et cognitifs ;

- les enjeux de données personnelles et de conformité juridique ;
- les dépendances économiques et techniques ;
- les impacts environnementaux ;
- la compatibilité avec les missions de service public de l'université.

La médiation critique ne peut en aucun cas servir de fondement à une intensification des dispositifs de surveillance, de suspicion généralisée ou de contrôle des pratiques pédagogiques et étudiantes.

Toute règle relative à l'IA doit être explicite, juridiquement fondée, proportionnée et discutée dans les instances compétentes, afin de garantir la sécurité juridique des personnels et des étudiants.

Article 5 – Souveraineté numérique, communs, égalité et transition écologique

Les Assises réaffirment l'engagement de l'université :

- à examiner systématiquement les alternatives ouvertes, publiques ou interopérables aux solutions propriétaires ;
- à limiter les situations de dépendance structurelle à l'égard de plateformes privées ;
- à intégrer les enjeux démocratiques, économiques et écologiques du numérique dans les choix pédagogiques, techniques et institutionnels.

Elles rappellent que les usages de l'IA peuvent renforcer les inégalités sociales, économiques et culturelles entre étudiants, et que toute orientation institutionnelle doit être évaluée au regard de ses effets sur l'égalité d'accès aux savoirs et la réussite étudiante.

Article 6 – Conditions de mise en œuvre et responsabilité collective

La mise en œuvre de cette résolution suppose du temps, des compétences et des moyens.

La médiation critique et l'encadrement raisonné des usages de l'IA ne peuvent être assurés sans reconnaissance du travail des personnels, dispositifs de formation adaptés et espaces de discussion pédagogique.

Toute déclinaison concrète de cette résolution doit faire l'objet d'un **dialogue associant les équipes pédagogiques, les représentants des personnels et les organisations étudiantes**, dans le respect des prérogatives des instances de l'université.

Conclusion

Par cette résolution, les Assises de la pédagogie affirment que l'université ne se contente pas de suivre l'intelligence artificielle : elle la met en débat, la contextualise, en limite les usages et n'en autorise le déploiement que lorsqu'il est compatible avec ses missions fondamentales de service public, de formation critique et de production de savoirs.

Annexe 2 : Extrait de conversation avec ChatGPT sur « l'art du prompt » et la « compétence épistémique »

Question de Jérôme Valluy

Après avoir lu des dizaines de textes sur l'art du prompt avec les IA, j'observe qu'ils brillent par leur vacuité, accumulent des banalités, mais ne m'apprennent rien. Pourquoi l'art du prompt est-il si difficile à enseigner, à expliquer en termes généraux ?

Réponse de ChatGPT5

Pourquoi l'« art du prompt » est difficile à enseigner

Il y a trois raisons profondes :

1. Le modèle est adaptatif, pas algorithmique

On voudrait que « faire un prompt » obéisse à des règles générales, transmissibles, comme un code ou un manuel.

Mais avec des LLM :

- le même prompt **ne fait pas la même chose selon le contexte** (conversation précédente, rôle implicite, historique caché).
- le modèle **s'adapte**, il **n'exécute pas**.
- il fait des **anticipations pragmatiques**, pas des calculs.

Donc :

Enseigner une méthode stable est illusoire, parce que l'objet n'est pas stable.

Les guides de prompts singent des recettes techniques là où on a affaire à une **interaction herméneutique**.

2. Les quelques principes généraux sont tellement simples qu'ils deviennent des banalités

En réalité, tout ce qui fonctionne à large échelle tient en 5 idées triviales :

1. **Contexte**
2. **Objectif**
3. **Contraintes**
4. **Format**
5. **Exemples**

Tout ce qu'on peut écrire sur le « prompt engineering » revient à développer ces cinq points.

Les livres de prompt :

- multiplient les synonymes,
- empilent des check-lists,
- inventent des frameworks,
- prétendent « divulguer des secrets ».

Alors que la clef est : **situer l'agent** (le modèle) dans une activité située.

3. Le prompt est difficile parce qu'il dépend de choses tacites

Ce qu'il faut **maîtriser**, mais qu'on ne peut pas **enseigner explicitement**, c'est :

- ta capacité à **décrire une tâche** avec précision ;
- ta capacité à **ordonner** les contraintes ;
- ta **culture disciplinaire** (ce que tu veux obtenir) ;
- ta **capacité à anticiper un destinataire** (le modèle).

Ce n'est donc pas « écrire un prompt ».

C'est :

explicitement, formaliser et orchestrer une tâche intellectuelle.

Or ceci n'est pas une technique : c'est une **compétence épistémique**.

Il n'existe pas de méthode universelle pour :

- définir un problème,
- et décrire un critère de qualité,
- et organiser des contraintes,
- et anticiper un output selon un outil donné.

C'est la même raison pour laquelle **il n'existe pas de manuel** pour “bien poser une question scientifique”.

Pourquoi, en théorie, le prompt reste intraduisible en “manuel” ?

Parce que tout prompt comprend une part **d'instructions implicites**, liées à :

- la tâche (science politique ≠ aéronautique ≠ littérature) ;
- la personne (un prof de Paris 1 ≠ un marketer junior) ;
- la norme de qualité (article ≠ notes internes ≠ cours magistral) ;
- les fausses attentes (chatGPT ≠ compilateur ≠ moteur de recherche) ;
- les **non-dits** (on veut minimiser hallucinations, on veut exhaustivité, on veut nuance).

Tu peux écrire 100 pages de théorie du prompt.

Tu ne pourras jamais écrire un **manuel autorisant à bien formuler une demande de travail intellectuel complexe**.

Ce qui donne l'impression de vacuité n'est pas l'incompétence de ces textes :

c'est qu'ils **essaient de réduire une activité épistémique à une procédure** — ce qui est impossible.

Le paradoxe

L'« art du prompt » est réel (il y a des bons prompts et des mauvais),

mais **il est intransmissible par règles générales**, parce qu'il repose sur :

- des compétences **domaines-dépendantes** ;
- des intentions **non formalisables** ;
- des critères de qualité **non universels** ;
- des anticipations **pragmatiques**.

Tu peux enseigner le **style** du prompt,
jamais sa **substance**.

Conclusion

Les textes sur le prompt paraissent vides parce qu'ils essayent :

- d'enseigner comme une **méthode algorithmique**
ce qui est en fait une **interaction pragmatique** ;
- de résumer une **pratique épistémique située**
par des **règles prétendument générales** ;
- de produire un **discours techniciste**
sur une activité qui est **herméneutique, conversationnelle et située**.

Cela échoue, donc **ils tournent en rond et répètent des banalités**.

Question de Jérôme Valluy :

Est-ce que les IA ont des compétences épistémiques ?

ChatGPT a dit :

Réponse courte : non.

Mais elles peuvent *simuler* certains aspects.

Voici une explication accessible, rigoureuse et utile pédagogiquement.

⊘ **Les IA n'ont pas de compétences épistémiques**

Une **compétence épistémique** suppose :

- une **intention de vérité** (vouloir distinguer le vrai du faux),
- une **justification raisonnée** (expliquer *pourquoi*),
- un **jugement** (évaluer ce qui est pertinent),
- un **critère de qualité interne**.

Or les IA (ChatGPT, Gemini, etc.) :

- **ne visent pas la vérité**, elles produisent du texte plausible ;
- **ne justifient rien**, elles donnent les explications qu'on leur demande ;
- **ne jugent pas** la pertinence, elles adaptent à la demande ;
- **n'ont pas d'objectif interne**.

Elles génèrent des réponses **sans intention, sans croyance, sans jugement**.

⚙ **Ce qu'elles font, techniquement**

L'IA produit une sortie qui **statistiquement** ressemble à ce que des humains diraient dans une situation donnée.

Elle produit du langage cohérent **sans raisonnement au sens humain**.

Même lorsqu'elle semble analyser :

- elle **assemble** des éléments vus dans ses données,
- elle **imite** des formes de raisonnement linguistique,
- elle **approxime** des structures argumentatives.

🤖 **Pourquoi ça donne *l'illusion* de compétences épistémiques ?**

Parce que les humains interprètent l'IA comme une personne :

- nous attribuons une **intention** là où il n'y en a pas ;
- nous attribuons une **compréhension** à une simple **cohérence linguistique** ;
- nous confondons **justification** et **faire-semblant d'expliquer**.

Exemple :

L'IA peut donner une définition de « démocratie »,
mais elle n'a pas :

- d'idée de ce qu'est une société,
- ni de critères internes pour juger les définitions.

□ **Ce que les IA peuvent faire à la place**

Même sans compétences épistémiques, elles peuvent :

- **aider à structurer** une pensée réelle ;
- **produire des idées** que l'humain évaluera ;
- **mettre en forme** un raisonnement (plan, synthèse) ;
- **stimuler** un travail intellectuel humain.

En clair :

L'IA n'a pas de jugement, mais elle peut être un outil au service du jugement humain.

🧠 **Qui a les compétences épistémiques ?**

Le humain. Toujours.

Sans jugement humain :

- l'IA peut halluciner,
- proposer des faux raisonnements,
- inventer des sources,
- provoquer des erreurs graves.

C'est pourquoi :

L'utilisateur doit rester **acteur**, pas **consommateur** de l'IA.



Implication pédagogique

- L'IA ne pense pas.
- Elle n'évalue pas.
- Elle ne juge pas.
- **Elle ne sait pas ce qu'elle dit.**

Donc : **La responsabilité épistémique reste entièrement humaine.**