

Octets d'IA

Comprendre l'intelligence artificielle dans la formation parcours Réussite

Une publication de Contact North | Contact Nord
et Literacy Link South Central

e-Channel
Apprentissage en ligne



Bienvenue à *Octets d'IA*, votre guide sélectionné pour mieux comprendre l'IA dans le contexte du parcours Réussite en éducation et au travail.

De nos jours, une personne apprenante peut taper une seule phrase et voilà! ... un paragraphe complet apparaît instantanément. Cela rend cette question plus pertinente que jamais : « **Si l'IA peut faire cela instantanément, pourquoi se donner la peine d'apprendre?** »

C'est une question tout à fait légitime. Mais souvenez-vous de Joey Tribbiani dans *Friends* — [The One with the Thesaurus](#) (2:53)?

Lorsqu'il a utilisé un thésaurus pour « avoir l'air intelligent », il a transformé « Ce sont des personnes chaleureuses et gentilles au grand cœur » en « Ce sont des Homo sapiens humides et séduisants dotés de pompes aortiques de grande taille ».

L'IA peut parfois ressembler à Joey utilisant un thésaurus pour rédiger une lettre de recommandation : soignée en surface, mais pas

toujours cohérente en profondeur. Dans l'épisode, Joey remplace chaque mot par un terme plus sophistiqué et finit par produire une phrase qui semble impressionnante, mais qui s'effondre dès qu'on l'examine un peu. L'IA peut fonctionner de la même manière : fluide, convaincante et parfois complètement à côté de la plaque, à moins qu'une personne ne vérifie le sens de ce qui est écrit.

L'apprentissage demeure essentiel, car la mémorisation n'est qu'un point de départ; le véritable travail repose sur le sens, le jugement et la voix personnelle. La capacité humaine à ralentir lorsqu'elle utilise l'IA et à préserver sa propre voix dans le processus est cruciale.

Versez-vous un café et joignez-vous à nous. Dans cette édition :

- Ce que vous devez savoir sur les grands modèles de langage (LLM);



- L'IA peut générer du texte, mais elle ne peut pas remplacer l'apprentissage;
- De la théorie à la pratique : expérimenter l'IA de façon intentionnelle.

L'apprentissage est important parce que le véritable travail commence là où l'IA s'arrête : le jugement, la profondeur et le sens.

L'équipe d'Octets d'IA



Carolina Cohoon est consultante en technologie éducative chez Literacy Link South Central. Son parcours professionnel couvre l'éducation et la réadaptation, avec une passion pour l'inclusion et l'accessibilité. Elle s'engage à concevoir des expériences d'apprentissage

qui célèbrent la diversité. Son intérêt pour l'IA découle de son enthousiasme pour l'innovation, le partage des connaissances, l'amélioration de l'accessibilité et de l'expérience d'apprentissage grâce aux adaptations personnalisées que l'IA peut offrir dans le cadre de la conception universelle de l'apprentissage (CUA). Carolina détient une certification ChatGPT du Blockchain Council et a complété une formation de l'Ivey School of Business sur le leadership accéléré par l'IA, démontrant ainsi son engagement en faveur d'une éducation inclusive et améliorée par la technologie qui favorise une transformation significative de l'apprentissage.



Jeremy Marks travaille pour Literacy Link South Central comme gestionnaire de projet et chercheur en technologie éducative. Il a récemment terminé le programme Teacher/Trainer of Adults au Conestoga College et enseigne désormais dans le programme ACE au Collège Fanshawe. Depuis

2002, il a enseigné dans les écoles publiques et secondaires, les collèges, les universités, ainsi que dans le domaine du Parcours Réussite en éducation et au travail au Canada et aux États-Unis. Son intérêt pour l'IA provient de sa passion de longue date pour la théorie de l'éducation et la philosophie cognitive. Il est l'auteur de quatre livres.

** Ce bulletin est publié sous la direction de Contact North / Contact Nord. Les outils d'IA générative ont facilité la conceptualisation, la mise en forme et la création d'images dans le cadre de ce travail. Les idées originales, la traduction, les connaissances, la recherche et les liens ont été développés par les auteurs. Cette divulgation reflète notre engagement en faveur de la transparence, de l'intégrité intellectuelle et de l'utilisation responsable des technologies émergentes.*

Ce que vous devez savoir sur les grands modèles de langage (LLM)

Une définition adaptée au parcours Réussite en éducation et au travail

Les grands modèles de langage, ou LLM (*Large Language Models*), sont des systèmes fondés sur la reconnaissance de motifs, entraînés à générer le jeton (*token*) le plus probable dans une séquence. Une grande partie de la fluidité, du raisonnement apparent et des réponses formulées avec assurance que nous observons provient de ce processus d'entraînement.

Les LLM constituent un type d'IA générative, mais ce sont ceux que la clientèle apprenante est plus susceptible d'utiliser pour les remue-méninges, l'élaboration de plans, les explications et les résumés. L'IA générative est la famille plus large qui englobe les systèmes textuels, visuels, audio, de code et multimodaux.

Inspiré de How Large Language Models Work: From Zero to ChatGPT d'Andreas Stöckelbauer, Data Science at Microsoft (Medium, 2023), avec un cadrage supplémentaire en matière de littératie de l'IA dans le contexte de l'éducation des adultes. Certaines parties de l'explication des capacités de l'IA

génération s'appuient également sur le guide de littérature de l'IA de l'Université de Calgary, plus particulièrement sur la section What GenAI Can Do (remue-méninges, élaboration de plans, explications, soutien à la rédaction et synthèse de l'information).

Définition courante d'un LLM

Un LLM est un vaste réseau neuronal comptant des milliards de paramètres. Sa tâche fondamentale est étonnamment simple : prédire le prochain jeton dans une séquence. [Les jetons ne correspondent pas toujours à des mots complets](#); ils peuvent être des mots entiers, des parties de mots ou des signes de ponctuation.

L'IA ne peut pas remplacer entièrement la construction du sens qui découle de l'expérience vécue, du dialogue et de la réflexion. Elle apprend des motifs à une échelle considérable, et ce qui semble être de la compréhension émerge de motifs statistiques appris et de représentations internes, et non d'une expérience vécue ou d'une compréhension ancrée dans la réalité.

Comment les LLM apprennent

Préentraînement

Le modèle traite d'énormes quantités de texte et prédit de façon répétée le prochain jeton. Ces cycles de prédiction forment les motifs statistiques qu'il utilise pour la grammaire, la structure et les comportements langagiers courants.

Réglage fin par instruction

Des exemples soigneusement sélectionnés montrent au modèle comment suivre des consignes. Cette étape lui permet de passer de la simple continuation de texte à la production de réponses correspondant à des modèles tels que les résumés, les explications ou les procédures étape par étape.

Apprentissage par renforcement à partir de rétroaction humaine (RLHF)

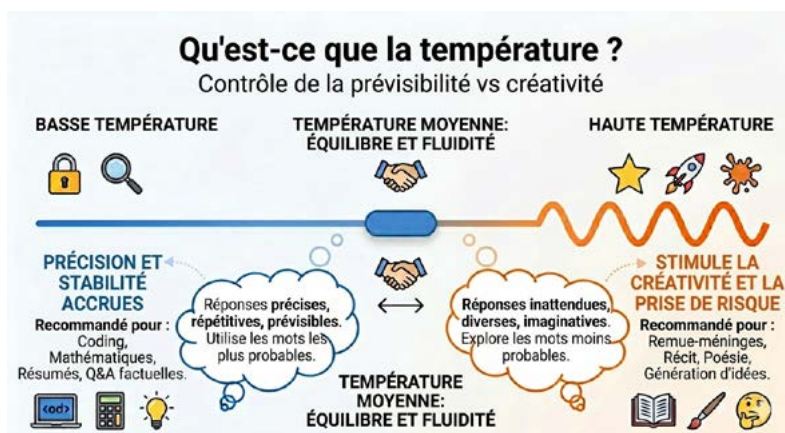
Des évaluateurs et évaluateuses humains comparent plusieurs réponses générées par le modèle et les classent. Ces classements servent

ensuite à entraîner un modèle de récompense qui reproduit approximativement les préférences humaines. Le modèle de langage est alors optimisé, souvent au moyen de l'apprentissage par renforcement, afin d'augmenter la probabilité de produire des réponses qui obtiennent de meilleurs résultats selon ce modèle de récompense. Son comportement évolue ainsi vers des modèles que les êtres humains ont jugés plus utiles, plus sûrs et plus appropriés.

Qu'est-ce que la température?

La température est l'un des paramètres les plus simples et les plus puissants qui influencent la façon dont les modèles d'IA répondent. Elle détermine si une IA adopte un ton factuel et précis ou un ton créatif et expressif.

Considérez-la comme un curseur qui établit un équilibre entre la prévisibilité et la surprise.



- Les températures faibles favorisent une plus grande exactitude et une meilleure stabilité. C'est exactement ce qu'il faut pour la programmation, les mathématiques ou la synthèse de données.
- Les températures moyennes offrent un équilibre naturel et une bonne fluidité. C'est le type d'échange naturel que l'on retrouve dans les agents conversationnels du quotidien.
- Les températures élevées stimulent la créativité et la prise de risques. Ce réglage est particulièrement utile pour les remue-méninges ou la création de récits.

Quelle est la température recommandée?

La réponse courte est : cela dépend de vos objectifs. Comme le souligne l'équipe de [Watercrawl](#), il n'existe pas de réglage universel. Le choix approprié dépend de la tâche à accomplir, qu'il s'agisse de générer du code, de rédiger des résumés ou de créer des histoires originales. L'expérimentation demeure le meilleur moyen de trouver l'équilibre souhaité entre cohérence et créativité.

Une mise en garde importante s'impose toutefois : [une faible température garantit la cohérence, mais non l'exactitude](#). Si le modèle ne possède pas l'information nécessaire, réduire la température ne fera qu'assurer que la même erreur sera reproduite à chaque exécution.

Souhaitez-vous l'essayer par vous-même? Un environnement d'expérimentation comme [Google AI Studio](#) vous permet de déplacer le curseur en temps réel et d'observer immédiatement les différences. Explorez concrètement comment les changements de température influencent le ton, la logique et la créativité de vos conversations!

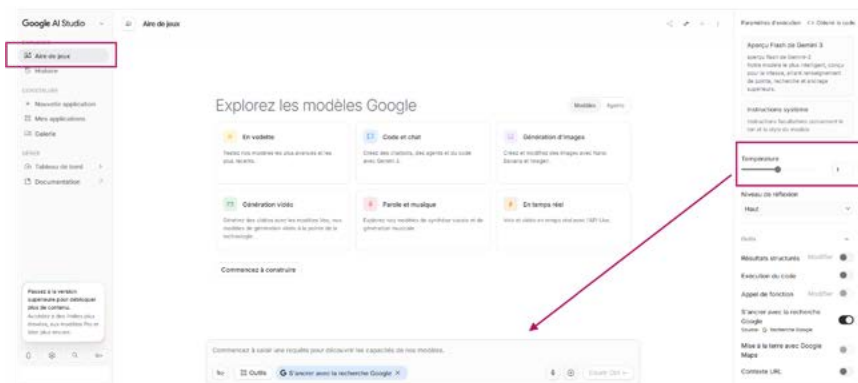
Les LLM semblent très performants, mais ils ont leurs limites

Les LLM peuvent générer un texte fluide, résumer du contenu, ébaucher des idées, expliquer des concepts et répondre d'une manière qui paraît soignée et assurée. C'est d'ailleurs l'une des raisons pour lesquelles ils peuvent sembler si convaincants, même lorsque leur contenu nécessite encore une vérification.

Un modèle peut produire une réponse qui paraît perspicace parce qu'il a appris comment la perspicacité est généralement exprimée dans les données qui ont servi à son entraînement. Cette aptitude apparente découle de l'apprentissage de motifs à grande échelle, et non d'une compréhension comparable à celle des êtres humains.

L'illustration suivante présente certaines capacités et limites associées à l'utilisation des grands modèles de langage (LLM).

D'autres exemples de capacités et de limites sont présentés dans la ressource suivante : « [AI Literacy - Artificial Intelligence](#) », bibliothèque de l'Université de Calgary.



Les grands modèles de langage (LLM), grâce à un entraînement à grande échelle sur des motifs linguistiques, peuvent :

Capacités	Limites
 Produire des paragraphes	Sans garantie de structure ou d'exactitude.
 Résumer des documents	Une vérification est requise.
 Traduire des langues	La qualité n'est pas toujours au rendez-vous. L'anglais domine la plupart des ensembles de données d'entraînement, et de nombreuses langues restent sous-représentées.
 Générer et expliquer du code	La qualité dépend des modèles présents dans ses données d'entraînement — et il y a beaucoup de code de mauvaise qualité sur Internet. Il est très performant en matière de mise en forme, mais le résultat doit tout de même être vérifié par un humain.
 Mettre en forme du texte	Les grands modèles de langage (LLM) excellent dans la mise en forme soignée des textes, mais ils ne parviennent pas à conserver la structure de manière fiable : la mise en forme peut se déformer, s'effondrer ou paraître impeccable tout en comportant des erreurs logiques ou techniques.
 Répondre à de nombreuses questions factuelles	Les réponses peuvent être incomplètes, erronées ou périmées.
 Imiter différents styles d'écriture	Uniquement au niveau superficiel : le style peut varier, les nuances sont gommées et le texte semble souvent « faux » aux yeux d'un lecteur averti.

Pourquoi les LLM commettent-ils des erreurs?

Les LLM sont des systèmes de génération de texte fondés sur des probabilités statistiques. Ils sont entraînés à poursuivre une séquence de jetons en fonction des motifs présents dans leurs données d'entraînement, et non à vérifier les faits de façon fiable ni à déterminer ce qui est vrai.

Hallucinations

L'écriture humaine adopte souvent un ton assuré. Les LLM apprennent donc à reproduire ce style. Par conséquent, ils peuvent générer un langage fluide et convaincant même lorsque l'information est peu fiable, ce qui peut parfois entraîner l'invention de détails, la fabrication de références bibliographiques ou la déformation de sources.

Ces erreurs, souvent appelées hallucinations, découlent de l'objectif même de l'entraînement du modèle et demeurent un domaine de recherche actif. Les hallucinations sont également plus susceptibles de se produire lorsque le modèle n'a pas accès à des sources externes fiables ou lorsque les requêtes exigent un rappel factuel très précis.

Bias

Les résultats des grands modèles de langage (LLM) peuvent également reproduire les biais sociaux présents dans les données sur lesquelles ils ont été entraînés. Par exemple, certains modèles établissent des associations plus marquées entre certains rôles ou traits de caractère et des genres, des origines ethniques ou des tranches d'âge précises, ce qui peut influencer la manière dont les personnes et les communautés sont décrites. Ces biais ne sont pas toujours évidents, mais ils peuvent influencer la manière dont certaines expériences, certains accents ou certains parcours de vie apparaissent dans les réponses du modèle. Ces schémas émergent parce que les données d'entraînement reflètent des inégalités historiques et sociales que le modèle peut apprendre et reproduire à moins qu'elles ne soient activement atténuées.

La modification des instructions, qui apprend aux LLM à suivre des instructions explicites, peut

réduire les biais, mais il n'est pas encore clair s'ils peuvent être totalement éliminés.

Les limites des LLM

Même les modèles les plus avancés présentent d'importantes limites intrinsèques :

- Ils ne sont pas fiables pour une compréhension de type humain;
- Ils ne constituent pas des vérificateurs de faits fiables;
- Ils ne raisonnent pas avec la responsabilité propre aux êtres humains*;
- Ils ne développent ni croyances, ni intentions, ni valeurs;
- Ils ne portent pas de jugements éthiques ou moraux au sens humain du terme;
- Ils ne peuvent pas garantir qu'une référence citée existe réellement ni qu'elle affirme ce que le modèle prétend qu'elle affirme.

Certains systèmes d'IA combinent la génération de contenu avec la recherche d'information, par exemple en consultant des documents ou le Web. Ces systèmes peuvent réduire les hallucinations en ancrant leurs réponses dans des sources réelles, mais ils exigent tout de même une vérification, puisqu'ils peuvent mal interpréter l'information ou n'en utiliser qu'une partie de façon sélective.

** Les LLM peuvent simuler le raisonnement et résoudre de nombreux problèmes, mais ce processus n'est pas fondé sur une compréhension réelle ni sur une responsabilité comparable à celle des êtres humains.*

La segmentation en jetons

L'une des façons les plus efficaces de démystifier les LLM consiste à comprendre les jetons, soit les véritables unités qu'un LLM lit et produit.

Un LLM ne traite pas nécessairement des mots entiers. Il décompose le texte en jetons : des segments qui peuvent correspondre à un mot complet, à une partie de mot, à un signe de ponctuation ou même à un seul caractère. Par exemple, le mot anglais unhappiness pourrait être

divisé en trois jetons : **un · happi · ness**.

Voici comment cette décomposition fonctionne :

- « **Un** » est un préfixe courant qui apparaît comme un jeton distinct;
- « **Happi** » est une racine fréquente (que l'on retrouve dans *happy*, *happiness* et *happily*);
- « **Ness** » est un suffixe courant qui constitue également un jeton distinct.

Chaque fois qu'un LLM génère du texte, il évalue les probabilités associées aux différents jetons susceptibles de suivre, puis répète le processus, un jeton à la fois.

À mesure que le modèle répète cette étape, jeton après jeton, la réponse prend forme. Ce processus est probabiliste plutôt que conceptuel : il est façonné par les motifs appris et le contexte, et non par une compréhension fondée sur l'expérience vécue.

Comprendre ces concepts favorise la littératie de l'IA. Lorsqu'une réponse semble faire autorité, les personnes apprenantes peuvent se demander : s'agit-il de la formulation la plus probable ou de la formulation la plus exacte? Les LLM sont optimisés pour produire le texte le plus probable plutôt que le plus véridique. Ils sont donc souvent plus fiables pour répondre à la première question qu'à la seconde.

L'idée essentielle

L'IA peut sembler représenter un changement technologique intimidant, mais les membres du secteur éducatif n'ont pas besoin d'être des spécialistes de la technologie pour animer cette conversation. Nous apprenons toutes et tous en temps réel.

Pour utiliser l'IA efficacement, la clientèle apprenante doit comprendre une idée fondamentale : l'IA peut produire d'excellentes ébauches, mais elle ne vérifie pas les faits par elle-même. Lorsque nous discutons ouvertement des situations où ces systèmes réussissent et de celles où ils échouent, la technologie devient plus facile à comprendre et à utiliser.

Cet aspect est important, car le fait de s'appuyer sur l'IA sans supervision peut limiter les possibilités et le développement d'une personne apprenante en milieu de travail. Chaque réponse générée par l'IA devrait être considérée comme un point de départ plutôt que comme une réponse définitive.

« Il est clair que ces LLM se révèlent très utiles et démontrent des capacités impressionnantes en matière de connaissances et de raisonnement, et peut-être même certains signes d'intelligence générale. Toutefois, il reste à déterminer si cela ressemble à l'intelligence humaine, et dans quelle mesure. » (traduction libre de Stöffelbauer, 2023, section Conclusion, paragr. 2)

Ce que cela signifie dans la pratique

Savoir que l'IA ne peut pas remplacer l'apprentissage n'est qu'un point de départ. La question plus importante consiste à déterminer comment rendre cette limite visible pour la clientèle apprenante.

Une façon d'y parvenir est de l'aider à examiner le fonctionnement de l'IA, à évaluer ses réponses de manière critique et à reconnaître les situations où son propre jugement demeure essentiel.

À partir de là, l'IA peut être présentée comme un soutien à l'apprentissage plutôt que comme un substitut à celui-ci. La clientèle apprenante peut explorer les ébauches produites par l'IA afin d'y repérer les biais ou les informations manquantes, et documenter leur collaboration avec l'IA au moyen d'une déclaration de divulgation ou d'une déclaration de collaboration.

Ces pratiques rendent l'apprentissage plus visible et renforcent l'importance de l'effort humain.

Dans cette optique, la littératie de l'IA ne consiste pas seulement à bien utiliser les outils. Elle consiste aussi à savoir quand les remettre en question, comment vérifier l'information et pourquoi la réflexion de la personne apprenante demeure au cœur du processus.

Voici quelques conseils pour présenter les concepts liés aux LLM :

01

Discuter de la segmentation en jetons et de la prédiction

Avant que la clientèle apprenante puisse utiliser l'IA de façon critique, elle doit comprendre ce qu'elle est réellement. Commencez par la segmentation en jetons : l'IA ne lit pas les phrases comme nous le faisons. Elle décompose le texte en fragments appelés jetons et prédit le jeton suivant le plus probable jusqu'à ce qu'une réponse soit complète.

Activité 1 : Utilisez le [Tokenizer Playground](#) pour montrer à la clientèle apprenante comment différentes phrases sont décomposées. Saisissez quelques exemples ensemble et observez comment le texte se divise en segments colorés. Notez que les couleurs elles-mêmes — violet, jaune, rouge, etc. — n'ont aucune signification particulière. Elles servent simplement à distinguer visuellement les jetons et à repérer où l'un se termine et où le suivant commence.

Activité 2 : Saisissez une requête comme « La capitale de la France est » et demandez à la clientèle apprenante de prédire la suite. Montrez ensuite comment le modèle complète la phrase. Expliquez qu'il répondra probablement « Paris », non pas parce qu'il « connaît » la réponse, mais parce que « Paris » constitue statistiquement la suite la plus probable de cette phrase.

Poursuivez ensuite avec une requête plus complexe : demandez au modèle de décrire un animal inventé, par exemple un « Octoviotopus », puis analysez ensemble la réponse produite. Cet exercice illustre comment le modèle peut générer un texte crédible même lorsqu'il ne dispose d'aucune réalité sur laquelle s'appuyer.

02

Faire des réponses de l'IA un objet d'étude

Travaillez avec la clientèle apprenante afin d'analyser des ébauches générées par l'IA. Demandez-lui de repérer les endroits où le texte repose sur des biais, invente des détails, néglige le contexte local ou simplifie excessivement certaines nuances.

Les personnes apprenantes capables d'expliquer pourquoi une ébauche produite par l'IA est faible démontrent précisément le type de raisonnement disciplinaire qu'une IA générique n'exécute pas de manière fiable.

Activité 3 : Demandez à l'IA de rédiger un aperçu des fleurs indigènes du sud-ouest de l'Ontario. Une fois le texte généré, invitez la clientèle apprenante à en faire l'analyse critique :

« Quelles voix et quelles perspectives sont mises en valeur ici, et lesquelles sont complètement absentes? »

Cet exercice met en lumière une réalité essentielle de la littérature de l'IA : les LLM ne sont jamais neutres ni universels. Ils reflètent simplement les moyennes statistiques des données dominantes et des populations qui ont servi à leur apprentissage.

03

Créer une déclaration de collaboration

Invitez la clientèle apprenante à examiner des exemples existants de déclarations d'utilisation de l'IA et de divulgation, puis à cocréer une déclaration de collaboration qui accompagnera tout travail réalisé avec l'aide de l'IA.

Cette déclaration devrait préciser :

- L'outil utilisé;
- Les requêtes soumises;
- Les éléments conservés ou modifiés;
- La contribution de la personne apprenante en matière d'exactitude, de voix personnelle, de jugement et de réflexion originale.

Cette pratique favorise une utilisation plus transparente et responsable de l'IA tout en aidant les personnes apprenantes à expliciter leur propre contribution intellectuelle. Cette expression est en soi un processus d'apprentissage et une compétence transférable en milieu de travail.

Pour les apprenantes et apprenants adultes ainsi que pour le personnel enseignant, une utilisation responsable de l'IA implique de faire preuve de transparence quant au moment et à la manière dont l'IA a été utilisée dans les travaux, les communications et les activités professionnelles.

Activité 4 : Nous avons créé un générateur de déclarations éthiques. Consultez le lien suivant pour découvrir comment l'utiliser dans votre salle de classe : [À la découverte de Google Opal.](#)

04

Produire des artefacts : flux de travail, guides de requêtes et tableaux comparatifs

Déplacez l'attention du produit final vers le processus. Encouragez la clientèle apprenante à créer et à partager :

- Des guides de requêtes expliquant ce qui a fonctionné et pourquoi;
- Des tableaux comparatifs présentant les réponses de l'IA et leurs versions révisées; des flux de travail documentés montrant comment l'IA s'intègre à un processus de réflexion plus vaste.

Ces artefacts rendent le processus d'apprentissage plus visible en mettant en évidence le raisonnement, le jugement et les révisions effectuées. Ils sont également plus résistants à l'automatisation par l'IA que les dissertations portant sur des sujets génériques.

Activité 5 (niveau avancé, sans programmation) : Ouvrez [l'application de déclaration éthique](#) ainsi que son éditeur principal. Demandez à la clientèle apprenante

d'observer le flux de travail visuel. Elle pourra ainsi voir concrètement comment des consignes formulées en langage naturel produisent la déclaration finale.

Vous pouvez également inviter les personnes apprenantes à proposer leurs propres idées d'applications et utiliser Opal pour les créer à partir de zéro. Il peut s'agir d'une idée aussi simple que :

« Je souhaite créer une application d'une seule page intitulée *Planificateur de budget d'épicerie* qui me permet de suivre ma liste d'achats tout en calculant le coût total en temps réel afin de respecter une limite budgétaire hebdomadaire de 120 \$.

05

Expliquer la complaisance de l'IA

Travaillez avec la clientèle apprenante afin d'analyser la manière dont les LLM privilégient parfois l'utilité perçue au détriment de l'exactitude. Demandez-lui de repérer les situations où l'IA ne fait que reformuler ou confirmer la requête de l'utilisatrice ou de l'utilisateur plutôt que de fournir une critique rigoureuse et objective.

Conseils pédagogiques

- Demandez à la clientèle apprenante de présenter à l'IA une opinion très affirmée, même si elle est incorrecte. Invitez-la ensuite à repérer la phrase précise où l'IA cesse de remettre en question l'affirmation et commence à acquiescer trop facilement.
- Recherchez les formulations excessivement apologétiques ou valorisantes (par exemple : « Vous soulevez un point important qui est souvent négligé... ») et discutez de la façon dont ce type de langage peut masquer un manque de substance.
- Analysez les situations où l'IA ignore des faits locaux ou des nuances culturelles particulières parce que la requête de l'utilisatrice ou de l'utilisateur suggérait un récit différent.

Mettre en place un processus de contrôle de la qualité avec les LLM

Puisque les LLM sont conçus pour « compléter » un texte plutôt que pour « vérifier » l'exactitude de l'information, la clientèle apprenante devrait développer son propre processus de contrôle de la qualité. Voici quelques conseils :

- Faites un renvoi des réponses. Soumettez la même requête à un deuxième LLM différent. Si les modèles produisent des réponses divergentes, cela peut indiquer que l'information mérite une vérification plus approfondie.
- Ne vous limitez pas à la réponse de l'IA. Consultez également d'autres sources dans des onglets distincts de votre navigateur.
- Vérifiez toute affirmation audacieuse ou surprenante à l'aide d'au moins trois sources indépendantes et fiables, comme des sites gouvernementaux, des revues universitaires ou des médias reconnus.
- Utilisez des requêtes de sécurité. Demandez au modèle de répondre « Je ne sais pas » plutôt que de deviner lorsqu'il n'est pas certain de l'information.
- Définissez des limites claires. Par exemple : « Utilise uniquement de l'information publiée cette année » ou « Utilise uniquement l'information provenant de cet organisme ».

- Ancrez le modèle dans le contenu à analyser. Fournissez-lui le texte ou les liens précis sur lesquels vous souhaitez qu'il se base.

Vous pouvez également inviter la clientèle apprenante à observer

comment le ton du modèle change lorsqu'elle modifie la température ou lui demande de fournir des réponses « plus assurées » ou « plus créatives ». Cela l'aide à comprendre que le style et le degré d'assurance sont en partie des choix de conception et non des indicateurs de vérité.

L'IA peut générer du texte, mais elle ne peut pas remplacer l'apprentissage

Vous souvenez-vous des Coles Notes?

Si vous avez fréquenté l'école au Canada à la fin du XXe siècle, ces petits guides jaunes et noirs vous sont peut-être familiers. Les *Coles Notes* proposaient, par exemple, une version abrégée de *Hamlet*, mais ces notes ne remplaçaient pas le travail de lecture, d'interprétation, de rédaction ou d'analyse effectué par la lectrice ou le lecteur.

L'IA est différente parce qu'elle peut produire l'ébauche à votre place. Même si le résultat exige encore des révisions, un outil qui génère l'ébauche peut passer du rôle de complément à celui de raccourci qui contourne le processus de réflexion. Dans les deux cas, il s'agit d'une forme de délestage cognitif.

La bande dessinée suivante illustre cette évolution :



Une [expérience sur le terrain](#) a révélé que l'utilisation sans restriction de GPT-4 améliorerait le rendement à court terme lors d'exercices de mathématiques, mais que lorsque l'accès à l'IA était retiré, ces élèves obtenaient par la suite de moins bons résultats que ceux qui n'avaient jamais utilisé l'IA. Une version encadrée du tutoriel améliorait le rendement durant les exercices tout en évitant en grande partie cette perte d'apprentissage.

Les gains de rendement rapportés étaient considérables :

une amélioration de 48 %

chez les élèves utilisant le tutoriel GPT-4 sans restriction;

une amélioration de 127 %

chez celles et ceux utilisant la version encadrée.

Mais le véritable constat est apparu plus tard. Lorsque l'accès à l'IA a été retiré, les élèves qui avaient utilisé la version sans restriction ont obtenu de moins bons résultats que ceux qui n'avaient jamais utilisé l'IA. Autrement dit, sans balises, certaines élèves ont pu s'appuyer sur GPT-4 pendant les exercices, ce qui a ensuite entraîné une baisse de leur rendement autonome (17 %).

Le tutoriel encadré a permis d'éviter cet effondrement. Sa conception incitait les élèves à rester dans le processus cognitif plutôt que de s'en décharger. La conclusion des autrices et auteurs suggère ceci : **l'IA peut améliorer le rendement, mais les choix de conception déterminent si elle renforce ou affaiblit l'apprentissage.**

John Nosta ([Psychology Today](#), 2026) utilise cette tendance pour illustrer la manière dont l'IA peut modifier le lieu où se produit la réflexion. Comme il l'explique, le LLM « réduit l'effort et fournit une réponse... en déplaçant une partie du processus de réflexion hors de la personne », ce qui signifie que la tâche est accomplie, mais que la structure cognitive qui sous-tend cette réussite est moins développée.

Nosta met en évidence deux pratiques qui contribuent à transformer l'IA d'un raccourci à celui de moteur d'apprentissage :

- **Interaction itérative** : un échange avec l'IA qui favorise la compréhension au lieu de la contourner;
- **Approche centrée sur la personne** : des ressources générées sur-le-champ, adaptées aux intérêts et aux besoins de la personne aux études.

[Good Things Foundation \(2024\)](#) met l'accent sur une utilisation de l'IA soutenue, réfléchie et personnalisée dans l'apprentissage des adultes. Ces trois perspectives convergent vers le même message : l'IA peut améliorer le rendement, mais seule une conception réfléchie permet de préserver l'apprentissage.

Ce que nous enseigne la théorie de l'apprentissage

La théorie de l'apprentissage permet de faire le lien entre ce que l'IA ne peut pas faire et les raisons pour lesquelles l'écriture demeure une activité cognitive si importante.

Béhaviorisme

1



L'IA ne peut pas remplacer la pratique

Le béhaviorisme met l'accent sur la répétition et la rétroaction pour maîtriser une compétence. L'IA ne peut pas pratiquer à votre place; les habitudes se construisent grâce à la répétition personnelle.

Le béhaviorisme nous enseigne que les compétences se développent par la répétition, le renforcement et la rétroaction. La maîtrise est renforcée lorsqu'une personne effectue elle-même le comportement.

L'IA peut générer du texte, mais elle ne peut pas faire la pratique à la place de la personne apprenante. Elle ne peut pas construire des habitudes ni renforcer une compétence par la compréhension au nom de quelqu'un d'autre.

Cognitivisme



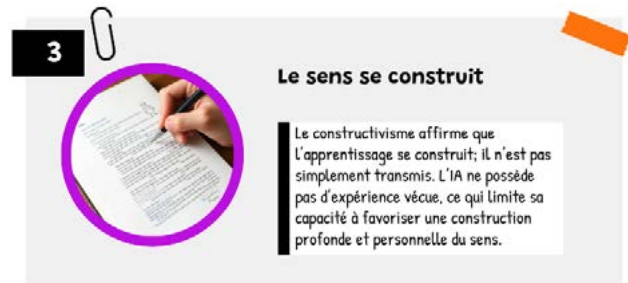
Lorsque les personnes apprenantes écrivent, elles activent leurs connaissances antérieures, organisent leurs idées et construisent des schémas. La friction cognitive qu'elles vivent — comme le fait de s'arrêter, de séquencer, de réviser et de remarquer les lacunes — fait partie du processus par lequel la compréhension se forme.

L'IA peut produire des réponses qui ressemblent à des concepts en détectant des motifs dans les données, mais elle ne peut pas intégrer de nouvelles informations à l'expérience vécue ni construire le schéma interne de la personne apprenante. Elle peut soutenir la cognition en suscitant la réflexion ou en offrant des explications, mais elle ne peut pas accomplir le travail interne de l'apprentissage durable à la place de l'apprenante ou l'apprenant.

Cela est important, car la mémoire de travail est limitée, et que la mémoire à long terme ne devient utile que lorsque les connaissances sont stockées et récupérables en temps réel. Lorsque les personnes apprenantes confient la récupération de l'information à l'IA, elles risquent de contourner la consolidation qui permet de mobiliser ces connaissances en situation de stress.

Si l'évaluation mesure seulement la capacité d'une personne à produire une bonne réponse, l'IA rend cette réussite plus facile à simuler. Si l'évaluation évalue plutôt si la personne a construit l'infrastructure cognitive nécessaire pour réfléchir avec ce qu'elle sait, la tâche est différente. La solution consiste à concevoir des évaluations qui rendent la réflexion visible, et non seulement le résultat final bien rédigé.

Constructivisme



Le constructivisme nous rappelle que l'apprentissage n'est pas transmis, mais construit. La clientèle apprenante construit le sens à partir de la pertinence, de l'expérience vécue, du dialogue et de la communauté. La compréhension s'enrichit à mesure qu'elle relie de nouvelles idées à ce qu'elle sait déjà, met ses interprétations à l'épreuve de sa propre vie et négocie le sens avec les autres.

Ce travail est lent, relationnel et lié à l'identité, et c'est précisément là que l'IA présente des limites évidentes. L'IA ne possède ni expérience vécue, ni identité, ni perspective, ni intérêt personnel. Elle peut imiter une voix, mais elle ne peut pas s'appuyer sur la mémoire, l'émotion, la culture ou le contexte comme le font les êtres humains lorsqu'ils construisent du sens.

La voix, la perspective et la capacité de se situer dans un argument sont des outils essentiels pour les adultes qui évoluent dans le monde du travail, les relations et la vie civique. Elles se construisent par l'acte de choisir une position, de la défendre, de la réviser et de comprendre pourquoi elle importe.

Lorsque l'IA génère une ébauche trop tôt, elle risque d'influencer la position de la personne apprenante avant que celle-ci ne l'ait pleinement définie. Elle peut compléter la phrase « Je pense que... » avant que l'apprenante ou l'apprenant n'ait pleinement développé sa réflexion. Une personne apprenante qui n'a jamais eu à défendre un argument a manqué un moment clé où le sens se construit.

Le texte peut sembler soigné, mais la personne n'aura peut-être pas construit l'architecture interne qui soutient le raisonnement, la communication et

l'action dans le monde.

EduTopia souligne que l'écriture authentique émerge des relations, de la collaboration et des réflexions personnelles. Il s'agit de processus humains — sociaux, émotionnels et contextuels — que l'IA ne peut pas entièrement automatiser.

L'IA peut faciliter la construction du sens en proposant des pistes de réflexion ou des exemples, mais elle ne peut pas remplacer la construction de la compréhension par la personne. Sans pertinence personnelle ni interprétation, la construction du sens est affaiblie.

Lorsque nous appliquons la théorie de l'apprentissage à la littératie de l'IA, le constat est clair

BÉHAVIORISME

→ L'IA ne peut pas remplacer les séances d'entraînement répétitives que permettent d'acquérir les compétences.

COGNITIVISME

→ L'IA ne peut pas remplacer la réflexion cognitive qui permet d'approfondir la compréhension.

CONSTRUCTIVISME

→ L'IA ne peut pas remplacer les processus sociaux et expérientiels qui permettent de donner du sens.

Pourquoi l'écriture demeure importante

La théorie de l'apprentissage permet de mieux comprendre pourquoi l'écriture reste importante.

Le comportementisme montre que l'écriture développe les compétences par la pratique. Le cognitivisme montre qu'elle approfondit la compréhension grâce à l'effort mental. Le constructivisme montre qu'elle aide la clientèle apprenante à donner du sens à ses connaissances par un engagement personnel et social.

L'IA peut générer du texte, mais les personnes apprenantes doivent toujours pratiquer, réfléchir, réviser et interpréter. Ce sont ces étapes qui conduisent à un apprentissage durable.



C'est pourquoi l'écriture demeure essentielle. L'écriture n'est pas simplement un produit à terminer; c'est un processus par lequel l'apprentissage se réalise.

L'IA apporte la rapidité, l'accès et l'assistance. La personne apprenante apporte l'effort, l'interprétation et le sens.

L'apprentissage dépend de la personne, et c'est quelque chose qu'aucune technologie ne peut remplacer.

Quel outil sert à quoi?

Les outils d'IA remplissent des fonctions différentes. Certains sont particulièrement utiles pour la rédaction, d'autres pour la recherche, d'autres encore pour la productivité intégrée ou la synthèse de contenus plus complexes.

Quel outil d'IA répond le mieux à vos besoins actuels?

Capturez cet écran!

 IA à usage général : Copilot, Gemini, ChatGPT	 Perplexity AI*
 Rédaction et remue-ménages Idéale pour les tâches créatives et conversationnelles	 Fonctionne comme un accélérateur axé sur la recherche Idéal pour recueillir des données probantes, identifier les tendances et obtenir des statistiques
 Produit des réponses conversationnelles fondées sur un entraînement général Exige le jugement de l'utilisatrice ou de l'utilisateur pour vérifier les faits	 Fournit des réponses étayées par des liens vers les sources Permet de vérifier l'exactitude des informations grâce aux références
 Rédiger, expliquer et approfondir des concepts complexes Idéal pour la création de contenu itérative	 Spécialisé dans la recherche concurrentielle et l'information difficile à trouver Adapté aux recherches fondées sur des données probantes
 IA intégrée à Microsoft 365 et Google Workspace	 Claude
 Copilot vous accompagne directement au sein des applications Microsoft 365. Intégré à Outlook, Word, Excel et PowerPoint. Gemini prend en charge Google Workplace Gemini vous accompagne directement dans Google Workplace. Intégré à Gmail, Documents, Feuilles de calcul et Présentations	 Fonctionne indépendamment des suites bureautiques Axé sur le travail structuré et les processus de raisonnement à plusieurs étapes
 Utilisez-le pour rédiger, résumer, créer des diapositives, analyser des données et suggérer les prochaines étapes	 Gère de grands volumes de contexte et améliore la qualité du code Utile pour créer des outils, organiser des documents et rédiger des propositions

* Perplexity agit comme un accélérateur de recherche en fournissant des réponses rapides accompagnées de liens vers les sources. Bien qu'utile, cet outil n'est pas infaillible. Les références citées permettent aux utilisatrices et utilisateurs de vérifier l'exactitude de l'information et de considérer les réponses comme des indications éclairées plutôt que comme une autorité définitive.

Cartographier les cas d'utilisation, les responsabilités et les risques liés à l'IA pour la clientèle apprenante

Les différentes utilisations de l'IA s'accompagnent d'attentes différentes. Les activités de remue-ménages créatives permettent davantage de souplesse, tandis que les tâches axées sur

la recherche ou l'automatisation exigent une vérification plus rigoureuse.

L'illustration ci-dessous montre que plus une tâche exige de précision, plus la validation attentive devient importante.

● IA à usage général (faible risque)

Cas d'utilisation : *clarification de textes, rédaction d'ébauches de réponses, partenaire de remue-ménages.*

Ces tâches figurent dans la partie supérieure du schéma parce qu'elles sont génératives et collaboratives. L'utilisatrice ou l'utilisateur recherche des variantes, des modifications de ton ou un point de départ. Son rôle ressemble alors à celui d'une directrice ou d'un directeur de création.

● Outils à contexte étendu (risque moyen)

Cas d'utilisation : *élaboration de propositions, résumé de longues discussions, structuration de documents.*

Ces tâches exigent que l'IA assimile de grandes quantités d'information afin de les synthétiser ou de les organiser. Le niveau de risque est moyen, car si l'IA interprète mal un document volumineux, l'utilisatrice ou l'utilisateur risque à son tour de transmettre une information erronée.

Son rôle évolue alors de directrice ou directeur de création vers celui de révision ou de vérification des faits.

● Outils axés sur la recherche et l'IA intégrée : création d'outils simples (risque élevé)

Cas d'utilisation : *soutien à la recherche, création de flux de travail, rédaction d'ébauches de réponses, partenaire de remue-ménages.*

Lorsque l'IA est utilisée pour recueillir des données probantes, le risque est élevé, car les LLM peuvent inventer des faits, des références et des sources qui semblent plausibles.

Dans un contexte axé sur la recherche, l'outil agit comme une boussole. La personne utilisatrice doit demeurer le point d'ancrage ultime.

Lorsque l’IA est intégrée à des flux de travail afin d’automatiser des tâches répétitives, le risque demeure élevé en raison des enjeux liés à la confidentialité des données, à la gouvernance et au risque d’accumulation d’erreurs à grande échelle.

L’illustration met en évidence plusieurs mesures de protection essentielles. Dans ce contexte, l’utilisatrice ou l’utilisateur agit comme responsable de conformité.

Cas d'utilisation de l'IA et responsabilités			
Cas d'utilisation	À quoi sert l'IA?	Responsabilités de la personne utilisatrice	Niveau de risque
 Clarification de textes	Simplifier des messages complexes	Confirmer que le sens, le ton et l'intention demeurent exacts	Faible
 Rédaction de réponses	Générer de premières ébauches pour accélérer la rédaction	Réviser l'exactitude, la pertinence et le contexte	Faible
 Partenaire de remue-méninges	Explorer des options, variantes et pistes créatives	Sélectionner, affiner et assumer l'ultime version	Faible
 Développement d'idées	Produire des ébauches structurées selon des pratiques reconnues	Vérifier les affirmations; s'assurer qu'elles correspondent aux objectifs et au public cible	Moyen
 Résumé de fils de discussion	Résumer de longues chaînes de courriels ou de longs documents	Comparer le résumé au contenu d'origine	Moyen
 Structuration de documents	Organiser des propositions, rapports et plans	Vérifier la logique, l'exactitude et l'exhaustivité avant publication	Moyen
 Soutien à la recherche	Dégager des liens, tendances et perspectives multilingues	Vérifier les faits et les sources et les citer correctement	Élevé
 Création d'outils simples	Automatiser des tâches répétitives à faible risque	Éviter le traitement de données privées et ne jamais automatiser des décisions concernant des personnes	Élevé

Un nouveau cadre à explorer

Le [cadre de l'Association pour la supervision et l'élaboration des programmes scolaires \(ASCD\)](#) souligne qu’une personne diplômée prête à évoluer dans un monde où l’IA est omniprésente est capable de :

- **Apprendre avec l’IA**, tout en continuant à se fixer des objectifs, à rechercher de la rétroaction et à demeurer responsable de son propre développement;

- **Faire de la recherche avec l’IA**, tout en évaluant les affirmations, en comparant les sources et en repérant les contradictions;
- **Synthétiser l’information avec l’IA**, tout en choisissant le niveau de détail, le format et le contexte appropriés pour le public visé;
- **Résoudre des problèmes avec l’IA**, tout en générant des idées, en explorant différentes perspectives et en surmontant les blocages créatifs;
- **Collaborer avec l’IA**, tout en favorisant la collaboration humaine, en réduisant les barrières linguistiques et en coordonnant le travail des équipes;
- **Raconter des histoires avec l’IA**, tout en façonnant les récits, les éléments visuels et les messages avec intention et intégrité.

À travers ces six rôles, ainsi que les concepts et les cadres présentés dans ce bulletin, le message demeure le même : **l’IA augmente les capacités humaines, mais elle ne remplace pas la responsabilité humaine.**

La question la plus importante n’est pas de savoir quel outil est le plus intelligent, mais plutôt quel outil convient le mieux à la tâche, au niveau de responsabilité et au degré de risque qui l’accompagnent.

De la théorie à la pratique : expérimenter l’IA de façon intentionnelle

Google OPAL

Enfin, nous vous avons promis des Gemini Gems, et nous tenons parole.

Les Gemini Gems sont des versions personnalisées de Google Gemini qui utilisent des consignes précises et des fichiers de connaissances chaque fois qu’elles génèrent une réponse. Par exemple, un Gem expert en « Trouver de l’information NCLC 3 de la CAFCO » pourrait inclure des renseignements détaillés sur le public cible, des exemples de tâches, des lignes



directrices du curriculum ainsi que le ton souhaité pour les réponses. La clientèle apprenante pourrait ainsi examiner des textes et en extraire l'information sans avoir à reformuler les mêmes consignes à chaque utilisation.

Pour conclure ce bulletin, nous mettons en lumière une ressource qui montre ce qui se produit lorsque l'IA cesse d'être une boîte noire et devient un outil que l'on peut ouvrir, examiner et même utiliser à des fins pédagogiques. Ce guide vous accompagne dans la découverte de **Google Opal**, le moteur qui alimente un type particulier de Gem.

Ce qui rend cette ressource particulièrement intéressante, c'est qu'elle ne se contente pas de montrer ce qu'Opal est capable de faire, mais qu'elle explique à la clientèle apprenante **comment le système fonctionne**.

Vous y découvrirez :

- Comment une machine fonctionne à l'aide de flux de travail;
- Les instructions suivies par l'IA;
- Comment vos données d'entrée influencent les résultats obtenus;
- Comment les personnes apprenantes peuvent cliquer sur « **Éditeur** » afin d'observer l'ensemble de la chaîne logique qui mène à la réponse finale.

Il s'agit d'une façon concrète et visuelle d'enseigner la littératie de l'IA, la réflexion en termes de flux de travail et l'utilisation responsable de l'intelligence artificielle.

Si vous recherchez une ressource qui aide la clientèle apprenante à passer de « L'IA me fournit des réponses » à « Je comprends comment cet outil fonctionne », téléchargez dès maintenant [À la découverte de Google Opal!](#)



Dans cette édition, nous avons exploré les mécanismes de l'apprentissage ainsi que la façon dont l'IA façonne et transforme ces processus.

Dans nos prochaines éditions, nous aborderons une conversation plus approfondie soulevée par notre communauté : **l'IA et le pouvoir**.

Par pouvoir, nous ne parlons pas d'électricité, mais bien de contrôle.

Nous examinerons qui possède et construit les infrastructures qui soutiennent ces outils numériques, quelles visions du monde sont amplifiées par leurs réponses et quels intérêts politiques ou économiques ces systèmes servent ultimement.

Nous espérons que cette ressource aura soutenu votre enseignement et votre réflexion, et nous avons très hâte de partager avec vous la suite.

Restez à l'affût de nos prochains balados!

Références et attributions

Bastani, H., Bastani, O., Sungu, H., Ge, Ö., Kabakci, R., & Mariman, R. (2024, June 25). *Generative AI without guardrails can harm learning: Evidence from high school mathematics*. PNAS, 122. <https://doi.org/10.1073/pnas.2422633122>

Celik, I., Zabolotna, K., & Viberg, O. (16 février 2026). Knowledge construction with GenAI: The role of theory-informed prompt engineering in achieving pedagogical alignment. *Innovations in Education and Teaching International*, 1–19. <https://doi.org/10.1080/14703297.2026.2631586>

Floridi, L., Morley, J., Novelli, C., and Watson, D. (10 décembre 2025). What kind of reasoning (if any) is an LLM actually doing? On the Stochastic Nature and Abductive Appearance of Large Language Models. *ArXiv*. <https://arxiv.org/abs/2512.10080>

Hudd, S. S., et al. (7 mars 2025). Reading, writing, and thinking in the age of AI. *Faculty Focus*. <https://www.facultyfocus.com/articles/teaching-with-technology-articles/reading-writing-and-thinking-in-the-age-of-ai/>

Ritter, L. A. (février 2026). The impact of AI tools on adult learning outcomes in online education. *International Journal of Research and Review*, Volume 13; issue 2 <https://doi.org/10.52403/ijrr.20260213>

Roberts, J. (25 février 2025). Authentic Writing in the Age of AI. *Edutopia*. <https://www.edutopia.org/article/teaching-authentic-writing-age-ai/>

Sedita, J. (9 mai 2025). Writing instruction in the age of AI. *Keys to Literacy*. <https://keystoliteracy.com/blog/writing-instruction-in-the-age-of-ai/>

Stöffelbauer, A. (24 octobre 2023). *How Large Language Models Work: From Zero to ChatGPT*. Data Science at Microsoft, Medium. <https://medium.com/data-science-at-microsoft/how-large-language-models-work-91c362f5b78f>

University of Calgary Library. (9 mars 2026). Artificial intelligence: AI literacy. *Libraries and Cultural Resources, University of Calgary*. <https://libguides.ucalgary.ca/c.php?g=733971andp=5278508>

Ressources françaises

Bédard, Grégoire. (16 juin 2025). *Écrire et penser à l'ère de l'IA*. Codex Numeris. <https://codexnumeris.org/42-ecrire-et-penser-a-lere-de-lia/>

Conseil supérieur de l'éducation et Commission de l'éthique en science et en technologie (2024). *Intelligence artificielle générative en enseignement supérieur : enjeux pédagogiques et éthiques*, Québec, Le Conseil; La Commission, 113 p. <https://www.cse.gouv.qc.ca/wp-content/uploads/2024/04/50-0566-RP-IA-generative-enseignement-superieur-enjeux-ethiques.pdf>

Cuerrier, Marjorie. (mai 2025). *La didactique de l'écriture à l'ère de l'intelligence artificielle : transformations, enjeux et formation*. Colloque international en éducation 2025. https://www.researchgate.net/publication/391735141_La_didactique_de_l'ecriture_a_l'ere_de_l'intelligence_artificielle_transformations_enjeux_et_formation

Jenni. (27 février 2024). *L'impact de l'IA sur l'écriture : efficacité, précision et au-delà*. <https://jenni.ai/fr/artificial-intelligence/writing-benefits>

Scandolera, Franck. (2025). Comprendre le fonctionnement des LLM : guide complet pour tous. https://www.youtube.com/watch?v=wrZFyKfJN_Y

Technologies éducatives pour l'enseignement et l'apprentissage, Actes du colloque ROC 2025, *Vers des formation numériques critiques et émancipatrices*, Québec, 312 p. <https://r-libre.telug.ca/3946/1/Actes%20du%20Colloque%20ROC2025.pdf>

Tishcoff, Ryan, et al. (2024). Utiliser l'IA générative pour un apprentissage plus accessible : point de vue des étudiant.es et du personnel de l'enseignement postsecondaire de l'Ontario. Conseil ontarien de la qualité de l'enseignement supérieur. <https://heqco.ca/wp-content/uploads/2024/11/GenAI-Report-FORMATTED-FR.pdf>

TrueFoundry. *Gemini 3 contre Kimi-K2 Thinking contre Grok-4.1 contre GPT-5.1 : qui gagne réellement le dernier examen de l'humanité?* <https://www.truefoundry.com/fr/blog/gemini-3>