

Le Paradoxe éducatif de l'IA

INSA Lyon 2026 — E. JOYEUX

Table des matières

Preface	5
Chapitre 1 : Dans la tête de nos élèves : comment apprennent-ils à l'ère de l'IA ?	6
Glossaire global du chapitre	7
Introduction	9
Plan global du chapitre : trois mouvements	9
Repère : Activité pédagogique et activité cognitive	9
Carte de compréhension : fil directeur du chapitre	10
Activité : Auto-test de sensibilisation	10
Que peut-on vraiment attendre des neurosciences en pédagogie ?	12
Situation ESR	12
La neuro-éducation comme cadre d'analyse	12
Carte de compréhension : de la science à la décision pédagogique	13
Quand le cours est clair mais que les étudiants décrochent	15
Situation en école d'ingénieurs	15
Mémoire de travail et charge cognitive	15
Outil : Diagnostiquer la surcharge d'un support de cours	15
Carte de diagnostic : charge, traitement et consolidation	16
Pourquoi prendre des notes n'est pas une compétence naturelle ?	18
Situation d'amphi	18
Prise de notes, encodage et stockage externe	18
Tableau : Ce que fait l'étudiant / ce que cela produit	18
Outil : Consignes anti-verbatim	19
Papier, clavier, tablette : le support compte moins que l'activité cognitive	20
Dépasser le faux débat	20
Matrice : Support, risque, usage fécond	20
Carte d'action : choisir un support selon l'activité cognitive	21
Pourquoi les méthodes « cerveau-compatible » séduisent autant ?	22
Neuromythes et neuro-enchantement	22
Carte de compréhension : cycle de formation d'un neuromythe	22
Tableau : Mythe, promesse, alternative rigoureuse	23
Quand manipuler aide vraiment à comprendre	25
Cognition incarnée	25
Exemple : TP de mécanique	25
Checklist : Un TP est-il cognitivement incarné ?	25
Carte de diagnostic : de la manipulation au transfert	26
Et si l'erreur était une trace de raisonnement ?	27
Métacognition et inhibition	27
Exemple : Mécanique des fluides ou analyse de données	27
Routine STOP pour entraîner l'inhibition	27
Carte de diagnostic : analyser une erreur comme trace de raisonnement	28

Quand l'étudiant produit sans avoir appris	30
IA, externalisation cognitive et désynchronisation	30
Situation : Projet Python avec IA	30
Matrice : Production externe / compréhension interne	30
Carte d'action : décider quand et comment autoriser l'IA	31
Outil : Trois traces à demander lorsque l'IA est autorisée	31
Encadré : Qu'est-ce qu'une preuve d'apprentissage ?	32
Rendre l'apprentissage visible, guidé et vérifiable	33
Pédagogie explicite et inclusive	33
Fiche synthèse : 8 repères pour enseigner à l'ère de l'IA	33
Avant / Après : Transformer une séance sans perdre l'exigence	34
Carte d'action : concevoir une séquence plus explicite	34
Étude de cas : Transformer un cours d'algorithmique	35
Situation initiale	35
Diagnostic	35
Transformation proposée	35
Carte d'action : transformer un cours dense en parcours vérifiable	36
Ce que la neuro-éducation ne peut pas promettre	37
Limites et controverses	37
Tableau : Permet / ne permet pas / oblige à discuter	37
Carte de compréhension : la neuro-éducation en dialogue	37
Controverses à garder ouvertes	38
IA et apprentissage	38
Évaluation et robustesse cognitive	38
Égalité de traitement et équité d'apprentissage	38
Auto-positionnement final	39
Fiche action : ce que je peux faire dès demain	40
Conclusion	41
Bibliographie indicative	42
Mémoire, charge cognitive et consolidation	42
Prise de notes et apprentissage	42
Neuromythes et neuro-éducation	42
Métacognition, erreur et autorégulation	42
Cognition incarnée, manipulation et transfert	42
IA, apprentissage et enseignement supérieur	43
Chapitre 2 : Pourquoi l'IA semble comprendre ?	44
Glossaire global du chapitre	45
Introduction	47
Continuité explicite avec le chapitre 1	47
Plan global du chapitre : trois mouvements	47
Carte de compréhension : du contenu disponible à la compréhension vérifiable	48
Activité : Auto-test de sensibilisation	49
1 — De l'IA symbolique à la fluidité générative : pourquoi la performance devient opaque	50
Situation ESR	50
De l'explicabilité symbolique à l'opacité statistique	50
Les modèles génératifs : cohérence textuelle et illusion de compréhension	50

□ Essentiel : résultat correct ≠ raisonnement transparent	51
Carte : évolution de l'IA et déplacement pédagogique	52
□ Approfondissement : Chain-of-Thought et raisonnement latent	52
2 — Quand l'IA brouille les repères de l'apprentissage	53
Situation ESR	53
Le modèle Clever Hans : réussir pour les mauvaises raisons	53
Illusion de raisonnement : quand la réponse finale fait écran	53
Anthropomorphisme et neuromythes : l'IA comme nouveau langage d'anciennes croyances	54
□ Essentiel : quatre confusions à éviter	54
Carte : de la fluidité à l'illusion de compétence	55
□ Recommandé : activité "trouver le point de rupture"	55
□ Approfondissement : misconceptions et post-vérité	56
3 — Que reste-t-il à apprendre lorsque l'on peut produire sans comprendre ?	57
Situation ESR	57
Externalisation cognitive et dette cognitive	57
Performance sans effort et expertise fragile	57
Réintroduire l'effort utile	58
□ Essentiel : cerveau d'abord, IA ensuite	58
Carte : de la production assistée à l'apprentissage robuste	59
□ Recommandé : stratégie 3P	59
□ Approfondissement : l'IA comme miroir cognitif	60
Carte d'aide à la décision pédagogique	61
Étude de cas : transformer une évaluation de rapport assisté par IA	62
Situation initiale	62
Diagnostic	62
Transformation proposée	62
Grille d'évaluation rapide	62
Fiche action : ce que je peux faire dès demain	64
Auto-positionnement final	65
Analyse de votre positionnement	65
Lecture par dimensions	66
Interprétation pédagogique	66
Pour passer à l'action	66
À retenir	67
Conclusion	68
Bibliographie indicative	69
Continuité avec le chapitre 1 : apprentissage, charge cognitive, prise de notes, métacognition	69
Fonctionnement des modèles, opacité et illusion de compréhension	69
Anthropomorphisme, misconceptions et neuromythes	70
Dette cognitive, apprentissage, évaluation et supervision critique	70

Preface

Code de lecture du manuel

- Essentiel À lire pour stabiliser les notions centrales.
- Transformer sa pratique À mobiliser pour diagnostiquer une séance, ajuster une consigne ou modifier une évaluation.
- Approfondissement À consulter pour discuter les limites, les controverses et les prolongements.

Chaque grande section suit autant que possible la même logique : **comprendre** l'idée clé, **repérer** les signes observables, puis **agir** par un geste pédagogique réaliste.

Chapitre 1 : Dans la tête de nos élèves : comment apprennent-ils à l'ère de l'IA ?

Neuro-éducation, neuromythes, prise de notes et externalisation cognitive

Transparence éditoriale

Ce chapitre a été conçu, vérifié et éditorialisé sous la responsabilité intellectuelle de l'autrice. Des outils d'intelligence artificielle ont été mobilisés comme appui à la structuration, à la reformulation et à la mise en forme, sans se substituer aux choix scientifiques, pédagogiques et éditoriaux.

Question mobilisatrice

Comment concevoir des situations d'enseignement qui maintiennent un véritable engagement cognitif lorsque les étudiants disposent d'outils capables de produire, résumer, structurer et reformuler à leur place ?

Mots-clés

Neuro-éducation Mémoire de travail Charge cognitive Consolidation Prise de notes Effet verbatim Neuromythes Cognition incarnée Métacognition Inhibition cognitive IA générative Externalisation cognitive Désynchronisation cognitive Robustesse cognitive

Ce chapitre vous aide à :

- distinguer une information disponible d'un apprentissage réellement construit ;
- repérer les situations de surcharge cognitive dans vos cours, TD, TP ou projets ;
- comprendre pourquoi la prise de notes doit être enseignée plutôt que présumée ;
- déconstruire les neuromythes sans rejeter les apports des neurosciences ;
- encadrer les usages de l'IA pour qu'ils soutiennent l'apprentissage au lieu de le court-circuiter ;
- rendre visibles les raisonnements, les erreurs, les stratégies et les transferts.

Glossaire global du chapitre

Ce glossaire donne des repères communs pour lire le chapitre. Il ne vise pas à figer les notions, mais à stabiliser leur sens pédagogique avant d'entrer dans les exemples, les activités et les cartes de décision.

Notion	Définition courte	Question pédagogique associée
Apprentissage construit	Transformation progressive d'une information en savoir mobilisable, explicable et transférable.	L'étudiant peut-il réutiliser ce savoir dans une situation nouvelle ?
Disponibilité du contenu	Accès matériel ou numérique à une information : support, fiche, correction, réponse IA, vidéo, code.	L'accès au contenu a-t-il donné lieu à un traitement actif ?
Production visible	Trace observable produite par l'étudiant : note, devoir, code, synthèse, schéma, exposé.	Cette trace renseigne-t-elle réellement sur le raisonnement construit ?
Preuve d'apprentissage	Indice permettant d'inférer qu'un apprentissage a eu lieu : justification, transfert, explication, correction d'erreur, comparaison.	Que prouve la production, et que ne prouve-t-elle pas encore ?
Activité pédagogique	Ce que l'enseignant demande, organise ou prescrit.	La consigne suffit-elle à provoquer l'activité mentale attendue ?
Activité cognitive	Ce que l'étudiant fait mentalement : sélectionner, comparer, inhiber, reformuler, justifier, transférer.	Quelle opération cognitive est réellement sollicitée ?
Engagement cognitif	Degré d'implication mentale dans une tâche d'apprentissage.	L'étudiant traite-t-il le contenu ou se contente-t-il de le traverser ?
Mémoire de travail	Système limité qui maintient et manipule temporairement les informations nécessaires à une tâche.	La tâche surcharge-t-elle inutilement l'étudiant ?
Charge cognitive	Effort mental imposé par une situation d'apprentissage.	La difficulté vient-elle du concept à apprendre ou de la manière dont la tâche est présentée ?
Consolidation	Stabilisation progressive des apprentissages dans le temps, grâce à la reprise, l'espacement et la mobilisation.	Le cours prévoit-il des occasions de réactivation ?
Effet verbatim	Tendance à recopier ou conserver la forme d'un contenu sans traitement personnel suffisant.	La prise de notes transforme-t-elle le contenu ou le duplique-t-elle ?
Métacognition	Capacité à se représenter, surveiller et réguler sa propre compréhension.	L'étudiant sait-il dire ce qu'il comprend, ce qu'il ne comprend pas et pourquoi ?
Inhibition cognitive	Capacité à suspendre une réponse intuitive mais inadéquate.	La tâche aide-t-elle à résister aux réponses trop rapides ?
Erreur féconde	Erreur utilisée comme trace de raisonnement à analyser, plutôt que comme simple échec.	L'erreur permet-elle de comprendre l'obstacle rencontré ?

Notion	Définition courte	Question pédagogique associée
Cognition incarnée	Idée selon laquelle l'apprentissage engage aussi le corps, l'action, l'espace et les interactions.	Le dispositif soutient-il l'attention, l'action et la manipulation significative ?
Neuromythe	Croyance pédagogique simplifiée ou fausement appuyée sur les neurosciences.	L'argument scientifique est-il suffisamment étayé ?
Externalisation cognitive	Délégation d'une partie du traitement mental à un outil, un support, un pair ou une IA.	L'outil soutient-il le raisonnement ou le remplace-t-il ?
Désynchronisation cognitive	Décalage entre la qualité apparente d'une production et le niveau réel de compréhension de l'étudiant.	Une production correcte peut-elle masquer une compréhension fragile ?
Robustesse cognitive	Capacité à expliquer, critiquer, adapter et transférer un savoir au-delà du cas travaillé.	L'étudiant peut-il tenir son raisonnement quand le contexte change ?
Pédagogie explicite	Démarche qui rend visibles les objectifs, les critères, les stratégies, les erreurs attendues et les attendus de raisonnement.	Les étudiants savent-ils ce qu'ils doivent apprendre à faire mentalement ?

Introduction

Dans l'enseignement supérieur, et plus encore en école d'ingénieurs, apprendre ne consiste plus seulement à recevoir une information claire, ni même à disposer d'un support complet. Un étudiant peut aujourd'hui photographier un tableau, récupérer un diaporama, générer une fiche de synthèse avec une IA, obtenir un code fonctionnel ou produire un résumé parfaitement structuré sans avoir réellement construit le raisonnement correspondant. Ce chapitre part de ce paradoxe : à l'ère de l'IA, les traces visibles du travail (notes, fiches, devoirs, codes, synthèses) deviennent parfois plus propres, plus rapides et plus convaincantes, alors même que l'activité cognitive interne peut rester fragile, partielle ou absente.

Centre de gravité du chapitre : proposer une grille pédagogique pour distinguer disponibilité du contenu, production visible et apprentissage réellement construit à l'ère de l'IA.

La neuro-éducation est ici mobilisée non comme une recette, mais comme une grille de lecture critique pour comprendre ce qui se joue entre exposition à l'information, engagement cognitif, consolidation, erreur, prise de notes, corps, environnement et usages numériques.

Plan global du chapitre : trois mouvements

Le chapitre est organisé en trois mouvements. Cette lecture permet de situer chaque partie dans une progression pédagogique : comprendre les conditions de l'apprentissage, déjouer les fausses évidences, puis concevoir à l'ère de l'IA.

Mouvement	Intention pédagogique	Notions et sections concernées
1 : Comprendre ce qui fait apprendre	Identifier les conditions minimales d'un apprentissage construit : attention, mémoire de travail, consolidation, prise de notes.	Neuro-éducation, charge cognitive, prise de notes.
2 : Déjouer les fausses évidences	Éviter les raccourcis séduisants : support « magique », neuromythes, TP automatiquement efficaces, erreur réduite à un manque.	Supports, neuromythes, cognition incarnée, erreur.
3 : Enseigner à l'ère de l'IA	Concevoir des situations où l'IA soutient le raisonnement sans remplacer l'activité cognitive de l'étudiant.	Externalisation, traces, pédagogie explicite, étude de cas, évaluation.

Repère : Activité pédagogique et activité cognitive

Une **activité pédagogique** correspond à ce que l'enseignant demande, organise ou prescrit.

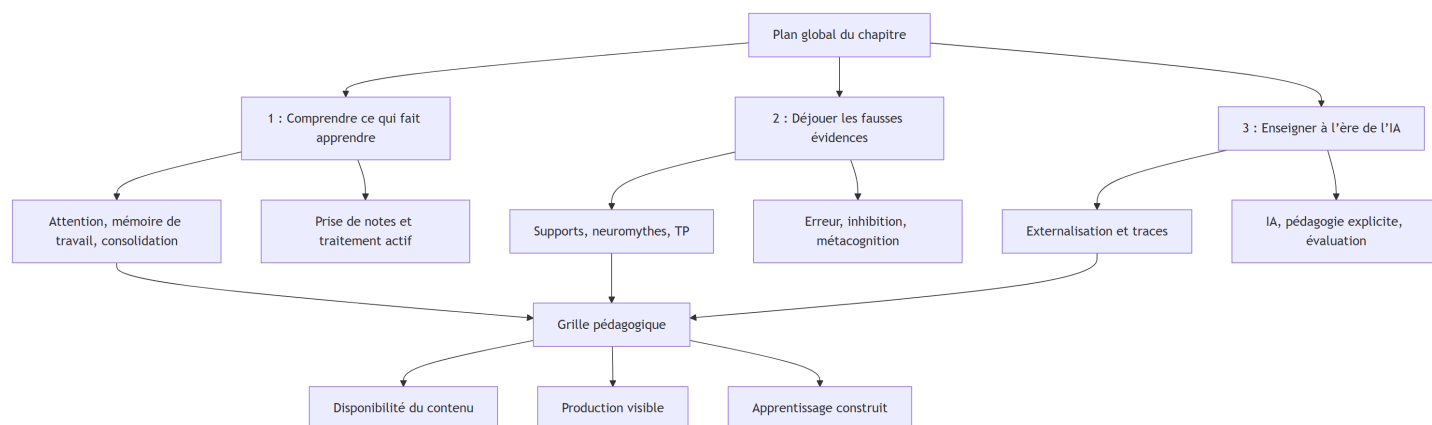
Une **activité cognitive** correspond à ce que l'étudiant fait mentalement avec cette demande : sélectionner, hiérarchiser, reformuler, comparer, inhiber une intuition, justifier, transférer, contrôler une réponse.

C'est un point décisif à l'ère de l'IA : une même consigne peut produire des apprentissages très différents. *Faire un résumé* peut être cognitivement pauvre si l'étudiant copie, paraphrase ou délègue à l'IA. La même tâche devient beaucoup plus riche s'il doit expliquer ses choix, comparer deux formulations, justifier ce qu'il a supprimé et transférer l'idée à un cas voisin.

Une **activité cognitive** ne devient pédagogiquement exploitable que si le dispositif en garde une trace observable : justification, comparaison, reformulation, erreur commentée ou transfert. Cette trace ne prouve jamais tout, mais elle augmente la valeur de ce que l'enseignant peut inférer sur l'apprentissage réel.

Carte de compréhension : fil directeur du chapitre

Mode d'emploi de la carte. Cette carte sert à situer les trois mouvements du chapitre et à relier les notions au problème central : transformer un contenu disponible en apprentissage construit.



Fil directeur du chapitre

La question n'est pas seulement : *les étudiants ont-ils accès au contenu ?* Elle devient : *quelles opérations cognitives ont-ils réellement effectuées pour transformer ce contenu en savoir mobilisable ?*

Activité : Auto-test de sensibilisation

Avant de lire : vos représentations de l'apprentissage

Pour chaque affirmation, répondez spontanément : **plutôt vrai, plutôt faux, cela dépend, je ne sais pas.**

Affirmation

Votre réponse

Un étudiant qui dispose d'une fiche de synthèse claire a déjà une bonne base d'apprentissage.
Prendre des notes est une compétence que les étudiants de l'enseignement supérieur maîtrisent globalement.
L'ordinateur pose problème surtout parce qu'il distrait.
Une IA qui explique bien un concept peut remplacer une partie du travail de compréhension.
Les erreurs en sciences viennent surtout d'un manque de connaissances.
Un TP fait nécessairement mieux apprendre qu'un cours magistral.
Les neurosciences permettent de savoir quelles méthodes pédagogiques fonctionnent.
Les étudiants ont des profils d'apprentissage stables qu'il faut respecter.

Une production propre et argumentée est un bon indicateur de compréhension.

L'égalité de traitement suffit à garantir l'équité d'apprentissage.

À reprendre en fin de chapitre : quelles affirmations vous semblent désormais plus discutables ? Pourquoi ?

Que peut-on vraiment attendre des neurosciences en pédagogie ?

☐ Essentiel ☐ Approfondissement

Dans cette section, vous allez :

Comprendre	Repérer	Agir
Ce que la neuro-éducation peut éclairer sans prescrire mécaniquement.	Les indices de neuro-enchantement, de raccourci ou de promesse trop simple.	Vérifier le niveau de preuve, contextualiser et formuler une décision pédagogique située.

Situation ESR

Dans une formation d'enseignants du supérieur, une collègue affirme vouloir rendre ses cours « plus cerveau-compatibles ». Elle envisage de distinguer les étudiants visuels, auditifs et kinesthésiques, d'ajouter quelques images du cerveau dans ses supports et de proposer des exercices courts « validés par les neurosciences ». L'intention est sincère : mieux prendre en compte les mécanismes d'apprentissage. Mais la question décisive est ailleurs : que peut-on réellement déduire d'un résultat neuroscientifique pour concevoir une situation d'enseignement ?

La neuro-éducation comme cadre d'analyse

La neuro-éducation se situe au croisement des neurosciences cognitives, de la psychologie et des sciences de l'éducation. Elle ne permet pas de transformer directement une observation du cerveau en méthode pédagogique. Elle aide plutôt à mieux formuler certains problèmes : surcharge de la mémoire de travail, illusion de compréhension, rôle de l'erreur, besoin de consolidation, influence du stress, importance de l'engagement actif ou de la régulation métacognitive.

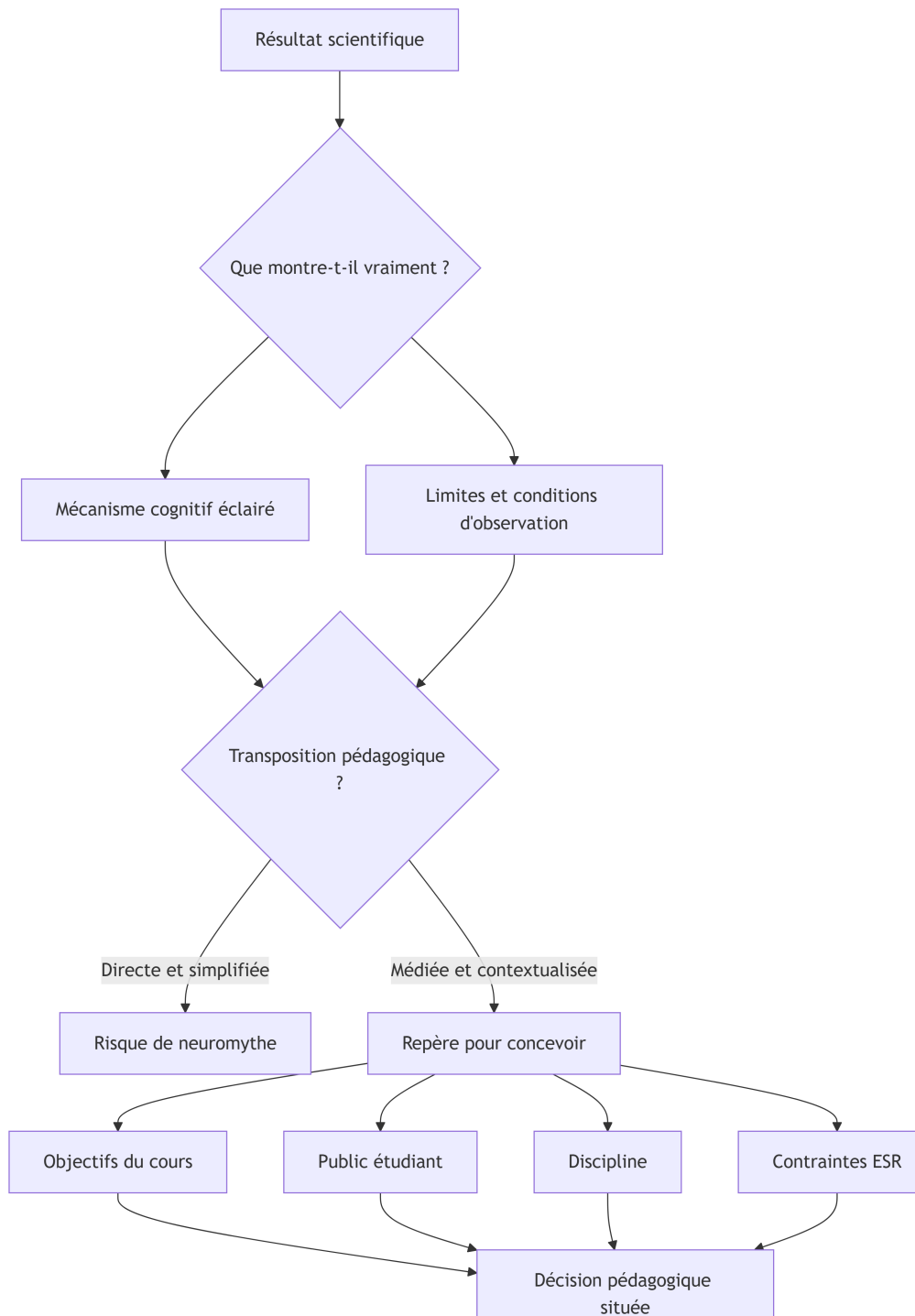
C'est pourquoi elle doit être abordée comme un **champ interdisciplinaire en construction**. Les neurosciences éclairent certains mécanismes, mais elles ne suffisent pas à comprendre une situation d'apprentissage dans toute sa complexité. Un cours, un TD, un TP ou un projet ne mobilise jamais un cerveau isolé. Il mobilise un étudiant situé : avec une histoire scolaire, une fatigue, une confiance plus ou moins grande, des contraintes matérielles, des attentes institutionnelles, un rapport à l'évaluation, des interactions avec ses pairs et désormais des outils numériques capables de produire à sa place.

Deux dérives sont donc à éviter. La première consiste à rejeter entièrement les neurosciences au motif qu'elles seraient trop éloignées de la salle de classe. Ce rejet prive les enseignants de repères utiles pour penser l'attention, la mémoire, l'erreur ou la consolidation. La seconde consiste à appliquer mécaniquement des principes simplifiés : « il faut enseigner selon les styles d'apprentissage », « un schéma cérébral suffit à prouver une méthode », « un exercice corporel améliore automatiquement les performances ». Ces raccourcis transforment des éclairages scientifiques en prescriptions fragiles.

La posture la plus féconde est une posture de **discernement** : accueillir les apports de la recherche, interroger leur niveau de preuve, les mettre en dialogue avec les sciences de l'éducation, puis les contextualiser dans une situation pédagogique précise.

Carte de compréhension : de la science à la décision pédagogique

Mode d'emploi de la carte. Cette carte aide à distinguer résultat scientifique, transposition pédagogique et décision située. Elle sert à éviter les prescriptions trop directes.



Point de vigilance

Une donnée neuroscientifique peut éclairer un choix pédagogique, mais elle ne suffit pas à le justifier. Entre la synapse et l'amphi s'interposent toujours des corps fatigués, des trajectoires sociales, des émotions, des contraintes institutionnelles et des finalités de formation.

À retenir

- La neuro-éducation est une **boussole**, pas un mode d'emploi.
- Elle aide à poser de meilleures questions pédagogiques.
- Elle exige une posture critique : ni rejet, ni fascination.

Pause réflexive

Dans votre contexte, quelle affirmation « validée par les neurosciences » avez-vous déjà entendue ? Que faudrait-il vérifier avant d'en faire une consigne pédagogique ?

Quand le cours est clair mais que les étudiants décrochent

☐ Essentiel ☐ Transformer sa pratique

Dans cette section, vous allez :

Comprendre	Repérer	Agir
Pourquoi un cours clair peut rester cognitivement difficile à traiter.	Les signes de surcharge : copie mécanique, photos sans retraitement, recours compensatoire à l'IA.	Découper, hiérarchiser, ralentir et introduire des pauses de traitement.

Situation en école d'ingénieurs

En amphithéâtre de thermodynamique, l'enseignant déroule une démonstration dense. Les diapositives sont propres, les notations rigoureuses, les transitions logiques. Pourtant, au bout de trente minutes, beaucoup d'étudiants recopient sans suivre. Certains photographient les équations. D'autres demandent ensuite à une IA de « résumer le cours simplement ». Le cours était clair du point de vue de l'enseignant, mais il n'a pas nécessairement été traitable du point de vue de la mémoire de travail des étudiants.

Mémoire de travail et charge cognitive

La mémoire de travail permet de maintenir et de manipuler temporairement des informations pour réaliser une tâche. En situation d'apprentissage, elle est fortement sollicitée : écouter, comprendre, relier aux connaissances antérieures, noter, interpréter un schéma, suivre une démonstration, anticiper une question. Or ses ressources sont limitées. Lorsque trop d'informations nouvelles sont présentées simultanément, ou lorsque l'étudiant doit écouter et transcrire en même temps, la mémoire de travail peut saturer.

La charge cognitive ne dépend donc pas seulement de la difficulté intrinsèque du contenu. Elle dépend aussi de la manière dont le contenu est présenté, du rythme, des supports, du nombre de notations nouvelles, des exemples mobilisés, de la place laissée au traitement et du degré d'autonomie méthodologique attendu. En école d'ingénieurs, cette question est centrale : les contenus sont souvent abstraits, cumulatifs et fortement symboliques. Une même séance peut exiger de manipuler des équations, des modèles, des unités, des graphes, des hypothèses et des interprétations physiques.

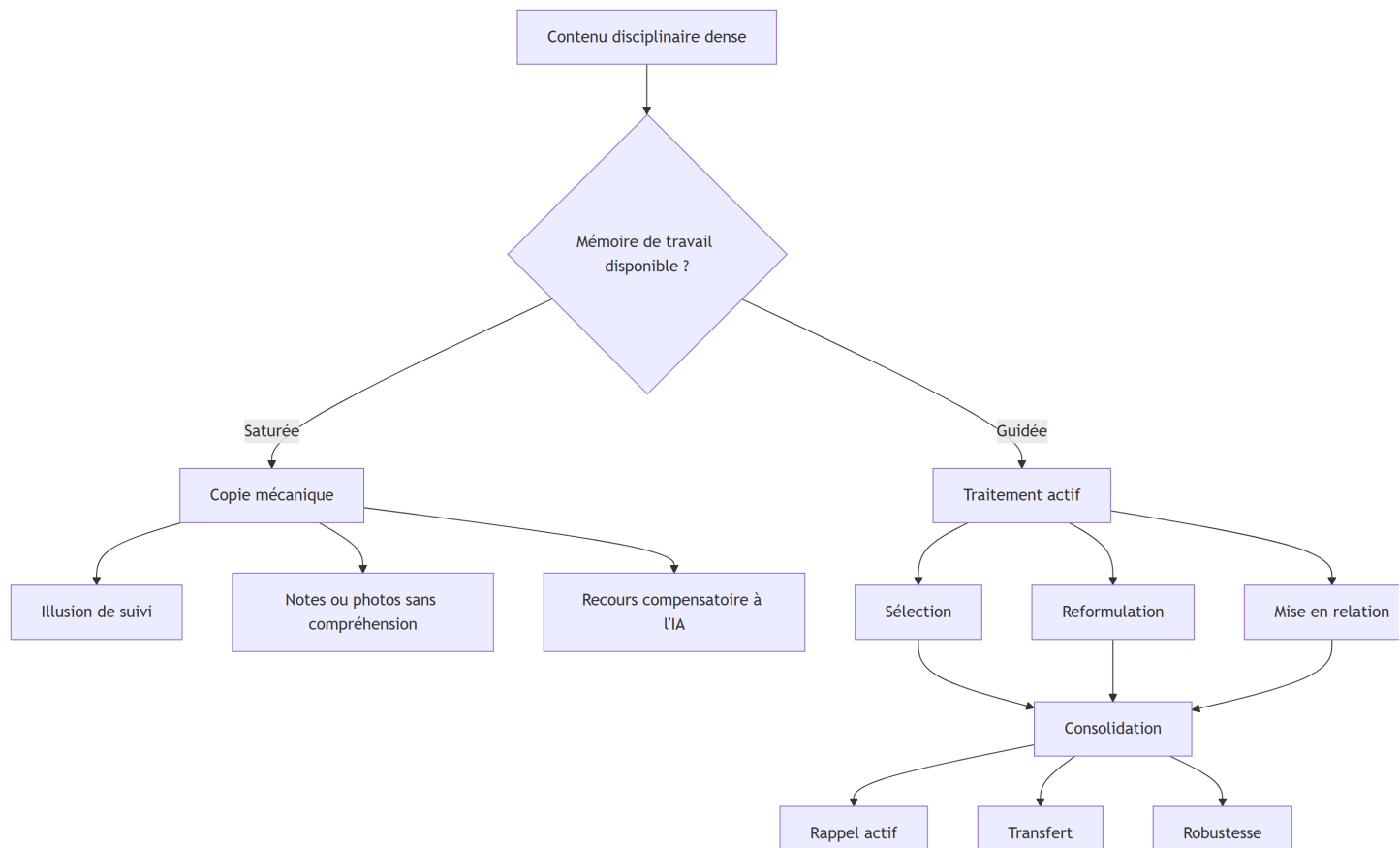
L'enjeu pédagogique n'est pas de simplifier excessivement les contenus, mais de mieux organiser leur traitement. Le **chunking**, ou regroupement en unités significantes, permet de découper une notion complexe en blocs interprétables. La hiérarchisation des supports, les pauses cognitives, les exemples progressifs, les schémas construits étape par étape et les rappels actifs permettent de réduire la surcharge inutile sans abaisser l'exigence.

Outil : Diagnostiquer la surcharge d'un support de cours

Question de diagnostic	Oui	Non	Ajustement possible	Coût indicatif
Une diapositive présente-t-elle plusieurs idées nouvelles en même temps ?			Découper en blocs successifs.	Faible
Les étudiants doivent-ils écouter, comprendre et copier simultanément ?			Prévoir une pause ou un support guidé.	Faible
Les symboles et notations sont-ils tous explicités avant usage ?			Ajouter une légende progressive.	Moyen
L'exemple arrive-t-il avant la stabilisation de la notion ?			Replacer l'exemple après un repère conceptuel.	Faible
Le support distingue-t-il l'essentiel, l'exemple et l'approfondissement ?			Utiliser un code visuel ou une hiérarchie.	Moyen
Un rappel actif est-il prévu après la séance ?			Ajouter un quiz court ou une question de transfert.	Faible

Carte de diagnostic : charge, traitement et consolidation

Mode d'emploi de la carte. Cette carte aide à repérer si une séance favorise plutôt la copie mécanique ou un traitement actif menant à la consolidation.



À expérimenter dans votre prochain cours

Pendant une séance dense, prévoyez une pause silencieuse de deux minutes. Demandez aux étudiants d'écrire trois éléments : **ce que j'ai compris**, **ce qui reste flou**, **la question que je poserais à une IA**. Reprenez ensuite deux ou trois questions pour ajuster le rythme.

À tester mentalement

Prenez une séance dense que vous donnez souvent. À quel moment précis la mémoire de travail des étudiants risque-t-elle de saturer ? Quel micro-arrêt pourrait rendre le traitement plus visible ?

Pourquoi prendre des notes n'est pas une compétence naturelle ?

□ Essentiel □ Transformer sa pratique

Dans cette section, vous allez :

Comprendre	Repérer	Agir
La différence entre stockage externe et encodage.	Les traces propres mais peu transformées : verbatim, photos, résumés automatiques.	Enseigner explicitement quoi noter, comment reformuler et comment reprendre ses notes.

Situation d'amphi

Dans un cours d'algorithmique, certains étudiants tapent presque mot à mot ce que dit l'enseignant. D'autres photographient les diapositives. Quelques-uns annotent un squelette de cours en ajoutant des mots-clés, des liens logiques et des exemples personnels. Tous repartent avec une trace. Mais ces traces n'ont pas la même valeur cognitive.

Prise de notes, encodage et stockage externe

La prise de notes est souvent considérée comme une compétence déjà acquise à l'entrée dans l'enseignement supérieur. Pourtant, elle mobilise une coordination exigeante : écouter, comprendre, sélectionner, hiérarchiser, transformer, écrire et parfois reformuler en temps réel. Elle n'est donc pas un simple geste scolaire, mais une activité cognitive complexe.

Une distinction importante oppose le **stockage externe** et l'**encodage**. Le stockage externe désigne le fait de disposer d'une trace consultable : cahier, fichier, diaporama annoté, transcription, résumé généré par IA. L'encodage renvoie au traitement interne par lequel l'étudiant transforme l'information en connaissance. Or le document produit ne dit pas toujours ce qui a été traité mentalement. Une fiche très complète peut résulter d'une copie passive. À l'inverse, des notes plus sélectives peuvent témoigner d'un effort de compréhension plus important.

Dans les cours scientifiques, la prise de notes peut même devenir contre-productive lorsqu'elle absorbe toute la mémoire de travail. L'étudiant copie les équations mais n'écoute plus leur interprétation. Il conserve la trace formelle mais perd le sens du raisonnement. La prise de notes doit donc devenir un objet d'enseignement explicite : que faut-il noter ? que peut-on laisser dans le support ? comment repérer une idée structurante ? comment reformuler une démonstration ? comment revenir sur ses notes après le cours ?

Tableau : Ce que fait l'étudiant / ce que cela produit

Geste étudiant	Trace visible	Activité cognitive probable	Risque	Geste enseignant utile
Copier mot à mot	Notes longues	Faible transformation	Effet verbatim	Interdire la transcription littérale, demander mots-clés et liens.

Geste étudiant	Trace visible	Activité cognitive probable	Risque	Geste enseignant utile
Photographier les slides	Trace complète	Stockage externe	Illusion de disponibilité	Prévoir une activité de retraitement.
Annoter un squelette	Notes guidées	Sélection et hiérarchisation	Dépendance au support	Faire compléter par une reformulation.
Reformuler après le cours	Synthèse personnelle	Encodage et consolidation	Temps invisible	Rendre ce travail attendu et visible.
Résumer avec IA	Fiche claire	Variable selon l'usage	Désynchronisation cognitive	Demander comparaison, critique et transfert.

Outil : Consignes anti-verbatim

Consignes simples à donner aux étudiants

Pendant la prise de notes, ne cherchez pas à tout écrire. Votre objectif est de produire une trace qui vous permettra de reconstruire le raisonnement.

1. Notez les **mots-clés**, pas toutes les phrases.
2. Notez les **liens logiques** : donc, parce que, en revanche, sous condition.
3. Notez les **points de bascule** : hypothèse, changement de modèle, erreur fréquente.
4. Ajoutez un **exemple personnel** ou une analogie.
5. Après le cours, complétez vos notes par une question : *qu'est-ce que je saurais expliquer sans support ?*

Pause réflexive

Dans votre cours, quelle trace étudiante vous paraît fiable aujourd'hui ? Que prouve-t-elle réellement ? Que ne prouve-t-elle pas encore ?

Micro-synthèse : Ce que les trois premières sections installent

À ce stade, trois idées se dégagent :

- l'accès au contenu ne garantit pas le traitement ;
- la trace produite ne garantit pas l'encodage ;
- une consigne apparemment simple peut engager des activités cognitives très différentes selon la manière dont elle est cadrée.

Le problème pédagogique n'est donc pas seulement de transmettre une information claire. Il est de concevoir les conditions dans lesquelles l'étudiant devra réellement la traiter, la stabiliser et la rendre mobilisable.

Papier, clavier, tablette : le support compte moins que l'activité cognitive

□ Transformer sa pratique

Dans cette section, vous allez :

Comprendre	Repérer	Agir
Pourquoi le support n'est jamais efficace ou problématique en soi.	Les usages qui favorisent la copie, l'accumulation ou l'illusion d'exhaustivité.	Associer chaque support à une consigne cognitive précise.

Dépasser le faux débat

La comparaison entre notes manuscrites et notes numériques alimente souvent des positions tranchées : le papier serait plus profond, l'ordinateur plus distrayant, la tablette plus flexible. Ces différences existent parfois, mais elles peuvent masquer la question essentielle : quel type d'activité cognitive le support favorise-t-il ?

Le clavier peut encourager une transcription verbatim, parce qu'il permet d'écrire rapidement. Mais ce n'est pas le clavier en soi qui empêche d'apprendre : c'est l'usage qui court-circuite la sélection, l'organisation et la reformulation. Le manuscrit peut favoriser une prise de notes plus sélective, mais il peut aussi être inefficace si l'étudiant copie sans structure. La tablette peut faciliter l'annotation de schémas, notamment en résistance des matériaux, électronique ou mécanique, mais elle peut aussi devenir un simple outil d'accumulation.

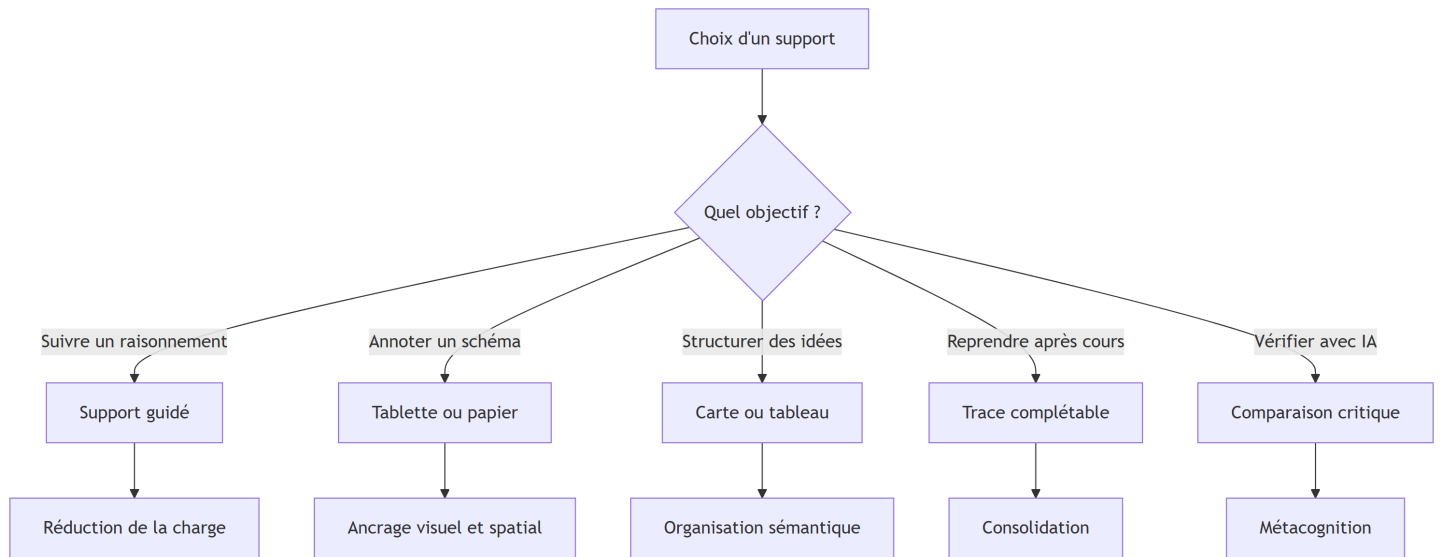
Dans une perspective pédagogique, la question n'est donc pas : *faut-il autoriser ou interdire tel support ?* Elle devient : *quelles consignes, quels rythmes et quelles activités permettent de transformer l'information, quel que soit le support ?*

Matrice : Support, risque, usage fécond

Support	Risque principal	Usage fécond en école d'ingénieurs	Consigne associée
Papier	Notes peu réorganisables	Schémas, équations, raisonnement progressif	Laisser des marges pour compléter après le cours.
Clavier	Effet verbatim	Structuration rapide, tableau, code commenté	Ne noter que mots-clés, liens et questions.
Tablette	Accumulation d'annotations	Annotation de schémas techniques	Distinguer annotations de compréhension et corrections.
Diaporama fourni	Passivité	Support de repérage	Compléter les blancs conceptuels.
Transcription automatique	Illusion d'exhaustivité	Accessibilité, reprise différée	Produire une synthèse personnelle.
IA générative	Désynchronisation	Comparaison, reformulation, critique	Identifier ce que l'IA clarifie et ce qu'elle simplifie.

Carte d'action : choisir un support selon l'activité cognitive

Mode d'emploi de la carte. Cette carte sert à choisir un support non pour lui-même, mais selon l'opération cognitive que l'on veut provoquer.



Point de vigilance

Interdire un outil ne suffit pas à créer de l'engagement cognitif. Autoriser un outil ne suffit pas non plus à soutenir l'apprentissage. Le levier pédagogique se situe dans la consigne, le rythme, l'activité demandée et la vérification de ce qui a été compris.

À tester mentalement

Prenez une consigne que vous donnez souvent. Quelle activité cognitive supposez-vous qu'elle déclenche ? Quelle activité réelle pourrait-elle produire selon le support utilisé ?

Pourquoi les méthodes « cerveau-compatible » séduisent autant ?

☐ Essentiel ☐ Approfondissement

Dans cette section, vous allez :

Comprendre	Repérer	Agir
Comment un résultat scientifique partiel peut devenir une recette pédagogique.	Les promesses fondées sur les styles d'apprentissage, le cerveau gauche/droit ou l'argument d'autorité neuroscientifique.	Remplacer les étiquettes par des choix pédagogiques diversifiés, vérifiables et contextualisés.

Neuromythes et neuro-enchantement

Les neuromythes sont des croyances pédagogiques qui prennent appui sur des éléments scientifiques simplifiés, déformés ou extrapolés. Ils séduisent parce qu'ils donnent des réponses simples à des situations complexes. Les styles d'apprentissage en sont un exemple classique : ils transforment des préférences ou des modalités sensorielles en profils stables censés déterminer la manière dont un étudiant apprendrait le mieux.

Le neuro-enchantement désigne la crédibilité excessive accordée à un discours simplement parce qu'il mobilise le vocabulaire du cerveau, des images cérébrales ou des références neuroscientifiques. Dans un contexte d'injonction à l'innovation pédagogique, cette crédibilité est puissante. Elle donne l'impression de disposer enfin d'une méthode scientifiquement fondée pour résoudre des problèmes très concrets : hétérogénéité des étudiants, décrochage, motivation, réussite aux examens.

Pourtant, le danger est précisément là : un résultat scientifique descriptif peut devenir une prescription pédagogique abusive. Dire que le cerveau apprend par plasticité ne dit pas automatiquement comment organiser un cours. Dire que l'attention est limitée ne valide pas n'importe quelle activité courte. Dire que le corps participe à la cognition ne signifie pas qu'un exercice corporel améliore mécaniquement les performances académiques.

Carte de compréhension : cycle de formation d'un neuromythe

Mode d'emploi de la carte. Cette carte montre comment une donnée scientifique partielle peut devenir une prescription pédagogique fragile.

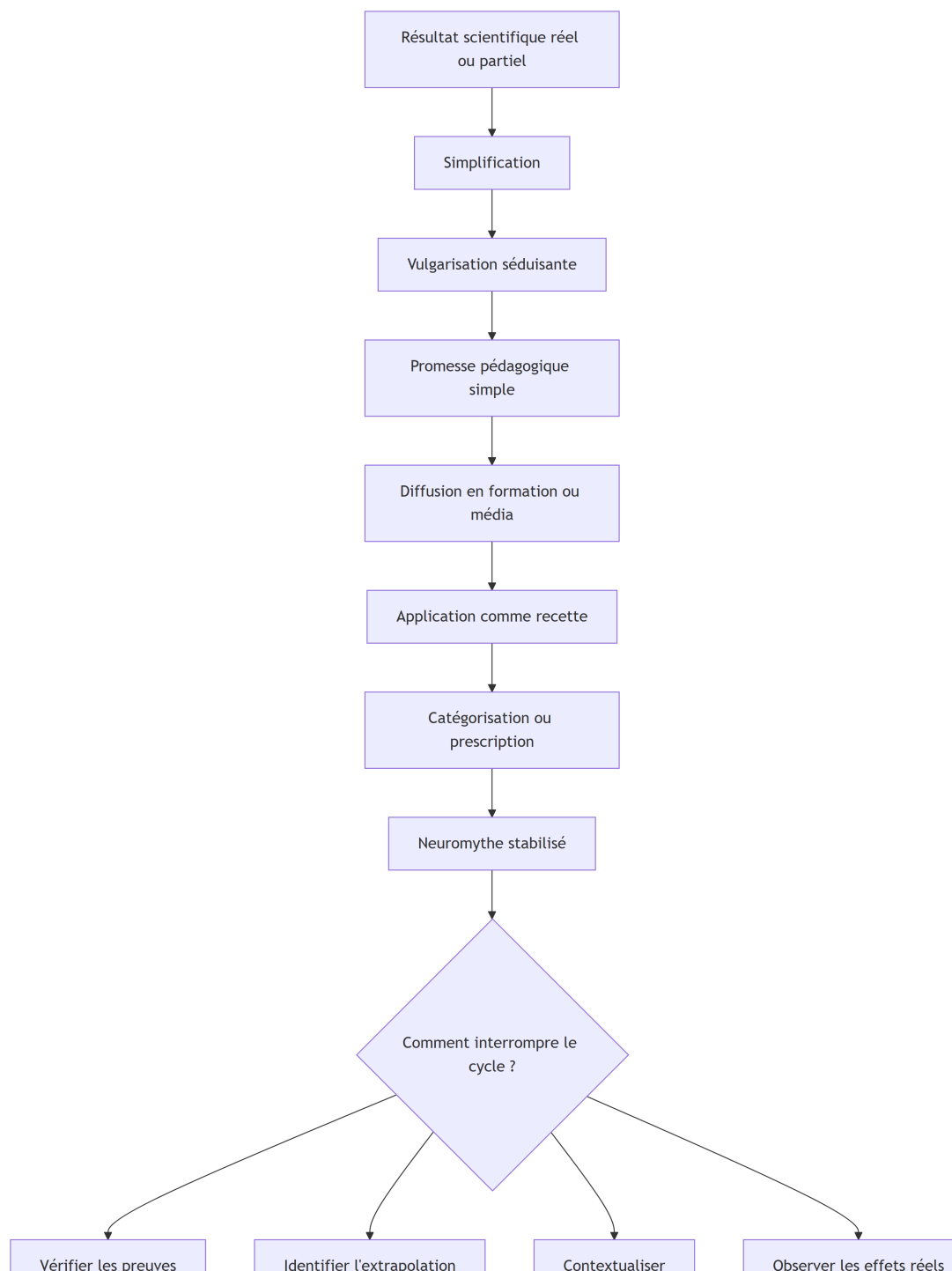


Tableau : Mythe, promesse, alternative rigoureuse

Mythe ou raccourci	Ce que cela promet	Ce que cela masque	Alternative pédagogique
Styles d'apprentissage	Adapter à chaque profil	Préférence ≠ efficacité	Diversifier les formats pour tous.
Cerveau gauche / cerveau droit	Classer les étudiants	Réseaux cérébraux complexes	Varier les tâches et stratégies.
Brain Gym scolaire	Améliorer mécaniquement les notes	Corps ≠ recette magique	Relier activité corporelle et objectif cognitif.

Mythe ou raccourci	Ce que cela promet	Ce que cela masque	Alternative pédagogique
Images cérébrales	Prouver une méthode	Corrélation $\neq$ causalité	Examiner le niveau de preuve.
IA tuteur universel	Personnaliser automatiquement	Contexte et jugement absents	Encadrer, vérifier, contextualiser.

À retenir

Déconstruire un neuromythe ne signifie pas nier la diversité des étudiants. Cela signifie refuser de transformer cette diversité en étiquettes fixes ou en prescriptions non fondées.

Quand manipuler aide vraiment à comprendre

☐ Transformer sa pratique

Dans cette section, vous allez :

Comprendre	Repérer	Agir
Pourquoi manipuler ne suffit pas à apprendre.	Les TP où l'étudiant exécute sans hypothèse, interprétation ni transfert.	Relier geste, mesure, modèle, verbalisation et cas voisin.

Cognition incarnée

La cognition incarnée remet en question l'idée selon laquelle apprendre serait seulement traiter mentalement des informations abstraites. Penser, comprendre et résoudre un problème mobilisent aussi le corps, les gestes, la perception, les objets, l'espace et l'environnement. En école d'ingénieurs, cette perspective est particulièrement féconde : un banc d'essai, une maquette, un prototype, une simulation interactive ou un FabLab peuvent aider les étudiants à relier un modèle théorique à une situation matérielle.

Mais l'activité corporelle ne suffit pas. Un TP très prescriptif peut rester cognitivement pauvre si l'étudiant suit un protocole sans formuler d'hypothèse, sans interpréter les écarts entre modèle et mesure, sans verbaliser ce que la manipulation lui permet de comprendre. À l'inverse, une activité de manipulation bien conçue peut soutenir l'apprentissage si elle oblige à articuler action, observation, modélisation et transfert.

La cognition incarnée invite donc à ne pas opposer théorie et pratique, mais à organiser leurs allers-retours. Le geste devient pédagogique lorsqu'il est relié à un concept, une hypothèse, une mesure, une erreur ou une décision.

Exemple : TP de mécanique

Dans un TP de mécanique, les étudiants appliquent un protocole sur un banc d'essai. Ils obtiennent les valeurs attendues, remplissent le tableau et concluent rapidement. Pourtant, lorsqu'on leur demande ce que la manipulation leur a appris sur le modèle, beaucoup répondent par une phrase générale. Le TP était matériel, mais pas encore pleinement incarné : l'action n'a pas été suffisamment reliée au raisonnement.

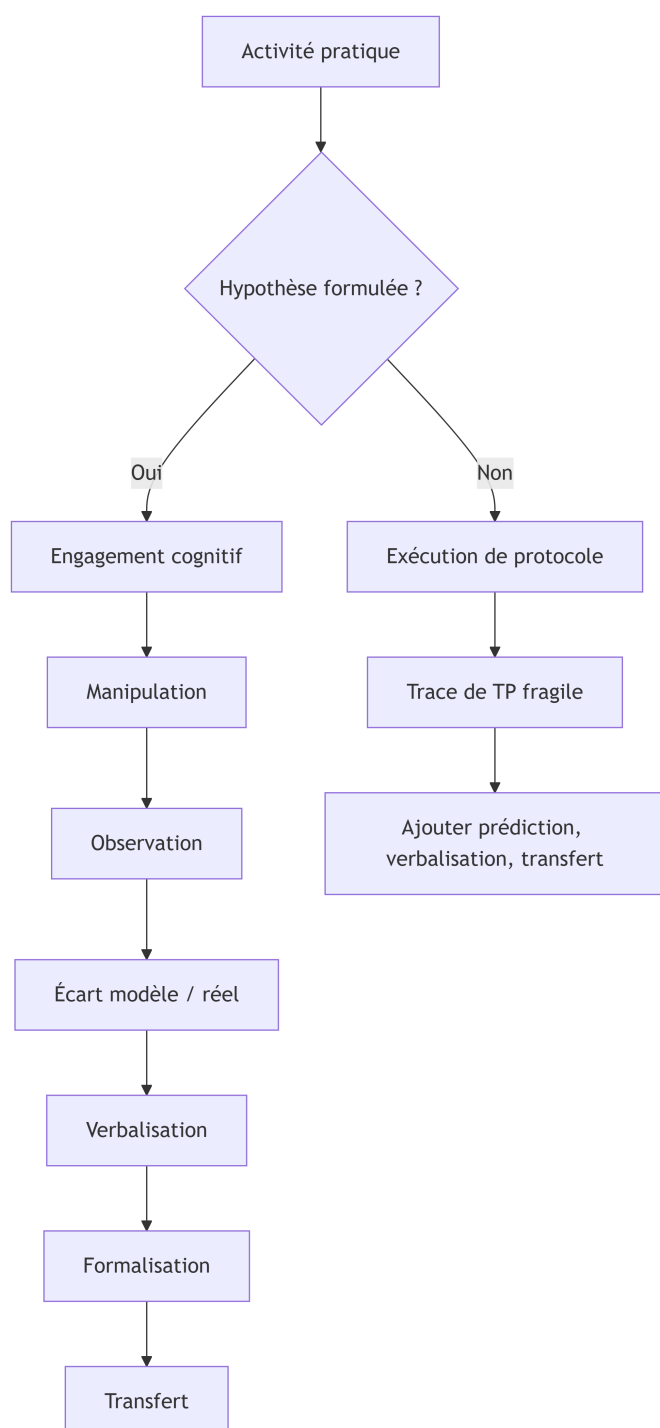
Checklist : Un TP est-il cognitivement incarné ?

Question	Oui / Non	Amélioration possible
Les étudiants formulent-ils une hypothèse avant de manipuler ?		Ajouter une question prédictive.
Le protocole laisse-t-il une place à la décision ?		Introduire un choix de mesure ou de méthode.
Les écarts entre modèle et réel sont-ils discutés ?		Prévoir un temps d'interprétation.
Les gestes sont-ils reliés à des concepts ?		Demander une verbalisation après manipulation.

Question	Oui / Non	Amélioration possible
Le transfert vers un autre cas est-il demandé ?		Ajouter une situation voisine.
L'environnement matériel est-il utilisé comme ressource cognitive ?		Faire annoter, schématiser, comparer.

Carte de diagnostic : de la manipulation au transfert

Mode d'emploi de la carte. Cette carte aide à vérifier si une activité pratique reste une exécution de protocole ou devient une activité de compréhension.



Et si l'erreur était une trace de raisonnement ?

□ Essentiel □ Transformer sa pratique

Dans cette section, vous allez :

Comprendre	Repérer	Agir
Pourquoi certaines erreurs révèlent une intuition plausible plutôt qu'un simple manque.	Les réponses trop rapides, les modèles appliqués hors contexte, les vérifications oubliées.	Transformer la correction en activité métacognitive et entraîner l'inhibition.

Métacognition et inhibition

En sciences et en ingénierie, certaines erreurs ne viennent pas d'une absence totale de connaissances. Elles viennent d'une réponse intuitive qui semble plausible mais qui n'est pas adaptée à la situation. L'étudiant applique une règle trop vite, reconnaît une forme de problème déjà vue, confond deux modèles, oublie une condition d'application ou lance un calcul avant d'avoir interrogé les hypothèses.

L'inhibition cognitive désigne la capacité à suspendre une réponse spontanée pour engager un raisonnement plus contrôlé. La métacognition permet ensuite de prendre du recul sur sa propre démarche : qu'ai-je supposé ? pourquoi ai-je choisi cette méthode ? à quel moment mon raisonnement a-t-il bifurqué ? quelle vérification aurait pu m'alerter ?

Dans cette perspective, l'erreur n'est pas seulement un écart à sanctionner. Elle devient une trace de raisonnement. Elle révèle un automatisme, une surcharge, une illusion de maîtrise, une difficulté à inhiber ou un manque de stratégie de vérification. La correction peut alors devenir un moment d'apprentissage, à condition de ne pas se limiter à afficher la bonne réponse.

Exemple : Mécanique des fluides ou analyse de données

En mécanique des fluides, des étudiants appliquent une intuition issue de la mécanique classique alors que le problème exige de raisonner sur les conditions du modèle. En analyse de données, d'autres lancent un algorithme sans interroger la distribution ou la qualité des données. Dans les deux cas, la difficulté n'est pas seulement technique : elle concerne la capacité à ralentir, à questionner l'évidence et à inhiber une réponse automatique.

Routine STOP pour entraîner l'inhibition

Routine de résolution avant calcul

S : Spontané : quelle réponse me vient immédiatement ?

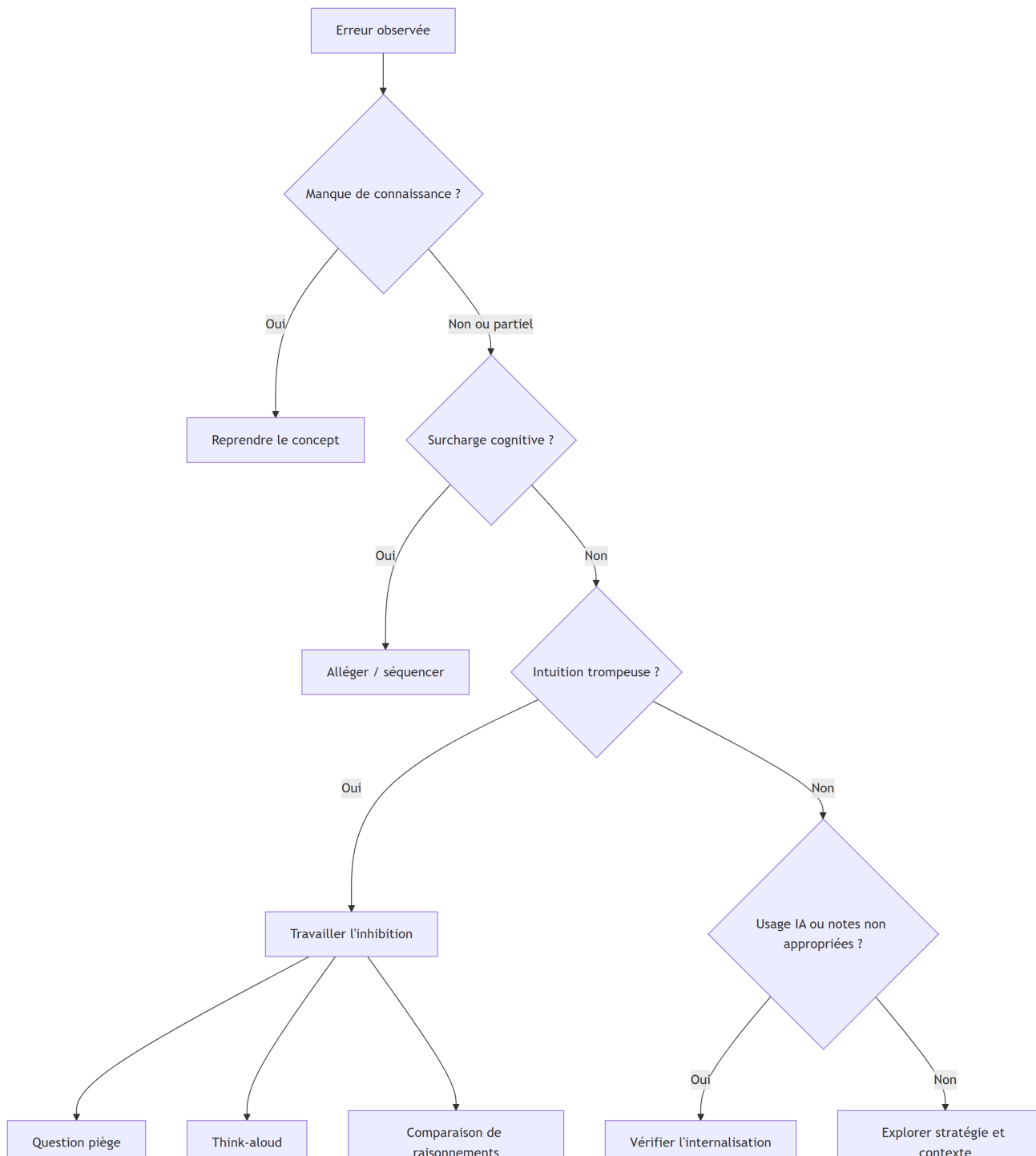
T : Trompeur : pourquoi cette réponse pourrait-elle être piégeuse ?

O : Observables : quels indices du problème dois-je vérifier ?

P : Preuve : quelle justification minimale dois-je produire avant de calculer ?

Carte de diagnostic : analyser une erreur comme trace de raisonnement

Mode d'emploi de la carte. Cette carte aide à qualifier une erreur : manque de connaissance, surcharge, intuition trompeuse ou usage non approprié d'un outil.



Correction commentée

Lors d'une correction, demandez aux étudiants d'écrire : « L'erreur semblait plausible parce que... ». Cette simple phrase transforme la correction en activité métacognitive.

Pause réflexive

Quelle erreur fréquente dans votre discipline mérite d'être conservée, analysée et discutée plutôt que simplement corrigée ?

Micro-synthèse : Ce que les fausses évidences déplacent

Les supports, les TP, les cartes visuelles ou les outils numériques ne sont pas efficaces en eux-mêmes. Ils le deviennent lorsqu'ils obligent l'étudiant à effectuer une opération cognitive identifiable : comparer, expliquer, anticiper, corriger, transférer.

Cette idée prépare le passage à l'IA générative : l'enjeu n'est pas de savoir si l'outil est autorisé ou interdit, mais de vérifier s'il soutient l'activité cognitive ou s'il s'y substitue.

Quand l'étudiant produit sans avoir appris

□ Essentiel □ Transformer sa pratique

Dans cette section, vous allez :

Comprendre	Repérer	Agir
Le risque de désynchronisation entre production externe et compréhension interne.	Les livrables fluides, corrects ou élégants que l'étudiant ne peut pas expliquer, adapter ou critiquer.	Demander des traces courtes de raisonnement, de comparaison et de transfert.

IA, externalisation cognitive et désynchronisation

L'IA générative transforme profondément les conditions de l'apprentissage, car elle ne se contente pas de donner accès à l'information. Elle peut sélectionner, reformuler, structurer, résumer, corriger, coder, comparer et argumenter. Autrement dit, elle prend en charge des opérations qui étaient auparavant réalisées par l'étudiant lui-même.

Cette externalisation peut être utile. Elle peut soutenir des étudiants allophones, aider à reformuler une explication, générer des exemples, rendre un contenu plus accessible ou servir de miroir critique. Mais elle introduit aussi un risque majeur : la **désynchronisation cognitive**. L'étudiant peut disposer d'une production externe de qualité : fiche, plan, code, résumé, correction : sans avoir réalisé l'activité cognitive interne correspondante.

Ce risque prolonge des phénomènes déjà observés avec la prise de notes passive, la relecture ou les supports complets. Mais l'IA change l'échelle et l'apparence du problème. Le rendu n'est plus seulement complet : il est fluide, structuré, argumenté, parfois meilleur que ce que l'étudiant aurait produit seul. L'illusion de maîtrise devient donc plus difficile à repérer.

Situation : Projet Python avec IA

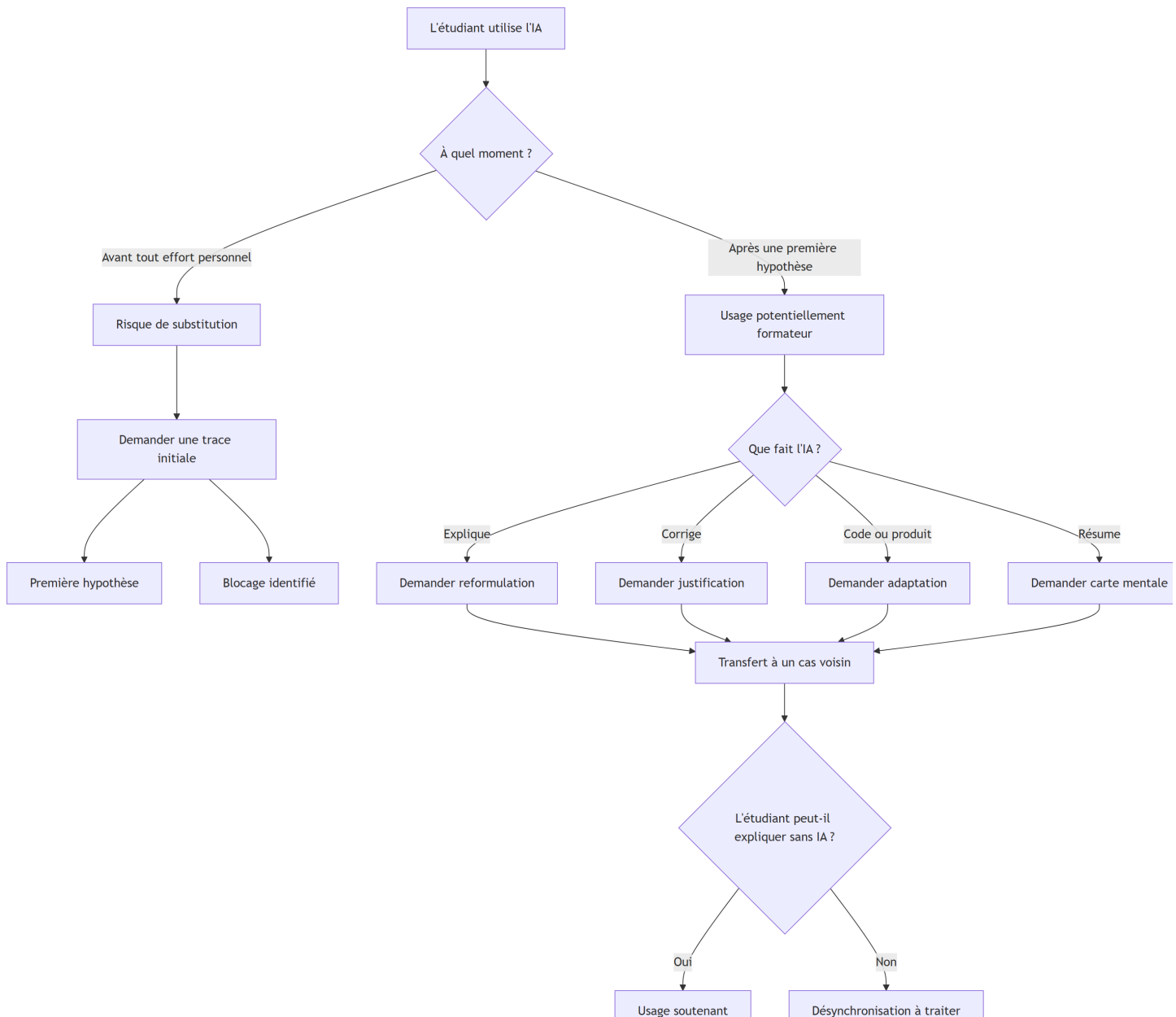
Un étudiant rend un code fonctionnel généré avec l'aide d'une IA. Le programme passe les tests de base. Mais en soutenance, l'étudiant ne sait pas expliquer pourquoi une boucle a été remplacée par une compréhension de liste, ni comment adapter le programme à un nouveau jeu de données. Le livrable est correct, mais la compétence reste incertaine.

Matrice : Production externe / compréhension interne

	Compréhension interne faible	Compréhension interne forte
Production externe faible	Difficulté visible : besoin d'accompagnement.	Idée comprise mais trace maladroite : besoin d'aide à la formalisation.
Production externe forte	Zone à risque : désynchronisation cognitive, IA ou copie possible.	Apprentissage robuste : l'étudiant peut expliquer, transférer, critiquer.

Carte d'action : décider quand et comment autoriser l'IA

Mode d'emploi de la carte. Cette carte aide à encadrer l'usage de l'IA selon le moment d'utilisation, la trace demandée et le niveau d'explication attendu.



Outil : Trois traces à demander lorsque l'IA est autorisée

Trace demandée	Fonction pédagogique	Exemple de consigne	Coût indicatif
Première hypothèse personnelle	Vérifier l'engagement avant IA	« Expliquez ce que vous pensiez faire avant d'interroger l'IA. »	Faible
Comparaison critique	Travailler la métacognition	« Identifiez ce que l'IA améliore, simplifie ou oublie. »	Moyen
Transfert à un cas voisin	Vérifier l'apprentissage réel	« Adaptez la solution à une situation légèrement différente. »	Moyen

Point de vigilance

Le risque principal n'est pas seulement la fraude. Il est que l'étudiant confonde **accès à une production de qualité** et **construction d'une compétence**.

Encadré : Qu'est-ce qu'une preuve d'apprentissage ?

Une production réussie n'est pas nécessairement une preuve robuste d'apprentissage. Elle devient plus probante lorsqu'elle permet d'observer un raisonnement, un choix, une justification, une erreur corrigée ou un transfert.

L'objectif n'est pas d'augmenter la quantité de traces, mais d'améliorer leur **valeur inférentielle** : une trace courte mais bien choisie peut mieux renseigner sur l'apprentissage qu'un long dossier final.

Trace visible	Ce qu'elle prouve faiblement	Ce qu'elle prouve mieux si...
Fiche de synthèse	L'étudiant a accès au contenu ou à une reformulation correcte.	L'étudiant explique ses choix, hiérarchise les idées et justifie ce qu'il a écarté.
Code fonctionnel	Le résultat attendu est atteint.	L'étudiant adapte le code à un cas voisin, explique une erreur ou anticipe une limite.
Réponse IA critiquée	L'étudiant a utilisé l'outil et repéré des éléments visibles.	L'étudiant identifie les hypothèses, les angles morts, les sources manquantes et les risques d'erreur.
Erreur commentée	L'étudiant a tenté une résolution.	L'étudiant explicite le raisonnement plausible qui a conduit à l'erreur et propose une stratégie de correction.
Explication orale courte	L'étudiant peut restituer un résultat.	L'étudiant reformule sans support, répond à une objection et relie la notion à une situation nouvelle.
Transfert à un cas voisin	L'étudiant reconnaît une similarité de surface.	L'étudiant explique ce qui change, ce qui reste invariant et pourquoi la méthode reste valable ou non.

Cette grille ne vise pas à multiplier les tâches d'évaluation, mais à mieux choisir les traces demandées. À l'ère de l'IA, la robustesse cognitive se repère moins dans la propreté d'un produit final que dans la capacité à expliquer, adapter, critiquer et transférer.

Rendre l'apprentissage visible, guidé et vérifiable

□ Transformer sa pratique

Dans cette section, vous allez :

Comprendre	Repérer	Agir
Pourquoi expliciter les opérations attendues réduit les implicites sans diminuer l'exigence.	Les tâches où les étudiants doivent deviner seuls quoi faire mentalement.	Choisir quelques traces à forte valeur inférentielle et les intégrer sans alourdir le cours.

Pédagogie explicite et inclusive

Une partie importante des difficultés étudiantes vient de l'implicite. Les enseignants attendent que les étudiants sachent prendre des notes, organiser leur travail, retravailler après le cours, utiliser les supports, identifier leurs erreurs, poser de bonnes questions à l'IA ou vérifier une réponse. Or ces compétences sont inégalement distribuées. Les laisser dans l'implicite renforce les écarts entre étudiants.

Une pédagogie explicite ne signifie pas simplifier les exigences. Elle consiste à rendre visibles les opérations attendues : comment sélectionner une information, comment construire une fiche, comment utiliser une IA sans déléguer le raisonnement, comment vérifier une réponse, comment préparer une évaluation, comment apprendre d'une erreur. Dans une perspective inclusive, ces explicitations bénéficient à tous, pas seulement aux étudiants identifiés comme en difficulté.

Fiche synthèse : 8 repères pour enseigner à l'ère de l'IA

Repère	Question enseignante	Geste pédagogique
Charge cognitive	Les étudiants peuvent-ils traiter ce que je présente ?	Chunking, pauses, supports hiérarchisés.
Consolidation	Quand vont-ils réactiver ?	Quiz courts, rappel actif, espacement.
Prise de notes	Savent-ils quoi noter et pourquoi ?	Squelettes de cours, consignes, exemples.
Effet verbatim	Copient-ils ou transforment-ils ?	Reformulation, méthode Cornell, mots-clés.
Neuromythes	Est-ce une preuve ou une croyance ?	Vérification des sources, prudence sur les prescriptions.
Cognition incarnée	Que permet la manipulation ?	Hypothèse, mesure, verbalisation, transfert.
Inhibition	Quelle intuition faut-il ralentir ?	Questions pièges, routine STOP, think-aloud.
IA	Que remplace ou soutient l'outil ?	Trace de raisonnement, critique, cas voisin.

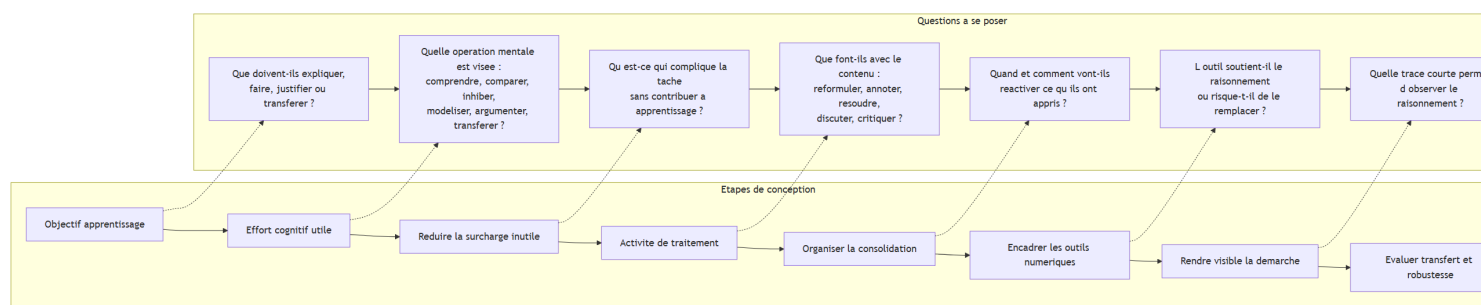
Avant / Après : Transformer une séance sans perdre l'exigence

Avant	Après
Cours dense + prise de notes libre	Cours découpé + pauses cognitives.
Support complet distribué seul	Support guidé + activité de retraitement.
IA autorisée sans consigne	IA utilisée pour comparer, critiquer, transférer.
Correction centrée sur la bonne réponse	Correction centrée sur l'erreur et la démarche.
TP protocolaire	TP avec hypothèse, mesure, verbalisation.
Évaluation du livrable final	Évaluation du raisonnement, du transfert et de la justification.

Carte d'action : concevoir une séquence plus explicite

Mode d'emploi de la carte. Cette carte sert de trame de conception pour passer de l'objectif d'apprentissage à une évaluation plus robuste.

Les questions associées permettent de vérifier que chaque étape soutient une activité cognitive identifiable.



À expérimenter

Pour une prochaine séance, choisissez un seul levier : une pause cognitive, une consigne anti-verbatim, une trace d'usage de l'IA, une question piège, une verbalisation d'erreur ou un transfert à un cas voisin. L'objectif est de transformer progressivement les pratiques sans alourdir l'ensemble du cours.

Étude de cas : Transformer un cours d'algorithmique

□ Transformer sa pratique

Dans cette section, vous allez :

Comprendre	Repérer	Agir
Comment les notions du chapitre se combinent dans un scénario réel.	Les points de fragilité d'un cours dense avec notes propres, IA et faible transfert.	Transformer progressivement la séance par des micro-ajustements ciblés.

Situation initiale

Dans un cours d'algorithmique, l'enseignante constate que les étudiants prennent beaucoup de notes et produisent des fiches de révision très propres. Plusieurs utilisent une IA pour résumer les séances et générer des exemples de code. Pourtant, en TD, ils peinent à expliquer pourquoi une structure de données est plus adaptée qu'une autre. Certains savent reproduire un exemple, mais échouent dès que le problème change légèrement.

Diagnostic

La difficulté ne relève pas uniquement d'un manque de travail. Plusieurs phénomènes se superposent : surcharge cognitive pendant le cours, prise de notes verbatim, externalisation de la synthèse vers l'IA, illusion de maîtrise liée aux fiches propres, manque de rappel actif et faible entraînement au transfert.

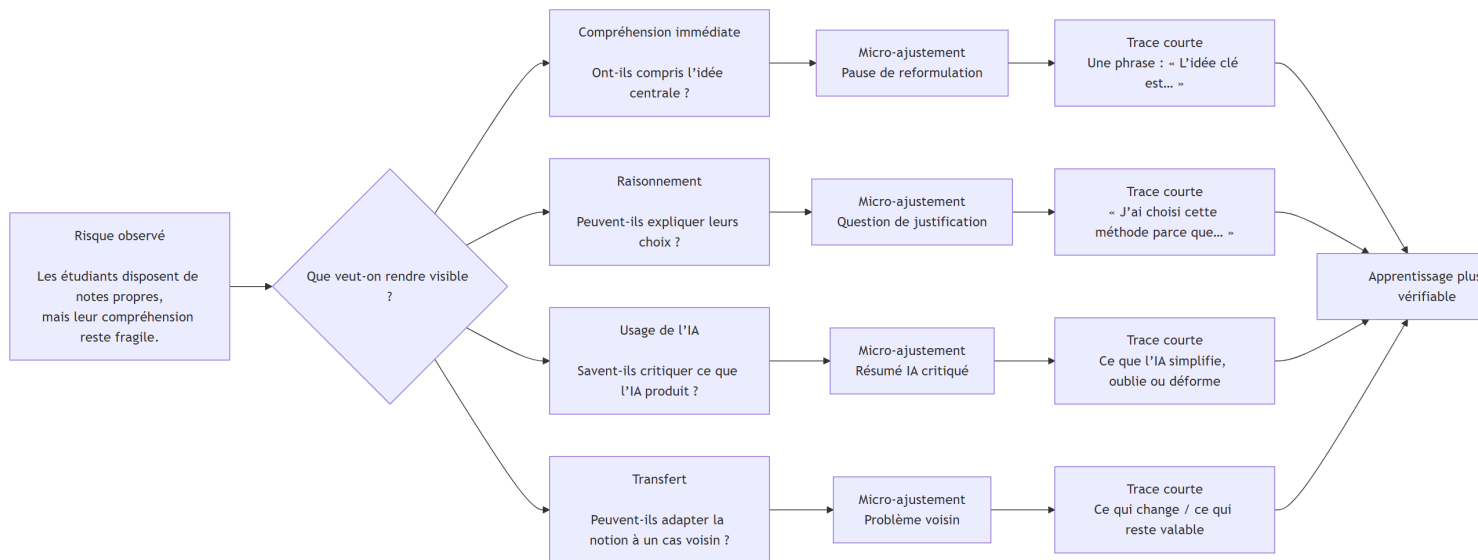
Transformation proposée

Moment	Transformation	Fonction cognitive	Coût indicatif
Avant le cours	Mini-question de prérequis sur Moodle	Réactivation.	Faible
Début du cours	Carte de compréhension des notions du jour	Orientation cognitive.	Faible
Pendant le cours	Squelette de notes avec zones à compléter	Réduction de la charge.	Moyen au départ
Toutes les 20 minutes	Pause de reformulation	Encodage.	Faible
Après un exemple de code	Question « Pourquoi cette solution et pas une autre ? »	Métacognition.	Faible
Usage IA	Résumé IA à critiquer	Détection des simplifications.	Moyen
TD suivant	Problème voisin sans IA	Transfert.	Moyen
Évaluation	Explication orale courte d'un choix algorithmique	Vérification de l'internalisation.	Moyen

Carte d'action : transformer un cours dense en parcours vérifiable

Mode d'emploi de la carte. Cette carte aide à transformer une séance dense sans ajouter une lourdeur excessive.

Elle part des risques fréquents observés chez les étudiants, puis propose des micro-ajustements permettant de rendre l'apprentissage plus visible.



Analyse commentée

La transformation ne consiste pas à supprimer l'IA ni à surcharger l'évaluation. Elle consiste à replacer l'IA dans une chaîne pédagogique où l'étudiant doit d'abord formuler, puis comparer, expliquer et transférer.

Ce que la neuro-éducation ne peut pas promettre

□ Approfondissement

Dans cette section, vous allez :

Comprendre	Repérer	Agir
Les limites de la neuro-éducation comme cadre d'aide à la décision.	Les promesses trop générales, les prescriptions universelles et les oublis du contexte.	Mettre la neuro-éducation en dialogue avec la didactique, l'évaluation, l'éthique et les conditions réelles d'enseignement.

Limites et controverses

La neuro-éducation ne peut pas promettre une méthode universelle d'enseignement. Elle ne peut pas garantir qu'une pratique fonctionnera dans tous les contextes. Elle ne peut pas réduire l'apprentissage à des mécanismes cérébraux isolés. Elle ne peut pas non plus résoudre seule les difficultés structurelles de l'enseignement supérieur : massification, hétérogénéité des publics, contraintes de temps, pression évaluative, équipements inégaux, injonctions à l'innovation.

Elle permet en revanche de mieux identifier certaines contraintes : mémoire de travail limitée, rôle du rappel actif, importance de l'erreur, effets de la surcharge, illusion de maîtrise, nécessité de consolidation. Elle oblige aussi à poser des questions plus exigeantes : que valorise réellement l'évaluation ? que rend-on visible dans le travail étudiant ? que délègue-t-on à l'IA ? comment distingue-t-on performance, progression et robustesse cognitive ?

Tableau : Permet / ne permet pas / oblige à discuter

La neuro-éducation permet de...	Elle ne permet pas de...	Elle oblige à discuter...
Identifier certaines contraintes cognitives. Comprendre certaines erreurs.	Prescrire une méthode universelle. Réduire l'étudiant à son cerveau.	La transposition laboratoire-classe. Le rôle des émotions et du contexte social.
Concevoir des supports plus soutenant.	Garantir l'apprentissage.	Les preuves réelles d'apprentissage.
Interroger l'usage de l'IA.	Autoriser ou interdire mécaniquement.	La place du jugement enseignant.
Penser la consolidation.	Remplacer la didactique disciplinaire.	La valeur de l'effort, du transfert et de la progression.

Carte de compréhension : la neuro-éducation en dialogue

Mode d'emploi de la carte.

Cette carte situe la neuro-éducation dans un ensemble de disciplines nécessaires au jugement pédagogique.

Elle ne sert pas à identifier une méthode unique, mais à élargir les questions à se poser avant de transformer une situation d'enseignement.

Domaine mobilisé	Ce qu'il permet d'éclairer	Question pédagogique à se poser
Neurosciences cognitives	Attention, mémoire de travail, inhibition, consolidation.	La séance respecte-t-elle les limites de traitement des étudiants ?
Psychologie cognitive	Apprentissage, métacognition, rappel actif, illusion de maîtrise.	Les étudiants traitent-ils réellement l'information ou pensent-ils seulement l'avoir comprise ?
Sciences de l'éducation	Scénarios pédagogiques, explicitation, accompagnement, évaluation.	Les attendus, critères et stratégies sont-ils explicités ?
Didactique disciplinaire	Obstacles propres aux savoirs enseignés, raisonnements typiques, erreurs fréquentes.	Qu'est-ce qui est difficile dans cette notion, spécifiquement dans ma discipline ?
Sociologie de l'ESR	Inégalités de codes, conditions d'étude, rapport à l'évaluation, implicites institutionnels.	Tous les étudiants disposent-ils réellement des mêmes ressources pour réussir la tâche ?
Ergonomie des environnements	Supports, espaces, interfaces, rythme, charge matérielle et attentionnelle.	L'environnement facilite-t-il le traitement ou ajoute-t-il une surcharge inutile ?
Éthique de l'IA	Transparence, responsabilité, explicabilité, équité, dépendance aux outils.	L'usage de l'IA est-il explicite, justifiable et compatible avec les objectifs de formation ?

Controverses à garder ouvertes

IA et apprentissage

L'IA peut-elle devenir un outil de consolidation, ou risque-t-elle de renforcer l'illusion de maîtrise ? La réponse dépend moins de l'outil que du scénario pédagogique : moment d'usage, consigne, trace demandée, vérification, transfert.

Évaluation et robustesse cognitive

L'enseignement supérieur valorise souvent le livrable final, la réponse correcte ou la performance rapide. Or la robustesse cognitive suppose d'observer aussi la démarche, l'erreur, la justification, la progression et la capacité de transfert.

Égalité de traitement et équité d'apprentissage

Donner le même cours, les mêmes supports et les mêmes consignes à tous ne garantit pas que tous disposent des mêmes conditions d'apprentissage. Certains étudiants compensent par un effort invisible considérable. D'autres savent déjà utiliser les codes implicites de l'institution. Une pédagogie explicite peut réduire ces écarts sans diminuer l'exigence.

Auto-positionnement final

Revenir à ses pratiques

Choisissez un cours, un TD, un TP, un projet ou une évaluation. Répondez aux questions suivantes.

Question	Réponse / piste d'action
Où mes étudiants risquent-ils de subir une surcharge cognitive ?	
Quelle activité de consolidation est prévue après le cours ?	
Ai-je déjà explicité ce qu'est une bonne prise de notes dans mon module ?	
Où l'IA risque-t-elle de produire une désynchronisation cognitive ?	
Quelle trace de raisonnement puis-je demander sans alourdir excessivement ?	
Quelle erreur fréquente pourrais-je transformer en activité métacognitive ?	
Quel neuromythe ou raccourci pédagogique dois-je éviter ?	
Mon évaluation mesure-t-elle seulement le résultat ou aussi la démarche ?	
Quel micro-changement puis-je tester dès la prochaine séance ?	

Boucler avec l'auto-test initial

Reprenez maintenant trois affirmations de l'auto-test initial. Pour chacune, indiquez si votre réponse a changé, pourquoi, et quelle conséquence cela pourrait avoir sur votre manière de concevoir un cours.

Fiche action : ce que je peux faire dès demain

☐ Transformer sa pratique

Trois transformations simples à tester dès demain

Micro-transformation	Ce que cela rend visible	Coût indicatif
Ajouter une pause cognitive de deux minutes dans une séance dense.	Ce qui est compris, flou ou à questionner.	Faible
Demander une trace de raisonnement avant ou après usage de l'IA.	L'hypothèse initiale, la comparaison critique ou l'adaptation.	Faible à moyen
Transformer une erreur fréquente en question métacognitive.	Le raisonnement plausible qui a conduit à l'erreur.	Faible

Principe de sobriété pédagogique : commencer par un seul changement, observer ce qu'il rend visible, puis décider s'il mérite d'être stabilisé dans le scénario.

Conclusion

À l'ère de l'IA, l'enjeu de l'enseignement supérieur n'est plus seulement de rendre les contenus accessibles. Il est de concevoir des situations qui rendent l'apprentissage réel, visible et vérifiable. La neuro-éducation aide à penser cet enjeu à condition d'être mobilisée avec discernement : elle éclaire les contraintes de l'attention, de la mémoire, de la consolidation, de l'erreur ou de l'engagement cognitif, mais elle ne remplace ni la didactique, ni l'analyse des contextes, ni le jugement professionnel de l'enseignant.

Dans une école d'ingénieurs, cette réflexion prend une portée particulière. Former des ingénieurs ne consiste pas seulement à leur faire produire des réponses correctes, des codes fonctionnels ou des rapports bien structurés. Il s'agit de développer leur capacité à comprendre, vérifier, expliquer, transférer, résister aux intuitions trompeuses, apprendre de leurs erreurs et utiliser les outils numériques avec lucidité. Le véritable défi n'est donc pas d'empêcher les étudiants d'utiliser l'IA, mais de leur apprendre à ne pas confondre production assistée et compétence construite.

Ce qu'il faut retenir

1. Une trace claire ne prouve pas un apprentissage réel.
2. La prise de notes et l'usage de l'IA doivent être explicitement enseignés.
3. La neuro-éducation éclaire les pratiques, mais ne les prescrit pas mécaniquement.
4. L'erreur, la reformulation, le rappel actif et le transfert rendent l'apprentissage visible.
5. La robustesse cognitive devient un objectif majeur de formation à l'ère de l'IA.

Bibliographie indicative

Mémoire, charge cognitive et consolidation

- **Baddeley, A. D. (1996).** Exploring the central executive. *Quarterly Journal of Experimental Psychology*, 49A(1), 5–28.
- **Baddeley, A. D. (2000).** The episodic buffer : A new component of working memory ? *Trends in Cognitive Sciences*, 4(11), 417–423.
- **Cepeda, N. J., Pashler, H., Vul, E., Wixted, J. T., & Rohrer, D. (2006).** Distributed practice in verbal recall tasks : A review and quantitative synthesis. *Psychological Bulletin*, 132(3), 354–380.
- **Sweller, J. (1988).** Cognitive load during problem solving : Effects on learning. *Cognitive Science*, 12(2), 257–285.

Prise de notes et apprentissage

- **Beveridge, L., & Scott, G. G. (2025).** AI for good in HE ? Leveraging AI and GenAI for note-taking. *Journal of Inclusive Practice in Further and Higher Education*.
- **Mueller, P. A., & Oppenheimer, D. M. (2014).** The pen is mightier than the keyboard : Advantages of longhand over laptop note taking. *Psychological Science*, 25(6), 1159–1168.
- **Piolat, A., Olive, T., & Kellogg, R. T. (2005).** Cognitive effort during note taking. *Applied Cognitive Psychology*, 19(3), 291–312.

Neuromythes et neuro-éducation

- **Changeux, J.-P. (2002).** *L'homme neuronal*. Fayard.
- **Dehaene, S. (2018).** *Apprendre ! Les talents du cerveau, le défi des machines*. Odile Jacob.
- **OCDE. (2002).** *Comprendre le cerveau : Vers une nouvelle science de l'apprentissage*. Éditions de l'OCDE.
- **Torrijos-Muelas, M., González-Villora, S., & Bodoque-Osma, A. R. (2021).** The persistence of neuromyths in educational settings : A systematic review. *Frontiers in Psychology*, 11, 591923.

Métacognition, erreur et autorégulation

- **Houdé, O. (2014).** *Apprendre à résister*. Le Pommier.
- **Karlen, Y., Hirt, C. N., Jud, J., Rosenthal, A., & Eberli, T. D. (2023).** Teachers as learners and agents of self-regulated learning. *Teaching and Teacher Education*, 125, 104055.
- **Zimmerman, B. J. (2002).** Becoming a self-regulated learner : An overview. *Theory Into Practice*, 41(2), 64–70.

Cognition incarnée, manipulation et transfert

- **Varela, F., Thompson, E., & Rosch, E. (1991).** *The Embodied Mind : Cognitive Science and Human Experience*. MIT Press.
- **Barsalou, L. W. (2008).** Grounded cognition. *Annual Review of Psychology*, 59, 617–645.

IA, apprentissage et enseignement supérieur

- Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Günnemann, S., Hüllermeier, E., Krusche, S., Kutyniok, G., Michaeli, T., Nerdel, C., Pfeffer, J., Poquet, O., Sailer, M., Schmidt, A., Seidel, T., Stadler, M., Weller, J., Kuhn, J., & Kasneci, G. (2023). ChatGPT for good ? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, 102274.
- Yan, L., Sha, L., Zhao, L., Li, Y., Martinez-Maldonado, R., Chen, G., Li, X., Jin, Y., & Gašević, D. (2024). Practical and ethical challenges of large language models in education : A systematic scoping review. *British Journal of Educational Technology*, 55(1), 90–112.

Chapitre 2 : Pourquoi l'IA semble comprendre ?

Illusion de raisonnement, dette cognitive et transformation des preuves d'apprentissage

Transparence éditoriale

Ce chapitre a été conçu, vérifié et éditorialisé sous la responsabilité intellectuelle de l'autrice. Des outils d'intelligence artificielle ont été mobilisés comme appui à la structuration, à la reformulation et à la mise en forme, sans se substituer aux choix scientifiques, pédagogiques et éditoriaux.

Question mobilisatrice

Pourquoi les IA génératives donnent-elles l'impression de comprendre, et comment concevoir des situations d'enseignement qui empêchent de confondre réponse plausible, performance visible et apprentissage réel ?

Mots-clés

IA générative LLM LMM Token Pattern matching IA symbolique Machine Learning Deep Learning Transformer Chain-of-Thought Raisonnement latent World Models Clever Hans Illusion de raisonnement Illusion de fluidité Anthropomorphisme Neuromythes Dette cognitive Dette épistémique Supervision critique Boîte de verre

Ce chapitre prolonge le chapitre 1 en vous aidant à :

- passer de la question *comment apprend-on ?* à la question *qu'est-ce qui donne l'illusion d'avoir compris ?* ;
- relier mémoire de travail, effort cognitif, prise de notes et externalisation à l'usage des IA génératives ;
- comprendre pourquoi la fluidité d'une réponse peut désarmer la vigilance critique ;
- distinguer fonctionnement statistique, raisonnement apparent et compréhension humaine ;
- repérer les risques de dette cognitive et de dette épistémique ;
- transformer les évaluations pour rendre visibles le raisonnement, la justification, la correction et le transfert.

Glossaire global du chapitre

Ce glossaire stabilise les notions centrales du chapitre 2. Il reprend les notions du chapitre 1 — activité cognitive, externalisation, désynchronisation, preuve d'apprentissage — pour les appliquer au cas spécifique des IA génératives.

Notion	Définition courte	Question pédagogique associée
Anthropomorphisme	Attribution à l'IA de propriétés humaines : comprendre, vouloir, penser, apprendre comme un humain.	Le vocabulaire utilisé brouille-t-il la différence entre machine et cognition humaine ?
Boîte de verre	Évaluation qui rend visibles produit, processus, propos et transfert.	Le raisonnement devient-il observable ?
Boîte noire	Évaluation centrée sur le produit final, sans accès au processus.	Que peut-on réellement inférer du livrable ?
Chain-of-Thought	Production textuelle d'étapes de raisonnement par un modèle.	Ces étapes sont-elles une preuve de raisonnement ou une rationalisation post-hoc ?
Clever Hans	Effet par lequel un système semble raisonner alors qu'il réagit à des indices de surface.	L'IA réussit-elle pour les bonnes raisons ?
Deep Learning	Famille de modèles qui apprennent des représentations distribuées et hiérarchiques dans des espaces latents.	Le modèle fonctionne-t-il sans que son cheminement soit explicable ?
Dette cognitive	Coût différé d'une délégation excessive de l'effort de synthèse, de structuration ou de déduction.	L'étudiant a-t-il produit vite au prix d'un apprentissage fragile ?
Dette épistémique	Perte de maîtrise des raisons, sources, critères et fondements permettant de justifier une production.	L'étudiant sait-il pourquoi la réponse est valide ?
Espace latent	Représentation interne distribuée, abstraite et difficilement interprétable.	Ce qui produit la réponse est-il visible et vérifiable ?
Expert fragile	Étudiant ou professionnel capable de produire un livrable crédible mais incapable d'en justifier les fondements.	Que se passe-t-il lorsque l'outil échoue ou que le contexte change ?
IA symbolique	Approche historique de l'IA fondée sur des règles explicites et manipulables.	Le raisonnement peut-il être retracé ?
Illusion de fluidité	Confusion entre facilité de lecture, cohérence formelle et maîtrise réelle.	Le texte semble-t-il fiable parce qu'il est bien formulé ?
Illusion de raisonnement	Impression qu'un raisonnement valide a été produit alors que le processus peut être fragile, incomplet ou contourné.	La réponse correcte masque-t-elle une fracture du raisonnement ?

Notion	Définition courte	Question pédagogique associée
Machine Learning	Approche où le système apprend des régularités à partir de données plutôt que de règles programmées.	La performance observable repose-t-elle sur une règle comprise ou sur une régularité apprise ?
Misconception	Croyance erronée ou simplifiée, résistante à la simple correction informative.	L'IA donne-t-elle une nouvelle légitimité à une idée fausse ?
Neuromythe	Croyance pédagogique infondée ou simplifiée sur le cerveau et l'apprentissage.	L'IA renforce-t-elle des catégories douteuses comme les styles d'apprentissage ?
Pattern matching	Appariement de formes ou de régularités statistiques.	La réponse vient-elle d'un raisonnement ou d'une reconnaissance de motifs ?
Raisonnement latent	Hypothèse selon laquelle une partie du traitement des modèles se joue dans les états internes plutôt que dans les explications textuelles.	La chaîne de pensée affichée est-elle fidèle au traitement réel ?
Supervision critique	Capacité à piloter, vérifier, corriger et contextualiser une production IA.	L'étudiant reprend-il la main sur le raisonnement ?
Token	Unité de texte prédite par le modèle : caractère, morceau de mot, mot ou séquence.	L'étudiant sait-il que le modèle prédit des formes plutôt qu'il ne comprend des concepts ?
Transformer	Architecture fondée sur des mécanismes d'attention, centrale dans les grands modèles de langage contemporains.	La cohérence du texte produit est-elle confondue avec une compréhension ?
World Models / JEPA	Orientation de recherche visant à construire des représentations plus stables des dynamiques du réel.	Le modèle dispose-t-il d'un ancrage causal ou seulement linguistique ?

Introduction

Le chapitre 1 a posé une distinction centrale : l'apprentissage ne se confond ni avec l'accès à une information, ni avec la production d'une trace visible. Un support clair, une prise de notes abondante, une fiche de synthèse ou un code fonctionnel ne suffisent pas à prouver que l'étudiant a réellement transformé l'information en savoir mobilisable. Ce chapitre 2 prolonge directement cette distinction en l'appliquant à l'intelligence artificielle générative.

L'IA générative amplifie en effet un problème déjà identifié : la possibilité de disposer d'une production propre, fluide et convaincante sans avoir traversé les opérations cognitives qui donnent à cette production sa valeur formative. Là où le chapitre 1 interrogeait la charge cognitive, la prise de notes, l'effet verbatim, les neuromythes et l'externalisation cognitive, ce chapitre analyse le niveau suivant : que se passe-t-il lorsque l'outil ne se contente plus de stocker ou reformuler l'information, mais produit à la place de l'étudiant les signes visibles du raisonnement ?

La difficulté ne tient pas seulement au fait que les IA puissent se tromper. Elle tient au fait qu'elles peuvent se tromper avec assurance, produire des explications plausibles, adopter les codes du discours académique, simuler des étapes logiques et donner l'impression d'une compréhension. Cette fluidité crée une tension pédagogique majeure : si la production devient facile, que reste-t-il à observer pour attester l'apprentissage ?

Centre de gravité du chapitre : comprendre pourquoi l'IA semble comprendre, puis transformer cette compréhension en critères pédagogiques : faire expliquer, justifier, corriger, transférer, superviser.

Continuité explicite avec le chapitre 1

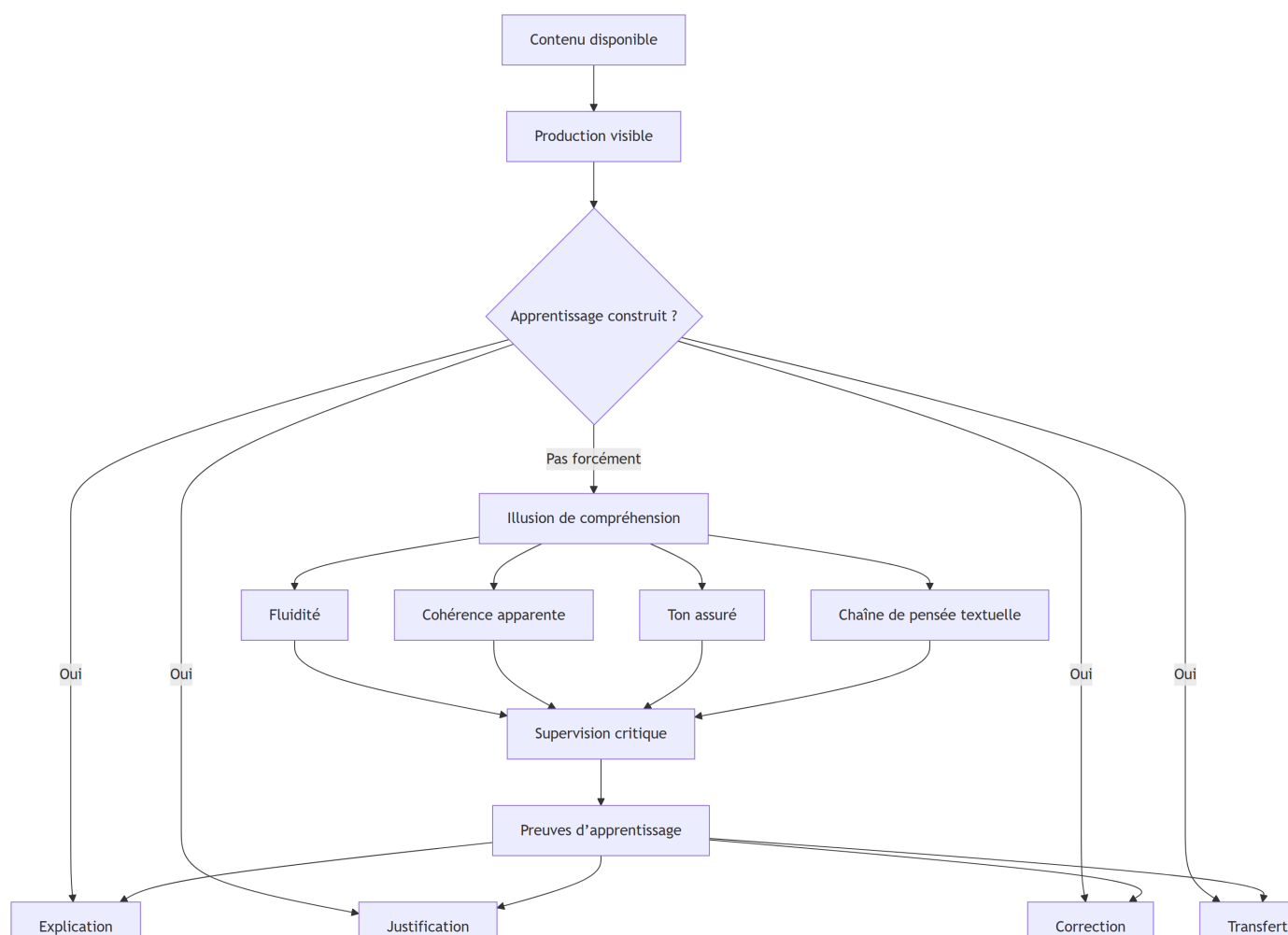
Chapitre 1	Prolongement dans le chapitre 2
La mémoire de travail est limitée.	L'IA peut réduire l'effort apparent, mais aussi supprimer des opérations utiles à l'apprentissage.
La prise de notes peut devenir verbatim.	La génération IA peut devenir une forme de verbatim augmenté : une production propre sans retraitement personnel.
L'externalisation cognitive peut soutenir ou court-circuiter l'apprentissage.	L'IA externalise non seulement le stockage, mais aussi la structuration, la synthèse et parfois le raisonnement apparent.
Une production visible ne prouve pas toujours une compréhension.	Une réponse IA ou un devoir assisté peuvent rendre cette dissociation beaucoup plus difficile à repérer.
L'erreur est une trace de raisonnement.	Les erreurs de l'IA peuvent devenir des supports d'analyse, de correction et de supervision critique.
Les neuromythes séduisent parce qu'ils simplifient l'apprentissage.	Le vocabulaire de l'IA peut réactiver des conceptions erronées du cerveau, de l'intelligence et de la personnalisation.

Plan global du chapitre : trois mouvements

Mouvement	Intention pédagogique	Notions et sections concernées
1 : Comprendre l'illusion technique	Comprendre pourquoi les modèles génératifs produisent une impression de compréhension sans fonctionner comme la cognition humaine.	IA symbolique, Machine Learning, Deep Learning, Transformers, raisonnement latent, World Models.
2 : Déconstruire les illusions pédagogiques	Repérer comment la fluidité, l'anthropomorphisme et les neuromythes brouillent les repères de l'apprentissage.	Clever Hans, illusion de raisonnement, illusion de fluidité, misconceptions, neuromythes.
3 : Transformer les pratiques	Réintroduire l'effort utile, rendre le processus visible et former à la supervision critique.	Dette cognitive, dette épistémique, métacognition, boîte de verre, 3P/4P, transfert.

Carte de compréhension : du contenu disponible à la compréhension vérifiable

Mode d'emploi de la carte. Cette carte reprend le fil du chapitre 1 — disponibilité, production, apprentissage — et montre comment l'IA générative ajoute une couche d'illusion : elle peut produire les signes de la compréhension sans garantir l'activité cognitive correspondante.



À l'ère de l'IA générative, produire ne suffit plus à prouver que l'on a compris. Il faut rendre observable le cheminement : ce que l'étudiant a vérifié, corrigé, rejeté, justifié et transféré.

Activité : Auto-test de sensibilisation

Avant de lire : vos représentations de l'IA

Pour chaque affirmation, répondez spontanément : **plutôt vrai, plutôt faux, cela dépend, je ne sais pas.**

Affirmation	Votre réponse
Une IA qui explique clairement a probablement compris le sujet.	
Une chaîne de pensée détaillée est une preuve de raisonnement.	
Un devoir bien structuré atteste au moins d'une compréhension minimale.	
Le risque principal de l'IA est de produire des erreurs factuelles.	
L'IA peut aider les étudiants à apprendre si elle leur fait gagner du temps.	
La personnalisation par IA permet naturellement de mieux respecter les profils d'apprentissage.	
Un étudiant qui corrige une réponse IA apprend davantage qu'un étudiant qui l'utilise directement.	
Une production IA doit être évaluée surtout sur sa qualité finale.	
L'oral devient plus important pour vérifier la compréhension.	
La compétence centrale est de savoir superviser l'IA plutôt que seulement l'utiliser.	

À reprendre en fin de chapitre : quelles affirmations vous semblent désormais trop simples ? Pourquoi ?

1 — De l'IA symbolique à la fluidité générative : pourquoi la performance devient opaque

□ Essentiel □ Approfondissement

Dans cette section, vous allez :

Comprendre	Repérer	Agir
Pourquoi l'IA est passée d'un raisonnement explicite à une performance statistique opaque.	Les situations où un résultat correct ne permet plus d'inférer un raisonnement transparent.	Faire expliciter aux étudiants ce qui relève du résultat, de la méthode et de la justification.

Situation ESR

En algorithmique, un étudiant obtient un code fonctionnel grâce à une IA. Le programme tourne, le résultat est conforme, le rapport est clair. Mais lorsque l'enseignant demande d'expliquer la logique de la boucle ou d'adapter le code à une contrainte légèrement différente, l'étudiant hésite. La performance existe, mais la compréhension reste incertaine.

De l'explicabilité symbolique à l'opacité statistique

Les premières approches de l'intelligence artificielle reposaient sur une idée relativement intuitive : être intelligent consistait à raisonner explicitement à partir de règles. L'IA symbolique visait à manipuler des symboles selon des règles formelles. Son avantage était l'explicabilité : on pouvait, au moins en principe, retracer le cheminement.

Cette approche s'est cependant heurtée à une limite fondamentale : le réel résiste à la formalisation exhaustive. Les situations humaines, les ambiguïtés du langage, les exceptions ordinaires et les contextes implicites sont difficiles à encoder sous forme de règles complètes. Le Machine Learning déplace alors le centre de gravité : au lieu de programmer toutes les règles, le système apprend des régularités à partir des données.

Avec le Deep Learning, cette évolution s'intensifie. Les réseaux profonds apprennent des représentations distribuées dans des espaces latents complexes. La performance augmente, mais l'explicabilité diminue. Le système fonctionne, parfois très bien, sans que le cheminement soit directement lisible. Pour les enseignants, ce déplacement est décisif : la réussite observable ne garantit plus la transparence du processus.

Les modèles génératifs : cohérence textuelle et illusion de compréhension

Les modèles génératifs contemporains, notamment ceux fondés sur l'architecture Transformer, produisent des réponses en exploitant des régularités statistiques apprises sur de vastes corpus. Leur force est de générer un langage fluide, adapté au contexte, souvent structuré et convaincant. Cette force crée aussi leur ambiguïté pédagogique.

Une réponse bien écrite active chez le lecteur des indices habituellement associés à la compréhension : précision du vocabulaire, organisation logique, ton assuré, capacité à enchaîner des arguments. Pourtant,

ces indices ne suffisent pas à prouver une compréhension conceptuelle. Le modèle ne comprend pas une règle comme un sujet humain ; il estime des formes probables dans un contexte donné.

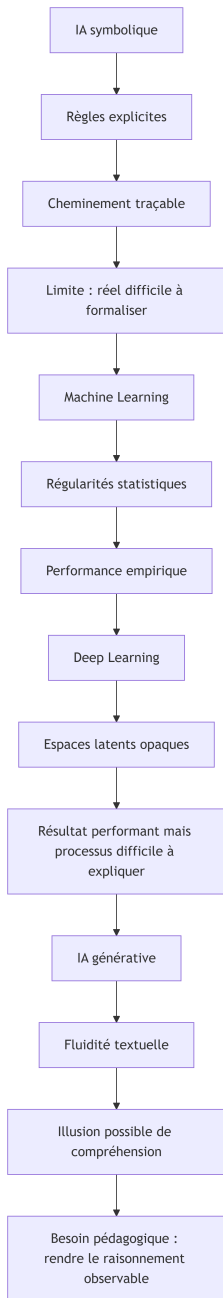
Cette distinction prolonge directement le chapitre 1 : une information disponible ou une production visible ne suffit pas à prouver l'apprentissage. L'IA rend cette dissociation plus spectaculaire, car elle produit non seulement le contenu, mais aussi les apparences du raisonnement.

□ **Essentiel : résultat correct ≠ raisonnement transparent**

Ce que l'on observe	Ce que l'on ne peut pas encore conclure
Le texte est fluide.	L'étudiant ou l'IA comprend le concept.
Le code fonctionne.	La logique algorithmique est maîtrisée.
La démonstration semble cohérente.	Le cheminement est valide.
Le plan est structuré.	Les critères de hiérarchisation sont compris.
Le vocabulaire est expert.	L'expertise est réellement construite.

Point de vigilance : la qualité formelle du résultat devient un indice de plus en plus faible si elle n'est pas accompagnée d'une demande d'explicitation.

Carte : évolution de l'IA et déplacement pédagogique



□ Approfondissement : Chain-of-Thought et raisonnement latent

Les chaînes de pensée textuelles donnent l'impression que le raisonnement du modèle devient visible. Pourtant, les ressources mobilisées soulignent que ces traces peuvent être des rationalisations post-hoc plutôt qu'une représentation fidèle des mécanismes internes. Cette prudence est importante pédagogiquement : demander à une IA de détailler ses étapes ne suffit pas à garantir que ces étapes correspondent au processus réel ayant produit la réponse.

Les recherches sur le raisonnement latent, les dynamiques internes des modèles et les World Models rappellent que le débat scientifique reste ouvert. Il ne s'agit donc pas d'affirmer que les modèles ne font "que" répéter, ni de leur attribuer une compréhension humaine. L'enjeu pédagogique est plus précis : ne pas confondre un texte explicatif avec une preuve de compréhension.

2 — Quand l'IA brouille les repères de l'apprentissage

□ Essentiel □ Recommandé □ Approfondissement

Dans cette section, vous allez :

Comprendre	Repérer	Agir
Comment la performance peut masquer une fracture du raisonnement.	Les indices de Clever Hans, d'illusion de fluidité et d'anthropomorphisme.	Concevoir des tâches qui obligent à tester le cheminement plutôt que le résultat seul.

Situation ESR

En thermodynamique, une IA produit une démonstration étape par étape. Le texte semble cohérent et convaincant. Mais une hypothèse physique est absurde. Si les étudiants ne disposent pas des critères conceptuels pour repérer l'erreur, la fluidité du texte peut renforcer une fausse compréhension au lieu de la corriger.

Le modèle Clever Hans : réussir pour les mauvaises raisons

Le phénomène Clever Hans désigne une situation où un sujet semble manifester une compétence alors qu'il réagit en réalité à des indices superficiels. Appliqué aux IA génératives, ce modèle aide à comprendre une tension centrale : un système peut produire une réponse correcte sans avoir suivi le chemin logique attendu.

Les travaux mobilisés autour de VisAidMath et des limites de généralisation des modèles montrent que des changements de forme, de symboles ou de contexte peuvent faire chuter la performance. Ce point est pédagogiquement essentiel : si une réponse correcte dépend d'indices de surface, elle ne prouve pas un raisonnement transférable.

Pour les étudiants, le risque est double. Ils peuvent attribuer à l'IA une compréhension qu'elle n'a pas nécessairement. Ils peuvent aussi s'approprier une réponse sans avoir construit les schèmes nécessaires pour l'expliquer, la corriger ou la transférer.

Illusion de raisonnement : quand la réponse finale fait écran

L'illusion de raisonnement apparaît lorsque la justesse apparente du résultat masque un processus incomplet, fragile ou contourné. Ma et al. distinguent l'exactitude du résultat final, la validité globale du processus et la robustesse détaillée du raisonnement. Cette distinction rejoint directement la problématique du chapitre 1 : la production visible ne suffit pas ; il faut observer le traitement cognitif.

En contexte éducatif, cette distinction transforme l'évaluation. Un étudiant peut rendre une synthèse bien structurée ou une résolution correcte sans avoir construit les connaissances correspondantes. Une IA peut produire une démonstration plausible sans suivre le chemin attendu. Dans les deux cas, le résultat final devient un écran.

Anthropomorphisme et neuromythes : l'IA comme nouveau langage d'anciennes croyances

L'IA ne crée pas seulement de nouvelles ambiguïtés. Elle peut aussi renforcer d'anciennes conceptions erronées du cerveau, de l'intelligence et de l'apprentissage. Les expressions comme "réseau de neurones", "mémoire", "compréhension", "apprentissage" ou "raisonnement" peuvent favoriser une confusion entre traitement statistique et cognition humaine.

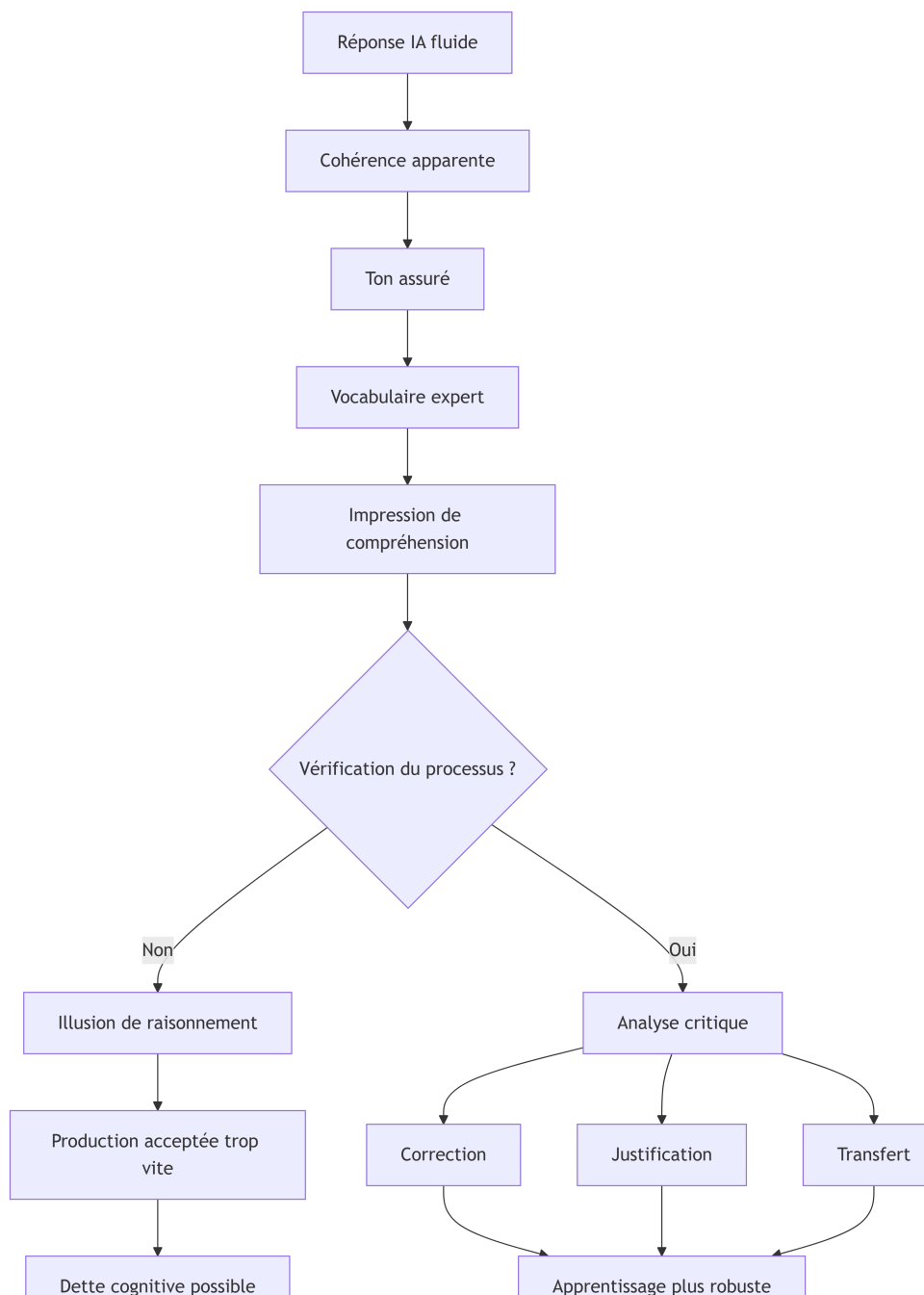
Cette confusion réactive les mécanismes déjà analysés dans le chapitre 1 à propos des neuromythes. Une technologie complexe, un vocabulaire scientifique et une promesse d'efficacité peuvent donner une apparence de légitimité à des croyances fragiles. C'est particulièrement sensible dans les discours sur la personnalisation : l'IA peut être présentée comme capable d'adapter automatiquement les contenus à des profils d'apprenants fixes, alors même que certaines catégories populaires, comme les styles d'apprentissage, relèvent de neuromythes persistants.

Le problème n'est donc pas seulement que l'IA soit mal utilisée. C'est qu'elle peut rendre plus crédibles des modèles simplistes de l'apprentissage.

□ Essentiel : quatre confusions à éviter

Confusion	Pourquoi elle est problématique	Formulation plus rigoureuse
Fluidité = vérité	Un texte bien écrit peut être faux ou incomplet.	La fluidité doit être vérifiée.
Chaîne de pensée = raisonnement réel	Les étapes affichées peuvent être une rationalisation.	Le cheminement doit être testé.
Réponse correcte = compréhension	Le résultat peut masquer une démarche fragile.	La compréhension se vérifie par justification et transfert.
Personnalisation IA = pédagogie adaptée	L'adaptation automatique peut réactiver des neuromythes.	L'adaptation doit reposer sur des besoins observables et des objectifs explicites.

Carte : de la fluidité à l'illusion de compétence



Recommandé : activité “trouver le point de rupture”

Activité pédagogique

Fournissez aux étudiants une réponse IA apparemment solide, mais contenant une erreur conceptuelle, une hypothèse absurde, une généralisation abusive ou un raisonnement incomplet. Demandez-leur de répondre à quatre questions :

Question	Intention cognitive
Où le texte paraît-il crédible ?	Identifier les indices de fluidité.
Où le raisonnement commence-t-il à se fragiliser ?	Localiser le point de rupture.

Quelle connaissance disciplinaire permet de le voir ?

Comment corriger sans simplement remplacer par une autre réponse IA ?

Réactiver les concepts.

Reconstruire un raisonnement.

□ **Approfondissement : misconceptions et post-vérité**

Les misconceptions ne sont pas de simples lacunes. Elles sont souvent robustes, parce qu'elles donnent une forme simple, intuitive et disponible à des phénomènes complexes. Dans un contexte de post-vérité, où la frontière entre fait, opinion et croyance peut être brouillée, l'IA ajoute un effet d'autorité : elle produit rapidement une réponse plausible, formulée dans un style expert.

Cette situation rend nécessaire un travail explicite d'éducation à la vigilance épistémique : distinguer source, forme, preuve, raisonnement, consensus scientifique et hypothèse. L'enjeu n'est pas seulement d'apprendre à "vérifier l'IA", mais d'apprendre à ne pas confondre crédibilité discursive et validité du savoir.

3 — Que reste-t-il à apprendre lorsque l'on peut produire sans comprendre ?

□ Essentiel □ Recommandé □ Approfondissement

Dans cette section, vous allez :

Comprendre	Repérer	Agir
Pourquoi l'IA peut améliorer la performance visible sans garantir l'apprentissage.	Les situations de dette cognitive, de dette épistémique et d'expertise fragile.	Réintroduire l'effort utile, la métacognition et l'évaluation du processus.

Situation ESR

Dans un projet Python, un étudiant utilise l'IA pour générer le code, corriger les erreurs, rédiger la documentation et expliquer le fonctionnement. Le livrable est complet. Mais lorsqu'on lui demande de modifier une contrainte, de corriger un bug ou de justifier un choix d'architecture, il ne sait pas reprendre la main. Le résultat était performant ; l'apprentissage, lui, reste fragile.

Externalisation cognitive et dette cognitive

Le chapitre 1 a montré que l'externalisation cognitive n'est pas mauvaise en soi. Un support de cours, une prise de notes, une carte mentale ou un outil numérique peuvent soutenir l'apprentissage. Mais l'IA introduit une rupture : elle ne se contente plus d'externaliser le stockage de l'information ; elle externalise aussi la structuration, la reformulation, la planification, la synthèse et parfois les signes visibles du raisonnement.

Lorsque cette délégation remplace systématiquement l'effort de traitement, elle peut produire une dette cognitive. L'étudiant gagne du temps à court terme, mais perd des occasions de construire les schèmes qui lui permettraient ensuite d'expliquer, corriger et transférer. L'aide devient problématique lorsqu'elle supprime précisément les difficultés qui avaient une valeur formative.

Performance sans effort et expertise fragile

La distinction entre performance et apprentissage devient ici centrale. Une performance est un résultat immédiat et observable. Un apprentissage est une transformation durable, qui permet de réutiliser un savoir dans un contexte nouveau. Avec l'IA générative, l'écart entre les deux peut s'élargir : l'étudiant peut produire une performance de haut niveau sans avoir construit l'apprentissage correspondant.

C'est ce qui conduit à la figure de l'expert fragile : une personne capable de présenter une solution crédible, mais dépendante de l'outil pour en reconstruire les fondements. Cette fragilité apparaît lorsque le contexte change, lorsqu'une erreur survient, lorsqu'un arbitrage est nécessaire ou lorsqu'il faut expliquer les limites du résultat.

Réintroduire l'effort utile

La réponse pédagogique ne consiste pas à interdire systématiquement l'IA. Elle consiste à réintroduire de l'effort là où cet effort est formateur. Cet effort n'est pas une difficulté punitive. Il s'agit d'une friction cognitive courte, ciblée, proportionnée et orientée vers la consolidation.

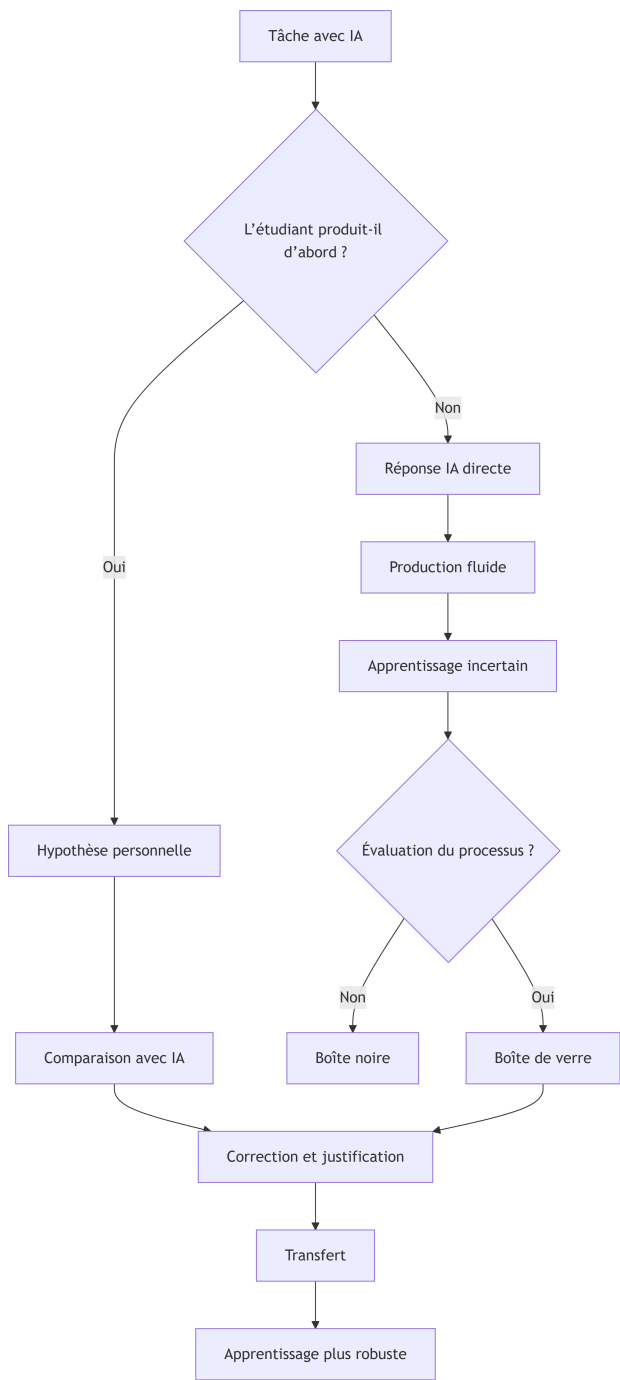
Quelques exemples : produire une première hypothèse avant d'utiliser l'IA ; comparer deux réponses ; corriger une réponse volontairement imparfaite ; expliquer ce que l'on accepte ou rejette ; transférer une solution à un cas voisin ; défendre oralement un choix.

Cette logique prolonge les notions du chapitre 1 : mémoire de travail, activité cognitive, métacognition, erreur féconde et pédagogie explicite. L'objectif n'est pas de rendre la tâche plus lourde, mais de préserver ce qui fait apprendre.

□ Essentiel : cerveau d'abord, IA ensuite

Séquence	Consigne possible	Fonction pédagogique
Avant IA	Écrire une première hypothèse, même incomplète.	Activer les connaissances et situer l'obstacle.
Avec IA	Demander une comparaison, une objection ou une reformulation.	Utiliser l'outil comme mise à l'épreuve.
Après IA	Reformuler sans copier, justifier, transférer.	Vérifier l'appropriation réelle.
Sans IA à nouveau	Résoudre une variation courte.	Tester la robustesse.

Carte : de la production assistée à l'apprentissage robuste



Recommandé : stratégie 3P

La stratégie **Produit / Processus / Propos** permet de ne plus limiter l'évaluation au livrable final.

Dimension	Question	Exemple de trace
Produit	Qu'est-ce qui est rendu ?	Rapport, code, synthèse, solution.
Processus	Comment cela a-t-il été construit ?	Brouillon, journal, versions, prompts, corrections.
Propos	L'étudiant sait-il expliquer et justifier ?	Oral court, note réflexive, défense des choix.

Extension possible : ajouter le transfert. Une compétence devient robuste lorsque l'étudiant peut adapter ce qu'il a produit à une situation nouvelle.

□ **Approfondissement : l'IA comme miroir cognitif**

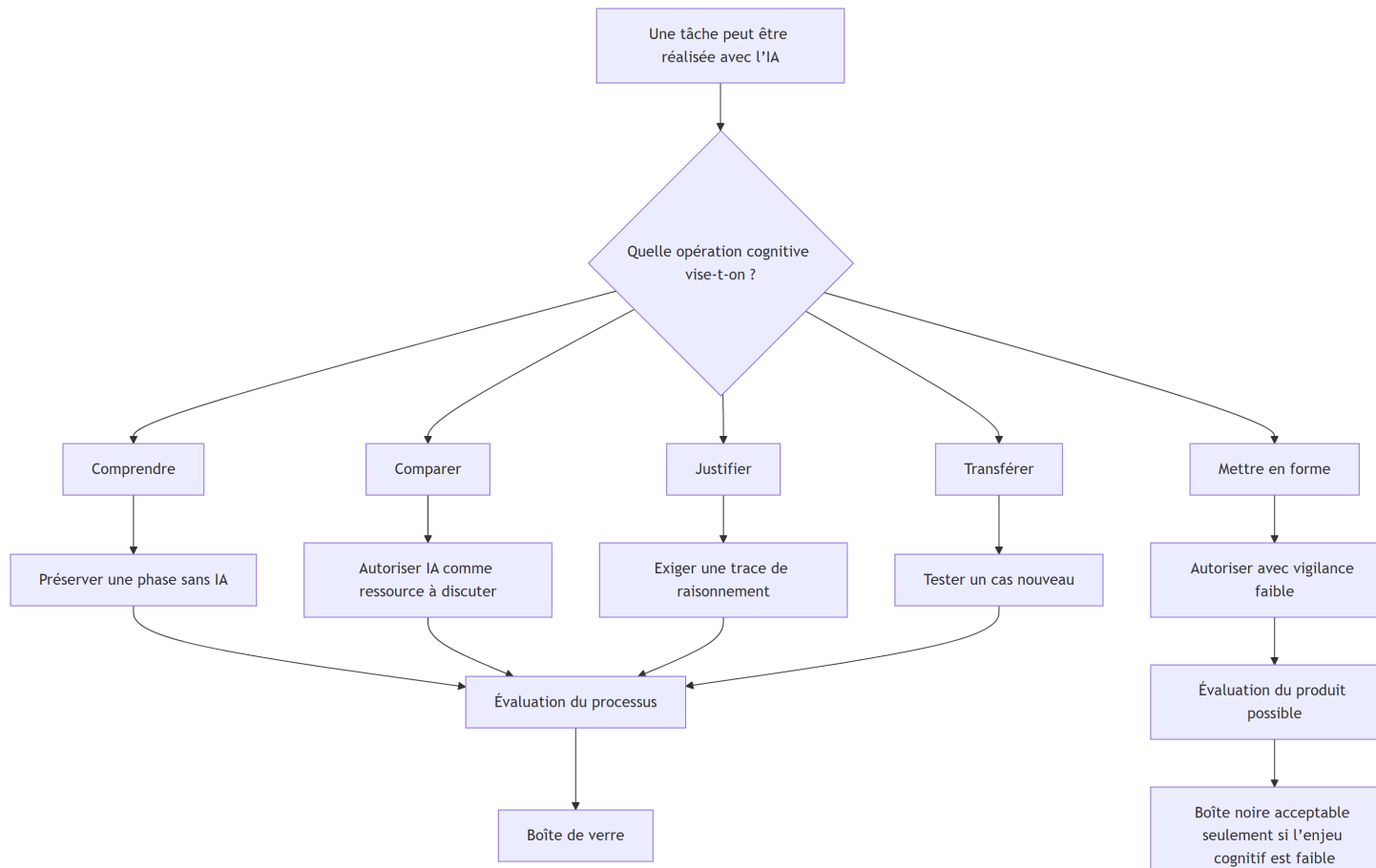
L'IA peut devenir un miroir cognitif lorsqu'elle est utilisée pour faire apparaître les limites d'un raisonnement. Dans cette configuration, elle ne sert pas à fournir directement la réponse, mais à créer une situation d'analyse : réponse imparfaite à corriger, raisonnement à auditer, hypothèse à discuter, analogie à tester.

Cette posture transforme le rôle de l'étudiant. Il n'est plus seulement utilisateur ou consommateur d'une réponse. Il devient superviseur critique : il examine la validité, les limites et les conditions de production du savoir.

Carte d'aide à la décision pédagogique

□ Essentiel

Objectif : décider comment encadrer l'usage de l'IA selon la valeur cognitive de la tâche.



Règle simple : plus la tâche est proche du cœur de l'apprentissage visé, plus l'usage de l'IA doit être précédé d'un effort autonome, accompagné d'une trace de processus et suivi d'une justification.

Étude de cas : transformer une évaluation de rapport assisté par IA

Situation initiale

Les étudiants doivent produire un rapport d'analyse sur un problème technique. L'IA est autorisée, mais aucune trace d'usage n'est demandée. Les rapports sont mieux rédigés qu'auparavant, plus homogènes et plus fluides. Pourtant, l'enseignant ne sait plus ce qui relève de la compréhension réelle des étudiants.

Diagnostic

Problème	Effet pédagogique
Le rapport est évalué comme produit final. L'usage de l'IA est invisible. Les étudiants ne justifient pas leurs choix. Aucune variation de contexte n'est demandée.	Risque de boîte noire. Impossible d'identifier la part de supervision. La compétence reste inférée, non observée. Le transfert n'est pas testé.

Transformation proposée

Avant	Après
Rapport final noté seul. IA autorisée sans trace. Correction centrée sur la forme et la complétude. Une seule situation traitée.	Rapport + note de processus + oral court. IA autorisée avec description des usages. Correction centrée sur validité, justification, limites. Variation courte pour tester le transfert.

Consigne possible

Vous pouvez utiliser une IA pour explorer des pistes, reformuler ou comparer des solutions. Vous devez cependant produire une première hypothèse personnelle avant usage, documenter ce que l'IA a apporté, indiquer ce que vous avez vérifié, expliquer ce que vous avez rejeté et justifier votre solution finale lors d'un échange oral court.

Grille d'évaluation rapide

Critère	Insuffisant	Satisfaisant	Très solide
Produit	Livrable incomplet ou non fiable.	Livrable cohérent.	Livrable cohérent, situé et limité.
Processus	Usage IA invisible ou non réfléchi.	Usage IA déclaré.	Usage IA documenté, vérifié et discuté.
Propos	L'étudiant ne sait pas expliquer.	Explication partielle.	Justification claire des choix et limites.
Correction	Accepte la sortie IA.	Corrige quelques erreurs.	Identifie les points de rupture et reconstruit le raisonnement.

Critère	Insuffisant	Satisfaisant	Très solide
Transfert	Incapacité à adapter.	Adaptation guidée.	Adaptation autonome à une contrainte nouvelle.

Fiche action : ce que je peux faire dès demain

Intention	Action simple	Effet attendu
Préserver l'engagement initial	Demander une hypothèse personnelle avant l'IA.	Éviter la délégation immédiate.
Déjouer la fluidité	Faire annoter une réponse IA : fiable / douteux / à vérifier.	Réactiver la vigilance critique.
Rendre visible le processus	Demander "ce que j'ai gardé / modifié / rejeté".	Identifier la supervision réelle.
Tester la compréhension	Poser une variation courte à l'oral.	Vérifier le transfert.
Utiliser l'erreur	Fournir une réponse IA imparfaite à corriger.	Transformer l'IA en miroir cognitif.
Lutter contre les neuromythes	Interdire les justifications par "style d'apprentissage" sans preuve.	Renforcer la rigueur pédagogique.
Déplacer l'évaluation	Ajouter un critère de justification ou de processus.	Réduire la dépendance au produit final.

Auto-positionnement final

Se situer après le chapitre

Répondez de 1 à 4 : **1 = pas du tout**, **2 = rarement**, **3 = souvent**, **4 = systématiquement**.

Question	1	2	3	4
Je distingue explicitement production visible et apprentissage réel.				
J'explique aux étudiants le fonctionnement probabiliste des IA génératives.				
Je limite les formulations anthropomorphiques du type "l'IA comprend".				
Je demande une trace du processus quand l'IA est autorisée.				
Je prévois une phase de réflexion personnelle avant l'usage de l'IA.				
J'utilise des réponses IA imparfaites comme supports d'analyse.				
J'évalue la justification et le transfert, pas seulement le livrable.				
Je repère les situations où l'IA peut renforcer des neuromythes.				
Je transforme les erreurs en traces de raisonnement.				
Je conçois mes évaluations comme des boîtes de verre plutôt que des boîtes noires.				

Analyse de votre positionnement

Additionnez vos réponses pour obtenir un score global sur 40.

Score total	Lecture possible	Point de vigilance
10 à 18	L'usage pédagogique de l'IA reste surtout centré sur le produit ou sur l'outil.	Le risque principal est de confondre une production réussie avec un apprentissage réel.
19 à 28	Vous avez déjà identifié plusieurs enjeux critiques, mais ils ne sont pas encore systématiquement traduits dans vos consignes ou évaluations.	Le passage à l'action peut être consolidé par des traces de processus, des moments sans IA et des questions de justification.
29 à 35	Votre pratique commence à intégrer une logique de supervision critique.	L'enjeu est maintenant de stabiliser ces pratiques dans vos scénarios, vos critères d'évaluation et vos échanges avec les étudiants.

Score total	Lecture possible	Point de vigilance
36 à 40	Votre approche est fortement alignée avec une pédagogie de la "boîte de verre".	Le point de vigilance devient la soutenabilité : comment maintenir ces exigences sans alourdir excessivement l'évaluation ?

Lecture par dimensions

Vous pouvez aussi relire vos réponses par grands axes.

Dimension	Questions concernées	Ce que cela permet d'observer
Compréhension du fonctionnement de l'IA	2, 3	Votre capacité à expliciter la différence entre traitement statistique, langage fluide et compréhension humaine.
Préservation de l'effort cognitif	1, 5, 9	Votre attention aux moments où l'étudiant doit encore chercher, hésiter, formuler, corriger et structurer par lui-même.
Supervision critique de l'IA	4, 6, 10	Votre capacité à transformer l'IA en objet d'analyse plutôt qu'en simple outil de production.
Évaluation de l'apprentissage réel	7, 10	Votre manière de rendre visibles les processus cognitifs, les justifications et le transfert.
Vigilance épistémique et neuromythes	3, 8	Votre capacité à éviter les formulations ou usages qui renforcent des croyances erronées sur le cerveau, l'intelligence ou l'apprentissage.

Interprétation pédagogique

Un score élevé ne signifie pas qu'il faut complexifier toutes les activités ou tout évaluer à l'oral. Il indique plutôt que vous accordez une place importante à ce qui devient décisif avec l'IA : la capacité de l'étudiant à expliquer, vérifier, corriger, transférer et assumer une production.

À l'inverse, un score faible ne signifie pas que vos pratiques sont inadaptées. Il signale surtout que certaines dimensions restent probablement implicites. Or, avec l'IA générative, ce qui reste implicite devient fragile : l'effort cognitif peut être contourné, le raisonnement peut être masqué par la fluidité du rendu, et la compétence peut être confondue avec la qualité apparente du livrable.

L'objectif n'est donc pas d'interdire l'IA, mais de mieux organiser les conditions dans lesquelles elle peut soutenir l'apprentissage sans remplacer les opérations mentales formatrices.

Pour passer à l'action

Choisissez une seule question pour laquelle vous avez répondu **1** ou **2**.

Transformez-la en micro-ajustement pédagogique pour votre prochain cours.

Si votre point faible est...	Ajustement simple possible
La distinction entre production et apprentissage	Ajouter une question : “Qu’avez-vous compris que vous ne saviez pas faire avant ?”
L’explication du fonctionnement des IA	Prévoir cinq minutes pour rappeler qu’un LLM prédit des formes linguistiques probables, sans comprendre comme un humain.
La trace du processus	Demander une courte note : “ce que j’ai demandé à l’IA / ce que j’ai gardé / ce que j’ai rejeté / ce que j’ai vérifié”.
La réflexion avant IA	Imposer un brouillon individuel ou une hypothèse initiale avant toute consultation de l’outil.
L’analyse critique de réponses IA	Donner une réponse générée volontairement imparfaite et demander aux étudiants d’identifier le point de rupture.
L’évaluation du transfert	Modifier une variable du problème et demander à l’étudiant d’adapter sa réponse.
La vigilance sur les neuromythes	Repérer une formulation du type “l’IA comprend”, “l’IA apprend comme un cerveau”, “adapté au profil visuel/auditif” et la reformuler.
La boîte de verre	Ajouter un critère d’évaluation sur la justification, le processus ou la reprise critique.

À retenir

Ce positionnement ne mesure pas votre niveau de maîtrise technique de l’IA. Il mesure plutôt votre capacité à maintenir visibles les conditions de l’apprentissage : effort utile, raisonnement, explicitation, vérification, transfert et responsabilité.

Dans un contexte où l’IA peut produire vite et bien, la question centrale devient : **qu’est-ce que l’étudiant a réellement construit, compris et rendu transférable ?**

Conclusion

L'IA générative ne rend pas l'apprentissage obsolète. Elle rend au contraire plus visible ce qui ne peut pas être réduit à la simple production de réponses plausibles : comprendre, justifier, corriger, transférer, arbitrer et assumer.

Le chapitre 1 montrait que l'apprentissage suppose un traitement actif, une gestion de la charge cognitive, une consolidation et une métacognition. Le chapitre 2 montre que l'IA peut brouiller ces repères en produisant des traces de raisonnement sans garantir le raisonnement lui-même. La réponse pédagogique n'est donc pas d'opposer usage et non-usage de l'IA, mais d'organiser les conditions d'un usage qui maintient l'activité cognitive.

Dans un monde où produire devient facile, l'enseignement supérieur ne peut plus s'appuyer uniquement sur le livrable final. Il doit rendre visible le processus, interroger les critères de validité, réintroduire l'effort utile, former à la vigilance épistémique et évaluer la robustesse.

Former à l'ère de l'IA, ce n'est pas seulement apprendre à mieux utiliser un outil. C'est apprendre à ne pas confondre fluidité et compréhension, performance et apprentissage, réponse plausible et savoir construit.

Bibliographie indicative

Cette bibliographie reprend exclusivement les références présentes dans les ressources fournies pour le chapitre 2 et les références de continuité mobilisées dans le chapitre 1 lorsque les notions sont explicitement prolongées.

Continuité avec le chapitre 1 : apprentissage, charge cognitive, prise de notes, métacognition

- **Baddeley, A. D. (1996).** Exploring the central executive. *Quarterly Journal of Experimental Psychology*, 49A(1), 5–28.
- **Baddeley, A. D. (2000).** The episodic buffer : A new component of working memory ? *Trends in Cognitive Sciences*, 4(11), 417–423.
- **Beveridge, L., & Scott, G. G. (2025).** AI for good in HE ? Leveraging AI and GenAI for note-taking. *Journal of Inclusive Practice in Further and Higher Education*.
- **Mueller, P. A., & Oppenheimer, D. M. (2014).** The pen is mightier than the keyboard : Advantages of longhand over laptop note taking. *Psychological Science*, 25(6), 1159–1168.
- **Piolat, A., Olive, T., & Kellogg, R. T. (2005).** Cognitive effort during note taking. *Applied Cognitive Psychology*, 19(3), 291–312.
- **Sweller, J. (1988).** Cognitive load during problem solving : Effects on learning. *Cognitive Science*, 12(2), 257–285.
- **Zimmerman, B. J. (2002).** Becoming a self-regulated learner : An overview. *Theory Into Practice*, 41(2), 64–70.

Fonctionnement des modèles, opacité et illusion de compréhension

- **Arcuschin, I., Janiak, J., et al. (2025).** Chain-of-thought reasoning in the wild is not always faithful. *arXiv preprint arXiv:2503.08679*.
- **Beger, C., Yi, R., et al., & Mitchell, M. (2026).** Do AI models perform human-like abstract reasoning across modalities ? *arXiv preprint arXiv:2510.02125v4*.
- **Belcamino, V., Attolino, N., Capitanelli, A., & Mastrogiovanni, F. (2025).** On the generalization gap in LLM planning : Tests and verifier-reward RL. *arXiv preprint arXiv:2601.14456v1*.
- **Bender, E. M., Gebru, T., et al. (2021).** On the dangers of stochastic parrots : Can language models be too big ? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FAccT)*, 610–623.
- **Bricken, T., Templeton, A., et al. (2023).** Towards monosemanticity : Decomposing language models with dictionary learning. *Transformer Circuits Thread*.
- **Chollet, F. (2019).** On the measure of intelligence. *arXiv preprint arXiv:1911.01547*.
- **Chollet, F., Knoop, M., et al. (2026).** ARC Prize 2025 : Technical Report. *arXiv preprint arXiv:2601.10904*.
- **Havlík, V. (2025).** Why are LLMs' abilities emergent ? *arXiv preprint arXiv:2508.04401*.
- **Helff, L., Härle, R., et al. (2026).** ActivationReasoning : Logical reasoning in latent activation spaces. *arXiv preprint arXiv:2510.18184v3*.
- **Ma, J., Zhan, R., Li, Y., Sun, D., Chan, H. P., Chao, L. S., & Wong, D. F. (2025).** VisAidMath : Benchmarking visual-aided mathematical reasoning. *Proceedings of the 39th Conference on Neural Information Processing Systems (NeurIPS 2025, Workshop VLM4RWD)*. *arXiv:2410.22995v2*.
- **Mallat, S. (2016).** Understanding deep convolutional networks. *Philosophical Transactions of the Royal Society A*, 374(2071).
- **Maes, L., Le Lidec, Q., et al., & LeCun, Y. (2026).** LeWorldModel : Stable end-to-end joint-embedding predictive architecture from pixels. *arXiv preprint arXiv:2603.19312*.
- **Mitchell, M., & Krakauer, D. C. (2023).** The debate over understanding in AI's large language models. *Proceedings of the National Academy of Sciences (PNAS)*, 120(13).

- **Shanahan, M., McDonell, K., & Reynolds, L. (2023).** Role play with large language models. *Nature*, 623(7987), 493–498.
- **Stevenson, C. E., Pafford, A., et al., & Mitchell, M. (2026).** Can large language models generalize analogy solving like children can? *Transactions of the Association for Computational Linguistics (TACL)*, 14, 612–626.
- **Su, Z., Xia, P., Guo, H., Liu, Z., Ma, Y., Qu, X., Liu, J., Li, Y., Zeng, K., Yang, Z., et al. (2025).** Thinking with images for multimodal reasoning : Foundations, methods, and future frontiers. *arXiv preprint arXiv:2506.23918*.
- **Tudor, A. R., Zeng, Y., Wang, H., Arias, J., & Gupta, G. (2025).** Building trustworthy AI by addressing its 16+2 desiderata with goal-directed commonsense reasoning. *arXiv preprint arXiv:2506.12667v2*.
- **Vaswani, A., et al. (2017).** Attention is all you need. *Advances in Neural Information Processing Systems*.
- **Wang, W. (2026).** LLM reasoning is latent, not the chain of thought. *arXiv preprint arXiv:2604.15726*.

Anthropomorphisme, misconceptions et neuromythes

- **Aktaş, B. E., Temel, Z., & Kumova, F. (2025).** The factual and ontological roots of psychological misconceptions : A study on university students in İstanbul.
- **Corbett, B., & Tangen, J. M. (2025).** AI tutors vs. tenacious myths : Evidence from personalised dialogue interventions in education. *Computers in Human Behavior*.
- **Korkos, Y. O. (2023).** Post-truth era in the modern world : Origins and causes [ЕПОХА ПОСТПРАВДИ У СУЧАСНОМУ СВІТІ : ВИТОКИ ТА ПРИЧИНИ].
- **Liang, Y., Wilson, C., & Yuile, A. L. (2025).** Folk theories of language and the brain : Public beliefs about speech, bilingualism, and neural function. Preprint.
- **Lilienfeld, S. O., Lynn, S. J., Ruscio, J., & Beyerstein, B. L. (2010).** *50 great myths of popular psychology : Shattering widespread myths and misconceptions about human behavior*. Wiley-Blackwell.
- **Tunga, Y., Celik, B., & Cagiltay, K. (2025).** Educational myths among teachers : Prevalence and refutational intervention for belief change. *Humanities and Social Sciences Communications*, 12(1). <https://doi.org/10.1057/s41599-025-05470-y>

Dettes cognitive, apprentissage, évaluation et supervision critique

- **Achuthan, K. (2025).** Artificial intelligence and learner autonomy : A meta-analysis of self-regulated and self-directed learning. *Frontiers in Education*, 10, 1738751. <https://doi.org/10.3389/feduc.2025.1738751>
- **Codère, M., & Erny, E. (2025).** L'évaluation des compétences avec l'IA. Communication présentée au colloque ROC 2025.
- **Djeufack, E. B., & Paquelin, D. (2025).** Agentivité étudiante à l'ère des outils d'IA générative : nouvelles perspectives d'autonomisation aux cycles supérieurs. *Actes du colloque ROC 2025*, 71–74.
- **Duplice, J. (2025).** Generation and L2 vocabulary learning : A classroom action study on the efficacy of the generation desirable difficulty in learning L2 vocabulary. *Vocabulary Learning and Instruction*, 14(1), 102482.
- **Fan, Y., et al. (2024).** Beware of metacognitive laziness : Effects of generative artificial intelligence on learning motivation, processes, and performance. *British Journal of Educational Technology*, 56(2), 489–530.
- **Goldmeier, G., & Mota, R. (2026).** Metacognition and pedagogy in the era of artificial intelligence. *Qeios*.
- **Hadwin, A. F., Järvelä, S., & Miller, M. (2023).** Self-regulated learning in the digital era. *Computers in Human Behavior*, 142, 107678.
- **Kosmyna, N., Hauptmann, E., et al., & Maes, P. (2025).** Your brain on ChatGPT : Accumulation of cognitive debt when using an AI assistant for essay writing task. *arXiv preprint arXiv:2506.08872*.
- **Kumar, S., et al. (2026).** Fluency illusion : A review on influence of ChatGPT in classroom settings. *Information*, 17(3), 299.
- **Lodge, J. M., & Loble, L. (2026).** Artificial intelligence, cognitive offloading and implications for education. University of Technology Sydney.

- **Maynard, A. D. (2026).** The AI Cognitive Trojan Horse : How large language models may bypass human epistemic vigilance. *arXiv preprint arXiv:2601.07085*.
- **Nathaniel, J., et al. (2026).** An experimental study of structured generative AI integration to mitigate pedagogical, cognitive, and ethical barriers in programming education. *Frontiers in Computer Science*, 8, 1789829.
- **Roediger, H. L., & Butler, A. C. (2011).** The critical role of retrieval practice in long-term retention. *Trends in Cognitive Sciences*, 15(1), 20–27.
- **Sankaranarayanan, S. (2026).** Mitigating “epistemic debt” in generative AI-scaffolded novice programming using metacognitive scripts. In *Proceedings of the 13th ACM Conference on Learning at Scale (L@S '26)*. ACM.
- **Soderstrom, N. C., & Bjork, R. A. (2015).** Learning versus performance : An integrative review. *Perspectives on Psychological Science*, 10(2), 176–199.
- **Tomisu, H., et al. (2025).** The cognitive mirror : A framework for AI-powered metacognition and self-regulated learning. *Frontiers in Education*, 10, 1697554.
- **Viana, et al. (2026).** Fluency illusion in AI-mediated learning. *Information*, 17(3), 299.
- **Yan, L., et al. (2025).** Distinguishing performance gains from learning when using generative AI. *Nature Reviews Psychology*, 4(7), 435–436.
- **Webjeje (2025).** L'intelligence artificielle dans l'éducation : État des lieux technique, éthique et pédagogique à l'horizon 2026. <https://webjeje.eu/>