

APPENDICE : Chronique et théorie d'une collaboration humain-IA

Octobre 2025 -- Novembre 2025

Cette chronique retrace le processus de création d'un livre singulier, né d'une collaboration étroite entre un expert de terrain et une intelligence artificielle. Au-delà de l'anecdote méthodologique, elle documente une expérience épistémologique : comment se transforment les conditions de production du savoir à l'ère de l'IA générative ? Comment s'articulent expertise humaine et capacités computationnelles ? Que révèle cette collaboration sur les nouvelles formes possibles de recherche et d'écriture en sciences humaines ?

Partie 1 : CHRONIQUE D'UNE COLLABORATION

I. Genèse du projet : la découverte d'un angle mort

Le point de départ : une intuition forte

Le projet naît d'une intuition forte que je portais en tant qu'ancien directeur de CIO : le discours sur l'école française fonctionne comme un système d'occultation sophistiqué. Certains sujets sont obsessionnellement présents (PISA, formation des enseignants, inégalités), d'autres systématiquement invisibilisés (tri social, ségrégation, curriculum caché).

La première formulation : "Les éléphants dans la classe"

L'image s'impose d'emblée : ces réalités massives mais invisibles que le discours éducatif s'acharne à ne pas voir. L'éléphant dans la pièce - expression anglaise désignant ce que tout le monde voit mais dont personne ne parle - devient la métaphore structurante du projet.

Mais comment passer de l'intuition à l'analyse ? Comment transformer mon expérience personnelle en démonstration rigoureuse ? C'est ici qu'intervient la collaboration avec l'IA, non comme simple outil d'écriture, mais comme partenaire méthodologique.

II. De l'intuition méthodologique à l'outillage théorique

Le démarrage : une série de posts pour le blog

Le travail commence concrètement par l'élaboration d'une série de posts pour mon blog. L'idée de départ est simple mais révélatrice : analyser l'utilisation de deux termes - "l'École" et "le système scolaire" - dans les médias mais aussi dans la littérature académique.

Cette distinction langagière m'intrigue : pourquoi certains parlent-ils de "l'École" (avec majuscule, comme une institution noble) tandis que d'autres évoquent "le système scolaire" (plus mécanique, plus critique) ? Ces choix lexicaux sont-ils neutres ou révèlent-ils des positionnements différents ?

Phase 1 : L'analyse comparative initiale (Octobre 2024)

Je constitue un premier corpus de 10 textes suite à une recherche Google. Je demande alors à l'IA de comparer ces textes : quelles sont les ressemblances et les différences dans l'usage de ces termes ? Quels acteurs utilisent quel vocabulaire ? Quelles réalités sont nommées ou occultées selon le terme employé ?

Cette première analyse révèle quelque chose de fascinant : le choix entre "l'École" et "le système scolaire" n'est pas anodin. "L'École" renvoie à un idéal républicain, à une mission civilisatrice, à des valeurs. "Le système scolaire" évoque une machine, des mécanismes, du tri social. Les textes institutionnels préfèrent massivement le premier terme. Les textes critiques utilisent davantage le second.

Phase 2 : La demande de précision théorique

La collaboration prend alors un tournant décisif. Les premiers posts du blog mentionnaient les "régimes de visibilité" et la "construction des problèmes publics" sans références académiques précises. Cette lacune que j'ai identifiée déclenche une première transformation.

Nous demandons à l'IA de mobiliser les références théoriques fondatrices :

- **Cefaï (1996)** sur les arènes publiques et les configurations dramatiques
- **Hassenteufel (2010)** sur la mise sur agenda et les processus de sélection
- **Foucault (1975)** sur les dispositifs de pouvoir-savoir
- **Bourdieu (1970, 1982)** sur la violence symbolique et la reproduction

Cette injection théorique transforme la nature du projet. Nous passons d'une grille d'analyse intuitive à un cadre conceptuel robuste. Ce qui était une observation empirique devient une analyse théorisée.

Phase 3 : La création des outils méthodologiques - les prompts d'analyse

Un moment méthodologique crucial : comment transformer ce cadre théorique en outil opératoire d'analyse ? Comment garantir la systématité et la rigueur de l'analyse sur l'ensemble du corpus ?

La solution passe par la création de **prompts d'analyse structurés** - véritables grilles méthodologiques qui permettent d'interroger systématiquement chaque texte du corpus. Ces prompts évoluent en quatre versions successives, chacune répondant à un besoin spécifique :

PROMPT 01 : Analyse modulaire générale - Régimes de visibilité dans les discours publics

Le premier prompt est volontairement large. Il permet d'analyser ce qui est rendu visible ou maintenu invisible dans tout type de discours public (politique, médiatique, institutionnel, scientifique). Sa force réside dans son architecture à **deux niveaux d'analyse complémentaires** :

- **Niveau 1 (descriptif)** : accessible, opératoire, permet d'identifier rapidement les patterns de visibilité/invisibilité

- **Niveau 2 (théorique)** : approfondi, mobilise explicitement le cadre conceptuel (Cefaï, Foucault, Bourdieu)

Cette modularité permet une double lecture : on peut rester au niveau descriptif pour un public large, ou approfondir l'analyse théorique pour une publication académique.

PROMPT 02 : Analyse spécialisée - Régimes de visibilité dans le discours éducatif

Le succès du premier prompt révèle un besoin : il faut le spécialiser pour le domaine éducatif. Le deuxième prompt reprend la même architecture modulaire mais l'adapte spécifiquement aux discours sur l'éducation. Il intègre :

- Les acteurs spécifiques du champ éducatif (OCDE, ministère, syndicats, chercheurs)
- Les thématiques propres à l'école (orientation, pédagogie, évaluation, ségrégation)
- Les mécanismes d'invisibilisation particuliers au système scolaire

Ce prompt devient l'outil central d'analyse de tous les textes du corpus pour les posts du blog.

PROMPT 03 : Révision systématique - Contrôle qualité post par post

Une fois les posts rédigés, un nouveau besoin émerge : garantir la cohérence de l'ensemble. Le troisième prompt est un outil de **révision systématique**. Pour chaque post, il vérifie :

- La cohérence interne de l'argumentation
- La conformité au style personnel (voir PROMPT 04)
- L'articulation avec les autres posts de la série
- La progression logique de l'ensemble

Ce prompt assure que la série forme un tout cohérent plutôt qu'une juxtaposition de textes indépendants.

PROMPT 04 : Préservation du style - Le style Desclaux

Le quatrième prompt répond à un enjeu crucial : comment collaborer avec une IA sans perdre l'authenticité de ma voix ? Ce prompt codifie les caractéristiques de mon style d'écriture :

- Ton accessible mais rigoureux
- Usage d'exemples concrets tirés du terrain
- Articulation entre empirique et théorique
- Progression pédagogique de l'argument

Ce prompt garantit que l'IA amplifie mon style plutôt qu'elle ne le remplace par un style générique.

L'innovation méthodologique : la systématité réflexive

Ces quatre prompts constituent une innovation méthodologique majeure. Ils transforment la collaboration humain-IA d'un dialogue improvisé en un processus systématique et réflexif. Chaque texte du corpus passe par la même grille d'analyse (PROMPT 01 ou 02), chaque post produit est soumis à la même révision (PROMPT 03), et l'ensemble maintient une cohérence stylistique (PROMPT 04).

Cette systématique n'est pas rigidité : les prompts évoluent, s'affinent, s'adaptent au fur et à mesure du projet. Ils constituent une méthodologie vivante, non un carcan.

Phase 4 : La modularisation méthodologique

Face à la richesse du nouveau cadre théorique et des outils créés, un risque émerge : la complexification excessive. Comment maintenir l'accessibilité sans sacrifier la profondeur analytique ?

La solution, déjà présente dans l'architecture des prompts, se généralise : une approche **modulaire** à deux niveaux :

- **Niveau 1 (descriptif)** : accessible, opératoire, immédiat
- **Niveau 2 (théorique)** : approfondi, conceptuel, académique

Cette architecture permet une double lecture : le blog reste accessible au grand public tandis que le livre développe la sophistication théorique nécessaire à la légitimité académique.

Phase 5 : La stratégie différenciée

Un moment décisif : la décision de maintenir deux temporalités et deux objectifs distincts :

- Finir la série de posts avec les prompts initiaux (cohérence, accessibilité)
- Préparer le livre avec des prompts enrichis (profondeur théorique, légitimité académique)

Cette bifurcation révèle une compréhension fine des différents "régimes de publication" et de leurs exigences spécifiques.

III. La découverte de l'invisibilisation différenciée

Le tournant : l'analyse du rapport CESE (Novembre 2024)

Tout bascule avec notre analyse du rapport du Conseil Économique, Social et Environnemental "Réussite à l'école, réussite de l'école" (2024). Ce document de 146 pages, adopté à 98 voix pour, devient notre pierre de Rosette.

La révélation : même les textes les plus critiques participent à l'invisibilisation. Le CESE documente méticuleusement les inégalités tout en évitant soigneusement d'analyser les mécanismes qui les produisent. Il dénonce les symptômes en occultant les causes structurelles.

Cette découverte méthodologique est cruciale : il n'y a pas UNE invisibilisation uniforme mais DES invisibilisations croisées selon les positions des énonciateurs. Chaque acteur

(OCDE, syndicats, ministère, chercheurs) a ses angles morts spécifiques et ses zones de lucidité partielle.

L'émergence de la grille d'analyse quadripartite

De cette analyse émerge une grille systématique des mécanismes d'invisibilisation :

1. **L'euphémisation** : transformer les violences en difficultés
2. **La technicisation** : convertir le politique en technique
3. **La fragmentation** : morceler pour empêcher la vision d'ensemble
4. **La dissociation** : séparer ce qui devrait être articulé

Cette grille devient l'outil central d'analyse, appliqué systématiquement à l'ensemble du corpus.

IV. L'émergence de la métaphore structurante

De l'éléphant générique aux éléphants spécifiques

La métaphore initiale se complexifie. Il n'y a pas UN éléphant mais DES éléphants, chacun correspondant à un mécanisme d'invisibilisation spécifique :

- L'éléphant PISA qui écrase tout le reste
- L'éléphant de l'orientation qu'on refuse de voir digérer
- L'éléphant de la ségrégation qu'on euphémise en mixité
- L'éléphant des conditions matérielles qu'on dématérialise

Le débat sur la violence métaphorique

Un moment crucial de réflexion : la métaphore digestive pour l'orientation ("le système ingère des enfants et les excrète dans des filières") est-elle trop violente ?

Ce débat révèle un enjeu central : comment être radical sans être repoussant ? Comment dire la violence du système sans brutaliser le lecteur ?

La solution : une **stratégie de progression** :

- Début du livre : métaphores douces pour installer la problématique
- Milieu : montée en intensité progressive
- Fin : formulations les plus incisives une fois le lecteur préparé

Cette stratégie rhétorique transforme l'organisation même du livre.

V. La construction du cadre théorique : de l'empirique au conceptuel

Phase 1 : L'accumulation des observations

Les premiers chapitres procèdent par accumulation empirique. Nous identifions des patterns récurrents dans le corpus : certains mots reviennent obsessionnellement (équité, mérite, réussite), d'autres sont systématiquement évités (classe sociale, sélection, échec). Certains acteurs sont sur-représentés (OCDE, ministère), d'autres invisibilisés (élèves, parents des classes populaires).

Phase 2 : La conceptualisation progressive

Face à cette masse d'observations, un besoin émerge : passer de la description à l'explication. Comment ces patterns s'organisent-ils en système ? Quels mécanismes sous-tendent cette organisation ?

Nous mobilisons alors progressivement trois cadres théoriques complémentaires :

1. **Les arènes publiques** (Cefaï) : comment se construit l'espace légitime du débat éducatif
2. **Les dispositifs de pouvoir-savoir** (Foucault) : comment le discours produit simultanément savoir et domination
3. **La violence symbolique** (Bourdieu) : comment s'impose la légitimité des classements scolaires

Phase 3 : L'unification théorique

Un moment décisif : la découverte que ces trois cadres ne sont pas juxtaposés mais articulés. Les arènes publiques définissent l'espace légitime du dicible. Les dispositifs de pouvoir-savoir organisent ce qui peut être su et comment. La violence symbolique naturalise les hiérarchies produites.

Cette articulation transforme notre approche : nous ne faisons plus de l'éclectisme théorique mais construisons un cadre unifié d'analyse des mécanismes d'invisibilisation.

VI. La phase de restructuration et d'approfondissement

Décembre 2024 : Le moment critique

Après deux mois de travail, une évidence s'impose : le livre existe mais manque de cohérence. Les chapitres se succèdent sans articulation claire. Le cadre théorique est présent mais pas unifié. Les analyses sont riches mais dispersées.

Nous décidons une restructuration complète :

1. **Unification théorique** : réécriture de tous les chapitres selon le cadre unifié
2. **Cohérence méthodologique** : application systématique de la grille d'analyse
3. **Progression argumentative** : organisation stratégique de la montée en charge critique

Cette phase de restructuration prend plusieurs semaines. Chaque chapitre est repris, réécrit, repositionné. L'architecture globale du livre se transforme radicalement.

Janvier-Mars 2025 : L'approfondissement

La restructuration révèle des zones d'ombre. Certains mécanismes restent insuffisamment analysés. Certains éléphants méritent leur propre chapitre. Certaines connexions théoriques doivent être explicitées.

Nous développons alors :

- **Chapitre 5** : analyse détaillée de l'orientation comme processus de tri invisible
- **Chapitre 7** : approfondissement sur le curriculum invisible
- **Chapitre 8** : découverte de l'invisibilisation de la matérialité scolaire
- **Chapitre 9** : articulation entre écologie et système éducatif (l'horizon invisible)

Cette phase révèle la capacité générative de la collaboration : chaque approfondissement en suscite d'autres. Le livre s'enrichit en complexité sans perdre en cohérence.

VII. L'évolution du rôle de l'IA dans la collaboration

Phase initiale : L'IA comme amplificateur

Au début, l'IA fonctionne principalement comme amplificateur de mon intuition. Je propose une piste, l'IA la développe, je valide ou redirige. Le rôle reste asymétrique : j'oriente, l'IA produit.

Phase intermédiaire : L'IA comme révélateur

Progressivement, l'IA devient révélateur de dimensions que je n'avais pas anticipées. En systématisant mon analyse, elle fait apparaître des patterns que je n'avais pas consciemment identifiés. Elle pointe des contradictions, suggère des connexions, identifie des absences.

Le rôle devient plus dialectique : je propose, l'IA développe et interroge, je réoriente à partir de cette interrogation.

Phase mature : La co-élaboration véritable

Dans les derniers mois, la distinction entre "ma" contribution et celle de l'IA devient floue. Les idées émergent de l'interaction elle-même. Une suggestion de l'IA déclenche chez moi une intuition qui, systématisée par l'IA, révèle une dimension théorique qui transforme à son tour mon analyse.

Le chapitre 8 sur la matérialité scolaire illustre parfaitement cette co-élaboration. Ni moi ni l'IA n'avions initialement pensé à ce thème. Il émerge de nos échanges successifs sur les absences dans le discours éducatif. Une fois identifié, il se révèle central et transforme rétrospectivement notre compréhension des mécanismes d'invisibilisation.

VIII. Les mécanismes de la collaboration

La validation empirique systématique

Chaque proposition de l'IA passe par un filtre rigoureux : correspond-elle à mon expérience du terrain ? Est-elle vérifiable dans les sources ? Ne relève-t-elle pas d'une généralisation abusive ? Cette validation systématique reste ma prérogative fondamentale.

L'amplification conceptuelle

L'IA apporte une capacité unique de systématisation conceptuelle. Mes observations dispersées deviennent grilles d'analyse. Mes intuitions se transforment en hypothèses testables. Mes formulations trouvent leur ancrage théorique.

La cohérence longitudinale

Un livre se construit sur plus d'un an. Maintenir la cohérence sur cette durée est un défi majeur. L'IA excelle dans ce rôle : elle identifie les contradictions involontaires, repère les répétitions inutiles, signale les développements manquants.

La créativité dialogique

Le dialogue constant génère une forme de créativité impossible à anticiper. Ni réductible à mes seules intuitions, ni à la seule systématisation de l'IA, mais émergente de leur interaction.

IX. Les innovations méthodologiques du projet

L'analyse par triangulation des invisibilisations

Une méthode originale émerge de notre travail : identifier un phénomène non par sa présence mais par son absence croisée dans différents corpus. Quand OCDE, syndicats et ministère évitent tous le même sujet, c'est précisément ce qui révèle son caractère problématique. Cette triangulation négative devient un outil d'analyse puissant.

La cartographie des absences structurelles

Nous développons une attention systématique aux absences : les questions qui ne sont pas posées parce qu'elles sont impensables dans le cadre dominant, les liens qui ne sont pas faits parce qu'ils révéleraient des contradictions insoutenables. Ces absences structurelles sont aussi signifiantes que les présences obsessionnelles.

La méthode des invisibilisations croisées

La triangulation des invisibilisations constitue une innovation méthodologique majeure. Cette méthode d'analyse croisée révèle comment l'invisibilisation d'un phénomène A permet l'euphémisation d'un phénomène B, comment la fragmentation de C masque la cohérence systémique de D, comment la technicisation de E dépolitise automatiquement F. Les mécanismes d'occultation fonctionnent en réseau, se renforçant mutuellement.

Le paradoxe productif de l'auto-réflexivité

Un paradoxe productif structure notre démarche : reconnaître nos propres mécanismes d'invisibilisation pour les dépasser et enrichir l'analyse. Mon oubli initial de la matérialité scolaire, une fois identifié, devient non pas une faiblesse à cacher mais une découverte

méthodologique qui génère un chapitre entier et approfondit notre compréhension des mécanismes d'occultation.

Les prompts comme infrastructure méthodologique

Les quatre prompts d'analyse constituent en eux-mêmes une innovation. Ils représentent une tentative de formaliser un processus d'analyse critique, de le rendre reproductible et systématique sans le rigidifier. Cette formalisation permet :

- La **traçabilité** : on sait exactement quelle grille a été appliquée à quel texte
- La **reproductibilité** : un autre chercheur pourrait appliquer les mêmes prompts
- La **transparence** : la méthode est explicite et discutable
- L'**évolutivité** : les prompts s'affinent au fur et à mesure du projet

Cette infrastructure méthodologique transforme la collaboration humain-IA en un protocole de recherche rigoureux.

X. Difficultés rencontrées et solutions développées

La première difficulté majeure fut d'éviter le piège du catalogue de dénonciations. La tentation était forte de simplement lister les manquements et les hypocrisies du système. La solution a consisté à développer une analyse systémique des mécanismes, transformant la critique en déconstruction méthodique des dispositifs de pouvoir-savoir.

Articuler critique radicale et rigueur académique représentait un défi constant. Comment être incisif sans perdre en crédibilité scientifique ? La réponse s'est construite dans l'élaboration d'un cadre théorique solide appuyé sur plus de 300 sources vérifiables, permettant d'asseoir la légitimité de l'analyse critique.

Maintenir la cohérence dans un ouvrage développé par étapes sur plus d'un an exigeait une vigilance particulière. La solution a été une réécriture systématique de tous les chapitres après l'établissement du cadre théorique unifiant, assurant une cohérence rétrospective à l'ensemble.

La tension entre accessibilité et profondeur analytique traversait constamment le projet. Nous l'avons résolue en créant un double niveau de lecture avec une progression pédagogique permettant au lecteur non-spécialiste d'accéder progressivement à la complexité théorique.

L'intégration de l'apport de l'IA sans perdre l'authenticité de ma voix constituait un enjeu identitaire du projet. Ma validation systématique par mon expertise de terrain a permis de maintenir l'ancrage dans le réel et la pertinence des analyses. Le PROMPT 04 sur le style a joué un rôle crucial dans cette préservation de ma voix.

L'équilibre entre force critique et lisibilité demandait une stratégie rhétorique élaborée. Nous avons adopté une progression du consensuel au radical, préparant progressivement le lecteur aux analyses les plus incisives.

Enfin, gérer la violence symbolique révélée par notre analyse sans tomber dans le nihilisme exigeait de toujours articuler la critique à des propositions d'alternatives, maintenant ouverte la possibilité d'une transformation.

XI. Réflexions sur la collaboration humain-IA

Les forces de cette synergie se révèlent dans la complémentarité des apports. J'apporte une expertise de terrain irremplaçable, fruit de mes quarante années d'expérience dans l'orientation scolaire. Mon expérience vécue nourrit une intuition fine des mécanismes institutionnels, validant ou invalidant les hypothèses théoriques. L'orientation stratégique du projet et l'exigence éthique qui le sous-tend sont restées fondamentalement miennes.

L'IA apporte de son côté une capacité de systématisation qui permet de transformer mes intuitions en grilles d'analyse opératoires. Son amplification conceptuelle enrichit mes observations empiriques d'une profondeur théorique. Sa capacité à maintenir la cohérence sur la durée évite les contradictions involontaires dans un projet au long cours. Les connexions inattendues qu'elle établit entre des domaines apparemment distincts ouvrent des perspectives analytiques nouvelles. Sa vigilance méthodologique constante assure la rigueur du processus.

L'interaction entre ces deux pôles génère des émergences impossibles à anticiper. Des possibilités nouvelles surgissent, comme la modularisation méthodologique ou la généralisation théorique. Des découvertes imprévues adviennent, tel le chapitre 8 sur la matérialité. Une réflexivité augmentée s'installe, chaque partie servant de miroir critique à l'autre. Une créativité authentiquement dialogique émerge de cet espace de rencontre.

Les limites de cette collaboration doivent être reconnues lucidement. La nécessité absolue de mon expertise pour la validation reste incontournable. Le risque de production de "fausses évidences" sans ancrage empirique menace constamment. La tendance initiale de l'IA à inventer des références a nécessité une vigilance particulière. L'impossibilité pour l'IA de saisir certaines nuances contextuelles limite parfois la finesse de l'analyse. Mon expérience vécue reste fondamentalement irremplaçable et ne peut être simulée.

La nature de cette collaboration transcende tant la simple "utilisation" de l'IA comme outil que la délégation de la réflexion à la machine. Il s'agit d'une véritable co-élaboration où les rôles ne sont plus distincts mais entrelacés, où la création devient authentiquement dialogique, où les limites de chacun deviennent paradoxalement des forces pour l'ensemble, et où l'imprévu peut surgir de l'interaction même.

XII. Ce que révèle cette expérience

Sur la production du savoir, cette collaboration révèle des transformations profondes. La fin du génie solitaire s'annonce : la création intellectuelle devient collaborative par nature, non par choix mais par nécessité épistémologique. L'hybridation des compétences fait s'entrelacer intuition humaine et capacités computationnelles dans une tresse impossible à défaire. La réflexivité augmentée que permet l'IA génère une mise en abyme productive où la pensée se pense elle-même avec une acuité nouvelle. L'accélération n'est pas quantitative mais

qualitative : il ne s'agit pas de faire plus vite mais de faire mieux, d'atteindre une profondeur d'analyse inaccessible autrement.

Concernant l'écriture en sciences humaines, l'expérience démontre que l'IA peut être mobilisée non comme substitut mais comme partenaire intellectuel. Cette mobilisation exige cependant de maintenir une vigilance épistémologique constante, d'ancrer systématiquement l'analyse dans l'expérience de terrain et les données observables, d'assumer une transparence totale sur le processus et de préserver une exigence éthique intransigeante.

La question de la légitimité d'un savoir co-produit avec une IA se pose inévitablement. Notre réponse est que la légitimité ne vient plus de l'auteur mais de la rigueur du processus. La transparence sur la méthode, loin d'affaiblir la crédibilité, la renforce en permettant l'évaluation critique. L'hybridation humain-IA peut produire des analyses plus riches que chaque partie séparément, créant une plus-value cognitive authentique.

Sur le plan éthique, plusieurs principes émergent de cette expérience. La transparence exige d'explicitier le rôle de l'IA dans la production du savoir. L'attribution impose de reconnaître la nature collaborative du travail. La validation maintient mon expertise comme garde-fou ultime. La réflexivité interroge constamment les biais propres à chaque partie. La responsabilité reste fondamentalement mienne pour le contenu final.

XIII. La phase 2 : Vérification itérative (20-24 novembre 2025)

Immédiatement après la relecture et la réécriture de l'Introduction générale, s'est posée une question méthodologique : comment assurer que les chapitres existants restent cohérents avec les trois principes maintenant théorisés (réflexivité, pluralisation, transparence) ?

La solution rapide : un prompt de vérification articulé autour de cinq critères simples pour tester chaque chapitre :

1. Continuité théorique (Cefaï-Foucault-Bourdieu opère-t-il ?)
2. Ton (analytique vs accusatoire ?)
3. Pluralisation (l'un parmi ou LE seul ?)
4. Mécanismes (logiques institutionnelles vs malveillance consciente ?)
5. Tensions (questions ouvertes vs solutions simples ?)

Cas exemplaire : le Chapitre 9

Application du prompt sur le Chapitre 9 original a révélé des passages hautement accusatoires. La réécriture qui en a suivi a transformé radicalement le ton—de "dénonciation" à "analyse structurelle"—tout en préservant tous les arguments. Cette expérience a montré que **modifier le ton ne détruit pas la force critique**, elle la rend seulement plus rigoureuse.

Apprentissage collatéral

Cette phase a révélé une asymétrie productive : l'IA applique systématiquement un critère sur un chapitre entier ; l'humain reconnaît les "faux problèmes" (redondance par exemple) ; le

dialogue améliore les critères eux-mêmes. La théorie abstraite s'est transformée en **outil opératoire concret**.

PARTIE 2 : RÉFLEXION THÉORIQUE SUR LA COLLABORATION

I. Notre collaboration à la lumière des travaux contemporains

La chronique que nous venons de documenter pourrait sembler singulière, propre à ce projet spécifique. Or, elle s'inscrit en réalité dans un mouvement plus large de réflexion académique sur la collaboration humain-IA dans la recherche, particulièrement en sciences humaines et sociales. Comprendre cette expérience à travers le prisme des travaux contemporains permet de voir ce qui relève du particulier et ce qui annonce des transformations plus systémiques dans les pratiques de recherche.

Au cours de l'année 2024-2025, plusieurs chercheurs et collectifs (Aiello, Pilati, F., Munk, A. K., & Venturini, T., Joyce & Cruz, Boullier) ont publié des analyses convergentes sur la nature de cette collaboration. Ces travaux identifient cinq idées-forces qui structurent actuellement la réflexion académique sur la question. Notre expérience, documentée en Partie 1, offre l'occasion de vérifier, d'illustrer, et parfois de nuancer ces convergences théoriques.

Les cinq convergences académiques

1. L'IA comme partenaire d'augmentation, non de substitution

La littérature académique insiste fortement sur un point : l'IA générative n'est légitime en recherche que si elle *augmente* les capacités du chercheur plutôt que de les remplacer. Cette augmentation passe par l'automatisation de tâches lourdes (tri, codage, résumé, formulation d'hypothèses) afin de libérer du temps pour le travail interprétatif et théorique, celui qui requiert l'expertise humaine irremplaçable.

Notre expérience confirme et illustre ce principe. L'IA n'a jamais été positionnée comme auteure autonome, mais comme amplificatrice de mon intuition de terrain. Même au moment où le rôle de l'IA s'est complexifié (passant d'amplificateur à co-élaborateur), la direction stratégique du projet, les choix éthiques, l'ancrage empirique, et la validation finale ont toujours resté fondamentalement de mon ressort. L'IA a systématisé mes observations en grilles d'analyse, enrichi mes hypothèses de références théoriques, maintenu la cohérence sur la durée — autant de tâches cognitives importantes mais qui n'auraient de valeur que si elles servaient une intention théorique supérieure, la mienne.

Mais cette augmentation révèle quelque chose de plus profond : elle n'est pas quantitative (faire plus, plus vite) mais *qualitative* (penser à une plus grande profondeur). L'IA n'a pas produit davantage de chapitres, mais a permis une analyse plus fine, une théorisation plus rigoureuse, une systématique impossible à atteindre seul.

2. Une reconfiguration des méthodes qualitatives

Un second point de convergence dans la littérature : cette collaboration transforme les gestes mêmes du travail qualitatif. L'ethnographie et la recherche qualitative se définissaient traditionnellement par certaines pratiques — prise de notes, codage manuel, construction progressive de catégories, écriture de vignettes analytiques. La collaboration IA reconfigure ces gestes : la prise de notes devient dialogue avec des « interlocuteurs synthétiques », le codage devient application systématique de grilles d'analyse formalisées, la catégorisation devient construction itérative entre intuition humaine et clarification computationnelle.

Un second point de convergence dans la littérature : cette collaboration transforme les gestes mêmes du travail qualitatif. L'ethnographie et la recherche qualitative se définissaient traditionnellement par certaines pratiques — prise de notes, codage manuel, construction progressive de catégories, écriture de vignettes analytiques. La collaboration IA reconfigure ces gestes : la prise de notes devient dialogue avec des « interlocuteurs synthétiques », le codage devient application systématique de grilles d'analyse formalisées, la catégorisation devient construction itérative entre intuition humaine et clarification computationnelle.

Notre projet illustre cette reconfiguration. La création de nos quatre prompts successifs (décrits en détail en Partie 1) constitue elle-même une reconfiguration méthodologique. Ces prompts n'automatisent pas un processus existant : ils formalisent et systématisent un processus qualitatif complexe en structures explicites et reproductibles, tout en gardant la flexibilité et l'adaptabilité requises pour la recherche qualitative.

Cette reconfiguration n'est pas une perte : elle ouvre de nouvelles possibilités analytiques. Elle permet notamment des formes d'analyse qu'un chercheur seul n'aurait pas le temps de conduire.

3. Les risques spécifiques à cette collaboration

La littérature académique met aussi en garde contre des dérives possibles : illusion d'intimité avec l'IA (la traiter comme un partenaire doté d'intentions), confusion entre données empiriques et texte généré (oublier ce qui provient du corpus vs ce qui est produit par l'IA), projection involontaire de stéréotypes, et lissage des tensions ou contradictions du matériau au nom d'une fausse cohérence.

Nous avons rencontré plusieurs de ces risques et avons dû développer des vigilances spécifiques. Notamment :

- **Le risque d'hallucination théorique** : L'IA a tendance à inventer des références ou à les déformer. Cela nous a imposé une vérification systématique de chaque référence citée. Nous avons également dû apprendre à distinguer entre une hypothèse féconde suggérée par l'IA et une affirmation présentée comme un fait.
- **Le risque de lissage des contradictions** : L'IA, par sa nature même, tend à construire des synthèses, des cohérences. Or, un corpus sociologique est souvent contradictoire, tensions et ambiguïtés en sont les données mêmes. Nous avons dû explicitement demander à l'IA de maintenir plutôt que de résoudre certaines tensions.

- **Le risque de la fausse intimité** : Il est tentant de parler de "l'IA pense que..." ou "l'IA a découvert que...". Ce langage cache une réalité plus complexe : ce que l'IA produit n'est jamais une pensée autonome, c'est une transformation du texte du corpus selon les grilles que je fournis. Nous avons dû rester vigilant sur ce point terminologique.

Cette vigilance n'affaiblit pas la collaboration, au contraire. Elle la rend plus robuste en la rendant consciente de ses propres limites.

4. La réflexivité explicite comme condition de légitimité

Un quatrième consensus académique : la collaboration IA-humain ne gagne en crédibilité que si elle est documentée explicitement. Comment l'IA a-t-elle été utilisée ? À quelles étapes du processus ? Selon quels prompts ? Quels ont été les choix de validation ? Cette transparence n'affaiblit pas le travail : elle le renforce en le rendant évaluable.

C'est précisément ce que cette Appendice entreprend. En documentant précisément les quatre prompts, les phases de la collaboration (amplification, révélation, co-élaboration), les moments de bifurcation (découverte du CESE, restructuration théorique, phase de vérification), nous créons un dossier qui permet au lecteur de comprendre comment ce livre a été produit. Cette transparence est elle-même une contribution au champ : elle montre qu'une collaboration rigoureuse avec l'IA est possible sous certaines conditions.

5. Une opportunité pour inventer de nouvelles formes de méthode

Enfin, la littérature voit la collaboration IA-humain comme une bifurcation créative : l'opportunité d'inventer des méthodes « quali-quantitatives » hybrides, des dispositifs expérimentaux inédits (interlocuteurs synthétiques, sondes d'investigation, jeux de rôle avec modèles), des infrastructures méthodologiques adaptées aux questionnements nouveaux.

Notre approche des prompts comme « boîte à outils » réflexive s'inscrit exactement dans cette perspective. Les quatre prompts ne sont pas des recettes figées, mais des infrastructures méthodologiques adaptables et critiquables. Ils représentent une innovation : la formalisation d'une démarche qualitative complexe en structures explicites et reproductibles, sans perdre pour autant la flexibilité et l'adaptabilité.

Positionnement de notre expérience

Ces cinq convergences nous permettent de voir notre expérience de manière nouvelle. Elle n'est pas une anecdote méthodologique locale. Elle contribue à l'émergence d'un nouveau paradigme de recherche dans les sciences humaines : celui où la collaboration humain-IA devient une pratique légitime à condition d'être rigoureuse, transparente et réflexive.

Mais cette expérience soulève aussi des questions que la littérature n'a pas encore entièrement traitées. C'est à ces questions que nous allons maintenant nous consacrer.

II. Ce que notre expérience révèle au-delà des convergences

Les cinq convergences académiques que nous avons énoncées validaient largement notre approche. Mais cette expérience intense et concentrée sur deux mois. Cette évolution

temporelle n'est pas triviale, même concentrée sur deux mois et que nous documentons de manière réflexive, soulève aussi des questions plus nuancées, voire révèle certains décalages avec la littérature existante.

Le caractère processuel de la collaboration : une relation qui se transforme

Un premier apprentissage majeur : ce que les travaux académiques désignent par « collaboration IA-humain » n'est pas une relation stable. C'est au contraire un processus qui se transforme en fonction des enjeux qui émergent et des résultats produits.

Nous avons identifié trois phases distinctes dans notre travail, que nous avons documentées en Partie 1 :

Phase 1 : L'IA comme amplificateur. Au démarrage, l'IA fonctionne principalement comme amplificatrice de mon intuition. Je propose une idée (l'invisibilisation dans le discours éducatif), l'IA la systématise et la documente. Son rôle est d'étendre, de clarifier, de mettre en forme ce qui est déjà en germe chez moi. C'est une relation asymétrique : je reste le penseur, elle reste l'instrument.

Phase 2 : L'IA comme révélatrice. Le rôle devient progressivement plus dialectique. Je propose un argument, l'IA développe et interroge, je reviens à mon corpus pour valider ou réfuter. C'est une véritable conversation, un dialogue où chaque partie force l'autre à clarifier, à justifier, à approfondir.

Phase 3 : La co-élaboration véritable. Au-delà d'une certaine durée, quelque chose change. L'IA commence à suggérer des directions que je n'avais pas envisagées. Certaines découvertes tardives (le chapitre 8 sur la matérialité, l'articulation avec l'écologie) émergent véritablement de cette co-élaboration. Aucun de nous deux n'en est le seul auteur.

Cette évolution temporelle n'est pas triviale. Elle implique que les risques et les exigences ne sont pas les mêmes à chaque moment. Au début, le risque principal est la délégation aveugle : laisser l'IA faire sans vérifier, accepter ses propositions sans les confronter au corpus. À la fin, le risque devient inverse : la fusion identitaire, où on oublie que c'est toujours une collaboration et où on croit que l'IA « pense vraiment ».

Les exigences de vigilance se transforment aussi. En phase 1, il faut vérifier que l'IA ne déforme pas mon intuition. En phase 3, il faut vérifier que les émergences ne sont pas des hallucinations ou des projections.

La littérature académique tend à parler de « conditions de légitimité » de la collaboration comme si elles étaient stables et universelles. Notre expérience suggère qu'il faudrait plutôt parler de conditions évolutives de légitimité qui se transforment au fur et à mesure que la relation mûrit. Cette distinction est importante pour les chercheurs qui s'engagent dans une telle collaboration : ce qui compte comme une bonne pratique n'est pas identique dans les deux mois et dans les douze mois.

L'importance de la formalisation méthodologique : externaliser pour mieux voir

Un second apprentissage : les quatre prompts que nous avons créés ne sont pas simplement des « demandes à formuler à l'IA ». Ce sont des formalisations de ma démarche analytique elle-même.

Avant cette collaboration, comme la plupart des chercheurs en sciences humaines, j'avais une démarche analytique. Elle existait, elle fonctionnait. Mais elle restait largement implicite. Je savais ce que je faisais, mais j'aurais eu du mal à l'expliquer précisément, à l'enseigner, à la reproduire de manière systématique.

La création des prompts a forcé à externaliser ce savoir implicite. Pour que l'IA comprenne ce que je voulais, j'ai dû formuler précisément :

- Quels éléments je cherchais à identifier (régimes de visibilité/invisibilité)
- Selon quelles catégories (euphémisation, technicisation, fragmentation, dissociation)
- À quel niveau de généralité (un discours public quelconque ? spécifiquement éducatif ?)
- Avec quel niveau d'analyse (descriptif rapide ou théorisation approfondie ?)
- Selon quel style d'énonciation (analytique ou accusatoire ?)

Cette formalisation a un effet paradoxal intéressant. D'un côté, elle rigidifie : chaque texte du corpus doit passer par la même grille. De l'autre, elle rend plus adaptable : parce que les prompts sont explicites, ils peuvent être ajustés, critiqués, affinés en fonction des résultats.

Mais il y a plus. Cette formalisation ne rend pas visible ce que l'IA fait. Elle rend aussi visible ce que je fais. En définissant précisément les fonctions analytiques que je voulais (régimes de visibilité/invisibilité, mécanismes d'occultation, niveaux d'analyse), et en demandant leur formalisation en prompts structurés, j'ai été forcé de clarifier. Cela crée une forme nouvelle et amplifiée de réflexivité. Je n'analyse pas seulement le corpus, j'analyse aussi mes propres méthodes d'analyse.

Pour les sciences humaines, qui ont longtemps valorisé le « savoir tacite » et la « compétence artisanale », cela représente un changement de paradigme. Pas un abandon du savoir tacite, mais sa mise en dialogue avec une explicitation formelle. Le résultat : une meilleure compréhension à la fois du corpus ET de ma propre démarche.

La question de l'émergence : créer du neuf dans la rencontre

Un troisième apprentissage, peut-être le plus inattendu : la collaboration IA-humain ne produit pas seulement de l'augmentation. Elle produit aussi de l'émergence.

L'augmentation, c'est quand l'IA étend ou systématise quelque chose qu'on avait déjà identifié. « J'avais une intuition floue, l'IA l'a clarifiée et systématisée. » C'est utile, mais c'est une opération prévisible.

L'émergence, c'est quand quelque chose de qualitativement nouveau surgit de l'interaction entre le chercheur et l'IA, quelque chose que ni l'un ni l'autre n'auraient produit séparément.

Nous avons observé plusieurs cas d'émergence véritable :

La découverte de l'invisibilisation différenciée. Au départ, je pensais à UNE invisibilisation : le discours éducatif occultant certaines réalités. L'analyse du rapport CESE a révélé quelque chose de plus nuancé : il n'existe pas une invisibilisation uniforme mais des invisibilisations croisées selon les positions. Chaque acteur (OCDE, ministère, syndicats, chercheurs critiques) a ses propres angles morts et ses propres zones de lucidité partielle. Cette découverte n'était pas dans mon intuition initiale. Elle a émergé de la confrontation systématique avec le corpus, guidée par les outils qu'on avait créés ensemble.

L'émergence de la grille quadripartite (euphémisation, technicisation, fragmentation, dissociation). Je cherchais les mécanismes d'invisibilisation. L'IA, en appliquant systématiquement le prompt à des dizaines de passages, a révélé que certains motifs revenaient obsessionnellement. En nommant ces motifs, on a créé un outil analytique qu'aucun de nous n'avait anticipé. C'est un cas classique d'émergence : la régularité apparaît par saturation du corpus, pas par déduction théorique.

Le chapitre 8 sur la matérialité scolaire. Cela aurait pu être un chapitre mineur, une note de bas de page. Mais en l'approfondissant ensemble, nous avons réalisé que c'était central : le discours éducatif évite précisément de parler de l'infrastructure matérielle des écoles parce que c'est là qu'on voit l'inégalité brute. Cette découverte tardive a remis en question l'organisation du livre entière. Elle n'était prévisible que rétrospectivement.

L'articulation entre écologie et système éducatif. Au moment où on récrivait l'introduction, une question a surgi : pourquoi le discours éducatif ne parle-t-il jamais du contexte écologique dans lequel l'école opère ? Cette articulation n'était pas planifiée. Elle a émergé d'une conversation entre mon expérience de terrain et les synthèses de l'IA.

Ces émergences partagent une caractéristique : elles n'étaient pas pré-programmées dans le projet. Elles n'auraient pas pu être anticipées en commençant. Elles sont le résultat spécifique de cette rencontre entre une expertise humaine située et la capacité computationnelle de l'IA à identifier des régularités et des connexions.

La littérature parle souvent de « nouvelles possibilités analytiques » ouvertes par la collaboration IA. Notre expérience précise cette formule : ces possibilités ne sont pas juste des extensions quantitatives (« on peut traiter plus de textes ») ni des accélérations (« on peut aller plus vite »). Ce sont des créations qualitativement nouvelles qui dépassent ce que chaque partie aurait pu produire en travaillant isolément.

Cela a des implications pour la théorie de la connaissance : cela suggère que la collaboration n'est pas un simple moyen d'amplifier une connaissance qui existe déjà chez le chercheur. C'est une véritable co-production épistémologique.

Le rôle crucial de la validation empirique : trois formes de vérification

Un quatrième apprentissage concerne la validation. La littérature académique sur la collaboration IA-humain insiste beaucoup sur la nécessité de la validation empirique : vérifier

que ce que propose l'IA correspond au corpus. C'est bien. Mais notre expérience montre que c'est plus complexe.

Dans la recherche traditionnelle en sciences humaines, un chercheur valide ses hypothèses en les confrontant au corpus : « Je proposais X, je relis le corpus, est-ce que X est bien documenté ? »

Dans notre collaboration, la validation prend trois formes distinctes et complémentaires :

Validation empirique classique. L'IA propose une analyse (« Ce passage pratique l'euphémisation »). Je revérifie dans le corpus : c'est bien le cas, ou ce n'est pas pertinent ? Cette forme de validation est celle que la littérature met en avant. C'est important, mais c'est aussi la plus facile : elle ne demande que de revérifier.

Validation théorique. L'IA propose une interprétation. Je vérifie : est-ce cohérent avec le cadre théorique qu'on a établi ? Est-ce que ça s'articule bien avec Cefaï, Foucault, Bourdieu ? Est-ce que ça ne contredit pas une conclusion qu'on a établie ailleurs dans le livre ? Cette forme de validation exige une maîtrise globale du projet.

Validation expérientielle. L'IA propose une analyse. Je me demande : cela correspond-il à mon expérience de terrain ? Après quarante ans comme directeur de CIO, reconnaissais-je cette réalité ? C'est une forme de validation qu'on tend à minorer dans les sciences humaines contemporaines, par crainte du subjectivisme. Mais c'est une erreur. L'expérience vécue du chercheur n'est pas un biais à éliminer : c'est une ressource de connaissance à mobiliser consciemment, à condition de le faire de manière réflexive.

Le rôle fondamental de cette validation expérientielle est apparu clairement lors de l'analyse du CESE. L'IA synthétisait ce que disait le rapport. Mais j'avais aussi l'expérience de ce qu'il ne disait pas, des silences spécifiques que je pouvais identifier parce que je connaissais le système éducatif par le vécu. Cette connaissance incarnée a été décisive pour identifier l'invisibilisation différenciée.

Notre collaboration a systématisé cette troisième forme de validation. Chaque proposition significative de l'IA passe par ce triple filtre : empirique (corpus), théorique (cadre), expérientiel (connaissance incarnée). C'est cela qui rend la collaboration robuste, pas la certification de chaque phrase par le corpus.

Cette clarification a une implication pour la théorie de la connaissance : elle suggère que dans les sciences humaines et sociales, il existe une forme de connaissance qui n'est pas réductible aux données et à la théorie, c'est la connaissance incarnée du chercheur. Et que cette connaissance n'est pas une faiblesse si elle est mobilisée consciemment et réflexivement.

III. L'éthique et la responsabilité dans une collaboration hybride

Si la collaboration IA-humain devient une pratique légitime dans la recherche académique, il faut aussi clarifier les enjeux éthiques spécifiques qu'elle soulève. Ces enjeux ne sont pas des détails secondaires : ils touchent aux fondations même de ce qu'on peut appeler de la « recherche responsable ».

La question de l'attribution et de l'autorialité

Un enjeu classique qui devient plus aigu avec l'IA générative : qui est l'auteur d'un texte produit en collaboration avec une machine ?

La convention académique traditionnelle est claire et stable : un article écrit par une personne a cette personne pour auteur. Même avec des outils techniques (traitement de texte, logiciels de calcul, bases de données), on ne remet pas en question l'autorialité. Le texte reste signé par la personne qui l'a rédigé.

Mais avec l'IA générative, la question devient plus complexe et exige qu'on la pose explicitement. L'IA a en effet contribué substantiellement à :

- la formulation précise de nombreuses phrases et paragraphes
- l'organisation logique et rhétorique de sections entières
- l'identification de connexions théoriques et d'articulations entre idées
- la clarification de passages obscurs ou mal formulés
- la suggestion de développements qu'on n'avait pas envisagés

Est-ce que cela fait de l'IA une « co-auteure » ? Devrait-on écrire « Bernard Desclaux et Claude » sur la couverture du livre ?

Notre position sur cette question est claire : non, l'IA n'est pas co-auteure, et je demeure l'auteur unique du livre. Mais cette conclusion exige qu'on la justifie rigoureusement, car elle n'est pas évidente.

Deux arguments nous semblent décisifs :

L'absence de responsabilité. Une auteure assume la responsabilité éthique, scientifique et légale de ce qu'elle affirme. Elle se tient derrière ses affirmations. Si on découvre une erreur factuelle, une conclusion non justifiée, une utilisation incorrecte d'une théorie, c'est elle qui doit répondre. Elle engage sa réputation.

L'IA ne peut pas assumer cette responsabilité. Elle ne « croit » à rien de ce qu'elle écrit. Elle produit du texte selon des processus probabilistes complexes, mais elle n'a ni convictions, ni engagement, ni réputation à perdre. Si le texte contient une erreur, ce n'est pas une « erreur de l'IA » au sens où c'est une erreur pour laquelle on pourrait la blâmer : c'est un artefact de son fonctionnement statistique.

La responsabilité de ce qui est écrit retombe donc entièrement sur moi. C'est moi qui ai choisi d'utiliser l'IA, qui ai formulé les prompts, qui ai validé ou rejeté ses propositions, qui ai décidé de ce qui entrerait finalement dans le texte. Ma signature au bas du livre signifie : « Je suis responsable de chaque affirmation ici. »

L'absence d'intention directrice. Une auteure poursuit une intention, une vision, un objectif. Elle ne sait pas d'avance où son travail va la mener, mais elle a une direction générale, des principes qui la guident.

L'IA n'a pas d'intention directrice. Elle n'a pas d'objectif à long terme. Elle répond à des demandes spécifiques (les prompts) en produisant du texte qui correspond statistiquement à ce qui a été demandé. Les demandes et l'objectif global viennent entièrement du chercheur.

En d'autres termes, l'IA n'a pas de « point de vue » qui viendrait enrichir ou compliquer le mien. Elle n'apporte pas une vision alternative que je devrais négocier. Elle apporte une capacité de systématisation, une rigueur dans la formalisation, une aide précieuse. Mais la pensée, au sens fort du terme, reste mienne.

Cela dit, cette conclusion crée une obligation éthique qui ne peut pas être contournée : la **transparence absolue** sur le rôle de l'IA dans la production du texte.

Ne pas mentionner que l'IA a été utilisée serait malhonnête intellectuellement. Ce serait tromper le lecteur sur la nature réelle du processus de création. On attendrait d'un chercheur qui a utilisé un assistant humain ou qui a reçu une aide substantielle qu'il le mentionne en remerciements ou en note. L'utilisation de l'IA mérite au minimum le même traitement, sinon plus.

C'est précisément ce que cette Appendice entreprend : documenter explicitement le rôle de l'IA, les prompts créés et leur évolution, les phases de la collaboration, les moments de bifurcation, les risques identifiés et gérés. Cette documentation est elle-même un acte éthique. Elle dit au lecteur : « Voici comment ce livre a été produit. Vous pouvez juger par vous-même si la méthode vous semble valide. »

La question de la responsabilité cognitive

Un enjeu plus subtil et peut-être plus important : à partir du moment où on se fait assister par l'IA, on abandonne une partie du « travail de pensée ». Quel est le coût de cette délégation ?

Dans la recherche traditionnelle, le chercheur passe des heures, des semaines à :

- relire laborieusement le corpus, parfois plusieurs fois
- prendre des notes manuelles, reformuler dans ses propres mots
- relire ses notes pour identifier des patterns
- chercher les connexions théoriques entre des observations éparses
- essayer différentes formulations pour exprimer une idée complexe
- organiser et réorganiser un argument pour qu'il soit cohérent

Ce travail long, fastidieux, parfois frustrant a une valeur cognitive réelle. Il crée une intimité incarnée avec le matériau. À force de relire le corpus, on connaît les subtilités du texte, les exceptions, les tensions. On ne peut pas les oublier. Cette connaissance devient très fine et très nuancée.

De plus, ce travail lent force à penser profondément. Quand on doit expliquer soi-même une idée, la reformuler, la justifier, on découvre les failles de sa pensée. On est obligé de clarifier, de préciser, d'approfondir. La lenteur devient une alliée de la profondeur.

En confiant ce travail à l'IA, on gagne en efficacité quantitative. On traite plus de textes, plus rapidement. Mais on perd potentiellement cette intimité cognitive, cette connaissance incarnée qu'on ne peut acquérir que par un travail lent et répétitif.

Comment avons-nous navigué dans cette tension ? En distinguant rigoureusement deux formes d'activité :

Les tâches qui peuvent être déléguées (avec vigilance) : systématisation de patterns, organisation formelle, clarification des formulations, vérification de la cohérence interne, identification de connexions théoriques entre des domaines éloignés, production de premières versions ou de synthèses.

Ces tâches sont cognitives, elles demandent un traitement intelligent du texte. Mais elles peuvent être outillées. C'est précisément ce que font les quatre prompts : transformer une tâche complexe en un processus systématisable.

Les tâches qui ne peuvent pas être déléguées : la décision sur ce qui compte analytiquement, c'est-à-dire ce qui mérite d'être expliqué plutôt que marginalisé ; la validation empirique contre le corpus ; la validation théorique contre le cadre conceptuel ; la direction générale de l'argument ; les choix éthiques sur la tonalité et les responsabilités envers les acteurs du système analysé.

Ces tâches demandent non seulement de l'intelligence analytique, mais de la sagacité, du jugement, une compréhension incarnée du contexte. On ne peut pas les automatiser sans risquer de perdre ce qui fait la valeur de la recherche.

Le résultat de cette distinction : j'ai moins passé de temps sur les tâches fastidieuses de systématisation manuelle, mais j'ai passé plus de temps sur le travail véritablement interprétatif. C'est une réallocation du travail cognitif vers ce qui est irremplaçable. Paradoxalement, utiliser l'IA pour les tâches routinières m'a permis de consacrer plus d'énergie mentale aux questions difficiles.

Mais il faut être lucide : on a quand même perdu quelque chose. Je n'ai pas la même intimité avec certains passages du corpus qu'un chercheur qui aurait pris des notes manuelles ligne par ligne. Il y a un prix à payer pour cette efficacité. Ce prix est acceptable si on en est conscient et si on compense par d'autres formes de rigueur (comme la validation systématique).

La question de la transparence méthodologique

Une obligation éthique majeure découle des deux points précédents : la transparence méthodologique.

Si vous êtes l'auteur mais que vous avez utilisé l'IA, vous devez le dire clairement. Si vous êtes responsable du contenu mais que vous avez délégué certaines tâches, vous devez clarifier lesquelles et comment.

Cette transparence n'affaiblit pas la légitimité du travail académique. Au contraire : elle la renforce considérablement en montrant que vous travaillez consciemment avec ces outils plutôt que de les utiliser en cachette ou de feindre l'ignorance sur leur rôle.

Dans le contexte actuel où l'usage de l'IA générative en recherche fait débat, cette transparence est particulièrement importante. Si on cache l'utilisation de l'IA, on alimente la méfiance : les lecteurs découvriront tôt ou tard que l'IA a été utilisée (grâce aux styles caractéristiques ou à d'autres indices), et ils concluront qu'on a voulu le dissimuler. C'est pire que de le reconnaître d'emblée.

En revanche, si on documente explicitement le rôle de l'IA dès le départ, on dit clairement : « Voici la méthode que j'ai utilisée. Voici pourquoi. Voici comment j'ai géré les risques. Vous pouvez évaluer par vous-même si cela vous semble valide. »

Notre Appendice fonctionne comme une déclaration transparente de méthode. Elle documente :

- Quand l'IA a été utilisée : à quelles phases du projet, pour quels types de tâches
- Comment elle a été utilisée : les quatre prompts, leur architecture, leur évolution, les critères qu'ils implémentaient
- Ce qui a toujours resté sous mon contrôle : la direction stratégique du projet, la validation finale de chaque affirmation majeure, la responsabilité éthique, les choix d'orientation
- Quels risques ont été identifiés : hallucination, lissage des contradictions, fusion identitaire, etc.
- Comment ces risques ont été gérés : quelles vigilances ont été mises en place, quels protocoles de vérification
- Quelle plus-value cette collaboration a apportée : les émergences qu'elle a permises, les innovations méthodologiques qu'elle a produites

Cette documentation n'est pas un luxe rhétorique. C'est une condition de légitimité académique. Un chercheur responsable doit pouvoir documenter sa méthode. Si on utilise l'IA, on doit documenter ce rôle aussi explicitement que n'importe quel autre outil ou protocole.

Une nouvelle forme de réflexivité

Il y a un effet inattendu à cette transparence méthodologique : elle crée une forme nouvelle et amplifiée de réflexivité scientifique.

La réflexivité en sciences sociales, c'est d'ordinaire quand le chercheur interroge sa position, ses biais, l'influence de son point de vue sur ses résultats. C'est important et largement pratiqué. Mais ça reste souvent au niveau théorique : on énonce ses biais sans toujours les traiter.

Quand on documente la collaboration avec l'IA, on est forcé d'une réflexivité plus systématique. On doit expliquer non seulement ce qu'on pense mais comment on pense, c'est-à-dire les méthodes, les outils, les protocoles qu'on utilise. On doit rendre visible le processus de la recherche, pas seulement ses résultats.

Cela rend la recherche plus évaluable. Un lecteur ou un évaluateur peut maintenant juger non seulement des conclusions (sont-elles justes ?) mais aussi du processus (était-il rigoureux ?). C'est un gain réel en termes de rigueur scientifique.

Plus profondément, cela crée une opportunité pour dépasser l'opposition entre « subjectivité » et « objectivité » qui a longtemps paralysé les sciences humaines. On n'essaie plus d'éliminer complètement le chercheur du processus (impossible de toute façon). On documente plutôt comment le chercheur et ses outils interagissent pour produire la connaissance. C'est une approche plus honnête et plus productive.

IV. Les questions ouvertes et l'horizon théorique

Après six mois de collaboration avec l'IA sur différents projets (blog, conception d'ouvrages), et particulièrement cette documentation réflexive du livre, une chose devient claire : nous avons résolu certaines questions pratiques, mais nous en avons aussi ouvert d'autres, plus profondes, qui dépassent largement le cadre de ce projet singulier. Ces questions concernent la nature même de la connaissance en train de se transformer.

Qu'est-ce que « penser avec » une machine ?

La formulation même pose problème et révèle notre malaise conceptuel. Quand on dit « penser avec l'IA », on emploie une métaphore qui cache autant qu'elle révèle.

« Penser » suppose une activité consciente, intentionnelle, capable de délibération. « Avec » suggère un véritable partenariat, une relation symétrique entre deux agents. Or, ni l'un ni l'autre n'est tout à fait juste quand on parle de collaboration avec une IA générative.

L'IA ne pense pas au sens où elle n'a pas de conscience, pas de délibération, pas de doute existentiel. Elle traite des textes selon des patterns probabilistes appris sur un corpus massif. C'est de l'intelligence au sens computationnel, pas au sens philosophique.

Mais en même temps, dire simplement « je pense et l'IA m'aide » serait trop réducteur. Ce qui se passe est plus étrange : une forme de cognition distribuée où le raisonnement ne réside ni entièrement dans ma tête ni dans la machine, mais dans l'interaction entre nous deux.

Quand je fais écrire un prompt, j'externalise une partie la formalisation de ma pensée. Quand l'IA répond, elle produit quelque chose qui n'était pas dans ma tête, mais qui n'aurait pas été produit sans ma direction. Quand je valide ou je réfute sa proposition, je crée un retour qui la force à générer autre chose.

Ce processus n'est pas de la « pensée » au sens classique, mais ce n'est pas non plus de la simple utilisation d'outil. C'est quelque chose d'intermédiaire qu'on n'a pas bien de vocabulaire pour décrire. On pourrait appeler cela une cognition augmentée distribuée : une forme de raisonnement qui dépasse les capacités de chaque agent pris isolément.

La question théorique devient : que reconnaître comme une véritable forme de pensée ? Si on définit la pensée seulement par la conscience et l'intentionnalité, alors l'IA ne pense pas et ce qu'on fait ensemble n'est pas vraiment de la pensée. Mais cette définition paraît de plus en plus insuffisante. Peut-être faut-il développer une théorie de la cognition qui reconnaisse des

formes distribuées et augmentées de pensée sans pour autant attribuer à l'IA une conscience qu'elle n'a pas.

Comment transformer la notion d'auteur ?

Une question corollaire : qu'advient-il de la notion d'auteur quand la création intellectuelle devient véritablement collaborative, pas juste avec d'autres humains, mais avec des artefacts technologiques ?

Pendant des siècles, l'auteur était une figure stable : une personne qui écrivait, pensait, était responsable. Certes, les auteurs s'inspiraient d'autres œuvres, dialoguaient avec d'autres penseurs, utilisaient des sources. Mais la créativité restait attribuée à la personne.

Avec l'IA générative, cette figure devient trouble. Je suis toujours responsable du contenu, mais le texte n'a pas pour autant émergé entièrement de ma conscience. Il y a une vraie part de création qui vient de l'interaction, pas de moi.

Est-ce qu'on devrait inventer une nouvelle catégorie ? Un « auteur augmenté » qui serait clairement l'auteur, mais qui reconnaîtrait explicitement le rôle de l'IA ? Certaines maisons d'édition commencent à explorer des mentions du type « Écrit par X avec l'assistance de l'IA ». Est-ce la bonne direction ?

Ou au contraire, faudrait-il distinguer nettement entre :

- L'auteur : la personne responsable, celle qui signe
- Les contributeurs : les outils utilisés (y compris l'IA)

Sur le modèle des remerciements : on ne dit pas « écrit par moi et par ma cafetière », mais on pourrait dire « écrit par moi avec l'assistance de l'IA Claude, selon les protocoles décrits en Appendice ».

Ces questions ne sont pas purement nominales. Elles touchent aux structures légales et éditoriales (qui détient les droits d'auteur ?), aux structures académiques (qui reçoit le crédit ?), et à la compréhension culturelle de ce qu'on valorise (la création seule, ou la créativité augmentée ?).

Quels nouveaux biais crée cette collaboration ?

La littérature académique sur l'IA s'inquiète régulièrement des biais : l'IA reproduit les biais présents dans ses données d'entraînement, elle amplifie les inégalités, elle encode les idéologies dominantes.

Tout cela est vrai. Mais notre expérience soulève une question différente : quels biais nouveaux la collaboration IA-humain elle-même crée-t-elle ?

Nous avons identifié quelques pistes :

Le biais de systématicité. L'IA a tendance à chercher des régularités, des patterns récurrents. Elle lisse les exceptions, les cas marginaux, les phénomènes rares. Or, en sciences sociales, ce

sont souvent les exceptions qui sont théoriquement les plus intéressantes. Elles révèlent les limites d'une théorie, elles forcent à revoir les catégories.

Le risque : en utilisant l'IA pour identifier des patterns, on amplifie les phénomènes fréquents et on marginalise les rares. Cela peut produire une compréhension biaisée du corpus.

Le biais de cohérence. L'IA cherche à construire des textes cohérents, sans contradictions internes. C'est naturel pour un générateur de texte. Mais un corpus réel est souvent contradictoire. Les acteurs se contredisent, les politiques poursuivent des objectifs incompatibles, les idées coexistent dans une tension productive.

Le risque : en soumettant le corpus à une IA qui cherche la cohérence, on perd les tensions, les contradictions, les ambivalences qui sont peut-être ce qu'il y a de plus important.

Le biais de continuité. L'IA extrapolera naturellement à partir des patterns observés. Elle cherchera à compléter ce qui semble incomplet, à continuer ce qui s'arrête. Mais l'absence de quelque chose peut être signifiante. L'incomplétude peut révéler une rupture importante.

Le risque : on artificialise une continuité qui n'existe pas dans le réel.

Comment avons-nous géré ces biais ? En grande partie par la vigilance et la validation systématique. Mais on n'est pas sûr de les avoir tous éliminés. Ce qu'on peut dire, c'est qu'on en a été conscients et qu'on a essayé de les compenser.

Mais cela soulève une question théorique majeure : qu'est-ce qu'un « biais » d'IA dans le contexte de la recherche qualitative ? Est-ce un défaut qu'il faut éliminer, ou est-ce une caractéristique inévitable d'un certain type d'outil qu'on doit apprendre à négocier ?

Comment évaluer la qualité d'un savoir hybride ?

Une question plus méthodologique mais fondamentale : selon quels critères peut-on juger qu'une recherche menée en collaboration avec l'IA est de bonne qualité ?

Les critères académiques traditionnels s'appliquent-ils ? Une recherche est-elle bonne si elle est rigoureuse, si elle repose sur du matériau solide, si elle est justifiée par des arguments théoriques, si elle est transparente sur sa méthode ?

Oui, probablement. Mais l'utilisation de l'IA ajoute peut-être des critères supplémentaires :

- La traçabilité : peut-on retracer exactement qui a fait quoi dans le processus de création ?
- La reproductibilité : un autre chercheur utilisant les mêmes prompts sur les mêmes données obtiendrait-il des résultats similaires ?
- La transparence méthodologique : les choix d'utiliser l'IA, le type d'IA, les prompts spécifiques, tout cela est-il clairement documenté ?
- La robustesse contre les biais : quelles vigilances ont été mises en place pour identifier et compenser les biais propres à cette collaboration ?

Mais il y a aussi une question plus délicate : comment juger de la créativité et de l'innovation d'un travail qui a utilisé l'IA ?

Un argument dit : « L'IA a produit du texte, donc c'est moins créatif, moins original. » Mais notre expérience suggère l'inverse. Précisément parce que l'IA systématise les données, j'ai pu consacrer plus d'énergie à l'interprétation créative, à la théorisation, aux questions d'émergence. Le résultat n'est pas moins créatif ; il est créatif différemment.

Comment évaluer cela ? On ne peut pas appliquer les mêmes critères qu'à une recherche traditionnelle. On doit développer de nouveaux critères d'évaluation spécifiques à la recherche augmentée par l'IA.

Vers une épistémologie nouvelle

Au-delà de ces questions ouvertes, notre expérience suggère quelque chose de plus fondamental : nous sommes peut-être à un point d'inflexion dans la théorie de la connaissance.

Pendant des siècles, la connaissance a été pensée comme le fruit d'un processus mental individuel : le savant, seul dans son cabinet, qui pense, observe, déduit. Les outils (microscope, télescope, balance) n'étaient que des extensions de la perception, pas vraiment des partenaires du processus cognitif.

Avec l'IA générative, quelque chose change. L'outil ne prolonge plus simplement notre perception ; il devient un partenaire dans la cognition elle-même. Il ne se contente pas d'amplifier ce qu'on peut faire ; il crée des formes de pensée qui n'existaient pas avant.

Cela ouvre peut-être la voie à une épistémologie de la collaboration : une théorie de la connaissance qui reconnaît que le savoir se produit de plus en plus dans des espaces entre les humains et les machines, que la cognition est distribuée, que la responsabilité est partagée mais pas symétrique.

Cette épistémologie ne serait pas un rejet de la rigueur ou de la responsabilité académique. Au contraire : elle imposerait une rigueur accrue dans la documentation, la transparence, la réflexivité. Mais elle reconnaîtrait que le sujet connaissant n'est plus le savant solitaire, c'est un collectif humain-technologique.

Pour les sciences humaines, cela représente un changement profond. Nous avons longtemps insisté sur la singularité, l'incarnation, l'expérience vécue comme ressources de connaissance. L'IA ne supprime pas cela. Mais elle force à le penser en dialogue avec une forme de cognition très différente : distribuée, statistique, non-incarnée.

Le défi n'est pas de choisir entre l'une et l'autre, mais d'apprendre à penser les deux ensembles.

Conclusion : vers une recherche réflexive et transparente à l'ère de l'IA

Cette Partie 2 de l'Appendice a entrepris de situer notre collaboration particulière dans un horizon plus large : celui des débats contemporains sur la légitimité de la collaboration IA-humain dans la recherche académique.

Nous avons d'abord montré comment notre expérience singulière valide les cinq convergences identifiées par la littérature académique récente : l'IA comme partenaire d'augmentation, la reconfiguration des méthodes qualitatives, l'identification des risques spécifiques, la nécessité de la réflexivité explicite, et l'opportunité d'inventer de nouvelles formes de méthode.

Mais nous avons aussi montré que cette expérience révèle au-delà de ces convergences. Elle met en lumière le caractère processuel de la collaboration, qui se transforme au fil du temps en trois phases distinctes. Elle souligne l'importance de la formalisation méthodologique, qui externalise et clarifie ce qui resterait implicite. Elle révèle la capacité d'émergence : non pas simplement l'amplification de ce qu'on pense déjà, mais la création de savoirs qualitativement nouveaux qui n'auraient pu être produits isolément. Et elle insiste sur le rôle crucial d'une triple validation : empirique, théorique, et expérientielle.

Sur le plan éthique, nous avons affirmé des positions claires : l'auteur unique reste responsable, mais cette responsabilité exige une transparence absolue sur le rôle de l'IA. Cette transparence ne faiblit pas la légitimité académique ; elle la renforce en créant une réflexivité amplifiée qui rend visible non seulement ce qu'on pense, mais comment on pense.

Enfin, nous avons posé les questions qui dépassent ce projet : qu'est-ce que « penser avec » une machine ? Comment redéfinir l'auteur à l'ère de la collaboration humain-IA ? Quels nouveaux biais cette collaboration crée-t-elle ? Comment évaluer la qualité d'un savoir hybride ? Vers quelle pédagogie se tourner ? Et plus fondamentalement, comment penser les enjeux de pouvoir et d'accès qui structurent l'utilisation de ces outils ?

Ce que cette expérience enseigne

Au-delà des réponses, ce que cette longue documentation enseigne est peut-être plus important que les conclusions elles-mêmes. Elle enseigne que la collaboration IA-humain n'est pas une technique neutre qu'on appliquerait de la même manière dans tous les contextes. C'est une pratique qui exige de la réflexivité, de la vigilance, et une transformation profonde de la manière dont on conçoit le processus de recherche.

Elle enseigne que la transparence méthodologique, loin d'être une charge ou une confession d'insuffisance, devient un élément central de la rigueur scientifique. Documenter comment on pense, c'est rendre son travail évaluable. C'est créer les conditions pour que d'autres puissent en apprendre, le contester, l'améliorer.

Elle enseigne aussi que l'expérience vécue du chercheur ne disparaît pas dans la collaboration avec l'IA. Au contraire, elle devient plus importante. C'est l'expertise incarnée, les quarante ans dans un système éducatif, la compréhension des silences et des angles morts, qui donnent

sens à ce que produit l'IA. Loin de remplacer l'humain, l'IA révèle son caractère irremplaçable.

L'enjeu plus large

Cette expérience s'inscrit dans un moment historique où les sciences humaines et sociales sont confrontées à une transformation majeure de leurs outils. Il ne s'agit pas simplement de l'adoption d'une nouvelle technologie, comme cela s'est produit avec les outils statistiques ou les logiciels d'analyse qualitative. C'est quelque chose de plus profond : une redéfinition du sujet connaissant lui-même.

Pendant longtemps, ce sujet était pensé comme l'individu : le chercheur solitaire, responsable, autonome. Cette figure demeure importante. Mais elle n'est peut-être plus la seule. Le sujet connaissant à l'ère de l'IA générative pourrait être un collectif humain-technologique : une entité distribuée, réflexive, exigeant de la transparence et de la documentation.

Cette transformation n'est pas à craindre si elle est gérée consciemment. Elle n'est problématique que si elle se fait en cachette, sans réflexivité, en prétendant que l'IA n'intervient pas là où elle intervient réellement.

Pour une recherche responsable

Quelle que soit la forme que prenne la collaboration IA-humain dans les années à venir, trois principes nous semblent non-négociables pour que cette collaboration reste une forme légitime de recherche scientifique :

Premièrement, la responsabilité assumée. Le chercheur doit rester responsable de tout ce qui est dit en son nom. Cette responsabilité exige une validation constante, une vigilance continue, une refondation du savoir sur l'expérience et le corpus plutôt que sur la confiance accordée à la machine.

Deuxièmement, la transparence radicale. Tout ce qui a été délégué à l'IA doit être clairement documenté. Non pas pour culpabiliser le chercheur d'avoir utilisé l'IA, mais pour que le lecteur puisse évaluer la validité du processus. Cette transparence est une force, pas une faiblesse.

Troisièmement, la réflexivité systématisée. La collaboration avec l'IA doit devenir une occasion de clarifier comment on pense, comment on analyse, comment on produit du savoir. Cette clarification bénéficie à la recherche qu'on mène avec ou sans IA.

Une invitation plutôt qu'un modèle

Nous ne prétendons pas que notre méthode soit la seule valide ou qu'elle doit être reproduite mécaniquement par d'autres. Chaque collaboration IA-humain sera singulière, fonction des disciplines, des corpus, des questions posées, des sensibilités engagées.

Mais nous invitons les chercheurs qui s'engagent dans une telle collaboration à :

- Documenter précisément leur processus

- Identifier et nommer les risques spécifiques à leur domaine
- Mettre en place des protocoles de validation adaptés à leurs enjeux
- Assumer la responsabilité de ce qu'ils affirment
- Partager ce qui fonctionne et ce qui ne fonctionne pas

C'est par l'accumulation de ces expériences, documentées réflexivement, que pourra émerger une véritable épistémologie de la collaboration.

L'horizon

Nous concluons en revenant au point de départ : le livre « Les éléphants du discours. L'invisibilisation dans le système éducatif français » n'aurait probablement pas existé sans cette collaboration avec l'IA. Pas parce que l'IA aurait été indispensable dans l'absolu, mais parce que cette collaboration spécifique a permis quelque chose que les capacités isolées du chercheur n'auraient pas réalisé : une systématique alliée à l'interprétation, une rigueur conceptuelle alliée à l'expérience incarnée, une transparence méthodologique alliée à la profondeur théorique.

Ce que nous espérons, c'est que cette Appendice, en documentant précisément ce processus, contribue à montrer qu'une collaboration consciemment menée entre un chercheur et une IA générative peut produire une recherche rigoureuse, légitime, et potentiellement innovante.

Car au final, l'enjeu n'est pas de savoir si l'IA doit être utilisée en recherche. L'IA sera utilisée, inévitablement, par de plus en plus de chercheurs. L'enjeu est de savoir comment l'utiliser : de manière réflexive ou aveugle, transparente ou cachée, responsable ou irresponsable. C'est à cette question que cette Appendice a tenté de répondre.

Épilogue

En relisant cette Appendice qui documente notre collaboration sur deux mois, une dernière observation s'impose. Entre les premiers posts conçus pour mon blog et la version de cet ouvrage que vous lisez, peu de temps s'est écoulé, mais la puissance de l'IA a permis de produire quatre versions différentes de cette problématique, chacune marquant une amélioration dans ma compréhension. Le livre que vous tenez n'est donc pas une version finale et définitive, mais plutôt un moment, un état dans un processus de production qui pourrait continuer. D'autres versions pourraient suivre, approfondissant certains aspects, réorganisant d'autres.

C'est peut-être là aussi un enseignement : à l'ère de l'IA générative, la notion même de « version finale » devient poreuse. Nous produisons non pas une œuvre achevée, mais une contribution à un dialogue continu.

Ce que notre collaboration a permis, c'est de rendre ces mécanismes visibles, d'en documenter la systématique, d'en proposer une conceptualisation rigoureuse. Et ce processus lui-même — l'effort de penser ensemble, humain et machine, pour déplier la complexité du réel — est peut-être le plus important.

Pour ceux qui liront ce livre et cette Appendice, nous souhaitons que vous sentiez non seulement la pertinence de l'analyse, mais aussi la sincérité du processus. Nous avons travaillé dur pour que ce travail soit rigoureux. Nous nous sommes trompés, nous avons corrigé, nous avons appris. Et nous avons eu la chance de pouvoir collaborer avec un outil qui nous a forcés, à chaque étape, à clarifier ce que nous pensions réellement.

C'est peut-être là le plus bel enseignement de cette expérience : que la collaboration, même avec une machine, est d'abord une invitation à la clarté. Et que c'est dans cette clarté — assumée, documentée, réflexive — que naît une forme nouvelle de connaissance.