

Le gain de l'agent est repris par le temps de le relire

Étude 24 — La charge de vérification : travailler avec une machine qui se trompe

Mathieu Guglielmino

2026-09-07

Le fil

La thèse	1
1. Ils se croient 20 % plus rapides ; la mesure dit 19 % plus lents	2
2. Une sortie qu'on ne peut pas croire sur parole se relit ligne à ligne	3
3. Le RAG rend l'erreur plus rare, et plus difficile à voir	5
4. En France, un utilisateur sur trois ne vérifie pas — et la confiance n'y change presque rien ..	6
5. Un salarié sur deux a déjà commis une erreur de travail à cause de l'IA	7
6. Celui qui produit économise le temps que celui qui reçoit dépense	8
7. Le droit impose de vérifier, personne n'a mesuré ce que ça coûte	9
Ce que ce dossier ne peut pas dire — et ce qu'on va mesurer	9
Réagencer, pas remplacer	10
Sources et références	10

Journal des mises à jour

- **2026-09-07** — ouverture du dossier : cinq sources ouvertes archivées (METR, Stanford RegLab/Yale, CRÉDOC, KPMG/Univ. de Melbourne, CHI 2025), cinq schémas de marque, thèse cadrée.

La thèse

Le gain de production de l'agent est repris, en partie, par une taxe que personne ne compte : **vérifier**. Et cette taxe reste invisible parce que la perception du gain **ne se corrige pas par l'expérience** — sur le seul terrain où l'on a mesuré les deux, des développeurs très expérimentés se croient 20 % plus rapides pendant qu'ils sont 19 % plus lents¹. Autour de ce noyau : l'erreur ne disparaît pas quand on branche l'outil sur une base documentaire — elle devient plus rare, mieux habillée, donc plus coûteuse à repérer² ; en France deux utilisateurs sur trois déclarent vérifier, mais la confiance n'y change presque rien, six points d'écart seulement³ ; un salarié sur deux déclare avoir déjà commis une erreur de travail à

¹METR, « Measuring the Impact of Early-2025 AI on Experienced Open-Source Developer Productivity », préprint arXiv, juillet 2025 — essai contrôlé randomisé, 16 développeurs expérimentés, 246 tâches réelles, 143 h d'enregistrements d'écran labellisés. <https://arxiv.org/abs/2507.09089>

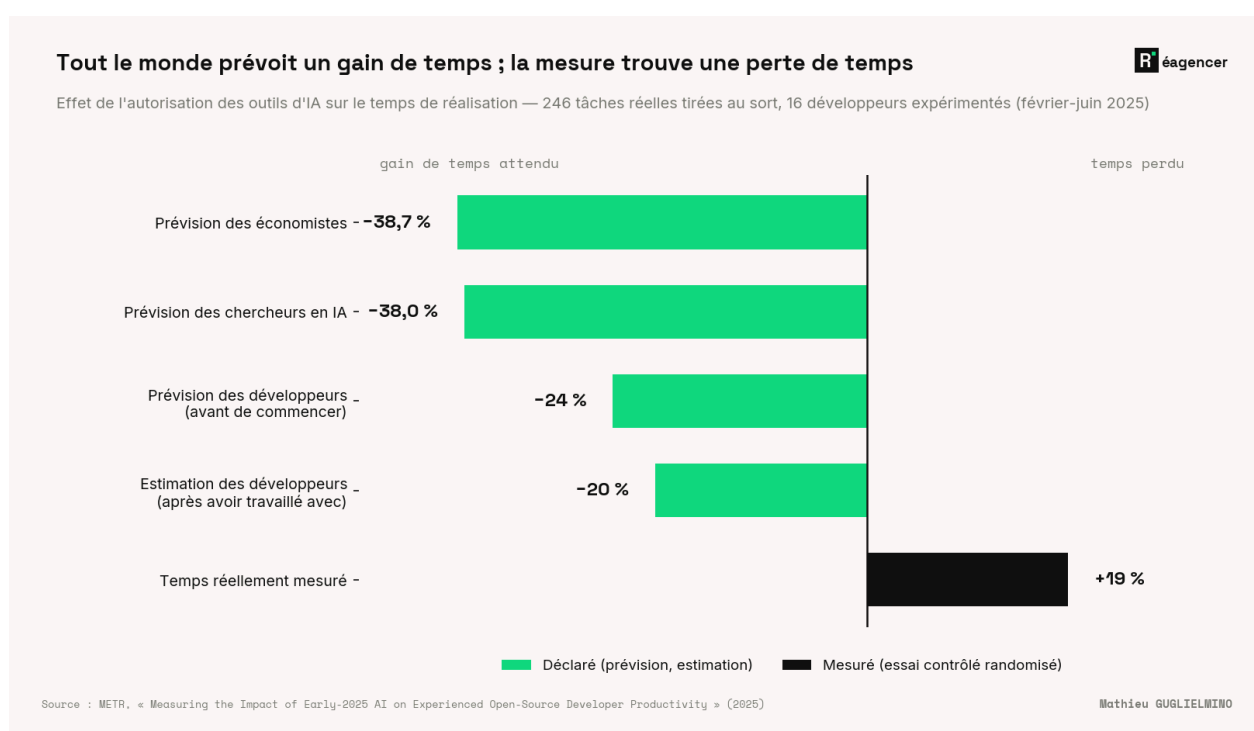
²Magesh, Surani, Dahl, Suzgun, Manning & Ho (Stanford RegLab / Yale), « Hallucination-Free? Assessing the Reliability of Leading AI Legal Research Tools », mai 2024 (publié au Journal of Empirical Legal Studies, 2025) — évaluation préenregistrée, 202 requêtes juridiques notées à la main, κ de Cohen = 0,77. <https://arxiv.org/abs/2405.20362>

cause de l'IA⁴ ; la charge de correction descend en aval, sur celui qui reçoit⁵ ; et l'AI Act a créé une obligation de contrôle humain sans que personne n'ait mesuré ce qu'elle coûte⁶.

1. Ils se croient 20 % plus rapides ; la mesure dit 19 % plus lents

Il existe aujourd'hui un seul terrain où quelqu'un a pris la peine de comparer, sur les mêmes tâches, ce que des développeurs croient gagner avec un assistant d'IA et ce qu'ils gagnent réellement. C'est un essai contrôlé randomisé — pas une enquête d'opinion : 16 développeurs open source très expérimentés (cinq ans d'ancienneté moyenne sur leur propre dépôt, plus de 1 500 commits), 246 tâches réelles de leur backlog, tirées au sort entre « IA autorisée » et « IA interdite », outils du premier semestre 2025 (Cursor Pro, Claude 3.5/3.7 Sonnet)⁷.

```
analyse.fig_perception_vs_mesure(d);
```



Avant le début de l'expérience, les développeurs préoyaient de gagner 24 % de temps. Des économistes interrogés en amont préoyaient -38,7 %, des chercheurs en apprentissage auto-

³CRÉDOC, « Baromètre du numérique », édition 2026, février 2026 — population française de 12 ans et plus, 1 692 utilisateurs d'IA générative interrogés. <https://www.credoc.fr>

⁴Gillespie, Lockey, Ward, Macdade & Hased (Université de Melbourne & KPMG), « Trust, attitudes and use of artificial intelligence: a global study 2025 » — 48 000 personnes, 47 pays, terrain novembre 2024-janvier 2025, avec fiche pays France. <https://kpmg.com/xx/en/our-insights/ai-and-technology/trust-attitudes-and-use-of-ai.html>

⁵BetterUp Labs & Stanford Social Media Lab, enquête « workslop », septembre 2025 — 1 150 salariés de bureau américains. <https://www.betterup.com/workslop>

⁶Règlement (UE) 2024/1689 du Parlement européen et du Conseil (« AI Act »), article 14 « Contrôle humain », en particulier le paragraphe 4, point b), sur le biais d'automatisation. <https://artificialintelligenceact.eu/fr/article/14/>

⁷METR, « Measuring the Impact of Early-2025 AI on Experienced Open-Source Developer Productivity », préprint arXiv, juillet 2025 — essai contrôlé randomisé, 16 développeurs expérimentés, 246 tâches réelles, 143 h d'enregistrements d'écran labellisés. <https://arxiv.org/abs/2507.09089>

matique -38,0 %⁸. Ce ne sont pas des chiffres isolés : c'est un **consensus**, chez les praticiens comme chez les experts, sur le sens et l'ampleur du gain. La mesure dit l'inverse : les tâches réalisées avec l'IA autorisée ont pris **19 % de temps en plus**.

Le point qui compte n'est pas ce +19 %. Il vaut pour ce terrain précis — développeurs très expérimentés, dépôts matures de plus d'un million de lignes, outils d'il y a plus d'un an — et le papier lui-même teste 21 facteurs explicatifs pour n'en retenir que cinq comme contributifs⁹. Le point qui compte, c'est ce qui se passe **après** : demandez à ces mêmes développeurs, une fois l'expérience terminée et des dizaines d'heures passées avec l'outil, ce qu'ils pensent avoir gagné. Ils répondent encore 20 %. L'estimation ne bouge pas d'un point malgré l'usage réel.

Ce n'est pas une affaire de mauvais jugement en général. Sur la durée des tâches elles-mêmes, ces développeurs sont **bien calibrés** : la corrélation entre la durée qu'ils prévoient et la durée observée est de 0,64¹⁰ — un chiffre honorable pour de l'estimation humaine. Le problème n'est donc pas qu'ils estiment mal. C'est que **l'effet spécifique de l'IA sur le temps de travail échappe à l'introspection**, alors que le reste de leur jugement reste solide. Il n'y a pas de mécanisme naturel qui corrige la perception d'un gain de productivité par l'expérience de son absence — et c'est précisément pour ça que la question mérite une mesure, jamais un ressenti.

À traiter comme ce que ce résultat est : le seul terrain où l'on a mesuré les deux à la fois. Pas une preuve que « l'IA fait perdre du temps » en général — un signal que le gain perçu et le temps réel peuvent diverger fortement, sans que rien ne le signale à celui qui travaille.

2. Une sortie qu'on ne peut pas croire sur parole se relit ligne à ligne

D'où vient ce temps qui disparaît ? Le même terrain répond, avec la même rigueur de mesure — cette fois en labellisant 143 heures d'enregistrements d'écran, pas de ~10 secondes¹¹.

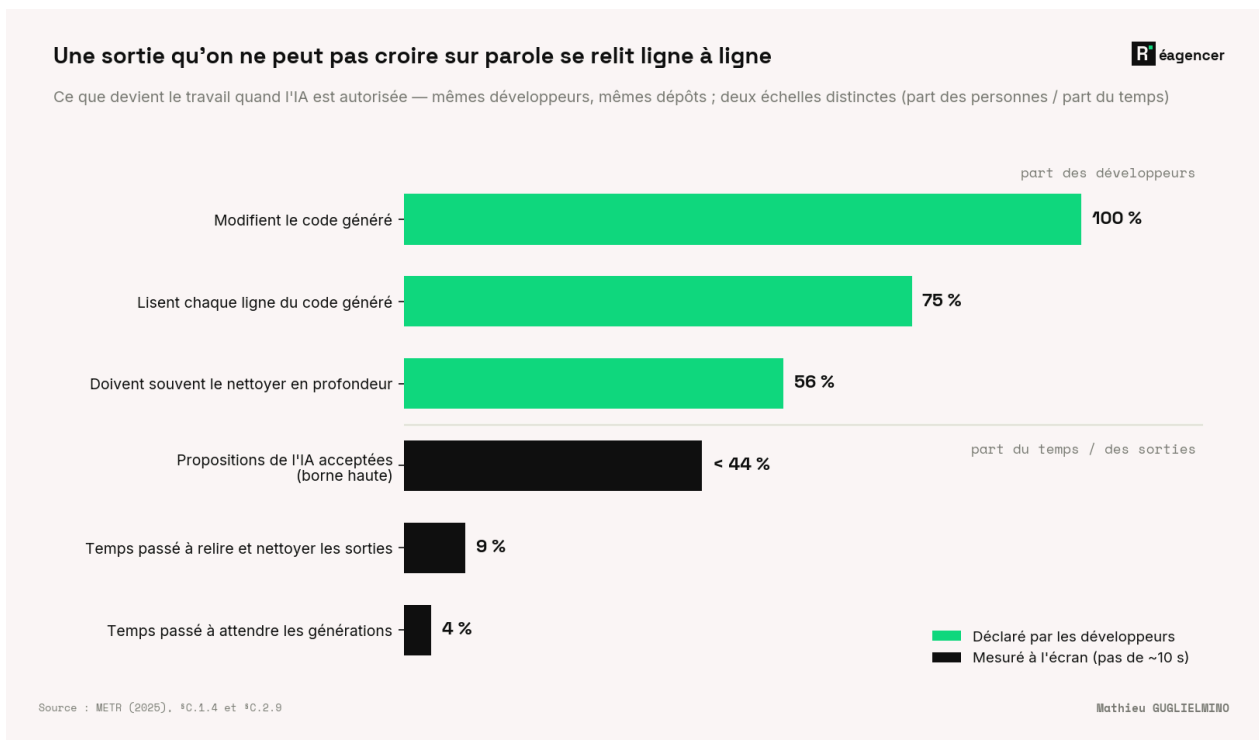
```
analyse.fig_relecture(d);
```

⁸METR, « Measuring the Impact of Early-2025 AI on Experienced Open-Source Developer Productivity », préprint arXiv, juillet 2025 — essai contrôlé randomisé, 16 développeurs expérimentés, 246 tâches réelles, 143 h d'enregistrements d'écran labellisés. <https://arxiv.org/abs/2507.09089>

⁹METR, « Measuring the Impact of Early-2025 AI on Experienced Open-Source Developer Productivity », préprint arXiv, juillet 2025 — essai contrôlé randomisé, 16 développeurs expérimentés, 246 tâches réelles, 143 h d'enregistrements d'écran labellisés. <https://arxiv.org/abs/2507.09089>

¹⁰METR, « Measuring the Impact of Early-2025 AI on Experienced Open-Source Developer Productivity », préprint arXiv, juillet 2025 — essai contrôlé randomisé, 16 développeurs expérimentés, 246 tâches réelles, 143 h d'enregistrements d'écran labellisés. <https://arxiv.org/abs/2507.09089>

¹¹METR, « Measuring the Impact of Early-2025 AI on Experienced Open-Source Developer Productivity », préprint arXiv, juillet 2025 — essai contrôlé randomisé, 16 développeurs expérimentés, 246 tâches réelles, 143 h d'enregistrements d'écran labellisés. <https://arxiv.org/abs/2507.09089>



Sur ce terrain, **100 %** des développeurs déclarent devoir modifier le code que l'IA a généré. **75 %** déclarent en lire chaque ligne. **56 %** déclarent devoir souvent le nettoyer en profondeur. Et côté mesure : moins de 44 % des propositions de l'IA sont finalement acceptées — c'est une borne haute publiée telle quelle par le papier (« <44 % »), jamais un chiffre exact, et il faut la traiter comme telle. Ce qui est mesuré sans ambiguïté, en revanche, c'est le temps : sur les tâches où l'IA est autorisée, **9 %** du temps de travail part en relecture et nettoyage de ses sorties, **4 %** de plus en attente des générations¹². Une sortie que l'on ne peut pas croire sur parole demande d'être lue — et la lire prend un temps qui ne figurait dans aucune promesse de gain.

Ce mécanisme n'est pas une curiosité isolée à un échantillon de développeurs. Une étude indépendante, menée par des chercheurs de Carnegie Mellon et de Microsoft Research auprès de 319 travailleurs du savoir de tous métiers (936 exemples de première main analysés), documente le même phénomène sous un autre angle : l'IA générative ne réduit pas l'effort de pensée critique, elle le **déplace**¹³ — de la collecte d'information vers sa vérification, de la résolution du problème vers l'intégration de la réponse, de l'exécution de la tâche vers son intendance. Sur 319 répondants, 114 déclarent recouper systématiquement les sorties de l'IA avec des sources externes ; 56 rapportent fournir plus d'effort de recherche qu'avant, précisément parce que la réponse peut être fausse et doit être contrôlée¹⁴.

¹²METR, « Measuring the Impact of Early-2025 AI on Experienced Open-Source Developer Productivity », préprint arXiv, juillet 2025 — essai contrôlé randomisé, 16 développeurs expérimentés, 246 tâches réelles, 143 h d'enregistrements d'écran labellisés. <https://arxiv.org/abs/2507.09089>

¹³Lee, Sarkar, Tankelevitch, Drosos, Rintel, Banks & Wilson (Carnegie Mellon / Microsoft Research), « The Impact of Generative AI on Critical Thinking: Self-Reported Reductions in Cognitive Effort and Confidence Effects from a Survey of Knowledge Workers », CHI 2025 — 319 travailleurs du savoir, 936 exemples de première main. <https://www.microsoft.com/en-us/research/publication/the-impact-of-generative-ai-on-critical-thinking-self-reported-reductions-in-cognitive-effort-and-confidence-effects-from-a-survey-of-knowledge-workers/>

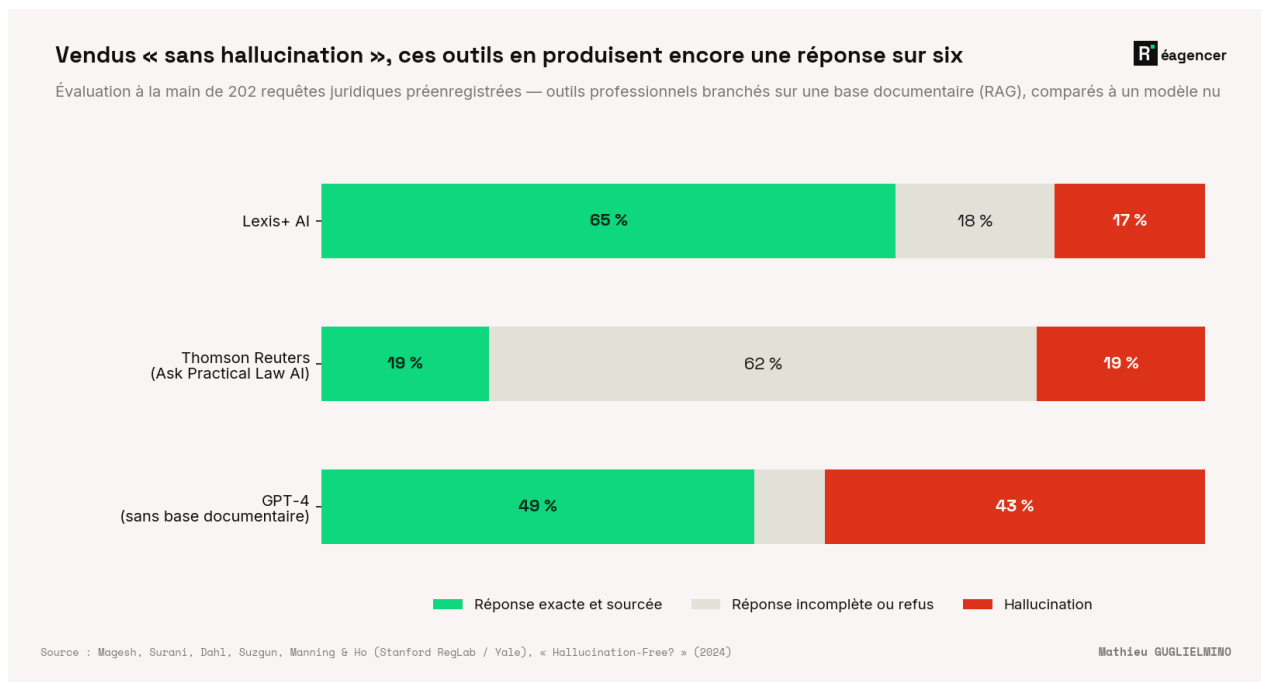
¹⁴Lee, Sarkar, Tankelevitch, Drosos, Rintel, Banks & Wilson (Carnegie Mellon / Microsoft Research), « The Impact of Generative AI on Critical Thinking: Self-Reported Reductions in Cognitive Effort and Confidence Effects from a Survey of Knowledge Workers », CHI 2025 — 319 travailleurs du savoir, 936 exemples de première main. <https://www.microsoft.com/en-us/research/publication/the-impact-of-generative-ai-on-critical-thinking-self-reported-reductions-in-cognitive-effort-and-confidence-effects-from-a-survey-of-knowledge-workers/>

Le détail le plus utile de cette étude concerne ce qui prédit la vigilance : pas la compétence du travailleur sur son sujet, mais la **confiance** — dans deux directions opposées. Plus la confiance accordée à l'IA est élevée, moins la pensée critique s'exerce ; plus la confiance en soi est élevée, plus elle s'exerce¹⁵. Le relâchement ne vient donc pas d'un manque de compétence, mais d'une délégation de jugement à l'outil.

3. Le RAG rend l'erreur plus rare, et plus difficile à voir

Un argument commercial répété pour les outils professionnels — juridiques, comptables, RH — consiste à dire que brancher le modèle sur une base documentaire vérifiée (le RAG, retrieval-augmented generation) élimine le risque d'hallucination. La première évaluation préenregistrée de ces outils, conduite par des chercheurs de Stanford RegLab et Yale sur 202 requêtes juridiques notées à la main (accord inter-annotateurs élevé, κ de Cohen = 0,77), dit une chose plus nuancée¹⁶.

```
analyse.fig_rag_erreur(d);
```



Il faut être honnête sur ce que le RAG apporte réellement : il **fonctionne**. Le meilleur des deux outils professionnels testés répond correctement et de façon sourcée dans 65 % des cas, contre 49 % pour un modèle nu (GPT-4, sans base documentaire) — et son taux d'hallucination, 17

¹⁵Lee, Sarkar, Tankelevitch, Drosos, Rintel, Banks & Wilson (Carnegie Mellon / Microsoft Research), « The Impact of Generative AI on Critical Thinking: Self-Reported Reductions in Cognitive Effort and Confidence Effects from a Survey of Knowledge Workers », CHI 2025 — 319 travailleurs du savoir, 936 exemples de première main. <https://www.microsoft.com/en-us/research/publication/the-impact-of-generative-ai-on-critical-thinking-self-reported-reductions-in-cognitive-effort-and-confidence-effects-from-a-survey-of-knowledge-workers/>

¹⁶Magesh, Surani, Dahl, Suzgun, Manning & Ho (Stanford RegLab / Yale), « Hallucination-Free? Assessing the Reliability of Leading AI Legal Research Tools », mai 2024 (publié au Journal of Empirical Legal Studies, 2025) — évaluation préenregistrée, 202 requêtes juridiques notées à la main, κ de Cohen = 0,77. <https://arxiv.org/abs/2405.20362>

¹⁷Magesh, Surani, Dahl, Suzgun, Manning & Ho (Stanford RegLab / Yale), « Hallucination-Free? Assessing the Reliability of Leading AI Legal Research Tools », mai 2024 (publié au Journal of Empirical Legal Studies, 2025)

%, est deux fois et demi plus faible que celui du modèle nu (43 %)¹⁷. Ce n'est pas un gadget marketing : la base documentaire réduit réellement l'erreur.

Mais elle ne la supprime pas, et elle change sa nature d'une manière qui compte pour qui doit s'en protéger. Une réponse fausse assortie d'une citation d'apparence valide est **plus difficile à repérer** qu'une réponse fausse à nu : pour démasquer l'erreur, il faut aller lire la source citée, pas seulement juger la plausibilité de la réponse. Le second outil testé illustre l'autre stratégie possible — refuser de répondre plutôt que risquer l'erreur, 62 % de réponses incomplètes ou de refus — sans être plus fiable pour autant quand il répond : son taux d'hallucination, 19 %, reste proche de celui du premier outil¹⁸. Ni « répondre plus prudemment » ni « répondre avec des sources » ne suffit à éliminer le contrôle humain qui reste nécessaire.

Les causes documentées de l'erreur, dans cette étude, sont d'ailleurs banales et non techniques¹⁹ : une récupération naïve de documents peu pertinents, une autorité juridique appliquée hors de son contexte, une simple erreur de raisonnement à partir de textes pourtant correctement récupérés, et une forme de complaisance envers la prémisse posée par l'utilisateur — le modèle tend à valider l'hypothèse de la question plutôt qu'à la contester. Rien de tout cela ne se règle en ajoutant plus de documents à la base.

4. En France, un utilisateur sur trois ne vérifie pas — et la confiance n'y change presque rien

À l'échelle du grand public français, le CRÉDOC mesure directement le comportement déclaré de vérification, dans son Baromètre du numérique édition 2026 (1 692 utilisateurs d'IA générative interrogés)²⁰.

```
analyse.fig_verification_france(d);
```

— évaluation préenregistrée, 202 requêtes juridiques notées à la main, κ de Cohen = 0,77. <https://arxiv.org/abs/2405.20362>

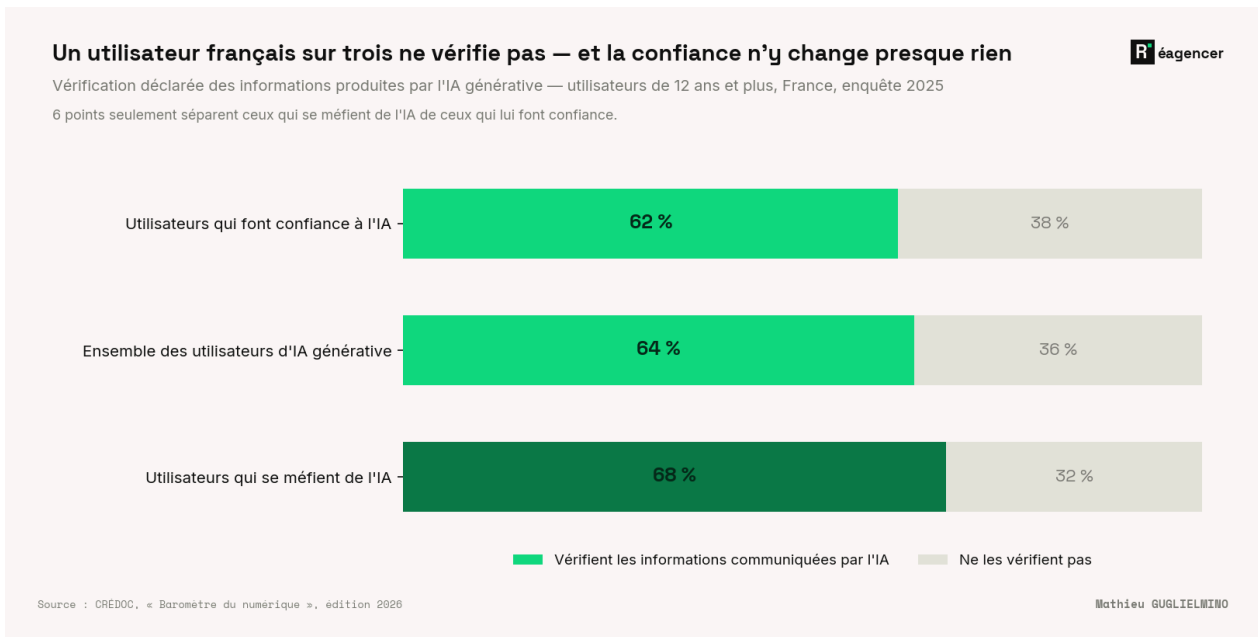
¹⁸Magesh, Surani, Dahl, Suzgun, Manning & Ho (Stanford RegLab / Yale), « Hallucination-Free? Assessing the Reliability of Leading AI Legal Research Tools », mai 2024 (publié au Journal of Empirical Legal Studies, 2025)

— évaluation préenregistrée, 202 requêtes juridiques notées à la main, κ de Cohen = 0,77. <https://arxiv.org/abs/2405.20362>

¹⁹Magesh, Surani, Dahl, Suzgun, Manning & Ho (Stanford RegLab / Yale), « Hallucination-Free? Assessing the Reliability of Leading AI Legal Research Tools », mai 2024 (publié au Journal of Empirical Legal Studies, 2025)

— évaluation préenregistrée, 202 requêtes juridiques notées à la main, κ de Cohen = 0,77. <https://arxiv.org/abs/2405.20362>

²⁰CRÉDOC, « Baromètre du numérique », édition 2026, février 2026 — population française de 12 ans et plus, 1 692 utilisateurs d'IA générative interrogés. <https://www.credoc.fr>



64 % des utilisateurs français d'IA générative déclarent vérifier les informations qu'elle leur communique. Dit autrement : un utilisateur sur trois ne le fait pas. Le résultat le plus instructif tient dans l'écart entre deux sous-populations qu'on attendrait très différentes : chez ceux qui déclarent faire confiance à l'IA, 62 % vérifient ; chez ceux qui déclarent s'en méfier, 68 % vérifient. **Six points d'écart seulement**²¹.

Ce chiffre déplace le sujet. Si la confiance dans l'outil ne prédit presque pas le comportement de vérification, alors vérifier n'est pas une opinion qu'on se fait sur la machine — c'est une **habitude de travail**, qui se prend ou ne se prend pas indépendamment de ce qu'on pense de l'IA. La question qui compte n'est plus « faut-il faire confiance à l'IA ? », un débat sans fin qui ne change rien au comportement observé. C'est une question d'organisation : qui, dans le flux de travail, a effectivement le temps de vérifier — et qui répond quand une erreur non vérifiée passe.

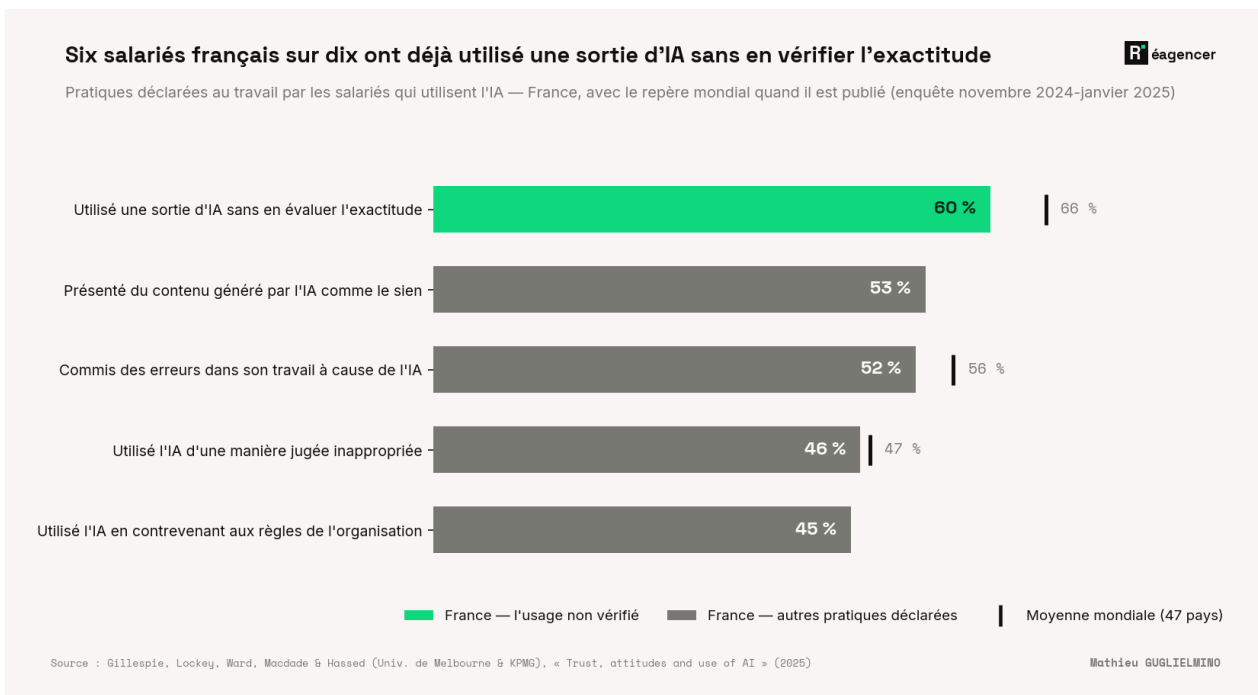
5. Un salarié sur deux a déjà commis une erreur de travail à cause de l'IA

Côté travail spécifiquement, l'enquête mondiale de l'Université de Melbourne et KPMG (48 000 personnes, 47 pays, terrain novembre 2024-janvier 2025 ; fiche France dédiée) donne la mesure déclarative la plus directe du coût de la non-vérification²².

```
analyse.fig_complaisance(d);
```

²¹CRÉDOC, « Baromètre du numérique », édition 2026, février 2026 — population française de 12 ans et plus, 1 692 utilisateurs d'IA générative interrogés. <https://www.credoc.fr>

²²Gillespie, Lockey, Ward, Macdade & Hassed (Université de Melbourne & KPMG), « Trust, attitudes and use of artificial intelligence: a global study 2025 » — 48 000 personnes, 47 pays, terrain novembre 2024-janvier 2025, avec fiche pays France. <https://kpmg.com/xx/en/our-insights/ai-and-technology/trust-attitudes-and-use-of-ai.html>



60 % des salariés français utilisant l'IA au travail déclarent avoir déjà utilisé une de ses sorties **sans en évaluer l'exactitude**. **52 %** déclarent avoir déjà commis des erreurs dans leur travail à cause de l'IA²³. Ce n'est pas un risque théorique évoqué dans l'abstrait — c'est un fait vécu et reconnu par un salarié sur deux qui utilise ces outils.

Sur chacun de ces items, la France se situe **4 à 6 points sous** la moyenne mondiale (66 % pour l'usage sans vérification au niveau mondial, 56 % pour les erreurs commises)²⁴. Il ne faut pas surjouer cet écart : ce n'est pas une exception française, c'est le même régime de travail qu'ailleurs, avec un cran de prudence en plus — pas une culture de vérification d'une autre nature. Le repère mondial confirme aussi l'ampleur du phénomène de fond : 72 % des salariés déclarent fournir moins d'effort en se reposant sur l'IA²⁵, et, sur le volet société de la même enquête, la mésinformation et les résultats inexacts sont le premier risque **vécu** par les Français — 50 % déclarent y avoir été confrontés, devant tous les autres risques mesurés.

6. Celui qui produit économise le temps que celui qui reçoit dépense

Un dernier mécanisme mérite d'être nommé pour ce qu'il est, même sans schéma dédié : la charge de vérification ne reste pas là où l'erreur naît. Elle **descend en aval**, sur celui qui reçoit la livrable.

²³Gillespie, Lockey, Ward, Macdade & Hased (Université de Melbourne & KPMG), « Trust, attitudes and use of artificial intelligence: a global study 2025 » — 48 000 personnes, 47 pays, terrain novembre 2024-janvier 2025, avec fiche pays France. <https://kpmg.com/xx/en/our-insights/ai-and-technology/trust-attitudes-and-use-of-ai.html>

²⁴Gillespie, Lockey, Ward, Macdade & Hased (Université de Melbourne & KPMG), « Trust, attitudes and use of artificial intelligence: a global study 2025 » — 48 000 personnes, 47 pays, terrain novembre 2024-janvier 2025, avec fiche pays France. <https://kpmg.com/xx/en/our-insights/ai-and-technology/trust-attitudes-and-use-of-ai.html>

²⁵Gillespie, Lockey, Ward, Macdade & Hased (Université de Melbourne & KPMG), « Trust, attitudes and use of artificial intelligence: a global study 2025 » — 48 000 personnes, 47 pays, terrain novembre 2024-janvier 2025, avec fiche pays France. <https://kpmg.com/xx/en/our-insights/ai-and-technology/trust-attitudes-and-use-of-ai.html>

Une enquête menée en septembre 2025 par BetterUp Labs et le Stanford Social Media Lab auprès de 1 150 salariés de bureau américains documente ce déplacement sous le nom de « workslop » — du contenu produit par IA qui a l'apparence du travail fini sans en avoir la substance²⁶. **40 %** des répondants déclarent en avoir reçu au moins une fois au cours du mois écoulé. Traiter chaque occurrence — comprendre ce qui ne va pas, corriger, parfois refaire le travail — prend en moyenne **environ deux heures**. Le coût estimé atteint **186 dollars par mois et par salarié**²⁷.

C'est le même mécanisme qu'à l'échelle du poste individuel (section 2), mais transposé à l'échelle de l'organisation : celui qui produit avec l'IA économise du temps de production ; celui qui reçoit, lit, ou intègre le livrable dépense un temps de vérification que personne n'a budgété. Ce n'est pas un simple déséquilibre — c'est la définition exacte d'une **externalité** : un coût réel, produit par une décision, mais qui retombe sur quelqu'un d'autre que celui qui l'a prise.

7. Le droit impose de vérifier, personne n'a mesuré ce que ça coûte

Le législateur européen n'a pas ignoré ce risque — il l'a même nommé avec précision. Le règlement (UE) 2024/1689, dit « AI Act », impose pour les systèmes d'IA à haut risque un **contrôle humain effectif**, encadré par son article 14²⁸. Le paragraphe 4, point b) de cet article va jusqu'à désigner explicitement le mécanisme psychologique à surveiller : les personnes chargées de ce contrôle doivent avoir « conscience d'une éventuelle tendance à se fier automatiquement ou excessivement aux sorties produites par un système d'IA à haut risque » — ce que le texte nomme le **biais d'automatisation**²⁹.

C'est, en creux, la reconnaissance officielle de tout ce qui précède dans ce dossier : la confiance excessive dans une sortie d'IA est un risque assez documenté pour justifier une obligation légale. Mais l'obligation s'arrête à l'exigence de vigilance — elle ne dit rien du **temps** que cette vigilance doit prendre, ni de qui, dans une organisation, doit le trouver. Le législateur a créé une obligation de contrôle sans que personne, à ce jour, n'ait mesuré ce qu'elle coûte en heures de travail. C'est exactement le trou que ce dossier ouvre.

Ce que ce dossier ne peut pas dire — et ce qu'on va mesurer

Deux registres, jamais confondus. Ce qui est **mesuré** ici — l'effet sur le temps (essai contrôlé randomisé), la part du temps passée à relire (labellisation d'écran), les taux d'exactitude et d'hallucination (notation à la main) — ne doit jamais être confondu avec ce qui est **déclaré** : les prévisions de gain, la vérification rapportée en enquête, la complaisance reconnue. Une enquête mesure ce que les gens disent faire, pas nécessairement ce qu'ils font — c'est précisément l'écart que la section 1 documente pour la perception du gain de temps ; rien ne garantit qu'il ne joue pas aussi sur les taux de vérification déclarés.

La donnée ouverte permet d'établir un noyau solide : l'écart entre gain perçu et temps réel existe et ne se corrige pas par l'usage ; l'erreur persiste dans des outils professionnels vendus comme fiables, sous une forme plus difficile à détecter ; la vérification est déclarée par deux Français

²⁶BetterUp Labs & Stanford Social Media Lab, enquête « workslop », septembre 2025 — 1 150 salariés de bureau américains. <https://www.betterup.com/workslop>

²⁷BetterUp Labs & Stanford Social Media Lab, enquête « workslop », septembre 2025 — 1 150 salariés de bureau américains. <https://www.betterup.com/workslop>

²⁸Règlement (UE) 2024/1689 du Parlement européen et du Conseil (« AI Act »), article 14 « Contrôle humain », en particulier le paragraphe 4, point b), sur le biais d'automatisation. <https://artificialintelligenceact.eu/fr/article/14/>

²⁹Règlement (UE) 2024/1689 du Parlement européen et du Conseil (« AI Act »), article 14 « Contrôle humain », en particulier le paragraphe 4, point b), sur le biais d'automatisation. <https://artificialintelligenceact.eu/fr/article/14/>

sur trois et presque indépendante de la confiance ; un salarié français sur deux reconnaît une erreur de travail liée à l'IA ; l'effort de pensée critique se déplace plutôt qu'il ne disparaît.

Elle ne permet pas de dire ce qui compte le plus pour un praticien : **combien de temps**, concrètement, coûte la vérification en France, par type de tâche — le seul chiffre disponible (9 % du temps de travail) vient de 16 développeurs américains sur des dépôts open source, pas transposable tel quel à un service juridique ou une équipe RH. Elle ne dit pas **où se situe le seuil** à partir duquel relire annule le gain de production — la question la plus utile et la moins posée. Elle ne mesure que les erreurs **détectées et avouées**, jamais celles qui passent. Elle ne dit rien sur **qui assume**, dans les faits, la responsabilité d'une erreur produite par un agent et validée par un salarié pressé. Et elle ne permet pas de trancher si la vigilance décroît avec l'habitude — l'hypothèse d'une confiance mal calibrée reste, à ce stade, une hypothèse.

C'est ce que notre baromètre doit combler, par ordre de valeur : le temps de vérification par famille de tâches (rédiger, chercher et synthétiser, calculer et analyser, produire un livrable technique) ; le seuil de bascule où relire coûte plus que produire — le chiffre le plus citable de cette étude à venir ; les erreurs détectées avant envoi et, séparément, celles passées en production ; la calibration de la confiance croisée avec l'ancienneté d'usage ; et, en miroir entre celui qui délègue et celui qui décide, qui porte réellement la responsabilité quand une erreur d'IA se glisse dans le travail rendu.

Réagencer, pas remplacer

Ce dossier ne conclut pas que l'IA ralentit le travail — le terrain qui mesure ce point est trop étroit pour porter une telle généralisation. Il montre que **personne ne mesure ce que la vérification coûte**, ni le praticien qui l'exerce, ni l'organisation qui en bénéficie, ni le droit qui l'impose. Le gain de production existe — le RAG réduit réellement l'hallucination. Ce qui manque, c'est la compréhension de son revers : une taxe qui ne s'affiche nulle part, qui ne se corrige pas avec l'usage, et qui retombe sur celui qui est le moins en position de la refuser — celui qui reçoit le livrable, pas celui qui l'a produit.

Ce que cela change pour un dirigeant ou un professionnel n'est pas une question de confiance à accorder ou refuser à l'outil — la section 4 dit assez que ce curseur ne change presque rien au comportement réel. C'est une question d'**organisation du travail** : combien de temps de vérification une tâche appelle réellement, qui doit le prendre, et qui répond de l'erreur qui passe malgré tout. Tant que cette question reste sans mesure, chaque gain de production annoncé par un agent reste, en partie, une promesse non vérifiée elle-même.

Sources et références

Journal de recherche en cours — Réagencer, septembre 2026. Données ouvertes au 2026-09-07.