

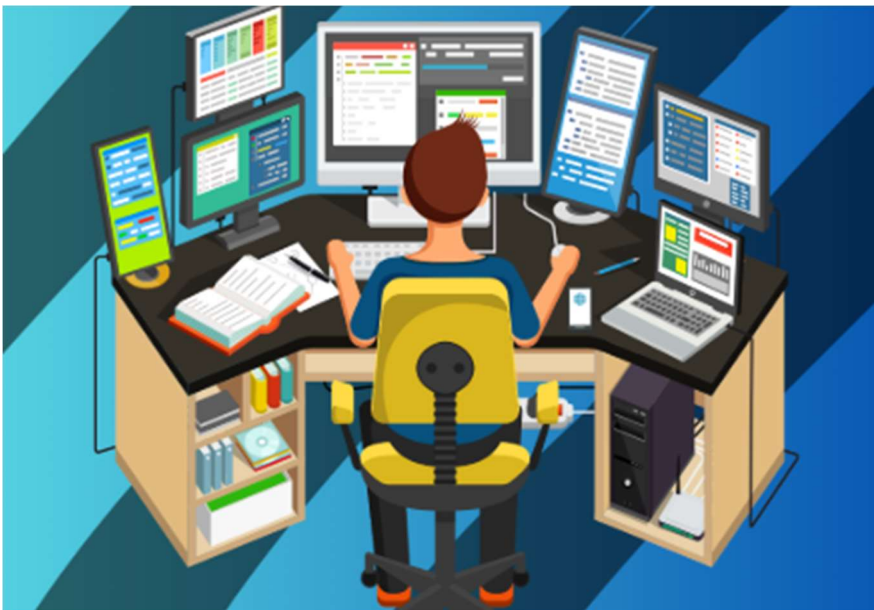
Les chercheurs avertissent que les industries de l'IA sont en train de submerger la science d'études inutiles, face à l'explosion des publications automatisées basées sur des données accessibles mais elles-mêmes pas nécessairement fiables scientifiquement.

Le 14 mai 2025 à 17:14, par [Bruno](#)



Le rapport de l'université du Surrey soulève une question cruciale : l'intégrité de la connaissance scientifique est compromise par la prolifération d'articles générés par l'IA, souvent superficiels et méthodologiquement fragiles. Ces « usines à papier », profitant de bases de données facilement accessibles, produisent en masse [des études biaisées](#), privilégiant des analyses simplistes à un seul facteur pour des problèmes de santé complexes. Cette pratique, amplifiée depuis 2021, inonde les revues, y compris celles évaluées par les pairs et menace de polluer le corpus scientifique avec des conclusions trompeuses. Si l'évaluation par les pairs reste un rempart, son efficacité est mise à mal par le volume croissant de ces publications et les limites inhérentes au système (évaluateurs non spécialisés, manque de rigueur).

Cette crise révèle aussi des enjeux systémiques : la marchandisation de la recherche, illustrée par des revues prédatrices, et l'instrumentalisation des données par certains acteurs (comme la Chine, devenue majoritaire dans ces publications). Les solutions proposées, vigilance accrue des éditeurs, encadrement des accès aux données, pointent vers une nécessaire réforme. Cependant, elles se heurtent à un paradoxe : l'IA, outil potentiel d'innovation, devient un vecteur de désinformation lorsque détournée par des logiques productivistes ou idéologiques. Ce phénomène s'inscrit dans un paysage plus large où le « slop » IA brouille les frontières entre réalité et fiction, exigeant une réponse collective pour préserver la crédibilité de la science.



Cette pratique connaît une croissance exponentielle depuis 2021, submergeant les revues scientifiques y compris celles soumises à évaluation par les pairs. Alors que seulement quatre articles de ce type étaient publiés annuellement entre 2014 et 2021, leur nombre est passé à 33 en 2022, 82 en 2023, et atteignait déjà 190 à la mi-octobre 2024. Cette inflation soudaine met en lumière les limites du système actuel d'évaluation scientifique, dont les mécanismes de contrôle apparaissent dépassés face à ce déluge.

L'étude révèle également un changement notable dans la géographie de ces publications. La part des chercheurs chinois parmi les auteurs principaux est passée de 8% avant 2021 à 92% entre 2021 et 2024. Cette concentration géographique, combinée à la prédominance des approches monofactorielles, accroît le risque de pollution du corpus scientifique par des conclusions erronées, particulièrement pour des sujets complexes comme la dépression ou les maladies cardiovasculaires.

Matt Spick, coauteur de l'étude, dénonce ces « fictions scientifiques » qui, sous couvert de données publiques, contournent les exigences méthodologiques fondamentales. « Ces articles ont l'apparence de la science mais ne résistent pas à un examen rigoureux », explique-t-il, pointant du doigt la combinaison dangereuse entre l'accès facilité aux bases de données et les capacités des grands modèles de langage. Cette situation engorge les revues scientifiques et dépasse les capacités des évaluateurs, menaçant à terme la crédibilité de l'ensemble de la recherche.

Monétisation, dragage de données et IA : vers une crise de la rigueur scientifique

Le phénomène s'inscrit dans un contexte plus large de marchandisation de la recherche scientifique, où certaines revues prédatrices monnayent la publication sans garantir la qualité des travaux. Les chercheurs identifient deux pratiques particulièrement préoccupantes : l'utilisation systématique d'analyses monofactorielles inadaptées à des problèmes complexes, et le « dragage » de données consistant à sélectionner arbitrairement des sous-ensembles pour confirmer des hypothèses préétablies. Le « dragage de données » est une pratique statistique qui consiste à explorer et analyser un ensemble de données de manière répétée, sans hypothèse préétablie, afin de trouver des corrélations ou des modèles qui pourraient sembler significatifs, mais qui en réalité seraient simplement le résultat du hasard.

Les entreprises spécialisées dans la falsification scientifique se multiplient, produisant des articles sur commande contre rémunération. Grâce aux avancées de l'intelligence artificielle (IA), ces contenus sont de plus en plus difficiles à détecter. Des outils sophistiqués permettent en effet de générer automatiquement du texte et des images convaincants, imitant le style et les données de véritables publications.

Face à ce risque, plusieurs grands éditeurs scientifiques ont pris des mesures. [Certains ont interdit ou limité l'usage de ChatGPT dans les articles soumis](#), craignant l'insertion de contenus erronés ou plagiés dans la littérature académique. Tandis que quelques chercheurs ont tenté de désigner le chatbot comme coauteur, des revues comme Science ont choisi d'interdire toute utilisation directe de son texte dans les manuscrits. Springer Nature, de son côté, accepte l'assistance de l'IA pour la rédaction, mais impose une transparence totale sur son usage et rejette l'idée de lui attribuer la qualité d'auteur.

Ces décisions interviennent alors que le débat sur la place de l'IA dans la production intellectuelle s'intensifie, notamment après les controverses liées à son emploi dans les médias comme CNET. De nombreux experts estiment que ChatGPT pourrait bouleverser durablement le secteur éditorial, en particulier dans les domaines facilement automatisables comme le journalisme sportif ou financier. Ce chatbot, développé par OpenAI, est capable de rédiger des textes complexes à partir de sources accessibles en ligne, posant de nouveaux défis en matière d'authenticité et de fiabilité.

« Les organisations informelles, voire illégales, capables de produire de faux articles avec des données de plus en plus crédibles vont proliférer grâce à l'IA », alerte Jennifer Byrne, biologiste moléculaire et spécialiste de l'intégrité scientifique à l'Université de Sydney.

Face à cette situation, l'équipe du Surrey propose plusieurs mesures correctives. Elles incluent un renforcement de la vigilance des éditeurs, un meilleur encadrement de l'accès aux données, et l'obligation de justifier toute analyse partielle des jeux de données. Ces propositions visent à rétablir des garde-fous sans pour autant entraver l'innovation ou restreindre indûment l'accès aux données.

Ce phénomène illustre le paradoxe de l'IA dans la recherche : outil potentiel d'avancées majeures, elle devient aussi un vecteur de désinformation lorsqu'elle est détournée par des logiques productivistes. La situation appelle une réponse collective de la communauté scientifique pour préserver les fondements mêmes de la connaissance, alors que les frontières entre recherche authentique et "science-fiction" deviennent de plus en plus floues.

L'impact des API et des outils d'analyse standardisés sur la recherche

La quantité de données biologiques à la disposition des chercheurs a considérablement augmenté ces dernières

années, ce qui a multiplié les possibilités de recherche axée sur les données. À mesure que davantage d'informations deviennent disponibles dans des formats prêts pour l'intelligence artificielle, la recherche, lorsqu'elle est effectuée conformément aux meilleures pratiques - devrait devenir plus rapide et plus reproductible. La grande disponibilité de ces ensembles de données peut toutefois poser de nouveaux problèmes, en facilitant la production de manuscrits de bout en bout, à grande échelle, avec l'aide de l'IA. Il s'agit d'une pratique qui peut être adoptée par les usines à papier, définies par le groupe de travail United2Act Research comme des organisations clandestines qui fournissent des manuscrits de mauvaise qualité ou fabriqués à des clients payants.

L'ancienneté de certaines banques de données a conduit à la création de bibliothèques R et Python qui fournissent, entre autres, des outils automatisés de recherche, d'extraction et d'analyse, offrant des flux de travail standardisés et améliorant la reproductibilité. Ces outils, ainsi que d'autres environnements de codage et bibliothèques largement utilisés, peuvent contribuer de manière significative à la production rapide de résultats et aux publications qui s'ensuivent. La capacité des chercheurs à automatiser le processus d'extraction des données par le biais d'une interface de programmation d'application (API ; conformément aux lignes directrices FAIR selon lesquelles les données doivent pouvoir être récupérées par identifiant à l'aide d'un protocole de communication normalisé), permet le transfert des données directement vers des environnements d'apprentissage automatique, ce qui facilite l'exploration rapide et complète des données.

La possibilité d'extraire des données via une API directement dans des environnements d'apprentissage automatique tels que R ou Python peut transformer la productivité, le nombre d'hypothèses pouvant être testées n'étant limité que par l'accès informatique, mais cela peut aussi comporter des risques. L'accent mis sur les analyses à facteur unique peut être particulièrement problématique, étant donné la nature multifactorielle de nombreuses maladies, ainsi que la difficulté de différencier les prédicteurs qui sont spécifiques à un état de santé de ceux qui sont communs à différents types de maladies.

En outre, la possibilité de générer un grand nombre de modèles d'apprentissage automatique permet d'étudier rapidement et a posteriori d'autres hypothèses, au cas où la principale hypothèse a priori ne serait pas confirmée (une forme d'émission d'hypothèses après que les résultats sont connus, ou HARKing). Grâce à l'accès facile à l'informatique, il est possible d'effectuer une recherche étendue pour toute combinaison d'indicateur, d'état de santé, de cohorte et de fenêtre temporelle qui produit une valeur p faible. Si le dragage des données est un phénomène bien décrit, les pipelines d'accès direct à l'IA peuvent rendre les pipelines de recherche basés sur des formules plus productifs qu'ils ne l'étaient auparavant. Ce gain de productivité devrait être particulièrement intéressant pour les papeteries.

Vers une gouvernance éthique de l'accès aux données scientifiques à l'ère de l'IA

L'étude de l'université du Surrey souligne que la multiplication des recherches basées sur des analyses à facteur unique accroît significativement le risque d'introduire des conclusions erronées dans la littérature scientifique. Cette approche réductionniste, appliquée à des phénomènes complexes, fausse la compréhension de problèmes de santé multifactoriels comme la dépression ou les maladies cardiovasculaires, pourtant reconnus comme résultant de multiples causes interdépendantes.

Face à cette dérive, les chercheurs proposent plusieurs mesures correctives. Ils recommandent notamment que les comités de rédaction considèrent systématiquement les études monofactorielles sur des sujets complexes comme des signaux d'alerte nécessitant un examen particulièrement rigoureux. Cette vigilance accrue permettrait d'identifier plus facilement les travaux problématiques avant publication.

Le rapport préconise également un meilleur encadrement de l'accès aux bases de données scientifiques.

L'instauration de clés API individuelles et de numéros d'application, à l'image du système déjà mis en place par la UK Biobank, pourrait limiter les pratiques de dragage de données. Chaque publication devrait ainsi inclure un identifiant vérifiable attestant d'un usage légitime des données.

Une autre proposition consiste à imposer l'analyse de l'ensemble des données disponibles, sauf justification

méthodologique solide pour se limiter à un sous-ensemble. Cette mesure viserait à prévenir les biais de sélection qui faussent souvent les résultats des études générées massivement. « Notre objectif n'est pas de restreindre l'accès aux données ou d'interdire l'usage de l'IA, mais d'instaurer des garde-fous essentiels », précise Tulsi Suchak, auteure principale de l'étude.

La situation n'est pas sans précédent. L'an dernier, l'éditeur Wiley avait déjà dû retirer 19 revues scientifiques de sa filiale Hindawi, compromises dans la publication massive d'articles produits par des usines à papier utilisant l'IA. Ce cas illustre l'ampleur du phénomène et la nécessité d'une réponse coordonnée de la communauté scientifique...