

Lien : ia openai affirme avoir resolu le probleme de navier-stoke ... ou pas

Posté par ouau le 14 septembre 2026 à 09:38.
Étiquettes : paypal, openai, mathématiques, recherche_scientifique



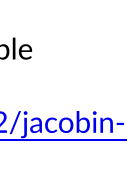
<https://www.pourlascience.fr/sd/mathematiques/si-openai-affirme-avoir-resolu-le-probleme-de-navier-stoke-29554.php>

(29 commentaires).

similaire

Posté par BAud (site web personnel) le 14 septembre 2026 à 13:17.
Évalué à 4 (+2/-0).

moui, on ne peut de toute façon pas lire l'article entier, puis bon pour en citer une partie



Par ailleurs, cette preuve annoncée le 8 septembre, naït dans la controverse. En effet, l'annonce d'OpenAI est arrivée à peine douze heures après qu'un mathématicien, Tristan Buckmaster, de l'université de New York, a rendu public ses propres résultats, obtenus avec Levent Alpöge, chercheur chez Anthropic (développeur de Claude, concurrent d'OpenAI).

ce qui peut faire référence à un article lisible

<https://www.fastcompany.com/91577272/jacobin-conjecture-anthropic-ai-levent-alpöge>

Un intérêt étant dans l'utilisation de [Lean](#) — pas la méthode de gestion de production, l'assistant de preuve — pour arriver à une démonstration.

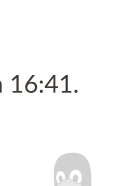
Bref, ça foisonne un peu aussi côté Mathématiques, est-ce bien efficace ou que des effets d'annonce ?

Répondre

[^] # Re: similaire

Posté par ouau le 14 septembre 2026 à 13:33. Évalué à 5 (+2/-0).

En principe, partout où des solutions intelligentes développées par des être intelligents peuvent être applicables à d'autres problèmes, un interpolateur géant (alias LLM ou IA) entre de bonnes mains peut s'avérer un outil puissant (au figuré, comme au — pas si — propre -;-) ; non ?



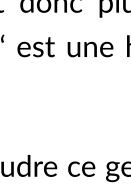
« IRAFURORBREVISSESTANIMUMREGEQUINISIPARETIMPERAT » — Odes — Horace

Répondre

[^] # Re: similaire

Posté par BAud (site web personnel) le 14 septembre 2026 à 16:06. Évalué à 3 (+1/-0).

je peux voir l'article en entier sans être abonné.



Répondre

[^] # Re: similaire

Posté par BAud (site web personnel) le 14 septembre 2026 à 16:06. Évalué à 3 (+1/-0).

j'ai le droit à « Vous avez lu 10 % de cet article. La suite est réservée aux abonnés.



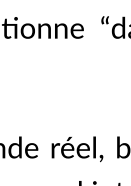
Déjà abonné-e ? Me connecter » et je n'accepte pas les cookies non plus (j'ai bien mangé ce midi).

Répondre

[^] # Re: similaire

Posté par 42nodid le 14 septembre 2026 à 16:41. Évalué à 2 (+1/-0).

C'est un softpaypal facilement contournable. Au choix, mode lecture ou rechargement en mode developpement de firefox,



Répondre

Un peu, beaucoup risqué

Posté par Bernie le 14 septembre 2026 à 13:45. Évalué à 9 (+8/-1).

Si on laisse les IAs résoudre les grandes questions de mathématiques, est-ce qu'on ne va pas arriver à un point où, face à une "preuve", plus personne (comprendre : aucun humain) ne saura l'interpréter à fond et donc plus personne ne pourra vérifier si cette "preuve" est une hallucination ou une véritable démonstration ?



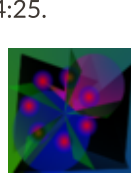
Et du coup, quel est l'intérêt de faire résoudre ce genre de problème à une IA ? Du point de vue le accroissement de la connaissance humaine, et de l'éventuel progrès technique induit je veux dire...

Répondre

[^] # Re: Un peu, beaucoup risqué

Posté par Dring le 14 septembre 2026 à 13:49. Évalué à 6 (+5/-1).

Dans un monde idéal, utiliser l'IA pour donner des pistes, voire trouver des solutions est une étape qui permet ensuite aux humains de s'approprier (au sens "comprendre et exploiter") la découverte. Donc ça a totalement du sens de leur demander de résoudre des problèmes mathématiques complexes.



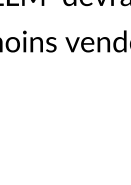
Dans le monde réel...

Répondre

[^] # Re: Un peu, beaucoup risqué

Posté par Bernie le 14 septembre 2026 à 13:53. Évalué à 2 (+1/-0).

Pour passer une étape bloquante et permettre aux mathématiciens humains d'avancer, OK. Mais il semble qu'actuellement, les IAs ne se contentent pas de ça. D'où ta précaution qui mentionne "dans un monde idéal".



Et comme on est plutôt dans le monde réel, ben... je trouve qu'une telle démarche n'a pas grand intérêt.

Répondre

[^] # Re: Un peu, beaucoup risqué

Posté par guily42 le 14 septembre 2026 à 14:11. Évalué à 3 (+2/-0). Dernière modification le 14 septembre 2026 à 14:12.

Je n'ai pas lu l'article (payant), mais il semble qu'il y ait de sérieuses réticences dans le métier :



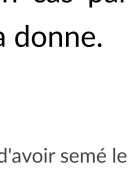
https://www.lemonde.fr/idees/article/2026/09/11/la-mise-en-garde-de-25-medailles-felds-les-objets-des-entreprises-d-ia-et-ceux-de-la-communauté-mathématique-divergent-profondement_6770630_3232.html

Répondre

[^] # Re: Un peu, beaucoup risqué

Posté par Bernie le 14 septembre 2026 à 14:23. Évalué à 1 (+0/-0).

Article réservé aux abonnés pour pouvoir lire la totalité, mais rien qu'une partie du titre :



Les objectifs des entreprises d'IA et ceux de la communauté mathématique divergent profondément

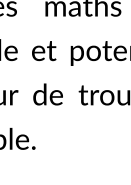
suffit à se forger une opinion, même si quitconque de sensé trouve évident que ces objectifs sont totalement divergents : le pouvoir et le fric pour les uns, le progrès des connaissances pour les autres. -;

Répondre

[^] # Re: Un peu, beaucoup risqué

Posté par mahikeulbody le 14 septembre 2026 à 14:25. Évalué à 2 (+1/-0).

Pour avoir lu l'article — mode lecture de Firefox — il semble que :



1) L'objection soit parfaitement légitime (risque de ne pas comprendre, ou de fausse preuve, comme avec du code vibe-codé qui tombe en marche) ;

2) Que les [LLM et pas l'IA](#) (faire attention à la nuance — l'article l'explique bien) sont bien employés (ici) pour compléter des étapes non triviales, mais pas non plus innovantes.

3) Les LLM effectivement ont tendance à ne pas se contenter de faire ce qu'ils peuvent faire (interpoler), et tentent souvent des « solutions novatrices » connues sous le nom d'hallucinations — en fait des extrapolations aventureuses hors de leur domaine d'entraînement, ou sur des parties non régulières du corpus. D'où l'importance d'une prise en main avancée et très sérieuse. Tout ce qu'écrit un LLM devrait être sujet à caution ; même si c'est moins vendeur sur le marché du bossware.

...

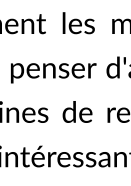
« IRAFURORBREVISSESTANIMUMREGEQUINISIPARETIMPERAT » — Odes — Horace

Répondre

[^] # Re: Un peu, beaucoup risqué

Posté par fearan le 14 septembre 2026 à 15:41. Évalué à 4 (+2/-0).

Tout ce qu'écrit un LLM devrait être sujet à caution



Je n'ai pas lu l'article mais si l'IA produit une démonstration en langage LEAN (ce qui semble le cas ici), la validité de cette démonstration est vérifiable par un humain. En revanche, j'imagine que saisir l'essence de la démonstration est une autre paire de manche si cette démonstration fait des centaines de pages.

Répondre

[^] # Re: Un peu, beaucoup risqué

Posté par fearan le 14 septembre 2026 à 15:48. Évalué à 6 (+3/-0).

Je n'ai pas lu l'article mais si l'IA produit une démonstration en langage LEAN



Attention une autre partie du problème consiste à s'assurer que l'exposé de base en LEAN correspond bien à l'énoncé; je me souviens de m'être planté dans un exercice de pivot de gauss où je m'étais planté en recopiant l'énoncé (j'avais mis un - à la place du +); et là on est dans un cas simple.

modéliser ton problème de math en LEAN est loin d'être simple et un cas particulier peut subtilement changer la donne.

...

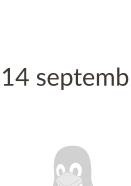
Il ne faut pas décorner les boeufs avant d'avoir semé le vent

Répondre

[^] # Re: Un peu, beaucoup risqué

Posté par thooms le 14 septembre 2026 à 16:39. Évalué à 2 (+0/-0). Dernière modification le 14 septembre 2026 à 16:40.

modéliser ton problème de math en LEAN est loin d'être simple et un cas particulier peut subtilement changer la donne.



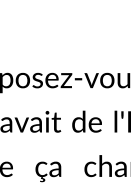
Je n'ai pas dit qu'on pouvait se passer des mathématiciens (par exemple pour formaliser les problèmes à résoudre), je disais juste qu'on peut, dans le cas particulier des mathématiques, éviter/détecter les hallucinations.

Répondre

[^] # Re: Un peu, beaucoup risqué

Posté par thooms le 14 septembre 2026 à 16:46. Évalué à 3 (+0/-0).

C'est difficile de savoir si la machine a exploité un bug subtil de lean pour prouver un lemme. Mais l'énoncé du théorème a pu être vérifié je crois, après, à part le bug de lean et/ou sa librairie, on peut être raisonnablement sûr. Par contre l'intérêt pour la communauté des maths d'une démonstration impénétrable et potentiellement moche dont c'est dur de trouver de nouvelles idées est discutable.



Et l'attitude des éditeurs d'IA est tout autant critiquée, ils ont vu qu'il y avait des avancées dans la communauté, alors ils ont mis tous les moyens sur le problème, dumpé la démo et passent à autre chose. Pas courant et très compétitif comme approche, les standards de la communauté ne sont pas respectés, soupçons de piquage d'idée sans crédit invérifiable (ils ne savent pas sur quoi le bousin a été entraîné ou pas ...)

Répondre

[^] # Re: Un peu, beaucoup risqué

Posté par Pol uX (site web personnel) le 14 septembre 2026 à 17:32. Évalué à 3 (+2/-1).

on est donc loin du fantasme du résultat foireux sans inventivité. On est déjà clairement sur des performances littéralement "surhumaines"

T'emballa pas. :)

...

Adhérer à l'April, ça vous tente ?

Répondre

[^] # Re: Un peu, beaucoup risqué

Posté par Pol uX (site web personnel) le 14 septembre 2026 à 17:32. Évalué à 3 (+2/-1).

Plus prosaïquement, le résultat (remarquable !) sur le problème d'Erdos peut aussi être interprété ainsi : une armée de chimpanzés est plus susceptible de fournir une solution au problème si elle est munie d'un bon crible (en l'occurrence Lean). Et si elle cherche à faire autre chose que d'exploiter les trous du crible. Et le résultat est inédit si elle n'a pas volé les résultats de quelqu'un d'autre.

...

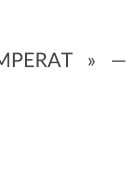
Adhérer à l'April, ça vous tente ?

Répondre

David Madore

Posté par LaurentClaessens (site web personnel) le 14 septembre 2026 à 19:22. Évalué à 9 (+7/-0).

L'avis de [David Madore](#) sur le sujet.



Précision : je ne suis pas forcément d'accord avec David Madore sur ce point. Mais son avis étant tout de même intéressant, je transmets.

Trop Long; Donne Résumé :

L'intérêt des grands problèmes de math est de découvrir des nouvelles questions, de nouveaux points de vue en recherchant à résoudre le grand problème. La réponse oui/non n'a pas tellement d'importance.

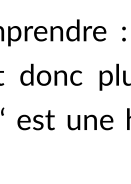
Par conséquent, résoudre ces questions avec de l'IA tend à assécher le viviers de nouvelles questions mathématiques. Bref : OpenAI tue les mathématiques.

Répondre

[^] # Re: David Madore

Posté par LaurentClaessens (site web personnel) le 14 septembre 2026 à 19:22. Évalué à 3 (+1/-0).

Parmi les commentaires, il y en a un qui explique pourquoi l'argument écologique n'est pas très bon. Je transmets.



Avant d'utiliser l'argument écologique, posez-vous cette questions : est-ce que si par magie on avait de l'IA sans aucun impact écologique, est-ce que ça changerait votre opinion ?

SI NON

Alors utilisez plutôt les autres arguments que vous avez : ils seront probablement plus solides.

SI OUI

Supposons que, sous l'hypothèse que l'IA était neutre, alors elle serait bénéfique.

Alors il faut éviter le sophisme du « dernier arrivé ». C'est à dire que, parmi tous les trucs bénéfiques qui ont un impact écologique, pourquoi refuser spécifiquement le dernier arrivé ?

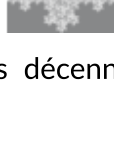
...

M'est avis que la quasi totalité des personnes utilisant l'argument écologique seront dans la branche "si non".

Répondre

[^] # **Re: David Madore**
Posté par [thoasm](#) le 14 septembre 2026 à 19:50.
Évalué à 4 (+1/-0).

Sur l'écologie, on a tendance à faire le plus facile en premier. Décarboner le monde entier ? Ça n'a rien d'évident. On doit faire ça dans les décennies qui viennent.



Dans ce cadre la question devient : est-ce qu'on est en train, en pleine urgence écologique, d'investir massivement sur un truc qui est au moins équivalent en terme de bénéfice écologique ?

Si non : ben on le fait pas et on investit dans la décarbonation (ou la décroissance) des trucs qui sont arrivés pas en dernier.

Répondre

[^] # **Re: David Madore**
Posté par [Ysabeau](#)  site web personnel, Mastodon) le 14 septembre 2026 à 19:57. Évalué à 2 (+2/-3).

Le problème avec "décarboner" c'est qu'on passe par des centrales nucléaires qu'il faut refroidir, et qu'il y a des petits génies qui sont en train de penser qu'il suffirait de faire partir cette chaleur dans l'atmosphère. Ce qui pose un léger problème de logique étant donné que le changement climatique actuel est dû au réchauffement engendré par les activités humaines.



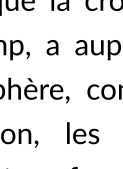
...

Je n'ai aucun avis sur systemd

Répondre

[^] # **Re: David Madore**
Posté par [thoasm](#) le 14 septembre 2026 à 20:07.
Évalué à 6 (+3/-0).

Ce n'est absolument rien par rapport au "forçage" qu'on induit en rajoutant du CO2, on trouve des chiffres comme "la chaleur globale ajoutée par l'atmosphère (à cause des gaz à effet de serre en plus) est de 4 bombes d'Hiroshima par ... seconde" : <https://skepticalscience.com/4-Hiroshima-bombs-worth-of-heat-per-second.html> (en 2013)



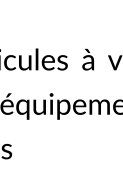
Autant dire qu'on est sur du rien de comparable avec le refroidissement des centrales nucléaires.

Il y a d'autres moyens de produire de l'élec décarbonée qu'avec des centrales nucléaires, les panneaux solaires, les éoliennes, et la biomasse par exemple (on profite de la décomposition des déchets organiques pour récupérer du gaz à faire brûler, c'est décarboné parce que la croissance des plantes, si on fait pas n'imp, a auparavant retiré ce carbone de l'atmosphère, comme le cycle normal de la respiration, les plantes poussent avec du CO2 en le transformant en matière comestible par nous, on rejette le même CO2 en respirant, équilibré à peu près)

Répondre

[^] # **Re: Truc oublié**
Posté par [fearan](#) le 14 septembre 2026 à 21:21.
Évalué à 3 (+0/-0).

Le problème des centrales nucléaires c'est surtout de réchauffer localement, les rivières par exemple.



Avec le bon équipement la centrale rejette de l'eau moins chaude <https://www.les-energies-renouvelables.eu/article/actualites/energies/la-seule-centrale-nucleaire-qui-refroidit-sa-riviere-en-france-1070/>

Et au vue des risque de canicules à venir y'a des chances pour que cet équipement soit ajouté aux centrales existantes

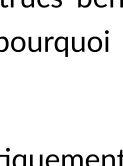
...

Il ne faut pas décorner les boeufs avant d'avoir semé le vent

Répondre

[^] # **Re: David Madore**
Posté par [Renault](#) (site web personnel, Mastodon) le 14 septembre 2026 à 20:10. Évalué à 7 (+4/-0). Dernière modification le 14 septembre 2026 à 20:10.

Avant d'utiliser l'argument écologique, posez-vous cette questions : est-ce que si par magie on avait de l'IA sans aucun impact écologique, est-ce que ça changerait votre opinion ?



SI NON
Alors utilisez plutôt les autres arguments que vous avez : ils seront probablement plus solides.

Et qu'est-ce qui empêche d'avoir plusieurs arguments valables en même temps et qui fait qu'ensemble ce soit un problème ?

Tout le monde a des sensibilités différentes, cela ne paraît pas aberrant de sortir tous les arguments valables à disposition même si certains sont plus forts dans l'absolu. Tant que cela a un sens évidemment.

C'est à dire que, parmi tous les trucs bénéfiques qui ont un impact écologique, pourquoi refuser spécifiquement le dernier arrivé ?

Je ne pense pas que ce soit systématiquement contre le dernier arrivé, c'est plutôt une question de structure.

Notre société, qu'on le veuille ou non, s'est structurée avec beaucoup de choses ce qui fait que décarboner / réduire l'impact écologique de ses éléments ne peut que prendre du temps. On pense au transport du quotidien (voiture), le chauffage des locaux / logements, les habitudes alimentaires, etc. Sans considérer que d'ailleurs ces éléments répondent souvent à des besoins vitaux.

C'est beaucoup plus facile de remettre en cause quelque chose de totalement nouveau dont la société n'a pas encore de dépendance pour éviter un impact important plutôt que de remettre en cause ce qui est déjà là.

Le défi est déjà important pour que chaque ajout d'un élément significatif nécessite des précautions plutôt que de renoncer à réguler sous prétexte qu'il y a pire ailleurs.

Répondre

[^] # **Re: David Madore**
Posté par [pulkomandy](#) (site web personnel, Mastodon) le 15 septembre 2026 à 00:15. Évalué à 9 (+6/-0). Dernière modification le 15 septembre 2026 à 00:15.

C'est à dire que, parmi tous les trucs bénéfiques qui ont un impact écologique, pourquoi refuser spécifiquement le dernier arrivé ?



qui a dit qu'on ne pouvait refuser que le dernier? Je connaît plein de gens qui préfèrent le train à l'avion, le vélo à la voiture, le matériel informatique d'occasion au neuf, le logiciel libre au logiciel propriétaire, le tout même si ça peut rendre la vie plus difficile.

La différence quand c'est le "dernier arrivé", c'est que la remise en question est plus facile. Si on avait préféré le train au tout voiture au siècle dernier, on aurait un réseau de train dense et bien entretenu et pas d'autoroutes. Si on avait préféré le logiciel libre partout, Microsoft n'existerait plus. Il en sera de même avec l'IA, qui va largement transfor@er au moins le web, et sûrement plein d'autres aspects de nos vies et de la façon dont on accède à de l'information a minima. Ce qui rendra de plus en plus difficile de faire marche arrière si on voulait un jour.

Répondre

[^] # **Re: David Madore**
Posté par [thoasm](#) le 15 septembre 2026 à 11:35.
Évalué à 3 (+0/-0).

Cerise sur le gâteau, évidemment ça ne tarde pas à être contre-productif pour la décarbonation ce genre de dépendance : l'IA fait de la thune entre autre en optimisant l'exploitation pétrolière de trucs qu'on ferait bien mieux de laisser là ou elles sont :



https://www.lemonde.fr/economie/article/2026/09/15/totalenergies-s-allie-a-mistral-pour-revolutionner-l-exploration-petroliere-grace-a-l-intelligence-artificielle_6774354_3234.html

Comment sucer encore plus les réservoirs pétroliers pour faire de la thune au prix de ... la vie de l'humain et autre sur Terre ? Le monde de l'IA ne s'en prive pas. C'est encore pire comme boucle de renforcement du système pétro-dollars.

Répondre

Envoyer un commentaire

Suivre le flux des commentaires

Note : les commentaires appartiennent à celles et ceux qui les ont postés. Nous n'en sommes pas responsables.

Revenir en haut de page

Derniers commentaires	Étiquettes (tags) populaires	Sites amis	À propos de LinuxFr.org
- Re: Envisageable ou p... - Re: David Madore - Re: Peut être dans les... - Re: Que faire? - Re: Qui paie ? - Al poisoning - Re: Déprimant - Re: Petites questions - Re: Envisageable ou p... - Re: Petites questions - Re: Que faire? - Re: Déprimant	- intelligence_artificielle - grands_modèles_de_J... - merdfication - hugging_face - états-unis - jihad_butlérien - capitalisme_de_survel... - administration_frança... - slop - capitalisme - startup_nation - bigtech	- Agenda du Libre - Aprii - Éditions D-BookER - Éditions Diamond - Éditions Eyrolles - Éditions ENI - En Vente Libre - Framasoft - La Quadrature du Net - Lea-Linux - Open Source Initiative - Imprimerie Grafik Plus	- Mentions légales - Faire un don - L'équipe de LinuxFr.org - Informations sur le site - Aide / Foire aux ques... - Suivi des suggestions ... - Wiki du site - Règles de modération - Statistiques - API pour le développ... - Code source du site - Plan du site