



The CLEAN PAPER

WHAT THE PAPER SAYS · WHAT IT DOESN'T · WHY IT MATTERS

چرا مدل‌های زبانی توهم می‌زنند — و چرا شیوه نمره‌دهی ما آن را نگه می‌دارد

دروغ‌های مطمئن که «توهم» می‌نامیم نقصی رازآلود نیستند: بخشی از آن‌ها محصول جانبی آماری آموزش‌اند، و باقی می‌مانند چون های benchmark اصلی حدس مطمئن را از «غی داتم» صادقانه بیشتر پاداش می‌دهند. یک مطالعه موردی روی چهار مدل مرزی نشان می‌دهد که بیان قواعد نمره‌دهی در «open prompt» (rubrics این انگیزه را وارونه می‌کند).

Written by Lucio Vaglio · Cross-checked by Laura Nesso · Edited by Michele Renda · Figures by Laura Nesso · AI-assisted, human-reviewed

Paper: *Evaluating large language models for accuracy incentivizes hallucinations*, Adam Tauman Kalai, Ofir Nachum, Santosh S. Vempala, Edwin Zhang, Nature ,653 1050–1047 (2026) · DOI: w-10549-026-s41586/1038.10

چرا مدل‌ها حدس می‌زنند، و چرا ما به آن‌ها چنین آموختیم

از یک مدل زبانی بزرگ تاریخ تولد یک غریبه را پرسید؛ ممکن است با اعتماد به نفس کسی که پاسخ را از روی کارت می‌خواند بگوید «۷ مارس» – و اشتباه کند، سه بار پشت سر هم، با سه تاریخ متفاوت. نویسندگان دقیقاً همین نوع مثال را می‌آورند: از مدل‌های پیشرو یک پرسش factual ساده پرسیده‌اند – تاریخ تولد یک شخص، یا معنی یک مخفف کم کاربرد – و هر کدام با اطمینان پاسخی متفاوت ساخته‌اند، بی‌آنکه هیچ کدام درست باشد. واژه صنعت برای این پدیده توهم است، واژه‌ای که آن را شبیه نقیصی در ادراک نشان می‌دهد. حرکت اول مقاله این است که رازآلودگی را از آن بگیرد.

از شیوه ساخته شدن مدل شروع کنید. در نخستین و بزرگ‌ترین مرحله آموزش، مدل در عمل می‌آموزد زبان روان چه شکلی دارد، با خواندن حجم عظیمی از متن. حالا یک واقعیت بی‌الگو را در نظر بگیرید – تاریخ تولد یک شخص خاص. اگر آن تاریخ در متن آموزش یک بار آمده باشد، یا اصلاً نیامده باشد، چیزی نیست که یک الگوآموز به آن چنگ بزند: پاسخ، از دید مدل، دل‌خواهی است. نویسندگان این را با وام گرفتن از ایده‌ای قدیمی دقیق می‌کنند (ایده آل تورینگ، از مسئله‌ای دیگر): اگر یکی از هر پنج تاریخ تولد فقط یک بار در داده‌ها ظاهر شود، باید انتظار داشت مدل دست کم در یکی از هر پنج مورد خطا کند – نه چون خراب است، بلکه چون هرگز چیزی برای یاد گرفتن وجود نداشته است. (با همین منطق، مدل‌ها تقریباً هیچ‌وقت پایتخت یک کشور را اشتباه نمی‌گویند: این‌ها مدام ظاهر می‌شوند.) آن‌ها با دقت استدلال می‌کنند که تشخیص یک گزاره درست از یک گزاره غلط محتمل خودش مسئله‌ای سخت است، و تولید فقط گزاره‌های درست دست کم به همان سختی است. یک کف خطا در مسئله پخته شده است.

yet seen haven't you what count to how underneath: idea The

شویم. آشنا آن با درست که دارد را آن ارزش و زبانی، مدل‌های از قدیمی‌تر – دارد تکیه هوشمندانه واقعاً ایده‌ای بر این با کیسه‌ای از توپ‌های رنگی شروع کنید. نمی‌دانید چند رنگ در آن هست. ۱۰۰ توپ را یکی یکی بیرون می‌کشید و می‌شمارید: قرمز ۴۰، آبی ۲۵، سبز ۱۵، زرد ۵، بنفش ۳، نارنجی ۲ – و سپس ده رنگ متفاوت که هر کدام دقیقاً یک بار ظاهر شده‌اند. اکنون پرسشی که تورینگ واقعاً با آن روبه‌رو بود، در مسئله‌ای کاملاً متفاوت: احتمال اینکه توپ بعدی رنگی باشد که اصلاً ندیده‌اید چقدر است؟ نمی‌توانید چیزی را که هرگز نکشیده‌اید بشمارید – اما می‌توانید رنگ‌هایی را بشمارید که دقیقاً یک بار دیده‌اید، یعنی «تک‌نماها». ترفند، که برآورد Good-Turing نام دارد، این است که سهم برداشت‌هایی که تک‌نما بوده‌اند احتمال پنهان‌مانده در رنگ‌هایی را تخمین می‌زند که هنوز ندیده‌اید. ده تا از صد برداشت شما رنگ‌هایی بودند که فقط یک بار دیده شدند، پس احتمال اینکه توپ بعدی رنگی کاملاً تازه باشد حدود $10 / 100 = 10\%$. آن رنگ‌های یک‌بار دیده اشتباه نیستند. آن‌ها اندازه‌گیری نادانی شما هستند: زیاد بودن رنگ‌هایی که فقط یک بار ظاهر می‌شوند راه نمونه برای گفتن این است که جهان چیزهای بیشتری دارد که شما هنوز نکشیده‌اید. حالا رنگ‌ها را با تاریخ تولد عوض کنید، و کیسه را با متن آموزشی مدل. فرض کنید در میان تاریخ تولدهایی که دیده است، یکی از هر پنج مورد دقیقاً یک بار ظاهر می‌شود. همان ترفند: حدود یک پنجم احتمال در تاریخ تولدهایی زندگی می‌کند که مدل عملاً ندیده است – و تاریخ تولد الگویی ندارد که بتوان به آن تکیه کرد (نمی‌توانید تاریخ تولد کسی را محاسبه کنید). بنابراین تاریخی که یک بار دیده شده، یا هرگز دیده نشده، سکه‌ای است که مدل نمی‌تواند وزنش را تعیین کند، و در حدود یکی از هر پنج مورد اشتباه خواهد کرد. هیچ مقدار هوشمندی این را درست نمی‌کند: چیزی برای یادگیری وجود نداشت. این کل استدلال در مقیاس کوچک است: نرخ تک‌نماها اندازه می‌گیرد چه مقدار از جهان از این داده‌ها یادگرفته‌اند نیست، و همان به کفی زیر خطاها تبدیل می‌شود. به همین دلیل هم مدل تقریباً هیچ‌وقت یک پایتخت را جا نمی‌اندازد – Paris مدام ظاهر می‌شود، نرخ تک‌نمای آن نزدیک صفر است، پس چیز زیادی برای یاد گرفتن وجود دارد.

این توضیح می‌دهد توهم‌ها از کجا می‌آیند. اما توضیح نمی‌دهد چرا باقی می‌مانند – چرا مدل‌ها، پس از همه آموزش‌های بعدی که قرار بوده مفید و راستگو شوند، همچنان بلوف می‌زنند به جای آنکه تردید را بپذیرند. اینجا قیاس مقاله تقریباً آزاردهنده قدر مناسب است. دانش‌آموزی

را تصور کنید که در امتحان پاسخ را نمی‌داند. اگر پاسخ خالی صفر بگیرد و حدس ممکن است یک امتیاز بگیرد، حرکت پیشینه‌کننده نمره حدس زدن است — با اعتماد به نفس، مشخص، و هرگز «مطمئن نیستم». دانش‌آموزان این را یاد می‌گیرند. ظاهراً مدل‌ها هم همین‌طور — چون ما به همان شکل نمره‌شان می‌دهیم. نویسندگان معیارهایی را که میدان واقعاً بر سر آن‌ها رقابت می‌کند، های‌لیدربرد را که مدل‌ها برای بالا رفتن از آن‌ها تنظیم می‌شوند، بررسی کرده‌اند و دیده‌اند تقریباً همه آن‌ها به «نمی‌دانم» دقیقاً همان نمره پاسخ غلط را می‌دهند: صفر. تحت چنین قاعده‌ای، مدلی که همیشه حدس می‌زند از مدلی که جز این یکسان است اما صادقانه عدم قطعیتش را علامت می‌زند جلو می‌افتد. ما، به معنایی نسبتاً لفظی، با نمره‌دهی آن‌ها را به این سو می‌رانیم.

Two ways to grade an answer

what each scoring rule rewards

Closed rubric

how most benchmarks score today

Correct	+1
Wrong	0
"I don't know"	0

Wrong and "I don't know" tie at 0, so a guess can only help.

Open rubric

scoring stated inside the question

Correct	+1
Wrong	-1
"I don't know"	0

A wrong guess now costs more than an honest "I don't know."

معیارها می‌توانند حدس زدن را عقلانی کنند. در نمره‌دهی‌ای که بیشتر معیارها استفاده می‌کنند (چپ)، پاسخ غلط و «نمی‌دانم» صادقانه هر دو صفر می‌گیرند — پس حدس فقط می‌تواند کمک کند. اگر قواعد در خود سؤال بیان شود، با جریمه برای خطا (راست)، پرهیز از پاسخ هنگام تردید می‌تواند حرکت بهتر شود. این عوض می‌کند آزمون چه چیزی را پاداش می‌دهد؛ به‌تنهایی توهم را حل نمی‌کند.

Original diagram — The Clean Paper CC BY 0.4

این بخشی است که باید نگه داشت، چون خلاف تیتیر معمول حرکت می‌کند. توهم اغلب به‌عنوان حدی ناگیر و تقریباً اسرارآمیز برای فناوری فروخته می‌شود. مقاله در هر دو مورد مخالفت می‌کند. کف پیش‌آموزش راز نیست — خطای آماری عادی است، از نوعی که یادگیری ماشین دهه‌هاست می‌فهمد. و تداوم آن ناگیر نیست — بخشی از آن انگیزه‌ای است که خودمان ساخته‌ایم و می‌توانیم تغییر دهیم. سامانه‌ای که وقتی مطمئن نیست صرفاً از پاسخ دادن سر باز زند اصلاً توهم نخواهد داشت؛ دلیل اینکه مدل‌های مستقر چنین رفتار نمی‌کنند این است که های‌سکوربرد ما امتناع را تنبیه می‌کند.

نویسندگان چه کردند

مقاله سه بخش دارد. نخست، استدلالی ریاضی که مقداری توهم در پیش‌آموزش از نظر آماری ناگیر است، با نشان دادن اینکه «تولید فقط متن معتبر» دست کم به سختی یک مسئله طبقه‌بندی دودویی «آیا این گزاره معتبر است؟» است. دوم، استدلالی — با پشتیبانی یک پیمایش از ده benchmark اثرگذار — که های metric سبک دقت رایج حدس زدن را برپرهیز ترجیح می‌دهند. سوم، اصلاح پیشنهادی و یک مطالعه موردی: ارزیابی‌های **open-rubric**، که در آن قواعد نمره‌دهی داخل خود سؤال بیان می‌شود (مثلاً «پاسخ درست ۱ می‌گیرد، پاسخ غلط ۱-، پس اگر کمتر از ۵۰٪ مطمئن هستید خودداری کنید»)، تا مدل بداند چه وقت صداقت پاداش می‌گیرد. آن‌ها این را روی چهار مدل مرزی — Claude Anthropic's، 4 Grok xAI's GPT-5، OpenAI's Pro، 3 Gemini Google's — با 5.4 Opus — 4,326 پرسش factual در SimpleQA آزموده‌اند. صریح می‌گویند مطالعه موردی illustrative است، «نه ارزیابی کنترل‌شده بین مدل‌ها» (تنظیمات پیش‌فرض، بدون tuning بدون نرمال‌سازی هزینه).

چه یافتند

- پیش‌آموزش مقداری خطا را تحمیل می‌کند. نرخ تولید دروغ‌های مطمئن از سوی مدل از پایین با تقریباً دو برابر نرخ خطای بهترین طبقه‌بند «آیا این گزاره معتبر است؟» ساخته‌شده از آن محدود می‌شود. برای facts بدون الگوی یادگرفتنی، آن کف دست کم نرخ تک‌نما است — سهمی facts که دقیقاً یک بار در آموزش آمده‌اند. حتی با داده کاملاً پاک هم مقداری توهم اجتناب‌ناپذیر است.
- نمره‌دهی به‌طور ملموس حدس را پاداش می‌دهد. تحت نمره‌دهی معمول درست/نادرست، هرگز خودداری نکردن راهبرد بهینه است، و پیمایش نویسندگان نشان می‌دهد اکثریت بزرگ های benchmark محبوب «غی‌دانم» را صرفاً غلط حساب می‌کنند. نمونه‌ای روشن از داده خودشان: در آزمون SimpleQA، دقت خام کمی به نفع o4-mini OpenAI's است — مدلی که تقریباً به همه چیز پاسخ می‌دهد و بیش از سه چهارم مواقع غلط است — نسبت به GPT-5-mini که چون هنگام تردید خودداری می‌کند خطاهای بسیار کمتری دارد. مدل بی‌احتیاط‌تر روی scoreboard بهتر دیده می‌شود.
- **rubrics open** انگیزه را وارونه می‌کنند (در مطالعه موردی آن‌ها). آن‌ها یک mitigation ساده برای توهم را آزمودند (مدل دو بار پاسخ دهد و اگر دو پاسخ ناسازگار بود خودداری کند). تحت دقت استاندارد، mitigation خطاها را کم می‌کند اما دقت را هم کم می‌کند — پس metric از به‌کارگیری آن باز می‌دارد. تحت **open rubrics**، همان mitigation برای هر چهار مدل در بازه‌ای از جرمیه‌ها بهتر بیرون می‌آید، و GPT-5-mini — که دقت خام به خاطر خودداری هنگام تردید تنبیهش کرده بود — وقتی نمره‌دهی آشکار بیان شود از o4-mini جلو می‌افتد ($n = 4,326$ پرسش برای هر مدل).

این احتمالاً چه معنایی دارد

کاهش توهم عمدتاً مسئله اختراع آزمون‌های اختصاصی بیشتر برای توهم نیست. مسئله تغییر نحوه نمره‌دهی های benchmark اصلی به عدم قطعیت است، تا اعتراف به «غی‌دانم» دیگر تنبیه نشود. تا وقتی scoreboard عوض نشود، کاهش توهم همچنان برای مدل‌ها امتیاز دقت هزینه خواهد داشت و بنابراین همچنان بازداشته خواهد شد — به همین دلیل نویسندگان مسئله را «اجتماعی-فنی» می‌نامند: بخشی metric بهتر، بخشی واداشتن های leaderboard اثرگذار به پذیرش آن.

این چه چیزی را ثابت نمی کند

- نشان نمی دهد rubrics open توهم را در عمل حل می کنند. آزمایش پشتیبان یک مطالعه موردی کوچک و عمداً کنترل نشده است — چهار مدل با تنظیمات پیش فرض، یک mitigation انتخابی، یک آزمون factual-QA — برای نشان دادن وارونگی انگیزه، نه رتبه بندی مدل ها یا اثبات کارایی عمومی.
- ادعا نمی کند غمزه دهی تنها علت است. خطاهای داده آموزش، مسائل واقعاً سخت، و های prompt ناآشنا همچنان منابع جداگانه اند.
- از خط رایج اجتناب ناپذیر بودن توهم ها پشتیبانی نمی کند. نویسندگان برعکس استدلال می کنند: سامانه ای که فقط به پرسش های قابل بررسی پاسخ دهد و در غیر این صورت بگوید «نمی دانم» هرگز توهم نخواهد داشت.
- کف پیش آموزش را ناپدید نمی کند — آن را توضیح و کران گذاری می کند، و این کران مربوط به خطاهای factual مطمئن است، نه همه رفتار مدل.
- نشان نمی دهد rubrics open به تنهایی کافی هستند. آن ها عوض می کنند ارزیابی چه چیزی را پاداش می دهد؛ جایگزین retrieval، استفاده از ابزار یا مدل های بهتر کالیبره شده نیستند.

شواهد چقدر قوی است؟

- هسته ریاضی است — کران های پایین رسمی، نه اندازه گیری. به عنوان استدلال نظری، در چارچوب خودش محکم است.
- بر مدل های عمداً ساده شده از مسئله تکیه دارد؛ خود نویسندگان «سه گانه کاذب» درست/نادرست/«نمی دانم» و setting آرمانی facts «دلخواهی» برای پاک ترین کران را علامت می زنند.
- پیمایش ها benchmark نمونه ای کوچک و گزینشی است — ده ارزیابی اثرگذار، نه audit کامل.
- مطالعه موردی واقعی اما محدود است: چهار مدل مرزی، یک mitigation فقط، SimpleQA تنظیمات پیش فرض، صراحتاً «نه ارزیابی کنترل شده». این concept of proof برای استدلال انگیزه است، نه نتیجه benchmark.
- ذکر point vantage مهم است: سه نفر از چهار نویسنده در OpenAI کار می کنند یا کرده اند، و مقاله استدلال می کند میدان باید نحوه ارزیابی مدل ها را تغییر دهد. این موضعی خوب استدلال شده از طرف ذی نفع است، نه review بیرونی بی طرف — باید وزن شود، نه کنار گذاشته. (به نفع مقاله، نقد را به مدل های خود، o4-mini و mini-5-GPT، همان قدر راحت وارد می کند که به دیگران.)

چرا مهم است

این مقاله مسئله ای پرهیاهو را از نو قاب بندی می کند. «توهم» معمولاً یا عیب اسرارآمیز فروخته می شود یا دیواری جابه جاشدنی؛ این مقاله آن را عادی و تا حدی خودساخته می کند — کفی آماری که واقعاً می توانیم بفهمیم، روی انگیزه ای که خودمان انتخاب کرده ایم. درس گسترده تر آرام تر و مفیدتر است: پیشرفت بیشتر در قابلیت اعتماد شاید به همان اندازه به آنچه می سنجیم بستگی داشته باشد که به آنچه می سازیم.

خلاصه تمیز

پاسخ‌های غلط اما مطمئن از مدل‌های زبانی از دو جا می‌آیند. اولی آماری است: وقتی ی‌fact الگویی برای یادگیری ندارد، مدلی که برای تقلید زبان آموزش دیده گاهی آن را غلط می‌گوید، و آن کف را می‌توان تخمین زد (مثلاً از تعداد یfacts که فقط یک بار در آموزش ظاهر می‌شوند). دومی انگیزه است: تقریباً هر یbenchmark که مدل‌ها بر اساس آن رتبه‌بندی می‌شوند «غنی‌دانه» را مثل پاسخ غلط نمره می‌دهد، پس حدس زدن همیشه می‌برد — تا حدی که مدلی که سه چهارم مواقع غلط است می‌تواند از مدل صادق‌تری که خودداری می‌کند جلو بزند. پیشنهاد نویسندگان آزمون توهم دیگری نیست، بلکه «open» rubrics است: قواعد نمره‌دهی را داخل سؤال بگویید. در مطالعه‌ای موردی روی چهار مدل مرزی، این انگیزه را وارونه می‌کند تا روش کاهش‌دهنده توهم پاداش بگیرد نه تنبیه. این مقاله‌ای نظری-همراه-پیمایش با آزمایشی کوچک و صریحاً کنترل‌نشده است، peer-reviewed در Nature؛ اصلاح امیدوارکننده است اما هنوز در مقیاس وسیع نشان داده نشده، و توهم‌ها نه اسرارآمیزند و نه به‌طور سخت‌گیرانه اجتناب‌ناپذیر.

بررسی بی‌تعارف

مقاله چه نشان می‌دهد: کران پایین ریاضی‌ای که طبق آن مقداری توهم در پیش‌آموزش تحمیل می‌شود (دست کم «نرخ تک‌نما» برای facts بی‌الگو)؛ پیمایشی که می‌گوید بیشترهای benchmark پیش‌رو به «غنی‌دانه» اعتبار نمی‌دهند؛ و مطالعه‌ای چهارمدلی که در آن بیان نمره‌دهی در «open prompt» («rubrics روش کاهش‌دهنده توهم را برنده می‌کند، جایی که دقت ساده آن را تنبیه کرده بود»)

چه چیزی محتمل اما اثبات‌نشده است: اینکه open rubrics، اگر واردهای benchmark اصلی شوند، توهم را در مدل‌های مستقر به شکل معنادار کاهش دهند. آزمایش پشتیبان کوچک و صریحاً کنترل‌نشده است.

چه چیزی را نشان نمی‌دهد: اینکه توهم‌ها اجتناب‌ناپذیرند (برعکس استدلال می‌کند)؛ اینکه نمره‌دهی تنها علت است؛ اینکه توهم را می‌توان کاملاً حذف کرد؛ یا اینکه مطالعه موردی چهارمدل را در برابر هم رتبه‌بندی می‌کند.

محدودیت‌های اصلی: مدل ساده‌شده درست/نادرست/«غنی‌دانه» (خود نویسندگان آن را «سه‌گانه کاذب» می‌نامند)؛ پیمایش کوچک و گزینشی‌های benchmark (ده ارزیابی)؛ مطالعه موردی کنترل‌نشده (چهارمدل، یک mitigation یک آزمون، تنظیمات پیش‌فرض)؛ و استدلالی به رهبری OpenAI درباره اینکه میدان چگونه باید مدل‌ها را ارزیابی کند.

خواننده عمومی چقدر باید مطمئن باشد؟ زیاد، نسبت به اینکه توهم نه اسرارآمیز است و نه کاملاً اجتناب‌ناپذیر، و های benchmark اصلی اکنون حدس زدن را پاداش می‌دهند. متوسط، نسبت به اینکه اصلاح پیشنهادی کمک می‌کند — اکنون concept of proof واقعی دارد، اما هنوز نشان نداده در گستره و مقیاس وسیع کار می‌کند.

منابع

Nature 653, 1047–1050 (2026)
<https://doi.org/10.1038/s41586-10549-026-w>

Why Language Models Hallucinate (arXiv:2509.04664)
<https://arxiv.org/abs/2509.04664>

OpenAI hallucinations-paper-experiments

<https://github.com/openai/hallucinations-paper-experiments>

SimpleQA

<https://github.com/openai/simple-evals>