



The CLEAN PAPER

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

Language models hallucinate क्यों करते हैं — और हमारी grading उन्हें ऐसा करते रहने के लिए क्यों प्रोत्साहित करती है

जनि आत्मवशिवसीझूठे उत्तरोकोहम “hallucinations” कहते हैं, वे कई रहस्यमय glitch नहीं हैं। कुछ training की statistical by-product हैं; और वे इसलिए बने रहते हैं क्योंकि mainstream benchmarks ईमानदार “मुझे नहीं पता” की तुलना में confident guess को अधिक reward करते हैं। चर frontier models पर case study देखी है कि prompt में scoring rules सफ़ लखि देना (“open rubrics”) इस incentive को उलट सकता है।

Written by Lucio Vaglio · Cross-checked by Laura Nesso · Edited by Michele Renda · Figures by Laura Nesso and Nova Vela · 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, 1047–1050 (2026) · DOI: 10.1038/s41586-026-10549-w

मॉडल अनुमति क्यों लगते हैं — और हमने उन्हें ऐसा करना क्यों सिखाया

किसी large language model से किसी अनजान व्यक्ति की जन्मतथि पूछिए। वह “7 मार्च” ऐसे आत्मविश्वास से बता सकता है जैसे काँड़ से पढ़ रहा हो — और लगता है तीन बार तीन अलग तरीके देकर तीनों बार गलत हो सकता है। लेखक ठीक ऐसा ही उदाहरण देते हैं: प्रमुख मॉडलों से किसी व्यक्ति की जन्मतथि या किसी दुर्लभ acronym का पूरा रूप पूछा गया और हर मॉडल ने आत्मविश्वास से अलग उत्तर गढ़ दिया, कोई सहिष्णुता उद्योग इसे hallucination कहता है, जिससे यह किसी धारणा संबंधी खरबी जैसा सुनई देता है। शोधपत्र का पहला कदम इस रहस्य को समझना बनता है।

शुरुआत इस बात से करें कि मॉडल बनता कैसे है। प्रशिक्षण के पहले और सबसे बड़े चरण में वह बहुत विशाल पठ-संग्रह पढ़कर, मटे टैर पर, धारणा वह भाषा के पैटर्न सीखता है। अब ऐसी जगह कर लें जिसके पीछे कोई पैटर्न नहीं है — किसी खास व्यक्ति की जन्मतथि। यदि वह तरीख प्रशिक्षण-पठ में एक बार आई या बिल्कुल नहीं आई, तो pattern learner के पास पकड़ने के लिए कुछ नहीं है: मॉडल की दृष्टि से उत्तर मनमा है। लेखक इसे एक पुराने विचार — अलग समस्या के लिए Alan Turing द्वारा इस्तेमाल विचार — से गणितीय रूप देते हैं। यदि पाँच में से एक जन्मतथि डेटा में केवल एक बार देखी देती है, तो मॉडल से कम से कम पाँच में से एक ऐसी जन्मतथि गलत हमें की अपेक्षा करनी चाहिए — इसलिए नहीं कि मॉडल टूटा है, बल्कि इसलिए कि सीखने के लिए पर्याप्त सूचना की भी नहीं है। इसी तरह से देश की राजधानी जैसे बातें मॉडल शायद ही गलत करता है, क्योंकि वे पठ में बार-बार आती हैं। लेखक समझती हैं कि सत्य कथन को विश्वसनीय लेकिन असत्य कथन से अलग पहचाना स्वयं कठिन समस्या है, और केवल सत्य कथन उत्पन्न करना कम से कम उतना ही कठिन है। इसलिए त्रुटि की एक न्यूनतम सीमा प्रशिक्षण में ही मौजूद रहती है।

मूल विचार: जो आपने अभी तक देखा ही नहीं उसे कैसे गिनें?

यह तरह एक सचमुच चतुर विचार पर आधारित है — भाषा मॉडलों से कहीं पुराना और अलग से समझने लायक।

रंगी गेंदों से भरे थैले की कल्पना करें। आपको नही पता उसमें कितने रंग हैं। आप एक-एक करके 100 गेंदें निकालते हैं और गनिते हैं: लाल 40, नीली 25, हरी 15, पीली 5, बैंगनी 3, नारंगी 2 — और फिर दस अलग रंग ऐसे हैं जो केवल एक-एक बार मिले।

अब वह सबल जिसे Turing ने एक अलग समस्या में पूछा था *अगली गेंद का रंग ऐसा हमें की संभावना कितनी है जो आपने अभी तक एक बार भी नहीं देखा?* जिसे कभी न किला ही नहीं उसे सीधे गनि नहीं सकते। लेकिन जो रंग **सर्वा एक बार मिले हैं** — singletons — उन्हें गनि सकते हैं। **Good-Turing estimation** नाम की तरकीब कहती है कि नमूने में singletons का अनुपात लगभग उस संभावना का अनुमान देता है जो अभी तक न देखे रंगों में छिपी है। 100 में 10 नकिली गेंदें ऐसे रंगों की थी जो केवल एक बार मिले, इसलिए अगली गेंद का बिल्कुल नया रंग हमें की संभावना लगभग $10 / 100 = 10\%$ है।

एक बार मिले रंग *गलत थिँन* हैं। वे आपकी अपनी अज्ञानता का माप हैं: यदि बहुत-से रंग केवल एक बार देखे, तो नमूना आप को बता रहा है कि दुनिया में ऐसे और रंग भी हैं जिन्हें आपने अभी तक न किला ही नहीं।

अब रंगों की जगह जन्मतथियाँ रखिए, और थैले की जगह मॉडल का प्रशिक्षण-पठ। मल लें कि देखी गई जन्मतथियों में **पाँच में से एक केवल एक बार देखी देती है।** वही तरकीब कहती है कि लगभग पाँचवाँ हिस्सा ऐसी जन्मतथियों की संभावना है जिन्हें मॉडल ने प्रभावी रूप से कभी ठीक से नहीं देखा। जन्मतथि में अनुमान लगाने लायक कोई समस्त पैटर्न भी नहीं होता। इसलिए एक बार या कभी न देखी गई तरीख पर मॉडल के पास सही उत्तर तक पहुँचने का भरोसा मंद आधार नहीं है, और लगभग **पाँच में से एक** पर गलती अपेक्षित है। केवल अधिक चतुर तरह से इसे ठीक नहीं किया जा सकता वहाँ सीखने के लिए पर्याप्त सूचना थी ही नहीं।

यही पूरे तरह का छोटो रूप है: singleton rate बताती है कि इस डेटा से दुनिया का कितना हिस्सा सीखना संभव नहीं है, और वही त्रुटियों के नीचे एक **न्यूनतम सीमा** बन जाता है। इसी वजह से राजधानी जैसे तथ्य बहुत कम चूकते हैं — *Paris* जैसे नाम लगता आते हैं, उनकी singleton rate लगभग शून्य है और सीखने के लिए बहुत कुछ मौजूद है।

यह बतता है कि hallucination पैदाकहाँसे होती है। लेकिन यह नहीं बतता कि बड़ के प्रशिक्षण के बमजुद वह बचीक्योर होती है — मॉडल संदेह मगने के बजय आत्मवश्वास से अनुमन क्योलगता है। यहाँ शोधपत्र की उपमा असहज रूप से सटीक है। परीक्षामें ऐसे छत्र की कल्पना कीजिए जिसे उत्तर नहीं पता। यदि खली छेड़ने पर शून्य अंक और अनुमन लगने पर सही हमें कीथेड़ी भी संभवना है, तो अंक अधिकतम करने की रणनीति है: अनुमन लगाना — आत्मवश्वास से, स्पष्ट उत्तर दो कभी “मुझे नहीं पता” मत कहो। छत्र यह सीखते हैं। मॉडल भी सीखते हैं, क्योंकि हम उन्हें लगभग इसी तरह अंक देते हैं। लेखकों ने उन benchmarks और leaderboards को देखा जिन पर क्षेत्र प्रतस्पर्धा करता है और जिनके अनुसार मॉडल optimize किए जाते हैं। लगभग सभी में “मुझे नहीं पता” को गलत उत्तर के समान ही अंक मिलता है: शून्य। ऐसे नियम में हमेशा अनुमन लगने वाला मॉडल उस बिल्कुल समान मॉडल से बेहतर स्कोर करेगा जो अनिश्चिति हमें पर ईमानदारी से रुक जाता है। कभी-कभी अर्थ में, हमारी scoring ही उन्हें अनुमन लगने के लिए प्रेरित करती है।

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."

Benchmarks अनुमन लगाना तर्कसंगत बना सकते हैं। समस्त scoring में (बाएँ) गलत उत्तर और ईमानदार “मुझे नहीं पता” दोनों को शून्य मिलता है, इसलिए अनुमन लगने से केवल फायदा हो सकता है। यदि प्रश्न में scoring सफ़ लखी जाए और गलत उत्तर पर दंड हो (दाएँ), तो अनिश्चिति हमें पर जवब न देना बेहतर रणनीति बन सकती है। इससे परीक्षण का प्रोत्साहन बदलता है; अपने-आप hallucination की समस्या हल नहीं होती।

Original diagram — The Clean Paper CC BY 4.0

यही हिसाब रखने लायक है, क्योंकि यह आम सुखीके वपरीत जाता है। Hallucination को अक्सर तकनीक की अनविध्य, लगभग रहस्यमय सीमा की तरह पेश किया जाता है। शोधपत्र दमोबत्तो को चुनौती देता है। Pretraining की न्यूनतम त्रुटि कोई रहस्य नहीं — यह समस्त संख्यकीय त्रुटि है, जिस तरह की चीज machine learning दशकों से समझती है। और बड़ में उसका बने रहना भी पूरी तरह अनविध्य नहीं — उसका एक हिसा हमारी बर्नई हुई incentive structure से आता है, जिसे बदला जा सकता है। यदि कोई प्रणाली अनिश्चिति हमें पर केवल उत्तर देने से मनाकर दे, तो वह उस स्थिति में hallucinate नहीं करेगी। तैनात मॉडल ऐसालगता नहीं करते क्योंकि हमारे scoreboards अक्सर उस ईमानदार refusal को दंडित करते हैं।

लेखकोने क्याकिया

शोधपत्र के तीसरे हिस्से हैं। पहला गणितीय तर्क है कि pretraining के दौरान कुछ hallucination संख्यात्मक रूप से अनविद्य है: “केवल वैध पष्ठ उत्पन्न करना” कम से कम binary “क्या यह कथन वैध है?” classification जटिलीकठिन समस्या है। दूसरा तर्क, दस प्रभावशाली benchmarks के सर्वेक्षण से समर्थित, यह है कि मुख्यधारा के accuracy metrics जब न देने की तुलना में अनुमान लगाने को प्रोत्साहित करते हैं। तीसरा एक प्रस्तावित सुधार और उसका छोटा case study है: **open-rubric evaluations**, जिनमें scoring नियम प्रश्न के भीतर ही साफ़ लिखे जाते हैं — जैसे “सही उत्तर +1, गलत -1; यदि 50% से कम भरोसा हो तो जवाब न दें” — तब मॉडल जवाब सके कि ईमानदार अनिश्चितता को इनमें मेलिगा। इसे उन्होंने चार अग्रणी मॉडलों — Google Gemini 3 Pro, OpenAI GPT-5, xAI Grok 4 और Anthropic Claude Opus 4.5 — पर SimpleQA के 4,326 factual questions के साथ आजमाया। वे स्पष्ट लिखते हैं कि यह केवल उदाहरणत्मक case study है, “मॉडलों के बीच नियंत्रित मूल्यंकन नहीं” — default settings, कोई tuning नहीं और लगत का normalization नहीं।

उन्हें क्या मिला

- **Pretraining कुछ त्रुटि अनविद्य बनाता है।** मॉडल द्वारा आत्मविश्वास से गलत तथ्य उत्पन्न करने की दर की एक नचिली सीमा है, जो मॉडल पर उससे बने सर्वोत्तम “क्या यह कथन वैध है?” classifier की त्रुटि दर से जुड़ी है। जिन तथ्यों में सीखने योग्य पैटर्न नहीं वहाँ यह सीमा कम से कम *singleton rate* — प्रशिक्षण में केवल एक बार दिखाई देने वाले तथ्यों का अनुपात — जटिली होती है। पूरी तरह साफ़ डेटा में पर भी कुछ factual hallucination से बचना संभव नहीं होगा।
- **Scoring अनुमान लगाने को प्रोत्साहित करता है।** समग्र सही गलत अंकन में कभी जवाब न छोड़ना सर्वोत्तम रणनीति बन जाती है, और लेखकों के सर्वेक्षण में लोकप्रिय benchmarks की बड़ी बहुमत “मुझे नहीं पता” को बस गलत मानती है। उनका स्पष्ट उदाहरण SimpleQA से है: raw accuracy OpenAI के o4-mini को थोड़ा बेहतर दिखाती है — जो लगभग हर सवाल का जवाब देता है और तीस-चौथे से अधिक बार गलत होता है — GPT-5-mini की तुलना में, जो अनिश्चितता हमें पर अधिक बार रुकता है और इसलिए बहुत कम गलतियाँ करता है। अधिक जोखिम लेने वाला मॉडल scoreboard पर बेहतर दिखाई देता है।
- **Open rubrics incentive को पलट सकती हैं — कम से कम इस case study में।** लेखक सरल hallucination-mitigation आजमाते हैं: मॉडल से दोबारा उत्तर लें और दोहरा उत्तर अलग हो तो जवाब न दें। समग्र accuracy में इससे गलतियाँ कम होती हैं, लेकिन accuracy score भी घटता है — यही metric इस उपग्रह को अपनाने से हतोत्साहित करती है। Open rubrics में penalties की अलग-अलग मात्रा पर यही उपग्रह चार मॉडलों के लिए बेहतर निकलता है; और GPT-5-mini, जिसे raw accuracy ने अधिक abstain करने के लिए दंडित किया था, scoring साफ़ बताए जाने पर o4-mini से आगे निकलता है। हर मॉडल के लिए $n = 4,326$ सवाल थे।

इसका शायद क्या मतलब है

Hallucination घटाने का बड़ा हिस्सा शायद और अधिक hallucination-specific tests बनाने में नहीं बल्कि मुख्यधारा के benchmarks में अनिश्चितता को अंक देने का तरीका बदलने में है, तब “मुझे नहीं पता” कहना दंडित न हो। जब तक scoreboard नहीं बदलता, hallucination कम करने से मॉडल को accuracy points खोने पड़ सकते हैं और इसलिए ऐसी रणनीतियाँ अपनाने का incentive कमजोर रहेगा। इसी वजह से लेखक समस्या को “socio-technical” कहते हैं: कुछ हिस्सा बेहतर metric का है, और कुछ यह मनवसे का कि प्रभावशाली leaderboards उसे अपनाएँ।

यह क्या सक्षमता नहीं करता

- यह नहीं देखता कि open rubrics वास्तविक दुनिया की hallucination समस्या हल कर देती हैं। सहस्रक प्रयोग छोटा और जटिल बनकर **uncontrolled case study** है — चार मॉडल, default settings, एक चुना mitigation, एक factual-QA test — जिसका उद्देश्य *incentive बदलना* देखना है, मॉडलों की ranking या वास्तविक प्रभावशीलता सक्षमता करना नहीं।
- यह नहीं कहता कि scoring ही एक मात्र कारण है। Training data की त्रुटियाँ, सचमुच कठिन समस्याएँ और अपरचित prompts अलग स्रोत बने रहते हैं।
- यह लेख प्रिय दल का समर्थन नहीं करता कि hallucinations अनिवार्य हैं। लेखक उल्टा तर्क देते हैं: ऐसी प्रणाली जो केवल जाँचे जा सकने वाले सवालों का उत्तर दे और बक़ी पर “मुझे नहीं पता” कहे, hallucinate नहीं करेगी।
- यह pretraining की न्यूनतम त्रुटि **ग़ायब नहीं करता**। शोधपत्र उसे समझता और सीमा देता है; वह सीमा आत्मविश्वास से दिए गए factual errors के बारे में है, मॉडल के हर व्यवहार के बारे में नहीं।
- यह नहीं देखता कि open rubrics अकेले पर्याप्त हैं। वे evaluation में मलिन वाले reward को बदलती हैं; retrieval, tool use या बेहतर calibrated models की जगह नहीं लेती।

सक्षम कतिनामजबूत है?

- मुख्य आधार गणति है — औपचारिक lower bounds, न कि केवल मझे गए correlations। सैद्धांतिक तर्क अपनी धारणाओं के भीतर मजबूत है।
- यह समस्या को जटिल बनकर सरल किए गए मॉडल पर टिका है। लेखक स्वयं हर उत्तर को सही गलत या “मुझे नहीं पता” में बाँटने को “false trichotomy” कहते हैं, और सबसे सख्त सीमा के लिए “arbitrary facts” वाली आदर्शकृत परिस्थिति स्वीकार करते हैं।
- Benchmark survey **छोटा और चुना हुआ नमूना** है — दस प्रभावशीलता evaluations, पूरे क्षेत्र का exhaustive audit नहीं।
- Case study **वास्तविक लेकिन सीमित** है: चार frontier models, एक mitigation, केवल SimpleQA, default settings और स्पष्ट धारणा कि यह “controlled evaluation” नहीं है। यह incentive तर्क का proof of concept है, benchmark result नहीं।
- दृष्टिकोण का स्रोत बताना भी उचित है: चार में से तीनों लेखक **OpenAI** में काम करते हैं या कर चुके हैं, और शोधपत्र क्षेत्र को model evaluation बदलने की सलाह देता है। यह हतिधारक पक्ष से आया सुवचरित तर्क है, न कि पक्ष बहरी समीक्षा नहीं — इसे तौलना चाहिए, केवल इसी कारण खराज नहीं। श्रेय की बात है कि शोधपत्र अपनी आलोचना To4-mini और GPT-5-mini जैसे OpenAI मॉडलों पर भी उतनी ही आसानी से लागू करता है।

यह क्यों मग्न रखता है

यह बहुत hype वाली समस्या को नए ढंग से रखता है। “Hallucination” को अक्सर या तो रहस्यमय खरबी या अचल दीवार की तरह पेश किया जाता है; यह शोधपत्र इसे सघन और आंशिक रूप से हमारी अपनी बनई समस्या बनाता है — एक संख्यकीय न्यूनतम सीमा जो समझा जा सकती है, और उसके ऊपर एक incentive जो हमने चुना है। बड़ा सबक शंत लेकिन अधिक उपयोगी है: reliability में आगे की प्रगति उतनी ही इस बात पर निर्भर हो सकती है कि हम क्या मानते हैं जितनी इस पर कि हम क्या बनाते हैं।

संक्षेप में

भाषामॉडलों के आत्मवश्र्वस से दिए गए गलत तथ्य दोस्तों से आते हैं। पहला संख्यिकीय है: जिस तथ्य के पीछे सीखने लयक कोई पैटर्न नहीं भाषाकीनकल सीखने वाला मॉडल कभी-कभी उसे गलत बताएगा। इस न्यूनतम त्रुटि का अनुमान लगाया जा सकता है, उदाहरण के लिए प्रशिक्षण में केवल एक बार दिखाई देने वाले तथ्यों के अनुपत्ति से। दूसरा incentive है: लगभग हर ऐसा benchmark जिस पर मॉडलों की ranking होती है, “मुझे नहीं पता” को गलत उत्तर के बराबर अंक देता है, इसलिए अनुमान लगाना हमेशा फायदे का सौदा है — इतना कि तीस-चौर से अधिक बार गलत हमें वलामॉडल भी अधिक ईमानदारी से abstain करने वाले मॉडल से बेहतर raw score पासकता है। लेखकों को प्रस्तुत कोई नया hallucination test नहीं बल्कि “open rubrics” है: scoring नियम प्रश्न में ही सफ़ लखें। चार frontier models पर छोटे case study में इससे incentive बदल गया और hallucination घटने वाला तरीका दंडित हमें के बजाय पुरस्कृत हुआ। यह theory + survey वाला शोधपत्र है, जिसके साथ छोटा और स्पष्ट रूप से uncontrolled experiment है। प्रस्तुत आशाजनक है, लेकिन अभी बड़े पैमाने पर सद्धि नहीं और मुख्य तर्क यह है कि hallucination न रहस्यमय है, न पूरी तरह अनविह्य।

बनाबढ़ाचढ़कर ज्ञान

शोधपत्र क्या दिखाता है: गणतिम lower bound जिसके तहत pretraining में कुछ hallucination अनविह्य है — patternless facts के लिए कम से कम singleton rate के स्तर पर; survey जिसमें अधकिंश प्रमुख benchmarks “मुझे नहीं पता” को कोई लाभ नहीं देते; और चार मॉडलों का case study जिसमें prompt के भीतर scoring सफ़ लखने — open rubrics — पर hallucination घटने वाला तरीका जितता है, जबकि सन्नरण accuracy उसे दंडित करती थी।

क्या संभव है, पर सन्नति नहीं कि open rubrics को मुख्य धारा के benchmarks में अपनाने से तैनात मॉडलों की hallucination अर्थपूर्ण रूप से घटेगी। सहायक प्रयोग छोटा और स्पष्ट रूप से uncontrolled है।

यह क्या नहीं दिखाता कि hallucinations पूरी तरह अनविह्य हैं — शोधपत्र उल्टा तर्क देता है; कि grading अकेला कारण है; कि hallucination पूरी तरह खत्म की जा सकती है; या कि case study चार मॉडलों की निष्पक्ष ranking है।

मुख्य सीमाएँ: सही गलत “मुझे नहीं पता” वाला जमबूझकर सरल मॉडल, जिसे लेखक “false trichotomy” कहते हैं; दस evaluations वाला छोटा curated benchmark survey; uncontrolled case study — चार मॉडल, एक mitigation, एक test, default settings; और model evaluation कैसे हो इस पर OpenAI से जुड़े लेखकों का तर्क।

समस्त पठक को कतिना भरोसा रखना चाहिए? इस बात पर उच्च कि hallucination न रहस्यमय है, न सख्त अर्थ में पूरी तरह अनविह्य, और वर्तमान mainstream benchmarks अक्सर अनुमान लगाने को पुरस्कृत करते हैं। प्रस्तुत सुधार बड़े पैमाने पर कतिना मदद करेगा इस पर मध्यम भरोसा रखें — वस्तुतः proof of concept है, व्यापक प्रदर्शन अभी नहीं।

स्रोत

Nature 653, 1047–1050 (2026)

<https://doi.org/10.1038/s41586-026-10549-w>

Why Language Models Hallucinate (arXiv 2509.04664)

<https://arxiv.org/abs/2509.04664>

hallucinations-paper-experiments (OpenAI)

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

SimpleQA

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