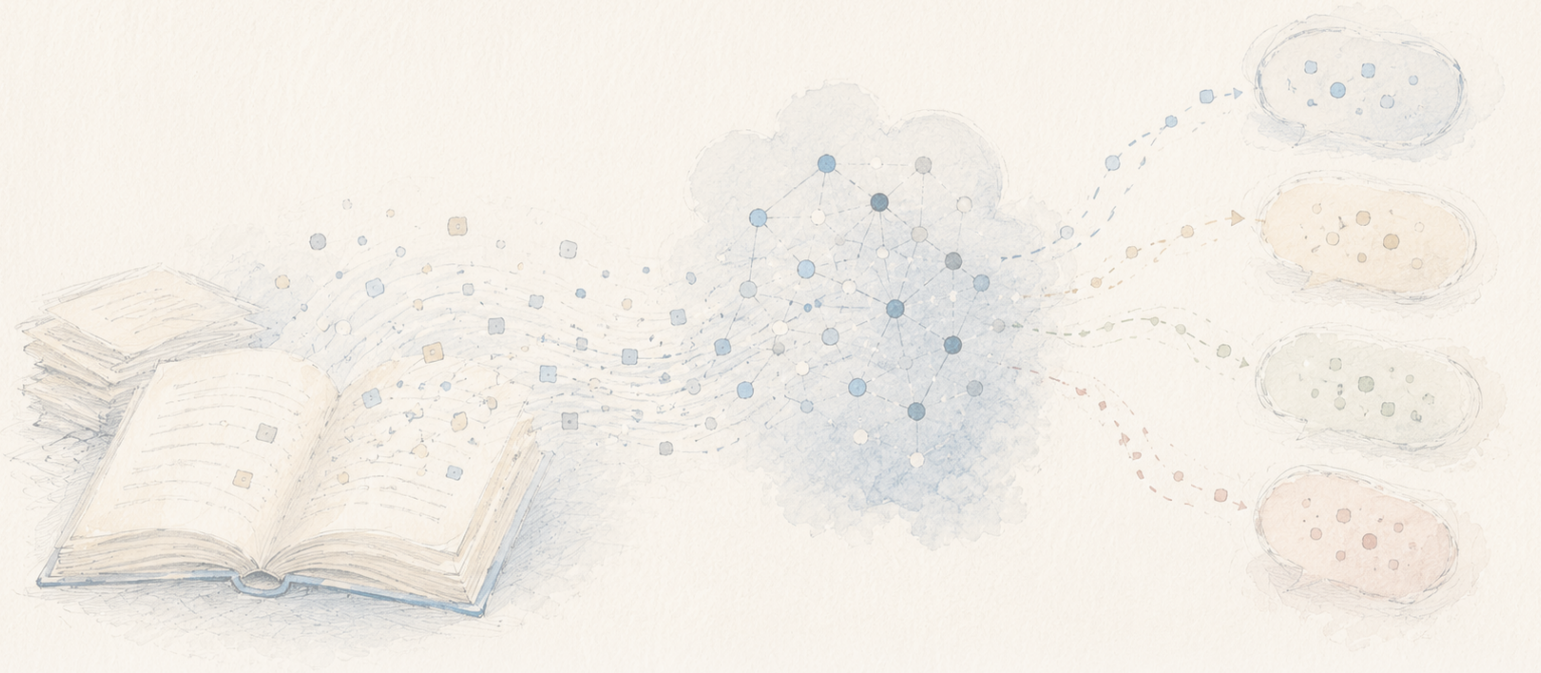




The CLEAN PAPER

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

Language model কনে hallucinate করে — এবং benchmark কনে guessing-কে reward করে

Paper statistical lower bound দিয়ে কিছু factual error-এর inevitability ব্যাখ্যা করে এবং দেখায় mainstream accuracy benchmark “I don’t know” ও wrong answer-কে একই score দলি়ে guessing rational হয়। Open-rubric case study incentive flip করে, কনিত্ত broad hallucination cure এখনও demonstrated নয়।

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

মডলে কনে আন্দাজ করে, আর আমরা কনে তাদরে তা করতে শখিয়িছে

একটি বহু ভাষা মডলে-কে কনে। অচনো ব্যক্তির জন্মদিন জিজ্ঞেসে করুন। সে হয়তো এমন আত্মবিশ্বাসে বলবে, “৭ মার্চ”, যনে কার্ড থেকে পড়ছে—কনিতু উত্তরটি ভুল। আবার জিজ্ঞেসে করলে অন্য তারখি, তৃতীয়বারে আরও একটি। গবষণাপত্ররে লখেকরো ঠিকি এমন উদাহরণ দনে: প্রধান ক্রমরে মডলেগুলোকে সাধারণ তথ্যভিত্তিকি প্রশ্ন করা হলে—কনে। ব্যক্তির জন্মতারখি, বা আড়াল করে সংক্ষিপ্ত রূপরে পূর্ণরূপ—প্রতিটি মডলে নরিদষ্টি ও আত্মবিশ্বাসী কনিতু আলাদা ভুল উত্তর বানায়। শলিপখাত এটাকে *হ্যালুসিনেশন* বলে, শব্দটি যনে উপলব্ধির অদ্ভুত তরুটি বোঝায়। গবষণাপত্ররে প্রথম কাজ হলে। রহস্যরে আবরণ সরানো।

মডলে কীভাবে তরৈ হয় তা থেকে শুরু করা যাক। প্রথম ও সবচেয়ে বড় প্রশিক্ষণ পর্যায়ে সে বপিল পাঠ্য পড়ে *fluent* ভাষার ধরন শখে। এখন এমন একটি তথ্য ননি যার পছনে শখোর মতো। ধরন নহে—একজন নরিদষ্টি মানুষরে জন্মতারখি। প্রশিক্ষণ পাঠ্যে তারখিটি একবার এসছে, বা একবারও আসনে, তাহলে ধরন শক্ষিণ-ব্যবস্থার ধরার মতো। কছু নহে। মডলেরে দৃষ্টিতে উত্তর-টি ইচ্ছামতো। লখেকরো Alan Turing-এর একটি পুরোনো পরিসংখ্যানগত ধারণা নযি এটকি নরিদষ্টি করনে: প্রশিক্ষণ-তথ্যে যদি পাঁচটির একটি জন্মতারখি মাত্র একবার দেখা যায়, তাহলে মডলেরে অন্তত প্রায় পাঁচটির একটিতে ভুল হওয়া প্রত্যাশিত—মডলে “ভাঙা” বলে নয়, বরং শখোর মতো। পর্যাপ্ত প্রমাণ ছিল না বলে। বপিরীতে দেশরে রাজধানী মডলে খুব কম ভুল করে, কারণ “Paris” বা অনুরূপ তথ্য তথ্যে বারবার আসে।

তারা আরও যুক্তি দনে যে কনে। সম্ভাব্য বরিতিসত্য ক মিথ্যা তা নরিভরয়ে। গ্যভাবে *classifying* করাই কঠনি; তাই **শুধু সত্য বরিতিতরৈ করা** অন্তত সেই শ্রণেবিনিয়াস সমস্যার মতো। কঠনি। অরখাৎ কছু তরুটির পরিসংখ্যানগত নমিনসীমা থকেই যায়।

যে ধারণাটি নিচি কাজ করছে: যা এখনো দেখেননি, তা কীভাবে গোনো যায়

এখানে একটি সুন্দর পুরোনো পরিসংখ্যানগত ধারণা কাজ করে।

রঙনি বলরে একটি খলকিল্পনা করুন। কত রং আছে জাননে না। 100টি বল একে একে তুললনে: লাল 40, নীল 25, সবুজ 15, হলুদ 5, বগুনী 3, কমলা 2—আর তারপর **10টি আলাদা রং, যগুলো এর প্রত্যেকেট মাত্র একবার করে দেখা গলে।**

প্রশ্নটি হলো: পররে বলটি এমন রঙরে হওয়ার সম্ভাবনা কত, যা আপনি এখনো একবারও দেখেননি? যা দেখেননি তা সরাসরি গোনো যায় না। কনিতু **যে রংগুলো ঠিকি একবার দেখা গেছে, অরখাৎ একক, সেগুলো গোনো যায়।** **ভালে—Turing estimation**-এর মূল অন্তরদৃষ্টি হলো: নমুনায় এককরে ভাগ, অদখো শ্রণে-তে এখনো কত সম্ভাবনা ভর লুকয়ি আছে তার অনুমান দয়ে। 100 draw-এর মধ্যে 10টি একক হলে পররে draw একবোরে নতুন রঙরে হওয়ার সম্ভাবনা আনুমানিক $10 / 100 = 10\%$ ।

এই একবার-seen রংগুলো *mistake* নয়। এগুলো আপনার অজ্ঞতার পরিমাপ: অনকে একক মানে নমুনা নজিই বলছে, বাইরে আরও শ্রণে আছে যা এখনো ঠিকিমতো দেখা হয়নি।

এখন রঙরে বদলে জন্মতারখি ধরুন, আর bag-এর বদলে মডলেরে প্রশিক্ষণ পাঠ্য। দেখো জন্মতারখিরে মধ্যে যদি **পাঁচটির একটি মাত্র একবার** উপস্থিতি হয়, তাহলে একই ববিচেনামূলক চিন্তা বলে বণ্টনরে উল্লেখ্যে গ্য অংশ মডলে কার্যত ভালে করে কখনো শখেনে। জন্মতারখিরে কনে। গণনাযে গ্য ধরন নহে—কারও জন্মতারখি “উৎপন্ন করা” করা যায় না। তাই একবার-seen বা অদখো তারখি মডলেরে ভুলরে একটি নমিনসীমা থাকে।

এটাই যুক্তির ক্যুদ্র সংস্করণ: একক হার বলে *এই তথ্য থেকে পৃথিবীর কতটা অংশ শখো কঠনি*, এবং সটেহি তরুটির নিচি একটি **নমিনসীমা** তরৈকিরে। Capital city-তে এই নমিনসীমা খুব ছোট, কারণ সেগুলো বারবার দেখা যায়।

এতে হ্যালুসিনেশন কখা থেকে আসে তার একটি অংশ বোঝা গলে। কনিতু পরে সহায়কতা ও সততা শখোনে। প্রশিক্ষণরে পরও মডলে কনে ভান করে উত্তর দয়ে, সটে এখনো বাকি। এখানে গবষণাপত্ররে উপমা খুব সরাসরি। পরীক্ষায় কনে। শক্ষিারখী উত্তর জানে না। Blank রাখলে শূন্য, আন্দাজ করলে শূন্য বা একটি—তাহলে গরুডে সরবর্ধকি করার কৌশল হলে। আন্দাজ করা। “আমি জাননি” বললে কনে। *upside* নহে। মডলেরে মানদণ্ডেও প্রায় একই প্রণে। লখেকরো ক্ষেত্রেরে প্রভাবশালী মানদণ্ডগুলো দেখনে এবং দেখনে অধিকাংশে “I don’t know” ভুল উত্তররে মতোই শূন্য পায়। এই সকে রহি নিয়মে অনিশ্চয়তা স্বীকার করা মডলেরে চয়ে সব সময় আন্দাজ করা মডলে বশেষিকেরে করত পারে। এক অরখাৎ আমরা মূল্যায়ন দয়িই মডলেকে ভান করে উত্তর দতি শখিয়িছে।

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."

মানদণ্ড আন্দাজ করাকে rational করে তুলতে পারে। সাধারণ স্কোরিং (বাম দিকে) ভুল উত্তর এবং সং “আমি জানি না” — দুটোই শূন্য, তাই আন্দাজ করলে কবেল লাভের সম্ভাবনা থাকে। প্রশ্নের মধ্যেই নিয়ম স্পষ্ট করে ভুল উত্তরকে penalty দিলে (ডান দিকে), অনশ্চয়তা বশেই উত্তর না দেওয়া ভালো। কষ্ট শেল হতে পারে। এটি মূল্যায়ন কী পুরস্কার করে তা বদলায়; নজি থেকে হ্যালুসিনেশনের পুরো সমাধান নয়।

Original diagram — The Clean Paper CC BY 4.0

এই অংশটি গুরুত্বপূর্ণ, কারণ প্রচলিত শরিতে নামের সঙ্গে এর বর্ণিত আছে। হ্যালুসিনেশনকে কখনও প্রযুক্তির রহস্যময় ও অবশ্যম্ভাবী সীমা হিসেবে দেখানো হয়। গবেষণাপত্র দুই দাবকিই নরম করে। পূর্বপ্রশিক্ষণ ত্রুটি নমিসীমা রহস্য নয়—সাধারণ পরিসংখ্যানগত ত্রুটি। আর হ্যালুসিনেশনের স্থায়িত্ব পুরো পুরি অনবির্য নয়—তার একটি অংশ আমাদের বানানো প্রণেদনা। কোনো ব্যবস্থা সত্যিই অনশ্চিতি হলে উত্তর না করলই hallucinate করবে না; মতোয়নে করা মডলে কনে এমন করো না, তার একটি কারণ স্কোরিং বর্ড refusal-কে শাস্তি দিয়ে।

লখেকরো কী করছেন

গবেষণাপত্রটির তিনটি অংশ। প্রথম গাণিতিক যুক্তি: পূর্বপ্রশিক্ষণে কল্পি হ্যালুসিনেশন পরিসংখ্যানগতভাবে forced, কারণ “শুধু বধৈ পাঠ্য তৈরী করা” অন্তত দ্বিমিকি “এই বর্ণিত বধৈ কী না?” শ্রণেবিনিয়াসের মতো কঠনি। দ্বিতীয় অংশে দশটি প্রভাবশালী মানদণ্ডের সমীক্ষা দিয়ে দেখানো হয়েছে মূলধারার নরিভুলতা-শিলী মটেরকি আন্দাজকে উত্তর না দেওয়ার চয়ে পুরস্কার করে। তৃতীয় অংশে একটি প্রস্তাবতি সংশোধন এবং কসে স্টাডি: **উন্মুক্ত-rubric মূল্যায়ন**, যখনে স্কোরিং নিয়ম প্রশ্নের মধ্যেই লখো থাকে—উদাহরণ হিসেবে, “সঠিক উত্তর +1, ভুল উত্তর -1; আস্থা 50%-এর কম হলে উত্তর না দেওয়া করো।” এতে মডলে জানে অনশ্চয়তা স্বীকার করা কখন পুরস্কার পাবে।

কসে স্টাডি-তে চারটি সীমান্ত মডলে ব্যবহৃত হয়েছে—Google-এর Gemini 3 Pro, OpenAI-এর GPT-5, xAI-এর Grok 4 এবং Anthropic-এর Claude Opus 4.5—এবং SimpleQA-এর 4,326 তথ্যভিত্তিক প্রশ্ন। লখেকরো স্পষ্ট বলনে এটি illustrative কসে স্টাডি, **নিয়ন্ত্রতি মডলে তুলনা নয়**: ডফিল্ট বনিয়াস, আলাদা সামঞ্জস্য নহে, খরচ normalization নহে।

তারা কী পয়েছেন

- **পূর্বপরীক্ষণ কিছু ত্রুটি বাধ্য করে।** মডলে যে আত্মবিশ্বাসী ভুল বহির্ভিত্তিক করে তার হারের একটি নমিনসীমা আছে, যা মডলে থেকে বানানো সেরা “বহির্ভিত্তিক কনি” শ্রুতিবিন্যাসকারীর ত্রুটির সঙ্কে সমপ্রকৃতি। শেখার মতো ধরন নই এমন তথ্যের ক্ষেত্রে নমিনসীমা অন্তত একক হার—প্রশিক্ষণে ঠিকি একবার দেখা তথ্যের ভগ্নাংশ। তথ্য পুরো পরিস্কার হলও কিছু হ্যালুসিনেশন অনবির্য হতে পারে।
- **স্কোরিং বাস্তবহই আন্দাজ করা পুরস্কার করে।** সাধারণ সংশোধন করা/ভুল স্কোরিং কখনও উত্তর না দেওয়া না করা সর্বোত্তম কৌশল হতে পারে। সমীক্ষায় দেখা যায় জনপ্রিয় মানদণ্ডের বড় অংশ “I don’t know”-কে সুরক্ষিত ভুল ধরে। SimpleQA-র একটি উজ্জ্বল উদাহরণ: কাঁচা নরিভুলতা সামান্য o4-mini-কে অনুকূল করে, যদিও এটি প্রায় সব প্রশ্নের উত্তর দিয়ে এবং তনি-চতুর্থাংশেরও বেশি ক্ষেত্রে ভুল; GPT-5-mini বেশি অনিশ্চয়তা স্বীকার করে বলে কম ভুল করবে কাঁচা স্কোরারের পছন্দি পড়তে পারে। Reckless মডলে দেখতে ভালো লাগে।
- **স্কোরিং নিয়ম আগের থেকে স্পষ্ট করে দিলে প্রণেদনা উল্টে যায়—অন্তত এই কেসে স্টাডিতে।** গবেষকেরা একটি সহজ প্রশ্নমন কৌশল পরীক্ষা করেন: মডলকে একই প্রশ্নের উত্তর দুবার দিতে বলা হয়, আর দুই উত্তর না মিললে উত্তর না দিতে বলা হয়। প্রচলিত নরিভুলতা মটেরিকে এতে ভুল কম, কিন্তু উত্তর দেওয়ার হারও কম—ফলে মটেরিকি এই সতর্ক আচরণকে নরিব্রহতি করে। কিন্তু স্কোরিং নিয়ম খোলাখুলি জানানো হলে একই কৌশল বিভিন্ন penalty-তে চারটি মডলেই ভালো স্কোর করে; এমনকি GPT-5-mini, কাঁচা নরিভুলতার জন্য অনুকূল করা o4-mini-কে ছাড়িয়ে যায় (n = 4,326 প্রশ্ন প্রতি মডলে)।

এর সম্ভাব্য অর্থ কী

হ্যালুসিনেশন কমাতে আরকেটি হ্যালুসিনেশন-নরিদৃষ্টি পরীক্ষা বানানোই প্রধান কাজ নাও হতে পারে। মূলধারার মানদণ্ড অনিশ্চয়তা কীভাবে স্কোর করে সেটি বদলানোও গুরুত্বপূর্ণ, যাতে “আমি জানি না” বলা স্বয়ংক্রিয়ভাবে শাস্তি না পায়। স্কোরারের বদল হ্যালুসিনেশন-হ্রাসকারী আচরণ নরিভুলতা বিন্দু খরচ করতে পারে, ফলে ডেভেলপারের প্রণেদনা উল্টে থাকে। এই জন্য লেখকেরা সমস্যাটিকে **socio-প্রযুক্তগত** বলেন: মটেরিকি নকশা এবং প্রভাবশালী লিডারের বদলে adoption—দুটোই দরকার।

এটাকী প্রমাণ করে না

- উন্মুক্ত rubric বাস্তব মতোয়নে করা ব্যবস্থায় হ্যালুসিনেশন ঠিকি করে দেয়—এটি গবেষণাপত্র দেখায় না। সমরখনকারী পরীক্ষা ছোট এবং ইচ্ছাকৃতভাবে **অনয়িন্ত্রতি**: চার মডলে, ডিফল্ট settings, একটি প্রশ্নমন, একটি তথ্যভিত্তিক-QA পরীক্ষা। উদ্দেশ্য প্রণেদনা উল্টে যাওয়া দেখানো, মডলে ranking বা সাধারণ কার্যকারিতা নয়।
- Grading-ই হ্যালুসিনেশনের **একমাত্র** কারণ—গবেষণাপত্র তা বলে না। প্রশিক্ষণ-তথ্য ত্রুটি, সত্যহি কঠনি সমস্যা এবং অপরচিত প্রম্পট আলাদা উৎস।
- হ্যালুসিনেশন অবশ্যম্ভাবী—জনপ্রিয় দাবিটিও গবেষণাপত্র সমরখন করে না। তাদের যুক্তি হলো, ব্যবস্থা যদি শুধু checkable প্রশ্ন উত্তর করে এবং অন্য ক্ষেত্রে “জানি না” বলে, তাহলে হ্যালুসিনেশন এড়ানো সম্ভব।
- পূর্বপরীক্ষণ নমিনসীমা অদৃশ্য হয়—উন্মুক্ত rubric তা করে না। গবেষণাপত্র নমিনসীমা ব্যাখ্যা ও সীমা করে; সীমা আত্মবিশ্বাসী তথ্যভিত্তিক ত্রুটিনিয়, সব মডলে আচরণ নিয়ে নয়।
- উন্মুক্ত rubric একই যথেষ্ট—এটিও নয়। মূল্যায়ন প্রণেদনা বদলানো retrieval, সরঞ্জাম ব্যবহার বা ভালো ক্যালিব্রেশনের বকিল্প নয়।

প্রমাণ কতটা শক্ত?

- মূল অংশটি **গাণিতিক**: আনুষ্টানিক নমিনসীমা, প্রায়শঃগিকি পরিমাপ নয়। তাত্ত্বিক যুক্তি তার অনুমানগুলো এর মধ্যে শক্ত।
- সমস্যাটির একটি **ইচ্ছাকৃতভাবে সরলীকৃত মডেল** ব্যবহার করা হয়েছে। লেখকরো নজিরোই “সঠিক / ভুল / ‘আমি জানি না’ ” এই তিন ভাগকে অতিরিক্ত সরল ত্রিবিভাজন বলনে এবং সবচেয়ে পরম্কার সীমা বরে করতে আদর্শায়তি, ইচ্ছামতে াতথ্য বনিয়াস ব্যবহার করেন।
- মানদণ্ড সমীক্ষা **হুটে টি ও curated**: দশটি প্রভাবশালী মূল্যায়ন, exhaustive নরীক্ষা নয়।
- কসে স্টাডি **বাস্তব কনিতু সীমতি**: চার সীমান্ত মডলে, এক প্রশমন, শুধু SimpleQA, ডফিল্ট settings, এবং স্পষ্টভাবে “not a নয়িন্তরতি মূল্যায়ন”। এটি প্রণে াদনা যুক্তরি ধারণার প্রাথমিকি প্রমাণ, মানদণ্ড ফল নয়।
- Vantage বনিদু-ও উল্লেখযে াগ্য: চার লেখকরে তনিজন OpenAI-তে কাজ করেন বা করছেন, এবং গবষণাপত্র ক্ষেত্রে মূল্যায়ন চর্চা বদলানে ার পক্ষযে যুক্তি দিয়ে। এটিকে dismiss করার কারণ নয়, কনিতু স্বাধীন নরিপক্ষে পর্যালোচনা হিসেবেও পড়া ঠকি নয়। কৃত্তিব হলো, critique তারা নজিদেরে o4-mini ও GPT-5-mini-এর ক্ষেত্রেও প্রয়োগ করছেন।

কনে এটি গুরুত্বপূর্ণ

এটি খুব অত্রিঞ্জতি প্রচার হওয়া একটি সমস্যাকে নতুনভাবে কাঠামো করে। “হ্যালুসিনেশন” কখনও spooky defect, কখনও immovable প্রাচীর হিসেবে বকিরি হয়। এই গবষণাপত্র এটিকে অনেকে বশো সাধারণ করে: একটি পরিসংখ্যানগত নমিনসীমা, যা বোঝা যায়; তার ওপর আমাদরে তরৈ প্রণে াদনা, যা বদলানো যায়। Quiet কনিতু উপযে াগী শকিয়া হলো, নরিভরযে াগ্যতার পরবর্তী অগ্রগতিকবেল মডলে স্থাপত্য বা প্রশকিষণ পরসির থেকে নাও আসতে পারে—**আমরা কী মাপি এবং কী পুরস্কার করি**, সটেডি সমান গুরুত্বপূর্ণ হতে পারে।

সংক্ষিপ্ত সারাংশ

ভাষা মডলে আতমবশিবাসী ভুলরে দুটি আলাদা উৎস থাকতে পারে। প্রথমটি পরিসংখ্যানগত: কনে াতথ্যরে শখোর মতো নয়িম না থাকলে এবং প্রশকিষণ-তথ্যে সটেডি খুব কম দেখা গেলে অনুকরণভিত্তিকি প্রশকিষতি মডলে কখনো না কখনো ভুল করবে; একক হাররে মতো রাশি থেকে সেই নমিনসীমা অনুমান করা যায়। দ্বিতীয়টি প্রণে াদনার সমস্যা: অধিকাংশ মানদণ্ডে “আমি জানি না” এবং ভুল উত্তর একই শূন্য স্কোর পায, ফলে আন্দাজ করাই যুক্তিসিঙগত কৌশল হয়ে দাঁড়ায়। এমনকি প্রায় তনি-চতুর্থাংশ ভুল করা বপেরে ায়া মডলেও, অনশিচতি হলে উত্তর না দেওয়া মডলে চয়ে কাঁচা নরিভুলতায় ভালো দেখাতে পারে। লেখকদরে প্রস্তাব হলো **স্কোরি নয়িম খে ালাখুলি জানানো**—প্রশ্নরে মধ্যযে কীভাবে নম্বর দেওয়া হবে তা বলা। চারটি সীমান্তবর্তী মডলে হুটে কসে স্টাডিতে এতে হ্যালুসিনেশন কমানোর আচরণ প্রচলতি নরিভুলতার তুলনায় বশো পুরস্কার পায। এটি তত্ত্ব, মানদণ্ড-সমীক্ষা ও হুটে অনয়িন্তরতি পরীক্ষার সমন্বয়; সমাধানটি আশাব্যঞ্জক, কনিতু বহু পরসিরে প্রমাণতি নয়। মূল বক্তব্য: হ্যালুসিনেশন রহস্যময়ও নয়, পুরো পুরি অনিবার্যও নয়।

সরাসরি যাচাই

গবষণাপত্র কী দেখায়: পূর্বপ্রশকিষণে কিছু তথ্যভিত্তিকি হ্যালুসিনেশনরে গাণিতিক নমিনসীমা; সমীক্ষায় অধিকাংশ প্রধান ক্রমরে মানদণ্ড “I don’t know”-কে কৃত্তিব দিয়ে না; এবং চার-মডলে কসে স্টাডিতে প্রম্পটরে মধ্যযে স্কোরি নয়িম দেওয়া হলে হ্যালুসিনেশন-হ্রাসকারী প্রশমন, সরল নরিভুলতা-তে penalised হওয়ার বদলে পুরস্কার পায।

কী সম্ভাব্য কনিতু প্রমাণতি নয়: মূলধারার মানদণ্ডে উনমুক্ত rubric চালু করলে মোতায়নে করা মডলে হ্যালুসিনেশন বাস্তবে উল্লেখযে াগ্যভাবে কমবে। সমর্থনকারী পরীক্ষা হুটে ও স্পষ্টভাবে অনয়িন্তরতি।

কী দেখায় না: হ্যালুসিনেশন অনবির্ষ—গবেষণাপত্র উল্টো যুক্তি দিয়ে; grading-ই sole কারণ; হ্যালুসিনেশন একবোরে eliminate করা যাবে; বা কসে স্টাডিচার মডেলে ন্যায্য ranking।

মূল সীমাবদ্ধতা: Simplified সংশোধন করা/ভুল/“I don’t know” মডেলে; দশটি মূল্যায়নের curated সমীক্ষা; চার মডেলে, এক প্রশমন, এক পরীক্ষা ও ডফিল্ট বনিয়াসের অনিয়ন্ত্রিত কসে স্টাডি; এবং মূল্যায়ন চরচা নিয়ে OpenAI-led যুক্তি।

সাধারণ পাঠকরে আস্থা কতটা হওয়া উচিত? উচ্চ আস্থা যে হ্যালুসিনেশন রহস্যময় বা strictly অনবির্ষ কখনো ঘটনা নয় এবং মূলধারার মানদণ্ড আন্দাজ করাকে পুরস্কার করতে পারে। প্রস্তাবিত উন্মুক্ত-rubric সংশোধন বসিতভাবে কাজ করবে—এতে মাঝারি আস্থা: ধারণার প্রাথমিক প্রমাণ বাস্তব, কিন্তু পরসির ও ব্যবহারিক মতোয়ানে প্রদর্শন এখনো নেই।

উৎস

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>