



The CLEAN PAPER

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

Ìdí tí language model fi ń hallucinate — àti ìdí tí ọ̀nà tá a fi ń fún wọn ní máàkì fi ń mú un dúró

Àwọn èké tí a ń pè ní “hallucination” tí model ń sọ pẹ̀lú ìdánílójú kì í ẹ̀ ẹ̀ àbùkù àdìtú: àwọn kan jẹ̀ èsì statistiki tí ìkòni, wọn sì ń dúró nítorí pé benchmark pàtàkì ń san èrè fún ìmẹ̀fò pẹ̀lú ìdánílójú ju “Mì ọ̀ mọ̀” tó jẹ̀ òtítọ̀ lọ. Case study lórí frontier model méréin fi hàn pé síso ọ̀nà fífúnni ní máàkì nínú prompt (“open rubric”) yí iwúrí náà padà.

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, 1047–1050 (2026) · DOI: 10.1038/s41586-026-10549-w

Ìdí tí àwọn model fi n méfò, àti ìdí tí a fi kò wọn láti ẹ̀ bẹ̀

Béèrè lówó large language model nípa ojó ibí àjèjì kan, ó lè dáhùn “March 7” pèlú idánilójú bí ẹnì tó n kà á lórí káàdì — kí ó sì sàšìṣe lẹmẹta ọtòtòtò pèlú ojó mẹta tó yàtò. Àwọn ònkòwé fúnni ní irú àpẹrẹ yìí gan-an: wọn béèrè ibéèrè òtítò lasán lówó àwọn model ašáájú — ojó ibí èniyàn kan, tàbí ìtumọ acronym kan tí a kò mò púpò — ọkọọkan wọn sì dá èsì mǐ sílẹ̀ pèlú idánilójú, tí kò sí èyikéyì tó tọ. Ọrọ tí ilé-iṣe n ló fún èyí ni *hallucination*, tó mú un dà bí àšìṣe nínú ohun tí a rí. Ìgbésẹ̀ àkókó iwé náà ni láti yọ àdìitú kúrò nínú rẹ.

Bèrẹ̀ pèlú bí a ẹ̀ kò model. Ní ìpele ìkòni àkókó tó sì tóbi jù lọ, ó kẹkọ̀, ní kókó, bí èdè tó n Ẹ̀n dáadàa ẹ̀ rí nípa kíka ọ̀pọ̀lọ̀pọ̀ ọ̀rọ̀. Wá mú òtítò kan tí kò ní àpẹrẹ lẹyìn rẹ̀ — ojó ibí èniyàn pátó kan. Bí ojó náà bá fara hàn nínú ọ̀rọ̀ ìkòni lẹ̀ẹ̀kan, tàbí kò fara hàn rárá, kò sí ohun tí akẹ̀kọ̀-àpẹrẹ̀ lè dì mú: láti ojú iwòye model, èsì náà jẹ̀ ohun tí kò ní ilànà. Àwọn ònkòwé sọ èyí di pípé nípa lílo èrò àtijọ̀ kan (ti Alan Turing, láti inú ìṣòro mǐ): bí ojó ibí kan nínú márùn-ún bá fara hàn lẹ̀ẹ̀kan ẹ̀so nínú dátà, a yẹ kí a retí pé model yòò sàšìṣe ó kéré tán lórí ọ̀kan nínú márùn-ún — kì í ẹ̀ nítorí pé ó bàjẹ̀, ẹ̀gbọ̀n nítorí pé kò sí ohun kan nǐbẹ̀ láti kọ. (Pèlú ìrònú kan náà, model fẹ̀rẹ̀ má sàšìṣe lórí olú-ilú orílẹ̀-èdè: wọn máa n fara hàn ní gbogbo ìgbà.) Wọn jiyàn pèlú ìṣọra pé yíya gbólóhùn òtítò sọ̀tò kúrò nínú èké tó ẹ̀e gbà jẹ̀ ìṣòro tó nira fúnra rẹ̀, àti pé pípèsè gbólóhùn òtítò *nìkan* ó kéré tán le tó ìṣòro yẹn. Ìpele àšìṣe kan wà nínú ètò náà láti ibèrẹ̀.

Èrò tó wà lábẹ̀ rẹ̀: bí a ẹ̀ n ka ohun tí a kò tǐ rí

Èyí dúró lórí èrò ọ̀lọgbọ̀n gidi kan — tó dàgbà ju language model lọ, tó sì yẹ kí a mò dáadàa.

Bèrẹ̀ pèlú àpò bọ̀lù aláwò. O kò mò iye àwò tó wà nínú rẹ̀. O fa 100 jáde, ọ̀kọ̀ọ̀kan lẹ̀ṣeṣe, o sì ka wọn: pupa 40, búlúú 25, ewé 15, yélò 5, aláwò elése àlùkò 3, ọ̀sàn 2 — lẹyìn náà **àwò mẹwàá ọ̀tòtòtò tí ọ̀kọ̀ọ̀kan fara hàn lẹ̀ẹ̀kan ẹ̀so.**

Wá wo ibéèrè tí Turing dojú kọ ní tòótò, lórí ìṣòro tó yàtò: *kí ni ànfààní pé bọ̀lù tó kàn jẹ̀ àwò tí o kò tǐ rí rárá?* O kò lè ka ohun tí o kò tǐ fa jáde — ẹ̀gbọ̀n o lè ka àwọn àwò tí o tǐ rí lẹ̀ẹ̀kan ẹ̀so, àwọn “singleton.” Ọ̀gbọ̀n náà, tí a n pè ní **Good–Turing estimation**, ni pé ìpín àwọn ìfà rẹ̀ tó jẹ̀ singleton n fọ̀jú díwọn probability tó sì fara pamọ̀ sínú àwọn àwò tí o kò tǐ rí. Mẹwàá nínú ọ̀górùn-ún ìfà rẹ̀ jẹ̀ àwò tí o rí lẹ̀ẹ̀kan ẹ̀so, nítorí náà ànfààní pé bọ̀lù tó kàn jẹ̀ àwò tuntun pátápátá jẹ̀ nńkan bí **10 / 100 = 10%.**

Àwọn àwò tí a rí lẹ̀ẹ̀kan náà kì í ẹ̀ àšìṣe. Wọn jẹ̀ ọ̀ṣùwọn àìmọ̀ tìrẹ̀: bí ọ̀pọ̀ àwò bá fara hàn lẹ̀ẹ̀kan, àpẹrẹ̀ náà n sọ fún ọ̀ pé ayé ní púpò sí i tí o kò tǐ fa jáde.

Wá fi ojó ibí rọ̀pò àwò, kí o sì fi ọ̀rọ̀ ìkòni model rọ̀pò àpò náà. Ká sọ pé, nínú àwọn ojó ibí tí ó rí, **ọ̀kan nínú márùn-ún fara hàn lẹ̀ẹ̀kan ẹ̀so.** Ọ̀gbọ̀n kan náà: nńkan bí idá márùn-ún probability wà nínú àwọn ojó ibí tí model kò rí ní tòótò — ojó ibí kò sì ní àpẹrẹ̀ tí a lè gbára lé (o kò lè *ṣíró* ojó ibí ẹnì kan). Nítorí náà, ojó tí a rí lẹ̀ẹ̀kan, tàbí tí a kò rí rárá, dà bí owó tí model kò lè fi iwọn sí, yòò sì sàšìṣe lórí nńkan bí **ọ̀kan nínú márùn-ún** wọn. Kò sí ọ̀gbọ̀n tó lè tún èyí ẹ̀: kò sí ohun nǐbẹ̀ láti kọ.

Èyí ni gbogbo àrìyànjiyàn nàà ní kùkùrú: iye singleton n díwòn iye ayé tí a kò lè kọ *látí inú dátà yìí*, èyí sì di **ààlà isàlẹ̀** fún àṣìṣe. Ó tún jẹ ìdírí tí model fẹrẹ́ má ẹ ṣàṣìṣe olù-ìlú kan — *Paris* fara hàn ní gbogbo ìgbà, iye singleton rẹ sún mọ ọ̀dọ, nítorí nàà ọ̀pọ̀lọ̀pọ̀ ohun ló wà látí kọ.

Èyí ṣàlàyé ibi tí hallucination ti wá. Kò ṣàlàyé ìdírí tí wòn fi yè — ìdírí tí àwọn model, lẹyìn gbogbo ìkóni tó tẹlẹ́ tí a ẹ látí mú wòn wúlò kí wòn sì jẹ olóòtítọ́, ẹ ní ẹ bí ẹni pé wòn mọ dípò kí wòn jẹwọ àní-àní. Nihìn-ín, àfiwé ìwé nàà fẹrẹ́ jẹ ohun tó n kó ènìyàn ní ìtìjù. Foju inú wo akẹkọ̀ọ̀ nínú ìdánwò tí kò mọ èsì. Bí ààyè ọ̀fọ bá gba ọ̀dọ, tí ìmẹfò sì lè gba ọ̀kan, ìgbésẹ̀ tó n mú máàkì pọ̀ jù ní látí mífò — pẹ̀lú ìdánilójú àti kíkún, kì í ẹ “Mi ọ́ dájú.” Àwọn akẹkọ̀ọ̀ n kọ èyí. Ó hàn pé àwọn model nàà n kọ ọ — nítorí pé bá yìí nàà ni a ẹ n fún wòn ní máàkì. Àwọn ònṅkọwé ṣàyèwò àwọn benchmark tí ẹka nàà n díje lórí gan-an, àwọn leaderboard tí a n ẹtò model látí gòkè, wòn sì rí i pé fẹrẹ́ gbogbo wòn n fún “Mi ọ́ mọ” ní máàkì kan nàà pẹ̀lú èsì àṣìṣe: ọ̀dọ. Lábẹ̀ ọ̀fin yẹn, model tó máa n mífò ní gbogbo ìgbà yòò borí model tó jọ ọ ní gbogbo ọ̀nà mii sùgbọ̀n tó n sọ àní-àní rẹ ní ọ̀títọ́. Ní ìtumọ̀ tó fẹrẹ́ jẹ tààrà, a n fi máàkì kọ wòn látí ẹ é.

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.”

Àwọn benchmark lè mú ìmẹfò jẹ ohun tó bójú mu. Lábẹ̀ ọ̀nà fífúnni ní máàkì tí ọ̀pọ̀ benchmark n ló (òsì), èsì àṣìṣe àti “Mi ọ́ mọ” tó jẹ ọ̀títọ́ n gba ọ̀dọ — nítorí nàà ìmẹfò lè ẹ̀rànwọ̀ nikan. Sọ àwọn ọ̀fin nàà nínú ibéèrè, pẹ̀lú ìjìyà fún àṣìṣe (ọ̀tún), ìkọ̀ látí dáhùn nígbà àní-àní sì lè di ìgbésẹ̀ tó dára jù. Ó n yí ohun tí ìdánwò n san èrè fún padà; fúnra rẹ, kò yanjú hallucination.

Original diagram — The Clean Paper CC BY 4.0

Apá yìí ló yẹ ká pa mọ, nítorí pé ó lòdì sí àkọlẹ́ tí a sàbà rí. Wòn sàbà máa n ta hallucination gégé bí ààlà tí kò ẹ́e yẹ, tó fẹrẹ́ jẹ ohun ìjìnlẹ́ ti technology. Ìwé nàà tako isọrọ̀ méjèèjì. Ààlà

ìsàlẹ̀ pretraining kì í ẹ̀ ẹ̀ àdìtú — àṣìṣe statistiki lasán nì, irú èyí tí machine learning tí mò fún ọ̀pọ̀ ọ̀dún. Ìdúró rẹ̀ pé náà kò sì ẹ̀ dandan — ní apá kan, iwúrí tí àwa fúnra wa kọ̀ nì, a sì lẹ̀ yí i padà. Ètò kan tó kàn kọ̀ láti dáhùn nígbà tí kò dájú kò ní hallucinate rára; ìdí tí àwọn model tí a n ló kò fi ẹ̀ bẹ̀ ẹ̀ nì pé scoreboard wa nì fi iṣẹ̀ jẹ̀ ìkọ̀ láti dáhùn.

Ohun tí àwọn ònkòwé ẹ̀

Ìwé náà ní apá mètá. Àkókó, àríyànjìyàn mathematiki pé hallucination diẹ̀ jẹ̀ dandan ní statistiki nígbà pretraining, nípa fífi hàn pé “pèsè ọ̀rọ̀ tó fẹ̀sẹ̀ múlẹ̀ nìkan” ó kéré tán le tó ìṣòro classification alakomeji “sẹ̀ gbólóhùn yí fẹ̀sẹ̀ múlẹ̀?” Èlẹ̀kẹ̀jì, àríyànjìyàn — tí iwádìí benchmark olókíkí mewaà ẹ̀ àtìlẹ̀yìn fún — pé àwọn metric accuracy tó wọ̀pọ̀ nì san èrè fún ìmẹ̀fò ju ìkọ̀ láti dáhùn lọ. Èlẹ̀kẹ̀ta, àtúnṣe tí wọn dábàà àtì case study tó dán an wò: ìdánwò **open-rubric**, níbi tí a ti sọ ọ̀nà fifúnni ní máàkì nínú ibeèrè fúnra rẹ̀ (fún àpẹ̀rẹ̀, “èsì tó tọ̀ gba 1, èyí tó ẹ̀sàṣe gba -1, nítorí náà má dáhùn bí ìdánílójú rẹ̀ kò bá tó 50%”), kí model lè mọ̀ ìgbà tí a nì san èrè fún ọ̀títọ̀. Wọn dán èyí wò lórí frontier model méréin — Gemini 3 Pro ti Google, GPT-5 ti OpenAI, Grok 4 ti xAI àtì Claude Opus 4.5 ti Anthropic — ní lílo àwọn ibeèrè ọ̀títọ̀ 4,326 ti SimpleQA. Wọn sọ ní kedere pé case study náà jẹ̀ àpẹ̀rẹ̀, “kì í ẹ̀ ìdánwò tó ní ìṣàkóso láàárín model” (ètò default, kò sí tuning, kò sí ìṣòkan owó).

Ohun tí wọn rí

- **Pretraining nì fipá mú àṣìṣe diẹ̀.** Iye tí model fi nì sọ èké pẹ̀lú ìdánílójú ní ààlà ìsàlẹ̀ tó jẹ̀ (ní àfojúsùn, ilọpo méjì) iye àṣìṣe classifier “sẹ̀ gbólóhùn yí fẹ̀sẹ̀ múlẹ̀?” tó dára jù lọ tí a kọ̀ láti inú rẹ̀. Fún àwọn ọ̀títọ̀ tí kò ní àpẹ̀rẹ̀ tí a lè kọ̀, ààlà ìsàlẹ̀ náà ó kéré tán jẹ̀ *ie singleton* — ipín àwọn ọ̀títọ̀ tí ó fara hàn lẹ̀kàn ṣoṣo nínú ìkòni. Hallucination diẹ̀ kò ẹ̀e yẹ kòdà pẹ̀lú dátà tó mò pátápátá.
- **Fífúnni ní máàkì nì san èrè fún ìmẹ̀fò — ní kedere.** Lábẹ̀ fífúnni ní máàkì tó nì sọ pé ó tọ̀ tàbí ó ẹ̀sàṣe, àìkọ̀ láti dáhùn láé nì strategy tó dára jù, iwádìí àwọn ònkòwé sì rí i pé ọ̀pọ̀ jù lọ benchmark olókíkí nì ka “Mi ọ̀ mọ̀” sí àṣìṣe lasán. Àpẹ̀rẹ̀ tó hàn gbangba láti ọ̀dọ̀ wọn: lórí ìdánwò SimpleQA, raw accuracy diẹ̀ ẹ̀ *ojú rere* sí o4-mini ti OpenAI — tó dáhùn fẹ̀rẹ̀ ohun gbogbo tí ó sì ẹ̀sàṣe ju ìdáméta-méréin lọ — ju GPT-5-mini, tó ẹ̀ ẹ̀sàṣe diẹ̀ púpọ̀ nítorí pé ó kọ̀ láti dáhùn nígbà àní-àní. Model tó nì ẹ̀ ewu jù ló dà bí ẹ̀ni tó dára jù lórí scoreboard.
- **Open rubric yí iwúrí padà (nínú case study wọn).** Wọn dán ọ̀nà ríṣẹ̀rùn láti dín hallucination kù wò (jẹ̀ kí model dáhùn lẹ̀mẹ̀jì, kí ó sì kọ̀ láti dáhùn bí èsì méjèjè kò bá mu). Lábẹ̀ accuracy tó wọ̀pọ̀, ọ̀nà náà dín àṣìṣe kù *sùgbón ó tún dín accuracy kù* — nítorí náà metric náà nì dẹkun lílò rẹ̀. Lábẹ̀ open rubric, ọ̀nà kan náà borí fún gbogbo model méréin lábẹ̀ oríṣíríṣi iṣẹ̀; GPT-5-mini — tí raw accuracy fi iṣẹ̀ jẹ̀ nítorí ìkọ̀ láti dáhùn nígbà àní-àní — sì borí o4-mini nígbà tí a sọ ọ̀nà fifúnni ní máàkì ní gbangba (n = 4,326 ibeèrè fún model kòkọ̀n).

Ohun tí èyí ẹ́ẹ́ ẹ́ kí ó tùmò sí

Dídín hallucination kù kì í ẹ́ ní pàtàkì nípa dídá ìdánwò mǐ tó jẹ àkànṣe sí hallucination sílẹ̀. Ó jẹ nípa yíyí bí àwọn benchmark pàtàkì ẹ́ ní fún àní-àní ní máàkì, kí jíjẹwọ̀ “Mi ò mò” má bàa tún gba ìjyà. Títí scoreboard yóò fi yí padà, dídín hallucination kù yóò máa ná àwọn model ní máàkì accuracy, a ó sì máa dẹkun rẹ́ — ìdí nìyẹn tí àwọn ònṣẹ̀wé fi pe ìṣòro náà ní “socio-technical”: apá kan jẹ metric tó dára sí i, apá kejì sì jẹ mímú àwọn leaderboard tó ní ipa gba á.

Ohun tí èyí kò fi hàn

- Kò fi hàn pé **open rubric** n tún hallucination ẹ́ ní ayé gidi. Ìdánwò tó ẹ́ àtìlẹ̀yìn fún un jẹ case study kékeré tí a mọ́mọ́ ẹ́ **láìsí ìṣàkóso** — model méréń ní ètò default, ọ̀nà ìdínkù kan, ìdánwò factual-QA kan — láti fi *iyípadà iwúrí* hàn, kì í ẹ́ láti to model ní ipò tàbí fi ipa gbogboogbo múlẹ̀.
- Kò sọ pé **fífúnni ní máàkì ní ìdí kan ṣoṣo**. Àṣìṣe nínú dátà ìkọ̀nì, ìṣòro tó nira gidi, àtí prompt tí model kò mò ẹ́ jẹ orísun mǐràn.
- Kò ẹ́ àtìlẹ̀yìn fún ọ̀rọ̀ olókíkí pé hallucination kò ẹ́ẹ́ yẹ. Àwọn ònṣẹ̀wé n jiyàn ní ọ̀díkẹjì: ètò kan tó dáhùn ìbéèrè tó ẹ́ẹ́ ṣàyèwò nikan, tó sì sọ “Mi ò mò” nígbà mǐràn, kò ní hallucinate láé.
- Kò mú ààlà ìsàlẹ̀ pretraining tán — ó ṣàlàyé rẹ́, ó sì fi ààlà sí i; ààlà náà kan àṣìṣe ọ̀títọ̀ tí a sọ pẹ̀lú ìdánílójú, kì í ẹ́ gbogbo ìhùwàsí model.
- Kò fi hàn pé **open rubric** tó fúnra rẹ́. Ó n yí ohun tí ìdánwò n san èrè fún padà; kì í ẹ́ arọ̀pò retrieval, lílo tool, tàbí model tí calibration rẹ́ dára sí i.

Báwo ni ẹ́rì náà ẹ́ lágbara tó?

- Ìpílẹ̀ náà jẹ **mathematiki** — ààlà ìsàlẹ̀ àṣe, kì í ẹ́ ọ̀ṣùwọ̀n. Gégé bí àrìyànjiyàn theory, ó fẹ́ṣẹ̀ múlẹ̀ lábé àwọn àfojúsùn tirẹ̀.
- Ó dúró lórí **àwòṣe ìṣòro tí a mọ́mọ́ mú rọ̀rùn**; àwọn ònṣẹ̀wé fúnra wọ̀n tọ́ka sí “false trichotomy” tí pípa gbogbo èsì sí tó tó, tó ṣàṣìṣe, tàbí “Mi ò mò”, àtí ipò “òótọ̀ tí kò ní àpẹ̀rẹ́” tí wọ̀n lò fún ààlà tó mò jù.
- Ìwádíí benchmark jẹ **àpẹ̀rẹ́ kékeré tí a yàn** — ìdánwò olókíkí mẹ́wàá, kì í ẹ́ audit gbogbo rẹ́.
- Case study náà jẹ **gidi sùgbón ó ní ààlà**: frontier model méréń, ọ̀nà ìdínkù kan, SimpleQA nikan, ètò default, tí a sì sọ ní kedere pé “kì í ẹ́ ìdánwò tó ní ìṣàkóso.” Ó jẹ proof of concept fún àrìyànjiyàn iwúrí, kì í ẹ́ èsì benchmark.
- Ó yẹ ká sọ ibi tí iwọ́ye náà ti wá: mẹ́ta nínú ònṣẹ̀wé méréń n ṣìṣe tàbí ti ṣìṣe ní **OpenAI**, iwé náà sì n jiyàn pé ẹ́ka náà yẹ kí ó yí ọ̀nà tó fi n dán model wò padà. Èyí jẹ ìdúró tí a ṣàlàyé dáadàa láti ọ̀dọ̀ ẹ̀nì tó ní ipa nínú ọ̀rọ̀ náà, kì í ẹ́ àyèwò ọ̀de aláìṣojúsàájú — ohun tó yẹ ká

wọn, kì í ẹ̀ ẹ̀ ká kò. (Ohun tó dára ní pé iwé náà fi àríwísí kan àwọn model tiwọn, o4-mini àti GPT-5-mini, gégé bí ó ẹ̀ kan àwọn mǐ.)

Ìdí tí ó fi ẹ̀ pàtàkì

Ó tún ọ̀nà tí a fi ń wo ìṣòro tí a ti polówó rẹ̀ gan-an ẹ̀. Wọn sáà mǎa ń ta “hallucination” bí àbùkù àjẹ̀jẹ̀ tàbí odi tí kò ẹ̀ ẹ̀ kọ́jǎ; iwé yí mǔ un jẹ̀ ohun lasán tí a sì dá sílẹ̀ ní apá kan — ààlà ìsàlẹ̀ statistiki tí a lẹ̀ lóye, tó jókòó lórí iwúrí tí àwa fúnra wa yàn. Ẹ̀kọ́ tó gbòòrò jẹ̀ díẹ̀díẹ̀, ó sì wúlò jù: ìlọ́síwájú sí i nínú reliability lẹ̀ sinmi lórí *ohun tí a ń díwọn* gégé bí ó ti sinmi lórí ohun tí a ń kọ́.

Àkótán mímọ́

Àwọn èsì èké tí language model ń sọ pẹ̀lú ìdánílójú wá láti ibi méjì. Àkọ́kọ́ jẹ̀ statistiki: nígbà tí òtítọ́ kan kò ní àpẹ̀ẹ̀ẹ̀ láti kọ́, model tí a kọ́ láti fara wé èdè yóò ẹ̀ ẹ̀ ẹ̀ nígbà mǐràn, a sì lẹ̀ fojú díwọn ààlà ìsàlẹ̀ yẹ̀n (fún àpẹ̀ẹ̀ẹ̀, láti iye òtítọ́ tó fara hàn lẹ̀ẹ̀kan ẹ̀ ẹ̀ ẹ̀ nínú ìkọ̀nì). Ẹ̀lẹ̀ẹ̀kejì ni iwúrí: fẹ̀rẹ̀ẹ̀ gbogbo benchmark tí a fi ń to model ní ipò ń fún “Mi ò mọ́” ní mǎàkì kan náà pẹ̀lú èsì ẹ̀ ẹ̀ ẹ̀, nítorí náà ìmẹ̀fò mǎa ń borí — tí tí dé ibi tí model tó ẹ̀ ẹ̀ ẹ̀ ní ìdámẹ̀ta-mẹ̀rin àkókò lẹ̀ ju èyí tó jẹ̀ olóòtítọ́ sí i tó sì kọ́ láti dáhùn lọ. Àbá àwọn ònṁkọ́wé kì í ẹ̀ ìdánwò hallucination mǐ, bí kò ẹ̀ “open rubric”: sọ ọ̀nà fifúnni ní mǎàkì nínú ìbẹ̀èrè. Nínú case study lórí frontier model mẹ̀rin, èyí yí iwúrí padà kí ọ̀nà tó ń dín hallucination kù gba èrè dípò ìjìyà. Ó jẹ̀ iwé theory-pẹ̀lú-iwádíí, pẹ̀lú ìdánwò kékeré tí a sọ pé kò ní ìṣàkóso, tí a ẹ̀ peer review rẹ̀ nínú *Nature*; àtúnṣe náà ní ìlérí sùgbọ̀n a kò tǐ fi hàn pé ó ń ẹ̀ ẹ̀ ní iwọn ńlá, àwọn ònṁkọ́wé sì jìyàn pé hallucination kì í ẹ̀ àdìitú tàbí ohun tí kò ẹ̀ ẹ̀ yẹ̀ pátápátá.

Àyẹ̀wò láìsí ọ̀rò asán

Ohun tí iwé náà fi hàn: Ààlà ìsàlẹ̀ mathematiki tí ó mǔ kí hallucination díẹ̀ jẹ̀ dandan nígbà pretraining (ó kéré tán “iye singleton” fún òtítọ́ tí kò ní àpẹ̀ẹ̀ẹ̀); iwádíí tó rí i pé ọ̀pọ́ benchmark ẹ̀ ẹ̀ ẹ̀ kò fún “Mi ò mọ́” ní mǎàkì; àti case study model mẹ̀rin níbi tí síṣọ́ ọ̀nà fifúnni ní mǎàkì nínú prompt (“open rubric”) tí mǔ kí ọ̀nà dídín hallucination kù borí, níbi tí accuracy lasán ti fi ìjìyà jẹ̀ ẹ̀.

Ohun tó ẹ̀ ẹ̀ gbà sùgbọ̀n tí a kò fidí rẹ̀ mǐlẹ̀: Pé open rubric, bí a bá fi sínú àwọn benchmark pàtàkì, yóò dín hallucination nínú model tí a fi sílò kù ní ọ̀nà tó ní ìtumọ́. Ìdánwò tó ẹ̀ àtílẹ̀yìn fún un kéré, a sì sọ ní kedere pé kò ní ìṣàkóso.

Ohun tí kò fi hàn: Pé hallucination kò ẹ̀ ẹ̀ yẹ̀ (ó jìyàn ní òdìkejì); pé fifúnni ní mǎàkì ni ìdí kan ẹ̀ ẹ̀; pé a lẹ̀ pa hallucination rẹ̀ ráúrú; tàbí pé case study náà ń to model mẹ̀rin sí ipò.

Àwọn ààlà pàtàkì: Àwòṣe tó mòṣmò mú ìṣòro rọrùn sí tó tọ/tó ṣàṣiṣe/“Mi ò mò” (àwọn òhàkòwé pè é ní “false trichotomy”); ìwádìí benchmark kékeré tí a yàn (ìdánwò mọ̀wàá); case study tí kò ní ìṣàkóso (model méréin, ọ̀nà ìdínkù kan, ìdánwò kan, ètò default); àti àríyànjìyàn tí OpenAI darí nípa bí ẹ̀ka náà ṣe yẹ kí ó máa dǎn model wò.

Ìgbẹ̀kẹ̀lé mélòó ni olùkà gbogboogbo yẹ kí ó ní? Ìgbẹ̀kẹ̀lé gíga pé hallucination kì í ṣe àdìtú tàbí ohun tí kò ṣeé yẹ pátápátá, àti pé benchmark pàtàkì n san èrè fún ìméfò ní bá yí. Ìgbẹ̀kẹ̀lé àárín pé àtúnṣe tí a dábàá yóò ran lówó — proof of concept gidi wà bá yí, ṣùgbón kò tū sí ìfihàn pé ó n ṣiṣe káàkiri àti ní ìwọ̀n nlá.

Àwọn orísun

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>