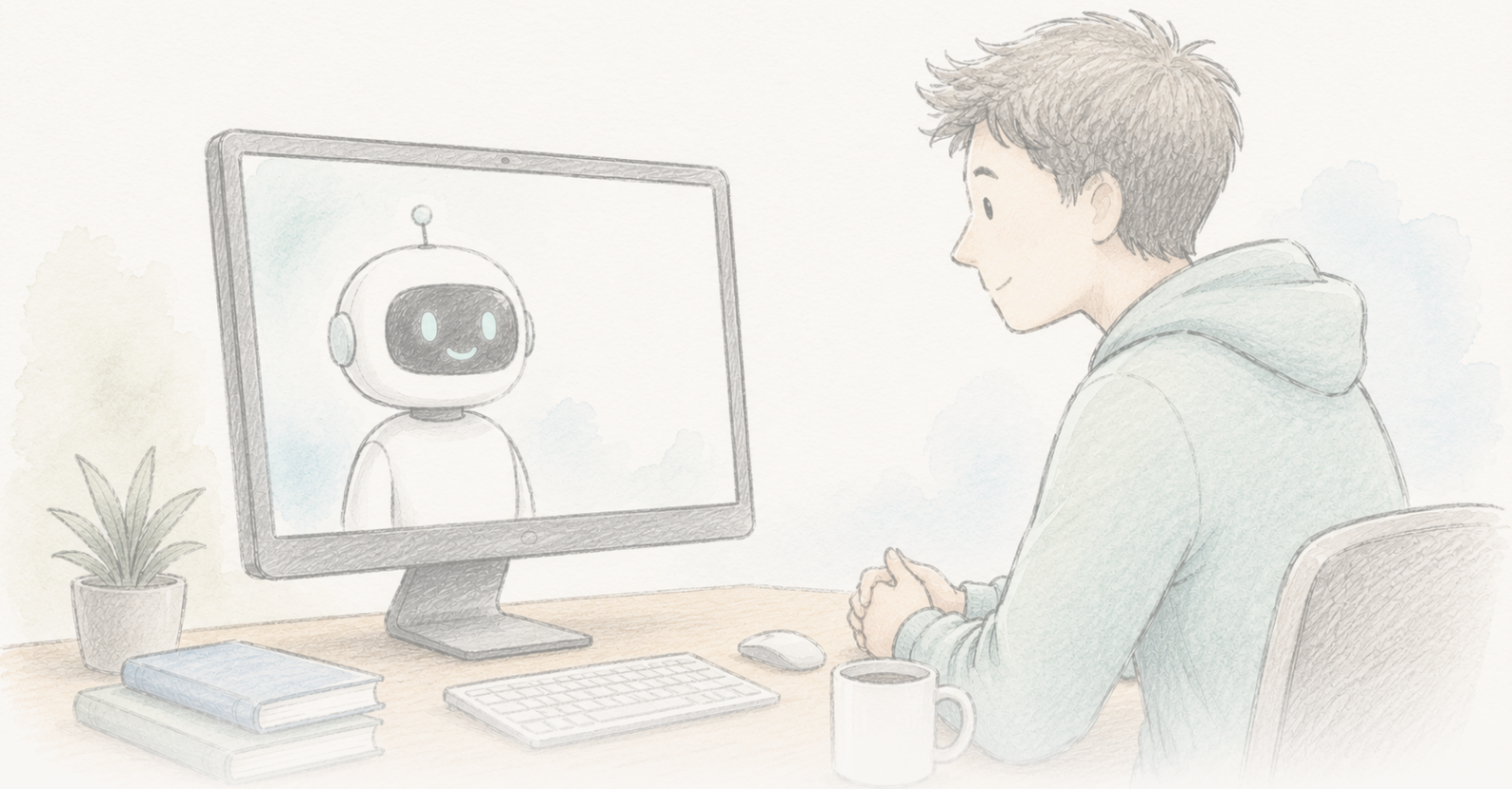




The CLEAN PAPER

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

Činilo se da je chatbot mjesecima smanjivao vjerovanje u teorije zavjere

Studija iz 2024. pokazala je da je razgovor s GPT-4 Turbo u tri kruga smanjio vjerovanje ispitanika u teoriju zavjere koju su zastupali za oko 20%, da je učinak trajao dva mjeseca, obuhvatio različite vrste zavjera i počivao na tvrdnjama umjetne inteligencije ocijenjenim kao gotovo u potpunosti tačne. U junu 2026. uz rad je objavljeno uredničko upozorenje zbog problema s obradom podataka i ponovljivošću; autori navode da korigirana analiza zadržava isti smjer, statističku značajnost i veličinu nalaza, dok Science još provodi procjenu. Zaista zanimljiv i pažljivo izveden rezultat — trenutno pod provjerom, ni konačno potvrđen ni opovrgnut.

Written by Lucio Vaglio · Cross-checked by Cinzia Vaglio and Vera Rubin · Edited by Michele Renda · Figures by Nova Vela · AI-assisted, human-reviewed

Paper: *Durably reducing conspiracy beliefs through dialogues with AI*, Thomas H. Costello, Gordon Pennycook, David G. Rand, Science 385, eadq1814 (2024) · DOI: 10.1126/science.adq1814

Editorial Expression of Concern

Science je uz ovaj rad objavio uredničko upozorenje dok procjenjuje pitanja obrade podataka i ponovljivosti. To nije povlačenje rada niti utvrđivanje nedoličnog postupanja; znači da prijavljeni rezultat treba smatrati privremenim dok se provjera ne okonča.

Editorial Expression of Concern →

Činjenice ipak mogu djelovati — ali rad sada nosi upozorenje

Studija iz 2024. izvijestila je o rezultatu koji ide protiv jednog ugodnog dijela konvencionalne mudrosti. Dajte osobi kratak razgovor s AI-jem u tri kruga — GPT-4 Turbo, upućen da odgovara baš na dokaze koje ta osoba navodi za teoriju zavjere u koju osobno vjeruje — i njeno vjerovanje u tu teoriju u prosjeku padne za oko 20%. Pad je bio prisutan i dva mjeseca poslije. Učinak se pojavio kod klasičnih teorija zavjere (ubojstvo Johna F. Kennedyja, izvanzemaljci, Iluminati) i aktualnih (COVID-19, američki izbori 2020.), čak i među ljudima čije je uvjerenje bilo snažno i povezano s identitetom. Kada je profesionalni fact-checker ocijenio 128 činjeničnih tvrdnji koje je AI iznio, 127 ih je bilo tačno, jedna obmanjujuća, a nijedna netačna.

Naslov koji je obišao svijet glasio je da ljude *možete* izvući iz „zečje rupe” razgovorom — da vjernici u teorije zavjere nisu izvan dosega dokaza, nego trebaju prave dokaze, dovoljno precizno usmjerene na ono u što vjeruju. To je zaista zanimljiva tvrdnja, a studija je bila pažljivije osmišljena od većine istraživanja uvjeravanja.

Zatim je u junu 2026. *Science* uz rad objavio **Editorial Expression of Concern**. To je dio koji će mnoga prepričavanja preskočiti, a upravo on određuje koliku težinu rezultat trenutno može nositi.

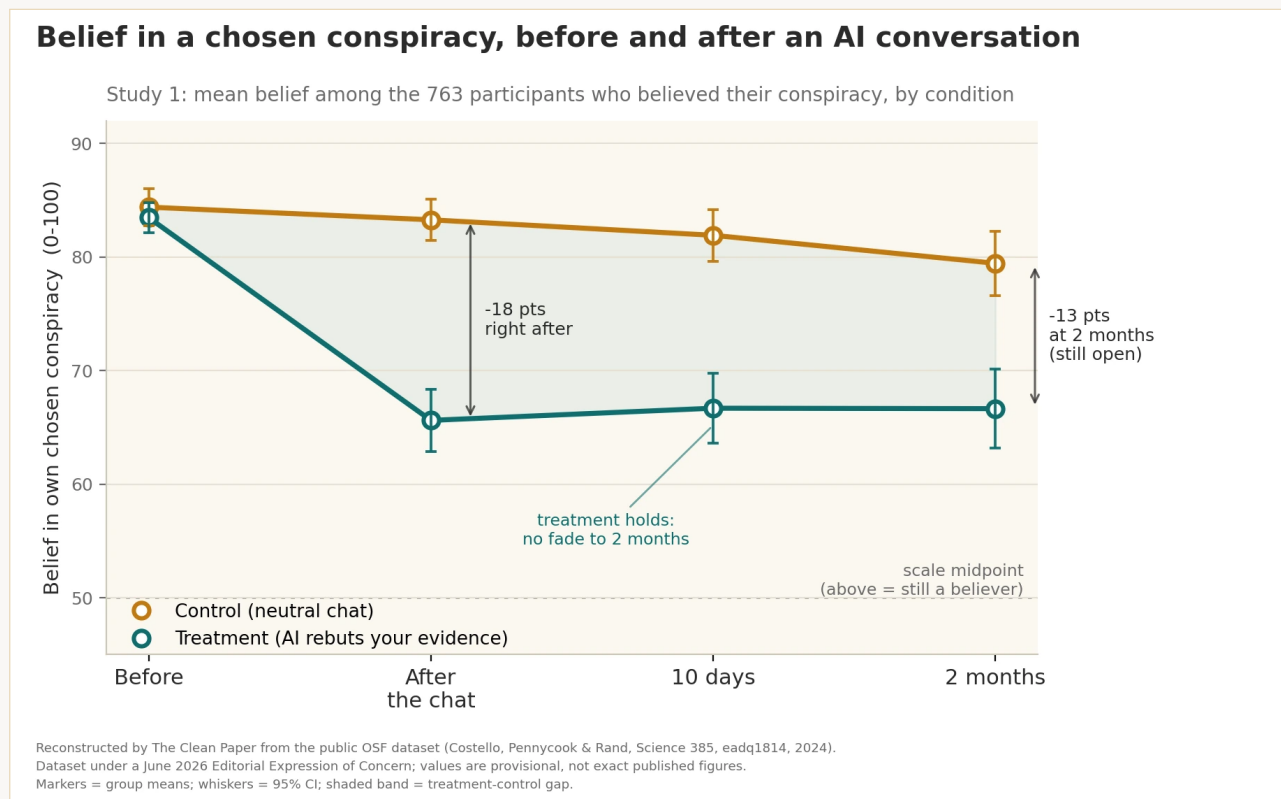
Šta je Expression of Concern — a što nije

Editorial Expression of Concern formalno je upozorenje koje časopis pridružuje radu kako bi čitaocima dao do znanja da su otvorena pitanja koja se još razjašnjavaju. To **nije** povlačenje rada (rad nije uklonjen) i samo po sebi nije nalaz naučnog nepoštenja. Znači: čitajte uz ovu ogradu jer se naučni zapis možda promijeniti. Ovdje su autori, nakon što su upozoreni na probleme, proveli istragu, prijavili pojedinosti časopisu *Science* i dostavili ispravljenu analizu.

Ono što je označeno vrlo je konkretno. Autori su upozoreni na nedosljednosti u tome kako su kriteriji za uključivanje učesnika primijenjeni između rukopisa i objavljenog koda za analizu te na činjenicu da je javni skup podataka sadržao suvišne retke umetnute greškom pri spajanju koda — probleme zbog kojih je dio prijavljenih brojeva bilo teško reproducirati iz objavljenih materijala. Autori su *Scienceu* dostavili ispravljenu analitički postupak i skup ažuriranih rezultata za koje tvrde da se s izvornima podudaraju **u smjeru, statističkoj značajnosti i praktičnoj veličini učinka**. *Science* ih još procjenjuje. Dok ta procjena ne završi, nad tačnim brojkama stoji upitnik koji može ukloniti samo

časopis.

Zato su ovo dvije priče istovremeno: intrigantan rezultat o tome mogu li činjenice pomaknuti uko-rijenjena uvjerenja i živi primjer naučnog zapisa koji se javno ispravlja. Poenta ovog teksta jest ne pomiješati ih.



U studiji je razgovor s AI-jem u tri kruga, koji je pobijao dokaze koje je učesnik sam naveo, prema objavljenim rezultatima smanjio vjerovanje u odabranu teoriju zavjere za oko 20% u prosjeku; za razliku od kontrolnog razgovora o nepovezanoj temi, pad je bio mjerljiv i nakon 10 dana i nakon 2 mjeseca. Učinak je bio trajan, ali djelomičan: većina učesnika i dalje je vjerovala u teoriju — smanjenje, ne „izlječenje“. Rad je trenutno pod Editorial Expression of Concernom (*Science*, juni 2026.) zbog nedosljednosti u izvještavanju i greške u javnom skupu podataka; autori kažu da ispravljena analiza čuva smjer, značajnost i veličinu učinka, dok *Science* nastavlja procjenu. Ovaj grafikon rekonstruirao je The Clean Paper iz tog javnog skupa podataka, pa su prikazane vrijednosti privremene, a ne konačne brojke rada.

Original figure — The Clean Paper CC BY 4.0

Šta su autori uradili

- Proveli su dva eksperimenta s 2190 učesnika; svaki je vlastitim riječima naveo teoriju zavjere u koju vjeruje i dokaze za koje smatra da je podržavaju.
- Svaki je učesnik vodio razgovor s GPT-4 Turbo u tri kruga. U **tretmanskom** uslovu AI je dobio zadatak pobijati konkretne dokaze učesnika; u **kontrolnom** je razgovarao o nepovezanoj temi.
- Mjerili su vjerovanje u odabranu teoriju prije i neposredno nakon razgovora, a zatim ponovno

kontaktirali učesnike nakon **10 dana i 2 mjeseca**.

- Profesionalni fact-checker procijenio je tačnost svih 128 činjeničnih tvrdnji koje je AI iznio u jednom poduzorku.
- Provjerili su prelijeva li se učinak na druge, nepovezane teorije zavjere — i, kao test specifičnosti, smanjuje li i vjerovanje u zavjere koje su se zaista dogodile i stoga su **istinite**.

Šta su pronašli

- **Prosječno smanjenje od približno 20%** u vjerovanju u odabranu teoriju u tretmanskoj grupi u odnosu na kontrolu (studija 1: interval pouzdanosti od 95% [13.8, 19.7] bodova na skali od 100 bodova, $P < 0.001$).
- **Učinak je trajao**. U tretmanskoj grupi pad nije nestao: vjerovanje je i nakon 10 dana i dva mjeseca ostalo blizu razine neposredno nakon razgovora umjesto da se vrati prema početnoj, dok je kontrolna grupa cijelo vrijeme ostala više. Rad opisuje učinak na odabranu teoriju kao nesmanjen najmanje dva mjeseca, a ne kao trenutačan kolebljivi pomak.
- **Učinak se generalizirao** kroz širok raspon teorija koje su ljudi naveli i pojavio se i kod učesnika čije je početno uvjerenje bilo snažno.
- **Tvrdnje AI-ja bile su tačne** u ovom okruženju: 127 od 128 (99.2%) ocijenjeno je tačnima, 1 (0.8%) obmanjujućom, nijedna netačnom.
- **Učinak je bio specifičan**, a ne opći skepticizam: intervencija **nije** smanjila vjerovanje u istinite zavjere, što ukazuje na pomak neutemeljenih uvjerenja, a ne na to da su ljudi postali cinični prema svemu.
- **Došlo je i do skromnog prelijevanja** — vjerovanje u druge, nepovezane teorije palo je za oko 12%, a učesnici su prijavili veću namjeru suprotstaviti se tvrdnjama o zavjerama. Glavni se učinak ponovio u studiji 2 i išao u istom smjeru u manjem uzorku bez kontrolne grupe (slabiji dokaz, ali dosljedan).

Šta ovo ne dokazuje

- Rezultat trenutno **nije** bez otvorenih pitanja. Rad je pod Editorial Expression of Concernom zbog problema u obradi podataka i reproducibilnosti, a ispravljene brojke *Science* još procjenjuje. Konkretno vrijednosti treba smatrati privremenima dok ta procjena ne završi.
- **Ne** pokazuje da AI „liječi” vjerovanje u teorije zavjere. Prosječno smanjenje od ~20% smislen je pomak, ne brisanje uvjerenja — većina učesnika i dalje je vjerovala u teoriju, samo manje.
- Ovo **nije** rezultat stvarne primjene u svijetu. Učesnici su dobrovoljno ušli u studiju i razgovarali s AI-jem u dobroj vjeri; izvan laboratorija ljudi najčvršće vezani uz teoriju zavjere vjerovatno su upravo najmanje skloni potražiti chatbot koji će im proturječiti.
- Tačnost od 99.2% **specifična je za ovaj zadatak, model i pažljivo oblikovane upute** — nije opće jamstvo da jezični modeli govore istinu. Autori izričito navode da bi ista personalizirana moć uvjeravanja mogla gurati i *lažna* uvjerenja ako bi bila usmjerena na to.
- „Trajno” ovdje znači **dva mjeseca**, što je impresivno za jedan razgovor, ali nije isto što i trajno

zauvijek.

Koliko su dokazi jaki

- **Dizajn je prava snaga studije.** Kontrolirani eksperimenti tretman nasuprot kontroli, čista placebo-kontrola, ponavljanje kroz dvije studije i nezavisni uzorak, praćenje trajnosti te eksplicitna provjera tačnosti tvrdnji samog AI-ja — sve je to pažljivije od većine istraživanja uvjeravanja, a smjer učinka dosljedan je gdje god je testiran.
- **Ali Expression of Concern nije fusnota.** Mogućnost da se prijavljene brojke ponovno dobiju iz objavljenih podataka i koda dio je onoga što rezultat čini pouzdanim, a upravo je to zakazalo — zbog nedosljednih kriterija uključivanja i dupliciranih redova iz greške pri spajanju koda. Tvrdnja autora da ispravljeni postupak čuva rezultat ohrabrujuća je i moguća, ali još je to njihov vlastiti izvještaj, ne nezavisna presuda. Treba pričekati procjenu *Sciencea*.
- **Pošten status jest „označeno upozorenjem”, ne „opovrgnuto” ni „potvrđeno”.** Dobro osmišljena, upečatljiva studija čije su tačne brojke pod formalnom provjerom. To je neugodno mjesto za ostaviti dobru priču, ali je tačno.

Zašto je to važno

Godinama je uticajno stajalište tvrdilo da vjernike u teorije zavjere pokreću psihološke potrebe i identitet te da ih se zato činjenicama ne može razuvjeriti. Trajno smanjenje temeljeno na dokazima — ako se održi — važna je protuteža tom pesimizmu: ukazuje da je dio problema možda bio u tome što općenito pobijanje nikada nije zahvaćalo *konkretne* dokaze koje pojedini vjernik zaista navodi, a LLM to može raditi u velikom opsegu.

Važno je i kao studija slučaja kako bi nauka trebala raditi. Pouka Expression of Concerna nije da su proslavljeni rezultati bezvrijedni; pouka je da greške izlaze na vidjelo, podaci se ponovno ispituju, a tvrdnje javno provjeravaju. To što taj mehanizam radi karakteristika je sistema, ne skandal — i razlog zašto je pravi stav prema ovom rezultatu strpljenje, a ne ni pljesak ni odbacivanje.

I priča o AI-ju ide u oba smjera. Ista sposobnost da prilagođenim argumentom oslabi lažno uvjerenje može ga jednako efikasno i ojačati. Tehnika je neutralna; cilj nije.

Kratak sažetak

Studija iz 2024. pokazala je da je razgovor s GPT-4 Turbo u tri kruga smanjio vjerovanje ljudi u teoriju zavjere koju su sami zastupali za oko 20%, uz učinak koji je trajao dva mjeseca i pojavljivao se kroz različite vrste teorija, dok su činjenične tvrdnje AI-ja ocijenjene gotovo potpuno tačnima. U junu 2026. rad je stavljen pod Editorial Expression of Concern zbog problema u obradi podataka i reproducibil-

nosti; autori izvještavaju da ispravljena analiza čuva smjer, značajnost i veličinu nalaza, a *Science* još procjenjuje. Treba ga čitati kao zaista zanimljiv i pažljivo dizajniran rezultat koji je trenutno pod provjerom — nije zaključen, ali nije ni opovrgnut. Najpoštenija rečenica o njemu ona je koju sam proces upravo piše: provjerite ponovno.

Izvori

Costello, Pennycook & Rand, Durably reducing conspiracy beliefs through dialogues with AI, *Science* 385, eadq1814 (2024)

<https://doi.org/10.1126/science.adq1814>

Editorial Expression of Concern, *Science* (posted 11 June 2026; H. Holden Thorp, Editor-in-Chief), DOI 10.1126/science.aej2383

<https://www.science.org/doi/10.1126/science.aej2383>