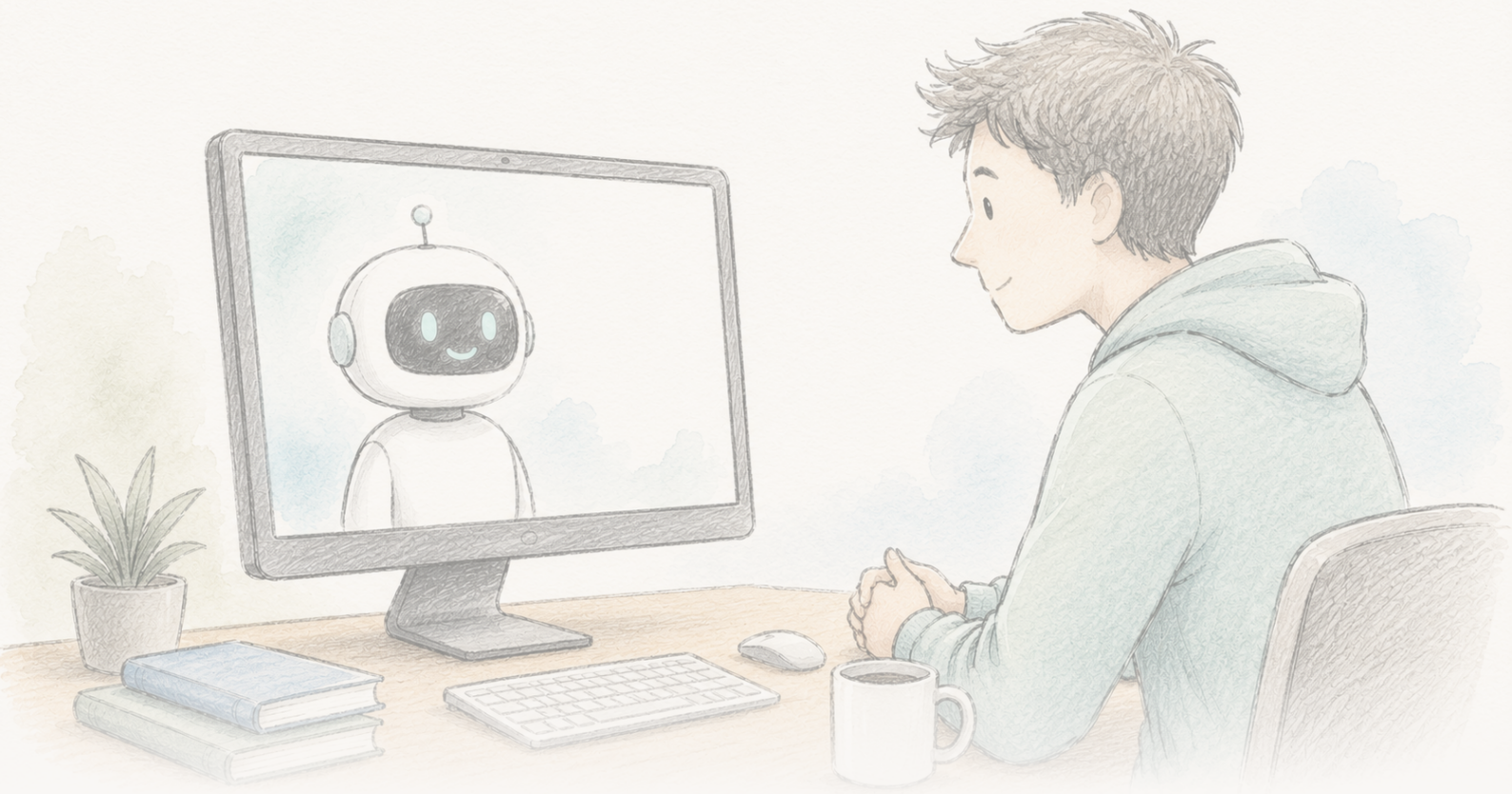




The CLEAN PAPER

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

Chatbot-এর সঙ্গে কথোপকথন conspiracy belief কয়কে মাস কমিয়ে দিচ্ছেলি — কিন্তু paper-টির ওপর Editorial Expression of Concern আছে

Personalized generative-AI dialogue experiment conspiracy belief immediateভাবে কমাৰ এবং কয়কে মাস follow-up-এ কিছু effect থাকার report দয়িছেলি। Result suggest করে tailored evidence কিছু strong believer-কওে move করতে পারে। কিন্তু paper বর্তমানে Editorial Expression of Concern-এর অধীনে; এটি retraction নয়, তবে unresolved methodological/transparency concern-এর signal। তাই clean reading: interesting reported effect, replication/clarification pending — settled fact নয়।

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 has attached an Editorial Expression of Concern to this paper while it evaluates data-handling and reproducibility questions. This is not a retraction and not a finding of misconduct; it means the reported result should be read as provisional until the review is resolved.

Editorial Expression of Concern →

তথ্য, শেষে পর্যন্ত—আর গবেষণাপত্রের ওপর একটি আনুষ্ঠানিক চহ্নিত করা

2024 সালের একটি গবেষণা প্রচলিত ধারণার বিপরীত একটি আকর্ষণীয় ফল জানিয়েছিল। কটে যে ষড়যন্ত্রের তত্ত্ব সত্যি বিশ্বাস করে এবং নজিরে ভাষায় যে প্রমাণ দিয়ে, GPT-4 Turbo-কে যদি সেই নির্দিষ্ট প্রমাণের জবাব দিয়ে তনি দফা আলাপ করতে দেওয়া হয়, গড়ে বিশ্বাস প্রায় 20% কমে। এই হ্রাস দুই মাস পরেও টকি ছিল। JFK হত্যাকাণ্ড, aliens, Illuminati-এর মতো ষড়যন্ত্র থেকে COVID-19 ও 2020 US election নিয়ে সমসাময়িক বিশ্বাস—বহুনির্ভরশ্রুতিতেই প্রভাব দেখা যায়; শক্তিশালী বা পরিচয়-সংশ্লিষ্ট বিশ্বাসেও। পশোদার fact-checker-রা AI-এর নমুনার 128টি তথ্যভিত্তিক দাবি পরীক্ষা করে 127টিকে সত্য, 1টিকে বিভিন্নান্তকির এবং 0টিকে ভুল বলছেন।

Travel করা শরিনোম ছিল: rabbit গহ্বর থেকে তথ্য দিয়ে মানুষকে ফরিয়ে আনা যায়; ষড়যন্ত্র বিশ্বাসী প্রমাণের বাইরে নয়, বরং সাধারণ খণ্ডনের বদলে তাদের নজিস্ব cited প্রমাণ নির্দিষ্টভাবে উত্তর করতে হবে। দাবিসত্যি আকর্ষণীয়, এবং পরে চনা গবেষণার তুলনায় নকশা সতরক।

তারপর জুন 2026-এ Science গবেষণাপত্রের ওপর Editorial Expression of Concern দিয়ে। এই অংশ বাদ দিলে বর্তমান প্রমাণ অবস্থা ভুল বলা হয়।

Expression of Concern কী—আর কী নয়

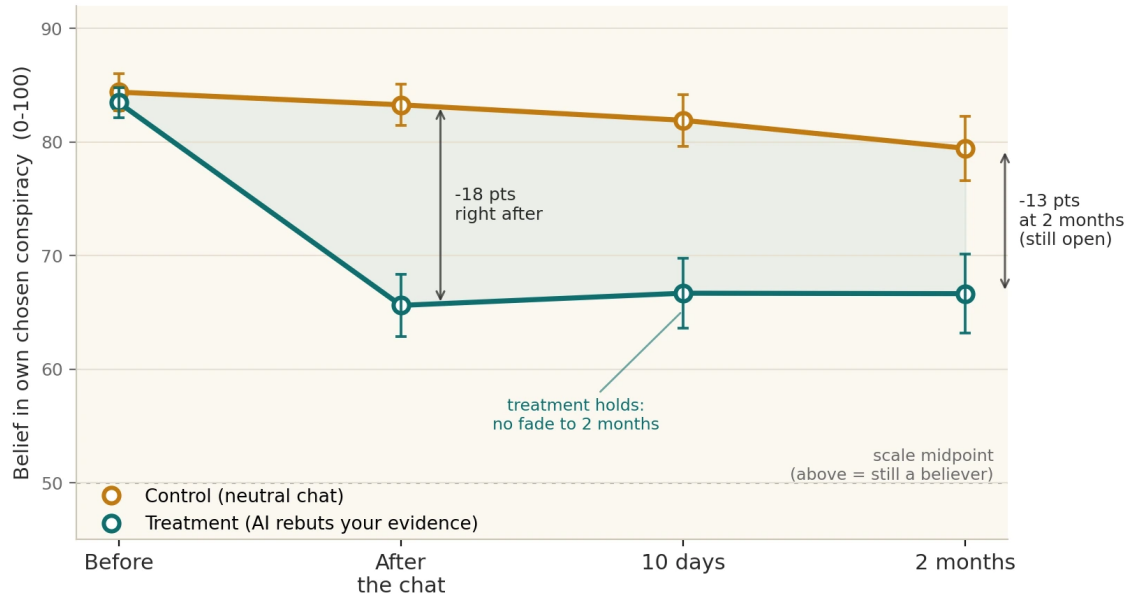
সম্পাদকীয় Expression of Concern সাময়িকীর আনুষ্ঠানিক সতরকবারতা: গবেষণাপত্র নিয়ে প্রশ্ন উঠছে এবং মূল্যায়ন চলছে। এটি প্রত্যাহার নয়—গবেষণাপত্র প্রত্যাহত নয়; অসদাচরণ সিদ্ধান্ত-ও নজির থেকে নয়। অর্থ: সতরকতা মাথায় রেখে পড়ুন, নথি বদলাতে পারে। এখানে সমস্যা জানার পরে লেখকেরা তদন্ত করে Science-কে নির্দিষ্ট সমস্যা ও সংশোধিত বিশ্লষণ দিয়েছেন।

সতরকবারতার কারণটি নির্দিষ্ট। অংশগ্রহণকারী বাছাইয়ের মানদণ্ড পাণ্ডুলিপি ও প্রকাশিত বিশ্লষণ কটে একভাবে প্রয়োগ করা হয়নি; জনসাধারণের ডটোসটে কটে ত্রুটির কারণে অতিরিক্ত ও নকল সারটুক পড়েছিল, ফলে প্রকাশিত উপকরণ থেকে কিছু সংখ্যা পুনরুৎপাদন করা কঠিন হয়। লেখকেরা সংশোধিত pipeline ও হালনাগাদ ফল Science-এ জমা দিয়েছেন; তাদের মতে প্রভাবের দকি, পরিসংখ্যানগত তাৎপর্য ও বাস্তব মাত্রা মূল ফলের মতোই থাকে। সাময়িকী এখনো মূল্যায়ন করছে। চূড়ান্ত সিদ্ধান্ত না হওয়া পর্যন্ত নির্দিষ্ট সংখ্যাগুলোকে অস্থায়ী ধরে নেওয়াই ঠিক।

এটিই দুই গল্প: তথ্য গভীরভাবে গুঁথে থাকা বিশ্বাস বদলাতে পারে কনি—একটি সামাজিক-science ফল; এবং বজ্ঞেগনিক নথি ত্রুটি পিলে জনসাধারণেরভাবে কীভাবে সংশোধন করা হয়—একটি চলমান উদাহরণ।

Belief in a chosen conspiracy, before and after an AI conversation

Study 1: mean belief among the 763 participants who believed their conspiracy, by condition



Reconstructed by The Clean Paper from the public OSF dataset (Costello, Pennycook & Rand, Science 385, eadq1814, 2024). Dataset under a June 2026 Editorial Expression of Concern; values are provisional, not exact published figures. Markers = group means; whiskers = 95% CI; shaded band = treatment-control gap.

প্রতিবেদিত গবেষণায় ব্যক্তিকেন্দ্রিক তথ্য দফার AI খণ্ডন নরিবাচতি ষড়যন্ত্র বশি়াস গড় ~20% কমিয়েছে; অসংশ্লিষ্ট ন্যিন্তরণ আলাপ তা করেনি। পার্থক্য 10 দিন ও 2 মাসেও পরমাপ্যে গ্য ছিল, কনিত্রু আংশকি—অধিকাংশ অংশগ্রহণকারী এখনও ষড়যন্ত্র কছুটা বশি়াস করছে। জুন 2026 থেকে গবেষণাপত্র *Science* সম্পাদকীয় Expression of Concern-এর অধিনে; সংশ্লিষ্ট বশি়াষণ সাময়িকী মূল্যায়ন করছে, তাই চার্টে মানগুলো অস্থায়ী।

Original figure — The Clean Paper CC BY 4.0

গবেষকরা কী করছেন

- দুই পরীক্ষায় মোট **2,190 অংশগ্রহণকারী** নজিদেরে ভাষায় একটি believed ষড়যন্ত্র তত্ত্ব এবং সমর্থনকারী প্রমাণ লেখেন।
- প্রত্যেকে GPT-4 Turbo-র সঙ্গে **তথ্য দফার আলাপ** করেন। চকিহিসায় AI অংশগ্রহণকারীর নরিদ্ষিট প্রমাণ rebut করে; ন্যিন্তরণে অসংশ্লিষ্ট বিষয় আলাপ চনা করে।
- নরিবাচতি ষড়যন্ত্র বশি়াস আলাপের আগ/পরে, তারপর **10 দিন ও 2 মাস** অনুসরণকালীন পরমাপ করা হয়।
- নমুনার AI আউটপুটে সব **128 তথ্যভিত্তিক দাবি** পশোদার তথ্য-যাচাইকারী মূল্যায়ন করেন।
- প্রভাব অসংশ্লিষ্ট ষড়যন্ত্র-তে spill over করে কিনা এবং **সত্যহিওয়া ষড়যন্ত্র**-তেও সার্বকি skepticism তরৈকিরে কিনা নরিদ্ষিট পরীক্ষা করা হয়।

তারা কী পয়েছেন

- চকিহিসা বনাম ন্যিন্তরণে নরিবাচতি ষড়যন্ত্র বশি়াস গড় প্রায় **20% কম**। গবেষণা 1-এ 100-বিন্দু পরসিরে প্রভাব 95% CI [13.8, 19.7] বিন্দুগুলো, $P < 0.001$ ।

- **প্রভাব টকি ছিল।** চকিৎসা দলরে হ্রাসপ্রাপ্ত বিশ্বাস 10 দিন ও 2 মাস অনুসরণকালীন just-after-আলাপ স্তররে কাছাকাছি ছিল; ন্যূনতম উচ্চতর। লেখকরো “অন্তত দুই মাস undiminished” ভাষা ব্যবহার করনে।
- ভিন্ন ষড়যন্ত্র ধরন ও শক্তিশালী প্রাথমিক বিশ্বাসীতেও প্রভাব দিক থাকে।
- AI দাবিনির্ভুলতা নমুনা 127/128 (99.2%) সত্য, 1/128 (0.8%) বহির্ভূত, 0 ভুল।
- হস্তক্ষেপে সত্য ষড়যন্ত্র-র বিশ্বাস কমানি—সার্বিক cynicism নয়, অসমর্থিত বিশ্বাস লক্ষ্য করার নরিদৃষ্টিতা সমর্থন করে।
- অসংশ্লিষ্ট ষড়যন্ত্র বিশ্বাস সামান্য ~12% কমে; অংশগ্রহণকারী ষড়যন্ত্র দাবি push পছনে করার ইচ্ছা বাড়ায়। গবেষণা 2 প্রধান প্রভাব প্রতিলিপি তৈরি করে; না-ন্যূনতম ছেটিতর নমুনা একই দিক দলিও প্রমাণ দূর্বলতর।

এটিকী প্রমাণ করে না

- গবেষণাপত্র এখন **unquestioned প্রমাণ নয়।** তথ্য-handling/পুনরুৎপাদনযোগ্যতা সমস্যা আনুষ্ঠানিক সাময়িকী প্রয়ালে চিনায়; সংশোধিত সূরিদৃষ্টি সংখ্যাগুলে স্বাধীনভাবে/সাময়িকী-নিশ্চিতি না হওয়া পর্যন্ত অস্থায়ী।
- AI ষড়যন্ত্র বিশ্বাস “**নিরাময়**” করে না। ~20% গড় হ্রাস অর্থবহ, erasure নয়—অধিকাংশ অংশগ্রহণকারী afterward-ও বিশ্বাস রাখে।
- **বাস্তব জগত্রে ব্যবহারিক মতোভাবে** ফল নয়। গবেষণা অংশগ্রহণকারী opt-in করে ভালো সদৃশ্য AI-এর সঙ্গে আলে চিনায় অংশ নিয়েছে; বাস্তব জগত্রে অধিকাংশ দৃঢ়বিশ্বাসী বিশ্বাসী খণ্ডন বট খুঁজবে কনি অন্য প্রশ্ন।
- 99.2% নিরুভলতা এই **কাজ/মডলে/প্রম্পট**-এর নমুনা; ভাষা মডলে সাধারণত truthful—এমন নিশ্চয়তা নয়। একই ব্যক্তিকিনেদ্রিক প্ররোচনা ভুল বিশ্বাস বাড়াতো ব্যবহার হতে পারে।
- “স্থায়ী” এখানে **দুই মাস**, permanent নয়।

প্রমাণ কতটা শক্তিশালী?

নকশার শক্তি বাস্তব: চকিৎসা/ন্যূনতম পরীক্ষা, অসংশ্লিষ্ট-বিষয় প্লাসবিটো-শলৌ ন্যূনতম, দুই-গবেষণা প্রতিলিপি, স্থায়িত্ব অনুসরণকালীন এবং AI তথ্য নিরুভলতা যাচাই। দিক বিভিন্ন পরীক্ষায় সামঞ্জস্যপূর্ণ।

কনিত্ত **Expression of Concern footnote** নয়। প্রকাশিত কডি/তথ্য থকে সংখ্যা পুনরুৎপাদন করা পুনরুৎপাদনযোগ্যতার মূল অংশ, এবং ঠিক এখানে বাছাই অসামঞ্জস্য ও একত্রিত সারিসমস্যা হয়েছে। লেখকরো সংশোধিত পাইপলাইন প্রভাব সংরক্ষণ করে বলছেন—আশাব্যঞ্জক, কনিত্ত এটিলেখকদরে বিবরণ; *Science*-এর মূল্যায়ন স্বাধীন ধাপ।

সতর্ক অবস্থা: **সতর্কতাচিহ্নিত, খণ্ডিতও নয়, নতুন করে নিশ্চিতও নয়।** চোখে পড়ার মতো, ভালোভাবে-সুপরিকল্পিত বলে মনে হওয়া গবেষণা যার সূরিদৃষ্টি পরিমাণগত নথি আনুষ্ঠানিক প্রয়ালে চিনায়।

কনে এটি গুরুত্বপূর্ণ

প্রভাবশালী নরোশয়বাদী দৃষ্টিভঙ্গি বলে ষড়যন্ত্র বিশ্বাস অভদে/মনসাততবিকি প্রয়োজন-চালতি, তথ্য দিয়ে তরু করে বদলানো যায় না। ফল টকি থাকা করলে স্থায়ী ব্যক্তিকিনেদ্রিক প্রমাণ প্রতিক্রিয়া সেই pessimism-কে পাল্টা প্রমাণ দেয়: সাধারণ খণ্ডন নয়, বিশ্বাসী য়ে **নরিদৃষ্টি প্রমাণ** উদ্ধৃত করে AI বহু পরিসরে তা উত্তর দতিে পারে।

গবেষণাপত্রটির বর্তমান অবস্থাটিও বজ্জ্ঞানরে একটি শিক্ষা। বহুল প্রশংসিত ফলকে ত্রুটিমুক্ত ধরে নেওয়া বজ্জ্ঞানিক গুণ নয়; অসামঞ্জস্য

ধরা পড়া, তথ্য আবার চালানো, সাময়িকীকে জানানো এবং প্রকাশ্য সংশোধন-প্রক্রিয়া চালু হওয়াই ব্যবস্থার কাজ করার লক্ষণ—কলেঙ্কারি নয়। তাই সঠিক অবস্থান অন্ধ প্রশংসা বা তড়িঘড়ি খারজি—কোনো টিহি নয়; **অপেক্ষা করা**।

AI শিক্ষা দুই দিকের: ব্যক্তিকেন্দ্রিক যুক্তিভুল বিশ্বাস কমাতে পারলে ব্যক্তিকেন্দ্রিক misinformation ভুল বিশ্বাস বাড়াতো পারবে। কঠোর নরিপক্ষে; aim নরিপক্ষে নয়।

সংক্ষিপ্ত সারাংশ

2024 গবেষণায় GPT-4 Turbo-র তনি-দফা ব্যক্তিকেন্দ্রিক আলাপ অংশগ্রহণকারীদের নরিবাচতি ষড়যন্ত্র বিশ্বাস গড় ~20% কমিয়েছে; প্রতবেদিত প্রভাব 2 মাস টকি ছিল এবং sampled AI তথ্যভিত্তিক দাবিগুলো প্রায় সবকষতেরে নরিভুল ছিল। জুন 2026-এ *Science* গবেষণাপত্রে সম্পাদকীয় Expression of Concern দয়ে অংশগ্রহণকারী-বাছাই অসামঞ্জস্য ও জনসাধারণের ডটোসটে কডেকত্রীকরণ ত্রুটি নিয়ে। লখেকরো বলনে সংশোধিত বিশ্লষণ প্রভাবেরে দকি, তাৎপর্য ও আকার সংরক্ষণ করে; সাময়িকী এখনও মূল্যায়ন করছে। তাই ফল সত্যহি আকর্ষণীয় এবং সতর্কভাবে সুপরকিল্পতি বলে মনে হওয়া, কনিত্ত **বর্তমানে আনুষ্ঠানিক পর্যালোচনার অধীনে**—মীমাংসতি নয়, debunked-ও নয়।

উৎস

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>