

Human-AI dialogue research from the team of David Rand, Gordon Pennycook, and Tom Costello

- Durably reducing conspiracy beliefs through dialogues with AI [Science 2024](#) [[NYTimes write up](#)] [[MIT Tech Review Op Ed](#)] [[X thread](#)] [[BlueSky thread](#)] [[9min presentation](#)] [[45min presentation](#)] [[Experimental materials](#)] [[Browse the conversations](#)] [[Try DebunkBot yourself!](#)]
- Just the facts: How dialogues with AI reduce conspiracy beliefs [Working paper](#) [[BlueSky thread](#)] [[X thread](#)]
- Dialogues with Large Language Models reduce conspiracy beliefs even when the AI is perceived as human [PNAS Nexus 2025](#) [[BlueSky thread](#)]
- Reducing belief in conspiracy theories as they unfold using large language models [Working paper](#) [[BlueSky thread](#)] [[X thread](#)]
- Short dialogues with AI reduce belief in antisemitic conspiracy theories [Working paper](#) [[BlueSky thread](#)] [[LinkedIn thread](#)] [[20min presentation](#)] [[Try DebunkBot yourself](#)]
- Addressing climate change skepticism and inaction using human-AI dialogues [Working paper](#) [[BlueSky thread](#)] [[Try the bot yourself](#)] [[Browse the conversations](#)]
- Personalized Dialogues with AI Effectively Address Parents' Concerns about HPV Vaccination [Working paper](#) [[BlueSky thread](#)] [[Twitter thread](#)] [[Try the bot yourself](#)]
- Conversations with a large language model improve attitudes toward Muslims and Islam without harming attitudes toward Jews [Working paper](#)
- Increasing the effectiveness of charitable giving using human-AI dialogues [Working paper](#)
- Persuading Voters using Human-Artificial Intelligence Dialogues [Nature 2025](#) [[BlueSky thread](#)] [[X thread](#)] [[Atlantic write-up](#)]
- The Levers of Political Persuasion with Conversational AI [Science 2025](#) [[BlueSky thread](#)] [[X thread](#)] [[Podcast interview](#)]
- Deep canvassing using AI [Working paper](#) [[BlueSky thread](#)] [[Browse the conversations](#)]
- LLMs as Scalable Tools for Interactive Consumer Behavior Experiments: Comparing Persuasion Strategy Effectiveness [Working paper](#)
- It's the Thought that Counts: Evaluating the Attempts of Frontier LLMs to Persuade on Harmful Topics [Working paper](#)
- How Malicious AI Swarms Can Threaten Democracy [Working paper](#)
- A brief therapeutic conversation with AI decreases intentions to exclusively seek human therapy [Working paper](#)

Check out our other work related to understanding and combating misinformation here:

https://docs.google.com/document/d/1k2D4zVqkSHB1M9wpXtAe3UzbeE0RPpD_E2UpaPf6Lds/edit?usp=sharing