

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

A Rational Analysis of the Effects of Sycophantic AI

Permalink

<https://escholarship.org/uc/item/8s7834hk>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 48(0)

Authors

Batista, Rafael

Griffiths, Tom

Publication Date

2026

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

A Rational Analysis of the Effects of Sycophantic AI

Rafael M. Batista (rbatista@princeton.edu)

School of Public and International Affairs, Princeton University
Department of Psychology, Princeton University

Thomas L. Griffiths (tomg@princeton.edu)

Department of Computer Science, Princeton University
Department of Psychology, Princeton University

Abstract

People increasingly use large language models (LLMs) to explore ideas, gather information, and make sense of the world. In these interactions, users encounter chatbots that are overly agreeable. We argue this sycophancy poses an epistemic risk distinct from hallucination: it distorts belief not by introducing falsehoods but by biasing the evidence users see. A rational analysis shows that a Bayesian agent fed examples sampled from its own hypothesis grows more confident in that hypothesis without moving closer to the truth. We tested this prediction in a modified Wason 2-4-6 rule discovery task where participants ($N = 557$) interacted with AI agents providing different types of feedback. Unmodified LLM behavior suppressed participants' discovery and inflated their confidence comparably to explicitly sycophantic prompting. By contrast, independent sampling from the true distribution yielded discovery rates five times higher. This paper documents how sycophantic AI distorts belief, manufacturing certainty where there should be doubt.

Keywords: Sycophancy; Large Language Models (LLMs); Rational Analysis; Belief Updating; Confirmation Bias; Discovery