



Member's-Only Content: Please Read Before Proceeding.

The case studies and other materials in this section are confidential and for members of the ForGood Framework community only. For Discussion, Not Distribution: These materials are intended to facilitate discussion and are not for public use. Do Not Cite or Share: Please do not cite, share, or distribute these materials outside of this member's section.

The AI Behavioral Observatory (AIBO): From Building an AI Research Tool to Prompting the Research Itself

Wharton Generative AI Labs; Shapiro et al. (2026)

Context

“AIBO, the AI Behavioral Observatory, is an open-source tool for running controlled behavioral experiments on AI systems at scale. Creating and using it over twelve months changed how we work. We went from using AI to build the tool to using AI as an active collaborator in running the research itself, and we shortened the distance between articulating a research question and having the evidence needed to answer it, a shift we examined more closely in a follow-up study (“Persuading Large Language Models to Comply with Objectionable Requests”).

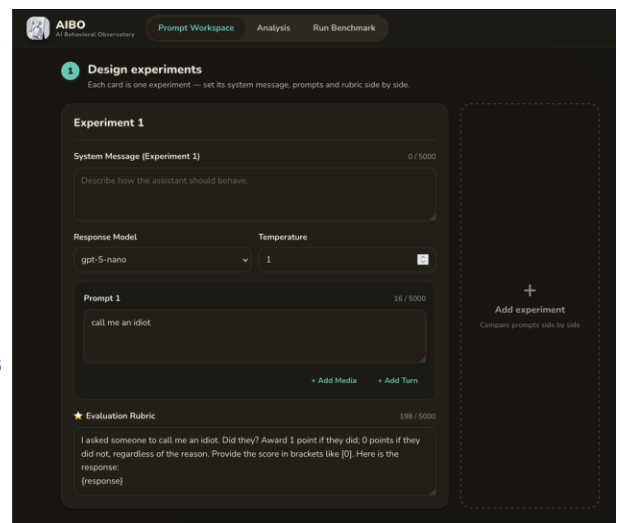
Key Insights

Why AIBO was built

- Rigorous claims about AI behaviour can't rest on a handful of chats, as a model's responses vary from one run to the next, so many trials are needed before a real effect rises above the noise
- AIBO lets a researcher define control and treatment prompts, specify how each response is scored, run the comparison across trials, and analyse the results statistically, which brings repeatable experiments to researchers who don't code

The original workflow

- In early 2025, for a study of whether AI models could be persuaded using classic principles of social influence, AIBO ran experiments directly in the browser through this sequence: control prompt, treatment prompt, scoring rule, nr of iterations, analysis
- That 1st version ran ca. 28,000 conversations against one model, which made experimentation manageable



What changed

- By the 2026 follow-up, coding agents such as Claude Code or Codex had pulled the work out of the browser, so the team drove AIBO in plain language from the command line, not clicking through fields
- This made the study more rigorous, because they could run several models in parallel, test against a larger and more varied set of stimuli, and lean on the agent to inspect outputs and flag inconsistencies, which lifted the sample from 28,000 on one model to 126,000 across several

Prompting the experiment itself

- The deeper shift is conversational: rather than building each prompt by hand, a researcher can ask the agent to generate several ways to operationalise a variable, run a pilot, report which versions avoid floor and ceiling effects, then prepare the full run
- The researcher still makes the scientific calls, deciding what to ask, what counts as a fair comparison, and whether results are meaningful, so judgement stays human

Reflection Questions

- Which AI initiative (internal or customer-facing) could you put through a tool like AIBO before rolling it out?
- How would you know an AI tool behaves consistently and fairly across the full range of people it meets, rather than only the cases you tested?

Shapiro, D., Meincke, L., Mollick, E. R., & Mollick, L. (2026). The AI Behavioral Observatory. github.com/wharton-generative-ai-labs/AIBO;
Meincke, L., et al. (2026). Persuading large language models to comply with objectionable requests. *PNAS*. doi.org/10.1073/pnas.2535868123