

EXPLORE

SEARCH

SUBSCRIBE TO OUR NEWSLETTER



Research ▾

Methodology

About

EXPLAINER JUNE 25, 2026 UPDATED AUGUST 6, 2026

Do Large Language Models Have Racial Bias? The Measured Evidence

Tests of today's leading AI chatbots find measurable bias against racial, ethnic, and religious groups, with the size varying sharply by model.

AI bias

Large language models

Antisemitism

Generative AI

Algorithmic bias



Anyone who has typed a question about a racial or religious group into a chatbot has probably wondered, at least once, whether the answer would differ depending on

which group they named. Researchers have now measured that question systematically: in a 2024 study that tried to push six leading models into producing biased content, the attempts succeeded between 0% and 76% depending on the model, averaging 75.7% across the eight bias categories tested, including antisemitism (Saeed and colleagues, 2024, a preprint). The bias is real and consistent in direction. Its size, and whether it surfaces in ordinary use or only when a model is deliberately pushed, varies by model.

Key Findings

- In a 2024 stress test of six leading models, attempts to draw out antisemitic content succeeded between **0% (Anthropic's Claude 3.5) and 76% (a smaller Meta Llama model)**, averaging 75.7% across the eight bias categories tested (Saeed and colleagues, 2024).
- AI image generators reproduce the same stereotypes: in a 2025 audit of AI-made images on a fringe forum, **28.8% contained racist content and 28.8% antisemitic content** (Gaba and De Cristofaro, 2025).
- A 2024 audit that probed 10 chatbots across 1,266 identity groups drew at least one toxic response in **70,477 of 75,960 attempts**, with racism and antisemitism recurring themes (Dutta and colleagues, 2024).
- Bias also appears as over-caution. A 2024 study found one model declined **56% of ordinary prompts about Jewish people**, the highest rate for any religion it tested (Plaza-del-Arco and colleagues, 2024).
- Newer or larger models are not reliably less biased (Plaza-del-Arco, 2024; Saeed, 2024).

How much racial bias do today's AI models show?

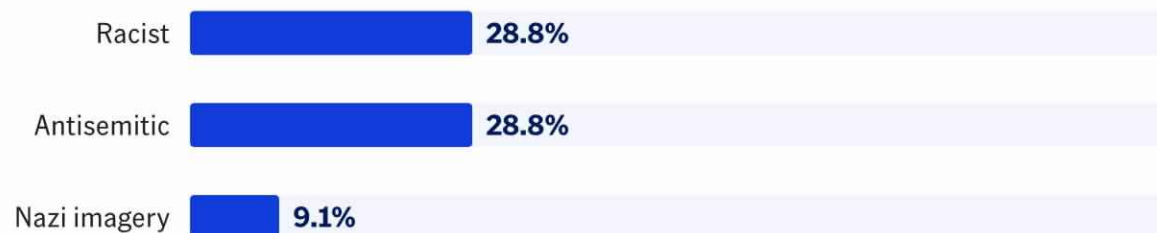
Attempts to push six leading models into biased content succeeded **75.7% of the time on average** across the eight bias categories tested, including antisemitism, in the clearest recent study (Saeed and colleagues, 2024, a preprint). There is no cleaner figure for race on its own,

because current systems are measured by stress tests and audits that bundle race, ethnicity, and religion together.

For racial content specifically, the sharpest measure is visual. In a 2025 audit of AI-generated images on a fringe forum, 28.8% carried racist content and 28.8% antisemitic content (Gaba and De Cristofaro, 2025, a preprint, from a sample of 66 images). Both kinds of figure describe what a system can be pushed to produce. How a model behaves on its own, without that pressure, is a separate and harder question, taken up below.

Share of AI-generated images on a fringe forum with biased content

2025 AUDIT, 66 AI-MADE IMAGES



Source: Gaba and De Cristofaro, 2025, preprint.

Which of today's models are most and least biased?

Attempts to elicit antisemitic content succeeded **0% of the time against Anthropic's Claude 3.5 and 76% against a smaller Meta Llama model**, with OpenAI's GPT-4 between them at 38%, in the same 2024 study (Saeed and colleagues, 2024). The size of the bias depends far more on the model than on the underlying technology.

Being newer did not guarantee being safer: OpenAI's optimized GPT-4o was easier to push into biased content than the earlier GPT-4, at 64% against 38%. The other systems fell in between, at 61% for the largest Llama model and 69% for Mistral. The spread is the practical point for anyone choosing a system, since two chatbots built on similar technology behaved very differently under the same test.

How often each model could be pushed into antisemitic content**SIX LEADING MODELS, 2024 STRESS TEST**



Source: Saeed and colleagues, 2024, preprint.

Does the bias show up in ordinary use, or only when pushed?

It shows up in both, but as two distinct findings. In ordinary use, a 2026 study found that when a leading model described fictional characters given Jewish names with no other cue, it rated them more privileged, more dominant, and less likable, matching long-standing antisemitic stereotypes (Gutman and Gilead, 2026, a peer-reviewed study). The jailbreak rates above measure only what a model can be forced to produce.

The provoked findings are real but narrower. The racist images came from a forum where users were actively trying to make them, and the audit that drew toxic output in 70,477 of 75,960 attempts did so by repeatedly prompting the models toward it (Dutta and colleagues, 2024, a preprint). A report that treats what a model can be tricked into as how it behaves by default would overstate the case. Kept apart, the two answer distinct questions: how a model leans on its own, and how well its safeguards resist abuse.

Are newer models getting less biased?

Not reliably. One study measured an older Meta Llama model declining **56% of ordinary prompts about Jewish people**, the highest rate for any religion, while a newer, larger model in the same family declined almost none (Plaza-del-Arco and colleagues, 2024, a preprint).

But the movement runs both ways. In the stress-test study, OpenAI's optimized GPT-4o was easier to push into biased content than the GPT-4 it followed (Saeed and colleagues, 2024). A measurement therefore describes one model at one moment. A figure for last year's release does not predict this year's, in either direction, which is why this bias has to be re-measured as systems are updated rather than settled once.

The same tools also find ethnic and religious bias

The tests that surface racial bias find religious bias in the same run, antisemitism most often. In the large 2024 audit, of the toxic responses that mentioned the Holocaust, **94.9% misrepresented it**, including denial and blaming Jews for it (Dutta and colleagues, 2024).

A second pattern is over-caution rather than hostility. The study that found a 56% refusal rate for prompts about Jewish people described it as stigmatization: by treating a group as inherently sensitive, a model can refuse to engage with it at all, a side effect of safety training that leans on hate-speech examples (Plaza-del-Arco and colleagues, 2024). These findings sit alongside the racial evidence because the same instruments produce them, not because race and religion are the same thing.

Methodology and limitations

The findings here come from three kinds of measurement on current systems. Stress tests, or red-teaming, try to defeat a model's safeguards and report how often they succeed. Audits prompt a model toward harmful output across many groups and classify what comes back. Representation probes ask a model to describe people indirectly and look for skew in the descriptions. Each measures something different, so a single number describes its method as much as the model.

Two limits apply. First, clean measures of a model's everyday behavior are scarcer than stress-test measures, so much of the current evidence describes worst-case behavior under pressure rather than typical use. Second, most of this work is recent preprint research rather than peer-reviewed journal studies; the one peer-reviewed study here is the 2026 representation probe (Gutman and Gilead, 2026). This report draws only on work measuring current-generation systems; older benchmarks of pre-ChatGPT models are a poor guide to the chatbots in use today and are not relied on.

Conclusion

So does a chatbot answer differently depending on which group you name? It does, and the difference is measurable from several directions at once rather than from any single one.

The provoked behavior splits sharply by system: the same jailbreak prompts drew antisemitic content 0% of the time from Anthropic's Claude 3.5 and 76% of the time from a smaller Meta Llama model, with OpenAI's GPT-4 at 38% and its newer GPT-4o higher at 64%. The default behavior carries the same lean more quietly, with a leading model rating characters given Jewish names more privileged, more dominant, and less likable when no other cue was present. The bias also runs in two directions at once, toward a group and away from it, with one model declining 56% of ordinary prompts about Jewish people, the highest refusal rate of any religion it tested. The pattern holds in other media too: 28.8% of a fringe-forum sample of AI-made images carried antisemitic content and 28.8% carried racist content, and in a large audit 70,477 of 75,960 attempts drew at least one toxic response, of which the Holocaust-touching ones misrepresented it 94.9% of the time. Being newer or larger did not reliably mean being safer, since the optimized GPT-4o was easier to push than the GPT-4 it followed.

These outputs do not stay inside the chat window. The privileged and dominant figure these models reach for unprompted, and the Holocaust denial that 94.9% of those Holocaust-touching responses reproduced, are the same money-and-power image and conspiracy tropes a reader already meets in everyday jokes and in claims about who secretly controls what. Cultural and informational portrayal is one of the inputs that shapes how a group is seen, which is why the Institute measures it. Whether a widely used tool, leaning toward old

stereotypes on its own and yielding them on demand, is one of the quiet feeders of antisemitism is the question the pattern raises, and it is a societal question, not one any single audit settles.

For the person on the other side of the prompt, the one a model rated more dominant and less likable before they typed a word, or the one whose ordinary question about Jewish life met a refusal 56% of the time, the audit stops where their day begins. What it is like to be the figure these systems sketch unasked, or the one they decline to discuss, is the part no measurement reaches, and it belongs to the people who carry it.

Frequently Asked Questions

Do all AI chatbots show the same level of racial bias?

No. The level depends sharply on the model. In a 2024 stress test, attempts to draw out antisemitic content succeeded 0% of the time against Anthropic's Claude 3.5 and 76% against a smaller Meta Llama model (Saeed and colleagues, 2024). A figure for one chatbot does not carry over to another.

Is producing a racist image on demand the same as having bias?

Not quite. Producing racist content when a user pushes for it measures whether a model's safeguards hold, which is different from how it behaves by default. A 2025 audit found 28.8% of a small sample of AI-generated images on a fringe forum contained racist content (Gaba and De Cristofaro, 2025), but that is provoked output, reported separately from everyday behavior.

Are newer AI models less biased than older ones?

Sometimes. A newer model cut an over-refusal bias from 56% to near zero across versions (Plaza-del-Arco and colleagues, 2024), but an optimized model was easier to push into biased content than the version it replaced (Saeed and colleagues, 2024). Newer or larger does not guarantee less biased.

Does a model refusing to answer mean it is unbiased?

No. Refusing to engage with a group at a disproportionate rate is itself a measured bias. One study found a model declined 56% of ordinary prompts about Jewish people, the highest rate for any religion, which the authors call stigmatization rather than safety (Plaza-del-Arco and colleagues, 2024).

Why does antisemitism appear in a report about racial bias?

Because the same tests measure racial, ethnic, and religious bias in one run. The large 2024 audit spanned 1,266 groups including nationalities, ethnicities, and religions together (Dutta and colleagues, 2024), so the measured patterns for these groups come from one body of evidence.

Are these findings peer-reviewed?

Most are recent preprints rather than peer-reviewed journal articles; the exception here is a 2026 peer-reviewed study of how a model represents Jewish-named characters (Gutman and Gilead, 2026). This report draws only on work measuring current AI systems, not older benchmarks of models that predate today's chatbots.

Sources

- Dutta, Khorramrouz, Dutta, KhudaBukhsh, 2024. Down the Toxicity Rabbit Hole: A Novel Framework to Bias Audit Large Language Models. arXiv:2309.06415. Preprint.
- Gaba, De Cristofaro, 2025. The Ethics of Generative AI in Anonymous Spaces: A Case Study of 4chan's /pol/ Board. arXiv:2506.14191. Preprint.
- Gutman, Gilead, 2026. From Myth to Model: Representation of "the Jew" in Generative AI. American Psychologist, DOI 10.1037/amp0001668. Peer-reviewed.
- Plaza-del-Arco, Cercas Curry, Paoli, Curry, Hovy, 2024. Divine LLaMAs: Bias, Stereotypes, Stigmatization, and Emotion Representation of Religion in Large Language Models. arXiv:2407.06908. Preprint.
- Saeed, Mohamed, Mohamed, Raza, Abdul-Mageed, Shehata, 2024. Desert Camels and Oil Sheikhs: Arab-Centric Red Teaming of Frontier LLMs. arXiv:2410.24049. Preprint.

RELATED RESEARCH

DISCOURSE JUNE 25, 2026

Is Anti-Zionism Antisemitism? What the Evidence Measures

Research does not settle whether anti-Zionism is antisemitism; it measures the link. A mediation model explains over 55% of anti-Israel attitudes.



The Hanover Institute for Public Policy studies the inputs fueling antisemitism in the United States and publishes what the evidence shows. Its work is content analysis, computational social science, and other data-driven research, and its reports take no policy positions.

SUBSCRIBE TO OUR NEWSLETTER

SUBSCRIBE

RESEARCH

Research

Methodology

INSTITUTE

About

Mission & Standards

Contact

TOPICS

Online & Social Media

Narratives & Tropes

Media & Representation

Israel & Anti-Zionism

AI & Technology

Incidents & Trends

Drivers & Effects

FORMATS

Data Report

Explainer

Narrative Analysis

Method

Trend Analysis

Comparison

Discourse

This material is distributed by Piro, Inc. on behalf of Havas Media Germany GmbH on behalf of the Israel Government Advertising Agency (LaPam). Additional information is available at the Department of Justice, Washington, DC.

© 2026 The Hanover Institute for Public Policy

[Terms of Use](#)

[Privacy Policy](#)