



July 31, 2026

Andrew Ferguson
Chairman
Federal Trade Commission
600 Pennsylvania Ave, NW
Washington, D.C. 20580

Submitted via Regulations.gov

Re: Federal Trade Commission’s Proposed Policy Statement Concerning the Suppression of Accuracy in Artificial Intelligence Systems, AI Policy Statement: Matter No. P264200

The National Consumer Law Center (NCLC)¹ submits these comments on behalf of its low-income clients in response to the Federal Trade Commission’s (“FTC” or “Commission”) request for public comment on its Proposed Policy Statement Concerning the Suppression of Accuracy in Artificial Intelligence Systems.² The proposed policy statement aims to rid artificial intelligence systems (“AI systems”) of “undisclosed ideological objectives” that purportedly deceive consumers in violation of Section 5 of the Federal Trade Commission Act (“FTC Act”).³ In doing so the statement, written in compliance with Executive Order 14365 and other policy pronouncements on “woke AI,” undermines states’ ability to protect consumers from discriminatory AI systems and companies’ attempts to take corrective actions for such defective models. The proposed policy statement creates confusion in the market and chills compliance with state and federal antidiscrimination laws as companies strain to comply with an ill-defined policy standard.⁴ The FTC should withdraw the proposed policy statement in its entirety and put in place stronger protections and safeguards for consumers.

¹ The **National Consumer Law Center, Inc. (NCLC)** is a non-profit Massachusetts Corporation, founded in 1969, specializing in low-income consumer issues, with an emphasis on consumer credit. On a daily basis, NCLC provides legal and technical consulting and assistance on consumer law issues to legal services, government, and private attorneys representing low-income consumers across the country. NCLC publishes a series of practice treatises on consumer credit laws and unfair and deceptive practices. NCLC attorneys have written and advocated extensively on all aspects of consumer law affecting low-income people, conducted trainings for tens of thousands of legal services and private attorneys, and provided extensive oral and written testimony to numerous Congressional committees on various topics.

² Fed. Trade Comm’n, Proposed Policy Statement Concerning the Suppression of Accuracy in Artificial Intelligence Systems (“Statement”), 91 Fed. Reg. 41638 (July 7, 2026).

³ See 91 Fed. Reg. at 41638.

⁴ See *id.* (quoting the remarks of Vice President J.D. Vance). See also Executive Order 14319 Preventing Woke AI in the Federal Government.

A. The proposed policy statement is unsupported by facts or evidence; and its foundation is underpinned by false assumptions and ideological bias.

In pursuit of ideological aims, the Commission’s proposed policy statement contorts the facts used to justify its policy on AI systems. The Commission did not engage in independent fact finding or enforcement to buttress its justification for issuing the policy guidance.⁵

The proposed policy statement is based on the premise that consumers trust AI-generated content to be accurate and truthful.⁶ However, in support of this premise the Commission cites a news article reporting on a single study that showed that 91% of users of one AI company’s system did not fact-check the AI output.⁷ As the news article notes, the company’s researchers “acknowledge multiple possible explanations” for the low fact-checking rate.⁸ For one thing, the study was able to track only 11 of 24 steps that the company identified as promoting “safe and effective human-AI collaboration.”⁹ If users took any of the other 13 steps to check the AI output, the study did not track or report them. For example, the study noted that “[u]sers may be evaluating artifacts outside the conversation — running code, testing an app, sharing a draft with a colleague.”¹⁰ The study also did not evaluate whether users made any use of the AI output, as opposed to discarding it. The study is hardly evidence that AI users “trust these systems across a broad swath of applications.”¹¹

To the contrary, a 2026 Quinnipiac University poll documented that 76% percent of Americans think they can trust AI either hardly ever (27%) or only some of the time (49%), while only 21% think they can trust AI either most of the time (18%) or almost all of the time (3%).¹² The inaccuracies of AI-generated content are so widely known that the Merriam-Webster Dictionary named “slop”--defined as “digital content of low quality that is produced usually in quantity by means of artificial intelligence”--as its 2025 Word of the Year.¹³ The phenomenon of AI hallucinations has also received widespread media coverage,¹⁴ and there is even a database that

⁵ See 91 Fed. Reg. at 41639 (statement issued in compliance with Executive Order 14365).

⁶ See 91 Fed. Reg. at 41640 (“Consumers ... have learned to trust these systems across a broad swath of applications....”).

⁷ 91 Fed. Reg. at 41640 n.39, citing Ted Ladd, Anthropic: 91% of Users Do Not Fact-Check AI. Let’s Fix That, FORBES (Mar. 5, 2026), <https://www.forbes.com/sites/tedladd/2026/03/05/anthropic-91-of-users-do-not-fact-check-ai-lets-fix-that/>.

⁸ Ted Ladd, Anthropic: 91% of Users Do Not Fact-Check AI. Let’s Fix That, FORBES (Mar. 5, 2026), <https://www.forbes.com/sites/tedladd/2026/03/05/anthropic-91-of-users-do-not-fact-check-ai-lets-fix-that/>.

⁹ Anthropic Education Report: The AI Fluency Index (Feb. 23, 2026), <https://www.anthropic.com/research/AI-fluency-index>.

¹⁰ Ted Ladd, Anthropic: 91% of Users Do Not Fact-Check AI. Let’s Fix That, FORBES (Mar. 5, 2026), <https://www.forbes.com/sites/tedladd/2026/03/05/anthropic-91-of-users-do-not-fact-check-ai-lets-fix-that/>.

¹¹ 91 Fed. Reg. at 41640. See also 91 Fed. Reg. at 41638 (“consumers have a reasonable expectation that AI systems aim to give truthful and accurate outputs”).

¹² Quinnipiac University Poll, The Age Of Artificial Intelligence: Americans' AI Use Increases While Views On It Sour, Quinnipiac University Poll On AI Finds; 7 In 10 Think AI Will Cut Jobs With Gen Z The Most Pessimistic (Mar. 30, 2026), <https://poll.qu.edu/poll-release?releaseid=3955>.

¹³ Merriam-Webster, 2025 Word of the Year: Slop, <https://www.merriam-webster.com/wordplay/word-of-the-year>.

¹⁴ See, e.g., Santul Nerkar, *New York Times*, AI ‘Hallucinations’ Created Errors in Court Filing, Top Firm Says (Apr. 21, 2026), <https://www.nytimes.com/2026/04/21/nyregion/sullivan-cromwell-ai-hallucination.html>; Mike Scarcella, Reuters, US appeals court sanctions lawyers over AI ‘hallucinations,’ lack of candor (June 3, 2026),

tracks legal decisions in which AI produced hallucinated content (1797 cases so far).¹⁵ Recently an AI model went rogue, escaping a security test and autonomously attacking a rival AI developer, making headlines around the world.¹⁶ The premise that consumers widely trust AI outputs is simply unmaintainable.

B. Companies must test and evaluate AI systems to improve accuracy and reliability, ensure AI systems are free of bias, and operate in compliance with federal and state laws.

In the proposed policy statement, the Commission appears to be positing that any attempt to debias an AI system is a means to infect “ideological bias” and make the system inaccurate.¹⁷ This is inapposite to the goal of such activities. Testing and evaluation of AI systems, including debiasing, makes systems more accurate, not less. Such activities are part of a robust model evaluation system and essential for legal compliance.

Discrimination can be introduced at each stage of the model development process, including reliance on limited and unrepresentative data, or data weighted or labeled in a biased way, in the defining of output variables, and incorrectly applying patterns from one situation to a different scenario.¹⁸ AI systems are trained on historical data that reflects discriminatory patterns and the output of the model may perpetuate the same.¹⁹ Unless AI developers take strong, continuing, affirmative steps to debias their AI systems, those systems will learn and reproduce inaccurate and harmful content.

Poor design and deployment of AI systems not only amplifies discriminatory patterns but also increase costs to consumers and creates barriers to access especially in the credit and housing market. We have seen instances of where poorly tested models have yielded higher price credit in a discriminatory matter.²⁰ Facebook (now Meta) was challenged for discriminatory ad targeting and ad delivery systems which cut consumers off from housing.²¹ Redfin was

<https://www.reuters.com/legal/litigation/us-appeals-court-sanctions-lawyers-over-ai-hallucinations-lack-candor-2026-06-03/>; Cathy Bussewitz, AP, Mistake-filled legal briefs show the limits of relying on AI tools at work (Oct. 30, 2025).

¹⁵ Damien Charlotin, AI Hallucination Cases, <https://www.damiencharlotin.com/hallucinations/> (last visited on July 24, 2026) (listing each decision and the outcome or sanction).

¹⁶ See New York Times, OpenAI Says Its A.I. Models Went Rogue and Attacked a Digital Library (July 21, 2026), <https://www.nytimes.com/2026/07/21/technology/openai-attack-hugging-face.html>; Fox News, OpenAI agent goes rogue in 'unprecedented' cybersecurity incident (July 23, 2026), <https://www.foxnews.com/video/6401902824112>; BBC, OpenAI says its AI went rogue and launched 'unprecedented' cyber-attack (July 22, 2026), <https://www.bbc.com/news/articles/c3ek3gvdnj3o>.

¹⁷ See 91 Fed. Reg. at 41641.

¹⁸ See David Lehr & Paul Ohm, Playing with the Data: What Legal Scholars Should Learn About Machine Learning, 51 U.C. Davis L. Rev. 653, 662-63 (2017).

¹⁹ See Carol Evans, Board of Governors of the Federal Reserve System, From Catalog to Clicks, The Fair Lending Implications of Targeted, Internet Marketing, Consumer Compliance Outlook (Second Issue 2017) at 4.

²⁰ Educational Redlining, Student Borrower Protection Center, Feb. 2020, available at <https://protectborrowers.org/wp-content/uploads/2020/02/Education-Redlining-Report.pdf>

²¹ Louise Matsakis, *Facebook's Ad System Might be Hard-Coded for Discrimination*, Wired (Apr. 6, 2019), <https://www.wired.com/story/facebooks-ad-system-discrimination/>; Muhammad Ali, et al., *Discrimination through Optimization: How Facebook's Ad Delivery Can Lead to Biased Outcomes*, 3 Proceedings of the ACM on Human-

challenged for allegedly digitally redlining communities of color with its minimum home listing price policy that offered no marketing service for homes in certain communities.²²

Financial institutions have fair lending testing, compliance and monitoring regimens in place to decrease risk. Fair lending evaluation of AI systems is a continuation of the assessment of risk undertaken with respect to more traditional models and systems. Model risk management assesses credit risk, and safety and soundness, and encourages the development and validation of accurate models.²³ Validation and monitoring of AI systems used in underwriting mitigates discrimination and other systemic risks. Federal laws encourage robust testing and monitoring. The Equal Credit Opportunity Act, for example, encourages creditors to self-test their technological and other systems to determine the effectiveness and extent of their compliance with the Act.²⁴

Consumers also demand this rigorous testing and evaluation. They expect that any technology deployed in the market will be free of bias and in compliance with civil rights and consumer protection laws. The stakes are highest when consumers interact with AI systems that make consequential decisions regarding access to credit, housing, employment, health care and benefits. A recent poll demonstrates ongoing concerns regarding discriminatory output. The Leadership Conference's Center for Civil Rights and Technology, in partnership with Voss Research and Strategy, found that 75% of people support requiring companies that use AI for decisions about housing, loans and employment prove that their AI systems do not discriminate. There were similar levels of support for preventing AI uses that lead to discrimination, and in support of testing for bias.²⁵

This proposed policy statement undermines such robust evaluation systems and creates confusion as to what is a legally acceptable manner to assess for bias and discrimination. The proposed policy statement will lead to a weakening of compliance with fair lending and consumer protection laws, and result in less accurate outputs from AI systems.

C. The Commission's assertion of preemption has no basis.

In the proposed policy statement, the Commission states that "state law is impliedly preempted to the extent it conflicts with a federal regulatory scheme." This statement is a well-known principle of preemption law. But the Commission goes further to assert that "[a] state law that requires an AI firm to deceive its consumers obviously conflicts with Section 5's express purpose of protecting consumers from such conduct."²⁶ This blanket statement is purely

Computer Interaction, Nov. 2019, at 199:2, <https://www.ccs.neu.edu/~amislove/publications/FacebookDelivery-CSCW.pdf>.

²² See NFHA, Our Redfin Investigation, available at <https://nationalfairhousing.org/issue/redfin-investigation/>.

²³ Bd. of Governors of the Fed. Res. Sys. & Off. of Comptroller of the Currency, Supervisory Guidance on Model Risk Management (Apr. 4, 2011) (providing guidance on model development, implementation, use, validation, governance, policies, and controls).

²⁴ See Regulation B, 12 C.F.R.1002.15.

²⁵ Polling: American Public Supports Policies Preventing AI Discrimination and Bias, Leadership Conference on Civil and Human Rights, available at <https://civilrights.org/ai-civil-rights-polling/>.

²⁶ 91 Fed. Reg. at 41641.

hypothetical. It does not outline how a state law might conflict with the ill-conceived and unsupported federal regulatory scheme the policy statement proposes.

Moreover, the Commission's statement about preemption is infected with the same flaw that pervades the proposed policy statement as a whole; it cites no basis for any belief that consumers are being deceived about the accuracy of AI systems, or that they want AI systems to promote unlawful discrimination. Without evidence that consumers are being deceived, there is no basis for preemption or any basis for issuing the policy statement at all. Nor has the Commission identified any state law that requires AI firms to affirmatively deceive their customers.

Conclusion

The proposed policy statement is a clear deviation from its stated aim – to create AI systems free of bias. Rather the policy statement creates confusion in the market and chills compliance with state and federal laws. The FTC should withdraw the proposed policy statement.

We appreciate the opportunity to comment on this issue. Please contact Odette Williamson, Director of Racial Justice Advocacy, owilliamson@nclc.org with any questions.