



LAWYERS' COMMITTEE FOR
CIVIL RIGHTS
U N D E R L A W

July 31, 2026

Andrew Ferguson
Chairman
Federal Trade Commission
600 Pennsylvania Ave, NW
Washington, DC 20580
Submitted via Regulations.gov

RE: AI Policy Statement, Matter No. P264200

Dear Chairman Ferguson,

On behalf of the Lawyers' Committee for Civil Rights Under Law, a nonpartisan, nonprofit civil rights organization, and the undersigned eight organizations that promote and protect civil and human rights, we submit this comment 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 ("proposed policy statement").¹

We urge the Commission to withdraw the proposed policy statement. It originates from partisan priorities, not from Commission findings on how artificial intelligence ("AI") and AI regulation impact consumers, or from expertise garnered through careful study, investigation, or law enforcement.² Instead of prioritizing consumer protection, it prioritizes global dominance and corporate interests³ from the flawed and scientifically unsupported perspective that regulation stifles innovation; it naively assumes that AI developers and deployers are primarily motivated to provide the best and most "accurate" outputs to consumers.⁴ It ignores well-documented harms resulting from biased, discriminatory, and inaccurate AI systems while seeking to undermine laws that would address those harms. And, it does so based upon flawed and unsubstantiated assumptions. Finally, the proposed policy statement would create confusion detrimental to consumers and businesses alike, that ultimately would result in the rollback or weakening of critical civil rights protections.

¹ *Policy Statement Concerning the Suppression of Accuracy in Artificial Intelligence Systems*, 91 Fed. Reg. 41,638 (proposed July 7, 2026) [hereinafter *Statement*].

² *Id.* at 41,639 (noting that Executive Order 14365 directs the Commission to issue the policy statement). Notably, the executive order did not grant the Commission discretion but declared that the "policy statement must explain circumstances under which State laws that require alterations to the truthful outputs of AI models are preempted by the Federal Trade Commission Act's prohibition on engaging in deceptive acts or practices affecting commerce." Exec. Order No. 14365, *Ensuring a National Policy Framework for Artificial Intelligence*, § 7, 90 Fed. Reg. 58,499, 58,500 (Dec. 11, 2025).

³ *Statement*, *supra* note 1, at 91 Fed. Reg. 41,638-39.

⁴ *Id.* at 41,640.

I. The Proposed Policy Statement Fails to Acknowledge Extensively Documented and Ongoing Algorithmic Harms and Biases

The proposed policy statement raises speculative concerns that consumers using AI systems would be deceived as an indirect result of state AI regulation while ignoring and dismissing the ways in which AI systems actually harm consumers.⁵ As the proposed policy statement recognizes, AI systems “are now used across nearly every sector of the economy.”⁶ Not only does this economy-wide use include instances in which consumers seek recommendations or information from generative AI systems, it also includes the deployment of AI systems *on* consumers. Consumers, often without their awareness or consent, are subjected to AI systems that surveil them, determine their worthiness for employment, housing, and credit, and make other risk-based predictions, recommendations, or decisions about them.⁷

Scientific consensus exists “that artificial intelligence (‘AI’) can exacerbate bias and discrimination.”⁸ As numerous well-documented instances illustrate, both generative and automated decision-making (“ADM”) AI systems have operated in inaccurate, discriminatory, and harmful ways, with Black and brown communities often bearing the brunt. Because racism and other biases are age-old problems affecting the data on which AI systems rely, the notion that AI systems are somehow immune from that phenomenon defies logic. AI systems are trained on data shaped by historical bias and discrimination⁹ and also may reflect the biases of their developers.¹⁰ This well-documented bias, itself a source of inaccurate AI outputs, occurs in the absence of, not because of, protections against AI bias and discrimination. Examples of bias in the context of ADM abound. For instance, such AI systems have discriminatorily denied or

⁵ The Commission’s turn of phrase “so-called ‘equity,’” *Statement, supra* note 1, at 91 Fed. Reg. 41,641, demonstrates a disturbing contempt for civil rights and equality.

⁶ *Statement, supra* note 1, at 91 Fed. Reg. 41,638.

⁷ FUTURE OF PRIV. FORUM, UNFAIRNESS BY ALGORITHM: DISTILLING THE HARMS OF AUTOMATED DECISION-MAKING (2017), <https://fpf.org/wp-content/uploads/2017/12/FPF-Automated-Decision-Making-Harms-and-Mitigation-Charts.pdf>.

⁸ *Affirming the Scientific Consensus on Bias and Discrimination in AI*, <https://www.aibiasconsensus.org/> (last visited July 30, 2026).

⁹ Atin Jindal, *Misguided Artificial Intelligence: How Racial Bias is Built Into Clinical Models*, 2 BROWN J. HOSP. MED. 38021 (2022).

¹⁰ Ankit Singh, *Should Developers Care About AI Bias?*, MEDIUM (Oct. 16, 2025), <https://medium.com/@ankitsingh1583/should-developers-care-about-ai-bias-13459aac267c>.

imposed conditions on access to housing¹¹ and credit.¹² The same holds true in the context of employment,¹³ education,¹⁴ healthcare,¹⁵ and criminal justice.¹⁶

Nor has generative AI, which garners much of the proposed policy statement's focus, been immune from bias and discrimination. For example, a study of an AI image generator demonstrated bias, showing that the tool disproportionately underrepresents certain demographic groups in ways inconsistent with actual labor demographics.¹⁷ Research also has shown that ChatGPT exhibited linguistic bias¹⁸ and amplified gender stereotypes.¹⁹ Further, AI systems that provide recommendations to consumers often are infected with various biases or motivations aside from truth and accuracy, including popularity bias,²⁰ and outputs to encourage continued engagement with the platform rather than what the Commission believes is the "truthful."²¹

AI systems have also harmed consumers by enabling mass surveillance, which should concern the Commission given its role as a privacy regulator and its recognition that "AI technologies should not be used to silence or censor lawful expression or dissent."²² AI-driven surveillance systems have enabled invasion of consumer privacy in a plethora of ways, including through facial recognition technology,²³ other types of automated biometric systems, license plate readers, and tools created to scrape social media to identify individuals whose ideology the government opposes. Further, some of these technologies, such as facial recognition technology, have been shown to operate at lower rates of accuracy on certain demographic groups, including

¹¹ See Emmanuel Martinez & Lauren Kirchner, *The Secret Bias Hidden in Mortgage-Approval Algorithms*, AP NEWS (Aug. 25, 2021), <https://apnews.com/article/lifestyle-technology-business-race-and-ethnicity-mortgages->

¹² See Robert Bartlett et al., *Consumer-Lending Discrimination in the FinTech Era* (NAT'L BUREAU ECON. RSCH., Working Paper No. 25943, 2019).

¹³ MIRANDA BOGEN & AARON RIEKE, UPTURN, HELP WANTED: AN EXAMINATION OF HIRING ALGORITHMS, EQUITY, AND BIAS 1 (2018), <https://www.upturn.org/static/reports/2018/hiring-algorithms/files/Upturn%20--%20Help%20Wanted%20-%20An%20Exploration%20of%20Hiring%20Algorithms,%20Equity%20and%20Bias.pdf>.

¹⁴ Todd Feathers, *College Prep Software Naviance Is Selling Advertising Access to Millions of Students*, THE MARKUP, <https://themarkup.org/machine-learning/2022/01/13/college-prep-software-naviance-is-selling-advertising-access-to-millions-of-students> (last updated Jan. 14, 2022).

¹⁵ Natalia Norori et al., *Addressing Bias in Big Data and AI for Health Care: A Call for Open Science*, 2 PATTERNS 100347 (2021).

¹⁶ Aaron Sankin et al., *Crime Prediction Software Promised to Be Free of Biases. New Data Shows It Perpetuates Them*, THE MARKUP (Dec. 2, 2021), <https://themarkup.org/prediction-bias/2021/12/02/crime-prediction-software-promised-to-be-free-of-biases-new-data-shows-it-perpetuates-them>.

¹⁷ Leonardo Nicoletti & Dina Bass, *Generative AI Takes Stereotypes and Bias from Bad to Worse*, BLOOMBERG (June 9, 2023), <https://www.bloomberg.com/graphics/2023-generative-ai-bias/>.

¹⁸ Eve Fleisig et al., *Linguistic Bias in ChatGPT: Language Models Reinforce Direct Discrimination*, in PROCEEDINGS OF THE 2024 CONFERENCE ON EMPIRICAL METHODS IN NATURAL LANGUAGE PROCESSING 13541 (2024).

¹⁹ Leander Gierbach et al., *A Large Scale Analysis of Gender Biases in Text-to-Image Generative Models*, ARXIV (Mar. 30, 2025), <https://arxiv.org/pdf/2503.23398>.

²⁰ Himan Abdollahpouri, *Popularity Bias in Ranking and Recommendation*, in PROCEEDINGS OF THE 2019 AAAI/ACM CONFERENCE ON AI, ETHICS, AND SOCIETY 529 (2019).

²¹ *Statement*, *supra* note 1, at 91 Fed. Reg. 41,638.

²² *Id.* at 41,639.

²³ *Rite Aid Banned from Using AI Facial Recognition After FTC Says Retailer Deployed Technology without Reasonable Safeguards*, FED. TRADE COMM'N (Dec. 19, 2023), <https://www.ftc.gov/news-events/news/press-releases/2023/12/rite-aid-banned-using-ai-facial-recognition-after-ftc-says-retailer-deployed-technology-without>.

Black and Asian people and women, putting members of those groups at heightened risk of being misidentified or targeted.²⁴

II. The Proposed Policy Statement Lacks Substantiation and Relies on Flawed Assumptions

Instead of tackling real and documented harms caused by AI systems, the proposed policy statement seeks to address the hypothetical harm that consumers using AI systems may somehow be deceived if AI developers or deployers comply with AI regulation. The Commission provides no substantiation that compliance with AI regulation has actually resulted in deception. Additionally, the Commission's theory relies on the flawed assumption that consumers "trust" AI and that compliance with AI regulation would require suppression of accurate outputs.

A. The Proposed Policy Statement Does Not Substantiate Its Claims that AI Regulation Leads to Deception

The proposed policy statement does not cite any documented instance in which a consumer using an AI system was deceived due to a platform or developer's attempt to comply with civil rights protections or state AI regulations. Significantly, the only state AI regulation named in the proposed policy statement is the Colorado AI Act, which was repealed in large part and has yet to go into effect.²⁵ Moreover, the proposed policy statement lacks grounding in Commission expertise or experience; it does not cite to any Commission study finding that AI regulation causes deception or any law enforcement actions that the Commission brought claiming that consumers suffered deceptive or unfair practices due to efforts to comply with state AI regulation. In fact, it contravenes past Commission study and findings regarding the risks that AI poses to consumers,²⁶ as well as those of other federal agencies.²⁷ Commission policy should be grounded in facts and expertise that the Commission has garnered through study and law enforcement, not pro-business speculation supported only by the worn trope that regulation stifles innovation.

²⁴ NAT'L ACADS. OF SCIS., ENG'G & MED., FACIAL RECOGNITION TECHNOLOGY: CURRENT CAPABILITIES, FUTURE PROSPECTS, AND GOVERNANCE (2024).

²⁵ Devin Culbertson, *Colorado AI Law Reset: Discrimination Duty Dropped, Disclosure Takes Its Place*, TECH TIMES (July 1, 2026), <https://www.techtimes.com/articles/319420/20260701/colorado-ai-law-reset-discrimination-duty-dropped-disclosure-takes-its-place.htm>.

²⁶ FED. TRADE COMM'N, POLICY STATEMENT OF THE FEDERAL TRADE COMMISSION ON BIOMETRIC INFORMATION AND SECTION 5 OF THE FEDERAL TRADE COMMISSION ACT (2023), https://www.ftc.gov/system/files/ftc_gov/pdf/p225402biometricpolicystatement.pdf; FED. TRADE COMM'N, A LOOK BEHIND THE SCREENS: EXAMINING THE DATA PRACTICES OF SOCIAL MEDIA AND VIDEO STREAMING SERVICES (2024), https://www.ftc.gov/system/files/ftc_gov/pdf/Social-Media-6b-Report-9-11-2024.pdf; *Rite Aid Banned*, *supra* note 23.

²⁷ U.S. Commission on Civil Rights Releases Report: *The Civil Rights Implications of the Federal Use of Facial Recognition Technology*, U.S. COMM'N ON C.R. (Sept. 19, 2024), <https://www.usccr.gov/news/2024/us-commission-civil-rights-releases-report-civil-rights-implications-federal-use-facial>; REHA SCHWARTZ ET AL., NAT'L INST. OF STANDARDS & TECH., NIST SP 1270, TOWARDS A STANDARD FOR IDENTIFYING AND MANAGING BIAS IN ARTIFICIAL INTELLIGENCE (2022), <https://nvlpubs.nist.gov/nistpubs/SpecialPublications/NIST.SP.1270.pdf>; PRIV. & CIV. LIB. OVERSIGHT BD., USE OF FACIAL RECOGNITION TECHNOLOGY BY THE TRANSPORTATION SECURITY ADMINISTRATION: STAFF REPORT (May 2025), [https://documents.pclob.gov/prod/Documents/OversightReport/90964138-44eb-483d-990e-057ce4c31db7/Use%20of%20FRT%20by%20TSA,%20PCLOB%20Report%20\(5-12-25\),%20Completed%20508,%20May%2019,%202025.pdf](https://documents.pclob.gov/prod/Documents/OversightReport/90964138-44eb-483d-990e-057ce4c31db7/Use%20of%20FRT%20by%20TSA,%20PCLOB%20Report%20(5-12-25),%20Completed%20508,%20May%2019,%202025.pdf).

B. The Policy Statement Misconceives How AI Works, How AI Would Impact Consumers, and Consumer Expectations

The proposed policy statement relies on misconceptions to establish its theory that AI regulation would deceive consumers as to the objectives of an AI system, starting with the premise that consumers “overwhelmingly trust that the system is designed to provide an answer tailored specifically to their objectives.”²⁸ This conclusion contradicts data showing wide-ranging consumer distrust of AI.²⁹ Such research includes studies by commercial actors,³⁰ government entities,³¹ and non-government organizations.³²

Regardless, the Commission should not rely on unsubstantiated assumptions about whether consumers generally trust AI as the basis for its enforcement policy. Whether an act or practice deceives consumers in violation of Section 5 depends not on the Commission’s general and misguided sense of consumer sentiment but on a contextual analysis of a specific claim.³³ Disregarding this well-established precept, the proposed policy statement makes a bizarre leap; it concludes that consumers would take general representations that AI systems are “useful and reliable,”³⁴ assist in solving problems,³⁵ or are “helpful”³⁶ to mean that AI, unlike all other goods and services, would operate in a vacuum from applicable laws and regulations. It also ignores the possibility that to many consumers, an AI system may not be “useful,” “helpful,” “accurate,” or “reliable,” if it discriminated against them, produced discriminatory outputs, or functioned less accurately because of the color of their skin or other physical attributes.

Furthermore, the Commission’s theory that AI regulation would suppress “accurate” outputs lacks factual and scientific grounding. Generative AI models are not designed to provide the truth; they predict words or word sequence based on patterns. With respect to ADM AI systems, the notion that AI regulation would result in less accurate outputs presumes that those systems can accurately predict outcomes and behavior in the first place, even though scientific research has shown that they cannot reliably forecast human behavior.³⁷ It also leads to the bizarre conclusion, for example, that it would be deceptive for a company to attempt to solve for

²⁸ *Statement*, *supra* note 1, at 91 Fed. Reg. 41,640.

²⁹ *What Americans Really Think About AI Algorithms: Public Confidence and Transparency in Government*, CORNELL BROOKS PUB. POL’Y (Nov. 12, 2025), <https://publicpolicy.cornell.edu/masters-blog/what-americans-really-think-about-ai-algorithms-public-confidence-and-transparency-in-government/>.

³⁰ Clifton Mark, *Most Americans use AI but still don’t trust it*, YOUTGOV (Dec. 9, 2025), [https://yougov.com/en-us/articles/53701-most-americans-use-ai-but-still-dont-trust-it-;](https://yougov.com/en-us/articles/53701-most-americans-use-ai-but-still-dont-trust-it-) *see also* Julie Ray, *Americans Express Real Concerns About Artificial Intelligence*, GALLUP NEWS (Aug. 26, 2024), <https://news.gallup.com/poll/648953/americans-express-real-concerns-artificial-intelligence.aspx>.

³¹ *What Americans Really Think*, *supra* note 29.

³² David Girling & Deborah Adesina, *Artificial authenticity: are NGOs risking their reputation when using AI-generated imagery?*, BOND (Mar. 23, 2026), <https://www.bond.org.uk/news/2026/03/artificial-authenticity-are-ngos-risking-their-reputation-when-using-ai-generated-imagery/>.

³³ FTC Policy Statement on Deception, *appended to Cliffdale Assocs., Inc.*, 103 F.T.C. 110, 174 (1984).

³⁴ *Statement*, *supra* note 1, at 91 Fed. Reg. 41,640 n.35.

³⁵ *Id.*

³⁶ *Id.*

³⁷ Seb Murray, *How to spot real value in AI — and avoid the snake oil*, MIT SLOAN (June 23, 2025), <https://mitsloan.mit.edu/ideas-made-to-matter/how-to-spot-real-value-ai-and-avoid-snake-oil>.

a tenant screening system that approves white applicants for rental housing but denies similarly situated Black applicants.

III. The Proposed Policy Statement Creates Uncertainty and Risk for Businesses and Consumers

The proposed policy statement is irresponsible. It muddies the waters rather than providing clarity as policy statements should, to the detriment of consumers and critical civil rights protections. It takes the unprecedented position that compliance with state law could place a company in the crosshairs of the FTC. Thus, businesses will have to reassess risk and compliance and may no longer be able to assume that they can comply with state law without risk. It raises the possibility that compliance with not only AI regulations but also longstanding civil rights laws, such as those that prohibit disparate impact discrimination could expose a company to FTC liability. This confusion will chill compliance with civil rights and consumer protection laws and incentivize businesses to roll back compliance, endangering consumers given the critical role that states have played in protecting consumer and civil rights. Crucial state protections cannot be preempted by threat of FTC enforcement, including based upon the unclear standards and unsupported and flawed rationale set forth in the proposed policy statement.

The lack of specificity and legal grounding in the proposed policy statement compounds the uncertainty. It states, “the Commission believes AI companies that steer the outputs of their AI systems toward unexpected objectives, and away from the objectives set by or reasonably expected by users are likely to deceive.” No clear standard exists for what constitutes an “unexpected objective,” especially if the Commission has determined that compliance with laws could be “unexpected.” Also unclear is the threshold for determining that compliance with an applicable law has crossed the line from “expected” to “unexpected.”

The Commission’s proposed policy statement represents a clear deviation from its stated mission to protect Americans from anticompetitive, unfair, and deceptive business practices—in the first place—without unduly burdening legitimate business activity— as an ancillary. Instead, it endorses industry claims and untested practices, while resting on unsupported assumptions. Consequently, the undersigned organizations oppose this policy statement and urge the FTC to adopt policies that fulfill its chartered mission to protect Americans. We appreciate the opportunity to provide comments on this issue. Please contact, Leah Frazier, Director, Digital Justice Initiative, at lfrazier@lawywerscommittee.org, should you have any further questions.

Sincerely,

/s/ Leah Frazier

Leah Frazier, Director

Digital Justice Initiative

Lawyers’ Committee for Civil Rights Under Law

SIGNERS

Asian Americans Advancing Justice (AAJC)

Common Cause

Hispanic Federation

National Action Network

National Consumer Law Center (on behalf of its low-income clients)

The Leadership Conference on Civil and Human Rights

The Legal Defense Fund

United Church of Christ Media Justice Ministry