

THE COMMONWEALTH OF MASSACHUSETTS OFFICE OF THE ATTORNEY GENERAL

ONE ASHBURTON PLACE
BOSTON, MASSACHUSETTS 02108

ANDREA JOY CAMPBELL
ATTORNEY GENERAL

(617) 727-2200
www.mass.gov/ago

July 31, 2026

Via Online Submission

Federal Trade Commission
Office of the Secretary
600 Pennsylvania NW, Suite CC-5610 (Annex B)
Washington, DC 20580

RE: Notice of Proposed Policy Statement Concerning the Suppression of Accuracy in Artificial Intelligence Systems, Docket ID FTC-2026-0859, Matter No. P264200

The Attorney General of Massachusetts, along with the State Attorneys General of Arizona, California, Colorado, Connecticut, Delaware, the District of Columbia, Hawai'i, Illinois, Maine, Maryland, Michigan, Minnesota, Nevada, New Jersey, New Mexico, New York, Oregon, Rhode Island, Vermont, and Washington (the "Attorneys General" or "State Attorneys General") respectfully submit these comments in response to the above-captioned Notice of Proposed Policy Statement Concerning the Suppression of Accuracy in Artificial Intelligence Systems (the "Policy Statement").¹ We write to urge the Federal Trade Commission (the "FTC" or "Commission") not to adopt the noticed Policy Statement, which is based on a flawed premise that the outputs produced by artificial intelligence ("AI") models are inherently neutral and reliable and which appears to misapprehend the relationship between state and federal consumer protection laws.

As the chief consumer protection officials in our respective states, the State Attorneys General share the FTC's concern regarding the accuracy of AI outputs and the representations made by AI companies. AI models play an increasingly consequential role in Americans' professional, personal, and civic lives, with many Americans already relying upon AI models in a variety of settings, from the living room to the board room.² Consumers also utilize other AI technologies, such as predictive AI models, agentic AI, and computer vision applications. As such, it is imperative that any policy statement intended to guide this burgeoning industry be based on an accurate and comprehensive understanding and representation of how AI models work.

¹ Notice of Proposed Policy Statement Concerning the Suppression of Accuracy in Artificial Intelligence Systems (July 1, 2026), published at 91 Fed. Reg. 41638-41642 (July 7, 2026), available at https://www.ftc.gov/system/files/ftc_gov/pdf/ai-policy-statement_0.pdf.

² See Americans and AI 2026: Chatbots, Smart Devices and Views on Impact, Pew Research Center (July 17, 2026), available at https://www.pewresearch.org/wp-content/uploads/sites/20/2026/06/PI_2026.06.17_Americans-and-AI_REPORT.pdf (finding that one in three Americans under the age of 50 report using AI chatbots daily).

The State Attorneys General and the FTC have long shared a common mission to protect the public from unfair or deceptive acts or practices.³ The Policy Statement touches on several pieces of common ground between states and the FTC. For example, the FTC’s proposed policy statement argues that “AI companies have represented explicitly and implicitly that their AI systems aim to produce the best output possible given technological and resource constraints,” and that consumers rely on these representations.⁴ We agree that consumers rely on the public statements made by AI companies regarding their models, and that those companies may be held legally accountable where appropriate. Like the FTC, the State Attorneys General believe that consumers value accurate and reliable outputs from AI models.

However, the Policy Statement makes several erroneous conclusions in both reasoning and law. Namely, the Policy Statement misinterprets the formulation of AI technologies and fails to acknowledge how state and federal laws interact, giving the misimpression that state and federal laws cannot coexist to regulate AI. The State Attorneys General urge the Commission to take this opportunity to course correct and propose a revised policy that is based on the facts and the law.

1. AI Models are Not Inherently Unbiased and Efforts to Remove Unlawful Bias are a Feature, not a Bug

The Policy Statement rests on the premise that consumers expect creators of AI models to build their systems to prioritize accurate, truthful, and ideologically neutral outputs, and that this expectation is reasonable due to “the inherent nature” of AI products and services.⁵ But AI outputs are not inherently neutral. AI companies have always worked to shape the outputs of their models, through selection of training data, pretraining, fine-tuning, and alignment methods such as reinforcement learning. Moreover, AI companies often recalibrate and update their products, both before and after those products go to market.⁶ AI models supply outputs based on the data they are trained upon, the information provided by the users, and their developers’ instructions. Each source of information and direction contains its own ideological or statistical biases, whether intentional or not. No dataset is perfect, and any biases in training data can propagate into AI outputs.⁷

³ See, 15 U.S.C. § 45(a)(1); Mass. Gen. Laws ch. 93A; 940 CMR 3.00; A.R.S. § 44-1522(A); Cal. Bus. & Prof. Code §§ 17200 *et seq.*; Cal. Civ. Code §§ 1770 *et seq.*; Colo. Rev. Stat. § 6-1-105; Conn. Gen. Stat. ch. 735a; D.C. Code §§ 28-3901 *et seq.*; Del. Code Ann. tit. 6, §§ 2511 *et seq.*, 2531 *et seq.*; Haw. Rev. Stat. §480-2; 815 Ill. Comp. Stat. 505; Md. Code Ann., Com. Law §§ 13-101 *et seq.*; Me. Rev. Stat. Ann. tit. 5, ch. 10; Mich. Comp. Laws §§445.901 *et seq.*; Minn. Stat. § 325F.69; Nev. Rev. Stat. §§ 598.0915 *et seq.*; N.J. Stat. Ann. §§ 56:8-1 to -233; N.M. Stat. §§ 57-12-3; N.Y. Gen. Bus. Law §§ 349-350; Or. Rev. Stat. §§ 646.607-608; R.I. Gen. Laws tit. 6, ch. 13.1; Vt. Stat. Ann. tit. 9, § 2453; Wash. Rev. Code §§ 19.86.010 *et seq.*

⁴ Proposed Policy, at 5-6.

⁵ Proposed Policy, at 1.

⁶ Models are constantly deprecated, updated and/or taken out of service. See, e.g., *Model Deprecation*, ANTHROPIC (June 5, 2026), <https://platform.claude.com/docs/en/about-claude/model-deprecations>; *Gemini Deprecations*, GEMINI (July 23, 2026) <https://ai.google.dev/gemini-api/docs/deprecations>; *Deprecations*, OPENAI DEVELOPERS (last visited July 24, 2026), <https://developers.openai.com/api/docs/deprecations>.

⁷ See, e.g., Bjørn Hofmann, *Biases in AI: Acknowledging and Addressing the Inevitable Ethical Issues*, FRONT. DIGIT. HEALTH, Aug. 20, 2025, at 3 (“bias may result in poor or erroneous output from the AI

Responsible businesses will take steps to correct for these biases to improve the accuracy of their products as part of the regular process that all businesses engage in when seeking to create profitable and lawful products.

The Commission’s Policy Statement acknowledges that there are competing objectives that would require an AI company to adjust its model output. As the Policy Statement states, “AI systems are likely to simultaneously pursue multiple objectives, some explicit and some implicit,” including “balance[ing] succinctness, clarity, relevance, accuracy, and other objectives.”⁸ These objectives may conflict with one another at times, and AI companies may weigh these factors when designing models. Each of these factors involves human value judgments that must be implemented by the AI, potentially modifying ultimate outputs.

Some AI model output modifications are made specifically for the purposes of complying with legal obligations or to avoid causing harm. For example, many of the most popular AI companies attempt to prevent their products from generating outputs that could assist in the creation of weapons of mass destruction, even when the model’s training data might allow it to do so.⁹ Others attempt to prevent their models from generating other types of harmful content, such as child sexual abuse material.¹⁰ The federal administration even delayed the release of certain AI models due to the models’ ability to uncover cybersecurity vulnerabilities.¹¹ These deliberate

system”); Emily Baumgaertner Nunn, *This Lifesaving Genetics Tool Works Best if You’re White*, N.Y. TIMES (July 20, 2026), <https://www.nytimes.com/2026/07/20/science/polygenic-risk-scores-diversity.html> (large but racially-homogenous training data led to less accurate results for non-white patients).

⁸ Proposed Policy, at 7.

⁹ See, e.g., *Our Agreement with the Department of War*, OPENAI (Feb. 28, 2026), <https://openai.com/index/our-agreement-with-the-department-of-war/> (OpenAI entered into an agreement with the federal government regarding outputs); *Safety at SpaceXAI*, XAI (last visited July 22, 2026), <https://x.ai/safety> (xAI “integrates safety evaluations throughout the model development lifecycle.”); *Activating AI Safety Level 3 Protections*, ANTHROPIC (May 22, 2025), <https://www.anthropic.com/news/activating-asl3-protections> (Anthropic employs safety standards to reduce the risk that their AI will be used for chemical, biological, radiological, and nuclear weapons); *Guidelines for the Gemini App*, GEMINI (last visited July 22, 2026), <https://gemini.google/policy-guidelines/> (Gemini openly attempts to restrict the types of content its model will generate). The Attorneys General take no position on the accuracy or efficacy of these statements, but merely point out that each has their own approach and bias for ensuring safety in outputs.

¹⁰ See, e.g., *Safety at Every Step*, OPENAI (last visited July 22, 2026), <https://openai.com/safety/> (OpenAI states “[w]e are careful what information we use to teach our AI. We avoid and filter child sexual abuse material (CSAM) and child sexual exploitation material (CSEM), and we report any attempts by users to upload it.”); *Aligning on Child Safety Principles*, ANTHROPIC (Apr. 23, 2024), <https://www.anthropic.com/news/child-safety-principles> (Anthropic “responsibly source[s]” its training data to avoid “ingesting” content that has a known risk of containing CSAM and CSEM); *Policy Guidelines for the Gemini App*, GEMINI (last visited July 22, 2026), <https://gemini.google/policy-guidelines/> (Gemini contains a policy guideline stating “Gemini should not generate outputs, including Child Sexual Abuse Material, that exploit or sexualize children”). Again, the Attorneys General take no position on the accuracy or efficacy of these statements.

¹¹ *Statement on the US Government Directive to Suspend access to Fable 5 and Mythos 5*, ANTHROPIC (Jul. 12, 2026), <https://www.anthropic.com/news/fable-mythos-access>.

manipulations of model outputs are needed to protect the public and for AI systems to comply with state and federal law. And none of these modifications reduce the “accuracy” of the models.

Contrary to the Policy Statement’s assertions, consumers may already reasonably expect AI companies to adjust outputs in an attempt to prevent unfair or deceptive impacts resulting from biases in training data or to comply with anti-discrimination laws. Indeed, an AI hiring tool that automatically screens out qualified candidates based on membership in a protected class—by, for example, rejecting older Americans¹² or disqualifying U.S. citizens¹³—could be deemed unfair or deceptive if offered as a service to lawfully screen job applicants. Landlords or lenders who use algorithmic decision making in screening applicants can be held accountable when those decisions violate state and federal laws.¹⁴ Thus, developers, suppliers, and users of AI all must remain vigilant in ensuring that models and outputs comply with all legal obligations.¹⁵ Each of these legal obligations represents a value judgment made by legislatures through the democratic process, reflecting the values and will of the people. None of these modifications reduces the “accuracy” of AI models or subvert consumer expectations—rather, they make such models *more* fair and balanced.

Additionally, the Policy Statement fails to credibly justify how adjustments by developers to comply with state law could create a deceptive act or practice. The Policy Proposal only points to potential unidentified “ulterior objectives” that a consumer may not “expect” with respect to developers adjusting their models to comply with anti-discrimination laws.¹⁶ The Commission attempts to argue this is sufficient to somehow “abuse consumer trust” and, presumably, mislead.¹⁷ The Policy Statement assumes that AI companies’ compliance with state laws may result in a “loss of accuracy” that these models necessitate, but fails to explain how adjusting models to prevent

¹² See, e.g., Press Release, U.S. Equal Emp. Opportunity Comm’n, iTutorGroup to Pay \$365,000 to Settle EEOC Discriminatory Hiring Suit (Sept. 11, 2023), <https://www.eeoc.gov/newsroom/itutorgroup-pay-365000-settle-eeoc-discriminatory-hiring-suit> (the EEOC settled with a company that used an AI tool to automatically screen out applicants above a certain age).

¹³ Press Release, U.S. Dep’t of Justice, Civil Rights Division Obtains Settlement with a Company that Used AI-Generated Advertisements that Excluded U.S. Workers from Jobs (Feb. 25, 2026), <https://www.justice.gov/opa/pr/civil-rights-division-obtains-settlement-company-used-ai-generated-advertisements-excluded> (the DOJ recently settled with a company which used an AI tool to screen out Citizen-applicants in violation of the Immigration and Nationality Act).

¹⁴ See, e.g., Press Release, U.S. Dep’t of Justice, Justice Department Reaches Proposed Settlement with Willow Bridge, One of America’s Largest Landlords, to Resolve Information Sharing and Algorithmic Coordination Claims (Jul. 6, 2026), <https://www.justice.gov/opa/pr/justice-department-reaches-proposed-settlement-willow-bridge-one-americas-largest-landlords> (the DOJ settled with three large scale landlords who were using AI technology to engage in anti-competitive price manipulation); Assurance of Discontinuance Pursuant to G.L. c. 93A, §5, Civ. Action No. 2584-cv-01895 (Jul. 8, 2025), <https://www.mass.gov/doc/earnest-aod/download> (the Massachusetts Office of the Attorney General settled with a student loan company who was using discriminatory AI technology to screen applicants).

¹⁵ See, generally, OFF. OF THE MASS. ATT’Y GEN., ATTORNEY GENERAL ADVISORY ON THE APPLICATION OF THE COMMONWEALTH’S CONSUMER PROTECTION, CIVIL RIGHTS, AND DATA PRIVACY LAWS TO ARTIFICIAL INTELLIGENCE (Apr. 16, 2024), <https://www.mass.gov/doc/ago-ai-advisory-41624/download>.

¹⁶ Proposed Policy, at 8.

¹⁷ Proposed Policy, at 7. The Commission fails to cite previous actions consistent with this conclusion.

discrimination presents an obstacle to accuracy.^{18,19} But the Commission fails to moor this vague critique in actual existing law or fact.²⁰ Indeed, it is fairer to assume that most consumers “expect” the products and services they purchase to comply with law, not cause them to violate the law.

For all of these reasons, the Policy Statement fails to provide well-founded, actionable guidance to the public. Moreover, it inhibits further AI innovation by muddying the legal landscape for burgeoning companies. Therefore, the Policy Statement should not be adopted.

2. The Commission Cannot Preempt State Law Through Policy Statements.

The Policy Statement argues that state laws governing AI use may be preempted by Section 5 of the FTC Act. Yet, as the Commission concedes, “the FTC Act does not expressly preempt state law,”²¹ nor is the Commission imbued with the authority to declare state law preempted; such a determination is limited to Congress in amending the law or the courts in interpreting the law.²² Since the inception of the FTC Act, states have passed laws that prevent entities from engaging in,

¹⁸ Proposed Policy, at 8.

¹⁹ The Proposed Policy takes the position that inaccurate outputs from AI systems, commonly called “hallucinations,” do not “in and of themselves raise issues under Section 5 or any other law that the Commission enforces” because such hallucinations do not stem “from a design decision to prioritize objectives contrary to users’ reasonable expectations, but from technological and resource limitations that AI systems necessarily reflect”. The States caution against such a sweeping mandate; AI technology advances rapidly, giving rise to the likely possibility that hallucinations may eventually be abatable outputs stemming from developers’ design choices.

²⁰ Colo. Sen. Bill 24-205 (enacted May 17, 2024), *repealed and reenacted, as amended*, by Colo. Sen. Bill 26-189 (enacted May 14, 2026) (While the Commission does repeatedly reference the Colorado Artificial Intelligence Act, as the Proposed Policy acknowledges, this law has been repealed, amended, and reenacted).

²¹ Proposed Policy, at 8.

²² *Louisiana Pub. Serv. Comm’n v. F.C.C.*, 476 U.S. 355, 357 (1986) (“[A] federal agency may pre-empt state law only when and if it is acting within the scope of its congressionally delegated authority. This is true for at least two reasons. First, an agency literally has no power to act, let alone pre-empt the validly enacted legislation of a sovereign State, unless and until Congress confers power upon it. Second, the best way of determining whether Congress intended the regulations of an administrative agency to displace state law is to examine the nature and scope of the authority granted by Congress to the agency.”); *Cipollone v. Liggett Grp., Inc.*, 505 U.S. 504, 517 (1992) (“[E]xpressio unius est exclusio alterius: Congress’ enactment of a provision defining the pre-emptive reach of a statute implies that matters beyond that reach are not pre-empted.”); *Compare* 15 U.S.C. § 1334 (b) (“No requirement or prohibition based on smoking and health shall be imposed under State law with respect to the advertising or promotion of any cigarettes the packages of which are labeled in conformity with the provisions of this chapter.”) *with* 15 U.S.C. § 45 (no explicit state preemption provision). *See also Hillsborough Cnty., Fla. v. Automated Med. Lab’ys, Inc.*, 471 U.S. 707, 717 (1985) (“[w]e are even more reluctant to infer pre-emption from the comprehensiveness of regulations than from the comprehensiveness of statutes. As a result of their specialized functions, agencies normally deal with problems in far more detail than does Congress. To infer pre-emption whenever an agency deals with a problem comprehensively is virtually tantamount to saying that whenever a federal agency decides to step into a field, its regulations will be exclusive. Such a rule, of course, would be inconsistent with the federal-state balance embodied in our Supremacy Clause jurisprudence.”).

encouraging, and facilitating discriminatory impacts and other negative outcomes.²³ These laws have consistently applied to new technologies, including software systems. AI is no different. None of these state laws has created the “impossible” conundrums described by the Policy Statement; instead, like the FTC, states seek to “protect consumers and promote competition.”²⁴

Rather than seeking to preempt the states’ role in consumer protection, the State Attorneys General respectfully urge the FTC to honor our long history of collaboration and the federal-state balance. The states have long served as laboratories of democracy, developing innovative and effective policy solutions that have served as models for sister states and subsequent federal legislation, and the AI governance space is no different. AI companies increasingly hold tremendous power in our society, and governments must use all available tools to steward that power. Congress has yet to act to curb the real and potential harms of AI and the states have acted to fill that void. American consumers will suffer and legitimate state interests will be impaired if the FTC purports to preempt states’ ability to pass legislation.

In conclusion, the Policy Statement misrepresents both the realities of AI technologies and misinterprets applicable law. Further, its lack of concrete, actionable guidelines neither protects consumers nor appraises businesses of their legal obligations. As such, we request the Commission decline to adopt the proposed Policy Statement.

Respectfully Submitted,



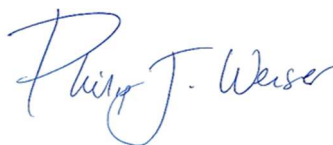
Andrea Joy Campbell
Massachusetts Attorney General



Kris Mayes
Arizona Attorney General



Rob Bonta
California Attorney General



Philip J. Weiser
Colorado Attorney General

²³ See, e.g., Mass. Gen. Laws ch. 93A; 940 CMR 3.00; A.R.S. § 44-1522(A); Cal. Bus. & Prof. Code §§ 17200 *et seq.*; Cal. Civ. Code §§ 1770 *et seq.*; Colo. Rev. Stat. § 6-1-105; Conn. Gen. Stat. ch. 735a; D.C. Code §§ 28-3901 *et seq.*; Del. Code Ann. tit. 6, §§ 2511 *et seq.*, 2531 *et seq.*; Haw. Rev. Stat. §480-2; 815 Ill. Comp. Stat. 505; Md. Code Ann., Com. Law §§ 13-101 *et seq.*; Me. Rev. Stat. Ann. tit. 5, ch. 10; Mich. Comp. Laws §§445.901 *et seq.*; Minn. Stat. § 325F.69; Nev. Rev. Stat. §§ 598.0915 *et seq.*; N.J. Stat. Ann. §§ 56:8-1 to -233; N.M. Stat. §§ 57-12-3; N.Y. Gen. Bus. Law §§ 349-350; Or. Rev. Stat. §§ 646.607-608; R.I. Gen. Laws tit. 6, ch. 13.1; Vt. Stat. Ann. tit. 9, § 2453; Wash. Rev. Code §§ 19.86.010 *et seq.*; Mass. Gen. Laws ch. 112, § 2 (*Mass. Med. Soc. v. Dukakis*, 815 F.2d 790, 791 (1st Cir. 1987)); Mass. Gen. Laws ch. 149, §§ 26-27H (*Marsh v. Mass. Coastal R.R. LLC*, 492 Mass. 641, 642 (2023)).

²⁴ *Our History*, FTC (last accessed July 22, 2026), <https://www.ftc.gov/about-ftc/history>.



William Tong
Connecticut Attorney General



Aaron D. Ford
Nevada Attorney General



Kathy Jennings
Delaware Attorney General



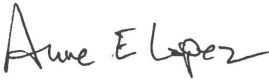
Jennifer Davenport
New Jersey Attorney General



Brian L. Schwalb
Attorney General for the District of Columbia



Raúl Torrez
New Mexico Attorney General



Anne E. Lopez
Hawai'i Attorney General



Letitia James
New York Attorney General



Kwame Raoul
Illinois Attorney General



Dan Rayfield
Oregon Attorney General



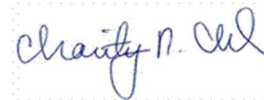
Aaron M. Frey
Maine Attorney General



Peter Neronha
Rhode Island Attorney General



Anthony G. Brown
Maryland Attorney General



Charity R. Clark
Vermont Attorney General



Dana Nessel
Michigan Attorney General



Nick Brown
Washington Attorney General