

**In the Matter of:
THE FEDERAL TRADE
COMMISSION’S PROPOSED POLICY
STATEMENT CONCERNING THE
SUPPRESSION OF ACCURACY IN
ARTIFICIAL INTELLIGENCE
SYSTEMS**

Docket: FTC-2026-0859

**COMMENTS OF
PUBLIC KNOWLEDGE, FIGHT FOR
THE FUTURE, AND THE ELECTRONIC
FRONTIER FOUNDATION**

July 31, 2026

I. INTRODUCTION

Public Knowledge, Fight for the Future, and the Electronic Frontier Foundation (the “Joint Commenters”) appreciate the opportunity to respond to the Commission’s Proposed Policy Statement Concerning the Suppression of Accuracy in Artificial Intelligence Systems.¹

Public Knowledge is a nonprofit organization that promotes freedom of expression, an open internet, and access to affordable communications tools and creative works. Fight for the Future is a nonprofit advocacy organization that works to defend free expression, privacy, and democratic control of technology, and to oppose mass surveillance and discriminatory automated systems. The Electronic Frontier Foundation is the leading nonprofit organization defending civil liberties in the digital world; EFF champions user privacy, free expression, and innovation through impact litigation, policy analysis, grassroots activism, and technology development.

Each of the Joint Commenters has urged this Commission to police genuine deception in technology markets, and each has as consistently opposed government efforts to dictate what private speakers may say. Those commitments are not in tension. They are why we oppose this Proposed Policy Statement and urge the Commission to withdraw it.

We have three primary comments for the Commission. First, the Proposed Policy Statement violates the First Amendment: it would install the Commission as an arbiter of which AI outputs are sufficiently accurate, a viewpoint-based judgment operating as a prior restraint on speech that the proposed disclaimer regime compounds rather than cures. Second, it exceeds the Commission’s authority under Section 5 and cannot preempt state AI laws that Congress has not displaced. Third, its vagueness and its stated provenance establish the conditions for improper jawboning of AI developers and deployers.

¹ *Policy Statement Concerning the Suppression of Accuracy in Artificial Intelligence Systems*, 91 Fed. Reg. 41,638, 41,641 (July 7, 2026) (hereinafter “Proposed Policy Statement”).

II. THE PROPOSED POLICY STATEMENT VIOLATES THE FIRST AMENDMENT AND UNLAWFULLY RESTRICTS FREE EXPRESSION

This section proceeds in five parts. Part A explains how AI systems are actually built, and why no model is an unmediated representation of truth: every system reflects subjective design choices at every stage, and there is no "unaltered" baseline to which the Commission could direct developers to return. Part B establishes that those choices, the outputs they produce, and users' receipt of those outputs are each protected by the First Amendment. Part C explains why the Proposed Policy Statement is a viewpoint-based prior restraint, and why its disclaimer regime compels speech rather than curing the defect. Part D shows that the Commission's theory of deception fails on its own terms. Part E identifies alternative transparency and disclosure measures that are more likely to pass First Amendment scrutiny than the Proposed Policy Statement.

A. AI Systems Are Varied in Function and Design and LLMs Are Not Objective Representations of Truth.

A cursory review of some of the technical features of AI systems, and Large Language Models (LLMs) specifically, will support a more detailed analysis of the Proposed Policy Statement and the inherently flawed concepts implicit in its framing.

First, AI systems vary considerably in their function and design.² While the Commission's statement acknowledges that AI is "an umbrella term covering a universe of different tools and systems," it gives very little consideration to how these differences impact how to apply a concept like "accuracy" in those varied contexts. Accuracy is easy to understand for some applications of AI, such as AI classifiers working in medical diagnosis, where there is labeled data and clear ground truth. The concept becomes more fuzzier as decisions become more subjective, like when AI is used to evaluate job applications. For generative AI models such as LLMs, which appear to be the Commission's focus, it becomes incoherent. Individual facts within an LLM's output may be accurate or inaccurate, but there is no clear way to apply the concept of accuracy to the model's behavior as a whole, just as one cannot say whether a search engine's algorithm is "accurate." Both involve creating methods that generate a response from a near-unlimited list of possible options based on an educated guess about what the user is looking for, where there is no single correct output, nor a single accurate model that can produce one.

To understand this in more detail, it's important to consider what LLMs ultimately represent. LLMs are not objective representations of ground truth about the world. They are representations of statistical patterns in language.³ While the FTC describes AI models as working to "distill all

² See, e.g., IBM, *Types of Artificial Intelligence*, <https://www.ibm.com/think/topics/artificial-intelligence-types> (last visited July 31, 2026); Salesforce, *Types of AI*, <https://www.salesforce.com/artificial-intelligence/types-of-ai/> (last visited July 31, 2026).

³ Cole Stryker, *What are large language models (LLMs)?*, IBM, <https://www.ibm.com/think/topics/large-language-models> (last visited July 31, 2026) ("LLMs work as giant

of human knowledge,” the text corpora LLMs are trained on “reflects only some use of language, and only in a very limited way. Humans use language in the context of the world they live in, and even an exceedingly large text corpus can reflect only part of this use due to the lack of worldly context.”⁴ In the strictest technical sense, “accuracy” for an LLM reflects its ability to match the probability distribution of the language the model was trained on.⁵ That kind of statistical fidelity does not inherently reflect the ground truth state of the world, nor does it necessarily satisfy the FTC’s main but ill-defined performance measure of matching consumer expectations. In fact, pretrained models can exhibit strange and unexpected behaviors, and on their own are not suitable for accomplishing many useful tasks.

In order to be capable of meeting user objectives and performing various practical tasks, a model must go through further processing after its initial training, in a stage known as “post-training.” Post-training includes teaching many of the basic behaviors that make AI systems suitable for different purposes, like a core personality, answering questions, conversational turn-taking, and how to structure different kinds of outputs.⁶ Post-training uses a variety of techniques, including Reinforcement Learning from Human Feedback (RLHF) which involves collecting human preferences about which outputs to prefer.⁷ This process of refinement requires judgement and choice from developers and encodes human values and preferences, which has legal implications discussed further below, but the critical fact to understand is that there is no way to complete post-training and develop a useful AI system *without* making choices—there is no purely factual, objectively truthful stage, perfectly mechanical stage in the LLM “lifecycle.”

B. AI Development, Deployment, and Use All Involve Protected First Amendment Interests.

The Proposed Policy Statement considers an AI system as if it were assessing the calibration of a mechanical measuring instrument that can be clearly right or wrong. As Section II(A) explains, modern AI systems are not simple fact summarizers, and they are the products of a long chain of human judgments including what a model reads in training, what it is optimized toward, and how

statistical prediction machines that repeatedly predict the next word in a sequence. They learn patterns in their text and generate language that follows those patterns”)

⁴ Christoph Durt & Thomas Fuchs, *Large Language Models and the Patterns of Human Language Use*, in *Phenomenologies of the Digital Age* 106, 115 (Marco Cavallaro & Nicolas De Warren eds., 1st ed. 2024), <https://doi.org/10.4324/9781003312284-7>.

⁵ Philip Resnik, *Large Language Models Are Biased Because They Are Large Language Models*, 51 *Computational Linguistics* 885, 888 (2025), https://doi.org/10.1162/coli_a_00558.

⁶ Davide Testuggine, *A Primer on LLM Post-Training*, PyTorch Blog (Aug. 26, 2025), <https://pytorch.org/blog/a-primer-on-llm-post-training/> (“a model coming out of pre-training is often bad at understanding that it should stop talking after a while and will blabber on forever, kind of like a Google autocomplete box.”)

⁷ David Rozado & Philip E. Tetlock, *Assessing Political Bias in AI Systems: A Framework for Disentangling Viewpoint Preferences from Epistemic Integrity*, 55 *Theory & Society*, art. no. 56, at 7 (2026), <https://doi.org/10.1007/s11186-026-09719-6>. (“Even when [human participants in RLHF] are drawn from a broad and demographically representative cross-section of society, the resulting outputs [of the AI system] are still likely to reproduce the prevailing norms, assumptions, and biases of that society at a particular historical moment.”)

it is instructed to behave. Those judgments are expressive. The resulting product is expressive. And the people who use it hold their own constitutional interest as listeners in what it outputs they receive. A policy that proposes to supersede those judgments is a policy that restrains protected speech, and it must be evaluated as such.

This conclusion follows from settled doctrine applied to a new medium, which is what the First Amendment has always demanded. “[W]hatever the challenges of applying the Constitution to ever-advancing technology, ‘the basic principles of freedom of speech and the press, like the First Amendment’s command, do not vary’ when a new and different medium for communication appears.”⁸

The Commission does not dispute that the choices it targets are expressive. The Statement is addressed to nothing else: what a developer trains a model to do, which objectives it prioritizes, and whether it has “steer[ed]” outputs toward “unexpected objectives.”⁹ The Commission has correctly identified the conduct at issue and drawn the wrong conclusion about its constitutional status.

1. AI developers and deployers exercise protected expression and editorial judgment through their AI systems.

Software code is protected speech under the First Amendment.¹⁰ As is the creation and dissemination of information, since facts “are the beginning point for much of the speech that is most essential to advance human knowledge and to conduct human affairs.”¹¹ Courts have also protected algorithmically generated results that reflect a provider’s judgments about relevance and quality.¹² Protection does not require a “narrow, succinctly articulable message.”¹³ And the Supreme Court has already grappled with questions of interactivity when it extended full protection to video games, whose expressive output is generated in response to, and varies with, user inputs.¹⁴

Most directly and recently, in *Moody v. NetChoice* the Supreme Court held that social media platforms engage in protected expression when they “make choices about what third-party speech to display and how to display it,” because they “include and exclude, organize and prioritize—and in making millions of those decisions each day, produce their own distinctive compilations of expression.”¹⁵ The Court grounded that holding in a line of authority running back through *Miami Herald Publishing Co. v. Tornillo* and *Hurley v. Irish-American Gay,*

⁸ *Brown v. Entertainment Merchants Assn.*, 564 U.S. 786, 790 (2011) (quoting *Joseph Burstyn, Inc. v. Wilson*, 343 U.S. 495, 503 [1952]).

⁹ Proposed Policy Statement.

¹⁰ *Junger v. Daley*, 209 F.3d 481, 485 (6th Cir. 2000).

¹¹ *Sorrell v. IMS Health Inc.*, 564 U.S. 552, 570 (2011); see *Bartnicki v. Vopper*, 532 U.S. 514, 527–35 (2001).

¹² *Zhang v. Baidu.com Inc.*, 10 F. Supp. 3d 433, 435–38 (S.D.N.Y. 2014).

¹³ *Hurley*, 515 U.S. at 569.

¹⁴ *Brown*, 564 U.S. at 790–91, 798.

¹⁵ *Moody v. NetChoice, LLC*, 603 U.S. 707, 716 (2024).

Lesbian and Bisexual Group of Boston, and emphasized that “[t]he principle does not change because the curated compilation has gone from the physical to the virtual world.”¹⁶ Curating and arranging expression is protected editorial activity, and that protection extends to entities that make those choices at scale and by algorithm.¹⁷

A series of judgements about what to include and exclude, and that ultimately organizes and prioritizes information, produces “distinctive compilations of expression,” and a law that “profoundly alter[s] the platforms’ choices about the views they will, and will not, convey” cannot survive First Amendment scrutiny.¹⁸

That an AI system is offered as a product changes nothing. Books, films, and video games are products. The question is not whether an item is sold in commerce, but whether the government’s action is reaching for the commercial acts of the offerer or the content of the product.¹⁹ The Commission retains ample authority over AI companies’ data practices, security representations, billing, and marketing claims, but it may not regulate the substance of what their systems say.

2. Users of AI systems hold independent First Amendment interests.

The right to receive information is an independently protected First Amendment right separate from the speaker’s rights, does not turn on the nature of the source, and extends to material the government considers false.²⁰ The government cannot indirectly target disfavored speech by regulating upstream, because pressure applied to an intermediary restricts the public’s access and not merely the intermediary’s.²¹ In addition, users of conversational AI systems are not mere passive recipients. Users can supply custom instructions, individual prompts, and additional source material. They then evaluate the outputs, and may re-prompt to redirect and further engage, with more of their own expression being added to the context for the AI system through this process of interaction.²² As in *Moody*, AI users seek to compose unique expression through this process of curation, discrimination, and judgment. The Proposed Policy Statement would displace their judgment with that of the Commission.

¹⁶ *Id.*; see *Miami Herald Publishing Co. v. Tornillo*, 418 U.S. 241 (1974); *Hurley v. Irish-American Gay, Lesbian & Bisexual Group of Boston, Inc.*, 515 U.S. 557 (1995).

¹⁷ *Id.*

¹⁸ *Moody*, 603 U.S. at 716, 731–32.

¹⁹ See Corbin Barthold, *AI + 1A: Why the First Amendment Protects Artificial Intelligence* (Mar. 23, 2026), <https://ssrn.com/abstract=6496380> (“Even if you call the LLM a ‘product,’ it is a product whose entire function is to emit speech, which we all have a First Amendment right to receive.”).

²⁰ *Martin v. City of Struthers*, 319 U.S. 141 (1943); *Lamont v. Postmaster General*, 381 U.S. 301 (1965); *Stanley v. Georgia*, 394 U.S. 557 (1969); *Virginia State Board of Pharmacy v. Virginia Citizens Consumer Council, Inc.*, 425 U.S. 748 (1976); *United States v. Alvarez*, 567 U.S. 709 (2012).

²¹ *Smith v. California*, 361 U.S. 147, 153–54 (1959).

²² See Inyoung Cheong, *Conversational AI and Human-Centered First Amendment*, 32 MICH. TECH. L. REV. 201 (2026), available at: <https://repository.law.umich.edu/mtlr/vol32/iss2/5>.

C. The Proposed Policy Statement Is a Viewpoint-Based Prior Restraint on Speech.

1. *The First Amendment constrains the government, not private persons.*

The Free Speech Clause “constrains governmental actors and protects private actors.”²³ The Proposed Policy Statement’s operative vocabulary—what is the “best output,” how the outputs of models are “steered,” and what “ideological objectives” developers have—describes private editorial choices. As established above, those choices are protected First Amendment activity. That expression must be protected *from* government interference, it is not the proper *target* of Commission scrutiny.

The Trump Administration has been confusingly inconsistent on this point, however the Commission invokes the White House’s National Policy Framework for Artificial Intelligence which states the proposition that AI technologies “should not be used to silence or censor lawful expression or dissent.”²⁴ The same Framework provides that Americans should be able to seek redress from the Federal Government for agency efforts to censor expression on AI platforms or to dictate the information those platforms provide. Only the second position describes a legally cognizable policy, because only the second constrains *government* conduct.

2. *The government may not install itself as the arbiter of truth.*

The core purpose of the First Amendment is to ensure that the government cannot “prescribe what shall be orthodox in politics, nationalism, religion, or other matters of opinion.”²⁵ An asserted interest in truthful discourse cannot by itself sustain a restriction on speech; holding otherwise would confer “a broad censorial power unprecedented in this Court’s cases or in our constitutional tradition.”²⁶ The Commission’s powers against deception extends to the pursuit of specific false statements tied to concrete commercial injuries. It may not announce a general federal interest in the accuracy of speech and enforce that policy against the people and companies that produce it.

Consider the same policy addressed to newspapers. An agency announces that newspapers have represented, explicitly and implicitly, that they aim to report the truth; that readers reasonably rely on those representations; and that a newspaper subordinating accuracy to an undisclosed ideological objective must either change its editorial practices or carry a prominent and persistent disclaimer of agency-determined sufficiency. Such a policy would be an outrage under our system of protected free expression, and no one would regard that as simple consumer protection. Substituting an AI model for a newspaper masthead does not, and cannot, allow the Commission to upend the constitutional order.

²³ *Manhattan Community Access Corp. v. Halleck*, 587 U.S. 802, 804 (2019).

²⁴ Proposed Policy Statement at 2 (citing The White House, National Policy Framework for Artificial Intelligence [Mar. 20, 2026]).

²⁵ *West Virginia State Board of Education v. Barnette*, 319 U.S. 624, 642 (1943).

²⁶ *Alvarez*, 567 U.S. at 723.

3. *“Accuracy” or “best output” are not content-neutral standards.*

Government “has no power to restrict expression because of its message, its ideas, its subject matter, or its content.”²⁷ Viewpoint discrimination is “an egregious form of content discrimination,” forbidden where “the specific motivating ideology or the opinion or perspective of the speaker is the rationale for the restriction.”²⁸

A federal accuracy standard is content-based on its face, since it cannot be applied without examining what a system said. Even administered without any particular political bias or distinct viewpoint from the government, a blanket accuracy or neutrality standard would be tantamount to forbidding an AI developer or deployer from designing a system in accordance with the wide range of possible views and perspectives that companies, researchers, and private persons may wish to encode in an AI system—that is still viewpoint discrimination under the First Amendment.

However, the Proposed Policy Statement is clear that a neutral application is not even the intention. The Proposed Policy Statement warns of a developer training a model “to correct what the developer believes are ‘historical injustices’ in the facts,” and of a company that might “suppress accuracy and interpose other objectives, such as so-called ‘equity.’”²⁹ This is not the establishment of a politically neutral (though still prohibited) policy of enforced neutrality; the Commission identifies one side of a contested political debate and marks it as disfavored.

4. *The Proposed Policy Statement’s provenance confirms a policy of specific viewpoint discrimination in line with the political objectives of the Trump Administration.*

The Statement is not a product of the Commission’s independent judgment about consumer harm. It exists because Executive Order 14365 directed the Chairman to issue it.³⁰ That Order builds on Executive Order 14319, which enumerates specific ideas the Administration considers objectionable: “critical race theory, transgenderism, unconscious bias, intersectionality, and systemic racism.”³¹ As applied to models the federal government itself procures, that Order may well be permissible; the government may generally decide what speech it endorses and adopts through its procurement. However, extending this into the private marketplace through Section 5 enforcement, an enumeration of disfavored concepts is a blatant example of viewpoint discrimination. Improper speech-suppressive motive can itself establish a First Amendment

²⁷ *Police Department of Chicago v. Mosley*, 408 U.S. 92, 95 (1972).

²⁸ *Rosenberger v. Rector and Visitors of Univ. of Virginia*, 515 U.S. 819, 829 (1995); see *Perry Ed. Assn v. Perry Local Educators’ Assn*, 460 U.S. 37, 46 (1983).

²⁹ Proposed Policy Statement at 7, 8.

³⁰ Exec. Order No. 14365, “Ensuring a National Policy Framework for Artificial Intelligence,” § 7, 90 Fed. Reg. 58499, 58500 (Dec. 11, 2025),

<https://www.whitehouse.gov/presidential-actions/2025/12/eliminating-state-law-obstruction-of-national-artificial-intelligence-policy/>.

³¹ Exec. Order No. 14319, “Preventing Woke AI in the Federal Government,” § 2(b), 90 Fed. Reg. 35389, 35389 (July 23, 2025)

<https://www.whitehouse.gov/presidential-actions/2025/07/preventing-woke-ai-in-the-federal-government/>.

violation,³² and a desire to “exclude the expression of certain points of view from the marketplace of ideas” is “plainly illegitimate.”³³

These are bedrock principles of free expression, and must be vigilantly guarded against a Commission of any composition. A future Commission could invoke this same authority to declare that a model’s treatment of climate science, vaccine safety, or American history constituted a “suppression of accuracy” that deceived consumers, and could demand design changes or disclaimers accordingly. The analysis would be identical, and so would our opposition. Granting a government agency the power to certify some outputs as accurate and others as ideologically distorted and therefore punishable is impermissible under the First Amendment.

5. The government may not intervene to rebalance the marketplace of ideas.

Whether or not AI systems currently reflect a skewed distribution of viewpoints, federal correction is not an available remedy. It is “no job for government to decide what counts as the right balance of private expression—to ‘un-bias’ what it thinks biased, rather than to leave such judgments to speakers and their audiences.”³⁴ Government may not restrict the speech of some “in order to enhance the relative voice of others,”³⁵ and the Court has held that line even against the claim that the marketplace of ideas had become “a monopoly controlled by the owners of the market.”³⁶ Concentration in the AI industry is a genuine concern and a reason for competition enforcement that is well within the Commission’s powers, but it is not a license for a policy statement of this kind that directly targets protected speech.

6. The Proposed Policy Statement would operate as a prior restraint on speech and the proposed disclaimer regime does not cure this and is impermissible compelled speech.

a. The Statement restrains speech in advance rather than penalizing it after the fact.

Prior restraints bear “a heavy presumption against [their] constitutional validity” and are “the most serious and the least tolerable infringement on First Amendment rights.”³⁷ A restraint operating before speech occurs suppresses it without the safeguards of an adversarial proceeding, and the expression it deters never surfaces to be defended.³⁸ The Supreme Court has held that nonbinding commission determinations—like this Proposed Policy Statement—can have the effect of a prior restraint on speech even where they lack the force of law or do not set up a

³² *Heffernan v. City of Paterson*, 578 U.S. 266, 273–74 (2016).

³³ *City Council v. Taxpayers for Vincent*, 466 U.S. 789, 804 (1984).

³⁴ *Moody*, 603 U.S. at 741.

³⁵ *Buckley v. Valeo*, 424 U.S. 1, 48–49 (1976) (per curiam).

³⁶ *Tornillo*, 418 U.S. at 251, 258.

³⁷ *Bantam Books, Inc. v. Sullivan*, 372 U.S. 58, 70 (1963); *Nebraska Press Ass’n v. Stuart*, 427 U.S. 539, 559 (1976).

³⁸ *Near v. Minnesota*, 283 U.S. 697 (1931); *Freedman v. Maryland*, 380 U.S. 51, 58–60 (1965).

formal system of censorship.³⁹ The Commission's "threat of invoking legal sanctions and other means of coercion, persuasion, and intimidation" is sufficient.⁴⁰

b. The proposed disclaimer regime is not a safety valve, but rather operates with the prior restraint to unconstitutionally compel speech from AI developers and deployers.

The Proposed Policy Statement—like the “Preventing Woke” Executive Order that preceded it—appears to offer an escape valve, perhaps recognizing both the clear invalidity under the First Amendment and the near impossibility of compliance with legal standards that require an objective evaluation of truthfulness or accuracy. The Proposed Policy Statement proposes that a clear and conspicuous disclosure about the “ideological objectives” that the system prioritizes could cure consumer confusion and therefore remove liability for deception.⁴¹ The Proposed Policy Statement specifies no content, form, or threshold of sufficiency, requiring only that it be "prominent," that it not be buried in fine print, and that it grow more "persistent and prominent" the further it "cuts against the reasonable expectations that users would take away from other contexts."

This alternative is cold comfort. One who must establish the lawfulness of his own expression under an unascertainable standard "must steer far wider of the unlawful zone than if the State must bear these burdens."⁴² Where First Amendment freedoms are at stake, "[p]recision of regulation must be the touchstone."⁴³ The absence of precision here is not plausibly an oversight. A standard no developer can satisfy with confidence yields the risk-averse, over-compliance the Commission would otherwise have to obtain by rule, with a record, subject to challenge.

Additionally, what the disclosure compels is not fact but a profession of belief—an announcement that the company has subordinated accuracy to an ideological aim. Compelled affirmations of belief occupy some of the least defensible ground in First Amendment law, whether the affirmation demanded is a flag salute, a state motto, a religious test, or a loyalty oath whose imprecision leaves the speaker guessing at his obligations.⁴⁴

There are clear exceptions for some commercial speech, but *Zauderer* generally permits compelled disclosure of "purely factual and uncontroversial information," and *NIFLA* confirms the exception goes no further.⁴⁵ As discussed further in sections below, there may be disclosures that are both compellable and helpful to users that fall into that category—a model's training

³⁹ *Bantam Books*, 372 U.S. at 67

⁴⁰ *Id.*

⁴¹ Proposed Policy Statement at 9.

⁴² *Speiser v. Randall*, 357 U.S. 513, 526 (1958).

⁴³ *NAACP v. Button*, 371 U.S. 415, 438 (1963).

⁴⁴ *Barnette*, 319 U.S. at 642; *Wooley v. Maynard*, 430 U.S. 705, 714–15 (1977); *Torcaso v. Watkins*, 367 U.S. 488, 495–96 (1961); *Baggett v. Bullitt*, 377 U.S. 360, 372 (1964); *Keyishian v. Board of Regents*, 385 U.S. 589, 603–04 (1967).

⁴⁵ *Zauderer v. Office of Disciplinary Counsel*, 471 U.S. 626, 651 (1985); *National Inst. of Family & Life Advocates v. Becerra*, 585 U.S. 755, 768 (2018).

sources, system prompt, or benchmark results might potentially be factual in that sense—but that is not the proposal the Commission has advanced.

D. The Commission’s Theory of Deception Is Incoherent.

Section III below further explains why Section 5 creates no conflict with state law, and why the Statement’s account of generalized consumer expectations is unsupported by any record. Three points bear noting here, as related to the First Amendment analysis.

1. The Proposed Policy Statement targets the wrong object and inverts the remedy.

Section 5 reaches representations about products, not the content of expressive products. The Commission’s own AI enforcement record honors that line. It has proceeded against companies for misrepresenting the accuracy of an AI content-detection tool, the assertion that a chatbot can replace a lawyer, the efficacy of AI-powered facial recognition and weapons screening, and the earnings claims made for a conversational AI product.⁴⁶ Each concerned a claim made to consumers about what a product could do. That is the correct scope of Section 5 authority as applied to AI, and it is meaningful consumer protection work.

The Proposed Policy Statement inverts the ordinary course of enforcement. The ordinary rule is that a company may not describe its product inaccurately, and the cure for a false claim is to stop making that false claim. Here, the Commission proposes a policy that a company which has described its product in a particular way must thereafter build the product to the Commission’s interpretation of that specification or else publish a Commission-approved disclaimer. That converts a truthful advertising requirement into a design mandate over expressive output.

2. The Proposed Policy Statement’s own citations on consumer expectations refute its premise.

The Proposed Policy Statement asserts that consumers “have no basis to believe that AI systems aim to produce outputs that are distorted by undisclosed ideological objectives,” and that the contrary expectation is “eminently reasonable.”⁴⁷ It cites no survey, no study of consumer understanding, and no complaint data. The omission is conspicuous because the Statement recites, pages earlier, that the Commission’s deception authority “requires the entity to substantiate the representation it makes about its products or services with sufficient evidence.”⁴⁸

The Proposed Policy Statement’s footnotes then supply its own refutation. It quotes xAI marketing Grok as a “truth-seeking AI companion for unfiltered answers,” and Anthropic

⁴⁶ Consent Order, In re Workado, LLC, No. C-4822 (F.T.C. Aug. 21, 2025); Consent Order, In re DoNotPay, Inc., FTC Docket No. C-4812 (F.T.C. Jan. 14, 2025); Consent Order, In re Intellivision Techs. Corp., No. C-4809 (F.T.C. Jan. 8, 2025); Stipulated Order, FTC v. Evolv Techs. Holdings, Inc., No. 1:24-cv-12940-PGL (D. Mass. Nov. 26, 2024); Compl., FTC v. Air Ai Techs., Inc., No. 2:25-cv-03068-SMB (D. Ariz. Aug. 25, 2025).

⁴⁷ Proposed Policy Statement at 1, 7.

⁴⁸ *Id.* at 4; cf. *Skidmore v. Swift & Co.*, 323 U.S. 134, 140 (1944) (an agency pronouncement’s weight depends on “the thoroughness evident in its consideration” and “the validity of its reasoning”).

marketing Claude as “trained . . . to be safe, accurate, and secure.”⁴⁹ Those are not the same promise. Developers compete on precisely these differences and advertise them to consumers, who select among them accordingly. The Commission has assembled the evidence that these systems carry distinguishable and openly declared commitments, and then concluded that consumers could have no basis to suspect as much.

3. Even demonstrated deception would not authorize the remedy contemplated.

If an AI developer genuinely misled consumers about how its system was designed, the government may be able compel commercial disclosure of “purely factual and uncontroversial information,” and no further.⁵⁰ Compelled statements that require a regulated party to characterize contested matters receive heightened scrutiny—requiring a company “to publicly condemn itself” is at least equally if not more constitutionally offensive than a prohibition on particular speech.⁵¹ What the Proposed Policy Statement contemplates is not a factual disclosure about a product’s characteristics or capabilities: it is a company’s account of its own ideological objectives, at a prominence and frequency the Commission will evaluate after the fact. The First Amendment protects “both the right to speak freely and the right to refrain from speaking at all,”⁵² and the government may not “coopt an individual’s voice for its own purposes.”⁵³

This is not an argument against transparency. As stated above and as Section II(E) further explains, we support meaningful disclosure obligations for AI systems. The distinction is between disclosing facts—perhaps things like what data a system was trained on, what its already written system prompt instructs, how it performed against a published peer-reviewed benchmark—and compelling a company to characterize its own values as somehow in opposition to “accuracy.” The first informs consumers, the latter hijacks the developer’s sacrosanct freedom of expression.

E. Transparency and Disclosure Are Important Goals That Can Be Compatible With the First Amendment.

While the FTC has no business establishing a policy to judge and control the outputs of AI systems, there is a legitimate consumer protection interest in transparency. Better public understanding of how AI systems function and what goals they prioritize can help consumers be better informed and promote a healthy marketplace founded on user and deployer choice and a more vital ecosystem for AI development and deployment overall. Auditing, evaluation, and transparency requirements can implicate the First Amendment by compelling speech, but there are versions of these policies that will likely pass constitutional scrutiny.

⁴⁹ Proposed Policy Statement at 6 fn.35.

⁵⁰ *Zauderer*, 471 U.S. at 651; *NIFLA*, 585 U.S. at 768–69.

⁵¹ *National Assn. of Manufacturers v. SEC*, 800 F.3d 518 (D.C. Cir. 2015).

⁵² *Wooley v. Maynard*, 430 U.S. 705, 714 (1977); see *Barnette*, 319 U.S. at 642.

⁵³ *303 Creative LLC v. Elenis*, 600 U.S. 570, 592 (2023).

The standard of scrutiny for disclosure requirements ultimately depends on the nature of the speech and what is being disclosed.⁵⁴ Technology transparency requirements may fall into one of three categories here: lenient rational-basis review for disclosures that are part of a broader regulatory system not principally concerned with speech, the more permissive *Zauderer* standard for factual commercial disclosures, and heightened scrutiny for mandates that are “inextricably intertwined” with and burden protected expression.⁵⁵

Some of this doctrine remains unsettled, as recent cases have raised questions about the scope of *Zauderer*’s precise scope and requirements.⁵⁶ It may not be the appropriate test for certain mandated AI disclosures. But it may be possible to create transparency requirements that pass constitutional muster by relying on purely factual disclosures that are designed for consumer understanding, tied closely to the product or service at issue, and not contingent on developers adopting the government’s framing of a contested judgment.⁵⁷ Several proposals in this realm may be constitutional, and while none of these are proposed by the Commission, we offer these as counterpoints to demonstrate that consumer and user needs can be met without onerous curtailment of speech.

One simple option is to require “model cards,” which function similar to nutrition labels, disclosing specified technical information about a model in a clear understandable format. This solution has already been adopted by industry to a fair degree, although the design and content of model cards is not standardized among those who have adopted them.⁵⁸ If appropriately tailored, a model card requirement may be a constitutional alternative to meet the need of clear information for consumers.⁵⁹

Another possibly valid approach, though slightly more ambitious, is to require disclosure of commercial AI models’ system prompts. This would go directly to the heart of the Commission’s

⁵⁴ See generally, Valerie C. Brannon, Cong. Rsch. Serv., R45700, Assessing Commercial Disclosure Requirements Under the First Amendment (2019), <https://www.congress.gov/crs-product/R45700>. (reviewing the status of commercial disclosure requirements under the First Amendment)

⁵⁵ See Becca Branum, *First Amendment Tech Transparency Roadmap*, Center for Democracy & Technology (Feb. 13, 2025), <https://cdt.org/insights/first-amendment-tech-transparency-roadmap/>.

⁵⁶ See Bahrad A. Sokhansanj & Mackenzie Arnold, *First Amendment Questions for AI Transparency Laws*, Lawfare (June 22, 2026) <https://www.lawfaremedia.org/article/first-amendment-questions-for-ai-transparency-laws> (“AI transparency laws... are arriving at a moment when at least some courts appear to be rethinking how to assess mandated corporate disclosure.”)

⁵⁷ *Id.* (“[T]raining data transparency laws may be on firmer footing when they call for disclosure of facts that help consumers evaluate a model, such as its training data sources. These laws are harder to defend when they call for interpretive judgments instead, such as assessments of possible bias.”)

⁵⁸ See Weixin Liang et al., *What’s Documented in AI? Systematic Analysis of 32K AI Model Cards*, arXiv (2024), <https://arxiv.org/abs/2402.05160>. (“Most AI models with a substantial number of downloads provide model cards, although with uneven informativeness.”)

⁵⁹ See Matt Perault, *AI Model Facts: Transparency that Works for Little Tech*, Andreessen Horowitz (Apr. 22, 2025), <https://a16z.com/ai-model-facts-transparency-that-works-for-little-tech/>.

concern about developers steering their outputs to prioritize certain objectives: that is what a system prompt does by definition.⁶⁰ The Administration itself has recognized the value of system prompt transparency, with the Woke AI Executive order including an option to “permit vendors to comply with the [order’s] requirement... to be transparent about ideological judgments through disclosure of the LLM’s system prompt.”⁶¹ There is also strong public support for this kind of transparency, with 89% of participants in a 2026 study expressing support for some level of system prompt transparency.⁶² Critically for First Amendment purposes, the information in a system prompt already exists. Disclosure does not require a developer to take sides in a contested debate, swear to the views it holds, or create any additional speech, and may be less likely to burden protected expression, thereby mitigating some of the key concerns underlying the compelled speech doctrine. The mandatory disclosure of the system prompt may therefore survive constitutional scrutiny, even if the words and content of the prompt itself is protected speech.⁶³ How these instructions ultimately function to shape model behavior may bring some controversy, but their content is unambiguous.⁶⁴

Another transparency measure may be to require disclosure of performance on an industry-accepted standard suite of benchmarks and evaluations. Benchmarks may be a suitable required component of model cards or other disclosure regimes once they are well-established as empirically grounded and accessible.⁶⁵ The current state of research does not appear to have reached this point, as current meta-analysis of existing benchmark methods reveals several persistent flaws around how much benchmarks actually measure their intended targets and how much their results are shaped by factors unrelated to model performance.⁶⁶ These problems

⁶⁰ See Sunil Ramlochan, *System Prompts in Large Language Models*, Prompt Engineering Institute (Mar. 8, 2024), <https://promptengineering.org/system-prompts-in-large-language-models/>. (“At their core, system prompts are a set of instructions, guidelines, and contextual information provided to AI models before they engage with user queries. These prompts act as a framework, setting the stage for the AI to operate within specific parameters and generate responses that are coherent, relevant, and aligned with the desired outcome.”)

⁶¹ E.O. 14319, Preventing Woke AI in the Federal Government (July 23, 2025)

⁶² Anna Neumann et al., *Who Controls the Conversation? User Perspectives on Generative AI (LLM) System Prompts*, in Proceedings of the 2026 CHI Conference on Human Factors in Computing Systems, art. no. 378, at 12 (2026), <https://doi.org/10.1145/3772318.3791726>.

⁶³ *Zauderer*, 471 U.S. 626, 651.

⁶⁴ See Anna Neumann et al., *Prompt Governance? On Governing Technologies Governed by Natural Language*, in Proceedings of the 2026 ACM Conference on Fairness, Accountability, and Transparency 6466 (2026), <https://doi.org/10.1145/3805689.3806763>. (observing “a fragmented research landscape advancing divergent and sometimes contradictory claims about what system-level instructions can achieve, under what conditions, and with what reliability”)

⁶⁵ Matt Perault, *AI Model Facts: Transparency that Works for Little Tech*, Andreessen Horowitz (Apr. 22, 2025), <https://a16z.com/ai-model-facts-transparency-that-works-for-little-tech/>.

⁶⁶ See, e.g., Andrew M. Bean et al., *Measuring What Matters: Construct Validity in Large Language Model Benchmarks*, arXiv, at 10 (2025), <https://arxiv.org/abs/2511.04703>.

(“We found that the operationalisation of abstract phenomena was often insufficient, with definitions being missing or contested. Tasks were frequently taken from pre-existing data sources without adjustments to ensure that they were representative of the target phenomenon. Statistical testing was also rarely performed. About half of the reviewed articles did discuss the validity of their benchmark, but nearly every paper had weaknesses in at least one area”); Maria Eriksson et al., *Can We Trust AI Benchmarks? An Interdisciplinary Review of Current Issues in AI Evaluation*, arXiv (2025), <https://arxiv.org/abs/2502.06559>.

suggest that First Amendment concerns would likely be implicated if they were made into requirements today. However, efforts in the field are continuing, with recent research working to elicit what makes a good benchmark⁶⁷ and standardize evaluation reporting to enable cross comparisons.⁶⁸ As the field develops, these methods may eventually produce enough consensus and reliability for certain commonly accepted benchmarks to be included in transparency mandates. One place where benchmarks may function well in the meantime is by providing substantiation for claims that developers already voluntarily make about their models, rather than standing as independent disclosure requirements. For example, researchers have proposed reframing model cards as structured “safety artifacts” that make a model’s safety claims and supporting evidence explicit, allowing benchmark and evaluation results to be re-evaluated to ensure the evidence continues to hold.⁶⁹ Benchmarks used this way would not be compelling new speech from developers, but rather helping to substantiate the claims they already make in their advertising, mirroring the best-faith interpretation of the core concern expressed by the Commission in the Proposed Policy Statement.

III. THE PROPOSED POLICY STATEMENT EXCEEDS THE COMMISSIONS LEGAL AUTHORITIES AND DOES NOT AND CANNOT PREEMPT STATE AI LAWS

The Proposed Policy Statement’s actual preemption analysis is cursory: “Although the FTC Act does not expressly preempt State law, State law is impliedly preempted to the extent it conflicts with a Federal regulatory scheme.” It then adds: “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.”⁷⁰ The Commission appears to hope that the word “obviously” will stand in for legal analysis. It does not. It has identified no state law that requires deception, and no source of legal authority for the Commission to act in this area.

The Joint Commenters do not dispute the obvious proposition that compliance with state law is not a defense to a violation of federal law. The question is whether there is any relevant federal law to begin with. There is not. The Administration’s preference for a “minimally burdensome” national AI policy is no substitute for Congressional action.⁷¹

(surveying technical, systemic, and ideological problems shown in AI benchmark papers over the past ten years, concluding “the field of quantitative benchmarking is currently ill-suited to single-handedly (or primarily) provide the safety and capability assurances requested by policy makers”).

⁶⁷ Anka Reuel et al., *BetterBench: Assessing AI Benchmarks, Uncovering Issues, and Establishing Best Practices*, Stanford HAI Policy Brief (Dec. 11, 2024), <https://hai.stanford.edu/policy/what-makes-a-good-ai-benchmark> (introducing “a novel assessment framework for evaluating AI benchmarks based on 46 criteria across five benchmark life-cycle phases” developed through benchmark literature analysis and expert interviews)

⁶⁸ Ruchira Dhar et al., *EvalCards: A Framework for Standardized Evaluation Reporting* (2025), <https://arxiv.org/abs/2511.21695>.

⁶⁹ Calvin Thuan-Phong Khuu et al., *Model Cards for Responsible AI: Stop Carding, Start Modelling*, ICSE-NIER '26 (2026), <https://doi.org/10.1145/3786582.3786804>.

⁷⁰ Proposed Policy Statement at 9.

⁷¹ Exec. Order No. 14,365, §§ 1–2, 90 Fed. Reg. 58,499, 58,499 (Dec. 16, 2025).

A. Policy Statements Do Not Have the Force of Law.

Section 18 of the FTC Act distinguishes “interpretive rules and general statements of policy” regarding unfair or deceptive acts or practices from rules that “define with specificity” the practices that are unlawful.⁷² While a policy statement may tell the public how the Commission currently understands the law and expects to exercise enforcement discretion, it does not and cannot by itself impose any legal duties or preempt state laws.

The Supreme Court has repeatedly held that interpretations contained in “policy statements, agency manuals, and enforcement guidelines” lack the force of law.⁷³ They may be persuasive, depending on the quality of their reasoning, but that is all. Further, the agency’s position carries weight only according to its “thoroughness, consistency, and persuasiveness.”⁷⁴ After *Loper Bright*, moreover, a court must exercise its own judgment about the meaning of Section 5. Statutory ambiguity by itself does not constitute a delegation of authority to the agency.⁷⁵

The Commission may announce that it may bring a Section 5 case against a company that misleads consumers about how an AI product works, whether or not that company claims compliance with state law as a defense. It may then test that theory in court. What it cannot do is use a policy statement to determine in advance that a state law *is* preempted—as opposed to stating its non-binding view that it should be.

B. Section 5 Does Not Create the Conflict the Policy Statement Imagines.

Conflict preemption exists “when it is impossible to comply with both state and federal law” or “where the state law stands as an obstacle to the accomplishment of the full purposes and objectives of Congress.”⁷⁶ A policy disagreement between federal and state policymakers does not meet this test.

Section 5 cannot be reasonably read as dictating the kinds of substantive answer an AI system must produce. It prohibits an unfair or deceptive “act or practice” in commerce.⁷⁷ Under the Commission’s deception standard, that requires a material representation, omission, or practice likely to mislead a consumer acting reasonably under the circumstances.⁷⁸

⁷² 15 U.S.C. § 57a(a)(1)(A)–(B). Section 18 also provides that the Commission has no other authority under the FTC Act to prescribe rules concerning unfair or deceptive acts or practices. *Id.* § 57a(a)(2).

⁷³ *Christensen v. Harris County*, 529 U.S. 576, 587 (2000) (denying controlling weight to an agency opinion letter); see also *United States v. Mead Corp.*, 533 U.S. 218, 234–35 (2001).

⁷⁴ *Wyeth v. Levine*, 555 U.S. 555, 576–77 (2009) (giving the FDA’s preemption position weight only according to its thoroughness, consistency, and persuasiveness).

⁷⁵ *Loper Bright Enters. v. Raimondo*, 603 U.S. 369, 412–13 (2024).

⁷⁶ *Schneidewind v. ANR Pipeline Co.*, 485 U.S. 293, 300 (1988); see Proposed Policy Statement at 9 fn.49; *English v. Gen. Elec. Co.*, 496 U.S. 72, 79 (1990).

⁷⁷ 15 U.S.C. § 45(a)(1).

⁷⁸ *FTC Policy Statement on Deception*, appended to *Cliffdale Assocs., Inc.*, 103 F.T.C. 110, 174, 175–76, 182–83 (1984); see *POM Wonderful, LLC v. FTC*, 777 F.3d 478, 490 (D.C. Cir. 2015).

The Policy Statement instead depends on unsupported theorizing about generalized consumer expectations. It says AI companies have represented that their systems aim to produce the “best output possible.” It then warns that steering a system toward an unexpected output is likely deceptive unless a sufficiently prominent and persistent disclosure resets consumer expectations.⁷⁹

This is vague and unspecific and not responsive to the reality of how AI systems are deployed and designed to meet a mix of objectives. Developers determine what a product is for, and choose training data, fine-tune their models, write system instructions, and establish safety rules in line with those goals. A medical-support tool, an employment-screening system, and an educational chatbot should not behave identically. There is not necessarily a “best” answer or output, and even different users of the same system may have different expectations.

Certainly, if a company says its product does one thing while secretly designing it to do another, ordinary deception law may apply. But Section 5 does not establish the Commission’s preferred hierarchy among design goals. Nor does any record establish that consumers expect every AI system to follow the Commission’s priorities.

The Colorado law invoked by the Policy Statement is instructive. Beginning January 1, 2027, the law will cover automated decision-making technology used to materially influence consequential decisions in fields such as education, employment, housing, lending, insurance, health care, and public benefits. It expressly excludes AI systems not intended for these kinds of uses. The law requires developers to provide specified documentation and deployers to give notice to affected consumers. And its liability provision allocates fault in actions alleging unlawful discrimination under existing state law, with developer liability limited to uses the developer intended, documented, marketed, advertised, configured, or contracted for.⁸⁰ Describing this as a law that requires AI companies to “suppress accuracy” is itself inaccurate.

The same reasoning applies to other state AI requirements the Commission’s and the Executive Order’s theory could reach. Transparency rules do not require deception; they require that AI companies provide additional (accurate) information. Testing and evaluation can reveal whether a product works as represented. Audits can uncover hidden errors or discriminatory effects. Rules addressing consequential decisions regulate how a tool is used to allocate housing, employment, credit, health care and information, and similar important matters. State laws vary greatly and any given measure may be wise or unwise, well drafted or clumsy, and some applications may raise constitutional questions. But they do not conflict with Section 5.

⁷⁹ Proposed Policy Statement at 9.

⁸⁰ S.B. 26-189, ch. 131, §§ 1, 5, 2026 Colo. Sess. Laws 569, 570–82, 585 (to be codified at Colo. Rev. Stat. §§ 6-1-1701 to -1709), <https://leg.colorado.gov/laws/session-laws/SB26-189/131/download>. See *id.* at 570–71 (scope and exclusions), 575–80 (developer documentation and consumer notices), 581–82 (allocation of fault and limits on developer liability), 585 (no new private right of action, effective date, and applicability).

The Commission's theory also confuses the outcome of an AI system, with misrepresentation. A system may be inaccurate without having been deceptively marketed. It may be designed to reduce discriminatory outcomes, and fail at that task, without misleading anyone. It may be a chatbot designed to be entertaining, and fail at that, too. And a system may produce an answer the Commission considers "truthful" while still being sold through false claims. Further, some statements mix fact with opinion or context making the question of "misrepresentation" all the more challenging.⁸¹ This highlights the disconnect between the Commission's stated policy goals and the legal authority it invokes.

It bears emphasizing that an answer the Commission labels "accurate" is not necessarily true. For example, Joe Biden won the 2020 election.⁸² The Trump White House nevertheless says that former cybersecurity official Chris Krebs "falsely and baselessly denied that the 2020 election was rigged and stolen."⁸³ Krebs was right; the White House is not. No one answerable to President Trump, however, appears capable of conceding this. Recently, Trump's nominee for Director of National Intelligence, Jay Clayton, refused to answer when Senator Jon Ossoff asked who won. Ossoff finally asked, "Isn't it humiliating to be unable to answer this question, to have to indulge the president's delusions?"⁸⁴ It is difficult to see how this Administration and its appointees could be counted on as neutral arbiters of what is true, and what is not.

Nor is that the only example of this Administration's loose relationship to the truth. After reviewing 31 primary studies and five meta-analyses, the World Health Organization's vaccine-safety committee concluded that the best available scientific evidence shows that vaccines do not cause autism. But the Administration's CDC now rejects that conclusion: it says that "vaccines do not cause autism" is "not an evidence-based claim."⁸⁵

Putting this Administration—or any administration—in charge of a federal truth standard would be laughable if the stakes were not so serious. An official label is no truth criterion, and Section 5 cannot enforce the Administration's preferred version of reality.

⁸¹ See *Milkovich v. Lorain Journal Co.*, 497 U.S. 1, 19–21 (1990); *Gertz v. Robert Welch, Inc.*, 418 U.S. 323, 339–40 (1974).

⁸² *2020 Electoral College Results*, Nat'l Archives, <https://www.archives.gov/electoral-college/2020> (last visited July 23, 2026).

⁸³ *Memorandum on Addressing Risks from Chris Krebs and Government Censorship*, 2025 Daily Comp. Pres. Doc. 465, at 1 (Apr. 9, 2025).

⁸⁴ *Open Hearing: Nomination for the Honorable Walter "Jay" Clayton III To Be Director of National Intelligence*, S. Select Comm. on Intelligence (July 15, 2026), <https://www.intelligence.senate.gov/2026/06/10/closed-briefing-intelligence-matters-74/>; Bill Barrow, *Sen. Ossoff Fuels Reelection Campaign with Attacks on Trump*, Associated Press (July 18, 2026), <https://apnews.com/article/b35cc90ec52f00194333aa47e093e8bd>.

⁸⁵ *Statement of the WHO Global Advisory Committee on Vaccine Safety (GACVS) on Vaccines and Autism*, World Health Org. (Dec. 11, 2025), <https://www.who.int/news/item/11-12-2025-statement-gacvs-vaccines-autism> (reviewing 31 primary studies and five meta-analyses and concluding that the available high-quality evidence indicates that vaccines do not cause autism); *Autism and Vaccines*, Ctrs. for Disease Control & Prevention (July 22, 2026), <https://www.cdc.gov/vaccine-safety/about/autism.html>.

C. The FTC Has No Delegated Power To Turn Section 5 Into a General Override of State AI Law.

Even a binding federal regulation can preempt state law only when the agency acts within authority Congress delegated to it. As the Supreme Court put it, “an agency literally has no power to act, let alone pre-empt the validly enacted legislation of a sovereign State,”⁸⁶ unless Congress has conferred the relevant authority. Valid federal regulations can have preemptive effect, but *Fidelity Federal* and *Geier*, the leading cases on which the Commission might rely, involved binding regulations issued under clear delegated authority.⁸⁷ They do not give nonbinding guidance the same status.

The FCC’s experience with its Internet Policy Statement is a close precedent to this matter. In 2004, FCC Chairman Michael Powell articulated four “Internet Freedoms.” The following year, after Kevin Martin became Chairman, the Commission incorporated those principles into a formal policy statement. The FCC said the document offered “guidance and insight” into its approach and cautioned that it was “not adopting rules in this policy statement.”⁸⁸

Three years later, the Martin Commission tried to enforce those principles against Comcast. Acting on a complaint filed by Free Press and Public Knowledge, the FCC ruled that Comcast’s interference with peer-to-peer applications contravened federal policy, and ordered Comcast to make disclosures and submit to compliance requirements. The D.C. Circuit vacated the order because the FCC had not grounded its action in a statutorily mandated responsibility. As the court found, “Policy statements are just that—statements of policy. They are not delegations of regulatory authority.”⁸⁹ An agency cannot turn guidance of this kind into an enforceable obligation, much less a basis for preempting state law.

Schneidewind, the only case cited in the statement’s preemption discussion, does not help the Commission.⁹⁰ The case concerned a Michigan law that required interstate natural-gas pipelines to obtain state approval before issuing long-term securities. The Supreme Court held that the law was preempted, not because FERC had published a policy statement, but because the Natural Gas Act gave FERC comprehensive control over interstate wholesale gas rates and facilities,

⁸⁶ *La. Pub. Serv. Comm’n v. FCC*, 476 U.S. 355, 374 (1986).

⁸⁷ See *Fid. Fed. Sav. & Loan Ass’n v. de la Cuesta*, 458 U.S. 141, 153–59 (1982); *Geier v. Am. Honda Motor Co.*, 529 U.S. 861, 874–86 (2000). See generally David S. Rubenstein, *The Paradox of Administrative Preemption*, 38 Harv. J.L. & Pub. Pol’y 263, 276–83 (2015) (distinguishing a regulation with preemptive effect from an agency’s interpretation of a statute’s preemptive scope).

⁸⁸ Michael K. Powell, Chairman, FCC, *Preserving Internet Freedom: Guiding Principles for the Industry* 5–6 (remarks at the Silicon Flatirons Symposium, Feb. 8, 2004), <https://docs.fcc.gov/public/attachments/DOC-243556A1.pdf>; *Appropriate Framework for Broadband Access to the Internet over Wireline Facilities*, 20 FCC Rcd. 14,986, 14,987–88 ¶¶ 3–5 & n.15 (2005).

⁸⁹ *Formal Complaint of Free Press & Public Knowledge Against Comcast Corp.*, 23 FCC Rcd. 13,028, 13,052–60 ¶¶ 43–55 (2008); *Comcast Corp. v. FCC*, 600 F.3d 642, 654, 661 (D.C. Cir. 2010).

⁹⁰ Proposed Policy Statement at 9 fn.49.

including tools for examining pipeline financing.⁹¹ Section 5 looks nothing like that detailed regulatory scheme. It is a generally applicable prohibition on unfair or deceptive commercial practices. It does not assign the FTC comprehensive authority over AI design.

Congress knows how to give an agency preemptive authority, and it knows how to create a national regulatory scheme. It did neither here.

D. The Proposed Policy Statement Fails the Major-Questions Test.

The major-questions doctrine requires clear congressional authorization when an agency claims exceptionally consequential power: “We expect Congress to speak clearly if it wishes to assign to an agency decisions of vast ‘economic and political significance.’”⁹² Relevant considerations include the scope and consequences of the claimed authority, whether the agency is asserting a novel power, and whether it relies on general language in an older statute to regulate beyond its established role.⁹³

The proposed Policy Statement treats Section 5 as authority to establish a substantive federal standard for AI “accuracy” and use that standard to preempt state laws governing AI systems. This is a novel claim of broad authority over a general-purpose technology. Section 5 lacks a clear delegation authorizing the Commission to prescribe substantive standards for AI outputs or to preempt state law. Congress has not clearly delegated the authority the Commission claims.⁹⁴

The result would be the same if the Commission promulgated a rule under Section 18, because compliance with rulemaking procedures cannot expand the Commission’s delegated authority. Section 18 authorizes the Commission to define unfair or deceptive acts or practices with specificity; it does not authorize broad national standards for AI outputs or the preemption of state AI laws. A rule attempting the same result would therefore also violate the major-questions doctrine.⁹⁵

E. The Executive Order Cannot Supply the Missing Authority.

Executive Order 14365 directs the FTC Chairman to issue a policy statement explaining when state laws that require alterations to “truthful” AI outputs are preempted.⁹⁶

⁹¹ *Schneidewind*, 485 U.S. at 300, 305–10.

⁹² *Util. Air Regul. Grp. v. EPA*, 573 U.S. 302, 324 (2014) (quoting *FDA v. Brown & Williamson Tobacco Corp.*, 529 U.S. 120, 160 (2000)).

⁹³ See *West Virginia v. EPA*, 597 U.S. 697, 721–24, 732–35 (2022); *Biden v. Nebraska*, 600 U.S. 477, 500–07 (2023); *FDA v. Brown & Williamson Tobacco Corp.*, 529 U.S. 120, 159–61 (2000).

⁹⁴ See 15 U.S.C. § 45(a)(1); *Proposed Policy Statement*, 91 Fed. Reg. at 41,641–42; *West Virginia*, 597 U.S. at 732–35.

⁹⁵ See 15 U.S.C. § 57a(a)(1)(A)–(B), (a)(2); *West Virginia*, 597 U.S. at 732–35.

⁹⁶ Exec. Order No. 14,365, § 7, 90 Fed. Reg. 58,499, 58,500 (Dec. 16, 2025).

The President may supervise executive officers and set enforcement priorities. But this does not give an agency statutory authority Congress withheld.⁹⁷ No executive order can convert a nonbinding FTC policy statement into federal law, much less one with preemptive effect.

Executive Order 14365 recognizes these limits. Its general provisions state that the Order must be implemented consistent with applicable law and may not impair authority granted by law to an agency or its head.⁹⁸ Its structure also shows that its drafters understood this, despite Administration rhetoric to the contrary. Section 7 calls for an FTC policy statement. Section 8 separately calls for proposed legislation establishing a uniform federal AI framework that preempts conflicting state laws.⁹⁹ The request for legislation makes sense because an executive policy statement is not legislation and cannot supply the FTC with the authority it lacks.

For these reasons, the Commission should remove the Policy Statement's suggestion that it can determine the preemptive effect of Section 5 in the abstract. At minimum, it should make clear that the statement has no independent preemptive force, that preemption depends on a concrete conflict between a specific state-law duty and valid federal law, and that nothing in the statement displaces state law.

IV. THE PROPOSED POLICY STATEMENT ESTABLISHES THE CONDITIONS FOR IMPROPER JAWBONING OF AI DEVELOPERS AND DEPLOYERS

The Proposed Policy Statement advanced an underspecified standard, makes attempts at facial neutrality while also preferencing the Trump Administration's political viewpoints, and exists in a context of enforcement that has failed to be evenhanded. This combination creates the perfect toxic combination of conditions for informal, extra-legal influence on the industry through the looming threat of enforcement. This is textbook jawboning, defined as "the use of official speech to inappropriately compel private action."¹⁰⁰

The Proposed Policy Statement fails to meet constitutional muster or to find roots in Commission authorities as explained in the sections above, but despite these deficiencies, its mere existence will have a chilling effect on companies' speech and product design choices. It advances a vague and unspecified standard that casts doubt on state laws in a way that creates confusion and uncertainty. The Commission even concedes that it cannot locate the boundary of its own standard, acknowledging that "the exact line of what constitutes bias may be difficult to draw."¹⁰¹

⁹⁷ See *Youngstown Sheet & Tube Co. v. Sawyer*, 343 U.S. 579, 585–89 (1952); *id.* at 635–38 (Jackson, J., concurring); *Medellin v. Texas*, 552 U.S. 491, 523–32 (2008).

⁹⁸ Exec. Order No. 14,365, § 9(a)–(b), 90 Fed. Reg. at 58,501.

⁹⁹ *Id.* §§ 7–8, 90 Fed. Reg. at 58,500–01.

¹⁰⁰ Will Duffield, *Jawboning Against Speech*, Cato Inst. Pol'y Analysis No. 934 (Sept. 12, 2022), <https://www.cato.org/policy-analysis/jawboning-against-speech>.

¹⁰¹ Proposed Policy Statement at 6 fn.37.

An FTC investigation carries significant costs, and the threat of those costs can be enough to change company behavior—even when the underlying policy is weakly supported in the law.

It is difficult to read the Proposed Policy Statement as anything but an attempt to exploit these ambiguities to advance an agenda that it could not advance under the law. The Proposed Policy Statement is rooted, and directly references the “Preventing Woke AI in the Federal Government” Executive Order,¹⁰² whose stated purpose is to remove left-aligned ideologies from the AI products the federal government procures. And while that Executive Order claimed that the “the Federal Government should be hesitant to regulate the functionality of AI models in the private marketplace”, the FTC policy statement appears to do just that. It even references some of the same “diversity, equity, and inclusion” language that the “Preventing Woke AI” and other Executive Orders have used in the attempt to remove perceived left-aligned ideologies from the federal government.

This is not the first time this Commission has used its power to attempt to jawbone companies by putting a thumb on the scale in favor of partisan speech and ideals.¹⁰³ Last year the Commission posted a request for public comment that invited complaints of censorship by social media companies in an effort to “smash the censorship cartel.”¹⁰⁴ This public request for comment was transparently designed to unearth claims of “conservative censorship” based on Trump’s long standing belief that social media companies censor conservatives,¹⁰⁵ a belief echoed by many conservative politicians and public figures. Without the Commission actually needing to engage in enforcement or defend a pro-censorship policy in court, this pressure campaign seemed to work, as Meta later announced an overhaul of its content moderation to allow slurs and dehumanizing rhetoric targeting LGBTQ+ people.¹⁰⁶ Notably, there were no new protections or changes in policy to help marginalized communities, including LGBTQ+ communities, who face over-enforcement of content moderation policies.¹⁰⁷

This same double standard is interwoven in the Proposed Policy Statement, creating an anticipation of unequal action intended to nudge corporate policies to favor particular political viewpoints. This current unequal action runs along two axes: treating companies differently

¹⁰² E.O. 14319, Preventing Woke AI in the Federal Government (July 23, 2025)

¹⁰³ Lisa Macpherson, *The Conservative Political Playbook Driving the FTC Platform Censorship Inquiry*, Tech Policy Press (May 21, 2025), <https://www.techpolicy.press/the-conservative-political-playbook-driving-the-ftc-platform-censorship-inquiry>.

¹⁰⁴ Brendan Carr (@BrendanCarrFCC), X (Apr. 3, 2025), <https://x.com/BrendanCarrFCC/status/1907939965251764294>.

¹⁰⁵ Jason Abbruzzese, *Trump Echoes Conservative Claims That Social Media Companies Censor Conservatives*, NBC News (Aug. 18, 2018), <https://www.nbcnews.com/tech/tech-news/trump-echoes-conservative-claims-social-media-companies-censor-conservatives-n901876>.

¹⁰⁶ Christopher Wiggins, *Meta Embraces Dehumanizing Anti-LGBTQ+ Slurs*, The Advocate (Dec 03, 2025), <https://www.advocate.com/business/meta-embraces-dehumanizing-lgbtq-slurs>.

¹⁰⁷ Ángel Díaz & Laura Hecht-Felella, *Double Standards in Social Media Content Moderation*, Brennan Ctr. for Just. (Aug. 4, 2021), <https://www.brennancenter.org/our-work/research-reports/double-standards-social-media-content-moderation>.

based on their perceived politics and treating state laws differently based on the makeup of their state governments.

A. The Commission will enforce unequally among private companies based on political allegiance.

Along the first axis, the disparate treatment of Grok as compared to Anthropic paints a clear picture of partisanship that this policy statement is a part of. There is a long history of evidence that Elon Musk and/or others at his company have altered or steered the output of Grok contrary to consumers' reasonable expectations. This includes the “MechaHitler” update,¹⁰⁸ an update where Grok began repeatedly mentioning “white genocide” in South Africa in response to unrelated topics,¹⁰⁹ and an update where Grok would heap unrealistic praise on Musk.¹¹⁰ Even though these manipulations of Grok are directly on point to the harms the policy statement claims to be aimed at, there has never been any investigation or adverse action taken by the FTC or Trump Administration against xAI. Indeed, shortly after the “MechaHitler” incident the Department of Defense awarded a \$200m contract to xAI.¹¹¹

Compare this with the Trump Administration's treatment of Anthropic, a company that Trump declared to be “A RADICAL LEFT, WOKE COMPANY.”¹¹² The Administration’s feud with Anthropic has been widely documented, and at least part of that feud stems from the perception of Anthropic as left-aligned and employing “radical” democrats like Katie Moussouris.¹¹³ In response, the Trump Administration has ordered US agencies to stop using Anthropic,¹¹⁴ labeled

¹⁰⁸ Lisa Hagen, Huo Jingnan & Audrey Nguyen, *Elon Musk's AI Chatbot, Grok, Started Calling Itself "MechaHitler,"* NPR (July 9, 2025),

<https://www.npr.org/2025/07/09/nx-s1-5462609/grok-elon-musk-antisemitic-racist-content>.

¹⁰⁹ Dara Kerr, *Musk's AI Grok bot rants about 'white genocide' in South Africa in unrelated chats*, The Guardian (May 14, 2025), <https://www.theguardian.com/technology/2025/may/14/elon-musk-grok-white-genocide>.

¹¹⁰ Josh Taylor, *Elon Musk's Grok AI tells users he is fitter than LeBron James and smarter than Leonardo da Vinci*, The Guardian (Nov. 21, 2025),

<https://www.theguardian.com/technology/2025/nov/21/elon-musk-grok-ai-bias-ranks-richest-man-fittest-smartest>.

¹¹¹ Nick Robins-Early, *xAI announces \$200m US military deal after Grok chatbot had Nazi meltdown*, The Guardian (July 14, 2025), <https://www.theguardian.com/technology/2025/jul/14/us-military-xai-deal-elon-musk>.

¹¹² Donald J. Trump (@realDonaldTrump), Truth Social (Feb. 27, 2026), <https://truthsocial.com/@realDonaldTrump/posts/116144552969293195> (“THE UNITED STATES OF AMERICA WILL NEVER ALLOW A RADICAL LEFT, WOKE COMPANY TO DICTATE HOW OUR GREAT MILITARY FIGHTS AND WINS WARS! ... The Leftwing nut jobs at Anthropic have made a DISASTROUS MISTAKE trying to STRONG-ARM the Department of War, and force them to obey their Terms of Service instead of our Constitution. ... Anthropic better get their act together, and be helpful during this phase out period, or I will use the Full Power of the Presidency to make them comply, with major civil and criminal consequences to follow. WE will decide the fate of our Country — NOT some out-of-control, Radical Left AI company run by people who have no idea what the real World is all about.”)

¹¹³ Maria Curi, Marc Caputo, *"They screwed us": Personality clashes sent Anthropic's models offline*, Axios (June 15, 2026), <https://www.axios.com/2026/06/15/anthropic-white-house-fable-mythos>.

¹¹⁴ Matt O'Brien, Konstantin Toropin, *Trump orders US agencies to stop using Anthropic technology in clash over AI safety*, AP News (Feb 27, 2026), <https://apnews.com/article/anthropic-pentagon-ai-hegseth-dario-amodei-b72d1894bc842d9acf026df3867bee8a>.

them as a supply chain risk,¹¹⁵ and temporarily banned the use of Mythos 5 and Fable 5 by foreigners.¹¹⁶ The Department of War's actions against Anthropic resulted in a lawsuit, where amici like the Electronic Frontier Foundation pointed to the significant First Amendment violations involved in DoW's actions.¹¹⁷ Indeed, we expect Anthropic to be one of the companies targeted by this policy statement in furtherance of the attacks on Anthropic's rights to free speech. Therefore, there is no expectation that the Commission will be taking an even hand in the enforcement of the Proposed Policy Statement.

B. The Commission will enforce unequally against state laws based on political composition.

Along the second axis, Colorado and other laws from Democrat-run states are the ones clearly in the crosshairs of this policy statement. Colorado's Automated Decision-Making Technology Act, which replaced the Colorado Artificial Intelligence Act, is directly referenced in the policy statement. The bill represents a legitimate state interest in protecting Colorado residents from bias and errors from automated technology when used to make "consequential decisions" regarding things like employment, access to housing, health care, and insurance.¹¹⁸ This follows real and documented harms with these technologies, like when Amazon discovered an automated recruiting tool they were developing showed bias against hiring women.¹¹⁹ The Proposed Policy Statement appears to be a backdoor attempt to preempt this and other state AI laws, a long held objective of the Trump Administration,¹²⁰ by forcing companies to not comply with those regulations.

However, there is no expectation that the FTC will attempt to do the same with AI laws passed in states with conservative leadership. For example, thirty states have passed legislation regulating deepfakes in elections – including many conservative states. These laws may be addressing real risks to our democratic system, but they also raise significant questions about free speech.¹²¹

¹¹⁵ Maria Curi, *Anthropic Sues Pentagon Over Supply Chain Risk Label*, Axios (Mar. 9, 2026), <https://www.axios.com/2026/03/09/anthropic-sues-pentagon-supply-chain-risk-label>.

¹¹⁶ *Donald Trump's Blocking of Anthropic Is Capricious and Chaotic*, The Economist (June 14, 2026), <https://www.economist.com/briefing/2026/06/14/donald-trumps-blocking-of-anthropic-is-capricious-and-chaotic>.

¹¹⁷ Brief of *Amici Curiae* Electronic Frontier Foundation et al. in Support of Plaintiff, *Anthropic PBC v. Dep't of War*, No. 3:26-cv-01996 (N.D. Cal. Mar. 9, 2026), <https://www.eff.org/document/anthropic-v-department-war-et-al-amicus-brief>.

¹¹⁸ Colo. S.B. 26-189 (2026), <https://leg.colorado.gov/bills/SB26-189>.

¹¹⁹ Jeffrey Dastin, *Insight: Amazon Scraps Secret AI Recruiting Tool That Showed Bias Against Women*, Reuters (Oct. 10, 2018), <https://www.reuters.com/article/world/insight-amazon-scraps-secret-ai-recruiting-tool-that-showed-bias-against-women-idUSKCN1MK0AG/>.

¹²⁰ Exec. Order No. 14365, "Ensuring a National Policy Framework for Artificial Intelligence," § 7, 90 Fed. Reg. 58499, 58500 (Dec. 11, 2025), <https://www.whitehouse.gov/presidential-actions/2025/12/eliminating-state-law-obstruction-of-national-artificial-intelligence-policy/>.

¹²¹ Mark Rumold, *It's Not a Hoax: The Very Real Threat of Political Deepfakes Laws*, Electronic Frontier Foundation (Apr. 2020), <https://www.eff.org/deeplinks/2020/04/not-hoax-very-real-threat-political-deepfakes-laws>.

Trump himself is a prolific poster of AI-generated images,¹²² including those of his political rivals.¹²³ These laws could result in companies “steer[ing] their systems' outputs toward objectives other than those that consumers ask for” to comply. There is less concern that the FTC will attempt to interfere with these laws in Florida, Texas, or any of the other conservative states.

While we would raise the same constitutional and FTC authority concerns with this policy no matter how it was applied, the anticipated unequal enforcement heightens the concern that its effect will be to enable extra-legal jawboning and suppression of free expression. Opposition to this practice has increasing bi-partisan support, and should not be an aim of the Commission or Trump Administration. The JAWBONE Act, introduced by Senate Commerce Committee Chairman Ted Cruz and Senator Ron Wyden, finds that providers of speech-enabling artificial intelligence systems “have a right to independent editorial judgment,” and would create a cause of action against federal officials who coerce such providers to add, alter, or remove generated speech.¹²⁴ Aspiring to protect free expression in this area from government interference should be a paramount concern that unites all Americans.

Lastly, we believe that ultimately the Proposed Policy Statement weakens the Commission's effectiveness and ability to execute its mission to protect consumers. The Commission is a resource constrained institution tasked with regulating broad sectors of the economy against anticompetitive behavior and unfair and deceptive business practices. Its success is largely due to its reputation for rigor and record of winning the fights it picks. Traditionally, the Commission's policy statements and guidelines, like the Horizontal Merger Guidelines, have strengthened its position by being well grounded in the law and are often relied upon by judges in interpreting the law. The Proposed Policy Statement, and its legally baseless attempt to introduce the conditions for intimidation and jawboning, may serve this Commission's immediate short-term goals, but these types of partisan policy statements will erode its credibility and effectiveness. The Commission simply does not have the resources nor capacity to fight every case in court. Without its reputation with judges as being a trusted arbiter, and without its reputation with the business community as being a fearsome litigator, more companies will take their chances in court. If the Commission gains a reputation for bluster instead of action—or losing in court on blatant constitutional overreaches—then it will cease to be sufficient to keep the business community in check. This will be disastrous for consumers who rely on the Commission for

¹²² Kinsey Crowley, *Trump goes on AI post spree, showing 2028 hats, Iran war images*, USA Today (July 27, 2026), <https://www.usatoday.com/story/news/politics/2026/07/27/donald-trump-truth-social-posts-third-term-dylan-mulvaney/91060977007/>.

¹²³ Ariana Baio, *Trump posts wild new AI meme trolling Biden, Obama and Clinton with autopen, ice cream and cocaine*, The Independent (May 7, 2026), <https://www.independent.co.uk/news/world/americas/us-politics/trump-ai-truth-social-biden-obama-clinton-b2972243.html>.

¹²⁴ Justice Against Weaponized Bureaucratic Overreach to Networked Expression (JAWBONE) Act § 2, 119th Cong. (2026), <https://www.commerce.senate.gov/wp-content/uploads/2026/06/JAWBONE-Act-FINAL.pdf>; see Press Release, U.S. Senate Comm. on Commerce, Sci. & Transp., Cruz, Wyden Introduce Legislation to Guard First Amendment Speech Rights Against Government Jawboning (June 2026).

protection from real harms, and will disrupt the Commission's ability to effectively cultivate an environment of prosperous competition and innovation in the United States.

V. CONCLUSION

The Commission should withdraw the Proposed Policy Statement in its entirety and desist from unconstitutional efforts to regulate free expression, speculative adventures in state preemption, and efforts to intimidate or cajole AI developers into ideological alignment with the Trump Administration. Instead, we implore the Commission to focus on its core strengths and the mission for which it is so urgently needed—promoting structural market competition and protecting consumers from real unfair and deceptive acts and practices—in both the burgeoning and critically important AI industry and across the broader technology marketplace.

On behalf of the Joint Commenters:

Nicholas P. Garcia

Senior Policy Counsel, Public Knowledge
nick@publicknowledge.org

John Bergmayer

Legal Director, Public Knowledge
john@publicknowledge.org

Matt Lane

Senior Policy Counsel, Fight for the Future
matt@fightforthefuture.org

Archer Amon

Legal Intern, Public Knowledge
archer@publicknowledge.org