



July 30, 2026

Matter No. P264200
Docket No. FTC-2026-0859
Federal Trade Commission
600 Pennsylvania Avenue NW
Washington, DC 20580

**Re: Federal Trade Commission’s Proposed Policy Statement Concerning the
Suppression of Accuracy in Artificial Intelligence Systems**

Introduction

The Foundation for Individual Rights and Expression (“FIRE”) submits this comment in response to the Federal Trade Commission’s (“FTC’s”) Request for Comment on its Proposed Policy Statement Concerning the Suppression of Accuracy in Artificial Intelligence Systems (“Statement”).

FIRE is a nonpartisan nonprofit that defends the freedoms of speech and thought protected by the First Amendment, including the rights of artificial intelligence (“AI”) developers and users. FIRE has substantial expertise at the intersection of AI and the First Amendment, having testified before legislative bodies, participated in litigation as *amicus curiae*, submitted regulatory comments, and published scholarship addressing constitutional issues AI regulations can raise.

FIRE regularly analyzes and comments on proposed AI legislation and agency actions affecting expressive technologies, including issues involving compelled speech, anonymous access to AI, content moderation, editorial discretion, and the constitutional limits on government regulation

of AI systems. Because AI systems increasingly serve as tools for creating and accessing protected expression, FIRE has a strong interest in ensuring potential government action fully accounts for the First Amendment implications of regulating AI system designs and outputs.

This comment urges the Commission to consider the implications of the Statement for AI systems that communicate and exchange information. AI systems function as expressive tools, and their outputs reflect editorial judgments embedded by their developers through choices about training data, system prompts, model alignment, and other design decisions. As the Commission considers how Section 5 applies in this context, it must account for the constitutional interests implicated by regulating those editorial choices.

This comment further explains why the Statement appears to depart from the Commission's traditional deception framework. Rather than focusing solely on commercial representations by AI companies, the Statement makes broad assumptions about AI developers' objectives and consumers' expectations, then anticipates using those assumptions to evaluate the substance of AI-generated outputs. The Statement's approach risks running afoul of the First Amendment by requiring the Commission to assess editorial judgments about how AI systems select, organize, and present information.

Because AI remains in the relatively early stages of widespread public adoption, the Commission must take particular care that any regulatory standards adopted today do not inadvertently shape the long-term development of expressive AI systems in ways that are difficult to reverse.

I. The development and use of AI is protected by the First Amendment.

Generative AI systems, commonly known as large language models, are tools for communicating information and ideas. They answer questions, summarize sources, create arguments, draft essays, translate languages, and engage users in open-ended conversations across virtually every subject.

Human developers make decisions about what text to use during pretraining to fundamentally shape how the model understands and interacts with users. Developers select and curate training data, fine-tune models, craft system prompts, and implement safeguards to affect how AI systems prioritize and present information. The resulting outputs therefore reflect not only a sophisticated statistical prediction system but also the developer's editorial judgment.

Newspaper editors decide which stories to publish, search engines rank competing sources, publishers determine what to print, bookstores curate collections, and recommendation systems decide what content to elevate or limit. AI developers likewise make editorial judgments about how to select, organize, and present information to users.

The First Amendment protects these editorial decisions.¹ As the Supreme Court recently reaffirmed, government interference with editorial discretion “alters the content of the compilation” because it overrides “a private party’s expressive choices.”² And although applying those protections to new technologies may present novel questions, “‘the basic principles’ of the First Amendment ‘do not vary.’”³ The government therefore may not compel private entities to present views they would reject, control their editorial choices, or force them to associate with speech they prefer to exclude.⁴ Those same principles apply equally here.

II. The Statement departs from Section 5 by regulating protected expression.

A. The Statement treats contested policy judgments as legal premises.

The Statement does more than identify deceptive commercial practices prohibited by Section 5. It begins by asserting a series of broad propositions about the objectives of AI developers, the expectations of consumers, and the proper function of AI systems. These propositions are presented as though they are self-evident or universally accepted. They are not.

For example, the Statement repeatedly asserts that AI companies seek only to produce “the best output,” that consumers uniformly expect AI systems to pursue users’ objectives free from any “hidden agenda,” and that AI companies implicitly promise to “distill all of human knowledge” in service of users’ goals. It further suggests departures from these assumptions — such as pursuing objectives other than those consumers purportedly expect — may constitute deceptive conduct under Section 5.

The Statement’s propositions are neither statutory requirements nor principles of deception law. They are contestable policy judgments raised by the Commission about how AI systems ought to function and what consumers should expect from them. But the Supreme Court “has many times held, in many contexts, that 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.”⁵

Developers possess the First Amendment right to pursue a wide range of objectives in creating and calibrating their AI systems, including but not limited to factual accuracy, safety, privacy,

¹ *Moody v. NetChoice, LLC*, 603 U.S. 707, 728–29 (2024).

² *Id.* at 732.

³ *Id.* at 733 (quoting *Brown v. Ent. Merchs. Ass’n*, 564 U.S. 786, 790 (2011)).

⁴ *Id.* at 731–33.

⁵ *Moody*, 603 U.S. at 719 (noting “[t]hat principle works for social-media platforms as it does for others”).

legal compliance, civility, and protection against misuse. Others may seek to provide users with a unique worldview.⁶ The First Amendment protects each of these pursuits — and many more — just as it protects publications with social, literary, artistic, political, religious, journalistic, and scientific focuses as varied as the American populace.

The Statement’s declaration that AI models uniformly dedicate themselves to producing “the best output” is a presumption, not a truism. The right to determine what “the best output” may mean in any given context is a decision the First Amendment reserves to the developer, not the federal government, and the Statement’s suggestion to the contrary is dangerously misguided.

Likewise, consumers’ expectations are neither uniform nor fixed. Different users choose different AI systems precisely because they prefer different balances among competing objectives, just as different readers may choose to subscribe to different news outlets. Again, the First Amendment protects from government encroachment that freedom to choose where to share and receive information, and the Commission has no authority to recharacterize protected editorial decisions as “deceptive” based on a declaration that AI users monolithically desire to, as the Statement suggests, “feel that their values are being respected.” The Commission has no authority to define Americans’ “values.”

By treating these contested propositions as a baseline for assessing deception, the Statement risks transforming the Commission’s preferred conception of AI’s development, deployment, and utility into an enforceable legal standard. Whether an AI system has improperly “steered” its outputs would then depend not simply on whether a company made a false commercial representation *about* the product or service. Rather, it would rest on whether the Commission agrees with the developer’s editorial and design judgments, how the system balances competing objectives, and the outputs of the system. The First Amendment does not permit such a result.

B. The Statement requires the Commission to police protected editorial judgments.

As the Statement recognizes, Section 5 of the Federal Trade Commission Act prohibits “unfair or deceptive acts or practices in or affecting commerce.”⁷ Under the Commission’s longstanding deception framework, liability ordinarily turns on whether a representation, omission, or practice

⁶ Talkie, for example, is a large language model trained on data before 1930. Tom Hawking, *Talkie Is a ‘Vintage LLM’ Trained on Pre-1930 Data to Help Facilitate ‘Time Travel’*, GIZMODO (Apr. 28, 2026, at 14:35 ET), <https://gizmodo.com/talkie-is-a-vintage-llm-trained-on-pre-1930-data-to-help-facilitate-time-travel-2000751758>.

⁷ 15 U.S.C. § 45(a)(1).

is likely to mislead “the consumer acting reasonably in the circumstances, to the consumer’s detriment.”⁸

The Statement instead would require the Commission to evaluate whether the company has “steer[ed] the outputs of [its] AI systems toward unexpected objectives, and away from the objectives set by or reasonably expected by users.” As a result, the Commission would no longer examine only whether the product or service provider had adequate support for its commercial representations. It would also evaluate the content of AI-generated responses against variable and ever-changing user objectives to determine whether those responses were deceptive.

AI systems routinely answer questions involving disputed history, economics, law, politics, philosophy, and other subjects on which people disagree. Developers of those systems make constitutionally protected editorial choices that influence what information their system may include, omit, and prioritize in its outputs — and how to characterize it.

Although the Statement distinguishes ordinary “hallucinations” from deceptive conduct — explaining that liability may arise from misrepresentations about the likelihood of hallucinations rather than the hallucinations themselves — it abandons that same distinction when discussing alleged ideological steering.⁹ Instead of asking whether an AI company accurately represented the objectives or capabilities of its system, the Statement invites the Commission to evaluate the substance of the system’s outputs themselves. Without an objective standard distinguishing protected editorial judgment from actionable deception, lawful expression risks being treated as deceptive conduct simply because the Commission disagrees with a developer’s editorial choices.

The examples provided in the Statement further illustrate this point. AI systems that produce “ideologically motivated distortions,” reflect “historical injustices,” prioritize “equity,” avoid “politically inflammatory outputs,” or further developers’ or employees’ “political agendas” would be in violation of Section 5 under the Statement. But each of these determinations would require the Commission to adjudicate intensely contested political and social questions, and to punish developers that reach different conclusions from the government. That is an authority the First Amendment does not permit a federal agency to wield. It is bedrock First Amendment law that “[o]n the spectrum of dangers to free expression, there are few greater than allowing the

⁸ FTC Policy Statement on Deception, 103 F.T.C. 174, 183 (Oct. 14, 1983), *appended to, In re Cliffdale Assocs., Inc.*, 103 F.T.C. 110 (1984).

⁹ See Proposed Policy Statement Concerning the Suppression of Accuracy in Artificial Intelligence Systems, 91 Fed. Reg. 41,638, 41,641 n.46 (July 7, 2026) (“While a company may, for example, unlawfully deceive consumers if it misrepresents the likelihood of such hallucinations, the Commission does not believe that such hallucinations in and of themselves raise issues under section 5 or any other law that the Commission enforces.”).

government to change the speech of private actors in order to achieve its own conception of speech nirvana.”¹⁰

C. The Statement requires the Commission to scrutinize the content of AI-generated expression.

The Statement marks a significant expansion of the Commission’s deception authority. Rather than evaluating whether an AI developer adequately substantiated its commercial claims, its approach would require the Commission to evaluate the substance of AI-generated expression itself against what *the Commission* believes users expect or infer about the product or service.

But again, determining whether the substance of AI responses satisfy users’ “reasonable expectations” necessarily requires the Commission to choose among competing views about the best use of AI rather than determine whether a seller adequately substantiated a commercial representation.

Imagine a user asks an AI model, “Were COVID-19 school closures justified?” The user might then ask about public-health benefits, learning loss, mental-health effects, and the evolution of scientific recommendations. Determining whether those responses have been improperly “steered” requires the Commission to evaluate the editorial judgments reflected in how the AI system selected, prioritized, and presented competing information.

The Department of Justice identified the same problem when it challenged Colorado’s AI Act, arguing the law compelled developers to make “different editorial decisions regarding training data, responses to prompts, model constraints, and more — all to generate Colorado’s preferred expressive outputs.”¹¹ That is the same power the FTC attempts to leverage here. The Statement likewise requires the Commission to evaluate those editorial decisions in determining whether AI-generated responses have satisfied the Commission’s standards regarding truthfulness, accuracy, and consistency with users’ assumed objectives and expectations.

D. The Statement creates strong incentives to suppress protected expression.

The net effect of the foregoing is that the Statement may well force AI developers to anticipate which lawful design choices the Commission may later characterize as deceptive and censor their models accordingly. Choices involving system prompts, retrieval, ranking, moderation, alignment, and uncertainty all shape AI-generated responses, yet the Statement provides no clear

¹⁰ *Moody*, 603 U.S. at 741–42.

¹¹ United States of America’s Complaint in Intervention ¶ 10, *x.AI LLC v. Weiser*, No. 1:26-cv-01515 (D. Colo. Apr. 24, 2026).

standard for distinguishing permissible editorial decisions from choices that would lead to actionable deception.

Developers seeking to avoid liability would have to anticipate how the Commission will interpret contested concepts such as neutrality, objectivity, and ideological steering, which may evolve over time with changes in Commission policy and priorities. Faced with that uncertainty, developers would have incentives to modify their models — not merely their representations — to reduce the risk that future responses will be characterized by the government as deceptive.

Although the Statement contemplates that disclosures may mitigate “deception,” it also indicates that disclosures must be sufficiently intrusive to alter what the Commission has assumed are consumers’ expectations. That approach suggests that revised representations or disclosures alone may be insufficient, creating pressure to alter the underlying AI system itself.

Conclusion

Section 5 has long prohibited deceptive commercial practices, not the editorial judgments embodied in expressive tools like AI. The Statement would depart from that tradition by requiring the Commission to evaluate whether AI-generated expression reflects the objectives the Commission believes users expect an AI system to pursue. Because that inquiry is necessarily normative and necessarily reaches protected editorial discretion, the Commission should decline to interpret Section 5 in a manner that authorizes the regulation of lawful expression.

Respectfully submitted,



John Coleman
Legislative Counsel
Foundation for Individual Rights and Expression (FIRE)

██████████
Philadelphia, PA ██████████
██████████