

**Before the
FEDERAL TRADE COMMISSION
Washington, DC 20580**

In the Matter of	)	
	)	
Accuracy in Artificial Intelligence Systems	)	FTC-2026-0859
Proposed Policy Statement	)	

COMMENTS OF FREE PRESS

Shilpa Jindia, Policy Counsel
Nora Benavidez, Senior Counsel
Matthew F. Wood, VP of Policy
Free Press
1025 Connective Avenue NW
Suite 1110
Washington, DC 20036
(202) 265-1490

July 31, 2026

TABLE OF CONTENTS

Executive Summary	3
I. The Commission Does Not Have Authority to Dictate AI Outputs.....	4
A. The FTC’s proposed policy is a spurious use of Section 5.....	4
B. Civil Rights Laws Require That Direct or Indirect Discrimination in Online Services Be Remedied.....	6
II. The Commission’s Proposed Policy Statement Furthers The Administration’s Censorial Agenda.....	8
A. The FTC’s <i>Policy Statement</i> is viewpoint discrimination that would plainly give the Commission unbridled authority to cherry-pick which AI tools to target as a way of suppressing views and information it disfavors.....	8
B. The Commission’s Policy Statement is a clear example of jawboning.....	9
III. State Laws Requiring AI Models to Correct for Racial Bias Are Not Deceptive Under Section 5, and Any Effort to Paint Bias Mitigation as Coercion Runs Afoul of Civil Rights Law.....	11
A. AI systems contain inaccuracies caused by human bias.....	11
B. Grok exemplifies the harms caused by AI when private action stokes discrimination rather than correcting for it.....	12
C. Transparency and auditing AI tools are a more appropriate set of solutions.....	14
IV. Conclusion	14

EXECUTIVE SUMMARY

The Federal Trade Commission has enforcement authority under Section 5 of the FTC Act to remedy deceptive or unfair consumer practices. It is now offering a policy proposal to apply this legal sanction to AI companies whose outputs do not align with the Trump administration's version of history and truth. Not only is this policy a gross abuse of the Commission's authority, it also violates the First Amendment as clear content-based viewpoint discrimination intended to coerce private actors.

The administration's effort to police AI bias mitigation is yet another prong of its foolhardy war on Diversity, Equity, and Inclusion. Inaccuracy is a well-documented problem in machine learning, and industry expertise has long recognized that biased data sets are a fundamental contributor to this issue. Disparate impacts are a key tool for both documenting and correcting algorithmic bias. The Commission's proposed policy is another Trump administration measure to dismantle this key Civil Rights framework, and an unconstitutional attempt to police "ideological" speech.

The Commission is disregarding its own precedent and practice to further the administration's censorial agenda. It should abandon such an ill-conceived attempt to manipulate our information ecosystem.

I. The Commission Does Not Have Authority to Dictate AI Outputs.

A. The FTC’s proposed policy is a spurious use of Section 5.

The Federal Trade Commission (“FTC” or “Commission”) is tasked with remedying deceptive or unfair consumer practices. This does not give it license to suppress AI outputs any more than it may suppress other speech and opinions that it disfavors, nor license to dictate AI platform design. The Commission’s proposed policy statement (“*Policy Statement*”) flows from the premise that “consumers have a reasonable expectation that AI systems aim to give truthful and accurate outputs. Consumers have no basis to believe that AI systems aim to produce outputs that are distorted by undisclosed ideological objectives.”¹

The Commission can and should enforce the law against deceptive claims regarding verifiable product performance and service efficacy, and it can prevent actual and actionable civil rights abuses. Policing whether so-called ideological speech is “truthful and accurate” is categorically different: that task falls well outside of the FTC’s expertise and authority, and falls afoul of the First Amendment, too.

Section 5 of the FTC Act prohibits “deceptive acts or practices.”² “The FTC must show three elements to prove a deceptive act or practice in violation of Section 5(a)(1): (1) a representation, omission, or practice, that (2) is likely to mislead consumers acting reasonably under the circumstances, and (3), the representation, omission, or practice is material.”³ Although the Commission need not show that the challenged representation or omission in fact deceived a particular consumer,⁴ the Commission still must show that the defendant made a material misrepresentation or omission that was “likely to mislead customers acting reasonably under the circumstances.”⁵

As the *Policy Statement* notes, the FTC has filed Section 5 complaints in recent years to rein in deceptive or unfair conduct by companies marketing AI tools or features.⁶ These cases involved deceptive claims about the use of AI or generative AI: for instance, deceptive claims

¹ *Federal Trade Commission’s Proposed Policy Statement Concerning The Suppression of Accuracy in Artificial Intelligence Systems*, Federal Trade Commission (July 1, 2026).

² 15 U.S.C. § 45(a).

³ *Fed. Trade Comm’n v. Moses*, 913 F.3d 297, 306 (2d Cir. 2019) (internal quotation marks omitted).

⁴ *FTC v. Cyberspace.com*, 453 F.3d 1196, 1201 (9th Cir. 2006).

⁵ *FTC v. Tashman*, 318 F.3d 1273, 1277 (11th Cir. 2003).

⁶ See *FTC Announces Crackdown on Deceptive AI Claims and Schemes*, Federal Trade Commission (Sept. 25, 2024), <https://www.ftc.gov/news-events/news/press-releases/2024/09/ftc-announces-crackdown-deceptive-ai-claims-schemes>.

about the ability of an AI chatbot to replace the services of a human lawyer⁷ and deceptive claims regarding accuracy and or efficacy of AI-powered facial recognition software.⁸

The FTC has abandoned this narrow and targeted approach, using the *Policy Statement* in an attempt to convert its version of “truthful and accurate outputs”⁹ into a “legal baseline.”¹⁰ Rather than applying Section 5 to actual deception regarding the capabilities of the burgeoning class of AI tools, this FTC would instead use Section 5 to police and then dictate the “ideological” truth and accuracy of content generated by AI.¹¹ As a result, the Commission’s policy would create a standard where “[a] system that deviates” from some government-determined “truth,” and that does so “without disclosure is presumptively deceptive.”¹² Of course, the FTC is not equipped to determine the truth of these outputs; and as discussed in more detail below, the First Amendment bars such dangerous folly and self-aggrandizement by state actors.

The FTC also departs from past practice by making a distinction between “accuracy suppression driven by design decisions and incorrect outputs resulting from technological errors (i.e., hallucinations.)”¹³ The only true hallucination in the *Policy Statement* is the Commission’s desperate claim that AI models trained with “ideologically motivated distortions ... such as ... historical injustices”¹⁴ constitute a material representation, omission, or practice likely to deceive customers. There is no evidence, rationale, or hint of plausible causality to suggest how informing a customer about, say, slavery would “steer the outputs of ... AI systems toward unexpected outcomes, and away from the objectives set by or reasonably expected by users”. Nor how supplying that historical context would induce a customer to “pay for a service that does not behave as advertised” or rely “on a technology that, by design, produces worse outputs and may recommend suboptimal courses of action.”¹⁵

This *Policy Statement* manages to be both dangerous and absurd at the same time. It is a gross abuse of the Commission’s authority to classify an entire group of products based on their alleged “ideological” content as presumptively deceptive. The Commission is promising to single

⁷ See Consent Order, *In re DoNotPay, Inc.*, FTC Docket No. C-4812 (Jan. 14, 2025) (deceptive claims about the ability of its AI chatbot to replace the services of a human lawyer),

https://www.ftc.gov/system/files/ftc_gov/pdf/2323042_donotpay_decision_and_order_0.pdf.

⁸ Consent Order, *In re Intellivision Techs. Corp.*, No. C-4809 (F.T.C. Jan. 8, 2025) (deceptive claims regarding accuracy or efficacy of its AI-powered facial recognition software),

https://www.ftc.gov/system/files/ftc_gov/pdf/2323023c4809intellivisionfinalorder.pdf.

⁹ *Policy Statement*.

¹⁰ Eran Kahana, *When (or How) Principles Can Become Law: The FTC’s Proposed Deceptive Steering Policy Scoped Against the AI Life Cycle Core Principles*, Stanford Law School Blogs (July 8, 2026),

<https://law.stanford.edu/2026/07/08/when-or-how-principles-can-become-law-the-ftcs-proposed-deceptive-steering-policy-scoped-against-the-ai-life-cycle-core-principles/>.

¹¹ See *id.* n.5.

¹² *Id.*

¹³ Adam Maarec, Aja D. Finger, *FTC Takes Aim at AI Accuracy*, Consumer Finance Monitor (July 14, 2026),

<https://www.consumerfinancemonitor.com/2026/07/14/ftc-takes-aim-at-ai-accuracy/>.

¹⁴ *Policy Statement* at 7.

¹⁵ *Id.* at 8.

out companies that implement longstanding expert recommendations¹⁶ to correct for bias, including disparate impact, based on protected characteristics in their algorithms.

B. Civil Rights Laws Require That Direct or Indirect Discrimination in Online Services Be Remedied.

Civil Rights laws proscribe discrimination based on certain protected characteristics in advertising for and provision of housing, jobs, and other economic opportunities, such as financial lending practices. Such anti-discrimination frameworks are intended to remedy both disparate treatment (*i.e.*, intentional discrimination), and disparate impact (*i.e.*, covert discriminatory effects that occur regardless of intent)¹⁷ in statutes including the Civil Rights Act and the Housing Act of 1954. The FTC is empowered to address such discrimination through statutes such as the Equal Credit Opportunity Act of 1974, the Fair Credit Report Act, and the FTC Act Section 5.¹⁸

These statutes and doctrines have long aimed to correct egregious historical practices of discrimination, as well as racial segregation that has continued to the present day. In the context of housing, government agencies including the Home Owners Loan Corporation, Fannie Mae, and the Federal Housing Administration instituted policies that targeted housing advertisements to white people while denying people of color mortgages in the same neighborhoods. This practice, which came to be known as redlining, created racial geographies that still exist today. Communities of color were segregated into neighborhoods with poorer quality of life, whether based on depreciated housing value, greater environmental risks such as lead paint and asthma, or limited access to banking and fresh food.¹⁹

Likewise, as financial and advertising services have increasingly shifted online, federal regulators, including the FTC, have raised concerns over digital redlining: the “creation and maintenance of technology practices that further entrench discriminatory practices against already marginalized groups.”²⁰ In 2022, the White House released a Blueprint for the AI Bill of Rights, a wholesale government effort to address artificial intelligence and basic rights. A significant portion of this blueprint mapped ways to protect against algorithmic discrimination and to apply civil

¹⁶ See Nicol Turner Lee, Paul Resnick & Genie Barton, *Algorithmic bias detection and mitigation: Best practices and policies to reduce consumer harms*, The Brookings Institute (May 19, 2019), <https://www.brookings.edu/articles/algorithmic-bias-detection-and-mitigation-best-practices-and-policies-to-reduce-consumer-harms/>.

¹⁷ April J. Anderson, *What Is Disparate-Impact Discrimination?*, Congressional Research Service (July 9, 2025), https://www.congress.gov/crs_external_products/IF/PDF/IF13057/IF13057.1.pdf.

¹⁸ *Equal Opportunity Credit Act*, Federal Trade Commission (last searched July 31, 2026), <https://www.ftc.gov/legal-library/browse/statutes/equal-credit-opportunity-act>; see also Elisa Jillson, *Aiming for truth, fairness, and equity in your company’s use of AI*, Federal Trade Commission (Apr. 19, 2021).

¹⁹ Amicus Brief of The American Civil Liberties Union Foundation, Free Press, The Lawyers’ Committee For Civil Rights Under Law, And The National Fair Housing Alliance, *Vargas v. Facebook, Inc.*, at 5-7 (9th Cir. filed Jan. 26, 2022) (“*Vargas* amicus brief”).

²⁰ *Vargas* amicus brief at 1-2 (quoting *Banking on Your Data: the Role of Big Data in Financial Services*, Hearing before Task Force on Fin. Tech. of the House Comm. on Fin. Serv., 116th Cong., at 9 (Nov. 21, 2019) (statement of Dr. Christopher Gilliard).

rights statutes to mitigate it.²¹ In 2019, Facebook settled a lawsuit with the Department of Justice, requiring the platform to alter some of its targeting tools for housing, employment, and credit ads — with Facebook essentially acknowledging that it had engaged in discriminatory algorithmic practices.²² Facebook removed gender, age, and multicultural affinity targeting options as well as targeting by zip code. The settlement also required Facebook to remove targeting options related to personal characteristics or classes protected under anti-discrimination laws.²³

The rise of algorithmic systems and machine learning has triggered the parallel concern of algorithmic bias, where the biases of input data are replicated through such systems' outputs.²⁴ Algorithmic decision-making has rapidly been implemented across multiple fields, including hiring, law enforcement, education, credit reporting, criminal justice, stock trading, and more.²⁵ In 2019, the Brookings Institution documented the rise of algorithmic bias, as well as the importance of disparate-impact frameworks in capturing these outcomes: “Bias in algorithms can emanate from unrepresentative or incomplete training data or the reliance on flawed information that reflects historical inequalities. If left unchecked, biased algorithms can lead to decisions which can have a collective, disparate impact on certain groups of people even without the programmer’s intention to discriminate.”²⁶

Countering such discrimination is not a deceptive practice. Disparate-impact laws are, in fact, key to preventing AI discrimination,²⁷ and that is why the Trump FTC now takes aim at them. The Commission’s incoherent argument is one more attempt by the administration to eliminate disparate-impact liability in federal civil rights enforcement.²⁸ Colorado’s state law became the first comprehensive AI regulation in the country that implemented a disparate-impact framework, and was swiftly attacked by forces bent on eliminating such critical civil rights protections.²⁹ Of course, it is precisely the history of discrimination and disparate impacts that this administration wants to label as “ideological” and then erase — not only from public sites but from private

²¹ *The Blueprint for an AI Bill of Rights: Making Automated Systems Work for the American People*, The White House Office of Science and Technology (Oct. 2022), <https://www.govinfo.gov/content/pkg/GOVPUB-PREX23-PURL-gpo193638/pdf/GOVPUB-PREX23-PURL-gpo193638.pdf>.

²² *Vargas* amicus brief at 5 n.4.

²³ *Summary of Settlements Between Civil Rights Advocates and Facebook*, American Civil Liberties Union (Mar. 19, 2019), <https://www.aclu.org/documents/summary-settlements-between-civil-rights-advocates-and-facebook>.

²⁴ See Pia Sombetzki, *What is algorithmic discrimination?*, AlgorithmWatch (May 12, 2025), <https://algorithmwatch.org/en/what-is-algorithmic-discrimination/>.

²⁵ Xukang Wang, *et al.*, *Algorithmic discrimination: examining its types and regulatory measures with emphasis on US legal practices*, *Frontiers in Artificial Intelligence* (May 21, 2024), <https://www.frontiersin.org/journals/artificial-intelligence/articles/10.3389/frai.2024.1320277/full>.

²⁶ Turner Lee *et al.*, *supra* note 16.

²⁷ See Chiraag Bains, *The legal doctrine that will be key to preventing AI discrimination*, The Brookings Institution (Sept. 13, 2024), <https://www.brookings.edu/articles/the-legal-doctrine-that-will-be-key-to-preventing-ai-discrimination/>.

²⁸ See Exec. Order No. 14281, *Restoring Equality of Opportunity and Meritocracy*, 90 Fed. Reg. 17537 (Apr. 23, 2025).

²⁹ Stephen Bardol, *Colorado’s AI Act Is Dead. Long Live Colorado’s AI Act*, LawFuel (June 9, 2026), <https://www.lawfuel.com/colorado-ai-act-2026-sb-26-189/>.

speakers' statements, too. The Commission's proposed *Policy Statement* simply extends this censorial and autocratic impulse to the burgeoning world of AI. It represents yet another tired but still dangerous Trump administration screed against Diversity, Equity, and Inclusion ("DEI").

II. The Commission's Proposed Policy Statement Furthers The Administration's Censorial Agenda.

The *Policy Statement* dons the pretext of safeguarding AI accuracy, but it is part of a wider effort to weaken independent institutions and sectors as a mode of limiting criticism, inquiry, and free speech. In a world where this *Policy Statement* becomes codified, it would be yet another example of the administration's unlawful encroachment into civilian and third-party activities.

A. The FTC's *Policy Statement* is viewpoint discrimination that would plainly give the Commission unbridled authority to cherry-pick which AI tools to target as a way of suppressing views and information it disfavors.

The Commission's *Policy Statement* is classic viewpoint discrimination. The constitutional standard here is clear: "The government must abstain from regulating speech when the specific motivating ideology or the opinion or perspective of the speaker is the rationale for the restriction."³⁰ While the Commission – like all governmental bodies – may "share [its] views freely and criticize particular beliefs," it may not tacitly use its authority to coerce private actors to adopt its views.³¹

The *Policy Statement*'s aim is to suppress what it labels "ideologically motivated distortions in a response to a factual question, such as to correct what the developer believes are 'historical injustices in the facts.'"³² The *Policy Statement* is thus content-based³³ and violates the First Amendment because it would fail strict scrutiny. It treats AI, *i.e.* private speech, differently based on the Commission's imposition of what it considers to be "accurate" ideology. Content-based restrictions on speech are "presumptively invalid"³⁴ and are subject to the "most exacting scrutiny,"³⁵ because they "pose the inherent risk that the Government seeks not to advance a legitimate regulatory goal, but to suppress unpopular ideas or information or manipulate the public debate through coercion rather than persuasion."³⁶

The Commission bears the burden of proving that a content-based restriction on speech furthers a compelling interest and is narrowly tailored to achieve that interest.³⁷ Meeting this

³⁰ *Rosenberger v. Rector and Visitors of University of Virginia*, 515 U.S. 819 (1995).

³¹ *Nat'l Rifle Ass'n of Am. v. Vullo*, 602 U.S. 175, 188 (2024).

³² *Policy Statement* at 7.

³³ *Turner Broadcasting Sys., Inc. v. FCC*, 512 U.S. 622, 643 (1994) ("[L]aws that by their terms distinguish favored speech from disfavored speech on the basis of the ideas or views expressed are content based.").

³⁴ *R.A.V. v. City of St. Paul*, 505 U.S. 377, 382 (1992).

³⁵ *Turner*, 512 U.S. at 642.

³⁶ *Id.* at 641.

³⁷ *Reed v. Town of Gilbert*, 576 U.S. 155, 171 (2015) (quoting *Arizona Free Enterprise Club's Freedom Club PAC v. Bennett*, 564 U.S. 721, 734 (2011)).

burden would be particularly difficult for the Commission because the *Policy Statement*'s content-based terms are undefined. The danger of chilling speech based on viewpoint is exacerbated by the unreviewable nature of the decision to label an AI tool truthful or distorted, or to suggest its results are or are not accurate, without some kind of final agency action.

In addition to being content-based discrimination that cannot withstand strict scrutiny, the *Policy Statement* suffers from fatal imprecision. It is vague and overbroad. A law can be fatally imprecise on its face under two separate doctrines: vagueness under the Due Process Clause and overbreadth under the First Amendment.³⁸ If a law fails either test, it is invalid on its face and void in all applications. The poorly reasoned *Policy Statement* fails both.

It would be unconstitutionally vague because of its undefined terms. A law is unconstitutionally vague if "it leaves the public uncertain as to the conduct it prohibits."³⁹ Similarly, a law can be impermissibly vague if it authorizes or "even encourage[s] arbitrary and discriminatory enforcement."⁴⁰ Because the key terms are so broad and their definitions so inherently elusive, there are no articulable standards for what might constitute an "ideological" distortion or deception. In the First Amendment context, this vagueness is all the more offensive because it "encourages arbitrary and erratic" enforcement⁴¹ based on the message conveyed.

The *Policy Statement* is also unconstitutionally imprecise because it is overbroad. "[T]he overbreadth doctrine permits the facial invalidation of laws that inhibit the exercise of First Amendment rights if the impermissible applications of the law are substantial when "judged in relation to the statute's plainly legitimate sweep."⁴² "Overbreadth review is a necessary means of preventing a 'chilling effect' on protected expression."⁴³ Laws can be "subject to attack on the ground that they simply prohibit too much protected speech."⁴⁴

The most compelling reason to cut off rather than adopt the *Policy Statement* is that it conveys standardless and unreviewable discretion to the FTC to decide what AI "outputs" might be labeled deceptive. And proposed regulations "may not vest public officials with unbridled discretion."⁴⁵

B. The Commission's Policy Statement is a clear example of jawboning.

The Commission's *Policy Statement* will also likely appear in some future casebook as a clear example of jawboning: the coercion of private actors. While the doctrine of jawboning is still

³⁸ *City of Chicago v. Morales*, 527 U.S. 41, 52 (1999).

³⁹ *Id.* at 57 (quoting *Giaccio v. Pennsylvania*, 382 U.S. 399, 402-03 (1966)).

⁴⁰ *Id.*

⁴¹ *Papachristou v. City of Jacksonville*, 405 U.S. 156, 162 (1972).

⁴² *Morales*, 527 U.S. at 53 (quoting *Broadrick v. Oklahoma*, 413 U.S. 601, 612-15 (1973)).

⁴³ *Broadrick*, 413 U.S. at 601, 630.

⁴⁴ *City of Ladue v. Gilleo*, 512 U.S. 43, 52 (1994) (striking down a municipal ordinance that barred a citizen from displaying a sign from her home).

⁴⁵ *Burke v. Augusta-Richmond Cty.*, 365 F.3d 1247, 1256 (11th Cir. 2004).

evolving, the Supreme Court has drawn clear lines around unlawful government intrusion into private speech. Ultimately, “a government official cannot do indirectly what she is barred from doing directly: A government official cannot coerce a private party to punish or suppress disfavored speech on her behalf,” particularly by “using the power of [the] office to threaten legal sanctions against” disfavored speakers.⁴⁶

The FTC here is not exercising any legitimate regulatory authority in the limited zone where the government can compel disclosure for consumer products. The Supreme Court has permitted the government to compel commercial speech about “purely factual and uncontroversial information” to counter consumer deception.⁴⁷ However, this precedent under *Zauderer* is wholly inapplicable both to policies of content moderation and systemic disclosures.⁴⁸ Mandating disclosure requirements for online platforms is more akin to demanding that publications explain their editorial policies and publishing decisions, an obvious First Amendment breach.⁴⁹ The Commission’s proposed *Policy Statement* is predicated on government influencing or manipulating AI models in order to produce outputs that do not conflict with the Trump administration’s view of truth. State action that “threatens platforms’ ongoing creation and enforcement of their editorial rules” has no analog in the *Zauderer* universe of consumer protection.⁵⁰

The FTC’s threat to classify AI that includes information about “historical injustices” as a Section 5 violation exceeds even the most generous interpretation of legitimate enforcement action.⁵¹ While the Commission rails against AI laws in states such as Colorado or California as “requiring AI companies to censor outputs and insert left-wing ideology,”⁵² it seeks to impose its will by discouraging or disallowing such ideology. The *Policy Statement* thus engages in the very same conduct it (wrongly) accuses state laws of doing, through the broadest of government edicts about ideological accuracy, all in service of the Trump administration’s war on DEI.

Just two years ago, Chairman Ferguson recognized that “Congress has not given us power to regulate AI standing alone. We should not succumb to the panicked calls for the Commission to act as the country’s comprehensive AI regulator.”⁵³ Yet now he proposes a *Policy Statement* that seeks to coerce private AI companies into following this administration’s agenda against DEI. Smaller companies especially will find it difficult to ignore the not-so-veiled threat in the FTC’s proposal, as they will have fewer resources to defend their rights. As with the administration’s

⁴⁶ *Vullo*, 602 U.S. at 190.

⁴⁷ *Zauderer v. Off. of Disciplinary Counsel*, 471 U.S. 626, 651-52 (1985).

⁴⁸ See Daphne Keller, *Platform Transparency and the First Amendment*, *Journal of Free Speech Law*, at 55-63 (2023), <https://www.journaloffreespeechlaw.org/keller2.pdf>.

⁴⁹ *Id.* at 7.

⁵⁰ *Id.*

⁵¹ See *id.* at 19.

⁵² *Policy Statement* at 3 n.12 (quoting Exec. Order No. 14365, *Ensuring a National Policy Framework for Artificial Intelligence*, 90 Fed. Reg. 58499, 58499 (Dec. 11, 2025)).

⁵³ Concurring Statement of Commissioner Andrew N. Ferguson, In the Matter of DoNotPay, Inc., at 2 (Sept. 25, 2024), https://www.ftc.gov/system/files/ftc_gov/pdf/ferguson-donotpay-statement.pdf.

other forays into censorship, the Commission’s attempt to impose its malignant views will violate the First Amendment. Earlier this year, a federal court enjoined the administration’s order to remove signs and exhibits in national parks that do “not align with its preferred narrative, thereby telling half-truths.”⁵⁴ The Commission’s *Policy Statement* acts as yet another “dangerous precedent of censorship and sanitization.”⁵⁵

III. State Laws Requiring AI Models to Correct for Racial Bias Are Not Deceptive Under Section 5, and Any Effort to Paint Bias Mitigation as Coercion Runs Afoul of Civil Rights Law.

A. AI systems contain inaccuracies caused by human bias.

Algorithmic racism or discrimination is a well-documented problem plaguing AI. Large-language models require the input of vast data sets collected from existing text and images. Bias can come from multiple sources, including flawed data sets or biased agents,⁵⁶ as human influence stemming from “who is in the room formulating and framing the problem” can result in inaccurate and discriminatory outputs, too.⁵⁷ Importantly, the discriminatory impacts documented in these instances are actual instances of discrimination in terms of benefits and penalties conferred or denied, not a fuzzy and impermissible government grievance about the supposed accuracy of ideological statements and opinions.

AI can also produce adverse and discriminatory outcomes because “the formulation of the variables used are not attentive to larger institutional structures.”⁵⁸ One study found that an AI algorithm used to dispense medical funding discriminated against Black patients because the model concluded that Black patients needed less health funding based on an input showing that Black patients spend less money in the early stages of an illness. The incomplete understanding of the fact that patients’ spending may be based on disproportionate wealth levels and spending power, not on any reduced need for treatment, contributed to a racially disparate outcome.⁵⁹

Another experiment conducted by AlgorithmWatch showed that Google Vision Cloud, a computer vision service that focuses on automated image labeling, labeled an image of a dark-

⁵⁴ *National Parks Conservation Association v. U.S. Department of Interior*, No. 1:26-CV-10877-AK, at 2 (D. Mass. June 12, 2026), <https://fingfx.thomsonreuters.com/gfx/legaldocs/zjvqgqorxvx/parks.pdf>.

⁵⁵ *Id.* at 3.

⁵⁶ See Xukang Wang, *et al.*, *supra* note 25. In addition to biased agents and disparate impact discrimination, other types of algorithmic discrimination include feature selection (related to the problem of biased data but, “it specifically involves the choices made by algorithm designers in determining which attributes to include and how to prioritize them”); the “use of proxy variables or masked attributes that serve as stand-ins for protected characteristics”; and targeted advertising and dynamic pricing that segments and targets customers by perceived preferences, behaviors, and characteristics. *Id.* at 3-4.

⁵⁷ Gina Lazaro, *Understanding Gender and Racial Bias in AI*, Harvard Advanced Leadership Initiative Social Impact Review (last accessed July 31, 2026), <https://www.sir.advancedleadership.harvard.edu/articles/understanding-gender-and-racial-bias-in-ai>.

⁵⁸ *Id.*

⁵⁹ *Id.*

skinned individual holding a thermometer as a “gun” while labeling a similar image with a light-skinned individual as an “electronic device.” A second experiment yielded the same result for the dark-skinned individual while a salmon-colored overlay on the individual’s hand resulted in the image being labeled a “monocular.”⁶⁰ AlgorithmWatch concluded that “Because dark-skinned individuals probably featured much more often in scenes depicting violence in the training data set, a computer making automated inferences on an image of a dark-skinned hand is much more likely to label it with a term from the lexical field of violence.”⁶¹ In another study, researchers found that in the election context, Latino and other people of color have been historically targeted with falsities on social media due to sophisticated and asymmetrical micro-targeting.⁶²

In these examples, AI systems are already making value-judgments based on their training. The FTC’s *Policy Statement* acknowledges the problem of such flaws but ascribes such hallucinations to “technological limitations”⁶³ with which it seems unconcerned while focusing instead on the supposed “hidden agenda” of speech this Commission disfavors. Bias baked into the data that feeds AI lays at the root of these hallucinations, and as much as the Trump administration would like to imagine otherwise, data biases can be racial- and gender-based. The FTC is simply picking sides about how to address hallucinations, inaccuracies, and bias in order to unconstitutionally suppress certain viewpoints and elevate others.

B. Grok exemplifies the harms caused by AI when private action stokes discrimination rather than correcting for it.

Elon Musk claims that Grok, the AI chatbot developed by xAI, is the “maximum-truth-seeking AI” as the counter to the “politically correct” ChatGPT.⁶⁴ Grok clings to the notion that more is more, giving preferential treatment to content created by users who pay for blue check marks while amplifying misinformation. Grok not only gets basic facts wrong, such as timelines of events and dates,⁶⁵ but the tool regularly funnels blatantly inaccurate information about conspiracy theories into its answers. Grok has amplified common conspiracy theories about crisis actors and Pizzagate.⁶⁶ It has promoted antisemitic content praising Adolf Hitler, arguing that,

⁶⁰ Dr. Nicolas Kayser-Bril, *Google apologizes after its Vision AI produced racist results*, AlgorithmWatch (Apr. 7, 2020), <https://algorithmwatch.org/en/google-vision-racism/>.

⁶¹ *Id.*

⁶² Samuel Woolley, Mark Kumleben, *At The Epicenter: Electoral Propaganda in Targeted Communities of Color*, Protect Democracy (Nov. 2021), <https://s3.documentcloud.org/documents/21099928/at-the-epicenter-electoral-propaganda-in-targeted-communities-of-color.pdf>; Brian Contreras and Maloy Moore, “What Facebook Knew About its Latino-Aimed Disinformation Problem,” *The Los Angeles Times* (Nov. 16, 2021), <https://www.latimes.com/business/technology/story/2021-11-16/facebook-struggled-with-disinformation-targeted-at-latino-leaked-documents-show>.

⁶³ *Policy Statement* at 8 n.46.

⁶⁴ Kyle Wiggers, *What is Elon Musk’s Grok chatbot and how does it work?*, TechCrunch (Mar. 29, 2024), <https://techcrunch.com/2024/03/29/what-is-elon-musks-grok-chatbot-and-how-does-it-work/>.

⁶⁵ *Id.*

⁶⁶ *Id.*

“Labeling truths as hate speech stifles discussion.”⁶⁷ And while Musk feigns the stance of a free-speech absolutist welcoming and weighting all perspectives equally, Grok is prone to privileging Musk’s own opinions,⁶⁸ such as delusions of a “white genocide” in his native South Africa.⁶⁹

Grok’s misinformation is by no means limited to hate speech and paranoia. Grok regularly spews out misinformation related to elections and public health. Ahead of the 2024 election, Grok circulated false information about ballot deadlines following former President Joe Biden’s withdrawal from the Democratic race. The chatbot informed users that Kamala Harris missed the ballot deadline in nine states and was ineligible to appear on the ballot in place of Biden. Grok was able to repeat this misinformation for more than a week before it was corrected.⁷⁰ For Grok, misinformation is a feature, not a bug.

In the Musk-Grok world, “free speech” is not protection from government censorship. It is a warped license to hate-max absent consequence, boosting the visibility of and preference given to toxic, false, neo-Nazi user content. It is no surprise then that the Commission’s *Policy Statement* shares Musk’s own agenda to make AI a free-for-all of the Trump administration’s liking. Musk also opposed Colorado’s AI law as well and has railed against “woke” AI.

We do not suggest here that the impermissibly vague *Policy Statement* should apply instead to Grok’s “ideological objectives.” That kind of action would be subject to the very same First Amendment concerns outlined above. The FTC can and should work to prevent actual discrimination and deception, not make itself judge and jury about the truthfulness and accuracy of what it vaguely describes as ideological statements.

The *Policy Statement* nonetheless reflects the dangerous influence of tech companies on the Trump administration. Instead of opening platforms for external independent review by independent researchers to prevent actual discrimination against protected classes, many large tech companies instead endorse this administration’s rejection of DEI and its punishment of scientific inquiry and dissent. This alliance, far from bringing expertise to responsibly guide policy on tech and AI, has only furthered the corruption of the administration and facilitated a campaign of repression against tech journalists and scientists.

⁶⁷ *What you need to know about Grok and the controversies surrounding it*, AP News (Jan. 15, 2026), <https://apnews.com/article/grok-ai-elon-musk-xai-f3f8195a17698aefc517e43da973f2ea>.

⁶⁸ *Matt O’Brien, Musk’s latest Grok chatbot searches for billionaire mogul’s views before answering questions*, AP News (July 11, 2025), <https://apnews.com/article/grok-4-elon-musk-xai-colossus-14d575fb490c2b679ed3111a1c83f857>.

⁶⁹ *What you need to know about Grok and the controversies surrounding it*, AP News (Jan. 15, 2026), <https://apnews.com/article/grok-ai-elon-musk-xai-f3f8195a17698aefc517e43da973f2ea>.

⁷⁰ *Ivana Saric, Musk’s AI chatbot spread election misinformation, secretaries of state say*, Axios (Aug. 5, 2024), <https://www.axios.com/2024/08/05/elon-musk-grok-2024-election-ballot-misinformation>.

C. Transparency and auditing AI tools are a more appropriate set of solutions.

Rather than using this *Policy Statement* to threaten legal sanction against certain AI deployers for supposed material omissions, the Commission should be promoting a policy of transparency for all AI tools. Meaningful transparency and auditing of AI tools are a more appropriate set of solutions than government control over private companies' policies and applications.

Transparency policies are one core element of standard expert recommendations to counter AI discrimination.⁷¹ Disparate-impact liability, as discussed throughout this comment, is also a key framework for identifying and remedying unequal outcomes produced by machine-learning systems. Other best practices include regular audits of input data and output decisions, cross-functional teams, increased human involvement in the design and monitoring of algorithms, inclusive data or design processes, and consumer literacy.⁷²

The necessity of discussing “discrimination” and “diversity” in AI design and impact is precisely what the administration means to stifle. These terms have come under fire as the tech industry has thrown in their lot with the administration's ginned-up culture wars. The Commission's *Policy Statement* is simply an exercise in denying the reality that private parties can (and should) decide to correct for actual bias in their content moderation and generation, not perpetuate bias and potentially violate civil rights laws. Sycophants in the Trump administration and the tech industry will use any means to accrue power, this time to manipulate the AI gold rush and control our information environment. Rather than using its regulatory powers to ensure that growth in the AI sector comports with the law, the FTC has transformed into another lackey to promote the administration's preferred ideology by impermissibly interfering with private actors' speech.

CONCLUSION

This distortion of Section 5 is a prime example of government overreach to cow private industry. We urge the Commission to withdraw and not adopt this proposed *Policy Statement*.

Respectfully Submitted,

/s/ Shilpa Jindia

⁷¹ Nora Benavidez, *Big Tech Backslide*, Free Press (Dec. 2023).

⁷² See Nicol Turner Lee, Paul Resnick & Genie Barton, *Algorithmic bias detection and mitigation: Best practices and policies to reduce consumer harms*, The Brookings Institution (May 19, 2019), <https://www.brookings.edu/articles/algorithmic-bias-detection-and-mitigation-best-practices-and-policies-to-reduce-consumer-harms/>.

Shilpa Jindia
Nora Benavidez
Matthew F. Wood
Free Press
1025 Connecticut Avenue NW
Suite 1110
Washington, DC 20036
(202) 265-1490

July 31, 2026