

Board of Governors of the Federal Reserve System.

Erin Cayce,

Assistant Secretary of the Board.

[FR Doc. 2026-13690 Filed 7-6-26; 8:45 am]

BILLING CODE:P

FEDERAL TRADE COMMISSION

[File No. P264200]

Policy Statement Concerning the Suppression of Accuracy in Artificial Intelligence Systems

AGENCY: Federal Trade Commission.

ACTION: Proposed policy statement; request for comments.

SUMMARY: The Federal Trade Commission (“Commission”) is proposing a policy statement regarding the application of the prohibition on deceptive acts or practices in section 5 of the Federal Trade Commission Act to companies that market artificial intelligence (“AI”) systems.

DATES: Comments must be received on or before Friday, July 31, 2026.

ADDRESSES: Members of the public may file a comment online or on paper by following the instructions in the Comment Submissions part of the **SUPPLEMENTARY INFORMATION** section below. Write “AI Policy Statement; Matter No. P264200” on your comment and file your comment online at <https://www.regulations.gov/docket/FTC-2026-0859>. If you prefer to file your comment on paper, mail your comment to the following address: Federal Trade Commission, Office of the Secretary, 600 Pennsylvania Avenue NW, Mail Stop H-144 (Annex P), Washington, DC 20580.

SUPPLEMENTARY INFORMATION:

I. Summary

Artificial Intelligence (AI) is reshaping how Americans consume information, educate our children, and perform our jobs.¹ As they have marketed their remarkable breakthroughs to the public, AI companies have spent years representing explicitly and implicitly that their systems aim to produce the best output—output that faithfully and accurately achieves users’ stated

objectives and the built-in objectives that users expect in the AI system—that is possible within their technological and resource constraints. Because of these representations and the inherent nature of the products and services in question, 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.

Nonetheless, an AI company might be tempted to alter or steer the output of its systems contrary to consumers’ reasonable expectations for various reasons, including attempted compliance with a State law, such as Colorado’s recently revised Artificial Intelligence Act. But steering an AI system in this manner may deceive consumers in violation of section 5 of the FTC Act. That is true even if the deceptive steering is done in an effort to comply with State laws. Of course, a company may be able to avert potential deception by making truthful, non-misleading representations about the aims of its model. But such representations would need to make clear the AI company is prioritizing objectives different than those consumers requested or would otherwise expect.

II. Background

Artificial intelligence is an umbrella term covering a universe of different tools and systems now used across nearly every sector of the economy and in most people’s daily lives.² Regardless of whether one has in mind modern large language models, AI applications applying large-language models to particular purposes, or a not-yet-developed superintelligence, the hallmark of a successful AI system is an ability to accurately “make predictions, recommendations, or decisions” for its user consistent with a given objective.³ Their utility can then be judged by how well the solution matches consumers’ objectives.

On a fundamental level, these products and services solve problems for people. For many Americans, they are becoming part and parcel of daily life, relied upon for research, analysis, and advice on both personal and business issues. This development reflects consumers’ confidence that AI

companies are bound by the same rules as other American companies: they will deal with consumers in good faith and with no hidden agenda, and they will not work in the shadows to punish those who hold opinions contrary to theirs. In large part, that confidence itself reflects the basic fact that the United States is the global standard-bearer for AI technologies.⁴ That American companies dominate every layer of the AI ecosystem should not be taken for granted. Geopolitical rivals are investing heavily in this sphere, hoping to inject their own companies and values into the marketplace. American dominance is thus about more than winning some abstract race. It is about ensuring Americans can continue to feel their values are being respected as they interact with, and benefit from, an AI-powered economy.

Under the President’s leadership, the Trump-Vance Administration has taken decisive steps to sustain America’s global AI dominance⁵ by removing regulatory barriers that impede AI innovation.⁶ Excessive AI regulation would undermine American AI supremacy by deterring and suppressing the same ingenuity responsible for making American AI great.⁷ At the same time, however, as President Trump’s National Policy Framework for Artificial Intelligence recognizes, responsible AI innovation can co-exist with prudent guardrails. For example, AI technologies should protect children and empower

⁴ See J.D. Vance, Remarks by the Vice President at the Artificial Intelligence Action Summit in Paris, France (Feb. 11, 2025), www.presidency.ucsb.edu/documents/remarks-the-vice-president-the-artificial-intelligence-action-summit-paris-france [hereinafter “Vance AI Summit Remarks”].

⁵ See, e.g., AI Action Plan, *supra* note 1; E.O. 14318, Accelerating Federal Permitting of Data Center Infrastructure, 90 FR 35385, 35385 (July 23, 2025) (prioritizing the rapid and efficient buildout of AI data centers and the infrastructure that power them by easing Federal regulatory burdens); E.O. 14319, 90 FR at 35389 (noting agencies have an “obligation not to procure [AI] models that sacrifice truthfulness and accuracy to ideological agendas”); E.O. 14320, Promoting the Export of the American AI Technology Stack, 90 FR 35393, 35393 (July 23, 2025) (establishing a coordinated national effort to support the American AI industry by promoting the export of full-stack American AI technology packages).

⁶ See, e.g., E.O. 14148, Initial Rescissions of Harmful Executive Orders and Actions, 90 FR 8237, 8240 (Jan. 20, 2025), <https://www.govinfo.gov/content/pkg/FR-2025-01-28/pdf/2025-01901.pdf> (revoking E.O. 14110, Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence, 88 FR 75191 (Oct. 30, 2023)); E.O. 14179, Removing Barriers to American Leadership in Artificial Intelligence, 90 FR 8741, 8741–42 (Jan. 23, 2025), www.govinfo.gov/content/pkg/FR-2025-01-31/pdf/2025-02172.pdf.

⁷ E.O. 14179, 90 FR at 8741; Vance AI Summit Remarks, *supra* note 4.

¹ Executive Office of the President, *Winning the Race: America’s AI Action Plan* at 1 (July 2025), <https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf> [hereinafter “AI Action Plan”]; The Council of Economic Advisors, Executive Office of the President, *Artificial Intelligence and the Great Divergence* at 1 (Jan. 2026), <https://www.whitehouse.gov/wp-content/uploads/2026/01/Artificial-Intelligence-and-the-Great-Divergence-5.pdf>.

² See, e.g., E.O. 14319, Preventing Woke AI in the Federal Government, at sec. 1, 90 FR 35389, 35389 (July 23, 2025) (“Artificial intelligence . . . will play a critical role in how Americans of all ages learn new skills, consume information, and navigate their daily lives.”).

³ 15 U.S.C. 9401(3).

parents.⁸ And most relevant to this statement, AI technologies should not be used to silence or censor lawful expression or dissent.⁹ Importantly, President Trump's proposed approach is a national AI framework, protecting innovation and competition by providing national regulatory clarity and certainty and avoiding a balkanized or patchwork regulatory approach driven by the States—or, most dangerously, imposed by certain anti-innovation State governments on the rest of the country.

America's AI Action Plan and other critical executive actions strike that balance between accelerating innovation and protecting Americans,¹⁰ but anti-innovation States' recent efforts to regulate AI are concerning. The growing number of enacted and proposed State AI laws threatens to create a patchwork of regulatory regimes and compliance challenges for American companies.¹¹ For example, some States have enacted laws regulating the outputs of various AI models, ultimately requiring American companies to embed "ideological bias within [their AI] models."¹²

On December 11, 2025, President Trump signed Executive Order 14365, "Ensuring a National Policy Framework for Artificial Intelligence."¹³ The Executive Order builds on the Administration's prior work to encourage adoption of AI and remove regulatory barriers, aiming to "sustain and enhance the United States' global AI dominance" by establishing a "minimally burdensome national policy framework for AI—not 50 discordant State ones."¹⁴ The Executive Order

directs the Commission to issue this enforcement policy statement clarifying the application of section 5 of the Federal Trade Commission Act ("FTC Act") to AI models, and in particular, address how State laws requiring alterations to the accurate outputs of AI models can conflict with the requirements of the FTC Act.¹⁵

III. The Commission's Authority Under Section 5 of the FTC Act

Nearly ninety years ago, Congress gave the Commission the authority to protect consumers from "unfair or deceptive acts or practices in or affecting commerce."¹⁶ Importantly, Congress did not provide State-law safe harbors—even where a State's laws shape the provision of those products or services, companies must comply with section 5.

Section 5 of the FTC Act's prohibition on deceptive acts or practices secures consumers' right to truth in the marketplace.¹⁷ In our 1983 Policy Statement on Deception, we said "the Commission will find deception if there is a representation, omission or practice that is likely to mislead the consumer acting reasonably in the circumstances, to the consumer's detriment."¹⁸ This test, which "undergird[s] all deception cases," consists of three elements.¹⁹ "First, there must be a representation, omission or practice that is likely to mislead the consumer."²⁰ "Second, we examine the practice from the perspective of a consumer acting reasonably in the circumstances."²¹ And "[t]hird, the representation, omission or practice must be a material one," such that it is "likely to affect the consumer's conduct or decision with regard to a product or service."²²

Federal courts have adopted the same principles in assessing whether an act or practice is deceptive.²³ Courts have also held the Commission's deception

authority covers implied misrepresentations, half-truths,²⁴ and instances where only a "significant minority" of consumers are misled.²⁵ Moreover, courts have recognized the Commission's deception authority requires the entity to substantiate the representation it makes about its products or services with sufficient evidence.²⁶ At its core, the Commission's deception authority is based on "well-established principles of advertising law and common sense."²⁷

The Commission has applied the principles underlying its section 5 deception authority to address false or misleading claims in a multitude of factual circumstances and across a wide variety of products and services,²⁸ including new and evolving technologies.²⁹ And we have taken aim at companies that failed to provide adequate disclosures,³⁰ as well as

²⁴ See, e.g., *Novartis Corp. v. FTC*, 223 F.3d 783, 787 (D.C. Cir. 2000); *Fanning v. FTC*, 821 F.3d 164, 170–72 (1st Cir. 2016); *Kraft, Inc. v. FTC*, 970 F.2d 311, 322 (7th Cir. 1992).

²⁵ *ECM Biofilms, Inc. v. FTC*, 851 F.3d 599, 610 (6th Cir. 2017) (quoting *POM Wonderful, LLC v. FTC*, 777 F.3d 478, 490 (D.C. Cir. 2015)). A small proportion of consumers, such as 10%, can constitute a "significant minority." See, e.g., *Telebrands Corp.*, 140 F.T.C. 278, 325 (2005), *aff'd*, 457 F.3d 354 (4th Cir. 2006) (10.5% is significant); *Firestone Tire & Rubber Co. v. FTC*, 481 F.2d 246, 249 (6th Cir. 1973) (10–15% is significant).

²⁶ See, e.g., *POM Wonderful*, 777 F.3d at 490–91; *FTC v. Direct Mktg. Concepts, Inc.*, 624 F.3d 1, 8 (1st Cir. 2010); *FTC v. Nat'l Urological Grp., Inc.*, 645 F. Supp. 2d 1167, 1190 (N.D. Ga. 2008), *aff'd*, 356 F. App'x 358 (11th Cir. 2009).

²⁷ *FTC v. Com. Planet, Inc.*, 878 F. Supp. 2d 1048, 1083 (C.D. Cal. 2012), *aff'd in part*, 642 F. App'x 680 (9th Cir. 2016); see also *FTC v. Johnson*, 96 F. Supp. 3d 1110, 1142 (D. Nev. 2015) (rejecting "fair notice" argument as to section 5(a) liability and noting there is "extensive case law and guidance on what constitutes a deceptive act or practice under Section 5(a)").

²⁸ See, e.g., *FTC v. Fleetcor Techs., Inc.*, 620 F. Supp. 3d 1268, 1289–1313 (N.D. Ga. 2022) (fuel cards and discounts), *aff'd sub nom FTC v. Corpay, Inc.*, 164 F.4th 807 (11th Cir. 2026); *FTC v. D-Link Sys., Inc.*, No. 3:17–CV–00039–JD, 2017 WL 4150873, at *2 (N.D. Cal. Sept. 19, 2017) (data security and protections for IoT devices); *Pom Wonderful*, 777 F.3d at 490–500 (dietary supplements); *Com. Planet, Inc.*, 878 F. Supp. 2d at 1083 (online web pages and negative option plans).

²⁹ See, e.g., *Stipulated Order, FTC v. Amazon.com, Inc.*, No. 2:23–CV–00932–JHC (W.D. Wash. Sept. 25, 2025), Dkt. No. 535 (enrolling consumers into online shopping memberships without consumers' express informed consent); *Compl., United States v. Facebook*, No. 19–cv–2184 (D.D.C. July 24, 2019), Dkt. No. 1 (misrepresenting that users would have to "turn[] on" facial-recognition technology); *Compl., FTC v. New Consumer Sols., LLC*, No. 1:15–cv–01614 (N.D. Ill. Feb. 23, 2015), Dkt. No. 1 (misrepresenting ability to detect melanoma by analyzing pictures of consumers' skin); *Compl., In re Sears Holdings Mgmt. Corp.*, FTC Docket No. C–4264 (Sept. 9, 2009) (tracking consumers' "online browsing").

³⁰ See, e.g., *FTC v. LendingClub Corp.*, No. 18–CV–02454–JSC, 2020 WL 2838827, at *6 (N.D. Cal. June 1, 2020) (inadequate disclosures of loan

⁸ The White House, National Policy Framework for Artificial Intelligence (Mar. 20, 2026).

⁹ *Id.*

¹⁰ See *supra* note 6.

¹¹ See The White House, Fact Sheet: President Donald J. Trump Ensures a National Policy Framework for Artificial Intelligence (Dec. 11, 2025), <https://www.whitehouse.gov/fact-sheets/2025/12/fact-sheet-president-donald-j-trump-ensures-a-national-policy-framework-for-artificial-intelligence/> [hereinafter "AI National Policy Fact Sheet"] ("State legislatures have introduced over 1,000 different AI bills[.]"; E.O. 14319, 90 FR at 35389).

¹² See, e.g., E.O. 14365, Ensuring a National Policy Framework for Artificial Intelligence, 90 FR 58499, 58499 (Dec. 11, 2025) ("For example, a new Colorado law banning 'algorithmic discrimination' may even force AI models to produce false results in order to avoid a 'differential treatment or impact' on protected groups."); AI National Policy Fact Sheet, *supra* note 11 ("States such as California and Colorado are considering requiring AI companies to censor outputs and insert left-wing ideology in their programming."). Colorado has since materially revised the law referred to by this Executive Order, but the new version poses many of the same concerns. See *infra* note 43.

¹³ E.O. 14365, 90 FR at 58499.

¹⁴ *Id.*

¹⁵ *Id.* at 58500 (Section 7).

¹⁶ Wheeler-Lea Act, Public Law 75–447, 52 Stat. 111 (1938) (codified in relevant part at 15 U.S.C. 45(a)(1)).

¹⁷ *In re Int'l Harvester Co.*, 104 F.T.C. 949, 1056 (1984).

¹⁸ FTC Policy Statement on Deception ("Deception Policy Statement"), 103 F.T.C. 174, 183 (1984), (appended to *In re Clifdale Assocs., Inc.*, 103 F.T.C. 110 (1984)), <https://www.ftc.gov/legal-library/browse/ftc-policy-statement-deception>.

¹⁹ *Id.* at 175.

²⁰ *Id.*

²¹ *Id.*

²² *Id.* Material information affects a consumer's decision to purchase a product but may also affect other kinds of consumer conduct. *Id.* at 182 n.45 (citing the Commission's complaint in *Volkswagen of America*, 99 F.T.C. 446 (1982)).

²³ See, e.g., *FTC v. Pukke*, 53 F.4th 80, 104 (4th Cir. 2022); *FTC v. Moses*, 913 F.3d 297, 306 (2d Cir. 2019); *FTC v. E.M.A. Nationwide, Inc.*, 767 F.3d 611, 631 (6th Cir. 2014).

governmental and private standard-setting entities that sought to restrict truthful advertising.³¹

IV. Deceiving Consumers as to the Objectives of an AI System Is a Section 5 Violation

AI products and services are not exempt from section 5's long-established, and generally applicable reach. The Commission has already brought a number of recent enforcement actions to combat deceptive claims in connection with AI-powered scams, as well as material misrepresentations about the performance, efficacy, and characteristics of AI products.³² For instance, the Commission recently invoked well-established section 5 principles to address an AI company's misleading claims about the ability of its conversational AI tool to replace human customer service representatives.³³ Although the Commission is careful to avoid unduly burdening innovation in the AI industry,³⁴ the Commission will

origination fees); see also *FTC v. Cyberspace.Com LLC*, 453 F.3d 1196, 1200 (9th Cir. 2006) ("A solicitation may be likely to mislead by virtue of the net impression it creates even though the solicitation also contains truthful disclosures.").

³¹ See, e.g., *In re Am. Med. Ass'n*, 94 F.T.C. 701, 235 (1979) (restricting physician advertising).

³² See, e.g., Consent Order, *In re Workado, LLC*, No. C-4822 (F.T.C. Aug. 21, 2025) (deceptive claims regarding accuracy or efficacy of its AI-content detection products), https://www.ftc.gov/system/files/ftc_gov/pdf/ContentatScaleAI-DecisionandOrder.pdf; Stipulated Order, *FTC v. TheFBAMachine Inc.*, No. 24-cv-6635-JXN-LDW (D.N.J. July 31, 2025), Dkt. No. 162 (falsely guaranteeing consumers could make money operating online storefronts using AI-powered software); 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_o.pdf; 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; Stipulated Order, *FTC v. Evolv Techs. Holdings, Inc.*, No. 1:24-cv-12940-PGL (D. Mass. Nov. 26, 2024), Dkt. No. 2 (false claims about the efficacy of its AI-powered security screening system for detecting weapons).

³³ See Compl., *FTC v. Air Ai Techs., Inc.*, No. 2:25-cv-03068-SMB (D. Ariz. Aug. 25, 2025), Dkt. No. 1 (deceptive claims regarding efficacy and profitability related to the company's "conversational AI" product).

³⁴ In fact, the Commission has already taken steps to address these unnecessary burdens, for example, by setting aside the previous Commission's order against a generative AI company. See Order Reopening and Setting Aside Order at 5, *In re Rytr LLC*, FTC Docket C-4806 (Dec. 22, 2025), https://www.ftc.gov/system/files/ftc_gov/pdf/Rytr-Order.pdf ("Where actors use AI to violate the law or deceive consumers about the capabilities of their generative AI, they should be held accountable, as the FTC has done and will continue to do. But that is not the case here. Rushing in to impose aggressive law enforcement unsupported by facts or law is improper and is not in the public interest.").

continue to protect consumers and the market from deceptive or unfair business practices, including those relying on AI systems or technologies, by using well-established section 5 principles.

In marketing their products as problem-solving tools, AI companies have represented explicitly and implicitly that their AI systems aim to produce the best output possible given technological and resource constraints.³⁵ AI companies aggressively market themselves as building products and services that distill all of human knowledge to solve problems, large and small, in furtherance of users' given objectives.³⁶ This is because companies know their success depends on their

³⁵ See, e.g., OpenAI, *Why Language Models Hallucinate* (Sept. 5, 2025) ("At OpenAI, we're working hard to make AI systems more useful and reliable. Even as language models become more capable, one challenge remains stubbornly hard to fully solve: hallucinations."), <https://openai.com/index/why-language-models-hallucinate/>; Anthropic, *Claude Product Overview*, <https://claude.com/product/overview> (last visited May 22, 2026) ("Meet your thinking partner . . . Tackle any big, bold, bewildering challenge with Claude . . . The AI for problem solvers . . . Tackle your toughest work . . . Like an expert in your pocket . . . There's never been a worse time to be a problem, or a better time to be a problem solver . . . What problem are you up against?"); Anthropic, *How People Ask Claude for Personal Guidance*, <https://www.anthropic.com/research/claude-personal-guidance> <https://www.anthropic.com/research/claude-personal-guidance> (Apr. 30, 2026) ("Helpfulness is one of Claude's most important traits. Speaking with Claude should be akin to a conversation with a brilliant friend, one who will speak frankly to a person about their situation, providing information grounded in evidence."); X.ai, *Chat With Grok*, <https://x.ai/grok> (last accessed May 20, 2026) ("Grok is your truth-seeking AI companion for unfiltered answers with advanced capabilities in reasoning, coding, and visual processing.").

³⁶ See, e.g., Anthropic, *Claude*, <https://claude.com> (last visited May 22, 2026) ("What is Claude and how does it work? Claude is an artificial intelligence, trained by Anthropic using Constitutional AI to be safe, accurate, and secure—the trusted assistant for you to do your best work. . . . What should I use Claude for? If you can dream it, Claude can help you do it. Claude can process large amounts of information, brainstorm ideas, generate text and code, help you understand subjects, coach you through difficult situations, simplify your busywork so you can focus on what matters most, and so much more."); OpenAI, *ChatGPT Capabilities Overview* <https://help.openai.com/en/articles/9260256-chatgpt-capabilities-overview> (last visited Mar. 17, 2026) (touting ChatGPT's ability to solve problems); OpenAI, *Extracting Insights with ChatGPT Data Analysis*, <https://help.openai.com/en/articles/9213685-extracting-insights-with-chatgpt-data-analysis> (last visited Mar. 17, 2026) (claiming ChatGPT can identify missing or outlier data); OpenAI, *ChatGPT for Healthcare*, <https://help.openai.com/en/articles/20001046-chatgpt-for-healthcare> (last visited Mar. 17, 2026) ("ChatGPT for Healthcare can pull from millions of peer-reviewed studies, clinical guidelines, and public health sources" to "[g]et clinical answers with citations" that "perform[] better than human baselines across every role measured").

products' utility, and an AI system's utility is a function of its ability to accomplish the objectives users ask it to accomplish.³⁷

Consumers share this understanding of the purpose of AI systems and have learned to trust these systems across a broad swath of applications, from education to health, finance, relationships or a myriad of other uses.³⁸ In fact, according to one major AI developer, consumers accept AI system outputs without conducting any further fact checking over 90% of the time, even though these systems purport to only strive for, rather than guarantee, complete accuracy.³⁹ In other words, consumers ask AI systems for help answering deeply personal and important questions, and they overwhelmingly trust that the system is designed to provide an answer tailored specifically to their objectives.⁴⁰ And because of both the companies' marketing and the inherent value proposition of an AI system, this widely

³⁷ Similarly, an ability to make accurate predictions, correctly solve problems, or make good decisions hinges on AI systems operating without undisclosed ideological biases that contradict explicit and implicit claims of accuracy and neutrality. Though the exact line of what constitutes bias may be difficult to draw, forcing an AI system to prioritize a singular ideological objective over all else is on the wrong side. See Concurring Statement of Comm'r Andrew N. Ferguson, *In re Intellivision Techs. Corp.*, Matter No. 2323023, at 1 (F.T.C. Dec. 3, 2024), https://www.ftc.gov/system/files/ftc_gov/pdf/intellivision-ferguson-concurrence.pdf.

³⁸ Paulo Vargas, *Your Teen Is Probably Using AI for Homework*, Digital Trends (Mar. 2, 2026) ("More than half of U.S. teens ages 13 to 17 say they have turned to chatbots . . . for school tasks."), <https://www.digitaltrends.com/computing/your-teen-is-probably-using-ai-for-homework/>; Walt Williams, *Survey: Consumers Increasingly Turn to AI for Financial Advice*, ABA Banking Journal (Sept. 2, 2025) ("A majority of consumers are turning to artificial intelligence for financial advice . . ."), <https://bankingjournal.aba.com/2025/09/survey-consumers-increasingly-turn-to-ai-for-financial-advice/>; Theo Burman, *Americans Are Using AI To Diagnose Their Health Issues*, Newsweek (July 20, 2025) ("Both clinicians and patients are using artificial intelligence more and more to help diagnose illness and injuries . . ."), <https://www.newsweek.com/ai-healthcare-diagnosis-chatgpt-doctor-2100091>; *How People Ask Claude for Personal Guidance*, Anthropic (Apr. 30, 2026) (finding 12% of all conversations with the Claude program concerned relationship navigation, 6% concerned personal development, and 4% concerned spirituality), <https://www.anthropic.com/research/claude-personal-guidance>.

³⁹ Ted Ladd, *Anthropic: 91% of Users Do Not Fact-Check AI. Let's Fix That*, Forbes (Mar. 5, 2026), <https://www.forbes.com/sites/teidladd/2026/03/05/anthropic-91-of-users-do-not-fact-check-ai-lets-fix-that/>.

⁴⁰ Reasonable consumers would not expect, of course, that an AI system would output content that is clearly and unequivocally not protected by the First Amendment, such as child pornography. And nothing in this statement shall be construed to prohibit companies from imposing limits on model use to prevent cybersecurity attacks.

held consumer expectation is eminently reasonable.

Of course, AI systems are likely to simultaneously pursue multiple objectives, some explicit and some implicit. This will often be consistent with consumers' reasonable expectations. Users of an AI chatbot might, for example, reasonably expect the system to balance succinctness, clarity, relevance, accuracy, and other objectives in its attempt to produce the best output. And while truth and accuracy are in many cases implicit objectives in requests to AI systems, a user could request output that is intentionally inaccurate or that deprioritizes accuracy in favor of some other objective like entertainment.

These representations and baseline consumer expectations notwithstanding, AI companies might steer their systems' outputs toward objectives other than those that consumers ask for or expect. A company could be tempted, for example, to abuse consumer trust by training a model surreptitiously to produce ideologically motivated distortions in a response to a factual question, such as to correct what the developer believes are "historical injustices" in the facts.⁴¹ Or a company could be pressured by a State law to alter its technology's outputs. The original version of Colorado's Artificial Intelligence Act,⁴² for example, imposed a broad duty on AI companies to avoid output that might lead to disparate impacts in various contexts, including when a customer's foreseeable use of that output could itself create a disparate impact. And the revised version of that law explicitly provides AI companies can be held liable for discriminatory outcomes caused by their customers' use of their products.⁴³ It is predictable that an AI company might suppress accuracy and interpose other objectives, such as so-called "equity," to avoid liability under this law, but fail to disclose these ulterior objectives in order to hide the loss of accuracy they necessitate.⁴⁴ AI

companies may also be subject to public pressure to surreptitiously avoid politically inflammatory outputs, as well as the temptation to clandestinely modify their AI systems to further their own or their employees' political agendas.

In light of the above, the Commission believes AI companies that steer the outputs of their AI systems toward unexpected objectives, and away from the objectives set by or reasonably expected by users, are likely to deceive⁴⁵ consumers in violation of section 5 of the FTC Act.⁴⁶ AI companies have made both explicit and implicit representations that their systems aim to produce outputs that achieve users' objectives as faithfully and accurately as they can.⁴⁷ Those representations, absent adequate disclaimers or qualifications, are material as consumers rely on them and are likely to consider, in their decision as to which AI system to use or pay for, whether a particular system is designed to produce output accomplishing their objectives or to intentionally produce a worse output in service of another objective. Consumers may be induced to pay for a service that does not behave as advertised. Consumers may also be deceived into relying on a technology that, by design, produces worse outputs and may recommend suboptimal courses of action, not because of any technological or resource limitations, but because the AI developer's hidden agenda subverted consumers' objectives.

As is always the case, a company's motives for deceiving consumers are irrelevant to the application of section

from making hiring and other employment decisions based on merit and skill, their needs, or the needs of their customers because of the specter that such a process might lead to disparate outcomes, and thus disparate-impact lawsuits. . . . Disparate-impact liability imperils the effectiveness of civil rights laws by mandating, rather than proscribing, discrimination."'). Indeed, it seems likely a law restricting truthful speech because it might lead another person to commit disparate impact discrimination would not survive First Amendment scrutiny, but that is orthogonal to the FTC Act issues we address here.

⁴⁵ The Commission at this time takes no position on whether the practices discussed in this statement may also be unfair under the FTC Act.

⁴⁶ This conduct is distinct from the problem of incorrect output from AI systems, frequently called "hallucinations," that stem not from a design decision to prioritize objectives contrary to users' reasonable expectations, but from the technological and resource limitations AI systems necessarily reflect. While a company may, for example, unlawfully deceive consumers if it misrepresents the likelihood of such hallucinations, the Commission does not believe such hallucinations in and of themselves raise issues under section 5 or any other law the Commission enforces.

⁴⁷ See *supra* notes 35 & 36.

⁴⁸ Whether motivated by profit, shaping public opinion, or anything else, section 5 prohibits deceiving consumers. These prohibitions apply even when a company engages in a deceptive act or practice in order to comply with a State law. 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.⁴⁹ 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.

An AI company can, of course, take actions that shape consumer expectations in ways that make it unreasonable for consumers to believe the AI system is designed to achieve the objectives that consumers would otherwise expect. A company can, for example, clearly and conspicuously disclose that its systems are designed to produce outputs that prioritize certain objectives over what users request and otherwise expect. But such a disclaimer would have to be adequate to shift consumer expectations that would be based otherwise both on companies' explicit representations and on the inherent value proposition of AI systems as tools to solve human problems. An adequate disclaimer could not be buried in terms of service, for instance. It would have to clearly and conspicuously dispel the notion that the system is designed to give the best answer possible. Such a disclaimer would need to be prominent, and it is doubtful a one-time disclosure subsequently hidden away in fine print would suffice. The more the disclosure cuts against the reasonable expectations users would take away from other contexts, the more persistent and prominent it would need to be.⁵⁰ A prominent misrepresentation is unlikely to be remedied by a less prominent,

⁴⁸ *FTC v. Bay Area Bus. Council*, 423 F.3d 627, 635 (7th Cir. 2005) (deception); *In re Kraft, Inc.*, 114 F.T.C. 40, 122 (1991) (deception).

⁴⁹ *Schneidewind v. ANR Pipeline Co.*, 485 U.S. 293, 300 (1988) ("[S]tate law is pre-empted when it actually conflicts with federal law. Such conflict will be found when it is impossible to comply with both state and federal law, or where the state law stands as an obstacle of the full purposes and objectives of Congress." (internal quotations omitted)).

⁵⁰ *FTC v. Direct Mktg. Concepts, Inc.*, 624 F.3d 1, 12 (1st Cir. 2010) (explaining disclaimers "are not adequate to avoid liability unless they are sufficiently prominent and unambiguous to change the apparent meaning of the claims and to leave an accurate impression."); *Removatron Int'l Corp. v. FTC*, 884 F.2d 1489, 1497 (1st Cir.1989) ("[d]isclaimers or qualifications in any particular ad are not adequate to avoid liability unless they are sufficiently prominent and unambiguous to change the apparent meaning of the claims and to leave an accurate impression.").

⁴¹ See, e.g., Leo Briceno, *Anthropic's Moral Compass Architect Suggested AI Overcorrection Could Address Historical Injustices*, Fox News (Apr. 22, 2026), <https://www.foxnews.com/politics/anthropics-moral-compass-architect-suggested-ai-overcorrection-could-address-historical-injustices>.

⁴² Colo. Sen. Bill 24–205 (enacted May 17, 2024), *repealed and reenacted, as amended*, by Colo. Sen. Bill 26–189 (enacted May 14, 2026).

⁴³ Colo. Sen. Bill 26–189, 6–1–1707 (enacted May 14, 2026).

⁴⁴ See E.O. 14281, *Restoring Equality of Opportunity and Meritocracy*, at sec. 1, 90 FR 17537 (Apr. 23, 2025) ("Disparate-impact liability all but requires individuals and businesses to consider race and engage in racial balancing to avoid potentially crippling legal liability. . . . On a practical level, disparate-impact liability has hindered businesses

subsequent disclosure.⁵¹ Ultimately, it is the company’s responsibility to ensure compliance with section 5.

V. Comment Submissions

You can file a comment online or on paper. For the Commission to consider your comment, we must receive it on or before Friday, July 31, 2026. Write “AI Policy Statement; Matter No. P264200” on your comment. Your comment—including your name and your State—will be placed on the public record of this proceeding, including, to the extent practicable, on the <https://www.regulations.gov> website.

We encourage you to submit comments through the <https://www.regulations.gov> website. Postal mail addressed to the Commission will be subject to delay because of heightened security screening. If you prefer to file your comment on paper, write “AI Policy Statement; Matter No. P264200” on your comment and on the envelope, and send it via overnight service to: Federal Trade Commission, Office of the Secretary, 600 Pennsylvania Avenue NW, Mail Stop H–144 (Annex P), Washington, DC 20580.

Because your comment will be placed on the publicly accessible website at <https://www.regulations.gov>, you are solely responsible for making sure your comment does not include any sensitive or confidential information. In particular, your comment should not include sensitive personal information, such as your or anyone else’s Social Security number; date of birth; driver’s license number or other State identification number, or foreign country equivalent; passport number; financial account number; or credit or debit card number. You are also solely responsible for making sure your comment does not include sensitive

health information, such as medical records or other individually identifiable health information. In addition, your comment should not include any “trade secret or any commercial or financial information which . . . is privileged or confidential”—as provided by section 6(f) of the FTC Act, 15 U.S.C. 46(f), and FTC Rule 4.10(a)(2), 16 CFR 4.10(a)(2)—including competitively sensitive information such as costs, sales statistics, inventories, formulas, patterns, devices, manufacturing processes, or customer names.

Comments containing material for which confidential treatment is requested must be filed in paper form, must be clearly labeled “Confidential,” and must comply with FTC Rule 4.9(c). In particular, the written request for confidential treatment that accompanies the comment must include the factual and legal basis for the request and must identify the specific portions of the comment to be withheld from the public record. See FTC Rule 4.9(c). Your comment will be kept confidential only if the General Counsel grants your request in accordance with the law and the public interest. Once your comment has been posted on the <https://www.regulations.gov> website—as legally required by FTC Rule 4.9(b)—we cannot redact or remove your comment from that website, unless you submit a confidentiality request that meets the requirements for such treatment under FTC Rule 4.9(c), and the General Counsel grants that request.

The FTC Act and other laws the Commission administers permit the collection of public comments to consider and use in this proceeding, as appropriate. The Commission will consider all timely and responsive public comments it receives on or before Friday, July 31, 2026. For information

on the Commission’s privacy policy, including routine uses permitted by the Privacy Act, see <https://www.ftc.gov/site-information/privacy-policy>.

By direction of the Commission.

April J. Tabor,
Secretary.

[FR Doc. 2026–13628 Filed 7–6–26; 8:45 am]

BILLING CODE 6750–01–P

FEDERAL TRADE COMMISSION

Granting of Requests for Early Termination of the Waiting Period Under the Premerger Notification Rules

Section 7A of the Clayton Act, 15 U.S.C. 18a, as added by Title II of the Hart-Scott-Rodino Antitrust Improvements Act of 1976, requires persons contemplating certain mergers or acquisitions to give the Federal Trade Commission and the Assistant Attorney General advance notice and to wait designated periods before consummation of such plans. Section 7A(b)(2) of the Act permits the agencies, in individual cases, to terminate this waiting period prior to its expiration and requires that notice of this action be published in the **Federal Register**.

The following transactions were granted early termination—on the dates indicated—of the waiting period provided by law and the premerger notification rules. The listing for each transaction includes the transaction number and the parties to the transaction. The grants were made by the Federal Trade Commission and the Assistant Attorney General for the Antitrust Division of the Department of Justice. Neither agency intends to take any action with respect to these proposed acquisitions during the applicable waiting period.

EARLY TERMINATIONS GRANTED
[05/01/2026, 05/31/2026]

05/01/2026		
20251510		Garage Topco, LP; Cantaloupe, Inc.; Garage Topco, LP.
05/05/2026		
20261221	G	Bending Spoons S.p.A.; tractive GmbH; Bending Spoons S.p.A.
20261233	G	Quad-C Partners X, L.P.; Norwest Equity Partners X, LP; Quad-C Partners X, L.P.
20261245	G	Brookfield Oaktree Wealth Solutions Alternative Funds S.A.; Brookfield NERH Aggregator LLC; Brookfield Oaktree Wealth Solutions Alternative Funds S.A.
20261249	G	KKR North America XIV (Bloom) Blocker; Roark Capital Partners VI (T) LP; KKR North America XIV (Bloom) Blocker.
20261262	G	Peak Rock Capital Fund IV LP; BioRidge Pharma Holdco, LLC; Peak Rock Capital Fund IV LP.

⁵¹ See, e.g., Deception Policy Statement at *48, *supra* note 18 (“Depending on the circumstances, accurate information in the text may not remedy a false headline because reasonable consumer may glance only at the headline.”); see also *FTC v. Corpay, Inc.*, 164 F.4th 807, 836 (11th Cir. Jan. 6, 2026) (disclaimers or qualifying language may not correct a misleading impression when the

disclaimer “is small, ambiguous, or contradicted by the body of the ad”) (citing *FTC v. Cyberspace.com LLC*, 453 F.3d 1198, 1200–01 (9th Cir. 2006); *FTC v. Direct Mktg. Concepts, Inc.*, 624 F.3d 1, 12 (1st Cir. 2010) (explaining disclaimers “are not adequate to avoid liability unless they are sufficiently prominent and unambiguous to change the apparent meaning of the claims and to leave an

accurate impression”); *Removatron Int’l Corp. v. FTC*, 884 F.2d 1489, 1497 (1st Cir.1989) (“[d]isclaimers or qualifications in any particular ad are not adequate to avoid liability unless they are sufficiently prominent and unambiguous to change the apparent meaning of the claims and to leave an accurate impression.”).