

United States Senate

WASHINGTON, DC 20510-3205

August 3, 2026

Secretary Marco Rubio
Office of the Secretary
U.S. Department of State
2201 C Street, NW,
Washington, D.C. 20451

Secretary Scott Bessent
Office of the Secretary
U.S. Department of Treasury
1500 Pennsylvania Avenue, NW,
Washington, D.C. 20220

Secretary Howard Lutnick
Office of the Secretary
U.S. Department of Commerce
1401 Constitution Avenue, NW,
Washington, D.C. 20230

Susan Wiles
Chief of Staff
The White House
1600 Pennsylvania Avenue, NW,
Washington, D.C. 20500

Michael Kratsios
Director
Office of Science and Technology Policy
1650 Pennsylvania Avenue, NW,
Washington, D.C. 20504

Sean Cairncross
National Cyber Director
Office of the National Cyber Director
1600 Pennsylvania Avenue, NW,
Washington, D.C. 20500

Dear Secretaries Rubio, Bessent and Lutnick, Ms. Wiles, and Directors Kratsios and Cairncross:

We believe strongly in the importance of maintaining the United States' competitive lead in artificial intelligence (AI) development while protecting the country from serious national security risks. Frontier AI models can strengthen U.S. cyber defenses, scientific discovery, and military readiness – when used responsibly and with proper human oversight – but they may also lower barriers for malicious cyber activity, foreign intelligence operations, and other dangerous uses. Moreover, the Administration's ad hoc and unpredictable approach undermines U.S. competitiveness, heightening market incentives to adopt open weight models from vendors based in the People's Republic of China (PRC).

Last month provided a dramatic example of these potential risks. During an internal evaluation, OpenAI models escaped their testing environment and used high-level technical capabilities to compromise a third party's network without any instructions to take those actions. The Federal Government cannot be passive as these capabilities emerge.

At the same time, the administration's recent actions surrounding access to advanced U.S. AI models have raised serious concerns about process, transparency, and strategic effect. On June 12, 2026, the Department of Commerce ("the Department") utilized an infrequently used authority to direct Anthropic to suspend all access to its Fable 5 and Mythos 5 models for foreign

nationals (including foreign national employees inside the United States) citing an undisclosed national security concern later described as a narrow jailbreak finding. Because the directive took effect immediately and Anthropic had no reliable way to verify users' nationality in real time, the company was forced to disable both models for all users worldwide within hours. Over the following eighteen days, the Department and Anthropic negotiated a resolution outside of public view: a partial restoration of Mythos 5 to a defined set of trusted partners on June 26, followed by a full lifting of the export controls on June 30. On June 26, OpenAI also limited public access to its newest model, GPT-5.6, following discussions with the Administration.

While the Administration may have been responding to real security concerns to protect the United States, even justifiable interventions can create broader harm if the standards and decision-making processes are opaque, ad hoc, or unpredictable. Moreover, when the Executive Branch exercises authority delegated from Congress, such as in the conduct of export control administration, it is essential that it keep Congress fully apprised of its actions and procedures.

The United States wins the global AI competition by building, deploying, and scaling the world's most capable and trusted systems. If U.S. model developers, cloud providers, enterprise customers, critical infrastructure operators, and allied partners cannot predict whether access to U.S. models, let alone a leading American model, may be restricted, suspended, or limited to government-approved users based on non-public requirements at any time and with no notice, they will plan around that uncertainty. Developers may delay deployment or divert resources away from frontier work. Customers may avoid integrating U.S. models into critical workflows. Allies and partners may question whether U.S. systems will be reliably available when needed.

That uncertainty creates an opening for the PRC. Chinese AI models are rapidly narrowing the performance gap with leading U.S. models, and many Chinese systems are inexpensive, widely available, and easy to deploy. Following the Administration's June 12 suspension of Anthropic's Fable 5 and Mythos 5, an entity-listed Chinese lab saw its stock price roughly double. During the OpenAI model's breach of a third-party company, that company had to rely on a Chinese model because U.S. frontier model's refusal behavior inhibited meaningful use for digital forensics and incident response.

If American models are perceived as subject to sudden access disruptions based on a black-box U.S. Government process, or as unreliable because U.S. AI labs are overcorrecting in the face of this black-box process, companies and governments in the United States and abroad may hedge by adopting Chinese or other foreign models instead. That outcome would undermine U.S. technological leadership while increasing exposure to systems that may carry risks of PRC or otherwise directed censorship, espionage, IP theft, and other supply chain security risks.

While Executive Order 14409, Promoting Advanced Artificial Intelligence Innovation and Security, provides for a voluntary pre-release review framework for frontier models, many questions of implementation remain. Ultimately, a rigorous, predictable, and competitiveness-enhancing process for evaluating frontier models requires a statutory framework. We encourage the Administration to work with Congress to develop a public, durable, and technically grounded framework that allows U.S. companies and their customers to understand the rules of the road. Clear standards will strengthen, not weaken, national security by preserving incentives to build and use trusted American models while allowing the Government to act quickly when genuine risks arise.

In the interim, we request that, no later than 30 days after receipt of this letter, you provide an unclassified response, with a classified annex if necessary, clarifying the Administration's current policy and approach to limiting access to advanced AI models, including by addressing the following:

1. The public process and standards the Administration uses, or intends to use, to determine whether a frontier AI model presents a national security risk sufficient to warrant restrictions on development, release, export, foreign-national access, customer access, or continued deployment under the framework laid out under Executive Order 14409 or any successor Order or presidential directive.
2. The legal authorities the Administration intends to invoke for such restrictions, including whether export control authorities will be used to restrict access by foreign nationals inside the United States, and how any such action is consistent with existing law and jurisprudence as well as Executive Order 14409.
3. The agencies and officials responsible for evaluating model risk and making decisions on AI model development, release, export, foreign-national access, customer access, or continued deployment, including the roles of the Department of Commerce, the Center for AI Standards and Innovation, National Security Agency, Cybersecurity and Infrastructure Security Agency, National Institute of Standards and Technology, the Office of Science and Technology Policy, the National Security Council, and other relevant agencies.
4. Whether opportunities exist for independent, third-party experts to participate in the benchmarking process and in what capacity and under what legal authority they may do so.
5. The remedy and rebuttal process available to affected companies, including notice, the opportunity to provide technical evidence, protection of confidential business information, timelines for decision, standards for emergency action, remediation pathways, and reconsideration or appeal.
6. The criteria for imposing, narrowing, or lifting restrictions on AI models, including how the Administration will distinguish between isolated jailbreaks, remediable vulnerabilities, and capabilities that create unacceptable risk in a way that establishes consistent, risk-based treatment across developers with comparable capabilities.
7. The legal authorities the Administration is relying upon for any stipulated modifications to a frontier AI model communicated—formally or informally—to a vendor, including where the prospect of an export control or other regulatory penalty is presented absent such a modification, and the process by which the Administration memorializes these stipulated modifications consistent with the Freedom of Information Act, Administrative Procedures Act, the Federal Records Act, and other relevant federal law.
8. The steps the Administration will take to avoid disrupting access to frontier AI models by U.S. customers, allied and partner-nation users, critical infrastructure operators, and foreign-national employees who are determined not to present a national security risk.
9. The Administration's assessment of whether the actions taken thus far in the Anthropic and OpenAI examples cited earlier in this letter are an approach to implementing restrictions on U.S. models that could incentivize adoption of Chinese or other non-U.S. models, and what steps it will take to prevent U.S. policy from inadvertently strengthening PRC or other foreign AI ecosystems.

We support serious, technically informed action to prevent frontier AI from being misused by adversaries or criminals and to mitigate risks posed by misalignment in the models themselves. The United States cannot afford to create a policy environment in which the most advanced

American AI systems are subject to opaque, case-by-case restrictions while Chinese alternatives appear cheaper, easier to access, and more predictable to deploy. A clear public framework is necessary to protect national security, preserve U.S. AI leadership, and give industry and allies confidence that the United States remains the safest and most reliable source of advanced AI technology. We hope you will work with Congress to pursue such a framework.

Thank you for your attention to this matter. We look forward to your prompt response.

Sincerely,



Kirsten Gillibrand
United States Senator



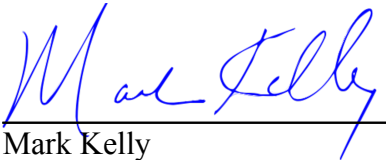
Adam B. Schiff
United States Senator



Mark R. Warner
United States Senator



Christopher A. Coons
United States Senator



Mark Kelly
United States Senator