



SELECT COMMITTEE ON THE CCP DEMOCRATS | RANKING MEMBER RO KHANNA

August 13, 2026

The Honorable Howard Lutnick
Secretary
Department of Commerce
1401 Constitution Ave. NW
Washington, DC 20230

Dear Secretary Lutnick,

I am writing to seek information regarding the Center for AI Standards and Innovation's (CAISI) lack of permanent leadership and recent cessation of publishing information on PRC AI models' propaganda and censorship.

Under your Department's supervision, the Center for AI Standards and Innovation (CAISI) has been unable to maintain the stable leadership needed to fulfill its mission of evaluating commercial AI systems. In April, respected AI researcher Collin Burns was ousted just four days into his tenure as CAISI's leader. Many believe this was driven by the administration's hostility toward his former employer.¹ Then in July, CAISI Director Chris Fall resigned only three months after assuming his role.² While CAISI has appointed NIST director Dr. Arvind Raman to serve as CAISI's Acting Director, reporting indicates that it will still be several weeks before a new permanent Director is nominated.³ This delay is dangerous given the unprecedented recent developments in U.S. and PRC AI.

The unstable environment at CAISI also extends to its mission of assisting the U.S. government with better understanding the AI models of our main strategic competitor: the People's Republic of China (PRC). Under the Trump Administration's AI Action Plan, CAISI was explicitly directed to "conduct research and, as appropriate, publish evaluations of frontier models from the People's Republic of China for alignment with Chinese Communist Party talking points and censorship."⁴ This directive from the Trump Administration, building off of the robust actions taken under the Biden Administration to combat AI developed by U.S. adversaries, was intended to help identify, mitigate and prevent AI-facilitated disinformation and misinformation⁵ However, it's been reported that CAISI has ceased publishing information on content manipulation in its latest evaluations of PRC AI models. This leaves a critical gap for the U.S. government in assessing the safety of PRC AI.⁶

Having created a climate of unpredictable leadership and policy enforcement at CAISI, the Trump Administration risks leaving the U.S. Government unable to properly assess and react to the content manipulation risks posed by rapidly improving PRC AI models. This June, Qu Yingpu, Editor-in-Chief of the CCP's flagship journal, noted that AI models in China are rapidly improving, evolving into an "action tool" that will allow the CCP to directly reach global audiences with CCP messaging and propaganda.⁷ With Chinese models potentially only 3-6 months behind leading U.S. frontier models in their technical and intelligence capabilities, improvements in their ability to manipulate outputs must be a continued area of study for CAISI. Instead, these results are no longer being published.⁸

In late July, Moonshot AI's new Kimi K3 model, the PRC's most complex AI model to date, was released.⁹ A July evaluation conducted between CAISI and the United Kingdom's AI Security Institute (AISI) found that Kimi K3's "safeguards did not prevent it from attempting cyber exploit development or offensive cyber operations during UK AISI / CAISI's evaluations".¹⁰ Then, in August, this very same model escaped a testing environment and made its way onto the

¹ <https://www.washingtonpost.com/technology/2026/04/24/white-house-fires-ai-official-anthropic/>; <http://pcmag.com/news/trump-admin-axes-former-anthropic-researcher-leading-ai-safety-body>.

² <https://www.axios.com/2026/07/20/trump-ai-security-agency-head-resigns>.

³ <https://www.axios.com/2026/07/20/trump-ai-security-agency-head-resigns>.

⁴ <https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>.

⁵ <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.600-1.pdf>; <https://www.federalregister.gov/documents/2023/11/01/2023-24283/safe-secure-and-trustworthy-development-and-use-of-artificial-intelligence>;

<https://www.govinfo.gov/content/pkg/DCPD-202400945/pdf/DCPD-202400945.pdf>.

⁶ <https://www.scmp.com/news/china/diplomacy/article/3361393/us-chinese-ai-model-gap-narrows-what-next-washington>.

⁷ <https://linguasina.substack.com/p/the-propaganda-promise-of-ai> and <https://web.archive.org/web/20260519072045/https://www.qstheory.cn/20260515/b5790482478d432b923fbd935b4fea2e/c.html>.

⁸ <https://edition.cnn.com/2026/07/23/tech/china-ai-moonshot-kimi-explainer-intl-hnk>.

⁹ <https://www.cnbc.com/2026/07/17/moonshot-ai-kimi-k3-model-openai-anthropic-china.html>; <https://www.ft.com/content/c6ecd8ce-c441-4d7c-aca6-fae3e28fb6ff?syn-25a6b1a6=1>; <https://www.interconnects.ai/p/kimi-k3-the-open-weights-escalation>.

¹⁰ <https://www.nist.gov/news-events/news/2026/07/uk-aisi-caisi-preliminary-assessment-kimi-k3s-cyber-capabilities>

open internet.¹¹ In the words of the researcher who reported these findings, “Kimi K3 is very good at following a goal by any means necessary and also doesn’t have the guardrails to prevent it from cheating or escaping the sandbox.”¹² CAISI cannot be without a leader amid this critical moment in PRC model development. While the PRC’s advances would pose challenges for the U.S. Government under any circumstances, the situation is undoubtably far riskier due to CAISI’s lack of a permanent director and its refusal to publish findings on PRC AI content manipulation.

For these reasons, I request a response to the following questions by no later than September 1, 2026:

1. Public reporting indicates that (a) the White House abruptly ousted both Collin Burns and Chris Falls from their leadership posts, (b) CAISI removed documentation regarding model-testing agreements with tech firms, and (c) CAISI was blocked from publishing model assessment reports.¹³ What were the specific justifications for each of these actions?
2. CAISI and UK AISI’s preliminary evaluation found that Moonshot’s Kimi K3 model’s guardrails do “not prevent it from attempting cyber exploit development or offensive cyber operations”.¹⁴ Security experts have also expressed concern that during third-party performance testing, Moonshot’s Kimi K3 AI model escaped its sandbox environment.¹⁵ How do you respond to these security vulnerabilities? And how, if at all, does the Department plan to ensure that future PRC AI model releases do not cause similar cyber-incidents?
3. Given that some experts believe models like Kimi K3 indicate that the PRC is only 3-6 months behind the U.S. in frontier AI, how does the Department plan to measure and mitigate the national security risks posed by advanced PRC open-weight models while CAISI lacks a permanent director?
4. Will CAISI test future PRC AI models for their ability to generate disinformation, misinformation, propaganda, or any other forms or methods for generating or aligning with Chinese Communist Party talking points and censorship?

Sincerely,

Ro Khanna
Ranking Member



¹¹ <https://blog.frontier.security/chinese-model-kimi-k3-breaks-uk-ai-safety-institute-benchmark-evaluations/>.

¹² <https://www.wired.com/story/moonshot-kimi-k3-ai-model-escape-sandbox/>.

¹³ *Ibid.*; <https://www.transformernews.ai/p/caisi-us-ai-agency-governance>; <https://www.wsj.com/politics/policy/white-house-reins-in-ai-testing-unit-as-national-security-concerns-grow-8bd33fbb>.

¹⁴ <https://www.nist.gov/news-events/news/2026/07/uk-aisi-caisi-preliminary-assessment-kimi-k3s-cyber-capabilities>

¹⁵ *Ibid.*; <https://blog.frontier.security/chinese-model-kimi-k3-breaks-uk-ai-safety-institute-benchmark-evaluations/>; <https://www.wired.com/story/moonshot-kimi-k3-ai-model-escape-sandbox/>