



SELECT COMMITTEE ON THE CCP

DEMOCRATS | RANKING MEMBER RO KHANNA

September 16, 2026

Mr. Yang Zhilin
Chief Executive Officer
Beijing Moonshot Technology Co., Ltd. d/b/a Moonshot AI
13th Floor, Building 1, JD Technology Building
No. 76 Zhichun Road, Haidian District
Beijing 100191
People's Republic of China

Dear Mr. Yang:

As Ranking Member of the House Select Committee on the Chinese Communist Party, I am writing to urge China's leading artificial intelligence (AI) companies to enter into a binding, enforceable, and verifiable agreement alongside America's leading AI labs to pace the frontier. The CEOs of Anthropic, OpenAI, and xAI do not agree on much, but earlier this month, they all came together and agreed that we should slow the rate of AI advancement until society can develop the necessary guardrails. It is time for Moonshot, Alibaba, and DeepSeek to do the same.

Some have cited the rapid advancement of AI models produced by your company and other leading labs in China as a reason why pacing the frontier would be a mistake. They argue that any strategic slowdown in AI development necessarily means ceding American AI leadership to Beijing. That perspective is distinctly misguided, in part because it presumes global pacing is not a possibility. Although our differences on many issues are vast, the United States and the People's Republic of China (PRC) share a mutual interest in preventing the loss of control or misalignment of advanced AI systems.

To that end, I have called on the governments of the United States, our democratic allies, and the PRC to urgently engage in negotiations to establish such baseline global standards. As Mr. Amodei of Anthropic recently wrote, we should "worry that in 6–12 months a swarm [of agents] could be capable of taking over the entire internet with a persistent botnet."¹ The fallout would be historic, and it would devastate the American and Chinese economies alike. Meanwhile, Mr. Altman of OpenAI has said "we could lose control of the future to AI."² Needless to say, both the U.S. and the PRC would suffer in a world where human interests diverge from the goals of misaligned superintelligence. Similarly, neither country would benefit from a biological or chemical attack facilitated by AI. These and other common interests should catalyze our leaders to negotiate an enforceable U.S.-PRC AI Treaty. These discussions should commence at the meeting between President Trump and General Secretary Xi next week. Unless the world adopts verifiable frontier pacing standards, including explicit limits on unchecked recursive self-improvement, I fear that this powerful technology will not always remain in service of humanity.

But if our political leaders do not pursue or fail to achieve such an agreement, it falls on leading AI companies themselves – wherever they may be domiciled – to ensure their decisions do not jeopardize the collective fate of humanity. The founders of Anthropic, OpenAI, and xAI have taken the first step

¹ <https://darioamodei.com/post/we-must-pace-the-frontier>

² <https://x.com/sama/status/2099352016988614852>

with their call to pace the frontier. Now it is your turn to act and echo their pledge. If all of the world's leading labs agreed to pace the frontier, the economic incentive to prioritize profit over safety would be rebalanced. It would also give our respective governments the time they need to negotiate and try to reach an agreement that binds the entire world to common sense safeguards.

There is no sugarcoating the fact that our countries have different values. Your government commits grave human rights violations against its people, and PRC AI models have played a significant role in facilitating those abuses. There is also no sugarcoating the fact that PRC AI models are used to the detriment of American and allied interests, including by helping to modernize the People's Liberation Army and supporting PRC intelligence agencies. In fact, Chinese law can compel such cooperation. However, in contrast, it is in the interest of both American and Chinese AI labs – as well as both the American and Chinese governments – that your companies refrain from practices that could result in humans losing control of AI or the development of rogue misaligned agent swarms. If both our countries' leading labs simultaneously agree on a more responsible path forward, it is more likely that both sides will actually follow through. We do not need to agree on much to agree that we should narrowly work together on this.

At a congressional hearing I am convening next Wednesday, I will urge President Trump and General Secretary Xi to pursue AI diplomacy. I will also demand that your company and the other leading Chinese AI labs adopt frontier pacing alongside your American peers. Accordingly, by no later than September 22, 2026, I request your written responses to the following questions:

1. For the benefit of all of humanity, will you agree to 'pace the frontier' and slow down your rate of AI advancement in order to ensure societal safeguards can catch up to technology?
2. If the leading U.S. frontier labs offer to enter into a binding, enforceable, and verifiable agreement with you to 'pace the frontier', including by adopting common sense safeguards like global limits on unchecked recursive self-improvement, would you execute such an agreement?
3. Pursuant to any such agreement, would you agree to neutral third party auditing and all-access inspections to ensure verifiable compliance with the frontier pacing agreement?
4. As they continue to advance, the risks posed by your AI models are not just a domestic PRC concern, but a concern for Americans and all of humanity. If you will not or would not agree to any of the commitments in the preceding questions, please describe what, if any, safeguards your company has to prevent the loss of control or misalignment of your AI models and future AI models.

Sincerely,



Ro Khanna
Ranking Member
House Select Committee on the CCP