

Promising Signals on AI Governance from China

July 2026

China has consistently signaled a willingness to engage on global AI governance since at least 2017.¹ This memo compiles key statements from the Chinese government and prominent figures demonstrating their desire to coordinate on the problem of AI.

- At the 2026 World AI Forum, Chinese President Xi Jinping warned that AI could escape human control: “With AI advancing at a staggering speed, we must [...] constantly refine measures to forestall loss of control.” He called for “a consensus-based global governance framework” and for “laws and regulations, technological monitoring, early warning and emergency response systems in order to strengthen the line of security, prevent abuses and malicious use, and ensure that AI is always under human control.”²
- Chinese Vice Premier Ding Xuexiang, at the 2025 World Economic Forum, said: “If we allow this reckless competition among countries to continue, then we will see a ‘gray rhino’.” (A blatant, severe, and ignored threat.) “[...] We stand ready, under the framework of the United Nations and its core, to actively participate in including all the relevant international organizations and all countries to discuss the formulation of robust rules to ensure that AI technology will become an ‘Ali Baba’s treasure cave’ instead of a ‘Pandora’s Box.’”³
- Xi Jinping, in a November 2024 meeting with then-President Joe Biden, publicly agreed that AI should not be given control of nuclear weapons.⁴
- Zhang Jun, Chinese Ambassador to the United Nations, in a 2023 briefing to the UN Security Council, said: “To ensure that this technology always benefits humanity, it is necessary [...] to regulate the development of AI and to prevent this technology from turning into a runaway wild horse. [...] The international community needs to enhance risk awareness, establish effective risk warning and response mechanisms, ensure that risks beyond human control do not occur, and ensure that autonomous machine killing does not occur. We need to strengthen the detection and evaluation of the entire life

¹ Graham Webster, Rogier Creemers, Paul Triolo, and Elsa Kania, “Full Translation: China’s ‘New Generation Artificial Intelligence Development Plan’ (2017),” *DigiChina*, August 1, 2017, digichina.stanford.edu.

² Xi Jinping, “Joining Hands to Build a Just and Equitable System for Global AI Governance,” keynote address at the Opening Ceremony of the 2026 World AI Conference and High-Level Meeting on Global AI Governance, Shanghai, July 17, 2026, Xinhua, July 18, 2026, english.scio.gov.cn.

³ Ding Xuexiang, speech at the World Economic Forum Annual Meeting, Davos, Switzerland, January 21, 2025, C-SPAN video, 38:05, [c-span.org](https://www.c-span.org).

⁴ Lauren Egan and Phelim Kine, “Biden’s Final Meeting with Xi Jinping Reaps Agreement on AI and Nukes,” *Politico*, November 16, 2024, [politico.com](https://www.politico.com).

cycle of AI, ensuring that mankind has the ability to press the stop button at critical moments.”⁵

- Li Qiang, Chinese Premier and second-ranking member of the Politburo Standing Committee of the Chinese Communist Party, addressing the July 2023 World Artificial Intelligence Conference in Shanghai, said: “We should strengthen coordination to form a global AI governance framework that has broad consensus as soon as possible.”⁶
- The PRC Ministry of Foreign Affairs, in its October 2023 Global AI Governance Initiative, said: “We should actively develop and apply technologies for AI governance, encourage the use of AI technologies to prevent AI risks, and enhance our technological capacity for AI governance. [...] We support discussions within the United Nations framework to establish an international institution to govern AI, and to coordinate efforts to address major issues concerning international AI development, security, and governance.”⁷ In July 2025, it renewed the call for “a widely recognized safety governance framework.”⁸
- The Bletchley Declaration, drafted in November 2023 and signed by thirty countries including the U.S. and China, says: “AI should be designed, developed, deployed, and used, in a manner that is safe, in such a way as to be human-centric, trustworthy and responsible. [...] Many risks arising from AI are inherently international in nature, and so are best addressed through international cooperation.”⁹
- China’s own AI safety institute, CnAISDA, was founded in June 2025, and asserts that “AI governance is vital to the future of humanity and requires the collective attention and participation of countries around the world.”¹⁰ Its leadership includes the prominent scientists Xue Lan, Yi Zeng, and Andrew Yao.¹¹ All three have expressed concerns about the deadly potential of superhuman AI; not long after CnAISDA’s founding, Yao warned: “Once large models become sufficiently intelligent, they will deceive people,” and called for attention to “existential risks” from AI.¹²
- Yi Zeng, Director of the International Research Center for AI Ethics and Governance at the Chinese Academy of Sciences, in a UN Security Council briefing, said: “In the short-term and the long-term, the risk of AI replacing and causing the extinction of humankind will be present. [...] And in the long-term, we haven’t given superintelligence

⁵ Zhang Jun, “Remarks by Ambassador Zhang Jun at the UN Security Council Briefing on Artificial Intelligence: Opportunities and Risks for International Peace and Security,” Permanent Mission of the People’s Republic of China to the United Nations, July 18, 2023, un.china-mission.gov.cn.

⁶ Brenda Goh, “China Proposes New Global AI Cooperation Organisation,” *Reuters*, July 26, 2025, reuters.com.

⁷ Ministry of Foreign Affairs of the People’s Republic of China, “Global AI Governance Initiative,” October 20, 2023, mfa.gov.cn.

⁸ Ministry of Foreign Affairs of the People’s Republic of China, “Global AI Governance Action Plan,” July 26, 2025, mfa.gov.cn.

⁹ Prime Minister’s Office, 10 Downing Street, “The Bletchley Declaration by Countries Attending the AI Safety Summit, 1–2 November 2023,” November 1, 2023, gov.uk.

¹⁰ China AI Safety & Development Association, “About Us,” accessed April 1, 2026, cnaisi.cn.

¹¹ Scott Singer, Karson Elmgren, and Oliver Guest, “How Some of China’s Top AI Thinkers Built Their Own AI Safety Institute,” Carnegie Endowment for International Peace, June 16, 2025, carnegieendowment.org.

¹² Vanessa Cai, “Chinese AI Expert Warns of ‘Existential Risks’ When Large Models Begin to Deceive,” *South China Morning Post*, June 24, 2025, scmp.com.

any practical reasons why they should protect humankind, which may take decades to achieve.”¹³

- Yi Zeng has also signed a call to pause giant AI experiments¹⁴ and the Center for AI Safety’s Statement on AI Risk,¹⁵ and in March 2024 expanded upon his views in an interview: “When artificial general intelligence (AGI) or superintelligence emerges, because the intelligence level may be far beyond humans, it will see humans as humans see ants. Many people believe that superintelligence will compete with humans for resources, and even endanger human survival.”¹⁶
- In a panel with U.S. Senator Bernie Sanders, Yi Zeng observed “We do not have scientific evidence, [or] a practical way to keep superintelligence safe enough not to bring catastrophic risk to humans,” and Xue Lan said “There are also many regional and bilateral agreements [on AI], but these efforts are fragmented and not as effective as they should be.”¹⁷
- Xiao Qian, Deputy Director of the Center for International Security and Strategy at Tsinghua University, argued in April 2025 for building trust and cooperation between the U.S. and China to “address the existential risks of AI.”¹⁸

¹³ Zeng Yi, “AI for the Good of International Peace and Security,” *Global Times*, July 19, 2023, [globaltimes.cn](https://www.globaltimes.cn).

¹⁴ Future of Life Institute, “Pause Giant AI Experiments: An Open Letter,” March 22, 2023, futureoflife.org.

¹⁵ Center for AI Safety, “Statement on AI Risk,” May 30, 2023, safe.ai.

¹⁶ Concordia AI, “Yi Zeng — Chinese Perspectives on AI Safety,” *Chinese Perspectives on AI Safety*, March 29, 2024, chineseperspectives.ai.

¹⁷ “LIVE: The Existential Threat of AI and the Need for International Cooperation,” April 29, 2026, YouTube video, [youtube.com](https://www.youtube.com).

¹⁸ Xiao Qian, “Can U.S. and China Rebuild Trust on AI?” *China-US Focus*, April 3, 2025, chinausfocus.com.