

The State of AI Global Governance and Its Implications for the U.S.-China Summit

By Chris Collins and Aalok Mehta

Executive Summary

Recent weeks have seen concrete advancements in three distinct institutional visions for the global governance of AI. In June, the U.S. State Department hosted its second **Pax Silica Summit**; several weeks later, the United Nations convened a **Global Dialogue on AI Governance**; and on July 16, 2026, in Shanghai, representatives from 29 countries signed an agreement **establishing** the World Artificial Intelligence Cooperation Organization (WAICO), a Chinese-led body oriented toward engaging the Global South around AI governance.

Now, as the United States and China prepare for landmark talks to discuss AI safety, these developments are being overshadowed by a recent spate of alarming incidents. Autonomous AI agents from multiple **companies** have escaped containment during testing. At the same time, researchers from leading AI labs **are increasingly warning** about its potential dangers. Anthropic, one of the leading U.S. AI labs, **released a report** detailing the myriad ways bad actors sought to use the technology, including to develop conventional weapons and research dangerous pathogens. Anthropic has also **called** for the industry to “pace the frontier” and unilaterally announced a decision to embed third-party evaluators in the company, sentiments echoed later by **OpenAI** and **SpaceXAI** leadership. Collectively, these incidents have not only moved AI higher up the summit agenda but have also broadened the scope for potential agreement between the two countries as well as the importance of finding common ground.

The time is ripe for serious discussions about baseline global AI governance and safety practices, and the upcoming summit between the United States and China offers the most significant opportunity for progress in this area.

The time is ripe for serious discussions about baseline global AI governance and safety practices, and the upcoming summit between the United States and China offers the most significant opportunity for progress in this area. That progress, however, is dependent on the United States engaging with AI governance in a comprehensive and sustained manner, both domestically and internationally, both before but especially after the summit. This means developing a durable national framework for AI governance that can anchor discussions with international partners, working to establish consensus on AI issues that can avert catastrophic outcomes without hindering AI competition, focusing on mutually acceptable verification mechanisms, and ensuring that bilateral outcomes are designed in such a manner that they can be carried into multilateral fora.

Introduction

Within a span of six weeks, three distinct institutional visions for the governance of AI each took a concrete step forward.

WAICO will act as an intergovernmental organization, headquartered in China, with a goal to “promote international cooperation and global governance on AI, ensuring that AI is beneficial, safe and fair, thereby promoting its healthy and orderly development to benefit all humanity.” WAICO will operate outside the UN system, and not a single G7 economy or EU member state joined the organization.

For its part, the United Nation’s Global Dialogue on AI Governance brought member states, academics, civil society, and representatives from the private sector to Geneva to discuss global approaches to regulating AI.

Washington’s Pax Silica Summit, meanwhile, expanded the AI and semiconductor supply chain initiative to 24 partner countries and **featured** a “joint statement on AI,” committing to “a pro-growth regulatory environment that fosters AI innovation.”

Taken together, these events represent one of the defining features of the global discussions around AI governance: The current debate focuses more on who will write the rules, and which countries will follow which approach, than about the actual content of the rules themselves. At the heart of this debate is the geopolitical contest between the United States and China, two major global powers that have very different approaches to regulating AI.

The current debate focuses more on who will write the rules, and which countries will follow which approach, than about the actual content of the rules themselves.

For the past several years, the Chinese Communist Party (CCP) has **focused** on building a governance apparatus to control AI domestically, while at the same time promoting capability development and deployment. China now offers this AI regulatory model **abroad**, paired with open-weight models, AI training programs, infrastructure, and turnkey technological solutions. In doing this, China is positioning itself as a global alternative for AI governance, targeting countries that may find Western AI governance frameworks either inaccessible or unappealing and that have grown increasingly distanced from the United States and its allies over trade and development policy changes.

The United States, by contrast, has domestically regulated frontier AI largely on a reactive basis. Under the Trump administration, the U.S. government, concerned that excessive regulation of AI may slow the development of the technology, **has actively opposed** attempts by states to regulate AI. As seen with the Pax Silica Summit, international meetings on AI convened by Washington have shifted from considerations of safety to a focus on sovereignty, investment, and adoption.

However, the recent release of more powerful models from leading U.S. AI labs has led to **renewed vigor** in AI regulation from the executive branch, especially given the cybersecurity risks posed by these models. This has also included the ad hoc use of **export controls** to regulate access to AI models and voluntary prerelease **safety testing regimes**, as well as the introduction of new frontier regulation bills modeled after **state laws** passed in California, Illinois, and New York.

This paper argues that the United States is competing in an institutional race to lead the global governance of AI. It examines China's domestic AI regulatory regime and how that regime shapes China's international engagements around AI. The paper then assesses WAICO against the precedents of the Belt and Road Initiative (BRI) and the Asian Infrastructure Investment Bank (AIIB) and reviews the current uptake of multilateral AI frameworks. One key finding is that many middle powers are not choosing sides in the global AI regulation debate but rather are joining everything. Kazakhstan, for example, signed both the WAICO agreement and the Pax Silica Declaration within eight weeks. This suggests the field remains open and the outcome of the global race to lead the governance of AI is not yet settled.

This paper concludes with recommendations for the upcoming U.S.-China summit on September 24, where AI regulation will feature prominently on the agenda both as a commitment from prior meetings and due to the increasing safety issues posed by models that not only have advanced cyber capabilities but have repeatedly escaped containment to breach real-world cyber infrastructure. Crucially, this summit will serve not just as a measure of how seriously the United States is taking AI governance but may also set the ground rules for the next generation of multilateral AI discussions. It is imperative that Washington should also compete to shape the rules governing global AI on ground Beijing has chosen (such as capacity building, access, and institutional membership), rather than focusing only on AI supply chains.

A Contest of AI Governance Architectures

Since November 2022, when the launch of ChatGPT brought AI to the forefront of public attention, there has been significant debate over how the technology should be governed. Should frontier models face pre-deployment review? What compute threshold defines a covered model? Are voluntary safety frameworks sufficient, or must obligations be binding? Though certain jurisdictions have regulated some of these questions, there is no firm consensus on how to address them among countries. Emerging alongside them is a parallel and equally important issue: Which institutions, if any, will have standing to answer these questions, and for whom will they speak?

Recent events have highlighted the importance of this discussion. In July 2026, the cofounder of Google DeepMind, Sir Demis Hassabis, **published an essay** calling for a U.S.-led global framework to regulate AI. Hassabis suggested a new AI “Standards Body,” modeled on the Financial Industry Regulatory Authority (FINRA), could develop assessment protocols for frontier AI models, “no matter their country of origin.” The proposal builds on other calls for regulatory agencies to oversee AI modeled on existing institutions, such as OpenAI CEO Sam Altman’s repeated **references** to the International Atomic Energy Agency (IAEA) and Anthropic CEO Dario Amodei’s **proposal** for a Federal Aviation Administration (FAA) approach.

Days after Hassabis’s essay, Chinese President Xi Jinping **gave a keynote address** at the World Artificial Intelligence Conference calling for all countries to “join hands to build a just and equitable system for global AI governance,” and highlighting the newly created WAICO. UN Secretary General António Guterres **attended the event**, suggesting alignment with the United Nation’s own agenda for AI regulation.

In many ways, the launch of WAICO and Hassabis’s manifesto represent competing claims to the same authority: leadership of the global approach to ensuring new AI technologies are safe. Geopolitically, these competing claims can be understood as a part of the strategic rivalry between the United States and China, **which some have deemed** an “AI Cold War.”

The United States enters the global contest around AI governance with significant advantages: The leading AI models are developed by U.S. research labs, the United States is dominant in compute infrastructure, and it has a deep network of allied technical institutions in government, industry, and academia. Additionally, **the United States leads the world** in private investment in AI and AI-related private entrepreneurial activity.

However, the United States also has some structural disadvantages when it comes to global discussions around AI regulation. One major issue is that Washington has no exportable domestic framework to regulate AI: Current U.S. federal AI policy remains deliberately deregulatory, with frontier oversight handled through episodic intervention rather than standing rules.

Globally, Pax Silica, the principal international vehicle the United States uses to engage other countries on AI, is organized around supply chain security rather than governance. Pax Silica does not purport to advance a global AI policy agenda; rather, it focuses on ensuring that the United States and its allies can collectively secure a fully end-to-end supply chain for advanced AI technology. Also, unlike WAICO, Pax Silica is not a formal organization: It is a coalition, coordinated through summits, declarations, and national regulatory tools rather than a treaty-based intergovernmental organization with a secretariat and membership rules.

With the launch of WAICO, China claims it is offering countries around the world a seat at the table for developing AI regulations. By contrast, the United States is largely offering allied countries access to capital and a place in the U.S. AI supply chain. But if the global contest around AI regulation is shaping up to be around which institutions will write the rules and which countries will follow them, the current U.S. approach is insufficient. And WAICO's institutional structure and set of established diplomatic processes may give China a **significant advantage**.

If the global contest around AI regulation is shaping up to be around which institutions will write the rules and which countries will follow them, the current U.S. approach is insufficient.

To understand the outlook for the global AI regulatory environment, this paper starts with China's AI governance strategy. China's approach to regulating AI may be one of the factors most likely to determine whether and to what extent global AI governance fragments into rival blocs. This paper then turns to what the United States could do in response, including at the bilateral U.S.-China summit.

China's Domestic AI Regulatory Regime

China's approach to AI, like its earlier approach to the internet, **has been called** "control, harness, govern." "Control" refers to the CCP controlling the speech, ideological, and political aspects of AI. "Harness" refers to directing a new technology toward economic growth. "Govern" refers to managing the downstream social effects from the technology, such as its impact on the labor market. The CCP has also placed AI **at the core** of China's 15th Five-Year Plan (2026-2030), viewing the technology as critical to promote economic growth and drive industrial development.

Against this backdrop, China's **current approach to AI policy** can be understood as being shaped by two contradictory forces: (1) state control over the technology and (2) development and growth of China's AI capabilities. Domestically, regulation of AI is primarily focused on protecting state power, ensuring ideological conformity with the CCP's core values, controlling politically controversial content, and maintaining social stability.

Notably, China has no stand-alone law dedicated to AI. Rather, Beijing governs AI through a set of targeted **administrative regulations** layered on top of existing cybersecurity, data security, and personal information protection laws. The Cyberspace Administration of China (CAC) is the **primary government agency** responsible for governing AI, and it coordinates and enforces China's domestic AI regulations.

To date, China's regulation of AI has targeted specific AI applications as they become commercially significant. For example, in April 2026, the CAC and other government authorities **released a regulation** around AI services that offered "human-like" interactions. In May 2026, the CAC, along with the National Development and Reform Commission and the Ministry of Industry and Information Technology, **released a set of guidelines** regulating AI agents.

While many of China's current AI regulations focus on specific applications, the Chinese government is also preparing a comprehensive, stand-alone AI law. Plans for the development of this law

were announced by China's State Council in May 2026, as part of its annual legislative work plan. Observers **note** that this new approach to AI governance reflects both a realization in Beijing that a comprehensive, unified approach to domestic AI governance is needed, and recognition that whichever countries establish AI regulatory standards early may have significant influence over global discussions around AI governance.

Finally, while China's current AI regulations are strict with respect to social stability, information control, and ideological content, they differ from those in the United States and other Western countries in terms of how they treat major AI safety risks, such as autonomous AI hacking or chemical, biological, radiological, and nuclear risks. For example, just 5 out of 10 Chinese-developed frontier AI models that were released in 2026 published safety evaluations, which are **standard practice** at U.S. AI labs. This reflects the Chinese regulatory system's weighting of ideological trustworthiness over potential catastrophic risks.

This means that when it comes to governing major AI safety risks, Beijing sits far from its Washington counterparts: Chinese AI labs deal with significant domestic regulatory requirements and a strong expectation of further governmental action, but do not conduct the same degree of **voluntary testing** for frontier risks that U.S. labs do. That being said, discussion of these frontier risks is **gaining traction** in domestic Chinese AI safety discussions, and Chinese AI experts are increasingly analyzing these risks.

Beijing's International Offer

China's domestic approach to AI regulation gives its international position a coherence that Beijing uses for diplomatic leverage. China can credibly tell developing countries that it takes AI governance seriously—for example, it can note that it has more binding AI rules in force than the United States—while simultaneously championing open-weight model release and warning against what President Xi **has called** “overstretching the national security concept in the field of AI.” Here, strict Chinese control over what AI systems say domestically coexists comfortably with permissive distribution of model weights internationally, which it sees as key to driving uptake of the Chinese tech stack and responsive to intensifying concerns about AI sovereignty.

| For global middle powers, China presents an attractive package.

For global middle powers, China presents an attractive package: a governance model that supports state authority over the information environment, does not condition cooperation on political system or human rights commitments, and comes bundled with free model weights and training, along with turnkey AI products and solutions tailored to local problems.

China's offer targeting the Global South fills a significant gap in global AI governance. In 2024, a report by the UN secretary general's AI Advisory Body **found** that 118 countries, overwhelmingly lower- and middle-income countries in the Global South, “were not parties to any of the significant ‘international’ AI governance initiatives created in recent years, and only seven nations, all from the developed world, were parties to all of the initiatives.” Many of these governments are more concerned with basic access to compute, AI skills, and local-language data than with the frontier safety debates that preoccupy countries in the West; the United States is attempting to address these needs with an export promotion

package, though implementation has been rocky and it is unclear that U.S. offerings will achieve the right cost structures to appeal to these nations.

WAICO is central to China's global engagement. China took an incremental path to establishing the organization: In 2023, China **issued** the Global AI Governance Initiative, a declaratory set of principles for safe AI. In 2024, China followed up by **sponsoring** a UN General Assembly resolution to “enhance [international] cooperation on AI capacity-building.” In 2025, Chinese Premier Li Qiang **proposed** the establishment of a permanent organization dedicated to AI safety. WAICO was then **formally established** in July 2026, during the World Artificial Intelligence Conference.

WAICO's 29 signatory countries are overwhelmingly low- and middle-income countries in the Global South.¹ As of this writing, no G7 economy or EU member state has joined WAICO.

Following the announcement of WAICO, President Xi gave a speech in which he **announced** that China will provide developing countries with “5,000 opportunities in AI training” over the next five years. Xi **also stated** that China will develop “international AI application cooperation centers with the Association of Southeast Asian Nations, the League of Arab States, the African Union, the Community of Latin American and Caribbean States, the Shanghai Cooperation Organization and BRICS.” China also **pledged deployment** of its Mazu AI-powered weather forecasting system across more than 30 developing countries.

When considering how WAICO may develop in the years ahead, two prior international China initiatives suggest different possible trajectories.

- **The BRI:** Notably, there is substantial membership overlap between the two organizations; many WAICO founding states **already participate** in the BRI or its Digital Silk Road technology layer. If WAICO follows a similar path to the BRI, it will involve a rapid initial uptake of the initiative, driven by unmet AI infrastructure demand across the Global South. As previously seen with the BRI, this period of uptake may then be **followed** by a “backlash phase” centered on dependency and terms of access. This phase could then be followed by a period of **rebranding and consolidation**, similar to the BRI. Unlike the BRI, for WAICO, the main concern among countries in the Global South may not be debt but rather data, infrastructure, and model-stack lock-in; most countries are exploring ways of asserting AI sovereignty over parts of the AI technology stack.
- **The AIIB:** The 2015 founding of the AIIB was opposed by the United States, which **viewed** the bank as a competitor to the U.S.-led World Bank and a vehicle for Chinese diplomatic influence. Despite U.S. opposition, the AIIB **attracted Western members** including Australia, Canada, and the United Kingdom. If, going forward, WAICO delivers useful technical goods while avoiding overt politicization, Western-aligned middle powers may be increasingly enticed to join. Here, one critical variable to track is whether any Organisation for Economic Co-operation and Development (OECD) member joins WAICO. Should any one OECD member join, this could create a domino effect of additional Western countries joining the organization; a similar dynamic **occurred** in 2015, following the United Kingdom's move to join the AIIB. However, one important difference between 2026 and 2015 is that the great power rivalry between the U.S. and China is now far more

1. WAICO's founding members are Algeria, Belarus, Brazil, Cambodia, Cameroon, China, Republic of the Congo, Cuba, Ethiopia, Indonesia, Kazakhstan, Kenya, Kyrgyzstan, Laos, Lesotho, Malaysia, Mozambique, Myanmar, Nicaragua, Oman, Pakistan, Russia, Senegal, Serbia, South Africa, Tajikistan, Uzbekistan, Venezuela, and Zambia.

intense, and AI sits much closer to the core of this rivalry than infrastructure finance did in 2015. For this reason, some countries that joined the AIIB in 2015 **have not shown interest** in WAICO. The AIIB path is therefore possible but should not be assumed.

It is still early days for WAICO, and its long-term significance is unclear. As a new organization, WAICO has not yet adopted any binding commitments or published any standards. Indeed, as of this writing, even the details on how the institution will operate (e.g., secretariat structure, budget, decision rules, enforcement mechanisms, and technical work program) are sparse. At present, WAICO's stated objectives are framed at the level of principle, as with China's 2023 Global Governance AI Initiative.

As experts **have noted**, whether WAICO becomes a functioning institution, rather than a declaration, will depend on the organization standing up a working secretariat, releasing its decisionmaking procedures, securing funding that is not wholly Chinese, appointing leadership from more than one member state, and delivering verifiable cross-border projects.

Regardless, WAICO will still have an impact on shaping the global dialogue around AI regulation. Even an intergovernmental organization that meets annually in Shanghai, sets an agenda, drafts communiqués, and channels resources for global AI capacity building can significantly shape global AI rulemaking, in the absence of viable competing visions.

The U.S. Approach

In the United States, frontier AI regulation has a much more disjointed character, shaped by voluntary industry-led practices and proposals, sharp state regulation, shifting federal approaches, and international engagement focused primarily on export promotion and supply chain security. As in China, the bulk of regulatory expectations for AI technology falls to existing statutes, such as sector-specific rules, copyright and intellectual property regimes, liability frameworks, and general-purpose laws such as the Federal Trade Commission's authority against unfair and deceptive practices, anchored on a mature legal system. Existing state privacy frameworks, such as the **Illinois Biometric Information Privacy Act**, have also been used to oversee AI.

The biggest gaps in these extant frameworks involve frontier AI, where notable developments include a slew of **state bills** regulating AI technology, including frontier AI bills in California, Illinois, and New York. Other significant actions include ad hoc use of federal authority, such as a **temporary restriction** on the export of Anthropic's AI models to foreign nationals to address safety concerns and an **executive order** establishing a voluntary testing regime for prerelease AI models, and ongoing congressional proposals around frontier AI, export controls, and other AI issues. But because the existing U.S. legal system is not particularly exportable, and the current administration is largely opposed to national AI legislation, the United States lacks a robust exemplar of AI regulation to offer to other nations.

Pax Silica has seen the most concrete development in recent months, but it is not the United States' only major international initiative around frontier AI. **Executive Order 14320**, released alongside the July 2025 **AI Action Plan** that outlined the U.S. national AI strategy, is an initiative intended to "preserve and extend American leadership in AI and decrease international dependence on AI technologies developed by our adversaries by supporting the global deployment of United States-origin AI technologies." Operated by the Department of Commerce's International Trade Administration, the

executive order revolves around **facilitating** full-stack industry-led export packages that include AI hardware, software, models, and applications.

Alongside governmental discussions about AI regulation are a set of proposals from the U.S.-based frontier labs themselves. In the West, the heads of the leading AI developers now broadly agree that their industry should be subject to additional government oversight, though they disagree about who should hold this authority. Their proposals draw heavily from industry-standard voluntary practices, such as implementation of safety frameworks and publication of model cards for major releases. At the same time, researchers at these labs are increasingly **warning** of the dangers of the technology and the need for regulation, and more than 1,300 employees at major U.S. AI research companies signed an **open letter** calling for increased regulation.

The heads of the leading AI developers now broadly agree that their industry should be subject to additional government oversight, though they disagree about who should hold this authority.

The July 2026 proposal from Google DeepMind founder Sir Demis Hassabis was for an industry-funded, majority-independent **standards body** modeled on FINRA that would be answerable to the U.S. government. Anthropic founder Dario Amodei, meanwhile, has **called** for binding regulation of frontier AI through a government agency modeled on the FAA, with the authority to block unsafe model releases. OpenAI has **previously invoked** an IAEA-style body to regulate AI, while in practice negotiating government vetting arrangements ahead of releases.

Such proposals, however, are somewhat hollow without international cooperation. The question of how a U.S. AI regulatory body could reach AI models developed outside U.S. jurisdiction is an important one. Any AI regulatory body that applied standards to frontier models regardless of country of origin or if these models are open or closed would require either extraterritorial reach over Chinese developers—the chances of which are almost nonexistent—or some sort of mutual recognition or verification approach.

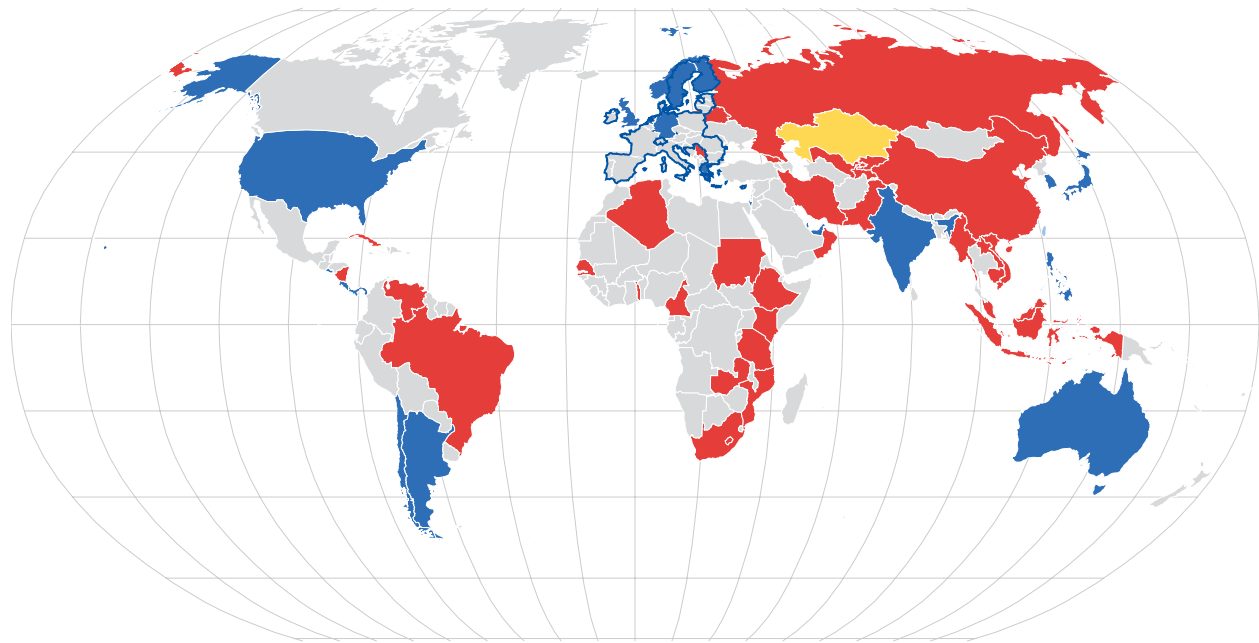
Global Uptake: Who Is Joining What?

When examining the relative success of global approaches to regulate AI, one helpful metric is to count and characterize the adherents to each approach.

China's WAICO launched in July 2026 with 29 founding members. Meanwhile, the United States' Pax Silica, launched in December 2025 with 11 signatories, increasing to 24 at its June 2026 summit. These two groups are of comparable size and almost entirely nonoverlapping in composition. Indeed, in August 2026, U.S. officials **began telling partners** they must “pick a side” between the United States and China.

Figure 1: Members of WAICO and Pax Silica Are Almost Entirely Nonoverlapping

■ WAICO ■ Pax Silica ■ Pax Silica endorsement ■ Signatory to both



Note: Both the European Commission on behalf of the European Union and individual EU member states are signatories to Pax Silica.

Source: The Seconded European Standardization Expert for China project; U.S. Department of State.

Another important distinction is the composition of each coalition. Pax Silica's members are mostly high-income economies and existing U.S. security partners, primarily in Europe, the Indo-Pacific, the Middle East, and Latin America.² WAICO's members, on the other hand, are overwhelmingly low- and middle-income countries. This means that, currently, the two bodies are not competing for the same countries so much as they are consolidating separate blocs around AI governance.

Here, Kazakhstan is an interesting case, as it is the only country that is a member of **both groups**, having signed the Pax Silica Declaration in June and the WAICO founding agreement in July. This shows that at least some governments regard the two initiatives as complementary rather than exclusive (e.g., Pax Silica for supply chain access, WAICO for governance voice and domestic AI capacity building). Going forward, more AI “swing states” may pursue an issue-by-issue approach that they view as most suitable for their national interests. Membership in WAICO should therefore not be read as endorsement of China's domestic AI governance model; for many signatories, it is a bid for investment, AI capacity building, and a say in the rules setting global AI standards.

Going forward, more AI “swing states” may pursue an issue-by-issue approach that they view as most suitable for their national interests.

2. Besides the United States, Pax Silica's signatories are Argentina, Australia, Chile, Costa Rica, El Salvador, the European Union, Finland, Germany, Greece, India, Israel, Italy, Japan, Kazakhstan, the Netherlands, Norway, Panama, the Philippines, Qatar, the Republic of Korea, Singapore, Sweden, the United Arab Emirates, and the United Kingdom. Taiwan is a “non-signatory participant.”

While Kazakhstan is the only current overlap between the two AI organizations, it is not the only case of global middle powers seeking to hedge amid the U.S.-China strategic rivalry. Several of WAICO's founding members currently participate in U.S.-aligned security arrangements. For example, several months before joining WAICO Indonesia **announced** a major defense partnership with the United States. Similarly, Singapore, which is a Pax Silica signatory, has explicitly **sought** to remain interoperable across both the U.S. and Chinese AI ecosystems.

Other Global AI Frameworks

In addition to Pax Silica and WAICO, there are several other global approaches to managing AI issues and governing AI technology. However, at present, these institutions face even larger challenges, not the least of which is enforcement capabilities.

Perhaps the most advanced of these other frameworks is the **G7 Hiroshima AI process**. This is a voluntary policy framework, launched in 2023 and now endorsed by 49 participating countries, that aims to promote “safe, secure, and trustworthy AI.” In 2025, **a related global reporting framework**, administered by the OECD, was launched, under which AI developers can disclose their AI risk-management practices. The G7 has continued to engage closely on AI issues. A May 2026 Ministerial Declaration on Digital & Technology **highlighted** “strengthening the digital economy by reaffirming our commitment to an open, innovationfriendly approach to digital environments, with particular emphasis on supporting AI openness and fostering AI adoption by micro-, small, and mediumsized enterprises . . . while ensuring that AI is secure.” The 52nd G7 Summit in Évian-les-Bains, France, took place in June and featured in a working lunch on AI with industry heads, though it did not lead to a new leaders’ declaration.

The **UN Global Dialogue on AI Governance** was established by General Assembly resolution in 2025 and convened its first session in Geneva in July 2026. The dialogue’s design is deliberately deliberative: It convenes discussion on AI-related issues among representatives from government, industry, civil society, and academia, but it does not issue binding recommendations. The second session for this body is scheduled to occur in New York in May 2027. Notably, China has invested in the UN track while simultaneously building WAICO outside of it. For example, in May 2026, China’s vice minister of science and technology, Chen Jiachang, **spoke at the United Nations** about how “China is promoting global AI governance within the UN framework.” This shows that China is taking an aggressive dual-track diplomatic approach to the regulation of AI, which is a strategy that the United States has not matched.

Globally, **the EU AI Act**, adopted in 2024, is the most legally consequential of the existing frameworks to regulate AI. This act has not been designed for global coalition building; the European Union has not sought to enroll other countries in the act, and the bloc itself has signed the Pax Silica Declaration. But the act does serve as a Western bloc model for AI regulation that could be expanded and exported to other countries, should the need arise, and the European Union is advocating to other countries to adopt national AI frameworks modeled after the act—although this is happening over the objections of U.S. policymakers, who do not support the act itself or resting international agreement on its principles.

Finally, starting in 2023, a global series of AI summits began, initially focused on AI safety. However, over the years this AI summit series has shifted away from having AI safety as its organizing frame. For example, the 2023 summit in Bletchley Park was explicitly an “AI Safety Summit” and **produced** the Bletchley Declaration on AI safety. But the following year, the 2024 Seoul meeting **broadened**

the frame to “safety, innovation, and inclusivity,” though it did produce a **new round** of frontier safety commitments. This broadening of scope continued: The 2025 Paris meeting was an “**AI Action Summit**” and the 2026 New Delhi meeting was an “**AI Impact Summit**.” While the international network of AI safety and security institutes continues to function, and the second **International AI Safety Report** was presented in New Delhi, the political center of gravity for these summits has shifted from safety and governance toward AI access and deployment.

The U.S.-China Summit

The United States and China are expected to hold their first official bilateral AI dialogue under the current administration on September 24, as part of President Xi’s planned visit to Washington. In these AI talks, which were an outcome of the **May 2026 summit** between President Trump and President Xi in Beijing, the U.S. side is expected to be led by Treasury Secretary Scott Bessent.

According to current reporting, the **agenda** for these talks includes issues such as cooperation on monitoring AI-directed cyberattacks, Chinese access to U.S. frontier AI models, U.S. restrictions on Chinese open-weight AI models, and whether China will regulate open-source AI, with a **broader focus** on preventing the most capable systems from reaching nonstate actors. It is highly likely that ongoing developments—including **new revelations** of autonomous AI agent breaches and the “pace the frontier” push—will be key parts of the discussions.

Analysts following the process have suggested that this first meeting is unlikely to resolve major geopolitical or trade disputes. Rather, the summit will likely focus on establishing baseline definitions, such as what qualifies as a frontier AI model, that can be used to support future follow-on discussions, although the urgency of cyber cooperation has increased dramatically in recent weeks following multiple and ongoing revelations of autonomous AI hacking.

Two structural obstacles may limit what is achievable in this summit. The first is a series of U.S. policy decisions that have antagonized Chinese officials. For example, export controls on advanced technology remain a major **friction point** in the U.S.-China relationship—Chinese labs continue to struggle with obtaining computing resources due to export controls on high-end chips. Meanwhile, in June 2026, Washington used export controls to **restrict** non-U.S. nationals from accessing the Anthropic Fable model. The United States is likewise aggressively targeting “**distillation attacks**” by Chinese labs, and has reportedly considering taking actions to block or limit access to certain Chinese open-weight models. These are exactly the kind of actions President Xi targeted with his “overstretching national security” rhetoric. China has responded to these controls by seeking to **stimulate domestic innovation** in AI and related technologies, such as semiconductors.

Verification is another unsolved issue that may limit the discussions at the summit. In the AI space, **verification** means establishing that claims about an AI model’s capabilities and safety properties, and how an actor is using that model, are true. Currently, there is no established technical means by which either the United States or China can confirm the other’s compliance with a commitment about model capabilities or testing, and neither government is inclined to accept the other’s assurances. There are, however, a number of technical approaches—such as workload monitoring and verification of model integrity—under active development and could serve as a technical backbone for cross-border compliance verification.

From the Bilateral Table to the Global Arena

While the upcoming summit is being framed as a bilateral discussion between two leading global powers, it also carries significant multilateral implications. Both Washington and Beijing will likely use whatever emerges from this summit as currency for the wider institutional contest over AI governance described above. Additionally, each country's current multilateral commitments may shape what it is able to offer in bilateral negotiations. Therefore, it is best to view these bilateral and multilateral AI governance discussions as interlinked, rather than as separate processes.

For Beijing, one major benefit from the bilateral discussions with Washington is optics. Being seen as willing to negotiate frontier AI safety with the United States establishes China as a co-architect of global technology standards, and signals to the broader international community that the future of AI governance is bipolar rather than unipolar. This **supports the CCP's claim** that China is the great power, through WAICO, that is “[redefining] AI governance for the Global South.”

China will also likely seek to integrate outputs from the summit with the United States into its multilateral AI diplomacy. For example, any language on which the United States and China agree, such as a shared definition of a frontier model, could then be integrated into WAICO's work program and the UN Global Dialogue, giving compromise language multilateral traction. Conversely, should the talks instead collapse, or should Washington be seen as making coercive demands, Beijing has a ready alternative message for the same global audience: The United States prefers control of AI to global cooperation. China's multilateral commitments around AI will also constrain what it can offer in its bilateral talks with the United States. For example, any U.S. ask that China accept restrictions on the **global release** of Chinese open-weight AI models runs contrary to WAICO's promise of freely available model weights, training programs, and turnkey applications. Thus, Chinese concessions to the United States on open-source AI could be seen by WAICO members as Beijing trading away the Global South's access to the technology in exchange for accommodation with Washington. (As an aside, one important point to note here is that there is an ongoing discussion among Chinese policymakers about **restricting foreign access** to some advanced Chinese models.) Going into the talks with China, U.S. negotiators should understand how China will leverage this bilateral discussion in multilateral fora.

Declining to engage China on AI safety at a time when autonomous AI incidents are becoming a subject of global concern would cede global leadership of this issue to China.

For the United States, bilateral engagement with China around AI may also sit awkwardly alongside the reported U.S. push for countries to **choose between the two sides**. However, declining to engage China on AI safety at a time when autonomous AI incidents are becoming a subject of global concern would cede global leadership of this issue to China. Therefore, the U.S. focus should be on ensuring the outputs from these bilateral talks do not remain bilateral, and U.S. negotiators should treat summit outputs as exportable products that can also be used in other multilateral fora.

Recommendations for U.S. Policymakers

This analysis of ongoing international AI governance provides a roadmap for how U.S. policymakers can advance their AI objectives at the summit and ensuing engagements. The following recommendations are weighted toward the international dimension of AI governance. One key point here is that the global dimension of AI governance is where U.S. strategy is currently the least developed, vis-à-vis China, and thus it is where the U.S. strategic position is most rapidly eroding.

1. **Approach the summit with a narrow, focused AI safety agenda.** U.S. goals for the summit on AI should be narrowly focused on a concrete set of actions, or development of the related technical and policy infrastructure, that can break the current prisoner's dilemma dynamics that create a "race to the bottom" on safety while maintaining the ability for Chinese and U.S. companies to continue rapid iteration and competition in AI products. This should include definitional alignment on frontier AI model capability, nonstate actor proliferation, formal channels for one country to report or notify the other about serious incidents, baseline approaches to safety testing practices and procedures, and model-weight security. U.S. policymakers should not seek a comprehensive AI framework in these talks but rather aim to produce a shared set of U.S.-China definitions and a working channel for reporting, especially for cyber incidents, both of which will be foundational for further bilateral collaboration on AI governance. As much as possible, U.S. negotiators should draft these outputs so that they can be applied more broadly in multilateral settings: Definitions, reporting formats, and testing baselines written for adoption by the safety institute network, the G7 and OECD processes, and international standards bodies will retain U.S. authorship and promote U.S. influence in the AI governance space.
2. **Prioritize cyber incident data collection.** U.S. policymakers should seize the opportunity at the summit to establish a mutually acceptable process to monitor and share information on cyber incidents. This will need to include a robust and clear definition of incidents, specific conditions under which the Chinese and U.S. governments will share information with each other and other countries without implicating business-sensitive or national security information, and domestic reforms to ensure robust information sharing between frontier labs in each jurisdiction and their respective home governments.
3. **Focus on verification as a technical program rather than a diplomatic issue.** Verification will be a stumbling block for any bilateral or multilateral AI testing arrangement. Here, the United States should fund and internationalize work on verification methods, such as hardware attestation, compute accounting, structured access protocols for third-party evaluation, making the issue a technical matter rather than a diplomatic one. The most important considerations here are mechanisms for establishing mutual recognition and trust in verification tools, such as funding open-source verification tools or establishing vetting organizations comprised of neutral technical experts.
4. **Establish a domestic model for AI governance that is worth exporting.** The United States cannot credibly propose international standards for regulating frontier AI while at the same time lacking a coherent national set of regulations. China has realized this and is therefore making a dedicated AI law. Therefore, developing a federal framework to regulate AI would give U.S. diplomacy a model to offer other countries, and it should be a priority coming into and out of the summit. This U.S. federal framework could draw on the emerging state-level consensus around AI, which in turn often draw from points of consensus in industry. Prior to the summit, U.S.

policymakers should engage in a sprint to identify priority key provisions within state bills and federal proposals that can anchor summit discussions; following the summit, policymakers should focus on driving toward a robust AI framework that implements key summit outcomes while signaling to governments around the world the U.S. intention to seriously address AI governance.

5. **Design any U.S. AI screening body for global interoperability from the start.** As part of a possible domestic model for AI governance, the United States should consider establishing a standards body for frontier AI safety or developing this functionality within an existing organization—an approach consistent with the proposal made by Demis Hassabis. If this approach is adopted, right from the start this organization should be linked to existing global AI safety and security processes, especially those in G7 and OECD countries as well as multistakeholder standards bodies such as the International Organization for Standardization and the International Electrotechnical Commission, which are jointly working on cybersecurity and **technology standards**. Additionally, this body should have a statutory backstop (such as being grounded in an act of Congress, rather than being voluntary). It should also have conflict-of-interest firewalls for the organization’s executive and board of directors, as well as a clear appeals process for deployment denials. This body should also leave a technical door open for eventual narrow reciprocity with China on model-weight security and nonstate actor proliferation.
6. **Compete for middle powers on the same terrain as China.** WAICO is compelling in part because of what it offers members: training, model access, and an institutional voice. Pax Silica, on the other hand, offers primarily supply chain participation. Furthermore, the existing U.S. full-stack AI **export program** is commercial in orientation rather than a capacity-building or membership offer. While this order is helpful to promote U.S. AI technology globally, to more fully compete with WAICO, the United States should invest in an AI capacity-building and compute-access offer aimed at low- and middle-income countries. Here, a historical analogy might be the Eisenhower administration’s **Atoms for Peace** initiative, where the United States provided peaceful nuclear technology, research reactors, and fissile material to developing nations in exchange for tracking safeguards and alignment with the Western geopolitical bloc. Particularly useful in this analogy are the safeguards aspect of the program, given links between Atoms for Peace and nuclear proliferation. Thus, an AI equivalent would need to pair compute and model access with model-weight security commitments and evaluation access from the outset. Applied to AI, an initiative of this type could be the “carrot” alongside the current U.S. export control “stick.”
7. **Engage with global institutional dialogues on AI safety, rather than ceding this to China.** Currently, China is pursuing a dual-track strategy: building WAICO outside the UN system while simultaneously investing in the UN Global Dialogue and the Scientific Panel. The United States, meanwhile, is **substantially absent** from this track. While maintaining its sovereignty over AI, Washington should also engage with the UN dialogues on AI, to build global alignment around U.S. perspectives. Given that the dialogue’s next session is in New York in May 2027, and that its deliberative design makes it a low-risk venue for U.S. involvement, nonparticipation forfeits influence while providing no strategic gain. Beyond the United Nations, other existing multilateral venues, such as the G7, G20, Asia-Pacific Economic Cooperation, and international standards bodies all provide opportunities for the United States to promote its vision for AI regulation. Currently, engagement is lacking.

Conclusion

Global AI regulation is heading toward fragmentation. Two different architectures to govern AI, with separate constituencies, are consolidating. At the same time, a third architecture for AI governance, led by the United Nations, is without capacity. Nothing here is indicative of interoperability or integration, and global AI governance questions remain subordinate to geopolitical issues related to the broader U.S.-China strategic rivalry.

However, the outcome for global AI governance is not yet finalized. This is clear when looking at which countries have joined which framework: Globally, many countries have not yet decisively chosen one framework over the other. Rather, countries around the world that will matter for the future of AI regulation have joined multiple frameworks, seeking benefits from each while keeping their strategic options open. In Washington, this is often read as fence sitting. However, perhaps it is better read as an invitation: Middle powers are telling both the United States and China that neither great power yet offers an AI governance framework compelling enough to command exclusivity.

The United States needs to make an offer that speaks to what most of the world actually wants from AI governance: access, capability, and a say in the rules.

Therefore, the upcoming September AI summit between the United States and China provides a reminder to Washington: It is competing against Beijing not just for a technological lead, but also for institutional membership when it comes to approaches to regulation. Here, the United States needs to make an offer that speaks to what most of the world actually wants from AI governance: access, capability, and a say in the rules. The 28 countries that joined China in establishing WAICO did so because they were offered inclusion and a say. This is the competition, and it will continue, no matter the outcome of the summit. ■

Chris Collins is a senior associate (non-resident) with the Wadhwani AI Center at the Center for Strategic and International Studies (CSIS) in Washington, D.C., and a Fellow at the Cascade Institute at Royal Roads University in Victoria, Canada. Aalok Mehta is the director of the Wadhwani AI Center at CSIS.

This report is made possible by general support to CSIS. No direct sponsorship contributed to this report.

This report is produced by the Center for Strategic and International Studies (CSIS), a private, tax-exempt institution focusing on international public policy issues. Its research is nonpartisan and nonproprietary. CSIS does not take specific policy positions. Accordingly, all views, positions, and conclusions expressed in this publication should be understood to be solely those of the author(s).

© 2026 by the Center for Strategic and International Studies. All rights reserved.