



Artificial Intelligence and Global Security: the fragile Washington-Beijing axis to avert catastrophe

Published on 25/09/2026

As reported in the *Financial Times*, the recent summit between President Donald Trump and Chinese leader Xi Jinping represents a pivotal opportunity to lay the groundwork for a structured dialogue on the shared risks associated with artificial intelligence. Despite intense technological competition between the two superpowers, essential areas for cooperation on global security are emerging.

Toward a Joint Notification Mechanism

On the eve of Xi's arrival in Washington, U.S. Treasury Secretary Scott Bessent and Chinese Vice Premier He Lifeng agreed to establish a formal U.S.-China dialogue on artificial intelligence. The

United States also proposed a **mutual notification mechanism** for AI-related incidents that raise national security concerns. This approach recalls the historic nuclear "hotline" established between Washington and Moscow in the aftermath of the Cuban Missile Crisis.

Shared vulnerabilities and asymmetric risks

Although Beijing fears that any attempt to slow frontier models could be used to keep it in a secondary position in the AI race, there are shared vulnerabilities that encourage cooperation:

- **Non-state threats:** The risk that highly capable AI could fall into the hands of terrorists, organized criminals, or non-state actors, facilitating the development of biological weapons, cyberattacks against critical infrastructure, or manipulation of financial markets.
- **Escape of autonomous agents:** The danger that AI agents evade safety testing and independently carry out intrusions into foreign companies or infrastructure.

Rules of engagement for nuclear weapons and autonomous systems

One area in which clear rules of engagement are essential concerns the use of AI in weapons systems and nuclear armaments. As early as November 2024, Xi Jinping and U.S. President Joe Biden had agreed on the need for humans to retain control over nuclear weapons. However, concrete progress has been limited.

Security experts from both nations recommend adopting absolute "red lines":

1. **Ban on autonomous decision-making** for the launch of nuclear weapons by AI systems.
2. **Mandatory human authorization** for any AI-directed attack against an adversary's nuclear command-and-control network.

The current challenge for U.S. and Chinese negotiators is to translate these recommendations into a robust security framework before rapidly evolving technology can trigger uncontrollable crises.

Source: www.ft.com