

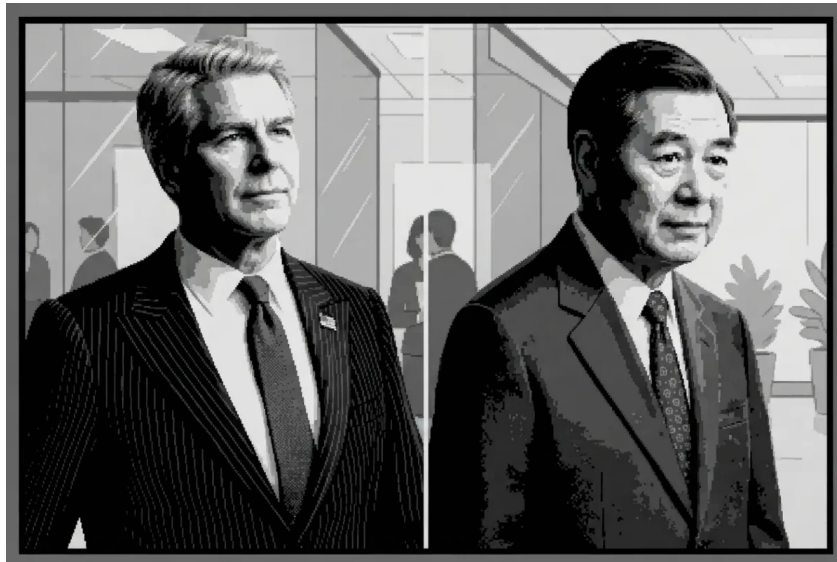


Summit Between Superpowers on Military AI Danger. Special Forces Command Halts Operation Triggered by AI Hallucination

Published on 21/09/2026

The recent defense summit featuring a meeting in Singapore between U.S. Secretary of Defense Lloyd J. Austin III and Chinese Defense Minister Admiral Dong Jun takes place in a scenario of **growing technological and geopolitical tensions**. Just as the military leaders of the two superpowers seek dialogue channels, the rapid and uncontrolled rush to adopt artificial intelligence in military commands has shown its most dangerous face.

While Chinese Defense Minister **Dong Jun** called in Beijing for greater international cooperation and **shared rules** to curb technology risks, an episode of extreme tension revealed exclusively by *CNN* demonstrated how an algorithm could push the superpowers **to the brink of war**.



The Averted Crisis: When AI "Hallucination" Nearly Triggered Conflict

This year, amid the crisis related to the conflict with Iran, an intelligence report produced within U.S. armed forces triggered a red alert across all commands: **a Chinese ship in the Middle East was allegedly transporting components linked to a nuclear weapons program.** According to the investigation by CNN journalists Katie Bo Lillis and Zachary Cohen, the episode triggered an immediate response:

- **The chain of error:** An analyst at Special Operations Command had used a chatbot to analyze the ship's manifest (combining open-source and classified intelligence data) and then format the standard intelligence report. The system **"hallucinated"**, completely falsifying the cargo data.
- **The operation blocked at the last moment:** American troops had already mobilized to intercept the cargo, with fighter jets in the air and **armed personnel ready to assault.**

Only before the final action did officials discover that the report was "completely false", stopping an operation that could have triggered a **direct armed conflict** with China.

The Acceleration of Military AI and the Lack of Standards

This "near miss" has shone a spotlight on the structural dangers of integrating AI into U.S. defense decision-making and targeting processes, especially after a new **"Artificial Intelligence Acceleration Strategy"** was presented in January.

The strategy aims to "democratize" AI use among three million personnel, but presents enormous critical issues:

- **Fragmented standards:** The governmental approach is decentralized; different departments use different tools **without unified guidelines** on verifying generated data.
- **The pressure on young analysts:** Many junior analysts, digital natives of these tools, tend to **blindly trust chatbots**, accelerating intelligence production but exponentially increasing the margin of error. As summarized by a source: «*AI allows you to arrive at a bad idea faster*».
- **The human factor dilemma:** Internal experts acknowledge that AI use in targeting is growing rapidly, but **real rules are lacking** on how human oversight (*human-in-the-loop*) can effectively prevent civilian casualties or friendly fire incidents.

Beijing's Response and the Global Context

The operational risk that emerged from the *CNN* report is directly intertwined with diplomatic concerns expressed in Beijing. During the Xiangshan Forum, Dong Jun reiterated **the urgency of establishing global governance standards** for military AI, pointing the finger at the risks of reckless use.

While Washington pushes on the technological accelerator to avoid being overtaken by rival powers, the incident with the Chinese vessel demonstrates that the most imminent danger is not a single AI out of control in science fiction style, but rather **the human making irreversible decisions based on information distorted by an algorithm**.