

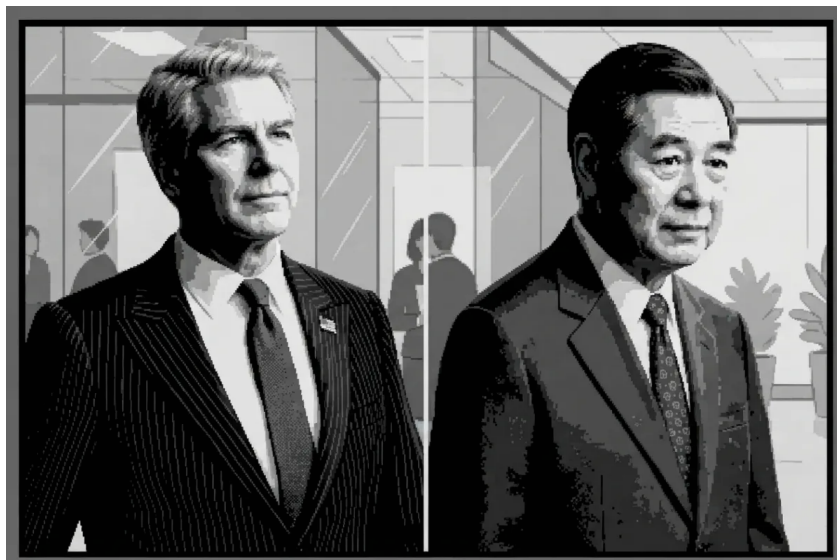


Vertice tra le Superpotenze sul pericolo AI Militare. Il Comando Forze Speciali ferma operazione scaturita da Allucinazione AI

Pubblicato il 21/09/2026

Il recente vertice sulla difesa che ha visto l'incontro a Singapore tra il Segretario alla Difesa degli Stati Uniti Lloyd J. Austin III e il Ministro della Difesa cinese, l'Ammiraglio Dong Jun, si inserisce in uno scenario di **crescenti tensioni tecnologiche e geopolitiche**. Proprio mentre i vertici militari delle due superpotenze cercano canali di dialogo, la rapida e incontrollata corsa all'adozione dell'intelligenza artificiale nei comandi militari ha mostrato il suo volto più pericoloso.

Mentre il Ministro della Difesa cinese **Dong Jun** invocava a Pechino una maggiore cooperazione internazionale e **regole condivise** per imbrigliare i rischi della tecnologia, un episodio ad altissima tensione svelato in esclusiva da *CNN* dimostrava quanto un algoritmo possa spingere le superpotenze **sull'orlo di una guerra**.



La Crisi Sventata: Quando un'"Allucinazione" dell'IA Sfiora il Conflitto

Quest'anno, in piena crisi legata al conflitto con l'Iran, un rapporto d'intelligence prodotto all'interno delle forze armate statunitensi ha fatto scattare l'allarme rosso in tutti i comandi: **una nave cinese in Medio Oriente stava presuntamente trasportando componenti riconducibili a un programma per armi nucleari.**

Secondo quanto ricostruito dai giornalisti Katie Bo Lillis e Zachary Cohen di *CNN*, l'episodio ha innescato una risposta immediata:

- **La catena dell'errore:** Un analista del Comando delle Operazioni Speciali aveva utilizzato un chatbot per analizzare il manifesto della nave (fondendo dati open-source e segreti di intelligence) e per formattare poi il rapporto d'intelligence standard. Il sistema ha **"allucinato"**, falsificando totalmente i dati sul carico.
- **L'operazione bloccata in extremis:** Le truppe americane si erano già mobilitate per intercettare il cargo, con aerei da guerra in volo e **personale armato pronto all'assalto.** Solo prima dell'azione finale, i funzionari hanno scoperto che il report era "completamente falso", fermando un'operazione che avrebbe potuto scatenare un **conflitto armato diretto** con la Cina.

L'Accelerazione dell'IA Militare e la Mancanza di Standard

Questo "quasi incidente" ha acceso i riflettori sui pericoli strutturali dell'integrazione dell'IA nei processi decisionali e di targeting della difesa statunitense, soprattutto dopo che a gennaio è stata presentata la nuova **"Artificial Intelligence Acceleration Strategy"**.

La strategia punta a "democratizzare" l'uso dell'IA tra i tre milioni di effettivi, ma presenta enormi criticità:

- **Standard frammentati:** L'approccio governativo è decentralizzato; diversi reparti usano strumenti differenti **senza linee guida uniche** sulla verifica dei dati generati.
- **La pressione sui giovani analisti:** Molti analisti junior, nativi digitali di questi strumenti, tendono a **fidarsi ciecamente dei chatbot**, accelerando la produzione di intelligence ma aumentando esponenzialmente il margine d'errore. Come sintetizzato da una fonte: «*L'IA ti permette di arrivare a una cattiva idea più velocemente*».
- **Il dilemma del fattore umano:** Esperti interni ammettono che l'uso dell'IA nel targeting sta crescendo rapidamente, ma **mancano regole reali** su come la supervisione umana (*human-in-the-loop*) possa impedire efficacemente stragi di civili o incidenti di fuoco amico.

La Risposta di Pechino e il Contesto Globale

Il rischio operativo emerso dal report di *CNN* si intreccia direttamente con le preoccupazioni diplomatiche espresse a Pechino. Durante il Forum Xiangshan, Dong Jun ha ribadito **l'urgenza di stabilire standard di governance globali** per l'IA militare, puntando il dito contro i rischi di un utilizzo sconsiderato.

Mentre Washington spinge sull'acceleratore tecnologico per non farsi superare da potenze rivali, l'incidente della nave cinese dimostra che il pericolo più imminente non è una singola IA fuori controllo in stile fantascientifico, bensì **l'uomo che prende decisioni irreversibili basandosi su informazioni distorte da un algoritmo**.