

Investigating the Harmful Implications of Generative AI in the Military Field

Mohammed Yasser¹

¹Delhi Private School Dubai

Publication Date: 2025/09/29

Abstract: Generative Artificial Intelligence (AI) has rapidly emerged as both a technological innovation and a global security concern. Its application in the military domain raises unique ethical, legal, and strategic challenges. This paper examines harmful implications of generative AI in warfare, supported by published data from surveys and policy reports. Ipsos (2023) found that 69% of respondents globally are concerned about autonomous weapons, and 73% are concerned about surveillance misuse. Similarly, a UK government survey highlighted that 45% of respondents fear job displacement, 35% worry about loss of human creativity, and 34% fear losing control over AI. Meanwhile, Brookings (2018) found that only 30% of respondents support AI use in warfare, but support increases to 45% if adversaries are already using AI-based weapons. These statistics reflect widespread societal concern about the destabilizing consequences of AI militarization. This paper analyzes these concerns across five domains: misinformation, autonomous weapons, accountability, bias in decision support, and adversarial vulnerabilities. It argues that generative AI may exacerbate risks of misinformation campaigns, unlawful targeting, biased decision-making, and loss of accountability, demanding urgent international regulation.

Keywords: *Generative AI, Military Applications, Autonomous Weapons, Misinformation, Ethics, Security Risks.*

How to Cite: Mohammed Yasser (2025) Investigating the Harmful Implications of Generative AI in the Military Field. *International Journal of Innovative Science and Research Technology*, 10 (9), 2046-2048.
<https://doi.org/10.38124/ijisrt/25sep1357>

I. INTRODUCTION

The acceleration of artificial intelligence research has generated tools that can automate tasks once thought to be exclusive to human intelligence. Generative AI, including large language models (LLMs) and deep learning systems capable of creating synthetic text, imagery, video, and audio, has found applications in diverse sectors such as education, healthcare, finance, and entertainment. However, its entry into military systems presents profound risks to international security and humanitarian law.

In recent years, militaries have begun experimenting with AI-powered systems for intelligence gathering, battle simulation, and autonomous targeting. The dual-use nature of generative AI makes it particularly dangerous—it can improve simulations and training environments but also produce convincing deepfakes for disinformation campaigns. For instance, during the Russia-Ukraine conflict, deepfakes of Ukraine's president calling for surrender were widely circulated before being debunked. Such incidents illustrate how generative AI can disrupt public trust and military morale.

Survey evidence shows broad public anxiety. Ipsos (2023) reported that nearly 70% of global citizens worry about AI-enabled autonomous weapons, while 73% are concerned about surveillance misuse. The UK Government's

Wave 3 AI Attitudes survey found that 34% of respondents fear losing control of AI systems altogether, highlighting deep societal unease. Moreover, Brookings (2018) demonstrated that public support for AI in warfare is conditional—only 30% support it outright, though this rises to 45% if adversaries are already deploying AI weapons. These findings suggest that acceptance of AI militarization is rooted in fear of being left vulnerable, rather than genuine trust in the technology.

This paper focuses on harmful implications of generative AI in the military, emphasizing five core domains: misinformation, autonomous weapon risks, accountability gaps, biased decision-making, and adversarial vulnerabilities. The analysis also integrates published survey data, academic research, and case studies to demonstrate the urgency of global governance and safeguards.

II. METHODOLOGY

This research adopts a qualitative synthesis methodology, analyzing secondary data from academic articles, policy reports, and survey results published by trusted organizations including Brookings, Ipsos, and government databases. The study categorizes harmful implications of generative AI into thematic areas and supplements analysis with visualizations based on published statistics. The methodology is not experimental but

interpretive, highlighting relationships between technological capabilities and associated risks.

III. RESULTS

The following figures and tables summarize published data on public perceptions of generative AI risks in military applications.

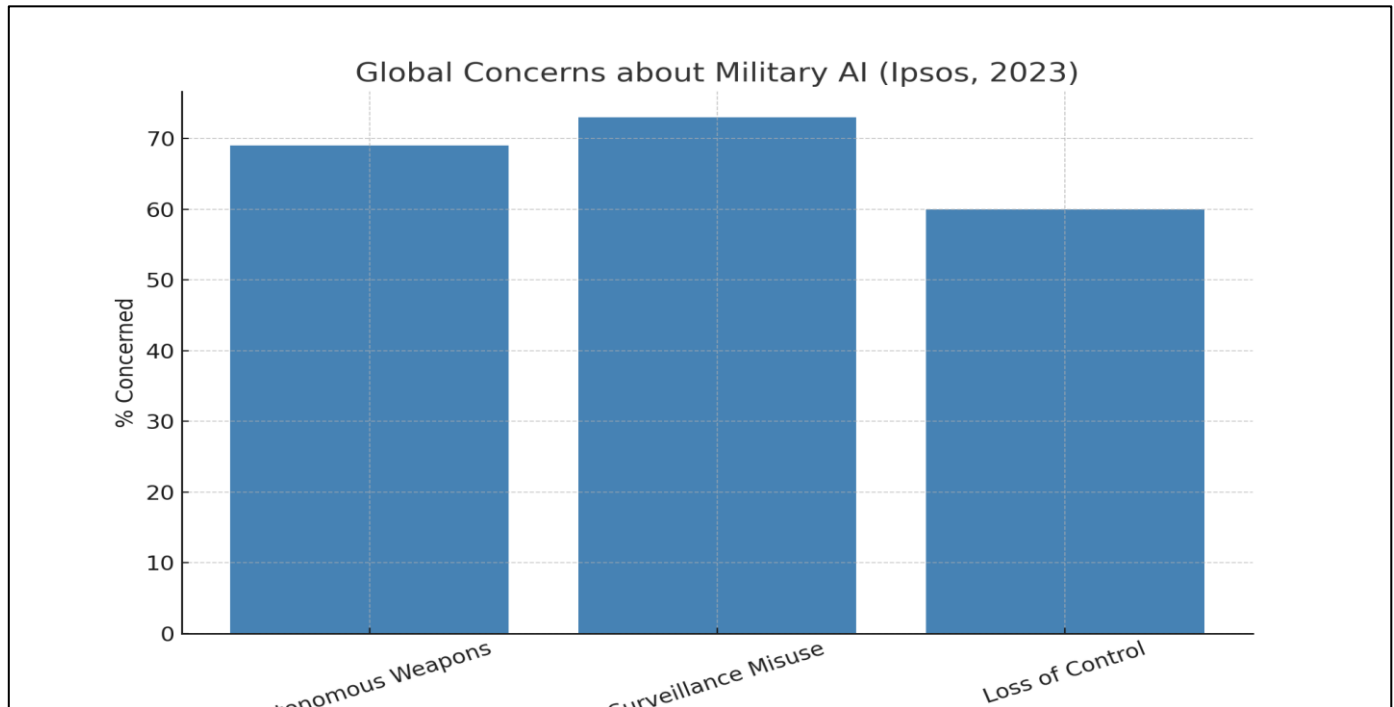


Fig 1 Global Concerns About AI-Enabled Military Threats (Ipsos, 2023).

Table 1 Presents Comparative Survey Results on AI Risks and Warfare Applications.

Risk / Scenario	Ipsos (2023) Global %	UK Gov (2023) %	Brookings (2018) %
Autonomous Weapons Misuse	69	-	30 baseline / 45 if adversary uses AI
Surveillance Misuse	73	-	-
Loss of Control	60	34	-
Job Displacement	-	45	-
Bias/Unfair Outcomes	-	14	-

IV. DISCUSSION

The data reveal widespread public concern regarding AI in military systems. Ipsos (2023) findings show nearly three-quarters of respondents fear misuse of AI in surveillance, aligning with ethical concerns raised in academic literature. Brookings (2018) demonstrates that support for AI weapons is reactive rather than proactive, reflecting a classic security dilemma—states may feel forced to adopt AI militarization simply because adversaries do.

Comparing across surveys, it becomes clear that misinformation and loss of control are central themes. The UK survey highlights public anxieties around creativity loss and bias, issues less visible in the Ipsos data but critical for democratic accountability. These results emphasize that risks are not only technical but deeply social, involving trust, legitimacy, and international stability.

The evidence suggests that unchecked generative AI in military contexts could destabilize both domestic societies

and global order. To mitigate such risks, policymakers should prioritize international treaties banning fully autonomous lethal weapons, require human-in-the-loop systems for targeting, and invest in robust detection mechanisms for synthetic media.

➤ Harmful Implications

• Misinformation & Psychological Warfare

Generative AI amplifies the speed and realism of propaganda. Deepfakes can produce fabricated orders, false casualty footage, or fake endorsements from leaders—tools perfectly suited for psychological warfare and influence operations. States and non-state groups can weaponize synthetic content to sow confusion in civilian populations and among military units, complicating crisis management and escalating tensions.

• Autonomous Weapon Risks

Generative models integrated into sensor-processing or decision pipelines can introduce novel failure modes. An AI

that (even partially) generates or suggests targeting priorities may misclassify civilians as combatants, or be fooled by spoofed sensor data. When humans rely on AI-generated recommendations without rigorous oversight, the chances of tragic errors rise.

- *Accountability and Legal Challenges*

International Humanitarian Law presumes human decision-making in lethal force. Delegating lethal choices—or even influential advice—to generative systems blurs responsibility. Determining legal liability becomes complex: software developers, system integrators, commanders, and operators could all bear overlapping culpability, making post-incident justice and remediation difficult.

- *Biased Decision Support*

Generative models encode biases present in training data. If military decision-support systems are trained on skewed historical records, they may reproduce discriminatory patterns or culturally blind assumptions, leading to poor strategic choices in diverse theatres.

- *Adversarial Vulnerabilities*

Generative models are not immune to poisoning or adversarial attacks. Opponents can manipulate input data or craft adversarial examples that degrade performance in surveillance, target recognition, or communications—turning AI from an advantage into a liability.

V. CONCLUSION

Generative AI presents transformative potential but poses significant dangers when applied in the military field. Published surveys show overwhelming global concern about AI-enabled autonomous weapons, surveillance misuse, and loss of human control. Case studies such as the use of deepfakes in modern conflicts already illustrate its power to destabilize societies. Without regulation, generative AI risks exacerbating misinformation campaigns, undermining international law, and eroding accountability in warfare. This paper argues for urgent implementation of global safeguards, transparency standards, and cooperative treaties to prevent harmful military applications of generative AI.

REFERENCES

- [1]. Ipsos (2023). Halifax International Security Forum Report: Public Perceptions of AI. <https://www.ipsos.com/en-th/halifax-report-2023-ai>
- [2]. UK Government (2023). Public Attitudes to Data and AI Tracker Survey, Wave 3. <https://www.gov.uk/government/publications/public-attitudes-to-data-and-ai-tracker-survey-wave-3>
- [3]. Brookings Institution (2018). Brookings Survey on AI in Warfare. <https://www.brookings.edu/articles/brookings-survey-finds-divided-views-on-artificial-intelligence-for-warfare>
- [4]. MDPI (2019). Artificial Intelligence Applications in Military Systems. *Electronics*, 10(7), 871. <https://www.mdpi.com/2079-9292/10/7/871>

- [5]. Slattery, S. et al. (2024). The AI Risk Repository: Taxonomies of Risks. arXiv preprint arXiv:2408.12622. <https://arxiv.org/abs/2408.12622>
- [6]. RAND Corporation. "Generative Artificial Intelligence Threats to Information Ecosystems." RAND, 2024. <https://www.rand.org/pubs/perspectives/PEA3089-1.html>
- [7]. ICCT. "The Weaponisation of Deepfakes." ICCT Policy Brief, 2023. <https://icct.nl/sites/default/files/2023-12/The%20Weaponisation%20of%20Deepfakes.pdf>
- [8]. U.S. Department of Defense. "Contextualizing Deepfake Threats to Organizations." CSI Deepfake Threats, 2023. <https://media.defense.gov/2023/Sep/12/2003298925/-1/-1/0/CSI-DEEPFAKE-THREATS.PDF>
- [9]. Brookings. "Deepfakes and International Conflict." Brookings, 2023. <https://www.brookings.edu/articles/deepfakes-and-international-conflict/>
- [10]. Carnegie Endowment. "Understanding the Global Debate on Lethal Autonomous Weapons Systems." 2024. <https://carnegieendowment.org/research/2024/08/understanding-the-global-debate-on-lethal-autonomous-weapons-systems-an-indian-perspective?lang=en>