

THE COST OF TEMPO

Military AI, the Acceleration of Targeting, and Compliance with International Humanitarian Law



By **Mohamed Badr
Er-rahaoui**

Under the mentorship of:

Abdelhak Bassou

Senior Fellow, Policy Center for the New South

Driss Alaoui Belghiti

Senior International Relations Specialist, Policy
Center for the New South

About the Mentorship Program

The Mentorship Program, jointly organized by the Policy Center for the New South (PCNS) and the Youth Policy Center (YPC), was launched to support master's students in developing policy-oriented research and writing skills, and to strengthen their contribution to public policy debates.

Implemented from March to June 2026, the pilot edition brought together six selected master's students from universities across Morocco, each paired with a mentor duo composed of a PCNS Senior Fellow and a PCNS Expert. The policy briefs developed during the program covered key themes including geopolitics and conflict dynamics, international cooperation, energy and sustainable development, macroeconomic and trade policies, labor market and human capital, and equality and inclusion.

Through methodological workshops, individual mentorship, and continuous feedback, participants were guided in transforming their research ideas into concise, evidence-based policy briefs. The program fostered intergenerational dialogue and knowledge exchange between emerging researchers and experienced policy experts, while providing participants with practical tools to define policy-relevant problems, structure their analysis, formulate actionable recommendations, and communicate their research to policy audiences.

This policy brief is the result of the mentorship and research process undertaken as part of the program's first cohort. It reflects the program's ambition to amplify youth-led policy research while nurturing a new generation of researchers committed to evidence-based policymaking.

For inquiries, please contact: mentorship@policycenter.ma

List of Abbreviations:

AI	: Artificial Intelligence
AWS	: Autonomous Weapon System(s)
CENTCOM	: United States Central Command
HOTL	: Human-on-the-Loop
IDF	: Israel Defense Forces
IHL	: International Humanitarian Law
ISR	: Intelligence, Surveillance, and Reconnaissance
LLM	: Large Language Model
MSS	: Maven Smart System
OODA	: Observe, Orient, Decide, Act (loop)

Executive Summary

Artificial intelligence has been progressively integrated into military domains because it accelerates the targeting cycle: by compressing the interval between observation and action, it allows for larger target selection and engagement of a far larger number of targets than would be possible through human workflows alone. In Ukraine, Gaza, and Iran, AI-enabled systems have created an asymmetry in operational tempo and scale between forces that use them and those that do not.

However, this advantage does not by itself determine the outcome of wars yet and is not without a cost. The same acceleration that makes military AI effective also undermines the human judgment on which compliance with International Humanitarian Law (IHL) depends. Biased and unrepresentative training data can lead to misidentification. Moreover, large language models (LLMs) can assert false information, and, above all, when the window for the human review of AI-generated target lists is compressed to seconds, oversight risks becoming more formal than substantive. Opacity, automation bias, and an emerging accountability gap compound the problem.

The central policy question is, therefore, not whether AI-enabled targeting systems should be adopted, but how to preserve accountability and restraint in the pursuit of speed. This brief recommends: (1) at the level of deployment, a guaranteed minimum human-review time, independent pre-deployment testing, systematic logging of targeting decisions, and a genuine capacity for operators to contest an algorithmic recommendation; (2) at the normative level, embedding these standards in IHL and multilateral frameworks; and (3) at the industry level, traceability obligations on the firms that build these systems.

Introduction

Military AI is attractive to military planners for one reason above all: its acceleration of the targeting process. This brief argues that the same acceleration strains human judgment, on which compliance with IHL depends. For most of the modern era, military advantage rested on mass, firepower, and industrial capacity. Today, the formula for a military victory is shifting toward the speed of data processing and decision-making; the militarization of AI sits at the very center of this shift. Thus, the rapid integration of AI into military operations has meant that an “algorithmic arms race” is no longer only a theoretical warning, but a practical reality, tested and developed in real time across the battlefields of Ukraine, Gaza, and Iran, which also serve as data-generation environments.

However, the relationship between material or technological superiority and the outcome of the conflict is not entirely deterministic and should not be overstated. Factors such as terrain, geography, human resources, politico-economic will, and tolerance of casualties have proved decisive in offsetting military capability disparities. The Vietnam War, the Soviet and American invasions of Afghanistan, and the war between Israel and Hezbollah, are commonly cited cases of conflicts in which technological edge could not secure political victory despite providing operational tempo.

This brief argues that the same accelerated targeting that makes military AI more appealing to military planners is also what strains the human control that sustains compliance with International Humanitarian Law (IHL). It proceeds to do so in three parts. The first section, *Faster Than Judgment: Military AI, the Acceleration of Targeting and the Strain on Human Control*, examines how AI compresses the targeting cycle and produces an asymmetry in the tempo and scale of the process. The second section, *The Cost of Acceleration: Data Quality, Human Judgment, and Compliance with International Humanitarian Law*, shows how the pace and opacity of AI-enabled targeting strains data quality, human control, and legal assessments needed for civilian protection. The third section, *Constraining the Pursuit of Speed: Safeguards across Deployment, Norms, and Industry*, sets out recommendations to restrain the pursuit of speed within enforceable limits.

01. Faster Than Judgment: Military AI, the Acceleration of Targeting, and the Strain on Human Control

The strategic competition in military AI, led primarily by the United States and China, is driven by the perception that falling behind technologically risks putting the nation in a state of strategic disadvantage (Güneysu, 2024). This competition is also driven by the promise of unmatched speed and tempo on the battlefield, compressing the “kill chain”—the observe, orient, decide, and act (OODA)—from hours to minutes, even seconds (Bovet, 2025).

This advantage of AI is most expressed in the intelligence, surveillance, and reconnaissance (ISR) domain, where its ability to process large data volumes at rates that are far faster than those of human operators, effectively reduces the effects of the “fog of war” (Faris, 2026).

The war in Ukraine illustrates the appeal of AI military integration. In a conflict swarmed by drones, satellite imagery and field reporting, the challenge is not gathering intelligence and information, but who processes it first (Yatsyk & Zharova, 2024). Ukraine's cloud-based DELTA system uses applications such as Veza and Avengers to classify Russian vehicles and equipment by analyzing thousands of simultaneous drone video feeds in an active combat zone (Matei & Reed, 2025). Additionally, the Maven Smart System (MSS), employed by the US XVIII Airborne Corps, has fused classified intelligence, commercial satellite imagery, and open-source data to geolocate Russian command posts behind the lines, enabling targeted artillery strikes (Russo, 2026). In each case, AI systems function as accelerators that generate operational advantage by shortening the interval between observation and action (King, 2024).

The same dynamic is largely visible in target-generation processes. In Gaza, the Israel Defense Forces (IDF) are documented to have used systems like "Gospel" to automate the identification process of hostile physical infrastructure, and "Lavender" which utilized machine learning to analyze the behavioral patterns of a large share of the population in Gaza (Faris, 2026). Reports indicated that the Lavender system was generating kill lists numbering in the tens of thousands (Russo, 2026), compressing the review process to seconds (Faris, 2026). During operation Epic Fury, US forces used the MSS to fuse up to 179 distinct data feeds across the space, air, land, and sea domains (Mande & Allen, 2026), to reportedly strike 1,000 targets in the opening day of the conflict (CENTCOM, 2026).

The scale of the resulting asymmetry is illustrated by the US Army's Scarlet Dragon exercise, in which a targeting cell of roughly 20 personnel using the MSS reportedly matched the output of a unit of around 2,000 operating without it, which is historically the same size as required by a time-critical targeting cell during operation Iraqi Freedom (Probasco, 2024).

However, faster and more efficient target-generation does not equate to a flawless process. Its use in Ukraine, Gaza, and Iran has raised many questions about safety and reliability.

02. The Cost of Acceleration: Data Quality, Human Judgment, and Compliance with International Humanitarian Law

Prioritizing operational tempo over safety testing exposes civilians and infrastructure to risks related to the quality of the data and the degree of human control. Data used to train those models are not free from bias, as these datasets reflect historical inequalities and prejudice based on race, gender, or class, and in a targeting context, these skews have operational consequences. If the data used to train the AI associates behavior with threats or overrepresents certain groups, neighborhoods, or movement patterns, the system can end up not only reproducing but also amplifying the "biased" associations, flagging individuals as hostile not because of their actions, but because of their resemblance to prior examples used to train the AI (Nuredin, 2025). Such is the case for predictive systems that assess the physical effects of weapons while relying on a "one-size-fits-men" approach, failing to account for varying body types or physical disabilities (Blanchard and Bruun, 2024).

The use of LLMs and generative AI models raises an issue of its own: hallucination (Probasco et al., 2025). Crucially, LLMs have a documented tendency of confidently asserting factually incorrect information and fabricating false rationales to satisfy the users' expectations. Clearly, in a battlefield scenario, this creates the risk of misidentification (Probasco et al., 2025).

To manage these risks, militaries mandate a "human-on-the-loop" (HOTL) arrangement, in which a human operator verifies and authorizes the strike on the target suggested by the AI (Longpre et al., 2022). However, as human operators are flooded with thousands of AI-generated target lists, they face severe cognitive overload and begin to place excessive trust in the automated system over their critical judgment; a phenomenon called automation bias (Nuredin, 2025). But when the review process of the AI-generated target lists is compressed to minutes or seconds, required substantive control risks being reduced to formal validation (Faris, 2026).

A further issue is that of opacity. The AI's neural networks act like a "black box", lacking transparency, that makes it practically impossible for a human operator to understand why the AI made the choice of flagging a target, or to conduct the case-by-case assessments of proportionality and distinction at the required level for compliance with IHL (Nuredin, 2025; UNIDIR, 2025).

As target lists grow faster and larger than operators can meaningfully review, it becomes harder for operators to remain in the loop and sustain substantive authority over each strike. As tempo rises, the conditions for human control erode, granting systems more autonomy to identify and engage targets without human oversight (Faris, 2026). As a result, this erosion of human oversight introduces new risks such as uncontrolled escalation: when deployed in a new environment, these systems can be highly brittle. An AI system can interpret a warning shot as a hostile action and retaliate before the human operator can intervene and de-escalate (Simmons-Edler et al., 2025).

Additionally, the deployment of autonomous weapon systems (AWS) creates the risk of an accountability gap: it would prove difficult to assign legal responsibility to an AI system committing a war crime in the absence of human oversight (Amoroso & Tamburrini, 2020).

03. **Constraining the Pursuit of Speed: Safeguards across Deployment, Norms, and Industry**

If militaries continue to pursue speed and operational advantage, constraint must operate where leverage exists: at the state deployment level, at the level of norm-setting, and finally, at the level of the firms that build said systems.

The proposed measures that follow are intended to be complementary and not alternatives to one another.

At the deployment level:

- Regular external testing for bias, error rates, and failure modes, rather than the assessments usually provided by the deploying force (Blanchard & Bruun, 2024).

- Realistic condition testing before deployment.
- Operators to be provided with a sufficient review window, information, and authority to contest and override an AI-proposed action (Nuredin, 2025).

At the international norm level:

- Continued norm-development grounded in IHL, to uphold human control and accountability as long-term safeguards (Milanovic, 2026).

At the level of the development firms:

- Creating investment funds, using public, multilateral or philanthropic capital, to acquire equity in these firms in order to exercise ownership rights and press for compliance by design, ensuring traceability, auditing, and rejection of fully autonomous weapons development. This applies established responsible-investment practice to the sector, on the logic that holding equity, rather than divesting, secures a seat at the table.
- Developers of military AI should be subject to traceability obligations, requiring documentation of design, data sources, and testing procedures sufficient to enable external audit independently from the deploying party (Taddeo et al., 2021).

Bibliography

- Amoroso, D., & Tamburrini, G. (2020). Autonomous weapons systems and meaningful human control: Ethical and legal issues. *Current Robotics Reports*, 1(4), 187–194. <https://doi.org/10.1007/s43154-020-00024-3>
- Blanchard, A., & Bruun, L. (2024). *Bias in Military Artificial Intelligence*. Stockholm International Peace Research Institute. <https://doi.org/10.55163/CJFT9557>
- Bovet, P. (2025). *Exploring Decision Advantages: Improving Speed, Precision and Efficiency in Military Targeting by Applying Artificial Intelligence* [Swedish Defence University]. <https://doi.org/10.62061/leiy7136>
- CENTCOM. (2026). [Operation Epic Fury fact sheet — OPERATION EPIC FURY-FIRST 24 HOURS.]. <https://media.defense.gov/2026/Mar/03/2003882611-1/-1/0/OPERATION-EPIC-FURY-FIRST-24-HOURS.PDF>
- Faris, L. (2026). Algorithmic targeting in the Iranian–Israeli confrontation: Technical realities, legal thresholds, and the boundaries of human control (Version 2). *F1000Research*, 14, Article 1200. <https://doi.org/10.12688/f1000research.169794.1>
- Güneysu, G. (2024). Lethal Autonomous Weapon Systems and Automation Bias. *Yuridika*, 39(3), 395–424. <https://doi.org/10.20473/ydk.v39i3.55188>
- King, A. (2024). Digital Targeting: Artificial Intelligence, Data, and Military Intelligence. *Journal of Global Security Studies*, 9(2), ogae009. <https://doi.org/10.1093/jogss/ogae009>
- Longpre, S., Storm, M., & Shah, R. (2022). Lethal autonomous weapons systems & artificial intelligence: Trends, challenges, and policies. *MIT Science Policy Review*, 3. <https://doi.org/10.38105/spr.360apm5typ>

- Mande, M., & Allen, G. C. (2026, June 2). What is the Maven Smart System, and what does it do? Center for Strategic and International Studies. <https://www.csis.org/analysis/what-maven-smart-system-and-what-does-it-do>
- Matei, S. A., & Reed, K. P. (2025, December 17). Mission command's asymmetric advantage through AI-driven data management. U.S. Army War College Press. <https://publications.armywarcollege.edu/News/Display/Article/4361748/mission-commands-asymmetric-advantage-through-ai-driven-data-management/>
- Milanovic, M. (2026). AI and the Commission and Facilitation of International Crimes: On Accountability Gaps and the Minab School Strike. *EJIL: Talk! (Blog of the European Journal of International Law)*.
- Nuredin, A. (2025). The Judgment of Algorithms: The Transformative and Discriminatory Impact of AI-Based Decision-Making Systems on Judicial Processes, Cybersecurity, and Individual Liberties. [Congress proceedings]. Faculty of Law, International Vision University.
- Probasco, E. S. (2024). Building the Tech Coalition: How Project Maven and the U.S. 18th Airborne Corps Operationalized Software and Artificial Intelligence for the Department of Defense. Center for Security and Emerging Technology. <https://cset.georgetown.edu/publication/building-the-tech-coalition/>
- Probasco, E. S., Toner, H., Burtell, M., & Rudner, T. G. J. (2025, April). AI for military decision-making: Harnessing the advantages and avoiding the risks (Issue Brief). Center for Security and Emerging Technology. <https://cset.georgetown.edu/wp-content/uploads/CSET-AI-for-Military-Decision-Making.pdf>
- Russo, A. (2026). Does AI Affect the Democratic Conduct of War? Analyzing US and Israeli Military AI Deployment. *Global Policy*. <https://doi.org/10.1111/1758-5899.70117>
- Simmons-Edler, R., Dong, J., Lushenko, P., Rajan, K., & Badman, R. P. (2025). *Military AI Needs Technically-Informed Regulation to Safeguard AI Research and its Applications*. [NeurIPS].
- Taddeo, M., McNeish, D., Blanchard, A., & Edgar, E. (2021). Ethical principles for artificial intelligence in national defence. *Philosophy & Technology*, 34, 1707–1729. <https://doi.org/10.1007/s13347-021-00482-3>
- United Nations Institute for Disarmament Research [UNIDIR]. (2025). Artificial intelligence in the military domain and its implications for international peace and security: An evidence-based road map for future policy action. <https://unidir.org/publication/artificial-intelligence-in-the-military-domain-and-its-implications-for-international-peace-and-security-an-evidence-based-road-map-for-future-policy-action/>
- Yatsyk, O., & Zharova, L. (2024). Applications of AI in Military Decision-Making and Strategic Planning. *Business, Economics, Sustainability, Leadership and Innovation*, (12), 71–78. <https://doi.org/10.37659/2663-5070-2024-12-71-78>



Youth Policy Center