

AI, INTERNATIONAL LAW, AND THE MINAB TEST

The war in Iran highlights the U.S. and Israeli militaries' reliance on military artificial intelligence. Yet the bombing of a school in Minab, which killed 170 civilians, did not set off any alarms.



Portraits of victims reportedly killed in a U.S.-Israeli airstrike on a residential building in Tehran on April 13, 2026. Photo: © AFP

By Louis Perez (for The Conversation France), 27 April 2026

The armed conflict against Iran launched on February 28 by Washington and Tel Aviv was quickly dubbed the “first AI war”. This assertion is, in fact, misleading in several respects. Not only has AI already been used extensively in recent conflicts, notably by Israel in Gaza, but more broadly AI, as a digital means of data processing and analysis, has a long history with armed conflicts, with technical foundations dating back to World War II.

Admittedly, the situation in Iran is distinguished by the unprecedented level of sophistication of these systems and by the military's unprecedented reliance on them. It also differs from the conflict in Gaza in that, this time, AI has been deployed against a state in the context of a high-intensity war. And never before have states communicated so openly about their use of these systems. It is this and the dramatic consequences of certain strikes that raise questions about the compatibility of these practices with international law.

USE OF AI IN THE IRAN WAR

Israel's use of AI in its war against Hamas was revealed by the newspaper +972. This outlet told what many experts had suspected for several years. In the context of the conflict in Iran, however, it was the US authorities themselves who announced their use of AI.

Indeed, US military forces admitted to using AI systems to compile and sort a list of targets at lightning speed. This process reportedly led to more than 1,000 strikes, described as highly precise, during the first 24 hours of the conflict. They reportedly used the Maven Smart System, a joint project utilizing Palantir's AI surveillance and data collection software, coupled with the generative AI system Claude, developed by Anthropic.

However, on the first day of the war, one of the US strikes targeted a school in Minab, killing about 170 civilians, mostly children. The United States acknowledged its responsibility for this strike, which it presented as a mistake. The school was indeed located near a Revolutionary Guards naval base. It had previously been an integral part of the same complex before being separated from it. It was therefore outdated information that reportedly led to authorization of the strike.

Such a mistake is not trivial. Many media outlets and NGOs quickly established the link between the school and the naval base. It was thus argued that the US military likely targeted this building based on outdated data by blindly following a recommendation from an AI system without conducting the necessary verification.

AI AND THE LAW OF ARMED CONFLICTS

To what extent are the use of AI to carry out these strikes and the error committed lawful under international law?

It should first be noted that AI is not prohibited per se by the law of armed conflict (International Humanitarian Law, IHL). At present, no legal rule specifically addresses the question of its lawfulness. Nevertheless, the issue does not exist in a legal vacuum. The general rules of IHL apply to the conduct of hostilities, regardless of the means and methods employed.

One of these rules is the principle of distinction, according to which only military targets may be attacked, while civilians and civilian objects must be spared. Directly targeting a school such as the one in Minab, in the absence of any military objective within it, therefore constitutes a clear violation of this principle. It is, however, unlikely that the US military had the deliberate intention of destroying the

school as such. As noted, this is more likely a case of mistaken target identification, possibly linked to an AI system trained on outdated data from the time when the building was still attached to the naval base.

Consequently, the violation relates more to the principle of precaution. This principle stipulates, in particular, that parties to the conflict must do everything practicable to verify that the targets to be attacked are indeed military targets. In this case, the US military does not appear to have carried out the necessary verifications to establish that the target was a school. A basic verification, such as that conducted by some media outlets, could have quickly dispelled any doubt.

During the war in Gaza, it was reported that Israeli soldiers sometimes had only 20 seconds to validate a target, which raises questions about the practical possibility of effectively adhering to this principle. Concerns regarding military AI often centre on the issue of autonomy and the risk that a system might identify and engage a target on its own; this is the crux of the issue with lethal autonomous weapons systems. This example still demonstrates that formally maintained human control may be merely nominal if the operator lacks both the time and the critical judgment necessary to evaluate an algorithmic recommendation.

On the Iranian side, it should be noted that the precautionary principle was not respected either. This principle not only imposes obligations on the attacker but also requires the attacked party to take certain passive precautions: in particular, the parties must keep civilians and civilian objects away from military targets. In this case, converting a building at a naval base into a school while keeping it in the immediate vicinity of the rest of the military complex deliberately exposed this civilian facility to the risks associated with the conduct of hostilities.

WHAT ARE THE LEGAL, POLITICAL, AND MORAL RESPONSIBILITIES?

Individual responsibility: such an attack does not constitute a war crime.

While the attack constitutes a violation of IHL, it is unlikely that any US military personnel will be convicted for such acts. Beyond issues of jurisdictional competence, the main obstacle is that neither the violation of the principle of precaution nor the errors leading to violations of IHL constitute war crimes under international criminal law.

The material act is clearly established, but the element of intent – that is, the will to commit the offence – is lacking. The current international criminal liability regime does not recognize liability for negligence in this context. This pragmatic approach could nevertheless evolve. On the one hand, if algorithmic targeting errors become more frequent, it will become increasingly difficult to argue that the error was “reasonable”, and the deliberate use of a system known to have flaws could imply a form of indirect intent to target civilians. On the other hand, the law could evolve in the future to punish military personnel who, through their negligence, cause the death of civilians.

The liability of Ai companies: a standoff between economic and political powers.

Another point of concern relates to the role of private companies specializing in AI, which today hold the lion's share of the technological expertise deployed on the battlefield. These companies could be held liable when they develop faulty systems, but beyond this liability, a fundamental moral and political question arises regarding the sale of AI technologies for military purposes.

Just before the United States entered the war Anthropic, the developer of the Claude system, opposed unrestricted cooperation with the Pentagon, particularly regarding autonomous weapons, citing its ethical commitments and the technical reliability limitations of its systems for the intended uses. The Pentagon then accused Anthropic of treason, although its systems continue to be used by the military.

Other companies in the sector, such as OpenAI, Google, Amazon, and Microsoft, appear to be collaborating unreservedly with the military, establishing themselves de facto as real defence contractors. It is interesting to note that companies, normally driven by profit, sometimes have more qualms about this issue than certain states, which are nevertheless supposed to safeguard the public interest.

State responsibility: accountability and preventing future violations.

States that develop and use military AI bear a special responsibility. In this case, the United States incurs international responsibility for committing an internationally wrongful act. This responsibility will, admittedly, be difficult to hold to account in practice. But beyond that, a responsibility emerges that is both legal and political. Under Article 1 common to the Geneva Conventions, states are indeed obligated to respect and ensure respect for IHL. However, the development of military AI tends to undermine this respect, and may even encourage and conceal violations of the law.

Various mechanisms could curb this phenomenon, such as training military personnel on the specificities of AI systems, developing rules of engagement tailored to AI, technical safeguards ensuring the reliability and transparency of systems, as well as regular testing and evaluation. Several international initiatives call for integrating such measures into new legal instruments. Yet political will is lacking, particularly among states at the forefront of the development and use of military AI.

US Secretary of Defence Pete Hegseth appears to be acting in the opposite direction. He recently dismissed military legal advisors whom he considered obstacles to the proper conduct of hostilities and described the rules of engagement as stupid. More broadly, the United States opposes any international legal regulation of military AI. AI thus appears to be both a driving force and a harbinger of a profound erosion of IHL.

Jacques Lacan said: "The real is when you hit yourself." The Minab accident is a dramatic event that confirms the risks military AI experts have been warning about for several years and that should have elicited far more reaction.

In reality, this seems to have been overshadowed by other considerations perceived as more urgent and more visible in the context of this war, starting with the nuclear risk. The Minab accident did not

serve as the wake-up call expected to prompt states to agree on a specific legal framework applicable to military AI. It remains to be seen whether such a wake-up call is still possible or even desirable.

This article, slightly modified and translated by Justice Info, is republished from [The Conversation France](#) under a Creative Commons licence. Read the [original article](#).



LOUIS PEREZ

Louis Perez holds a doctorate in public law and is a postdoctoral researcher at the Centre Thucydides at Paris Panthéon-Assas University. He is the author of a thesis on “International Legal Regulation of Military Artificial Intelligence”.