

Informal Exchanges on Artificial Intelligence in the Military Domain and its Implications for International Peace and Security

15 – 17 June 2026 Geneva, Tempus

Session 5: "The Life Cycle of AI in the Military Domain (Part 2)"

16 June, 10:00-11:30am

Notes: Prof. Ingvild Bode, Center for War Studies, SDU

- I will speak about two issues in my brief intervention: (1) first, I will draw attention to an under-recognized dimension of human judgment and control – namely human agency; and (2) I will demonstrate how thinking about this shapes concrete implementation pathways at the lifecycle stages of deployment and use

1. Human agency & human-machine interaction [slide 1]

- Let me speak about an under-recognized dimension of human judgment and control – namely human agency, which I understand as the intentional capacity to take actions and affect outcomes.
- Exploring human agency recognizes that there is a two-sided relationship between humans and AI systems – especially at the deployment and use stage of the lifecycle.
- Human agency has important implications for the ability to exercise control of AI systems, particularly with respect to relevant legal and ethical frameworks.
- **In the international debate about AI in the military domain, we have seen the emergence of a requirement for the sufficient capacity to exercise human agency as a possible condition for compliance with international humanitarian law.**
- The GGE on LAWS rolling text, for example, stipulates specific measures to implement human judgment and control as well as to promote human responsibility and accountability in the use of LAWS (GGE on LAWS 2026, para. 20).
- While this is not explicitly mentioned, these stipulations appear to be designed to ensure that human judgment and control, especially in the context of deployment and use, **does not turn nominal.**
- The notion of nominal human control describes a situation where humans appear to be exercising important legal, safety, and ethical functions in the deployment and use of AI systems but actually lack the time, information, or cognitive capacity to perform these functions.
- Here, humans engage in cognitive off-loading and over-rely on AI for complex tasks. This process is linked to “automation bias”, the over-reliance on the AI’s output and a failure to catch its mistakes.
- The simple presence of humans at the moment of deployment and use is therefore not sufficient to guarantee human judgment and control.
- This is perhaps best illustrated through considering the case of AI decision support systems. In the case of AI DSS, human operators have the final say over the verification and validation of targets.
- But this simple presence does not guarantee the exercise of human judgment and control – it does not consider how the human agency by those humans-in-the-loop is exercised and affected by various dynamics of human-machine interaction.

- Building on this, I argue that the risk of nominal human control can be countered if we integrate human agency into our thinking about governing AI systems in the military domain.

2. Implementation pathways [slide 2]

- In the second part of my intervention, I want to raise some concrete implementation pathways arising out of thinking about how the exercise of human agency conditions human control and judgment.
- The two implementation pathways that I sketch come from an international research project that I led, the European Research Council funded AutoPractices project, where we worked with a diverse set of stakeholders to co-produce a best practices toolkit around strengthening human agency when using AI systems in the military domain.¹
- **(1) Integrate feedback, communication and contestation mechanisms that connect human actors at the deployment/use stage with those at earlier lifecycle stages.**
- Ongoing feedback loops and communication channels among all participants at various life-cycle stages help people understand how AI systems work and how these are used in specific contexts.
- This improves transparency and strengthens human accountability.
- Human actors at the deployment and use stage should have the opportunity to cross-check the outputs of AI systems with other sources of data and intelligence.
- This cross-checking exercise helps humans to critically assess and potentially challenge AI outputs. It also safeguards against deskilling and cognitive erosion.
- These contestation mechanisms would raise human user awareness of known and potential limitations and aid them in understanding when the system is not working as planned.
- The impact of this practice ultimately rests on maintaining an organizational culture that is open to feedback and exchange.
- **(2) Include consistent documentation to transmit knowledge, concerns, and risks.**
- This should include establishing documentation channels to transmit knowledge, concerns or risks across the chain of command.
- Maintaining such consistent records and documentation is vital to ensure transparency of how responsible individuals or teams use AI systems for certain functions.
- What both of these pathways highlight is also that the deployment and use stage cannot be thought of in isolation from the earlier lifecycle stages. The two pathways I sketched require throughlines of documentation and feedback loops.

¹ The AutoPractices Project, *Strengthening Human Agency in the Military Domain: Best Practices Toolkit for Policymakers, Developers, and Users of AI Systems* (Center for War Studies, University of Southern Denmark, 2026), <https://doi.org/10.5281/zenodo.17671715>.