

The Impact of AI on Decision-Making in High-Stakes Environments

Giuseppe Galetta¹, Stefania De Simone², & Massimo Franco²

¹ Central Administration, Human Resources Area, University of Naples Federico II, Naples, Italy

² Department of Political Science, University of Naples Federico II, Naples, Italy

Correspondence: Giuseppe Galetta, Central Administration, Human Resources Area, University of Naples Federico II, Naples, Italy. E-mail: giuseppe.galetta@unina.it

Received: March 30, 2026

Accepted: April 23, 2026

Online Published: June 6, 2026

doi:10.5539/ijbm.v21n4p48

URL: <https://doi.org/10.5539/ijbm.v21n4p48>

Abstract

This study examines how artificial intelligence (AI) is transforming decision-making processes in organizations operating in high-stakes environments. Specifically, it analyzes the impact of AI-based Decision Support Systems (AI-DSSs) on governance, accountability, and performance. Military cases are used as analogies to highlight the acceleration of decision-making cycles and the challenges of maintaining meaningful human control over AI decisions. The research is based on a thematic analysis of 388 academic, military, and institutional sources, processed by *NVivo* software and analyzed through Grounded Theory methodology. Observational data from ongoing conflicts are used as an empirical context and complemented by SWOT analysis to assess risks and benefits. The analysis identifies three key constructs: decision acceleration, artificial agency, and meaningful human control. Findings show that AI acts as a ‘decision amplifier,’ increasing data processing speed, while introducing risks such as excessive dependence, loss of human control, bias, hallucinations, and accountability gaps. This outlines a bounded rationality decision acceleration paradigm, where speed tends to prevail over precision: AI-DSSs increase organizational information processing capacity and compress the OODA loop, shifting human decision-making from direct deliberation toward validation of machine-generated options. The study highlights the urgent need to preserve human control in critical decision-making processes, both in military and business contexts. It proposes a human-machine cooperative model capable of balancing efficiency with human oversight, avoiding full delegation to AI. Finally, the research emphasizes the importance of developing appropriate governance and compliance mechanisms to regulate AI use in high-stakes environments, ensuring accountability, respect for human values, and the prevention of conflict escalation.

Keywords: AI governance, artificial agency, decision acceleration, high-stakes decision-making, OODA loop, meaningful human control

1. Introduction

The adoption of artificial intelligence in firms is rapidly transforming how strategic and operational decisions are made under uncertainty. For managers, this transformation raises fundamental questions: How does AI affect the distribution of decision rights within organizations? What governance mechanisms ensure accountability when algorithms accelerate or even anticipate decision outcomes? How can firms preserve meaningful human control while exploiting AI-driven speed advantages?

These issues resonate strongly with the management literature on information processing, governance, and human-machine collaboration: the study situates its contribution within the historical waves of AI adoption in management research and practice (Cristofaro & Giardino, 2025), showing why the current high-velocity human-AI teaming moment requires fresh theorization about decision acceleration and oversight. The historical arc of AI in management highlights a cyclical pattern, and the contemporary moment emphasizes the need for responsible AI frameworks and integrated approaches that balance theoretical rigor with practical relevance.

To investigate these challenges, this paper draws on insights from military domain, where decision acceleration has long been critical to survival. In fact, the weaponization of artificial intelligence is rapidly transforming the logic of war, bringing profound changes to military decision-making and leadership (Baboş, 2021; Farnell & Coffey, 2024; Michel, 2024).

In this ever-evolving scenario, military doctrines – such as Mission Command or the OODA loop – offer rich

illustrations of how AI alters the tempo and structure of high-stakes decisions. Yet rather than focusing on military policy per se, our analysis translates these lessons into organizational practice, showing how firms can adapt governance, escalation protocols, and accountability structures in response to AI adoption. In doing so, we build a conceptual framework around three key constructs—decision acceleration, artificial agency, and meaningful human control—specifying their managerial relevance in terms of observable processes and outcomes.

An important clarification: in this paper, the term ‘Mission Command’ refers specifically to the doctrine of the U.S. Army (FM 6-0) (2022). In fact, the U.S. Marine Corps emphasizes *Maneuver Warfare*, articulated by Lind (1985), which shares similarities but differs in practice. We acknowledge these doctrinal differences explicitly to ensure precision in translation to organizational analogies.

Until a few years ago, the decentralized military command-and-control model (i.e. Mission Command) was supported by human skills, qualities, judgment, and values such as leadership, troop motivation, mutual trust between commanders and soldiers, clear understanding of orders, threat assessment, timeliness of decisions, instinct, and intuition.

Today, artificial intelligence poses new challenges to the traditional Military Decision-Making Process (MDMP) and chain of command leadership by enhancing human decision-making through algorithms and predictive models able to control the situation without human intervention. These systems enable a rapid understanding of the scenario and the formulation of fast and effective decisions (situational awareness).

This work analyzes artificial intelligence applications in ongoing conflicts, examining how AI can be tactically and strategically used in military planning and decision-making. Furthermore, we highlight how human decision-making processes can become more accurate, efficient, and targeted, while maintaining meaningful human control over AI-based Decision Support Systems (AI-DSSs). Additionally, we will address ways to reduce the risks of biases and hallucinations, which remain dangerous sources of uncertainty when using AI-driven decision-making models (Probasco et al. 2025).

Finally, after analyzing the risks and opportunities associated with AI weaponization, we suggest what seems to be the best possible path forward: fostering human-machine teaming. While it is still too early to speak of replacing human command with AI agents on the battlefield, the rapid development of artificial intelligence, coupled with the growing use of lethal autonomous weapon systems (LAWS), and concerns raised by the scientific community (including AI developers), highlights troubling prospects. All this leads to an essential question for organizational leaders: could a firm ever rely on the command of a machine?

In this paper we propose an alternative to the risks posed by the advent of technological singularity in the military domain (Kurzweil, 2022). In fact, in the context of the evolutionary analysis of military command models, we describe the transition to a different AI-assisted command logic. This logic requires human-machine collaborative practices and governance frameworks of cyborg ethics, ensuring meaningful human control over AI and clear accountability in military operations, incorporating military rules and doctrines into AI reasoning processes to create a truly adaptive automated command-and-control system, aiming to reduce the risk of AI instability in military decision-making cycle in war scenarios, where it is vital not to make mistakes.

Therefore, we hope for the development of a human-machine relationship model in the military, that does not subordinate humans to AI control but offers decision makers a cooperative framework, where the ‘human factor’ does not completely disappear behind artificial decisions made by machines (Goldfarb & Lindsay, 2022).

The theoretical logic of the paper is therefore not that AI simply makes decisions faster. Rather, AI changes the architecture of boundedly rational decision-making. Under conditions of uncertainty, time pressure, and information overload, human decision makers cannot evaluate all possible courses of action exhaustively. AI-DSSs extend the organization’s information-processing capacity by filtering signals, generating options, and ranking courses of action. At the same time, this extension alters the OODA loop: observation becomes data fusion, orientation becomes algorithmic sensemaking, decision becomes human validation of machine-generated alternatives, and action becomes a feedback-rich process that continuously updates the next cycle. This integrated view allows the paper to connect military decision tempo with general management questions of governance, accountability, and responsible delegation, which can also be applied to the business and managerial context.

2. Positioning the Contribution into the ‘Waves’ of Managerial AI

As previously mentioned, this study follows the historical waves of AI adoption by firms, as outlined by M. Cristofaro and P. L. Giardino (2025). The ‘wave’ metaphor, as specified by the authors, “*captures the dynamic*

and cyclical evolution of AI's integration in theory and practice, illustrating periods of rapid advancement followed by phases of stagnation and reassessment" (Cristofaro & Giardino, 2025).

Based on this line of management and organizational studies, our research is situated at the intersection between the 'fourth wave' of AI adoption—characterized by the transformation driven by big data and deep learning—and the emerging 'fifth wave,' which emphasizes responsible collaboration between humans and machines.

While previous waves focused on the automation of structured tasks and routine decisions, the current era, driven by data and machine learning, introduces ultra-high-speed decision-making cycles and an unprecedented perception of artificial agency within complex organizations, in high-risk and non-routine organizational contexts, such as the military.

Our contribution extends the managerial debate beyond mere information processing efficiency, focusing on new governance challenges, focusing on the high-risk end: the decision acceleration paradigm and the erosion of meaningful human control. By analyzing how AI is reshaping the chain of command in high-stakes contexts, the article answers the fifth wave's call for a new theory of AI governance and decision-making ethics, ensuring that the push for technological efficiency does not compromise human accountability and autonomy within organizations.

This positioning clarifies the article's theoretical contribution. Descriptively, military cases demonstrate that AI is already being used to accelerate data collection, target assessment, forecasting, and operational coordination, especially in high-stakes situations. Theoretically, it is argued that these practices constitute a still weak and imperfect organizational mechanism, which could prove risky: AI-DSSs relieve decision makers of the cognitive load of information processing and option generation, determining greater speed and effectiveness in the decision-making process, but significantly reducing the oversight and meaningful human control, exception management, and accountability. The concept of the decision-acceleration paradigm thus links the OODA loop to bounded rationality (Simon, 1947; 1982; 1991): speed is greater because AI reduces the search and processing load, but the quality of judgment and the ethicality of decisions can diminish if human actors are excluded from the orientation and decision-making phases.

3. Purpose and Method

This study adopts an *NVivo*-assisted thematic analysis combined with Grounded Theory procedures. The purpose is to explain how AI-DSSs affect high-stakes decision-making by increasing decision speed, expanding artificial agency, and challenging meaningful human control.

Since in the military field there is not yet a defined theoretical framework capable of embracing such a changeable object of research, characterized by unpredictable developments, and very difficult to classify due to the technological advances speed (especially when applied to the use of lethal force), we chose to base the study on the principles of Grounded Theory (GT).

In fact, at the present stage, we believe it is more realistic to rely on empirical data that emerge daily from the battlefield, where commanders operate 'boots on the ground': through military bulletins, journalistic investigations, research center reports, and OSINT sources, we can know when AI-driven decisions are implemented by the chains of command. This qualitative data can be more effectively compared through Computer Assisted Qualitative Data Analysis Software (CAQDAS), such as *NVivo*.

3.1 Research Questions

The analysis addresses three research questions: (RQ1) How do AI-DSSs reshape the OODA-based decision cycle in high-stakes environments? (RQ2) How AI, by expanding information processing capacity, increases the effects of bounded rationality while creating new risks of over-reliance and loss of judgment? (RQ3) What governance mechanisms are required to preserve meaningful human control when AI-DSSs move from decision support toward autonomous or semi-autonomous agency?

3.2 Search Protocol and Corpus Construction

The corpus was built through a structured search conducted across academic, military, and institutional sources. The academic databases included ScienceDirect, ProQuest, IEEE Xplore, ACM Digital Library, EBSCO, and ArXiv. Military sources included DTIC (*Defense Technical Information Center*), USMA Libraries (*U.S. Military Academies Libraries*), AUL LibGuides (*Air University Library Guides*), AULS (*Army University Library System*). The search was limited to English-language documents published between 2000 and 2024, with later sources included only where they were already part of the accepted reference set and directly relevant to the

emerging debate on military AI and human-machine teaming.

Furthermore, through a parallel investigation conducted on major government and institutional online platforms, we integrated the bibliographical corpus with some of recently declassified government documents, such as U.S. Department of Defense memoranda and reports, as well as other reliable sources in policy domain (e.g., reports by authoritative international think tanks like RAND Corporation, CNAS, RUSI, CSET, HCSS, Finabel, UNIDIR, or ICRC).

This was necessary because open search is often constrained by data limitations due to secrecy concerns. However, it is appropriate to underline that, without privileged access or security clearances, it is difficult to determine the precise role of AI-DSSs in use-of-force situations, as many studies are still classified or covered by military secrecy, and funded by the military-industrial complex through 'black budgets'. Despite these challenges, valuable insights have been obtained, allowing us to draw up an overall framework for understanding the role of AI-DSSs in military domain.

The search strings combined three groups of terms using Boolean operators.

- The first group captured keywords related to military decision-making, such as: military decision-making, command and control, mission command, network command, network-centric warfare, OODA loop.
- The second group included keywords related to AI-based technology in the military: military AI, AI-DSS, decision support systems, autonomous systems, lethal autonomous weapon systems.
- The third group of keywords concerned governance and human factors: meaningful human control, human-machine teaming, situational awareness, algorithmic bias, trust in automation, responsible AI, accountability, AI governance.

Searches were run using combinations, such as: ('AI' OR 'artificial intelligence' OR 'AI-DSS') AND ('military decision-making' OR 'OODA loop' OR 'command and control') AND ('meaningful human control' OR 'human-machine teaming' OR 'algorithmic bias').

3.3 Inclusion and Exclusion Criteria

Documents were included when they met at least one of the following criteria: (a) they addressed AI or automation in organizational or military decision-making; (b) they discussed command-and-control, OODA, Mission Command, Network-Centric Warfare, or JADC2 in relation to decision tempo; (c) they analyzed human-machine teaming, meaningful human control, trust, accountability, or bias in high-stakes AI use; or (d) they provided official, policy, or technical evidence on AI-DSSs in military or comparable high-risk contexts.

Documents were excluded when they were not in English, lacked accessible full text, dealt with purely technical AI performance without decision-making implications, focused on low-stakes consumer applications, duplicated already included records, or consisted of opinion pieces without an identifiable analytical or empirical basis. After de-duplication and screening of titles, abstracts, executive summaries, and full texts, 388 sources were retained for thematic synthesis.

3.4 NVivo Coding Procedure

The 388 documents were imported into *NVivo* as PDF files and coded through three iterative stages.

- First stage: open coding identified text segments referring to decision speed, information overload, OODA stages, bounded rationality, trust, reliability, bias, human control, artificial agency, and accountability.
- Second stage: axial coding grouped related open codes into broader categories, comparing military and organizational sources through constant comparison.
- Third stage: selective coding organized the categories around the core concept of decision acceleration, linking it to artificial agency and meaningful human control.

Frequency queries were used to identify the relative salience of themes, but theme frequency was not treated as proof of causal importance. Instead, frequency counts were used as a descriptive indicator to guide deeper interpretive analysis.

A codebook was maintained throughout the process. For each node, the codebook specified: node name, operational definition, inclusion rules, exclusion rules, example passages, and links to theoretical anchors.

Table 1 reports the thematic nodes and operational definitions, while Table 2 links the theoretical constructs to the corresponding bodies of literature and *NVivo* traceability codes. Analytical memos were created after each coding round to document why codes were merged, split, renamed, or elevated into theoretical constructs.

3.5 Inter-coder Reliability and Credibility Checks

To improve credibility, coding was conducted through an iterative inter-coder process. A pilot subset of documents was independently coded by the research team, after which coding discrepancies were discussed and the codebook was refined. A second coding round applied the revised codebook to the full corpus.

Reliability was assessed through agreement checks on the presence or absence of core thematic nodes and through consensus meetings for ambiguous passages. Because the corpus includes heterogeneous academic, policy, military, and institutional sources, the reliability procedure was used to strengthen interpretive consistency rather than to reduce the analysis to a purely quantitative content analysis. Given the heterogeneity of the corpus, reliability was operationalized as structured consensus coding supported by a shared codebook, agreement checks on core thematic nodes, and analytical memos, rather than as a standalone quantitative content analysis.

Table 1. NVivo-assisted thematic synthesis across 388 sources

Key theme identified by NVivo	Operational definition (for method replicability)	Analytic role	Frequency (%)
Decision acceleration (AI-driven MDMP)	References to AI-enabled reduction of time required for observation, orientation, decision, or action	Core category / theoretical construct	18.4
Situational awareness	References to the AI-supported perception, integration, and interpretation of battlefield or organizational information	Descriptive finding	15.2
Risks and benefits	Balanced references to operational advantages and vulnerabilities of AI-DSSs	Descriptive finding	12.8
Ethical and legal issues	References to IHL, accountability, responsibility, compliance, proportionality, or liability	Governance context	11.0
Evolution of AI-driven command models	References to Mission Command, Network Command, NCW, JADC2, OODA/DOODA, or command-and-control transformation	Contextual mechanism	10.5
Artificial agency/AI autonomy	References to autonomous or semi-autonomous option generation, prioritization, selection, or execution	Theoretical construct	9.7
AI-DSS reliability	References to robustness, explainability, validation, adversarial attacks, hallucinations, or error propagation	Boundary condition	8.6
Meaningful human control (MHC)	References to human oversight, intervention, override authority, final responsibility, and auditability	Governance moderator	7.4
Trust in AI-DSS	References to calibrated trust, over-reliance, halo effect, automation bias, or responsible reliance	Boundary condition	6.4

Source(s). Own elaboration.

Table 2. Theoretical constructs, main sources, role in the model, and traceability

Theoretical construct	Main reference sources	Role in the conceptual model	Traceability codes by NVivo
OODA-based decision cycle	Boyd (1995); Coram (2002); Lind (1985); Johnson (2022a)	Explains the temporal sequence through which AI compresses observation, orientation, decision, and action	Nodes A07, A12, A35
Bounded rationality	Simon (1947; 1982; 1991); March & Simon (1958); Tushman & Nadler (1978)	Explains why AI-DSSs are adopted to compensate for cognitive limits, information overload, and time pressure	Nodes A22, A35, A71
Decision Acceleration Paradigm	Tushman & Nadler (1978); Bennis & Nanus (1985); U.S. Department of Defense (2023)	Core category: AI increases information-processing capacity and makes speed a central source of advantage	Nodes A12, A35
Artificial Agency	Bandura (1997); Gugerty (2006); Johnson (2022c)	Explains the delegation of option generation, prioritization, or action-selection functions to AI-DSSs	Nodes A57, A89
Meaningful Human Control (MHC)	RAND (2024a, 2024b); Saidi (2022); Davis, 2024; human-AI teaming literature	Governance moderator that preserves human scrutiny, override authority, and accountability	Nodes A112, A198
Trust and reliability	Mayer (2023); Probasco et al. (2025); RAND (2024b)	Boundary condition affecting whether AI recommendations are scrutinized, contested, or passively accepted	Nodes A144, A203

Source(s). Own elaboration.

3.6 Role of Observational Material and SWOT Analysis

Observational material from ongoing conflicts is used as contextual evidence to illustrate the relevance of AI-DSSs in high-stakes environments. It is not used to infer direct causal effects or to generalize across all military or corporate contexts.

SWOT analysis (Table 3) is therefore positioned as a post-synthesis support: it organizes strengths, weaknesses, opportunities, and threats emerging from literature, identifying risks and benefits of implementing AI-DSSs in high-stakes environments such as in the military, but does not constitute a separate Grounded Theory coding stage, namely it is not presented as part of the GT-engine but as a heuristic to visualize managerial implications. This distinction addresses the risk of overgeneralization by separating descriptive observations from the paper's theoretical contribution.

Table 3. SWOT analysis with source-based justification and interpretative status

Dimension	Output with reference sources	Interpretative status
Strengths	Decision acceleration (Scharre, 2018); improved situational awareness (Allen, 2017); predictive capabilities (Velazquez, 2024); reduction of selected human errors (Horowitz, 2019)	Descriptive synthesis of reported benefits; not evidence that AI improves all decisions
Weaknesses	Loss of human oversight (Bode et al., 2025); ethical and legal concerns (Gunawan et al., 2022); bias and hallucination risks (RAND, 2024b); dependence on technology (Bode, 2025)	Boundary conditions limiting the decision acceleration paradigm
Opportunities	Innovation in military strategy (Clark et al., 2020; Vold, 2025); human-machine teaming (Daugherty & Wilson, 2018); cost efficiency (Defense Science Board, 2016); enhanced training (Hussen et al., 2025)	Managerial implications requiring governance design before adoption
Threats	Escalation risks of conflicts (Kissinger et al., 2021); cybersecurity vulnerabilities (Girei et al., 2025); ethical and legal dilemmas (Rowe, 2022); regulatory challenges (Meerveld et al., 2025)	Risks that constrain generalization and require meaningful human control

Source(s). Own elaboration.

4. Theoretical Foundations and Conceptual Model

The theoretical framework integrates three complementary key constructs: (1) the OODA loop; (2) bounded rationality; (3) information-processing theory. Each construct explains a different part of AI's impact on high-stakes decision-making. The OODA loop (Observe-Orient-Decide-Act) explains the temporal structure of decisions under competition and uncertainty (Boyd, 1995; Johnson, 2022a).

Bounded rationality explains why human decision makers cannot fully process the volume, speed, and ambiguity of information in such environments (Simon, 1947; 1982; 1991). Information-processing theory (IPT) explains why organizations adopt AI-DSSs to increase processing capacity and reduce uncertainty (Tushman & Nadler, 1978). Together, these theories support a cohesive model in which AI-DSSs compress decision cycles while redistributing agency and accountability between humans and machines.

4.1 The OODA Loop as the Temporal Structure of AI-driven Decisions

The OODA loop provides the process architecture of high-stakes decisions. In the traditional model, human commanders or managers observe the environment, orient themselves by interpreting signals, decide among available courses of action, and act. AI-DSSs transform each stage.

In observation, they aggregate sensor feeds, reports, open-source intelligence, and operational data. In orientation, they classify patterns, generate predictions, and produce algorithmic sensemaking. In decision, they recommend, rank, or automatically select courses of action. In action, they create feedback loops that update the next cycle. The key theoretical implication is that AI does not merely add speed to a pre-existing human process; it reconfigures the process by changing who or what performs each OODA function.

4.2 Bounded Rationality as the Human Constraint Intensified by AI

Bounded rationality explains why AI-DSSs become attractive in high-stakes environments. Human decision makers face cognitive limits, incomplete information, time pressure, and uncertainty; as a result, they satisfy rather than optimize. AI-DSSs expand the feasible search space by processing data at machine speed and by generating options that would be unavailable to unaided human cognition. However, the same mechanism can intensify bounded rationality in a new form.

When AI outputs become too complex, too fast, or too opaque to scrutinize, human decision makers may no longer deliberate over the underlying problem but instead deliberate over whether to trust the machine. The locus of bounded rationality shifts from the battlefield or organizational (or market) environment to the human-machine interface.

4.3 Information-processing Theory and the Decision Acceleration Paradigm

Information-processing theory links the previous two elements. Organizations perform effectively when their information-processing capacity matches environmental uncertainty. AI-DSSs increase this capacity, but they also create a speed-based governance challenge.

When the organization's competitive or operational advantage depends on completing the OODA loop faster than an adversary or competitor, speed may be prioritized over deliberative precision. We define this condition as the decision acceleration paradigm: an organizational state in which AI-enabled processing speed becomes a central source of advantage while simultaneously increasing the need for human oversight, accountability, and control.

The military context, characterized by extreme volatility, uncertainty, complexity, and ambiguity (VUCA), represents an environment with massive information needs (Bennis & Nanus, 1985; Barber, 1992; Sehgal, 2024). AI-DSSs become essential mechanisms to increase processing capacity, mirroring the strategic necessity for leadership agility in firms operating in complex global markets, which are rapidly evolving towards a new model of 'artificial agility'.

Military doctrines like the OODA loop illustrate this focus on tempo. In an organizational context, the goal of accelerating the OODA loop is to achieve the 'information dominance' by running the decision cycle faster than competitors, thereby securing a strategic advantage. The managerial relevance lies in understanding how this speed-first approach impacts the quality of human orientation and decision phases—the very stages where human judgment and ethical oversight are required (U.S. Department of Defense, 2023).

4.4 Artificial Agency

Artificial agency refers to the capacity of an AI system to exercise autonomous or semi-autonomous decision-making based on situational inputs, learned patterns, and pre-programmed objectives.

This concept, originally put into practice in 1956 by Allen Newell and Herbert A. Simon with the *Logic Theorist* program (Newell & Simon, 1956; Gugerty, 2006), is drawn from the psychological theory of human agency (Bandura, 1997), which describes the human capacity to act promptly and transformatively in response to real-world situations.

In a firm setting, artificial agency manifests when AI-DSSs are delegated authority for tasks that were previously human-mediated, such as automated trading, credit scoring, or real-time supply chain optimization.

While AI offers immense computational power, removing human agency raises critical questions about corporate accountability and the risk of algorithmic behaviors (e.g., hallucinations, or emergent untested actions) that diverge from organizational intent. For management scholars, understanding artificial agency is crucial for designing governance structures that allocate responsibility between human supervisors and autonomous systems.

Within the proposed conceptual model, artificial agency increases when AI-DSSs move from information retrieval to option generation, prioritization, target selection, or execution. This does not mean that the system possesses human judgment or moral responsibility. Rather, it means that practical control over parts of the decision cycle is delegated to technical systems.

The managerial and military problem is therefore not whether the machine is morally responsible, but how responsibility is allocated when the machine shapes the options humans perceive as available.

4.5 Meaningful Human Control (MHC) as a Governance Moderator

Meaningful human control (MHC) is the governance mechanism that prevents decision acceleration and artificial agency from becoming uncontrolled delegation: it represents the necessary framework for human oversight, intervention, and ultimate responsibility over AI-driven decisions, particularly those involving high stakes, ethical considerations, or legal liability. It is the counterpoint to artificial agency, ensuring that the delegation of decision rights remains within acceptable ethical and organizational bounds.

MHC is not just about the ability to hit a physical 'stop button' but involves embedding human judgment—concerning organizational values, ethics, and legal compliance—into the design and operational loop of the AI system (RAND, 2024a; 2024b).

The lack of MHC can lead to a state of AI over-reliance ('halo effect') where commanders or managers passively validate the machine's suggestion rather than actively scrutinizing the quality and rationale of the output. For management, establishing MHC is essential for implementing a Responsible AI framework, which ensures that

accountability remains anchored in the firm's leadership and management structure.

In the proposed conceptual model, MHC moderates the relationship between AI-DSSs and decision outcomes. It requires more than a formal human approval step. It requires sufficient time for scrutiny, access to explanations or confidence indicators, clear escalation thresholds, documented accountability, training against automation bias, and the authority to override machine recommendations. Without these conditions, the human actor may remain nominally present but substantively absent from the orientation and decision stages of the OODA loop.

Table 4. Conceptual model of AI-driven decision acceleration

Environmental conditions	Bounded-rationality mechanism	AI-DSS function	OODA effect	Governance implication
Uncertainty, time pressure, information overload	Human actors cannot exhaustively search, interpret, and compare all options	Data fusion, pattern detection, prediction, option generation	Observe and Orient are compressed; Decide shifts toward validation	Meaningful human control must preserve scrutiny, override authority, and accountability
Competitive/adversarial pressure	Need for fast satisficing under risk	Real-time recommendations and automated prioritization	Loop speed becomes a source of advantage	Risk of over-reliance, halo effect, and loss of deliberative precision
High ethical/legal stakes	Human judgment remains necessary for values and responsibility	Decision support rather than unrestricted autonomous command	Human-machine teaming balances speed and oversight	Escalation protocols, audit trails, responsible AI officers, and training

Source(s). Own elaboration.

4.6 Conceptual Model

Table 4 summarizes the conceptual model. Environmental uncertainty and information overload create bounded-rationality constraints. AI-DSSs respond to these constraints by expanding information-processing capacity. This expansion compresses the OODA loop and produces decision acceleration. As AI-DSSs take on option-generation or action-selection functions, artificial agency increases.

Meaningful human control moderates this process by keeping human judgment, values, and legal accountability embedded in the loop (Davis, 2024). The expected outcome is not full automation, but human-machine teaming: AI accelerates observation and analysis, while humans retain responsibility for orientation, validation, escalation, and final accountability in high-stakes decisions.

As illustrated, the model of decision acceleration links bounded rationality and OODA to AI-DSSs: AI expands processing capacity, compresses the decision cycle, increases artificial agency, and requires meaningful human control as a governance moderator.

5. Implementing AI into Military Decision-Making

5.1 The impact of AI on the Logic of Ongoing Conflicts

For over a decade, the integration of AI into military decision-making is emerging as a crucial issue in the current geopolitical and strategic landscape (Moisescu et al., 2010).

Technological advances in AI, alongside contemporary conflicts, are pushing governments and private companies of several countries to invest heavily in AI-based military technologies. The goal is to strengthen military capabilities and gain a competitive advantage over adversaries.

Some AI research organizations have established collaborations with defense-sector companies (e.g., *Palantir Technologies*, and *Anduril Industries*), prompting debate about mission orientation and ethical responsibilities.

However, the militarization of AI raises a series of complex ethical, legal, and political issues that require in-depth analysis by experts. Deploying these systems in warfare raises serious concerns about the loss of significant human control over decision-making, given the high risk of civilian casualties and collateral damage. AI could significantly lower the threshold for triggering a new global conflict, with destabilizing consequences for international security.

The rapid evolution of AI technology complicates the creation of adequate regulatory frameworks to regulate its use in military contexts, posing further challenges to policymakers (Goztepe, Dizdaroğlu & Sağiroğlu, 2015). According to authoritative scholars, facing these challenges requires promoting an informed and multidisciplinary public debate on the risks and benefits of AI militarization.

The international community, not only military organizations, should work together to develop clear ethical and legal standards that effectively regulate the use of AI for military purposes. This includes creating a shared regulatory framework, binding on individual states, to ensure that such technology is employed legally and ethically, in full compliance with humanitarian principles and international law.

Although the impact of AI on modern conflicts is still in its infancy and the long-term implications remain uncertain, technology is already transforming the nature of warfare. AI is profoundly transforming decision-making processes, introducing a new model of military command-and-control that leverages rapid information processing and adaptability to unpredictable and ever-changing war scenarios.

The integration of AI-DSSs into traditional decision-making models enables faster decision-making cycles, improving efficiency, accuracy, and speed of complex information analysis. These systems provide chains of command with essential support in highly complex and hazardous decision-making environments, where uncertainty and risk are significant variables. They optimize military decisions and operations on the battlefield.

5.2 A Brief Overview of AI Weaponization

The rise of artificial intelligence as a pervasive force in modern society is profoundly altering the nature of warfare (Glonek, 2024; Nadibaidze et al., 2024). The integration of AI models into weapon systems and military strategies is radically transforming the logic of modern conflict, with significant implications for global security, international law, and military ethics. This process has undergone an incredible acceleration due to the Russia-Ukraine war, which are currently expanding limits and capabilities of AI as a decision support tool for the armed forces (Robinson et al., 2023).

Military interest in integrating AI into decision-making processes is primarily driven by two needs: accelerating the decision-making cycle and making decisions more effective compared to traditional command models. Historically, organizational and military decision-making has always been based on the commanders' experience and intuition, as well as on an accurate analysis of available information. Today, this information is becoming increasingly complex, surpassing human processing capabilities.

The main advantage of AI-DSSs in the military is their ability to analyze huge volumes of data from different sources in real time, identifying recurring patterns and hidden schemes to support and accelerate command decisions. These systems can generate automated decisions, as in the case of autonomous weapon systems.

Modern military decision-making is based on a circular command-and-control model known as the OODA loop (Observe-Orient-Decide-Act), and AI developers had to implement this model into AI-DSSs.

The OODA loop is often attributed to USAF Colonel John R. Boyd (1995), but scholarship shows multiple interpretations. We cite Coram's biography of Boyd (2002), and Lind's *Maneuver Warfare Handbook* (1985) to clarify doctrinal lineage, especially with reference to U.S. Marine Corps doctrine. Our use of OODA emphasizes its organizational implications for decision tempo, not a single doctrinal interpretation.

As is known, the OODA loop includes four procedural stages:

- 1) Observation: collecting heterogeneous data related to the mission;
- 2) Orientation: processing the collected data to understand the strategic environment and outline possible scenarios, attributing meaning to the available data (sensemaking);
- 3) Decision: commanders determine a course of action to achieve the desired outcome;
- 4) Action: commanders test selected hypotheses by interacting with the operational environment and use feedback mechanisms to adapt and execute the chosen course of action (COA).

The main goal of the OODA loop is to complete the decision-making process faster than the adversary, enabling faster reactions and the ability to interrupt the opponent's OODA cycle. The faster the OODA loop is executed,

the greater the advantage over the enemy.

However, commanders face another significant challenge in accelerating the decision-making cycle: reducing uncertainty. They must analyze huge volumes of information from different sources in limited time. In fact, uncertainty is the main factor that pervades the dynamic decision-making environment of military operations, characterized by unpredictability and sudden changes in scenarios, which inevitably influence commanders' decisions on the battlefield, determining the need for thorough training to make decisions in complex operational environments.

In combat, soldiers must operate under pressure, in conditions of constant stress and danger, where the 'human factor' plays a crucial role. The high volatility of variables prevents commanders from relying on training data sets for AI systems that are repeatable and applicable to all scenarios.

AI-driven OODA loop command-and-control models aim to reduce the level of uncertainty that commanders face in combat (i.e. the 'friction' factors), or the unknowns that compromise their effectiveness. These systems constantly acquire real-time situational data and contextual feedback (uncertainty modeling), allowing the AI to process and adapt information promptly, enabling effective sensemaking. High processing speed is essential to transform information into a real strategic advantage against the enemy. Any excessive delay would allow the adversary to change its strategy, negating the information advantage.

Unfortunately, in war, knowledge has a short lifespan. The rapid pace of action limits the amount of information that can be collected and processed in a timely manner. Therefore, the crucial challenge for military decision makers (analysts, senior officers, and field commanders) is to deal with uncertainties in a race against time.

To tackle these uncertainties in a highly turbulent and risky battle environment ('fog of war'), militaries have developed new data collection tools, integrating them into AI-DSSs. Examples include sensor feeds from drones and satellites, which collect massive amounts of raw data from any geographic point at any time; or head-mounted displays (HMD) worn by the troops, connected to remote tactical commands (Minguela-Castro et al., 2021).

As the volume of data collected from various sources has grown exponentially since the early 2000s, analysts cannot process all the information obtained from reports, sensors, videos, and feeds without AI support. The organizational and military decision-making cycle has therefore become a scalable system. Due to the increasing speed and complexity of warfare fueled by large-scale multidomain operations, simultaneous OSINT data streams, and fluid operational scenarios (such as those in densely urbanized conflict zones), AI-DSSs have now become indispensable to rapidly process the collected data. They provide relevant real-time information (payload) from the five warfare domains (land, sea, air, space, and cyberspace), with the aim of achieving information dominance over the enemy.

The goal of implementing AI in organizational and military decision-making is to produce faster, more accurate, responsive, and effective decisions regarding the use of force while protecting civilians and their settlements from the harm of military operations (Rasch et al., 2003). This objective is becoming increasingly relevant in urban warfare scenarios, where conducting operations without causing collateral damage is more difficult due to the civilian presence and the interconnected nature of military and civilian sites. The current conflict in the Gaza Strip and recent IDF (Israeli Defense Forces) operations in Lebanon have highlighted this difficulty.

In such operational scenarios, AI-DSSs can process large data volumes to improve commanders' situational awareness on the battlefield, enabling new approaches to combat (Lee et al., 2023). Furthermore, the need to anticipate and disrupt enemy decisions has led to improvements in predictive analytics systems, which predict adversary moves based on specific probability indicators. However, these systems remain inaccurate. Some systems are even trained to bluff by implementing psyops simulation models capable of developing stratagems and deception methods: these are an integral part of military strategies and can impose a strategic advantage over the enemy (game theory).

Unfortunately, one of the most controversial aspects of AI militarization is the development of autonomous weapon systems (LAWS), which use AI-DSSs capable of identifying, selecting and hitting targets without direct human intervention, based on automated analyses, as well as the acquisition of biometric data of possible targets, such as the *Lavender* system (powered by *Palantir Technologies*), currently used by the IDF to eliminate prominent members of Hamas during the Gaza War; or *The Gospel (Habsora)*, used for targeting infrastructures and buildings, which calculates the potential collateral damage. As is well known, the use of LAWS raises ethical and legal concerns related to the loss of meaningful human control over the machine's decisions and the risk of using such autonomous weapons in violation of international humanitarian law (IHL).

6. Evolution of Military Command Models: From Human to Artificial Decision

6.1 The Relevance of the Human Factor: Mission Command

To understand the extent of the transformation underway in organizational and military decision-making, driven by the advent of AI-based decision support systems, and the risks associated with reducing the human factor in today's conflicts, it is essential to start with the human-based command and control model.

This model, known as Mission Command (MC), or decentralized command, highlights the complexity of the human variables that AI developers must implement into AI-DSSs to adapt to dynamic and unpredictable conflict scenarios, in order to support command decisions quickly and effectively. Indeed, MC is a military command model that allows field commanders significant autonomy in the execution of orders (mission-type tactics), as long as such actions are in line with the strategic intent of the higher military command.

The MC model allows field commanders to interpret orders, 'disobey' specific instructions, and break rules, if it ensures the mission's success. In sudden or unexpected tactical situations, personnel involved are required to make decisions quickly and effectively. This command model emphasizes delegation of authority, encouraging personal initiative and responsibility among subordinates to achieve mission objectives, while adapting to changing operational scenarios.

Human qualities such as courage, intuition, and the ability to make timely and appropriate decisions are crucial to MC because they enable commanders to quickly adapt to changing circumstances, while maintaining cohesion and motivation of troops. In practice, the human factor plays a crucial role in military operations according to the MC model, contributing to success or failure on the battlefield: leadership, motivation, mutual trust, and clear understanding of orders form the pillars of MC, enabling commanders to manage operations flexibly and reactively.

However, the MC model also has some risks that need to be considered when deciding to abandon it in favor of a more advanced command model, such as AI-assisted command. Not all human capabilities can be replicated by AI-DSSs. Indeed, MC requires a high level of training and coordination, adherence to military doctrine, focused initiative, critical thinking, deliberative thinking, and trust and respect culture within the chain of command, qualities that are difficult to replicate by AI.

Nonetheless, processing large volumes of real-time information from multiple sources to make rapid decisions that can impact soldiers' lives is a demanding and burdensome task for any experienced commander. Under the pressures of combat, they may overlook critical data needed for decision-making. The informational complexity of current warfare scenarios requires a processing speed that is beyond human capabilities, which however could be more successfully delegated to an AI-DSSs (Fazekas, 2023; Otto & Mănescu, 2023).

After describing the complexity of the decision-making process and the importance of the human factor in this command model, it is even more legitimate to ask for organizational leaders: would an AI system be able to act and decide like a human commander?

6.2 The OODA-loop Evolution: Network Command

The remarkable advances in digital technologies over the past thirty years have led to the emergence of a new military command doctrine (or strategic theory): Network-Centric Warfare (NCW). This doctrine was introduced by the US Department of Defense in the 1990s during the first Gulf War, marking the first networked military operations and the prelude to the militarization of AI (Smith, 2001). Since then, the principles of NCW (or digital warfare) have gained global traction, incorporating a new command-and-control model, Network Command (NC), into military doctrines.

This doctrine relies on the fact that gaining an informational advantage (through geographically distributed computers and sensor networks) translates into a military advantage on the battlefield. The goal is to achieve information dominance by creating a 'system of systems' that integrates intelligence sensors, command and control systems and precision weapons connected via a network, allowing for complete situational awareness on the battlefield, faster target assessment and automatic deployment of geographically distributed weapon systems in the event of a threat, supporting chains of command in the decision-making process. The concept of *Full Spectrum Dominance* assumes that accurate and real-time information improves decision-making, command speed, and operational response effectiveness, enabling better situational awareness. This model of network-centric warfare, so-called C4ISR (Command, Control, Communications, Computers, Intelligence, Surveillance and Reconnaissance), has transformed traditional warfare into algorithmic warfare, but has required the adaptation of human decision-making to a new technological reality.

Today, military decision-making now involves collaboration between humans and machines to achieve timely and meaningful situational awareness (Grujdin, 2022). AI supports this process by using ontologies to share information between human and non-human participants. Achieving interoperability among diverse entities enables decision-making automation and AI-driven decisions. This evolved and technologically advanced OODA loop emphasizes rapid action to penetrate the enemy's decision-making process and gain a competitive advantage (Johnson, 2022a).

AI implementation into MDMP has significantly accelerated this cycle, allowing faster data analysis and more effective adaptation to battlefield changes. In the operational reality of combat, it is an adaptive process, consisting of a continuous cycle of actions and feedback based on the observation of the evolution of the situation on the battlefield, through which information is processed in real time and translated into operational decisions, but also protected from possible external intrusions through the creation of specific 'tactical bubbles' (*Organic Design for Command and Control*). AI allows the analysis and processing of large amounts of data in real time, quickly adapting to unpredictability and surprise factors, generating predictive analysis on potential scenarios and calculating the probabilities of victory or defeat, projecting information dominance over the enemy.

The OODA loop then turns into DOODA loop (dynamic observe-orient-decide-act): 'friction' factors are filtered by sensors (from electronic sensors to human observations) that, by evaluating the resources and constraints of the available assets, allow to dynamically determine the command decisions, repeating the loop adaptively until the mission is accomplished, loss, or withdrawal.

All phases of the DOODA loop are associated with risk assessment and the level of uncertainty, but the complexity and diversity of tasks to be managed are so high that it is unlikely that in the near future there will be an AI capable of directly giving orders to troops on the ground.

To overcome this problem and increase the context capacity of an AI-DSSs, some authors have suggested the development of an 'option generator model', able to prevent strategic and ethical dangers (French & Lindsay, 2022), but we believe that the complexity and unpredictability of battle events are so high and complex that AI cannot provide 100% accurate decisions, totally replacing human decision-making by commanders, but only a statistical approximation, not free from bias and hallucinations. And, in the case of an 'option generator', AI would still have full control over the decision-making process.

AI modeling of military decision making provides the human decision maker with a high level of abstraction and a 'top-down view', improving battlefield success rates. However, information superiority must be safeguarded against enemy actions, which use the same AI-driven approach. Fully networked military forces can share real-time data to synchronize decisions and battlefield actions, sharing information in real time; this allows for rapid integration and synchronization between decision-making and action on the field, increasing readiness and decision effectiveness in battle.

However, network-based military operations management still presents some technical challenges, including interoperability between networks with different classification and security levels; and the reliability of GPS geolocation in low-coverage areas (such as inside forests, caves, bunker, buildings, urban bottlenecks, densely built-up areas, where many of the current war operations take place, such as in Ukraine and the Gaza Strip).

Therefore, human intervention is still necessary to determine the type of data that can or cannot be handled by AI-assisted communication systems, keeping in mind the need for rapid reconfiguration of network-centric operations management processes based on changing battlefield conditions.

Over-reliance on AI systems by military command centers and analysts, as well as the online distribution of false or misleading information, spread by the enemy through disinformation or misdirection operations, can distort the dataset analyzed by AI-assisted intelligence, amplifying the initial error by several orders of magnitude, propagating cascading errors and compromising military operations on the battlefield, according to dynamics typical of chaos theory.

In recent years, in fact, a new military doctrine has emerged that considers the above-mentioned risks, as well as the blurring of the boundaries between war and politics, partly criticizing the perspective of network-centric warfare, or *Fourth Generation Warfare* (4GW), which blends network-centric and traditional warfare.

This doctrine proposes an alternative, or rather hybrid, use of network-centric and traditional warfare, resorting to a more 'primitive' and human warfare, capable of displacing and confusing a hyper-technologized military intelligence apparatus, now accustomed to the use of AI systems, thus making actions on the battlefield more unpredictable and less linear and, precisely for this reason, more likely to succeed.

7. From Human Space to Artificial Agency in Battle

7.1 New Challenges to Traditional Decision-Making Models

Until the 2000s, most war conflicts took place in ‘human space’, a traditional, four-dimensional battlefield perceived by human senses. The advent of autonomous weapons (i.e. LAWS) has further separated warfare from human space, allowing these systems to detect, locate, identify, and attack targets without human intervention.

The rapid evolution of AI is progressively (and dangerously) replacing human control at all stages of the decision-making process, pushing chains of command to increasingly rely on AI-DSSs, raising ethical and legal concerns (Johnson, 2022c).

The distinction between an AI-assisted decision-making system and one that autonomously decides on attacks lies in its programming, which relies entirely on human input. Fully engineering the decision-making process removes human agency, leading to critical ethical and legal implications in combat actions.

This automation, although authorized by the central commands, introduces significant risks (bias, IHL violations, conflict escalations, etc.), and the natural tendency of humans to over-rely on AI systems. Known as the ‘halo effect’, this bias causes the operator to overestimate the AI’s analytical capabilities, leaving the entire decision-making process to the system (AI over-reliance).

Combat situations often force humans to disengage from their responsibilities to preserve themselves. In high-stress and dangerous scenarios, commanders may not have time to verify AI decisions, detect errors, or challenge outputs, risking becoming passive supervisors instead of active controllers of these systems. Further complications arise when AI systems autonomously adapt and modify their parameters through self-learning (deep learning).

It is therefore essential to maintain meaningful human control over AI systems to limit unpredictable outcomes and prevent collateral damage, by implementing specific measures in the machine’s reasoning mechanisms that preserve human judgment in the decision-making process, in order to achieve ‘responsible trust’ by human decision-makers.

A possible solution could be to accept a limitation of the speed of artificial decision-making systems, ensuring sufficient time for human deliberation, but ‘war is war’ and not all countries would be willing to accept it, just to ensure a strategic advantage over the enemy.

7.2 The Concept of Artificial Agency in the Military

The emulation of human decision-making by an artificial agent, facilitated by the modeling of organizational and military decision-making processes, re-proposes the concept of artificial agency (already illustrated in § 4.4).

In the military domain, artificial agency refers to the ability of an artificial agent or system to exercise autonomous decision-making, independent of human operators, based on AI algorithms. This eliminates the typical emotional responses of humans to stress and danger, typical of combat situations (human factor).

The greatest risk lies in the inability of AI systems to perform abstract reasoning (semantic gap). Unlike human decision-makers, AI-DSSs can detect objects they have been technically trained to recognize, but do not fully understand their meaning or context.

This lack of understanding poses a significant risk when AI systems are involved in decisions that require evaluative processes or value judgments, such as distinguishing between military and civilian targets, by incorrectly labeling a person or structure as a target, especially in the case of the use of lethal force by autonomous weapon systems.

Furthermore, AI-DSSs systems may exacerbate biases already present in their training data or even introduce new ones that developers may not have anticipated. These emergent behaviors could significantly alter decision-making processes during operations.

To ensure the responsible use of AI in military operations, it is essential to maintain meaningful human control over autonomous systems. This approach is critical to preserve accountability and minimize unintended effects.

7.3 AI-DSSs: Gray Zones and Information Processing Limits

Unfortunately, the huge volume of information from multiple sources leads to information overload for human operators. As the volume of data increases, it becomes more and more difficult for analysts to effectively organize and interpret it for decision making, often ignoring large portions of data to avoid cognitive overload.

Recent military research on cognitive overload (e.g., DARPA programs) parallels civilian studies of

information-processing limits in organizations (Tushman & Nadler, 1978). Integrating these literatures highlights how decision acceleration pressures manifest both in combat and in firms, where managerial attention is a scarce resource.

According to information processing theory (IPT), effective control requires a match between the information supply, information needs, and the user's processing capacity. Automated systems are deployed precisely to relieve the cognitive overload of human operators by managing the saturation of data in the infosphere and removing conflicting data, thus generating a coherent and prioritized view of the situation.

However, this comes at a cost: the procedure by which the AI system processes the information is invisible to the end users of the decision cycle (i.e. field commanders), leaving them in the dark about the decision mechanisms and returning an 'artificial' situational awareness overview elaborated by the machine. This means that commanders will receive an elaborate, often simplified and washed-out picture of the battlefield, with some elements emphasized and others de-emphasized. Some factors may even be completely omitted from the analysis, probably due to bugs in the AI system's programming or training data.

Since the decisions implemented by the autonomous system are still algorithmic decisions, the set of information loaded during the training phase of the AI underlying the decision-making system, even if entrusted to probabilistic models, would not be able to cover all the unpredictable and changeable aspects of reality, with the result of producing incorrect decisions, putting human lives at risk (Garcia, 2024).

For this reason, the most sustainable and less risky solution is to use a 'collaborative' logic (closer to the Network Command model): human-machine teamwork would allow the human agent to intervene when necessary to correct any AI-DSSs errors. Human judgment is needed to handle unpredictable variables that AI may not consider, especially in lethal scenarios, and should always be reviewed by human commanders to ensure compliance with legal and ethical standards.

Another concern is the 'decision conflict' between the human decision maker and the AI system. In fact, in the decisions of an autonomous decision system there is always a margin of unpredictability, which leads to a drastic decrease in the human being's ability to predict exactly what the machine's decision will be in a specific situation, making the understanding of the system's output uncertain: AI is a 'black box'.

In situations where human lives are at stake, such as planning combat missions or the use of lethal force, decisions made by an automated system should be considered simply as part of the information available from all sources. The most rational (and ethically sustainable) answer is to transform the role of the human operator from an active controller (subject to 'human' decision-making times) to a supervisor of the system, granting the system a safety function in case of error, thus creating a human-machine cooperative model able to preserve significant human control over the machine's decision-making cycle (Saidi, 2022).

8. Towards a Human-Machine Cooperative Decision-Making Model

8.1 Descriptive Findings from the Corpus and Insights from the Battlefields

The descriptive findings indicate that AI is currently used primarily as a 'decision amplifier'. It increases the speed, scale, and scope of data collection, pattern detection, prediction, and option generation, but it does not eliminate the need for human judgment.

The *NVivo* synthesis shows that decision acceleration and situational awareness are the most frequently discussed themes, while meaningful human control and trust in AI-DSSs receive comparatively less attention. This distribution should be interpreted descriptively: it shows where the literature places emphasis, not which variables are causally most important.

Observational data from ongoing conflicts are used to illustrate how AI-DSSs may accelerate target assessment, intelligence processing, and operational coordination. These cases should not be overgeneralized. They are context-specific illustrations of mechanisms identified in the literature, especially the shift from human search for information toward human validation of machine-generated outputs. The theoretical contribution lies in explaining this shift through the decision acceleration paradigm, not in claiming that all military or corporate organizations will experience AI adoption in the same way.

Battlefields offer a clear example of a high-risk environment, where the use of AI effectively responds to the need for rapid decision-making, but events suggest that such use dangerously increases the risk of errors and collateral damage.

At this stage of the 'wave', AI has not yet completely replaced human decision makers, but it is amplifying the speed, scale and scope of human decisions. The 'decision acceleration paradigm' highlights that processing

speed is preferred to the precision of traditional MDMP, as AI allows military operations to be optimized and made more efficient, although it entails risks due to the AI-DSSs imprecision, that cannot yet be avoided. Through AI, the role of the chain of command (as well as that of managerial leadership) shifts from the search for the decision to its validation.

This issue, which we found to be the most critical, can only be mitigated by slowing down the processing speed of AI systems to ensure that human operators remain involved in the decision-making loop. However, as observational data shows, the primary goal of achieving information dominance over the enemy is driven by the speed at which AI processes data. Slowing down the system to allow human involvement could undermine the advantages of faster data processing, as the enemy could exploit this ‘slowness’ to achieve its own information dominance by relying in turn on completely autonomous systems. As military commands become increasingly dependent on AI, the gradual abandonment of human control over MDMP seems inevitable.

8.2 Theoretical Contribution: from Decision Support to Governed Decision Acceleration

The paper’s theoretical contribution is the integration of OODA, bounded rationality, and information-processing theory into a model of AI-driven decision acceleration. AI-DSSs reduce some human cognitive burdens by expanding the organization’s information-processing capacity, but they also create new governance burdens because human actors must understand, contest, or validate machine-generated recommendations under time pressure. Thus, AI changes the nature of managerial and military judgment: the critical skill becomes not only choosing among options, but governing how options are produced, filtered, and presented.

8.3 Human-Machine Teaming as the Preferred Governance Response

The preferred response is a cooperative human-machine decision-making model. In this model, AI supports observation, data fusion, prediction, and option generation, while human decision makers must retain responsibility for interpretation, ethical judgment, escalation, and authorization in sensitive decisions. The model accepts that AI can legitimately accelerate parts of the OODA loop, but rejects full delegation where legal, ethical, or strategic consequences require accountable judgment. In practical terms, this requires documented decision rights, audit trails, thresholds for human review, training against automation bias, and organizational roles dedicated to responsible AI governance.

In the military domain, this model will allow military operations to benefit from AI efficiency and computational power, while ensuring that human values, ethics, responsibility, and respect for the laws of war remain integral to the decision-making process, especially in sensitive situations requiring the use of lethal force (van den Bosch & Bronkhorst, 2018; Konaev et al., 2021; Mayer, 2023; Nalin & Tripodi, 2023).

Human involvement will be essential to oversee and validate machine decisions: this human-machine cooperative model will improve operational outcomes without replacing human judgment, reducing the risks arising from the inherent instability of AI in high-stakes environments. The challenge will be to find a balance between AI efficiency and meaningful human control.

9. Conclusion and Managerial Implications

This study shows that AI-DSSs are transforming high-stakes decision-making by accelerating information processing, increasing artificial agency, and challenging meaningful human control. The theoretical framework explains this transformation through the interaction among OODA cycle, bounded rationality, and information-processing theory. AI-DSSs address bounded rationality by expanding the organization’s capacity to collect and process information, but they also create a new form of bounded rationality when human decision makers cannot fully inspect the speed, complexity, or opacity of machine-generated outputs.

As highlighted in the military, it is still too early to talk about trust in AI-DSSs: it is unwise for military commanders and leaders to rely entirely on machine decisions, especially when it comes to the use of lethal force.

Lessons drawn from a high-risk environment, such as the military, translate directly into core management imperatives for companies adopting the decision-acceleration paradigm. The primary challenge for organizational governance is not the technology itself but maintaining accountability and strategic oversight when operational decisions are delegated to systems exhibiting artificial agency.

The paper’s descriptive contribution is the mapping of recurring themes in the AI-DSSs literature and in conflict-related empirical material, taken as an example of high-risk environment: decision acceleration, situational awareness, risks and benefits, ethical and legal issues, AI-driven command models, artificial agency, reliability, meaningful human control, and trust. The paper’s theoretical contribution is narrower and more

specific: it proposes the decision acceleration paradigm as a mechanism explaining how AI-DSSs shift human decision-making from direct search and deliberation toward validation, supervision, and governance of machine-generated options.

The implications should therefore be interpreted cautiously. As already said, military examples are used as high-stakes analogies for organizations, not as evidence that firms and armed forces face identical decision environments.

Generalizable insight is not a claim about all AI use, but a governance principle: when AI compresses decision cycles, organizations must redesign accountability, escalation, and oversight mechanisms so that speed does not replace judgment.

Managers and military leaders should prioritize three actions:

- 1) Redesigning governance structures: firms must establish clear protocols for meaningful human control over automated decisions that impact legal compliance, brand reputation, or human safety. Ethical and regulatory challenges are, first and foremost, internal governance challenges that demand firm-level mechanisms. Therefore, it is necessary to create, in high-stakes decision-making environments, specific top figures, such as AI manager (or RAI –responsible AI officer– in the military), with the task of leading and supervising the implementation of AI-DSSs in complex organizations, but also training and qualifying AI-enabled DSS-operators, in order to build ‘augmented teams’ able to manage the transformative impact of the emerging ‘fifth wave’.
- 2) Developing hybrid leadership: leaders must be trained to manage decision conflicts, balancing human judgment with machine recommendations; recognize automation bias; interrogate AI outputs; mitigate the ‘halo effect’, and understand the ‘black box’ limitations. The value of human experience, critical thinking, and abstract judgment becomes amplified in managing strategic direction, while AI handles the tactical execution of complex, data-intensive tasks.
- 3) Fostering responsible teaming: a cooperative human-machine decision-making model, where AI supports but does not replace human judgment, offers the best way forward. This ensures that the pursuit of speed and efficiency does not erode organizational values or lead to unintended, escalating risks. As stated above, it is necessary to institutionalize responsible human-machine teaming through roles such as AI managers or responsible AI officers, who supervise AI-DSSs implementation, maintain the codebook of decision rules, and ensure that meaningful human control remains operational rather than symbolic.

The future of organizational and military decision-making will therefore depend less on whether AI can process information faster than humans - it already can - and more on whether organizations can govern this speed responsibly. A sustainable model is one in which AI accelerates observation and analysis, while humans retain accountable authority over interpretation, prioritization, escalation, and final decisions in high-stakes environments.

Informed consent

Obtained.

Ethics approval

The Publication Ethics Committee of the Canadian Center of Science and Education.

The journal and publisher adhere to the Core Practices established by the Committee on Publication Ethics (COPE).

Provenance and peer review

Not commissioned; externally double-blind peer reviewed.

Data availability statement

The data that support the findings of this study are available on request from the corresponding author. The data are not publicly available due to privacy or ethical restrictions.

Data sharing statement

No additional data are available.

Open access

This is an open-access article distributed under the terms and conditions of the Creative Commons Attribution license (<http://creativecommons.org/licenses/by/4.0/>).

Copyrights

Copyright for this article is retained by the author(s), with first publication rights granted to the journal.

References

- Adams, T. K. (2011). Future Warfare and the Decline of Human Decisionmaking. *The US Army War College Quarterly: Parameters*, 41(4), 57-71. <https://doi.org/10.55540/0031-1723.2600>
- Allen, G., & Taniel, C. (2017). *Artificial Intelligence and National Security*. Cambridge, MA: Harvard Kennedy School, Belfer Center for Science and International Affairs. Retrieved from https://www.belfercenter.org/sites/default/files/pantheon_files/files/publication/AI%20NatSec%20-%20final.pdf
- Baboş, A. (2021). Artificial Intelligence as a Decision Making Tool for Military Leaders. *Land Forces Academy Review*, 26-4(104), 269-273. <https://doi.org/10.2478/raft-2021-0034>
- Bandura, A. (1997). *Self-efficacy: The Exercise of Control*. New York, NY: W.H. Freeman.
- Barber, H. F. (1992). Developing Strategic Leadership: The US Army War College Experience. *Journal of Management Development*, 11(6), 4-12. <https://doi.org/10.1108/02621719210018208>
- Barrett-Taylor, R., & Karner, N. (2025). *AI Won't Replace the General: Algorithms, Decision-making and Battlefield Command*. London: The Alan Turing Institute. Retrieved from https://www.turing.ac.uk/sites/default/files/2025-09/turing_final_report_ai_wont_replace_the_general_2025.pdf
- Bennis, W. G., & Nanus, B. B. (1985). *Leaders. The Strategies for Taking Charge*. New York, NY: Harper & Row.
- Bode, I. (2025). *Human-Machine Interaction and Human Agency in the Military Domain*, Policy Brief No. 193. Waterloo, ON: Centre for International Governance Innovation (CIGI). Retrieved from https://www.cigionline.org/documents/3094/PB_no.193.pdf
- Bode, I., Nadibaidze, A., Watts, T., & Zhang, Q. (2025). Ensuring the exercise of human agency in AI-based military systems: concerns across the lifecycle. *Ethics and Information Technology*, 27(50). <https://doi.org/10.1007/s10676-025-09861-2>
- Boyd, J. R. (1995 [2018]). *A Discourse on Winning and Losing*. Maxwell Air Force Base, AL: Air University Press.
- Clark, B., Patt, D., & Schramm, H. (2020). *Mosaic Warfare: Exploiting Artificial Intelligence and Autonomous Systems to Implement Decision-Centric Operations*. Washington, DC: CSBA – Center for Strategic and Budgetary Assessments. Retrieved from https://csbaonline.org/uploads/documents/Mosaic_Warfare_Web.pdf
- Coram, R. (2002). *Boyd: The Fighter Pilot Who Changed the Art of War*. New York, NY: Back Bay Books – Little, Brown and Company.
- Cristofaro, M., & Giardino, P. L. (2025). Surfing the AI Waves: The Historical Evolution of Artificial Intelligence in Management and Organizational Studies and Practices. *Journal of Management History*. <https://doi.org/10.1108/JMH-01-2025-0002>
- Daugherty, P. R., & Wilson, H. J. (2018). *Human + Machine: Reimagining Work in the Age of AI*. Boston, MA: Harvard Business Review Press.
- Davis, J. L. (2024). Elevating humanism in high-stakes automation: experts-in-the-loop and resort-to-force decision making. *Australian Journal of International Affairs*, 78(2), 200-209. <https://doi.org/10.1080/10357718.2024.2328293>
- de Melo, C. M., Kim, K., Norouzi, N., Bruder, G., & Welch, G. (2020). Reducing Cognitive Load and Improving Warfighter Problem Solving with Intelligent Virtual Assistants. *Frontiers in Psychology*, 11. <https://doi.org/10.3389/fpsyg.2020.554706>
- Defense Science Board (DSB) (2016). *Summer Study on Autonomy*. Washington, DC: Department of Defense. Retrieved from <https://apps.dtic.mil/sti/pdfs/AD1017790.pdf>
- Farnell, R., & Coffey, K. (2024). AI's New Frontier in War Planning: How AI Agents Can Revolutionize Military Decision-Making. *Belfer Center for Science and International Affairs*. Retrieved from <https://www.belfercenter.org/research-analysis/ais-new-frontier-war-planning-how-ai-agents-can-revolutionize-military-decision>

- Fazekas, F. (2023). Mission Command and Artificial Intelligence. *Land Forces Academy Review*, 28(2), 69-79. <https://doi.org/10.2478/raft-2023-0010>
- French, S. E., & Lindsay, L. N. (2022). Artificial Intelligence in Military Decision-Making. Avoiding Ethical and Strategic Perils with an Option-Generator Model. In B. Koch, & R. Schoonhoven (Eds.), *Emerging Military Technologies. Ethical and Legal Perspectives* (pp. 53-74). Leiden, NL: Brill-Nijhoff. https://doi.org/10.1163/9789004507951_007
- Gaire, U. S. (2023). Application of Artificial Intelligence in the Military: An Overview. *Unity Journal*, 4, 161-174. <https://doi.org/10.3126/unityj.v4i01.52237>
- Garcia, D. (2024). Algorithms and Decision-Making in Military Artificial Intelligence. *Global Society*, 38(1), 24-33. <https://doi.org/10.1080/13600826.2023.2273484>
- Girei, A. A., Abraham, F., & Majekodunmi, A. O. (2025). Securing AI Models Against Adversarial Attacks in Military Surveillance Systems. *World Journal of Advanced Research and Reviews*, 27(2), 2119-2130. <https://doi.org/10.30574/wjarr.2025.27.2.3084>
- Glonek, C. (2024). The Coming Military AI Revolution. *Military Review*, 3, 88-99. Retrieved from <https://www.armyupress.army.mil/Journals/Military-Review/English-Edition-Archives/May-June-2024/MJ-24-Glonek/>
- Goldfarb, A., & Lindsay, J. R. (2022). Prediction and Judgment: Why Artificial Intelligence Increases the Importance of Humans in War. *International Security*, 46(3), 7-50. https://doi.org/10.1162/isec_a_00425
- Goztepe K., Dizdaroğlu, V., & Sağiroğlu, Ş (2015). New directions in military and security studies: artificial intelligence and organizational and military decision making process. *International Journal of Information Security Science*, 4(2), 69-80. Retrieved from <https://dergipark.org.tr/en/download/article-file/147972>
- Grujdin, I. (2022). Human-Machine Shared Situational Awareness. *Land Forces Academy Review*, 26-4(108), 376-385. <https://doi.org/10.2478/raft-2022-0046>
- Gugerty, L. (2006). Newell and Simon's Logic Theorist: Historical Background and Impact on Cognitive Modeling. *Proceedings of the Human Factors and Ergonomics Society Annual Meeting*, 50(9), 880-884. <https://doi.org/10.1177/154193120605000904>
- Gunawan, Y., Aulawi, M. H., Anggriawan, R., & Putro, T. A. (2022). Command responsibility of autonomous weapons under international humanitarian law. *Cogent Social Sciences*, 8(1). <https://doi.org/10.1080/23311886.2022.2139906>
- Horowitz, M. C. (2019). When speed kills: Lethal autonomous weapon systems, deterrence and stability. *Journal of Strategic Studies*, 42(6), 764-788. <https://doi.org/10.1080/01402390.2019.1621174>
- Hussen, N. et al. (2025). A Review of Real-time Military Training Simulator Based on Improving War Scenario Using AI Tools. *Advanced Sciences and Technology Journal*, 2(2). <https://doi.org/10.21608/astj.2025.343743.1020>
- Johnson, J. (2022a). Automating the OODA loop in the age of intelligent machines: Reaffirming the role of humans in command-and-control decision-making in the digital age. *Defence Studies*, 23(1), 43-67. <https://doi.org/10.1080/14702436.2022.2102486>
- Johnson, J. (2022b). Delegating Strategic Decision-Making to Machines: Dr. Strangelove Redux? *Journal of Strategic Studies*, 45(3), 439-477. <https://doi.org/10.1080/01402390.2020.1759038>
- Johnson, J. (2022c). The AI Commander Problem: Ethical, Political, and Psychological Dilemmas of Human-Machine Interactions in AI-enabled Warfare. *Journal of Military Ethics*, 21-3(4), 246-271. <https://doi.org/10.1080/15027570.2023.2175887>
- Joshi S. (2025). The Role of Artificial Intelligence in Strategic Decision-Making: A Comprehensive Review. *Preprints.org*; 2025. <https://doi.org/10.20944/preprints202505.0047.v2>
- Kissinger, H., Schmidt, E., & Huttenlocher, D. (2021). *The Age of AI: And Our Human Future*. London: John Murray Press, Hachette UK.
- Konaev, M., Huang, T., & Chahal, H. (2021). *Trusted Partners. Human-Machine Teaming and the Future of Military AI*. Report. Washington, DC: CSET – Centre for Security and Emerging Technology. <http://doi.org/10.51593/20200024>

- Kurter, O. (2025). The use of artificial intelligence for decision-making process for strategic management. *OPUS–Journal of Society Research*, 22(2), 195-210. <https://doi.org/10.26466/opusjsr.1632110>
- Kurzweil, R. (2022). *The Singularity is Nearer. When We Merge with AI*. New York, NY: Viking Press.
- Lee, C. -E., Baek, J., Son, J., & Ha, Y. G. (2023). Deep AI Military Staff: Cooperative Battlefield Situation Awareness for Commander's Decision Making. *The Journal of Supercomputing*, 79, 6040-6069. <https://doi.org/10.1007/s11227-022-04882-w>
- Lind, W. S. (1985). *Maneuver Warfare Handbook*, Milton Park, Abingdon, UK: Taylor & Francis Inc.
- March, J. G., & Simon, H. A. (1958). *Organizations*. New York, NY: John Wiley & Sons.
- Mayer, M. (2023). Trusting machine intelligence: artificial intelligence and human-autonomy teaming in military operations. *Defense & Security Analysis*, 39(4), 521-538. <https://doi.org/10.1080/14751798.2023.2264070>
- Michel, A. H. (2024). *Decisions, Decisions, Decisions: Computation and Artificial Intelligence in Military Decision-Making*. Geneva: Technical Report, ICRC – International Committee of the Red Cross. Retrieved from <https://shop.icrc.org/decisions-decisions-decisions-computation-and-artificial-intelligence-in-military-decision-making-pdf-en.html>
- Meerveld, H., Peperkamp, L., Šafář Postma, M., & Lindelauf, R. (2025). Operationalising responsible AI in the military domain: a context-specific assessment. *Ethics and Information Technology*, 27(48). <https://doi.org/10.1007/s10676-025-09865-y>
- Minguela-Castro, G., Heradio, R., & Cerrada, C. (2021). Automated Support for Battle Operational–Strategic Decision-Making. *Mathematics*, 9(13), 1534. <https://doi.org/10.3390/math9131534>
- Moisescu, F., Boşcoianu, M., Prelipcean, G., & Lupan, M. (2010). Intelligent Agents in Military Decision Making. *Science & Military*, 5(1), 58-64. Retrieved from <https://www.proquest.com/scholarly-journals/intelligent-agents-military-decision-making/docview/607928597/se-2>
- Nadibaidze, A., Bode, I., & Zhang, Q. (2024). *AI in Military Decision Support Systems: A Review of Developments and Debates*. Odense, DK: Center for War Studies, University of Southern Denmark. Retrieved from <https://www.autonorms.eu/ai-in-military-decision-support-systems-a-review-of-developments-and-debates/>
- Nalin, A., & Tripodi, P. (2023). Future Warfare and Responsibility Management in the AI-based Military Decision-Making Process. *Journal of Advanced Military Studies*, 14(1), 83-97. <https://doi.org/10.21140/mcu.j.20231401003>
- Newell, A., & Simon, H. A. (1956). *The Logic Theory Machine. A Complex Information Processing System*. Santa Monica, CA: The RAND Corporation. Retrieved from https://archive.org/details/bitsavers_randiplP86ineJul56_3534001/mode/2up
- Otto, S., & Mănescu, G. (2023). Will Artificial Intelligence (AI) Replace a Human Commander in the Army?. *Scientific Bulletin*, 28(1-55), 79-87. <https://doi.org/10.2478/bsaft-2023-0009>
- Petrović, P., Martins Dias, M. T., Denieul, M., Botella Berenguer, V. V., & Grasso, N. (2024). *The Role of AI Decision-Making for Land-Based Operations*. Brussels: Finabel – European Land Force Commanders Organisation. Retrieved from <https://finabel.org/the-role-of-ai-in-decision-making-for-land-based-operations/>
- Probasco, E., Toner, H., Burtell, M., & Rudner, T. G. J. (2025). *AI for Military Decision-Making: Harnessing the Advantages and Avoiding the Risks*. Washington, DC: CSET – Center for Security and Emerging Technology. Retrieved from <https://cset.georgetown.edu/publication/ai-for-military-decision-making/>
- RAND Corporation (2024a). *Strategic competition in the age of AI. Emerging risks and opportunities from military use of artificial intelligence*. Santa Monica, CA: RAND Corporation. Retrieved from https://www.rand.org/content/dam/rand/pubs/research_reports/RRA3200/RRA3295-1/RAND_RRA3295-1.pdf
- RAND Corporation (2024b). *Exploring Artificial Intelligence Use to Mitigate Potential Human Bias Within U.S. Army Intelligence Preparation of the Battlefield Processes*. Santa Monica, CA: RAND Corporation. Retrieved from https://www.rand.org/content/dam/rand/pubs/research_reports/RRA2700/RRA2763-1/RAND_RRA2763-1.pdf

pdf

- Rasch, R., Kott, A., & Forbus, K. D. (2003). Incorporating AI into Military Decision Making: An Experiment. *IEEE Intelligent Systems*, 18(4), 18-26. <https://doi.org/10.1109/MIS.2003.1217624>
- Robinson, E., Egel, D., & Bailey, G. (2023). *Machine learning for operational decisionmaking in competition and conflict: A demonstration using the conflict in Eastern Ukraine*. Santa Monica, CA: RAND Corporation. Retrieved from https://www.rand.org/pubs/research_reports/RRA815-1.html
- Rowe, N. C. (2022). The comparative ethics of artificial-intelligence methods for military applications. *Frontiers in Big Data*, 5. <https://doi.org/10.3389/fdata.2022.991759>
- Saha, S., Shakil, M., Yasmin, R., Chanda, A., & Ferdous, Z. (2024). Artificial Intelligence–Driven Decision Making and Its Impact on Strategic Performance in Organizations. *Journal of Artificial Intelligence General Science (JAIGS)*, 6(1), 649-661. <https://doi.org/10.60087/jaigs.v6i1.444>
- Saidi, I. (2022). The Weaponization of Artificial Intelligence in the Military: The Importance of Meaningful Human Control. *MAS Journal of Applied Sciences*, 7(2), 357-363. <https://doi.org/10.52520/masjaps.v7i2id187>
- Scharre, P. (2018). *Army of None: Autonomous Weapons and the Future of War*. New York, NY: W.W. Norton & Company.
- Sehgal, V. (2024). Artificial Intelligence and Military Decision Making: Revisiting OODA Loop Framework. *International Journal For Multidisciplinary Research*, 6(4). <https://doi.org/10.36948/ijfmr.2024.v06i04.26389>
- Simon, H. A. (1947). *Administrative Behavior: A Study of Decision-Making Processes in Administrative Organizations*. New York, NY: Macmillan.
- Simon, H. A. (1982). *Models of Bounded Rationality*. Cambridge, MA: The MIT Press.
- Simon, H. A. (1991). Bounded Rationality and Organizational Learning. *Organization Science*, 2(1), 125-134. <https://doi.org/10.1287/orsc.2.1.125>
- Trif, R.-C., & Dumitrascu, O. (2024). Leveraging Artificial Intelligence for Enhanced Decision-Making in Management: Bibliometrics and Meta-Analysis. *Annals of the “Constantin Brâncuși” University of Târgu Jiu, Economy Series*, 2, 178-188. Retrieved from <https://ideas.repec.org/a/cbu/jrnlec/y2024v2p179-188.html>
- Tushman, M. L., & Nadler, D. A. (1978). Information Processing as an Integrating Concept in Organizational Design. *The Academy of Management Review*, 3(3), 613-624. <https://doi.org/10.2307/257550>
- U.S. Army (2022). *FM 6-0. Commander and Staff Organization and Operations*. Washington, DC: Department of the Army.
- U.S. Department of Defense (2023). *Data, Analytics, and Artificial Intelligence Adoption Strategy. Accelerating Decision Advantage*. Washington, DC: Department of Defence. Retrieved from https://media.defense.gov/2023/Nov/02/20033333300/-1/-1/1/DOD_DATA_ANALYTICS_AI_ADOPTION_STRATEGY.PDF
- van den Bosch, K., & Bronkhorst, A. (2018). *Human-AI Cooperation to Benefit Military Decision Making*. Brussels: NATO Headquarter. Retrieved from <https://publications.sto.nato.int/publications/STO%20Meeting%20Proceedings/STO-MP-IST-160/MP-IST-160-S3-1.pdf>
- Velazquez, L. E. (2024, May). *Harnessing Deep Learning for Enhanced Military Simulations*. Paper presented at MODSIM World 2024, Norfolk, VA. Retrieved from https://www.modsimworld.org/papers/2024/MODSIM_2024_paper_8.pdf
- Vold, K. (2025). Augmenting military decision making with artificial intelligence. *Cambridge Forum on AI: Law and Governance*, 1(e48), 1-13. <https://doi.org/10.1017/cfl.2025.10022>

Copyrights

Copyright for this article is retained by the author(s), with first publication rights granted to the journal.

This is an open-access article distributed under the terms and conditions of the Creative Commons Attribution license (<http://creativecommons.org/licenses/by/4.0/>).