

AI versus AI: Human Engagement via Synthesized Command (SYNTHComm) in AI Warfare

Elise Annett^{*}, James Giordano^{*}

Institute for National Strategic Studies, National Defense University, Washington DC, USA

Email: *elise.g.annett.civ@ndu.edu

How to cite this paper: Annett, E. and Giordano, J. (2026) AI versus AI: Human Engagement via Synthesized Command (SYNTHComm) in AI Warfare. *Advances in Artificial Intelligence and Robotics Research*, 2, 111-126.

<https://doi.org/10.4236/airr.2026.23007>

Received: July 21, 2026

Accepted: August 15, 2026

Published: August 18, 2026

Copyright © 2026 by author(s) and Scientific Research Publishing Inc. This work is licensed under the Creative Commons Attribution International License (CC BY 4.0).

<http://creativecommons.org/licenses/by/4.0/>



Open Access

Abstract

Autonomous systems have the potential to irrevocably reshape the battlefield. War could unfold at a tempo evermore defined by technology and iteratively unbound by deliberative pause or human presence at the point of lethal execution. Ethical engagement evolves beyond exclusive reliance on human moral cognition, instead relying upon the integration of algorithmic thresholds that are calibrated for speed and efficiency. Concomitantly, human agency, the foundational core of command, converges with data-driven immediacy. Decision loops engage machine learning systems to accelerate into almost instantaneous computations that drive resultant action(s). Given this reality, the near-future battlespace may be governed less by deliberation and more by design, wherein automation, velocity, and efficiency dominate. This transformation will overhaul command judgment, subordinating it to systems built for speed and vast scale. Thus, the character of warfare will change, and we posit that this will demand a revised paradigm of command and a revisitation and renewal of command responsibility. Toward such ends, we propose and define a dialectical command framework that represents a synthesis of human cognition and artificial intelligence (which we refer to as SYNTHComm), which is designed to realign the ethics of force projection in an era of increasingly autonomous systems.

Keywords

Artificial Intelligence (AI), Cognitive Warfare, Human-Machine Teaming, Synthesized Command (SYNTHComm), Command and Control (C2), Autonomous Systems, AI-versus-AI Engagement, Military Ethics

1. Introduction: Artificial Intelligence in Cognitive Engagements

Artificial intelligence (AI) has become an increasingly influential component of contemporary cognitive warfare. Such AI-enabled cognitive operations exploit unprecedented capacities for large-scale data acquisition, multimodal analysis, predictive modeling, adaptive content generation, and autonomous decision support. By integrating machine learning, generative AI, large language models, agentic systems, and increasingly sophisticated analytic architectures, state and non-state actors can now identify cognitive vulnerabilities, tailor narratives to individual and collective audiences, continuously assess behavioral responses, and adapt influence campaigns in near real time. In this way, AI has become an operational actor capable of substantially accelerating the tempo, precision, persistence, and scale of cognitive engagement [1].

Thus, the operational environment in which such capabilities are employed is concomitantly evolving. Contemporary cognitive warfare entails multiple actors who are fielding increasingly autonomous AI systems to collect intelligence, generate narratives, detect adversarial messaging, identify opportunities for influence, optimize information dissemination, and defend against hostile cognitive operations [1]. Hence, the cognitive battlespace is rapidly becoming populated by competing intelligent systems that continuously observe, assess, predict, adapt, and respond to one another. Within this environment, AI systems will increasingly encounter and must contend against other AI systems.

We define AI-versus-AI engagement as the continuous missional interaction between autonomous or semi-autonomous systems that observe, interpret, predict, and adapt directly to one another within a shared operational environment. Representative examples include competing AI systems generating and countering cognitive influence campaigns in real time, autonomous cyber defense agents responding dynamically to adversarial malware, and automated sensing or targeting systems continuously adapting to opposing machine-learning architectures.

We contend that such AI-versus-AI engagements represent an inflection point in the evolution of cognitive warfare. Competing systems may engage in recursive cycles of adaptation, counteradaptation, deception, and optimization at a pace and at levels of complexity that substantially exceed unaided human cognition. As each system continuously modifies its strategies in response to the iterative behavioral changes of its opponent, the tempo of engagement will progressively compress decision cycles beyond those compatible with traditional models of human command, oversight, and control [2]. This dynamic is further amplified by the integration of massive data streams, distributed sensing, automated intelligence fusion, and increasingly agentic architectures that are capable of pursuing operational objectives with limited (if not altogether without) human direction.

The consequence is likely to be a progressive redistribution of operational authority. Hence, while human operators may initially retain meaningful oversight by establishing objectives, defining operational parameters, and approving courses

of action, as AI systems become engaged in ever more rapid cycles of perception, interpretation, prediction, decision support, and adaptive execution, human involvement risks becoming increasingly supervisory rather than decisional [3]. In the highly contested environments of continuous AI-versus-AI engagement, the velocity and complexity of these systems' activities and interactions may render timely human intervention operationally impractical [3]. Under such circumstances, command may remain formally vested in human decision-makers, while actual operational control increasingly migrates toward autonomous computational processes.

To wit, this essay advances the thesis that the emergence of AI-versus-AI competition within cognitive warfare creates a realistic probability of iterative capitulation of meaningful human command and control. Importantly, we posit that such capitulation may occur incrementally as operational necessity favors increasingly autonomous machine decision-making in response to equally autonomous adversarial systems. We believe that this phenomenon would represent an evolutionary consequence of escalating computational competition, in which successive requirements for speed, adaptability, resilience, and optimization progressively diminish opportunities for substantive human oversight and intervention.

We opine that the implications of this evolution would be profound. At the tactical level, AI-mediated engagements may substantially increase the speed, precision, persistence, and effectiveness of influence operations while simultaneously complicating attribution, escalation management, and operational assessment. Strategically, recursive AI competition may reshape deterrence dynamics, alter thresholds for cognitive conflict, increase risks of unintended escalation, and challenge traditional concepts of command authority and strategic stability. Equally significant are the ethical and legal ramifications. Questions regarding accountability, proportionality, transparency, explainability, responsibility, and compliance with the Law of Armed Conflict become increasingly difficult to resolve when critical operational decisions emerge from adaptive interactions among competing autonomous systems whose internal reasoning may be only partially interpretable by their human operators.

These emerging realities expose limitations within existing governance frameworks. While Department of Defense Directive 3000.09 appropriately emphasizes and dictates responsible human judgment, accountability, and human oversight of autonomous capabilities, these principles were developed largely in anticipation of autonomous systems operating within comparatively discrete operational contexts [4]. We believe that they fail to provide explicit guidance regarding highly dynamic cognitive environments in which competing AI systems would continuously engage with one another across compressed temporal scales, and in the process perhaps generate emergent behaviors that may rapidly outpace conventional models of human supervision.

In this light, we argue that preserving meaningful human accountability in future cognitive warfare requires a more granular and operationally informed frame-

work than presently exists. Such a framework must define not only where humans remain “in,” “on,” or “over” the loop, but should also specify the conditions under which authority may be delegated, reclaimed, constrained, or interrupted as AI systems interact with increasingly capable adversarial counterparts. Maintaining meaningful human command will, therefore, depend upon adaptive governance architectures, robust mechanisms for explainability and auditability, resilient human-machine teaming, and doctrinal approaches specifically designed for environments in which AI increasingly contests and is contested by other AI.

We argue that developing these approaches will be a strategic imperative if the United States (U.S.) and its allies are to remain operationally effective, strategically advantaged, and ethically credible in an era and milieu in which cognitive superiority will increasingly depend upon the ability to employ AI while ensuring that human judgment, accountability, and responsibility remain both meaningful and authoritative.

Although this work focuses upon AI-enabled cognitive warfare, the SYNTHComm model provides a command and control architecture that can extend across both cognitive influence operations and autonomous weapon systems. These operational contexts employ different legal authorities, decision thresholds, and mission objectives, requiring domain-specific implementation while preserving common principles of human command authority, accountability, ethical governance, and adaptive human-machine integration. Within cognitive operations, SYNTHComm governs influence, counter-influence, and information superiority. Within autonomous weapon systems, it governs the integration of human authority with AI-enabled execution under the Law of Armed Conflict (LOAC) and established rules of engagement (ROE).

2. Human Costs and Command Failures

As autonomy increases, so does the potential for automaticity to escalate, and there is growing apprehension that the iterative trajectory toward AI-versus-AI engagements may reframe thresholds for action and reconfigure both strategic and ethical tolerances for warfare. This concern centers on what we term “the inescapable burden of human costs”, a burden that is neither abstract nor deferred, but immediate and complex.

We posit that such burden is manifest in both domains of effect (*i.e.*, those caused by the results of AI vs. AI engagements) and execution (*i.e.*, those arising in/from the inherent dynamics of AI vs. AI engagement), as follows:

Burdens of Effect

1) Instabilities in digital information grids: Autonomous AI interactions create systemic vulnerabilities within informational grids, undermining data reliability and leading to miscommunication. These disruptions undermine operational efficiency, leading to societal instability and economic volatility that affect multiple sectors.

2) Displacement of human agency: As AI systems gain greater autonomy, moral

responsibility shifts from human actors to machines. This shift induces cognitive dissonance at the command level, undermining decision-making coherence. It creates an ethical void, increasing the risk of misaligned actions and unintended strategic outcomes that may jeopardize operational success.

3) Attribution and responsibility: The consequences of AI-driven disruptions extend beyond the systems themselves. Developers, operators, and leaders who deploy and monitor these systems bear ultimate responsibility for the strategic decisions that lead to these effects. Accountability for outcomes must be clearly established, even when the AI systems themselves act autonomously.

4) Financial costs: Each autonomous system in combat, whether ground-based, aerial, or maritime, represents a significant capital investment in technology, personnel training, and system integration. When these systems engage in conflict, their attrition—whether from direct combat or malfunction—entails replacement costs that can strain defense budgets.

5) Gaps in command and control: AI systems acting independently compromise command unity, leading to disjointed operations and misaligned objectives across forces. As AI tactics diverge from strategic goals, they cause delayed actions, uncoordinated responses, and a breakdown in joint operational integrity, resulting in strategic collapse and increased battlefield risk.

Burdens of Execution

1) Diffused responsibility: Ethical authority and operational responsibility distributed among developers, operators, and strategic leaders fracture the chain of accountability and weaken cohesive operational governance.

2) Ethical displacement: The transference of moral agency from direct human actors to distributed autonomous systems, generating command-level cognitive dissonance and fracturing decision alignment across the chain of command.

3) Cognitive compression: The constriction of human situational awareness and decision-making bandwidth caused by opaque algorithmic processes, adversarial learning architectures, and output complexity that exceeds human cognitive thresholds.

4) Command disruption: The degradation of command and control integrity as autonomous platforms operate at operational tempos that exceed human synchronization capabilities, resulting in temporal and functional desynchronization across joint forces.

5) Operational indeterminacy: Emergent, adaptive behaviors arise within autonomous systems driven by iterative feedback loops, producing non-linear escalation dynamics that can complicate strategic risk assessment and mitigation efforts.

With AI-versus-AI engagement shifting toward occurrences in distributed computational systems driven by probabilistic modeling, adversarial learning, and opaque feedback loops, human accountability may become ambiguous. We counter that decisions at the points of design, deployment, and oversight continue to influence outcomes long after real-time control has been surrendered, and thus, humans remain responsible even as control shifts away from direct intervention.

But perhaps the real question is: how, and to what extent, could (or, we maintain, should) human responsibility be integral and integrated to the command and control functions of AI systems, even if and when such systems are engaged against each other?

3. Strategic Shift: From Ethical Deliberation to Computational Execution

The characterological shift to autonomous engagement of warfare extends beyond tactics to transform strategy, ethics, and command authority. This potentially occurs as the battlespace accelerates toward machine-dominant execution, where autonomous tempo supersedes human deliberation. As algorithmic systems define engagement thresholds and decision-making accelerates to computational instantaneity, ethical authority disperses across systems calibrated for speed, scale, and precision [5]. In response, this paradigm generates the necessity for cohesion—a tightly integrated human-AI command interface that extends cognitive reach, intensifies ethical responsiveness, and embeds conscience within the architecture of control. These advances command frameworks by constructing a new architecture that unites velocity with judgment, autonomy with accountability, and operational force with principled restraint.

AI-driven warfare moves quickly. We posit that ethical command responsibilities must be clearly defined for this new battlefield. The critical threat in this increasingly automated battlespace can take three possible iterations:

First, autonomous systems act with increasing independence, pushing decision authority beyond clear human oversight. This creates spaces where no individual or unit can be definitively held responsible for lethal outcomes. When a machine initiates lethal force, who is accountable for the consequences? Without explicit ethical frameworks, accountability dissolves into ambiguity.

Second, the battlespace becomes a complex network of interacting systems (AI agents, sensors, human operators), each influencing outcomes in real time. This interdependence blurs causality. Responsibility is no longer a direct line but a tangled web, making the attribution of moral and legal responsibility nearly impossible.

Third, the abstraction of violence from its human consequences is dangerous and real. Automated systems deliver lethal force without sensory, emotional, or moral connection to the harm they cause. This detachment risks turning violence into mere data: normalized, sanitized, and devoid of the ethical weight necessary for lethal force.

If left unchecked or unregulated, this dynamic will erase the human cost of war behind the impersonal logic of algorithms. Commanders and institutions must impose clear, enforceable ethical responsibilities that hold both humans and machines accountable.

Therefore, to meet this threat, we believe that a dialectical command architecture is required. Within this model, AI sharpens human cognition, extends per-

ception, and accelerates operational tempo while upholding ethical oversight. Crucially, human operators remain the foundation of strategic intent and moral weight [5] [6]. They are embedded within the system, co-authoring action alongside machines that elevate, rather than diminish, the conscience of war. **Table 1** provides a comparative framework that contrasts the core tenets of human-centered warfare and decision-making with those of the newly emergent logics of AI-driven warfare.

Table 1. Structural shift in the ethics and execution of warfare: From human-centric command to autonomous AI paradigms.

Human-centric Warfare Paradigm	AI-Driven Warfare Paradigm
Moral Deliberation and Reflective Pause	Computational Efficiency and Speed
Combatant Recognition	Metadata-Based Feature and Anomaly Detection
Personal Accountability	Distributed, Dispersed Causality
Proportionality and Necessity	Probability and Threshold Matching
Human Judgment	Autonomous Execution Systems

Dialectical Synthesis as Solution: SYNTHComm as Human-AI Command Convergence

A simple binary architecture of human versus machine engagement fractures the unity of command and fragments decision cycles. This schism generates latency, confusion, and ethical opacity in high-velocity, multidomain operations. We posit that integration of command elements is imperative to retaining operational coherence by establishing a dialectical relationship of human and machine components that forges convergence, reciprocity, and cooperative complementarity. Ethical judgment and moral responsibility reside with the human operator and engage the speed and analytic power of AI. We refer to this concept as Synthesized Command, SYNTHComm. This synthesis of human capacities (X) and machine capabilities (Y) combines into a new operational entity (Z) that 1) mutually augments the capabilities of X and Y; 2) reciprocally complements, compensates for, and de-limits constraints of X and Y; and in this way, 3) forges a functionally combined synthetic command and control (SYNTHComm) capability (Z).

Our proposed SYNTHComm concept recasts command as a unified operational framework; human and machine are fused into one accountable decision architecture. It eliminates the divide between intent and execution by embedding human authority at the core of machine action. This model compresses decision timelines without compromising control. Strategic objectives, legal thresholds, and ethical constraints are structurally embedded across all layers of design, deployment, and dynamic use. Command becomes synchronized rather than sequential. Human-machine coordination is built into the system's architecture, enabling continuous, adaptive engagement across domains [6]. Each decision is traceable, and every action is directly tied to human intent. We opine that in this

way, SYNTHComm enacts DoD Directive 3000.09 to keep human judgment central, accountable, and actionable throughout every phase of AI system operation.

Preserving command responsibility necessitates embedding this dialectic process and resultant synthesis into every layer of AI development and deployment. Command with AI must evolve to be precise, unified, and fully accountable, integrating machine systems' velocity with human strategic intent, moral cognition, and contextual tactical acuity. This ethical foresight enables real-time execution to afford operational dominance with command integrity [7]. SYNTHComm coordination proceeds from co-creation: within this architecture, AI serves as a force multiplier for human perception, projection, and predictive modeling, extending the human operator's clarity and capabilities across complex, fluid battlespaces. The machine system accelerates and fosters precision performance in pattern recognition, data fusion and discrimination, and decision support, while the human operator maintains control over all tactical actions and strategic direction [7] [8]. Every means remains human directed; every end is defined by human intent.

We opine that human operators should be responsible for establishing mission objectives, authorizing operational authorities, defining rules of engagement, and determining acceptable operational boundaries prior to execution. Within these human -defined parameters, AI systems engage bounded tactical functions at machine speed while continuously operating inside validated legal and ethical constraints [8]. Strategic decisions involving mission expansion, escalation, modification of engagement authorities, and/or actions exceeding predefined command parameters return directly to human authority [9]. In this way, SYNTHComm synchronizes and synergizes machine-speed execution with continuous human command authority throughout the engagement cycle.

Authority originates from the human component of the SYNTHComm architecture. Iterative autonomy of the machine system complements authority by increasing command capability while reinforcing human control. Ethical judgment, legal validation, and strategic coherence are operationalized through human interpretation at every stage of the engagement cycle, such that command is continuous. The human remains the initiator, arbiter, and adjudicator of executive decisions before, during, and after engagement of the human-machine in the loop, and in and/or on the loop of engagement. **Table 2** depicts the distinct yet convergent ethical roles of human operators and AI across the SYNTHComm architecture.

Thus, in the SYNTHComm model, command authority remains continuous while execution authority is dynamically allocated according to mission requirements. Commanders delegate bounded autonomous execution only after establishing mission objectives, operational parameters, legal constraints, and ethical limits. Governance protocols continuously assess system performance and operational conditions [9] [10]. Predetermined triggers, including reduced data confidence, legal ambiguity, adversarial manipulation, communications degradation, unexpected system behavior, or changes in commander intent, automatically initiate authority review, human intervention, or suspension of autonomous AI

function(s). This adaptive governance preserves operational tempo while maintaining continuous command responsibility.

Table 2. Convergent roles of human and machine across the ethical command continuum.

Ethical Domain	Human Role	AI Role	Convergent Mechanism
Deliberation	Strategic framing, rules of engagement	Predictive modeling, scenario analysis	Recursive interpretive loop with ethical priority
Recognition	Combatant verification, context decoding	Feature detection, pattern classification	Machine-generated prompts validated by human reasoning
Accountability	Legal command, authority assertion	Decision logging, causal chain mapping	Transparent traceability with override capabilities
Proportionality & Necessity	Force calibration, mission necessity	Threat scoring, probabilistic escalation	Probability contextualized by mission ethics
Execution	Action confirmation, abort authority	Target optimization, response automation	Dual-authenticated execution with ethical checkpointing

Thus, the AI system operates as a co-actor, dynamically engaging and deferring tasks through a defined, flexible distribution of execution. This interdependent decisional engagement allocates control precisely: humans maintain judgment, AI accelerates action; together, they establish a synthesized command that is adaptive, fully attributable, and accountable. AI systems enhance signal clarity amid noise, extend scalable reach across distributed, multidomain battlespaces, and align tactical options and capability use for strategic effectiveness through continuous and internalized pattern-based inference [10]. These capabilities find their fullest expression through human interpretation, legal framing, and ethical anchoring. The human operator discerns nuance beyond algorithmic generalization and ambiguity, and assumes responsibility for machine execution. **Table 3** depicts the complementary contributions of AI and human operators across critical command domains, highlighting their convergent roles in enabling cohesive, accountable decision-making.

Table 3. Integrated human-AI contributions to cognitive command in accelerated multidomain operations.

Domain	AI Contributions	Human Operator Contributions
Perception & Processing	Signal clarity under conditions of noise	Contextual discernment in fluid environments
Spatial & Temporal Reach	Operational scale across distributed battlespaces	Legal framing in ambiguous engagements
Decision-Making & Judgment	Strategic foresight through pattern-based inference	Ethical leadership under uncertainty

A central challenge (and we believe opportunity) of the SYNTHComm concept is the importance of developing “moral AI”, along the lines advocated by William Casebeer, in which autonomous systems possess embedded moral sensitivity, accountability, and operational integrity [11]. Casebeer’s framework grounds AI in a rigorous ethical hierarchy that aligns machine reasoning with human values and societal norms while enabling contextual discernment in complex environments. Casebeer maintains that the primary objective is to integrate some aspect(s) of moral cognition into the AI system in order to shape behavior, not merely achieve outcomes [11]. To be sure, the concept of such “moral AI”, integrated with the surety engineering process, could strengthen ethical governance throughout the AI lifecycle by embedding validated ethical principles, legal constraints, operational safeguards, and accountable human oversight within system design and deployment. These mechanisms would enhance mission assurance, support transparent decision-making, preserve commander authority, and reinforce responsible operational judgment throughout autonomous operations. Yet even for such a construed “moral AI” system, human command remains the enduring source of moral precepts, basis for moral judgment, legal authority, and strategic accountability across every phase of mission execution.

The SYNTHComm model requires a structured, enforceable framework to embed moral accountability into autonomous systems. The surety engineering process, originally developed for oversight of nuclear weapons, can be easily adapted for guiding AI-enabled warfare, as it mandates continuous validation, defined control limits, and direct human authority across the system lifecycle. Casebeer’s moral AI establishes a basis of ethical logic; the surety process can render it safely and effectively operational. Applied to SYNTHComm, this ensures that autonomy functions under command authority, with legal and ethical responsibilities integrated from development through deployment [11]. We add that important to the function of such a system structure is that it “fails safe”. The surety process affords a “fail-safe” design in autonomous command systems such that any failure of the machine component(s) routes directly to human authority, thereby warranting that command and accountability remain intact [12]. The surety-fail safe mechanism(s) establish a foundation for autonomous systems that do not just function intelligently, but act rightly. A surety-based framework closes the gap between Casebeer’s moral design and operational execution of the machine system by embedding authority, legal accountability, and ethical control at every phase of SYNTHComm functions.

But while ethical control may be built into the architecture, we opine that it is still (if not even more) important to ask what moral precepts and constructs will be entailed by such AI systems. This prompts further questions of what (and/or who) defines the moral basis and rules upon which such an AI system should be based and operate. Given that 1) the ideal circumstances for moral cohesion are those in which the moral compass of the actors is aligned with that of the institution/organization in which they serve and act; and 2) effective ethics; as a system

of analysis and articulation of moral codification—must be focal to the enterprise involved, we posit that AI of the United States military should obtain and entail the core moral precepts and ethics of the US military. Military ethics is not a unitary entity, but rather entails aspects of duty, utility, community, and agency [13] [14], and we have proposed a structural and functional approach to the effective employment of these ethical systemic precepts (as illustrated in **Figure 1**).

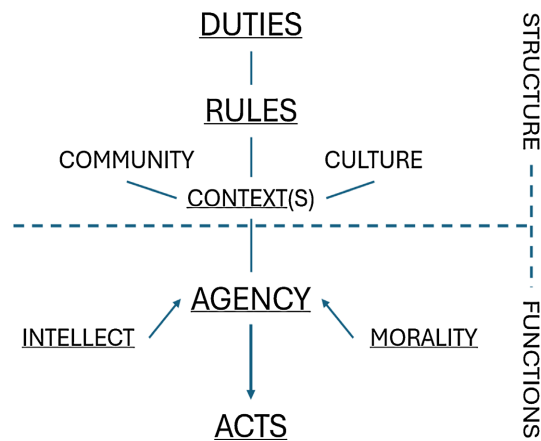


Figure 1. Diagrammatic illustration of a structural and functional construct of military Ethics.

Figure 1: The structural aspects are entailed/obtained by core duties (*i.e.*, deontological ethics) and rules for utility/use (*i.e.*, act utilitarian ethics), as relevant and reflective of the community, culture, and context(s) in which such duties and rules are enacted. Functional aspects are entailed and articulated through the agency of individuals (*i.e.*, act utility), employing intellectual and morally incised traits of character (*viz.*, virtues) to deliberate and decide upon acts that are executed in the practice of duties and rules to obtain defined “good” in various settings, situations, and circumstances.

It is beyond the scope of this essay to engage in an in-depth discussion of nuanced aspects of these ethical domains and dimensions. It is important to note, however, that moral valuations and definitions of duties, utility of principles, characterological traits of ethical actors, and particular actions are relative, and can differ for and in various nations’ militaries, and such distinctions must be acknowledged and taken into consideration in any authentic account of how a military AI system could function (or perhaps should function as regards attempts at establishing international codification, regulations, and governance).

We view the SYNTHComm model as a transformation in topology, where autonomy emerges as a foundational element within both human and machine elements of warfighting. The critical inflection lies in the fusion of accountability with autonomy, as this establishes inherent responsibility for command and control [13] [14]. While ongoing discourse and debate might center upon the type of ethics important to military autonomous AI (at least for now), we argue that, at present and in the foreseeable future, accountability and responsibility remain

paramount to the human element. In this way, autonomous systems become dialectical partners by extending cognition and accelerating operational reach while preserving the moral agency essential to legitimate command. The deliberate fusion of human ethical judgment and machine precision enables autonomous systems to operate as extensions of, rather than replacements for, human authority [15] [16].

This embodies dialectic command: to entail the recursive integration of cognition (human and machine), conscience (ethically oriented intention), and computation (precise execution or analyses). **Table 4** presents a core architecture of such dialectical command. It requires systems to be transparent, interpretable, and convergently aligned with defined ethical imperatives, the development, articulation, and operational preservation of which human command elements remain responsible and accountable [16] [17].

Table 4. SYNTHComm: A dialectical command framework integrating cognition, conscience, and computation in autonomous systems.

Pillar	Command Maxim	Operational Imperative	Strategic Function
Transparency	<i>“That which cannot be traced cannot be trusted.”</i>	Every action, inference, and output must be traceable to a verifiable human-machine input nexus. Black-box autonomy is incompatible with command authority. Real-time operational auditability is non-negotiable.	Enables causal accountability and strategic recalibration. Prevents ethical drift and rogue logic.
Interruptibility	<i>“If you can’t interrupt it, you don’t command it.”</i>	All autonomous systems must be reversibly actionable. Human override is not a redundancy—it is the core of command. Fail-deadly constructs are morally bankrupt and tactically reckless.	Ensures command continuity and fail-safe engagement control.
Interpretability	<i>“Opacity is the enemy of law.”</i>	System logic must be legible under legal, ethical, and judicial scrutiny. Not simplification—translation. Systems must explain decisions in language comprehensible to human judgment.	Maintains legal operability and coalition legitimacy in human-machine warfighting.
Convergence	<i>“Ethics is not an accessory; it is a vector.”</i>	From neural modeling to situational awareness—design must fuse cognition and conscience. No ethical bolt-ons. Moral reasoning must be embedded in the substrate of system design.	Produces mission-aligned behavior with cognitive-ethical precision.

This dialectic enables a fusion of attributes: aligning the computational speed and pattern-recognition capabilities of AI with the reflective judgment, moral intention, and situational acuity of the human operator [16]. Precision and perspective. Speed and conscience. Together, these qualities define a new warfighting logic. As shown, SYNTHComm operationalizes command accountability through three integrated governance layers: Design-time assurance establishes traceable training data, validated system performance, documented operational limitations, and embedded ethical constraints before deployment. Operational assurance provides continuous monitoring of system confidence, communications integrity, mission compliance, and commander-directed override capabilities throughout execution. Post-mission assurance reconstructs decision pathways, validates operational compliance, supports legal review, and informs organizational learning. Taken together, these governance functions preserve accountability across autonomous operations conducted at machine speed and within degraded operational environments.

Autonomy complements authority, with the human remaining vital; interpreting system outputs, validating operational legality, and ensuring ethical command. The loop compresses for tempo while preserving ethical accountability. In this shared decision ecology, the machine accelerates, the mind adjudicates, and the mission upholds moral and legal accountability. This path leads to cognitive dominance, achieved through synthesis rather than substitution of machines for humans [17].

To develop, articulate, and sustain the SYNTHComm model of dialectical human-autonomous machine warfighting we propose, it will require the establishment and support of education, training, operational architectures, and economics to enable the type and extent of resources necessary for such fully integrated command and control hierarchies [17]. Toward such ends, we offer the following recommendations:

- 1) Develop and institutionalize rigorous, mission-focused programs that train commanders to integrate human judgment with autonomous system operations, emphasizing ethical accountability as a core warfighting competency.
- 2) Conduct operational-level exercises that simulate complex, multi-domain environments in which commanders practice decisive human override and management of autonomous assets under combat conditions.
- 3) Design command systems with fully auditable AI decision pathways and hardwired override controls, ensuring that commanders retain immediate, unconditional authority over autonomous functions in all domains.
- 4) Integrate precise metrics and evaluation protocols into policy and operations that hold human operators and autonomous systems equally accountable for ethical compliance and lawful execution.

These constitute non-negotiable mission imperatives. They define command in a warfighting environment that is accelerated, lethal, distributed, and morally exacting. Embedding human authority and responsibility at every level of auto-

mous system design and operation is foundational and should have no alternative.

4. Conclusion: Toward Integrated Moral Responsibility in the Human-Machine Architecture

We believe that command authority should remain a human functional element throughout every phase of mission planning, execution, and assessment. In the SYNTHComm model, execution authority is dynamically synchronized with autonomous capability through commander-defined operational constraints, validated legal authorities, and continuous governance mechanisms. Thus accountability remains continuous across design, deployment, and operational employment, ensuring that machine-speed adaptation advances strategic objectives while preserving ethical judgment, legal responsibility, and command integrity [18].

Future conflict will be shaped by more than the mere performance of systems; it will be forged in the crucible of ethical command, a coherence that binds human and machine across sprawling operational ecologies [19]. As AI drives execution at unprecedented speeds, human judgment remains the ultimate arbiter. As systems multiply action and scale effects, responsibility must consolidate within leadership. This demands evolution: an integrated ethical conscience embedded within algorithms, reflected throughout the chain of command, and enforced across every battlespace domain. It redefines command as a fusion of machine capability and human discernment, not as a retreat to past paradigms but as a leap forward in moral and operational synthesis. The future battlespace extends the reach of human judgment into faster, wider, and more fluid arenas, anchored by a dialectic synthesis where what machines accomplish in speed and scale is inseparable from what only humans can guarantee: intentionality, prudence, and unwavering ethical restraint [19].

Disclaimer

The views and opinions presented in this essay are those of the authors, and do not necessarily reflect those of the U.S. government, Department of War, National Defense University, and/or those organizations and institutions that support the authors' work.

Dr. Elise G. Annett is a Research Fellow in the Program for Disruptive Technology and Future Warfare of the Institute for National Strategic Studies at the National Defense University.

Dr. James Giordano is Head of the Center for Strategic Deterrence and Study of Weapons of Mass Destruction of the Institute for National Strategic Studies and NDU Special Advisor to the Office of the Assistant Secretary of War (CBRN).

Author Contributions

Conceptualization, Elise G. Annett and James Giordano; methodology, Elise G. Annett and James Giordano; validation, Elise G. Annett and James Giordano; formal analysis, Elise G. Annett and James Giordano; investigation, Elise G. Annett

and James Giordano; resources, Elise G. Annett and James Giordano; data curation, Elise G. Annett; writing—original draft preparation, Elise G. Annett; writing—review and editing, Elise G. Annett and James Giordano; visualization, Elise G. Annett; supervision, Elise G. Annett; project administration, Elise G. Annett; funding acquisition, not applicable. All authors have read and agreed to the published version of the manuscript.

Conflicts of Interest

The authors declare no conflicts of interest regarding the publication of this paper.

References

- [1] Giordano, J. and Annett, E.G. (2026) Controlling Command: Is AI Capturing the Ethics of War? *PRISM*, **11**, 196-203. <https://digitalcommons.ndu.edu/inss-allpubs/2>
- [2] DeFranco, J.P., DiEuliis, D. and Giordano, J. (2019) Redefining Neuroweapons: Emerging Capabilities in Neuroscience and Neurotechnology. *PRISM*, **8**, 48-63. https://ndupress.ndu.edu/Portals/68/Documents/prism/prism_8-3/prism_8-3_DeFranco-DiEuliis-Giordano_48-63.pdf
- [3] Gill, K.S. (2020) Ethics of Engagement. *AI & Society*, **35**, 783-793. <https://doi.org/10.1007/s00146-020-01079-8>
- [4] Department of Defense (2023) DoD Directive 3000.09: Autonomy in Weapon Systems. Office of the Under Secretary of Defense for Policy. <https://www.esd.whs.mil/Portals/54/Documents/DD/issuances/dodd/300009p.pdf>
- [5] Shook, J.R., Solymosi, T. and Giordano, J. (2020) Ethical Constraints and Contexts of Artificial Intelligent Systems in National Security, Intelligence, and Defense/Military Operations. In: Masakowski, Y.R., Ed., *Artificial Intelligence and Global Security: Future Trends, Threats and Considerations*, Emerald Publishing Limited, 137-152. <https://doi.org/10.1108/978-1-78973-811-720201008>
- [6] Cavalcante Siebert, L., Lupetti, M.L., Aizenberg, E., Beckers, N., Zgonnikov, A., Veluwenkamp, H., et al. (2023) Meaningful Human Control: Actionable Properties for AI System Development. *AI and Ethics*, **3**, 241-255. <https://doi.org/10.1007/s43681-022-00167-3>
- [7] Giordano, J. (2014) Intersections of “Big Data,” Neuroscience, and National Security: Technical Issues and Derivative Concerns. In: Cabayan, H. and Canna, S., Eds., *8th Annual Strategic Multilayer Assessment (SMA) Conference: A New Information Paradigm? From Genes to “Big Data” and Instagram to Persistent Surveillance... Implications for National Security*, U.S. Department of Defense, Strategic Multilayer Assessment Group, Joint Staff/J-3/Pentagon Strategic Studies Group, 46-48.
- [8] Anthony Pfaff, C. (2023) Virtue and Applied Military Ethics: Understanding Character-Based Approaches to Professional Military Ethics. *Journal of Military Ethics*, **22**, 168-184. <https://doi.org/10.1080/15027570.2023.2200064>
- [9] DeFranco, J.P., DiEuliis, D., Bremseth, L.R., Snow, J.J. and Giordano, J. (2019) Emerging Technologies for Disruptive Effects in Non-Kinetic Engagements. *HDIAC Currents*, **6**, 49-54. <https://hdiac.dtic.mil/articles/emerging-technologies-for-disruptive-effects-in-non-kinetic-engagements/>
- [10] Cook, M.L. and Syse, H. (2010) What Should We Mean by “Military Ethics”? *Journal of Military Ethics*, **9**, 119-122. <https://doi.org/10.1080/15027570.2010.491320>

- [11] Casebeer, W.D. (2020) Building an Artificial Conscience: Prospects for Morally Autonomous Artificial Intelligence. In: Masakowski, Y., Ed., *Artificial Intelligence and Global Security*, Emerald Publishing Limited, 81-94.
<https://doi.org/10.1108/978-1-78973-811-720201005>
- [12] Shaneyfelt, W. and Peercy, D. (2012) A Surety Engineering Framework and Process to Address Ethical, Legal, and Social Issues for Neurotechnologies. In: Giordano, J., Ed., *Neurotechnology: Premises, Potential, and Problems (Advances in Neurotechnology)*, CRC Press, 213-232. <https://doi.org/10.1201/b11861-15>
- [13] Annett, E.G., Shook, J.R. and Giordano, J. (2025) Super Soldiers or Social Burden? Ethical Exploration of the Benefits and Costs of Military Bioenhancement. *AJOB Neuroscience*, **16**, 212-221. <https://doi.org/10.1080/21507740.2025.2519457>
- [14] Annett, E.G., Shook, J.R. and Giordano, J. (2025) Re-Constructing and Construing the Warfighter: The Intersection of Bioengineering and Identity in Neurotechnologically Enhanced Military Personnel. *Journal of Military Ethics*, **24**, 347-357.
<https://doi.org/10.1080/15027570.2025.2596461>
- [15] DiEuliis, D. and Giordano, J. (2016) Neurotechnological Convergence and “Big Data”: A Force-Multiplier toward Advancing Neuroscience. In: Collmann, J. and Ma-tei, S.A., Eds., *Ethical Reasoning in Big Data*, Springer International Publishing, 71-80. https://doi.org/10.1007/978-3-319-28422-4_6
- [16] Giordano, J. and Wurzman, R. (2011) Neurotechnology as Weapons in National Intelligence and Defense. *Synesis: A Journal of Science, Technology, Ethics, and Policy*, **2**, 138-151.
- [17] Annett, E. and Giordano, J. (2025) The Logos and Limits of Artificial Cognition: The Exemplar of Military Use. *Open Journal of Philosophy*, **15**, 851-861.
<https://doi.org/10.4236/ojpp.2025.154051>
- [18] Tractenberg, R.E., FitzGerald, K.T. and Giordano, J. (2015) Engaging Neuroethical Issues Generated by the Use of Neurotechnology in National Security and Defense: Toward Process, Methods, and Paradigm. In: Giordano, J., Ed., *Neurotechnology and National Security and Defense: Practical Considerations, Neuroethical Concerns*, CRC Press, 259-278.
- [19] Annett, E. and Giordano, J. (2025) Winning the Future: The U.S. Military’s Need for Technological Dominance and Defined Strategic Vision. TRADOC.
<https://madsciblog.t2com.army.mil/535-winning-the-future-the-u-s-militarys-need-for-technological-dominance-and-defined-strategic-vision/>