

The Kill Chain Is Not a Theory of Victory

By: [Brian Lampert](#), VP of Sales and Business Development

The Panacea Trap

Don't mistake a targeting system as the blueprint for military AI.

Palantir's [Maven Smart System \(MSS\)](#) may be one of the most consequential AI-enabled capabilities the Department of War has fielded to date. It can fuse data, visualize the battlespace, compress targeting workflows, and help commanders move from detection to action at machine speed. Those are not small things. In a fight, latency kills. The force that can see, decide, and engage before the enemy can react owns a real advantage.

At the architecture level, MSS is a data-integration and targeting platform. Computer vision flags objects in sensor feeds, a fusion layer correlates them into tracks, and a language-model layer from Palantir's AIP summarizes activity and compares strike options, so an operator can move from detection to action [inside one system that used to take eight or nine](#).

But wars are not won by kill-chain optimization alone. They are won by commanders and frontline leaders who understand context, assess risk, adapt plans, and connect tactical actions to operational and strategic purpose.

MSS is not the problem. In many ways, it is proof that the Department can field AI-enabled software at operational scale. The problem is the temptation to treat an AI-enabled kill-chain accelerator as a comprehensive answer to military AI and automation. Warfighting is not simply a targeting problem. The Department needs AI that helps commanders reason from facts to meaning, from meaning to options, from options to risk, and from risk to decision. The panacea trap is the assumption that a faster kill chain is the same as solving military AI for warfighting.

The Kill Chain Wins Fights

Improving the kill chain is important.

I spent too many years chasing terrorists and insurgents with disconnected tools, too few analysts, and too many systems that did not talk to each other. Data arrived by email, appeared in chat rooms, sat buried in databases, or remained trapped inside siloed applications that did not talk to each other. Humans were the integration layer. It was slow, methodical, and exhausting. We made it work and captured hard-won knowledge in scattered Microsoft OneNote pages, briefing slides, and shared drives. Institutional memory was a patchwork of tribal knowledge.

And that was against an adversary that, while dangerous, did not represent the speed, scale, or technical sophistication of a near-peer nation-state. This is why sensor-to-shooter matters. A commander who can see the current operational and intelligence picture, identify a fleeting target, understand available effects, and move from detection to action faster than the enemy can react has a meaningful edge.

This is exactly the kind of problem AI and automation should attack. Humans should not be the copy-and-paste function between incompatible systems, spending their valuable time searching for facts that software can retrieve, correlate, and present in seconds. That is the promise of systems like MSS. They attack a real and painful problem: too much data, too many systems, too little time, and too few humans with enough cognitive bandwidth to make sense of it all.

But the Kill Chain Does Not Win Wars

The same speed that compounds a sound plan multiplies a flawed one, and a commander owns that risk.

The same experience that taught me targeting also taught me its limits. The most well-oiled targeting and strike machine cannot win a war absent coherent operational and strategic decisions. A well-run engagement can still be disconnected from the broader purpose of the campaign. This is where the Department needs to be careful. The kill chain is optimized for speed, synchronization, and action. It cannot, by itself, determine whether those decisions are wise.

This is not a minor distinction. If the underlying plan is sound, speed compounds the advantage. It allows commanders to seize fleeting opportunities, impose dilemmas, protect friendly forces, and act before the enemy can adapt. But if the underlying plan is wrong, a faster kill chain does not create decision advantage. It accelerates error. A bad plan carried out slowly is bad. A bad plan carried out at machine speed is worse.

That is the risk of automation without operational understanding. Automation compresses time. Compressed time can reduce space for skepticism, dissent, red-teaming, and commander's judgment. If the assumptions are right, that compression creates tempo. If the assumptions are wrong, it exponentially scales the mistake.

The argument is to accelerate more than the kill chain. We need AI that helps commanders understand the operational and strategic context, see options and decision points, surface risks and tradeoffs, rapidly plan and replan, and assess whether actions are producing the intended effects. Those systems must challenge assumptions, not just route approvals. It needs tools that can help a staff ask whether the plan remains valid, not just whether the next target can be serviced. It needs AI that can support restraint as well as action, adaptation as well as tempo, and judgment as well as speed.

That is where humans remain essential. As I argued in [“Humans are for Edge Cases,”](#) AI should absorb routine, repeatable, lower-risk tasks so warfighters can focus on the ambiguous and consequential decisions that require judgment. The Department should field systems that make warfighters faster. But it should not confuse faster targeting with better warfighting. The future of military AI cannot stop at finding, fixing, and finishing targets. It must help commanders observe, orient, decide, and act under conditions of uncertainty. The kill chain wins fights. It does not, by itself, win wars.

Context Turns Data Into Military Knowledge

A clean common operating picture can masquerade as understanding.

A sensor can detect movement. A system can plot that movement on a map. An algorithm can correlate the movement with other tracks. A targeting agent can match the appropriate fires platform to the target. But none of that, by itself, explains what the enemy is trying to accomplish.

Ten vehicles moving north is data. A feint, a withdrawal, a spoiling attack, a logistics displacement, or the opening move of a larger encirclement is military knowledge. Choosing among courses of action is human judgement. The first can be sensed. The second must be interpreted. The third must be decided.

We are not doing enough to help commanders move from observation to meaning: Why does it matter? What should we do about it? How does this affect our goals and objectives? The Medium essay [“From Battlefield Data to Military Knowledge”](#) frames this problem well: command is not optimized merely for speed; it is concerned with coherence of intent and the interpretation of military meaning.

Context turns information into meaning. It connects observations to doctrine, terrain, timing, logistics, enemy patterns, political constraints, weather, morale, deception, and the commander’s intent. Context is what allows a staff to distinguish between a retreat and a trap, between an opportunity and a distraction, between a target that is merely visible and a target that is operationally decisive.

[I have argued before](#) that common operational and intelligence pictures are indispensable because they show the “what,” “where,” and “when” of the battlefield. But they can also create a dangerous illusion of understanding when leaders mistake clean visualization for comprehension. The map can show that something happened. It cannot always explain why it happened, whether it matters, or how it should change the plan.

Sensemaking is where today’s AI opportunity becomes most interesting. This is the work AI should increasingly help staff perform. Not by replacing judgment, but by expanding the commander’s aperture before judgment is applied. A good measure of success is whether staff hours move from aggregation to judgment. AI can generate hypotheses, compare activity against known patterns, surface anomalies, summarize relevant reporting, identify contradictions, test assumptions, and present competing interpretations. It can help a staff move faster from “something happened” to “here are

plausible explanations, here is the evidence for each, here is what we should collect next, and here are the decisions this may force.”

The Department needs to move faster from sensor to shooter. But the next leap in military AI is not simply a faster path from detection to engagement. Because the decisive advantage in war will belong to the force that can turn data into meaning faster than the enemy can turn meaning into action.

Targeting Is One Function of an Operational AI Stack

Targeting is one function. We need an AI-powered stack to holistically support warfighting.

A system of record stores what happened. Command also needs systems of work, where intent is expressed, options are built against it, and action is governed as it happens.

To be clear, MSS is already more than a simple map or targeting workflow. Palantir and others describe it as an AI-enabled data integration and decision-support system, and its capabilities are expanding into planning, analysis, and third-party integrations. The question is whether the Department should allow any single platform’s expansion to become its default theory of AI-enabled command.

The Department needs an operational AI stack that should include systems that represent commander’s intent, surface assumptions, reason about enemy behavior, generate options beyond fires, visualize risk, maintain running estimates, interoperate across vendors, and operate at the edge. In practical terms, that means at least eight capabilities beyond dynamic targeting:

- **Commander’s intent representation:** AI systems should understand the mission, end state, operational approach, constraints, and risk tolerance.
- **Assumption surfacing:** AI should identify known and hidden assumptions, show which ones are most fragile, and alert staffs when new data undermines them.
- **Enemy intent and deception analysis:** AI should generate competing hypotheses about enemy behavior, providing probabilistic scoring and identifying black swans.
- **COA generation beyond fires:** A real operational assistant should propose options across all the functions of warfighting, providing decision space and options.
- **Risk and tradeoff visualization:** AI should enable commanders to see risks and tradeoffs side by side, ensuring decisions are as informed as possible.
- **Running estimates:** AI should help staffs ask: Did the action create the intended effect? Did it change enemy behavior? Did it advance the campaign? Did it create new risk?
- **Open, multi-vendor extensibility:** The Department should ensure other AI powered tools and systems can interoperate through open interfaces.

- **Edge resilience:** Any AI-enabled architecture must degrade gracefully under denied, degraded, intermittent, and limited connectivity.

Sensemaking Requires Different Primitives Than Targeting

When you choose one of these systems you are buying primitives, and primitives set the ceiling on what it can ever do.

Sensemaking is the work of turning observation into military knowledge, and it depends on the objects a system can represent. A free-form prompt cannot be governed or audited, because the model cannot reliably separate a directive from a description. A typed, governed workflow can, because intent and constraints enter as structured inputs with known meaning.

Turning observation into knowledge also requires resolving entities and their relationships at runtime, and scoring how informative each relationship is, since a “seen near” link tells you less than a confirmed unit association. That requires runtime entity-relationship modeling and dynamic ontologies, paired with passage-level classification tagging and citations, so a generated hypothesis carries its evidence and markings.

Representing intent, assumptions, and uncertainty as structured objects costs engineering effort and some flexibility. In national security that constraint is the feature, because it makes the reasoning inspectable and the system accreditable. A platform that cannot represent intent, assumptions, and uncertainty as first-class objects cannot help a staff reason about a fight. It can only help them act inside one faster.

Authority, Portability, and the Edge Are Architectural Choices

By the time control, portability, and edge operation are features you bolt on, the architecture has already decided them.

Authority over carried-out action is a property of the architecture, not a policy laid over it. It is enforced through [human checkpoints on consequential steps](#), access control inherited from the acting user, and attribution and audit on every action, in a form an accrediting official can inspect. Removing open-ended agency in favor of governed workflows is what contains the [excessive-agency and rogue-agent failure modes OWASP flags for agentic systems](#). Anyone building agentic systems for the national security community should enforce this through guided approvals, role-based and attribute-based access control, and end-to-end audit emission, which also makes a running estimate honest, since an effect no one can audit is an effect no one can trust.

Open, multi-vendor extensibility is an operational requirement. [After a 2026 dispute the Department of War moved off a single AI provider](#), and a senior defense official said [it would never again rely on one model](#). An operational AI stack treats the model as a substitutable, governed input, so swapping it carries no loss of governance, and mission

data never leaves the customer's environment when the model does. The stack should be model-agnostic and connect to existing systems through open interfaces.

Edge resilience is the requirement most tools quietly fail, because most assume cloud access and stable connectivity. Centurion by Legion Intelligence runs governed workflows in denied, degraded, intermittent, and limited conditions, manages identity across nodes, and moves data between central and forward nodes. The architecture that carries a commander's intent into governed action across existing systems, keeps a human in command of every consequential step, and still works when the network fails, is the one Legion Intelligence is building.

A Theory of Victory Starts After the Kill Chain

The kill chain wins fights, and any serious military AI strategy has to make that chain faster, more resilient, and more lethal. The Department should do that. But it cannot mistake faster targeting alone for better warfighting. An AI system that finds targets faster is valuable. An AI stack that helps commanders understand context, challenge assumptions, generate options, assess risk, maintain running estimates, and decide under uncertainty can be decisive.

Use MSS. Improve it. Integrate it. Push it to the edge where it helps warfighters win. But do not confuse the kill chain with a theory of victory.

Continue the series

Follow Legion Intelligence on LinkedIn: [linkedin.com/company/legion-intel](https://www.linkedin.com/company/legion-intel)

Explore all Command Papers: [legionintel.com/command-papers](https://www.legionintel.com/command-papers)

Ready to talk: <https://www.legionintel.com/contact>