



Shee Atiká
GOVERNMENT SERVICES

The Right Human in the Loop

Building an AI-Competent Force for Agentic Systems



Bob Oshel
VP, Technology
bob.oshel@sheeatikagov.com

INTRODUCTION

About the Speaker



Bob Oshel

VP of Technology
Shee Atiká Government Services

Currently owns technology strategy across a multi-entity portfolio and the internal R&D program that moves early-stage concepts to fielded, marketable capability.



20+ Years in DoD & Enterprise Infrastructure

Cloud and security architecture for the DoD and Fortune 500 enterprises



Program Track Record

USAF CloudOne, DISA milCloud, SAIC, Lockheed Martin



Hands-on with Agentic AI Today

Architecting multi-agent orchestration, zero-trust agent workflows, and confidential computing



Compliance Depth

CMMC, ITAR, IL4/IL5 — the governance lens for today's discussion

AGENDA

Session Roadmap

- 1 The Shifting Question** From “a human in the loop” to the right human in the loop
- 2 The Amplifier Effect** How agentic power magnifies human strengths — and weaknesses
- 3 Key Personnel Is a Competency** The hybrid profile: mission context + technical stack + AI behavior
- 4 Provenance, Governance & Trust** Tying AI-assisted decisions to data, models, approvals, and policy
- 5 Building the Force** AI fluency vs. literacy, role-based competency models, training pathways

Plus key takeaways, a Monday-morning action list, and Q&A.

OUTCOMES

Key Learning Objectives



Explain why the right human in the loop is essential — and how AI power amplifies both good and bad human guidance



Identify competency characteristics for AI-critical roles: AI Supervisors, AI Product Owners, AI Risk Officers



Describe how to maintain provenance, governance, and trust by tying decisions to data, models, approvals, and policy/RMF constraints



Differentiate AI fluency from basic AI literacy — and training that builds V&V and governance discipline, not just tool familiarity

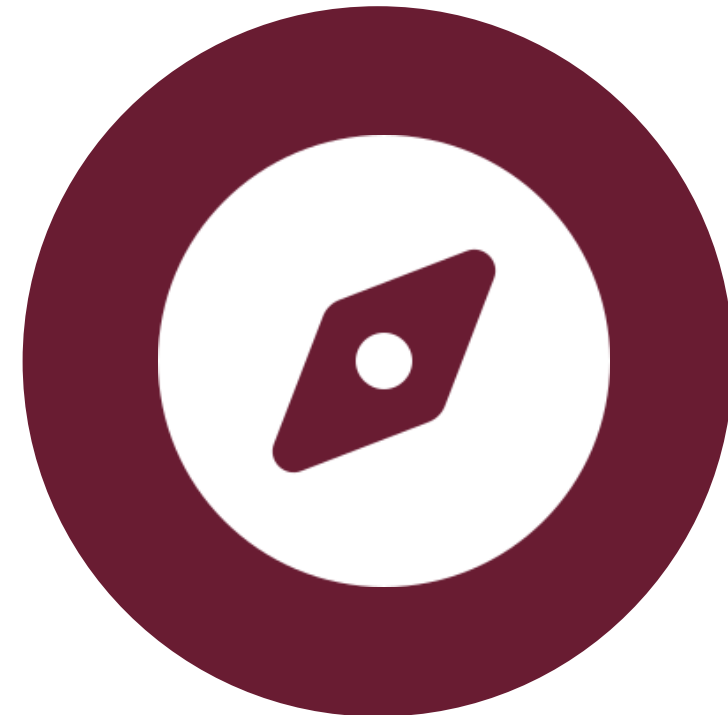


Apply a practical workforce framework to recognize who belongs in the loop and design progressive training pathways

PART 1

The Shifting Question

Agentic AI is already operating at scale. The human-in-the-loop box is checked. So why doesn't that feel like enough?



PART 1 | THE SHIFTING QUESTION

Agentic AI Is Already on Mission



Cyber Defense

Agents triage alerts, correlate sensor data across enclaves, and propose containment actions at machine speed.



Mission Planning

Agents decompose commander's intent into courses of action, logistics estimates, and synchronized task lists.



Software Modernization

Agents generate, refactor, and test code across legacy systems — producing artifacts that ship to production.

Common thread: these systems don't just answer questions — they produce decisions and artifacts that someone must sign.

PART 1 | THE SHIFTING QUESTION

Not the Chatbot Your Policy Was Written For

Chat Assistant

- Responds to one prompt at a time
- Suggests — the human performs every action
- Human drives each step and sees all the work
- Errors are visible at the point they happen

Agentic Mission Partner

- Reasons, plans, and acts across multi-step workflows
- Uses tools: queries data, writes code, files artifacts
- Human supervises intent and outcome — not every step
- Errors can compound silently across the chain

Mental model: a brilliant intern who never sleeps — not a search box. Extraordinary leverage — if someone is actually supervising it.

The Question Has Shifted

Old: ~~“Do we have a human in the loop?”~~

Now: “Do we have the *right* human in the loop?”



A human signature on a flawed output is not oversight. It is liability with a name attached.

PART 2 | THE AMPLIFIER EFFECT

The Amplifier Effect

Agentic AI doesn't average out human judgment — it multiplies it. Strengths scale. So do weaknesses.



PART 2 | THE AMPLIFIER EFFECT

Same System, Opposite Outcome



“Is it right – and can I prove it?”

The AI fluent supervisor

- Clear intent & explicit constraints
- Disciplined review of the reasoning
- Evidence demanded before approval

Amplified out: a disciplined, transparent force multiplier



**AGENTIC
AMPLIFIER**



“Is it finished?”

Untrained human in the loop

- Vague guidance & unstated assumptions
- Hidden bias in framing and data
- Weak controls, rubber-stamp review

Amplified out: polished, confident, scaled-up error

PART 2 | THE AMPLIFIER EFFECT

A Tale of Two Supervisors

Same tasking: the agent ships a 200-rule firewall change to kill a critical CVE.

"Is it finished?"

Untrained

Accepts the change set at face value — it looks complete and well-formatted

Approves without asking which sources or assumptions drove the plan

Change breaks a mission system; no record ties inputs to the decision



Result: outage, broken traceability, and eroded trust in every future AI output

"Is it right?"

AI fluent

Challenges the reasoning: "Show the affected systems and why each rule changed"

Demands evidence — verifies against config baselines and threat intel sources

Approves a corrected set; provenance links data, model, and approval



Result: faster than manual work, fully auditable, trust in the system grows

The Hybrid Profile: Three Fluencies, One Person

Not a billet, not a title, not seniority. A definable, trainable, testable set of skills — and a new hybrid profile.



The Hybrid Profile: Three Fluencies, One Person



Mission Context

Knows what right looks like: the operational stakes, constraints, and second-order effects of a decision



Technical Stack

Understands the systems the agent touches — data sources, integrations, and where artifacts land



AI Behavior

Knows how models fail: hallucination, sycophancy, brittle confidence, and drift under ambiguity

The intersection is the job — and the unicorn. A subject-matter expert alone rubber-stamps. A generic “AI operator” lacks mission judgment. The right human holds all three fluencies at working depth.

PART 3 | KEY PERSONNEL IS A COMPETENCY

What the Right Human Actually Does

-  **Frames Intent & Constraints**
Gives the agent a clear objective, boundaries, and explicit no-go conditions before work starts
-  **Challenges Reasoning, Not Just Conclusions**
Interrogates how the agent got there — assumptions, sources, and decomposition, not just the answer
-  **Interprets Confidence Signals**
Reads uncertainty for what it is — and treats fluent confidence as a claim to verify, not proof
-  **Detects Hallucinations**
Spot-checks citations, data points, and system references against ground truth
-  **Demands Evidence**
“Show your work” is a standing order: sources, tool calls, and intermediate artifacts on request
-  **Preserves Provenance**
Keeps the chain intact from inputs to decision — so every approval can be reconstructed later

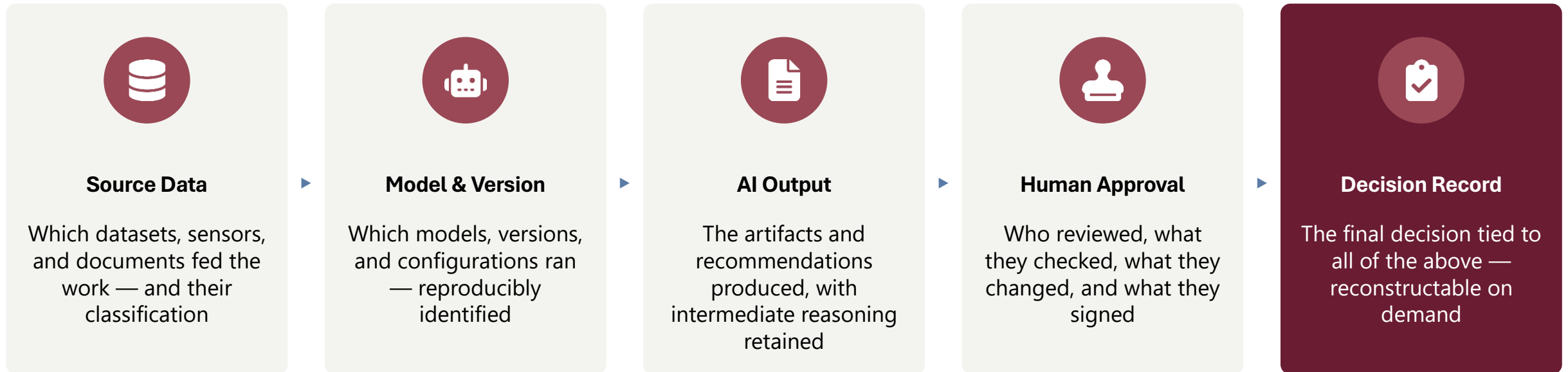
Provenance: The Unbroken Chain

If you can't reconstruct how a decision was made, you can't defend it. Agentic speed makes the paper trail harder — and more essential.



PART 4 | PROVENANCE, GOVERNANCE & TRUST

Provenance: The Unbroken Chain



One broken link — an unlogged model swap, an untracked prompt edit — and the decision is indefensible. For the record: “trust me, bro” is not a provenance record. Neither is “the AI said so.”

Governance: Existing Guardrails, New Strain Points



RMF & Accreditation Continuous-ATO thinking fits agents better than point-in-time approval — but agent behavior changes with every model update, prompt change, and tool grant



Policy & Ethics Constraints Constraints must be machine-enforceable where possible and human-verified where not — written into agent instructions, not just binders



Audit & Oversight Traditional audit assumes human-paced artifacts; agentic volume demands sampling strategies and automated provenance capture



Continuous Monitoring Supervision doesn't end at deployment — drift, new failure modes, and changed data sources require a standing review cadence

Trust is an Operational Property, Not a Feeling



Traceable

Every decision can be walked back to its data, model, and approvals



Reproducible

Rerun the chain and you can explain — or recreate — the result



Accountable

A named, qualified human owns each approval, on the record

PART 5 | BUILDING THE FORCE

AI Fluency

AI 101 will not produce AI supervisors. Fluency is a different curriculum — role-based, progressive, and disciplined.



AI Literacy is Not AI Fluency

AI Literacy

The broader force — everyone

- What AI is and where it helps
- Safe, appropriate everyday use
- Basic prompt hygiene & data handling
- Knowing when to escalate

AI Fluency

Designated AI key personnel

- Directing multi-agent workflows against mission intent
- Verification & validation as standing discipline
- Reading confidence, catching hallucination, demanding evidence
- Governance, provenance, and risk ownership

Fluency Training: Discipline, Not Tool Familiarity



V&V as Reflex

Structured verification drills against ground truth — graded on what the trainee catches, not what they produce



Adversarial Exercises

Briefings seeded with landmines: planted hallucinations, biased framing, broken provenance — who steps on one?



Governance Reps

Practice tying real decisions to RMF, policy, and ethics constraints — until the paperwork instinct is automatic



Role-Specific Scenarios

Cyber supervisors train on cyber workflows; planners on planning agents — fluency is contextual, not generic

Grading standard: a fluent supervisor is measured by the errors they catch, not the outputs they accept.

Framework Step 1: Role-Based Competency Models



AI Supervisor

Is this output right?

Directs and approves agent work day-to-day

- Frames intent & constraints
- Challenges reasoning, verifies evidence
- Owns each approval on the record



AI Product Owner

Is this capability worth running?

Owns the agentic capability over its lifecycle

- Prioritizes use cases vs. mission value
- Manages model/tool change & versioning
- Ensures training & manning for it



AI Risk Officer

Can we defend this?

Independent check on governance & compliance

- Maps workflows to RMF / policy / ethics
- Audits provenance chains by sampling
- Tracks incidents, drift, lessons learned

PART 5 | BUILDING THE FORCE

Framework Step 2: The Skill Map



Mission Expertise

Deep operational context for the domain the agents serve



Data Literacy

Knows the data: sources, quality, classification, and limits



Prompt & Intent Engineering

Translates mission intent into precise agent direction and constraints



Verification & Validation

Systematic output checking against ground truth and requirements



Governance & Risk

RMF, policy, ethics, and provenance discipline applied in-flow



AI Behavior & Limits

Failure modes, confidence reading, and knowing when not to use the agent

Each role weights these differently — the map is the common language; the weighting is the role design.

Framework Step 3: Training Pathways

1 Foundational Literacy

The entire force

- What agents are, safe use, escalation
- Delivered broadly: CBTs, commander's calls
- Goal: a force that uses AI safely

2 Applied Practitioner

Frequent AI users by career field

- Role-relevant workflows & prompt discipline
- Supervised practice on real (low-risk) tasks
- Goal: productive, careful power users

3 AI Key Personnel

Designated supervisors, POs, risk officers

- Adversarial V&V, governance reps, certification
- Recurring evaluation — a rating, not a badge
- Goal: qualified final authority over agents

Applying This in Your Unit (Monday Morning)

-  **1. Inventory** List every place agentic AI already touches your workflows — including the unofficial ones
-  **2. Name the humans** For each, identify who approves outputs today. Are they qualified — or just available?
-  **3. Assess against the map** Score candidate key personnel against the six skill domains; gaps become the training plan
-  **4. Designate & empower** Make key personnel formal: written designation, authority to reject outputs, protected training time
-  **5. Build the pathway** Stand up the three tiers — start small with your Tier-3 candidates and expand from results

REVIEW

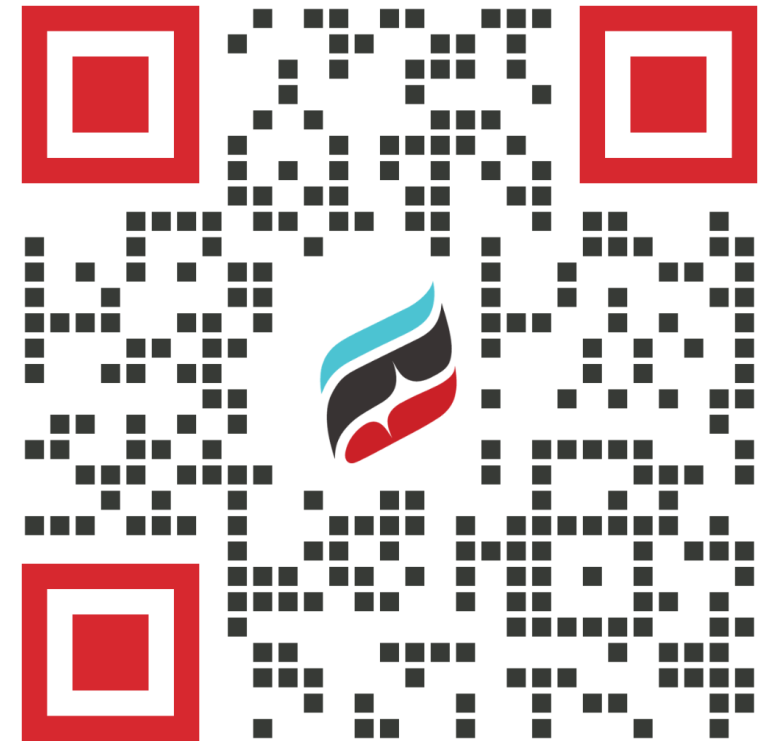
Five Things to Take Home

- 1 Agentic AI amplifies its human - the supervisor who asks “Is it right?” scales judgment; the one who asks only “Is it finished?” scales error with a confident finish.
- 2 “Key personnel” is now a competency: mission context + technical stack + AI behavior, in one person, at working depth.
- 3 Provenance is the discipline that makes decisions defensible — data, model, output, approval, record, unbroken.
- 4 Fluency ≠ literacy: power users need adversarial, role-specific training — not more AI 101.
- 5 **Mission success depends on the right human, prepared and empowered, as the final authority over intelligent machines.**

Questions & Discussion

Three things you can do this week:

- Inventory agentic AI use in your unit — official and unofficial
- Name your candidate AI key personnel and score them against the skill map
- Brief leadership on fluency-vs-literacy before the next training budget cycle



Bob Oshel
VP, Technology
bob.oshel@sheeatikagov.com