

The Right Human in the Loop

Building an AI-Competent Force for Agentic Systems

Agentic AI amplifies whoever directs it. A prepared supervisor turns these systems into disciplined, auditable force multipliers; an unprepared one scales error with a confident finish. **"Key personnel" is now a competency, not a billet** — definable, trainable, and testable.

1 THE HYBRID PROFILE THREE FLUENCIES, ONE PERSON, AT WORKING DEPTH

Mission Context

Knows what right looks like: operational stakes, constraints, and second-order effects.

+

Technical Stack

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

+

AI Behavior

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

2 SUPERVISOR BEHAVIORS USE AS AN ASSESSMENT RUBRIC

- **Frames intent & constraints.** Clear objective, boundaries, and explicit no-go conditions set before work starts.
- **Interprets confidence signals.** Treats fluent confidence as a claim to verify, not as proof.
- **Demands evidence.** "Show your work" is a standing order: sources, tool calls, intermediate artifacts on request.
- **Challenges reasoning, not just conclusions.** Interrogates assumptions, sources, and task decomposition.
- **Detects hallucinations.** Spot-checks citations, data points, and system references against ground truth.
- **Preserves provenance.** Keeps the chain intact from inputs to decision so any approval can be reconstructed later.

3 THE PROVENANCE CHAIN

- 1 **Source data** — datasets, sensors, documents, classification
- 2 **Model & version** — models, versions, configuration, reproducibly identified
- 3 **AI output** — artifacts and recommendations, intermediate reasoning retained
- 4 **Human approval** — who reviewed, what they checked, changed, and signed
- 5 **Decision record** — the decision tied to all of the above, reconstructable on demand

TRUST TEST

Is the decision traceable, reproducible, and accountable to a named, qualified human?

Every link runs under policy, ethics, and RMF constraints. One unlogged model swap or untracked prompt edit breaks the chain.

4 ROLE * SKILL MAP EVERY ROLE NEEDS ALL SIX DOMAINS, WEIGHTED DIFFERENTLY

ROLE	MISSION EXPERTISE	DATA LITERACY	PROMPT & INTENT ENGINEERING	VERIFICATION & VALIDATION	GOVERNANCE & RISK	AI BEHAVIOR & LIMITS
AI Supervisor Directs and approves agent work day-to-day; owns each approval on the record.	●●●	●●○	●●●	●●●	●●○	●●●
AI Product Owner Owns the agentic capability across its lifecycle: use cases, change control, manning.	●●●	●●●	●●○	●●○	●●○	●●○
AI Risk Officer Independent check on governance and compliance; deliberately outside the throughput chain.	●●○	●●●	●○●	●●●	●●●	●●●

●●● **primary** — certify and evaluate recurringly · ●●○ **working depth** — train and refresh · ●○● **awareness** — literacy sufficient. Small units may combine roles, but the functions must all exist and the conflicts must be visible.

5 PROGRESSIVE TRAINING PATHWAY FLUENCY IS A DIFFERENT CURRICULUM, NOT MORE AI 101

TIER 1

Foundational Literacy

The entire force

- What agents are, safe use, escalation
- Delivered broadly — CRTs, commander's calls
- Basic prompt hygiene and data handling

Goal: a force that uses AI safely

TIER 2

Applied Practitioner

Frequent AI users by career field

- Role-relevant workflows and prompt discipline
- Supervised practice on real, low-risk tasks
- Introduction to V&V and escalation criteria

Goal: productive, careful power users

TIER 3

AI Key Personnel

Designated supervisors, POs, risk officers

- Adversarial V&V drills — planted errors, biased framing, broken provenance
- Governance reps against RMF, policy, ethics
- Certification with recurring evaluation

Goal: qualified final authority over agents

Design rule: grade supervisors on the errors they catch, not the outputs they accept — and make Tier 3 a rating you maintain, not a badge you earn once. Agent behavior changes with every model update, tool grant, and prompt revision.

6 APPLY IT IN YOUR UNIT FIVE STEPS, IN ORDER

1 INVENTORY

Map every place agentic AI already touches your workflows — including the unofficial ones.

2 NAME THE HUMANS

Identify who approves each output today. Are they qualified — or just available?

3 ASSESS

Score candidates against the six skill domains. The gaps are your training plan.

4 DESIGNATE & EMPOWER

Written designation, authority to reject outputs without friction, protected training time.

5 BUILD THE PATHWAY

Stand up the three tiers. Start small with Tier 3 candidates; let results make the case.