

The AI Operating Model for Teams

Ed Schaefer · Version 1.0 · August 2026 · edward-schaefer.com/templates/ai-operating-model-for-teams · Free to share

As of: _ (each team stamps its own date)

As AI absorbs the mechanical work of building software, what's left is the part that still has to be human: the judgment. Someone has to go first. Then mark the route.

The Three Shifts

1. From IN-THE-LOOP to ON-THE-LOOP — steer the flow, don't gate it.
2. From TRUST to VERIFICATION — reliability from checking, not model magic.
3. From DISCOVERY WHILE BUILDING to DISCOVERY BY DESIGN — put back the learning AI skips.

The tier scale

Org level: the scale itself

- Tier 1 — just do it. Small, low-risk, reversible. No ask.
- Tier 2 — do it, tell me, let me green-light. A little input, then proceed.
- Tier 3 — quick exchange first. A small decision, or a manual step I need to take.
- Tier 4 — a real conversation. Evaluate options, maybe a research spike, maybe I deploy or do something by hand.
- Tier 5 — fully human. Standing up infrastructure, anything irreversible, anything I have to own. We walk through it together.

Autonomy granted = earned-trust × blast-radius. High trust and low blast radius → full auto. Low trust or high blast radius → I'm engaged.

(From [Human-in-the-Loop Is a Coaching Problem](#).)

Team level, repeatedly: which tier the work sits at

Trust moves both directions: a trusted agent still stops at Tier 5, a brand-new one gets watched even at Tier 1. Set a tier once and never revisit it, and that's the Static Ladder, the named failure mode.

The Four Values, as the tie-break

Use these to break ties when priorities pull against each other, and read them with the line below.

Outcomes over output. Earned autonomy over standing permission. Human judgment where it matters over human hands on everything. Verified work over trusted work.

There's value on both sides. When they pull against each other, choose the left.

The standing practices

- **Delivery-state visibility.** Make delivery-state visible; pay down comprehension debt; never rubber-stamp. Comprehension debt is what builds up when shipped work outruns what your people can still account for.
- **Attention and dispatch capacity.** Ask each of your people what their own dispatch limit is, and what "done" means for the work their agents hand back.
- **Owned durable substrate.** Adopt AI as change management, on an owned and durable substrate, matched to your risk. The substrate is the wiring you own, so it survives the models changing out from under you.
- **Corrections land in the setup, not the chat.** The agents aren't the part that improves. The setup and the person running it are.

Who answers for the outcome

An agent can do the work, but a named person still has to own what happens because of it.

Named owner for this team: _ (fill in if your organization works that way; skip if it doesn't)

Revisit this page: quarterly, and after any agent incident. _ (your team's own trigger, if different)

This page is the short version. The detail lives elsewhere: [field guide](#) for the full canon, the [pre-flight checklist](#) before a rollout, [Definition of Verified](#) for sign-off, the [incident runbook](#) when something breaks.

This is a fit question, not a recipe. Where your team actually sits against this list, and what that's costing, is what the [Engineering Ways of Working Diagnostic](#) is built to find.