

Safety systems, not paperwork

An AI agent with real access is a new hire you onboarded and skipped every safeguard for. Governance decides what it is allowed to do. **Controls are how it stays inside that line once it is live and moving fast** – the same things you give a person on day one: scoped access, a paper trail, a probation, an undo, and a manager who would notice.



THE FIVE CONTROLS ————— each one is a thing you'd give a new hire

- I

Access LEAST PRIVILEGE

The narrowest access that lets it do the job, and nothing beyond. If it only needs to read, it should never be able to write. **Broad access isn't power, it's blast radius.**
- 2

Audit trail A RECORD OF WHAT IT DID

Every meaningful action logged in plain language, so on Monday you can answer what it did and when. **A log nobody can read is the same as no log.**
- 3

Rollback THE UNDO BUTTON

Know how you'd undo an action before you grant the access. If you can't cleanly reverse it, it probably shouldn't be automated yet.
- 4

Monitoring SOMEONE WOULD NOTICE

Tripwires you set on purpose: a metric moving too far, too many actions too fast, a decision it's unsure about. **Silence isn't safe, it usually means nobody's watching.**
- 5

Human-in-the-loop THE PROBATION

Start reviewed closely, then loosen the leash as trust is earned. Never the reverse.

HOW MUCH LEASH: THE AUTONOMY DIAL ————— a dial, not a switch



The tier you grant must match the controls you've actually built. **The classic failure is level-3 autonomy on level-0 controls** – let it run free with no real audit, no rollback, nothing watching. Start low, climb one level at a time as the controls underneath earn the climb.

Monday morning YOUR FIRST MOVE

Take one AI agent you already run – or one you're about to switch on – and answer the five onboarding questions **out loud**: what exactly can it touch, where's the log of what it did, how would you undo its last action, what would trip an alarm, and who is actually reviewing it right now?

Whichever one you can't answer cleanly is an unsupervised employee with real access. **Fix that control before more autonomy, before the next use case.**

OPERATOR INSIGHT · LAYER 4

Controls aren't bureaucracy - they're operational safety systems. So a fast mistake doesn't quietly become a slow disaster.

"Your AI maturity will never exceed your operational maturity."

Free

Score your own controls – the AI-Ready Operations Assessment runs all five layers as a readiness check, on a scale from "ad hoc" to "AI-ready". vihrenlabs.gumroad.com/l/rmkje