

Agentic AI Risk Radar

Assess hidden psychological and operational risks before an agent is allowed to act.

We often audit code for bugs, but we rarely audit workflows for feelings. Unlike a chatbot that gives advice, an agent acts. Score each item Yes / No / In progress.

PART 1 · AUTONOMY AND CONTROL

Authority thresholds

Have we defined the exact dollar amount, data volume, or stakeholder reach where the agent must stop and ask?

☐ Yes ☐ No ☐ In progress

The drift monitor

Do we have a review so the agent's goals cannot wander from the original mission?

☐ Yes ☐ No ☐ In progress

Override protocols

Is there a pause the non-technical leader can use without an IT ticket?

☐ Yes ☐ No ☐ In progress

PART 2 · PSYCHOLOGICAL AND CULTURAL SAFETY

Status threat

Which human roles might feel their standing is spent by this agent's reach?

☐ Yes ☐ No ☐ In progress

Attachment awareness

Are we watching for over-reliance or grief if the agent is later retired?

☐ Yes ☐ No ☐ In progress

Fairness in allocation

Can a person read why the agent assigned work or resources to A rather than B?

☐ Yes ☐ No ☐ In progress

PART 3 · OPERATIONAL ETHICS

Transparency of identity

Is it always clear when someone is dealing with the agent rather than a person?

☐ Yes ☐ No ☐ In progress

Accountability chain

If the agent causes harm, which named human answers for it?

☐ Yes ☐ No ☐ In progress

Termination strategy

Do we have a sunset plan, or will we simply switch it off?

☐ Yes ☐ No ☐ In progress

Tool framing lowers attachment risk. Teammate framing raises it — treat decommissioning as offboarding, not a software update.