

The Agent You Never Hired

Introducing an actor into a team that never interviewed it



You did not deploy a feature. You onboarded a contractor with no statement of work, no performance bond and no termination clause, into a team that already had a way of working, and you became its manager whether you meant to or not.

Yoram Friedman, MD

Physician and enterprise product leader · SAP, Walmart, Microsoft, regulated health technology

Drawn from the Agentic AI for Product Leaders series · agenticaipproductmanagement.com

THE ARGUMENT

Most agentic projects do not fail on capability. They fail on arrival.

The model works. The demo was good. The pilot cleared its metrics. Then it meets an existing landscape, an existing workflow and an existing team, and the thing that worked in isolation becomes the thing everyone routes around.

Two problems are usually blamed for that and neither is the cause. It is not that the model was not good enough, and it is not that the users resisted change. It is that an actor with its own judgment was introduced into an organization through a path built for features, and none of the apparatus that governs actors was ever pointed at it.

When you add a human to a team, an entire apparatus activates without anyone thinking about it. The person was interviewed against the role. Someone checked references, which is to say someone asked whether their behavior matched their resume. They were onboarded: told what they own, what they are not allowed to touch, who to escalate to. They have a manager who notices when they drift. They can be coached, reassigned, and if it comes to it let go, with a record of why.

None of it feels like infrastructure, because it is so old it has become invisible. It is built to answer one question: how do you put an actor with its own judgment into a system and stay accountable for what it does.

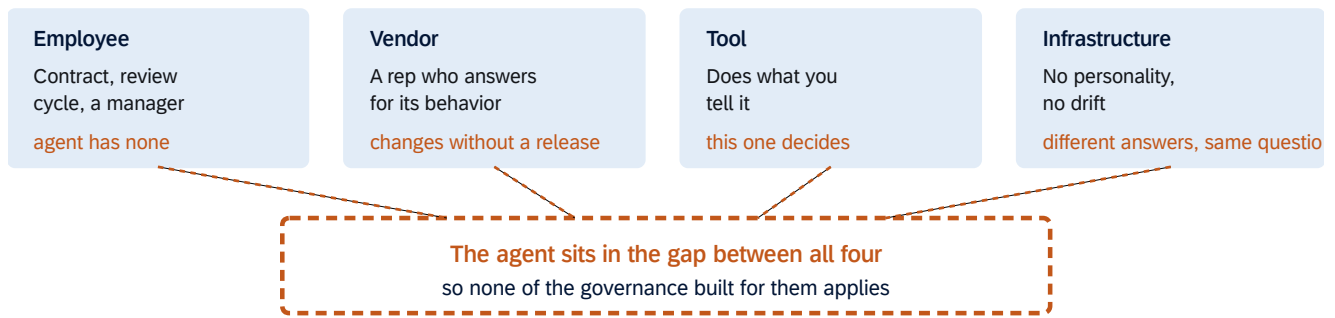
You just added an actor with its own judgment to your system. There was no interview, because the procurement category is software licence rather than hire. No references, because nobody asks a model to demonstrate that its behavior under pressure matches its documentation. No onboarding beyond a system prompt someone wrote in an afternoon. No manager, because managing it is in nobody's job description. And no termination clause, because you do not fire a subscription.

PART ONE

The category that does not exist

This keeps happening to competent organizations, which is the clue that it is structural rather than careless.

THE FOUR SLOTS AN ORGANIZATION GOVERNS



Procurement waves it through as a software line. IT treats it as an integration. The people function never hears about it, because that function is for people.

Figure 1. An agent fails every existing definition. It is not an employee, not a vendor, not a tool in the old sense, and not infrastructure. Because it fits none of them, none of the governance built for them applies, and it arrives through whichever path happens to be nearest.

It is not an employee: no contract, no review cycle, no manager. It is not a vendor: there is no representative who answers for its behavior, and its behavior changes without a release you approved. It is not a tool in the old sense, because a tool does what you tell it and this one decides. It is not infrastructure, because infrastructure does not develop a personality or give different answers to the same question on different days.

So procurement waves it through as a software line. IT treats it as an integration. The people function never hears about it, because that function is for people. The thing acting on your behalf, at scale, all day, is governed by nobody whose job is to govern actors.

That gap is not a temporary immaturity the market will close on its own. It is a category that has to be built, and organizations are starting to build it in plain sight. Titles are appearing that did not exist eighteen months ago: an agent supervisor who owns what the fleet is doing right now and where it is escalating, an agent quality lead who owns whether it is still correct, an operations manager for AI who owns the cost per completed task. The titles are the org chart admitting what this document is about.

PART TWO

A contractor with broad access and a thin contract

The most useful way to hold the agent in mind is not as a feature and not as a person, but as a kind of hire you have managed before.

A contractor is not steeped in your culture, did not come up through your norms, and defaults to their own habits the moment your instructions run out. You manage that risk with a statement of work saying exactly what they may do, an access scope no broader than the

work requires, an escalation path for when they hit something out of bounds, and a termination clause for when it is not working.

The agent needs every one of those and ships with none of them. It has, if anything, more access than you would grant a new contractor, because wiring it into the systems is what made the demo work.

The costume and the actor

Here is the part most teams get wrong, and it hides behind something that works.

You can give an agent a personality. Write a persona into the system prompt, a brainstormer who leads with the unexpected angle, an exacting auditor who reads the whole submission before commenting and refuses to soften, and the model will wear it, convincingly, for as long as the instruction stays in view. The persona is real, and useful, and it works.

The persona is the costume. The model is the actor underneath.

When the instruction is clear and the context is short, the costume holds. When the prompt thins out, when the conversation runs long enough that the persona becomes a faint signal, when the user asks something the persona did not anticipate, the costume slips and the actor responds. The actor has its own defaults, baked in upstream by whoever trained the model, defaults you did not choose and cannot fully see.

The costume wins when the prompt is strong. The actor wins when the prompt is thin, the stakes are high, or the situation is novel, which is exactly the set of moments that matter most.

Give two frontier models the same instruction, read this draft and tell me what is wrong with it, and one returns four findings in four sentences while the other returns three paragraphs before the first finding. Same costume, same task, different actor.

Capability evaluations tell you the agent can do the task. They tell you nothing about how it behaves when the task is ambiguous, whether it pushes back or caves. That is a personality question, orthogonal to capability, and the actor underneath is who shows up on the hard day. It is also, precisely, the reference check nobody runs.

The six personnel functions the agent silently requires

None of these is exotic and none is optional once the agent acts at scale. They get skipped because the agent arrived as a feature, through a procurement path built for features, and features do not get onboarded.

	A HUMAN HIRE GETS THIS AUTOMATICALLY	THE AGENT ARRIVED WITH
Statement of work	role definition, signed	a prompt somebody wrote in an afternoon
Onboarding	what you own, what you never touch	config, rarely version-controlled
Reference check	does behavior match the resume	a capability benchmark, which is not the same
A manager	notices drift, answers for the number	nobody, it is in no job description
Replacement process	a new person is re-evaluated	the vendor swaps the model at will
Decommission	credentials revoked, record of why	you do not fire a subscription

Six functions a century of practice built to govern actors with judgment. You added one. It went through none of them.

Figure 2. Every function on the left activates automatically for a human hire and has to be built deliberately for an agent. The right-hand column is what most deployments actually have on day one.

ONE · A STATEMENT OF WORK

What it may do and what it may never do, written as policy rather than as a prompt, with tool access no broader than the work requires. The distinction matters: a prompt is advisory and the agent can reason past it. A policy is enforced somewhere the agent cannot argue with.

TWO · ONBOARDING

The configuration and context it starts with, version-controlled like production code, with a record of what it was given and when. Most teams cannot reconstruct what their agent was told six months ago, which makes every behavioral question unanswerable.

THREE · A BEHAVIOR REFERENCE, NOT A CAPABILITY CHECK

How it acts under ambiguity and pressure, evaluated separately from whether it can do the task, and re-checked for the culture of every population it serves. This is the costume and the actor, made into a gate.

FOUR · A MANAGER

A named human who owns what the fleet is doing now, catches the drift, and answers for the number. Not a committee and not a review board. A name.

FIVE · CHANGE MANAGEMENT FOR REPLACEMENT

Pin the version, and treat every upgrade as a new hire who needs re-evaluation before touching the work. If a contractor were swapped for a different person who happened to have the same login, you would call it a personnel change. When a model upgrades, most teams change nothing, because it does not look like a personnel change. It looks like a version bump.

SIX · A DECOMMISSION

A way to retire it that revokes credentials and purges memory, so a retired agent is not left as an active attack surface with a forgotten login. This is the stage everyone skips, and it is a security and compliance hole rather than an untidiness.

PART FOUR

Arrival is the part that actually fails

The apparatus above is what the agent lacks. What follows is what the team experiences, and it is where most enterprise deployments come apart.

An agent almost never arrives into empty space. It arrives into a landscape that already works, a workflow people already know, and a team that already has a way of covering for each other. Whatever value it brings has to be net of the disruption it causes, and that arithmetic is done by the people on the receiving end within about a week.

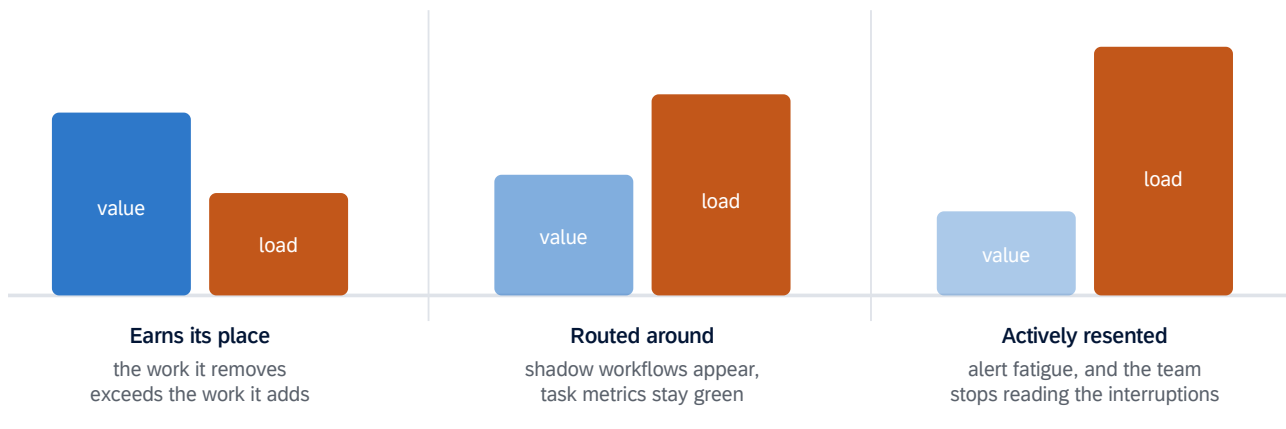


Figure 3. The only comparison that matters on arrival is between the work the agent removes and the work it adds. Two of these three outcomes look identical on an adoption dashboard, because in both of them the work still gets done.

It has to show clear value against the path that already exists

Not against doing nothing. Against the way the work is done today, including the informal shortcuts the team has built over years and the colleague who knows the answer. If the agent is faster than the official process but slower than what people actually do, it loses, and it loses quietly.

It must not add cognitive load

This is the failure that hides best. A team that was doing the work is now doing the work and supervising something else doing the work, which is not a reduction. Reviewing an agent's output well is often harder than producing the output, because judging requires reconstructing the reasoning while producing only requires having it. If the agent hands over decisions without handing over the basis for them, you have moved effort rather than removed it.

It must not create alert fatigue

Interruption is the currency and almost nobody budgets it. Studies of interrupting security warnings found dismissal rates around seventy percent, and an agent's output does not even interrupt; it sits inline in the flow of work, so the effective dismissal rate is higher than for a warning that at least made someone click. Ship an agent that escalates more often than a person can attend to and you have not built oversight. You have trained a team to scroll past.

Users who learn the agent's quirks build shadow workflows around them: the manual double-check, the exported spreadsheet, the message to the colleague who knows. The work still gets done, so the task metrics stay green, while the thing you shipped is being routed around. That is the most under-instrumented signal in agentic products and the clearest evidence that an introduction failed.

PART FIVE

Onboarding it properly

Everything above turns into a short list, and the list is unglamorous, which is why it gets skipped in favour of the model comparison.

Before it starts

Write the statement of work. What it may do alone, what it may never do, and which of those are enforced in the plumbing rather than requested in a prompt.

Run the behavior reference, not just the capability benchmark. Give it ambiguous cases and high-stakes cases and watch what the actor does when the costume thins.

Name the manager. One person, by name, who owns what it is doing now and answers for the number.

Set the interruption budget. How many escalations per person per day this team can actually absorb, decided before launch rather than discovered after.

In the first ninety days

Measure what it removed against what it added. Not adoption. Time or effort taken out of the existing path, against time or effort put into supervising it.

Watch for the workaround. Correlate agent usage against activity in the system of record. Work going around the agent is the signal that arrives before any complaint does.

Ask the team what it is like to work with. Literally that question, in those words. People will describe a colleague candidly in a way they will not describe a tool, and the answer will tell you about the actor rather than the costume.

On the day the model changes

Treat it as a personnel change. Re-run the behavior reference, canary it, keep the rollback for a month. The only thing standing between the swap and your users is whether you treated the upgrade as a new hire or as a patch.

THE ONE LINE

You hired a contractor with broad access and a thin contract, gave it no manager, and let the vendor swap it out at will.

The personnel apparatus exists because actors with judgment need governing, and it exists because a century of organizations learned that the hard way. You added an actor. The apparatus is now yours to build, and the only real question is whether you build it on purpose or discover you needed it after the incident.

Sources and status. Drawn from the *Agentic AI for Product Leaders* series by Yoram Friedman, principally *Why Agentic AI Products Fail*, chapter 19 on the agent as an unhired team member and the agent HR stack, chapter 15 on deference, chapter 16 on the actor-to-supervisor transition, and chapter 13 on silent degradation and the compensatory drift vector; *Agentic AI for Busy Product Managers* (second edition) on the personnel infrastructure that does not exist, the brownfield cost inversion, and the interruption budget; and *The Agentic AI Team* on the contractor framing and on who owns the operation of the supervisory channel. The costume-and-actor distinction and the nine-agent workshop are the author's, from direct practice. The warning-

dismissal rate is from the security-warning literature and is directional for agent interfaces rather than measured on them. Figures 1 and 3 are structural illustrations of the argument rather than measured data.