

The six dimensions

Before we write a line of agent code for anyone, we score their sales system on six things. This is that instrument, written out in full: what each dimension measures, why it gates agents specifically, what good looks like, and what to fix first.

It is meant to be used, not read. Run it on your own organization, in a room with the people who would own the result. You do not need us in the room to do it.

How to use this

- 1 Place your organization on each row of the matrix.** Be honest rather than aspirational. The value of the exercise is entirely in where it disagrees with the story you tell externally.
- 2 Expect an uneven profile.** Nearly every organization is agent-ready on two or three dimensions and absent on one. The absent one is the finding, and it is usually the one nobody owns.
- 3 Sequence, do not remediate everything.** The six are not equally weighted and they are not independent. The reading order is at the bottom of this page.

The maturity matrix

Six dimensions, four levels. Most organizations land in more than one column, which is the point: the shape of the profile tells you where to start.

DIMENSION	Absent No mechanism at all	Manual Happens only because a person holds it together	Assisted Partly systematic, real gaps remain	Agent-ready Automated, observable, and trusted
01 Systems connectedness	Scattered, nothing talks to anything	Many, stitched together by hand	Several, partly connected	Connected, syncing automatically
02 Admin load	Admin eats the week, selling happens in the gaps	Mostly admin and follow-up, roughly 20/80	Roughly balanced, 50/50	Mostly selling, roughly 80/20
03 CRM hygiene	The forecast is a fiction rebuilt by hand	Reps clean it up right before the pipeline review	Mostly current, some deals go stale	Current, every open deal has a real next step
04 Follow-up latency	Often never, leads go cold on their own	A few days, if someone remembers	Same day	Within minutes, automatically
05 Signal response	Nothing, we find out from a press release or never	Someone notices eventually, usually too late	A rep sees it and acts that week	Automated, the account surfaces with context attached
06 Traceability and approval	No way to tell, so we do not automate at all	We would see the result, never the reasoning	We could reconstruct most of it with some digging	Every action logged, and approved before it sends

The dimensions in full

01 Systems connectedness

How many systems hold your sales data, and do they reconcile without a person?

WHAT IT MEASURES

How many systems hold a fact about an account, and whether those facts reconcile without someone doing it. Count the CRM, email, calendar, call recordings, support tickets, product usage, billing, and the spreadsheets in between.

WHAT GOOD LOOKS LIKE

The CRM, email, calls, tickets, product usage, and billing land in one place on a schedule you control, and that place is infrastructure you own rather than a vendor's copy of your business.

WHY IT GATES AGENTS

An agent can only reason over what it can see. Point one at the CRM alone and it inherits every one of the CRM's blind spots. It will draft a confident follow-up to an account that filed three support tickets last week, because nothing told it otherwise. The most common cause of an agent that looks unreliable in production is not the model. It is an agent reasoning over a partial record and being right about the wrong thing.

WHAT TO FIX FIRST

Inventory before integration. List every system that holds a fact about an account and mark which ones are readable by API today. The gap list is your integration backlog, and it is almost always shorter than the room expects.

02 Admin load

How much of a rep's week goes to CRM admin and manual follow-up versus actually selling?

WHAT IT MEASURES

The split between administrative work (logging calls, updating stages, writing recaps, chasing next steps) and live conversations with buyers. Measured per rep, per week.

WHAT GOOD LOOKS LIKE

The post-call work is drafted before the rep reopens their laptop. Nobody is retyping into a CRM something they already said out loud on a call.

WHY IT GATES AGENTS

This is the dimension that funds the program. Admin load is the only item on this list a finance function will approve on its own merits, because it converts cleanly into hours and hours convert into loaded cost. If you cannot state the number, you cannot build the budget case, and the project stalls at the budget conversation rather than the technical one.

WHAT TO FIX FIRST

Measure it before you automate it. Take a two-week sample across five reps and count honestly. The baseline is what your pilot gets judged against, and a pilot without a baseline gets judged on vibes.

03 CRM hygiene

How current and trustworthy is your CRM?

WHAT IT MEASURES

Whether open deals carry a real next step and an honest close date, continuously, rather than in the hours before a pipeline review.

WHAT GOOD LOOKS LIKE

Hygiene is enforced continuously by something that never forgets, instead of reconstructed by people under deadline pressure.

WHY IT GATES AGENTS

Agents act on what the data claims. A close date three weeks in the past is not a neutral blank, it is an instruction, and an agent will follow it. An agent running on a dirty CRM nudges the wrong deals, skips the right ones, and burns the team's confidence inside two weeks. That trust gets spent once. Recovering it costs more than the cleanup would have.

WHAT TO FIX FIRST

This is the strongest candidate for a first agent, because it is the rare one that improves its own substrate: a nightly sweep for deals with no next step, no activity in two weeks, or a close date already past. It posts the list to the owner, and once the rule is approved it fixes the fields itself.

04 Follow-up latency

After a call or an inbound signal, how long until follow-up actually happens?

WHAT IT MEASURES

Elapsed time from a triggering moment to the thing that follows it: the recap email, the booked next step, the reply to a demo request or a pricing-page visit.

WHAT GOOD LOOKS LIKE

Same hour, every time, with the message drafted automatically and sent by a person until the team trusts it enough to stop reading every one.

WHY IT GATES AGENTS

This is where agents produce revenue rather than savings. Every other dimension gives time back. This one closes deals that were otherwise lost to silence, which is a different and much better line on a business case. It also has the cleanest before-and-after measurement of the six, which makes it the pilot most likely to survive scrutiny.

WHAT TO FIX FIRST

Instrument it. Most organizations have never measured median time-to-follow-up and are genuinely surprised by the number. The measurement alone changes behavior before a single agent ships, which makes it the cheapest intervention on this page.

05 Signal response

When a target account raises a round, hires an exec, or launches something, what actually happens?

WHAT IT MEASURES

Whether externally visible buying signals at your target accounts reach the person who could act on them, while they still matter.

WHAT GOOD LOOKS LIKE

The account surfaces on its own with the hook attached and the outreach already drafted, waiting for a rep to approve rather than to discover.

WHY IT GATES AGENTS

Nobody fails this one on purpose. It fails because it has no owner, and work without an owner never reaches a roadmap. It is also the dimension best suited to automation on first principles: it is pure sustained watching, which is the thing humans are worst at and machines are indifferent to.

WHAT TO FIX FIRST

Write the trigger inventory. Which events would genuinely change a rep's priorities this week? Five is plenty. A signal agent watching five well-chosen events beats one watching fifty, because the second one gets muted by week three.

06 Traceability and approval

If an automated action went out on your behalf, could you tell what it did, what data it read, and why?

WHAT IT MEASURES

Whether you can reconstruct what an automated action did, which records it read to decide, and who approved it. And whether a person could have stopped it first.

WHAT GOOD LOOKS LIKE

Every read, every action, and every approval is logged. Autonomous versus approval-gated is set per agent and per action by the business owner, and changing it does not require a deployment.

WHY IT GATES AGENTS

This is the dimension that decides whether the program ships at all. The other five determine whether agents will work. This one determines whether you will be permitted to run them. Security review, legal, and the executive who owns the customer relationship all ask the same question, and there is no version of the answer that starts with we would see the result. Programs that defer this build something impressive in pilot and then die at the enterprise readiness review, having spent the political capital and produced nothing deployable.

WHAT TO FIX FIRST

Decide the graduation policy on paper before the first agent, not after. Everything starts in approval mode: the agent drafts, a person clicks send. Then define what earns an action its autonomy, whether that is a volume of clean approvals, a rejection rate under a threshold, or a named owner's sign-off. Writing this down early is free. Retrofitting it is not.

Reading order: what to do with an uneven profile

The six dimensions do different jobs, and treating them as a checklist to complete in order is the most common way this exercise gets wasted.

FOUNDATIONS

Systems connectedness and traceability

These two gate everything else. One decides what an agent can see, the other decides whether you are allowed to let it act. Neither produces value on its own, which is exactly why both get deferred, and why deferring them is what kills agent programs at the readiness review rather than at the demo.

FASTEST RETURN

Admin load and CRM hygiene

These two pay for the program. They convert into hours and hours convert into a number a finance function will fund. Hygiene has the additional property of improving the substrate the other agents run on, so it compounds.

UPSIDE

Follow-up latency and signal response

These two produce revenue rather than savings, which is a stronger business case but a harder one to attribute cleanly. Run them once you have the credibility from the first two, and instrument the baseline before you start.

The mistake worth naming

Teams try to finish systems connectedness before starting anything. You do not need every system connected. You need the systems your agent reads, which is usually two or three, and you need the governance policy written down. A narrow agent running against complete data for one workflow beats a broad one running against partial data for ten, and it produces the evidence that funds the rest of the program.