

Applying First Principles to AI Readiness

"AI readiness" usually gets treated as a checklist (data warehouse? check. LLM API access? check. an AI strategy doc? check) rather than reasoned from the ground up. Breaking it down from first principles clarifies what actually matters.

Step 1: Strip the term to what it's actually asking

"Ready" is meaningless without an object. Ready *for what*? A company "ready" to use AI for customer support triage has almost nothing in common with one "ready" to use AI for medical diagnosis support. So the first move is: don't ask "are we AI-ready," ask "ready for this specific decision or task to be touched by AI."

Step 2: Ask what has to be true, independent of the word "AI"

Forget that it's AI for a second. Any tool — human or machine — inserted into a decision needs the same underlying conditions to actually help:

1. A real, well-defined problem exists. Not "we should have AI somewhere," but a specific decision or task that's currently slow, expensive, inconsistent, or bottlenecked.

Example: "Support tickets take an average of 14 hours to get a first response" is a well-defined problem you can point AI at. "We should have an AI strategy" is not — there's no task, no baseline, nothing to improve. A retailer that says "let's use AI for inventory" is still vague; "we're overstocking size-M in 30% of stores because reorder decisions are made off last month's average, not seasonality" is a problem AI can actually be pointed at.

2. Sufficient, honest input. The tool needs accurate, representative information about the thing it's deciding on. This is the true "data readiness" — but the fundamental question isn't "do we have a data warehouse," it's "does the information reflect reality closely enough for this decision."

Example: A hospital wanting AI to help triage patient intake has years of structured records — that's sufficient input. A startup wanting AI to predict "which customers will churn" but only logging signups, not usage or cancellation reasons, doesn't have what the task needs, no matter how good the model is. "We have a lot of data" isn't the same as "we have the right data" — a company might have millions of rows of sales data and still lack the one field (e.g., why a deal was lost) that the decision depends on.

3. A way to check the output before it matters. Someone or something needs to catch it when the tool is wrong, before the wrongness causes damage. This scales with stakes — low-stakes draft emails need almost no check; loan approvals need a lot.

Example: A marketing team using AI to draft social captions can just eyeball them before posting — low stakes, cheap to catch mistakes. A bank using AI to flag fraudulent transactions needs a human review step or a well-tested confidence threshold, because a false positive locks out a real customer

and a false negative lets fraud through. Skipping this step is how you get AI systems that confidently generate wrong legal citations that nobody catches until a judge does.

4. The organization can actually change what it does with the answer. If the AI produces a better answer but no process exists to act on it (approvals still take 3 weeks, nobody owns the output), nothing improves. This is usually the real bottleneck, not the tech.

Example: An AI model that predicts equipment failure a week in advance is useless if maintenance scheduling is locked to a rigid quarterly calendar that can't accommodate an early repair. Conversely, a sales team that can immediately act on an AI-flagged "this deal is at risk" alert — because reps already have autonomy to adjust their approach — gets real value from the same kind of prediction.

5. A feedback loop. Some mechanism to notice when it's working or not, and adjust. Without this, "readiness" decays — even a good setup goes stale.

Example: A recommendation engine that never learns whether suggested products were actually purchased will drift out of step with what customers want. A content moderation system with no channel for reviewers to flag missed violations or wrongful takedowns will keep making the same category of mistake indefinitely, even as the underlying content on the platform shifts.

6. Someone accountable for the outcome. Not the AI — a person or team whose job it is to own what happens when the tool is right, and especially when it's wrong.

Example: A company using AI to screen job applicants needs an HR owner who can explain and defend a rejection, not a shrug of "the algorithm decided." Compare that to a company where an AI-driven pricing tool sets prices with no one monitoring for discriminatory or erratic outcomes — when it eventually sets a price that damages customer trust or breaks a regulation, there's no one positioned to have caught it early or answer for it after the fact.

Step 3: Notice what falls out of this

- Infrastructure and headcount devoted to "AI" are not the fundamental unit — the **decision** is. A company with no dedicated AI team can be "ready" for a narrow, well-scoped task; a company with a huge AI budget can be completely unready if none of the above six hold.
- "Readiness" isn't binary or organization-wide. It's task-by-task. The right question is never "are we AI-ready" but "is *this* decision AI-ready."
- Most AI failures trace back to #3 or #4, not model quality — no verification loop, or no process to actually use the output.

Practical Takeaway

If assessing readiness for something specific, score the actual task against these six conditions honestly, rather than against a generic maturity checklist.