

INSIGHT PAPER 01

Operational AI Readiness

Why enterprise AI fails before implementation, and how leaders build the foundation for measurable value.

By Rachel Lane, Founder, CXaiS

The argument in one page

Boardrooms have stopped asking whether to invest in AI. They are asking how quickly they can show a return. Yet most enterprise AI work stalls before it reaches the people it was built for. The reason is rarely the model. It is whether the organisation was ready to run AI at scale before it started building.

This paper makes one argument, evidenced throughout. AI does not fail on technology. It fails on operational readiness. The organisations that succeed are not the ones with the best models. They are the ones that understood their workflows, their data, and their commercial case before they chose a platform.

It then sets out six dimensions that determine operational readiness, a maturity model leaders can use to place themselves honestly on the curve, and the sequence that turns readiness into commercial value. The framework can be applied independently. Nothing in the next ten pages requires a consultant to act on.



THE PROBLEM

Technology is not the constraint

The evidence on enterprise AI is unusually consistent, and unusually stark. In 2025, MIT's Project NANDA studied more than 300 deployments and found that 95 percent of enterprise generative AI pilots delivered no measurable impact on the profit-and-loss account. Only 5 percent created significant value. MIT was explicit about the cause. The divide was not driven by model quality. It was driven by approach.

Other researchers reach the same place by different routes. RAND found that more than 80 percent of AI projects fail, roughly twice the failure rate of conventional IT projects. S&P Global Market Intelligence found that 42 percent of companies abandoned most of their AI initiatives in 2025, up from 17 percent the year before. Gartner expects more than 40 percent of agentic AI projects to be cancelled by the end of 2027, on cost, weak value, and poor controls.

EXHIBIT 1 • THE FAILURE EVIDENCE

95%

of enterprise GenAI pilots delivered no measurable P&L impact.

MIT Project NANDA, 2025

80%

of AI projects fail, about twice the rate of non-AI IT projects.

RAND Corporation, 2024

42%

of companies abandoned most AI initiatives in 2025, up from 17%.

S&P Global Market Intelligence, 2025

Read those numbers alongside one fact. The technology has never been better. Models improve every quarter, budgets are approved, vendors are selected, and the work still does not reach production. When capability rises and outcomes do not, the explanation has to lie somewhere other than the capability.

It lies before go-live. It lies in the foundation the deployment was built on, and in whether anyone checked that foundation before committing the budget.

The winners fix the workflow first

McKinsey's own research points in the same direction. Only 7 percent of companies have fully scaled AI across their organisations, and more than two-thirds of high performers name data, not models, as the primary obstacle to enabling AI. In McKinsey's 2025 survey, the organisations reporting significant financial returns were about twice as likely to have redesigned their end-to-end workflows before they selected their modelling approach.

The winners fixed the workflow first, then chose the technology. The stalled majority did it the other way round.

The failure modes are predictable, and they repeat across sectors. Organisations automate an inefficient workflow, so AI performs the wrong work faster. They build on fragmented data, so a small flaw becomes a systemic one at machine speed. They deploy without clear ownership or governance, so no one is accountable when the numbers do not move. And they select a platform before defining the business outcome, so the technology arrives in search of a problem.

None of these is a technology failure. Each is a readiness failure. And each is visible before a single line of production code is written, provided someone looks in the right place.

Four assumptions worth challenging

Some of the most expensive AI decisions rest on assumptions that sound reasonable and do not survive contact with a real programme.

THE ASSUMPTION

WHAT WE OBSERVE

Better models produce better customer outcomes.

Better operating models do. A strong model on a broken workflow automates the breakage faster.

Data must be perfect before you start.

Data must be fit for the specific use case, governed, and traceable. Good enough is defined per workflow, not in the abstract.

AI is an IT project.

AI is an operating-model change that happens to use IT. Treated as an IT project, it stalls at adoption.

Vendor selection is the first decision.

It should be one of the last. Choose the platform after the workflow, the data, and the outcome are understood, not before.

THE FRAMEWORK

Six dimensions of operational readiness

Operational readiness is not a single measure. It resolves into six dimensions, each of which independently determines whether AI will perform in a given workflow. A programme can be strong on five and fail on the sixth. The value of assessing all six is that it surfaces the one that would otherwise have stalled the rollout after the money was spent.

EXHIBIT 2 · THE OPERATIONAL AI READINESS FRAMEWORK

01

Data Quality

Is the data fit, traceable, and governed for the decision it supports?

02

Process Maturity

Is the workflow understood as it actually runs, not as mapped?

03

Channel Architecture

Does context survive as the customer moves across channels?

04

Organisational Change

Is ownership named and the operating model ready to change?

05

Use-Case Prioritisation

Is the portfolio ranked on value against complexity, with a defensible sequence?

06

Commercial Alignment

Is the return expressed in the language the board already funds?

Dimension walk

01 Data Quality

Scope: Signal-to-noise, historical completeness, VoC reliability, integration accuracy.

Why it matters. AI amplifies whatever it is fed. On strong data it compounds value. On weak data it compounds error, and it does so faster than any manual process could.

Common failure pattern. Teams assume that because data is searchable it is usable. Searchability is not usability. A document can be correct in full and still produce a wrong answer if the wrong fragment is retrieved.

What good looks like. Data is traceable and fit for the specific decision it supports. The target is not perfect data. It is data that is good enough for this use case, governed, and defensible if an output is challenged.

The executive question. Can we trace an AI output back to the source that produced it?

02 Process Maturity

Scope: Service delivery model, knowledge management, first-call resolution, scalability.

Why it matters. Automating an immature process does not fix it. It industrialises it. The inefficiency scales with the technology.

Common failure pattern. The workflow that gets automated is the loudest one, not the one that creates value. Visibility is mistaken for priority.

What good looks like. The process is understood as it actually runs, including on its worst day, not as the process map claims it runs.

The executive question. Do we know how this workflow behaves under stress, not just in the demo?

03 Channel Architecture

Scope: Multi-channel analysis, friction points, unified journey mapping, channel prioritisation.

Why it matters. AI deployed into a fragmented channel estate inherits every disconnect in that estate. The customer experiences the seams, even when each channel works alone.

Common failure pattern. One channel is automated while the customer moves across five. Context is lost at every switch, and the AI looks worse than the humans it replaced.

What good looks like. The journey is mapped end to end before any single touchpoint is automated, so context survives the handoffs.

The executive question. Does our view of the customer survive a channel switch?

04 Organisational Change

Scope: Change-management maturity, leadership buy-in, skill gaps, communication readiness.

Why it matters. The best-designed deployment dies quietly if the people around it were not brought with it. Adoption, not accuracy, is where most rollouts are lost.

Common failure pattern. A technically sound system that no one uses, because ownership was never named and the operating model never changed.

What good looks like. Ownership is explicit, skills are in place, and the operating model evolves alongside the technology rather than after it.

The executive question. Who owns this workflow after go-live, and are they ready to own it?

05 Use-Case Prioritisation

Scope: Value-vs-complexity scoring, ROI framework, strategic alignment, technical feasibility.

Why it matters. Most programmes automate the use case that is most visible, not the one that returns the most. Sequence is where value is created or forfeited.

Common failure pattern. A portfolio chosen by enthusiasm rather than evidence, with no defensible reason for what comes first.

What good looks like. A ranked portfolio, scored on value against complexity, with a stated rationale for the sequence that a sceptical board can interrogate.

The executive question. Can we defend why this use case is first, and not third?

06 Commercial Alignment

Scope: Linking CX to P&L, revenue retention, customer lifetime value, acquisition economics.

Why it matters. Work that cannot be expressed in the board's language stays a cost line. It is the first thing cut when budgets tighten, regardless of its merit.

Common failure pattern. A business case that looks convincing on a slide and cannot be defended in a budget review, because the return was never modelled in P&L terms.

What good looks like. Commercial impact is expressed in the language the CFO already uses, before the build starts, not reconstructed afterwards to justify it.

The executive question. Can we state the return in numbers the finance function already trusts?

THE MATURITY MODEL

Placing yourself honestly on the curve

Readiness is not binary. It is a curve, and most enterprises are further down it than their AI ambition suggests. The following five levels describe what operational readiness looks like in practice. Leaders can use them to locate their programme before they commit to the next wave of spend.

EXHIBIT 3 · THE READINESS MATURITY MODEL

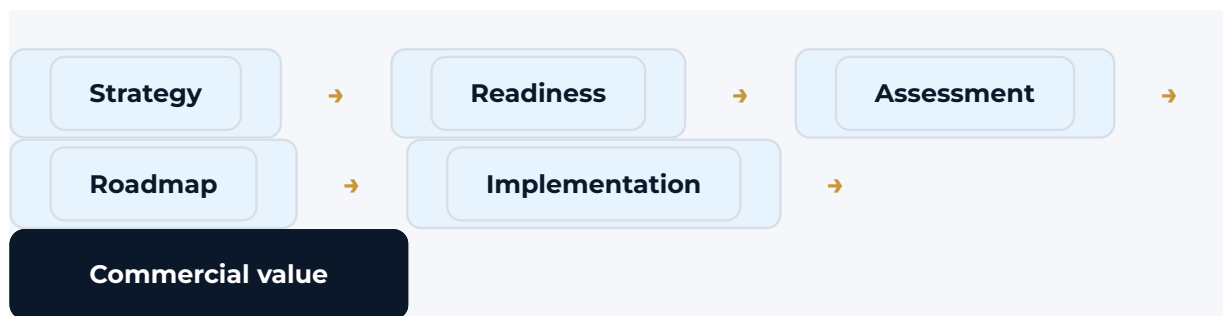
1	Explore	AI opportunities are being investigated, but there is no formal strategy, governance or readiness assessment.
2	Build	Foundations are being established. Initial pilots, governance and capabilities are emerging but remain inconsistent.
3	Operational	AI is delivering value in selected business areas with defined governance, repeatable processes and measurable outcomes.
4	Optimised	AI is embedded into business operations, continuously improved through data and governance, and delivering predictable value at scale.
5	Leading	AI is a strategic competitive advantage. The organisation innovates continuously, benchmarks externally and influences industry best practice.

Most enterprises sit at Explore or Build. The distance between Build and Operational is where value is won or lost, and it is an operational distance, not a technical one. No new model closes it. Only readiness does.

THE SEQUENCE

From readiness to value

The dimensions and the maturity model resolve into a single sequence. Each step depends on the one before it, and the order is not negotiable.



The point of the sequence is simple. Value sits downstream of readiness, and readiness sits upstream of technology. Run it in this order and AI has a foundation to perform on. Reverse it, and you rejoin the 95 percent.

APPLYING THE FRAMEWORK

From insight to decision

Leaders can use this framework independently. The six dimensions are a discipline any capable operator can run against their own programme, and the maturity model is a mirror any leadership team can hold up to itself. If reading this paper changes the order in which you make your next three AI decisions, it has done its job.

Where an organisation wants an independent read, this is the framework CXaiS applies. We assess readiness across the six dimensions, upstream of platform selection, and return a ranked portfolio with a clear verdict per workflow, expressed in commercial terms. The work is deliberately independent of any vendor, because the value of a readiness verdict depends entirely on it being honest. When we say a workflow is not ready for AI, that judgment is based on whether AI will perform there, not on whether anyone needs the deal to close.

See the value. Act on it.

To apply the framework to your own programme, start with an independent readiness read.

cxais.ai/contact

About and sources

ABOUT CXAIS

CXaiS is an independent AI in CX consultancy. We build the commercial foundation, ROI framework, and implementation plan that make AI investment in customer experience perform. We are practitioner-led, vendor-independent, and commercially accountable at every stage. We work upstream of deployment, on the readiness that determines whether the technology you buy will actually return.

ABOUT THIS SERIES

This is Insight Paper 01 from the CXaiS Institute. Forthcoming papers in the series: Workflow Cartography, on identifying the processes AI should automate first; Data Readiness, on building the foundation for enterprise AI; and The Value at Stake, on measuring the commercial impact of AI.

DISCLAIMER AND SOURCES

This paper is intended as an educational framework. Organisations can apply these principles independently or engage specialist support where appropriate.

Sources: MIT Project NANDA, The GenAI Divide: State of AI in Business (2025). RAND Corporation, on AI project failure rates (2024). S&P Global Market Intelligence, Voice of the Enterprise: AI and Machine Learning (2025). McKinsey and Company, AI data readiness: The key to scaling impact (2026), and McKinsey Global Survey on AI (2025). Gartner, on agentic AI project cancellations (2025). Figures are attributed to their original publishers and used for illustration.



cxais.ai · 2026