

PLAXONIC PERSPECTIVE · ENTERPRISE AI

Why 80% of Enterprise AI Projects Stall

Technical findings clearly explain why most implementations stall. This report examines these breakdown points alongside the specific blueprint followed by high-performing teams

Enterprise AI · Research Report



Contents



03	The Demo Works. The project Dies Anyway.
03	Why Enterprise AI Stalls
04	Reason 1: The Wrong Problem, Solved Well
04	Reason 2: The Data Was Never Ready
05	Reason 3: Chasing the Technology Instead of the Outcome
05	Reason 4: No Infrastructure to Reach Production
06	Reason 5 Asking AI to Do The Impossible
06	Reason 6: Treating a Capability Like a Project
07	Figures
08	Closing
09	Figures
10	Resources and References
11	About Plaxonic

The demo works. The project dies anyway.



There has never been more enthusiasm for enterprise AI, or more evidence that most of it does not work out. The headline figure is stark and well sourced: by RAND's estimate, **more than 80% of AI projects fail, roughly twice the failure rate of IT projects that do not involve AI**. The demo impresses. The pilot validates. And then the project quietly stalls somewhere between proof of concept and production.

It would be comforting to blame the technology, but the evidence points the other way. When RAND's researchers interviewed 65 experienced data scientists and engineers about why projects fail, **four of the five root causes they identified were organizational, not technical**, and 84% of industry interviewees pointed to leadership and problem-framing issues as the primary cause. The models are rarely the problem. What surrounds them is.

This report works through why enterprise AI stalls, grounded in the public research and in what we see building these systems. A note on the numbers up front, in the spirit of honesty this topic badly needs: the widely-quoted failure figures measure different things. We keep them separate throughout, and treat each source as authoritative only for what it actually measured.

WHY ENTERPRISE AI STALLS

- **The failure is organizational, not technical.** RAND found four of five root causes sit with leadership, problem framing, and data, not the algorithm. The AI usually works; the project around it does not.
- **The wrong problem is the most common failure.** Projects optimize a metric no one in the business actually needed, and succeed technically while failing entirely.
- **Data readiness decides the outcome before modeling begins.** The most-cited obstacle to enterprise AI is not talent or tooling. It is data that is siloed, low-quality, or untrusted.
- **"Pilot purgatory" is the norm, not the exception.** Validating in a controlled demo is easy. Surviving contact with production, at scale, with real users and real governance, is where most stall.
- **The meta-pattern is treating AI as an experiment to productionize later, instead of a production system from the start.**

The organizations in the successful minority are not the ones with the best models. They are the ones who framed the right problem, prepared the data, designed for production from day one, and treated AI as a capability to build rather than a tool to buy.

Reason 1 The wrong problem, solved well

The single most common reason enterprise AI stalls is also the least technical: the project set out to solve the wrong problem, or never defined the right one clearly enough. RAND found this at the top of its list, [business stakeholders often misunderstand or miscommunicate what problem needs to be solved, and teams deploy models optimized for the wrong metric or that do not fit the actual workflow](#).

This failure is insidious because it looks like success right up until it doesn't. The data-science team hits its accuracy target, the model performs beautifully in evaluation, and then it lands in the business and no one uses it, because it optimized something no one actually needed. The project was well executed and pointed in the wrong direction.

PLAXONIC'S VIEW

Start every AI initiative with a one-sentence problem statement and a business metric it must move, agreed by the people who own that outcome, before a single model is trained. The most valuable question at the start of an AI project is not "what can the model do?" but "what decision or workflow will be measurably better, and how will we know?" A model that optimizes the wrong number is not a partial success. It is a complete failure that took months to become visible.

Placeholder for a real Plaxonic example: an engagement where sharpening the problem definition up front (or inheriting one that wasn't) changed the outcome. Add a real, de-identified account here.

Reason 2 The data was never ready

If the wrong problem is the most common failure, unready data is the most fundamental. AI models are only as good as what they learn from, and enterprise data is routinely siloed, inconsistent, incomplete, or simply not trusted. RAND named the lack of adequate data as a core root cause; broader surveys put it at the very top of the obstacle list.

The evidence is consistent across studies. [Informatica's 2025 CDO survey found data quality and readiness cited as the leading obstacle to AI success, ahead of technical maturity and skills](#), and Gartner has forecast that a majority of AI projects will be held back specifically by inadequate, AI-unready data foundations. Notice what is not on those lists: the algorithm.

PLAXONIC'S VIEW

Audit the data before you build the model, and be willing to make the first phase a data project. The unglamorous work of connecting, cleaning, governing, and making data trustworthy is not a prerequisite to the AI work; it substantially is the AI work. Organizations that skip it are not moving faster, they are simply moving the failure to a later, more expensive stage, when a model trained on broken data produces results no one can trust.

Reason 3 Chasing the technology instead of the outcome

A third pattern is a bias toward the newest, most impressive technology rather than the approach that best fits the problem. RAND identified this directly, [projects fail when the organization focuses more on using the latest and greatest technology than on solving real problems for its intended users](#). The hype cycle makes this worse every year, as each new capability arrives promising to make the hard parts disappear.

It rarely does. Chasing sophistication inflates expectations, complicates delivery, and often applies a heavyweight solution to a problem a simpler approach would have solved more reliably. The excitement is real; the fit is not.

PLAXONIC'S VIEW

Choose the simplest approach that solves the problem, not the most advanced one that impresses the room. A boring, reliable solution in production beats a state-of-the-art one that never ships, every time. The maturity to say "you do not need the most advanced model here" is often worth more to a client than the ability to build it. The goal is a solved problem, not a showcase.

Reason 4 No infrastructure to reach production

Even when the problem is right, the data is ready, and the approach fits, projects stall at the last and hardest step: getting to production and staying there. RAND's fourth root cause is exactly this, [organizations lacking the infrastructure to manage their data and deploy completed models](#). The prototype that worked for one user in a notebook meets the reality of many users, real latency budgets, monitoring, and cost.

This is the heart of "pilot purgatory," and it is widely measured. [Gartner has projected that at least 30% of generative-AI projects will be abandoned after proof of concept](#), and MIT's 2025 research found that [roughly 95% of enterprise generative-AI pilots produced no measurable profit-and-loss impact](#). Those figures measure different things than RAND's 80%, but they all point at the same cliff: the drop between a working demo and a running system.

PLAXONIC'S VIEW

Design for production from the first line of code, not as a phase you bolt on later. The reason so many pilots die at the deployment cliff is that they were never built to cross it, the prototype code became the production code by default. Locking down the deployment realities early, latency, cost per call, throughput, an accuracy floor, monitoring, and recovery, is the difference between an AI product that ships and one that quietly bleeds effort. This is the core of how we work: production-ready from the start.

Reason 5 Asking AI to do the impossible

The fifth of RAND's root causes is the most honest one, and the one vendors mention least: sometimes AI fails because [the technology is applied to a problem that is simply too difficult for it to solve](#). AI is not a magic wand. Even the most advanced models cannot automate away every hard task, and pretending otherwise sets a project up to fail before it begins.

This is distinct from chasing shiny technology. Here the problem may be well defined and the data adequate, and the honest answer is still that the current state of the art cannot deliver what was promised. Recognizing that early saves months and budgets; discovering it late is one of the most expensive lessons in enterprise AI.

PLAXONIC'S VIEW

Part of good AI judgment is knowing what AI cannot yet do, and saying so. It is tempting, in a competitive market, to promise that AI can solve anything, but the more valuable partner is the one who will tell you when a problem exceeds what the technology can reliably deliver today, and propose a narrower version that will actually work. Honesty about limitations is not a weakness in an AI partner. It is the thing that separates the successful projects from the cautionary tales.

Reason 6 Treating a capability like a project

Underneath RAND's five causes sits a meta-pattern that practitioners describe again and again: [AI projects fail because they are treated as experiments to be productionized later, rather than as production systems from the start](#). The one-off project mindset, build it, launch it, move on, is a poor fit for AI, which needs ongoing data, monitoring, retraining, and ownership to keep delivering value after launch.

The organizations that succeed treat AI as a durable capability: the data discipline, the deployment infrastructure, the governance, and the people who can operate and improve the system over time. The ones that stall keep launching projects and wondering why the value does not last.

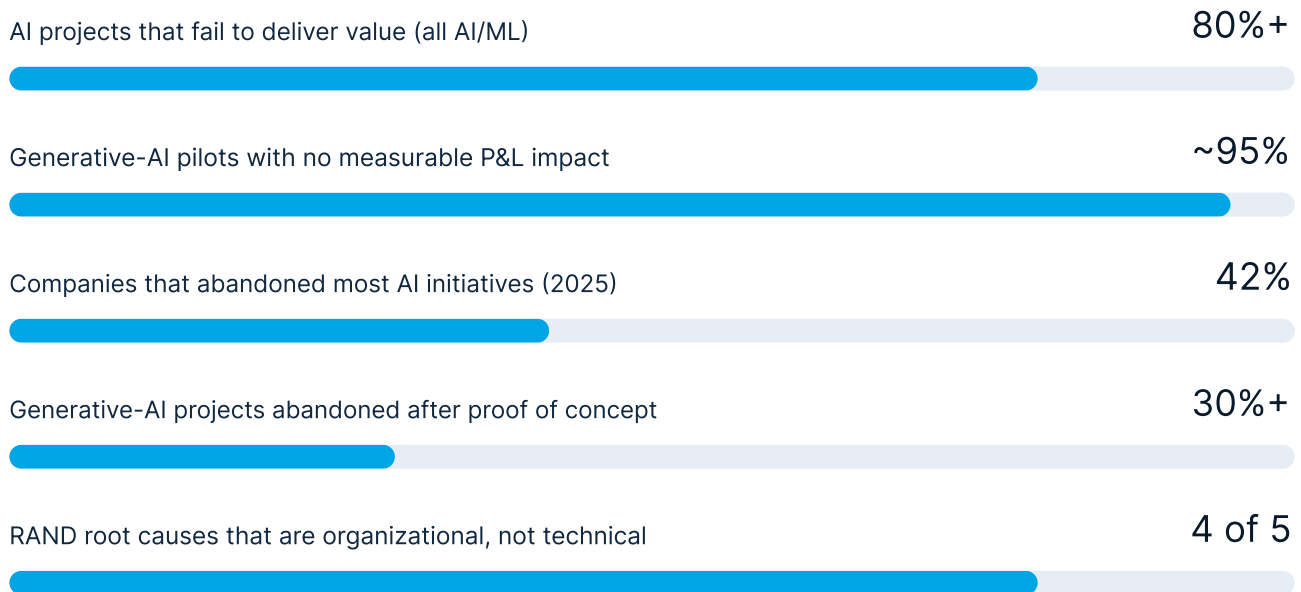
PLAXONIC'S VIEW

Build the capability, not just the model. The highest-leverage thing an enterprise can do is stop treating AI as a series of experiments and start building the durable foundation, data, deployment, governance, and skills, that lets every future initiative succeed faster. A good AI partner should leave you more capable of doing the next one yourselves, not more dependent on them for it. That is the difference between an AI demo and an AI capability that compounds.

The failure rates, kept honest



These figures are quoted together everywhere, and they measure different things. We keep them separate. They are not contradictory; they have different denominators, measured in different years.



Sources, each measuring a different thing: RAND, The Root Causes of Failure for AI Projects, 2024 (80%+ of all AI projects fail; four of five root causes organizational). MIT Project NANDA, The GenAI Divide, 2025 (~95% of generative-AI pilots no measurable P&L impact). S&P Global Market Intelligence, 2025 (42% abandoned most initiatives, up from 17% in 2024). Gartner, 2024 (30%+ of generative-AI projects abandoned after PoC). See References. These are not directly comparable with one another.

HOW TO READ THESE NUMBERS

- The 80% (RAND) covers all AI and machine-learning projects, and is qualitative, best read as "the large majority fail," not a precise rate.
- The 95% (MIT) is narrower and stricter: generative-AI pilots specifically, measured against P&L impact.
- The 42% and 30% measure abandonment, a different outcome again, over different periods.
- Quoting them as if they were one number is exactly the kind of imprecision this topic suffers from. The honest summary: most enterprise AI does not reach durable value, and the reasons are consistent.

Before you start: six questions that predict success



Almost every failure in this report traces to a question that was not asked, or not answered honestly, before the work began. Here is the sequence we run with clients before committing to an enterprise AI initiative.

- 1 What business outcome must this move, and who owns it? If you cannot state the problem in one sentence with a metric and a named owner, you are not ready to build. This alone prevents the most common failure.
- 2 Is the data real, reachable, and trusted? Audit it before modeling. If the honest answer is no, the first project is a data project, and that is progress, not delay.
- 3 Is this the simplest approach that works? Resist the pull of the most advanced option. Match the method to the problem, not to the hype cycle.
- 4 Have we designed for production from day one? Lock the deployment realities, latency, cost, throughput, an accuracy floor, monitoring, and recovery, up front, not after the pilot.
- 5 Is this problem actually solvable with today's AI? Be honest about the limits. A narrower problem that works beats an ambitious one that cannot.
- 6 Are we building a capability or buying a tool? Plan for the data, governance, and people who keep the system delivering after launch. AI that is not maintained does not stay valuable.



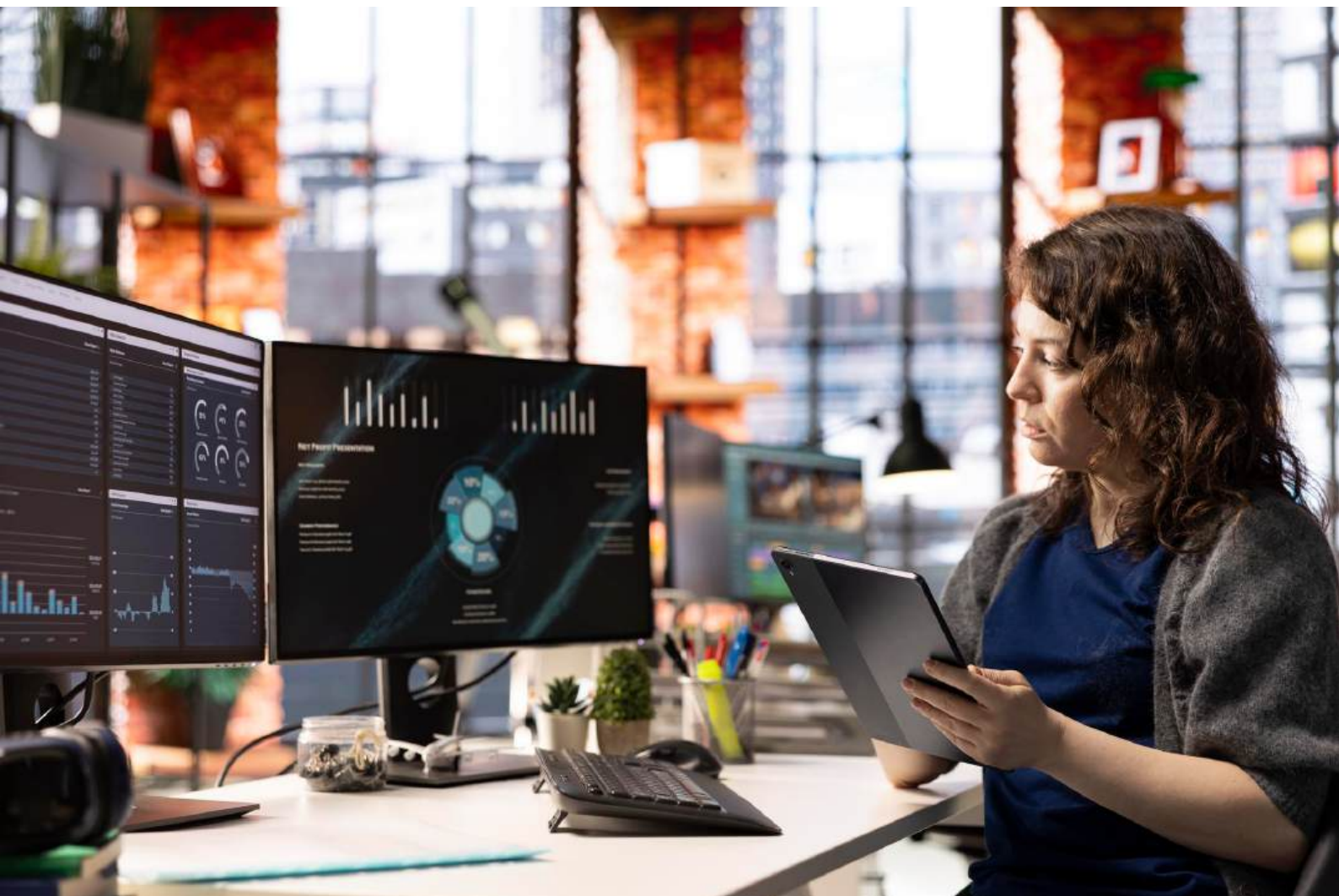
None of these six questions is about the model. That is the point. The research is unambiguous that enterprise AI succeeds or fails on the decisions made around the technology, not the technology itself, which is precisely why the failure rate has stayed high even as the models have become extraordinary.





IN CLOSING

The models have never been more capable, and most enterprise AI still stalls. That is not a paradox. It is the clearest possible signal that the bottleneck was never the technology. It is the problem you chose, the data underneath it, the honesty about what AI can do, and the discipline to build for production rather than for a demo. The 80% who stall are not short of talent or tools. They are short of those four things. The minority who succeed treat AI as something you build the capability to do well, again and again, not something you buy once and hope works. That is the whole difference, and it is entirely within reach.



Resources and references

Every quantitative figure in this report is traced below to a named third-party source. The failure-rate figures measure different populations, outcomes, and years, and are not directly comparable. Original sources should be treated as authoritative.



- 1 RAND Corporation, "The Root Causes of Failure for Artificial Intelligence Projects and How They Can Succeed" (RR-A2680-1), 2024. Based on interviews with 65 data scientists and engineers with 5+ years of experience. Source for the 80%+ failure figure (roughly twice the non-AI IT rate), the five root causes, and the finding that the majority of causes are organizational rather than technical. Cited throughout, and in Reasons 1–5.
- 2 MIT Project NANDA, "The GenAI Divide: State of AI in Business 2025," 2025. Reviewed 300+ public AI initiatives with interviews and surveys. Source for the finding that ~95% of enterprise generative-AI pilots produced no measurable P&L impact. Cited in Reason 4 and Figures.
- 3 S&P Global Market Intelligence, 2025 survey. Source for 42% of companies abandoning most AI initiatives in 2025, up from 17% the prior year. Cited in Figures.
- 4 Gartner, AI predictions, 2024. Source for the projection that at least 30% of generative-AI projects will be abandoned after proof of concept. Cited in Reason 4 and Figures.
- 5 Informatica, 2025 CDO Insights survey. Source for data quality and readiness being the leading cited obstacle to AI success, ahead of technical maturity and skills. Cited in Reason 2.
- 6 Practitioner analyses of RAND's findings (multiple, 2024–2026), including the observation that ~84% of industry interviewees cited leadership and problem-framing issues, and the meta-pattern that AI projects fail when treated as experiments to be productionized later rather than as production systems from the start. Cited in the introduction and Reason 6.

About Plaxonic



Plaxonic is a global AI-first engineering and product company. We help organizations turn ambition into working systems: intelligent, cloud-native, and automation-driven software built to perform in production, not just in a demo.

Our work spans the full journey of a modern digital system. We build production-ready AI that makes it into daily operations, engineer enterprise software that businesses depend on, and design scalable cloud platforms that hold up under real demand. Where legacy systems hold companies back, we modernize them.

Where operations are slow or manual, we streamline and automate them. With globally distributed teams, we bring together the engineering depth and range that complex problems require, working alongside our clients rather than at arm's length. The result is technology that is built for scale and measured by outcomes: systems that run reliably, operations that cost less to run, and digital products that move the business forward.

We build what actually ships, and we make sure it delivers once it does.



WEB www.plaxonic.com LINKEDIN linkedin.com/company/plaxonic

Disclaimer: The material in this document reflects information available at the time it was prepared, as indicated by the date on the front page. Conditions across enterprise AI are changing rapidly and the position may change. This content is provided for general information purposes only, does not take account of any reader's specific circumstances, and is not intended to be used in place of consultation with a qualified professional adviser. Plaxonic disclaims, to the fullest extent permitted by applicable law, any and all liability for the accuracy and completeness of the information in this document and for any acts or omissions made on the basis of such information.

Figures cited in this document are drawn from named third-party sources listed in the References section and have not been independently verified by Plaxonic. They measure different populations and outcomes across different years and are not directly comparable. A number of them are analyst findings or forecasts rather than measurements of current conditions, and forecasts are frequently revised. Sections marked "Plaxonic's View" are opinion, offered as perspective rather than as findings.

This document refers to marks owned by third parties. All such third-party marks are the property of their respective owners. No sponsorship, endorsement, or approval of this content by the owners of such marks is intended, expressed, or implied.

Copyright © 2026 Plaxonic. All rights reserved.