



BURAK GALİBA
Thoughts on AI Transformation

ESSAY · ENTERPRISE AI

AI in Name Only — How Enterprises Pick Flagship AI Projects That Aren't AI

Enterprise AI portfolios fail before day one — at selection, where every incentive rewards relabeling analytics as AI. Here's the test, and the fix.

AI STRATEGY

GOVERNANCE

LEADERSHIP

Burak Galiba · July 30, 2026 · 6 min read

<https://burakgaliba.com/blog/ai-in-name-only>



SCAN TO READ ONLINE

[68% of S&P 500 earnings calls cited "AI" in Q4 2025](#) — a ten-year high. And in BCG's latest survey, [61% of CEOs say their boards are rushing AI transformation](#).

Inside a large enterprise, those two numbers produce a committee. Business units keep running small pilots on their own products, and a central body selects a handful of flagship "AI projects" that get the real IT resources. That selection meeting is where the portfolio's fate is decided — and here's the contrarian read: the famous AI failure statistics are partly measuring projects that were never AI to begin with. The portfolio wasn't mismanaged. It was mislabeled, before day one, at selection.

I've written about [killing pilots](#) and [escaping pilot purgatory](#). Both are downstream problems. This one is upstream: walk into a typical flagship portfolio review and a majority of the "AI projects" on the slide are BI dashboards, data-warehouse builds, or plain workflow automation wearing an AI badge.

The selection committee is a market

Nobody in that room is lying. Everyone is responding rationally to their incentives — which is exactly why the outcome is so reliable.

- **Top management is buying keywords.** The board wants AI, the holding company wants AI, and executives pattern-match on terms from the reports they've read. What they actually want is savings; "AI" is the word the pressure arrives in.
- **IT leadership is selling certainty.** Asked to showcase that it can deliver AI, IT nominates the projects it already knows how to build — which are analytics and dashboard projects. A guaranteed delivery beats an honest experiment in every steering committee ever convened.
- **Individual IT leads are buying career polish.** "Led the enterprise's flagship AI initiative" is a sentence with a market value, and it reads the same on a CV whether the system contained a model or a pivot table.

The committee is a market, and the currency is the label, not the outcome. This isn't a competence problem — it's an incentive problem, and it matches what RAND found when it autopsied failed AI projects: [the leading root cause is stakeholders misunderstanding or miscommunicating what problem needs solving](#), followed closely by choosing projects for the shine of the technology rather than the problem.

And the label genuinely pays. Back in 2019, MMC Ventures found that [around 40% of European "AI startups" showed no evidence of using AI](#) — and that AI-labeled firms raised 15–50% more capital. Gartner now calls the vendor version "**agent washing**": of thousands of self-declared agentic AI vendors, [it estimates only about 130 are the real thing](#). The Guardian reports PR executives sending [roughly half of their AI press releases under duress](#), describing the claims as "Bikram yoga-level stretches." The same rebranding happens inside enterprises, one steering committee at a time. There's no published statistic for what share of internal enterprise "AI projects" contain no machine learning at all — nobody measures it, because nobody at the table benefits from measuring it.

A dashboard is a rear-view mirror

The test for the label is simpler than the committee makes it look. Ask one question of the proposed system: **does it make probabilistic predictions, decisions, or generations that change an action — or does it present historical data for a human to decide?**

If you can specify the output exactly in advance, it's software. Good software, possibly. Not AI. A dashboard is a rear-view mirror; a model is a steering input. Both are useful, but only one of them turns the car.

Apply the test to a typical flagship portfolio and watch it reclassify itself. The "AI-powered invoice matching" that's deterministic rules. The "intelligent claims platform" that's a workflow engine with a queue. The "AI forecasting initiative" that's a data warehouse and a set of charts a human squints at every Monday. Each one arrived at the committee with the word "intelligent" in the deck, which is the corporate equivalent of racing stripes.

The mislabeled projects are usually good projects

Here's the part the cynics get wrong: most of these relabeled projects deserve to exist. Data pipelines, semantic layers, clean dashboards — this is exactly the foundation real AI needs. Gartner predicts organizations will [abandon 60% of AI projects unsupported by AI-ready data](#), and reports 63% of organizations don't have AI-ready data practices. The enterprise that builds its data layer first is doing the right work.

The harm isn't the work. It's the label. Three things break when a data project ships under an AI flag:

First, it burns the AI budget's credibility — the flagship consumed the money and political capital reserved for AI, so the actual AI candidates queue behind it. Second, it sets P&L expectations the project was never designed to meet; a semantic layer doesn't produce the savings the board was promised, because it was never going to. Third — and worst — when it ships and the result is a dashboard, leadership concludes "AI doesn't deliver." The NANDA report's famous finding that [95% of enterprise GenAI pilots produce no measurable P&L impact](#) gets quoted as proof that AI fails. Some of that 95% never contained AI to fail. **Mislabeling is how you manufacture your own 95%** — and then cite it as the reason to stop.

Steve Blank named the genus years ago: [innovation theater](#), activity that signals transformation without changing what the company does. The flagship AI portfolio is its most expensive current production.

How to fix the selection

The fix isn't smarter committee members. It's changing what the committee's currency buys.

1. **Label honestly — and run two tracks proudly.** A funded "data & analytics" track next to a smaller, real "AI" track. The data projects stop needing the costume the moment they don't have to compete for the AI budget to survive.

2. **Require a written value hypothesis plus one paragraph titled "what makes this AI."** If the paragraph can't name the prediction, decision, or generation the system makes, the project goes to the other track — with full funding and zero shame.
3. **Put a business owner's metric on it, not IT's.** Cycle time, loss rate, cost per claim — a number that lives in someone's P&L. "Platform delivered on schedule" is an IT metric, and IT metrics are how relabeled projects declare victory.
4. **Select for measurability, not demo appeal.** The best flagship is usually a high-volume back-office decision with a baseline, not the use case that looks best on stage at the town hall.
5. **Reserve the right to reclassify.** A standing rule that any "AI project" can be relabeled mid-flight — keeping its budget but leaving the AI portfolio's scoreboard. Reclassification without punishment is what makes honest labels cheap.

None of this slows anyone down. It just makes the label stop paying more than the outcome — at which point the market at the table reprices itself.

Where to start

This week, ask one question of your flagship portfolio: *which of these projects could ship, exactly as specified, without a model anywhere in it?* Every project that passes is a good software project wearing the wrong badge — and everyone in the room quietly knows which ones they are.

If you want to ask it with more rigor, I built a free [AI Label Test](#) that runs one flagship project through 12 questions on exactly the two axes that matter: does the system genuinely make AI-shaped calls, and does it have a value case either way? The verdict is one of three — AI project, good software wearing the wrong badge, or keyword project — and it doesn't care what the project is called on the slide. It asks what the system decides, who owns the number, and whether the committee would still fund it without the label — which, it turns out, is everything the selection meeting forgot to ask.

Read the essay online at <https://burakgaliba.com/blog/ai-in-name-only>

I built a free AI Readiness Score — 20 questions across pilots, data, talent and governance, about ten minutes, no email required: <https://burakgaliba.com/readiness>

© 2026 Burak Galiba · <https://burakgaliba.com> · Views are my own.