

What You're Getting Wrong About AI Enablement

FIVE MISTAKES AND THE SEQUENCE THAT FIXES THEM | TL PARTNERS | FIELD NOTES 2025-2026; THIS WRITE-UP: JULY 2026 | THOMAS IGEME

I have spent the better part of two years inside enterprise AI enablement: the manager and leader programs at a publicly traded technology company rebuilding its management corps mid-restructuring, and at a late-stage consumer company whose engineers had already handed most of their coding to AI agents. Around them, advisory work across roughly fifteen PE-backed software companies and a running exchange with the CHROs doing this work elsewhere. Different companies, different tools, the same five mistakes.

The public numbers cannot even agree on how badly it is going. MIT's much-quoted study found 95 percent of enterprise GenAI pilots never touched the P&L; Wharton found three-quarters of enterprises reporting positive ROI the same season. Same world, measured differently; hold that thought. What nobody disputes: the median firm spends eleven dollars per employee per month on AI. Most companies have not seriously tried yet, and the ones that try still fail for reasons that are not technology problems and, less obviously, not skills problems.

MISTAKE ONE: YOU MADE AI THE STRATEGY

Half the executive teams I talk to now run a strategy deck with AI as the headline. If your investment thesis underwrites an AI-driven margin structure, fine: AI is in the strategy, priced and dated. Most companies mean something softer, "be an AI-first company," and people hear a vibe, not a direction. Vibes do not survive contact with a Tuesday morning. I watched one leadership team ask everyone for their most creative selves when what they wanted was 150 creative ways to move in one direction. Say the direction.

The fix is almost embarrassingly simple, and I have watched companies stay mired for months for lack of it. Find the enablement anchor: take the strategy you already have, line by line, and ask what AI changes about each line. What does it accelerate? What does it make cheap enough to finally do? And, the question nobody asks, what work does it expose as not worth doing at all? One page: every strategy line marked accelerated, newly affordable, or stopped, each with a named workflow and owner. De-prioritization is the first dividend of AI, available before a single workflow is rebuilt. Skip it and you end up where I found one client this spring: the AI-first operating model asserted rather than defined, the slogan set, the day-to-day work unchanged. An announcement, not a transformation.

MISTAKE TWO: YOU ARE TREATING ADOPTION AS A SKILLS PROBLEM

The budget tells the truth first. The median firm barely spends; the firms that do put nearly everything into models, platforms, and licenses, a split usually decided before the People team sees it. BCG's ratio has not changed in two years: AI value is 10 percent algorithms, 20 percent technology, 70 percent people and process. Budgets still run the other way.

Even the people budget buys the wrong thing, because skill is the barrier everyone diagnoses. It is the comfortable one: a skills gap implies a course, and a course can be procured. In the rooms I run, skill is the smallest of three barriers. The second is capacity and incentives: a support rep with forty tickets a day has no slack to experiment, and only 13 percent of people say they are rewarded for reinventing how work gets done when results are not guaranteed. Nobody redesigns their own job on unpaid time. The third, the one nobody names, is identity, and it binds hardest where accuracy is the job: finance, legal, operations, the senior ranks. The thing that makes you great with AI, the willingness to experiment and get it wrong, is the opposite of what made you great at your job. Where accuracy is the standard, being wrong does not just feel uncomfortable. It feels like failing. Underneath sits the question people rarely say out loud: if AI does that part, what am I for? People do not experience workflow change as workflow change. They experience it as identity change.

Courses can teach mechanics. They do not move identity. Three things do. A picture of excellence in the person's own function, produced live, on their own work. Safety to admit a gap, which only leaders create by going first: the general counsel automating part of his own workload in front of everyone beats any curriculum, and so does the CHRO who opens by admitting she had not touched the tools a year ago. And an accessible start: real work finished in the room, the first rep taken with a coach present. The data shows the same pattern: at companies pulling ahead, 88 percent of managers role-model AI daily; at laggards, 25 percent. Your people are not watching your enablement calendar. They are watching you.

MISTAKE THREE: YOU ARE ENABLING INDIVIDUALS AND LEAVING THE OPERATING LAYER ALONE

The default playbook, licenses plus prompt training, produces impressive individuals inside an unchanged company. An SVP I work with delivers, by her own account, three months of work in two and a half weeks. A self-report, the same instrument I am about to tell you to distrust, and I believe her anyway; both can be true, which is the problem with anecdotes. Individual lift is real, and it does not sum. Microsoft's 2026 research found organizational factors, culture, manager support, talent practices, account for more than twice the reported AI impact of individual factors. McKinsey finds the same shape from the other direction: the strongest predictor of AI reaching the P&L is whether workflows were fundamentally redesigned.

Both findings point at the operating layer: your managers and the machinery they run. It is the layer enablement skips, and it is holding the new unpriced work: checking AI-assisted output, coaching the transition, answering performance questions nobody has answered for them. At one client, the executive team hit it mid-calibration: how do you evaluate someone who collaborates with AI on everything? What is them and what is the model? If someone gets faster and simply stops, is that good performance? Nobody knew. That is not a training gap. That is an operating model that has not caught up with its own tools. You cannot hand managers that load without taking something off their plate; spend the de-prioritization dividend there first.

Everyone in your company is becoming an editor of work they did not author, and editing is a skill you have never hired for or trained. The old quality signals died overnight; a well-put-together deck used to tell you how many hours of thinking went in, and now it tells you nothing. The editor has to know what the author does not: what the data says, what is missing, where instinct disagrees. The people who get lost are the ones who simply produce more. And the redesign unit is the role, not the task: every leader runs their own job through the tools and returns with a picture of the future version, then repeats it for each role on the team.

The 2027 version of this mistake is already here. At the consumer company, engineering had handed most coding to agents, and the live questions were about none of this year's training: who reviews agent output, at what sampling rate; who edits the agents; what a span of control means when a growing share of a team's throughput is machine-produced. If your program teaches copilot chat while your 2027 budget buys agents, you are training for the wrong decade. The operating layer is where the agent boundary per function gets set, published, and revised. No tool sets it for you.

MISTAKE FOUR: YOU ARE MEASURING ACTIVITY AND CALLING IT EVIDENCE

Seats provisioned, weekly active users, self-reported hours saved: the standard dashboard, and it is theater. One CHRO whose program I rate highly told me plainly that the hours-saved numbers in her reporting were invented. Save the judgment for the ask: her board demanded a number that does not exist, and an invented number is what that demand produces, everywhere. The best evidence on self-report is a 2025 randomized trial of experienced developers: 19 percent slower with AI, believing they were 20 percent faster. It is why MIT finds 95 percent failure while Wharton finds three-quarters success. Nobody is measuring the same thing. You must decide what counts as evidence in your company.

What counts is small, named, and dated. Every leader leaves with a one-page experiment: the workflow being reimaged, the blocker being removed, what they are willing to be imperfect at, and the signal expected within thirty days. Commitments live where work lives: a goal in the performance system, a session on two calendars. No real performance

system? Then the QBR document and the board deck are the system. A commitment that is not in a system somebody checks is a preference. Follow-through is peer pairs at thirty and ninety days, with the People team owning the rhythm: a peer expecting to hear how it went, not an audit.

Proficiency is an observable working level: what the AI actually is in a person's work. Cleanup tool. Research assistant. Thought partner. Co-worker. The tells sit at the extremes: at the bottom, they could stop tomorrow and nothing would change; at the top, the question has flipped from "should I use AI for this" to "how would I do this without it." Managers read levels in the work, against anchors written per function, because level three in legal is not level three in engineering, and anything earnable by self-report will be gamed the day it touches a review. The distribution by function, quarter over quarter, is the leading indicator your dashboard is missing. It compounds, too: six months of sustained use measurably outperforms the first month, and adoption jumps only past roughly five hours of coached practice. Enablement is a curve, not an event. Budget for the curve.

What to show your board: five numbers survive scrutiny. Named workflows redesigned and running, a count short enough to list by hand. Cycle time or cost on those workflows, measured by the system that runs them, not by survey. The working-level distribution, moving right. Experiments shipped and experiments killed, both dated, because a portfolio with no kills is not being measured. And the share of managers carrying an AI outcome goal. Weekly actives can stay as a floor. It cannot be the story.

MISTAKE FIVE: YOU ARE PRETENDING THE ROOM IS NOT SCARED

AI stopped being neutral inside your company the day it became the public rationale for efficiency. In every post-restructuring room I have run, the exercise of sorting work into own, augment, automate lands on someone thinking: am I planning the next reduction? If the message was "adopt AI and you will be safe," and people were not safe, that credibility is spent, and no framework survives spending it.

The opposite of pretending: name the fears before the frameworks. Collect what keeps the room up at night, distill the themes on the spot, and have the most senior person present speak to them for ten unscripted minutes. That is the opening move, not the repair. Give the business context from the top, because when executives withhold it and hand the People team the session anyway, they have assigned an impossible brief. People call the result a lecture from HR. The lecture was written upstairs. And when you change the operating model, call it an experiment with a review date, because that is what it is.

Mandates deserve a cleaner verdict than they get. Shopify's 2025 memo, reflexive AI use as a baseline expectation, worked; Duolingo's near-identical message, led with a public announcement, bought a backlash and a walk-back. The difference was not the mandate. It was years of modeled behavior underneath it, enablement already in place, a message that ran internal-first. So split the mandate: standardize the toolset, mandate the outcome, invite the practice. Pressure without enablement has a measurable product: 40 percent of employees now report receiving workstop, output polished enough to pass and hollow enough to push the thinking onto the recipient. Meanwhile two-thirds of professionals admit using AI against policy, senior leaders worst of all. The demand exists. What is missing is permission with standards.

If the announcement already happened and the trust is already spent, the path back is the same sequence run humbly: re-anchor to the strategy in public, kill the surviving work with no remaining consumer so relief arrives as fact rather than promise, leaders take their first visible reps, one honest story about where freed capacity goes, measurement small enough to be believed, and a first goal cycle that is developmental rather than evaluative. The recovery is not a different program. It is the same program minus the theater. The work-structure half, the inventory and the nobody column, is worked as a sample in our AI reset brief.

THE SEQUENCE THAT WORKS

Clarity first, because exposure without direction is frightening. The executive team anchors AI to the strategy, decides what stops, and settles the story about freed capacity, growth or cuts, before anyone below them is asked to change. That story is a decision, not a communication. It is the WHAT, the value hypothesis everything else is scored against; enablement is a HOW intervention, and no amount of HOW rescues a missing WHAT.

Leaders second, rehearsing the narrative, not the plan. The last exercise I run with any leadership team is in the first person: what will you say to your team on Tuesday? Not what you will do. What you will say. If no executive will go first, borrow the most credible operator you have. If there is truly no one, stop: enablement against a leadership team that will not touch the tools is theater, and the room knows before the sponsors do.

Teams third: real work in the room, out with the thirty-day brief. Then the rhythm, owned. The People function runs the operating layer; the CIO owns the toolchain and the data boundary; where no real CHRO exists, the CEO names an owner, because an unowned rhythm is dead by October. Run it on the business calendar, not as an initiative with a launch week. For investors, one addition: treat post-reduction savings claims as unproven until a system of record says otherwise, and put the anchor exercise and the first cohort inside the 100-day plan.

WHAT THIS COSTS YOU TO IGNORE

Run the counterfactual on your own company. The spend is not the cost. The cost is the year: a stalled program while competitors compound down the learning curve, a workslop tax priced near two hundred dollars per employee per month, decisions moving at the speed of unread drafts, and the quiet exit of exactly the people you need; three-quarters say they would move for an employer that develops these skills. The people most worth keeping are the ones with options.

Five questions tell you whether the year is already leaking. Which line of our strategy does the AI work anchor to, and what did we stop doing? Who on the leadership team took a visible first rep this quarter? Which of our managers' questions about AI-assisted performance have answers? Show me one workflow where the before and after is provable without a survey. And what story did we tell about freed capacity, and was it true?

Across the two enterprise programs behind this paper, more than 80 percent of participating managers reached target proficiency on the working levels. The rest of the results are shared in client rooms, which is where results belong. The instruments exist: the anchor map, the thirty-day brief, the level rubric, the board page. We present the full approach, with the artifacts, at executive team and fund partner meetings: fifteen minutes of material, your questions after.

NOTES

Field material is drawn from TL Partners engagements and practitioner conversations, 2025-2026, anonymized; program results are shared at engagement level. External figures: MIT NANDA, The GenAI Divide: State of AI in Business (Aug 2025); Wharton GenAI Lab, Accountable Acceleration (Oct 2025); Ramp AI Index (Jun 2026); BCG, Closing the AI Impact Gap (Jan 2025) and AI at Work (Jun 2025), source of the five-hour coached-practice threshold; Microsoft, 2026 Work Trend Index (May 2026), which reports organizational factors accounting for more than twice the AI impact of individual factors (67 vs 32 percent) and the 13 percent rewarded-for-reinvention figure; McKinsey, The State of AI (Nov 2025); METR randomized developer trial (Jul 2025); BetterUp Labs and Stanford Social Media Lab on AI workslop, via Harvard Business Review (Sep 2025, Jan 2026), source of the 40 percent and per-employee cost figures; PagerDuty shadow-AI survey (Jun 2026), source of the policy-violation and would-move figures; Anthropic Economic Index (Mar 2026), source of the six-month learning-curve finding; manager role-modeling figures (88 vs 25 percent) from BCG, AI Transformation Is a Workforce Transformation (Feb 2026).