

Navigating AI Without a Map

A talk for technology CEOs · ~15 minutes

Uncertainty changes the decision framework

There is something I want to say at the outset that you may not hear very often, and it may be the most important thing I can tell you.

Nobody knows where this technology is going.

I don't mean that as a hedge or a disclaimer. I mean it as a precise, literal observation. The researchers don't know. The people spending fifty billion dollars a year on computers don't know. I have sat in rooms with some of the most technically sophisticated people in our industry, and the honest answer — once you strip away the confident projections and the fundraising narratives — is that we are watching a technology evolve in real time with no reliable map. The models that exist today are orders of magnitude more capable than what existed three years ago. What will exist three years from now is truthfully unknown, and the people who tell you otherwise are working from hope, not evidence, and they probably have an agenda.

One practical consequence of that rate of change is that it is entirely possible, and even common, to be making decisions based on a version of technology that no longer exists. I spoke recently with a senior engineering leader who had piloted AI coding tools about a year ago for a complex vertical market enterprise application. He concluded they were not sufficiently reliable, and moved on. His assessment was probably right at the time. But the tools available today are not the tools he evaluated, and the tools available a year from now will not be the tools available today. I told him he needed to go back and try them again, and not assume that what was true twelve months ago is still true now. That applies to your teams as well. Staying current with what AI can actually do — not what it could do when someone last checked — is not optional anymore. It changes how you operate.

I want to spend a few minutes on why that uncertainty matters practically, because I think it changes how you should be making decisions right now — about your internal operations, about your product, and about how you talk to your boards and your investors about it.

Start with the outcome, not the technology

Most of you have spent your careers making product decisions in a world where the outcome was knowable before you started. You decide to build a new feature. You scope it, resource it, execute it, ship it. The question was always cost and time, not whether it was possible. That predictability shaped how we planned, how we budgeted, how we reported progress, how we held teams accountable.

AI breaks that model, and the sooner you internalize that, the better the decisions you will make.

Here is the discipline I would encourage you to bring to every AI initiative — whether it is internal or product-facing, because the fundamental process is the same. Start not with the technology, but with the outcome. Pick something that would genuinely matter — not incrementally, but meaningfully. Reduce the time and effort to close your books by fifty percent. Improve conversion in your sales pipeline by ten percent. Cut customer onboarding time in half. The specificity and ambition of the outcome matters, because it sets the bar against which you will evaluate everything that follows. A five percent improvement in something that needed forty percent is not a success. It is a signal — that you may need a different approach, or that this particular problem may not be tractable with AI right now — and that is actually useful information if you treat it that way.

Once you have a meaningful outcome defined, the next question is whether you have the raw materials to go after it. Do you have the data? Is it the right data, and is it clean enough to be useful? Is the underlying process well-defined enough that a model can learn from it? These are questions you have to answer honestly before you commit significant resources, because unlike traditional software development, you can do everything right and still find out that the problem is not solvable with the tools available today. That is not a failure of execution. It is the nature of working at the frontier of a technology that does not yet have well-understood limits.

This is also where your sequencing decisions matter enormously. Before you build, evaluate a buy. If a vendor's off-the-shelf solution cannot get you toward your outcome, that is important and relatively cheap information to have before you commission a custom development effort. On the product side, before you invest in a distilled, cost-optimized model for a feature, prototype with a foundation model. Yes, it will be expensive at scale, but the prototype allows you to determine whether the capability is actually there and whether users respond to it quickly. Prove the thing works before you invest in making it efficient. Too many teams optimize something that was never going to work in the first place.

And as you move toward your outcome iteratively, resist the pull toward projects that look like AI work but are not connected to the outcome you defined. The pressure to have an AI story — for your board, for your customers, for your investors — is real enough that doing something visible can feel better than doing something right. Initiatives get launched that generate talking points but do not move numbers. Usage mandates get set that create activity metrics and change nothing in the business. That is not strategy, it is theater, and it consumes the organizational attention that could go toward something that actually matters. The test is simple: can you draw a direct line between what you are doing and a measurable business outcome that you and your board cares about? If you cannot draw that line, find a different, more impactful, problem.

Some of the things you point AI at will not reach the bar you set. The problem will turn out not to be tractable right now, or the data will not be there, or the outcome will plateau well short of being compelling. If you have held the bar high, you will encounter this, and how you respond to it matters. If it gets treated as a failure of the team, you will get a conservative response — people will stop proposing ambitious problems and start proposing safe ones, and nothing will change in the business. Separate your evaluation of the outcome from your evaluation of the work. A team that defined a meaningful problem, assessed the raw materials honestly, moved toward the outcome with discipline, and found that the problem was not solvable right now did something valuable. Make sure they know it.

There is one more thing worth saying here, and it connects directly to how aggressively you should be willing to move. The cost of writing software is falling rapidly — not to zero, but dramatically. What used to take a team months now takes days. That changes the economics of risk in a fundamental way. If iteration is cheap, you can take more swings. If a project does not reach the outcome you defined, the cost of walking away is lower than it has ever been — the sunk cost that tempts you to keep going is simply smaller. The organizations that understand this will move faster, run more experiments, and accumulate learning at a pace that organizations treating software costs as they were two years ago simply cannot match. The question becomes not whether you can afford to take more risk. The question is whether you can afford not to.

The demo problem is structural

I want to talk about something specific to the product side that I think is genuinely underappreciated, and that will resonate with anyone who has shipped an AI feature and watched the customer reaction not quite match the demo.

Traditional software features are, in an important sense, as good as they will ever be on the day you ship them. They do not change until you update the software. The customer's experience on day one is the experience, and it is consistent — every customer who uses that feature gets essentially the same thing.

AI features work the opposite way. They learn. Which means they improve over time. But it also means that the gap between what you demonstrated in a controlled environment and what a customer experiences in their first weeks with the product can be significant. The demo was real — the capability exists — but the system has not yet learned from this customer's data, this customer's patterns, this customer's edge cases. It needs time and usage before it approaches the quality that made everyone in the boardroom excited. You need to manage customers through an early experience that may be disappointing relative to their expectations, and you need to set those expectations honestly before they go live, not after.

There is a second thing that is less obvious. With traditional software, two customers using the same feature get the same behavior. With AI features, two customers can have materially different experiences with the same feature — because the quality of the output depends heavily on the quality and structure of their underlying data. A customer with clean, well-organized data will see better results than a customer whose data is sparse, fragmented or inconsistent, even though they are running the same model. The often stated assertion that you need both volume and variety of data to get great outcomes holds true. This is the fundamental nature of this technology. But it means your customer success motion, your support model, and your quality expectations all have to account for variability in a way that traditional software did not require. We are not used to a world where two customers have legitimately different experiences with the same feature and both are correct, but we need to get used to it and adapt to it.

This is the realm of the forward deployed engineer and why they have become so essential for AI forward companies - they bridge the gap between possibility and reality for customers.

Stop the long roadmap

Let me turn to something that connects back to where I started.

The uncertainty I described at the beginning is not just a constraint you are operating under. It is something you have to build your operating model around, and that has a very specific implication for how you plan.

If you have a detailed three-year AI roadmap, I would encourage you to look at it carefully and ask yourself what it is actually based on. The models available eighteen months ago could not do things that current models do routinely today. The models available eighteen months from now will almost certainly do things that current models cannot, in ways that are genuinely hard to predict even for the people building them. A detailed plan built against today's capabilities will systematically underinvest in what becomes possible. A plan built against speculative future capabilities will systematically overinvest in things that may not materialize when you assumed they would. Either way the plan is wrong in ways that will matter.

The only defensible planning horizon for specific AI initiatives right now is months, not years. Six months is reasonable. Twelve is the outer edge. Beyond that you are not planning, you are speculating with a precision that will give people false confidence and lead to real misallocation.

I understand the tension I just created. Your customers want roadmaps. Your teams want certainty. Your investors want to see a multi-year vision. The answer is to separate your strategic intent from your tactical plan. The intent can be stable and durable and outcome focused — we will use AI to fundamentally change our cost to serve, we will make AI central to our product experience in ways that drive retention and expansion. Hold your intent firmly. But the specific initiatives, the technology choices, the sequencing — those need to live on a short cycle and be revisited constantly as the underlying technology continues to evolve.

And it is not just the product roadmap that needs to change — it is the function that produces it. When development accelerates by an order of magnitude, the bottleneck shifts. It moves upstream, from engineering to product management. The traditional PM function — organized around prioritization ceremonies, backlog management, and roadmap coordination — was designed for a world where execution was the constraint. When execution is no longer the constraint, the entire function has to restructure around that reality. The organizations moving fastest will do this: concentrate product investment in a small number of people with genuine strategic judgment, deep domain knowledge, technical builder capability, or enterprise customer and market understanding — and letting AI handle the coordination and triage work that currently consumes most of the function's time. That is a separate and longer conversation, but the leaders who are rethinking their roadmaps should also be rethinking the organization that produces them.

Make decisions crisply, hold them lightly

The last thing I want to leave you with is something more personal than organizational.

Operating in conditions of rapid change and genuine uncertainty requires a specific kind of leadership discipline that most of us were not trained for. We were trained to make decisions, commit to them, and defend them. That served us well in a more stable world. It is a liability in this one.

The decisions you make about AI, such as which problems to pursue, which technology to bet on, which organizational changes to make, need to be made crisply but held lightly. Make the call, move forward, but keep the door open to new information that would cause you to revisit it. The enemy of that posture is ego, it is false certainty. When a decision owns the person who made it, new information stops being a useful signal and starts being a threat. That is how organizations end up defending positions that the evidence has already moved past. The goal is not to be right, it is to stay current with reality, and in a technology environment moving at this pace, those are not the same thing.

The second part of this can feel less intuitive. For decisions that are hard to reverse — architectural commitments, organizational structures built around a particular technology assumption, contracts that lock you in — make as few of them as possible, and make them as late as you can. This feels uncomfortable because it can look like indecision. It is not. It is the rational response to operating with incomplete information that is improving over time. Every week you wait on a consequential, hard-to-undo decision is a week in which the information available to you gets better, assuming you constantly and actively seek out that new information. Of course, that has always been true to some degree. With something progressing as rapidly as AI right now, it is dramatically true. Preserve your ability to change course for as long as you responsibly can.

What to do tomorrow

If you take nothing else from this conversation, here is what I would ask you to do coming out of this week.

Pick one outcome that would truly matter to your business — not a modest improvement, something that would change your trajectory — and ask honestly whether you have the raw materials to go after it. Not whether you have the technology. Whether you have the data, the process clarity, and the organizational readiness. Start there.

Pull out your AI roadmap and draw a line between what is stable strategic intent and what is a specific tactical bet. Anything on the tactical side with a horizon longer than twelve months deserves a hard look. The world it was planned against may already be different.

Find one initiative that is generating activity but not moving a number you care about. Have the conversation about redirecting that energy toward something with a clearer line to a business outcome. That conversation will be uncomfortable. Have it anyway.

And go back and look at something your team tried with AI 6-12 months ago and concluded was not ready. The assessment may have been right at the time. It may not be right anymore. The cost of checking is low. The cost of assuming the answer is the same is potentially much higher.

That is not a sign of strategic weakness. It is the only intellectually honest position available right now. And at our scale — companies your size, with the ability to make decisions and change course faster than large public enterprises ever could — it is a genuine advantage. The organizations that will do best with AI over the next several years are not necessarily the ones who made the best bets in 2025. They are the ones who built the organizational capability to learn quickly, course-correct cheaply, and recognize a good outcome when they find one.

That capability is entirely within reach for all of us.

Notes for revision

Section labels are orientation markers only — the talk is intended to flow as a single continuous argument without explicit signposting. Approximate spoken length at a fast pace: 16–18 minutes.