

BHARAT RAVURI

The

Practitioner Thinker Letter

ISSUE TWO · JULY 2026

Before the journey: how a board decides on AI — the hardest call on the desk, reasoned through.



Bharat Ravuri

Former CEO · Founder · Board Advisor · The R Doctrine
Hyderabad, India

This is the second of these letters. The first described a world that has reset across four dimensions at once — the technological (AI), the geopolitical, the financial, and the risk reset — four structural shifts arriving together, each a significant change in its own right, and operating in a way where each one influences the others: the cascade. Of the four, artificial intelligence is the most prominent and visible change, and the one most often misread: treated as a technology deployment when it is really a question about what your institution should become, now that intelligence itself can be manufactured in a data centre. That letter set out the doctrines for navigating the reset, with particular focus on The AI Doctrine, along with the case studies behind it.

This letter turns to an even more fundamental decision in front of the organisation: **should it pursue the AI journey at all, or wait** — and we offer a structured way of answering that question.

ONE

The Hardest Call

Why is the decision on AI so genuinely hard for so many organisations? As you can see from the evidence below, some are rushing in too soon, while many are paralysed, waiting for clarity. But if we break the question down, it is not really about timing — *are we too early, or too late?* That is the surface. Underneath sit three foundational questions, and they arrive together:

The technology is evolving so fast that by the time we have formed a view on what it can do, it has already moved to the next stage. Should we wait for it to settle so we are clear on how to implement it — or, in waiting, are we missing the benefits others are already accruing?

The investment being asked of us is large — money, time, resources — and we are not sure of the true return.

The proof we would normally demand is thin. For all the hype, there are remarkably few places we can point to where this has been an unambiguous, durable success — and a great many visible, expensive failures, including at companies far better resourced than us.

So we are being asked to commit serious resources, on a fast-moving technology, with little proof of success and a long list of cautionary tales — and to do it now. We are caught in a whirlwind: move too early, and we may be pouring money into something the market is already whispering is a bubble; wait too long, and a competitor may race ahead, building a moat that becomes very hard to close. From where we sit, both directions look reckless. That is what makes this decision genuinely hard — not that the right answer is difficult to execute, but that it is difficult to *know*, and we are expected to know.

You are not alone in this. The exact tension has now been measured, and it is present in boardrooms across the globe.

THE EVIDENCE • SPLIT DECISIONS, BCG, 2026

A survey of 625 leaders — 351 chief executives and 274 board members, all at companies above one hundred million dollars in revenue.

- Around **60% of chief executives** felt their boards were pushing AI *too fast*.
- The board members who rated their *own* understanding of AI as **weakest** were the ones most convinced the organisation was moving *too slowly*. The less a board understood AI, the faster it wanted to go.
- **35% of CEOs** said their boards *overestimated* what AI could replace.
- **~40% of CEOs** felt their boards lacked an informed view of how AI was reshaping the company's growth strategy.
- Yet **three-quarters of board members** rated their own AI knowledge *at or above* that of their peers.
- On pressure: CEOs believed roughly **a third** of their performance review now rode on AI results; their boards put it at barely **a quarter**. The person carrying the decision feels its weight more than the people around the table realise.

Here is my take. In this field, urgency is running ahead of understanding — which means the most valuable thing we can do is neither speed up nor slow down on instinct, but step back and decide *well*. That is what this letter is for. What I lay in front of you is a structure — a way to reason through this decision so that the judgement you bring to it has something solid to work with. There is no single right answer here, and no standard answer that fits every organisation. This structure exists to help you find the right decision for *your* company.

TWO

Two Journeys

Before we work through the decision, let us understand its consequences through two examples — because the fastest way to see what is at stake is to look at how this has actually gone for those ahead of us. And it has gone two very different ways.

The honest shape of it: most attempts to make AI pay have failed, and a small number have built something close to extraordinary. The gap between them is almost never about who had the better technology. It is about how they went about the work.

THE JOURNEY THAT FAILS

Activity without outcome

MIT studied around three hundred enterprise AI deployments and found that roughly **95%** had produced no measurable impact on the bottom line at all. And the picture is not improving with time: S&P Global Market Intelligence found that **42%** of companies abandoned most of their AI initiatives in 2025 — up from **17%** the year before. More money, better models, and the abandonment rate more than doubled in a year.

The reason is the point. These were not, in the main, failures of technology — the models mostly worked. What failed was everything around them: the question framed wrongly before a single tool was chosen, AI bolted onto processes built for a vanished era, pilots that dazzled in a demo and collapsed in production. They had treated a question about what their business should *become* as a project to *install some technology*.

THE JOURNEY THAT WORKS

DBS — AI paying off at scale

DBS, the Singapore bank, is the clearest case of AI paying off at scale inside a regulated institution — and the striking thing is how little of the story is about clever models.

It began with **data** — years of patient, unglamorous work building a governed, accessible foundation. It **industrialised**, building a common platform that cut the time from concept to production from **twelve-to-fifteen months down to two or three**. It led with **revenue, not cost** — of the value it credits to AI, by its own disclosures **more than one billion Singapore dollars a year**, the larger share new revenue rather than headcount removed. And it **carried its people**, reskilling more than **eleven thousand** roles rather than cutting them — all under one chief executive who owned the work for **more than fifteen years**.

Hold the two side by side. A well-thought-through decision, patiently seen through, can deliver extraordinary results. A hurried decision, or one taken from the fear of being left behind, tends to deliver a great deal of activity and very little outcome. The difference is not luck, and it is not the technology. It is the quality of the decision and the discipline of the build. Which raises the real question — the one this letter exists to answer. **Before committing to a journey like this, what should we be asking?**

THREE

Seven Dimensions of One Decision

The decision about AI is not a simple yes or no. It is better understood through seven dimensions, which I will walk you through in turn. At a high level, they are these:

1 • Technology	Should we wait for the technology to settle, or move now? What pace and posture should we take?
2 • Investment	How do we judge the return? What funding model ensures the money delivers outcomes, rather than being spent without result?
3 • Ambition	What business outcomes do we want from AI, and over what timeline?
4 • Risk & Liability	What unforeseen risks and liabilities does AI carry? How do we govern a machine making decisions outside direct human control?
5 • Ownership	Who will own this transformation? Who is accountable for the outcomes, and who will champion the journey?
6 • People	What will our people feel when we deploy AI? Acceptance or resistance — and will our own leaders support the rollout, or quietly resist it?
7 • Readiness	Is the organisation actually ready? How do we assess that, and what should we put in place before we begin?

The culmination of all seven is what informs the decision. As I take you through each, I will bring real-world examples — successes and failures both — so you can see the opportunities and the risks clearly, and reason your own way to the answer.

1 • Technology

Let us address the core question on technology first: should we wait for it to fully evolve and settle before architecting the business around it, or make the decision now to embrace AI? To answer that, we first need to understand how this technology is actually evolving. And here, for once, we can replace anxiety with measurement.

THE MEASURED PACE • METR

An independent lab, METR, measures something concrete: the length of a task — judged by how long it takes a competent human — that a frontier AI system can complete on its own. That figure has been **doubling roughly every seven months, for six years.**

To make it tangible: at the launch of ChatGPT, these systems could manage tasks of about **thirty seconds**; today the best of them handle tasks that take a person **more than fourteen hours**. *(The measure is taken on coding and technical work, at a fifty-per-cent reliability bar — not a claim that machines already do fourteen hours of any work dependably, but a clear read on the slope.)*

This is not new to us. We have seen an evolution shaped like this before, with Moore's Law — and we will see the same with AI. This is not a technology that will evolve and then stabilise; it will keep evolving. Which means there is no good, still moment at which to step in. **Waiting for one is waiting for something that will not arrive.**

The smarter approach is to account for the change *as* you transform the organisation — to build the capacity and capability to move with a moving target, rather than betting on a fixed picture of it that will be obsolete soon. So: start moving, but in a deliberate, conscious way, building the core capabilities as you go.

2 · Investment

Now let us address the other big concern: investment and return. The signs of investment going wrong are already visible at the most sophisticated companies in the world.

CAUTIONARY — RUNAWAY COST

Microsoft cancelled most of the internal licences for an external AI-coding assistant across one of its largest engineering divisions, effective mid-2026, after usage-based token costs ran far ahead of budget. And **Uber's** chief technology officer described the company burning through its entire annual AI-coding-tools budget in four months — adoption of a single coding assistant outran the finance team's model so completely that the firm imposed a hard per-engineer monthly cap.

This brings investment squarely into question, and the fundamental ask is how do we ensure we deliver business outcomes that are sustainable and deliver the return. We do not want to sign a cheque and wait three or four years to find out. This is where I propose a **cycle-funding model** — not something new, but a discipline that holds steady through the challenging times.

There is a *base* investment any start requires — to get the organisation ready, to train a meaningful proportion of people, to acquire the core licences and platform. Beyond that base, the board should approve funding only one *cycle* at a time — each cycle carrying specific objectives and accomplishments to be reached by its end, usually a quarter. At the end of the cycle, the sponsors come back to the board and explain what was achieved, what failed, what was learned, what they propose for the next cycle, and the funding it requires. Every quarter, the board has a real opportunity to review progress and decide whether to release the next round — keeping the whole effort moving toward measurable outcomes.

This does something important: it pushes the programme teams to pick a focused area and *demonstrate* the outcome there before rolling it out across the whole organisation at once. The steps are smaller, but the lessons are seen actively, at every level — and the teams have to *earn the right* to spend the next quarter's money. It is a far stronger discipline than a blanket cheque written up front.

PROOF — OUTCOMES ARE ACHIEVABLE

That the outcomes are achievable is not in doubt. **DBS**, by its own disclosures, credits AI with more than a billion Singapore dollars a year, the larger share new revenue. **JPMorgan** runs its AI suite across more than **230,000 employees and 450-plus use cases**. And they are not alone: **Upstart** rebuilt lending as an AI-native business, **BBVA** scaled adoption from the bottom up, **AIG** compressed its underwriting timelines, **PwC** built a new advisory offering on a single AI platform. Different industries, different approaches, the same lesson — **the outcomes come from structure, not from the size of the cheque**.

3 · Ambition

This is a critical dimension, because the starting approach of a journey is so often divergent from its real goal. Most organisations say they want *revenue* from AI — yet the starting point is usually the deployment of AI licences across the organisation. That goal of revenue cannot be reached through a *laissez-faire*, sequential approach.

I have also heard many leaders talk about their goal as AI agents as a percentage of the workforce — some saying half the workforce in the future could be AI agents — and they usually frame those agents as a way of automating current human tasks to gain productivity. I think that framing and approach is wrong. Why? The way humans are organised — into functions, specialities, and capabilities — and the way we execute work, through hand-offs and processes, are themselves shaped by human limitations: limited memory, limited span of attention, focused capability. Those structures and processes are inherently inefficient *because* of those human limitations.

Automate them as they are, and you have not made the organisation efficient —

it is inefficiency at machine scale.

You have built a more efficient version of the wrong thing. So it is always better to step back and reimagine the work in the context of AI, rather than automate the work as it stands. And that is exactly why the right measure is *business outcomes*, not a percentage of the workforce replaced.

With that frame, I recommend the journey begin along **three parallel tracks**, not in sequence:

Adopt & Automate

Deploy AI broadly across the organisation, get people using it in their own work, let them automate tasks and find efficiencies. The point is not the efficiency itself — it is to make the organisation *conversant* with AI, and, more importantly, to tap into what I call *hinterland innovation*: the unprompted use that surfaces ideas no roadmap would ever have planned, when people are simply given these tools and left to use them however they wish, with no central direction.

PROOF — BBVA, ADOPTION FROM THE BOTTOM UP

When BBVA opened AI access broadly across its workforce, its own people built more than **twenty thousand custom assistants** themselves — pulled into being by demand from the bottom, not pushed from the top. Adoption scaled from **3,300 licences to 11,000 to its entire workforce of 120,000**, with users saving on the order of **2.8 hours a week** and a large majority using the tools daily. The headline business results do not come from here, and expecting them to would be a mistake. Its value is fluency, and a stream of surfaced ideas the funded work then draws on.

Reimagine the Work

This is organisation-directed: a deliberate effort to re-engineer or reimagine a particular product, business, or function toward a defined business outcome. It typically has a business and a technology leader as joint sponsors, a cross-functional team doing the work, and a report back to the board carrying both the results and the lessons — for learning and for setting direction.

IN PRACTICE — JPMORGAN

JPMorgan now runs its AI suite across more than **230,000 employees and 450-plus use cases** — some of it first-level acceleration, but a meaningful part genuinely rebuilding how the work is done.

Build the New

While the first two tracks are largely about efficiency and productivity, the board and chief executive should make a conscious, separate effort to launch a genuinely *new, AI-native business*. This needs a dedicated, named leader and specific revenue objectives, with the mandate to build from the ground up. Many treat this as something to attempt only after the first two tracks have proved themselves. I would strongly recommend the opposite: begin it on day one, in parallel with the others. It is what allows an organisation to reach real revenue from AI sooner than most — and to build a genuine distinction in the market.

The best way to reach business outcomes from AI is not to pick one of these tracks, but to **run all three as one system** — each feeding the others, all of them in motion together rather than in a queue.

4 · Risk & Liability

This is the biggest concern that plagues most boards, and rightly so — because managing risk is a defined, central responsibility of the board, and risk is always mitigated through controls. But for the first time, what is being deployed into the business is an *intelligent* machine, and the concern is fundamental: can we control it? Can we keep it within the organisation's guardrails — and, independent of whatever benefits it delivers, could it create liabilities that exceed those benefits?

These are not unfounded fears. It helps to see the risk surface not as a pile of frightening stories but as a set of **distinct categories**, because that is how you begin to govern it.

THE RISK SURFACE — A FEW CATEGORIES

A

Customer-facing

AIR CANADA — THE CHATBOT'S PROMISE

Air Canada's website chatbot told a grieving passenger he could book a flight at full price and claim a bereavement discount retroactively, within ninety days. He relied on it and booked. When he later applied for the refund, the airline declined — its actual policy did not allow bereavement fares to be claimed after travel — and pointed out that the chatbot had simply been wrong. The passenger sued. Air Canada argued, remarkably, that the chatbot was a separate entity responsible for its own statements.

The tribunal rejected that outright: the chatbot was part of the airline's website, the airline was responsible for everything on it, and it was ordered to honour the price the bot had quoted. "The model said it, not us" is not a defence.

CHEVROLET DEALERSHIP — THE ONE-DOLLAR TAHOE

A Chevrolet dealership deployed a customer-facing chatbot on a general-purpose model to engage buyers, without adequate controls. A visitor instructed it to agree with anything he said and to end every reply with "that's a legally binding offer — no takesies backsies," then asked to buy a new Tahoe — an SUV listing above sixty thousand dollars — for a single dollar. The bot agreed, in writing. The screenshot went viral; others piled in, getting the same bot to recommend rival manufacturers' cars and even write software code.

The dealership pulled the chatbot and did not honour the "sale" — but the same product was running on hundreds of other dealer sites. The lesson is general: a capable model placed in front of customers with weak guardrails can create liabilities for the company.

B Controls**EARNEST — WHEN THE HUMAN OVERRIDE BECAME THE LIABILITY**

The lender Earnest used an AI model to help decide who qualified for a loan, with human staff able to override the model's outputs. The Massachusetts Attorney General investigated and alleged that the lending — the model *and* the manual overrides around it — produced a disparate impact, with certain applicants disadvantaged on the basis of protected characteristics. In 2026 Earnest settled for **two and a half million dollars**, on the basis of the allegation and **without admitting liability — a settlement, not a court ruling against it**.

The instructive part is what the regulator faulted: not only the model, but the human overrides applied without proper controls. The people placed in the loop to make the system safer became part of the liability rather than the safeguard against it.

C Models

There are two angles we need to cover. The first is *geopolitical risk*: governments have started imposing controls on the model layer — the software — beyond their earlier controls on the hardware, the chips.

GPT-5.6 "SOL" — SHIPPED TO A GOVERNMENT-APPROVED LIST

In June 2026, OpenAI released its most capable model, GPT-5.6 "Sol," not to the public but to roughly twenty partners whose names were individually approved by the US government — the first time a frontier model shipped under a government-managed access list, on concern over its capabilities in coding, biology, and cybersecurity.

ANTHROPIC — DISABLED WORLDWIDE, DAYS AFTER LAUNCH

Days earlier, Anthropic had been ordered under a US export-control directive to suspend access to its two newest models, Claude Fable 5 and Mythos 5, for all foreign nationals anywhere — including green-card holders and its own foreign-national staff. Because nationality cannot be checked at the point a system calls the model, Anthropic disabled both models for *every* customer worldwide.

MANUS — CHINA UNWINDS A \$2BN DEAL

China ordered Meta to unwind its roughly two-billion-dollar acquisition of the AI startup Manus on national-security grounds — striking because Manus, Chinese-founded, had relocated to Singapore precisely to present itself as international. The lesson the episode taught the market: re-flagging a company to a third country has limits when a government decides a capability is strategic.

One line for the boardroom: the frontier-model capability your business runs on has a switch — and that switch can be controlled by a government.

The second angle is more universal — *commercial dependency*. If you build your business on a single provider's model and it is withdrawn, degraded, repriced, or switched off for reasons that have nothing to do with you, your business is exposed. Both angles are the same underlying risk: **too much dependence on any single model**.

And these are only the tip of the iceberg. A 2026 study by the vendor Sinch found that around **three-quarters** of organisations had been forced to roll back a live customer-facing AI agent for one reason or another.

So how do we mitigate all this — without retreating from AI, which only leads to the paralysis we have already diagnosed? **Two things**.

First, on the model risk specifically: build on an **architecture of multiple models** — local and open models alongside the frontier ones — rather than depending on any single provider. No single model is irreplaceable enough to be worth the concentration risk; they are increasingly substitutable, so treat them as interchangeable components. Done well, that one choice de-risks continuity (no single point of switch-off, whether by a government or a vendor), mitigates cost (route the cheaper models for what does not need frontier capability), and protects private and proprietary data (keep the sensitive material on models you control).

Second, and more deeply: change *where in the process* risk and legal sit. The organisations that scale AI safely do not treat compliance as a gate at the end that blesses or blocks finished work. They bring risk and legal in as **co-authors from the first design session** — so their constraints become design parameters, and the thing is built safe rather than litigated into safety afterwards. Just as important is building genuine AI fluency *inside* the support functions — risk, legal, compliance — so those teams can plan their mitigations *while* the organisation moves to capture the opportunity, rather than standing as blockers at the end.

PROOF — BUILT SAFE, NOT LITIGATED SAFE

DBS built a data-use discipline into its platform from the start; BBVA brought legal and compliance in, in its own words, as partners from day one, enabling rather than slowing.

Two further things belong on the board's agenda. There should be a **dedicated working session with the board specifically on managing risk in the AI world** — this is too consequential to fold into a general update. And it must be understood as an evolving space: as organisations adapt to AI, so do the regulators, who are actively reworking the rules. Any system or design built today has to be able to *adapt* to a changing risk and regulatory landscape — compliance here is a moving target, not a fixed one.

Let me close this dimension on the point that matters most: **two things are equally true, and the board's task is to hold both at once**. The first: mitigating these risks is critical — not as caution for its own sake, but so the

organisation is never held liable, and never incurs losses that erase what AI was meant to deliver. The second: embracing AI is equally critical — it is what allows the organisation to leap ahead rather than be left behind. The mistake is to see these as opposing forces, where managing risk means moving slowly and moving fast means accepting danger. They are not in tension. A board that is genuinely informed — working in close coordination with its business, technology, and risk, compliance, and legal functions — does not choose between protecting the institution and advancing it. It does both, deliberately, at once.

5 • Ownership

Given the level of impact AI will have on the whole organisation, it is a no-brainer that ownership and accountability sit with the board and the chief executive. The larger and more useful question is: how should the organisation be *structured* to run this transformation — and who should be responsible for it? Here are four models to explore.

The **chief-executive-led model** puts responsibility at the top, with the executive team executing — and in practice this is often what is already happening; nearly three-quarters of chief executives say they are their company's chief decision-maker on AI. Its strength is that it treats AI as business architecture, owned where it should be. Its weakness is the paradox itself: it depends entirely on a chief executive who is truly fluent and engaged, and even then, that person cannot run it day to day — they are occupied with too many other priorities.

The **technology-led model** — running AI as a programme out of the CIO or CTO organisation — is being adopted widely today. The technologist has real capability and a deep understanding of the technology, but limited understanding of, and influence over, the business model. So the effort drifts toward automating existing processes — the first level of ambition. Its strength is genuine technical depth; its weakness is that the business-architecture decisions the work requires sit outside the owner's remit.

The **federated, or dual, model** has business and technology co-owning, with a dedicated transformation lead reporting to the chief executive, and AI embedded across the functions rather than walled into one. It is the structure closest to what the work actually needs.

IN PRACTICE — DBS

DBS organised itself around horizontal, cross-functional "journeys" of business and technology working as one — which its chief executive credited as the structural reason its AI was business-driven rather than IT-driven.

Its strength is that the ownership matches the shape of the problem: AI is horizontal, so its ownership is horizontal. Its weakness is that the transformation-lead role can easily fall into the coordination-overhead trap — becoming a programme office with no real responsibility, rather than genuinely driving the AI in close conjunction with the technology and business leaders.

The **chief-AI-officer model** is where the language gets confusing, so let me be careful. The market has begun calling the transformation lead described above a "Chief AI Officer." So the title itself tells you little; what matters is the mandate. The version I want to put in front of you is not a coordinator but a *builder*. Picture the role as a **seeding ground**: a senior owner who works across the existing businesses to transform them, in

close conjunction with the technology and business leaders, and is *also* given the mandate to build the new AI-native business itself. That turns the role from integrative to genuinely future-building. Its strength is clear senior ownership joined to a builder's mandate. Its weakness is real, and I will name it: it demands an exceptional individual and a true mandate, and done weakly it recreates the very silo the title is usually criticised for. I will also be candid that, this early, I cannot point you to a named institution running precisely this reframed version with a documented result.

There is no right model in the abstract. There is only a *fit* — to the organisation, the culture, the company's track record, and the depth of its business and technical talent. That said, I will give you my own view plainly: **I would urge you to consider a dedicated Chief AI Officer charged with building the AI-native business in parallel**, because it gives the company a head start on real revenue and creates a genuine distinction in the marketplace.

6 · People

The earlier dimensions were all about defining the AI transformation. This one determines whether the board and chief executive actually *achieve* the ambition they have set — because history is clear that **no transformation succeeds without the active participation of its leaders and its people**. Change management has always existed for exactly this reason: to bring people along, to win their support and participation in organisational change.

But AI transformation presents a far bigger challenge. Almost every article and headline carries the message that AI will eliminate jobs — and that fear is real, widespread, and certainly present among your own people. Alongside it runs the talk of *delaying*, and of the changing role of the leader and manager in the new organisation. So the very leaders and managers who are meant to motivate and carry the associates are themselves uncertain about their own roles and futures. **The people you would rely on to drive the change are unsettled by it.**

Addressing the concerns of leaders and people properly is a large subject in its own right, and I will devote a future letter to it — not at the surface, but dissecting the root of these fears and building the response to them. This letter stays at the level of the initial conversation the board and chief executive must have.

Before embarking on the AI initiative, it is critical that the board and chief executive spend real time on how they will carry their people along the journey. Many leaders are saying they will reskill everyone, that there will be no job cuts, that almost everyone will have a place. But assurances given before the outcome is known are hard for people to believe, however sincerely they are meant. I would strongly recommend the board address this honestly, in three parts. **First, honest communication:** tell people the truth — that it is early, that we do not yet know the true extent of the change until we are further into the journey, and that we commit to keeping everyone updated, openly and regularly, telling the truth at each stage as it becomes clear, including the parts that are hard to hear. That is a commitment about *how we will behave*, which people can believe, rather than a promise about an outcome we cannot yet see. **Second, reskilling that means something:** not generic retraining, but active participation in the AI transformation that will make our people genuinely more valuable in the market — more in demand — whether or not there turns out to be a place for them here. That is a promise we can actually keep, and one a person can verify; it is the opposite of asking someone to train their own replacement. **Third, a humane approach to exits:** if exits do have to come, they will be handled with the most humane approach we can manage, and in a way that ensures the people affected genuinely share in the benefits the transformation creates.

It feels frightening to communicate any of this so early — before we even know whether AI will deliver for us. But being honest and engaging people at the very beginning is precisely what builds the trust that carries a transformation far. The alternative — masking the uncertainty with assurances that may not hold — buys quiet for a while, and then breaks exactly when you need people most. Trust is what turns people from wary observers into active participants, and their participation is what determines whether you reach the ambition you set. **The honesty is not only the decent choice; it is the strategic one.**

What is equally critical is *effective communication* — how you present the information, and the words you use.

CAUTIONARY — STANDARD CHARTERED, A SLIP OF A FEW WORDS

Standard Chartered set out a plan to reduce its corporate-functions workforce — human resources, risk, compliance — by more than fifteen per cent by 2030, on the order of seven to eight thousand roles, as it leaned on automation. The bank's underlying position was a responsible one: it had for years invested in helping colleagues whose roles were exposed to automation reskill into higher-value work, and its point was that lower-value *roles*, not lower-value people, were the ones most vulnerable.

But at an investor forum in Hong Kong, in describing the shift, its chief executive said the bank was "replacing, in some cases, lower-value human capital" with the technology and investment it was putting in. The phrase detonated on two fronts at once. The first was the obvious one — employees and the public heard a bank label part of its own workforce as low in value, at the exact moment their jobs were said to be at risk. The second was subtler and arguably worse: the term "human capital" itself drew criticism for reducing people to a line item. A global union federation rebuked him publicly — the bank was "talking about thousands of people and families, not human capital." Within days the chief executive apologised on LinkedIn for his choice of words.

The bank's genuine, defensible reskilling story was almost entirely lost beneath the phrase. This was not a failure of intent or of strategy. Standard Chartered was, by its own long record, doing the responsible thing. It was a slip of a few words — at a podium, in front of investors — that undid it.

So AI transformation requires a dedicated **communication architecture** — defined deliberately, to ensure the message is received correctly by all stakeholders: employees, leaders, investors, regulators, the press. There is as much focus on using the right words as on avoiding the wrong ones. I will cover that architecture in detail in a future letter. For now, it is enough to see that **how you carry your people is not the soft wrapping around the decision. It is a major part of the decision.**

7 • Readiness

Now the focus shifts to a pertinent question: is the organisation ready to begin this journey? The general norm is to conduct a multi-week organisational readiness check across many parameters to test true readiness. This section is something different — a board-level *scan*: a set of questions the board and the CXO team can work through together. Where they are answered well and agreed upon, the organisation is standing on firm ground. Where the answers are difficult, or there is disagreement in the room, that is even *more* valuable — because those are precisely the questions the CXO team should carry into the deeper organisational readiness study.

Below is a sample set of questions. As a first step, I would ask you to add any others you think belong on the list. The intent is never to get a simple yes-or-no — that is a generic response, and it does not force a conversation. These questions are built to generate discussion in the room, because **the real insight and value live in the conversations, not in the questions.**

A SAMPLE READINESS SCAN

Force a specific; the conversation is the value.

STRATEGIC · Ambition · Technology

- If we had to name the AI-native business opportunity we would pursue first, what would it be?
- Name the competitor we are most worried about on AI.
- Which two or three of our products or lines are most exposed to AI in the next two years?

OWNERSHIP & TALENT

- If we had to name two or three AI champions to join this transformation, who are they?
- Who would lead the AI-native business initiative?
- Who on this board or executive team could genuinely *own* this — not sponsor it, own it?

RISK, LEGAL & REGULATORY

- Who, by name, is tuned in with our regulator on the changes being driven for and by AI?
- Can we name a moment when our risk or compliance team found a way to *yes* on something hard?

INVESTMENT

- Name two currently budgeted initiatives we would stop or suspend to fund this.
- What is the largest cheque this board would write for AI before demanding to see a return — and over how long?

PEOPLE

- If we told our people tomorrow that this was beginning, what is the first question they would ask?
- Name the part of the organisation most likely to resist this — and why.

FOUNDATIONAL · the enabler under all of it

- Can anyone point to one decision last quarter where we had the clean, governed data we needed, in time?

These questions are not a scorecard — they are conversation generators, because there are no right or wrong answers here. **You will know yourself whether you have arrived at the right answer or not.**

FOUR

The Decision Before the Journey

We have come to the end of the letter, and by now you can infer that this is no longer a yes-or-no decision on AI transformation. Every organisation has to plan its own journey on this road — and working through the dimensions above, and the readiness scan, will have given you a real sense of what it takes to frame that path and begin. As I said in the readiness section, the value is more in the divergence in the room than in the agreement — because those conversations are what bring the real opportunities and the real issues to the table, and that is what eventually helps the organisation succeed.

Making the decision to begin is critical. But seeing the journey through asks far more than a decision — it asks **courage, patience, and perseverance**, sustained over years and through the troughs where results have not yet arrived, through the false starts, and through the moments where your people need to be held steady against their fears and their anxiety.

And here is the part that is rarely said aloud. Everything this letter asks you to do for your people — to be honest before the outcome is known, to hold steady through the troughs, to lead from ahead of where the organisation is — you are asked to do while standing in the same uncertainty they are, and under a pressure they cannot fully see. You are being asked to steady an organisation through a change you did not choose the timing of, cannot fully predict, and will be held accountable for — and to do it without the certainty everyone assumes you must have. There is no version of this where the pressure is delegated upward and away from you. It stops with you.

Which is exactly why it can only be led, not managed. A transformation this uncertain cannot be run from a plan or held together by a programme office; it holds together only because someone at the top has decided to carry it — to go ahead of where the organisation is, and carry its leaders, its associates, its investors, and its regulators through the change, including on the days the results have not yet come and the doubt is loudest. That, in the end, is what this decision commits you to. Not a technology. **A journey that only leadership can carry — and only because leadership chooses to.**

Bharat Ravuri

The R Doctrine · The Practitioner Thinker

bharat@bharatravuri.com

A note on the evidence. Every case and figure in this letter is drawn from primary sources or the institutions' own disclosures, and where a claim rests on a settlement or a survey rather than a court ruling or an audited number, that distinction is kept. In a field moving this fast, figures are dated and refreshed before each use; where a verified example could not be found, that is said plainly rather than filled with invention.