

ARTIFICIAL INTELLIGENCE

Do You Own Your Enterprise Cortex? The AI Strategy Risk CEOs May Not See Coming.

By Aaron Arnoldsen, Rich Lesser, Djon Kleine, and Sanjeev Reddy

ARTICLE AUGUST 06, 2026 12 MIN READ

No CEO would build a mission-critical supply chain around a single supplier. This decade's hard experience—COVID-19 factory shutdowns, war-driven commodity shocks, semiconductor shortages that idled entire production lines—has taught companies to diversify critical inputs and design for failures that they can't predict. For anything strategic, the need for resilience outweighs the efficiency of a single source. That principle is already settled in how companies run

their operations, but the same thinking should also apply to AI as companies' decisions increasingly rely on it.

Today, providers across the AI ecosystem—from frontier model labs and hyperscalers to open-weight developers and specialist platforms—are racing to own as much of the stack as they can. Past technology waves have shown how difficult it can be to unwind dependencies once a vendor platform becomes essential to the daily operations of a business. At that point, technological lock-in takes hold.

Indeed, in our conversations with CEOs, we found that many are aware of the risks of overreliance on a single provider. As a result, in recent months, the top-of-mind question for informed CEOs has often shifted from, “Which model should we use?” to, “Are we committing too much, too soon to a single platform?” This question cuts to the heart of how to protect and strengthen an organization's unique identity as the role that AI plays in enterprise operations grows.

To preserve their organization's autonomy and flexibility to respond to a rapidly changing AI landscape, CEOs need to build a layered AI tech stack with a defined security perimeter around their most valuable internal knowledge. We call this the **enterprise cortex**—the brain of the company.

What is your organization's enterprise cortex? It's your IP, essential data, key business rules, proprietary information, and codified understanding of how processes work and how they link to your core business strategies, purpose, and values. These intellectual assets constitute the enterprise's most valuable internal knowledge, enabling it to thrive over time and maintain its distinctiveness versus the competition.

Why AI Takes Technological Lock-In to a New Level—Cognitive Lock-In

It's reasonable for a CEO to wonder, “Why do I need to worry about creating a protective layer around my organization's cortex if I have built privacy and ownership governance into the enterprise contracts I have signed with my platform and LLM providers?”

The answer is that in the AI era, a new reality amplifies the problem of technological lock-in: AI tools will increasingly become part of how the organization thinks and makes decisions.

As AI models and agents influence the way organizations solve problems, they can become inextricably linked to the organization. Over time, organizations risk becoming unduly dependent not just on a technology platform, but on an external source of intelligence. We call this phenomenon **cognitive lock-in**. This risk is not limited to proprietary frontier models. Open-weight deployments can reduce dependence on a provider while creating new dependencies around a particular checkpoint, tuning pipeline, serving infrastructure, or operating team.

Cognitive lock-in occurs when an organization thoroughly embeds its data and all of its operational context so deeply into a model, platform, or surrounding operating stack that changing any of them becomes prohibitively difficult. Strong contracts can protect your data, but exposure to an organization's data is only part of the issue. At least as important is the operational context, which includes key information—decision paths, legal rules, regulations, and any additional, unstructured yet valuable proprietary information such as surveys, standard operating procedures, and lessons learned from previous actions.

If all of that crucial operational context becomes interwoven with a particular model or architecture, the organization may believe that it's still making independent decisions. But it's making those decisions inside a technology provider's architecture that it doesn't wholly own and that it can't change to suit its immediate needs.

The risks are also more difficult to mitigate through existing or new contractual agreements. The language would need to account for interpretation, judgment, and ideas—all of which are difficult to define, monitor, and enforce in a world subsumed by AI, where they are often indistinguishable from the outputs of large language models (LLMs).

How can CEOs gauge whether their organization is drifting toward cognitive lock-in? A few signs are observable without a technical audit:

- When teams increasingly fail to explain “the why” because they are relying more on the model's reasoning than on their own human judgment
- When business leaders start voicing frustration that AI isn't meeting their actual needs
- When teams begin shaping strategy around what the model does well rather than what the business requires

These signs are strong indications that the model has started steering the enterprise rather than serving it.

At first glance, this situation may seem familiar. Organizations have seen similar patterns with ERP and SaaS platforms, where the need to accommodate the constraints of the technology reshaped processes. The crucial difference now is that the dependency is cognitive rather than operational. Instead of merely dictating how work gets done, the model shapes the enterprise's thinking to the point where switching it becomes too risky to attempt.

Cognitive lock-in need not result from misconduct by a provider. It can emerge from perfectly rational decisions by both the provider and the enterprise.

Architecting an AI-Transformation Tech Stack to Protect the Business

Avoiding cognitive lock-in does not mean rejecting vendor AI. Model providers, hyperscalers, and platform partners are producing extraordinary capabilities that companies can clearly benefit from. At the same time, every major AI platform—including model labs, hyperscalers, and data and software giants—is seeking to play a broader role in enterprise AI, extending into the enterprise cortex, the layer of the tech stack that houses the organization’s most valuable IP. The objective is to define the right boundary between vendor innovation and the enterprise’s cognitive core so that both can contribute what they do best.

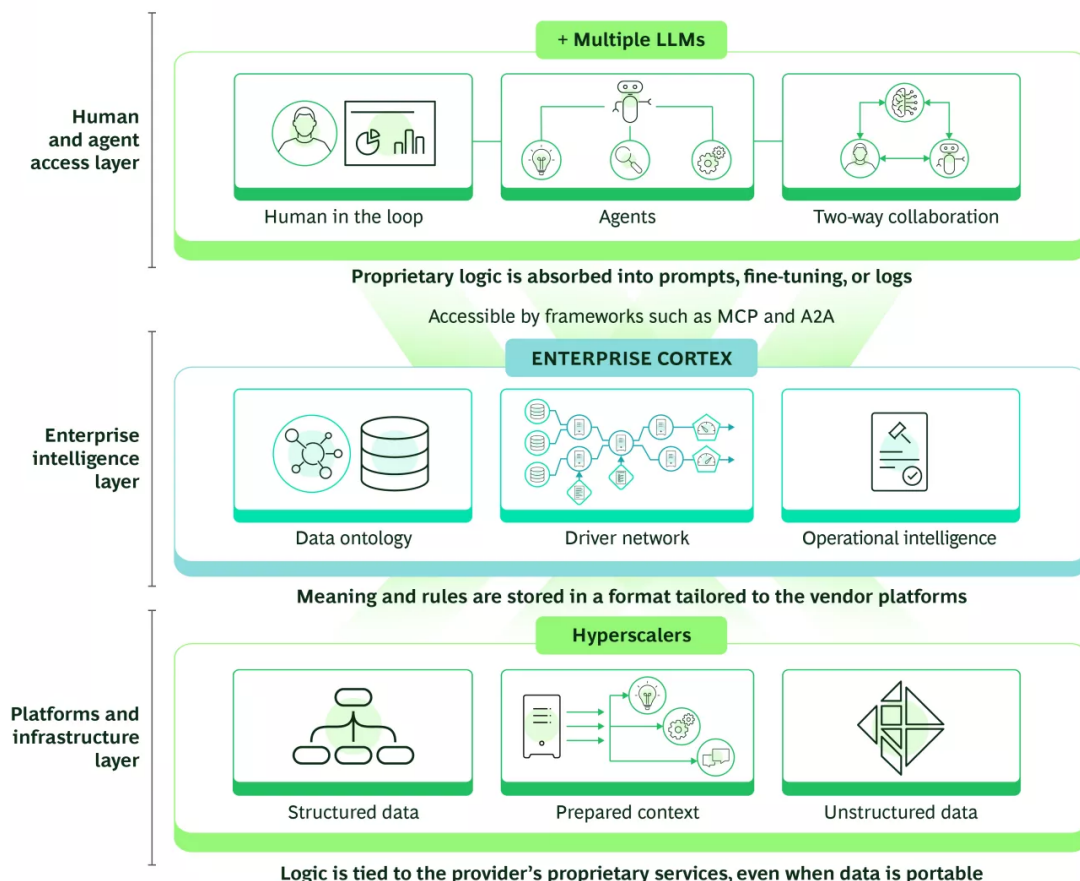
It is critical to note that organizations have never before invited such extensive access to their most valuable IP and internal knowledge, derived from their own insights and operations and from outside vendors alike. Companies have entrusted core information and business rules to vendor platforms for decades, but those systems largely execute predefined logic, holding your data and running your processes without interpreting, deciding, or generating judgment as AI models and agents do. What organizations are now exposing is not just the data and the rules, but the reasoning layer that sits on top of them—how the enterprise thinks, decides, creates, and acts. That is a new category of exposure, and it warrants careful thought about the possible repercussions.

The AI tech stack has three essential layers, each of which has a distinct boundary:

- **The Human and Agent Access Layer.** The top layer includes frontline management and workers, as well as the agents and applications through which they use AI. Those systems should be able to draw on a governed portfolio of models—from small language models (SLMs) and open-weight models to frontier LLMs—depending on the task.
- **The Enterprise Intelligence Layer or Enterprise Cortex.** The middle layer is the corporate brain layer, which supports and will increasingly drive the organization’s key decisions. It includes a network of your data ontology, rules, and operational intelligence.
- **The Platforms and Infrastructure Layer.** The third layer is where vendors, hyperscalers, or the enterprise itself hosts and serves models and data.

The middle layer—the enterprise intelligence layer—safeguards the enterprise cortex, and ensures that the system does not inadvertently share key intellectual property across layers. Common routing, evaluation, and fallback logic should make this model portfolio modular without exposing the cortex. (See the exhibit.)

The Three Layers of the AI Tech Stack



Source: BCG analysis.

Note: A2A = agent-to-agent; LLM = large language model; MCP = model context protocol.

As the following examples show, leading enterprises are already building this layer:

- **A global entertainment company** built an agentic platform that enables creatives and their managers to use AI to generate marketing communication assets at scale. Each creator's voice, brand guidelines, and personal preferences remain codified inside the company's cortex, and multiple technology providers' models sit on top, each drawing the context it needs from that single shared layer to do what it does best. The result: a frontier LLM scales the creator's unique ideas to millions of fans, while the sensitive material that defines them remains owned and protected.
- **A major retailer** uses AI to boost frontline associates' productivity. But LLMs answer on the basis of probabilities, and are prone to hallucinations, so an unguarded agent will cost you

more at the register than it will deliver in improved productivity. The enterprise cortex acts as a GPS over the company's own data and rules, keeping the model on the road and significantly improving the accuracy of its answers.

- **A global beauty company** encodes decades of marketing know-how—brand standards, retail expertise, hard-won decades of experimentation—directly into the enterprise cortex that its marketers work from. That knowledge stays exclusively the company's own, governed by the company and remaining portable across models. The payoff: the AI speaks in the brand's voice, not in a generic one, and the expertise behind it can't walk out the door.

How Organizations Can Create Boundaries Without Limiting the Value of AI

The challenge CEOs face is how to set boundaries for the AI models and agents without limiting the value it creates for the enterprise. Five principles can help them shape a tech stack that safeguards what makes the company truly unique, while enabling them to get the most out of LLMs.

Stay Clear-Eyed on Vendor Value and Ownership

To state this principle as simply as possible: Own the content, rent the containers, and buy or build the components from the best available, including a graph database and orchestration tooling. Tools are replaceable, but the corporate brain should never be.

The CEO must understand that protecting the enterprise cortex is not strictly an IT problem. Treat the tech stack as a governed enterprise asset that spans data, technology, risk, and the business itself. The CEO can designate a clear owner—a domain expert—who is responsible for validating the rules and defining a process for upkeep and scalability. Although you don't maintain the AI tech stack, you should view it as a critical business success pillar to uphold, just as you consistently scrutinize the health of the brand.

Set the Parameters for What Is Possible

LLMs are powerful because they can interpret language, summarize complexity, generate options, and reason through ambiguity. But an enterprise cannot run on reasoning alone.

The organization must build its own compliance constraints, permission structures, audit trails, business logic, and definitions that remain consistent from one prompt to the next. The agent should have the context to navigate to the right answer or workflow—routing work from task to task, recommending prices within defined parameters, approving exceptions, allocating supply, or triggering customer actions. Let the model handle language and reasoning, but let your enterprise cortex handle the rules, limits, and consequences.

Keep the Model and Platform Layers of Your AI Tech Stack Modular

The AI market will keep shifting as models improve, platforms change, and new agent frameworks emerge. Consequently, CEOs need to keep their AI-transformation tech stack as modular as possible. An organization should be able to swap models, tools, or platforms without redesigning its entire stack or workflow. Portability should extend across model classes, not just across technology providers. The operating principle is workload placement: use the smallest, least-expensive model that meets a task's quality, latency, and risk requirements, with a governed fallback to a stronger model—or a person—when it is unequal to the task.

Standards governing exactly this kind of portability are beginning to appear. Various open, standards-based protocols—including the Model Context Protocol (MCP), Agent-to-Agent (A2A), and Agent Communication Protocol (ACP)—have been emerging to ensure that models and agents can connect to external tools and to one another. Organizations can require support for open, standards-based context interfaces to their environments and in enterprise contracts to ensure that they remain accessible to any model that they may want to adopt in the future.

Move Fast Through Focused Execution

The choice that CEOs face is not whether to favor control or speed, but where to apply both control and speed first. The companies making the biggest strides don't accumulate isolated use cases. Instead, they start with a clear picture of what they want the platform to become, and then they reshape an entire high-value workflow from end to end, with a full pricing process, a service recovery journey, and a credit decisioning flow. They secure the operating logic behind the workflow, prove the financial payback, and then scale outward, decision by decision, as the speed compounds. Focusing means going deep on one workflow that matters, not shipping shallow features across many workflows.

Manage Models as a Portfolio—and Own Them Selectively

For a data-rich company, an open-weight model or fit-for-purpose SLM tuned to the company's own data can become a controlled, specialized asset. The distinction is task-shaped, not a blanket rule. Frontier models repay their cost on open-ended, low-volume, high-variety work—novel reasoning, synthesis across domains, and tasks whose shape you can't predict in advance. SLMs may prove attractive on the opposite profile—narrow, well-defined, high-volume tasks in the company's own language, run frequently enough for the token economics and easier governance to outweigh the raw capability you give up.

If a task is repetitive, bounded, and easier to train, it may be a candidate to bring in-house; if it's varied, exploratory, or rare, keep leveraging frontier capability. The model can become an asset. The enterprise cortex—together with the evaluation and routing logic surrounding it—makes it durable as your needs and model capabilities evolve.

Many companies are reaching or will soon reach a critical juncture in their AI journeys. The decisions that CEOs make at that inflection point may well determine whether AI strengthens what makes their enterprise distinctive or slowly erodes it.

Above all, the organization's unique identity—and the knowledge that underpins its competitive advantage—must stay within the protected center. The enterprise cortex must remain protected and autonomous, capable of expanding at the organization's pace and adapting as technology evolves.

Authors



Aaron Arnoldsen

Managing Director & Partner,
BCG X
Seattle



Rich Lesser

Global Chair
New York



Djon Kleine

Managing Director & Partner
San Francisco - Bay Area



Sanjeev Reddy

Managing Director & Partner
Boston



ABOUT BOSTON CONSULTING GROUP

Boston Consulting Group bridges the gap between ambition and outcomes for the world's leading companies and organizations. We are built for this era of unprecedented change — bringing strategic clarity rooted in over 60 years of deep domain knowledge, combined with applied AI shaped by our practitioners. BCG works shoulder-to-shoulder with CEOs across industries and geographies to deliver transformative impact at scale: stronger returns, transferred capabilities, and change that sticks. For more information, visit bcg.com.