

The Inference Placement Framework

An executive framework for placing AI reasoning across the enterprise

Jennifer Viens · Viens Strategic Advisory · July 2026

Executive Summary

For three decades, enterprise architecture has been organized around one question: where should the application run? AI adds a second: where should intelligence run? That question is now being asked across the industry, and rightly so; what has been missing is not the question but a disciplined way to answer it.

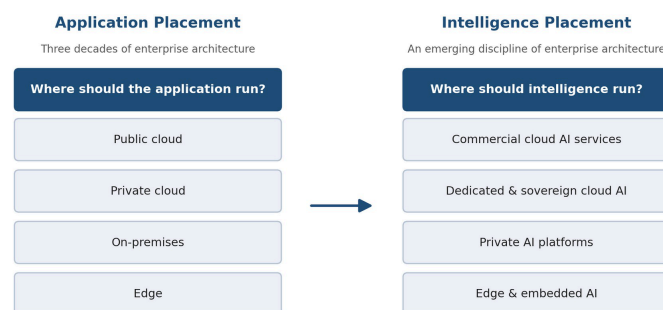
The framework in this paper is that instrument, and its unit of analysis is deliberately small: most architecture approaches place the AI workload; this one places each consequential inference step within it, because the answer differs by decision. Some decisions demand the lowest possible latency, some prioritize cost or resilience, some require the highest levels of trust and governance, and others need data or domain expertise that lives in only one place. Too many AI strategies begin by selecting a model or a platform, when the better starting point is the decision itself.

This paper argues that intelligence should be treated as a portfolio rather than confined to a single model, platform, or environment, and it offers a practical instrument: start from the decision, score it on nine dimensions, place it deliberately on a four-position spectrum, and govern the whole portfolio through one gateway.

1. The Question That Built Enterprise Architecture Has a Successor

The application question built a discipline, a market, and a generation of architects, and it never went away; nothing here retires it. AI simply adds a second placement question on top of it, and the second is harder than the first.

An application executes logic someone already wrote, so placing it is largely a matter of performance, cost, and control. An AI workload produces reasoning at run time, which means its placement decides not only how fast and how cheaply an answer arrives but what data the reasoning can see, what rules it operates under, and whether the organization can explain the answer afterward. A question with that many consequences deserves its own discipline, not a footnote in the application playbook.



The first question placed workloads. The second places reasoning.

2. Placement Is Already Happening, Just Not on Purpose

The distributed future is not a forecast. A 2026 survey of more than 1,100 IT decision-makers found that 78 percent of organizations now run inference themselves, hosting models in their own data centers, in private cloud, or on cloud capacity they control, rather than consuming AI only as a public API service. The same survey found that 77 percent name inference as their main AI activity, and that the typical enterprise has seven AI models in operation or under evaluation [1]. IDC projects that by 2028, seven in ten of the most AI-advanced enterprises will route work dynamically across many models rather than standardizing on one [2].

So most enterprises already own a multi-model, multi-environment intelligence estate, and that estate is an asset, not a problem: the buying and building of the past two years created real capability. The opportunity now is to run it deliberately. Placement choices made inside individual pilots and procurements harden into architecture over time, and the cheapest day to make them on purpose is before they do.

3. The Starting Point Is the Decision, Not the Model

Model selection is where most AI strategies begin, and it is an unstable anchor, because the model landscape turns over faster than any enterprise planning cycle. The model that leads the benchmarks this quarter is routinely displaced the next.

The stable unit is the decision the business needs to make: approve or deny, escalate or resolve, reorder or hold, alert or wait. Nearly every consequential decision now carries an inference step, and that inference step, not the application and not the model, is the unit of placement.

Every major platform provider now offers deployments across this spectrum, from public cloud services to sovereign regions, customer-operated infrastructure, and the edge. The market itself is conceding that placement will be plural. What buyers still need is a provider-neutral executive discipline for choosing among those environments at the level of each inference step. That is the gap this framework addresses.

Working backwards from the decision prevents the two most expensive mistakes in enterprise AI: over-securing low-stakes workloads until they die in review, and under-securing high-stakes workloads until they surface in an audit.

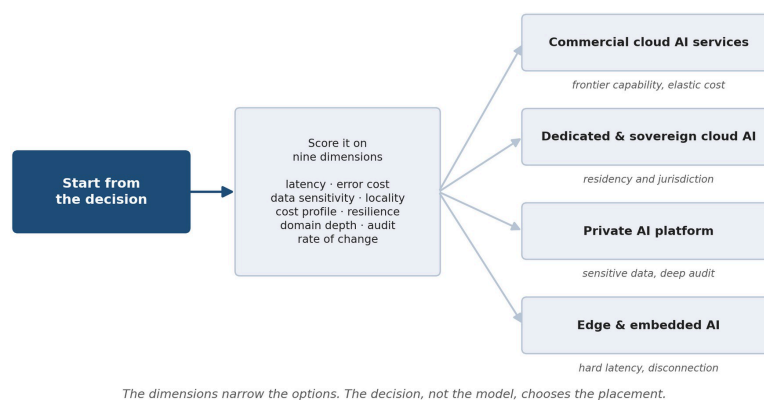
4. Every Decision Scores Differently on Nine Dimensions

The assessment translates business requirements into placement requirements. Each dimension is an executive question, answerable by the owner of the decision, not only by a technologist, and each answer narrows the viable placements.

Dimension	The executive question
Latency	How fast must the answer arrive for the decision to still matter?
Consequence of error	What happens when the answer is wrong, and who bears it?
Data sensitivity	Can the data behind this decision leave its current boundary?

Dimension	The executive question
Locality and jurisdiction	Does regulation fix where the data, and the reasoning over it, must sit?
Cost profile	Is the volume steady and high, or occasional and spiky?
Resilience	Must the decision still be made when the network is degraded or gone?
Domain depth	Does the decision need specialized knowledge a general-purpose model does not carry?
Auditability	Will you need to reconstruct how this answer was produced, months later?
Rate of change	How often will the underlying capability need to improve or be replaced?

The locality question deserves particular attention, because it is the one regulators are now answering for you. Gartner frames AI sovereignty as the need for authority over every layer of the AI stack while operating inside geographic boundaries [3], and predicts that by 2027, 35 percent of countries will be tied to region-specific AI platforms, leaving multinational enterprises to manage several platform partnerships, each with its own compliance and data-governance regime [4]. For a global enterprise, some placements will be chosen by geography whether the architecture anticipates it or not.



5. The Spectrum Runs from Commercial Cloud to the Edge

The four environments are not competitors. They are positions on a continuum of control, and a mature enterprise portfolio will occupy all four at once.

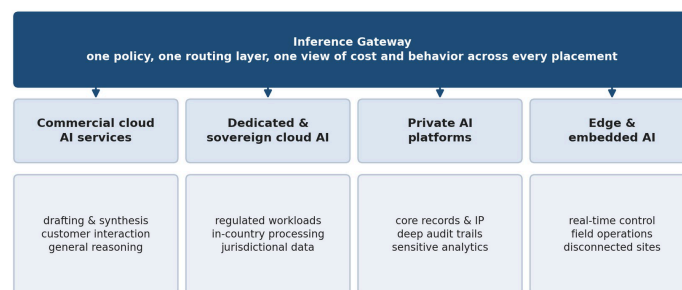
Commercial cloud AI services are often the fastest path to broad capability and scale. Frontier models, continuous improvement, and consumption economics arrive here first, at a pace privately run environments rarely match. Reflexively excluding this tier is itself a placement decision, made by default, and it strands low-sensitivity work that should move at commercial pace.

Dedicated and sovereign cloud AI delivers the same operating model inside a jurisdictional or compliance boundary, and it is the answer when locality, not capability, is the binding constraint.

Private AI platforms bring the reasoning to data that cannot move: core records, proprietary IP, and workloads where the audit trail must be owned end to end.

Edge and embedded AI is where latency or disconnection decides, running small, task-specific models at the point of action.

One process can legitimately span several placements. On a manufacturing quality line, defect detection runs on camera at the line because milliseconds matter, the supplier correspondence it triggers is drafted by a commercial cloud service because capability matters, and the warranty exposure analysis runs on the private platform because it reasons over proprietary claims data. One process, three placements, each earned.



Intelligence as a portfolio: every placement earns its position, and one gateway governs them as a single system.

6. One Gateway Governs the Whole Portfolio

Distribution without a control point is how a portfolio becomes a sprawl. The gateway is that control point: a single layer that routes each inference step to its assigned placement, applies one policy everywhere, meters cost across every environment, and logs behavior so the audit question is answerable in one place.

The gateway is also what makes the portfolio durable. When a better model appears, and it will, the enterprise changes a route, not an architecture.

7. Intelligence Is a Portfolio, and Portfolios Get Managed

The future of enterprise AI will not be defined by a single model or a single cloud. It will be defined by the ability to place intelligence where it delivers the best outcome, and just as we learned to architect applications across multiple environments, we will learn to architect reasoning across them. Intelligence placement will take its seat as a core discipline of enterprise architecture, with its own review question standing alongside the one we have asked for thirty years.

The exercise takes thirty minutes. List the ten most consequential decisions in your business that now carry an inference step. Run each through the nine questions, then compare where the inference runs today with where the scoring says it should run.

The gaps are your inference placement roadmap.

The question that closes each one is simple:
Where should this intelligence run?

References

1. F5, 2026 State of Application Strategy Report, survey of more than 1,100 IT decision-makers, May 2026.
2. IDC, 2026 AI and Automation FutureScape; see also IDC, "The Future of AI Is Model Routing," December 2025.
3. Gartner, Predicts 2026: AI Sovereignty, October 2025.
4. Gartner, Top Predictions for IT Organizations and Users in 2026 and Beyond, October 2025.