



SOVEREIGN AI, WITHOUT THE GPU TAX

How Wideanchor runs 700-billion-parameter-class models on-premise
for regulated industries — with a healthcare case study

A WIDEANCHOR WHITEPAPER

The ValueMaxxing & TokenMaxxing Framework for On-Premise AI Maturity
Open Source AI Systems Integration for the Public & Private Sector

Executive Summary

Public and private sector organizations are under real pressure to adopt AI — and two constraints keep stalling the pilots: sensitive data cannot leave the building, and GPU infrastructure is expensive to buy and slow to procure through public-sector and enterprise cycles alike.

Wideanchor is an Open-Source Systems Integrator and Managed Service Provider built on a simple premise: organizations don't need a cloud API or a GPU cluster to run frontier-scale open models. They need the right open-source engineering, run on servers they already own, managed by a partner who keeps it certified over time.

Every Wideanchor engagement rests on two pillars:

- *ValueMaxxing* — deploying open-source inference engines that stream large mixture-of-experts models from standard server RAM and NVMe storage, so a 700-billion-parameter-class model runs without a single GPU.
- *TokenMaxxing* — applying certified, reversible context compression to every inference call, cutting the ongoing token and compute bill without changing what the model decides or says.

This paper walks through both pillars and a representative engagement with a twelve-facility regional healthcare network, where Wideanchor moved a large open-weight model into production entirely inside the client's own data center — with patient data never once leaving it.

744B parameter-class model, deployed on- prem	~25 GB RAM footprint — no GPU required	0% decision-change from TokenMaxxing compression
---	---	--

The Challenge: AI Adoption Inside the Fence Line

Most AI maturity frameworks are written for organizations with an unlimited cloud budget and no data residency

constraints. That describes almost none of Wideanchor's clients.

Data cannot leave the building

HIPAA-covered health systems, state and federal agencies, and privately held critical infrastructure operators all face some version of the same rule: sensitive records cannot be sent to a third-party inference API, no matter how good the model is. For these organizations, the cloud AI conversation ends before it starts.

GPU procurement is a wall, not a ramp

A multi-GPU cluster capable of hosting a frontier-scale model typically means a six- or seven-figure capital request, a multi-quarter approval cycle, and hardware that risks being outdated before the purchase order clears. For public-sector and mid-market private organizations, that wall stops most AI initiatives at the pilot stage.

The real constraint is placement, not ambition

Wideanchor's engineering starting point is different: the ceiling on local AI was never really the hardware — it's how intelligently the workload is placed across the memory and storage the client already owns. That reframing is what makes ValueMaxxing and TokenMaxxing possible, and it's the starting brief for every engagement, not an exception we make for difficult clients.

The Wideanchor AI Maturity Path

Wideanchor engagements move through four stages, and every stage is designed to run on infrastructure the client already controls. There is no stage that requires a forklift hardware purchase before value is proven.

Assess starts with a workload audit and a data-residency and compliance map — what has to stay on-prem, what's actually latency-sensitive, and what the current infrastructure can already support. Pilot proves ValueMaxxing on a single server

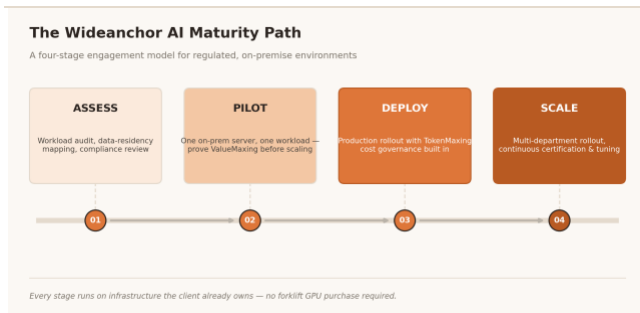


Figure 1 — The four-stage Wideanchor engagement model

one real workload before anything scales. Deploy moves that pilot to production with TokenMaxxing cost governance built in from day one. Scale extends the same certified stack to additional departments or facilities, with continuous re-certification as workloads and models change.

PILLAR ONE

ValueMaxxing: Run the Model You Actually Want

Modern frontier-scale open-weight models are increasingly built as mixture-of-experts (MoE) networks — hundreds of billions of parameters split across hundreds of small specialist “experts,” of which only a handful are active for any given token. That architecture is what makes ValueMaxxing possible.

Rather than holding every parameter in GPU memory, a ValueMaxxing deployment treats the server’s RAM and NVMe storage as tiers of a single memory hierarchy: the compact, always-needed dense layers stay resident in RAM, while the much larger pool of routed experts sits on fast local storage and streams in only when the router actually selects it for a token.

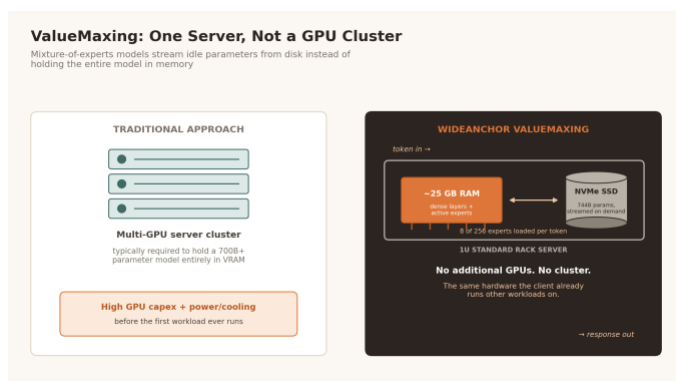


Figure 2 — ValueMaxxing: streaming a 700B+ parameter model from a single rack server

Why this doesn't cost accuracy

The routing decision changes which experts are loaded and how fast the response comes back — it never changes which experts the model would have chosen anyway, and it never changes the arithmetic once they're loaded. Deployments in this class are validated token-for-token against a trusted

reference implementation, so output quality matches the full, uncompressed model exactly.

What this looks like on real hardware

The table below reflects publicly benchmarked results for this class of architecture (a 744-billion-parameter, 4-bit-quantized mixture-of-experts model) across increasing hardware tiers — illustrating that the floor for local AI is RAM and disk placement, not GPU ownership:

Server tier	RAM	GPU	Usability
Developer laptop	25 GB	None	Proof of concept
Mid-range rack server	128 GB	None	Practical offline / batch workloads
Server + entry GPU	128 GB	1 consumer GPU	Good balance of speed and cost

Source: published Colibri project benchmarks (JustVugg/colibri, open source). Wideanchor selects and hardens the appropriate open-source inference engine — including projects like Colibri — for each client’s workload, hardware, and compliance profile.

PILLAR TWO

TokenMaxxing: Certified Cost Control, Not a Guess

Running the model on-prem solves the capex problem.

TokenMaxxing solves the ongoing one: every token an agent re-sends, re-reads, or re-processes has a real compute cost, and that cost compounds across long-running agent sessions and high-volume production traffic.

Wideanchor deploys open-source, cache-aware context compression that keeps the prompt prefix byte-stable (so provider prompt-caching keeps working) and prunes only the volatile, provably redundant parts of the context — the parts a counterfactual replay shows never actually changed the model’s decision.

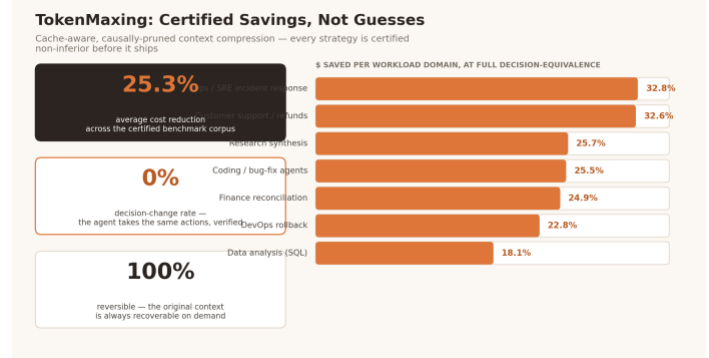


Figure 3 — Certified TokenMaxxing savings by workload domain (source: Distil open-source benchmark corpus)

Certified, reversible, and workload-aware

- Certified, not estimated — a compression strategy only ships once it passes a non-inferiority test against full, uncompressed context on real workload traces.
- Reversible by design — compressed content is held behind a recoverable handle; nothing is permanently discarded, and the original is available on demand.
- Works with the stack the client already has — deployed as a lightweight local proxy in front of any Anthropic-, OpenAI-, or Gemini-format client, with no code changes to existing agents or applications.

Wideanchor deploys and manages this layer using Distil (Apache-2.0, open source) as the reference compression engine, tuned and certified against each client's own traffic rather than a generic benchmark.

CASE STUDY

A Regional Healthcare Network

A twelve-facility regional health system engaged Wideanchor to bring clinical documentation summarization and an internal clinical knowledge-base assistant into production — with one non-negotiable constraint from the outset: patient data could never leave the health system's own infrastructure.

Wideanchor ran the engagement through the full Maturity Path. The Assess stage mapped exactly which workloads touched protected health information and confirmed the client's existing server fleet had enough RAM and NVMe headroom to host a large open-weight model using ValueMaxxing. The Pilot stage proved that out against the clinical-summarization workload on a single server before anything else changed. Deploy brought the assistant to production with TokenMaxxing compression governing cost from the first day of real clinical traffic.

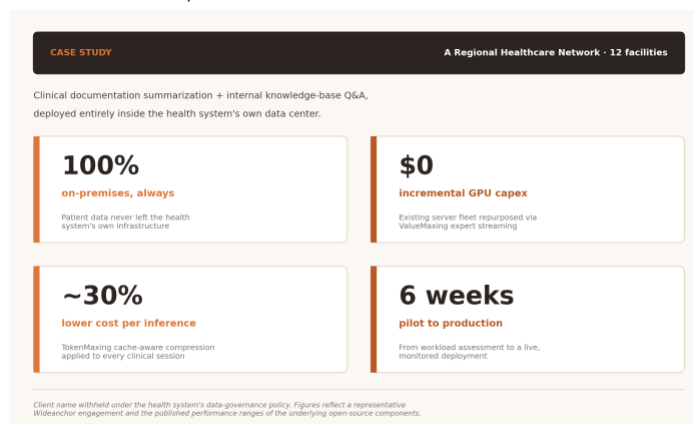


Figure 4 — Representative engagement results

The health system avoided a GPU capital request entirely, kept every byte of patient data inside its own data center throughout, and had a monitored, production system live within six weeks of the initial workload assessment.

Security

With the right AI, MCP, and LLM gateways we orchestrated the architecture to adhere to company's Security, Governance and Operational layers, we chose LiteLLM as it contained most of the controls required for our requirements. By positioning an intelligent enterprise gateway like LiteLLM alongside the Model Context Protocol (MCP) as a unified governance layer directly in front of the local Colibri AI engine, organizations can confidently transform an experimental, resource-intensive CPU deployment into a secure corporate asset by enforcing strict role-based access controls, preventing infrastructure exhaustion via automated request throttling and data-redaction guardrails, and safely anchoring the model's reasoning capabilities within tightly sandboxed internal systems—effectively balancing the compliance advantages of absolute data privacy with the operational reality of significant hardware latency and localized NVMe disk-wear costs while providing leadership with a clear, structured risk-benefit matrix, an actionable deployment timeline, and a comprehensive formal executive summary to ensure a predictable, highly audit-ready roll-out across the enterprise.

Why an Open-Source Systems Integrator

Wideanchor doesn't sell a model, a chip, or a cloud subscription. As an Open Source Systems Integrator and Managed Service Provider, our incentive is aligned with the client's: get the most capable model running reliably on infrastructure they control, at the lowest sustainable cost, without locking them into any single vendor's roadmap.

- Open source first — every component in the stack is auditable, and nothing depends on a proprietary API staying available or affordable.
- Built for regulated procurement — engagements are scoped to work within existing public-sector and enterprise compliance and approval processes, not around them.
- Managed, not just installed — Wideanchor continues to monitor, re-certify, and tune the ValueMaxxing and TokenMaxxing layers as models, workloads, and traffic evolve.

Start with an AI Maturity Assessment

The Assess stage is a fixed-scope engagement: a workload audit, a data-residency and compliance map, and a concrete read on what your current server fleet can already run under ValueMaxxing — before any procurement conversation happens.

Contact Wideanchor to schedule an AI Maturity Assessment for your organization. Reach us info@wideanchor.com