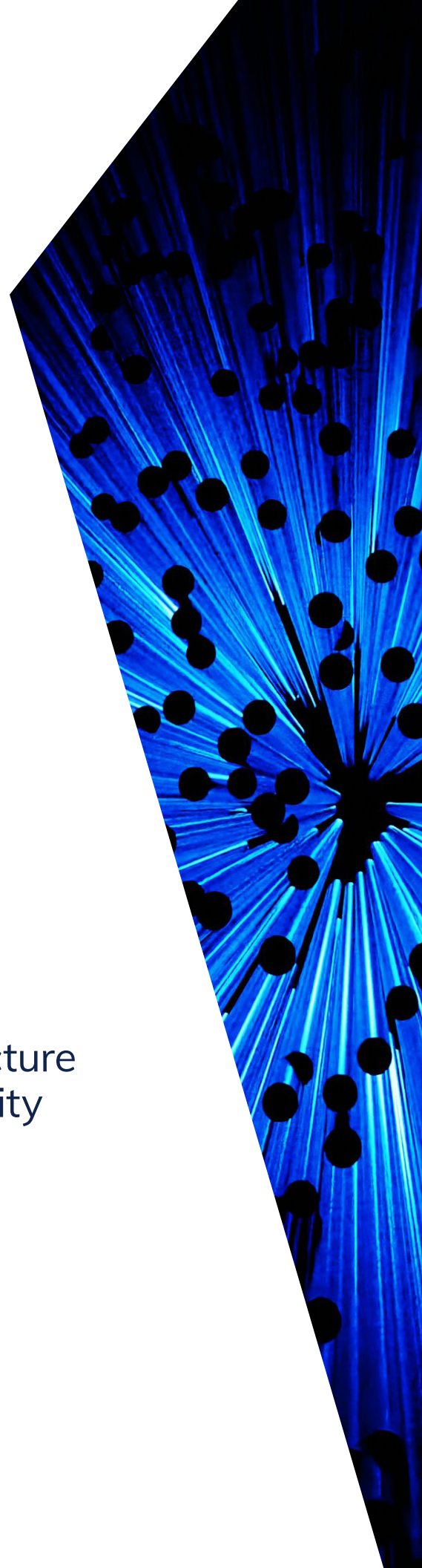




Operating the AI Factory

Governing Mission Resilience, Infrastructure
Efficiency, and Operational Accountability
Across Public Sector AI Infrastructure



Executive Takeaway

Public sector organizations are rapidly expanding AI factory infrastructure across cloud, on-premises, hybrid, sovereign, and air-gapped environments. Agencies are scaling AI-powered applications, analytics systems, and agent-driven workflows that support modernization initiatives, mission operations, citizen services, research environments, and public safety programs.

But many “AI observability” approaches simply extend traditional monitoring to GPUs by collecting infrastructure metrics, generating alerts, and reporting utilization.

That’s necessary, but not sufficient.

AI factories are not isolated infrastructure environments.

They are distributed execution systems spanning:

- AI applications and services
- Models and inference pipelines
- Kubernetes orchestration
- GPUs and CPUs
- Storage and networking
- Shared hybrid infrastructure
- Multi-agency operational environments

Most operational failures are not isolated hardware events.

They are:

- Wasted infrastructure investment
- Misallocated GPU capacity
- Infrastructure contention
- Slow training and inference
- Unstable workloads
- Degraded mission-critical and citizen-facing services
- Limited operational accountability across shared environments

The modern requirement is not GPU monitoring.

It is AI Factory Observability: a system-aware operational control plane that connects application behavior, orchestration, infrastructure conditions, workload execution, and operational impact into a unified operational model.

More telemetry does not create operational clarity. System-aware observability does.

Public Sector AI Operations Are Different

Public sector AI environments operate under constraints rarely seen in commercial organizations.

Teams must manage:

- Shared infrastructure across agencies and mission programs
- Budget accountability and procurement oversight
- Regulated, sovereign, and air-gapped environments
- Hybrid infrastructure spanning cloud and on-premises systems
- Long infrastructure lifecycle expectations
- Mission-critical service delivery requirements
- Operational resilience mandates

AI operational failures increasingly affect:

- Emergency response coordination
- Transportation modernization systems
- Public health analytics
- Research and higher education environments
- Citizen service platforms
- Defense and intelligence operations

As AI becomes a persistent public sector operating model, organizations require more than infrastructure monitoring.

They require system-aware operational governance across the entire AI factory.

The Market Misconception: “AI Observability” = GPU Monitoring

Most AI observability narratives focus on:

- GPU telemetry
- Utilization tracking
- Alerting and dashboards
- Component-level troubleshooting

This is traditional monitoring applied to GPU infrastructure.

A similar misconception exists at the AI application layer where request monitoring, latency tracking, or API visibility are treated as AI observability.

They are not.

Legacy observability tools surface symptoms at isolated layers.

AI Factory Observability identifies operational constraints across the full execution stack.

Traditional APM stops at application telemetry.

AI Factory Observability follows execution across:

- AI applications and services
- Kubernetes orchestration
- Infrastructure and shared resources
- GPU, CPU, storage, and network environments
- Cloud, on-premises, hybrid, and sovereign infrastructure

Without system-aware operational context, organizations can observe activity but cannot explain causality, operational risk, or mission impact.

The reality is that governing AI factories requires understanding three fundamental operational truths:

Reality #1: AI Infrastructure Is Now a Governance Problem

GPUs are expensive resources whether purchased on-premises or consumed through cloud providers.

For public sector organizations, the operational problem is rarely:

“Do we have metrics?”

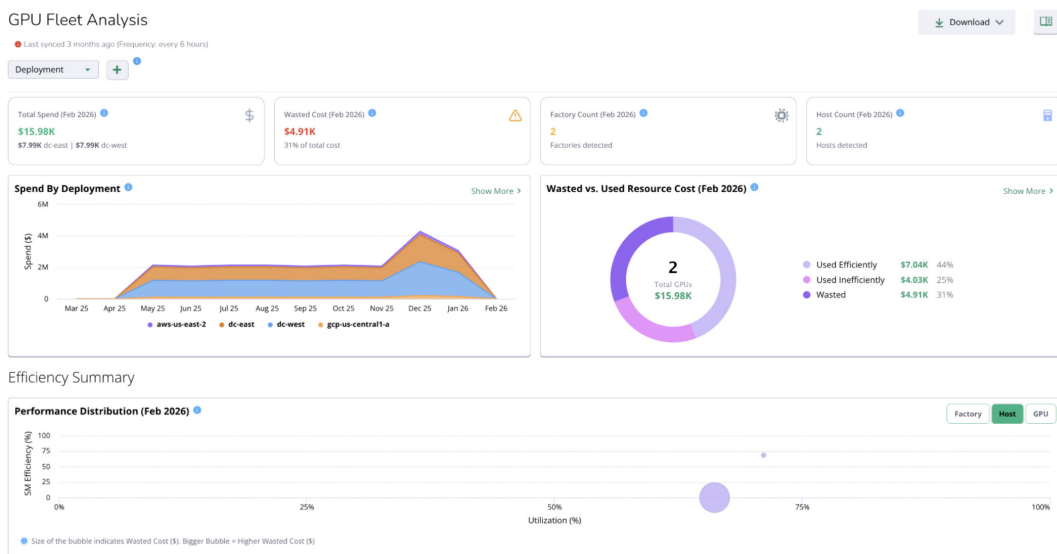
It is:

- Are scarce GPU resources allocated to the right mission workloads?
- Which agencies or programs consume the most infrastructure capacity?
- Are shared environments operating efficiently?
- Can infrastructure investments be justified with defensible operational data?
- Can organizations forecast future capacity and procurement requirements?

AI factory economics vary across environments. In cloud AI services, consumption models can directly affect operational cost. In sovereign and on-premises AI factories, operational efficiency depends more heavily on GPU allocation, orchestration behavior, workload prioritization, and infrastructure utilization.

In both cases, organizations must understand execution behavior across the system to govern AI operations effectively.

Figure 1 — AI Infrastructure Efficiency Is an Operational Governance Problem



AI factory infrastructure can appear healthy while delivering poor operational efficiency as scheduling delays, orchestration bottlenecks, data pipeline issues, and shared-resource contention leave GPUs underutilized even as infrastructure costs rise. In public sector environments, where AI infrastructure is often shared across departments and mission programs, this creates both operational and budgetary risk, making utilization efficiency, workload prioritization, and operational accountability critical. The operational question is no longer whether infrastructure is active, but whether AI infrastructure is producing measurable operational value.

Reality #2: AI Workloads Operate as Full Execution Systems

Many performance issues attributed to GPUs are actually caused by:

- Slow storage and data pipelines
- Network throughput bottlenecks
- Kubernetes orchestration constraints
- Shared infrastructure contention
- Unstable or misconfigured workloads
- Hybrid environment coordination failures

Applications are not isolated code.

AI applications operate as distributed execution systems spanning models, orchestration layers, infrastructure, storage, networking, and compute resources.

Observability must understand the behavior of the entire system behind every inference and training execution.

Without end-to-end operational context, teams may observe low utilization, failed requests, service degradation, or reduced mission responsiveness without understanding the underlying operational constraint.

This challenge becomes significantly more severe across:

- Hybrid cloud environments
- Distributed agency infrastructure
- Public safety operations
- Research and higher education systems
- Multi-vendor operational environments
- Secure and air-gapped infrastructure

AI Factory Observability requires unified operational context across:

- Applications and AI services
- Kubernetes orchestration
- Infrastructure and shared resources
- GPU, CPU, storage, and network environments
- Cloud, on-premises, hybrid, and sovereign infrastructure

Without this system-aware context, organizations cannot effectively scale or govern AI modernization initiatives.



Reality #3: Execution Behavior Defines Operational Outcomes

AI training and inference are execution-driven operational processes.

Organizations need to understand:

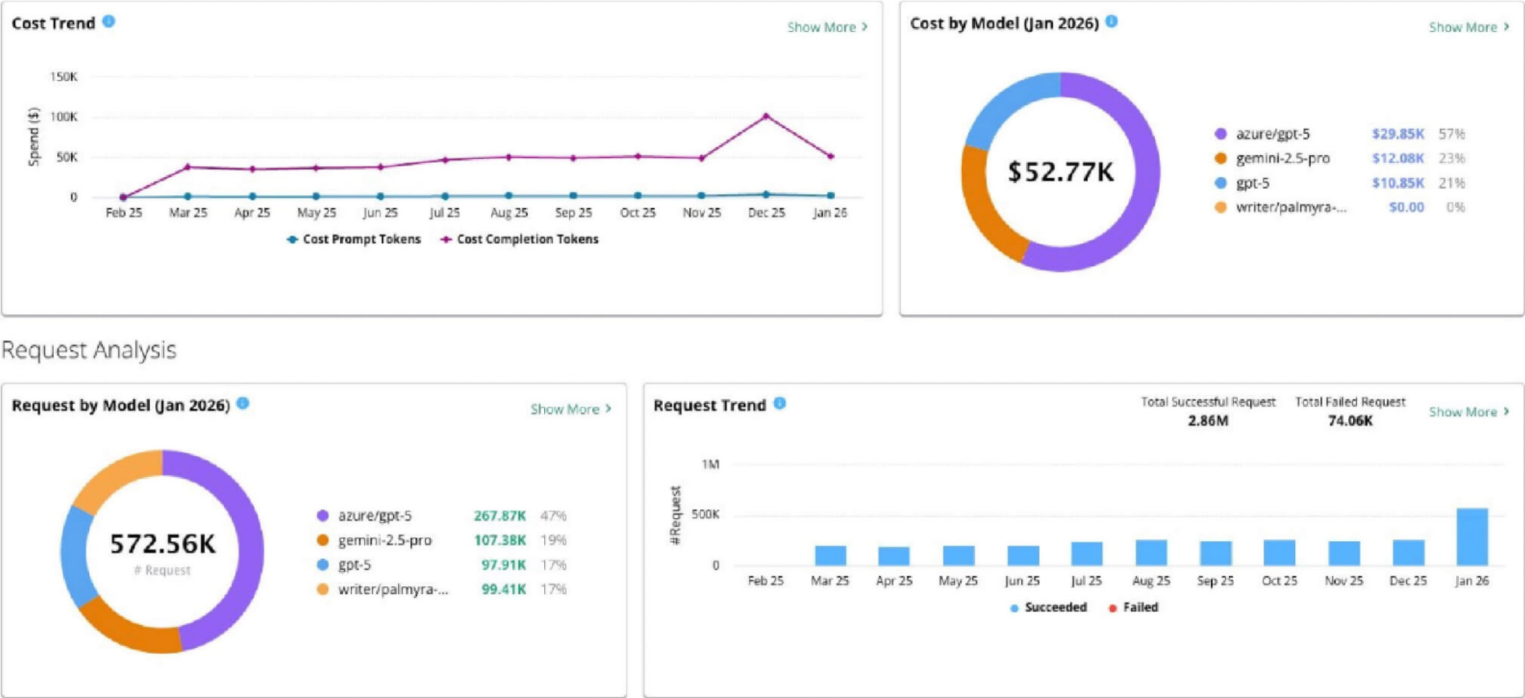
- How long jobs and requests run
- When and why failures occur
- Where bottlenecks emerge
- How workloads behave over time
- How infrastructure conditions affect mission outcomes
- Which workloads consume disproportionate resources

Monitoring infrastructure devices or requests alone does not explain:

- Operational efficiency
- AI service reliability
- Capacity consumption
- Infrastructure contention
- AI factory performance
- Mission impact

Execution-level analysis exposes how workloads, orchestration paths, and infrastructure conditions interact across the system, enabling organizations to move from reactive monitoring to operational causality.

Figure 2 — Execution-Level Intelligence for AI Factory Operations



AI cost, performance, and operational resilience are driven by workload execution behavior, not infrastructure metrics alone. Execution-level analysis exposes how workloads consume infrastructure capacity, where inefficiencies accumulate, why latency and failures occur, and how orchestration, infrastructure conditions, and AI request behavior affect operational outcomes. In cloud AI services, token consumption can directly influence cost and performance, while in sovereign and on-premises environments, execution efficiency and infrastructure utilization become the primary operational constraints.

Understanding these realities changes how organizations should operate AI factories. The result is a new operating model built on three core capabilities.

The Operating Model Shift: From Monitoring to AI Factory Observability

A modern AI factory operating model must deliver three capabilities simultaneously.

Key Capability #1: Infrastructure Governance and Capacity Planning

Organizations need operational clarity across AI infrastructure environments:

- Define operational economics for GPU consumption
- Improve infrastructure utilization
- Expose underutilization and resource waste
- Support showback and chargeback across agencies and mission programs
- Quantify capacity constraints and prioritization decisions
- Improve procurement planning and long-term forecasting

In cloud environments, organizations must manage variable consumption costs.

In sovereign and on-premises environments, organizations must maximize utilization of fixed, high-value infrastructure assets.

AI Factory Observability transforms AI infrastructure from opaque spending into governed operational systems.

Key Capability #2: System-Aware Observability Across the Full Stack

AI workloads operate across a connected execution stack:

- AI applications and services
- Kubernetes orchestration and scheduling
- Infrastructure including GPU, CPU, storage, and networking
- Hybrid and sovereign operational environments

AI observability must connect workload behavior to orchestration decisions and infrastructure conditions.

Without end-to-end operational context, teams cannot identify true constraints, understand downstream impact, or maintain resilient mission-critical services.

More dashboards do not solve this problem. System-aware operational intelligence does.

Key Capability #3: Execution-Level Intelligence and Operational Causality

Operations teams must measure and compare workload executions:

- Duration, failures, and error patterns
- Resource contention during execution
- Infrastructure conditions correlated to outcomes
- Utilization and scheduling efficiency
- Workload behavior across shared environments

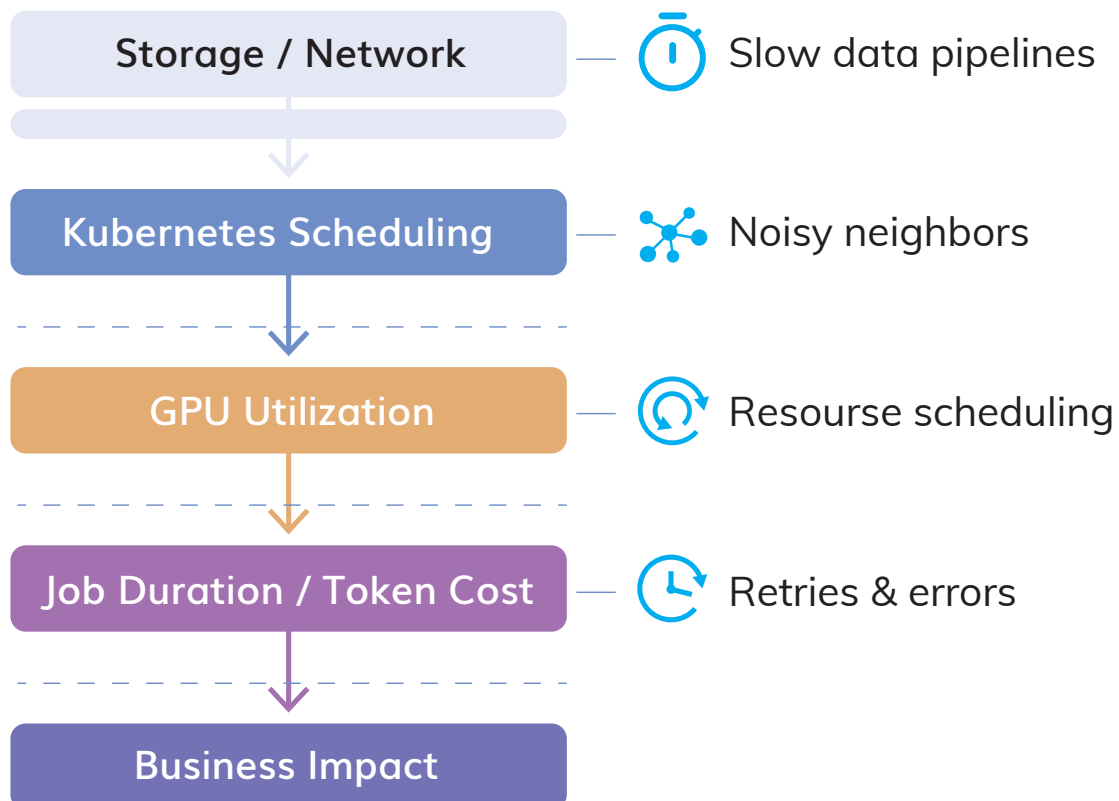
This is how organizations move from reactive troubleshooting to governed AI factory operations.

Modern AI factory operations also require intelligent operational reasoning across infrastructure, orchestration, workloads, and execution behavior.

Agentic AI systems can correlate telemetry across the stack, identify operational constraints, explain causality, accelerate decision-making, and improve mission resilience.

Figure 3 — AI Factory Outcomes Are Driven by Full-Stack Causality

Root Cause Analysis - Full Stack Context



In AI factory environments, operational outcomes are shaped by interactions across the stack, from storage and scheduling to orchestration behavior, infrastructure utilization, and workload execution. Without end-to-end operational context, organizations can observe symptoms but cannot explain cost, performance, reliability, mission impact, or public service degradation. Full-stack causality is what makes AI factory operations governable at scale.

Why This Matters Now

AI factory infrastructure has become a strategic operational concern for public sector organizations.

The economics are unforgiving.

GPU infrastructure, AI workloads, and modernization programs can quickly become major operational and budgetary challenges across both cloud and on-premises environments.

At the same time, the systems are too interconnected to troubleshoot manually.

When performance depends on orchestration, workload behavior, infrastructure utilization, and shared operational conditions, more dashboards do not create operational clarity.

Many public sector organizations now operate AI infrastructure under:

- Procurement oversight
- Multi-year funding constraints
- Shared agency resource models
- Resilience mandates
- Sovereign infrastructure requirements
- Pressure to demonstrate measurable modernization outcomes

Organizations that treat AI factories as extensions of traditional monitoring will continue to face fragmented visibility, infrastructure inefficiency, capacity contention, and operational risk across mission-critical services.

Organizations that treat AI factories as governed operational systems built on system-aware observability and operational causality will scale modernization initiatives with greater confidence, resilience, and operational efficiency.

Bottom Line

GPU monitoring explains what infrastructure devices are doing.

Application monitoring explains whether services are responding.

AI Factory Observability explains how the entire execution system behaves:

- Where capacity is constrained
- Why performance degrades
- How infrastructure conditions affect mission outcomes
- Which operational constraints impact mission-critical services
- How shared infrastructure affects modernization priorities
- What actions should be taken across infrastructure, orchestration, and workload execution

As AI becomes a persistent public sector operating model, organizations will require system-aware observability that moves operations teams from reactive monitoring to governed AI factory operations.

The organizations that succeed will be the ones that understand the system, prove operational causality, optimize infrastructure efficiency, and operate resilient AI environments across cloud, on-premises, hybrid, sovereign, and air-gapped infrastructure as one unified operational platform.

ABOUT VIRTANA

Virtana delivers the deepest and broadest observability platform for hybrid and multi-cloud, with full-stack AI observability spanning applications, services, data pipelines, GPUs, CPUs, networks, and storage. Powered by high-fidelity data and agentic AI, Virtana provides unmatched visibility across end-to-end IT services and AI workloads, correlating health, performance, cost, and user impact in real time.

With advanced event intelligence and autonomous insight generation, Virtana delivers clarity no other provider can match. Trusted by Global 2000 enterprises and public sector organizations, Virtana helps IT operations and DevOps teams reduce risk, strengthen resilience, improve efficiency, and modernize with confidence across multi-cloud, on-premises, and edge environments.

Learn more at virtana.com

info@virtana.com | **+1 408-579-4000** | **virtana.com**

©2026 Virtana. All rights reserved. Virtana is a trademark or registered trademark in the United States and/or in other countries. All other trademarks and trade names are the property of their respective holders.