



Enterprise AI Governance

CONFIDENTIAL

From Autonomous Agents to Governed Operations

RuntimeAI — AI Security, Control and Governance Platform

Audience: Enterprise AI Leaders, CISOs, Operations & Compliance Teams

Classification: Confidential — For Recipient Use Only

The Governance Gap in Enterprise AI

Every enterprise AI program follows the same arc. It begins with automation — rule-based, predictable, human-supervised. It progresses to AI-assisted decision-making, then to autonomous agents executing decisions at machine speed. At each stage, productivity gains compound.

The problem appears at Stage 4 and Stage 5 — when agents stop recommending and start acting.

At that point, the enterprise needs answers to four questions it often cannot answer:

- **Who** is this agent, and is it authorized to act on our behalf?
- **What** is it permitted to do — and what is explicitly off-limits?
- **What** data is it accessing, and is that exposure compliant?
- **What** did it do, and can we prove it under audit?

Without answers, autonomous AI is not an efficiency gain — it is an uncontrolled liability. RuntimeAI is the governance layer that answers all four questions, in real time, before any agent action executes.

The Four Pillars of Agent Governance

1. Agent Identity — Know Who Is Acting

Every AI agent in a governed environment must carry a verifiable, cryptographic identity. Anonymous agents — processes that take actions without traceable authorization — create accountability gaps no compliance framework can tolerate.

RuntimeAI assigns a persistent identity to every agent. That identity travels with every action the agent takes, across every system it touches. No agent acts without a verified identity. No action goes unattributed.

Capability	Business Outcome
Per-agent cryptographic credentials	Every action is attributed — no anonymous agents
Just-in-time entitlements	Agents receive only the permissions required for the current task
Instant identity revocation	A compromised or misbehaving agent is isolated in real time
Cross-system attribution	Agent actions are traceable across all enterprise systems — ERP, CRM, WMS, email, cloud

2. Policy Guardrails — Define the Boundaries of Autonomy

Autonomous agents operate at machine speed. So does RuntimeAI's policy enforcement. Every agent action — every system update, every outbound request, every data access — passes through the inline policy engine before it executes. Actions within policy proceed without delay. Actions outside policy are blocked, redirected, or escalated to a human reviewer.

Policy enforcement operates across three dimensions:

- Behavioral policy** — What is this agent allowed to do? Guardrails define permitted actions by agent type, workflow context, and data classification. A finance agent authorized to read accounts payable records has no authority to modify payment instructions — even if the underlying system permits it.
- Scope policy** — Where is this agent allowed to operate? System scope and organizational scope are enforced at the infrastructure level. An agent scoped to customer support operations cannot reach HR systems or financial records, regardless of what it is instructed to do.
- Temporal policy** — When is this agent allowed to act? Time-bound authorization ensures that agents operating in elevated-privilege windows are subject to tighter controls, with automatic de-escalation when the window closes.

3. Data Guardrails — Control What the AI Sees and Shares

Enterprise AI agents operate across systems that hold the organization's most sensitive data: customer PII, financial records, regulated health information, proprietary contracts. Without active guardrails, that data flows freely into AI model inputs and outbound requests — creating exposure that violates HIPAA, GDPR, SOC 2, EU AI Act, and CMMC requirements.

RuntimeAI enforces bidirectional data protection on every channel the AI layer uses:

- Inbound protection** — Sensitive data is identified and redacted or masked before it reaches the AI model. Agents receive the information they need to make a decision — and nothing they do not.
- Outbound protection** — Agent outputs and outbound requests are inspected before execution. Data that should not leave the organization is blocked at the point of egress — not after the fact.
- Context-aware classification** — The same customer record that a support agent is permitted to read may not be accessible to an operations agent. Classification follows the workflow, not just the data field.

Compliance frameworks mapped: HIPAA · GDPR · SOC 2 · EU AI Act · CMMC 2.0 · NIST AI RMF · ISO 42001

RuntimeAI does not pursue third-party certifications. It provides its own compliance documentation — CISA SSDF self-attestation, SOC 2 Trust-Service-Criteria control mappings (documentation, not a certificate), an independent-style security review, a threat model, architecture documentation, and a current SBOM — to support the customer's own assessment.

4. Audit Trail — Prove What Happened

At Stages 4 and 5, AI agents make consequential decisions on behalf of the enterprise. Every decision must be auditable — not just logged, but immutable, tamper-proof, and structured for compliance review.

RuntimeAI maintains a complete audit record of every agent action:

- Every policy evaluation: what the agent requested, what policy decided, and what the outcome was
- Every data access: which records were read, which were modified, and under what authorization
- Every system interaction: database writes, API calls, document updates, workflow transitions
- Every escalation: when an agent reached a policy boundary and deferred to a human reviewer

This record cannot be modified, deleted, or overwritten — by the agent, by an administrator, or by the platform. It is the compliance evidence layer for the autonomous enterprise.

Operational Safety: The Kill Switch

No governance framework is complete without an emergency stop. RuntimeAI provides a **sub-50ms kill switch** — the fastest agent termination capability in the market — operable at three levels:

Level	Scope	Use Case
Single agent	Terminate one agent instance	Anomalous behavior detected on a specific workflow

Level	Scope	Use Case
Agent group	Terminate all agents of a type	Pattern of policy violations across a class of agents
Full tenant	Terminate all AI operations	Security incident, compliance hold, or emergency stop

The kill switch is available from the dashboard, the RuntimeAI API, and as an automated policy trigger. When an agent's behavior crosses a defined risk threshold measured against its established baseline, the system invokes the kill switch automatically — without waiting for human intervention.

Extended Capabilities

AI Integration Fabric — Governed Middleware for All AI Traffic

As enterprise AI environments grow — adding models, tools, third-party integrations, and external data sources — the governance surface expands. RuntimeAI's **AI Integration Fabric** is the governed middleware layer at the center of this network. Every AI tool call, model request, and integration event flows through a single, policy-enforced control point, regardless of what is on either end.

Without AI Integration Fabric	With AI Integration Fabric
Each integration manages its own access controls	Single policy layer governs all integrations consistently
Security gaps appear as new tools are added	New tools inherit the governance framework automatically
Audit trail is fragmented across systems	All integration events log to one unified record
Shadow AI — unregistered tools accessing enterprise data	Every AI connection is registered, verified, and monitored

QuantoSign — AI-Governed Document Execution

As agents reach Stage 4 and Stage 5 autonomy, they initiate and route consequential documents: contracts, agreements, compliance records, approvals. QuantoSign is RuntimeAI's AI-governed electronic signature platform — extending the governance layer into the document execution step.

Every agreement initiated or routed by an AI agent carries a compliant, auditable signature trail. Policy-enforced signing workflows ensure no agent bypasses human sign-off where the business requires it. Every signature event is captured in the same compliance record as agent actions — creating an end-to-end governance chain from agent decision to executed document.

Compliance alignment: ESIGN Act · UETA · EU eIDAS · sector-specific requirements

Coding Agent Defense — Securing the AI Development Pipeline

AI coding agents have elevated access to enterprise systems: source code, configuration, deployment pipelines, infrastructure credentials. An ungoverned coding agent is one of the highest-risk entry points in the enterprise AI stack.

RuntimeAI governs coding agents with the same identity, policy, and audit framework applied to operational agents:

- **Scoped credentials** — coding agents receive time-bound, least-privilege access to the systems they need
- **Output inspection** — code and configuration changes are reviewed against policy before they are committed or deployed
- **Development audit trail** — every change initiated by a coding agent is logged to the immutable audit record, creating a complete development provenance chain for security and compliance review

Governance Across the AI Maturity Curve

RuntimeAI adds governance value at every stage of the AI maturity model — not just at Stage 4 and 5. Governance investment compounds as AI capability grows, rather than requiring a retrofit after autonomous operations are already in production.

Maturity Stage	Governance Capability Activated
Stage 1 — Manual	Foundational data classification; access controls baseline
Stage 2 — Automated	Audit trail for rule-based automations; data classification across connected systems
Stage 3 — AI-Assisted	Data guardrails on AI model inputs and outputs; policy review on AI-generated recommendations
Stage 4 — Governed Agents	Full agent identity, inline policy enforcement, kill switch, behavioral monitoring
Stage 5 — Autonomous	Automated policy triggers, continuous behavioral baseline, compliance evidence package

Deployment Flexibility

RuntimeAI deploys over any existing AI infrastructure with no rip-and-replace. It operates across:

Deployment Model	Details
SaaS (fully managed)	RuntimeAI hosts and operates the full platform — fastest time to governance
Hybrid	RuntimeAI control plane managed; agent traffic stays in customer environment
On-Premises (recommended for regulated data)	Full platform deployed inside the customer's private infrastructure — preferred posture so the customer's own controls and compliance boundary govern the data
Air-Gapped (recommended for regulated data)	Full offline deployment for classified, regulated, or critical infrastructure environments — no data leaves the boundary

Works with any cloud provider, any LLM, and any existing enterprise system stack. No vendor lock-in. For sensitive, regulated, or sovereign workloads, on-prem / air-gapped deployment is the recommended and preferred model.

Why RuntimeAI

Differentiator	What It Means
Sub-50ms Kill Switch	Terminates any rogue agent mid-execution before damage occurs
Inline policy enforcement	Every agent action evaluated before it executes — never pass-through
No rip-and-replace	Deploys over existing infrastructure in days, not quarters
Vendor-neutral	Works with any cloud, any LLM, any enterprise system
Air-gap ready	Full platform available for classified and critical infrastructure environments
Compliance automation	Auto-generates evidence packages — SOC 2, HIPAA, EU AI Act, CMMC 2.0, NIST AI RMF

Get Started

Free trial — up and running in 15 minutes: trial.runtimeai.io

Executive overview: www.runtimeai.io

Roshan Shaik

Founder & CEO, RuntimeAI

roshan@runtimeai.io

+1 650-868-7724

<https://calendly.com/roshan-runtimeai/30min>

RuntimeAI, Inc. · www.runtimeai.io · trial.runtimeai.io

Confidential — prepared for recipient use only