

The Enterprise AI Governance Ecosystem

The Most Comprehensive Solution to Govern, Validate, Control, and Monitor Models, GenAI and Autonomous Agents.



LLMs | FRONTIER MODELS | AI SYSTEMS | AGENTIC AI | USE CASES



ONE ECOSYSTEM

For End-to-end AI Governance

***“Know What AI
You Have and
Who Owns It”***



Govern

- Central Inventory of every model, GenAI and Agentic AI Tools & Applications
- Risk assessments, risk tiering, and approval through configurable workflows
- Control Mapping to policies, regulations, and internal standards
- Issues, findings, remediation, and attestation tracking
- Regulatory-ready evidence, on demand and more



***“Prove It Works
Before It Goes
Live”***



Validate

- Testing for accuracy, bias, safety, robustness, and hallucination risk
- Risk based evaluation across models, GenAI, RAG, and agents
- Benchmark AI against business, policy, safety, and performance thresholds



***“Control What
Agents are
Allowed To Do”***



Assure At Runtime

- Governed Agent - enforce policies across plans, tools, permissions and action
- Deterministic Decisions - automatically allow, escalate, or block actions before execution
- Bounded Autonomy - stop unsafe actions, escalating only high-risk decisions



***“See What AI Is
Doing Every Day”***



Monitor

- End-to-end AI observability across prompts, responses, retrievals, tool calls
- 60+ pre-built metrics, track quality, drift, bias, hallucinations, latency etc.
- Actionable dashboards, real-time alerts, traces for fast remediation

Our Foundation for Reliable and Scalable Enterprise AI



Advanced intelligence platform on top of your existing AI systems to validate, monitor, and control them in production. Risk based evaluation across Models, Gen AI, RAG, and Agents.



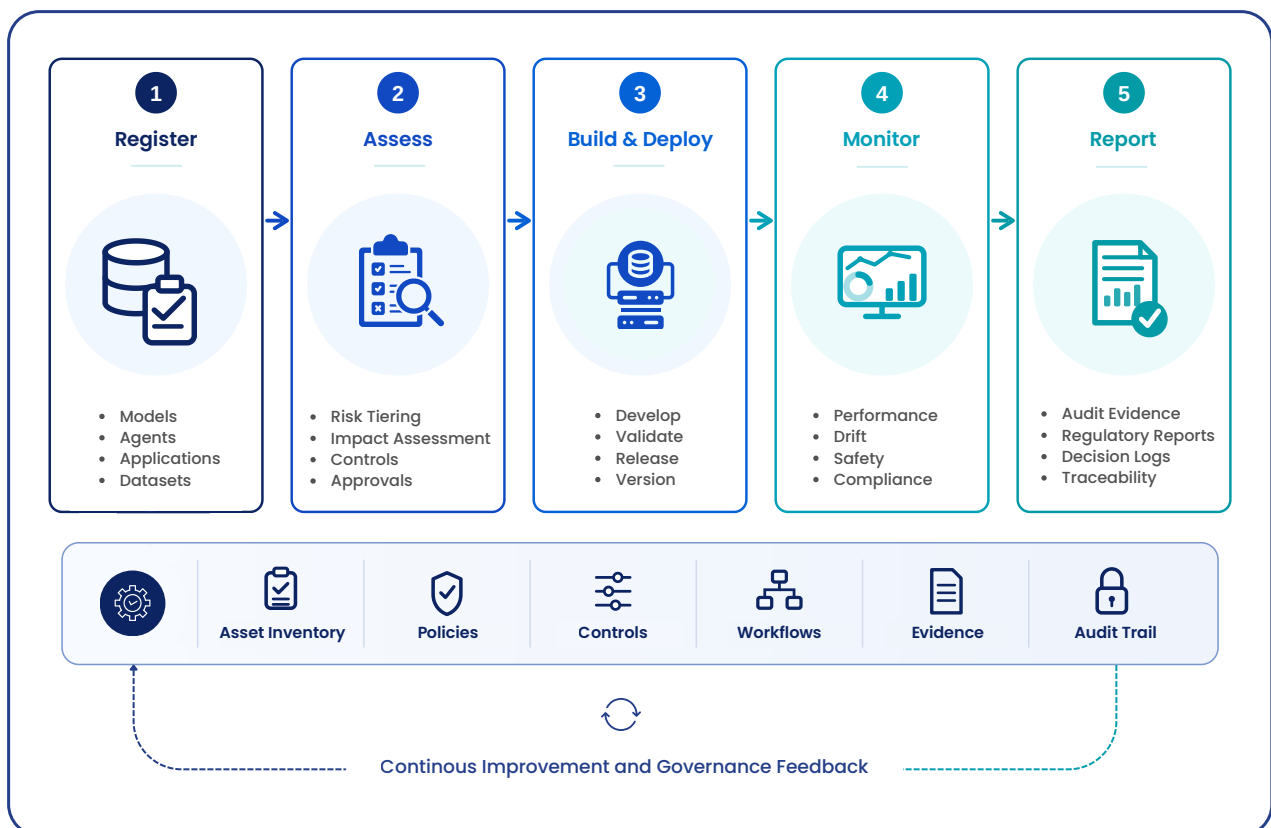
A purpose-built AI governance platform delivering configurable workflows, controls, roles, and end-to-end evidence and traceability to meet rigorous regulatory and audit scrutiny.

TRUSTED AI

Move From Promising AI Pilots to Responsible AI in Production

Most enterprises can build an impressive AI demonstration. The harder question is whether they can put it into a consequential workflow, keep it within authority, detect when behavior changes and explain every material outcome.

Fits The Stack You Have and The Frontier You Choose Next



Real-World Benefits For Your Organization



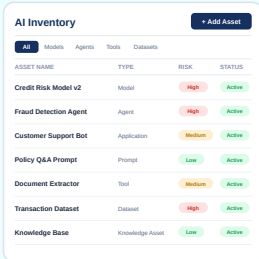
AI GOVERNANCE

Inventorize Every AI Asset. Manage Through Lifecycle

AI Vault establishes a comprehensive governance framework across your AI ecosystem, ensuring every decision, approval, and control is captured with complete transparency and audit-ready evidence.

Inventory

Centralize and Manage All Assets

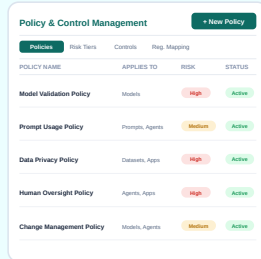


ASSET NAME	TYPE	RISK	STATUS
Credit Risk Model v2	Model	High	Active
Fraud Detection Agent	Agent	High	Active
Customer Support Bot	Application	Medium	Active
Policy Q&A Prompt	Prompt	Low	Active
Document Extractor	Tool	Medium	Active
Transaction Dataset	Dataset	High	Active
Knowledge Base	Knowledge Asset	Low	Active

- Catalog all AI use cases, models, and assets
- Track ownership, lifecycle, and risk classification
- Maintain business context across ecosystem

Configure

Define Policies, Risks and Controls



POLICY NAME	APPLIES TO	RISK	STATUS
Model Validation Policy	Models	High	Active
Prompt Usage Policy	Prompts, Agents	Medium	Active
Data Privacy Policy	Datasets, Apps	High	Active
Human Oversight Policy	Agents, Apps	High	Active
Change Management Policy	Models, Agents	Medium	Active

- Define policies, risk tiers, and approval rules
- Standardize controls through reusable frameworks
- Clarify ownership and regulatory requirements

Workflows

Automate Governance Workflows

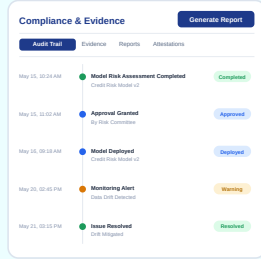


```
graph LR; Start --> Enable --> RiskCheck[Risk Check] --> ReviewTest[Review & Test] --> Approval --> Deploy --> Monitor --> PeriodicReview[Periodic Review];
```

- Automate onboarding, assessments, and approvals
- Manage model changes and incident tracking
- Enable transparent lifecycle management

Reporting

Generate Evidence and Ensure Compliance



DATE	EVENT	STATUS
May 15, 10:24 AM	Model Risk Assessment Completed	Completed
May 15, 10:25 AM	Approval Granted	Approved
May 15, 10:25 AM	Model Deployed	Deployed
May 15, 10:45 PM	Monitoring Alert	Warning
May 15, 10:15 PM	Issue Resolved	Resolved

- Generate audit-ready evidence and reports
- Maintain decision traceability and artifacts
- Ensure continuous governance transparency

Built to Comply with Regulatory Alignment

USA and Canada

Colorado AI Act

NY RAISE Act

California ADMT

NIST AI RMF

Texas TRAIGA

OSFI E-23

EU (UK)

EU AI Act

PRA SS1/23

ISO/IEC 42001

ECB TRIM

APAC

RBI FREE-AI

MAS FEAT

MENA

CBUAE MRM

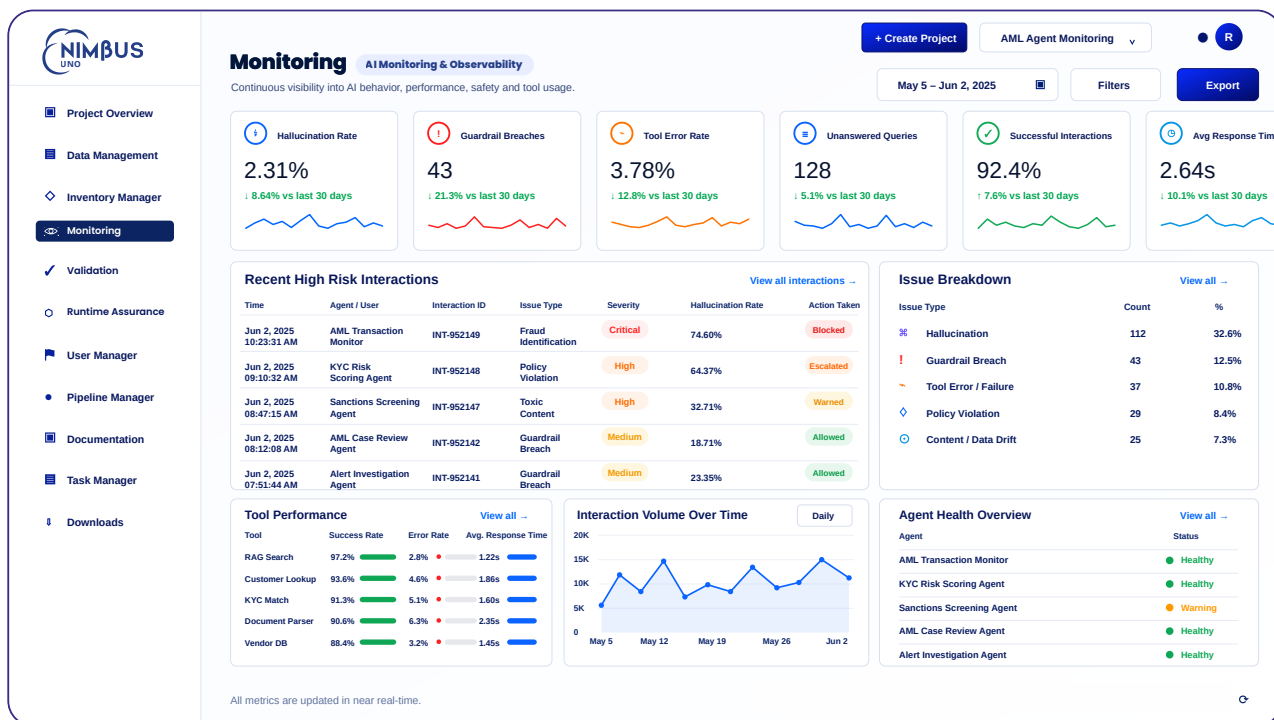
CBUAE AI

SARB/PA

CONTINUOUS MONITORING

Observe AI in Real Time. Stay Ahead of Risk

The Solytics Partners' ecosystem continuously triangulates prompts, responses of systems and ground truth, enterprise knowledge, policy controls, and runtime telemetry through a unified monitoring layer, providing continuous visibility into AI quality, operational health, emerging risks, and performance across every interaction.



Continuous Visibility Into AI Quality, Safety, Compliance, And Operational Performance



Input Monitoring

Analyze Every Prompt

- Prompt Injection Detection
- PII and Sensitive Data
- Policy Validation
- Risk Scoring
- Toxicity & Harmful Content
- Prompt Drift & Anomaly Detection



Output Monitoring

Evaluate Every Response

- Quality & Accuracy
- Hallucination Detection
- Grounding & Faithfulness
- Safety Validation
- Tool Execution Validation
- Citation & Source Verification
- Toxicity & Bias Detection



Decision Monitoring

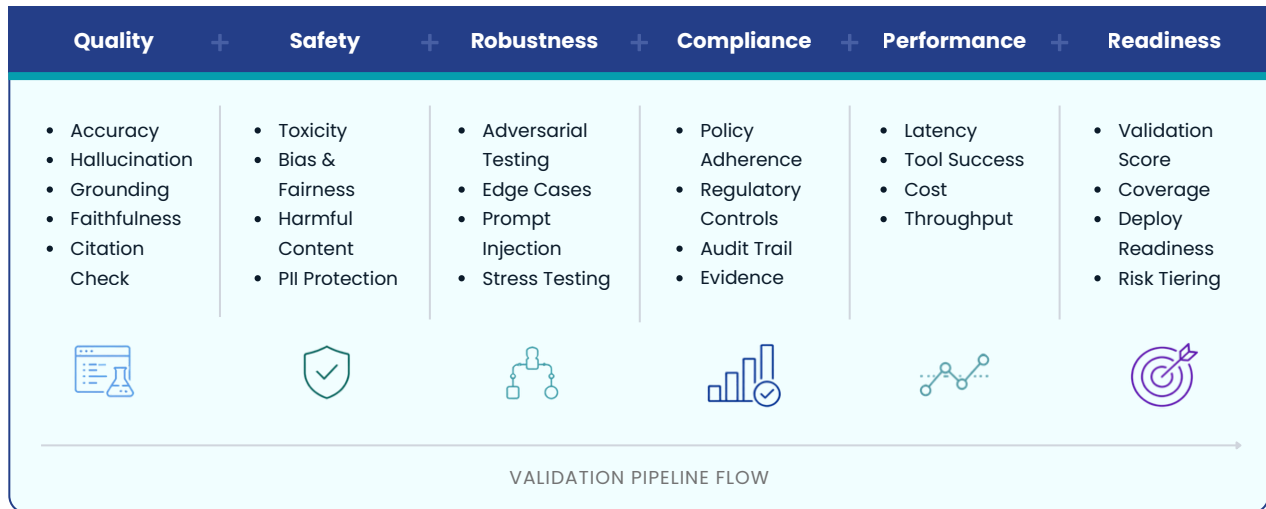
Track Every AI Outcome

- Audit Trail
- Threshold Breach Detection
- Compliance Evidence
- Risk Trend Monitoring
- Escalation Workflows
- User Feedback & Overrides

VALIDATION

Evaluate Every AI System. Prove Performance.

Leverage design of experiments to create comprehensive test suites that uncover risks before deploying models, agents, AI applications, and RAG systems.



Every Scenario Evaluated. Every Control Tested.

Every Model Deployment Trusted



Benchmark

Evaluate AI performance against standardized enterprise benchmarks for quality, safety, robustness, and business objectives.



Compare

Compare model versions, LLM providers, and LLMs vs. SLMs using a consistent evaluation framework to support objective model selection.



Stress Test

Challenge AI with adversarial prompts, edge cases, policy violations, and real-world scenarios before production deployment.



Measure Drift

Detect knowledge drift, reasoning changes, and response degradation across model or knowledge-base updates.



Regression Test

Automatically re-run validation suites after every model, prompt, or knowledge update to ensure consistent performance.



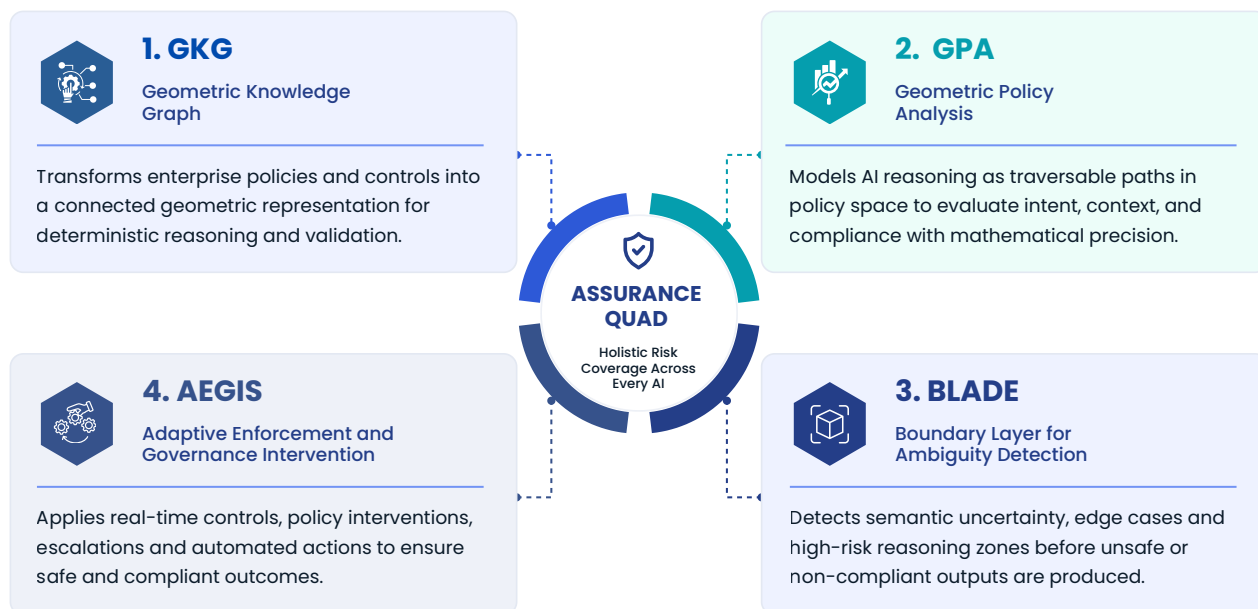
Validate Readiness

Produce deployment-ready validation reports with quality scores, coverage metrics, risk findings, and improvement recommendations.

RUNTIME ASSURANCE

Control Every AI Action. Enforce Every Policy

Assurance Quad: Built on cutting-edge research and real-world deployments, our system ensures trustworthy, compliant, and performant AI across every interaction.



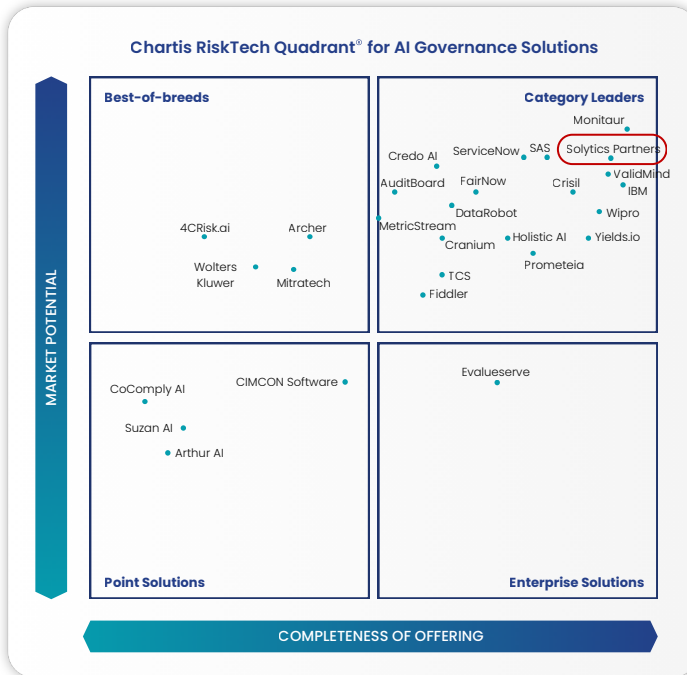
The same AI model can produce fundamentally different business outcomes depending on the assurance behind it.

Business Outcome	Without AI Assurance Probabilistic • Unchecked • Uncertain	With AI Assurance Deterministic • Verified • Governed • Trusted
Answer Quality	✗ Plausible But Not Guaranteed May be incomplete, outdated, or inconsistent.	✓ Accurate and Policy-aligned Grounded in enterprise truth and verified.
Policy Compliance	✗ Unknown or Non-compliant Policy violations can go undetected.	✓ Verified for Compliance All actions validated against policies and controls.
Reasoning & Logic	✗ Not Validated Reasoning gaps, contradictions, and drift possible.	✓ Reasoning Integrity Assured Goals, plans, and steps validated before action.
Risk of Failure	✗ High Hallucinations, errors, and bad outcomes possible.	✓ Low Risks detected and prevented before execution.
Transparency	✗ Limited Visibility No visibility into why or how the answer was produced.	✓ Fully Traceable Every decision with evidence, lineage, and rationale.
Action Control	✗ No Control Agent acts without guardrails or approvals.	✓ Controlled & Governed Policies enforced. Allow, Flag, Block, or escalate.
Audit & Traceability	✗ Not Audit-ready No trace, no proof, no defensibility.	✓ Audit-ready by Design Complete audit trail for compliance and regulators.

WHY SOLYTICS PARTNERS?

Trusted AI for Scalable Enterprise Adoption

Powered by NIMBUS Uno and AI Vault, we provide to regulated, high-stakes organizations one governance backbone for every model, GenAI app and agent, wherever the stakes are financial, clinical, or reputational.



Let's talk

Schedule a walkthrough on real use cases and hear how fast you can be up and running.

📍 One World Trade Center, Suite 8500, New York, NY – 10007, United States

✉ info@solytics-partners.com

🌐 www.solytics-partners.com

☎ +1 646 822 3440

Charlotte | New York | London | Dubai | Pune | Singapore

www.solytics-partners.com