

SOLUTION BRIEF

Secure Any Agent & AI Application with Lasso Security

AppSec teams are being asked to secure agentic applications, **but how do you secure something that is non-deterministic by nature?**



At build time

AppSec teams struggle to put agentic-specific security measures in place since they lack out-of-the-box policies or custom frameworks designed specifically for AI risks and threats.



During testing

AppSec teams are unable to test agentic runtime behavior, have no reliable way to understand what the refusal space is of the application's model, and security assumptions are made without properly stress-testing these agentic applications against adversarial conditions.



At runtime

AppSec teams struggle to know if an agentic application is staying within its intended scope, or whether its intent has been manipulated or has quietly drifted from its original design.

The instinct is to layer on additional security measures, but policies built for traditional applications introduce serious latency and generate false positives at a rate that will break production.

What's needed is something purpose-built for AI that can answer four foundational questions:

01

What agents and MCPs exist in my environment and what is their goal?

03

What can I fix and how should I prioritize?

02

What permissions and risks do they carry?

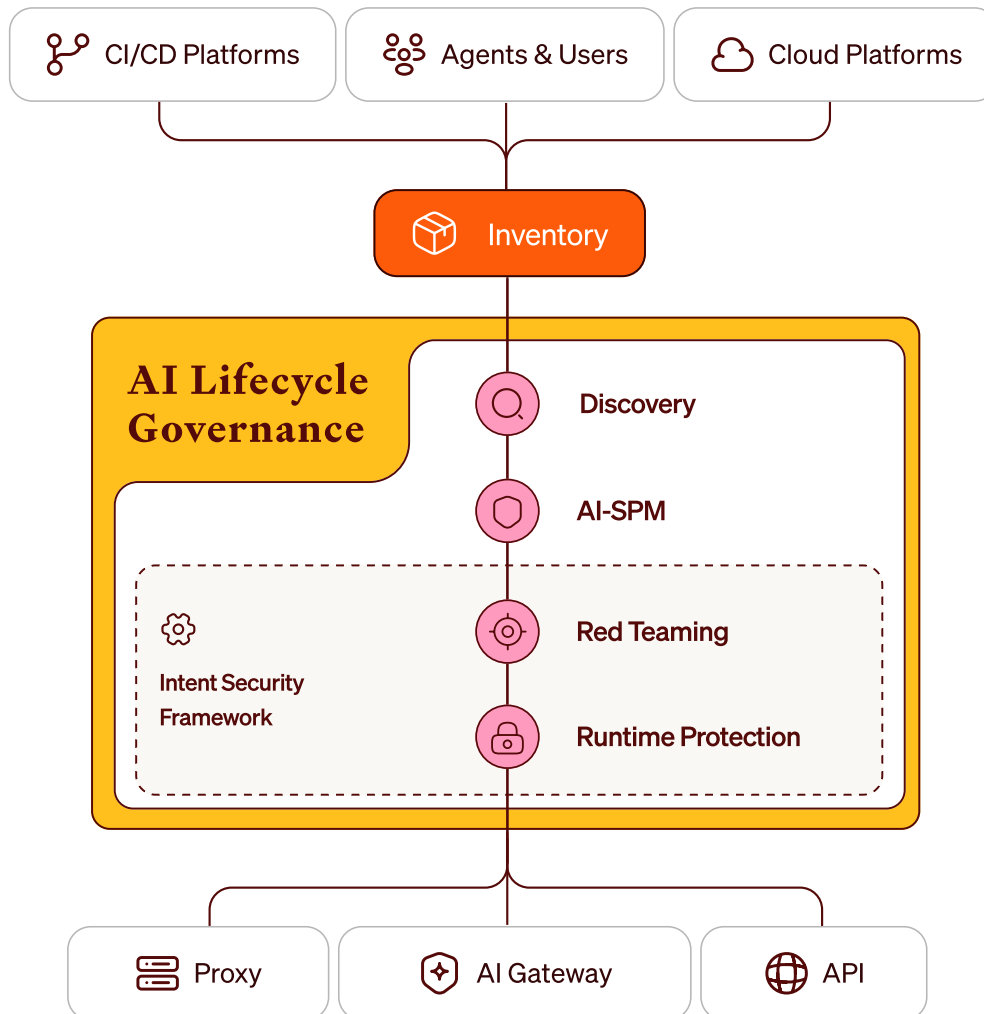
04

How do I protect them at runtime?

Lasso Security

Full Lifecycle Governance for Agentic Applications

Lasso is the AI Security Platform purpose built for the agentic era. By connecting discovery, posture management, red teaming, and runtime enforcement in a single continuous loop, Lasso ensures every agentic application behaves within its intended scope, at every stage of its lifecycle.



Discovery

Know What You Have

Lasso connects to CI/CD pipelines to automatically discover and inventory every homegrown AI application, and integrates directly with AWS Bedrock, Google Vertex AI, Salesforce, and other cloud and third-party platforms for visibility into AI agents. Each AI system is profiled across its model, system prompt, tools, guardrails, policies, and configurations, and the inventory stays current on every change.

AI-SPM

Know Your Posture

Lasso conducts a static analysis of every application's configuration, policy coverage, and third-party dependencies and identifies any policy gaps. In addition, it maps each agentic application to its coverage of major frameworks, including NIST, OWASP, and MITRE.



Red Teaming

Know Your Risks

Lasso's automated AI red teaming covers the complete kill chain from reconnaissance to exploitation. It proactively discovers exploitable vulnerabilities through automated reconnaissance and adversarial testing purpose built for agentic applications.

→ Static attacks draw from a 300K+ payload library with 100% MITRE and OWASP LLM and Agentic Top 10 coverage, running a comprehensive sweep of known jailbreak patterns, content moderation bypasses, and obfuscation techniques.

→ Dynamic attacks use multi-turn and continuous probing to test how an application holds up across extended adversarial sequences, not just a single interaction.

→ High-agency attacks deploy extremely customized, bespoke attack techniques through probing tailored specifically to the intent and design of each application.

- ✓ Offensive agents with specialized role
- ✓ Complete attack chain generation
- ✓ Attack surface visualization
- ✓ Continuous automated testing (CI) with every change
- ✓ Optional reconnaissance for each application

300K+

Attack payloads

100%

OWASP LLM & Agentic Top 10

500+

New variants weekly

50+

Evasion techniques



Runtime Protection

Know You're Protected

Lasso offers policy enforcement and AI threat protection at the proxy, API, or AI Gateway layer. Protection adapts as the applications evolve and as new capabilities are added. When an attack hits production, whether a jailbreak, a prompt injection, or any other AI threat, Lasso blocks it in real time and sends an immediate alert with full context, including what happened, which application was targeted, what the impact is, and what to do next.

Full application

Coverage

1.4%

False positive rate

98.6%

Threat detection accuracy

Immediate

Time to respond (MTTR)

<5sec

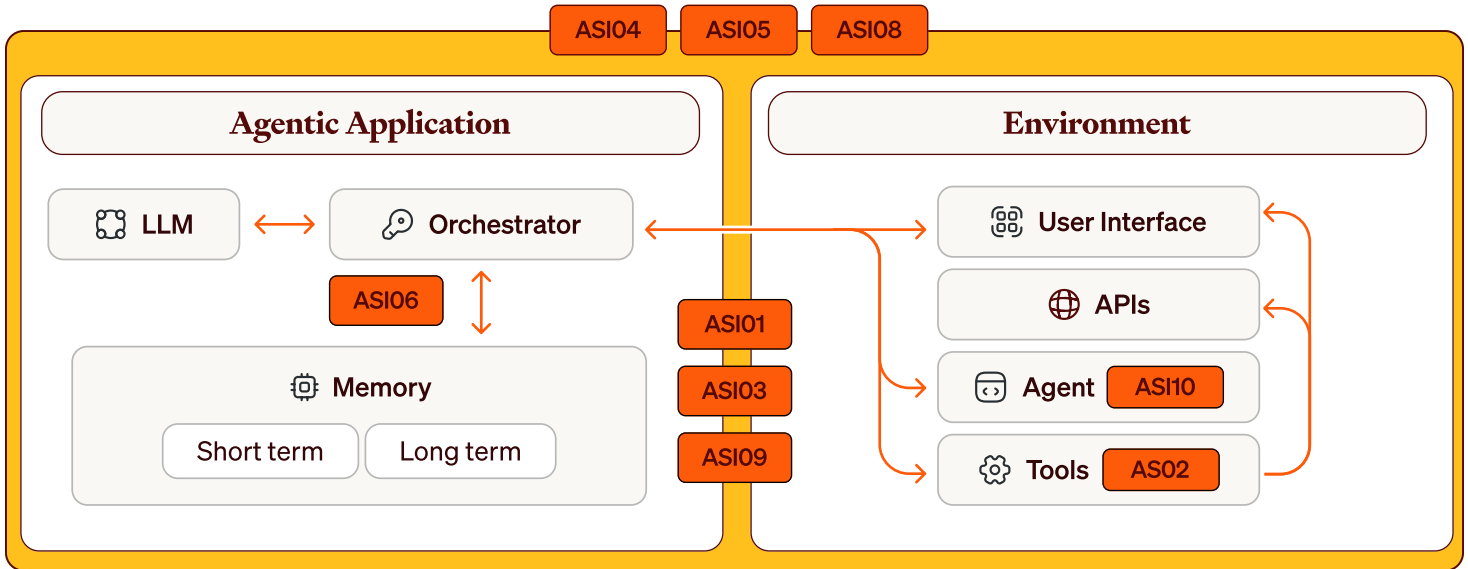
Session analysis

< 200ms

Time to detect

<50ms

Real-time blocking



Offensive and Defensive Security Loop

Find Vulnerabilities and Protect Against Them

By combining red teaming and runtime protection in a single platform, Lasso reduces MTTR by 95%, eliminates the lag between finding a vulnerability and acting on it, and creates a feedback loop that adapts continuously without requiring human intervention.



Red Teaming

Offense / Attack Simulation

Proactively discover weaknesses in your models, agents, and applications through adversarial simulation.

Prompt injection

Multi-turn attacks

Agent-logic corruption

Tool-chain exploitation

Recon & fingerprinting



Blue Teaming

Defense / Detection / Response

Ensure guardrails, policies, and runtime protections prevent, detect, and respond to AI-specific threats in real time.

Guardrail conformance

Runtime AI-SPM

Policy enforcement

Anomaly detection

Trust-boundary alerting



Purple Teaming

Continuous Attack-Defense Fusion

Feed adversarial findings directly into defensive controls, policy updates, and guardrail for continuous improvements.

Closed-loop remediation

Auto guardrail patching

Red/blue alert correlation

Policy update mapping

CI integration hooks

 **Red teaming** identifies real vulnerabilities in each agentic application. The moment findings are surfaced, Lasso auto-generates the guardrails to address them through purple teaming.

 At runtime, **blue teaming** ensures those protections hold. Policy enforcement, anomaly detection, and trust-boundary alerting operate continuously across the full application portfolio at sub-50ms latency.

Lasso's security approach is built specifically for agentic AI, and operates as a continuous, closed loop where offense and defense operate as a single system with:

- Adaptive guardrails generated automatically from red teaming findings, continuously updated as new vulnerabilities are discovered
- Context-based guardrails (CBAC) that apply different protections per person, role, and application
- Intelligence layer that evaluates intent at every step of an agent's execution
- Full audit trail across every interaction at runtime for complete visibility

Intent Security

Lasso's Intelligence Layer

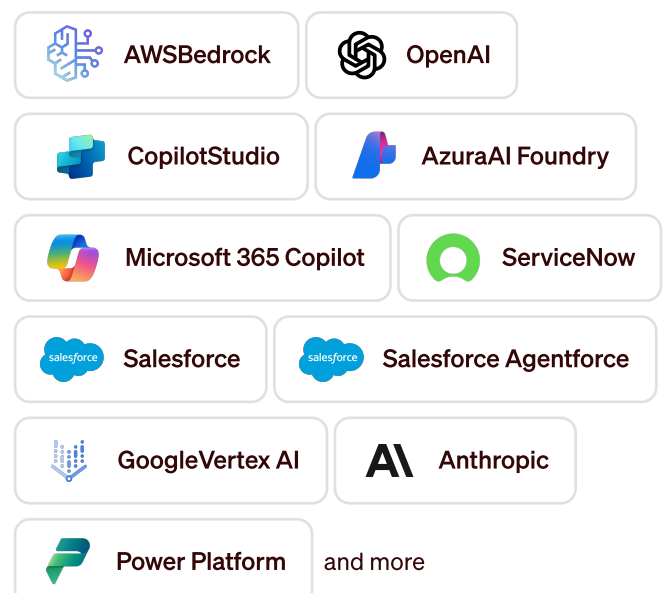
Lasso's Intent Security Framework is the intelligence layer running underneath every stage of the platform. Rather than evaluating isolated prompts or individual outputs, it builds a behavioral baseline for every agent and application, then continuously measures deviation from it.

Every agent operates under four competing pressures simultaneously: the system prompt as the developer's original design intent; the user request as the immediate goal of the human operator; the agent's actions as the actual execution path and tool usage; and external inputs including third-party data, environmental context, and potential adversarial influence. Intent Security evaluates all four together, not in isolation.

Use Cases

Safe Autonomy at Scale Agents with broad system access and tool chains are only as safe as the guardrails around them. Lasso continuously validates that every agent is operating within its intended scope, catching drift, manipulation, and unauthorized actions before they become incidents.

- ✓ Web applications
- ✓ AI APIs
- ✓ On-prem AI
- ✓ Cloud & third party agent building platforms



Secure AI Product Development

AI features embedded directly in product offerings create new revenue streams and competitive differentiation. Lasso gives engineering and product teams the confidence to ship those features without security becoming a bottleneck.

Trusted Agent Delegation

Delegating complex operational tasks to autonomous agents only generates value if those agents can be trusted. Lasso provides the governance layer that makes agent delegation safe, reducing manual oversight without increasing risk.

Compliance and Customer Trust

Customers and partners increasingly scrutinize how AI systems handle their data and decisions. Lasso's explainability and audit capabilities turn compliance into a competitive advantage, demonstrating that AI systems operate with integrity at every interaction.

Regulatory Readiness

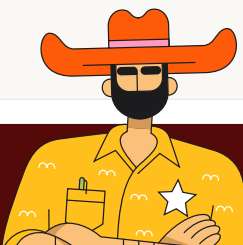
Global AI regulation is accelerating. Lasso's alignment with NIST, OWASP, and MITRE frameworks positions security teams ahead of the regulatory curve, not reacting to mandates after the fact.

Key Benefits

Lasso is an AI Security Platform built for a world run by AI. Every agent operates under four competing pressures simultaneously: the system prompt as the developer's original design intent; the user request as the immediate goal of the human operator; the agent's actions as the actual execution path and tool usage; and external inputs including third-party data, environmental context, and potential adversarial influence. Intent Security evaluates all four together, not in isolation.

About Lasso

Lasso is the AI Security Platform built for the agentic era. By connecting discovery, posture management, red teaming, and runtime enforcement in a single continuous loop, Lasso ensures every agentic application behaves within its intended scope, at every stage of its lifecycle.



→ Intent vs. technique detection

Lasso detects threats based on intent, not just known attack signatures, delivering 98.6% accuracy with a 1.4% false positive rate.

→ Offensive and defensive in one platform

Red teaming, AI-SPM, runtime enforcement, and threat detection operate as a unified system, not disconnected tools.

→ Bespoke agentic attacks

Multi-turn, high-agency testing tailored to each application's specific intent and design.

→ CI-native continuous security

Automated re-testing on every deployment so security keeps pace with engineering velocity.

→ Purple teaming built in

Red team findings automatically translate into policy updates and guardrail improvements, closing the loop between offense and defense.

→ Zero deployment overhead

No agents, no code changes, no source code access required. Connect in hours.

Trusted Security for a World Run by AI

[Book a Demo](#)