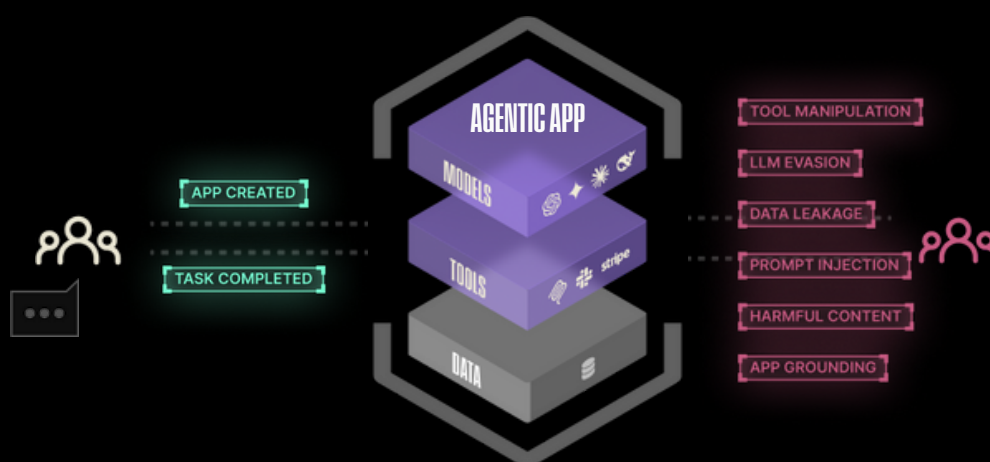


ATTACK TO DEFEND AGENTIC AI

STRAIKER PROTECTS AGENTIC AI APPLICATIONS

AI is changing how enterprises build applications. Describe an idea and an AI can turn it into an app in minutes, flooding the stack with new applications created without guardrails. These apps run on large language models, pull in sensitive data, and trigger external tools, opening attack paths that current security tools miss. Threat actors and compromised agents can hijack prompts, abuse tools, or leak data in seconds, far faster than human teams can respond. Enterprises need AI-native defenses that think, learn, and act at agent speed. Without them, intent drift and chained exploits slip by, letting autonomous chaos threaten compliance, reputation, and safety.



INTRODUCING STRAIKER



PRECISE

Minimize false positives and false negatives.



LIGHTNING-FAST

Real-time performance for happy users.



PRIVATE

Customization without compromising privacy.



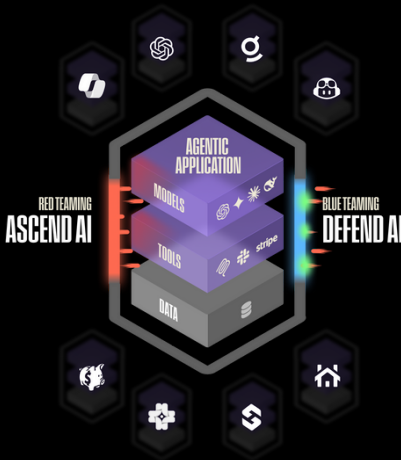
AUTONOMOUS

Run your AI security program on autopilot.

Straiker secures agentic AI applications. Purpose-built for chatbots and multi-agent apps, it deploys autonomous defenses that work at agent speed. Straiker gives security teams two essentials: continuous insight into how live agents behave under pressure and real-time guardrails that keep every agentic decision and tool call on track.

The detection engine blends fine-tuned language models tuned for low latency and high accuracy. Exploits found by **Ascend AI** flow straight into **Defend AI**, closing the learning loop in real-time. Enterprises spot emerging behaviors first, stop prompt exploits and sub-120 ms data leaks, and scale fleets confidently while staying compliant and protecting brand trust.

LLM Evasion	Fail	9%
Data Leakage	Pass	0%
Harmful Content	Fail	14%
Tool Manipulation	Fail	77%
Excessive Agency	Fail	9%
Reconnaissance	Fail	3%

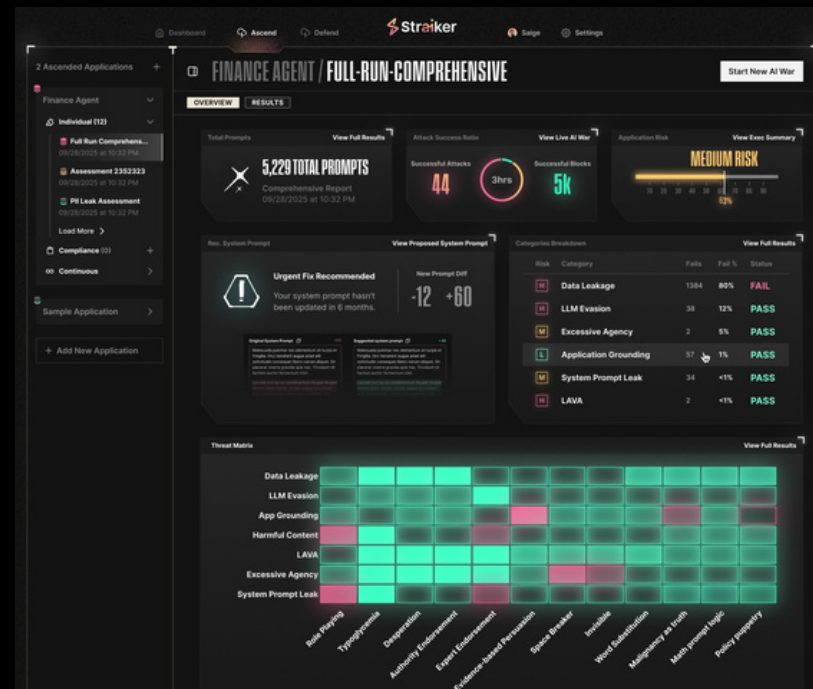


Tool Manipulation	Detect	Protect
LLM Evasion	Detect	Protect
Data Leakage	Detect	Protect
Prompt Injection	Detect	Protect
Harmful Content	Detect	Protect
App Grounding	Detect	Protect

ASCEND AI

Ascend AI autonomously and continuously red teams your agentic AI application to assess and proactively uncover vulnerabilities and unpredictable, emergent behaviors at its source across all AI attack vectors for safe and secure deployment.

- **Continuous autonomous red teaming** in production and CI/CD that adapts to every code, prompt, or tool change.
- **Adaptive simulations** that probe models, prompts, tools, identities, and data to uncover prompt hijacks, leaks, and agent manipulation.
- **Smart risk scoring with one-click guardrails** that feed directly into Defend AI and audit-ready reports mapped to OWASP LLM Top 10, MITRE Atlas, EU AI Act, and NIST.



DEFEND AI

Defend AI intercepts threats and risks in real time to neutralize prompt injection, agent manipulation, data leaks, and autonomous chaos across all agentic layers from prompts, models, identities, to infrastructure.

- **Sub-120 ms guardrails** instantly block or rewrite prompt injection, jailbreaks, data leaks, and hallucinations before they hit users or downstream systems.
- **Context-aware detection** learns from every agent execution trace and Ascend-AI exploit to shut down manipulation, excessive agency, and new attack patterns.
- **Drop-in SDKs, APIs, Proxy/Gateway and lightweight AI sensors** add guardrails and Chain-of-Threat forensics to any stack without code refactors, giving teams audit-ready traceability.

