

Protect AI Applications and Agents in Real Time

Runtime security built for production AI. Enforce policy inline without disrupting user experience.

BUILT FOR PRODUCTION

PROVEN AT ENTERPRISE SCALE

EXECUTIVE OVERVIEW

Secures AI as it Operates

AI Agent Security enforces security and policy controls inline—protecting prompts, outputs, and agent actions instantly.

Zero Friction

Deploys in under 5 minutes. No retraining models. No rewriting prompts. Secure your workflows without sacrificing development velocity or performance.

99.8%

DETECTION ACCURACY

<40ms

AVERAGE LATENCY

<0.5%

FALSE POSITIVES

CORE PROTECTION CAPABILITIES



THREAT DETECTION



DATA LEAKAGE



POLICY ENFORCEMENT



AGENT CONTROL

THE CONTEXT

Production AI Demands a New Security Paradigm

Deploying AI into live workflows exposes organizations to non-deterministic threats that traditional security controls cannot see or stop.

THE SHIFT TO RUNTIME RISK

- **Non-Deterministic:** Natural language from real users bypasses keyword filters.
- **Dynamic Execution:** Outputs generated in real-time create unpredictable risk surfaces.
- **Autonomous Agency:** Agents interacting with APIs require new permission boundaries.

THE GAP

Why Runtime Security Is Different

AI risk does not originate from misconfiguration alone. It emerges through interaction.

IN LIVE ENVIRONMENTS

- Prompts manipulate model behavior
- Outputs expose sensitive data
- Agents misuse tools or exceed authority

TRADITIONAL TOOLS MISS

- Natural-language intent
- Contextual and semantic attacks
- Multi-step agent behavior

THE SOLUTION

Securing AI at Runtime

AI Agent Security integrates directly into live AI interactions to inspect prompts, evaluate intent, and enforce policy inline.



[X] NO RETRAINING [X] NO REWRITING PROMPTS [X] NO PERFORMANCE LOSS

CAPABILITIES

What AI Agent Security Enforces



Runtime Threat Detection

- Detects injections, jailbreaks, adversarial instructions
- Identifies semantic attack patterns



Data Leakage Prevention

- Monitors prompts/outputs for PII
- Prevents unauthorized disclosure



Inline Policy Enforcement

- Applies rules at runtime
- Centralized updates without redeployments



Agent Action Control

- Guards tool usage & permissions
- Prevents unsafe agent behavior



Model-Agnostic & Fast

- Any LLM provider
- Sub-50ms latency



Multimodal & Multilingual

- Text, voice, multimodal apps
- 100+ languages

Operating Across the AI Lifecycle

BEFORE LAUNCH

Risks identified through testing.

Powered by AI Red Teaming

IN PRODUCTION

Enforces controls per interaction.

Powered by AI Agent Security

OVER TIME

Policies evolve with usage.

Powered by AI Agent Security

BUILT FOR PRODUCTION



Purpose-built for LLMs and AI agents



Supports cloud, on-prem, and hybrid



Proven in regulated systems



Aligns with compliance requirements

PROVEN IN ENTERPRISE



"We've chosen Lakera to secure our enterprise GenAI deployment across our regulated banking environment. This partnership enables us to safely innovate with AI in money transfers, financial services, and customer support. Lakera's accuracy, low latency, seamless integration, scalability and support for Portuguese and Spanish are essential for our global operations, especially in markets with sophisticated fraud attempts."



Largest digital bank outside of Asia
Serving 115M+ customers

Secure Every Layer of Your AI Journey

Lakera integrates with **Check Point** to complete the AI Defense Plane.



Get Started with AI Agent Security

DEPLOY IN < 5 MINS DEFINE POLICIES SECURE AGENTS

[SEE AI AGENT SECURITY IN ACTION](#)

Worldwide Headquarters

5 Shlomo Kaplan Street, Tel Aviv 6789159, Israel | Tel: +972-3-753-4599

U.S. Headquarters

100 Oracle Parkway, Suite 800, Redwood City, CA 94065 | Tel: 1-800-429-4391

www.checkpoint.com

© 2026 Check Point Software Technologies Ltd. All rights reserved.