

AI Security for the Enterprise

Secure Every Layer of Your AI Journey

As organizations rapidly adopt AI, risk no longer lives in one place. It spans employee AI usage, production AI applications, and autonomous agents — creating a fragmented security surface that traditional controls cannot see or govern.

Lakera, a Check Point company, delivers end-to-end AI security through a unified platform built for how AI is actually used in the enterprise.

THE CHALLENGE

AI Has Changed the Security Model

AI risk doesn't originate from misconfiguration alone. It emerges through human interaction, probabilistic reasoning, and autonomous action.

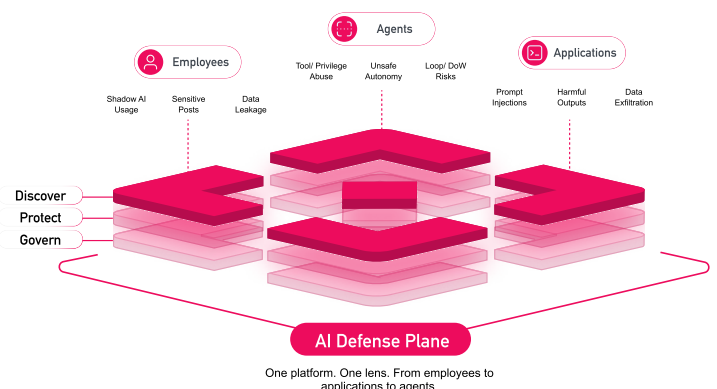
- ✗ Employees use hundreds of AI tools across browsers, desktops, and IDEs
- ✗ AI applications generate unpredictable outputs in real time
- ✗ Agents act on human authority and corporate systems
- ✗ Security teams lack visibility into who is using what, how, and with what risk

Without a unified approach, AI adoption outpaces security and governance.

THE SOLUTION

The AI Defense Plane

Lakera integrates with Check Point to form the AI Defense Plane — a single control plane for managing AI risk across the enterprise.



- ✓ **Unified visibility** into AI usage across the organization
- ✓ **Inline protection** where AI operates and decisions are made
- ✓ **Centralized governance** across people, apps, and agents

Comprehensive Protection

1



Workforce AI Security

GOVERN EMPLOYEE AI USAGE WITHOUT SLOWING PRODUCTIVITY

Know which AI tools are in use — and how they're being used. Security teams gain visibility and control — without blocking innovation.

- ✓ Discover AI tools across browser and device
- ✓ Break down activity by application & user
- ✓ Apply granular policies by app & data type
- ✓ Identify unapproved and Shadow AI usage
- ✓ Understand user intent to assess risk
- ✓ Prevent risky connections to corporate resources

2



AI Agent Security

PROTECT AI APPLICATIONS AND AGENTS AT RUNTIME

Secure AI where decisions are made. Inspect prompts, outputs, and agent actions inline to prevent prompt injection and unsafe autonomy.

- ✓ Inspect prompts, outputs, and agent actions inline
- ✓ Enforce policy without retraining models
- ✓ Prevent prompt injection & data leakage
- ✓ Sub-50ms latency, built for production

3



AI Red Teaming

EXPOSE AI FAILURE MODES BEFORE ATTACKERS DO

Test AI systems like a real adversary would. Identify vulnerabilities in reasoning, workflows, and tool usage before launch.

- ✓ Simulate real-world AI attacks & misuse
- ✓ Prioritize risk based on business impact
- ✓ Identify vulnerabilities in reasoning & tools
- ✓ Deliver actionable remediation guidance

Built for the Enterprise

SECURE AI ACROSS THE LIFECYCLE



Employee Usage

Govern employee AI usage and Shadow AI

WORKFORCE AI SECURITY



Before Launch

Identify risks and unsafe behaviors

AI RED TEAMING



In Production

Enforce security and policy per interaction

AI AGENT SECURITY



Over Time

Continuously govern evolving AI usage

AI DEFENSE PLANE

ENTERPRISE CAPABILITIES



Model-agnostic

Works with any LLM provider



Multimodal & Multilingual

Global coverage for text & voice



Flexible Deployment

Browser, desktop, and clientless



Proven at Scale

Trusted by regulated industries

The Enterprise AI Security Platform

One platform. One control plane. One source of truth for AI risk.

From employees to applications to agents, the AI Defense Plane enables organizations to adopt AI confidently.

GET STARTED