



Prompt Security

Secure the AI powering modern work —
without slowing the people building it.

As AI proliferates throughout organizations, security leaders face a new class of threats. These risks range from employees sharing sensitive data with AI tools to attackers manipulating models through prompt injection in customer-facing applications. Yet despite all the risks, AI unlocks immense value, and adopting it is fast becoming essential to business survival.

Prompt Security, A SentinelOne Company, empowers enterprises to embrace this new AI era with confidence. By providing the guardrails and governance needed to innovate safely, it enables companies to move fast, unlock the full potential of AI, and build trust with customers and employees alike – turning AI from a source of risk into a catalyst for growth.



Secure Workforce AI Adoption

Protect your organization from AI-associated risks by securing every interaction with AI applications without slowing adoption.



Safeguard AI-Powered Applications

Embed security and governance across every AI-enabled application, whether homegrown, integrated, or customer-facing.



Protect Data Everywhere

Ensure sensitive information stays private across all AI interactions by enforcing real-time data controls and adaptive privacy protections.



Enable Safe Innovation

Adapt to emerging threats, from autonomous agents to new AI attack vectors to enable your organization to embrace AI with confidence.

Prompt Security Addresses Key AI Risks

- ✓ **Shadow AI**
Identifies and monitors unsanctioned AI usage to eliminate blind spots.
- ✓ **Prompt Injection**
Detects and blocks malicious inputs designed to manipulate AI models.
- ✓ **Sensitive Data Disclosure**
Prevents confidential or regulated information from leaking into AI tools.
- ✓ **Harmful LLM Responses**
Shields users from being exposed to inappropriate, harmful or off-brand content generated by LLMs.
- ✓ **Insecure Agents**
Applies safeguards to AI agents to ensure safe automation at scale.
- ✓ **Jailbreak and Prompt Leaks**
Blocks attempts to override model safeguards or reveal hidden prompts.
- ✓ **Denial of Wallet/Service**
Detects abnormal usage and blocks it to prevent outages and excessive costs.

Gartner
Peer *Insights*.

“

Prompt Security's GenAI platform gave me clear visibility into AI-related risks and was surprisingly easy to integrate. The real-time monitoring and policy control made me feel more confident about securing generative AI use



BANKING, 50M-250M USD



Prompt for Employees

Enable your workforce to adopt AI without worrying about Shadow AI, data privacy, and regulatory risks.

- **Observability:** Detect and monitor all AI tools in use, eliminating Shadow AI risks and highlighting the riskiest apps and users.
- **Data Privacy:** Prevent data leaks through automatic anonymization and data privacy enforcement.
- **Risk and Compliance:** Establish and enforce granular department and user rules and policies.
- **Employee Awareness:** Coach your employees on the safe use of AI tools with non-intrusive explanations.



Prompt for AI Code Assistants

Adopt AI-based code assistants like GitHub Copilot or Cursor while safeguarding secrets, scanning for vulnerabilities, and maintaining developer efficiency.

- **Secrets and PII Protection:** Instantly redact and sanitize code.
- **Full Visibility and Governance:** Track AI usage across development cycles and flag potential privacy violations.
- **Broad Compatibility:** Integrates with thousands of AI tools and assistants and nearly 30 coding languages.



Prompt for Homegrown AI Applications

Unleash the power of AI in your homegrown applications without worrying about prompt injection, data leaks, or harmful responses.

- **Address AI Risks:** Secure your AI apps from prompt injection, jailbreaks, denial of wallet, RCE, and other exploits.
- **Data Protection:** Enforce automatic anonymization and data privacy enforcement to prevent data leaks.
- **Content Moderation:** Prevent user exposure to inappropriate, harmful or off-brand content generated by LLMs.
- **Visibility and Compliance:** Log and monitor inbound and outbound traffic from AI apps with full oversight.

AI Red Teaming

Expose prompt injection, jailbreaks, data poisoning, and more before you ship. Get risk-scored findings with remediation guidance, and a native path from pre-production vulnerability discovery to runtime protection.



Prompt for Agentic AI

Get real-time visibility, risk assessment, and policy enforcement across AI agents operating in your environment.

- **Agent and MCP Discovery:** Surface every AI agent, MCP server, and shadow deployment across SaaS apps, desktops, browsers, and developer tools, including those spun up without IT knowledge.
- **Continuous Risk Assessment:** Dynamically score agents based on permissions, behaviors, data access, and connected system exposure, prioritizing those with the broadest blast radius.
- **Policy Enforcement:** Apply least-privilege controls by user, agent, or action type to prevent unauthorized action chaining before it reaches your CRM, financial systems, or databases.
- **Audit Logging:** Maintain a complete, searchable record of every agent decision and interaction for compliance and forensic review.