



Securing AI Health Check

AI tools are delivering true business innovation acceleration, but they are moving faster than your business security controls

Service overview

Ensuring AI is securely adopted & emerging risks are managed protects your business whilst supporting innovation with accelerated efficiencies. Unfortunately, not all AI tools are adopted or used as securely as they should be.

CloudGuard's Securing AI provides an Agile & highly responsive range of best practice modular security services to enable your business to adopt & adapt AI securely.

Securing AI

Agentic AI Threat Detection

Enhanced AI Security User Awareness

Continuous Code, Agent & Data monitoring

Risk Management policies & rules

Greater visibility & data controls

AI Governance Framework

- **Readiness** – collaboratively test how your controls and processes perform against emerging AI risks
- **Responsiveness** – discover how your team detects, investigates and responds as a scenario unfolds
- **Resilience** – help you build lasting capability through direct knowledge transfer & adaptive governance

Customer Challenges

Shadow AI in businesses

85%

of organisations have employees using AI tools outside IT-sanctioned controls, exposing sensitive data with zero visibility.

Faster attack pattern & persistence growth

3×

AI agents, MCP servers, and autonomous workflows expand the attack surface at a rate traditional security controls were never designed to cover.

Operate without effective AI security policies in place

72%

of SMBs have deployed AI tools with no formal governance framework, leaving them exposed to data loss and prompt injection attacks.

Why this matters now

Your teams are already using AI, ChatGPT, Copilot, Gemini, and dozens of tools embedded in daily workflows. Most have access to internal data, emails, and systems. Few are governed. When an employee pastes client data into a public LLM, or an AI agent is granted excessive permissions, the damage doesn't wait for your next risk review.

Service outcomes

AI Governance Framework Health Check

Best Practice Secure AI adoption includes:

- Updated AI security policies,
- Enhanced AI & data access controls,
- Improved AI & Agent dev standards,
- Wider real time security monitoring & rules,
- AI-specific IR & BC plan.

Greater visibility & data controls

Best Practice Secure AI adoption includes:

- Reviewed CASB + Entra ID Conditional Access
- Ability to identify and govern every AI tool
- Detection of sanctioned or shadow AI in your business
- Validate & governance of NHI Agents, Autonomous Agents & decision making

Continuous Code, Agent & Data Security

Continually reviewed user & data guardrails across the business & partner landscape to assure data security, processing & security.

Risk Management Policies & Rules

Expanded coverage of security policies to include LLM-integrated apps, agentic pipelines & workflows, autonomous agents, MCP's, and connected AI services.

Enhanced AI Security User Awareness

Richer & more engaging AU user security awareness training & guidance.

Agentic AI Threat Detection

Enhanced security rules, logic, detections & responses covering Agents, agentic behaviours, unusual API patterns, lateral movements, user context OpenClaw personal assistance & AI interactions.

The Gaps

AI agents read files, send emails, execute code, and call APIs through protocols like MCP. A single misconfigured agent or compromised MCP server can move laterally across your environment faster than any human attacker.

Traditional NDR, EDR, and SIEM tools were never built to detect this.

- IDS/IPS and firewall rules cannot inspect AI agent actions or MCP server trust chains
- CASB catches sanctioned apps, not shadow AI running in browser extensions
- DLP misses context: sensitive data shared via prompt is not a file transfer
- Most IR playbooks have no AI-specific scenario, SOC teams lack a response path
- ENTRA ID identity & supporting conditional access policies to ensure consistency in Human & agentic accounts
- Ensure Data governance BEFORE activation as existing gaps will be amplified without controls & it is much harder to close these gaps
- Implement MCP and MCX standards and apply AI guardrails to code, agents, projects, API's, interfaces and maintain an inventory of these for lifecycle management
- Deploy AI-specific user awareness training, detection & response rules covering anomalous agent & non-human behaviours, usual API call patterns and permission escalations

The Security Confidence Gap is widest here.

Most organisations feel their AI tools are low risk because they trust the vendor. They're not assessing the prompt, the data, the agent permissions, or the MCP server trust chain. CloudGuard exists to close that gap.

FRAMEWORK COVERAGE

Mapped to OWASP Agentic AI Top 10, Microsoft Purview AI Hub, ISO 27001, and Cyber Essentials Plus, defensible evidence for audits and insurance.

Prompt Injection AA01 / OWASP	Excessive Permissions AA03 / OWASP	Data Exfiltration AA05 / OWASP	MCP Trust Chains AA07 / OWASP
Shadow AI Microsoft CASB	Copilot Governance M365 / Purview	Agent Monitoring Sentinel KQL	Incident Response NIST / ISO 27001

BUSINESS BENEFITS

- **Full AI visibility:** Know exactly what AI tools are in use, who has access, and what data they can reach, sanctioned and shadow.
- **Reduce AI-related risk:** Close the gaps traditional controls miss, prompt injection, agentic misuse, ungoverned data flows.
- **Compliance confidence:** Governance aligned to OWASP, ISO 27001, EU AI Act, NIS2 and Cyber Essentials. Evidence your auditor and insurer will accept.
- **Enable AI safely:** Controls that let your people use AI with confidence. Safe adoption beats blanket restriction.
- **24/7 agentic monitoring:** CloudGuard's SOC monitors AI-specific Sentinel detections continuously.

Securing AI Health Check

Uncertain where to start with Securing AI? In a single structured engagement, CloudGuard will assess your AI security posture, identify your highest-risk areas, and deliver a prioritised roadmap, including quick wins you can implement immediately. No commitment beyond the assessment. Full AIR delivered within two weeks.

WORKS WELL WITH

Red/Purple Teaming	PROTECT MDR	Entra ID / Identity
Purview / DLP	PROTECT+ Managed SOC	ENHANCE Secure Access



Enhanced Readiness

AI governance frameworks, shadow AI discovery, MCP security assessment, and prompt injection risk management, proactive controls before threats materialise.



Optimised Responsiveness

11 purpose-built Sentinel KQL detections for agentic AI threats, monitored 24/7 by CloudGuard's dual SOC in Manchester and Belfast.



Maximised Resilience

AI-specific incident response playbooks, post-incident review, and continuous service evolution as the AI threat landscape changes.

What CloudGuard Securing AI Delivers

AI Governance Framework Health Check

Five-document pack: security policy, access controls, dev standards, Sentinel KQL rules, and AI-specific IR plan.

Shadow AI Discovery & Control

CASB + Entra ID Conditional Access to identify and govern every AI tool – sanctioned or shadow – across your estate.

MCP Security Assessment

37 guardrails across 8 domains. Identify misconfigurations, excessive permissions, and trust chain vulnerabilities in MCP servers.

Prompt Injection Risk Management

CloudGuard PIRM framework (CG-PIRM-POL-001) covering LLM-integrated apps, agentic pipelines, and M365-connected AI services.

Agentic AI Threat Detection

11 Sentinel KQL rules covering anomalous agent behaviour, unusual API patterns, and lateral movement via AI toolchains.

AI Security Health Check

OWASP AA Top 10 assessment. Delivered as a CloudGuard Actionable Insight Report with maturity score and prioritised roadmap.

CloudGuard delivers enterprise-grade cybersecurity for growing businesses. Combining automation-led technology with UK-based experts, we provide 24/7 protection, without the complexity or cost of traditional platforms. Our fully managed approach simplifies detection, protection and improvement, so you can focus on growing with confidence.

[Get in touch](#)

