

Cloudflare AI Security Suite

在 AI 生命週期中控制資料和管理風險，藉此來保護 AI 互動。

自信地擴展 AI

應對 AI 風險需要現代安全性

您的團隊正在使用 AI 加速創新，但這也帶來了重大的安全風險。封鎖 AI 或新增複雜的單點解決方案這樣的傳統策略正在失效，因為它不僅會扼殺創新，還忽略了以下現實：

- 85% 的員工使用的 AI 工具未經 IT 部門審查。¹
- 93% 的員工承認在未經批准的情況下將公司資料輸入 AI。²
- 63% 發生資料外洩的組織未制定 AI 治理原則。¹

瞭解如何為團隊提供支援，並獲得 AI 的生產力優勢，同時保持業務所需的安全性和控制力。



確保主動式 AI 和 GenAI 安全的統一平台

Cloudflare 提供單一平台，讓您的組織能夠滿懷信心地採用 AI。我們協助網路安全領導者管理風險、技術團隊提高生產力以及平台工程師安全建置，確保所有人都能安全地共同創新。

開始使用 Cloudflare AI Security Suite

Cloudflare AI Security Suite 在統一平台上交付，用於確保工作團隊安全使用 AI 工具和面向公眾的應用程式。探索影子 AI、保護模型免遭濫用、保護主體存取，以及防止提示中的資料暴露，這樣，您的企業就能憑藉更輕鬆的可見度與更強的控制，來安全、高效地進行創新。



延伸可見度

涵蓋 AI 應用程式、API 端點和 AI 主體連線。



緩解風險

藉助防護機制、狀態控制和即時威脅緩解。



保護資料

透過為員工和 AI 主體提供提示防護和使用控制。



在 Cloudflare 上開發 AI 時，
已內建安全性。

使用 Cloudflare AI Security Suite 確保生成式 AI 和自主式 AI 通訊安全

透過控制工作團隊在生成式 AI 中使用的資料，並管理自主型主體帶來的安全風險，以保護所有 AI 通訊。



在 AI 生命週期中將安全納入設計，加速採用 AI



確保工作團隊安全使用 AI

實施 AI 使用控制和 AI 安全狀態管理 (AI-SPM)，以緩解風險並保護資料。



保護採用 AI 技術的應用程式

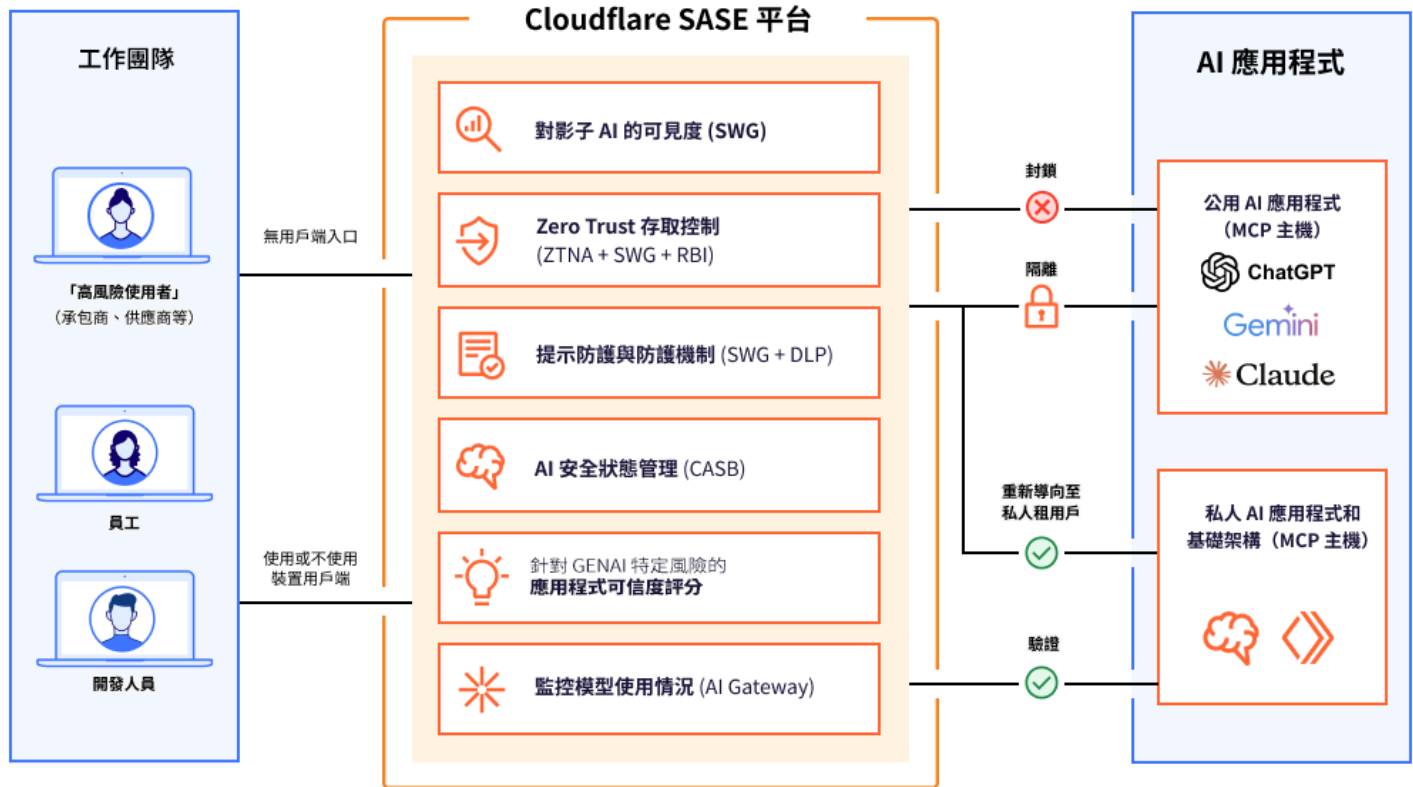
即時保護 AI 應用程式和 API，以防提示注入和資料外洩。



安全地建置 AI

憑藉整合式可觀測性、限速和內嵌 AI 防護機制，協助開發人員保護 AI 應用程式。

透過 Cloudflare 的 SASE 平台，確保工作團隊安全使用 AI 應用程式和工作負載



SSE 用於保護人類與 AI 之間的通訊

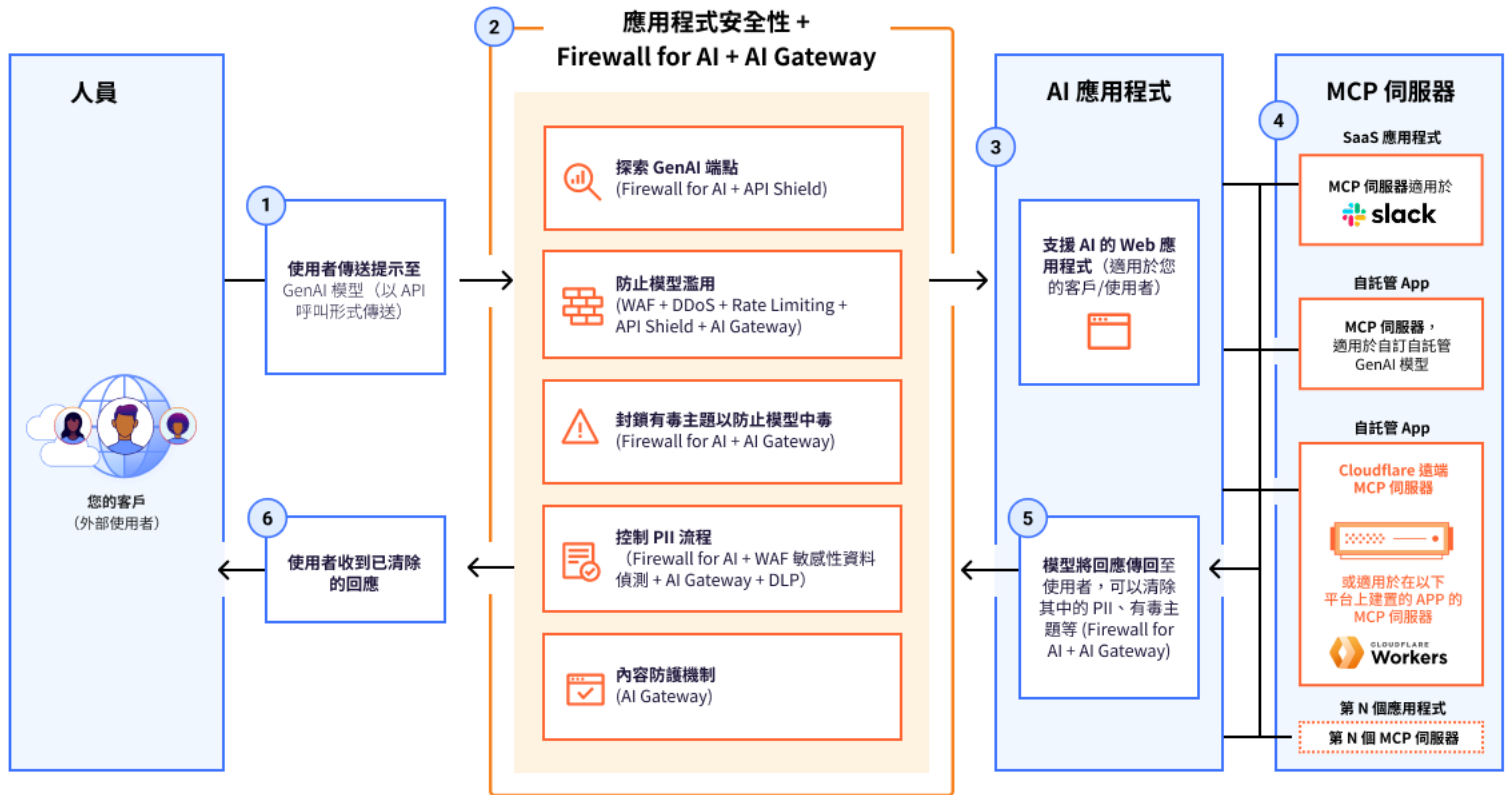
- **可見度：**透過內嵌流量檢查，探索並分析影子 AI 的使用。使用透明評分評估那些 AI 應用程式所帶來的風險。
- **存取控制：**封鎖、隔離、重新導向或允許使用者連線。針對每個應用程式強制執行基於身分的 Zero Trust 規則。
- **提示防護和防護機制：**基於意圖（例如，越獄嘗試、程式碼濫用、PII 請求）偵測並封鎖使用者提示。
- **資料安全：**透過針對 PII、原始程式碼等執行採用 AI 技術的資料丟失預防 (DLP) 偵測，阻止敏感性資料暴露。
- **AI 安全狀態管理：**透過 API 與 GenAI 工具整合，使用我們的雲端存取安全性代理程式 (CASB) 搜尋設定錯誤。現已支援 ChatGPT、Claude 和 Google Gemini。

MCP 伺服器入口網站用於保護 AI 與資源之間的通訊

- **可見度：**彙總所有模型情境通訊協定 (MCP) 請求記錄，以供稽核和分析。對每個 MCP 伺服器進行審查並核准，然後再將其新增至入口網站。
- **驗證：**基於身分驗證使用者對入口網站的存取。根據最低權限原則，設定對 MCP 伺服器的存取權限。
- **連線：**使用單一 URL 連接所有可存取的 MCP 伺服器，而非個別設定每個 MCP 伺服器。
- **統一管理：**對 AI 連線強制執行與對人類使用者相同的精細化存取原則。

注意：MCP 伺服器入口網站支援任何 MCP 伺服器，包括（但不限於）在 Cloudflare 上建置或部署的遠端 MCP 伺服器。此功能可作為 Zero Trust 網路存取 (ZTNA) 控制使用。

使用 Cloudflare 不受限於模型的內嵌安全性， 保護支援 AI 的應用程式和工作負載



使用應用程式安全性和 Firewall for AI 保護面向公眾的 AI：

- **探索 GenAI 端點：**自動探索 Web 資產中的所有 AI 模型和 API。
- **保護 AI 模型免遭濫用：**使用我們專門建置的 [Firewall for AI](#)，以阻止提示注入、模型中毒、過度使用，以及其他可能繞過傳統安全保護的威脅。
- **控制 PII 流程：**掃描使用者提示和模型回應，以 [阻止敏感性資料](#) 暴露，從而協助您保持合規性。
- **實施內容防護機制：**使用 Llama Guard 等整合模型，[封鎖不安全或有害的提示](#)。建立自訂 WAF 規則，以輕鬆封鎖或記錄可疑的 AI 互動。

使用開發人員平台和 AI Gateway 保護您建置的 AI

- **統一 AI 控制平面：**透過單一儀表板管理所有 AI 應用程式。路由請求、快取回應、控制成本並監控效能。
- **在邊緣保護認證：**在邊緣安全地儲存 API 金鑰和 [密碼](#)，防止用戶端暴露並簡化提供者之間的金鑰輪換。
- **實施內容安全防護機制：**自動識別並 [封鎖](#)/刪減提示和回應中的有害內容和 PII。

Cloudflare 是唯一可以同時確保公開和私人 AI 環境安全的廠商

實作適當的防護機制，以便充滿信心地採用 AI，從而確保安全措施能夠加速而不是阻礙創新。

- **統一的 AI 生態系統防護：**複雜的安全性堆疊會增加風險。使用單一平台保護資料，並確保在整個 AI 生命週期內實現合規性。
- **具備前瞻性的全球架構：**如今，憑藉後量子安全網路，可防範未來的挑戰。該安全網路能針對任何流量規模進行擴展、不斷適應新威脅，並且可針對新的使用案例進行程式設計。
- **採用 AI 技術的安全性：**我們採用 AI 技術的防禦系統可檢查提示和回應，以進行即時威脅偵測。
- **久經驗證的 AI 領導地位：**在生成式 AI 領域的前 50 大公司中，有 80% 的公司信任該平台，因此可在該平台上充滿自信地進行創新。
- **不受限於模型的部署：**安全控制適用於您環境中的所有 AI 模型，提供了一種統一的方法來管控 AI 部署。

客戶評價



全球排名第一的
求職網站
[詳閱案例研究](#)

識別及控制影子 AI
與 VPN 取代專案同時進行



保險科技公司
[詳閱案例研究](#)

隔離如 ChatGPT 等公用
GenAI 工具

以封鎖複製-貼上敏感性
資料



支援 AI 的 SaaS 公司

保護 PII
透過防止客戶將敏感性
資訊提交給面向公眾的
GenAI 端點



AI 驅動的金融科技公司

推斷成本降低了 95%

透過採用 Cloudflare 快取
並執行來自 AI 模型提供者
的回應

準備好討論您的 AI 安全性需求了嗎？

[與專家討論](#)

1. 2025 年 Manage Engine 研究：[來源](#)
2. 2025 年 IBM《資料外洩成本報告》：[來源](#)