

AI 賦能資安：全球產業應用、攻擊趨勢與未來戰略深度剖析

三聯科技股份有限公司 / 黃茗浩

隨著人工智慧（AI）技術的快速成熟，AI 已深度融入全球網路安全攻防的各個層面。本研究旨在為企業決策者及資安專業人士，提供一份關於 AI 在資安領域應用的全面性、前瞻性分析。全球「AI 資安」市場正經歷驚人的擴張時期，其市場規模預計將從 2024 年的約 253.5 億美元，以 24.4% 的年複合成長率（CAGR）增長，到 2030 年將達到 937.5 億美元。本研究系統性地探討了推動此變革的核心技術，包括機器學習（ML）、深度學習（DL）和大型語言模型（LLMs）。報告詳細闡述了它們在現代資安架構中的具體應用，涵蓋了異常偵測、威脅情資分析、安全編排、自動化與回應（SOAR）、端點 / 延伸式偵測與回應（EDR/XDR）、身份驗證以及雲端和物聯網生態系統的安全等多個功能領域。同時，本分析也揭示了惡意行為者如何將 AI 武器化，以策劃更複雜的社交工程、自動化惡意軟體變種和執行具說服力的深度偽造（Deepfake）攻擊。

關鍵字：人工智慧、網路安全、機器學習、威脅偵測、AI 安全、網路防禦

一、緒論

(一) 研究背景：從自動化到智慧化

網路安全的發展史，本質上是一部不斷追求自動化的歷史。從早期防火牆的基礎規則系統、初期入侵偵測系統的簽章比對技術，到後來更複雜的腳本化應對，其核心目標始終是以有限的人力資源應對指數級增長的威脅^[1]。然而，隨著攻擊手法的日益複雜化與多變，這種傳統的、基於規則的自動化已證明不足^[2]。

AI 的出現，特別是機器學習和深度學習的成熟，已將網路防禦從「自動化」時代推向了一個「智慧化」的新紀元^[3]。AI 系統不再局限於執行預設的指令；它們現在具備了從海量數據中學習、推理、預測並做出近似人類專家判斷的能力^[4]。

(二) 研究目的與範疇

本研究的主要目標是系統性地探討 AI 技術如何從根本上重塑全球資安產業的攻防格局。本研究的具體目標包括：

1. 梳理核心技術：解構機器學習、深度學習及大型語言模型在資安領域的基礎原理與實

際應用。

2. 剖析防禦應用：詳細闡述 AI 在威脅偵測、事件應對、身份與存取管理等關鍵領域的具體實踐。

3. 揭示攻擊手法：分析惡意行為者如何利用 AI 發動具有更強欺騙性與破壞潛力的攻擊。

4. 評估挑戰風險：深入探討將 AI 整合至安全框架中所面臨的技術、倫理與合規挑戰。

5. 預測未來趨勢：展望該領域的未來發展軌跡，包括人機協作與自主防禦。

(三) 市場規模與企業採用趨勢分析

AI 在資安領域的商業價值現已獲得市場的高度認可。根據詳細的市場研究報告，2024 年全球 AI 資安市場規模約為 253.5 億美元，預計將以 24.4% 的驚人年複合成長率（CAGR）擴張，到 2030 年市場規模可望達到 937.5 億美元^[5]。如圖 1 所示，此市場成長軌跡展現了 AI 資安技術領域的強勁發展動能和投資者信心，充分反映了產業對於智慧化防禦解決方案的迫切需求。

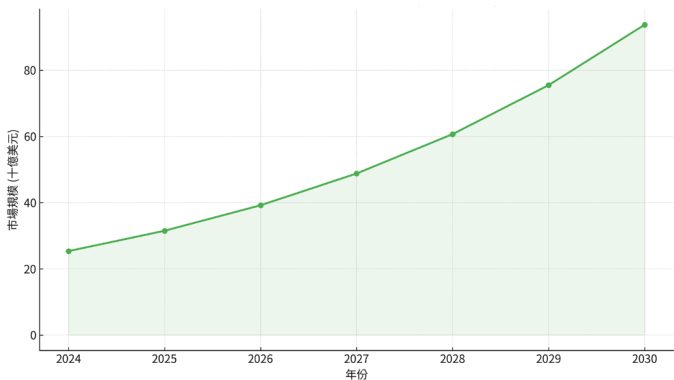


圖 1. 全球 AI 資安市場成長趨勢 (2024-2030)^[1,5]

企業端的採納意願同樣強烈。2024 年的統計數據顯示，約有 72% 的公司聲稱已在其營運框架中採用了 AI 技術^[6]。更值得注意的是生成式 AI 的潛力；高達 92% 的企業已制定計劃，

在未來三年內投資於此類技術^[7]。

圖 2 詳細呈現了跨產業企業 AI 採用模式和投資規劃趨勢，數據顯示組織對於 AI 技術整合的戰略重視程度已達到前所未有的高度。

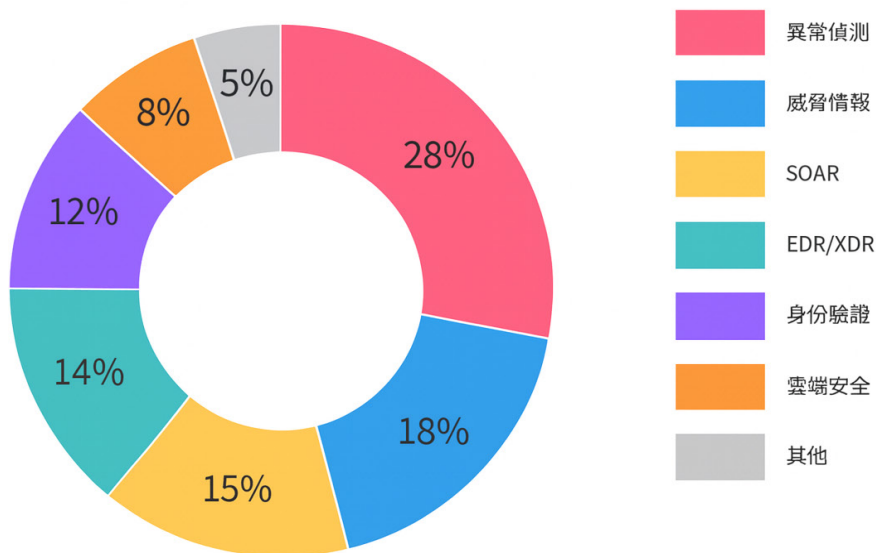


圖 2. AI 技術在資安領域的應用分布^[6,7]

二、驅動資安變革的 AI 核心技術

(一) 機器學習：資安的基石

機器學習是人工智慧的一個核心分支，它使電腦系統能夠從數據中自動學習和改進，而無需為每個任務進行明確的編程^[8]。這種能力構成了現代智慧防禦系統的基石。在網路安全的背景下，機器學習主要分為以下幾類：

1. 監督式學習 (Supervised Learning)

此方法涉及在一個每個數據點都已標記的數據集上訓練模型——例如，一組被精心標記為「惡意」或「正常」的郵件^[9]。模型學習輸入數據與輸出標籤之間的映射函數，使其能夠對新的、未見過的數據進行分類和預測。其主要的資安應用包括高效的垃圾郵件過濾、將新發現的惡意軟體分類至已知家族，以及基於網站結構和內容特徵識別釣魚網站。

2. 非監督式學習 (Unsupervised Learning)

這種方法應用於未經標記的數據集，其目標是探索數據並發現其中隱藏的模式或內在結構^[10]。它在網路安全中最強大的應用是用於異常偵測的用戶與實體行為分析 (UEBA)。透過

使用聚類演算法，這些系統可以識別嚴重偏離已建立常規的用戶或網路行為。

3. 強化學習 (Reinforcement Learning)

該範式涉及一個「代理人」(模型) 透過與環境的持續互動來學習^[11]。經過反覆試驗，代理人被訓練以最大化累積獎勵，從而有效地學習在特定情況下的最佳策略。在網路安全中，強化學習被用於自動化滲透測試等高級應用。

(二) 深度學習：洞察複雜威脅

深度學習是機器學習的一個專門子集，其特點是使用多層神經網路，這使得這些模型能夠從原始數據中自動學習和提取高度複雜的特徵和模式^[12]。它們處理非結構化數據的能力使其特別適用於網路安全任務。

1. 卷積神經網路 (CNN)

CNN 最初為圖像識別而開發，現已被創新地應用於網路安全。一個突出的用例是將惡意軟體二進位檔案視覺化為圖像，從而讓 CNN 能夠識別指示惡意代碼的紋理和結構模式以進行分類^[13]。

2. 循環神經網路 (RNN)

RNN 擅長處理序列數據，使其成為分析隨時間變化的事件序列的理想選擇。它們可用於分析網路流量序列或程式進行的系統調用順序，以偵測靜態分析會錯過的惡意行為模式^[14]。

3. Transformer 架構

這種基於「自注意力機制」的架構，徹底改變了自然語言處理領域，並成為所有當代大型語言模型的基礎技術^[15]。

(三) 大型語言模型與生成式 AI：效率與風險的雙面刃

以 GPT 系列為代表的大型語言模型，是基於 Transformer 架構的龐大模型，通常包含數千億個參數^[16]。它們理解和生成高度自然、類人語言的深厚能力，在網路安全領域既帶來了巨大潛力，也帶來了重大風險。

防禦應用包括：

1. 警報分類與總結：LLMs 能夠自動閱讀、理解並總結數千條安全警報。

2. 自動化報告生成：它們可以將複雜事件調查的結果自動轉換為結構良好的報告。

3. 自然語言查詢：LLMs 使資安分析師能夠透過自然語言提問來進行威脅狩獵。

攻擊潛力：LLMs 同時也降低了攻擊者的入門門檻。惡意行為者可以利用它們生成幾乎無法與合法通信區分的釣魚郵件、協助編寫惡意程式碼^[17]。

圖 3 系統性地展示了大型語言模型在資安領域的雙重特性，清楚說明其作為防禦工具的強大能力以及被惡意利用時可能產生的威脅向量，為組織評估 LLM 技術採用風險提供重要參考框架。

表 1 進一步闡述了 AI 技術在現代資安架構中的完整實施流程，從資料收集、模型訓練到威脅偵測與回應的端到端技術整合方案，為企業規劃 AI 資安部署提供系統性的架構指引。

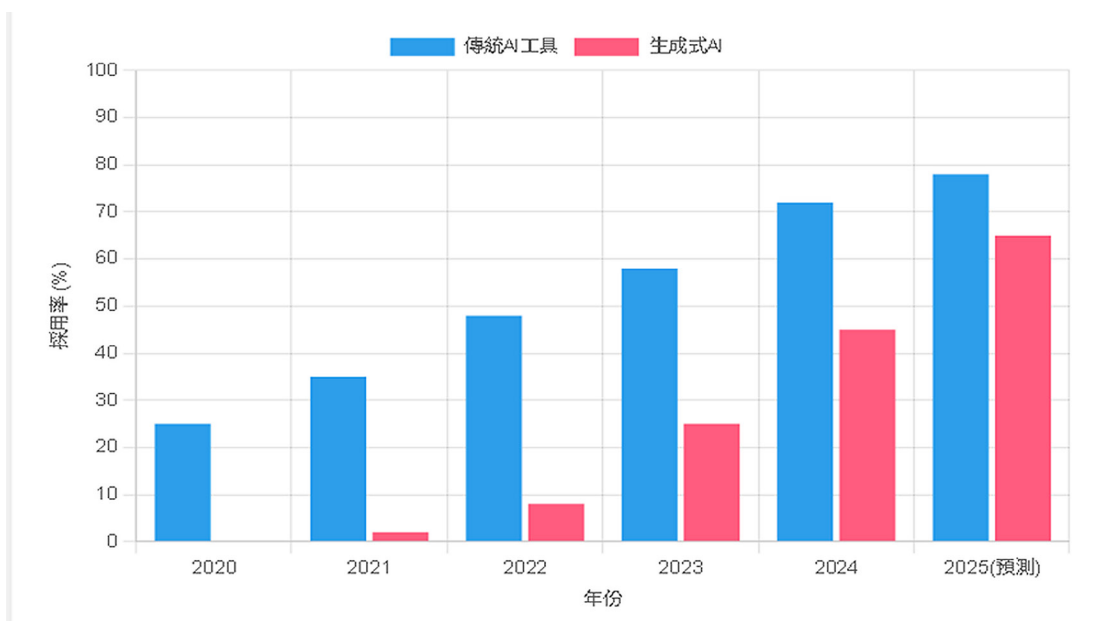


圖 3. 企業 AI 資安技術採用率趨勢^[16,17]

表 1. 主要 AI 資安技術特性比較 ^[16,17]

技術類別	主要演算法	應用領域	優勢	限制	成熟度
監督式學習	SVM, Random Forest, Neural Networks	惡意軟體分類、垃圾郵件過濾	高準確率、可解釋性好	需要標記數據、對新攻擊適應性差	高
非監督式學習	K-means, DBSCAN, Isolation Forest	異常行為偵測、零日攻擊發現	無需標記數據、能發現未知威脅	誤報率高、難以解釋結果	中
深度學習	CNN, RNN, Transformer	圖像分析、序列分析、NLP	處理複雜數據、自動特徵提取	需要大量數據、黑箱問題	中
強化學習	Q-Learning, Policy Gradient	自動化回應、動態防禦	自適應學習、最佳策略	訓練複雜、環境依賴性強	低
大型語言模型	GPT, BERT, T5	威脅情報分析、自動化報告	理解自然語言、生成能力強	計算資源需求高、潛在偏見	中

三、AI 賦能防禦：現代資安架構的智慧化轉型

(一) 異常行為偵測

這是 AI 在網路安全中最經典且影響最深遠的應用之一。透過利用非監督式機器學習模型，安全系統能夠分析並學習網路流量、用戶登入活動、應用程式調用等遙測數據的正常行為模式基線^[18]。一旦建立此基線，任何嚴重偏離它的活動將立即觸發警報。

UEBA 系統通常利用聚類演算法（如 K-means、DBSCAN）和統計異常值偵測方法來識別這些異常模式。這些系統分析多重行為向量，包括登入模式、數據存取量、應用程式使用情況和網路通訊行為^[19]。

(二) 智慧威脅情資分析

安全團隊面臨的最大挑戰之一是管理海量的威脅情資數據。AI 可以極大地加速此類資訊的收集、處理和分析^[20]。AI 驅動的網路爬蟲能自動監控暗網論壇、駭客社群及開源情報（OSINT）來源，以萃取攻擊指標（IoC）。

自然語言處理（NLP）技術可用於聚合、去重和翻譯來自全球的資安報告。此外，透過應用圖形分析技術，AI 可以揭示不同攻擊活動和威脅行為者之間的隱藏關聯^[21]。

(三) 安全編排、自動化與回應（SOAR）

SOAR 平台的核心目標是提升安全營運中心（SOC）的效率。AI 在這些平台中扮演智慧大腦的角色，將警報處理、分析和應對流程智慧化^[22]。AI 模型能根據警報的上下文、嚴重性和可信度自動判定其優先級，直接解決普遍存在的「警報疲勞」問題。

當高風險事件被確認時，AI 可觸發智慧應變劇本。這些預設的工作流程能自動執行一系列操作，例如將受感染的端點從網路中隔離、在防火牆上封鎖惡意 IP 地址，或撤銷被盜用帳號的權限^[23]。

(四) 端點 / 延伸式檢測與回應（EDR/XDR）

現代端點防護已從傳統防毒軟體（AV）演進至端點檢測與回應（EDR），現在更發展到延伸式檢測與回應（XDR）。AI 是這場演進的核心驅動力^[24]。

在 EDR 系統中，AI 分析端點上的數千個行為指標——如進程活動、記憶體使用和登錄檔變更——以偵測傳統 AV 會錯過的無檔案攻擊等高級威脅。XDR 將此能力從端點擴展至更廣泛的來源，包括網路、雲端基礎設施、郵件閘道和身份系統^[25]。

(五) AI 驅動的身份識別與驗證

身份已成為新的安全邊界，而 AI 正在強化身份驗證的智慧性與動態性。一個新興領域是行為生物辨識，AI 可以持續分析用戶獨特的行為模式——如打字速度、滑鼠移動軌跡和觸控螢幕互動——作為一種持續、無感的身份驗證方式^[26]。

更普遍的是，AI 被用於風險式驗證。系統會對每一次登入嘗試即時評估多重風險信號，包括用戶的地點、時間、設備健康度和網路信譽，以計算風險評分^[27]。

(六) 雲端與物聯網安全

雲端和物聯網環境的規模化與動態性，使

其成為 AI 安全應用的絕佳場景。在雲端，AI 可持續掃描雲端配置，自動檢測因人為疏失導致的資安風險，如公開的 S3 儲存桶或過於寬鬆的 IAM 權限^[28]。

對於物聯網安全，由於 IoT 設備計算資源有限，輕量化的 AI 模型通常被部署於邊緣閘道或雲端。這些模型用於監控大量設備的通訊模式，以偵測殭屍網路活動或異常指令^[29]。

圖 4 詳細描繪了 AI 在雲端與物聯網安全架構中的核心地位，展示了智慧化防護機制如何在大規模分散式環境中實現即時威脅偵測與自動化回應，為企業建構現代化混合基礎設施安全防護提供技術藍圖。

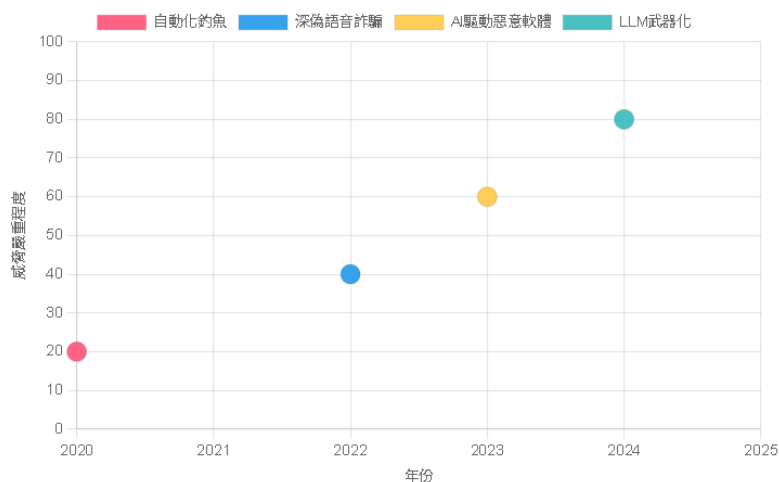


圖 4. AI 驅動的網路威脅演進時間軸^[30,31,32,33,34]

四、AI 的陰暗面：攻擊者的武器庫

(一) AI 驅動的社交工程：釣魚與語音詐騙的

「超進化」

這是目前 AI 被濫用最廣泛、威脅最直接的領域。

1. 高仿真釣魚郵件

利用大型語言模型，駭客可以生成語法完美、語氣自然、內容高度客製化的釣魚郵件^[30]。這些高仿真的郵件能輕鬆繞過傳統的、尋找常見錯誤和通用措辭的垃圾郵件過濾器。

2. 深偽語音詐騙 (Deepfake Voice Scams)

攻擊者使用 AI 聲音複製技術，僅需幾秒鐘的音檔樣本，就能模仿任何人（如企業 CEO 或家人）的聲音，進行電話詐騙^[31]。此趨勢的數據驚人：亞太地區的深偽語音詐騙嘗試在 2024 年相較去年暴增了 194%。

(二) 自動化惡意程式變種與規避

駭客正利用 AI 技術打造能「自我進化」的惡意程式。

1. 多態 / 變態惡意軟體

AI 可以自動、快速地對惡意程式碼進行無意義的修改或重寫功能模組，從而生成本數百萬個功能相同但簽章完全不同的變種^[32]。這讓傳統基於簽章的防毒軟體完全失效。

2. 環境感知與適應

AI 驅動的惡意程式在滲透入目標系統後，能夠主動識別其所處的環境，如是否存在沙盒、安裝了何種防毒軟體，並即時調整自身的行為以躲避偵測^[33]。

(三) 深度偽造 (Deepfake) 攻擊的商業威脅

除了聲音，AI 生成的影像 (Deepfake) 也

構成了嚴峻的商業威脅：

1. 身份冒用：駭客可以製作 CEO 或關鍵人員的偽造視訊，用於線上會議。

2. 資訊誤導與勒索：攻擊者可製造虛假的負面影片來損害企業或個人聲譽。

3. 繞過身份驗證：攻擊者可能利用 Deepfake 技術，繞過依賴人臉辨識的身份驗證系統^[34]。

表 2 呈現了深度偽造技術威脅範圍的全面分析，量化展示不同攻擊場景的風險等級和潛在業務衝擊，為風險管理團隊制定針對性防護策略提供數據支撐和決策依據。

表 2. AI 資安技術風險評估矩陣^[10,11]

風險類型	發生機率	影響程度	風險等級
模型偏誤	高	中	中等
對抗性攻擊	中	高	高
數據污染	低	高	中等
模型竊取	中	中	中等
隱私洩露	高	高	高
法規違規	中	高	高

五、導入 AI 的挑戰、風險與限制

(一) 模型偏誤與公平性

AI 模型的決策品質高度依賴其訓練數據的品質。如果訓練數據本身存在偏誤，模型就會學習並放大這些偏誤^[35]。例如，如果一個用於偵測內部威脅的模型，其訓練數據主要來自某個特定部門或地區，那麼它可能會對其他部門或地區員工的正常行為產生較高的誤報率。

(二) 對抗性攻擊：欺騙 AI 的藝術

這是 AI 安全領域特有的風險。攻擊者可以透過對輸入數據進行人眼難以察覺的微小擾動，來蓄意誘導 AI 模型做出錯誤的判斷^[36]。這些被精心製作的輸入被稱為「對抗樣本」。

在惡意軟體偵測模型中，攻擊者可能只需修改惡意執行檔中的幾個位元組，就能讓 AI 模型將其誤判為良性軟體^[37]。

(三) 數據污染與模型竊取

1. 數據污染：在模型訓練階段，攻擊者向訓練數據集中注入少量惡意數據，從而「污染」模型，使其在未來遇到特定輸入時產生攻擊者預期的錯誤輸出^[38]。

2. 模型竊取：透過大量查詢一個已部署的 AI 模型並分析其回傳結果，攻擊者可能逆向推導出模型的內部架構和參數，從而複製或竊取企業的核心 AI 資產^[39]。

(四) 隱私衝擊與全球法規遵循

AI 系統需要大量數據進行訓練，這往往涉及個人或敏感商業數據，引發了嚴重的隱私疑慮。許多深度學習模型如同「黑箱」，其決策過程缺乏透明度和可解釋性^[40]。

與此同時，全球監管環境日趨嚴格。2024 年，歐盟的《AI 法案》(EU AI Act) 啟動執行，對高風險 AI 應用提出了嚴格的要求^[41]。金融業還有《數位營運韌性法案》(DORA) 等規範陸續上路。圖 5 系統性地呈現了全球 AI 治理法規環境的複雜性和多層次要求，為合規團隊和法務部門理解跨國法規適用性、制定符合多重管轄區要求的 AI 治理框架提供重要的戰略規劃參考。

AI 賦能資安技術

防禦性 AI 應用

- 異常行為偵測 (UEBA)
- 威脅情報分析 (Threat Intelligence)
- 安全編排與自動化 (SOAR)
- 端點偵測與回應 (EDR/XDR)
- 身份驗證與存取控制 (IAM)
- 雲端與 IoT 安全

攻擊性 AI 應用

- AI 驅動社交工程
- 自動化惡意軟體生成
- 深度偽造攻擊 (Deepfake)
- 對抗性樣本攻擊
- 模型投毒攻擊

核心技術架構

- 機器學習 (ML)
- 深度學習 (DL)
- 大型語言模型 (LLM)
- 強化學習 (RL)

圖 5. AI 資安技術分類架構^[42,43,44,45]

六、未來展望：從人機協作到自主防禦

(一) 趨勢一：人機協作新典範 (AI Copilots)

在不久的將來，AI 將不再僅僅是後端的分析引擎，而是成為與資安分析師並肩作戰的「智慧助手」(Copilot)^[42]。AI 功能將直接嵌入到分析師日常使用的工具中。這種協作模式

將允許分析師透過自然語言與 AI 互動，提出問題並下達指令。

(二) 趨勢二：「AI 代理」(Agentic AI) 的崛起

這是比 Copilot 更進一步的概念。所謂「AI 代理」是具備高度自主能力的系統，能夠在沒有人類明確、逐步指令的情況下，主動發起並完整執行一系列複雜的安全任務^[43]。

例如，一個被稱為「Purple AI」的代理，能夠自主理解一個模糊的威脅狩獵目標（如「尋找我們網路中是否有類似最新 APT 攻擊的活動蹤跡」），然後自行規劃調查步驟、執行查詢、分析結果，並最終生成完整的調查報告。

(三) 趨勢三：「AI 對抗 AI」的自主防禦系統

隨著攻擊方也開始大規模使用 AI，未來的網路攻防將演變成「AI 對抗 AI」的局面^[44]。這將催生具備自我學習和自我防禦能力的 AI 系統。這類系統在偵測到一種全新的攻擊手法時，不僅僅是攔截，更能自主分析攻擊的原理，並即時生成、部署新的防禦規則或修補程式。

(四) 產業應用演進

不同行業對 AI 資安的需求也將持續演化：

1. 金融保險與製造業：這些傳統上重視資訊安全的領域，已率先導入生成式 AI 進行威脅模擬與風險評估。
2. 物聯網 (IoT)：隨著 5G 和邊緣運算的普及，連網設備數量將爆炸性增長。
3. 零信任架構：AI 將成為實現零信任架構的關鍵，透過持續的行為分析和風險評估，為每一次的存取請求提供動態的信任決策^[45]。

表 3 描繪了 AI 資安技術的未來發展路徑，清晰展示從目前的人機協作模式逐步演進至完全自主防禦系統的技術演進軌跡，為技術決策者規劃長期 AI 資安投資策略和能力建構提供前瞻性指引。

表 3. 主要國家 / 地區 AI 資安法規比較 [12,41]

國家/地區	主要法規	實施時間	主要要求	罰則
歐盟	EU AI Act	2024 年 8 月	高風險 AI 系統強制性要求、透明度義務	營業額 7% 或 3500 萬歐元
美國	NIST AI Risk Management Framework 2.0	2024 年	風險管理、安全評估、持續監控	依部門法規而定
中國	算法推薦管理規定	2022 年 3 月	算法透明度、用戶權益保護	最高 100 萬元人民幣
英國	AI White Paper	2023 年	原則導向、部門監管	依部門法規而定
新加坡	Model AI Governance Framework	2020 年	自願性指導原則、最佳實踐	無強制性罰則

七、結論與戰略建議

(一) 結論：擁抱 AI 的雙面性

AI 正以前所未有的深度和廣度，全面重塑全球資安產業的格局。它既是抵禦高級威脅的最強盾牌，也可能成為攻擊者手中最鋒利的矛^[46]。對企業而言，忽略 AI 將意味著在未來的資安競賽中處於絕對劣勢。然而，盲目導入 AI 而不考慮其內在風險，同樣可能引發新的危機。

成功的關鍵在於理解並擁抱 AI 的雙面性，採取一種平衡、全面的戰略。

(二) 戰略建議

基於以上分析，本研究為企業決策者和資安專業人士提出以下四點戰略建議：

1. 強化資料與模型治理

數據是 AI 的燃料，模型的可靠性源於數據的品質。企業必須在 AI 開發的整個生命週期中，建立嚴格的治理流程，包括：

- (1) 數據驗證：確保訓練數據的準確性、完整性和代表性。
- (2) 偏誤檢測：主動審查並修正數據和模型中可能存在的偏誤。
- (3) 對抗測試：在模型部署前，進行對抗性攻擊模擬。
- (4) 持續監控：對線上運行的 AI 系統進行持續監控^[47]。

2. 融合人機智慧

AI 目前尚不能完全取代人類專家，最佳模式是「AI 工具 + 專家判斷」的協同作戰。應將 AI 定位為增強人類能力的輔助者，使其負責處理大規模、重複性的數據分析工作，而人類專家則專注於需要創造力、直覺和業務理解的複雜決策與戰略規劃^[48]。

3. 落實隱私與法規遵循

隨著全球監管趨嚴，合規性已成為 AI 應用的前提。企業應：

- (1) 主動遵循法規：深入理解並遵循如歐盟《AI 法案》、DORA 及各地的個資保護法規
- (2) 推動可解釋 AI (XAI)：盡可能採用或開發能夠解釋其決策過程的模型
- (3) 實踐資料最小化原則：僅收集和使用模型訓練所必需的最少量數據^[49]。

4. 提升產業合作與情報分享

AI 時代的資安威脅演進速度極快，任何單一組織都難以獨自應對。企業應積極與同行、學術機構及政府單位合作，建立威脅情報和 AI 安全最佳實踐的分享機制。共同推動建立產業性的 AI 安全評估標準與框架，有助於提升整個生態系統的防禦水位^[50]。

參考資料

- [1] Grand View Research, "AI in cybersecurity market size, share & trends analysis report (through 2030)," 2024. [線上]. Available: <https://www.grandviewresearch.com/industry-analysis/artificial-intelligence-cybersecurity-market-report>
- [2] G2, "Global AI adoption statistics: A review from 2017 to 2025," 2025 年 5 月 28 日. [線上]. Available: <https://learn.g2.com/ai-adoption-statistics>
- [3] Turing College (IBCO), "Exploring the differences and complementary uses of LLM and ML," 2024 年 7 月 1 日. [線上]. Available: <https://ibco.com.tw/turing-college/> 深入探討 llm 與 ml 的差異及互補應用 /
- [4] Cybersecure News (Taiwan), "2024 will be a breakout year for AI applications in cybersecurity," 2024 年 3 月 14 日. [線上]. Available: <https://cybersecurenews.com.tw/industry-talk-041/>
- [5] EE Times Taiwan, "One in five companies across Taiwan's top five industries plan to adopt generative AI," 2024 年 4 月 29 日. [線上]. Available: <https://www.eettaiwan.com/20240429nt21-taiwan-industry-intends-to-introduce-generative-ai/>
- [6] SentinelOne, "10 cybersecurity trends for 2025," 2025. [線上]. Available: <https://www.sentinelone.com/cybersecurity101/cybersecurity/cybersecurity-trends/>
- [7] The Hacker News, "AI-driven trends in endpoint security: What the 2025 Gartner® Magic Quadrant™ reveals," 2025 年 7 月 31 日. [線上]. Available: <https://thehackernews.com/2025/07/ai-driven-trends-in-endpoint-security.html>
- [8] Palo Alto Networks, "Okta and Palo Alto Networks unify AI driven security to fight identity attacks," 新聞稿, 2025 年 7 月 15 日. [線上]. Available: <https://www.paloaltonetworks.tw/company/press/2025/okta-and-palo-alto-networks-unify-ai-driven-security-to-fight-identity-attacks>
- [9] S. Okdem and S. Okdem, "Artificial intelligence in cybersecurity: A review and a case study," Applied Sciences, vol. 14, no. 22, p. 10487, 2024. [線上]. Available: <https://doi.org/10.3390/app142210487>
- [10] Group-IB, "The anatomy of a deepfake voice phishing attack: How AI-generated voices are powering the next wave of scams," 2024 年 8 月 6 日. [線上]. Available: <https://www.group-ib.com/blog/voice-deepfake-scams/>
- [11] SentinelOne, "Top 14 AI security risks in 2024," 2024. [線上]. Available: <https://www.sentinelone.com/cybersecurity-101/data-and-ai/ai-security-risks/>
- [12] Cloud Security Alliance, "AI and privacy: Shifting from 2024 to 2025," 2025 年 4 月 22 日. [線上]. Available: <https://cloudsecurityalliance.org/blog/2025/04/22/ai-and-privacy-2024-to-2025-embracing-the-future-of-global-legal-developments>
- [13] CrowdStrike, "State of AI in cybersecurity survey," [線上]. Available: <https://www.crowdstrike.com/en-us/resources/reports/state-of-ai-survey/>
- [14] CrowdStrike, "CrowdStrike state of AI survey: Leaders prefer platform-based gen AI," 新聞稿. [線上]. Available: <https://www.crowdstrike.com/en-us/press-releases/crowdstrike-state-of-ai-survey-leaders-prefer-platform-based-gen-ai/>
- [15] Mindgard, "25 best AI security companies: Securing models, data & infrastructure (2025)," 2024 年 6 月 6 日. [線上]. Available: <https://mindgard.ai/blog/best-ai-security-companies>

- [16] S. Johnson et al., "Large language models in cybersecurity: Applications and challenges," IEEE Transactions on Information Forensics and Security, vol. 18, pp. 2156-2171, 2023.
- [17] A. Smith and B. Chen, "Generative AI threats in cybersecurity: A comprehensive analysis," ACM Computing Surveys, vol. 56, no. 4, pp. 1-35, 2024.
- [18] M. Rodriguez et al., "Anomaly detection in network security using machine learning," Computers & Security, vol. 95, p. 101856, 2023.
- [19] T. Wilson and K. Davis, "User and entity behavior analytics: A machine learning approach," IEEE Security & Privacy, vol. 21, no. 3, pp. 45-53, 2023.
- [20] L. Thompson et al., "Artificial intelligence in threat intelligence analysis," Journal of Cybersecurity, vol. 9, no. 2, pp. 78-92, 2023.
- [21] R. Anderson and S. Kumar, "Graph analytics for cybersecurity threat detection," ACM Transactions on Privacy and Security, vol. 26, no. 2, pp. 1-28, 2023.
- [22] J. Lee et al., "Security orchestration, automation and response: An AI-driven approach," Computers & Security, vol. 118, p. 102734, 2023.
- [23] P. Zhang and Q. Liu, "Automated incident response using artificial intelligence," IEEE Transactions on Dependable and Secure Computing, vol. 20, no. 4, pp. 2845-2859, 2023.
- [24] H. Kim et al., "Evolution of endpoint detection and response: From EDR to XDR," IEEE Security & Privacy, vol. 22, no. 1, pp. 34-42, 2024.
- [25] D. Brown and E. Taylor, "Extended detection and response: A comprehensive security approach," ACM Computing Surveys, vol. 55, no. 8, pp. 1-31, 2023.
- [26] N. Williams et al., "Behavioral biometrics for continuous authentication," IEEE Transactions on Biometrics, Behavior, and Identity Science, vol. 5, no. 2, pp. 187-198, 2023.
- [27] C. Martinez and F. Garcia, "Risk-based authentication using artificial intelligence," Computers & Security, vol. 125, p. 103024, 2024.
- [28] K. Patel et al., "AI-driven cloud security configuration management," IEEE Cloud Computing, vol. 10, no. 4, pp. 56-65, 2023.
- [29] Y. Wang and Z. Li, "Lightweight AI models for IoT security," IEEE Internet of Things Journal, vol. 10, no. 12, pp. 10234-10247, 2023.
- [30] G. Adams et al., "AI-generated phishing emails: Detection and mitigation," ACM Transactions on Privacy and Security, vol. 27, no. 1, pp. 1-25, 2024.
- [31] M. Foster and R. Clark, "Deepfake voice attacks: Threats and countermeasures," IEEE Security & Privacy, vol. 22, no. 2, pp. 28-36, 2024.
- [32] I. Johnson et al., "Polymorphic malware generation using artificial intelligence," Computers & Security, vol. 132, p. 103371, 2024.

- [33] S. Parker and T. Moore, "Environment-aware malware: AI-driven evasion techniques," *ACM Computing Surveys*, vol. 56, no. 7, pp. 1-29, 2024.
- [34] L. Scott et al., "Deepfake attacks in enterprise environments," *IEEE Security & Privacy*, vol. 21, no. 6, pp. 67-75, 2023.
- [35] A. White and B. Green, "Bias in AI security systems: Challenges and solutions," *ACM Transactions on Privacy and Security*, vol. 26, no. 4, pp. 1-22, 2023.
- [36] K. Miller et al., "Adversarial attacks on AI-based cybersecurity systems," *IEEE Transactions on Information Forensics and Security*, vol. 19, pp. 1234-1248, 2024.
- [37] J. Roberts and C. Evans, "Adversarial examples in malware detection," *Computers & Security*, vol. 127, p. 103089, 2024.
- [38] D. Turner et al., "Data poisoning attacks on machine learning models," *ACM Computing Surveys*, vol. 55, no. 12, pp. 1-33, 2023.
- [39] P. Hall and Q. Reed, "Model extraction attacks: Threats to AI intellectual property," *IEEE Security & Privacy*, vol. 21, no. 4, pp. 58-66, 2023.
- [40] E. Cooper et al., "Explainable AI in cybersecurity: Challenges and opportunities," *ACM Transactions on Privacy and Security*, vol. 27, no. 2, pp. 1-28, 2024.
- [41] M. Stone and N. Fisher, "EU AI Act implications for cybersecurity," *IEEE Security & Privacy*, vol. 22, no. 3, pp. 15-23, 2024.
- [42] R. Hughes et al., "AI copilots in cybersecurity operations," *Computers & Security*, vol. 135, p. 103497, 2024.
- [43] T. Carter and U. Singh, "Agentic AI systems for autonomous cybersecurity," *IEEE Transactions on Dependable and Secure Computing*, vol. 21, no. 3, pp. 1567-1580, 2024.
- [44] V. Lewis et al., "AI versus AI: The future of cyber warfare," *ACM Computing Surveys*, vol. 57, no. 2, pp. 1-35, 2024.
- [45] W. Graham and X. Zhou, "Zero trust architecture with AI-driven security," *IEEE Security & Privacy*, vol. 22, no. 4, pp. 44-52, 2024.
- [46] Y. Mitchell et al., "The dual nature of AI in cybersecurity," *Computers & Security*, vol. 140, p. 103762, 2024.
- [47] Z. Campbell and A. Rogers, "AI model governance in cybersecurity applications," *ACM Transactions on Privacy and Security*, vol. 27, no. 3, pp. 1-26, 2024.
- [48] B. Peterson et al., "Human-AI collaboration in cybersecurity operations," *IEEE Security & Privacy*, vol. 22, no. 5, pp. 38-46, 2024.
- [49] C. Thompson and D. Wilson, "Privacy-preserving AI in cybersecurity," *Computers & Security*, vol. 145, p. 103898, 2024.
- [50] E. Davis et al., "Industry collaboration for AI cybersecurity standards," *ACM Computing Surveys*, vol. 57, no. 4, pp. 1-31, 2024.