

出國報告（出國類別：開會）

國際資訊安全會議
(RSA Conference 2025)
出國報告

服務機關：財政部財政資訊中心

姓名職稱：謝明峯 副主任

派赴地區：美國（舊金山）

出國期間：114年4月25日至5月4日

報告日期：114年7月31日

摘要

RSA Conference 為全球資訊安全領域最具規模與影響力的年度盛會之一，本次與會主要在於瞭解最新資通安全發展趨勢、技術與應用，以強化機關在資安治理與防禦架構上的能力。

AI 正在重塑威脅樣貌與防禦策略。AI 的賦能與挑戰無所不在，它既是提升防禦效率的利器，也是攻擊者的幫兇。這場攻防競賽的天平，將取決於誰能更有效地駕馭 AI。

國家級駭客的普及化與策略轉變，值得警惕；尤其對臺灣的潛在威脅，更值得我們持續關注。

從 AI 驅動的企業變革、國家級威脅的演進，到雲端原生與非人類身分的資安挑戰，環環相扣，描繪出一個更複雜的資安新常態。

目次

摘要.....	1
目次.....	2
壹、 會議簡介、過程及參加目的.....	3
貳、 研習摘要.....	5
一、 The AI-Powered Enterprise: A CEO's Mandate for Strategic Supremacy (AI 驅動的企業：CEO 邁向戰略優勢的使命).....	5
二、 Five Keys to Unlocking the Potential of AI Agents (解鎖 AI 代理程式潛 力的五大關鍵).....	6
三、 Making America Safe Again Through Cyber Defense (透過網路防禦，讓美 國再次安全).....	8
四、 Under Siege: How APTs and Nation-States Are Coming for Everyone (圍 攻：進階持續性攻擊與國家級威脅全面來襲).....	9
五、 World on the Brink: War, Geopolitics and Cybersecurity (世界瀕臨邊 緣：戰爭、地緣政治與網路安全).....	11
六、 Cybersecurity Year-in-Review and The Future Ahead (網路安全年度回顧 與未來展望).....	13
七、 From Exploit to Exfil: Rethinking a Cloud-Native Ransom Attack (從漏 洞利用到資料外洩：重新思考雲原生勒索攻擊).....	14
八、 Software Supply Chain Security, AI, and Regulation (軟體供應鏈的安 全、人工智慧與監管).....	16
九、 Securing Non-Human Identities in the Age of AI Agents (在 AI 代理時代 的非人類身分識別).....	18
參、 心得及建議.....	21

壹、會議簡介、過程及參加目的

一、會議簡介

RSAC 2025 (RSA Conference 2025) 為全球資訊安全領域最具規模與影響力的年度盛會之一，本年以「眾聲齊聚，共築社群」(Many Voices. One Community) 為主題，呼籲全球資安社群攜手應對挑戰。

會議特色之一為數個場館內數十個活動同時舉行，由參加者各自選擇感興趣的主題參加。活動型式包含：主題演講 (Keynote)、技術研討 (Track Session)、教程 (Tutorial)、研討會 (Seminar)、創新沙盒 (Innovation Sandbox)、推廣活動 (Event) 及產品展覽 (Expo) 等。

二、會議過程

會議於本年 4 月 28 日至 5 月 1 日於美國加州舊金山莫斯康中心 (Moscone Center) 舉行，同時於 Moscone West、Moscone South、Moscone North、YBCA (芳草地藝術中心) 及 YBCA Blue Shield of California Theater (YBCA 加州藍盾劇院) 等 5 個場館進行。

本年規模創下新高，吸引了近 44,000 名與會者、730 位演講嘉賓及 650 家參展商。

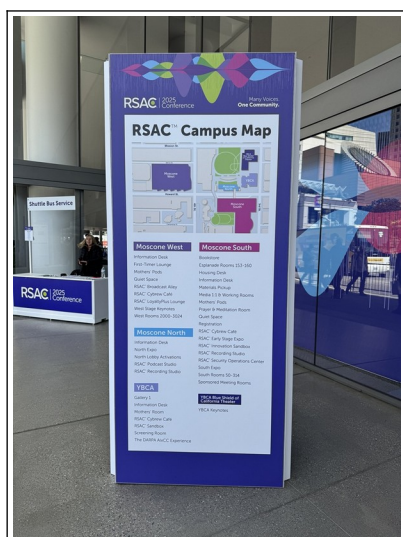


圖 1 及圖 2 RSAC 2025 的場館分布圖，及場館之一的 Moscone South。

三、參加目的

主要在於瞭解最新資通安全發展趨勢、技術與應用，以強化機關在資安治理與防禦架構上的能力，期能在現今嚴峻的環境中提升資安韌性。

本報告將聚焦於會議期間所參與之研習，摘要相關內容作為未來參考。

貳、研習摘要

一、The AI-Powered Enterprise: A CEO's Mandate for Strategic Supremacy (AI 驅動的企業：CEO 邁向戰略優勢的使命)

主持人：Cloud Security Alliance / President / Illena Armstrong

主講者：Procter & Gamble / Chief Technology Officer / Alan Boehme

重點摘要：

主講者曾在寶僑、ING、H&M 和可口可樂擔任高階職位。他認為，面對日益增長的全球性挑戰，企業需要根本性地改變利用技術的方式，而人工智慧（AI）是答案。他強調，前瞻的領導者現在就需要採取行動，成為由 AI 驅動的企業。

OpenAI 的出現立即引發廣泛地討論，這是自 1999 年 Netscape 帶來網路時代後未曾見過的情況。然而，許多 CEO 的反應是被動的，不確定如何應對。目前，全球只有不到 20% 的組織擁有涵蓋 AI 的明確策略，許多組織僅內部各自為政地嘗試。

AI 普及的主要障礙包括：

- (一) 部分高階主管對 AI 的恐懼或不理解，他們往往只將 AI 視為提高生產力或加快速度的工具，並尋找不採用的理由。
- (二) 組織內部的阻力，人資、法務和採購這三個部門常被視為阻礙的力量，因為他們將工作重點放在「保護」而非「賦能」，並傾向於維持現狀。
- (三) 將 AI 視為對現有基礎設施的「附加功能」，而非「典範轉移」。
- (四) 對 AI 的誤解，只談論技術本身而非其帶來的商業價值。

為了採用並成為 AI 驅動的企業，主講者提出以下建議：

- (一) 教育：教育內容應涵蓋業務風險和機會，而非僅限於技術或資安風險。教育不應只針對高階主管。安全團隊應站出來證明變革是安全的，並引領業務增長。
- (二) 領導力：管理團隊有責任向阻礙部門解釋不採用 AI 可能造成的損害，他們應該在轉型以提高未來競爭力方面發揮領導作用，從被動轉向主動。

- (三) 展示可能性：向業務部門展示技術能快速實現變革的能力。
- (四) 組織結構：可以考慮設立首席 AI 長 (Chief AI Officer) 或其他獨立的組織或子公司來推動 AI 專案，以規避內部阻力（如採購和人資的規則）。
- (五) 供應商：供應商應利用現有客戶關係的優勢，專注於提供業務價值和引領客戶進入下一代變革，而不是簡單地將 AI 作為附加產品銷售。
- (六) 從歷史中學習並加速：過去的技術轉變（如雲端）提供了經驗，但 AI 的發展速度會更快。
- (七) 整合與治理：需要預先建立強大的治理機制，包括數據、治理和道德委員會。開始培養內部人員（包括業務部門）的 AI 技能。將 AI 視為一個整合的流程。
- (八) 指標與策略：所有部門的目標必須與公司的業務增長、營收和利潤目標緊密相連，避免脫節。
- (九) 領導層換代：理解行動裝置和計算能力的年輕領導者採用 AI 的速度會加快。
- (十) 立即行動：抓住當前的機會，從被動轉向主動。透過教育、溝通和開放來推動變革。

主講者補充說明，部分使用 AI 的公司聲稱成本節省不到 10%，主講者認為這可能是公司只側重於成本削減，將 AI 視為類似於機器人流程自動化（RPA）的工具，且帳務處理可能導致 AI 採用數據被誇大。

二、Five Keys to Unlocking the Potential of AI Agents (解鎖 AI 代理程式潛力的五大關鍵)

講者 1：Microsoft / Dorothy Li

講者 2：Microsoft / Ryan Munsch

重點摘要：

AI 正在改變世界，就像電力的普及重寫了規則、帶來新的產業和工作方式一樣，AI 也處於一個類似的轉折點。AI 代理程式（agent）與傳統自動化不同，它們會推理、適應、預測，並與我們並肩工作。起初可能從小任務開始，但隨著時間推移，其影響將逐漸變得巨大。

為了充分釋放 AI 代理程式在資安領域（在其它領域也類似）的潛力，需要掌握五個關鍵：

（一）在使用者日常工作的地方提供工具：

情境切換會扼殺生產力，研究顯示情境切換會讓生產力下降，平均每次情境切換會浪費 9.5 分鐘。在資安領域，特別是在第一線處理警報的情境中，這可能是阻止攻擊和成為受害者之間的區別；尤其是資安領域的分析師，一天之中可能在不同系統間切換 12 到 15 次。關鍵是選擇能夠無縫整合到現有工作流程的代理程式，最好是將代理程式置於分析師的工作流程中。

（二）消除 IT 和資安工作中的繁瑣事務：

資安團隊常面臨大量的重複性任務和警報疲勞，例如，用戶提交的網路釣魚警報中超過 90% 是誤報，安全營運中心（SOC）頻繁接到的警報轟炸中 83% 可能是誤報。AI 代理程式擅長處理這類高投入、低腦力的任務，例如自動分類誤報、生成報告、日誌匯總等。這不僅是為了減少工作量，更為了通過自動化來讓資安團隊能夠專注於思考、深入調查和果斷行動。

（三）通過透明度建立信任：

代理程式必須展示其推理過程、引用來源、承認不確定性，並賦予用戶質疑、審計、驗證和調整的能力。可解釋性（Explanability）不是錦上添花，而是必不可少的功能；代理程式不應是黑箱，需要建立清晰的溝通和理解框架。可解釋的代理程式不僅建立信任，還能賦予專業人士快速且自信地行動的能力。人類應始終參與最終判斷。

（四）從被動安全轉向主動安全：

許多組織在識別和修補漏洞方面速度緩慢，有時需要多達 30 天或更久，這種延遲增加了風險；配置錯誤是攻擊者入侵的主要途徑之一，且常被忽視。代理程式具有幫助團隊在問題蔓延前發現潛在威脅的潛力，代理程式可以自動監控配置，收集情報，並針對複雜系統的漏洞修補和配置更改制定行動計畫，然後提交給團隊審查和執行。

（五）以同理心引導：

儘管代理程式帶來變革，但人類仍然是這場技術轉變的核心。設計 AI 代理程式時，必須以人為本，讓使用者感到被看見、被支持和被尊重。要讓實際使

用者參與設計、測試和反饋，並優先考慮建立人機協作的工作流程。好的代理程式應賦予使用者控制權，讓他們設定信心閾值、提供手動選項，並接受反饋，以幫助使用者做出更好的決策。

三、Making America Safe Again Through Cyber Defense (透過網路防禦，讓美國再次安全)

主持人：Dakota State University / President / José-Marie Griffiths

主講者：U.S. DHS / Secretary of Homeland Security / Kristi Noem

重點摘要：

美國國土安全部（DHS）部長表示，國土安全部的核心使命是確保美國的安全，並且以誠信的方式進行；她強調，網路安全即是國家安全，DHS 對涉及美國安全的人員、貨物及系統擁有管轄權，特別是關鍵基礎設施。

部長指出，世界面臨更多惡意行為者的威脅，他們不斷試圖滲透系統，利用技術對抗美國，網路戰爭和犯罪普遍存在，特別是攻擊關鍵基礎設施（如：水電等公用事業），其主要目的是為了控制它們，而非僅僅竊取數據。根據情報評估，中國仍然是美國最大的網路威脅，他們常從攻擊小型企業和地方政府這些「最薄弱環節」開始。

DHS 的目標是回歸其網路安全使命。部長指出 CISA（網路安全與基礎設施安全局，DHS 轄下機關）過去涉足部分非其應負責任，目前正在對 CISA 進行評估與調整，使其工作重點回到國會設定的網路安全任務。

部長強調需要打破政府內部的隔閡以及與情報機構之間的壁壘，加強跨部門和跨機構的合作，讓所有相關方都能參與解決方案的討論，並認為必須對惡意行為者追究責任，讓其承擔後果。

在未來計畫方面，DHS 將利用其採購能力推動「安全設計」(Secure by Design)，要求科技產品從一開始就內建安全性，客戶不應再為已包含在產品中的安全附加元件付費。

DHS 正檢視資源配置，計劃利用 AI 和先進加密技術來進一步保護資訊系統。

部長重申，DHS 和總統都高度重視網路安全，將其視為重要的國家安全責

任；她指出，DHS 轄下有 23 個機關構，職責廣泛，但網路安全貫穿所有領域。

四、Under Siege: How APTs and Nation-States Are Coming for Everyone (圍攻：進階持續性攻擊與國家級威脅全面來襲)

講者 1：Broadcom / Fellow, Symantec Threat Hunter Team / Eric Chien

講者 2：Broadcom / Vice President & General Manager / Jason Rolleston

重點摘要：

過去 5-10 年前，進階持續性攻擊 (APT) 與國家級威脅主要鎖定高端目標，如：政府、關鍵基礎設施及大型企業。然而，情況正在快速變化，現在這些威脅已廣泛影響所有產業及各種規模的組織；許多小型企業依賴「藉由隱匿來保持安全」，並認為規模小就不會被鎖定，但事實並非如此，數據顯示，85% 的攻擊目標是中型市場的跨產業組織。

主要國家級威脅來源及其特徵如下：

(一) 中國：

擁有眾多由國家支持的駭客團體，彼此相互獨立但共享情報和工具。他們極度活躍，攻擊範圍廣泛，主要動機是經濟間諜活動，例如竊取智慧財產權和談判策略。他們常將組織融入供應鏈的一部分，以達成最終目標。中國攻擊者傾向於重複使用工具，且不在乎是否被抓。值得注意的是，部分駭客團體會利用國家提供的工具兼差，為個人利益部署勒索軟體。

(二) 俄羅斯：

駭客團體與軍事情報機構合作，他們非常激進，導致顯著的外溢效應，波及中小型企業。SolarWinds 遭受 NotPetya 勒索軟體攻擊是著名案例，造成廣泛的全球影響。該國駭客團體的攻擊（如：ShockWorm）不一定複雜，但有效，經常透過大規模的釣魚攻擊，及利用被入侵的網路作為攻擊其他組織的跳板。

(三) 北韓：

動機經常純粹是經濟利益，他們針對加密貨幣組織、銀行（如：孟加拉國的銀行）、自動櫃員機等攻擊，並常使用勒索軟體。他們甚至建立「筆電農

場」(Laptop Farm)，由北韓人員於遠端工作，冒充美國公民賺取數百萬美元。北韓每年從這些活動中獲利數億甚至數十億美元。他們不關心政府或關鍵基礎設施，只關心金錢，但他們為了金錢可能使被攻擊的國家不穩定。

(四) 伊朗：

Seaworm 是該國自 2017 年起活躍的駭客團體，傳統目標是政府、國防、能源，最近則開始針對中東以外的公用事業、運輸、製造業等攻擊，包括中小型企業。他們常經由一般家庭中的路由器作為攻擊發起點，入侵後則使用滲透測試工具和 reg.exe 等合法工具竊取憑證，進行橫向移動。Seaworm 的顯著特點是大量使用合法、現成的工具，很少使用惡意軟體，使得傳統基於檔案特徵碼的防禦機制難以奏效。

應對這些威脅的防禦策略需要改變：

- (一) 傳統依賴「藉由隱匿來保持安全」的觀念已經過時。
- (二) 需要將大型企業使用的技術引入中小型市場，包括網路安全、端點偵測與響應 (EDR)、雲端安全和供應鏈安全。
- (三) 漏洞修補至關重要，尤其是在網路邊界的設備（如：虛擬私人網路 VPN、防火牆），許多中小型企業在這方面存在盲點。
- (四) 攻擊者大量使用合法工具進行「就地取材」(Living Off the Land) 攻擊，防禦需要基於行為分析，識別合法工具的惡意使用。自適應 (Adaptive) 防護技術可以根據環境基線，阻止特定工具的惡意使用，這能有效挫敗攻擊者。
- (五) 網路安全至關重要，包括針對內部流量（東西向）的深度封包檢測。
- (六) 端點偵測與響應 (EDR) 是必要的。
- (七) 新的趨勢是利用大型語言模型 (LLM) 來預測攻擊者的下一步行動。透過分析大量攻擊鏈 (Cyber Kill Chain) 數據，將事件視為「詞語」，攻擊鏈視為「句子」，AI 模型可以高準確度預測攻擊的下幾個步驟。這使得預防性封鎖成為可能。
- (八) 提升安全意識非常重要，威脅是真實存在。
- (九) 勒索軟體的滯留時間 (dwell time) 大幅縮短，過去平均約 4 天至一周，現在攻擊者常利用漏洞直接攻擊，在數小時內完成數據竊取和加密。

講者回答問題時提到，來自拉丁美洲的新興威脅，目前主要是大量的金融攻擊，針對當地的銀行流程和網站。但攻擊活動尚未達到像中國或前述國家的規模。

五、World on the Brink: War, Geopolitics and Cybersecurity (世界瀕臨邊緣：戰爭、地緣政治與網路安全)

主持人：Politico / Cybersecurity Reporter / Maggie Miller

主講者：Silverado Policy Accelerator / Dmitri Alperovitch

重點摘要：

內容主要為主講者闡述其著作《世界瀕臨邊緣》(World on the Brink) 的內容。

(一) 中國的網路威脅與臺灣議題

主講者處理中國的網路威脅已有 15 年，從早期針對美國進行的「極光行動」(Operation Aurora) 即已開始。然而，最近幾年與過去大相逕庭，尤其是「伏特颱風」(Volt Typhoon) (一個與中國政府相關的駭客組織) 的行動特別值得關注。這些行動被軍方稱為「戰場準備」，因為駭客組織入侵網路並持續潛伏，即使沒有太多有價值的資訊可竊取；但這些被入侵的網路對美國軍事行動、美軍及其盟友在全球關鍵區域的運作，以及破壞美國本土日常生活極有價值。最近還發現另一個組織「鹽颱風」(Salt Typhoon) 潛伏在美國電信網路中，並已入侵包括高層官員設備在內的許多裝置。

主講者認為中國一直在進行軍事建設和演習，而網路領域可能是重要的預警信號。中國非常專注於臺灣的通訊基礎設施，特別是海底光纖；中國的民用船隻頻繁「意外」切斷通往台灣的光纖，光是去年就發生了 8 起事件。

此外，主講者認為美軍依賴民用基礎設施（如：航空公司、鐵路、港口）運輸部隊、後勤和武器系統，這些設施正是前述行動的目標，目的是為了阻礙後勤運輸，爭取時間優勢。主講者另外提及，中國也可能攻擊美軍非保密資訊的專用 IP 網路 NIPRNet，來影響部隊戰備狀態。

(二) 俄羅斯的網路活動與教訓

在烏克蘭衝突之前，俄羅斯被認為是網路強國之一；然而，他們在戰爭中的表現並不理想，網路行動常顯得零散，似乎缺乏將網路行動整合的戰略思考。

從俄羅斯的行動可以學到一個教訓：利用網路進行威懾非常困難。俄羅斯試圖透過網路行動、縱火、暗殺等手段威懾西方國家停止援助烏克蘭，但沒有產生效果。

(三) 北韓的創新網路活動

主講者認為北韓是網路世界迄今為止最具創新性的參與者，他們走在如何利用網路達成戰略需求的尖端。北韓是第一個使用破壞性攻擊的國家，也是第一個利用加密貨幣和竊盜資助其導彈和核武計畫的國家。他們也是第一個利用遠端 IT 工作者進行網路活動的國家，這些遠端 IT 工作者進入美國數百家公司來獲取薪資，且隨時可能被用於間諜、勒索軟體或竊盜。

(四) 伊朗的網路能力與潛在衝突

伊朗過去幾年在網路行動方面進步非常大，一直針對沙烏地阿拉伯、以色列等磨練這些能力。主講者認為，如果伊朗的核設施成為攻擊目標，可以預期他們會在網路世界進行反擊。

(五) AI 在網路安全中的應用

所有國家級駭客都在使用 AI。北韓在利用 AI 方面似乎更成功，可能與其遠端 IT 工作者有關。AI 最直接的用途體現在社交工程上，它可以改進釣魚郵件的文法、製作深度偽造 (deep fake) 影音。AI 也被用於情報分析，協助總結資訊、識別趨勢和威脅。隨著未來技術進步，可能會出現更自動化或半自動化的攻擊，從而擴大行動規模；同時，AI 也將被深度整合到防禦技術中，維持著攻防之間的貓捉老鼠遊戲。

(六) 美國政府的應對與國際合作

主講者認為，美國現狀的網路安全系統需要進行重大改革，例如 CISA 的角色需要重新思考，並建議其應成為聯邦政府的資安長，為缺乏專業能力的部門提供服務。

在國際合作方面，儘管存在緊張關係（如關稅政策），但盟友基於永久的國家安全和經濟利益，主講者認為，無法與美國脫鉤，例如，「五眼聯盟」等

關係仍在正常運作。

六、Cybersecurity Year-in-Review and The Future Ahead (網路安全年度回顧與未來展望)

主持人：Silver Buckshot Ventures / Nicole Perlroth

主講者：Ballistic Ventures / Founder & General Partner / Kevin Mandia

重點摘要：

就數據來看：

- (一) 攻擊者主要通過幾種方式突破防線，漏洞利用 (Exploitation) 是最常見的方式。在 2023 年佔了 38%，這相較於 2022 年的 32% 有所上升，並與 2019 年之前主要依靠魚叉式網路釣魚 (spear fishing) 的情況有所不同。魚叉式網路釣魚在 2023 年佔 17%，而憑證被盜 (credentials stolen) 佔 10%。雖然具體比例有所波動，但主要入侵手段並沒有發生劇烈變化。
- (二) 然而，在勒索軟體攻擊中，入侵方式略有不同。暴力破解 (Brute forcing) 是勒索軟體攻擊中最主要的方式，攻擊者會對邊界電腦「轟炸」帳號和密碼進行嘗試。
- (三) 針對雲端基礎設施的入侵方式也完全不同，並且預計將會改變。目前已經出現了雲端原生入侵 (cloud-native intrusions)，意味著一些駭客天生就瞭解雲端環境，並且在雲端中運作得更有效率。
- (四) 主講者也列舉了一些最常見的漏洞。包括，Palo Alto Networks 的一項漏洞有超過 12 個駭客組織使用，其中一個勒索軟體組織 (ransom hub) 將其作為零日攻擊使用；Avanti 的兩個可以串聯攻擊的漏洞，被至少來自中國的五個不同組織使用；Fortinet 的一項漏洞被犯罪組織 (Finn 8) 發現並利用。一個值得注意的觀察，這些被利用的設備都無法安裝端點偵測與響應 (EDR)，這表明攻擊者正在針對防禦最薄弱的外部或邊緣攻擊。
- (五) 滯留時間 (dwell time) (從駭侵事件發生到被發現的時間) 過去曾趨於穩定，但在過去一年中，從平均 10 天增加到了 11 天。但主講者強調，這項指標可能偏差，因為他通常只在發生了事件後才會被聘請 (亦即他觀察到的數據會偏長)，而許多成熟的安全營運中心 (SOC) 能在更短時間內偵測到問題。

(六) 零日漏洞的數量每年都在上升，但在 2024 年的數量與前一年差不多。

主講者強調一個關鍵結論：資安事件很難避免，駭客每天都在進步，且攻擊者可以無限次嘗試，技術變革使得整體環境更加複雜。因此，資安專業人士必須提升自己的能力，並善用技術，例如 AI。AI 競賽正在進行中，它對進攻和防守都有利，問題在於誰能更好地利用它。

除了技術進步外，基本的資通安全衛生 (hygiene) (指良好的資通安全習慣) 也再次被強調。過去曾有人認為衛生對付不了頂尖駭客，但現在許多重大的資料外洩事件是利用既有的漏洞 (end day exploit)，而不是零日漏洞。因此，不斷修補漏洞和關閉面向網際網路的不必要通道至關重要。另外，身分安全 (identity security) 從未如此重要，幾乎每一個重大的資料外洩事件都與身分安全問題有關。

註：該場次中，主講者亦多次提及前文所提到「伏特颱風」(Volt Typhoon) 及「鹽颱風」(Salt Typhoon) 等事件，避免重複，此處不再摘要。

七、From Exploit to Exfil: Rethinking a Cloud-Native Ransom Attack (從漏洞利用到資料外洩：重新思考雲原生勒索攻擊)

主講者：Wiz / Director, Cloud Response / Yotam Meitar

重點摘要：

主講者分享了一個真實的雲端原生勒索攻擊，及隨後的事件應變過程，並表示，為了保護受害者身份，部分細節已被修改。

事件始於組織收到勒索信，多位高階主管都收到了類似的訊息，要求支付一筆巨額比特幣 (當時約 2000 萬美元)。攻擊者聲稱已竊取資料並加密檔案，並提供外洩資料的樣本作為證明。他們還提供了暗網網站連結，供受害者上傳被加密的檔案，以證明他們有能力解密。收到勒索信之前，組織並未察覺任何惡意活動，但該次攻擊直接導致業務中斷。

起初，組織感到較為安心，並基於下列兩個因素，決定不支付贖金：

(一) 資料備份在 Azure 雲端的不可變備份中，而主要的生產環境在 AWS 運行。組織認為有了不可變備份，就可以無需解密金鑰自行恢復。

(二) 攻擊者提供的樣本資料看起來並不包含高敏感資訊。

然而，好消息很快變成了壞消息。問題出在儲存於 Azure 雲端的備份，解開該備份的金鑰存放在 AWS 的金鑰管理系統 (HashiCorp Vault)，而進入該系統所需的關鍵檔案也被勒索軟體加密了，導致備份雖然完好無損，但無法取出使用。

面對困境，組織決定將進入金鑰管理系統所需的關鍵檔案上傳，請求攻擊者解密。這是一個有風險的決定，因為若被發現，攻擊者可能意識到組織無法取回備份並大幅提高勒索價格。最終，這個策略奏效，攻擊者傳回了被解密的關鍵檔案。組織成功取得關鍵檔案後，得以將 Azure 中的備份取回使用。

組織成功恢復了備份，並決定不支付贖金；但當截止日期過後，攻擊者公布了他們竊取的檔案，當中顯示，攻擊者擁有的資料遠比樣本多，並包含許多敏感的客戶個資。這凸顯了僅憑樣本評估洩漏範圍的風險。即使決定不支付贖金，與攻擊者進行接觸是推薦的做法（由談判專家進行，而非高階主管），這可以爭取時間用於調查和應變。

資安團隊通過回溯分析揭露了完整的攻擊鏈（細節略），並提出攻擊鏈每個步驟的預防和偵測方法：

- (一) 快速的漏洞識別和修補至關重要。且應對非正常活動進行分析，例如與特定機器不尋常通訊的 IP。
- (二) 防止賦予虛擬機 (VM) 過於寬鬆的權限。對罕見事件（如創建新使用者、罕見命令被執行）發出警報。
- (三) 偵測高權限操作，包含：新使用者被添加到管理員群組、高權限使用者首次執行聯邦登入 (federated login，指跨伺服器的登入)。在雲端中僅給予聯邦登入最低存取權限。
- (四) 禁用 AWS 的 IMDS (Instance Metadata Service，實例中繼資料服務) v1，並偵測 IMDS 罕見使用的命令。
- (五) 組織需要清楚知道哪些資料存在風險，且偵測 AWS 系統管理工具 (SSM) 使用不尋常的命令。

主講者並提出了三個關鍵的高層次結論：

- (一) 雲端原生的特定攻擊需要雲端特定的防禦：雲端環境與地端環境不同，攻擊面也不同。不能簡單地將地端環境的安全策略複製粘貼到雲端。
- (二) 快速風險補救的重要性：攻擊發生得非常快（本次攻擊的滯留時間不到兩天），必須快速、持續地識別風險並優先處理和修補。
- (三) 集中式監控：不同來源（如 AWS、Azure、OS 日誌、GitHub 日誌等）的所有數據需要整合到一個地方。這對於基於關聯性的有效偵測、反應速度以及調查至關重要。

八、Software Supply Chain Security, AI, and Regulation (軟體供應鏈的安全、人工智慧與監管)

講者 1：Futurum Group / VP / Mitch Ashley

講者 2：Sonatype / SVP / Tyler Warden

重點摘要：

該演講主要分享來自 Futurum Research 和 Sonatype 的研究報告及數據。

(一) AI 的廣泛採用與高層期望

我們正處於一個 AI 深入人心並實質採用的時代。全球的 CEO 正積極地在組織中針對 AI 進行試點專案，並對 AI 將如何影響業務寄予期望，儘管許多人尚不清楚具體影響是什麼。Futurum Group 一項針對 800 多位 CEO 的研究顯示，59% 的公司正在積極進行基礎性（初期試點）的 AI 專案；89% 的 CEO 承認利用 AI 的戰略重要性。這表明 AI 已成為組織高層關注並期望帶來效益的重點。

(二) 資訊長 (CIO) 和資安長 (CISO) 對 AI 的主要擔憂

85% 的受訪者將安全列為首要問題，緊隨其後的是資料合規性。這對資安業者來說是個好消息，意味著無需費力就能引起人們對 AI 安全的關注。然而，儘管需求增加，資安預算普遍面臨壓力，需要更有效地證明資安投入對專案的價值。

(三) AI 在軟體開發中的應用與挑戰

AI 正在影響軟體開發的各個環節。許多技術組織，包括資安和軟體開發團

隊，都期待利用 AI 來提高生產力。42% 的受訪者表示正在尋求在軟體開發和測試中使用 AI。然而，有報告顯示，在開發中採用 AI 的組織實際生產力反而下降，這可能與處於實驗和學習階段有關。目前，業界對 AI 在程式碼生成中的討論主要圍繞「速度」，而非「品質」。這是一個有趣的現象，但演講者認為，AI 真正的價值在於應用於工作流程和團隊，提升整體團隊的生產力，而非僅僅是個人效率。

AI 也正在快速融入現有的開發工具鏈和基礎設施。人們普遍在工作中使用 AI 工具，這導致了對「影子 AI」(Shadow AI，指使用不是組織批准的 AI 工具) 的擔憂。一個新興的趨勢是使用 MCP (Model Context Protocol，模型文本協定)，MCP 的快速普及也產生了新的資安和基礎設施的挑戰，例如需要維護、升級和修補漏洞。Kubernetes 作為主流的平台，也成為部署 AI 工作的常見選擇 (63% 的組織將 AI 工作容器化並部署在 Kubernetes 中)。

(四) 軟體供應鏈安全的現狀與挑戰

軟體供應鏈安全是一個關鍵的投資領域。主講者強調了開源在軟體開發中的重要性，但同時也指出了風險。Sonatype 的報告數據顯示，惡意開源軟體的數量持續增長，年增長率高達 150%，而 JavaScript (npm) 和 Python 套件的需求量也呈現顯著的年增長 (分別為 70% 和 80%)。儘管有漏洞修補，但持續的風險依然存在，例如 Log4j 漏洞在出現 3 年後，仍有 30% 的 Log4j 是易受攻擊的版本。雖然 SBOMs (軟體物料清單) 的使用有所增加，但軟體的透明度和可追溯性仍有不足，資訊產業仍待在這方面加強努力。

(五) AI 對軟體供應鏈的特定影響與新興威脅

特別是 AI 開源模型，給軟體供應鏈帶來了新的挑戰。目前最迫切的需求是提高可見度，了解組織正在使用哪些開源模型以及它們的來源，這與十幾年前圍繞開源的討論類似。除了傳統的程式碼漏洞，AI 還帶來了業務層面的漏洞和影響，例如，利用操縱數據或提示詞 (prompt)，可能導致 AI 在商業決策中出現偏頗，造成實際業務損失。

供應鏈的風險也延伸到了提示詞。提示詞不僅是我們輸入模型以生成內容或程式碼的指令，模型本身也會進行推理並生成、改進自己的提示詞。這形成了一個新的攻擊目標，攻擊者可以通過提示詞注入程式碼或繞過安全防護。未來的安全關注點將需要涵蓋從底層硬體到最上層的提示詞來源和生成過程。

(六) 制定 AI 政策和治理

對於制定 AI 政策，特別是針對大語言模型（LLM），目前各組織正在探索不同的方法。一些常見的做法包括：

- 在持續整合及持續交付（CI/CD）中阻止使用未經批准的模型。
- 定義允許使用的模型標準和來源。
- 利用網路可觀察性等技術，識別開發和生產環境中違反政策的 AI 使用。
- 建立沙箱環境，用於實驗未完全批准的模型。

(七) AI 管制政策

正圍繞幾個主要方面展開：

- 模型類型：區分商業通用模型、軟體中內嵌的 AI 功能以及用於內部開發實驗的開源模型。
- 技術決策：例如是否允許使用特定的技術框架（如 PyTorch vs. SafeTensor）。
- 數據類型：規定 AI 模型可以訪問哪些數據（內部 vs. 外部），以及是否允許在某些類型的內容上進行訓練。
- 應用程式分類：根據應用程式的類別或敏感度來限制 AI 的使用。

總體而言，AI 正以前所未有的速度改變軟體開發和安全領域，這帶來了巨大的潛力，但也引入了新的、複雜的供應鏈風險，需要業界和組織共同努力來解決。

九、Securing Non-Human Identities in the Age of AI Agents (在 AI 代理時代的非人類身分識別)

講者 1：Astrix Security / CEO / Alon Jackson

講者 2：Workday / Senior Director / Albert Attias

重點摘要：

該演講主要探討非人類身份（Non-Human Identity, NHI）的演變及風險，

特別是在 AI 代理程式 (agent) 日益普及的時代。

(一) NHI 的演變與定義

NHI 這個概念並非一開始就存在，它經歷了演變，最初稱為「第三方整合存取管理」，隨後轉向「應用程式對應用程式安全」，最終才有了 NHI 這個名稱，現已成為行業慣用的標準。

NHI 指的是所有獲得企業存取的非人類事物，這包括機器、服務、應用程式，特別是那些越來越具備 AI 能力的事物。例如：

- 獲得企業存取權限的第三方應用程式。
- 內部憑證，例如 SSH 金鑰、通行 token 等。
- 在自動化和 AI 環境中，日益自主運行的角色。

這些 NHI 正變得越來越獨立，它們在企業中扮演著類似「虛擬員工」的角色，例如幫助寫程式碼、寫電子郵件、做會議記錄等，而且無處不在，甚至在個人生活中也會出現協助我們完成任務的代理大軍。

(二) AI Agents 與 NHI 的關聯

特別值得注意的是代理式 AI (agentic AI)，它被視為多個 NHI 的結合。它是一種附有大型語言模型的自動化系統，代表用戶執行任務，這可能涉及到存取多個系統的金鑰。與傳統的機器人流程自動化 (RPA) 不同，AI 自動化具有「出乎意料」(unexpected) 的特性。AI 會自我改進，並可能以不同的方式達成相同的目標，這種不確定性需要以不同於過去的方式來管理。

(三) NHI 帶來的風險與挑戰

NHI 雖然能提升企業效率，但也帶來嚴重的安全風險。過去幾年發生的許多攻擊都與 NHI 有關。其中最常見的風險是權限過高，一個 API 金鑰就可能導致數百萬或數十億筆記錄的洩露；而不僅僅是資料洩露，還可能涉及操作風險，例如發送電子郵件、轉移資金或停止生產，這些甚至比資料洩露更有害。

此外，攻擊者可以利用權限過高的 NHI 進行橫向移動，例如，一個簡單的 AWS KMS 金鑰就可能成為攻擊的起點。

另一個嚴峻挑戰是缺乏有效的監控與治理。許多非人類身份，例如老舊的服務帳戶 (service accounts)，可能已經存在多年，所有者早已離職，但仍

保有存取權限，甚至沒有人知道它們的具體用途。在典型的人類離職流程中，相關的NHI 令牌或金鑰可能沒有被同步解除，這會留下不必要的安全風險。傳統上，企業對NHI 沒有完善的流程、生命週期管理和治理結構。

隨著AI 和自動化的普及，NHI 的數量正呈指數級增長。它們無處不在，其中一些權限過高。如何管理這些日益獨立的「虛擬員工」，追究其責任，是一個複雜的問題，因為它們不是可以被追究責任的法律實體。

參、心得及建議

一、心得

資訊安全攻防的典範已然轉移，過去我們熟悉的防禦邊界與攻擊手法，正被人工智慧（AI）與雲端原生環境徹底顛覆。會議中的多場演講，從不同面向共同指向一項事實：AI 正在重塑威脅樣貌與防禦策略。AI 的賦能與挑戰無所不在，它既是提升防禦效率、從海量數據中預測攻擊的利器，也是攻擊者用於擴大攻擊規模、自動化攻擊的幫兇。這場攻防競賽的天平，將取決於誰能更有效地駕馭 AI。

其次，國家級駭客的普及化與策略轉變，值得警惕。例如，「就地取材」（Living Off the Land）策略，大量使用合法工具進行惡意活動，使得傳統基於特徵碼的防禦機制難以奏效。另外，《世界瀕臨邊緣》（World on the Brink）一書中所提及對臺灣的潛在威脅，也值得我們持續關注。

再者，在《從漏洞利用到資料外洩》（From Exploit to Exfil）所揭示的雲端原生勒索攻擊案例，清楚地顯示了攻擊者利用雲端環境的權限設定不當與跨雲備份的依賴關係，製造出傳統地端環境思維無法預見的防禦死角。而《非人類身分識別》（Securing Non-Human Identities）中，更點出一項長期被忽視的巨大風險：數量龐大且權限過高的服務帳號、API 金鑰等「非人類身分」，正成為攻擊者橫向移動與權限提升的最佳跳板。

總結而言，本次會議所吸收的知識，從 AI 驅動的企業變革、國家級威脅的演進，到雲端原生與非人類身分的資安挑戰，環環相扣，共同描繪出一個更複雜、更動態、邊界更模糊的資安新常態。這也意味著我們的防禦思維，必須從被動回應轉向主動預測，從單點防禦轉向體系化的韌性建構。

二、建議

基於本次會議的觀察與心得，並針對當前的業務特性與挑戰，下列議題值得持續思考及研議，以強化本中心的資安防禦韌性：

（一）啟動「非人類身分」（NHI）的盤點與生命週期管理

可評估全面盤點所有系統、應用程式、服務之間的存取金鑰與服務帳號，

建立完整的 NHI 生命週期管理，確保其權限遵循「最小權限原則」，並在任務結束後能被確實移除或停用，防範其成為攻擊者的突破口。

(二) 強化軟體供應鏈的管理

本中心多數專案已導入 SBOMs (軟體物料清單)，可檢視是否已整合至日常程式庫進出館的管理流程，以確保在開發或程式改版階段即能阻擋已知的高風險漏洞，提升軟體供應鏈的安全。

(三) 深化雲端原生環境的安全組態與防禦策略

雲端勒索攻擊案例顯示，地端的安全策略無法直接複製到雲端。應持續關注有無公有雲環境適用的安全組態基準 (如：雲端版的 GCB)。此外，應整合不同公有雲專案的監控數據至單一安全營運中心 (SOC)，進行集中式監控與關聯分析，以實現對雲端原生攻擊的快速偵測與反應。

(四) AI 代理程式 (agent) 的治理框架

代理程式在 AI 賦能的狀況下，可能執行意料外的動作，也可能存取預期外的資料。未來本中心導入 AI 代理程式時，應同步評估其治理方式。