



林業及自然保育署
臺東分署

115年8月

公務機密 資訊安全維護



目錄

01

資安時事案例

- AI詐騙浪潮衝擊企業資安 深偽與社工手法致巨額損失
- 駭客用傳統文字加鹽手法，騙過新式AI郵件防護

02

個人資料保護法

- Pi拍錢包遭駭！官方證實350萬用戶個資外洩 補償出爐
- Q&A 60則個資法常識(每月一常識)

Q28 在什麼情況下，才可將個人資料做為其他之用？

03

生活中的資安

- Meta AI眼鏡遭貼「變態眼鏡」標籤用戶坦言不敢戴出門

04

數位學習

[養龍蝦熱潮！用戶遭OpenClaw報復？](#)



AI詐騙浪潮衝擊企業資安 深偽與社工手法致巨額損失



圖／示意圖

service@sunmedia.tw (商傳媒 SUN MEDIA)

2026年7月14日週二 上午11:46

商傳媒 | 葉安庭／綜合外電報導

隨著人工智慧（AI）技術快速發展，企業正遭遇一波由AI驅動的社交工程詐騙浪潮。這類高度鎖定的詐騙活動，利用AI生成逼真的文字、語音和影像，誘騙企業內部人員洩露敏感資訊、授權詐欺性交易或未經授權存取關鍵系統，其規模與複雜度均大幅攀升。

根據聯邦調查局（FBI）報告，2025 年美國網路犯罪總損失已超過 200 億美元，其中超過 30 億美元與商業電子郵件詐騙有關。資安公司 KnowBe4 於 2026 年 4 月發布的報告顯示，網路釣魚攻擊在過去六個月內增加了 17.1%。單次成功的攻擊不僅會導致資料外洩、監管調查與民事訴訟，更會對企業聲譽造成長期損害，2025 年美國資料外洩事件的平均成本高達 1,022 萬美元。

詐騙者日益運用深偽技術（透過AI合成高度擬真的音訊與影像），利用公開可得的聲音、筆跡等資訊來模仿組織內部或外部的聯繫對象。2025 年 5 月，聯邦調查局曾發布警示，指出有惡意行為者冒充美國政府高階官員，透過文字和AI合成語音管道，引導收件人點擊惡意連結。詐騙者常冒充高階主管、第三方顧問（如外部律師）或 IT 部門人員，利用員工傾向服從權威的心理，要求緊急或秘密的資金轉帳。

大型語言模型（LLMs，一種能理解並生成人類語言的AI模型）讓詐騙者能大規模產生潤飾精良、具說服力的詐騙電子郵件，並從領英（LinkedIn）等公開平台收集資料，量身打造攻擊內容。詐騙者甚至會要求目標員工簽署保密協議，以推進虛假的併購交易，並強調其保密性，阻止員工向高階主管或法務部門諮詢，藉此規避內部查核。

因應新興威脅，各國監管機構已強化對企業資安的審查。例如，對於有義務向美國證券交易委員會（SEC）揭露資訊的公司，SEC 的網路安全揭露規則要求在確定重大網路安全事件後，通常需在四個工作天內進行揭露。歐盟金融實體依《Digital Operational Resilience Act》規定，在確定為重大事件後的四小時內，以及從首次發現事件起的 24 小時內，均須啟動初步通報。同時，《一般資料保護規範》（GDPR）及其他隱私法規，也要求在資料外洩後通報主管機關及受影響個人。主管機關與執法單位日益期望企業採行基於風險的治理、分層技術控制、強健的身分驗證程序，以及適當的董事會與管理層監督。

為此，企業應定期實施強制性員工訓練，特別針對深偽辨識和AI生成通訊內容的識別。針對高階主管、財務人員及資訊科技（IT）員工等高價值被冒充目標，應提供更具針對性的訓練。企業亦應強化財務控制，實施多層級核准機制以進行資金轉帳，並對透過電子郵件收到的所有財務請求，透過次要且獨立的管道（如安全訊息平台或已知電話回撥）進行驗證。同時，應利用AI驅動的內容偵測技術，輔助傳統電子郵件過濾系統，並重新評估多因子驗證（MFA，除了密碼外，還需額外驗證步驟才能登入的方式）的實施方式。此外，確保網路保險（cyber insurance）涵蓋AI驅動的社交工程詐騙，並具備足以應對潛在損失的額度，也至關重要。企業應維持穩健的事件應變計畫，並定期透過桌面演練進行測試。

駭客用傳統文字加鹽手法，騙過新式AI郵件防護

資安業者Barracuda自4月以來偵測逾100萬次「文字加鹽」釣魚攻擊。駭客把惡意訊息藏在大量隱藏的無害文字中，干擾機器學習與LLM判斷，使釣魚信可能被誤認為正常郵件



iThome

文/陳曉莉|2026-07-20發表

資安業者Barracuda上週四（7/16）揭露，駭客正利用名為「文字加鹽」（Text Salting）的傳統垃圾郵件規避手法，干擾由機器學習及大型語言模型（LLM）驅動的郵件防護工具。Barracuda自今年4月以來已偵測到逾100萬次相關釣魚攻擊，主要偽裝成零售業者寄送的點數、禮品卡或限時兌換通知。

文字加鹽是指在惡意郵件中加入大量隨機且看似無害的文字，以稀釋獎勵、到期或信用卡等可疑詞彙所占的比例，降低郵件被判定為垃圾郵件或釣魚信的機率。這些填充內容可能是隨機故事、一般對話或專案筆記，並大量使用小狗、訓練、筆記、任務及書籍等無害詞彙。

為了不讓收件人察覺，駭客會透過CSS將填充文字的顯示範圍裁切至零、把文字移到畫面外，或將字體大小設為零。部分現代郵件防護工具已能強制顯示並分析這些隱藏文字，因此駭客開始同時採用多種隱藏技術；即使其中一種被識破，其他方式仍可能讓填充內容保持隱藏。

Barracuda指出，這種多年來用來躲避傳統垃圾郵件過濾器的手法，如今也能混淆部分AI郵件防護工具。部分LLM直接分析郵件文字及HTML原始碼，除非經過特別訓練或指示，否則未必知道哪些內容對收件人不可見。隱藏的無害文字因而可能改變模型對郵件意圖、用途及風險的判斷，使釣魚信遭錯誤歸類為正常郵件。

生成式AI也讓文字加鹽更容易大規模運作。攻擊者可以快速產生大量內容自然且彼此不同的填充文字，將短短數十字的惡意訊息藏進數百字的無害故事中，增加追蹤與偵測難度。

Barracuda建議企業採用多層式郵件防護，不應只依賴關鍵字或AI內容分析，而應同時檢查郵件結構、寄件者信譽、身分驗證結果、內嵌網址及HTML顯示方式，並比較收件人實際看到的內容與郵件原始碼，找出遭刻意隱藏的文字。

Pi拍錢包遭駭！官方證實350萬用戶個資外洩 補償出爐



PChome旗下第三方支付服務「Pi拍錢包」日前遭駭客組織Settra聲稱入侵。(示意圖／翻攝自pixabay)

CBC東森新聞

實習編輯 林予安 2026年7月17日週五 下午7:32

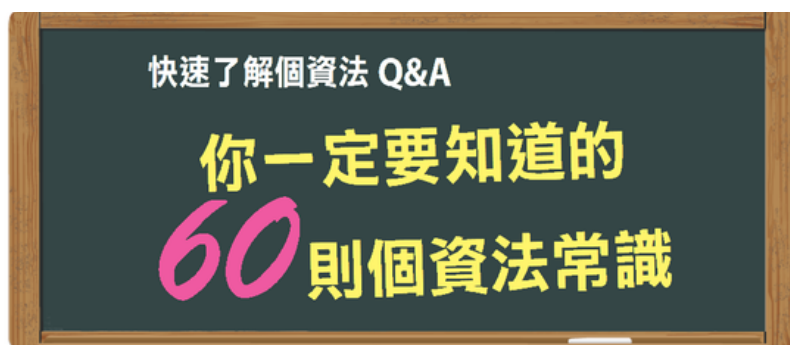
PChome旗下第三方支付服務「Pi拍錢包」日前遭駭客組織Settra聲稱入侵，引發外界關注逾350萬名用戶個資是否外洩。拍付國際今（17）日公布完整調查結果，證實確有部分會員資料遭非法存取，但強調外洩內容並非交易系統資料，而是內部員工日常作業檔案，涉及的資訊包含部分會員姓名、手機號碼及車牌號碼。官方也同步宣布補償方案，符合資格的受影響用戶可領取價值超過500元的超商飲品電子優惠票券。

駭客利用主機防火牆漏洞取得系統權限

拍付國際表示，經資安團隊與第三方專業鑑識機構調查，本次事件起因於駭客利用主機防火牆漏洞取得系統權限，隨後非法進入公司內部員工檔案伺服器，導致部分作業檔案遭到存取。進一步確認後發現，外流資料主要來自2019年及2022年的專案文件，內容包含部分會員姓名、手機號碼及車牌號碼。

不過拍付國際強調，此次事件並未影響Pi拍錢包App、官方網站及交易系統。包括用戶帳號密碼、交易紀錄、信用卡卡號、點數等重要資訊，皆儲存在獨立且具高度防護的系統內，經全面檢查後確認沒有遭駭客存取，也未發現交易資料外洩情形。對於此次個資事件，拍付國際表示深感歉意，坦言未能妥善保護部分會員資料，已完成必要的系統修補與資安強化措施，希望降低類似事件再次發生的風險。

每月一常識



新版個人資料保護法與過去有很大的不同，新法進一步擴大了個人資料的保護範圍，並且讓所有產業一體適用；新法甚至首度增加團體訴訟，而且違法的罰則也加重了，企業老闆要負更大的責任。接下來，我們以60個Q&A，快速帶你認識新版個人資料保護法

Q28 在什麼情況下，才可將個人資料做為其他之用？

A 企業在以下情況下，才可以將個人資料做為原訂目的之外使用：

1. 法律明文規定。
2. 為增進公共利益。
3. 為免除當事人之生命、身體、自由或財產上之危險。
4. 為防止他人權益之重大危害。
5. 公務機關或學術研究機構基於公共利益為統計或學術研究而有必要，且資料經過提供者處理後或蒐集者依其揭露方式無從識別特定之當事人。
6. 經當事人書面同意。

Meta AI眼鏡遭貼「變態眼鏡」標籤 用戶坦言不敢戴出門



Meta AI 眼鏡近期因偷拍事件及隱私疑慮備受爭議。
(圖／路透)

ETtoday新聞雲

ETtoday 的故事/2026/7/17

Meta 與 Ray-Ban 合作推出的 AI 智慧眼鏡近期因偷拍事件及隱私疑慮備受爭議，不少使用者坦言，如今已不敢在公共場合配戴，甚至擔心因此被貼上「變態」標籤，產品形象也因此受到衝擊。

根據《Engadget》報導，部分創作者利用 Meta AI 眼鏡內建鏡頭，在未經他人同意的情況下拍攝陌生人，尤其有男性創作者錄下搭訕女性的過程並發布至社群平台，引發外界對偷拍及個人隱私的質疑。

偷拍內容引發反感 「變態眼鏡」稱號流傳

報導指出，除了未經同意拍攝陌生人外，甚至曾出現有人利用秘密錄影內容向受害者索取金錢的案例，加深外界對智慧眼鏡的負面觀感。另一方面，Meta 過去曾因資料蒐集及隱私爭議受到批評，如今又積極將臉部辨識等 AI 技術導入智慧眼鏡，也讓不少民眾擔憂個人資料可能遭到濫用。

由於相關爭議持續延燒，不少網友甚至直接將 Meta AI 眼鏡稱為「變態眼鏡（Pervert Glasses）」，相關標籤也開始影響一般使用者。

使用者：戴出去怕被誤會

旅遊內容創作者 Danielle 表示，她原本對這款產品充滿期待，但如今已不願在公共場所使用。她認為，部分男性使用者的不當行為破壞了整體產品形象，因此自己也不希望別人因為看到她配戴智慧眼鏡而感到不自在。她更直言，如今這副眼鏡「就像昂貴的紙鎮」。

另一名自由影音製作人 Will Kujawa 則表示，自己原本打算購買 Meta AI 眼鏡，但看到大量網友留言後改變想法。他表示，不少人認為戴著這類眼鏡的人就像偷拍者或意圖不軌的人，也讓他意識到，「很多場合其實並不適合把攝影機直接戴在臉上」。

Meta 仍看好智慧眼鏡市場

儘管爭議不斷，Meta AI 眼鏡的銷售表現仍明顯優於早年的 Google Glass，其他科技公司也正積極投入智慧眼鏡市場。

Meta 近日也找來名人 Kylie Jenner 參與最新宣傳活動，執行長 Mark Zuckerberg 依然看好智慧眼鏡的發展，並認為這類裝置未來有望逐步取代智慧型手機，成為下一代主要運算平台。

不過，在偷拍事件、隱私爭議及負面觀感持續發酵之下，即使產品銷售持續成長，如何化解外界對智慧眼鏡的疑慮，仍是 Meta 未來推廣這類產品必須面對的挑戰。