

2025

第三季電子郵件安全觀察



ASRC

Spam Mail

Virus Mail

Malicious Mail



2025 年第三季，電子郵件攻擊呈現「手法更精緻、攻擊鏈更長、利用合法服務為跳板」的趨勢。

攻擊者大量採用 AI 生成式工具強化社交工程，使郵件文案更具說服力且在多語系場景下更難以辨識；同時，他們善用第三方服務，如短網址、雲端平台、電子簽章、以及資安廠商自己的置換連結機制等，來串接、轉向與掩飾攻擊路徑，使既有的靜態過濾與基於來源信譽的防護失效。針對資安廠商的社交工程試探也明顯增加，攻擊者以「少量多次」的方式對公開服務窗口埋伏惡意程式或蒐集情報，嘗試取得長期潛伏的後門。最後，AI 的普及不僅提升釣魚攻擊成功率，也使漏洞挖掘、自動化探針與隱蔽指令的濫用更有效率，導致端到端的電子郵件防護必須由單層規則升級為行為與情境感知、跨通道監控與資料外洩的整合防護體系。

以下為 ASRC 與中華數位科技在這一季的特殊觀察：

置換連結防護遭到濫用

為降低收件者誤點惡意連結的風險，許多郵件防護機制在郵件傳遞途中執行「置換連結 (URL Rewriting)」，將郵件內原始連結改寫為防護服務自有的檢查跳板，以便在使用者點擊時即時檢查最終目的連結的安全性，並記錄使用者與點擊時間，以利事後鑑識或封鎖。這種即時檢測機制在傳統釣魚攻擊中有效降低成功率，並提高事件追溯能力。

然而在 2025 年第三季出現明顯濫用的趨勢：攻擊者串接多個置換連結與合法跳轉服務（例如短網址、合法雲端或第三方追蹤域名），形成「多段轉址」的攻擊鏈。其操作邏輯與風險如下：

1. 串接防護跳板以躲避即時檢測

攻擊者先利用合法服務（或被入侵的服務）生成短網址或跳轉連結，再將這些連結放入釣魚郵件中。當收件者點擊時，第一個被檢查到的 URL 可能是某資安廠商或其它合法防護的置換域名，因其來源被視為「可信」，系統就不會進一步深度解析或標示為可疑，導致最終惡意目的地得以通過。

2. 繞過記錄與追蹤機制

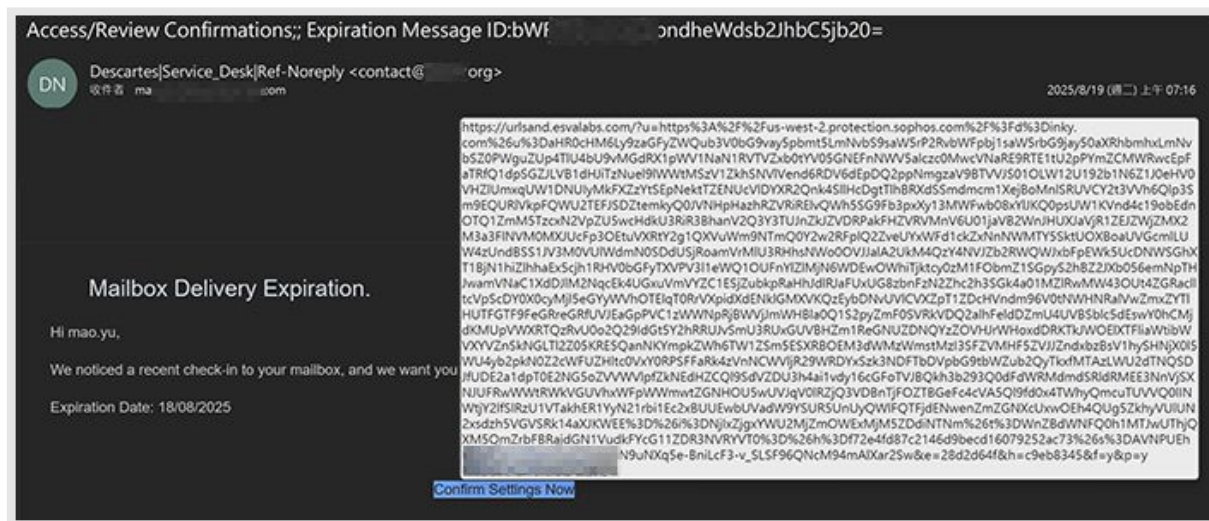
若多段轉址使中間某些 Click-tracking / 置換節點被系統視為正常流量，系統可能不會完整紀錄最終目的地或點擊者資訊，削弱事後鑑識與責任歸屬。

3. 利用合法資源作掩護

當轉址鏈包含受信任的第三方（例如廣告追蹤、電子簽章或大品牌雲端），攻擊行為顯得更「自然」，讓使用者與系統更難辨認其惡意意圖。

4. 自動化與規模化

攻擊者可以自動化生成大量多段轉址連結，配合 AI 編寫的人性化文案，顯著提升釣魚效率。



釣魚郵件的攻擊者嘗試串起不同防護的置換連結，並搭配縮址、轉址的功能，讓置換連結防護失效

潛在影響

防護閘道的「第一層檢查」被合法外殼所迷惑，導致「偽陰性」增加；事後鑑識資訊不完整，延遲事故回應與補救；以及使用者信任度下降（尤其當合法廠商的置換域被濫用時），導致品牌與服務信譽風險。

針對資安公司的社交工程攻擊與試探

第三季觀察到攻擊者特別將目標瞄準「資安公司」或其公開服務窗口，常見攻擊路徑與手法可區分為兩大類：

1. 長期潛伏式 Web / 服務窗口滲透

攻擊者嘗試以惡意程式感染服務窗口，成功感染後，以小量、頻繁的請求（通常透過 HTTPS / 443）取得後續惡意程式，目標是建立可長期維持的後門或定期蒐集目標主機資訊，並定期將資訊上傳到特定外部站點。攻擊行為刻意低調（低頻率、分散來源 IP、混淆 User-Agent），以避免被即時偵測系統標為異常流量。

2. 社交工程與商務洽談偽裝

以「購買服務」、「產品諮詢」或「技術合作」之名接觸業務承辦人，誘導其提供企業內部資訊、技術細節或測試存取權限。手段常結合精緻的語言、模擬的公司文件與偽造聯絡人資訊，單靠表面核查難以立即識破。

常見破綻

- 發信來源與真實性不一致

不少攻擊使用的郵件並非來自他們聲稱的公司域名或官方郵件流程，若針對郵件頭、來源 IP 與 SPF / DKIM / DMARC 進行核查，仍可發現破綻。

- 表單回應機制缺少驗證

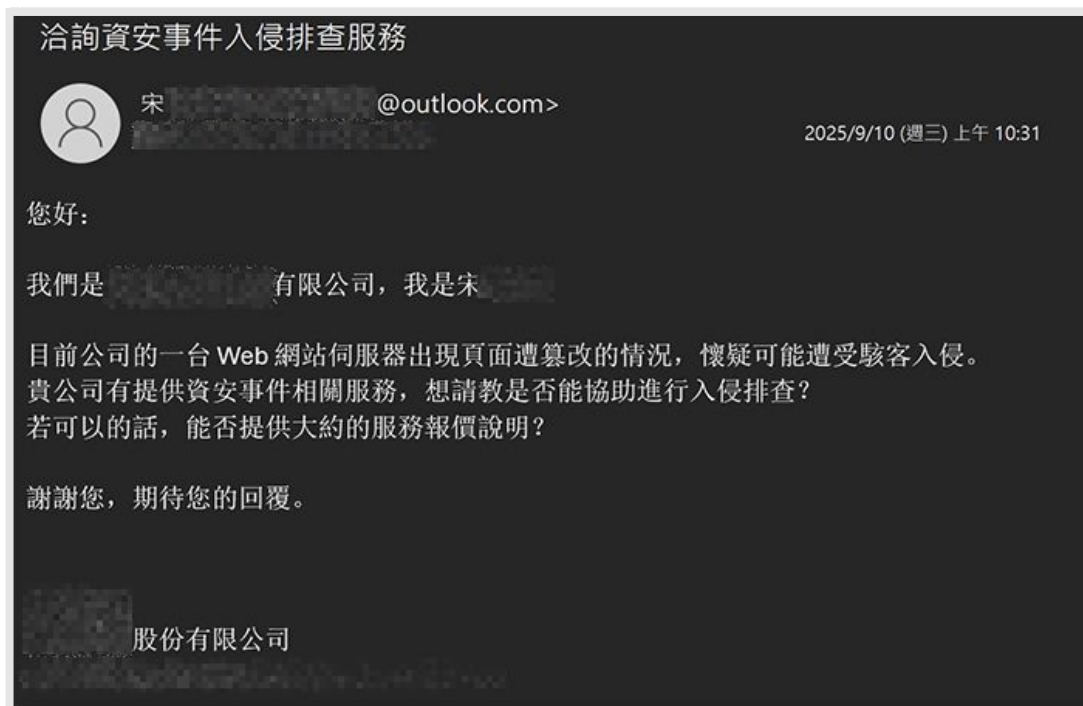
公司若以網頁表單作為第一接觸點，但未對回覆者進行強制驗證（例如電話回撥、企業郵件網域驗證或商業憑證），將提高被社交工程騙取資訊的風險。

- 內部資訊過度披露

公開的 FAQ、技術支援說明或產品文件若包含過多架構或技術細節，能被攻擊者快速收集並用於定向攻擊。



▮ 試圖在服務窗口的電腦上埋入可以長期潛伏並洩資的後門



▣ 試圖誘騙相關業務承辦人洩露過多的訊息或技術情報

AI 進化帶來的威脅

AI 與大型語言模型 (LLM) 在 2025 年下半年已廣泛被攻防雙方採用，對電子郵件資安的影響主要有三個面向：

1. 強化社交工程內容產出

AI 可生成高品質、針對特定組織或個人語氣與文化語境的郵件文案，包含合理的時間脈絡、專業術語與稱謂，有效提高釣魚與偽冒成功率。並且能自動化 A/B 測試郵件標題、內容與呼籲動作，快速優化可欺騙率。

2. 自動化漏洞發現與攻擊鏈組合

攻擊者利用 AI 加速漏洞掃描、解析郵件伺服器或附件的潛在弱點，並自動生成對應利用代碼或 payload。當 AI 結合自動化工具（如腳本、代理、多段轉址生成器）時，可以大規模產生變異化攻擊，使傳統簽名式防禦失效。

3. 對企業內部 AI 系統的濫用（prompt injection / 隱藏指令）

隨著企業導入 AI 助手處理郵件（如自動摘要、回覆建議、敏感資訊檢測），攻擊者可能在郵件正文中嵌入隱蔽指令（例如極小字體、白底白字、或特殊格式），誘使 AI 揭露敏感資訊或執行不當行為（稱為 prompt injection）。若 AI 的輸出未經適當的審核或上下文限制，可能成為內部資料外洩或錯誤自動化決策的來源。

電子郵件攻防邁入新階段

電子郵件攻防正進入「以合法性與自動化為盾與矛」的新階段。攻擊者大量利用 AI 生成社交工程與自動化工具，並利用合法第三方與置換連結的信任層來掩護惡意路徑；同時，資安供應鏈本身與對外窗口成為高價值目標。單一層級的靜態防禦（例如只靠 SPF / DKIM / 傳統過濾）已不足以應付這類複合、動態攻擊。

未來趨勢預測

- 多段轉址與合法服務濫用將更加普遍，防護會從「域名信譽」轉向「轉址鏈分析」與「行為得分」。
- 攻擊者對資安業者與業務窗口的試探會持續，促使資安公司本身採用更多“對抗式”自我測試與服務窗口硬化（hardening）。
- AI 相關的 prompt injection 與模型濫用將成主要攻擊向量，企業若不設限，AI 反而可能成為資料洩露的幫兇。

企業防護建議

- **強化身分與接觸驗證流程**
對外業務 / 客服 / 表單回應採用多因子驗證與實體回撥核實。
- **AI 使用原則與防護**
針對內部 AI 處理郵件的流程設計輸入淨化、輸出審核與最小授權。
- **釣魚演練與社交工程防禦訓練**
針對 VIP / 財務 / 客服進行定向演練與應變流程訓練。
- **建立業界協作與即時通報機制**
當發現被濫用的第三方或置換域名，應快速通報、分享入侵指標 IoCs 與同步封鎖。

關於ASRC 垃圾訊息研究中心

ASRC 垃圾訊息研究中心 (Asia Spam-message Research Center)，長期與中華數位科技合作，致力於全球垃圾郵件、惡意郵件、網路攻擊事件等相關研究事宜，並運用相關數據統計、調查、趨勢分析、學術研究、跨業交流、研討活動..等方式，促成產官學界共同致力於淨化網際網路之電子郵件使用環境。

更多資訊請參考 www.asrc-global.com



ASRC垃圾訊息研究中心