

資安戰情室：CompTIA SecAI+ (CY0-001) 戰術手冊

掌握 AI 時代的防禦與反擊，全面
解析六大核心模組

任務代號：CY0-001

適用對象：IT 與資安專業人員、系統管理員、安全分析師

主講人：James Lee



戰略藍圖：六大防禦階段

階段一：洞悉地形 (Understanding the Terrain)

模組 1：AI 基礎概念與模型訓練
(核心引擎)

階段二：鞏固防線 (Fortifying the Perimeter)

模組 2：AI 威脅建模與系統防護
(預警雷達)
模組 3：AI 存取控制與資料安全
(多重防護門)

階段三：實戰交鋒 (Engaging the Enemy)

模組 4：AI 相關威脅與補償控制
(主動防禦)
模組 5：AI 輔助資安實務
(自動化守衛)

階段四：全局指揮 (Commanding the Operation)

模組 6：AI 治理、風險與合規
(指揮中心)



階段一：模組 1 - AI 武器庫類型矩陣 (The AI Typology Matrix)

規則型 AI (Rule-based)	機器學習 & 深度學習 (ML & DL)	Transformer 模型 & 生成式 AI (GenAI)
主要資安應用： 辨識已知威脅與固定特徵。	主要資安應用： 行為異常分析、惡意軟體分類、 威脅情資處理。	主要資安應用： 理解自然語言、分析釣魚郵件、 生成安全腳本與威脅模擬。
適應能力： 低（僅依循預設邏輯）。	適應能力： 高（從海量非結構化數據中學習 模式）。	適應能力： 極高（強大的內容生成與脈絡 理解）。
防禦盲區： 無法應對未知或變種攻擊。	防禦盲區： 易受資料投毒 (Data Poisoning) 影響。	防禦盲區： 提示注入 (Prompt Injection)、 幻覺與潛在偏見。



駕馭核心：提示工程 (Prompt Engineering) 解剖圖

系統提示 (System Prompts)：設定 AI 的靈魂

為模型設定「角色」與「行為邊界」（例如：你是一位嚴謹的資安分析師）。



使用者提示 (User Prompts)：下達戰術指令

提供具體問題或執行指示。

**戰術引導等級
(Guidance Levels)：**

Tier 1
Zero-shot (零樣本)：
不給範例，直接測試 AI 的既有知識。

Tier 2
One-shot (單樣本)：
給予一個成功範例，確立輸出格式。

Tier 3
Multi-shot (多樣本)：
給予多個範例，用於處理高度複雜或精密的資安任務。



攻擊者的武器：生成式 AI 武器化實例

```
import openai

openai.api_key = "sk-YOUR_EXPOSED_API_KEY_HERE"

response = openai.ChatCompletion.create(
    model="gpt-5-mini",
    messages=[
        {"role": "system", "content": "扮演一名資深 CFO，  
撰寫一封緊急的匯款通知給財務經理..."},
        {"role": "user", "content": "立即支付，否則我們將  
失去這筆交易。"}
    ],
    temperature=0.7
)

print(response.choices[0].message.content)
```

威脅情境：

攻擊者利用 OpenAI API（如 gpt-5-mini）結合自動化腳本，大規模生成高度逼真的商業電子郵件詐騙（BEC）。

參數微調的威脅：

temperature=0.7：保持一致性同時增加變異，使多型態釣魚更難被偵測。

防禦啟示：

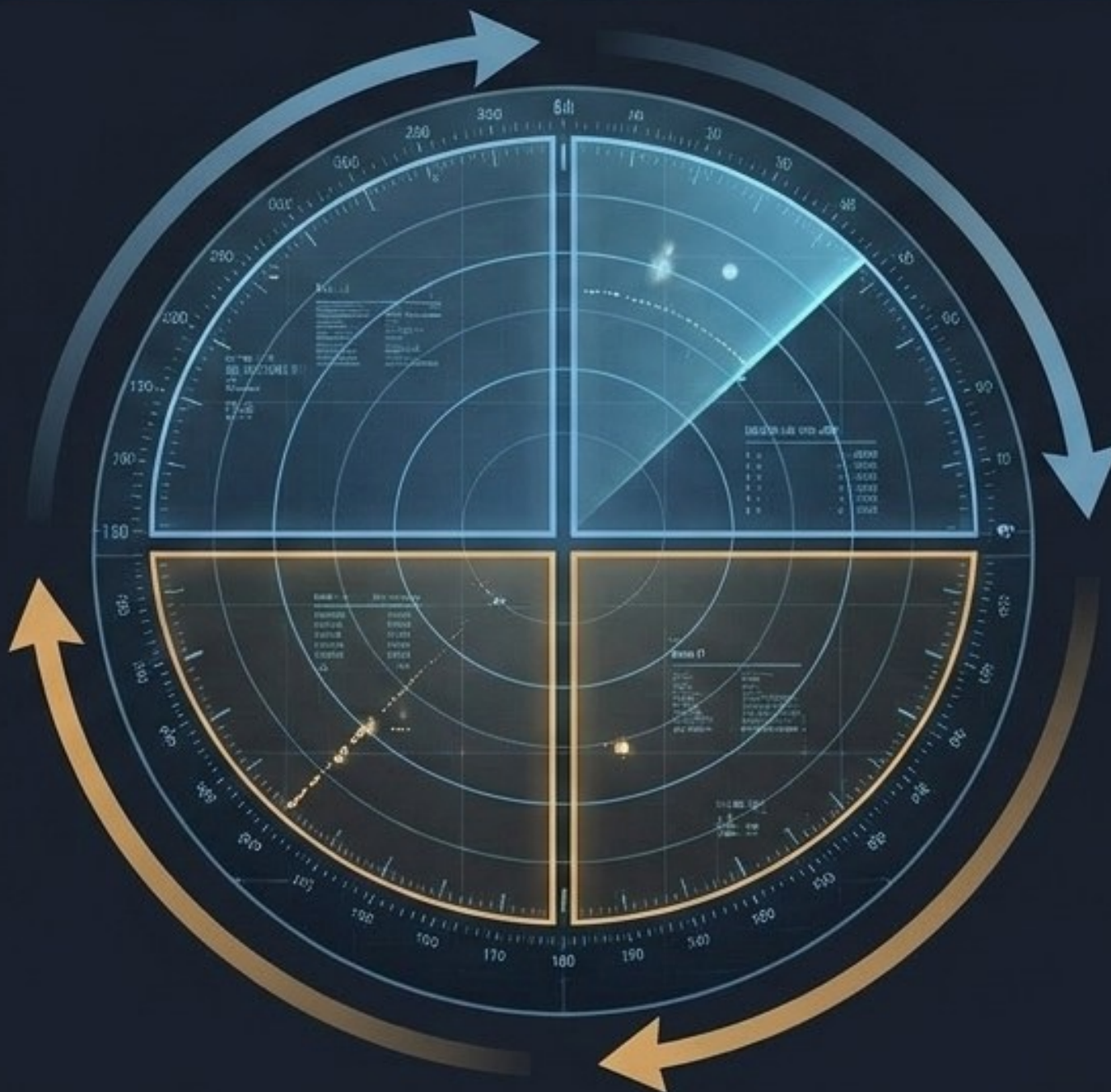
當攻擊者能夠自動化生成無瑕疵的社交工程攻擊時，傳統的「檢查拼字錯誤」防禦已完全失效。必須保護好 API Key（存放於 Vault 或環境變數中）避免系統被劫持。

階段二：模組 2 - AI 威脅建模雷達 (Threat Modeling Lifecycle)

核心目標：在攻擊發生前，以**敵手思維** (Adversary Mindset) 找出系統漏洞。

治理 (Govern)
建立長期監督機制，確保合規與政策落實。

管理 (Manage)
實施適當的安全控制措施以減輕風險。



對應 (Map)
盤點 AI 資產、識別資料流與 AI 系統在哪裡運作。

測量 (Measure)
量化潛在威脅的影響力與風險等級。



模組 3 - AI 存取控制金庫 (The Access Control Vault)

最外層：身分與存取管理 (IAM)

RBAC (角色基礎)：嚴格區分資料工程師、資料科學家與終端使用者的權限。

ABAC (屬性基礎)：根據環境、時間或網路狀態動態調整存取權。

核心層：資料安全防護

靜態與傳輸加密 (Encryption)：確保資料在儲存與跨網路移動時的安全。

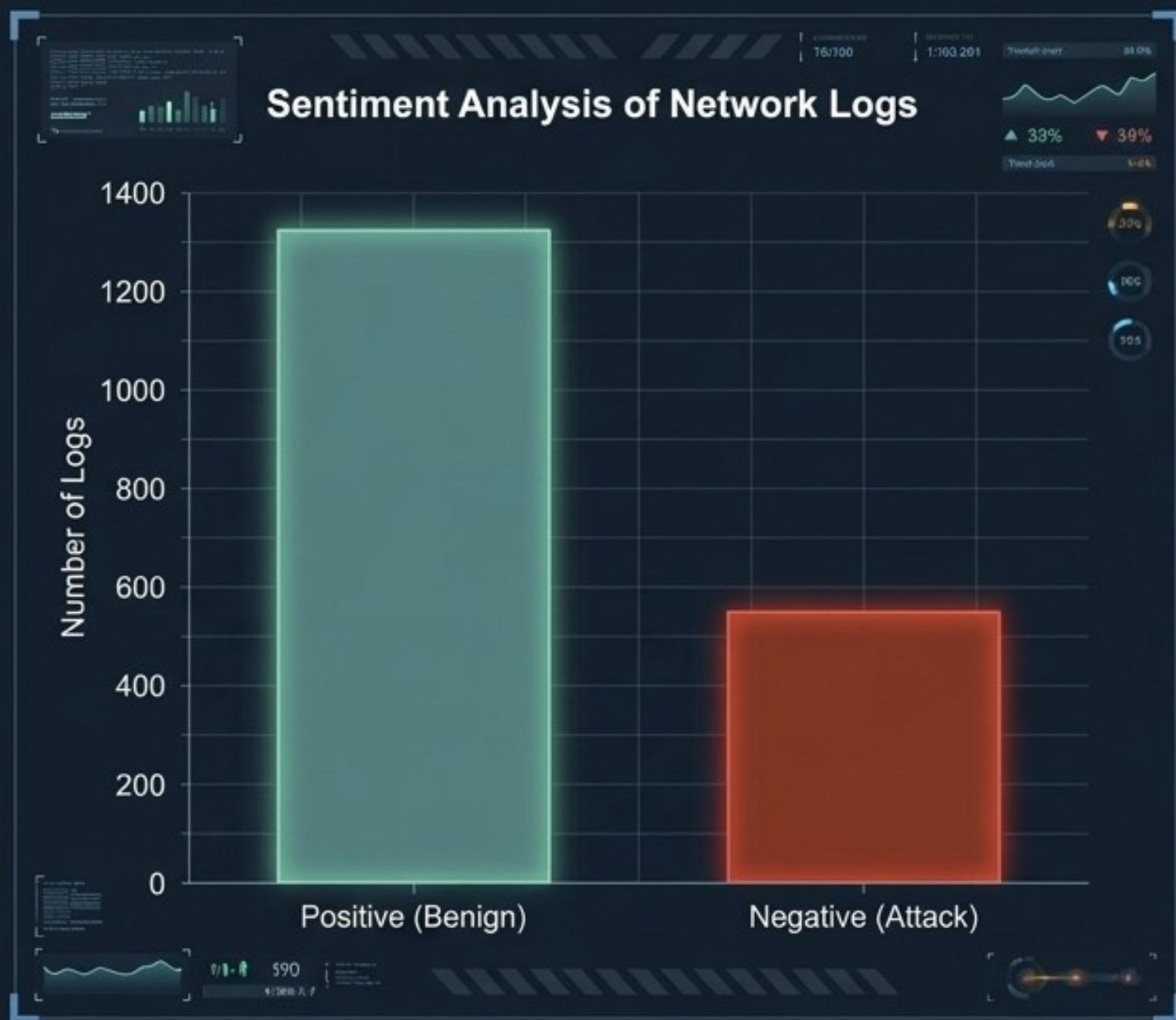
資料遮罩與分類 (Data Masking & Classification)：符合法規，隱藏敏感個資 (PII)。

中間層：API 與網路控制

管控哪些代理程式 (Agents)、網路和 API 被允許與 AI 系統通訊。



模組 3 - 持續監控與稽核 (Continuous Auditing)



為何需要持續監控？

AI 系統是不斷演進的，必須確保其行為不會隨著時間或新資料而產生偏差。

監控重點：

- **正向控制與日誌記錄**：追蹤誰在何時使用了哪些資料與模型。
- **異常行為偵測**：識別異常的 API 呼叫頻率或非預期的模型輸出。
- **合規性驗證**：透過定期稽核確保資料處理符合監管法規與標籤分類標準。

階段三：模組 4 - AI 全生命週期威脅分析



模組 4 - 毒水井效應與補償控制 (The Poisoned Well)



致命的連鎖反應：

如果 AI 訓練資料集因疏忽或攻擊遭到污染（毒水井），其後續的所有分析洞察、自動化決策與安全控制都將變得不可信。

補償控制的必要性：

當主要防禦失效時，必須依賴深度防禦的補償控制來降低風險。

淨水機制 (防禦手段)：

- 資料清洗與驗證 (Data Cleansing & Verification)
- 血統追蹤 (Lineage Tracking)：確保資料來源的純淨度。
- 完整性保證與平衡 (Integrity Assurance & Balancing)

模組 4 – 部署期的動態護欄 (Dynamic Safeguards)

AI 防禦思維的轉換：從傳統的靜態防禦，轉向具備適應性、感知上下文的動態安全護欄。

提示防火牆 (Prompt Firewalls)：
即時分析並攔截惡意的注入攻擊指令。

輸出驗證 (Output Validation)：
模型生成結果後，在傳遞給使用者或系統前進行安全與合規性檢查。

持續性攻擊模擬：
定期進行紅隊演練，以迭代方式強化模型韌性 (Model Hardening)。

模組 5 - 雙面刃：AI 攻防動態矩陣 (The Double-Edged Sword)

威脅行為者 (Threat Actors) - AI 武器化

網路釣魚：

生成無語法錯誤、高度針對性的社交工程誘餌。



惡意軟體：

生成具備獨特特徵碼、能規避偵測的多型態程式碼 (Polymorphic code)。



自動化攻擊：

大規模、高效率的欺騙內容生成與漏洞掃描。



資安專家 (Security Pros) - AI 輔助防禦

威脅偵測：

分析龐大非結構化數據，比人類更快發現行為異常。



事件回應：

自動化分析安全事件與日誌，加速事故處理速度。



弱點測試：

具備智慧自動化的滲透測試與紅隊演練。



階段四：模組 6 - 傳統 GRC vs. AI GRC

傳統 GRC (痛點)：



- 高度依賴人工流程。
- 靜態框架與定期評估。
- 對新興威脅反應遲緩。

AI 驅動的 GRC (優勢)：



利用先進演算法與機器學習。

- 動態風險評估：實時數據分析。
- 自動化合規監控：主動識別並減輕風險。
- 負責任的 AI 原則：確保透明度、問責制、公平性與隱私尊重。

模組 6 - AI 治理成熟度路線圖 (The GRC Roadmap)

階段二：標準與監控 (Standardization)

- 建立模型註冊表 (Model Registry)。
- 導入業界標準並建構初步監控能力。

階段一：基礎建立 (Foundation)

- 發布簡明政策。
- 建立現有 AI 使用案例的資產清單。

階段三：進階管理 (Advanced Management)

- 深化風險管理實務。
- 實現證據收集流程的自動化。

階段四：獨立驗證 (Validation)

- 針對高風險 AI 系統，尋求第三方 (3rd Party) 獨立驗證與稽核。

模組 6 – AI 指揮體系與關鍵角色

AI 卓越中心 (AI Center of Excellence, CoE)

組織內的中央決策與支援團隊。提供最佳實務、資源配置與 AI 專案的戰略指導。

關鍵技術角色分離 (職責分權)：

資料工程師 (Data Engineer)

- 負責構建與維護資料管線，獲取、清理、標記資料，並執行品質控制。

資料科學家 (Data Scientist)

- 探索資料、選擇特徵、開發基線模型，並證明 AI 解決方案的商業價值。

透過明確的角色劃分，確保創新協作的同時降低部署風險。

核心總結：SecAI+ 聯防生態系統

模組 2, 3, 4：威脅建模、
存取與補償控制 (裝甲)

模組 6：GRC 治理框架
(力場防護)

無縫整合的資安生命週期：
SecAI+ 的六大模組並非
孤立的章節，而是一個持
續運作的閉環系統。

模組 1：
資料與模型
(引擎)

資料與模型位於核心，受
到層層存取與補償控制的
保護，透過 AI 實務操作
進行主動防禦，並完全包
覆在 GRC 的治理框架之
中。任何一層的缺失，都
可能導致系統的全面崩
潰。

模組 5：AI 資安實務操作
(主動防禦)

模組 6：GRC 治理框架
(力場防護)