



## 責任 AI 政策

### 一、目的

本公司為因應人工智慧 ( Artificial Intelligence, AI ) 技術快速發展及其對營運、產品與社會之影響，建立涵蓋 AI 導入、開發、採購、使用、監督與退場之全生命週期治理原則與管理要求。

本政策旨在：

- 確保 AI 應用在促進創新同時有效控管風險
- 系統性管理資安、隱私、法遵及營運風險
- 落實公平性、透明性、可解釋性與倫理原則
- 強化企業治理與永續發展責任

### 二、適用範圍

本政策適用於本公司及全球各子公司，涵蓋所有 AI 相關活動，包括但不限於：

- 自行開發之 AI 系統
- 外購或第三方 AI 服務
- 生成式 AI ( Generative AI ) 與代理式 AI ( Agentic AI )

適用範圍涵蓋 AI 系統之資料處理、模型開發、測試、部署、營運與退場之全生命週期。

### 三、核心治理原則

本公司依據「風險導向 ( Risk-based )」與「以人為本 ( Human-centric )」原則，建立以下 AI 治理基準：

#### 1. 公司在 AI 使用與/或開發過程中，尊重數據隱私

於 AI 全生命週期中，應確保資料之合法性、正當性與最小化使用原則，並落實資料分類分級、存取控制與保護機制，以保障個人資料及敏感資訊之安全，並遵守適用之法規與公司內部資料治理規範，且僅限於核准之目的與範圍內使用。

#### 2. 在 AI 使用與/或開發過程中，保護系統的網路安全和資訊安全

AI 系統及其基礎架構應納入企業資安治理範圍，採取風險導向之技術與管理措施，包括：

- 身分與存取控管
- 資料加密與保護
- 弱點管理與持續監控
- AI 特有風險 ( 如 Prompt Injection、模型濫用 ) 防護
- 防止異常行為
- 避免不當內容輸出

以確保 AI 系統之安全性、完整性與可靠性。

#### 3. 在 AI 使用與/或開發過程中，避免潛在偏見

AI 系統應避免因資料偏差或模型設計造成不公平或歧視性結果，並透過測試、監控與持續改進機制，確保其決策具合理性與可信度。

#### 4. 保留人為決策機制

針對可能對個人、客戶或營運產生重大影響之決策，應導入「Human-in-the-Loop」機制，確保：

- AI 決策具備人員審核
- 最終決策由人員負責



以降低高風險自動化決策所帶來之不確定性。

#### 5. 確保 AI 系統的透明度，以及 AI 生成結果/決策的可解釋性

AI 系統應於合理範圍內揭露其用途、運作方式與限制，並提供必要之解釋機制，使使用者及利害關係人得以理解相關決策依據，並確保使用者得以辨識其互動對象為 AI 系統。

#### 6. 針對 AI 模型/工具產出的結果，建立明確的管理機制

本公司應明確定義 AI 治理相關角色與責任，包括治理單位、系統負責人及相關業務單位之職責分工，並建立以下管理機制：

- 設立或指定 AI 治理與風險控管之專責機制或單位
- 明確定義系統負責人 ( Owner ) 與使用責任
- 建立事件通報、調查與改善流程

確保 AI 應用之可追溯、責任歸屬與可持續改善。

#### 7. 針對 AI 的功能權限，定義明確邊界

AI 系統應依核准用途與範圍使用，不得逾越設計目的或未經授權之延伸應用。

公司應建立：

- AI 使用邊界管控
- 高風險使用情境審查機制
- 例外核准流程

以避免誤用與合規風險。

#### 8. 推動永續發展

AI 發展與應用應納入環境與資源效率考量，並將永續發展納入 AI 治理與風險管理機制中，於合理範圍內優先採用具能源效率之架構與技術，以降低運算資源消耗與碳排放。

#### 9. 禁止使用歐盟 AI 法案所規範之禁制型 AI 系統

本公司 AI 應用應符合各國相關法令及國際規範，並明確禁止使用：

- 操縱人類行為之 AI
- 利用弱勢族群或漏洞之 AI
- 社會評分系統
- 未經合法授權之生物辨識監控

並遵循《歐盟人工智慧法案 ( EU AI Act ) 》之禁制型 AI 規範，以確保不侵害基本人權與社會公平。

董事長 焦佑鈞

2026年7月