

群聯電子股份有限公司 人工智慧治理政策

前言

群聯電子股份有限公司（以下簡稱「本公司」）致力於推動負責任且可信賴之人工智慧，於人工智慧生命週期中遵循相關法規與國際治理原則，以確保其合規性與可信度。

適用範圍： 本政策適用於本公司各內部單位，現階段暫不延伸至供應商及合作夥伴；未來如業務需要（如委外開發、雲端 AI 服務採購），將於本政策改版時另行納入相關規範與供應商承諾條款。

一、承諾

1. 確保人工智慧發展符合法規與公司政策要求。
2. 確保人工智慧發展與公司永續目標及資安政策一致。
3. 強化員工人工智慧素養與責任意識。
4. 確保人工智慧應用尊重資料隱私。
5. 確保人工智慧系統開發生命週期符合資通安全。
6. 避免人工智慧應用存在潛在偏見。
7. 於人工智慧應用保留人為介入機制。
8. 確保人工智慧應用之透明度與可解釋性，並對 AI 產出內容及 AI 驅動決策予以明確標註。
9. 建立並落實人工智慧治理與問責機制。
10. 明確界定人工智慧之可執行及不可執行界線，並限制敏感人工智慧功能（如生物辨識、內部監控）之存取權限。
11. 兼顧低碳與永續發展，積極評估並推動人工智慧應用之能源效率與生態足跡管理。
12. 禁止使用或部署具操縱行為、剝削弱勢、社會評分或未經授權生物辨識監控之人工智慧應用。

二、治理架構

1. 本公司成立**人工智慧治理委員會**（以下簡稱「AI 治理委員會」），獨立於現有資安 / 風險管理委員會之外運作，由高階主管領導，成員涵蓋法務、資安、人力資源、業務代表等跨部門代表。
2. AI 治理委員會職責包括：審議本政策及其修訂、監督人工智慧治理架構之落實、審核高風險人工智慧應用場景之人為介入機制設計、定期審查人工智慧倫理與公平性管理成效，並將治理成效定期陳報高階管理層（或董事會）。
3. AI 治理委員會之運作應與公司永續發展、資料治理及資通安全政策整合，避免治理疊床架屋，並定期檢討與持續精進。

三、行動

（一）法規遵循

人工智慧發展持續對齊政府法規、法令及相關指引，並定期檢視其符合性與執行情形。

（二）教育訓練

規劃並推動人工智慧教育訓練與宣導機制：全體員工應接受基礎人工智慧倫理與風險意識教育；直接使用或操作人工智慧工具之單位（如人力資源招募團隊、資安 / 法遵團隊、研發、業務等）應接受加強訓練，強化風險意識、倫理認知及法遵能力。

（三）資料治理與隱私保護

導入並落實人工智慧資料治理與隱私保護機制，依據個人資料保護相關法規及本公司隱私權保護政策，於人工智慧之資料蒐集、訓練、開發、部署與維運各階段，遵循個資保護、目的限制、資料最小化、保存刪除及第三方資料合法使用原則。

（四）資通安全

於人工智慧系統開發生命週期落實安全設計、權限控管、加密、弱點管理、滲透測試、提示注入攻擊防護及事件應變等機制，以確保系統與資料之安全性、完整性與可用性。

（五）倫理與公平性

建立並執行人工智慧應用倫理與公平性管理機制，參考國際人工智慧治理原則（如 OECD AI Principles），採取偏誤識別、公平性測試、代表性資料檢核（稽核）及持續監測機制，避免人工智慧對特定族群造成不公平、歧視或偏頗之結果；招募／績效評估篩選類人工智慧應用應列為公平性稽核之優先對象，至少每年檢視一次。

（六）人機迴圈（Human-in-the-Loop）

確保人工智慧應用具備人機迴圈機制，對重大權益、高風險或不可逆之人工智慧決策，保留人工審核、介入、覆核、停止或覆寫機制，不得完全自動化執行。本公司現階段明訂以下場景須有人為審核並保留最終決定權，不得由 AI 自動執行最終結果：

- 人力資源招募、面試評分或績效評估篩選
- 門禁／生物辨識存取控制與內部監控類應用
- 資安或法遵事件之自動應變決策

（七）透明度與可解釋性

於規劃、導入及運用人工智慧系統時，以可理解方式揭露人工智慧之用途、能力、限制、風險與決策依據，並於適當情境明示使用者其正與人工智慧系統互動而非真人。對外或對內提供之 AI 產出內容，以及以 AI 輔助／驅動之決策結果（如招募篩選結果），應明確標註其為人工智慧產出或 AI 輔助決策。

（八）管理與問責機制

明確建立人工智慧之管理與責任機制：

- **AI 業務負責人：**由導入／使用人工智慧應用之各業務單位主管或指派代表擔任，負責該應用之業務合理性與影響評估。
- **模型負責人：**由技術系統整合處擔任，負責模型技術安全審查、效能與弱點管理。
- **資料負責人：**由資料擁有部門擔任，負責資料來源合法性、品質與隱私合規。
- **風險審查單位：**由法務處（或風險管理單位）擔任，負責法規遵循與風險審查。
- **資安審查單位：**由安全暨資源整合部擔任，負責資安風險審查。

並建立人工智慧事件通報、調查、矯正、責任追溯機制；為受人工智慧決策影響之同仁或第三方（如求職者），建立申訴／複核管道（得整合現有員工申訴或舉發機制），以識別、評估並降低相關風險。

（九）使用規範與防護機制

制定人工智慧使用規範及防護機制，落實人工智慧系統開發生命週期各階段管理，明確規範授權用途、能力邊界、禁止用途及例外升級審查機制；限制敏感人工智慧功能（如生物辨識、

行為監控)之存取權限僅限經授權人員;並導入風險控管與文件化管理,以確保可追溯性、合規性及防止功能濫用、擴張或非預期用途。

(十) 低碳與永續

於採用自有或第三方人工智慧模型、算力資源及資料中心時,評估其能源效率、碳排強度、水資源使用、運算優化及再生能源使用情形,優先選擇生態足跡較低之方案;引用公司現有 ESG / 永續報告書之碳排與能耗數據作為基準,並將人工智慧相關措施對永續發展成果之量化影響(如節能效益、算力效率提升)納入既有永續 KPI 揭露。

(十一) 禁止不當應用

建立檢核機制,禁止使用或部署具有操縱行為、剝削弱勢、社會評分或未經授權生物辨識監控之人工智慧應用。門禁 / 差勤等經授權之生物辨識用途得依現行規範繼續使用(並應告知員工使用目的與範圍),但禁止任何未經授權、隱密進行之員工行為監控、情緒 / 行為評分或社會評分類應用。

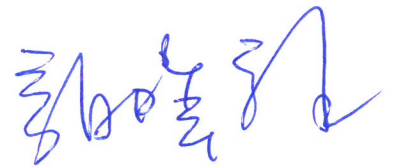
(十二) 維運監控

建立並執行人工智慧系統維運監控機制,持續監測模型效能、資料漂移、偏誤與安全事件;未來如政策範圍延伸至供應商,應強化供應商風險評估與合規審查。

四、政策生效與審查

本政策應每年或發生重大變更時,審查其適切性,並進行必要之修訂。本政策經 AI 治理委員會審議後,陳報核定後實施,修訂時亦同。

簽署人：_____



董事長

顏璋駁

2026 年 7 月