

樺漢科技股份有限公司

AI 人工智慧政策

第一條 制定目的

樺漢科技股份有限公司（以下簡稱「本公司」）致力於以創新科技推動產業升級、營運效率提升及永續價值創造。面對人工智慧（Artificial Intelligence, AI）技術之快速發展與廣泛應用，本公司秉持以人為本、誠信經營、風險治理、資訊安全、隱私保護、透明問責與永續發展原則，建立負責任人工智慧之治理框架，以確保 AI 之研發、導入、部署、使用、維運與管理，符合適用法令、社會倫理、公司治理標準及利害關係人期待。

本政策旨在作為本公司負責任人工智慧之最高治理原則，藉以提升技術可信度、保障利害關係人權益、降低風險，並支持企業創新、營運韌性與永續發展目標。

第二條 適用範圍

本政策適用於樺漢科技股份有限公司及持股 50% 以上之子（孫）公司，並涵蓋下列情境：

- 一、自行研發、共同研發、訓練、調校、部署或維運之 AI 系統、模型與應用。
- 二、採購、導入、整合 or 使用第三方 AI 工具、平台、模型、雲端服務與資料處理服務。
- 三、員工於營運管理、產品研發、資料分析、知識管理、流程自動化、客戶服務、內部協作及決策支援等情境中對 AI 工具之使用。
- 四、本公司向客戶提供之產品、服務、解決方案或內部流程中涉及 AI 技術之活動。

第三條 國際框架與法令遵循

本政策參照下列國際框架與法規精神制定：

- 一、OECD AI Principles：遵循國際間推動可信賴 AI 的核心原則，包含包容性增長、永續發展與福祉、以人為本的價值觀與公平、透明度與可解釋性、穩健性、安全性與保障，以及問責制。
- 二、UNESCO《Recommendation on the Ethics of Artificial Intelligence》：依循聯合國教科文組織所發佈之全球首個 AI 倫理建議書，保障人權、基本自由與尊嚴，促進環境與生態系統

之永續。

- 三、歐盟《Artificial Intelligence Act》（歐盟人工智慧法案）：參照其全球首部全面性 AI 監管法律之風險導向治理、透明揭露、人類監督、AI 素養與禁止用途之規範精神。
- 四、中華民國《人工智慧基本法》：遵循其所揭示之永續發展與福祉、人類自主、隱私保護與資料治理、資安與安全、透明與可解釋、公平與不歧視及問責等治理原則，以及高風險應用、風險分類、標示警語、責任歸屬與救濟機制之立法方向。

本公司於 AI 之研發、導入、部署、使用及管理過程中，除遵循本政策外，並應遵守中華民國及營運所在地之相關法令，包括但不限於《人工智慧基本法》、《個人資料保護法》及其他與智慧財產、營業秘密、勞動權益、消費者保護、資訊揭露、產業監理及各目的事業主管機關 AI 應用管理規範相關之法律及命令。

第四條 治理原則

一、 永續發展與以人為本

本公司推動 AI 應用，應兼顧營運效益、社會公益、環境永續、數位平權及人權保障，以支持人類自主、維護人格權及尊重民主法治與文化價值為前提，不得以 AI 取代必要之人類判斷、責任承擔與權利保障。

二、 Risk-Based 風險分級與高風險控管

本公司對 AI 應用採取風險導向管理原則，依其用途、資料類型、資料敏感度、影響對象、影響程度、可逆性、法令遵循要求及潛在損害，進行風險辨識、分級、審查與控管。

本政策所稱高風險 AI 應用，係指人工智慧之研發、導入、部署或使用，如其目的、資料類型、適用場景或輸出結果，可能對個人之生命、身體、健康、安全、隱私、基本權利、就業機會、財產利益，或對本公司重大營運、法令遵循、資訊安全及利害關係人權益造成重大影響者。

對高風險 AI 應用，本公司應採取較嚴格之事前評估、權責核准、人類監督、警示標示、持續監控、事件通報、例外升級及必要之限制或停用措施。

三、 資料隱私與資料治理

本公司於 AI 系統之研發、訓練、測試、部署及使用過程中，應妥善保護個人資料、客戶資料、員工資料、供應商資料、商業機密及其他敏感資訊，遵循合法、正當、必要、特定目的及資料最小化原則。

未經授權，不得將機密、敏感、受限制或依法不得外傳之資料輸入不適當之外部 AI 工具、平台或服務。涉及個人資料之 AI 應用，應依適用法令建立蒐集、處理、利用、保存、刪除及跨境傳輸之控管機制。

四、 維護資安與系統安全

本公司應以既有資訊安全管理制度為基礎，對 AI 系統及其相關資料、模型、介面、運作環境與第三方服務採取適當之安全控管措施，包括但不限於身分驗證、權限分級、資料保護、日誌留存、弱點管理、事件通報、備援、復原及供應鏈安全要求，以降低未授權存取、資料外洩、模型濫用、對抗攻擊、提示注入及其他資安風險。

五、 避免偏誤與不公平使用

本公司要求 AI 系統於設計、資料蒐集、訓練、測試、驗證、部署及使用過程中，應審慎檢信資料來源、資料品質、模型邏輯及應用情境，避免因資料偏差、設計缺陷或使用方式不當，導致對特定個人或群體產生不公平、不一致或歧視性結果。

對涉及員工、客戶、供應商或其他利害關係人權益之 AI 應用，本公司應加強偏誤風險辨識與檢視，並視需要採取人工覆核、調整、限制使用或停止使用等措施。

六、 重要決策須保留人工判斷與介入

本公司不允許 AI 在缺乏適當人為監督之情況下，單獨作成不可逆、重大、高風險或可能影響個人權益之決策。

凡涉及法令遵循、個資隱私、資訊安全、勞動權益、客戶權益、重大營運風險或其他高影響事項之 AI 應用，均應設置明確之人工審查、核准、介入、覆核、駁回、升級或停止機制。

七、 透明度、可解釋性與適當標示

本公司於適當情況下，應揭露 AI 之使用事實、用途、主要能力、限制、風險及責任分工；當使用者係與 AI 系統互動，或內容係由 AI 生成或輔助生成時，應依情境採取適當之資訊揭露、說明或標記。

對高影響用途，本公司應盡力提供足以支持理解、查核與覆核之說明基礎，不得將 AI 輸出視為不經驗證之最終結論。

八、 責任歸屬與可追溯性

本公司要求各 AI 應用均應有明確之責任單位及權責分工，包括業務/產品責任人、資料管理人、系統管理人及風險、法遵或資安審查責任。

本公司應視需要保留必要之文件、決策紀錄、審查紀錄、模型版本與異常處理紀錄，以支持責任歸屬、內部管理、法遵檢核、調查與改善。若 AI 應用造成錯誤、偏差、爭議、申訴、資安事件或其他不利影響，應依內部機制進行通報、調查、矯正及改善。

九、 明確界定 AI 可為與不可為之邊界

本公司應就 AI 使用建立明確授權範圍、用途界線、使用規範、能力限制與禁止事項。未經適當審查與核准，不得將 AI 用於超出原設計目的、超出授權範圍、欠缺法令依據或欠缺相稱控管之用途。

員工不得以 AI 從事偽造、誤導、侵權、洩密、操縱、規避內部控制、違反法令或其他不當行為。

十、 低生態足跡與環境永續

本公司關注 AI 對能源使用、運算資源、水資源及碳排放之潛在影響，並鼓勵於 AI 研發、部署、採購及使用過程中，納入能源效率、運算最佳化、設備資源使用效率、基礎設施永續性及供應商永續表現等考量。

在可行條件下，本公司優先考量具較佳能效、較低資源浪費與較高永續表現之模型、平台、資料中心或供應商方案。

十一、 禁止不當或高風險濫用

本公司不開發、不部署、亦不支持下列 AI 用途：

- (一) 以操縱、欺瞞、誤導或刻意扭曲方式實質影響個人行為之 AI 系統。
- (二) 利用年齡、身心障礙或社會經濟脆弱性進行不當影響、剝削或差別待遇之 AI 系統。
- (三) 以社會評分方式對個人進行不正當分類、排序或不合理差別待遇之 AI 系統。
- (四) 未經合法授權之即時遠端生物特徵辨識、監控或侵害隱私之 AI 應用。
- (五) 其他違反法令、侵害人權、侵害隱私、危害安全、資訊誤導、造假或顯失公平之 AI 應用。

十二、 AI 素養、教育訓練與能力建構

本公司應依不同角色、風險情境及應用範圍，持續推動 AI 倫理、資料治理、隱私保護、資訊安全、人機協作、輸出辨識與法遵要求等教育訓練與宣導，以提升員工及相關人員對 AI 之知識、風險意識與正確使用能力。

第五條 治理架構、責任單位與跨部門職掌分工

本公司建立分工明確之 AI 治理架構，由下列角色與對應之權責單位依職掌辦理，以落實跨部門協作與政策推動：

治理層級 / 角色	指定權責單位	具體管理職掌
一、最高監督層級	董事會	1. 負責核准本政策及其後續修訂。 2. 聽取管理階層之 AI 治理與重大風險管理進展報告。 3. 監督並指導公司重大 AI 技術之應用方向。
二、管理督導層級	總經理室 / 永續發展暨提名委員會 (ESG)	1. 負責推動本政策之落實，督導相關制度、資源配置與跨部門協作。 2. 定期檢視 AI Risk 與法遵稽核進度，向董事會報告重大進展與異常事件。
三、功能審查與支持單位	資訊管理處 (IT) / 資訊安全部 (Sec)	1. 負責 AI 系統之基礎設施、部署環境之技術審查。 2. 建立 AI 資安防禦機制，進行弱點管理、系統備援與日誌留存。
	法務與合規室 (Legal & Compliance)	1. 負責各類 AI 導入與應用之合規性檢視。 2. 協助研擬與供應商之合約與數據授權條款。
	個資保護小組 (DPO)	1. 負責審查涉及個人資料處理之 AI 流程。 2. 確保資料治理符合最小化原則與跨境傳輸規定。
	人力資源單位 (HR)	1. 辦理員工 AI 倫理與素養之教育訓練。 2. 確保 AI 在 HR 應用具備公平性與人為監督。
四、執行與使用單位	各業務、產品研發、專案與營運單位	1. 負責 AI 工具之日常風險控管、監控與記錄保存。 2. 發生異常時應立即人工介入並進行通報。

第六條 申訴、事件通報與改善

本公司就 AI 使用所衍生之疑慮、申訴、事故、偏差、資訊安全事件、隱私事件或其他不利影響，應建立適當之通報、調查、處理、矯正與改善機制；必要時，應採取暫停使用、限制使用、下架、修正、通知利害關係人或其他相應措施。

利害關係人如認為本公司 AI 應用或相關資料處理行為對其權益造成影響，得透過公司官網公開揭露之申訴與聯繫管道提出申訴或反映意見。本公司將依案件性質進行受理、調查、回應、矯正及追蹤改善。

官網申訴管道：<https://www.ennconn.com/tw/consumer-protection-and-privacy/>

第七條 政策檢討與修訂

本公司將依技術演進、法令更新、國際標準、利害關係人期待、營運需求與風險情勢，定期檢討本政策之適切性，並視需要進行修訂，以持續精進負責任人工智慧治理機制。

第八條 施行

本政策經董事會核准後施行，修正時亦同。