

責任人工智慧行動方案

方案目的與適用範圍

為呼應本公司《責任人工智慧治理政策聲明》，並推動人工智慧（AI）於工程服務及企業營運領域之負責任應用，本公司提出本行動方案，作為推動 AI 治理、風險管理與永續實踐之行動框架。

本方案旨在促進 AI 應用符合透明、可信賴、安全、永續與以人為本之治理精神，並作為本公司完善 AI 治理機制、風險管理措施及永續發展實踐之參考依據。其適用範圍涵蓋本公司於 AI 技術開發、導入、採購、使用、維運及管理等相关活動，本公司將依據應用情境之風險程度、技術成熟度及營運需求，採取彈性且相稱之管理措施。

AI 倫理與安全訓練

本公司持續推動 AI 教育訓練與知識宣導，透過實體課程、專題講座、工作坊及中鼎大學（CTCI University）等多元學習管道，提升同仁對 AI 技術應用、倫理、安全、風險管理及負責任使用之認知與能力。

集團創研中心於 2025 年共辦理 12 堂 AI 相關實體課程，累計 724 人次參與；為進一步擴大教育訓練之觸及範圍並提升學習便利性，本公司亦持續將相關課程數位化並上架中鼎大學，提供同仁更具彈性的學習方式。截至本行動方案初版發布前，相關 AI 數位課程累計點閱已超過 3,000 人次。未來，本公司將持續依 AI 技術發展、法規趨勢及實務需求，更新及充實相關訓練內容，涵蓋 AI 素養、AI 工具與應用、AI 倫理與治理、風險管理、資訊安全、資料隱私及資料保護等主題。

相關訓練紀錄及統計資料，得作為本公司推動 AI 素養提升與成效追蹤之依據，並依永續資訊揭露規劃，納入永續報告書或其他對外揭露資料，以展現本公司推動負責任 AI 應用及人才培育之成果。

AI 敏感功能之存取限制與權限管控

本公司重視企業機密、個人資料、客戶資訊、專案資料及其他敏感資訊之保護。針對涉及敏感資料或重要業務流程之 AI 應用（如：人臉辨識、監控系統或生物辨識相關應用等）應依循本公司既有資訊安全及權限管理原則，實施嚴格之存取控管與使用限制。

相關 AI 功能之使用，應依資料屬性、應用目的、業務需求及風險程度，採取適當之權限區分、授權管理或其他保護措施，以降低資料外洩、未授權使用、誤用或其他潛在資安風險。隨著 AI 應用發展，本公司亦將持續檢視相關管理方式之適切性，以支持安全且可信賴之 AI 應用環境。

AI 模型公平性與偏誤評估

本公司重視 AI 應用可能衍生之偏見、差別影響及公平性風險。於導入、採購、使用、自行開發、訓練或微調 AI 模型與相關服務時，應依應用情境、資料特性、使用對象及可能影響範圍，將公平性、偏見風險及倫理影響納入整體考量。

針對採用外部 AI 模型（例如 Google、OpenAI 等公司所提供的 AI 模型）之應用場景，本公司得視實務需求，檢視供應商提供之技術文件、安全報告、風險說明或其他相關資料，並於合理可行範圍內評估其對不同群體可能造成之影響。若發現不當偏誤或公平性疑慮，本公司應採取更換供應商或限制使用等相應改善措施，以維持公平、包容與可信賴之 AI 應用。

針對自行開發、訓練或微調之模型，本公司應檢視資料來源、代表性及品質，並依模型之應用目的、影響程度及風險高低，進行必要之偏誤與公平性管理。相關模型每年辦理一次評估 (Assessment)，評估內容視實際需要得包括稽核 (Audit)、偏誤檢測 (Bias Testing) 及公平性評估 (Fairness Evaluation) 等項目，以確認模型不致對不同群體產生不合理之偏誤或差異。若發現模型輸出存在不合理偏誤或差異，應採取資料調整、模型修正、限制使用或其他適當之改善措施，以維持公平、包容與可信賴之 AI 應用。

AI 模型偏移修正機制

本公司關注 AI 模型或 AI 服務於實際應用過程中，可能因模型版本更新、資料偏移或外部應用環境變化，而導致輸出品質下降、適用性差異或異常結果等現象。針對重要之 AI 應用，本公司應建立適當之觀察、使用者回饋或異常反映機制，並以每年檢視一次為原則；惟得依應用風險、使用頻率、業務影響程度、模型或服務版本更新情形及實際管理需求，調整檢視頻率與檢視方式，以偵測模型偏移 (Model Drift) 或效能退化情形。

當發現 AI 應用之輸出品質、穩定性或適用性偏離預期時，相關單位得採取必要之檢視、調整提示工程 (Prompt Engineering)、使用提醒、限制使用、調整應用方式、更新服務設定、向技術供應商反映並協作修正，或採取其他適當措施，以確保 AI 應用之可靠性與準確度。本公司將隨 AI 技術發展與應用經驗累積，逐步強化相關觀察、回饋與改善能力。

AI 生成內容之透明性揭露與標註機制

本公司重視 AI 生成或輔助產出內容之透明性。對於本公司提供或導入之 AI 系統、AI 服務或 AI 輔助功能，應視其應用情境、服務對象及內容用途，採取明確之標示 (Labeling)、說明或揭露方式。相關標示及說明應使使用者或利害關係人能理解其內容是由 AI 生成、整理或分析建議，以便其認知相關內容之性質與限制，並做出明智之判斷與選擇。

AI 生成或輔助內容之揭露方式，應依使用情境、服務對象、內容用途及可能影響程度採取相稱作法，以維持資訊傳遞之清楚性、專業性與可信賴性。

申訴程序與使用者權益保障

本公司重視 AI 應用可能對使用者、客戶、合作夥伴及其他利害關係人造成之影響。針對本公司提供或導入之 AI 服務，將依服務性質、使用對象及影響程度，指定適當之負責單位及處理窗口；並視實際情境，於相關網頁、系統介面或其他適當管道提供聯繫資訊，使相關人員或利害關係人得就 AI 應用結果、使用情形或相關疑慮提出意見或申訴 (Appeals)。

本公司 AI 應用相關意見或申訴，本公司員工可透過員工意見平台反映，或直接聯繫相關負責窗口；如對客戶、供應商或其他外部利害關係人產生影響，當事人可透過本公司對外聯絡信箱 (ctcigrp@ctci.com)、中鼎集團舉報網站，或洽相關業務負責窗口提出申訴。

中鼎集團舉報網站：<https://secure.conductwatch.com/ctci/?Pg=1&Lang=zh-TW>

相關案件由集團創研中心統籌受理，並依案件性質、影響程度及既有管理機制，交由適當權責單位審慎處理；必要時，得提送<資訊安全推動委員會>審議及決定。

必要時，得依公司內部規範、契約約定或相關法令要求進行後續檢視或改善，以維護利害關係人權益、公司信譽，並落實 AI 應用之問責機制與社會責任。

生態足跡倡議

本公司認同 AI 技術之發展與應用，應兼顧節能及生態足跡之影響，並於合理可行範圍內，提高能源使用效率，降低不必要之運算及能源消耗。

於採購或使用外部 AI 服務時，本公司將依實際應用需求評估其適用性，並將供應商在節能、生態足跡、淨零等永續面向的表現納入考量，以兼顧 AI 應用效益與環境永續。

如涉及自建或自主管理 AI 運算資源，本公司將於合理可行範圍內，優先考量採用能源使用效率較佳之硬體設備。目前本公司自有 AI 運算資源均設置於公司總部大樓內，其相關用電屬公司整體能源使用之一部分，並納入公司既有之能源管理及減碳管理範疇，如 IT 機房採用高效率空調。本公司已設定科學基礎減碳目標並通過 SBTi (Science Based Targets initiative) 審查，並導入 ISO 50001 能源管理系統，以持續推動能源使用效率及減碳管理。

於模型開發、訓練及應用過程中，亦將透過適當之模型選型、模型壓縮或量化、及演算法優化等方式，減少不必要之運算及能源消耗。

量化永續影響

本公司重視 AI 技術對企業營運效率、工程服務品質及永續發展目標之潛在貢獻。於 AI 應用推動過程中，將依應用場景、資料可取得性、技術成熟度及衡量方法之可行性，持續探索 AI 導入後對作業效率、資源使用、能源管理、風險降低及服務價值提升等面向之影響。

對於具備明確應用場景與可衡量基礎之 AI 應用，本公司得於合理可行範圍內，建立量化指標 (KPIs)，例如碳排放量、能源使用量、資源節約或風險降低等指標，以評估其導入後之實際效益。相關量化結果將作為後續 AI 策略優化之參考，並得視實務需求及揭露規劃，於永續報告書或其他適當之對外揭露文件中說明，以數據驅動方式展現 AI 應用對環境與社會之正面價值。