

議題結論報告

團隊名稱:駭未來 Hire Future

討論議題:是對手還是隊友？AI 世代的職場再設計

辦理時間:114 年8月2日(星期六)

辦理地點:台大新生教學館201號教室

活動緣起：

AI 的發展勢如破竹，各國政府已意識到這股浪潮既是機遇也是挑戰。為了在鼓勵創新的同時確保社會安全與倫理，全球主要經濟體紛紛開始制定相關法律與政策。AI 的發展已無法被「阻止」，因為它所帶來的效率提升與資源重分配，已讓企業與政府難以抗拒這股技術紅利。然而，正因為這波浪潮來得如此迅速，各國更積極透過立法與制度設計來「治理」這股力量，試圖在創新與風險之間找到平衡。

2024 年，歐盟率先通過全球第一部《AI 法案》(EU AI Act)，建立起以「風險分級」為核心的監管架構，針對醫療、教育、招聘等高風險領域進行嚴格規範，並明確禁止部分 AI 應用，如社會信用評分與公共場所的人臉辨識。這代表歐洲已不再只是監管科技產品，而是將 AI 視為影響社會正義與民主制度的潛在力量。

美國則以行政命令與原則性法案作為回應。拜登政府於 2023 年頒布的《AI 行政命令》要求聯邦機構強化對 AI 使用的審查與透明度，並同步推動《AI 權利法案藍圖》，提出包括演算法不歧視、知情選擇權、保有人類決策權等五大原則，強調在創新過程中仍須守住基本人權。

台灣目前尚無完整 AI 專法，但已開始修訂個資法並草擬「生成式 AI 原則」，針對模型透明度、責任歸屬與風險揭露進行討論，顯示我們也正面對同樣的制度挑戰與倫理思辨。

從這些趨勢可以看出，AI 發展已不是要不要用的問題，而是如何規範、如何共存、如何問責的公共議題。立法正在追趕技術，治理正在追趕應用，社會也正在追趕對話。人類雖然無法擋下技術浪潮，但能夠決定這股力量的方向與邊界。

活動流程：

早上首先透過破冰卡牌，使參與者沉浸式體驗想像，在2030的世界中，AI究竟會如何與人類共事。反思自己對於AI的選擇偏好，是否反映了自己對於AI想像的界線。接著透過議題導讀，讓參與者了解「技術躍進週期縮短、組織的壽命、本次討論的 AI 應用範疇、對工作者的影響？、職場現況 & 挑戰」等相關內容。在早上充分吸收學習了AI知識後，下午則藉由願景工作坊的形式，將學員分成五組，模擬三個劇本中角色所面臨的困境，而後發想解方。議題導讀較偏向單向輸出的講座，願景工作坊則是小組成員互動多元觀點激盪，深入討論換位思考。

活動宗旨：

本活動為旨在鼓勵公民針對重大公共議題，提出多元觀點、未來想像、政策建議。期望透過願景工作坊與角色模擬，鼓勵不同世代、背景的參與者皆能在安全、包容的討論空間中，探索 AI 對未來職場培育與制度變革的影響，激盪出兼具公平、創新與人性關懷的未來職場設計。

活動背景設定在2030年，台灣已全面部署具備 AGI 能力的「智慧協作代理人 (AI Co-Agent)」，廣泛應用在科技、醫療、服務等各行各業。這些 AGI 能夠獨立決策、指派任務，甚至進行跨領域知識轉換與提案。政府已提出《人機共事基本法》，保障「人類工作者的再價值化與再協作權」，但社會對未來工作型態仍充滿焦慮與討論。

本活動預期實現三個目標：

- 一、理解工作者在 AI 競合下的機會與困境
- 二、探討未來職場中的人機協作倫理議題
- 三、共創可行性高的未來制度草案

一、現況或問題

議題現況：

AI進入職場對個人、組織、社會的影響？

AI快速發展與自動化技術的普及，全球勞動市場正經歷劇烈轉型。

根據麥肯錫《2023 未來工作》報告，2023至2027年間，約23%勞工將失去或轉換工作，整體職缺恐減少2%。

AI不僅重塑傳統職業樣貌，更對勞工適應力、職能培訓及產業政策帶來深遠挑戰，例如：遠距與彈性工時等新工作模式，使企業需要重新設計組織架構與人力資源策略。

同時，AI的運用也引發公平性、隱私與倫理風險，如何防止演算法偏誤與勞工權益受損，成為政策與法律制定的重要議題。

此外，政府、教育機構，該提供哪些資源及培訓，才能有效填補勞工技能落差，提升國家競爭力。

(一) 在2030當AGI 取代人類的某些能力，人如何重新定義工作價值？

第一組跟第二組拿到的劇本，從公司新進員工角度出發，藉由AI的幫助做出了比原先自己產出還好的工作成果，也因此順利通過了試用期。但也同時產生自我懷疑，對自己的專業是否還重要產生迷惘。劇本主要是在討論「人類工作價值在 AGI(通用人工智慧)普及後面臨的挑戰與不確定性」，許多過去被認為是人類獨有的能力，如複雜分析、創意發想和策略規劃，都正被機器以更高的效率和更低的成本所取代。而2030的AGI已經進化到有辦法比人類更出色地運用某些能力。根據劇本設定兩個小組有了以下四點關於這個議題會產生的困境：

1、侷限人類思考，致使大腦萎縮

決策、記憶、判斷、想像……這些曾經需要經年累月鍛鍊的大腦功能，正因為「不用而退化」。今天我們不再記住路線，只會開啟導航；不再思考答案，只會複製生成；不再等待靈感，而是輸入提示。人類的心智能力，逐漸被「外包」給機器。當我們習慣性地接受AI給出的答案，而不是去質疑、查證或從多個角度思考時，我們的大腦會失去辨別資訊真偽、建立邏輯鏈結的能力。

2、資料產出正確性

AGI 生成的內容高度可信，常讓人忽略其潛在的謬誤與偏見。由於AGI是基於現有的大數據進行訓練，它會自動學習並複製數據中的歷史偏見、錯誤資訊，甚至是人類社會中的不平等。例如，如果訓練資料中存在性別或種族歧視，AGI產出的內容也可能無意識地帶有這些偏見。當我們將這些經過AI美化過的資訊視為絕對正確時，不僅可能做出錯誤的判斷，更會進一步加劇社會中的既有問題。

3、對專業的不尊重

法律、醫療、設計、寫作……幾乎每一個過去需要長年累積經驗的專業，都開始被 AI 模型涉足，甚至做得更快、更便宜。當一個 AI 模型能設計建築、寫法律文件、擬定診療建議，人們開始質疑專業的意義。

當一個僅受過幾個月訓練的AI模型，就能在某些方面表現出媲美甚至超越數十年資深專家的能力，這會讓社會大眾對專業人士的辛勤付出、漫長養成與獨特經驗產生輕視感。人們可能因此認為，專業知識不再是稀缺資源，而是唾手可得的數據集合，專業人士也從不可或缺的「匠人」變成了隨時可被替換的「操作員」。

4、AI的自主性、進步速度，讓人產生自我懷疑

AGI 的進化速度遠超預期，更新快、版本多、功能日益精密。我們還來不及適應新的工具，下一輪淘汰已經開始。自我懷疑成為一種普遍狀態：不只是怕被取代，而是根本不確定自己的努力是否還有意義。更深層的焦慮來自於：當人不再能定義未來技術的方向，只能被迫追趕，它到底還是「工具」，還是正在主導我們生活與工作的他者？在這樣的環境中，人類該如何重建對自己的信任與信心？

(二) AI 與人為判斷出現衝突怎麼辦？

第三組及第四組拿到劇本二，主角身為公司的主管，在專案負責人上需要欽點人選，主角心中有兩個人選難以抉擇，最終在用人調度上相信AI根據數據分析後推薦的人選。然而在實際工作過程中，發現該人缺乏臨場反應的能力，懷疑自己當初是否應該選擇另一個人作為專案負責人。在2030年的時空背景之下，AI 模型快速給出結論、建議與預測，為人類社會的運作

待帶來更高速的效率，但也與人類產生了更多的衝突。在「人機共決」的時代，暴露出三個尚未解決的深層問題：

1、權責不明

當 AI 的建議與人類判斷相左時，若結果導致損失或災難，責任該由誰承擔？是採納建議的人？設計模型的開發者？還是根據不完整數據做出決策的 AI 本身？現行法律與制度對這類「人機共決」的責任歸屬仍處於模糊地帶。企業將錯推給演算法，個人推說只是依 AI 建議行事，AI 本身則「無法負責」。當判斷權愈來愈多交給機器，責任卻無人承接，我們正邁向一個道德與法律的灰色地帶。這不僅模糊了過失的界線，更可能侵蝕社會賴以維繫的信任基礎。

2、個資外洩

AI 所做的每一個判斷、每一次推薦，都建基於大量個人資料的收集與運算。然而，這些資料從哪裡來？被誰看見？在何時被用？我們很少知情，幾乎從未真正授權。在人與 AI 判斷出現衝突時，我們才驚覺：AI 憑什麼對我們「了如指掌」？而這份掌握是否早已超出我們能控制的範圍？人類的隱私、偏好與選擇，逐步被資料化、商品化，但我們卻無法真正決定它們的去向。

3、AI倫理

AI 不只處理事實，更在預測、推薦與建議的過程中，實質參與了「價值判斷」的層次。當 AI 建議一家公司裁撤某些員工、建議醫療系統優先照顧某些病患，它其實已經介入了「應不應該」的倫理抉擇。問題是，這些標準從哪裡來？是誰設定的？我們是否過早將選擇交給一個無感情、無責任的程式碼？當倫理與效率產生矛盾，人類是否還有最後說話的空間？

(三)如何設計一個公平的人機共事制度表讓資方、勞方雙贏

第五組拿到的劇本從資深人資長的角度出發，探討在與AI協作、AI同事、AI代理人的時代導入AI 協同人資系統，如何設計一個公平的人機共事制度表讓資方、勞方雙贏，以下是可能面臨到的三個問題：

1、人機合作的責任歸屬與績效分配

在一個專案中，AI扮演了關鍵角色。當專案成功時，該如何區分是個人的智慧，還是AI的運算能力所致？如果團隊成員過度依賴AI，且最終表現不佳，責任該歸咎於人，還是AI的侷限性？

這關係到獎金分配的公平性。一個能善用AI的員工，其產出可能遠高於傳統工作者，這是否意味著他應獲得更高的獎金？反之，如果一個員工的貢獻幾乎都來自AI，他是否還應該得到全額獎酬？我們需要思考如何將AI的貢獻納入考量，並建立一套公平的績效分配模型。

2、技能與薪酬制度的脫鉤

傳統的薪酬制度通常與專業技能的深度與廣度掛鉤。但在AI時代，許多專業技能，如數據分析、文案寫作，AI都能高效完成。這使得過去需要長年累積的專業知識，其「稀缺性」大幅下降。

如果專業知識不再是薪酬的核心依據，那麼我們該如何定價？未來的薪酬制度可能需要與「AI協作能力」、「批判性思考」、「情感智力」等軟實力掛鉤。企業可能需要投入更多資源來培訓員工這些新技能，並將其納入薪酬與晉升的考量。

3、創意的歸屬與智慧財產權

當一個設計師使用AI生成了新的Logo或一個作家利用AI創作了新的故事，最終的創意歸屬權該如何界定？是屬於設計師，還是屬於開發AI的公司，還是屬於AI本身？在獎勵制度中，「原創性」通常是高額獎金的來源。然而，在AI協作下，原創性的定義變得模糊。這不僅是法律層面的問題，更會影響公司內部對於創新獎勵的公平性與透明度。

二、結論或建議

(一) 重新檢討法規、修法

現行的勞動法、獎酬設計與責任歸屬機制，大多預設「人是唯一的創造者」，因此在面對AI深度參與工作流程時顯得捉襟見肘。

可能調整方向包括：

- 建立「AI 協作工作責任歸屬準則」，釐清 AI 建議與人類判斷之間的界線與責任邊界
- 修訂著作權法與智慧財產相關法規，界定 AI 生成內容的歸屬與利潤分配
- 規範企業在員工評估時，若使用 AI 協作工具，須揭露使用程度與原始貢獻內容

挑戰在於法律常落後於技術，若無跨部門協作與前瞻思維，制度很容易流於形式或被濫用。

(二) 提高公民 AI 素養

AI 工具普及不是終點，真正的關鍵在於人是否具備理解、判斷、運用與反思 AI 的能力。

可推動方式包括：

- 在中高等教育、企業訓練中納入「AI 素養教育」，不只是操作教學，更包含倫理、判斷與風險意識
- 鼓勵員工紀錄 AI 協作歷程（如 prompt、修正歷程、判斷依據），作為貢獻評估的新依據
- 培養「使用 AI 的反思力」，讓人們不只會問 AI，也能質疑 AI、補充 AI 所無法給的判斷

最終目的不是每個人都要變 AI 工程師，而是讓每個人都知道在 AI 介入的工作中，如何看見自己的角色與價值。

(三) AI 專業顧問公司健檢

當組織內部無法客觀評估 AI 導入後的個體貢獻差異時，引入外部協力成為一個中立選項。

可行模式：

- 由 AI 專業顧問公司或第三方機構定期進行「AI 導入後工作流程影響評估」
- 發展一套「人機共創貢獻指標 (Human-AI Collaboration Metrics)」，協助企業評估員工在 AI 協作中的獨特價值
- 協助企業識別「可被 AI 替代的重複流程」與「需保留人判斷的關鍵節點」，避免一刀切的效能導向錯誤

藉由外部單位的專業性與客觀性降低內部的爭議與偏見。

(四) 組織文化與個人心態重塑

當 AI 成為職場中不可或缺的協作夥伴，組織文化與個人心態也勢必要重新調整與重塑。企業不應只聚焦在技術導入，更需引導員工建立與 AI 共存的價值觀與角色認知。例如，從「AI 是否會取代我」的恐懼，轉向「我如何與 AI 共創價值」的合作思維。組織可以透過內部對話、工作坊與情境模擬，引導員工反思自身優勢與 AI 的界線，並鼓勵嘗試錯誤，降低「我不能輸給機器」的心理負擔。同時，領導者應以身作則，打造一種容許不確定、尊重人機協作的文化氛圍。唯有如此，人才才能在心態上完成「競爭對手 → 協作夥伴」的轉換，為組織帶來更長遠的適應力與創新力。