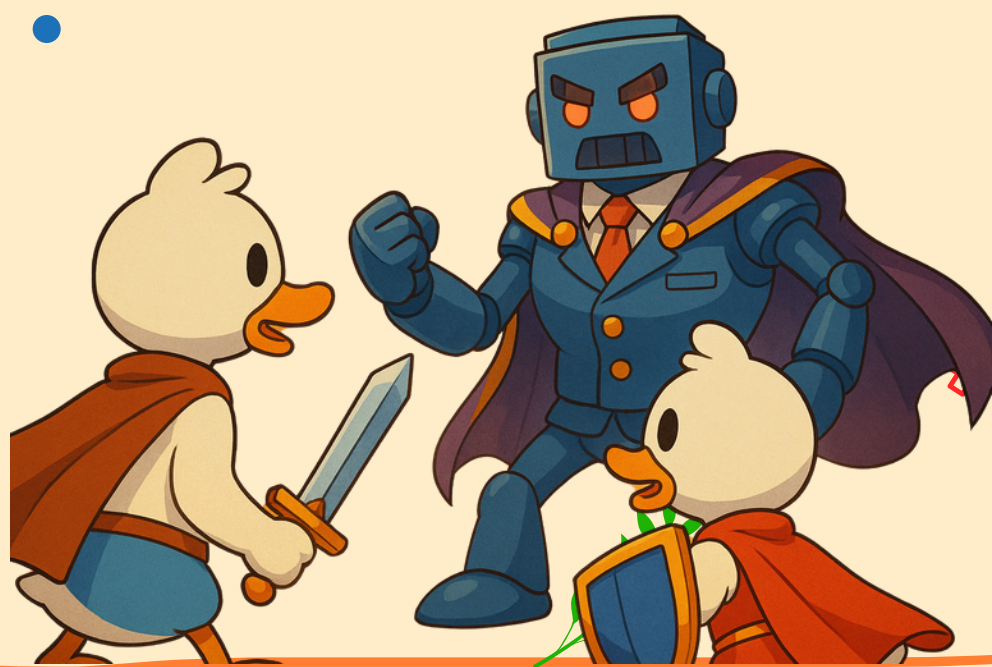


當虛擬主管襲來！ 上班誰說了算？

議題手冊



時間 | 2025年7月19日 09:30-16:40

地點 | 東吳大學城中校區2/23大會議室



教育部青年發展署
Youth Development Administration Ministry of Education





114年「青年好政-Let's Talk」計畫

目錄

一、計畫簡介	00
二、提案源起	01
三、團體介紹	04
四、講者介紹	05
五、活動流程	06
六、劇本	08
七、議題背景介紹	12
(一)機器管理人類？	12
(二)當AI工作時，是否有偏見？	17
(三)AI對工作自主性是否有影響？	22



一、計畫簡介

關於青年好政：

為鼓勵青年參與公共事務，提供青年與政府政策對話發聲管道，教育部青年發展署（下稱青年署）於2013年改制後，參考前身行政院青年輔導委員會的青年國是會議、青年政策大聯盟，以審議民主之精神辦理青年政策論壇。2015年青年政策論壇結合青年較常使用之新媒體（如Google網上論壇、臉書），建構跨越地理及時間限制的公共政策參與平臺。2017年則邀請縣市政府、大專校院及關注青年議題民間團體，共同辦理青年好政論壇，整合產、官、學資源，提供青年更多政策參與機會。

2018年起，青年署將青年好政論壇轉型為以青年為主體的「青年好政-Let's Talk」（下稱Talk）計畫，號召青年於全臺各地自主辦理公共議題政策審議討論；2019年Talk計畫則配合行政院政策，融入「開放政府」透明、參與、課責及涵容精神，讓青年在共聚、聆聽、交流過程中，培養思辨和公民參與行動力；2020年臺灣為加入全球「開放政府夥伴關係聯盟」（OGP），宣示推動「臺灣開放政府國家行動方案」，Talk計畫即被列為該行動方案有關「擴大公共參與機制」承諾事項之一。

二、提案緣起

議題現況及問題

AI已經在我們的日常生活中，扮演越來越重要的角色，我們的生活也藉由AI更顯得便利。但AI同時也已經進入到企業經營管理的各個層面，並被企業用來提升生產力與管理效率。這種變革對於勞動市場與勞動條件的影響，並非單純帶來正向發展。AI技術的應用不僅改變了工作型態，還塑造出新的勞動監控與管理模式，使得勞工的工作環境與權益面臨新挑戰。

對於青年勞動者而言，究竟是身處在一種什麼樣的勞動環境中？比起其它族群，青年是否因為較具備數位能力，而更能感受到安全舒適的工作環境？又或者只是在進入或者準備進入到職場時，就直接面對企業運用AI要求持續提高生產力的工作場所？

而因應AI出現的新型態勞動機會，青年勞動者能夠擁有更多自主性，藉此滿足自我的時間和規劃，又或者是面臨更嚴密的管理與勞動監控，被迫承受著更大的工作壓力？

二、提案緣起

議題之多元觀點 與公共重要性

針對AI對活絡勞動市場或者勞動控制的影響，社會上存在多種不同觀點，並涉及公共利益的重要討論。

企業觀點：

AI技術被視為提升管理效率與競爭力的關鍵工具。透過AI強化勞工管理，能夠減少人為錯誤、提升生產力，甚至改善工作環境。例如，智慧化調度系統可以讓工作更為精準，避免人力資源浪費。然而，這種「效率優先」的思維，可能與勞工的福祉產生衝突，因為強化監控與績效導向可能會導致工作壓力增加，甚至損害勞工的健康與權益。

勞工觀點：

AI的應用可能導致監控與管理措施過於嚴苛。例如，某些企業已經開始使用AI追蹤員工的工作狀況、螢幕使用時間，甚至透過影像分析技術監測勞工行為，這種做法可能進一步出現剝削的狀況，讓勞工的隱私與自主性受到侵害。此外，當AI系統根據歷史數據進行決策時，可能會無形中排除某些群體，導致求職或升遷的不公平現象，特別是對於青年、女性與其他弱勢群體的影響更為顯著。

二、提案緣起

議題之多元觀點 與公共重要性

政策與法律觀點：

AI的應用符合正義與公平原則。目前部分國家已經開始針對AI管理提出規範，例如歐盟正在2024年批准的《人工智慧法案》，試圖確保AI的應用不會侵犯基本人權。在台灣，目前針對AI監控勞工的相關法規仍然較為薄弱，如何在促進科技發展的同時，保障勞工權益，成為政府與社會必須共同面對的挑戰。

社會整體的角度：

AI的發展可能影響勞動市場的結構與就業機會。許多研究以及相關經驗指出，AI取代部分重複性高、標準化的工作，導致低技能勞工失業風險升高。然而，也有另一種觀點認為，AI將創造新的工作機會，勞工應透過學習新技能來適應這場變革。關鍵問題在於，這種轉型是否能夠順利進行，政府與企業是否能夠提供足夠的培訓資源，協助勞工應對職場變革。

三、團隊介紹



1984年5月1日「台灣勞工法律支援會」誕生，1988年更名為「台灣勞工運動支援會」，1992年「台灣勞工陣線」成立。

隨著台灣從威權到民主，勞陣也從解決個別勞資爭議、協助勞工籌組工會，到目前致力於勞動政策的批判與改革。

台灣長期以來沉溺於「低成本、低賦稅、自由化」的血汗經濟模式，造成集體低薪與過勞，未來勞陣將從造成血汗經濟的問題根源著手，提出新的經濟發展模式，以追求「勞動正義的經濟民主」與「分配正義的社會民主」為目標。

勞陣作了什麼？

- 推動提高基本工資、周休二日、大量解僱勞工保護；
- 催生最低工資立法、失業保險、勞保年金、職災勞工保護體系；
- 改革集體勞動三法，建構工會保護制度及產業民主之勞工董事；
- 參與公共托育、長期照護、社會住宅、稅制改革、全民健保等倡議行動。

四、講者介紹

邱羽凡

- 現職：國立陽明交通大學科技法律學院副教授。
- 相關經歷：德國哥廷根大學法學博士，專業領域為勞動法、民事法、勞動關係理論。
- 將於活動第一階段，分享有關「AI與勞動」相關重要見解，協助參與者對議題建立初步理解，進一步探索個人與議題的關聯性，為後續的劇本發展與願景凝聚奠定基礎。

黃志雯

- 現職：於科技業專責數位匯流政策與法規事務、台灣科技大學資訊管理系博士候選人。
- 相關經歷：曾任職政府所屬法人，協助行政院科技會報辦公室、行政院智慧國家推動小組，以及經濟部加速行動寬頻服務及產業發展推動小組，執行各項國家大型資通訊政策。
- 本次將於活動第二階段，分享相關政策與經驗如何在科技發展與勞動保障之間取得共存。

韓仕賢

- 現職：全國金融業工會聯合總會祕書長。
- 相關經歷：長期擔任勞資爭議調解委員，為勞動部認證資深調解人。擁有豐富協助工會與資方協商團體協約經驗。
- 將於活動第二階段，分享相關政策與經驗如何在科技發展與勞動保障之間取得共存。

五、活動流程

時間	流程	進行方式
09:00-09:30	報到	
09:30-09:45	開場	開場並說明願景工作坊進行流程
09:45-10:35	議題講座	邀請邱羽凡老師簡述 A I 與勞動的發展、衝擊與國際趨勢
10:35-10:45	情境概述	針對此次審議主題進行說明
10:45-11:45	小組討論I	邀請主持人擔任桌長，與小組成員對主題現況與相關問題進行討論
11:45-12:00	小組成果分享I	由各桌桌長，針對小組的討論進行報告；大場主持人確認分享內容。
12:00-13:30	午餐	小組自由交流

五、活動流程

時間	流程	進行方式
13:30-14:40	議題講座	邀請黃志雯與韓仕賢分享針對議題 民間與工會的因應方式
14:40-15:00	茶點時間	
15:00-15:05	討論前說明	說明劇本中 A I 運用時可能的情境
15:05-16:10	小組討論II	由桌長帶領與小組成員針對主題進行討論
16:10-16:30	小組分享II	由各桌主桌長，針對小組的討論進行報告；大場主持人確認分享内容。
16:30-16:40	賦歸	交流及大合照

六、劇本

2045年未來劇本：AI與AGI主導的時代

背景：

在2045年，通用人工智慧（AGI）驅動的監控系統無處不在，企業廣泛使用AI分析員工行為數據（如鍵盤輸入、語音情緒、動作）以評估績效與忠誠度，違規者可能被立即解雇。AI監控也延伸至公共領域（如智慧城市的安全監控），引發隱私與權益爭議。審議將聚焦AI監控的挑戰與應對策略。

一、產業變革與就業機會流失

1. 客服與行政支援工作

- 轉變：AI整合語音分析、表情辨識與客戶歷史資料，可即時偵測客服員的「服務品質不足」或「態度不佳」的情緒指標。
- 失業風險：基礎客服與行政助理幾乎被AI替代。

2. 教育與教學領域

- 轉變：AI分析教師表情變化、語音語速、情緒強度等，以預測學生是否會失去興趣或理解困難，並評分教學表現。或以AI監控教師學生互動：AI偵測教師對學生是否「一視同仁」、有無「偏頗對待」、是否使用「正向激勵用語」，資料進入升遷評鑑系統。

六、劇本

3. 媒體與內容創作產業

- 轉變：AI直接分析內容表現（點閱率、轉發率、情緒影響力），並自動建議主題、修稿與發佈時間。
- 失業風險：記者、編輯等角色轉為「AI輔助產出器」，決策權大幅被系統取代。

4. 製造業與物流產業

- 轉變：工廠工人與倉儲員工配戴感測器，AI實時監控疲勞度、效率與偏差，稍有異常即被調職或淘汰。
- 失業風險：中階技術工人與現場管理人員被逐步取代。

5. 金融與行銷分析

- 轉變：分析人員的決策會被AI重新評分，如投資失誤等可能構成解僱依據。如提案自動分析與排名：將所有行銷方案輸入系統，AI會比對歷史案列、消費反應、ROI預測，自動給分與調整建議。或是產生客戶互動評估：AI可分析業務員與客戶互動的語音、文字、表情，評估「說服力」、「信任感指數」、「轉換率預測」，數據會影響獎金與續聘。導致投資分析師與行銷人員幾乎無法再憑經驗與直覺做決策，AI建議被默認為「更客觀」、「更準確」，發揮個人風格的空間被壓縮，創意案若違反AI模式偏好可能直接被篩除。

六、劇本

二、社會與經濟後果

1. 心理健康危機與職場壓力上升

- 被全面監控與「數據驅動績效」的環境下，勞工普遍出現「數位壓迫感」，進而導致焦慮、抑鬱與職場冷感。
- 青年特別受到影響，因其進入職場即面對全面量化的KPI與行為評分。

2. 勞資權力失衡

- 雇主藉由AI進行「預測性人力管理」：提前預測跳槽傾向、忠誠風險、效率下滑等進行處理，導致工會與人權組織失去談判空間。
- 雇員則面對「隱形算法主管」，難以抗辯與申訴。

3. 經濟結構極端集中

- 控制AI系統的少數科技平台企業，掌握人力配置與就業機會的主導權。
- 就業市場出現「平台封建制」：個人不再是自由勞動者，而是以績效數據向平台租賃工作機會。

4. 民主與監督機制的弱化

- 傳統的公共監察與勞動保障機難以適用於深度演算法決策機制，並且技術日新月異，監督機制無法及時跟上。

六、劇本

討論環節

小組討論I：AI的未來想像與生活融入

子題：1.現在如何使用AI

2.未來使用場景與工作樣態想像

小組討論II：面對AI監控、包容性的就業與AI發展

子題：1.針對AI監控，作為從業者的個人因應方式

2.工會與民間單位因應方式

3.政策協助的可能性

(一)機器管理人類？

把人當機器管理

摘錄自：鱸魚（異類矽谷），將人關進鐵籠裡工作：亞馬遜的血汗專利，
20180925

要了解血汗專利的背景與角色，首先要先了解亞馬遜出貨中心運作的流程。一個商品從貨架到包裝好上卡車，一共有三個階段。第一個階段是找貨和取貨，第二個階段是掃描和包裝，第三個階段是出貨上卡車。取貨的階段最花費時間。

在已經自動化的出貨中心，找貨是由機器人取代。這些機器人其實只是全自動化高速運轉的載重轉盤。機器人只知道商品的位置，但無法辨認商品，因為很多商品，比方像衣服鞋子及日用品，歸類非常困難，而且沒有固定的形狀，必須要靠人來辨識撿選。所以取貨仍舊是靠人工。

取貨的部分是由機器人把大約一人寬七尺高的貨架，整個搬到取貨員面前，由人工撿選，機器人再把貨架送回原位。在已經自動化的出貨中心，一項商品從下單到出貨，只要花25分鐘。有機器人的出貨中心，一年可以節省兩千萬美金的營運成本。如果推廣到150個出貨中心，亞馬遜一年的營運成本就可以減少30億。這對任何企業來講都是極大的誘惑。

可是亞馬遜對這個25分鐘也許還是不滿意。因為整個貨架拿過來再送回去就為了一項商品，效率上不夠節約。他們要把它壓縮到更短，最好的方法就是讓機器人直接篩選。

七、議題背景資料

鐵籠中的操作員——人將成為機器的一部份

2018年9月7號那天《西雅圖時報》報導亞馬遜申請的「工作鐵籠」正式得到專利局認可，引發輿論的撻伐。不過在當天亞馬遜立刻就對外宣布，這純粹只是一項紙上談兵的研究。他們不會把這項專利付諸實行。

亞馬遜並沒有說明鐵籠設計的實際用途。不過從專利的圖解可以合理猜測得出來，他們是打算把取貨員關在機器轉盤上的鐵籠裡，跟著機器人移動。在抵達貨架之後，由籠中的操作員利用機械手臂撿取需要的貨品。這樣就不需要把整個貨架搬到一百公尺以外的工作枱撿選，完了再把貨架送回去歸位。這種人機合一的鐵籠設計，不需移動貨架而且可以一次拿多樣商品。

這種設計營造的工作關係是由人來補強機器的不足。可是如果我們停下來思考一下，也許會覺得毛骨悚然。因為人已經變成機器的一部分，而且改變了主從的關係。我們好像不再控制機器，而是被機器控制。我們的角色淪為只是幫助機器，做它所不能做的事。

為了績效而拋棄人性？

坦白講，如果這項專利是把人放在一個玻璃控制室內，而不是鐵籠內，亞馬遜所遭受的批評也許不會這麼嚴厲。無論亞馬遜在設計曝光之後，用千百個理由來自圓其說，都很難讓人信服。他們的理由是出貨中心到處都是機械手臂，把人放在鐵籠裡是為了操作員的安全。問題是為什麼要用鐵籠？鐵籠難道沒有羞辱的意義嗎？這問題媒體沒有質問，亞馬遜也沒有解釋。我相信這麼做是因為鐵籠成本最低。他們的本意並沒有要羞辱任何人，只是亞馬遜可以為了成本而不在乎員工的感受。這就是他們的文化。

七、議題背景資料

新聞媒體《晨間早報》(Morning Call)曾經揭露亞馬遜賓州出貨中心，在攝氏40度的炎夏寧可僱用救護車在場外待命，也不願意安裝冷氣，因為這樣比較划算。如果員工中暑，就由救護站人員先行處理，恢復後就得回去繼續工作。情況嚴重的就抬上救護車送到醫院。這是幾年前另一個為了成本而不在乎人性的例子。所以我合理地相信，這種成本決定一切的文化也左右了上面的鐵籠設計。

追蹤員工一舉一動的手環

2018年稍早，《西雅圖時報》揭露過亞馬遜引起極度爭議的另一項專利——那就是可以追蹤員工手臂位置，以及移動頻率的手環。如果手該動而沒有動，或移動速度不夠快而跟不上工作進度，手環就會自動發出震動的警訊。這使我想起在研究動物行為的實驗科學上，科學家們用發出微量電擊的頸環來控制動物的行為。

監控員工效率作為未來考績的參考是一回事，可是直接用手環振動來警告又是另一回事。如果監工在每天下班的時候用報表顯示給員工，告訴他們今天落後多少——即使同樣是用手環追蹤，這項專利的爭議性也許也不會這麼高。同樣地，這背後的意義是為了效率，我們可以用機器來管理人，警告人……是不是將來也會用來懲罰人？人與機器的界線到底在哪裡？

當然手環也可以追蹤你所有的一舉一動，包括手臂有多少時間不是用在工作上，你去了哪些地方，上了幾次廁所，每次花多少時間。這個手環還嚴重暴露了個人隱私問題。只有機器人不需要隱私，也不會在乎羞辱。這項設計無疑是把員工當作機器人來管理。想得更恐怖一點，如果哪天這兩項專利同時用在同一個人身上……

七、議題背景資料

完全不在乎「人」的感受

《西雅圖時報》訪問了亞馬遜在英國物流中心的幾位員工。員工的回答是他們早就覺得自己已經是機器人了，不同的是他們只是個有血肉的版本。如果你的手放在不該放的地方，隨身的裝置都會發出警告。公司規定每一位包裝員，一個小時必須處理一定數量的包裹。如果落後，裝置也會發出警告——包括上廁所耽擱太久。有員工抱怨上廁所次數太多，或時間過久都會被扣分處罰。扣分點數集滿到六點就會被開除。

這種裝置只是把人當作機器人來管理使用。就像一位匿名的員工所說的「員工對他們來說就是機器人。在科技還沒有辦法完全取代人之前，他們卻已經全面化把人當作機器來管理。」在亞馬遜，人與機器的界線似乎越來越模糊。

身為消費者，亞馬遜的確有不可抗拒的魅力，這種魅力已經大到足以讓社會大眾願意忽視他們對員工的不人道。每次看到媒體披露這些血汗內幕，大家在一陣震驚之後，也許隔天又回到網上繼續購物，沒有人想到要去抵制。如果血汗工廠出現在中國，美國人馬上走上街頭呼籲大家抵制血汗產品。我不知道這算不算是雙重標準，不過這也是亞馬遜最成功的地方。他們的顧客絕對至上原則，早就把全球兩億多客戶的心收買得服服貼貼。世界上只有對亞馬遜抱怨的員工，沒有對亞馬遜抱怨的消費者——這是他們背後靠著撐腰的武器。

另外亞馬遜的出貨中心都是蓋在美國中西部比較落後的鄉下。那些地方本來就沒有足夠的經濟規模或工作機會。這些員工在失業與不人道的管理文化之間，除了忍氣吞聲似乎沒有其他的選擇，這也給外界落了個願打願挨的口實。所以這樣的文化就沒有阻攔地持續下去。

七、議題背景資料

有了兩億客戶撐腰，有了為績效作藉口，對於幾十萬員工血汗管理的批評，亞馬遜一路走來都只是見招拆招。這是文化問題而不是技術問題，所以看來短期內也不會有太大的改變。

讓數據左右人性？

企業追求利潤與績效是天經地義的事，只是這個代價的底線到底在哪裡？企業的一切都以數據衡量——從盈收到效率到產能……無一倖免。只是不要忘記，人性不能，也永遠不應該放在天秤的另一端衡量。

思考爭點：

AI是否有可能驅動新一波的「泰勒主義」（以科學化管理與工作標準化為核心，透過精細分工、時間控制與效率最大化來提高勞動生產力的管理制度）？包括為使效率與生產量極大化，透過即時數據分析與演算法決策，加強對勞動者的監控，進一步削弱勞動者的人道考量。

七、議題背景資料

(二)當AI工作時，是否有偏見？

改寫自：許慧瑩，AI偏差歧見的另類難題 – 使用者對AI風險評估工具的差異性互動

不少AI支持者確信，AI不但可增進決策的精確與準確性，因為少了人類的介入，可大幅減緩人類決策的偏差歧見（bias），提供較為客觀、公平的決定。但是，這個論述在實際應用案例上受到不少質疑。以美國為例，法院採用的風險預測與評估輔助工具 – COMPAS協助法官就嫌疑人的再犯風險做為量刑參考，引發許多爭議，其一為風險評估的結果帶有種族的偏差歧見。

Pro Publica在2016年的調查研究報告中指出，COMPAS的演算法在種族存在嚴重的偏差歧見，對於黑人嫌犯再犯風險的估算為白人嫌犯的兩倍之多。美國亦有研究發現，醫療院所與保險公司大量使用之健康照護資源分配的演算法，可能含有系統性偏差歧見，對於可能面臨類似病況的黑人病人與白人病人，演算法可能不會推薦（引導）給可改善複雜醫療需求的患者。後續追究可能的原因為，雖然有某些疾病在黑人群體比例較高，例如：糖尿病、貧血、腎衰竭、高血壓等，但總體而言，依據健康照護年度支出來看，若為患有同樣慢性病時，黑人比白人支付額少1,800美元，又演算法基礎於支出來劃分風險群組，而因此在健康照護的資源在黑白種族間產生了差異。

有論者主張，透過研究實驗並未發現COMPAS的演算法有偏差歧見的可能，另一派論者則認為，基於資料科學之輸入垃圾、輸出也是垃圾（garbage-in, garbage-out）的原理，若訓練資料早存有人類的偏差歧見，那麼演算法經學習與訓練後所提供的預測與建議，不可避免的也將或多或少含有偏差歧見（bias-in, bias-out）。

七、議題背景資料

針對演算法可能因為訓練資料而產生偏差歧見的預測建議，目前已有不少機器學習的學者專家，正著手於研究如何發現與調校偏差歧見，並提出因應個別議題的新興技術，例如：依據公平性指標將不同層級標註為不相容的工具、研發得以減少偏差歧見之新演算法等。

不過，就演算法的開發與訓練是否帶有偏差歧見，最終還是回到最原始的問題－資料、資訊、與知識。對於演算法來說，最完美的理想狀況為完全記錄可呈現人類行為的資料、透過人工或機器標註資料與資料間的關聯性或連結資料（不論是結構化或非結構化資料），並尋找相關行為模式所得之知識，才得以提供演算法之開發與學習的可能。然，就現況觀察，以人類法官量刑為例，法官最後心證的形成所基礎的資料，除了實體證據之外，尚包含被告的對話、反應、肢體動作等。法官的判決基本上只儘可能的說明心證的過程，但是還有一些心證形成的過程，並無法完全描繪說明、且記錄下來。最終，法官的判決不僅由有形的文字與資料間歸納而成，還有許多無形或感官上的因素，成就最終的決策。以現時非結構化之判決文字/鑑定書之文字表達，可能無法完全呈現法官/專家於法庭論辯環境所形成之最終心證/鑑定之過程。但是，如果沒有資料或足夠的資料可開發與訓練演算法，對於演算法來說，就是一個問題。有相關的資料才得以檢視，演算法的開發與訓練是否帶有偏差歧見。

除此之外，華盛頓郵報提出，有關檢視是否為開發或訓練資料造成亦帶有偏差歧見，光是要如何定義公平，以及與將公平資料化，就是一件頗為困難的挑戰。再者，即便有可能完整記錄人類的行為資料（含感官資料間的關聯性），適用到專業人員進行專業決策（例如：醫療、法律），差別對待或處置是否即可直接解讀為具有偏差歧見，可能尚無簡單明確的判斷方法。

七、議題背景資料

截至目前為止，演算法是否可做出比人類更準確、精確與公平客觀的判斷，進而於未來可能取代人類或專業，仍待證實，但至少現階段演算法仍可以做為人類進行決策時的輔助工具。最近，越來越多的公私部門逐漸採用演算法輔助評估、決策（在公務機關的處分、司法刑事偵察與法院系統、醫護專業之臨床決策輔助系統、公司員工聘僱的篩選），以刑事司法系統所使用之風險評估工具為例，風險評估工具僅扮演著輔助法官進行決策的工具提供預測建議給法官參考，是以，風險評估工具的預測建議是否正確與公平，最多僅可能是影響法官作最終決策的間接因素，最終決策者還是人類。

Ben Green 與 Yiling Chen 在〈Disparate Interactions: An Algorithm-in-the-Loop Analysis of Fairness in Risk Assessments〉一文中，透過實驗探討風險評估工具對人類決策的影響。他們認為，討論風險評估工具的公平性時，關鍵不在工具本身是否帶有偏見，而是其預測建議如何影響使用者（如法官）的最終決策。研究以刑事司法中「審判前釋放」為例，透過 Amazon Mechanical Turk 平台進行實驗，觀察一般人在有無工具預測建議情境下的判斷差異。結果顯示，風險評估工具的預測確實會影響人類判斷，即使實驗對象非法律專業者，仍可反映工具與使用者互動對決策產生實質影響。

該研究指出，風險評估工具並未如預期減少自由裁量，反而轉移至法官如何解讀與採納工具建議的過程。這種「人機間差異性互動」（disparate interaction）可能在不知不覺中滲入人類決策，特別是當工具的預測建議已內含偏誤時。作者強調，若要理解AI對決策過程的真實影響，不能僅關注演算法技術本身，更應觀察人類在使用工具過程中的行為模式與採納程度，特別是在高度專業領域如司法與醫療。

七、議題背景資料

最後，Ben Green & Yiling Chen提出「在決策過程納入演算法」（Algorithm-in-the-Loop, AITL）的概念，主張將演算法視為社會實踐的一部分。這意味著我們應重視使用者與AI的互動，而非僅看演算法的輸出結果。隨著AI工具日益滲透公共與專業決策場域，人機互動的社會效應（如依賴、信任、責任轉移）將變得更重要，未來研究須進一步揭示演算法如何影響專業判斷與決策責任，並提出相應的制度設計與倫理規範。

有學者將人機互動不同的行為模式加以分類為：偏好傾向接受機器預測（亦即較不信任人為預測）一稱為信任演算法（algorithm appreciation），以偏好傾向不接受機器預測（亦即較為信任人為預測）一稱為厭惡演算法（algorithm aversion）。即便有這些行為現象，但究竟，人類決策者對於演算法所提供的預測建議，在人機互動之差異性對待/處置/歧視偏見的背後原因是甚麼？使用者採納演算法之預測建議與採納的程度，在一般使用者與專業使用者是否會有所差異；專業使用者差異性對待/處置做出判斷背後的原因為何；如果人類將大幅仰賴系統作出的專業性預測建議而做出最終決策時，雖然表面上仍是由人類作出判斷，但演算法的建議對於人類決策過程是否有邏輯論理的引導（暗示）的效果；再加上演算法的黑箱效應（black box effect），是否會使專業人員失去專業自主？

當人類採納演算法的預測建議，但卻無法理解其原因，演算法還是輔助工具，人類還有決策權嗎？另外，傳統被動靜態系統與主動動態系統，在做為輔助工具時，是否會產生不同的結果，亦或作為專家決策系統要為不同之管制？在使用輔助工具時，使用者決策後的責任感，是否會有不同之心路歷程？

七、議題背景資料

思考爭點：

- AI判斷是否有可能出現隱形偏見，影響勞動環境？包括AI可能學習過去數據中的偏見，進一步影響勞動條件與機會，特別是青年身為脆弱勞動者，恐將進一步受到衝擊。
- 應對AI勞動控制如何面對？包括從技術、政策與工會角色出發，確保勞動者的自主權與更公平的演算法機制，但同時又不因過度限制而降低技術創新的動力。

七、議題背景資料

(三)AI對人類判斷的影響

改寫自：Min，【還有 700 天準備】AGI 若在 2027 年實現，人類如何因應「智慧被超越」的衝擊？，20250430

多名人工智慧業界知名人士，近日聯合發表了《AI 2027》產業發展預測，為外界帶來2027 年人工智慧發展狀況及影響推論的詳細路線圖。在《AI 2027》預測中最值得注意的部分，是專家認為「通用人工智慧」(AGI) 將在 2027 年實現，而比 AGI 更強大的「人工超級智慧」(ASI)，更是會於 AGI 誕生後的幾個月內登場，整體發展速度超乎外界預期。

無論是科學研究或創造性工作，專家認為 AGI 在幾乎所有認知任務上，都將可以展現出跟人類相互匹敵，甚至是有所超越的能力。當 AGI 能夠主動去適應環境、執行常識推理與自我完善，就會進而演變為 ASI，即電腦系統的智慧大幅超越人類，解決我們至今仍然無法理解的問題。而《AI 2027》的推論仍然基於假設，對此AI界有支持和質疑截然不同的意見。

AI 自動化引發失業潮，2 年緩衝期遠遠不足

譽為 AI 教父的圖靈獎得主 Geoffrey Hinton指出 AGI 最快 2028 年就會達成。當外界對於 AGI 的到來，同時感到興奮和恐懼之時，人類將很難避免「更聰明的 AI」所帶來的各種衝擊與影響。專注於 AI 領域的外媒記者 Jeremy Kahn 就指出，假如 AGI 真的在未來幾年內實現，社會勢必會引發失業潮，畢竟許多企業、組織都會開始轉向自動化。然而，僅僅要用剩下的 2 年時間，為 AGI 所造成的衝擊做出準備，這樣的寬限期顯然不夠讓個人與企業提前適應，包含客戶服務、內容創作、程式編寫和資料分析等產業，可能會因此面臨劇烈的動盪，假如全球經濟在轉變的過程中出現衰退，那麼因 AI 而出現的失業者，他們所面臨的壓力將會更大。

七、議題背景資料

當把思考「外包」給 AI，人類剩下什麼？

另一方面，假設 AGI 的出現並沒有引起失業潮，人類仍然會開始對自身存在的意義提出疑惑。人類可以透過「思考」行為，找到自身存在的確定性，即使自己被感官欺騙，或是被他人誤導，只要擁有思考的事實，就能夠證明本身的存在。

Gary Grossman 表示，笛卡兒的「我思故我在」使人類以思想者的身分，成為現代世界的中心人物，然而如果機器也就是 AI 可以開始思考，或者說看起來像是在思考，而人類又把思考任務「外包」給人工智慧，那麼這對於現代人的自我價值來說，究竟將意味著什麼？近來亦有一份國外研究表明，當人們在工作上嚴重依賴 AI 時，他們的批判性思考就會減少，久而久之就會導致本應保留的認知能力，因此出現了退化。

開始為最好、最壞的情況都做好準備

愛德曼人工智慧卓越中心全球負責人 Gary Grossman 認為，如果人工智慧發展正如《AI 2027》所預測，即 AGI 與 ASI 在未來幾年或之後不久出現，人類必須迅速處理 AI 對工作、安全，以及人類自身所帶來的各種影響。同時人類也無法否認，AI 在加速創新探索、減少工作痛苦及擴展人類原有能力方面，擁有難以比擬的非凡潛力。

《AI 2027》對於 AGI 的各種預測，也許會出現失誤或問題，但其可信度已經足以讓人類展開行動，以身為企業、政府與社會的一份子的角度，為各種最好、最壞的情況做好準備。

七、議題背景資料

從企業、政府到個人，面對 AGI 能怎麼做？

對於企業而言，這意味著商業領導人要開始對 AI 安全領域發起投資和研究，同時關注組織彈性，透過 AI 增強人類員工的優勢；對於政府來說，主管機關得加速發展 AI 監管框架，以解決模型安全性評估等當前問題，以及盡早釐清 AI 所長期擁有的風險。

至於個人方面，Gary Grossman 提醒我們應該擁抱終身學習，並專注於那些獨特的人類技能，例如創造力、情緒智商和做出複雜判斷，同時跟 AI 工具之間發展健康的工作關係，避免遭到 AI 削弱主體能動性。

Gary Grossman 強調，對於 AI、AGI 及 ASI 進行抽象辯論的時代已經過去，人類要開始為 AI 轉型做好具體準備；我們的未來不僅是由演算法所主導的社會，更會是從現在這一刻開始，由人類所選擇且堅持的價值觀所共同塑造。

思考爭點：

- AI是否將取代人類判斷，技能退化與職業自主性消失？包括AI決策如何減少勞動者的判斷空間，勞動者是否僅成為執行者，而喪失自主性，亦或勞動者的技能是否會因為長期依賴AI而退化。

七、議題背景資料

附錄：其他參考資料

1. ILO，公司如何利用人工智慧招聘、監控和裁員及有組織的勞工可以如何反擊

<https://reurl.cc/O52o4R>

2. 天下雜誌，亞馬遜用AI自動解雇900人 你的電腦也被公司監控了嗎？

<https://www.cw.com.tw/article/5097612>

3. 中央研究院法律學研究所 許慧瑩，AI偏差歧見的另類難題 – 使用者對AI風險評估工具的差異性互動

<https://infolaw.iias.sinica.edu.tw/?p=2261>

4. CMoney，工作鐵籠、監控手環，連上廁所時間都被控制…亞馬遜的2項血汗專利！為了成本漠視員工的感受，就是他們的文化

<https://www.cmoney.tw/notes/note-detail.aspx?nid=184444>

5. Adecco，台灣未來勞動力調查報告：人工智慧對當代工作者的挑戰

<https://reurl.cc/nqQa2D>

6. KPMG，解密歐盟人工智慧法案

<https://kpmg.com/tw/zh/home/insights/2024/04/decoding-the-eu-artificial-intelligence-act.html>

7. AGI 若在 2027 年實現，人類如何因應「智慧被超越」的衝擊？

<https://buzzorange.com/techorange/2025/04/30/agi-in-2027-ai-predictions/>

8. 最容易被AI機器人取代的10大職業，「這工作」竟被比爾蓋茲、馬斯克點名

<https://www.businessweekly.com.tw/careers/blog/3018528>

9. 數位人權刻不容緩：公民團體對台灣政府提出9大訴求

<https://www.amnesty.tw/node/21604>

10. 生成式AI下的勞資關係

<https://voicetank.org/20250324-2/>

11. AI不會成為影視終結者？好萊塢編劇工會明文規管AI，美國勞資首開先例

https://global.udn.com/global_vision/story/8664/7524465

12. 數位科技驅動服務業勞動市場轉型之現況與因應

作者：吳慧娜；劉光哲

《臺灣勞工季刊》77期(2024 / 03) P.66-70

13. 數位時代勞工資訊隱私權保護問題初探

作者：顏雅婷(Ya-Ting Yan)

《勞動及職業安全衛生研究季刊》28卷4期(2020 / 12) P.60-69

七、議題背景資料

附錄：其他參考資料

14. AI科技的發展對勞動市場、組織工作及訓練資源分配的影響

作者：黃文柔(Huang, Wen-Rou)

《台灣教育研究期刊》 6卷2期 (2025 / 03) P.219-237

15. 論數位時代下工作場所改變之挑戰及對我國啟示－以美國電傳工作加強法案為例

作者：許雲翔

《臺灣勞工季刊》 62期 (2020 / 06) P.21-30

16. 新聞機器人為誰「勞動」？自動化新聞學引入新聞產製的影響及論述

作者：劉昌德(Chang-De Liu)

《中華傳播學刊》 37期 (2020 / 06) P.147-186

17. 《血汗AI：為人工智慧提供動力的隱性人類勞工》

作者：James Muldoon, Mark Graham, Callum Cant 出版社：大塊文化

18. 《技術陷阱：從工業革命到AI時代，技術創新下的資本、勞動力與權力》

作者：Carl Benedikt Frey 出版社：八旗文化

Let's Talk

參與青年問卷



114年「青年好政-Let's Talk」計畫