

NDC-MIS-110-001（委託研析報告）

**公部門機器演算法應用
之制度調適與路徑分析
（結案報告）**

**國家發展委員會編印
中華民國 111 年 3 月**

NDC-MIS-110-001（委託研析報告）

**公部門機器演算法應用
之制度調適與路徑分析
（結案報告）**

受委託單位：數位治理研究中心

計畫主持人：陳敦源

協同主持人：廖洲棚、黃心怡、張濱璿、王千文

研究助理：陳郁函、郭彥廷、張瀛心、邱家妤

國家發展委員會編印

中華民國 111 年 3 月

目次

目次	I
表次	III
圖次	VII
摘要	IX
ABSTRACT	XI
第一章 緒論.....	1
第一節 研究背景	1
第二節 研究目的與方法	1
第三節 名詞解釋	3
第二章 文獻回顧.....	9
第一節 數位轉型中演算法的革新與落差	9
第二節 演算法對於決策的影響與輔助	16
第三節 各國演算法應用之個案採擷.....	22
第四節 人機協作職能模式	56
第三章 調查分析方法	67
第一節 研究設計與方法	67
第二節 中央各部會應用演算法之現況調查	81
第三節 法務部保護司觀護人之實驗調查	86
第四章 調查分析發現	99
第一節 法規調適之跨國分析	99
第二節 資源配置之調查分析	107
第三節 才能管理之實驗分析	141
第四節 落差分析	164
第五章 結論與政策建議.....	185
第一節 研究結論	185
第二節 政策建議	190
參考文獻	193
附錄	209

附錄一、 中央各部會應用演算法之調查問卷.....	209
附錄二、 法務部保護司觀護人實驗調查問卷.....	220

表次

表 1：子計畫一演算法類型與美國數位治理的案例	10
表 2：子計畫一跨國資料分析之設計	15
表 3：各國對於演算法治理之價值原則以及組織規範架構比較表	50
表 4：傳統官僚與演算官僚的比較	57
表 5：歐盟電子能力架構（European e-Competence Framework 4.0）	60
表 6：人機協作職能地圖	65
表 7：子計畫一研究資料來源摘要	68
表 8：子計畫一深度訪談受訪者背景資料表	72
表 9：子計畫一深度訪談題綱	73
表 10：子計畫一第一場焦點座談參與者背景資料表	74
表 11：子計畫一第一場焦點團座談題綱	75
表 12：子計畫一第二場焦點座談參與者背景資料表	75
表 13：子計畫一第二場焦點團座談題綱	76
表 14：子計畫一第三場焦點座談參與者背景資料表	76
表 15：子計畫一第三場焦點團座談題綱	77
表 16：子計畫一第四場焦點座談參與者背景資料表	77
表 17：子計畫一第四場焦點團座談題綱	78
表 18：子計畫一第五場焦點座談參與者背景資料表	78
表 19：子計畫一第五場焦點團體座談題綱	79
表 20：子計畫一第六場焦點座談參與者背景資料表	79
表 21：子計畫一第六場焦點團體座談題綱	80
表 22：子計畫一調查構面、操作化定義及題號一覽表	85
表 23：子計畫一觀護人實驗前期規劃時程表	88
表 24：子計畫一觀護人實驗收案之時程表	90
表 25：第一階段問卷調查結果綜整表	110
表 26：第二階段調查受訪者基本資料一覽表	111

表 27：文官運用 AI 態度相關題項的 EFA 結果.....	113
表 28：文官心理資本相關題項的 EFA 結果.....	115
表 29：文官所屬機關擁有 AI 系統所需資料準備度相關題項的 EFA 結果	117
表 30：文官所屬機關組織管理現況相關題項的 EFA 結果.....	118
表 31：文官運用 AI 預期態度之多元迴歸分析結果.....	120
表 32：行政機關運用 AI 系統所需人才職能分析結果.....	120
表 33：人才職能應具備與實際具備程度落差分析表.....	121
表 34：人機協作部分訪談資料之彙整.....	123
表 35：問卷樣本基本資料整理.....	141
表 36：各組樣本基本資料統計.....	143
表 37：觀護人工作狀況與「檢察機關案件管理系統」使用現況調查..	145
表 38：實驗前觀護人對機器演算法導入再犯評估之期望調查.....	147
表 39：實驗前受試者對機器演算法支持程度調查.....	148
表 40：實驗各案例分配次數及比例.....	149
表 41：各組每輪修改答案比例之比較.....	154
表 42：給予實驗組受試者正確或錯誤判斷後，修改答案比例之差異比較	155
表 43：實驗後受試者對機器演算法導入之支持度調查.....	156
表 44：受試者對機器演算法導入工作之接受度調查.....	157
表 45：受試者創新意願之調查.....	159
表 46：受試者對工作流程自動化之需求調查.....	160
表 47：受試者公共服務動機之調查.....	161
表 48：各面向議題內容及次數彙整表.....	169
表 49：行政運作面向議題內容及次數彙整表.....	169
表 50：服務設計面向議題內容及次數彙整表.....	170
表 51：人才職能面向議題內容及次數彙整表.....	171
表 52：以聊天機器人支援公共價值的傳遞.....	174
表 53：外部介入資訊對受試者判斷影響整理.....	175

表 54：機器學習演算對公部門個體層次影響之理論視角落差分析	176
表 55：三種策略途徑下的落差分析	178

圖次

圖 1：AI 發展下的演算法分類	4
圖 2：子計畫一公部門機器演算法應用之制度分析理論架構	12
圖 3：兩個因子在不同公共任務上的決策影響	19
圖 4：以風險框架為基礎的演算法分類	20
圖 5：OECD AI 政策的生命週期	24
圖 6：OECD 各國治理組織設定盤點	25
圖 7：EU 值得信賴的 AI 治理框架	27
圖 8：EU 值得信賴的 AI 七大要件關係圖	28
圖 9：EU 落實值得信賴的 AI 的循環生態	29
圖 10：美國針對 AI 規劃之國家政策里程碑	34
圖 11：ICO 監理架構圖	39
圖 12：整合性演算法治理架構	58
圖 13：人機協作下跨領域技能	64
圖 14：子計畫一研究架構與設計	67
圖 15：子計畫一研究流程圖	71
圖 16：行政院組織架構圖	82
圖 17：調查問卷設計架構圖	84
圖 18：子計畫一實驗調查流程示意圖	88
圖 19：子計畫一實驗之操作頁面圖示	92
圖 20：子計畫一實驗組 Arm 1 之確認答案圖示	93
圖 21：傳統程式演算法與機器學習演算法之差異	100
圖 22：歐、美、英 AI 規範演進比較	103
圖 23：人才職能落差分析雷達圖	122
圖 24：受試者回答 B1「該個案犯罪行為的產生受到不良的社會影響的 程度為何？」之平均分數	151
圖 25：受試者回答 B2「個案在社會中受到排斥且缺乏對他人的關心， 產生與社會疏離的程度為何？」之平均分數	151

圖 26：受試者回答 B3「個案在自我特質上，因缺乏同理心與負面情緒，導致有問題解決技巧的認知不良情形的程度為何？」之平均分數.....	152
圖 27：受試者回答 B4「個案願意接受觀護與治療的配合程度為何？」之平均分數	152
圖 28：受試者回答 B5「請問個案的危險級別為何？」之平均分數	153
圖 29：受試者回答 B6「請問個案的再犯風險級別為何？」之平均分數 資料來源：本計畫。	153
圖 30：受試者回答 B7「針對個案之危險級別，觀護處遇評估為第幾級？」之平均分數	154
圖 31：落差分析模型.....	165

摘要

新興科技的發展越趨蓬勃，綜觀國內外皆有其相關領域結合新興科技使業務的執行更加有其效益，較廣為人知的新興科技為人工智慧的廣泛運用，如欲發展理想中的人工智慧，其前置作業需要演算法的運用讓機器學習發展。我國運用機器演算的數位服務於中央與地方政府中皆有嘗試，但推動的過程還是遇到許多挑戰，包含：數位成熟度不一、可以運用該技術的業務未盤點、相關依循法規與風險控管不明，面對諸多問題，政府部門需要針對法規調適、組織資源配置、數位職能養成進行系統性地評估，因此，為因應這股不可逆的趨勢，本計畫認為機器運算運用於政府數位服務，仍有賴前瞻性政策和制度規劃。

為深入探索此議題，本計畫重點有三：第一、蒐集3-5個國家將機器演算法應用於政府數位服務的政策、法制、組織數位成熟度等，對照臺灣目前現況進行比較分析，包括：正確性（accuracy）、公平性（fairness）、課責（accountability），以及對外溝通機制，例如：透明（transparency）、知情同意（informed consent）等；第二，透過個案研析，瞭解國內中央政府機關推行應用機器演算法的現況、衝擊，以及其對於組織或個人所面臨的挑戰；第三，參考國內外政府與個案對於機器學習的服務應用，從組織調適的角度，進行落差分析（gap analysis），並從中研提機器演算法應用於政府數位服務的具體改革路徑（path）建議，其中至少包括：法規的調適、資源的配置，以及能力的管理等。

針對上開重點，探討出目前研究上的三大部分的研究結果：第一、借鏡OECD、歐盟、美國、英國、新加坡的經驗，進行四大項目的比較，可發現規制層級的專責、領域的多元與差異性、執行流程的共通性，皆能成為檢視我國目前現況與發展該機制的借鏡；第二、從中央各部會應用演算法的調查中可發現，對於演算法應用於公共服務中是可行的，但執行上會礙於資源面，如：預算、人力職能、組織共識等影響成為推行阻礙；第三、人機協作的職能發展其應然與實然的狀況仍有差距，但該職能的發展確實不可或缺，需要以整體組織發展以及數位人才培育思維，使整體政府部門的認知有所共識，才有利於組織對於該技術應用的推動。

最後，本計畫目前得出三項研究建議，第一、短程目標：盤點各部會對於機器演算法運用的規劃與投入資源；第二、中程目標：適合各部會遵循的共通性原則與執行指引建置；第三、長程目標：公務人員的數位轉型素養養成與數位職能訓練。

關鍵詞：演算法、演算法治理、人工智慧、數位創新、政府數位轉型

Abstract

As the development of emerging technologies becoming more and more vigorous, the utilization of artificial intelligence is not only popular but inevitable for public sector organizations. For Taiwanese governmental organizations, there have been various attempts to adopt algorithm in digital services, but those experiences are still facing many challenges. In this research, an overall review of those experiences will be researched to build up a forward-looking institutional plan for government to be successful in moving into the Machine-Algorithm driven governing era.

In order to explore this issue in depth, this research plan to do three things. First, collect the policies, legal guidelines, and experiences of 3-5 countries which apply machine algorithms to digital services, and conduct a comparative institutional analysis against the current situation in Taiwan through four public values: Accuracy, Fairness, Accountability, and Transparency. Second, through an “inventory” survey on central government’s organizational and individual level, to understand ministries’ adopting efforts of machine algorithms. Thirdly, we also conduct a quasi-experiment about the machine-algorithm assisted decisions for parole officers to learn how the behavior changed under the adaptation. After all the above research efforts, we conduct a gap analysis and then propose specific reform paths for the application of machine algorithms to government digital services, including at least the adjustment of laws and regulations, the allocation of resources, and the management of individual talents.

The results of this research is as follows. First, borrowing from the experience of OECD, the EU, the United States, the United Kingdom, and Singapore, and comparing them against the case of Taiwan, it is clear that that the choice of regulatory agencies’ responsibilities, the implementing processes, and building up legal guidelines for adopting machine algorithms can all be important for the future development in Taiwan. Second, according to the “inventory” survey of ministries, the application of machine algorithms is welcomed but successfully rare, which is due to the feasibility, decision-making and implementation problems in which the applications are hindered by budget, human reactions, organizational consensus and ethical concerns. Third, there are practical gaps between the adaptation and development of human-machine collaboration, especially on the talent management and

building of organizational consensus on moving forward on the technological development.

Finally, this project has come up with three recommendations as path for future development. First, short-range goals: to take inventory of the public organization to encompass the difference between organization and come up with different plans. Second, middle-range goals: establishing suitable and common principles via guidelines for government agencies to promote and adopt machine-algorithm projects. Third, combined short, middle, and long-term goals: promoting digital transformation through resource reallocation and training of civil servants to be capable of handling the transformation.

Keywords : Machine-Algorithm 、 Algorithm Governance 、 Artificial Intelligence 、 Digital Innovation 、 Government Digital Transformation

第一章 緒論

第一節 研究背景

從實務上來看，德國政府2019年挹注五億資金推動人工智慧（Artificial Intelligence，簡稱「AI」）發展，應用於道路交通儀控系統、人臉辨識、犯罪回報及查緝自動化等；歐盟為促進AI科技應用，同時兼顧民眾隱私、正確性以及公平性，降低產業發展伴隨而來的新興風險，2020年提出《人工智慧白皮書》（White Paper on Artificial Intelligence）和《歐盟資料策略》（European Data Strategy），並成立法律發展組織負責研議AI和數位民主的影響評估。另外，澳洲政府於2019年提出《人工智慧白皮書》（Artificial Intelligence Roadmap），強調AI用以解決問題、發展經濟、改善人民的生活品質，並應重視資料治理、信任、透明、倫理及數位環境整備度等議題。另一方面，美國政府機關已有多元的《人工智慧使用手冊》（AI Toolkit）用以政策制定和執行，協助達成有關廉正（integrity）、效率（efficiency）和公平（fairness）的挑戰性任務，可稱為「新演算法治理」（New Algorithm Governance）。

我國運用機器運算的數位服務諸如：臺北市政府交通局的智慧化交通量調查分析系統、臺中水湳智慧城花博園區的自駕公車、健保署去識別化的醫療影像應用，另有AIoT建置的智慧生活服務系統、空汙通報與環保實地稽查的智慧環境監控等。然而，政府部門推動的過程還是遇到許多挑戰，首先，因為公共組織的數位成熟度（digital maturity）不一，領導團隊不清楚自己的組織可以與不可以發展機器運算的部分為何；再者，政府相關法制結構的密度甚高，新興的機器運算可能會給組織甚至廣大社會帶來的影響未明，政府組織面對其應用風險控制需要正視；最後，面對機器運算對傳統組織人力的替代效果，政府部門在人力（human resource）、能力（talent）與資源配置的調適（adaptation）也需要系統性地評估，因此，為因應這股不可逆的趨勢，本計畫認為機器運算運用於政府數位服務，仍有賴前瞻性政策和制度規劃。

第二節 研究目的與方法

基於上述說明，本計畫的主要目的與工作有三：

第一，蒐集3-5個國家將機器演算法應用於政府數位服務的政策、法制、組織數位成熟度等，對照臺灣目前現況進行比較分析，包括：正確性

（accuracy）、公平性（fairness）、課責（accountability），以及對外溝通機制，例如：透明（transparency）、知情同意（informed consent）等。

第二，透過個案研析，瞭解國內中央政府機關推行應用機器演算法的現況、衝擊，以及其對於組織或個人所面臨的挑戰。

第三，參考國內外政府與個案對於機器學習的服務應用，從組織調適的角度，進行落差分析（Gap Analysis），並從中研提機器演算法應用於政府數位服務的具體改革路徑（path）建議，其中至少包括：法規的調適、資源的配置，以及能力的管理等。

再者，本計畫主要採用質／量並重的研究方法，透過深度訪談、焦點座談、中央部會問卷調查、跨國比較與實驗研究等五種研究方法蒐集資料（詳細內容請參第三章圖14與表7），之後透過質化與量化資料的交叉分析，以發掘國內中央政府機關應用機器演算法於行政實務的現況與遭遇之法規面、資源面及人才面衝擊和挑戰，並嘗試就現況的掌握研擬應對策略，俾提供本計畫之委託單位參考。

接著，由於本計畫在概念的策略途徑包含三個層次：法規調適、資源配置、與才能管理，深度訪談與焦點座談在三個層次都有所應用，中央部會問卷調查主要針對資源配置與才能管理的部份，而跨國比較則專注於法規調適的層次，而觀護人的實驗研究，主要是在才能管理的層次，針對法務部觀護人進行實驗調查。當政府調查問卷特別是實驗研究的部分有兩個目的：第一、提供主管單位，公部門使用機器演算嵌入之數位公共服務產品、平臺與工具時的適用性建議；第二、預計能提供機器演算導入公部門機關時，以個案方式提供實務經驗中，組織與個人的主觀想法、態度以及在實驗情境下公務人員決策過程與結果的影響。將有助於該部門後續導入方案的設計與改善，也能提供其他公務機關對於機器演算治理中，組織與個人層次對於公共服務本質與內涵的再探。

最後，在此特別說明，本計畫進入中後期的階段，從初步的中央機關問卷與深度訪談當中得知，政府在演算法的應用上還處於初階的時期，其中得知「聊天機器人」（Chatbot）是較多機關曾經導入應用的個案，因此就特別為此舉辦焦點座談與深度訪談，這個隨著研究進程發現並且外加的個案焦點，主要目的仍然是要讓本計畫的落差與路徑分析的基礎，是來自於明確實務狀態下的真實實踐案例。整體來說，本計畫在政府推動演算法應用的初階時期進行整體的策略性研究，讓政府面對數位轉型的階段性發展，能有一個完整的回顧與前瞻的循證評估（evidence-based evaluation）的機會，並且能夠在實務的基礎上，針對未來包括新成立的數位發展部所要面對的各種關於演算法應用的公共政策研擬與執行，帶來重要的策略性資訊。

第三節 名詞解釋

一、 什麼是「演算法」(Algorithm)？

簡單來說，演算法是「在數量化的世界裡，由人們所設計出來，針對特定的問題所寫下的一個可操作的運算步驟，其目的是用來指揮電腦『自動』解決人們所想處理的特定的問題」。這個設計的步驟是由不同的數學運算所構成（比方說， $9 > 7$ 中的「大於」符號），其中含有特定判斷的邏輯功能（比方說，如果 $9 > 7$ ，那就 $x=7$ 中的「如果...那就; if... then」之判斷符號）。比方說，我們想比較並且找出下面七個數字中哪一個數字最小：

$$\{5, 9, 2, 8, 6, 1, 4\}$$

如果是人類的判斷，我們只要一眼看過去就知道數字“1”是該集合中最小的，不過，人類要如何教電腦作同樣的事，就是要設計一個演算法來尋求解方，比方說，人們的解法可能如下面四個步驟的演算設計：

步驟一：我們設一個數字空間，該空間只能放一個數字，先放進無限大(∞)。

步驟二：依序將每一個數字與空間數字比較，如果比空間數字小，就置換。

步驟三：顯示真實步驟。

- 1、 先比 5 與 ∞ ，5 比較小，接著將空間數字置換成 5。
- 2、 再比 9 與 5，9 比較大，空間數字不置換，仍然是 5。
- 3、 再比 2 與 5，2 比較小，接著將空間數字置換成 2。
- 4、 再比 8 與 2，8 比較大，空間數字不置換，仍保持 2。
- 5、 再比 6 與 2，6 比較大，空間數字不置換，仍保持 2。
- 6、 再比 1 與 2，1 比較小，接著將空間數字置換成 1。
- 7、 再比 4 與 1，4 比較大，空間數字不置換，仍保持 1。

步驟四：確認空間數字 1 是最小的，程序結束。

當然，這不是解決這問題的唯一演算法；比方說，設計者也可以從第一個數字開始與其他數字進行比較，如果一發現有比該數字小的數字，就判斷這個數字不是群組中最小的而刪除，接著一個一個比直到發現“1”這個數字比大家都小的時候，程序就停止。因此，不同的設計者可以建構不同的演算法，經過某種效能、效率、或是錯誤率等標準的檢驗，一個問題可能會有目前設計出最有效率或出錯率最低的演算法，這就是所謂演算法的「比較」或「標竿」研究 (Comparing and Benchmarking Algorithm;) (Beiranvand, Vahid, Warren & Yves, 2017)。

演算法是電腦語言的基礎，而 AI 的爆炸性發展，也直接促使演算法研究的蓬勃發展，一般來說，AI 的專家們對於應用一個又一個的演算法堆疊起來一個決策模型，需要不斷設計與比較最適於達成目標的演算法，也相信經過這樣的精進作為，有朝一日就可以發展出來勘比擬人類決策的機器人，這樣的想法是一種「手作的 AI」，不過，有更多的資料科學家相信，AI 成形的關鍵，是在於設計者找出了能夠產生「**人工智慧自產的演算法**」（AI-generating Algorithms, AI-Gas）（Clune, 2019）的 AI 演算法，讓機器自己的深度學習等同或是更優於人類，無論如何，AI 除了應用領域的周邊程序與實體機具以外，最核心的靈魂就是演算法的發展與（自我）進化，因此圖 1 將 AI 發展內演算法的類型¹描繪如下：

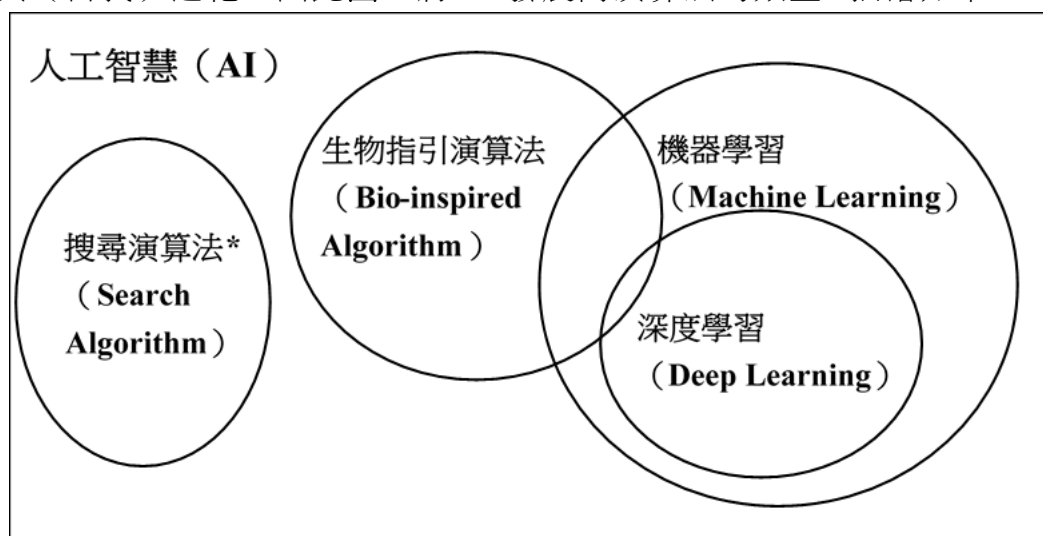


圖 1：AI 發展下的演算法分類

資料來源：Rishal Hurbans（2020）。

從上圖中可以看出，各種演算法的發展雖有其不同脈絡，但是都是 AI 發展的重要基礎，比方說，本節開始所舉的那個尋找一組數字中最小數字就是一個「搜尋演算法」，而生物指引的演算法類似用數學模式去發展出生物的演化或是生存的演算法，這樣的演算法類似生物力學對於航空與航海工程發展之影響一樣，也深深地影響 AI 最核心機器學習與深度學習，當然，不論是機器還是深度學習的設計，也都是以演算法為其基礎，因此，AI 發展機器學習（包含深度學習）的領域，也都是以演算法其核心，大略可以分為三種：監督式學習（supervised learning）、非監督式學習（unsupervised learning）、與強化學習（reinforcement learning）²。

比方說，「決策樹的演算法」（Decision Tree Algorithms）導引下的機器學習就是一種標準的監督式學習，將資料的性質標註清楚，並且設好

¹ Rishal Hurbans（2021）。AI algorithm families，2021 年 11 月 28 日，取自：<https://dev.to/rishalhurbans/ai-algorithm-families-2kf7>。

² Neelam Tyagi（2020）。Top 10 Machine Learning Algorithms，2021 年 11 月 28 日，取自：<https://www.analyticssteps.com/blogs/top-10-machine-learning-algorithms>。

決策流程，讓機器依照標註好的資料與決策流程去進行決策；而「集群演算法」（Clustering Algorithms）則是非監督式學習的常用方法，不事先標註資料，讓演算法從資料歸納特徵並且自行分類，雖然誤差有可能較大，但機器決策過程的自主性高，AI 的意涵足夠；最後，強化學習的演算法大多與人類腦神經網絡的結構有密切關係，常被用在影像辨識的「積疊腦神經網絡演算法」（Convolutional Neural Networks Algorithms, CNN），就是讓電腦自行辨識與分類影像，用函數相互影響並且能夠產生新函數的運算法，讓機器從影響當中自行歸納、分類、回饋定義、最終自行判斷物件為何的一種機器學習。

二、為何政府 AI 的治理需要從演算法著手？

演算法與 AI 這兩個當代的熱門關鍵字之間，到底有甚麼樣的關係？這個問題可以從資訊科學的角度來回答，不過，由於本計畫主要是從政府公共管理的角度出發，對於演算法與 AI 之間的關係，將從下面三個面向來討論。首先，我們將先討論演算法與 AI 的定義連接；再者，我們將從 Herbert Simon 的決策理論來討論演算法與 AI 的關係；最後，我們將討論演算法治理（Algorithm Governance, AG）從政府管制的角度來看，為何是切入未來 AI 管理的關鍵鑰匙。

（一）AI 是演算法所搭建起來的未來

AI 是一個人類創造自我的過程，目前雖然尚未達到最終目標（創造出一個與自己智能相當的機器人），但是不同面向的進展時有所聞，因此，早在 1950 年代電腦科學家艾倫·圖靈（Alan M. Turing）曾經提出一個「圖靈測試」（Turing Test）的區辨概念（demarcating concept），意指一個人要應用數理方法去測試電腦互動中的那個「他」是機器人還是人類；導致人類如想成功創造出與自己相同的機器人，必須通過圖靈測試，通過的機器人在聰明程度已與真人無異。綜括來看，不論是機器人智能的建構，還是測試機器人的建構，都需要不同程度演算法協助，才能達成這個自我創造的過程，因此，演算法類似樂高遊戲的積木塊，而 AI 則是用這些積木塊搭起來的智慧機器人。

對政府的運作來說，「自動化」的演算法應用概念早在 AI 一詞被高舉之前就已出現，1980 年代「辦公室自動化」（office automation）推動的過程中，比方說，將政府公文系統電腦化的改革過程中，也一樣是透過演算法來協助推動，特別是該系統中電腦與使用者的介面建構，演算法的應用更是不可少的核心（Konsynski, Bracker & Bracker, 1982），這樣看來，對於電子化政府常久以來的變革，早在 AI 因為網際網路、雲端、與大數據等資訊技術爆炸性發展以前，改革工具的基礎就是以演算法的應用為基礎的，當然，這些年大數據與 AI 的概念更加成熟的發展，也可以看作是演算法應用的一場革命，目前電子化政府向 AI 發展的重心，就是

演算法的進化，Naqvi & Munoz（2020:100）學者說這樣的革命是：

「一種本質上統計學，奠基於資料科學，但是應用演算法來(讓機器)學習達成任務（It is statistical in nature, based upon data science and driven by training algorithms that learn to accomplish tasks.）」。

更重要的，AI的發展對於政府所面臨的數位轉型有著決定性的影響，然而，早期的發展最重要的就是「人機協作」（Human-robot Collaboration, HRC; Bauer et al., 2008）的導入與應用，本計畫將在第二章文獻回顧的第四節當中，將 HRC 當作一種公務職能來進行討論。

（二）從專家系統到 AI 的演算法轉變

不過，我們如果回到前 AI 大爆炸的 20 世紀，人們對於政府自動化的期待是樂觀（一種科技樂觀論，陳敦源，2016），這其中包括兩個成份，一個是類似「科技程序正義」（technological due process）由法律學者 Danielle Keats Citron（2007）提出，另一則是提出專家輔助決策系統由諾貝爾經濟學家 Herbert Simon 提出。前者從國家行政諸多複雜的管制決策網絡中，尋找更加中立、正確與效率的「自動化行政國」（automated administrative state），這樣的期待從行政管制國家的角度當然就會引起諸多課責的問題；而後者著眼的是「決策的品質」是否得以經過機器輔助而提升的問題，這是國家制度以及政府組織追求公共利益的一個重要的切入點，對於政府數位轉型的當口，不論從行政國還是公共決策，都是十分重要的問題。

（三）演算法治理可以引導良善的 AI

事實上，經過 2016 年美國總統大選與英國脫歐公投過程中，一家應用資料技術影響選舉的 Cambridge Analytics 公司，以及臉書（facebook）社群網路平台演算法對於前述重大選舉公投活動的涉入，被認為對社會產生了所謂「網絡中立性」（Network Neutrality）（Wu, 2003）的影響問題³，有偏誤的演算法人工智慧（Algorithm AI）之應用，被認為是重要的罪魁禍首之一，這個社會普遍的想法使得 AI 中藉由演算法而產生的決策受到質疑，有必要對於 AI 的發展過程中演算法是否符合人們的價值期待（比方說，公平與效率等）進行課責的管制，這也是政府數位發展與治理的重要任務（Zarsky, 2016）。因此，有 Thierer、Sullivan 與 Russell 等美國學者（2017）認為 AI 對社會的負面影響有兩種方法禁行治理，首先，管制 AI 技術的研發與應用，但是這樣「絕禁」的作法會阻礙科技發展對社會

³ 網絡中立性簡單來說是一個網路平臺經營的原則，就是網路平臺公司必須平等對待網路上每一則溝通信息，不應該因其個人、內容、平臺、裝置、以及其他足以影響平臺公平應用性的原則，而這個爭議的焦點，很多時候是放在演算法之上，演算法表面看起來是可開放的篩選規則，但事實上它也會受到設計者本身偏誤（bias）的影響。

帶來的正面影響；而另外一種方法則是依循學者 I. Rahwan (2018) 提出「人類留在 AI 環境中」(human in the loop) 的「弛禁」策略，對 AI 的基本元素演算法進行管制，主要方法就是將社會規範與人類價值置入演算法，Wirtz、Weyerer 與 Sturm 等三位德國學者 (2020: 819) 以下列文字描繪這種管制的方法：

「政府與企業要依照從人類價值與社會規範而來的期待，建立管制目標，並且攜手建構相關規範與管制法令，這些規範與法令要通過 AI 演算法來落實，以期達到（管制 AI 發展負面效應）的政策目標（Government and industry are meant to work together to provide regulations and standards and represent the goals and expectations of society based on human values, ethics and social norms. The results are then implemented into an AI algorithm and evaluated towards the defined goals.）」

綜合來說，演算法是數學與資訊科學的基礎元素，並非所有演算法的目的是發展 AI，但是 AI 的發展不可能沒有演算法的領域，而 AI 的發展是一日千里，伴隨而來的公共價值被扭曲風險也是日益嚴峻，不過，作為需要資金與市場發展的 AI 產業，以及該技術攸關一國未來國際貿易優勢與國家安全的需求下，絕禁或是「先禁再說」的管制方式是不利於於國家競爭力發展的。因此，本計畫基於前面的論述，將本計畫的核心名詞定義為「人工智慧機器演算法」(AI Algorithm)，為了論述簡潔的方便，故本計畫之後都將該名詞簡稱為「演算法」(algorithm)，代表以 AI 為主要領域，包括在機器學習等關鍵領域中演算法的治理，是未來智慧政府發展息息相關的新興科技管理的領域。而「演算法治理」(Algorithm Governance)，意指與人工智慧機器演算法相關的政府與跨部門的機關，在其推動、管制、以及市場化過程中的策略管理相關議題，也是政府在 AI 時代的管制作為，可以切入並且建構有效治理框架與行政法制建構的一個起點與新領域⁴。

三、 其他名詞解釋

- (一) 落差分析：落差分析主要是針對研究對象的「實際」(reality) 與「期待」(expectation) 狀態之間差異的分析，最有名的就是「服務品質」的 PZB 模型中的五個落差分析 (SERVQUAL;

⁴ 本計畫面對 AI 與 Algorithm 這兩個詞的使用，基本上是「相互適用」(interchangable)，主要原因有三：(1) AI 這個詞過度受歡迎但範圍較不明確，但是 Algorithm 這個詞過度專業但卻定義更基礎，也更廣泛，但本計畫需要兩者的聯集；(2) 專業文件上只有少數使用演算法，大多是用 AI 來概括，因此，本計畫主要是用更基本的「演算法」來蓋括定義；(3) 由於專業文件上如果使用 AI 這個詞，並不代表就不是在談演算法，而本計畫也無法且無意去作正在發展中的領域之名詞正名的工作，因此，團隊會延用專業文件上的用法，放入本計畫報告當中，但是仍然是在「演算法」的討論範圍內。

Zeithaml et al., 1988)，這樣的分析也可以應用在法規調適、資源配置、與才能管理等面向的分析。

- (二) **路徑分析**：路徑 (path) 分析除了是統計學因果關係的分析方法外，它也是一種改革策略的時間序列分析與論述 (Johnson and Wanta, 1996)，舉例來說，AI 人工智慧在政府發展的先期建設，不演算法導入，而是資料治理 (data governance)，先資料治理再導入演算法就是一種政府數位轉型的路徑建議。
- (三) **制度分析**：制度理論是討論人與其環境限制的學說，制度分析最有名的就是諾貝爾經濟學獎得主 Ostrom (1986) 對於資源配置的制度分析與發展框架 (Institutional Analysis and Development, IAD)，主要討論在法規、資源、與能力有限的狀態下，組織要如何更有效地進行組織與管制對象的改革與發展。本計畫依照 Ostrom 的理論，將制度分析的框架應用在本計畫的「制度調適」(institutional accesement; Williamson, 1980) 研究，一共分為三個層次：法規調適、資源配置與才能管理，將分別於下面定義。
- (四) **法規調適**：法規調適意指廣義制度規範的改變與因應。從政府端面來說，對 AI 的新興科技，應該進行相關法規的調適，另一方面，政府推動法規變革的過程中，也應該從被管制者的角度進行「管制影響評估」(regulatory impact assesement, RIA; Radaelli, 2009)，本計畫主要從跨國比較進行法規調適的分析與建議。
- (五) **資源配置**：公共組織必須依賴資源而運作，如果從公共管理的理論來看 (Starling, 2005: 18)，資源管理的標地包括人力、財務與資訊等，而資源有稀缺性 (scarcity) 與排他性 (exclusivity) 本質，因此，從落差分析尋找增補資源的缺口外，在有限資源的前提下進行改革就需要對資源進行最有效的配置 (allocation)。
- (六) **才能管理**：公部門或其他部門面對數位轉型的新科際發展的需求，最核心的努力領域就是公部門人力「才能」(talent) 的發展與管理 (Charlwood, 2021)，包括政府演算法人力的需求如何滿足？政府長期新科技人力的組織與管理作為應該如何重新設計等問題，這也將是未來數位發展部的重點業務之一。

第二章 文獻回顧

第一節 數位轉型中演算法的革新與落差

一、政府數位轉型中的演算法創新革命

嚴格說起來，數位政府的領域是一個不斷「追新」的政策領域，過去這些年一路上被雲端、大數據、物聯網...等科技商業模式給拖行了這麼久，總是要面對「創新」(innovation)承諾的落實問題，特別是目前演算法這個概念充滿了科幻式的創新想像下，從推動政府數位轉型(digital transformation)的政府部門來看，越來越需要一個統整的科技概念來兜攬五花八門的科技創新與應用市場上的技術，以期能精準投入政府有限資源推動數位轉型；另一方面需要應用這個概念當作政府創新政策推動的跨域協調平臺，藉此政府數位轉型的政策將得以產出並且開始推進。因此，本計畫選擇「演算法」(algorithm)這個概念當作基礎，所謂演算法，根據綜合各家英文字典的描繪：「演算法是一組計算的規則或是程序，主要是人類設計給電腦遵守並且被用來協助決策與問題解決的基礎」，不過，這些規則也可以是開放且具創造性的，也就是說，雖然由人類設計，但是機器可以在一定的範圍內自我學習，也就是所謂的「機器學習」(machine learning)，這部分也是AI最重要且最令人期待。

2020年美國聯邦行政會議(Administrative Conference of the United States, ACUS)盤點142個聯邦政府機關使用AI的現況⁵，顯示有高達88%機關採用基礎AI技術，高達45%機關已採行AI或機器學習，惟過程中不確定其正確性與效率是否能夠得到保障；而將演算法(Algorithm)應用於政府治理分為五大類、列出對應的案例，並針對具有影響力的案例進行個案研究，最後提出幾項公部門應用AI技術的建議方向。

當然，根據 Gil-Garcia 等學者(2016)的歸納，智慧政府的「智慧」(smartness)具有 14 個不同的面向，然而，非智慧系統一樣可以讓政府變得更智慧，比方說，實體的專家系統應用德菲法進行政府服務改善的集體決策，就不會用到 AI 的系統，因此，1970 年 Herbert Simon 所擘畫的

⁵ 美國聯邦行政會議(Administrative Conference of the United States, ACUS)，是國會通過「行政聯席會法」(The Administrative Conference Act)而建立於 1964 年的美國聯邦政府的獨立機關(independent agency)，其建立目的是通過各種研究與討論的方法，針對美國聯邦政府相關運作功能，包括管制方案、行政補貼與津貼、或是政府運作所需的執行與評估能力，進行有系統且前瞻的評估與建議，以達成改善政府的效率、公平、與正確性。相關資料請參：<https://www.acus.gov/administrative-conference-united-states-acus>。

AI 系統，仍然是以真人為主，而機器只是輔助真人決策的專家系統(expert system)而已⁶，雖應用到演算法，仍和目前的機器學習有極大差異，Konsynski、Bracker 與 Bracker 等美國學者(1982:10-11)認為 Simon 時代與目前強調機器學習有三大差異：其一，比起過去專家系統，目前 AI 演算法與結果之間的連結「難以洞悉」(inscrutable)；其二，目前 AI 演算是「非直覺」式(non-intuitive)邏輯，即便懂得演算法的專家也無法通盤理解並且複製；其三，Simon 的 AI 系統以專家為核心，現今演算法治理範圍中的政府運作，包括強化執行(enforcement)、監理分析與監控(regulatory research, analysis, and monitoring)、裁決(adjudication)、公共服務(public services and engagement)以及內部管理(internal management)如表 1，項目豐富且與公部門工作職能關係較密切，當然，這個差異的轉變，讓 AI 發展過程中產生負面影響的疑慮升高，政府被期待在其中扮演管制者的壓力越來越大。

表 1：子計畫一演算法類型與美國數位治理的案例

應用類型	實例
強化執行 (enforcement)	- 證券交易委員會、醫療保險和醫療補助中心以及國稅局用以決定執行標的之順序 - 海關和邊境保護局以及運輸安全管理局之人臉識別系統 - 食品安全檢查局用以預先告知食品檢測事宜
監理分析與監控 (regulatory, analysis, and monitoring)	- 消費者金融保護局對消費者投訴的分析 - 勞動統計局對工傷的敘述進行編碼 - 食品藥品管理局對不良藥物事件之分析
裁決 (adjudication)	- 社會安全局用於糾正裁決錯誤 - 專利商標局用於裁決專利和商標申請
公共服務 (public services and engagement)	- 郵政署的自動車和手寫辨識工具 - 住房暨城市發展部及美國公民暨移民服務局的聊天機器人 - 機構對法規制定相關意見的分析
內部管理 (internal anagement)	- 衛生及公共服務部用於協助制定採購決策 - 總務署用於確保聯邦招標的合法性

⁶ 專家系統根據 Herbert Simon 自己的設想，主要是程式設計師借重領域專家對於專業內涵的知識，比方說，病理學家或是政府中的查帳專家等，請這些人依照決策邏輯設立一組判斷規則，可以藉此用來在一個專家決策情境中，獲得一個合理的決策結果。

應用類型	實例
	國土安全部用於應對機構系統遭網路攻擊的工具

資料來源：美國聯邦行政會議《機器演算法應用於政府治理》報告。

最後，政府應用演算法進行的數位創新與轉型，一方面是值得期待的，但是另一方面卻也是令人恐懼的，因為，這些演算法仍然受到人類價值觀的直接影響（Diakopoulos, 2015），從公共價值（public values）的角度來討論演算法，可以看出公部門雖然已經從技術上大量引進演算法到行政運作的各個環節，但是仍然必須面對各種價值層面的挑戰（Veale, et al, 2018），學者Andrews（2019）就認為，政府演算法的應用會帶來四種治理風險，包括：「選擇誤差」（selection error）、「演算法違法」（algorithmic law-breaking）、「操控與對奕」（manipulation or gaming）與「演算法宣傳」（algorithmic propaganda）的問題，需要政府以適當的政策介入並且避免之。

二、 演算法創新的制度調適層次問題

演算法若要運用於公部門中，目前面臨的制度應有何調適？本計畫根據「理性選擇制度論」（Rational Choice Institutionalism, Shepsle, 2006；陳敦源，2019）的說法，強調以個人或可以清楚定義的團體，通稱為「個體」（agent）為分析單位，基於理性自利（self-interest）的人性假說，以演繹推論的方法（deductive logic），尋找公共事務運作之因果解釋。新興科技的發展，不管是演算法、大數據、AI等，人們都希望能運用這些新興科技來幫助組織發展、個人效益提升等，在私部門已倡議並運行多年的過程中，公部門也漸漸意識到，若公共服務需要更貼近民眾需求，同時，以更有效率、效能方式完成該項任務時，新興科技的協助為之重要。

反觀我國公部門運用演算法之數位服務，誠如研究背景與目的所述，上至中央政府下至地方政府都已開始嘗試試辦甚至運用於公共服務中。然而，推動過程總會面臨到數位成熟度不一、組織不清楚可發展演算法的部分等。面臨這些問題，政府有需要更有系統性盤點若該機關要導入演算法於公共服務中，現今人力、資源配置上應如何調適。

本計畫為了建構公部門演算法應用之制度調適，以2009年諾貝爾經濟學獎得主Elinor Ostrom（1986）提出制度分析（Institutional Analysis）的理論，作為本計畫的分析架構，從憲法選擇層次、集體選擇層次與個人選擇層次三個層面，建構出三種策略途徑，分別為法規調適（憲法選擇層次）、資源配置（集體選擇層次）與才能管理（個人選擇層次），試圖從

三種途徑建構出公部門應用演算法，其制度調適的架構，該制度分析理論（如圖2）對本計畫之研究具有導引作用。

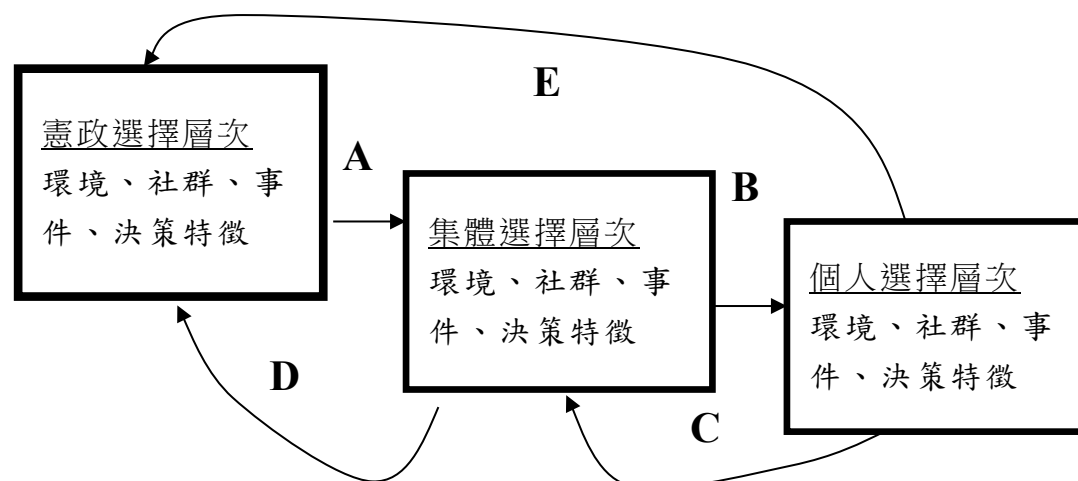


圖2：子計畫一公部門機器演算法應用之制度分析理論架構

資料來源：本團隊改繪自Ostrom（1986）。

為達成本計畫之研究目的，從制度分析架構發展出三種策略途徑，依據本計畫之主題，從相關文獻試圖整理出三項「制度調適」（institutional assessment; Williamson, 1980）分析：「法規調適」、「資源配置」與「才能管理」對於公部門運用演算法之定義。

第一、「法規調適」為公部門中常見之課題，尤其面對新興科技如雨後春筍般出現，公共服務如欲導入新興科技，事實上，法律規範的修法速度是無法趕上日益更迭的技術發展。因此，數位政府的治理過程中，如欲運用新興科技於公共服務中，首當其衝即是法規調適的問題，如：金融監理沙盒（Regulatory Sandbox）為使FinTech得以發展，在現行法規尚未完成修法階段，為先行試驗該技術的可行性，監理機關與受監理業者間透過積極的互動與對話模式，讓監理機關透過制定監理框架，於框架中較為彈性修改當前制度甚至是灰色地帶之法規。國發會也為因應數位科技發展，試圖降低法規適用之不確定性，予新創產業發展空間，成立「新創法規調適平臺」；另外，經濟部中小企業處發展條例第二條也點出法規調適制度的定義：

「本要點所稱中小企業法規調適制度，係要求行政機關對有關中小企業權益之法規，應注意其對中小企業經營影響其遵行能力，並應給予參與法規形成過程之機會，以建立公平、合理之中小企業法規環境。」

因此，若參考該法規調適制度的定義，以公部門運用演算法來看「法規調適」，即係檢討有關領域之法規於適用對象的影響，並給予適用對象參與法規形成調整過程之機會，為建立公平、合理、適當之法規環境，所

採行之各項因應調整與準備作為。透過探索相關類似個案的發展，從法制層面、資源層面與監理單位層面剖析，並且反觀我國之現況，試圖提出應調適之處。

第二、「資源配置」(resource allocation)根據國家教育研究院的定義為：「在經濟活動中，由於資源稀有性(scarcity of resource)的產生，造成可用資源的匱乏與不足。在資源稀有下，如何使既得的資源能夠發揮使用的效率，達成公平、合理原則，就要注意到資源有效分配的問題。資源分配者可以優先順序(priorities)或取捨(trade-off)等方式，來達成分配效果。」數位政府發展歷程中，資源配置是另一常見的課題。行政院為發展智慧國家政策，邁向國家的數位轉型，以強化國家科技發展整體規劃與提升備戰力為目標，確認組改增設數位發展部，其組改目的如下所述：⁷

「將業務統整，並統一事權、整體規劃數位發展推動相關資源，統籌產業、政府、社會與國民生活數位轉型的基礎工程及環境之整備，以及資安、資料治理、監理沙盒、人才培育與法制規範等業務，協助各機關落實數位治理。」

從上述對於組改目標的陳述，即是一種資源配置的重分配過程，將業務進行盤點，釐清權責與歸屬，並確立組織目標。數位發展部的目標為建構數位整備度並協助各機關落實執行，業務包含：組織結構的調整、資源的整備、法規的調適以及因應改革的人才培育。故「數位政府的資源配置」，在資源有限的情況下，進行該願景發展相關政策的配置，於數位治理過程中，釐清各機關的組織願景與目標，據此提出短期、中期、與長期之預算編製、人力編列與數位資源(如：資料管理、倉儲與檔案)配置的操作。

第三、「才能管理」(talent management)為人力資源發展中重要的管理議題，隨著全球化之影響，公私部門對於人力資源管理越趨重視，才能管理也漸被提起，對此該定義也逐漸達成共識。關於才能管理的定義從切入的角度不同而有所差異，分別為：1、才能管理是以典型的人力資源部門的實踐與職能為理念；2、才能管理是根據人力資源規劃和預測人力需求來定義；3、才能管理關注高績效與高潛能的人才，反之亦然，故為一種通則性的管理(Lewis & Heckman, 2006)。

受到新公共管理影響，公部門師法企業引進人力資源管理，才能管理的議題備受討論，其才能管理的定義為「建構系統性的管理模式，識別組織內存有哪些關鍵職位，這些職位對於組織之競爭力與優勢上具有不同

⁷ 鄭宏達(2021)。行政院明將拍板組改草案 數位發展部、科技部等先行，2021年3月24日，取自：
https://money.udn.com/money/story/7307/5341032?from=udnamp_storysns_fb&fbclid=IwAR0UxCQDGHvqQITJ8q8Cuyv4FxdWQeqvemFpRU14IAMZL83wikyVCrfLukU。

的貢獻，因此，需開發具有高潛力與績效的人才擔任該職位，並培育該人才；另外，也需要建構具有差異性的人力資源架構，適度安排組織內之非關鍵職位但仍各司其職的成員，以確保該員對組織的持續承諾」（Ariss, Cascio & Paauwe, 2009）。

針對上述對於三種策略途徑的闡述，其研究目的之一的探討，將從「法規調適」與「資源配置」途徑，從法制、資源與監理單位三方面分析國外案例，並針對我國運用演算法之現況進行比較分析；針對研究目的二，分析我國中央各部會機關以及法務部觀護人之個案，透過「資源配置」與「才能管理」，瞭解中央各部會機關其人力、預算與數位資源的配置，以及從人力資源發展的角度，剖析中央各部會機關不同的階層與組織內的分工，以及應培育之能力發展與培訓藍圖，另外，以法務部觀護人為個案，進行問卷實驗瞭解當演算法參與決策後，如何影響檢察機關觀護人進行案件評估的決策。最後，研究目的三為綜合研究目的一與二之研究發現，從三個策略途徑進行落差分析並提出具體改革之建議。基此，三層次的制度分析架構，有助於對於公部門如欲運用演算法於公共服務中，於法制、組織與才能三大途徑有更全面的剖析。

三、 政府演算法創新的落差與變革路徑分析

綜上所討論的政府應用機器演法的現實，如果政府的演算法應用是不可避免的變革方向，政府相關政策的扶持將會是政府數位轉型政策的重要環節，因此，本計畫從預評估與過程評估的角度，針對目前臺灣中央政府的演算法的應用現況進行落差分析（Gap Analysis, Flinders & Dommett, 2013）。舉例而言，績效落差（Appleby & Alvarez-Rosete, 2003），「落差」一詞是指「政策現況（status quo; S）」與「政策應該達到的境界（policy goal; G）」之間的差異，通常也是公共問題的「問題」一詞的定義，因此，本計畫將從憲政、集體與個人的三個層次（前一節當中的制度分析架構層次），尋找目前政府演算法應用的“G”與“S”之間的落差，包括國外法規制度（G）與我國法規制度（S）的落差、國內部會的標竿機關（G）與其他機關（S）的落差、以及制度預期（G）與行為人行為結果（S）之間的落差，進行現象與原因的分析。

接著，從前述的資料蒐集與分析當中，本計畫將從訪談與焦點座談的方法，蒐集各方意見，從策略管理的「策略路徑」（strategic path）的概念，對於我國政府運用演算法未來的變革，進行「路徑分析」（path analysis），主要是從前面一節的架構下，針對部會組織的機器運算之組織變革提出改革策略，配合法規調適的路徑與個人管理誘因的路徑，提出短、中、長期的改革建議，當然，這些研究成果將會引領我國政府在智慧政府的新時代中，一方面能夠獲取科技創新的效益，另一方面降低科技危害的風險，讓演算法在政府的應用可以順利並且有效益。比方說，就目前

蒐集到的文件顯示，我國在制度層次可以進行的策略落差與路徑分析的大概模式如下。

個人資料運用為例，因應科技的快速更迭，國外從對於影響個人權利的資料保護，漸漸轉移重視公共利益之權利的資料治理，在追求資料運用上的思維，會依照需求評估可帶來的利益，但在利益與權利的取捨上，可能會使個人資料保護的自主性有所退讓，進而產生演算法應用可能的落差(gap)。為使落差得以改善，經文獻檢閱發現，制度性的變革路徑(path)可透過幾種方式為之，如：設立具有監管功能的委員會，將資料運用的倫理層面納入考量，甚至將其建立具體規範；或是讓參與者透過滾動的方式，在過程中動態同意運用與否，並賦予退出權利（何之行，2021）⁸，上述皆為可參考之方式。

反觀國內，目前「個人資料保護法」有具體規範外，對於巨量資料的使用並無其他嚴格法規的規範。以「全民健康保險研究資料庫」之建置為例，起初中央健康保險署其資料蒐集之目的僅係在保費申報與稽核使用，後續委託國家衛生研究院所探討應如何建置該資料庫，但很顯然後續研究如運用該資料庫是已超出原目的外之使用，具有爭議。因此，在公共利益與個人資料保護上，如何使其二者取得平衡，是未來任何新興科技運用上的首要課題。

因此，關於演算法應用的法規調適研究，期待透過文獻資料的法規盤點、專家訪談、焦點座談等方式，蒐集各國之詳細規範運作與未來討論方向，與相關專家進行深入探討，希望提出未來可行之政策方向與法律體系建構規劃，如表2所示。

表 2：子計畫一跨國資料分析之設計

研究重點	研究主題	研究方法
法制	國外法律與我國法律（個人資料保護、研究法規、資料庫規範等）	文獻探討 深度訪談 焦點座談 與行政機關溝通會議
資源	資料庫建置、個人資料	
監理單位	行政機關、專門委員會（實質審查、倫理審查）	

資料來源：本計畫。

⁸ 何之行（2021）。死者健保資料再利用，釐清個資疑慮，2021年3月22日，取自：<https://udn.com/news/story/7266/5334140>。

第二節 演算法對於決策的影響與輔助

一、 演算法機制對決策的影響框架

AI中的技術機器之演算正以無法估計的速度成長演進，是公部門不可忽視的技術革新。處在資訊社會，演算法輔助人類決策的工具對人們而言非常具有吸引力，由機器學習技術產生的演算法，用來推薦、預測到所需要的人們來執行工作任務，藉由這種方式，人類讓機器學習過去的資料數據，用來訓練與精進系統，而逐漸讓演算法自己決定要執行的流程與任務執行。但一個值得公部門關注的核心問題是，當這些演算法預測的準確性若聲稱與人類相當時，此時人類的角色為何？如何讓人機協作的運作順利且符合組織需要，進而達到數位轉型的目標，是探討演算法機制對決策影響的主要實務貢獻。

早在1999年，Barth與Arnold (1999) 即預告AI與這類能協助決策能力的技術發展，將對公部門組織帶來質變。不單是組織層次的影響與變革，對於組織個人來說更是全新的挑戰，不僅包括態度認知的部分，也包括公務人員本身業務執行層次的行為層面。

他們預期未來系統將有能力可以獨立的學習、扮演自主代理人 (autonomous agents) 的角色、評估環境，並擁有思考價值、動機和情感的能力。但這也帶來了許多關於行政裁量權與組織決策的反思與困境，包括回應性 (responsiveness)、判斷力 (judgement) 和課責 (accountability) 等挑戰。

首先在回應性的部分，Barth與Arnold (1999) 提到，公共行政人員面臨的困境之一是行使裁量權，即便我們知道人與機構都有可能受到自身利益影響，但我們仍依賴著公務人員的「公共利益」來行使裁量。但事實上這種抉擇並不容易，即便他們沒有被自我意識與自身利益影響，但還是不免的存在既有偏見。公共政策中，決策都反映了價值選擇和決定做與不做的偏見，機器之演算所希望達到的人工智慧化，是希望能夠預測和瞭解潛在的價值和偏見，增加了回應決策的可能性。在判斷力的部分，大部分學者承認公共行政人員在執行法律必須擁有部分的合理判斷空間。換言之，官僚不可避免地必須行使裁量，過去傳統認為這部分是人類優於電腦之處，是只有人類才能行使的能力（因為機器不適合做最後判斷，也無法對非法律規範外的灰色地帶做決策），但也有些知識與常識是人類從很小的時候就開始學習且未包含在參考書中不證自明的事實，要做出正確的判斷，不僅需要運用嘗試與經驗，有時還需要同情心。這也是目前演算法研究人員期望創造出與人類思維合而為一的機器。

最後在課責的部分，Barth與Arnold (1999) 認為機器之演算所希望發展的終極AI有可能讓公民認為公職人員更應該對公共服務與政策負起責

任，畢竟這些系統都是政府所主導開發的。另外，就演算法治理而言，釐清每個「責任點」（responsibility points）變得更為重要，明定從政策制定到前線的作業人員，在每個環節下須由誰負起責任，以落實課責。而當發生錯誤時，造成疏失的人員應負起實質的法律、政治或行政等相關責任。資訊科技與新興機器之演算或許更能不費力追蹤流程，達到監督的作用。在公民參與變成政府治理常態的今日，公民對政策制定的影響仍決定於是否具有足夠的專業與獨立評估能力。隨著通信技術的進步，擁有能夠瞭解、近用演算法系統的公民或許可提高對公共政策問題的評估能力，AI與機器學習的技術或許也可以幫助部分缺乏專業知識和資源有限的公民對於政府決策進行評估，算是一部分的好處。

但學者也提醒此發展將可能存在風險。當中就以演算法工具的不透明性最受到批評，意即許多演算法內部的運作具有非公開的特性，機器在學習的過程，也很難讓人們知道演算法是如何運作的，如何學習的。但Vaithianathan（2017）認為，即便工程師將演算法公開，可能仍無法達到實質「透明」，原因在於演算法的機制之間的交互作用會產生複雜性（從而產生不透明）。且也需避免官僚將疏失歸咎於「系統失靈」（systemic failure）本身（Anderson, 2009; O'Brien, 2012; Leon & Orriols, 2019），而難以達到真正的課責（Guerin, 2018）。

綜合上述，機器學習演算法已逐漸視為當代商業模式的核心（Zuboff, 2015）。演算法改變公共服務的供給形式，儘管諸多文獻談及AI與機器演算能夠為公部門組織帶來低成本、效率與效能提升等優勢，但學者們也認為使用它們存在潛在的風險。在設計與開發的階段，開始有學者提出不同的降低風險的見解，如強化演算法的解釋性、透明度與輔助客觀性等。

二、 演算法的解釋性

首先，機器學習演算在嘗試達到超級智慧的能力之前，更被學者與實務界認為需要解決AI系統背後複雜模型與演算所產生的「黑盒子」（black box）問題。因應而生，科學社群（包括資訊工程）學者紛紛起而發展「解釋性人工智慧」（Explainable Artificial Intelligence, XAI），意即開發能夠解釋與詮釋學習模型的AI。當大量資料被運用來不斷訓練複雜系統時，雖然逐漸增加了預測的能力（predictive power），但卻可能也失去解釋如何得出結果之內部機制的能力。

從社會科學角度觀之，這樣的解釋性是相當必須的。Oswald（2018）的文章從行政法（規則）中的「自然正義」（natural justice）概念出發，談及演算法需承擔合理性和適用性的責任，以及程序公正的必要性，主張公共組織利用演算法做決策時必須受到程序和知識的控制。該作者認為演算法納入決策的過程中也容易產生「懷疑真實性」的風險，演算法不可

能用較低的標準來提供藉口或辯護。當人們使用演算法時，我們可能無法理解演算法產生出結果過程中評估的「要點」（因為我們不知道演算法是如何運作），而這樣的不透明可能就因此會妨礙到他人的權利。

然而Oswald（2018）也認為我們對演算法要求更高標準的期望需要被調整。演算法輔助的優點在於，我們可以在考慮所有相關因素，而不考慮任何不相關的因素下，追求所有案件的處理方式，只要條件一致，就應達到相同的效果，以達到消除偏見的目的。但近來機器學習演算法許多流派，即使是傳統監督式的分類演算，也可能出現黑盒子模式，以機器所判斷出的隨機機率預測，用來作為人類決策依據。但學者們也承認，人類大腦永遠無法像機器一樣，不但無法直接對龐大數據資料進行記憶與快速處理，更遑論要排除個人既有文化與社會化的諸多影響。因此從行政法原則上重新探討，我們對演算法輔助決策的要求應非達到解釋性的更高標準，而是希望它更適用於組織結構與人機協作的實務。

換言之，公部門決策者在應用機器學習演算法時，必須要能夠對該機器的預測結果進行解釋，最主要的原因是希望能透過後推（如：逆向工程），理解整體的決策過程是否存在缺陷，或是決策者本身是否會受到該缺陷的影響。Polson & Scott（2018）評論說：機器可以根據他的編碼假設進行預測，但只有人類可以檢查這些假設。Van Kleek（2018）等人也將「解釋演算法」的過程稱為「sense-making」，將其呈現為預測或是風險是一種邏輯的過程。但這並非表示參與決策階段的所有人都必須為此提出解釋，而是應該要根據不同使用者以及演算法所呈現出的結果來提供不同的特性和解釋。

如圖3主要體現出不同行政官僚決策行為的即時性以及對於人民權益影響程度的高低，藉以回應演算法解釋性可能產生缺陷的不同面向原因。Oswald（2018）提出了以下的標準基礎，來說明行政決策行為潛在缺陷存在的可能原因和影響，包括：「決策的即時性」和「決策效果的嚴重性」（也就是對個人自由權利的影響程度）。以情況C為例，在警察處理的家庭暴力中，我們要求的是官員迅速決定是否將暴力者逮捕，在此情形下演算法工具所提供的詳細訊息或風險評估可能對決策過程並不會太有幫助。反觀情況A，觀護人在處理假釋案件時，該個案所對社會安全的危害程度，將是決策過程最為看中的因素，也因此解釋了大部分觀護人的判斷行為。回到與機器學習演算法對公部門決策影響的關係討論上，究竟機器學習演算法的預測，多大程度解釋了人機協作中，組織與個人在決策過程中的結果？除了演算法本身產生的技術性邏輯外，演算法與行政官僚互動過程中，因為業務差異所造成的解釋缺陷，也是值得探索的重點。

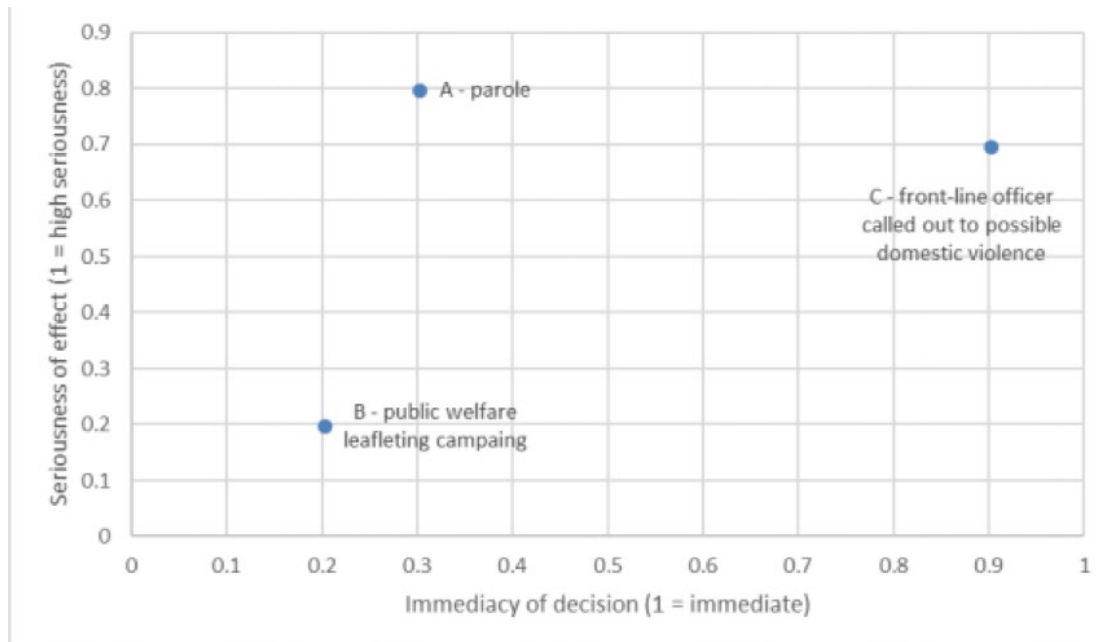


圖 3：兩個因子在不同公共任務上的決策影響

資料來源：Oswald (2018)。

三、演算法的透明度

與可解釋 AI 有某種程度相關，但同樣被學者們重視的透明度，也是演算法能否達到公共性與課責的重要因素。學者 Peeters (2020) 認為演算法透明度不足可能源自以下原因，包括：(1)運作程序受到專利保護。(2)演算法分析人類很多無法處理的資料。(3)演算法因不斷地機器學習作改變，使人類最後難以瞭解。因此，透明度問題可能因無法透徹檢視 (Danaher, 2016)，或受到人類對它認識的限制而受阻礙 (Zarsky, 2011: 293)。就此相關衍生的問題為透明度問題應被視作瞭解演算法的途徑，抑或是充分瞭解演算法運作的形式。隨著決策引用演算法協助逐漸成為公部門趨勢，Janssen、Hartog、Matheus、Ding 與 Kuk (2020) 透過實驗研究檢驗設計三種情境，以瞭解有無演算法協助，對個人決策制定的影響。此三類情境分別為：(a)控制組，沒有任何演算法的協助。(b)有商業專家規範之邏輯 (Business Rules, BR) 的協助，源自於專家系統，透過敘述性的邏輯方式詮釋和複製人類的行為，以模擬公部門實際的決策情形。(c)有機器學習演算法 (Machine Learning, ML) 的協助。在該研究的實驗中，受試者須在不同情境下做出正確的決策，在此 BR 及 ML 可能會提供正確或錯誤的建議，以觀察受試者是否會影響判斷，以及藉此瞭解 BR 和 ML 的限制。在實驗中發現演算法能夠幫助受試者作出正確的決策，對於有工作經驗或受過訓練的情況下，決策者較能察覺演算法是否判斷錯誤。然而，即便是受過良好訓練的人才也未必能偵測出全部演算法所產生的錯誤判斷，該研究認為演算法的透明度越高，則決策品質也會相對提升，另外決策制

定時選擇適當的演算法以及人才的培養都是提升課責性及透明度的關鍵因子。

四、 演算法的輔助客觀性

另一個能降低機器學習演算風險的原則，是強化其輔助客觀性。Bannister 與 Connolly (2020) 認為在建造演算法時，嵌入的資訊中所隱含的人類價值觀點或偏見都有可能造成機器作出判斷時造成歧視等負面影響。因此如何能夠偵測以及將各個風險所分類變成設計過程中不可或缺的環節。作者透過簡單的二維分類法區分運用在公部門中的演算法，並依此建構出一個風險管理框架，主要的框架邏輯在於如何確保演算法的輔助客觀性。第一個維度是將演算法運用的資料特性分為「主觀」(subjective) 或「客觀」(objective) 的資料。雖然我們所稱的「客觀」基本上也是建構在人類創造的準則上，然而仍能與其他僅能透過間接衡量的變項，如幸福、貧窮、忍受程度、趨避風險程度等作出區隔。簡單區分的方式為：在相同環境之下，機器設備針對客觀資料所給予的答案應為一致，而主觀資料則有可能相異；大多數能藉由機器偵測、蒐集的資料多為客觀資料，而主觀資料則包含各種建構、評估、價值判斷、人為預測等較為複雜的面向。然而就客觀資料認定上也存有問題，以非結構式資料（如照片、影片等）為例，演算法系統要如何不帶有偏差地詮釋這些資料（如臉部辨識）將是一大考驗。

其次，第二個維度是將演算法的運作分為「被動」(passive) 以及「主動」(active) 兩種模式。至今大多數的演算法皆為被動式，意即它們主要用於輔助人類制定決策，而並非主動行使行為。而隨著科技的進步，能夠自己執行決策的主動式演算法的應用也正迅速地成長，逐漸變成電子或數位治理 (E-governance or Digital Governance) 中受到關注的應用與議題 (Sharma, 2019; Plantera, 2017)。圖4將上面的分類方式作成二乘二矩陣形式，並進一步針對不同種類的演算法作說明。

	客觀資料 Objective	主觀資料 Subjective
被動式演算法 Passive	OP	SP
主動式演算法 Active	OA	SA

圖 4：以風險框架為基礎的演算法分類

資料來源：翻譯修改自 Bannister & Connolly (2020)。

(一) 客觀／被動 (Objective/ Passive, OP)

在清楚的規則架構下，基於數據、不同特徵作出，得在毋須人力介入之下自動作出判斷。這類的演算法廣泛應用在公部門之中，如退休金的給付，或是在學術研究上如線性規劃 (linear programming) 或西北角法 (North-West corner)。

(二) 客觀／主動 (Objective/ Active, OA)

此類演算法藉由感測系統判斷客觀事實，並採取實際措施，相關例子如飛行器的空速校準，或電梯的體重感測以及測速照相。

(三) 主觀／被動 (Subjective/ Passive, SP)

許多專家系統皆屬此類，透過人類的價值觀點或決策行為作出判斷，然而有時會隱含偏差，如 Google 的搜尋引擎。其他例子如醫療診斷，相較於人類，機器能夠針對特定疾病作出更為準確的診斷 (Budd, 2019)，然而這方面的應用須仰賴醫生過去的診斷發現，以供機器作學習。美國私人醫療保健資訊網站「WebMD」⁹則是較為簡單的應用，使用者能在網站上針對自己的病徵或需求得到客製化的答案。然而隨著這些系統的發展日趨複雜，人類是否將杜絕於機器決策之外將是一大考驗。

(四) 主觀／主動 (Subjective/ Active, SA)

此類演算法能透過學習人類的價值及行為模式而自己作決策，然而機器是否能考量各式複雜的面向以及資料 (如質性資料) 則是一大問題，故在應用的風險最高。相關例子如應徵工作者時需要面試者回答的問答題，或是少年事件處理案件。

從美國犯罪管理分析系統 COMPAS (Correctional Offender Management Profiling for Alternative Sanctions) 的使用案例來看，演算法應用於公共服務上有諸多限制及疑慮，以此訴訟案為例，COMPAS 運作過程的不公開將有損判決過程的透明度 (transparency)，而法官也無法解釋系統實際的演算邏輯。再者，由於機器判斷是基於過去的大量資料所作出的通盤性結果，而難以考量到具體個案，故該評估結果應視作政策或判決上的建議，而最終仍須由人為檢視並作出最終決定，尤其涉及到高複雜性或影響人民權利程度較高的業務，更不能僅仰賴演算法作決策。機器本身仍可能潛藏著人類社會具有的認知偏差 (cognitive bias)，導致分析與判斷上的偏誤 (bias)，進而造成歧視性行為 (discriminative behavior) 的出現。然而即便演算法的使用可能潛藏著許多危機，我們仍不可否認它所帶來的效益。「美國全國經濟研究局」(National Bureau of Economic Research) 曾研究與 COMPAS 功能與使用目的相近的演算法，在不考慮種族的變項

⁹ 美國醫療保健，網站連結：<https://www.webmd.com/>

之下，演算法本身能夠有效降低地方的犯罪率，以及適時糾正法官的錯誤（COMPAS, 2017）。

就此個案的發展與爭議來看，隨著資訊時代的來臨，透過資通訊科技（Information and Communication Technology, ICT）以及AI等應用日漸頻繁，無論是用以提升產能及效率，或是協助運算人類無法處理的複雜資訊，這些科技的使用確實有著改善人類生活品質的潛力，而公部門也能思考如何運用科技以推行品質更為優良的公共服務。然而公共利益所須顧及甚廣，包括：多元價值上的衡量、考慮各利害關係人（stakeholder）的立場觀點、是否符合程序正義、能否達到透明性（transparency）與課責性（accountability）、以及決策本身是否會造成偏見與歧視等，皆是考量的面向，故在科技的使用上應更為謹慎，避免非理性（non-rationality）的行為出現而導致政府失靈（Government Failure）。然儘管文獻研究中，已有諸多論文都點出此挑戰，卻甚少有研究採以實證方法，系統性的探討人機協作過程中，人與機器互動時，對公共服務與業務在流程、決策考量原因、與最終決策的關聯性。而此問題卻是在推動數位轉型時，無論在才能管理與組織數位流程再造時，不可或缺的關鍵因素。

第三節 各國演算法應用之個案採擷

AI隨著數據生成、儲存、分析相關科技的進步，越趨蓬勃發展，各國亦將AI技術設定國家重大發展目標，期透過AI運作，取得勞力成果品質以及整體社會利益的提升，而AI實則為一門以數據資料結合演算法運算的技術，AI系統的產出結果極大部分取決於核心的演算法的設定。誠如前一節演算法對於決策的影響與輔助所述，AI中的技術其演算法的成長為人類所無法掌握的速度前進，更是公部門不可忽視的技術革新。因此，接下來以演算法運用最多的AI新興科技，做更全面的具體探討。

公部門在數位轉型的過程中，初期或許僅係行政機關配合國家數位轉型政策而為，僅應用演算法進行大量資料分析，尚無所謂人工「智慧」的效果；但遠程而言，若能利用演算法的高功能運算能力，能更精準預測或分析國家政策方針，亦有助於服務民眾品質與公共利益之提升。因此，雖然公共治理尚未達以演算法治理的程度，目前多僅於演算法應用的初期階段，但在各國法制政策設計上，均已逐步走向以演算法系統的研發與終極應用作為規劃的方式。

因此，本部分以OECD、以及歐盟等重要國際組織所提出之AI政策原則出發，觀察美國、英國、新加坡等先進國家的政策方向與制度，盤點各國AI政策發展與組織規劃之過程，再從中梳理主要文件所提出AI政策推動之規制方式與重點原則，希望提供我國演算法發展政策上就組織規劃、

規範目的與規制內容之參考；其後，由於演算法之應用涉及資料之蒐集、處理、利用，因此計畫會進一步整理各國重點原則與資料蒐集、處理、利用之關聯，以回應我國公部門若採演算法為應用時，可能發生的問題及其制度調適與路徑分析。

一、國際組織所揭示的 AI 原則

為協助各國政府在制訂AI導入各項公私部門服務時的政策規範，各國國際組織均有提出相關指引甚至法規政策，作為各國進行各項規範時的具體參考依據，亦同時提出進行AI政策規劃時應考量之倫理基礎，可供各國政府制訂具本土化差異之法規時，共同應遵循之準則。本計畫整理經濟合作暨發展組織（OECD）以及歐盟提出的各項指引及法規，此揭準則不區分公部門或私部門之應用，而係針對國家整體的AI策略時應注意或遵循之事項。

（一）經濟合作暨發展組織（OECD）

OECD於2019年5月通過了《AI原則》（AI Principles），設定了實用和靈活的標準，以促進使用具有創新性和可信賴性、並尊重人權和民主價值觀的AI。其中包括有五項基於價值觀的原則（value-based principle）與給予執政者五項建議（recommendations for policy makers）：

1、價值觀原則（value-based principle）

- (1) 原則 1.1：包容性增長、持續的發展和福祉（inclusive growth, sustainable development and well-being）：強調值得信賴的 AI 能為所有人的整體增長和繁榮做出貢獻並推動全球發展的目標。
- (2) 原則 1.2：以人為本的價值觀和公平（human-centred values and fairness）：AI 系統的設計應尊重法治、人權、民主價值觀和多元性，並應有適當的保障措施以確保社會公平公正。
- (3) 原則 1.3：透明性和可解釋性（transparency and explainability）：確保人們瞭解何時會與他們互動並可以挑戰 AI 運算的結果。
- (4) 原則 1.4：穩定性和安全性（robustness, security and safety）AI 系統必須在整個生命週期中以穩健和安全的方式運行，並且應不斷評估和管理潛在風險。
- (5) 原則 1.5：課責性（accountability）：開發、部署或操作 AI 系統的組織和個人，應根據 OECD 基於價值觀的 AI 原則，對其維持系統正常運作負責。

2、給執政者的建議（recommendations for policy makers）

- (1) 原則 2.1：投資 AI 研發（invest in AI R & D）：政府應促進公共和私人對研發的投資，以刺激可信賴 AI 的創新。
- (2) 原則 2.2：為 AI 培育數位生態系統（foster a digital ecosystem for AI）：政府應藉由數位基礎設施和技術，以及數據和知識共享機制，來建立可近的 AI 生態系統。
- (3) 原則 2.3：為 AI 塑造有利的政策環境（shape an enabling policy environment for AI）：政府應該創造一個可供部署值得信賴 AI 的系統政策環境。
- (4) 原則 2.4：培養人才能力並為勞動力市場轉型做準備（build human capacity and preparing for labour market transformation）：政府應該讓人們具備 AI 的技能並支持勞工，以確保順利轉型。
- (5) 原則 2.5：可信賴 AI 的國際合作（foster international co-operation for trustworthy AI）：各國政府應跨界和跨部門合作，共享訊息制定標準，並致力於對 AI 進行負責任的管理。

OECD 並成立 AI 政策觀察站（The OECD AI Policy Observatory），從 2020 年 6 月起藉由各國專家組成的工作小組每月會議，以瞭解這些 AI 原則的導入以及 AI 政策所遇到的挑戰與好的典範，並於 2021 年 6 月出版了對於成員國導入 OECD AI 原則以及執行 AI 政策狀況的報告。在這份報告中，將 AI 政策的生命週期，分為 AI 政策設計、AI 政策實施、AI 政策監測、以及 AI 領域的國際和多方利害關係人合作等四個階段，並依此區分進行評估，參見圖 5。



圖 5：OECD AI 政策的生命週期

資料來源：翻譯修改自 OECD（2021）。

在政策設計階段，各國利用公共和包容性對話來增加對演算法的信賴，並需要跨政府部門的合作，因此可能是透過現有政府機制下的部門、或新設部門來進行，負責演算法政策執行的整合、演算法科技的影響評估、以及解決倫理議題。政策執行階段，則重視各國政府如何將OECD五項給執政者的建議具體落實，包括是否具體投入資金、是否建立開放的資料庫、是否對於現有政策以及法規架構進行配套調整以符合演算法研發到商品化過程的需求、是否已開始進行人才培育等等，也同時評估各會員國是否有相關政策執行的監視機制、以及與各國合作的情形。而為了使AI政策可以順利執行，治理上的組織設定，有些國家是用現有的部門、有些國家新設獨立部門、有的是成立專家諮詢小組、也有是透過一些獨立諮詢和監督機構所提供的意見來進行監管，參見圖6。



圖 6：OECD 各國治理組織設定盤點

資料來源：翻譯修改自OECD（2021）。

在法規架構（regulatory frameworks）部分，各國的政策制定者都瞭解數位轉型相關的監管挑戰，並以多種方式做出回應，從「觀望」到「實驗和學習」，再到徹底禁止數位商業模式等不同強度的處理方式，有些會員國是建立比較軟性的指引（guidance）、有些建立較硬性規定的成文規則。另外有些國家在立法前，透過沙盒實驗（Sandboxes）方式進行政策預評估，以瞭解演算法系統所可能帶來的影響看所需的監理環境，例如現已普遍應用在金融科技（FinTech）體系，也逐漸擴大應用的領域。

（二）歐盟

1、發展

歐盟多年來致力演算法的科研發展與跨領域合作，以提高整個歐盟的競爭力，同時確保發展與合作方向與歐盟價值觀一致。初期歐盟法律並未針對AI制訂具體法律框架，僅透過《一般資料保護條例》（General Data Protection Regulation, GDPR）和《執法指令》（Law Enforcement Directive）之規定給予基本規範。而AI系統之使用者，在「歐洲聯盟基本權利憲章」

（The EU Charter of Fundamental Rights）規範下，亦受平等原則的約束，基於種族、宗教、性別、年齡、殘疾和性取向等應受保護之理由而禁止歧視，歐盟並通過幾項主要的反歧視（平等）指令，這些指令作為歐盟所有成員國製定標準；但是，每個成員國都有責任制定具體的立法來實現這些目標。

自2016年起，歐洲議會就針對具有感知、學習能力（capacity to learn and sense）的機器，進行機器應用涉及民事法律規則之討論（Civil Law Rules on Robotics）。「歐盟執委會」（European Commission）下設之「AI-高階專家工作小組」（Artificial Intelligence-High Level Experts Group，簡稱AI-HLEG），於2019年4月8日發布其所研擬的一份《值得信賴的AI倫理指引》（Ethics Guidelines for Trustworthy AI，簡稱「AI倫理指引」），經採納公眾意見徵詢程序所收集之回饋意見完成修正的AI應用共通性指引，適用於所有受AI設計、開發、部署、實施、使用所影響利益相關者之相關問題，包括（但不限於）：企業、公共服務、政府單位、民間團體、學術機構、研究人員、個人、勞工及消費者等等。

其後歐洲議會於2020年針對相關技術發佈《AI、機器人和相關技術的倫理框架》（Framework of Ethical Aspects of Artificial Intelligence, Robotics and Related Technologies，簡稱「AI倫理框架」），歐盟執委會則於同年底發佈的《AI白皮書》（White Paper on AI，下稱2020 AI白皮書）提出了清楚的願景：「卓越與信任的AI生態系統」，為AI法規提案奠定基本輪廓與架構，主要著重於消費者保護與產品安全，亦回應歐盟2016年起對AI系統的立場；2021年4月進一步再提出《AI法律調和規則草案》（Proposal for a Regulation Laying Down Harmonised Rules on Artificial Intelligence; Artificial Intelligence Act, and Amending Certain Union Legislative Acts，下稱「AI規則草案」），旨在平衡「AI運用帶來的好處」與「AI對個人或社會帶來的潛在負面風險」，提供各會員國對AI監管機制採取協調一致的態度及方法，維護歐洲價值與公民基本權利。

2、重點原則

在歐盟GDPR規範下，針對資料控制者處理個人資料時，提供相關應遵循之規則和原則，以保護資料主體之權利，包含合法性（lawfulness）、透明性（transparency）、公平性（fairness）、正確性（accuracy）、資料最小化（data minimization）、目的和存儲限制（purpose and storage limitation）、保密性和課責（confidentiality and accountability）。

歐盟的2020 AI白皮書，將政策方針主要分成兩區塊：(1)建立優質的生態系統：仰賴跨部門、公私協力的共同行動；(2)值得信賴的生態系統：仰賴特別的規範框架，旨在充分尊重歐洲公民價值觀與權利的同時，得以快速、安全發展演算法技術。由於民眾對於演算法的監理框架，會擔心在

面對演算法決策的訊息不對稱時，造成自己的權利和安全風險，一般公司也擔心法律上的不確定性，需特別重視行為監督、技術安全、隱私保障、資訊透明、不歧視與公平、社會福祉以及責任制度。因此對於未來監理架構，必須從演算法應用的範圍（scope）出發，除了應釐清「資料」與「演算法」這兩個主要元素的意義，因為演算法的學習會受到開發者對於此二元素界定的拘束；也需要釐清所應用領域的法律規範，並且要訂定明確的標準。

白皮書中並整合了歐盟執行委員會於2019年的「AI倫理指引」，作為規範框架的基礎。在AI倫理指引中，AI-HLEG專家直接點明，「可信賴度」（trustworthiness）是在開發、部署及使用AI系統之重要先決條件。「值得信賴的AI」（trustworthy AI）包含3項要素，如圖7：(1)應為「合法的」（lawful）：應遵守所有適用的法律規則；(2)應為「合倫理的」（ethical）：必須尊重倫理原則與價值；(3)應為「穩定的」（robust）：不論從技術面與社會面觀點都應確保AI系統穩定以避免風險。

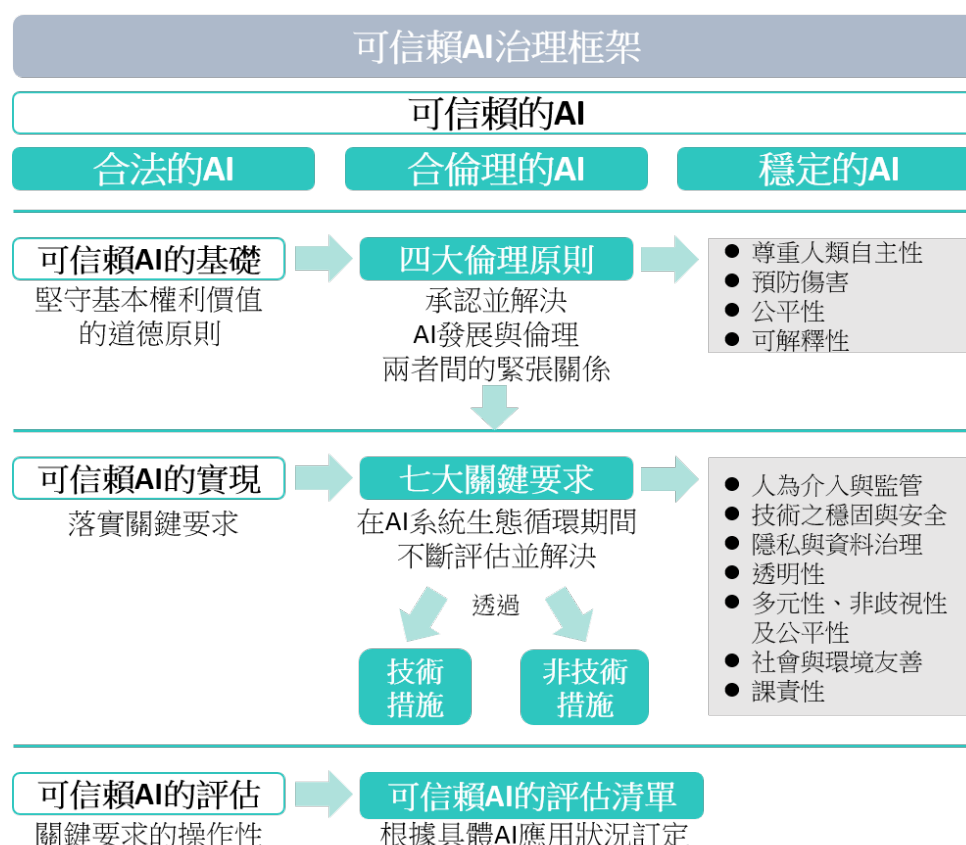


圖7：EU值得信賴的AI治理框架

資料來源：翻譯修改自European Commission AI-HLEG（2019）。

AI倫理指引所提出的3項構成要素的抽象概念下，值得信賴的演算法之基礎必須扣合於基本權利（fundamental rights）保護，從「歐盟基本權利憲章」的基本價值觀中，得到包含尊重自主、避免傷害、公平、可解釋性等四項倫理原則，作為歐盟AI倫理規範架構的4大原則。而要實現值得

信賴的演算法，尚包括七大要求：人為介入與監管（human agency and oversight）、技術穩定與安全（technical robustness and safety）、隱私與資料治理（privacy and data governance）、透明性（transparency）、多元性、非歧視性及公平性（diversity, non-discrimination and fairness）、社會與環境友善（societal and environmental wellbeing）、課責性（accountability）；七項關鍵每一項均同等重要，彼此相互支援，且於AI系統之整個生命週期中應加以落實及評估（如圖8）。

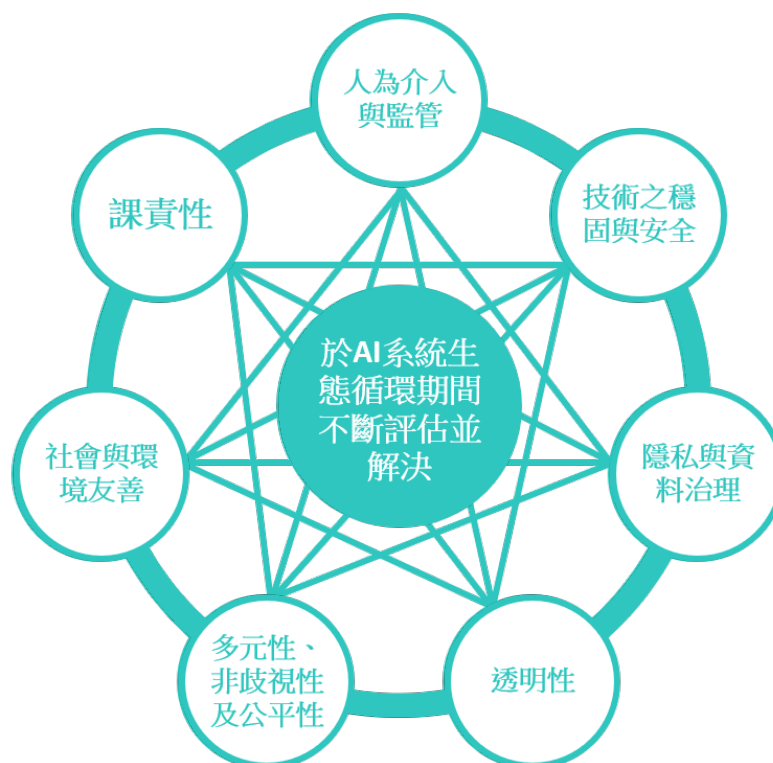


圖8：EU值得信賴的AI七大要件關係圖

資料來源：翻譯修改自European Commission AI-HLEG（2019）。

在AI倫理原則落實方面，AI倫理指引指出於將四大原則轉化為七大要求後，再透過實施合適的「技術性」或「非技術性」方法，將各要求確實導入演算法系統的整體生命週期中。一個具備自我學習能力的演算法系統，不斷接觸外在環境所給予的訊息，進一步適應環境或解決問題，使演算法不斷演進，因而在過程中需要持續不斷的評估與進行調適，才能達到值得信賴的演算法目標，其循環生態如圖9。

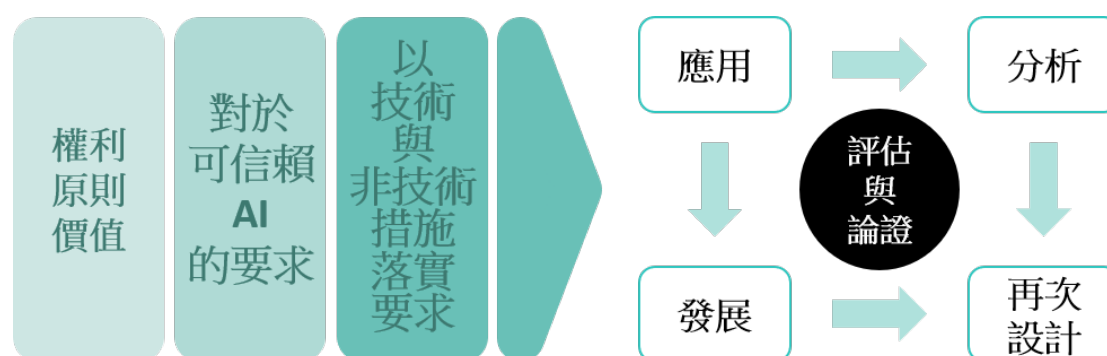


圖 9：EU 落實值得信賴的 AI 的循環生態

資料來源：翻譯修改自European Commission AI-HLEG（2019）。

在2021年「AI規則草案」，是以「風險基礎導向」(risk-based approach) 制定的，透過風險的程度將公共服務進行類型化，並思考分類管制的可能性，希望能避免對新技術發展造成不必要的限制或阻礙。依照AI系統對基本權利或價值所創造的風險程度，予以不同程度的監理方式，區分為三種類型，包括不可承受之風險 (unacceptable risk)：原則上禁止使用此類型演算法系統；高風險 (high risk)：需遵循各種風險管控措施並於進入市場後持續監控；低度風險 (low or minimal risk)：由演算法系統提供者自願建立行為準則 (codes of conduct)。其中針對高風險的服務類型，高風險AI系統必須符合一定之規範以及高密度的管制，才能作為公眾使用，從風險管理、資料管制、技術要求、紀錄溯源與保存、資訊透明、用戶所得理解、人為監督，以及系統準確與穩定性、網路安全等面向，都要能獲得確保。

在「AI規則草案」第六條 (Article 6) 規定，如果AI 系統作為產品的安全管制使用 (safety component)，或是另有法規特別規範的AI產品，便應視為高風險產品。「AI規則草案」特別另以附件 (Annex) 的方式列舉可能的類高風險型領域：

- 1、自然人的生物特徵識別與分類 (Biometric identification and categorisation of natural persons)：
 - (1) AI 系統擬使用於「即時 (real-time)」和「事後 (post)」的遠端自然人生物識別。
- 2、重大基礎設施的管理與營運 (Management and operation of critical infrastructure)：
 - (1) AI 系統擬作為交通管理與營運、以及和水、瓦斯、暖氣和電力的供應管理與營運的安全管制使用。
- 3、教育和職業培訓 (Education and vocational training)：
 - (1) AI 系統擬用於對自然人決定是否進入或如何分配至教育和職

業培訓機構。

- (2) AI 系統擬用於教育和職業培訓機構學生的評分，或用於入學考試評估之目的。

4、人力資源管理（Employment and workers management）：

- (1) AI 系統擬用於招募或挑選求職者。
- (2) AI 系統擬用於決定晉升、終止聘僱、任務分配、評估人員表現。

5、獲得必要私人或公共服務（Essential private services and public services）：

- (1) AI 系統擬用於評估人民取得、減少、撤銷或收回社會福利和救助的資格。
- (2) AI 系統擬用於評估自然人信用或建立其信用評分，但僅於小規模廠商自行使用者除外。
- (3) 調度（dispatch）或建立在緊急服務調度中的優先順序，包括消防和醫療援助。

6、執法行為（Law enforcement）：

- (1) 執法機關意圖使用 AI 系統評估形式犯罪或再犯罪的風險。
- (2) 執法機關意圖使用 AI 系統作為測謊儀和類似檢測情緒狀態之工具。
- (3) 執法機關意圖使用 AI 系統偵測被操控（manipulate）的偽造（deep fake）數據。
- (4) 執法機關意圖使用 AI 系統評估刑事犯罪過程中證據之可信性（reliability）。
- (5) 執法機關意圖使用 AI 系統，依據過去的犯罪行為對自然人進行特徵或人格分析，以預測刑事犯罪實際發生或再犯的可能性。
- (6) 執法機關意圖使用 AI 系統於刑事犯罪偵查、調查或起訴過程中。
- (7) 執法機關意圖使用 AI 系統對自然人進行犯罪分析，並進一步搜索來自不同來源或格式、相關或不相關的複雜大數據資料，以判斷數據中的隱藏關聯。

7、移民、庇護和邊境管制（Migration, asylum and border control）：

- (1) 主管機關意圖使用 AI 系統進行測謊和類似工具、或檢測人的

情緒狀態。

- (2) 主管機關意圖使用 AI 系統，針對將入境或已入境之人進行風險評估，包括安全、非法移民或健康之風險。
- (3) 主管機關意圖使用 AI 系統作為驗證旅行證件的真實性，以及透過文件之安全功能（security features）以發現偽造證件。
- (4) 主管機關意圖使用 AI 系統協助檢視對庇護、簽證和居留許可的申請人資格與申訴。

8、司法和民主程序（Administration of justice and democratic processes）：

- (1) AI 系統擬用於協助司法機關研究、理解事實，並進行法規適用。

另外「AI規則草案」在第10條針對資料治理部分，要求資料控制者於訓練、驗證和測試相關數據資料庫時，應遵循資料治理和管理之規定，包括：(1)相關設計之選擇；(2)資料收集；(3)資料準備之處理與操作，如註解、標誌、清除、濃縮和聚集；(4)相關假設之制訂，特別是應被衡量或具代表性之資料；(5)資料庫之可取得性（availability）、數量和適合性之事先評估；(6)針對可能存在之偏見進行審查；(7)審查任何可能之資料差距或不足與其解決方法。

第13條規範對於資料使用者應充分揭露系統運作訊息的透明性（transparency），尤其在高風險的AI系統，就有關資料控制者應向資料主體揭露之部分，該AI規則草案即規定應確保AI系統之設計和開發足夠透明，應確保資料主體知悉其個人資料正運用於該AI，且針對與資料主體相關、可近性和易於理解之訊息，應以適當之數位格式或其他完整、正確和清晰之形式呈現說明之。這些訊息應包括AI提供者及其授權代表之身份和聯繫方式，以及AI系統之特性、能力和性能限制。其內容包括預期目的、已經過測試和驗證，以及可預期的準確性、穩健性和網絡安全水準。同時亦應包括可能導致健康和基本權利的風險、人工監督措施、為解釋AI系統輸出而採取之技術措施、AI系統之預期壽命，以及任何必要之更新、維護和保養措施，以確保該AI系統得以正常運行之更新。

基於其中對於對基本人權的保護部分，在針對草案內容所為之解釋性備忘錄（explanatory memorandum）第3.5段，特別寫明若演算法技術將對歐盟基本權利憲章規定之基本權利產生不利影響，包括人性尊嚴、生活隱私、個人資料保護、非歧視、兩性平權、表現自由、有效救濟程序之保障，乃至於行政機關應遵守之原理原則、具體工作權、消費者保護等，應予以事前檢測與事後監督，以確保演算法技術應用時所可能產生的錯誤風險與偏差的最小化。對於創新支持措施部分，AI規則草案第五條（Title

V) 亦鼓勵會員國建立演算法監理沙盒 (Regulatory Sandbox) 機制，使創新演算法系統進入市場之前，能於可控環境、時間中，依明確計畫進行開發、測試與驗證。

二、 個別國家或區域演算法規範發展

(一) 美國

1、 發展

歐巴馬政府分別在2014年及2015年，因應新興科技成立之科技政策辦公室 (Office of Science and Technology Policy) 發布兩份關於大數據報告後，至川普執政，2016年5月再度針對大數據與演算法發佈相關報告。在這份報告中，白宮強調「大數據是客觀的」這個假說前提是錯誤的，並提醒勿因為過於依賴演算法，而忽略了其所附帶的歧視風險，有兩大原因：(1)若演算法資料選取偏差或不正確、不完整，本身就是導致歧視的原因；(2)若演算法設計不佳、個人化的推薦服務限縮用戶自主、決策系統的錯誤假設導致歧視、資料蒐集不足以代表特定人口。對於沒有直接參與大規模數據系統算法技術開發的人來說，這系統的最終產品可能會感覺像一個「黑盒子」，一個不透明的機器，它接受輸入、執行一些一般人通常不知道演算法運作的技術過程、再交付來自該過程卻無法解釋的輸出，而這些過程常被操作者視為機密或專有。

在2019年、2020年則分別由川普總統頒布13859、13960號行政命令 (executive order)。川普總統於2019年2月11日所簽署發布以「維持美國在AI方面之領導地位」為題的行政命令第13859號 (Executive Order on Maintaining American Leadership in Artificial Intelligence, 簡稱「13859命令」)，其提出之美國AI倡議 (The American Artificial Intelligence Initiative) 也就是川普政府所提出來的國家AI政策，其中提出五大原則，列出了五個關鍵領域：研發、釋放AI資源、建立AI治理標準、建立AI勞動力，以及國際協作與保護，作為美國維持其在AI研發與部署過程中所涵蓋的科學、技術及經濟等面向領導地位的重要指引；另針對行政部門制訂未來促進AI創新的六大策略目標，內容主要側重美國的科技優勢與國防安全。為了AI政策的執行和政策協調而成立「AI專責委員會」 (Select Committee on Artificial Intelligence)，至於策略性目標的推動除了促進AI研發、提供資料和資源近用、培訓AI世代勞動力、保障美國技術優勢以外，就行政機關對私部門發展AI應用的法規監管原則，也確立必須兼顧促進AI創新應用且可同時保障到人民自由隱私、收集公眾意見以增加信任、制定技術標準以支持可靠的、穩健的與值得信賴的AI技術系統。至於2020年的行政命令第13960號 (Executive Order on Promoting the Use of Trustworthy Artificial Intelligence in the Federal Government, 簡稱「13960命令」) 則延續13859

命令的政策內容，進一步提出適用於國安以外各大領域的一般性準則，以確保AI之用途符合國家的價值觀並有益於公眾利益，並透過跨部門、部會的協力合作，提升並普及AI技術與應用相關知識，以銜接現有法治框架，及培育下一代的AI研發人才。

2020年1月白宮科技政策辦公室，依據2019年2月由川普總統所簽發第13859號行政命令，與美國行政管理和預算局（Office of Management and Budget, OMB）、國家政策委員會（Domestic Policy Council）、國家經濟委員會（National Economic Council）共同以備忘錄形式公告《AI應用的規範指引》（Guidance for Regulation of Artificial Intelligence Applications），提出十點政府機構規範私部門對AI技術應用的法規或非法規管理方式的原則供相關聯邦機構參考。另外，政府機構於評估特定AI應用後，若確認現有的法規已經足夠、或建立新法規的利益無法合理化其成本，該機構可考慮不採取任何措施、或採取非法規的監理方式，例如發布特定產業之政策指引或規範架構、或提供先導型或相關實驗性計畫（pilot programs and experiments）、或由相關產業自願性建立共同標準（voluntary consensus standards）。

至於在立法部分，2019年底眾議院通過的《AI政府法案》（Artificial Intelligence Government Act），要求聯邦政府各部門在系統運行中盡可能使用AI，並設立能建議和促進聯邦政府開發AI的創新用途的單位，但此法案尚未獲參議院通過。2020年3月12日眾議院又提出《國家AI倡議法案》（National Artificial Intelligence Initiative Act of 2020），是川普政府時代的兩項行政命令的倡議後第一個化為具體的AI法案，建立國家AI倡議，以促進AI研究和機構間合作，並製定AI最佳實踐和標準，以確保美國在負責任的AI發展方面發揮領導作用；並由白宮科技政策辦公室（Office of Science and Technology Policy）和國家科學基金會（National Science Foundation）的政府官員，與來自學術界、政府和行業的12名成員，於2021年組成國家人工智慧研究資源工作組（National Artificial Intelligence Research Resource Task Force）。有關美國針對AI規劃之國家政策里程請參見圖10。

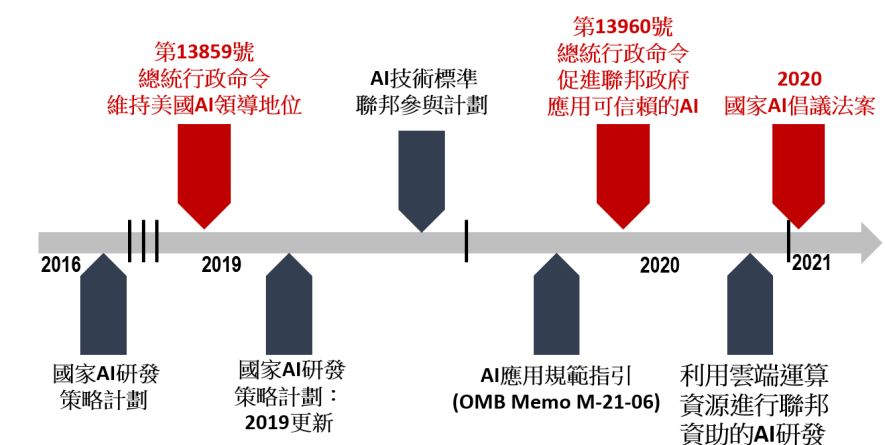


圖 10：美國針對 AI 規劃之國家政策里程碑

資料來源：翻譯修改自 OECD（2021）。

2、重點原則

2020 年 13960 命令中，提出了在政府中使用 AI 的原則。在聯邦政府中設計、開發、獲取和使用 AI 時，各機構應遵守以下原則：

- (1) 合法並尊重國家的價值觀 (lawful and respectful of our nation's values)：尤其重視隱私、人權與人民的自由。
- (2) 有目的性的和表現導向 (purposeful and performance-driven)：風險必須大於利益。
- (3) AI 應用必須「準確、可靠、有效」(accurate, reliable, and effective)。
- (4) AI 應用必須「安全、可靠且有彈性」(safe, secure, and resilient)。
- (5) 可以理解的 (understandable)：AI 應用的操作和結果必須可以理解。
- (6) 負責任和可追溯 (responsible and traceable)。
- (7) 定期監測 (regularly monitored)。
- (8) 透明 (transparent)：相關資訊必須揭露。
- (9) 可課責的 (accountable)：對於 AI 應用的正確使用和運作。

2020 年備忘錄形式的《AI 應用的規範指引》，提出了下列十點對於任何機關決定是否或如何使用 AI 時，制訂一般性或特定領域之「法規」(regulatory) 或「非法規」(non-regulatory) 監理措施時需注意的原則，並呼籲美國應領導國際社會制定有關 AI 技術的標準。雖然這份文件僅針對政府機構規範私部門對 AI 技術應用的法規或非法規管理方式加以要求，但其促進創新的思維仍可窺見，此揭原則對於公部門導入演算法治理時，亦有其參考之價值。

- (1) 公眾信任 (public trust in AI): 應藉由促進可靠、穩健和值得信賴的 AI 應用，來建立公眾對 AI 的信任。依據風險性質及避險工具，適當監管來保護個人隱私，並確保 AI 不會損害個人知情決定的權利。
- (2) 公眾參與 (public participation): 根據第 13563 號行政命令「改進監管和監管審查」，應納入公眾參與機制，尤其在 AI 可能涉及個人資訊使用的情形，有助於政府之責任減輕及改善監理成果，提升公眾信任。
- (3) 科學完整性和訊息品質 (scientific integrity and information quality): 政府機構對 AI 應用所採取之「法規」與「非法規」監理措施，在維持科學誠信的前提下，使政策決定及程序透明公開。
- (4) 風險評估和管理 (risk assessment and management): 監管措施應與風險評估結果一致，並以跨機構及跨技術領域之風險管理為基礎。
- (5) 收益與成本 (benefits and costs): 在考慮與 AI 應用程序開發和法規建置時，政府機構應在合法情況下仔細考慮全部社會成本、收益和資源分配的影響。
- (6) 使用彈性 (flexibility): 應採取績效導向、技術中立的靈活方法，且注意不會損害創新。
- (7) 公平與不歧視 (fairness and non-discrimination): 在一般情況下，AI 應用具可降低因人類所為主觀判斷所造成的歧視，卻也有可能將現實的偏見導入，而產生具歧視或偏頗的結果。
- (8) 揭露與透明度 (disclosure and transparency): 在考慮額外的揭露和透明措施之前，政府機構應仔細衡量現有或不斷變化的法律、政策和監管環境的合適性，並視具體情況判斷可帶來之潛在利益及危害。
- (9) 資料安全 (safety and security): 評估或導入 AI 政策時，政府應促進安全、可靠、按預期運行的 AI 系統的開發，並鼓勵在整個 AI 設計、開發、布建和運作過程中仔細考量安全問題。
- (10) 機構間協調 (interagency coordination): 各政府機構間彼此應協調、分享與 AI 政策之相關經驗，以確保在法規監理上的一致性及其可預測性。

從美國政策與立法的演進，可知美國一直在尋求開發AI的戰略和政策，對於保護公眾免受AI技術潛在風險、與鼓勵積極創新和競爭力之間取

得平衡。著重基於風險分級、顧及成本效益的AI監理方法，並在可能的情況下優先考慮寬鬆的非監理方法，並突破現有的監理障礙。2020年的13960命令要求聯邦政府各部門先行盤點AI使用的案例清單，作為後續政策制訂的基礎。

3、公部門應用個案研析：美國矯正犯人替代處罰管理分析COMPAS

美國「矯正犯人替代處罰管理分析」(Correctional Offender Management Profiling for Alternative Sanctions, COMPAS)，係由管理顧問公司 Northpointe, Inc.於 1998 年所創，透過機器學習演算法用以輔助美國司法部門作被告的再犯風險評估。目前美國數個地方司法體系已利用 COMPAS 作為法官判決的輔助工具，分別有紐約州(2007 年開始使用)、威斯康辛州(2012 年開始使用)、加州及佛州的布勞沃德郡(2008 年開始使用)。

為了能有效評估被告的再犯風險，COMPAS 由風險／需求評估、犯罪個體決定追蹤、治療紀錄追蹤、個體行為檢視、個體態度以及程式操作檢視等模組組成，其所設計的風險評估表共有 137 個問題，而這些問題能分別從被告填寫或過去的犯罪紀錄中獲得答案。風險評估表的問題內包含關於被告的身世背景，如被告的原生家庭與人際關係、犯罪經歷、或是被告原先的工作、休閒娛樂等。此外，在評估表中也透過李克特量表 (Likert scale) 要求被告瞭解被告身心狀態以及處世態度，相關問題如「在別人眼中，我看起來是冷漠且無情的」或「當人們犯輕微的罪或吸毒時，除了會傷害自己，並不會對社會造成危害」等。各項評估將對被告作評分，量測指數為 1 至 10 分，1 至 4 分為「低度風險」；5 至 7 分為「中度風險」；而 8 至 10 分為「高度風險」。COMPAS 能作的評估有三，以下作說明：

- (1) 假釋風險評估 (pretrial release risk scale)：用以評估個體獲假釋後違反報到規定或再犯的可能。
- (2) 再犯評估 (general recidivism scale)：根據被告過往犯罪、吸毒以及少年案件等記錄，評估釋放後再次犯罪的可能。
- (3) 暴力犯罪評估 (violent recidivism scale)：根據被告暴力行為紀錄，以及違法、工作或教育經驗以及初犯年齡，評估釋放後個體發生暴力犯罪行為的可能。

然而 COMPAS 演算法並非完美無瑕，相對地該系統被詬病對於不同膚色、族群的評估具有歧視性的對待，以下分別對 COMPAS 的爭議經典案例以及調查性非營利組織 Pro Publica 的報導作介紹。

(1) 經典案例：State v. Loomis

2013 年，Eric Loomis 因涉入飛車槍擊案，經 COMPAS 分析 Loomis 的過去犯罪紀錄以及風險評估調查問卷，認定該名嫌犯為再犯的高風險群，故而被判六年有期徒刑。Loomis 不服判決而上訴至威斯康辛州（Wisconsin）最高法院，認定法官仰賴 COMPAS 作判決有瑕疵，然而該上訴案為最高法院於 2017 年駁回。

a. Loomis 觀點

Loomis 不服法院判決的具體理由如下：

- (a) 違反量刑個人化：司法機關作量刑應針對個案的具體情節作判斷，然而 COMPAS 的分析方式為比對個人資料與歷史資料後作出判斷，僅能作出通盤性的判斷，而無法具體衡量個體的再犯可能。
- (b) 不符正當法律程序：法官應依據明確資訊（accurate information）認定嫌犯是否罪，然而如同許多演算法的缺點，COMPAS 的運作機制像是黑箱（black box），僅能檢視機器所判斷的結果，而無法詳知運作的過程。
- (c) 以性別作為判斷依據：Loomis 認為 COMPAS 已針對不同性別作出專屬的量表，故在分析時不應再將性別納入考量中。

b. Wisconsin 最高法院觀點

最高法院駁回上訴的理由則有三：其一為 COMPAS 所使用的資訊皆依據被告以及過去歷史資料提供，在審判的過程中若有任何不符或須解釋處應於當時提出；其二為 COMPAS 並非法官裁量的唯一基準，除了參考分析結果之外，法官尚會考慮個案的具體作為，及援引其他資料作判決；其三，法院認為考量性別因素並未產生任何歧視行為，再者 Loomis 無法舉證法院有基於此變項作判決。

(2) Pro Publica 的調查

此外，非營利組織 Pro Publica 於 2016 年發布有關 COMPAS 的調查報導 “How We Analyzed the COMPAS Recidivism Algorithm”。該組織花費兩年分析超過一萬名布勞沃德郡的被告，發現該系統所作出的判決對非裔族群較為不利。以下為 Pro Publica 的具體發現：

a. 錯判各族群之再犯風險機率不一

同樣為兩年內未再度犯案的被告，非裔族群被系統錯認再犯風險較高的機率幾乎為白人的兩倍（非裔族群：45%，白人 23%）；而白人被錯認再犯風險較低的機率遠高於非裔族群（非裔族群：28%，白人 48%。而

即便控制犯罪紀錄、未來再犯率、年紀、性別等變項，非裔族群仍比白人更容易被判定為再犯風險高的族群。

b. 錯判各族群之暴力犯罪機率不一

非裔族群被錯判再犯暴力事件的機率也比白人高兩倍，而白人被低估的機率比非裔族群多上63%。即便控制犯罪紀錄、未來再犯率、年紀、性別等變項，非裔族群被認定為再犯暴力事件高風險群的機率仍比白人高77%。

無論是再犯風險機率或是暴力犯罪機率的錯判，皆可發現非裔族群被認定為高度風險族群的機率遠大於白人，而認定COMPAS就族群的變項上同時發生了型一及型二錯誤，且造成了不必要的歧視行為，讓非裔族群成為系統偏誤底下的受害者，使COMPAS系統的公平性受到了質疑。

(二) 英國

1、發展

英國政府科學辦公室（Government Office for Science, UK）於2016年提出一份關於演算法的報告，認為政府應透過政策，提供演算法相關產業發展及推展新技術的有利環境，對產品安全性、隱私權保護和從業人員進行倫理教育並制定完善的管理制度，以提高民眾對演算法產品的信任、保護民眾的個人隱私、降低科技應用的可能風險。報告也建議政府應善用演算法提高公共服務效率，使資源達到最佳化分配。不過，政府將演算法應用在行政管理和服務時，必須確定技術使用是否善盡相關義務，避免危害民眾權益，具體措施包括採取嚴密的個人資料保護措施、提高執法人員的科技倫理教育、透過部分技術內容的公開降低演算法偏差可能導致的歧視、以及設計明確的課責制度。

2017年11月英國商業、能源及產業策略部（Department for Business, Energy and Industrial Strategy）發佈《產業策略白皮書》（Industrial strategy white paper），將AI列為未來產業發展的四大挑戰之一，並於2018年由英國商業、能源及產業策略部以及數位、文化、媒體暨體育部（Department for Digital, Culture, Media & Sport, DCMS）共同發佈《AI產業協議》（AI Sector Deal），並同時成立了三個組織：AI委員會（AI Council）、AI辦公室（Office for AI）、數據倫理與創新中心（Centre for Data Ethics and Innovation），以尋求企業、科研以及政府的三方合作；官方指引文件則分別發佈《AI路線圖》（AI Roadmap）與《公部門AI應用指南》（A Guide to Using Artificial Intelligence in the Public Sector）。

2、重點原則

英國 AI 委員會（AI Council）於 2021 年 1 月提出的獨立報告－《AI 路線圖》，回應業界、學界與社會關注。AI 路線圖將發展重點分成四個大項：研究、開發與新創（research, development and innovation）、技能和多樣性（skill and diversity）、數據、基礎設施和公共信任（data, infrastructure and public trust）、全國跨部門採用（national, cross-sector adoption）；而在《公部門 AI 應用指南》中，提及 AI 應用的相關計畫，應考慮包括與 AI 相涉之倫理與安全，例如數據品質（data quality）、公平性（fairness）、課責性（accountability）、隱私（privacy）、解釋性與透明性（explainability and transparency）以及系統建立與維護成本（cost）。

3、具體監理架構

英國資訊專員辦公室（Information Commissioner's Office, ICO）於 2019 年 3 月，提出《人工智慧稽核架構》（AI Auditing Framework）（ICO, 2019），並以「治理與課責（governance and accountability）」與「AI 特定風險領域（AI-specific risk area）」為架構的核心要素。前者是指機構有責任採取措施以符合資料保護要求；後者則針對部分 AI 特定領域可能出現的潛在資料保護風險，以及為了管理這些風險的適當措施（ICO, 2019）。由於 AI 之應用會增加資料保護風險，且由於 AI 應用對資料提供者也可能帶來損害，因此使用資料者須重新評估現有的治理機制與風險管控是否有效。ICO 所提出的監理架構圖 11 如下：

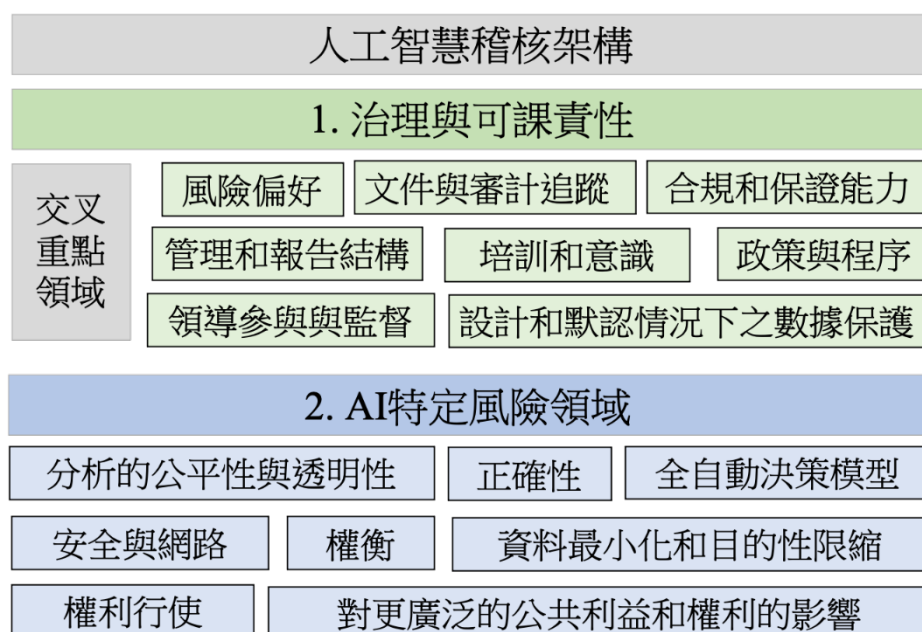


圖11：ICO監理架構圖

資料來源：翻譯修改自ICO（2019）。

此架構深受歐盟個資法（GDPR）的影響，尤其是對於可課責性此一重要原則。因此 ICO 在此架構中，提出了此架構所考量 AI 的八個特定風險領域（AI-Specific risk areas）：

- (1) 分析（profiling）的公平性與透明性：包括偏見和歧視問題、AI 應用的可理解性、以及 AI 決策對資料主體進行解釋。
- (2) 正確性（accuracy）：需涵蓋 AI 應用程序中所使用的資料以及衍生資料。
- (3) 全自動決策模型（fully automated decision making models）：包括基於不同人為介入程度的 AI 決策，以及對全自動決策模型的人工審查（human review）問題。
- (4) 安全與網路（security and cyber）：包括測試和驗證的挑戰、外包風險和重新識別（re-identification）的風險。
- (5) 權衡（trade-offs）：於優化 AI 模型時，應衡平不同限制要求（constraints）之挑戰（例如正確性與隱私）。
- (6) 資料最小化和目的性限縮。
- (7) 權利行使（exercising of rights）：包括個人的被遺忘權（right to be forgotten）、資料可攜性（portability）和個人資料的可取得權利。
- (8) 對更廣泛的公共利益和權利的影響：例如自由或結社、言論自由等（ICO 並強調此框架主要考量與資料保護立法有關之公共利益與權利，不代表更廣泛的公共政策目標）。

4、公部門應用個案研析：英國健康保險的健康資料使用 NHS Digital

(1) 英國健康資料收集使用

英國國家保健服務（National Health Service, NHS）是英國的公費醫療保險體系，相當於我國的全民健康保險，也是全球最大的單一保險人制度醫療體系。這一體系主要經費來自於稅收，並隸屬於英國衛生部的行政管轄，形成英國龐大的公醫體系，其醫療服務包含初級照護（General practice, GP）、住院治療、長期護理、眼科和牙科的服務，且大不列顛各成員國各有自己的 NHS 體系。英格蘭於 2012 年 3 月 27 日通過衛生和社會照護法（The Health and Social Care Act 2012），並成立衛生與社會照護資訊中心（The Health and Social Care Information Centre, HSCIC），作為醫療健康資料的專責機構，大幅影響過去病歷資料之蒐集、分享和分析方式。根據衛生和社會照護法之規定，HSCIC 若受到衛生部長（Secretary of State for Health）之指示，或照護品質委員會（Care Quality Commission,

CQC)、英國國家健康與臨床卓越研究院(National Institute for Health and Clinical Excellence, NICE)、醫院監管機構 Monitor 等命令要求之特定情況下，無需尋求病患同意，得直接從家庭醫師(GP Practice)端獲得病患個人機密資料(Personal Confidential Data, PCD)。2013年3月HSCIC獲NHS授權，於6月開始執行的care.data服務，藉由定期蒐集醫療照護過程之相關資料，針對病患所為之健康和社會照護資訊(例如住院、門診、意外事故和緊急救護記錄)進行延續性連結，以提供即時、正確的NHS治療和照護資訊給民眾、門診醫師和相關部門，進而達成支援病患為治療選擇、加強顧客服務、資訊透明性、優化成果產出、增加課責性以及驅動經濟成長等六項care.data之目標(資策會，2013)。然由於care.data係以一般科醫療收集服務(General Practice Extraction Service, GPES)為基礎，所收集之資料包括家族歷史、接種疫苗、醫師診斷、轉診記錄、生理指標及所有NHS處方。care.data在進行初級和次級資料連結時，會透過NHS號碼、生日、性別和郵遞區號，將這四項可識別資料之進行比對。因此，縱使care.data在涉及敏感性資料時會加以排除，但部分醫師、隱私專家和的社會團體仍對此項服務提出質疑，質疑care.data是否有充分告知病人、HSCIC所宣稱的匿名性是否足夠、此項服務對醫病關係的衝擊、該服務所宣稱的資料分享退出機制(opt-out)並未妥善等。

2016年7月HSCIC更名為NHS Digital，同時取消了飽受個人健康資料共享爭議之cara.data計畫。NHS Digital為英國醫療保健資料之全國託管機構，負責整個醫療和社會保健系統之資料和資訊的標準化、收集、分析、發布及共享。透過家庭醫師(GP Practice)收集的患者資料來支持各種研究和分析，以幫助運行和改進衛生和護理服務，NHS Digital除了取代GPES已經完成的工作之外，家庭醫學資料之規劃和研究服務亦有助於支持衛生和護理服務的規劃和調適、衛生和護理政策的制定、公共衛生監測和干預措施，並支持許多不同的研究領域，例如COVID-19對人群的長期影響、分析醫療保健不平等，以及研究和開發治療嚴重疾病的方法(NHS Digital, 2020)。

NHS Digital可能會從家庭醫療記錄中共享以下資料：

- a. 收集在英國家庭醫師診所註冊之所有患者，包括兒童和成人
- b. 在收集開始後死亡，且之前在英國家庭醫師診所註冊的患者

NHS Digital不會收集患者的姓名或地址，或任何其他可以直接識別患者身份的資料(例如NHS編號、出生日期)，可識別之資料都將由去識別軟體所生成之代碼替代，再與NHS Digital共享資料，此過程稱為假名化(pseudonymisation)。NHS Digital從患者家庭醫療記錄中收集結構化和編碼資料，其內容包括：

- a. 診斷、症狀、觀察結果、測試結果、藥物、過敏、免疫接種、轉診、回診和掛號資料，也包括有關身體、心理和生殖健康等資料。
- b. 關於性別、種族和性取向之資料。
- c. 有關治療過患者的工作人員之資料。

而不包括：

- a. 姓名和地址（除郵政編碼以特殊代碼形式保護外）。
- b. 書面筆記（例如與醫生和護士交談的細節）。
- c. 圖像、信件和文件。
- d. 年代久遠而不需要之資料（例如超過 10 年之藥物、回診和掛號紀錄）。
- e. 法律不允許家庭醫學共享之資料（例如有關人工受孕之治療或性別重置手術之訊息）。

(2) 英國健康保險資料使用之權利保障與行使

2017 年英國資訊專員辦公室 (Information Commissioner's Office, ICO) 認定 Google 旗下人工智慧部門 DeepMind 與英國國家醫療服務體系 (NHS) 的資料共享協議違反英國資料保護法。其後英國衛生部 (Department of Health and Social Care) 於 2018 年 5 月修正施行新「國家資料退出指令」(National data opt-out Direction 2018)，健康與社會照護相關機構得參考國家保健服務體系 (NHS) 10 月公布之國家資料退出操作政策指導文件 (National Data Opt-out Operational Policy Guidance Document) 規劃病患退出權行使機制。

該指導文件主要在闡釋英國病患退出權行使之整體政策，以及具體落實建議作法，主要為退出因應措施、退出權行使以及例外得限制退出權情形 (資策會，2018)。若未來病患表示欲退出國家資料共享者，相關機構應配合完整移除資料，並不得保留重新識別 (de-identify) 可能性。又因指令不溯及既往適用，修正施行前已合法處理提供共享之資料，不因此中止或需另行進行去識別化等資料二次處理。此外，病患得行使其退出權，且於退出後得重新加入國家資料共享體系；惟退出權之行使，係採整體性行使，即病患退出後不得選擇部分加入 (例如僅同意特定臨床試驗之資料共享)。

因此，若患者不希望出於自身護理以外之目的經 NHS Digital 與其他機構共享自身可識別資料，得透過線上註冊 Type 1 Opt-out、National Data Opt-out 或兩項皆可之方式行使其退出權 (NHS Digital, 2021)。無論選擇

哪種方式，患者之個人照護都不會受到影響。而在有正當法律理由時，NHS Digital 也能夠將該去識別軟體的唯一代碼轉換回得以在某些情況下直接識別患者的資料，意即病患之退出權行使仍非毫無限制，當基於當事人同意（consent）、傳染病防治（communicable disease and risks to public health）、重大公共利益（overriding public interest）、法定義務或配合司法調查（information required by law or court order）等四種情形之一者，健康與社會照護相關機構得例外限制之。

另一方面，根據資料保護法（Data protection law），NHS Digital 透過「透明度通知」（Transparency Notice）告訴患者他們是如何使用患者的個人資料，包括從其他組織收集，並與其他組織共享之個人資料，以及共享之原因（NHS Digital, 2021）。英國家庭醫學診所保存之家庭醫療記錄中，患者每天的資料透過規劃和研究來改善健康、護理和服務，幫助診所找到更好之治療方法，並改善患者護理。因此，NHS Digital 所收集之任何資料將僅用於健康和護理目的，而不會與營銷或保險公司共享。

如同歐盟個資法通過後之潮流，英國 NHS Digital 對於病人健康資料的使用，透過對於病人的充分告知資料收集使用之各項細節、以及當事人退出權的行使，以達到資料使用透明性的要求，並貫徹當事人自主權利的保障。臺灣全民健康保險研究資料庫即曾發生不允許行使退出權之爭議，目前正在司法院大法官會議研議中，預期會做成相關解釋，英國 NHS Digital 的運作，對於公部門利用各項既有資料前應有之作為，應可借鏡。

（三）新加坡

1、發展

新加坡在 2019 年底宣布推動國家 AI 策略（National Artificial Intelligence Strategy），並成立智慧國家與數位政府辦公室（Smart Nation and Digital Government Office, SNDGO），鎖定九大重點目標並設定滾動式佈署（AI deployment loop）以推動國家 AI 技術之發展；另一方面，新加坡個人資料保護委員會（The Personal Data Protection Commission, PDPC）發佈《AI 模型治理框架》（Model AI Governance Framework）因應 AI 應用所帶來的挑戰（PDPC, 2018）。其中滾動式佈署，新加坡已透過下列五項計畫在各種公共服務領域實際應用 AI 技術：

- （1）智慧物流計畫（Intelligent Freight Planning）
- （2）無縫高效市政服務（Seamless and Efficient Municipal Services）
- （3）慢性疾病預設與管理（Chronic Disease Prediction and Management）
- （4）客製化教育計畫（Personalized Education through Adaptive

Learning and Assessment)

(5) 通關計畫 (Border Clearance Operations)

2、重點原則

新加坡在監理上尤其重視個人資料的保護。其個資保護委員會(PDPC)在 2018 年所提出的報告指出, AI 的治理框架應納入下列兩項原則: (1)任何由 AI 所作成或協助作成之決策, 應具備可解釋性、透明性與公平性三項特質; 及(2)透過 AI 作成之決策, 應以人類為中心, 並遵守行善(beneficence)及不傷害(do no harm)原則。上開原則其實也是在生命倫理中所要求的基本原則。

2019 年底推動之國家 AI 策略(National Artificial Intelligence Strategy), 在願景與路徑(vision and approach)中, 分別以「AI 為服務人民而發展」、「積極面對 AI 應用造成之風險與治理問題」以及「與 AI 共利共榮」三個角度詮釋 AI 技術應以人為本(human-centric)。

其後 PDPC 發佈《AI 模型治理框架》(Model AI Governance Framework), 舉出兩項最高指導原則(Guiding Principle), 以促進 AI 技術的可信賴度以及應用時的理解:

- (1) 應用 AI 輔助決策之組織, 應確保決策過程具可解釋性(Explainable)、透明(Transparent)和公平(Fair)。
- (2) AI 的解決方案應以人為本(Human-centric), 首要考慮人民的利益、福祉與安全。

另於此治理框架的附件 A 彙整了 AI 的基本倫理原則, 供應用 AI 技術之組織於相關或必要情況下, 除了上述最高指導原則可斟酌參考外, 包括課責性(accountability)、正確性(accuracy)、可驗證性(auditability)、可解釋性(explainability)、公平性(fairness)、以人及其福祉為本(human centrality and well-being)、人權保障一致性(human rights alignment)、涵容性(inclusivity)、漸進性(progressiveness)、負責任、課責且透明(responsibility, accountability and transparency)、穩定與安全性(robustness and security)、可持續性(sustainability)。另特別強調此框架是不針對特定演算法(algorithm-agnostic)、不針對特定科技(technology-agnostic)、不針對特定部門(sector-agnostic)、以及不針對特定規模和商業模式(scale and business model-agnostic)。

2020 年 1 月, PDPC 又提出了第二版的 AI 模型治理框架(Model Framework), 除再次確立第一版中對於 AI 應用的兩項指導原則外, 並優化了第一版的架構, 提出從原則(Principle)到落實(Practice)的額外考量項目(IMDA & PDPC, 2020):

- (1) 內部治理的結構和措施：組織中有明確的角色和職責；有風險監理的標準作業程序；有完整的人員培訓。
- (2) 確定人類參與 AI 增強決策（AI-Augmented Decision-Making）的程度：AI 中有適當程度的人為參與；盡可能減少對個人造成傷害的風險。
- (3) 營運管理：最小化數據和模型中的偏差；基於風險導向的監測方法，例如可解釋性、穩健性和定期調整（regular tuning）。
- (4) 利害關係人（Stakeholder）的互動與溝通：讓用戶了解 AI 政策；盡可能允許用戶提供反饋；讓溝通更容易理解。

3、公部門應用個案研析：新加坡人工智慧國家策略

(1) 新加坡 AI 五項國家計畫（National AI Projects）

新加坡於 2019 年底宣布推動國家 AI 策略（National Artificial Intelligence Strategy），這是新加坡以國家主導邁向智慧國家的關鍵一步。這項策略揭示新加坡深入使用 AI 技術來進行經濟轉型，且不僅是採用技術，而是要從根本上重新思考商業模式並進行根本改革，以提高生產力並創造新的發展領域。國家 AI 策略中五項計畫之 AI 技術實際應用(SNDGO, 2019)：

a.智能貨運計劃智慧物流計畫（Intelligent Freight Planning）

物流公司透過 AI 規劃貨車到港口集裝箱堆場之路線和調度，減少等待時間和局部交通堵塞，使公司能夠提高業務利潤並增強其業務競爭力。物流公司所收集的數據亦有助於政府以資料進行的政策評估，並進行更好的城市規劃，以提高新加坡物流生態系統的整體效率、創建貨運生態系統，讓關鍵各方利用 AI 共享信息和協作，實現更高效的貨運和卡車運輸。

b.無縫高效市政服務（Seamless and Efficient Municipal Services）

居民透過聊天機器人得直接對各種機關提出反饋或疑問，聊天機器人之 AI 系統引導市民描述問題，並將問題提交予相對應負責之主管機關。居民無需煩惱應向哪個機關提出疑問，或該如何描述其問題。國家人工智慧辦公室（The National AI office）表示，AI 協助機關更瞭解當地居民以及各社區應如何規劃和建設能滿足居民需求之設施。另外亦透過 AI 分析數據來預測房屋問題，以優化房屋維護工作，例如預測建築物中下一次電梯可能故障時間，以提前進行維修、保養工作。

c.慢性疾病預設與管理（Chronic Disease Prediction and Management）

AI 可用於分析臨床數據、醫學圖像、健康行為和基因組資料，建立

個人化風險評估分級，該評估資料可幫助個人採取適當之預防措施，使護理團隊獲得更早、更具特定性之介入措施。目前初步發展眼部病變分析儀（SELENA+），使用深度學習系統分析視網膜照片，以檢測糖尿病眼病、青光眼和年齡相關性黃斑變性之三種主要眼部疾病。SELENA+可以像人類評分者一樣準確和快速地分析視網膜照片，同時該分析視網膜圖像方面之能力也將得到擴展，未來計畫將用於開發心血管疾病之預測風險評估模型。

d.藉由適應性學習和評估進行個別化教育計畫（Personalized Education through Adaptive Learning and Assessment）

2018 年 5 月推出「學生學習空間」（Singapore Student Learning Space, SLS），向全國學校系統中所有學生和教師提供 AI 在線學習平台。透過 AI 自動適應學習系統增強學生對 SLS 之學習體驗，該系統使用機器學習告訴每位學生對學習素材與活動之反應，並為每個學習者推薦一組循序漸進之學習路徑。新加坡政府優先於中小學試辦 AI 英語語言自動評分系統，教師能透過 AI 自動評分系統更有效地評估學生表現，包括評估學生開放式之回答，例如簡答題和論文，並提供快速反饋。

e.通關計畫（Border Clearance Operations）

政府希望加強邊境安全並改善旅客體驗，期能提供所有旅客全自動化出入境檢查。利用 AI 協助移民官員在旅客抵達檢查站前，透過收集和分析包括旅客提交之新加坡入境卡、航空公司提供之旅客信息，預先評估旅客風險狀況，制定分級化安全檢查。另一方面，新加坡政府亦將研究如何重新設計出入境手續，讓所有旅客都能藉由自動化通關設施享受安全、無縫的出入境體驗。

(2) 建構 AI 生態體系（Building the AI Ecosystem）

新加坡政府除了五大國家 AI 重點計畫項目之外，也建立整體 AI 的生態系統（AI Ecosystem），提升整體經濟中 AI 之創新與應用，並確立了五項生態系統推動因素（SNDGO, 2019）：

- a.三重螺旋合作夥伴關係（Triple Helix Partnership）：學界、產業界和政府間的三重螺旋合作夥伴關係，使 AI 解決方案之基礎研究與部署得以快速商業化。
- b.AI 人才培養：培養 AI 之人才和教育以解決與 AI 有關之工作角色需求，並推出工程師認證框架（Framework for artificial

intelligence (AI) engineers)¹⁰。

- c. 資料架構 (Data Architecture)：建立快速且安全之資料架構，使高品質的資料集得以跨越各種部門。
- d. 進步且值得信任的環境：此一環境對於測試、開發和部署 AI 解決方案十分重要。
- e. 跨國合作：與跨國研究人員、企業和政府開展國際合作，推動和支持人工智能的可持續發展。

隨著新加坡快速發展數位經濟與 AI，為能在促進技術和商業創新間取得適當平衡，建立值得信賴的生態系統及治理和監理制度至關重要，更需建立 AI 之責任以取得使用者之信任，同時保護使用者之利益。新加坡透過成立以產業導向之 AI 和資料道德使用諮詢委員會 (Advisory Council on the Ethical Use of AI and Data)，該委員會對於業界開發負責任 AI 及其商業化部署，政府可能將面臨到的政策方向或監理等問題提供相關建議，並針對不同行業和應用環境下所開發之 AI 提供相對應之業務行為準則或專業行為準則、策劃實現可解釋的 AI 之技術解決方案、由新加坡電算協會 (Singapore Computer Society's) 成立之 AI 倫理與治理指導委員會，發展 AI 道德與治理方面之培訓和認證，以管理 AI 解決方案及實施 AI 項目之專家、發布組織評估指南，以評估 AI 治理流程與 AI 模型治理框架之一致性。

(3) 新加坡對於資料保護的作為

新加坡資訊通信媒體發展局 (Info-communications Media Development Authority, IMDA) 及個人資料保護委員會 (PDPC) 亦於 2019 年 6 月宣布推出「可信任資料共享框架」(Trusted Data Sharing Framework)，以促進資料共享的信賴程度。資料之可信任性為數位經濟之基礎，可信任之資料集可為企業帶來更有效的訊息交換、實現共享資料資產，進而支持開發創新之解決方案，為不同市場定制配套服務和流程。對於使用者而言，共享數據之意願同時反映了對企業使用和保護數據能力之信心程度。而雖然企業組織已體認資料的價值，但資料資產共享常面臨以下挑戰 (IMDA and PDPC, 2020)：

- a. 缺乏資料共享之指南、方法與系統化方法；
- b. 與合作夥伴建立信任關係；

¹⁰ 一級認證 AI 工程師之資格，必須在 AI 相關職位上有至少一年工作經驗，且至少部署了一個價值至少 50 萬美元的 AI 項目。第二級和第三級的申請人，則須進一步證明其有能力管理整個組織中的 AI 工程團隊和多款 AI 應用程式之能力。三個級別的認證費用分別為 500 美元、800 美元和 1,000 美元。

- c.確保遵守個人資料保護法（Personal Data Protection Act，PDPA）；
- d.評估資料資產價值；
- e.資料共享可能導致業務競爭力下降或商業機密洩露。

可信任資料共享框架提供通用的「資料共享語言」（data-sharing language），協助企業組織建立一套可信任資料共享夥伴關係之系統化方法，結合目前 PDPC 指南中關於個人資料匿名化與共享內容，新增資料數據共享之資料評估指南、合約範本，透過可信任資料共享框架的指引與監管，利於企業組織在當地生態系及全球為數位經濟提供更好及更具信任性之資料傳輸（IMDA, 2019）。

（4）演算法之審查（Algorithm Audit）

在第二版的 AI 模型治理框架中，針對演算法的審查，提出了需要考量的因素。首先，為確認審查之合法性，必須是具有法律授權的監理機構，於必要之情形下，對於 AI 模型中演算法之實際運作進行審查；或者由 AI 技術提供廠商，於監理單位對其客戶提出合規要求時進行協助。審查報告可能會超出大多數被審查個人和機構的理解範圍，因此進行演算法審查可能需要具有技術專長的外部專家，且演算法審查所需的費用和時間應與從審查報告中獲得的預期收益相衡平。最終，演算法審查通常應該在此審查要能帶來明顯的好處相當明確的情況下才進行。

另外要進行演算法審查應該要考量以下因素：

- a. 進行演算法審查的目的：為促進 AI 的應用，演算法審查之前，建議考慮已經向個人、組織、企業以及監理機構提供的信息是否充分和可信（例如產品或服務描述、系統技術規範、模型培訓和選擇）記錄、數據來源記錄等等。
- b. 審查結果的目標對象：審查者（專家）應確認該審查的目標對象為何，以確認其報告內容所提供的訊息能確實解釋該 AI 的決策過程以及資料有效運用，以達到可受信賴 AI 的目的。如果是監理機關所需的審查報告，應先從資料可信性以及演算法的功能性著手，若對於 AI 的運作所使用訊息的真實性及完整性有所疑義，該審查報告應要能完整呈現 AI 運作的模型。
- c. 一般資料的課責性：受調查對象應能提出使用一般資料課責機制的相關資訊，包括開發階段如何採取保存資料的可溯源紀錄、如何確保資料正確性、如何減少資料中的既有偏差、基於不同目的進行資料分類、而且有定期檢視和更新資料。
- d. AI 模型中的演算法須是具有商業價值、可以影響市場競爭力

的信息：例如，該演算法須是營業秘密（**trade secret**）或可作為營業秘密的商業規則。如果考慮進行技術審核，還應考慮相應的保密措施（例如保密協議）。

三、 各國法規發展共同特色

由各國的資料顯示，目前關於應用演算法所使用的演算法研究與分析，多尚未有單一針對演算法應用管理法規，而是針對應用演算法之 AI 模式，提出指引以及處理原則。明文的法規範部分，目前所注重的仍在於個人資料保護之法規，歐盟的「一般資料保護規則」（**GDPR**）是目前最典型也最嚴格的規範，其所重視的「當事人自主原則」，也影響了各國關於法律制訂的思維。

然而，除了個人資料保護的角度外，也需要考慮透過巨量資料演算分析所帶來對於公共利益的好處。因此從政策形成與制訂而言，考量演算法的益處與挑戰的平衡，需要關注的面向應該需包括：提供公眾參與以增加數據共享的誘因、提高安全性和保護措施、更新演算法的監管方法、評估可接受的風險和做決策的倫理、以及考量其他政策需求的平衡，以避免在鼓勵創新與權利侵害中產生兩難。

也就是說，包括歐盟在內之先進國家政府，對於包含演算法充滿著高度不確定風險的新興科技發展，按例都會對該項科技在倫理（**ethic**）、法律（**legal**）、與社會（**social**）等三個面向上（所謂的 **ELSI**）所引發之相關問題或疑慮進行評估後，依據風險程度再決定下一步治理策略。就各國對於演算法治理之價值原則以及組織規範架構，整理如表 3。

至於個人資料的使用，依據本團隊訪談專家所獲，目前國際的觀念似有從個人化權利的資料保護（**Data Protection**），逐漸過渡到群體式公益權利的資料治理（**Data Governance**），資料使用的個案需求考量所帶來的利益，可能會使個人保護得絕對自主必須有部分退讓。因此，透過設立具有功能的資料倫理委員會（**Data Ethics Committee**）或是資料近用委員會（**Data Access Committee**）的監管，充分納入倫理考量甚至於建立規範的思考；或是給予參與者回饋機制給予滾動式動態同意機會，以及如英國衛生和社會照護法（**The Health and Social Care Act 2012**）賦予退出權（**opt-out**）的設計，都是可能的參考方式。

表 3：各國對於演算法治理之價值原則以及組織規範架構比較表

國家或組織	OECD	歐盟	美國	英國	新加坡
倫理或價值原則	OECD (2019)： 1.原則 1.1:包容性增長、持續的發展和福祉 2.原則 1.2:以人為本的價值觀和公平 3.原則 1.3:透明性和可解釋性 4.原則 1.4:穩健性和安全性 5.原則 1.5：課責性	AI-HLEG (2019)： 值得信賴的 AI 三要素： 1.合法的 (lawful) 2.合倫理的 (ethical) 3.穩定的 (robust) 四大原則： 1.尊重自主 2.避免傷害 3.公平 4.可解釋性	第 13960 命令 (2020)： 1. 合法並尊重國家的價值觀 2. 有目的性的和表現導向 3. AI 應用必須「準確、可靠、有效」 4. AI 應用必須「安全、可靠且有彈性」 5. 可以理解的 6. 負責任和可追溯 7. 定期監測 8. 透明 9. 可信賴的	公部門 AI 應用指南 (2018)： 1. 數據品質 2. 公平性 3. 課責性 4. 隱私 5. 解釋性與透明性 6. 系統建立與維護成本	AI 模型治理框架 (2018)： 1. 課責性 2. 正確性 3. 可驗證性 4. 可解釋性 5. 公平性 6. 以人及其福祉為本 7. 人權保障一致性 8. 包容性 9. 漸進性 10.負責任、課責且透明 11.穩定與安全性 12.可持續性
政策原則	OECD (2019)： 1.原則 2.1：投資 AI 研發	AI-HLEG (2019) 七大關鍵要求： 1.人為介入與監管	OSTP (2020)： 監理 AI 應用的十點原則：	AI Roadmap (2021) 四大重點： 1. 研究、開發與新創	PDPC (2018)： 1. AI 決策應具備可解釋性、透明性

國家或組織	OECD	歐盟	美國	英國	新加坡
	2.原則 2.2：為 AI 培育數位生態系統 3.原則 2.3：為 AI 塑造有利的政策環境 4.原則 2.4：培養人才能力並為勞動力市場轉型做準備 5.原則 2.5：可信賴 AI 的國際合作	2.技術穩固與安全 3.隱私與資料治理 4.透明性 5.多元性、非歧視性及公平性 6.社會與環境友善課責性	1. 公眾信任 2. 公眾參與 3. 科學完整性和信息品質 4. 風險評估和管理 5. 收益與成本 6. 使用彈性 7. 公平與不歧視 8. 揭露與透明度 9. 資料安全 10.機構間協調	2. 技能和多樣性 3. 數據、基礎設施和公共信任 4. 全國跨部門採用	與公平性三項特質。 2. AI 決策，應以人類為中心，並遵守行善及不傷害原則。
政府組織	建議整合協調的組織： 1. 現有的部門 2. 新設獨立部門 各國現有之治理監督的組織類型（如圖 6）： 1. 現有的部門 2. 新設獨立部門 3. 專家諮詢小組		1. AI 專責委員會（Select Committee on AI）。 2. 國家 AI 研究資源工作組（National Artificial Intelligence Research	1. 科學辦公室（Government Office for Science, UK）。 2. 2.AI 委員會（Office for AI）、AI 辦公室（Office for AI）、數據倫理與創新中心	1. 智慧國家與數位政府辦公室（Smart Nation and Digital Government Office, SNDGO）。 2. 個人資料保護委員會（The Personal Data

國家或組織	OECD	歐盟	美國	英國	新加坡
	4. 獨立諮詢監督機構提供意見		Resource Task Force)：含白宮科技政策辦公室 (OSTP) 和國家科學基金會 (NSF) 的官員及產學界專家。	(Centre for Data Ethics and Innovation) 。	Protection Commission, PDPC) 。
規範架構	建議： 1. 軟性指引 (guidance) 2. 硬性規定 (regulation) 3. 可進行沙盒實驗 (Sandboxes)	建議： 依據不同風險程度進行監理： 1.不可承受之風險：原則禁止使用 AI 系統。 2.高風險：需遵循風險管控措施並持續監控。 3.低度風險：自願建立行為準則 4.鼓勵會員國建立 AI 監理沙盒機制	行政命令 (含備忘錄)，目前正在立法中。	1. 產業策略白皮書 2. AI 路線圖 (AI Roadmap) 3. 公部門 AI 應用指南	國家 AI 策略 (National AI Strategy, 2019) 。

資料來源：本計畫。

四、我國現況

(一)我國 AI 整體規範概況

目前臺灣雖有「個人資料保護法」的明文規範，但缺乏實際的監理機關，對於巨量資料的使用並無其他嚴格法規的規範。在生醫研究上藉由事先取得檢體獲資料捐贈者同意，可以合法申請設立「人體生物資料庫」之方式對於檢體與資料進行研究使用；但如中央健康保險署委託國家衛生研究院所推動「全民健康保險研究資料庫」之建置，由於資料蒐集時其目的僅係為經常在保費申報與稽核使用，則後續研究已超出原目的外之使用，迭生爭議。

我國政府將2017年宣誓為「臺灣AI元年」，相繼於同年8月推出「AI科研戰略」，以及於隔年推動4年期的「臺灣AI行動計畫」；科技部與其轄下四個AI創新研究中心共同於2019年9月發佈「AI科研發展指引」，期以「以人為本」、「永續發展」及「多元包容」三大核心價值，指引科研人員於研究發展之際審慎關注相涉議題，以兼顧科研人員學術自由，以及維護人權與普世價值，亦為貫徹從倫理出發的基本思維。然而，此一「AI科研發展指引」是針對研究時所應注意之規範，僅為產品或服務開發之前階段，並非對於產品與服務開始上市或提供後之指引規範。

臺灣目前對於AI服務或產品的演算法規制方式，尚無法律層級的立法或提案。但針對AI所提供之公共服務以及轉型，政府過去有許多委託民間協助制訂標準或指引的計畫，例如國家發展委員會曾就公民網路參與輿情分析委託電子治理研究中心制訂「網路輿情整合式管理機制作業程序」（陳敦源、廖洲棚、黃心怡，2016）；亦曾針對公部門AI數位轉型曾委託電子治理研究中心制訂「數位轉型AI育成手冊」（曾冠球、廖洲棚、黃心怡、陳敦源、何宗武，2019），以作為政府相關政策制訂或作業程序的準則。對於政府針對具體化的AI運用之服務範例，最具代表性的應為「自動化投資理財顧問」，也就是「理財機器人」的規範，由信託暨顧問商業同業公會於2017年制訂「中華民國證券投資信託暨顧問商業同業公會證券投資顧問事業以自動化工具提供證券投資顧問服務（Robo-Advisor）作業要點」（以下簡稱自動投顧作業要點），並報請主管機關金融監督管理委員會（下簡「金管會」）函釋公告，並作為主管機關審查之標準，使得原性質可能僅為民間團體針對內部會員的約束規範，成為具有外部效力之行政規則。

(二)我國關於理財機器人之規範

1、理財機器人之運作特色

理財機器人，亦有稱為「機器人理財顧問」(Robo-Advisor, RA)或「自動化投資理財顧問」(Automated Investment)，是透過大數據分析及演算法為其運作基礎，依據金融市場中交易價格、交易量、財務資訊等各種數據資料，結合投資人自行輸入所得、年齡、投資目標、風險承受度等個人資料，經由數據分析、財務模組與決策樹(decision tree)建立相對中性之投資建議模組，搭配AI之機器學習(Machine Learning)，對投資人之風險屬性進行分類後，以電腦系統提供客戶投資組合建議與投資管理服務(曲建仲，2017¹²；Frankenfield, 2021)。期望藉由機器人的分析，即時調整投資組合之配置，降低投資風險並提升預期報酬率，並達到深入分析消費者行為與偏好，提供最符合消費者需求的財富管理服務之效果(陳安斌等，2018)。

事實上現階段理財機器人之發展規模，多數仍僅在運用AI之機器學習模式執行任務，尚難達到運用深度學習的模式。基於我國金融管制將證券業與投資信託顧問(投信顧)業分業經營，因而無法提供全方位「個人資產管理服務」。目前投信顧業運用理財機器人之方式，多僅為提供投資組合試算或推薦自家金融商品等行銷手段，故我國理財機器人僅能稱為「投資顧問服務自動化」的半自動顧問型機器人服務，並未達到完整以機器取代真人服務。根據金管會整理我國理財機器人之發展類型可概分為三種，其一以基金投資理財服務為主，例如富邦銀行、王道銀行等；其二為針對指數股票型基金(Exchange Traded Fund, ETF)之投資平台，例如元大投信AI 智能投資平台；其三為提供銀行理專挑選基金之量化系統，例如大拇哥投顧、瑞銀理財機器人(王儷玲，2019)。

2、我國現行有關理財機器人之法規

金管會於2017年開放「機器人投資顧問」業務，信託暨顧問商業同業公會於同年6月制訂了「中華民國證券投資信託暨顧問商業同業公會證券投資顧問事業以自動化工具提供證券投資顧問服務作業要點」(自動化投顧作業要點)，為理財機器人目前唯一明文規範。該要點全文共9點，除進一步定義理財機器人，具體規範可分為三大面向：其一為「資訊揭露」，業者應於客戶使用理財機器人前告知注意事項，使客戶確實瞭解相關內容，以及投資組合再平衡之告知；其二「規範金融業者之金融消費者保護行為」，業者應落實瞭解客戶(Know Your Customer, KYC)作業與建議投

¹² 曲建仲(2017)。翻轉人類未來的AI科技：機器學習與深度學習，科技新報，2021年04月40日，取自：<http://technews.tw/2017/10/05/ai-machine-learning-and-deep-learning/>。

資組合，避免利益衝突、確保公平客觀之執行；其三為「業者之內部治理要求」，業者對理財機器人演算法應有監管機制，以及專責委員會之監督責任。其他涉及包含金融消費者保護法在內之法規，對金融消費者之契約公平性、行銷廣告內容真實性、瞭解客戶、適合性原則以及忠實義務與說明義務訂有規範。但此揭規範主要以自然人理專以及業者為主要監管對象，尚非針對理財機器人本身。

較為特殊為，自動化投顧作業要點特別在第三條針對演算法之監管有獨立之條文規範。規範強調演算法（Algorithms）乃自動化投資顧問服務系統之核心，並反映投資顧問業者對市場分析與研究之邏輯，其設計之正確與否，影響運算的結果，攸關客戶權益。故業者對系統所運用之演算法，應有效進行監督與管理。此一對於演算法所扮演角色與性質的闡述，充分說明AI應用服務監管的核心在於演算法，對接受服務者的權益影響甚鉅。因此，自動化投顧作業要點特別要求業者內部對於演算法應設立包括期初審核以及定期審核的監管機制。在期初審核除對於預期結果的成效應予評估外，對於演算法使用之方法論必須確實理解，並掌握系統之假設、偏誤與偏好；而定期審核的重點在於對於系統使用模型於市場情況或經濟條件變化時的適當性評估，並定期測試與原先預期結果之差異；另外在第五條規定若演算法涉及挑選投資組合，便應針對演算法本身進行審核；演算法本身也必須是專責委員會監督以及告知客戶注意事項的主要內容項目。

3、理財機器人發展所面臨之困境

現階段理財機器人雖僅得提供半自動化之投資顧問服務，金管會仍察覺此服務模式尚有不少缺失，主要可分為以下四點（彭禎伶，2020）¹³：

- (1) 沒有訂定演算法模組參數審核之作業流程，以及對參數之調整或確認時，未敘明原因且未提供佐證資料供驗證其合理性；
- (2) 沒有訂定市場發生重大變動的情境與相關因應措施，或即使與客戶約定若「市場發生重大變動時」須提供因應的投資建議，但卻出現未對全數客戶檢視投資建議妥適性並提出建議的情事。
- (3) 對客戶在短期間查填 KYC 問卷內容有不一致或相互矛盾的情況，沒有建立適當處理機制。
- (4) 程式設計上，發現理財機器人對不同風險屬性之投資人，產出相同投資標的及投資比例之情況，顯未確實按照客戶不同

¹³ 彭禎伶（2020）。機器人理財也會出包 金管會糾出四大缺失，中時電子報，2021年10月20日，取自：<https://www.chinatimes.com/newspapers/20200310000253-260202?chdtv>。

風險屬性，提供差異化之投資理財建議。

自上開缺失可知，原先設定透過機器人處理，是希望減少人為主觀判斷造成的偏差；但是極有可能因為演算法程式撰寫之不合理，造成提出建議有妥適性之疑義，因此需要有完整機制進行監督驗證。在我國政策規範上，雖尚無具體法律層級對於AI監理的法規範，但如同理財機器人透過部分委託規範或團體內規，並進一步使其產生外部規制的效力，似為目前比較具體之做法。如何衡平個人權利與公眾利益，透過委員會的倫理審查或監管、個人退出權的賦予、或使用授權者之回饋機制，並進一步要求演算法方法論的透明性揭露，取得公眾信任，應為後續思考之方向。

第四節 人機協作職能模式

一、 人機協作理論基礎

隨著全球化發展，科技的興起與普及，網路的服務越來越便利，使得人們逐漸習慣於一種新型態的活動方式，隨之而來的新形態的工作職能也應有所改變，特別是演算法學習、大數據是第4代工業革命重要的突破，未來組織內部的團隊分工，不再僅由人類完成，還會包括AI的協助，人機協作（Human-robot Collaboration, HRC）現象將會越趨平常（Molitor, 2020）。而這樣的變革，Vogl, Ganesh, Seidelin與Bright等（2020）即曾以「演算型官僚」（algorithmic bureaucracy）替代以往「傳統型官僚」（traditional bureaucracy）稱之。演算型官僚是一種新興行政模式，該模式結合人、演算法系統所需之文件與數據等來落實公共服務價值，並克服傳統官僚的侷限性。

表4即為傳統官僚與演算官僚的比較。而在這樣新的官僚模式下，該文即提出六種人機社會技術互動（Socio-technical Interaction）關係，並歸類於兩大模式，分別是可單獨運作的自動代理，以及協助決策的預測輔助工具，在這兩大類六種互動方式中，演算法系統不僅可以取代人，也正在改變人與系統之間的社會技術關係，以及影響著公部門的工作方式（Vogl et al., 2020）。

表 4：傳統官僚與演算官僚的比較

面向	傳統型官僚	演算型官僚
組織	層級節制	協力
服務產出	程序：遵循規則系統	系絡：回應公民獨特需求
知識	專業分工	集體智慧
工具	儲存：文件紀錄被儲存在文件櫃	反饋：記錄可用來進行即時決策
價值	程序平等	結果平等

資料來源：翻譯修改自 Vogl et al. (2020)。

無獨有偶，Wirtz等人（2020）奠基於管制理論（Regulation Theory）及相關演算法挑戰的討論，提出整合性AI治理架構（AI Governance Framework）。此架構的基本結構係由五層所組成，包括1、AI應用/服務與科技層：由於AI技術、服務和應用有可能產生失靈現象，這一層代表了管制的目標；2、AI挑戰層：AI所造成的市場失靈，將展現在AI技術的外部影響和對社會構成的相關挑戰；3、AI管制過程層：為了應對可能的負面影響，需要一個管制過程來評估成本、效益，以及評估有無管制的可能結果；4、AI政策層：透過政策、法律等管制手段以防止或調整導致市場失靈的問題；5、AI協作治理層：有鑑於AI失靈將造成社會巨大影響及其潛在負面影響，受影響的多元利害相關人，以及公私部門的代表應支持整個管制過程提出建言。圖12即描繪了各層演算法治理架構。

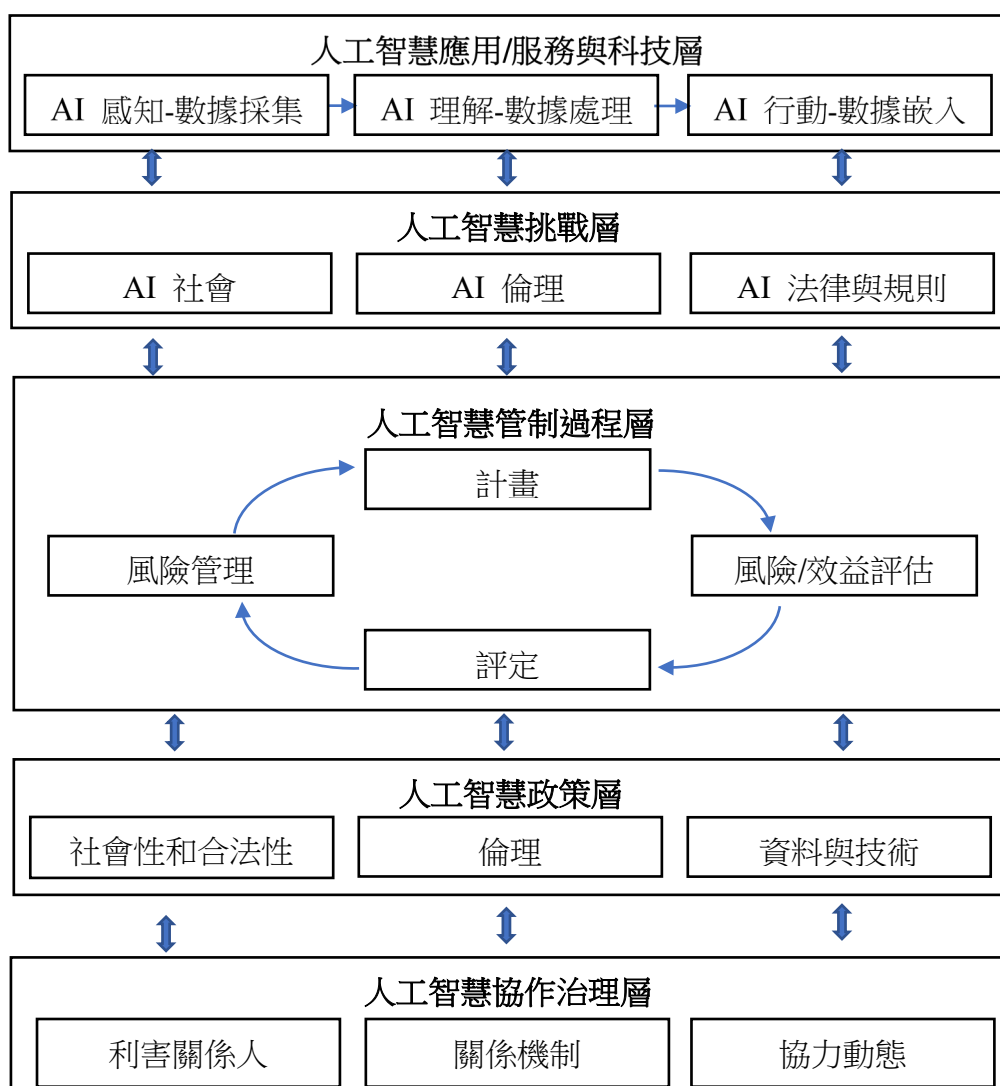


圖 12：整合性演算法治理架構

資料來源：翻譯修改自 Wirtz et al. (2020)。

第一層AI應用／服務和技術層，這一層界定AI的領域範圍，包括如何收集和處理數據以達既定結果。為能理解AI的系統過程，Wirtz等引用Bataller和Harris (2016) 的研究，提出感知、理解與對給定數據採取行動的概念。AI感知主要是通過搜索或整合已獲得的數據集或數據庫來實現，也可以通過外部感測器（如相機、觸覺傳感器或麥克風）來實現。第二層為AI挑戰層，管制的原因是因上述應用AI科技和提供服務時所帶來的挑戰。AI的挑戰分為社會、倫理和法律三部分，儘管這些領域是相互關聯的，但該框架將它們視為不同的部分，其中每個部分都需要對其潛在問題和不同形式的管制進行分析。第三層為AI管制過程層，該層主要透過四個階段來提供如何執行管制的細節說明。這四個過程依序分別是規劃、風險與成本效益評估的界定、實際透過資料蒐集以評定風險與成本效益、最後則是關於AI的風險管理行動。最後兩層則是與管制有關的政策制訂以及為了執行管制而涉入多元利害關係人概念（Cath, Wachter, Mittelstadt,

Taddeo and Floridi, 2017)。為有效執行管制政策，AI的管制不應僅止於政府的政策制定，仍包括私部門及非營利部門的涉入，藉由多元利害關係人以平衡彼此對AI在社會中運作所產生的挑戰與價值衝突，藉由協力機制彼此溝通對話，公私雙方彼此提出解決方案與可行的政策制定，提高社會對於運用AI的信任與接受度。

從上述討論引入 AI 所形成的研究框架來看，政府部門引進 AI 的變革過程，涉及「協力」(collaboration)此核心概念。協力是一個過程，在此過程中涉及多元參與者一起達到共同的目標。過去，人與機器之間在工作職場上各自有負責的工作範圍，而以協力為核心概念的人機協作，則強調彼此的直接互動，以及工作執行上的協力綜效，在這樣的工作場域中稱之為「人機協作」。也因此，機器的角色從過往工具性定位轉變為工作中的協力隊友(Warta, Kapalo, Best & Fiore, 2016)，此從團隊合作的角度出發，人與機器的演算法系統間形成一種夥伴關係，彼此貢獻自己的核心能力(Hoffman & Breazeal, 2004)。

演算法除了與人之間形成有別於以往角色的協力夥伴關係外，我們從上述治理框架中也可觀察到，公私部門引入演算於工作任務的執行，將涉及多種利害關係人的介入，這些關係來自於組織內部及外部。組織內部中涉及跨單位資料所轄權限範圍的共享協力關係；就組織外部而言，政府部門應用演算的影響，將展現在實際的社會場域中，因此社會中多元利害關係人都有可能成為演算法所影響的對象，為能讓演算法產生實質效益，則政府與多元利害關係人之間，彼此須能針對演算法的應用與影響有實質的對話與討論，藉以增進彼此對於演算法的信任與促進演算法能達到更貼近現實的實質效益。

為了適應這種新環境的工作方式，目前迫切需要新的勞動能力和技能，以及相對應的管理方式，逐步改變政府提供公共服務的方式以及與民眾互動的經驗(丁玉珍、林子倫，2020)，因此新的工作職能有建構的實質必要性，而上述所討論的協力概念，也將被涵容於本計畫所建構的職能當中。以下即針對國內外相關研究報告，進行有關人機協作所應具備職能的討論，並於最後提出一整合性的「人機協作職能模式」(Competency Models of Human-robot Collaboration)。

二、 文獻檢閱暨人機協作職能模式的提出

2019年5月OECD會員國在部長級數位經濟政策委員會中，通過了《人工智慧建議原則》(The Recommendation on Artificial Intelligence)。該建議書界定了五項價值原則，以確保AI的課責與可信度。此五項原則略說明如下：(1)AI應能促進成長、永續發展與福祉：各利害關係人應積極建立可信的AI，以追求對人有益的結果，如增強人類的能力與創造力、

促進少數群體的涵容性、減少經濟、社會、性別與其他的不平等；(2)AI 的設計應尊重法治、人權、民主價值和多樣性，並應包括適當的保障措施以確保社會公平性；(3)AI 系統應致力於透明化與負責任的揭露角色，確保人們能理解演算法的結果並可挑戰此結果。希冀受 AI 系統不利影響的個人，能夠根據簡易的資訊來挑戰其結果；(4)AI 系統必須在其整個生命週期中，以穩健、安全的方式運行，並應不斷評估和管理其潛在風險；(5)開發、建造或執行 AI 系統的組織與個人，應秉持上述原則對 AI 的運作來負責（OECD, 2019）。上各項原則主要在確保人機協作下的課責實踐，以及為有效執行人機協作所應具備跨部門合作的能力要求，與前此理論文獻所提及的協力概念似有所呼應。

前述《人工智慧建議原則》在 2019 年提出後，於 2021 年 3 月 OECD 即提出第一次執行評估報告。其中與人機協作職能相關者為報告中「AI Skills, Jobs and Labour Market Transformation」一章所示之內容。該報告以加拿大、新加坡、英國與美國的線上求職平臺資料進行分析發現，人機協作下為能有效與機器人協力執行任務，則人們必須具備某種程度的程式處理能力。同時該報告也指出，像是溝通、解決問題、創造力與團隊合作等技能，也隨著時間的遞移越顯現其重要地位。總體如報告所言，在工作職場中引入 AI 將改變過往的工作環境與內容，益發重視通用技能的趨勢，如組織、管理與溝通技巧，快速反應複雜環境變化的適應力也將特別重要（Grundke et al., 2018；OECD, 2017）。

歐盟注意到技術變革正在改變全球經濟，而為持續繁榮則新的數據技能將是歐洲工業的主要優先發展事項。事實上，全球人才競爭要求歐洲勞動力不斷提高技能，以提高就業能力並推動競爭力。而電子能力架構（European e-Competence Framework 4.0）即是針對 ICT 人員所建構的演算法治理能力框架，此架構共區分出三十種 ICT 專業角色。能力架構請參見表 5 所示，第一欄為該架構之目標，第二欄即為運用 AI 在各目標項下，所必需具備的 41 項基本能力。此四項構面包括規劃（plan）、建造（build）、執行（run）、賦能（enable）與管理（manage）。

表 5：歐盟電子能力架構（European e-Competence Framework 4.0）

目標	技 能	說 明
規劃	A1. 資訊系統與經營策略一致性	預測長期的業務需求，以改善組織流程的效率與效能。
	A2. 服務級別管理	界定適用的服務級別協議。
	A3. 業務計畫發展	處理業務或產品規劃的設計與結構，包括識別替代方案及可能的報酬。
	A4. 產品／服務規劃	分析和界定組織當前和目標狀態。

目標	技 能	說 明
	A5. 建構設計	考慮數據和資訊基礎設施的要求，具體、完善、更新並提供實施解決方案和服務的正式方法。
	A6. 應用設計	根據 IS 政策和使用需要建構執行模式。
	A7. 科技趨勢監測	監測資訊科技的最新發展趨勢。
	A8. 永續管理	評估資訊科技對於環境永續發展的影響。
	A9. 創新	針對產品或服務提供創新性的想法與解決方案。
	A10. 使用者經驗	應用人機互動原則創造直觀、易用、安全且有效的產品與服務。
建造	B1. 應用發展	根據民眾的需要設計與發展合適的應用。
	B2. 組件整合	整合軟硬體與子系統至現有或新創的系統中。
	B3. 測試	建構測試系統以瞭解資訊系統。
	B4. 解決方案佈署	遵循預先設立之標準落實既定規畫之干預以執行解決方案與服務。
	B5. 文件生產	藉由整合資訊與相關要求來產製文件。
	B6. 資訊科技系統工程	建構所需的網絡與連結、組件和接口。
執行	C1. 使用者支持	回應民眾的要求並記錄相關資訊。
	C2. 變革支持	評估與指導資訊科技的變革。
	C3. 服務傳輸	確保按照既定的服務級別協議提供服務。
	C4. 問題管理	管理事件和問題的生命週期。
	C5. 系統管理	監控資訊科技服務及其物理系統與硬體設施。
賦能	D1. 資訊安全策略發展	確保組織策略與文化能維護資訊使用安全免除內外部威脅。
	D2. 資訊科技品質策略發展	界定與改善正式策略以滿足民眾期望與組織績效。
	D3. 資訊教育與培訓	實施資訊科技培訓政策，以解決技能落差問題。
	D4. 採購	應用一致的採購程序
	D5. 銷售發展	為組織的產品與服務建立系統性的流程。
	D6. 數位行銷	瞭解數位行銷的基本原則。
	D7. 資料科學與分析	使用數據分析技術洞察複雜環境中的變化與挑戰。

目標	技 能	說 明
	D8. 契約管理	根據組織程序提供與協商契約制定。
	D9. 人力發展	診斷個人與團隊技能，評定技能間的需求與差距。
	D10. 資訊與知識管理	開發相關程序以辨識與組織相關的資訊與知識。
	D11. 需求辨識	積極瞭解內外部顧客的意見，並闡明他們的需要。
管理	E1. 預測發展	解釋市場需求並評估民眾對於產品或服務的接受度。
	E2. 計畫管理	執行變革計畫之計畫。
	E3. 風險管理	在風險管理政策的規範下，執行跨系統間的風險管理。
	E4. 關係管理	在多元利害關係人中發展業務關係。
	E5. 流程改善	測量現有或新的資訊流程方法的有效性。
	E6. 資訊科技品質管理	執行資訊科技品質政策以維持與促進產品與服務的提供。
	E7. 企業變革管理	評估數位轉型與變革的影響。
	E8. 資訊安全管理	根據資訊科技的策略與組織中相關的風險文化，管理與技術、人員、組織及相關威脅有關的資訊與系統安全政策。
	E9. 資訊系統治理	根據業務需求界定、部署與控制資訊系統、服務與資料管理。

資料來源：翻譯修改自European e-Competence Framework，
<https://ecfusertool.itprofessionalism.org/>。

英國政府網站也提供了幾點有關人機協作下的架構－《倫理、透明化與課責架構》（Ethics, Transparency and Accountability Framework）。該架構聚焦於人機協作下有關數據處理方面的倫理與課責要求，因此出現如數據倫理、數據保護法等專有名詞。英國政府所提出的七點架構將幫助政府部門更安全、永續與合乎道德地使用自動化或演算法決策系統。此七點分別是(1)測試以避免任何意外的結果；(2)提供所有使用者和公民公平的服務；(3)針對演算法的結果有明確責任歸屬；(4)安全的處理數據保護公民利益；(5)協助告知民眾瞭解演算結果對他們產生甚麼樣的影響；(6)確保人機協作能遵守法律規範；(7)監測演算系統並持續蒐集過程中所產製之數據與資料以作為日後查核所用。而根據歐盟和英國計算機協會的調查，人們對先進技術的監管存在著明顯的不信任。

該網站也提及未來人機協作中，演算法扮演兩種協助決策的角色，分

別是完全自動化決策(無人工判斷)和自動化輔助決策(輔助人工判斷)。完全自動化決策意味著完全自動而無需人工判斷，本質上較適用於經常重複、例行的業務；自動化輔助決策則是指自動化或演算法系統輔助人類進行判斷，此較為複雜且對公民具有較大的影響。歐盟於 2018 年所提出的《通用資料保護規則》(General Data Protection Regulation, GDPR) 第 22 條，也針對自動化決策有所規定，原則上禁止完全以自動化來處理個人資料，並產生對資料主體做出具有法律效果或重大類似影響的決策 (GDPR, 2021)。

從英國政府與歐盟 GDPR 第 22 條的規定，均強調人機協作中有關資料安全性，以及對於遵守倫理規範要求的重視。之所以著重上述資料處理的能力，係希冀透過這些規範來提升民眾對於演算法的信任度。對於資訊科技的信任，其實來自於是否有足夠的能力，讓民眾相信我們對於資料的處理，而其中最重要的即在於對資料處理的透明化要求 (Rossi, 2019)。Van Den Bosch, Eric (2017) 對於人機協作中提升信任度的研究也指出，除了對於資料安全性與倫理規範的要求外，社會中人際信任的文化風格也會轉嫁至人對於演算法的信任程度。其將文化區分為三種，分別是尊嚴型文化 (dignity, 如西歐、北美)、面子型文化 (face, 如東亞) 與榮譽型文化 (honor, 如中東與拉丁美洲)，研究結果發現尊嚴型文化對自動化與 AI 的相對信任度最高，榮譽型文化最低，而面子型文化居中。因此，為了能有效提升民眾對於演算結果的信任感，最根本的做法應該是在人機協作下具備對資料處理的安全性與倫理規範的能力要求，以落實資料處理的透明化價值。

前此從相關國際性組織與研究報告，討論有關人機協作的職能要求；反觀國內如鍾文雄 (2019) 也曾強調在全球科技發展浪潮下，組織成員應具備跨領域技能 (cross functional skills)，始能因應未來工作型態的轉變，此些技能包括社會技能 (social skills)、資源管理技能 (resource management skills)、系統技能 (systems skills)、技術能力 (technical skills) 與解決複雜問題的技能 (complex problem solving skills) 等。各能力構面及其所涵蓋的細部技能可參見圖 13 所示。鍾文雄於文中提出上述技能來自於 2018 年亞太區人力資源聯盟 (APFHRM) 論壇，然進一步探詢其文獻源頭可發現，此些技能應修改自「美國勞工部就業和培訓管理局」(U.S. Department of Labor/Employment and Training Administration) 資助「北卡羅萊納州商務部」(North Carolina Department of Commerce) 所成立的職業信息網站 (The Occupational Information Network, O*NET) 之資訊，以該網站內所涵蓋的技能為基本構面，選擇並增加與人機協作所必須具備之職能建構而成。

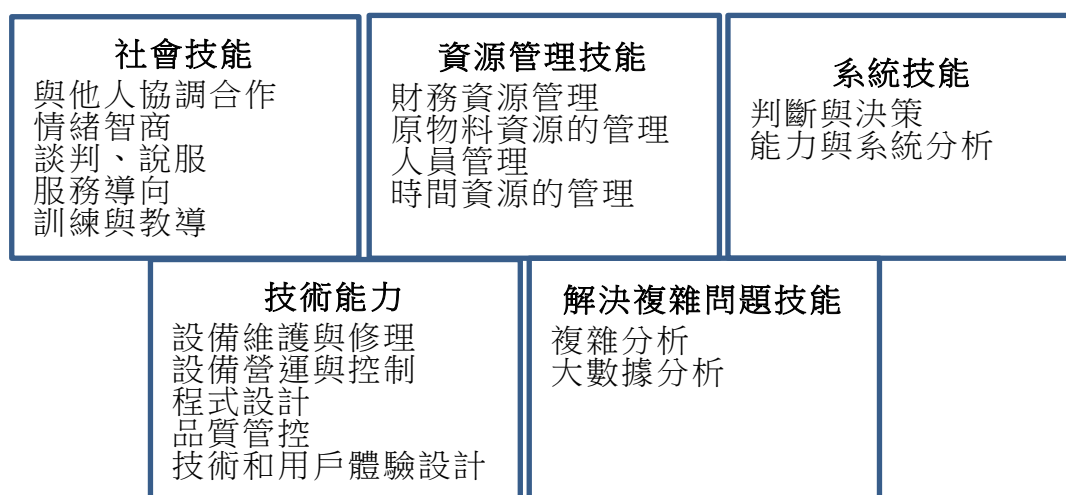


圖 13：人機協作下跨領域技能

資料來源：鐘文雄（2019）。

而張智奇、連瑋鑫（2021）也同樣提出了類似跨領域技能的看法，除了將職能區分為專業職能、核心職能、管理職能三者，此外，透過人機協作與智慧聯網生產職能需求調查，發現專業知識外，跨領域技能與態度相關之專業職能的重要性，內容包括了，重要知識之職能、跨知識與技能之職能（人際影響、資訊收集、協商談判、問題解決、組織規劃等）、跨技能與態度之職能（關係建立、人際影響、成效導向等）、態度之職能（彈性變通、自發性、自我挑戰）。從上述可知，在新型態的工作環境下，必須同時具備專業職能（包括系統面、技術面、處理問題面）、管理職能，以及核心職能（包括社會性、創新性、安全性等類似技能），才能因應數位轉型下的新工作模式。

有關職能中對於創新性的要求，除了前此在 OECD、歐盟與張智奇等（2021）有被提及外，謝天和（2019）所提出的演算法環境應具備的專業能力中，即認為前三重要者分別為「創新思維」、「相關資訊安全法令道德素養」、「解讀與分析數據」，其中以創新思維為最。換言之，在數據化、演算法的發展下，個人除了應具備基本解讀及分析數據的專業能力外，還需透過創新的思維模式，保持自身應變能力與新興人機合作方式，以面對未來創新科技的挑戰。

除此之外，在數位轉型下的政府治理模式，將改變以往公部門的職場運作環境，人機協作的運作將成為未來運作的主要工作型態，公部門導入演算法存在著許多挑戰與風險，各國也都紛紛提出演算法應變措施與相關法令，臺灣也不例外，並於 2018 年正式啟動了為期四年的演算法計畫，作為國家發展策略（丁玉珍、林子倫，2020），而人才的運用、資料的分析以及支援經費等都是演算法發展重要的部分，由此可見，為了發展人機協作的有效性，除了必須具備專業職能與核心職能外，相關管理資源的分配也必須作為輔助角色，始能協助落實人機協作的計畫執行。

私領域導入演算法較公部門來得快，然而發展速度也導致公私領域存有技術上的差距，政府部門將演算法整合到公共服務中尚未發展成熟（Mikhaylov et al., 2018），若是要將演算法成功導入公部門中，勢必要透過公私協力的方式，結合企業運作模式，將知識導入公部門，增加公務員對演算法技能的認知（丁玉珍、林子倫，2020），也因此呼應前此我們所提出的，人機協作的整體運作在各個階段其實均植基於協力的核心價值而來。以下本計畫將整合上述所提出對於人機協作的原理原則，與相關職能所提出的討論，相較於國際性組織較著重於資料安全性與倫理規範的重視，本計畫將擷取鍾文雄所提出的技能為基礎，延伸至美國勞工部網站蒐集各項技能的細部說明，並增加有關安全性與倫理規範構面，當然，歐盟的資訊架構能力也有其完整性，但因該架構係以資訊部門人員為對象所建構，與本文初期欲建構一能適用性較廣的人機協作職能的目的有別，同時，納入相關文獻對於創新性的能力要求。**表 6** 即為本計畫所整合提出的人機協作職能地圖。

表 6：人機協作職能地圖

職能	技能	細項	說明
核心職能	社會性技能 具備能與他人共事以達成目標之能力。	協調合作	根據他人的行動調整行動。
		情緒智商	瞭解他人的反應並理解他們為什麼會如此反應。
		談判能力	將其他人聚集在一起並試圖調和分歧。
		說服能力	說服他人改變主意或行為。
		服務導向	積極尋找幫助人們的方法。
		訓練與教導	教別人如何做某事。
	倫理與安全 確保政府安全、永續與合乎道德規範的使用 AI 系統。	避免非意圖結果	測試以避免任何意外的結果。
		確保公平服務	提供所有使用者和公民公平的服務。
		安全處理數據	安全的處理數據保護公民利益。
		遵守法律規範	確保人機協作能遵守法律規範。
	創新性 為提供新概念、想法、產品或服務設計創造性的解決方案。	具備創新思維	運用新穎和開放的思維來設想利用技術進步來滿足商業／社會需求或研究方向。

職能	技能	細項	說明
管理職能	資源管理技能 具備有效資源分配的能力	財務資源管理	確定如何花錢來完成工作，並核算這些支出。
		原物料資源的管理	獲得並確保正確使用完成某些工作所需的設備、設施和材料。
		人員管理	在員工工作時激勵、發展和指導他們，確定最適合工作的人選。
		時間資源管理	管理自己的時間和他人的時間。
		資訊與知識管理	瞭解要部署的適當工具來創建、提取、維護、更新和傳播業務知識，以便從資訊資產中獲利。
		契約管理	建立和維護供應商關係和定期溝通。
專業職能	系統技能 針對 AI 系統的輸出具備有判讀、監測與改善的能力。	判斷與決策	考慮系統各項運作方式的相對成本與效益，並選擇最合適者。
		能力與系統分析	決定系統該如何運作以及瞭解條件、操作和環境的改變將對系統產生甚麼影響。
	技術能力 對於 AI 系統的應用具備有設計、建立、運作與故障排除的能力	設備維護與修理	對設備進行日常維護，並確定何時需要進行什麼樣的維護。使用所需工具修理機器或系統。
		設備營運與控制	控制設備或系統的運作。
		程式設計	為各種目的編寫系統程式。
		品質控管	對產品、服務或過程進行測試和檢查以評估輸出品質與績效。
		技術與用戶體驗設計	產生或調整設備和科技以滿足民眾需求。
	解決複雜問題能力 能在複雜的現實環境中運用大數據分析解決新的、定義不明確的問題。	複雜性分析	識別複雜的問題並審查相關資訊以發展和評估方案並實施。
		大數據分析	使用和應用數據分析技術，例如數據挖掘、機器學習、規範性和預測性分析，以應用數據洞察來應對組織的挑戰和機遇。

資料來源：本計畫。

第三章 調查分析方法

第一節 研究設計與方法

針對上述議題，本計畫以三大方面著手，分別為跨國資料分析與法規調適、國內部會調查以及法務部觀護人實驗與調查。以本計畫主軸之制度分析「法規的調適」（憲法層次）、「資源的配置」（組織層次）、以及「能力管理」（個人層次）三層次，指出這三者的「落差」，建構出我們的改革路徑建議。

基於上述，本計畫採量化與質化的混合研究方法（Mixed Methods）方式為之，量化研究方面將進行二項調查，分別為：1、行政院部會演算法應用與推動調查；2、法務部觀護人的決策演算法實驗；質化研究分為三部分，分別為：1、文獻與理論的建構，以演算法應用為政府數位轉型的論述，於理論上制度調適的文獻蒐集、制度調適與落差分析的連接；最後，從制度調適、落差分析到制度改革路徑分析的操作文獻的整理；2、透過跨國資料蒐集，瞭解各國之發展並反思我國可借鏡之處，並以深度訪談與專家座談方式，瞭解我國相關實情。成為未來公部門應用演算法技術的參考。整體而言，本計畫之研究架構與設計可表達如圖 14，研究資料來源摘要如表 7 所述。

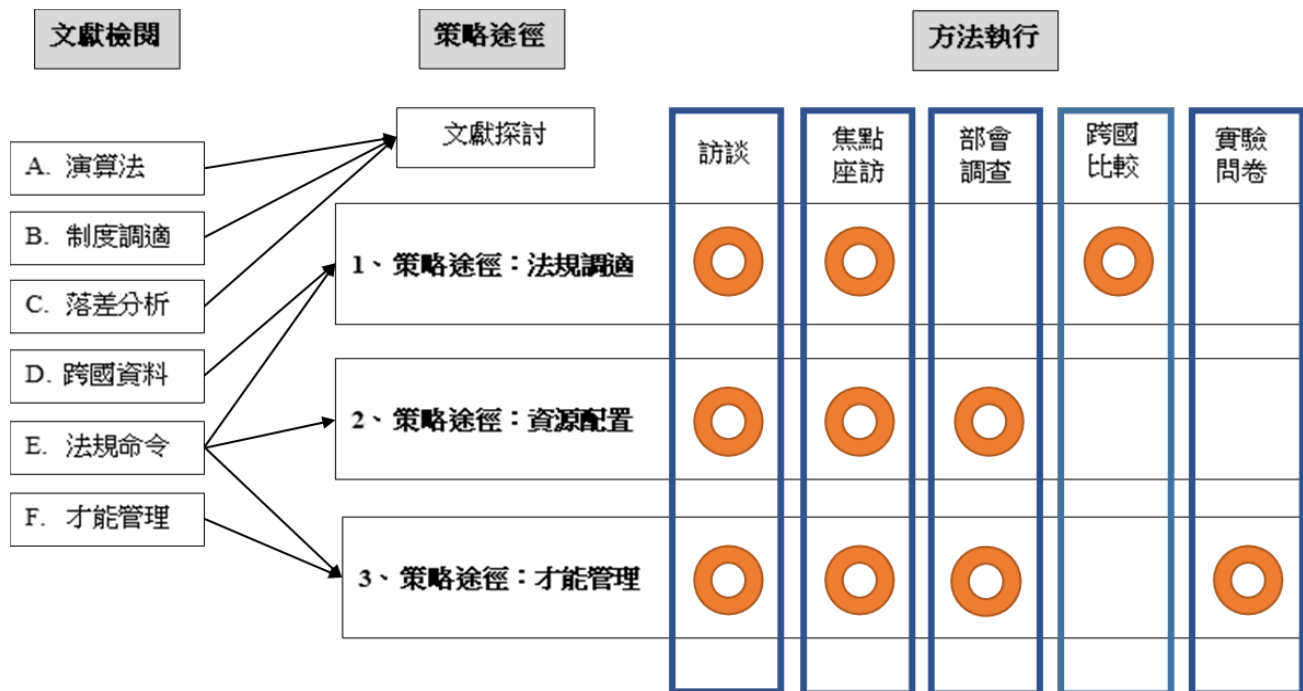


圖 14：子計畫一研究架構與設計

資料來源：本計畫。

表 7：子計畫一研究資料來源摘要

研究目的	資料搜集方法與來源
蒐集 3-5 個國家將機器演算法應用於政府數位服務的政策、法制、組織數位成熟度等，對照臺灣目前現況進行比較分析，包括：正確性（accuracy）、公平性（fairness）、課責（accountability），以及對外溝通機制，例如：透明（transparency）、知情同意（informed consent）等。	● 為發掘國外相關機器演算法之應用，本計畫初步針對美國、歐盟、英國以及新加坡等先進國家，透過以下途徑蒐集資料： 1. 英文期刊資料庫。 2. 學術機構、國際組織、國內外政府出版品。 3. 國際性數位政府研究機構之資料庫。 4. 彙整並比較分析機器演算法應用於公部門之相關研究文獻。 ● 除了透過上述的資料之分析，更透過深度訪談與焦點座談聚焦討論問題成因與可能解方，並試圖提出建議。特說明相對應的焦點座談場次為第三場焦點座談、第四場焦點座談、第五場焦點座談與第六場焦點座談。 1. 第三場焦點座談：針對國內外法規制度進行問題成因分析。（問題分析階段） 2. 第四場與第五場焦點座談：針對國內個案 Chatbot 技術於公部門應用之探討，分別從組織與個人層面面臨之問題探析，必且從中找尋未來規劃與執行上可解決之道。（問題分析階段、解方分析階段） 3. 第六場焦點座談：透過實務角度來檢視目前國外針對 AI 的法規規範與制度的建立，反思我國目前之規範現況與 AI 法治之建構。（解方分析階段、建議提出階段）

研究目的	資料搜集方法與來源
透過個案研析，瞭解國內中央政府機關推行應用機器演算法的現況、衝擊，以及其對於組織或個人所面臨的挑戰。	● 針對該研究目的，本計畫透過兩階段進行資料的蒐集與分析： 1. 質化與量化的混合研究方法：針對已應用機器演算法於行政實務之政府機關利害關係者、相關領域之學者專家以及中央政府機關資訊專責人力進行調查；另外，面對人機協作的未來趨勢，人力資源發展也是本計畫探究之重點之一。因此，本計畫先藉由深度訪談瞭解目前中央各部會推行機器演算法之現況與挑戰，藉此也舉辦第一次焦點座談（問題分析階段），邀集相關學者針對發放之問卷，來釐清問卷的適切性與後續可關注之處。 2. 三角定錨法（Triangulation）分析：將蒐集之質化與量化資料進行分析，以掌握當前中央政府機關推行機器演算法之現況，並進一步發掘機器演算法對於中央政府機關及政府資訊專責人力所帶來的挑戰。 ● 本計畫挑選法務部保護司作為研究個案，透過混合研究方法研究此個案。 1. 深度訪談與焦點座談（第二次焦點座談）：瞭解該單位組織與個人決策的實際狀況、工作流程與數位轉型應用的需求與想法。 2. 問卷實驗：根據訪談資料，設計較符合現實的實驗調查研究，將設計實驗情境、實驗介入與實驗前後問卷調查。
參考國內外政府與個案對於機器學習的服務應用，從組織調適的角度，進行落差分	● 綜上研究目的一與二之資料蒐集與分析，進一步進行落差分析，歸納研究結果與發現，透過深度訪談或焦點座談方

研究目的	資料搜集方法與來源
析 (Gap Analysis)，並從中研提機器演算法應用於政府數位服務的具體改革路徑 (path) 建議，其中至少包括：法規的調適、資源的配置，以及能力的管理等。	式，針對重要結果與發現進行探討，提出政策建議。

資料來源：本計畫。

一、 研究流程

本計畫進行流程 (如圖15) 可概分三個階段。第一階段為「問題分析」的階段，背景資料及文獻的蒐集與探討，採擷國外演算法的實際推動個案，歸納、分析我國政府未來演算法運用於公共服務或業務中其必要性、可能面臨的挑戰以及可行性，過程中也會透過焦點座談與深度訪談試圖釐清問題。該階段的焦點座談分別有第一場焦點座談針對中央各部會問卷之調查前的問題釐清、第二場焦點座談為法務部觀護人調查之問題探析、第三場焦點座談為針對國內外法規制度進行問題成因分析。

本計畫第二階段為「解方分析」的階段，部分焦點座談、深度訪談都會在這個階段開始實施。該階段藉由深度訪談，確立本計畫在第一階段透過文獻檢閱、他國經驗等釐清問題與困境，同時進一步補充文獻描述中所缺乏的部分。另外，針對中央各部會之問卷調查以及法務部觀護人之實驗，試圖從組織與個人層面瞭解我國目前對於演算法的認知、未來執行之挑戰等。該階段的焦點座談分別有第四場與第五場焦點座談，透過深度訪談與文獻資料來得知目前公部門Chatbot運用於公共服務的情形，進一步從國內個案中探討Chatbot技術運用之情形，各從從組織與個人層面面臨之問題探析，必且從中找尋未來規劃與執行上可解決之道。為問題分析階段與解方分析階段的綜合座談。

第三階段亦有討論焦點座談會，並在最後進行資料整理、分析。焦點座談會、深度訪談的結果，也將進一步透過逐字稿的整理予以歸納，並撰寫報告，最後提出研究結論與建議。該階段為第六場焦點座談，透過實務角度來檢視目前國外針對AI的法規規範與制度的建立，反思我國目前之規範現況與AI法治之建構，為解方分析階段與建議提出之座談。

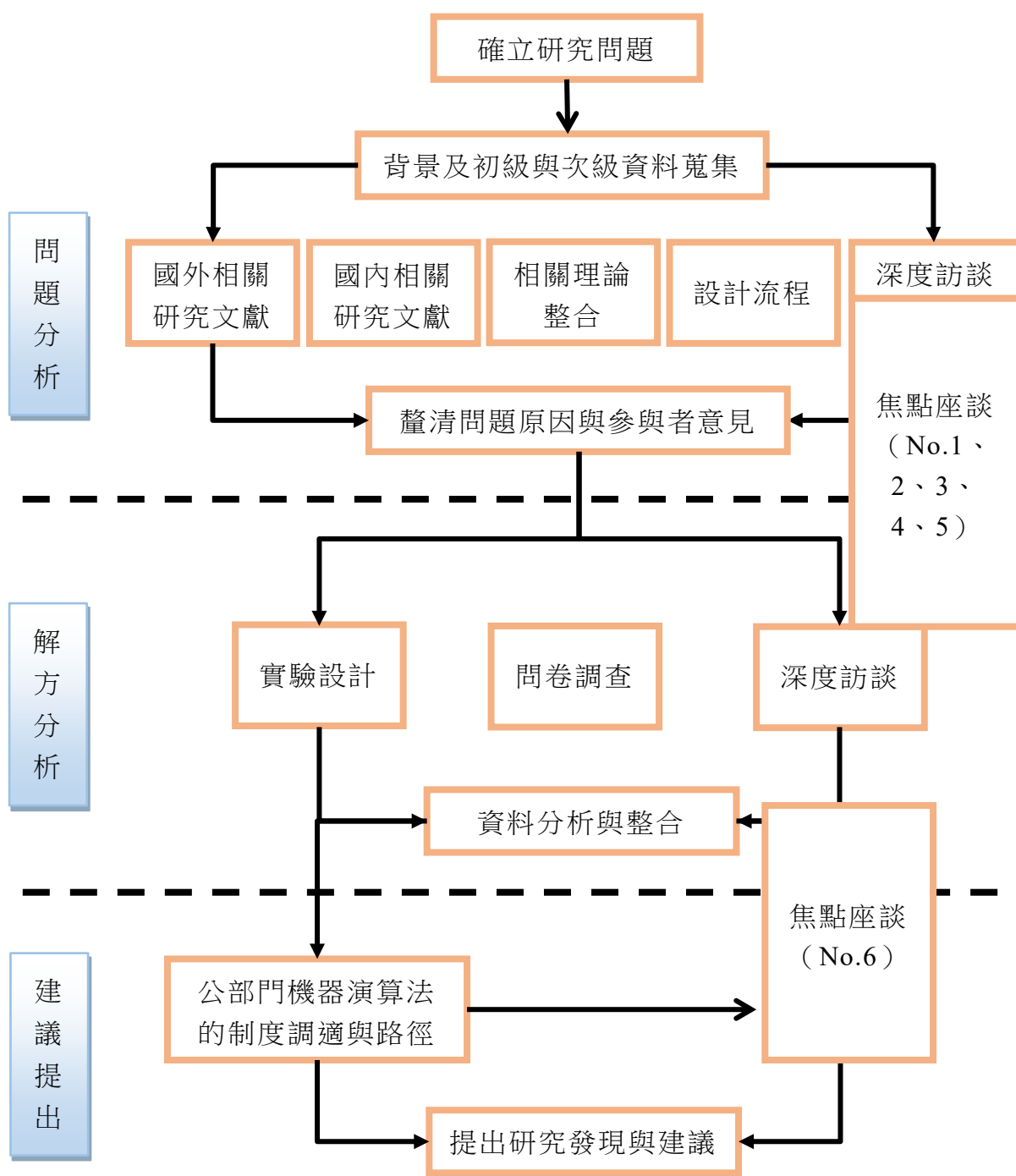


圖15：子計畫一研究流程圖

資料來源：本計畫。

二、 深度訪談、焦點座談

本計畫自110年5月開始，已進行國內外相關文獻檢閱，舉辦焦點座談5場、深度訪談6位，由本計畫研究團隊陳計畫主持人（D）、廖協同主持人（L）、黃協同主持人（H）、王協同主持人（W）與張協同主持人（C）訪談。

對於深度訪談的挑選，經由立意抽樣的方法，透過本計畫團隊的人際網絡關係、國發會資管處引薦之下，其領域來自政府機關（G1）、法律（E1、E3、E6）、相關研究領域學者（E2、E4、E5、E7），共計訪談8位受訪者，詳細深度訪談受訪者背景資料表如表8所示。為保護受訪者之權利，訪談前由研究團隊說明研究參與者之權利與本計畫對於訪談內容之運用，並徵求受訪者同意錄音，簽署研究參與者知情同意書後進行訪談。

表 8：子計畫一深度訪談受訪者背景資料表

序號	機關位階	代表性	性別	代碼	訪談時間
1	臺灣新北地方檢察署	觀護人	女	G1	5/27
2	中央研究院歐美研究所	副研究員	女	E1	6/28
3	東吳大學巨量資料管理學院	助理教授	男	E2	7/1
4	世新大學智慧財產暨傳播科技法律研究所	副教授	男	E3	7/6
5	國立師範大學科技應用與人力資源發展學系	助理教授	男	E4	8/19
6	國立中正大學企業管理學系	教授	男	E5	8/27
7	國立清華大學科技法律研究所	副教授	男	E6	12/3
8	中研院資訊處	處長	男	E7	12/10

資料來源：本計畫。

深度訪談以半結構式方式進行，根據本計畫之研究目的設計訪談題綱（如表9），以「國內外法規制度落差探討與法規調適對策建議」、「演算法的案例探討」、「法務部觀護人其工作流程瞭解」與「人機協作之職能模式建構」四大面向深入探討，目的為瞭解現今國外與我國公私領域對於演算法的想法與實際情形，從法規調適、資源配置與才能管理三角度去多元探討，進而思考公部門對於演算法之應用應如何進行制度的調適。訪談題綱的設計，並非會完全依照以下題目完全向受訪者請益，會依照受訪者的專業與擅長領域進一步挑選題目；訪談過程若受訪者提出之想法引發訪談者進一步的問題思考，會視情形開放式請益，使訪談內容更加聚焦。

表 9：子計畫一深度訪談題綱

構面	題綱
國內外法規制度落差探討與法規調適對策建議	1. 您認為公部門應用演算法治理（Algorithm Governance）內涵為何？例如，如何定義「演算法」、透明性之要求，及與 AI 在概念上的差異？ 2. 公部門應用演算法治理，除了資料處理與保護的相關問題以外，尚有哪些法律規範的問題應加以考量？ 3. 公部門應用演算法執行任務時，是否應受限制？須限制之主因為何？此限制與追求公共利益間應如何平衡？ 4. 公部門應用演算法之正面／負面影響為何？並就您的觀察，國內外是否已有公部門演算法運用之實例，若有，有哪些各國想解決的主要問題？ 5. 就國外公部門應用演算法之經驗，臺灣應用的可能性與範圍為何？ 6. 您認為臺灣在規劃公部門應用演算法治理時，目前法規範體制可能有哪些不足的面向？
演算法的案例探討	1. 請問您就「演算法相關的技術於行政業務革新」乙事，有無印象特別深刻的個人經驗或熟知的案例？若有可否介紹一下該案例的相關內容？ 2. 請問您如何看待前述提及案例的實施成效？（成功或失敗的部分皆可） 3. 從應用資訊科技進行行政革新的角度而言，您認為一個成功的革新應注意哪些面向以及關鍵的成功因素為何？
法務部觀護人其工作流程瞭解	1. 請問觀護人，一般在進行風險評估時，你們的工作流程為何？會需要哪些相關的參考資料輔助？ 2. 觀護人在面對不同類型參考資料（e.g., 保護管束人前科紀錄、本案判決書、在監服刑情形、家庭與訪談資料）的權重如何分配？ 3. 再犯風險評估與對應處置決策的關聯性為何？ 4. 在評估過程與決策中，會遇到哪些困難？（任何面向都可以） 5. 年資與經驗對評估與決策的是否有差異？
人機協作之職能模式建構	1. 依著科技變革導致工作方式產生人機協作的變化，在這樣的環境中，組織如何重新界定其策略與發展目標？必須考量哪些層面與因素？

構面	題綱
	2. 演算法智慧發展的改變，使得人機協作成為最有效率的工作方式，在人機協作下，人與機器各自扮演了甚麼樣的角色？我們將依據甚麼樣的判準，將工作任務分配給人？分配給演算法？甚麼時候又分配給人機協力來進行？ 3. 人機協作下，個人必須具備甚麼樣的背景、態度或者相關經歷？ 4. 從組織層次觀察人機協作模式的成效，則個人必須具備甚麼樣的核心職能？以及為順利完成工作所需的專業技能為何？其中，就人與演算法來說，兩者間的溝通與對話是發生在甚麼樣的時機點，以及應具備甚麼樣的能力？ 5. 就您個人所知，有沒有推薦的公、私部門中有那些人機協作的實際案例？您認為他們之所以運用成功或者仍具改善空間的原因為何？（成功的或者仍具有改善空間）

資料來源：本計畫。

除透過深度訪談瞭解對於演算法的見解與未來公部門如欲運用的建議，也以焦點座談方式，邀請相關領域學者專家，參與者需具備對於該座談主題相關之知識或經驗，針對座談討論議題提出意見與感受。本計畫已於110年7月15日執行第一場焦點座談，針對「中央各部會應用演算法創新公共服務之現況調查」之問卷架構與設計進行討論，此場座談邀請大學任職之學者共計5位，詳細第一場焦點座談參與者背景資料表如表10所示。

表 10：子計畫一第一場焦點座談參與者背景資料表

序號	機關位階	代表性	性別	代碼	座談時間
1	世新大學行政管理學系	教授	男	P1	7/15
2	臺北大學公共行政暨政策學系	副教授	男	P2	
3	臺灣大學政治學系	副教授	女	P3	
4	臺灣師範大學公民教育與活動領導學系	教授	男	P4	
5	淡江大學課程與教學研究所	助理教授	男	P5	

資料來源：本計畫。

第一場焦點座談聚焦於討論「中央各部會應用演算法創新公共服務之現況調查」的問卷架構與設計，擬定之題綱如表11所示：

表 11：子計畫一第一場焦點團座談題綱

研究目的	題綱
透過個案研析，瞭解國內中央政府機關推行應用演算法的現況、衝擊，以及其對於組織或個人所面臨的挑戰。	1. 本計畫之實驗流程設計有何需要注意之處？ 2. 針對本問卷設計之內容，請問有何需要修改增進之處？ 3. 其他開放式請益，或與會專家意見的補充。

資料來源：本計畫。

110年8月16日執行第二場焦點座談，針對「法務部保護司導入機器學習演算系統實驗調查設計規劃」之問卷架構與設計進行討論，此場座談邀請大學任職之學者以及法務部保護司觀護人代表共計6位，詳細第二場焦點座談參與者背景資料表如表12所示。

表 12：子計畫一第二場焦點座談參與者背景資料表

序號	機關位階	代表性	性別	代碼	座談時間
1	臺北大學公共行政暨政策學系	副教授	男	P2	8/16
2	臺北大學公共行政暨政策學系	教授	男	P6	
3	臺灣大學圖書資訊學系	副教授	男	P7	
4	臺北市立大學社會暨公共事務學系	副教授	男	P8	
5	法務部保護司	觀護人	女	G2	
6	法務部保護司	觀護人	女	G3	

資料來源：本計畫。

第二場焦點座談聚焦於討論「法務部保護司導入機器學習演算系統實驗調查設計規劃」的問卷架構與設計，擬定之題綱如表13所示：

表 13：子計畫一第二場焦點團座談題綱

研究目的	題綱
透過個案研析，瞭解國內中央政府機關推行應用機器演算法的現況、衝擊，以及其對於組織或個人所面臨的挑戰。其中，透過實驗調查，瞭解法務部保護司觀護人在有無演算系統的輔助下，對個案進行再犯評估風險的判斷及決策擬定後處理措施之差異，希冀能從中理解演算法應用對公部門實際的好處、顧慮、與決策困境等。	1. 本計畫之研究架構，其因果關係設計有何需要改進之處？針對取樣方法有何建議？ 2. 針對本問卷設計之內容，請問有何需要修改增進之處？ 3. 其他開放式請益，或與會專家意見的補充。

資料來源：本計畫。

110年9月8日執行第三場焦點座談，針對「國內外AI與演算法應用與法規制度落差探討」進行討論，此場座談邀請法律學界之學者以及實務專家共計6位，詳細第三場焦點座談參與者背景資料表如表14所示。

表 14：子計畫一第三場焦點座談參與者背景資料表

序號	機關位階	代表性	性別	代碼	座談時間
1	財團法人資訊工業策進會科技法律研究所	研究員	男	P9	9/8
2	政治大學法學院	副教授	女	P10	
3	臺灣大學法律學院	副教授	女	P11	
4	臺灣大學國家發展研究所	副教授	女	P12	
5	臺北醫學大學醫療暨生物科技法律研究所	副教授	男	P13	
6	博思法律事務所	律師	男	P14	

資料來源：本計畫。

第三場焦點座談聚焦於討論「國內外法規制度落差探討」，擬定之題綱如

表15所示：

表 15：子計畫一第三場焦點團座談題綱

研究目的	題綱
瞭解各國制度之設計，反觀我國公共服務如欲導入演算法，其應扮演之角色、規範、管制、權利保障與主責機關之規劃。	1. 政府公共服務導入演算法運用，其適當的角色為何？ 2. 政府公共服務導入演算法應用時，應採用何種層級之規範方式？ 3. 政府公共服務導入演算法應用時，是否應該以及如何將公共服務進行類型化，以作不同程度之管制？ 4. 政府公共服務導入演算法應用時，可能對人民權利可能產生之影響？以及可能影響哪些基本權利？ 5. 對政府公共服務導入演算法應用之協調與監理，應由何種類型之組織或單位負責？

資料來源：本計畫。

110年9月24日執行第四場焦點座談，針對「Chatbot技術應用的實務運作之探討」進行討論，此場座談邀請法律學界之學者以及實務專家共計5位，詳細第三場焦點座談參與者背景資料表如表16所示。

表 16：子計畫一第四場焦點座談參與者背景資料表

序號	機關位階	代表性	性別	代碼	座談時間
1	臺灣大學政治學系	教授	女	P3	9/24
2	空中大學管理與資訊學系	助理教授	男	P15	
3	臺北市政府資訊局系統研發中心	主任	男	G4	
4	中央銀行資訊處	前處長	男	P16	
5	意藍資訊股份有限公司	董事總經理	男	P17	

資料來源：本計畫。

第四場焦點座談聚焦於討論「Chatbot技術應用的實務運作之探討」，擬定之題綱如表17所示：

表 17：子計畫一第四場焦點團座談題綱

研究目的	題綱
透過個案研析，瞭解國內中央政府機關推行應用機器演算法的現況、衝擊，以及其對於組織或個人所面臨的挑戰。其中，以現行之 Chatbot 技術為主題，瞭解於公部門之實務運用，希冀瞭解政府組織運用該技術時應具備之條件、適用之公共服務領域、制度設計、組織運作與個人應具備之能力等要素之建構。	1. 請問目前 chatbot 技術有運用在政府的公共服務中嗎？有哪些應用類別？有無個人經歷過或是知曉的案例？若有可否介紹一下該案例的相關內容？ 2. 您覺得 chatbot 技術在政府運用，政府組織需要具備什麼條件？又在哪些公共管理領域居多？該技術如要輔助組織內成員的業務執行，運用的目的應該為何？ 3. 運用 chatbot 技術，需要懂得訓練 chatbot 的工作？請問你所知道的訓練方法有哪些？這些方法的優缺點又是甚麼？ 4. 從應用資訊科技進行行政革新的角度而言，您認為一個成功的 chatbot 導入革新，從制度設計、組織運作以及個人反映，其成功的要件為何？ 5. 其他開放式請益，或與會專家意見的補充。

資料來源：本計畫。

110 年 10 月 8 日執行第五場焦點座談（表 18），以 Chatbot 為例探討人機協作導入職能模型的建構，受訪者挑選則以產官學三面向的專家學者為主，包含產業界 AI 專家、政府單位應用 AI 服務者（尤其著重於具備承辦聊天機器人相關經驗）、以及 AI 相關領域學者，透過實務角度來檢視該職能地圖有無需要修改及更正之處，並期望透過學術專業的角度激盪出產官學界對於人機協作人才職能的建議。

表 18：子計畫一第五場焦點座談參與者背景資料表

序號	機關位階	代表性	性別	代碼	座談時間
1	高雄市政府研究發展考核委員會	專門委員	女	G5	10/08
2	國立臺灣師範大學科技應用與人力資源發展學系	助理教授	男	E4	
3	臺北市政府研究發展考核委員會話務管理組	組長	男	G6	
4	經濟部商業司	科長	男	G7	

序號	機關位階	代表性	性別	代碼	座談時間
5	詠鉉智能股份有限公司(Chimes AI)	執行長／ 博士	男	P18	
6	東吳大學資料科學系	助理教授	男	P19	

資料來源：本計畫。

第五場焦點座談聚焦於討論「以 Chatbot 為例探討人機協作導入職能模型的建構初探」擬定題綱如表 19：

表 19：子計畫一第五場焦點團體座談題綱

研究目的	題綱
透過個案研析，瞭解國內中央政府機關推行應用機器演算法的現況、衝擊，以及其對於組織或個人所面臨的挑戰。其中，以 Chatbot 為例探討人機協作導入職能模型的建構初探。	1. 在您執行 Chatbot 經驗中，除上述所提出的三項主要核心能力外，是否還有哪些核心能力是您認為也應具備的？為什麼？而在這些核心能力中，您個人認為在您承辦該項業務的執行過程裡，他們分別扮演了甚麼樣的角色？ 2. 在每項核心能力下，目前各自有其組成項目，就您個人執行 Chatbot 的經驗中，各項核心能力下是否有目前未提及或者需要調整的項目能力？為什麼？對於各項目的界定是否有需要調整之處？ 3. 承上，在每個項目下也有其組成的次項目，在您執行 Chatbot 的經驗中，各次項核心能力下是否有目前未提及或者需要調整的項目能力？為什麼？對於各項次項目的界定是否有需要調整之處？

資料來源：本計畫。

110 年 12 月 14 日執行第六場焦點座談（表 20），探討國內外法規制度落差探討與法規調適對策建議，受訪者挑選則以產官學三面向的專家學者為主，包含產業界 AI 專家以及科技法律相關領域學者，透過實務角度來檢視目前國外針對 AI 的法規規範與制度的建立，反思我國目前之規範現況與 AI 法治之建構。

表 20：子計畫一第六場焦點座談參與者背景資料表

序號	機關位階	代表性	性別	代碼	座談時間
1	詠鉉智能股份有限公司	執行長	男	P18	12/14

序號	機關位階	代表性	性別	代碼	座談時間
2	中央研究院法律學研究所資訊法中心	副主任	男	P20	
3	臺北醫學大學醫療暨生物科技法律研究所	教授	男	P21	
4	臺北智慧城市專案辦公室	主任	男	P22	
5	臺灣大學法律學系	副教授	男	P23	
6	世新大學法律學系	副教授	男	P24	

資料來源：本計畫。

第六場焦點座談聚焦於討論「國內外法規制度落差探討與法規調適對策建議」擬定題綱如表 21：

表 21：子計畫一第六場焦點團體座談題綱

研究目的	題綱
透過實務角度來檢視目前國外針對 AI 的法規規範與制度的建立，反思我國目前之規範現況與 AI 法治之建構。	1. 公部門應用機器演算法治理（algorithm governance）之內涵，除資料處理與保護外，尚有哪些法律問題應加以考量？ 2. 國外是否已有公部門機器演算法運用之實例？其中有哪些各國想解決的主要難題？其解決方法為何？是否有為此制訂相關法制規範？是否遇到技術領域或其他面向的反彈？臺灣公部門是否有借鏡之可能性如何？ 3. 臺灣公部門應用機器演算法的可能性與範圍為何？此類新興型態之行政治理與革新，對於傳統公部門依法行政的思維與行為模式，可能帶來什麼影響？現階段應用面有哪些難題及不足？ 4. 自法規制度及資訊技術之不同角度觀察，公部門於應用演算法之政策、制度或法律規範上，應給予何種程度之監管框架，始較能達成追求公共利益、人力發展及科技創新間之平衡？與對私部門應用機器演算法之規範制訂上，是否應有何不同考量？ 5. 臺灣目前面對公部門應用演算法治理之政策、制度或法律規範，可能有哪些不足或困難？（例如：是否宜

研究目的	題綱
	重新建構體制，或於數位資訊或科技技術上有遵循上之難處？）若產生權利侵害之風險時，法制規劃上應如何進行事前之預防、事中之控管、與事後之究責？又自資訊技術面以觀，是否可能做到？

資料來源：本計畫。

第二節 中央各部會應用演算法之現況調查

本計畫團隊於110年9月2日至10月15日間，辦理中央各部會應用機器演算法創新服務之現況調查，茲就調查的設計以及調查題項內容說明。本節將就本項調查設計重點依序說明如下：

一、 調查目的

為瞭解中央政府各部會應用機器演算法(AI)¹⁴創新公共服務的現況，本計畫團隊特就各部會在組織層面之行政運作實況，以及在個人層面之公務人員應用意願、預期風險或遭遇阻礙等面向進行問卷調查，俾蒐集相關經驗資料進行落差分析，以進一步研提政策建議。

二、 調查方式

本計畫團隊採用Surveymonkey問卷平台製作之網路問卷進行調查，調查分為兩階段，第一階段調查旨在瞭解中央各部會應用或規劃應用機器演算法(AI)創新公共服務的現況；第二階段調查旨在瞭解中央各部會公務人員對於行政機關應用或規劃應用機器演算法(AI)創新公共服務的看法。各階段調查皆採填答者匿名以及先邀請後發送問卷方式實施，先由本計畫團隊以國立政治大學名義函請中央各部會指派受訪人員，再由該人員填妥公文隨附之各階段受訪者聯絡資料表後回傳，計畫團隊接著再依據受訪者提供之電子郵件地址分別寄送第一階段及第二階段網路調查問卷並跟催後續填答情況，調查期間也安排研究助理回答受訪機關提問，以提高問卷的填答率。

三、 調查研究問題

第一階段調查以部會為分析單元；第二階段調查以個別公務人員為分析單元，本計畫團隊期透過兩階段調查回答以下問題。這部分調查研究

¹⁴ 本計畫為使受測者更能理解演算法之運用，以AI為例藉此協助受測者對於演算法運用的理解，進而使用演算法（如：AI）更能針對題目語意回答問題。

的問題意識，是從文獻回顧、期中報告前的深度訪談與焦點座談、以及跨國資料的分析中，歸納出研究團隊可以在我國中央政府可以進行的調查研究議題，這些議題如下所列：

(一)第一階段調查：

- 1、中央部會應用機器演算法(AI)創新公共服務可行嗎？可能優先推動的項目為何？
- 2、中央部會應用機器演算法(AI)創新公共服務可能遭遇之制度限制為何？
- 3、中央部會應用機器演算法(AI)創新公共服務所需的資源條件為何？包括預算、人力、資料需求等。

(二)第二階段調查：

中央部會公務人員對於應用機器演算法(AI)創新公共服務的態度、預期效能以及對可能遭遇問題的解決態度為何？

四、抽樣設計

由於缺乏中央各部會潛在受訪者之母體清冊，故兩階段調查之抽樣設計分別採用立意抽樣法以及滾雪球抽樣法。

- (一)第一階段調查採立意抽樣法，由政治大學發文函請行政院各部會(32個，如圖16所示)各指派1名所屬公務人員填答問卷，該人員需協助蒐集現況資料並以部會立場填答問卷。

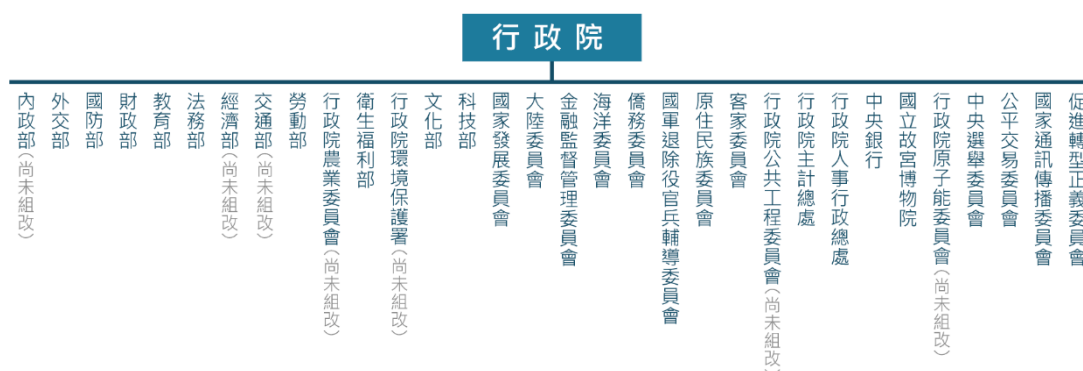


圖 16：行政院組織架構圖

資料來源：行政院官網：<https://www.ey.gov.tw/Page/62FF949B3DBDD531>，最後檢閱日期 2021/11/12。

- (二)第二階段調查採用滾雪球抽樣法，由第一階段部會指派人員推薦中央各部會的資訊、相關業務及人事單位科長級以上主管1名後，再由各主管推薦相關業務同仁2名參與調查。若該部會無與調查目的業務相關單位，但附屬三級機關有業務相關單位

者，可邀請最相關業務單位之科長級以上主管及該主管指定相關業務同仁至少 2 名受訪。

五、問卷設計

本計畫團隊依據調查研究目的以及調查研究問題，設計兩份調查問卷，第一階段調查問卷著重在瞭解中央各部會應用機器演算法創新公共服務的現況，調查焦點包括：部會應用 AI 系統的導入規劃、資訊人力規模、資訊預算規模、法規調適現況、創新公共服務類型以及服務所需資料來源等。第二階段調查著重在瞭解中央各部會文官對於應用機器演算法創新公共服務的態度，調查焦點包括：文官參與行政革新經驗、部會資料準備度、部會管理系統、文官心理資本、公共服務動機和人才職能等構面。本計畫團隊於 110 年 7 月 15 日邀請 P1、P2、P3、P4、P5（詳見第三章第一節表 11 名單）參加本調查設計的專家座談，討論問卷設計構想以及初版問卷內容。經採擷與會專家寶貴建議並修改問卷內容後，計畫團隊於同年 8 月 5 日至 12 日邀請國發會 8 位同仁參加問卷前測，並在 13 日辦理前測參與者座談會討論前測問卷內容，並依據前測座談會議意見修訂最終版問卷，茲將本調查問卷設計架構描繪如圖 17，各調查構面的操作化定義以及對應之正式調查問卷題號，整理如表 22 所示，正式調查問卷內容請參閱附錄四。

第一階段調查

機關現況：AI 系統導入規劃、人力規模、預算規模、法規調適現況、公共任務類型、應用領域

第二階段調查

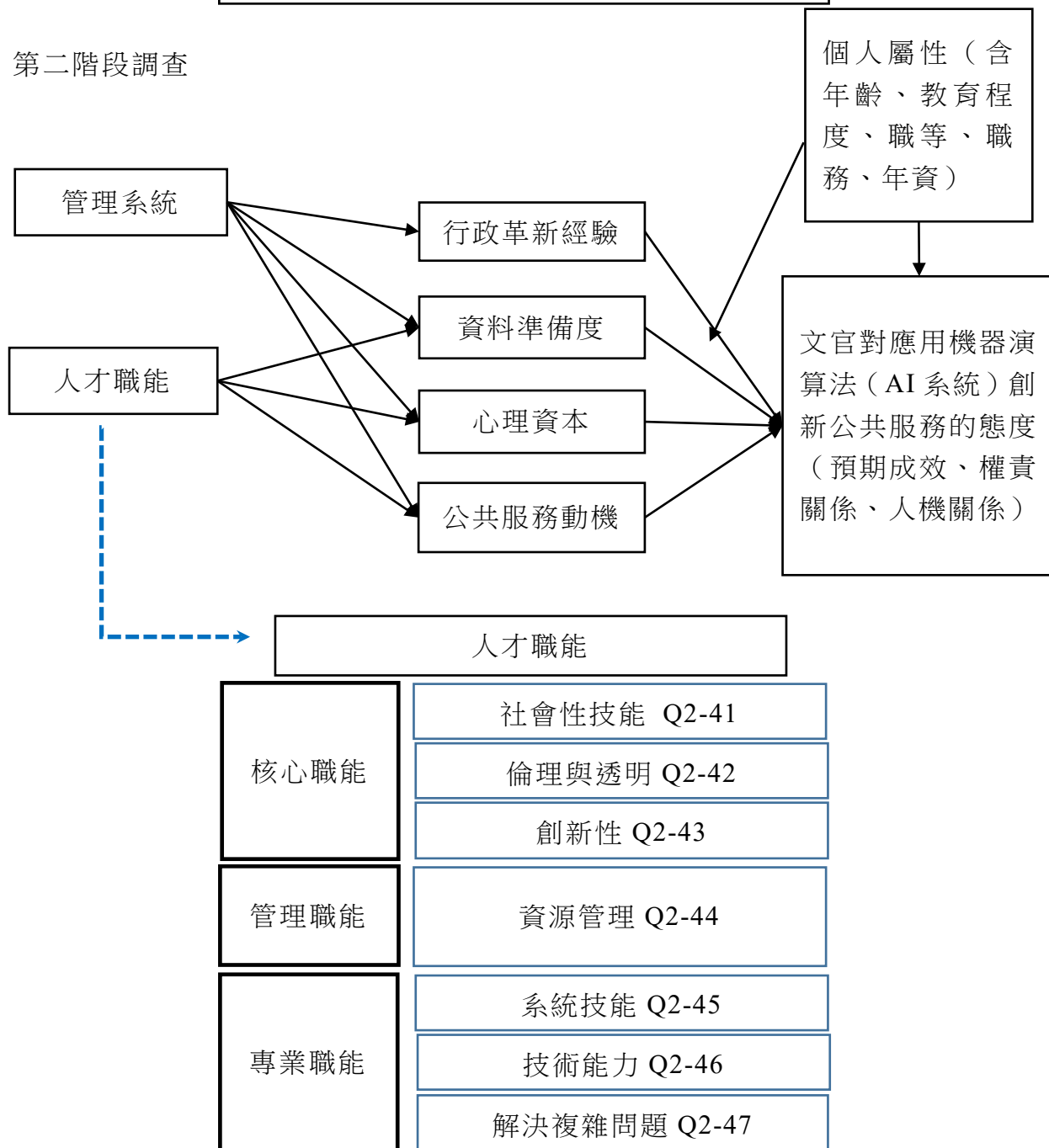


圖 17：調查問卷設計架構圖

資料來源：本計畫。

表22：子計畫一調查構面、操作化定義及題號一覽表

調查構面	操作化定義	題號
機關現況	本計畫將應用 AI 科技於公共服務的電腦資訊系統(Artificial Intelligent System, 以下簡稱 AI 系統) 定義為一種結合數位資料、機器學習演算法以及人類經驗而成的一種全自動化或半自動化輔助行政業務的電腦資訊系統。本調查針對中央部會及相當於中央二級機關之獨立機關，推動 AI 系統輔助行政業務的情況，包括系統導入及規劃、參與人力規模、預算規模、法規調適現況、應用行政業務類型及應用範疇。	1-1～1-27
文官運用機器演算法（AI）創新公共服務的態度	依據相關文獻的討論，本計畫從以下三個構面測量文官對機器演算法（AI）創新公共服務的態度，分別是預期成效、權責關係、人機協作關係。	2-1～2-12
心理資本	個人克服行政革新逆境的能力，包含四個次構面，自我效能感、希望、韌性、科技樂觀。	2-13～2-24
公共服務動機	個人從事公共服務的動機。	2-25～2-29
資料準備度	個人所屬機關擁有 AI 系統所需資料的資料準備度，包括資料完整度、機器可讀程度、資料正確性、資料近用度以及資料使用歷程紀錄完整度。	2-30～2-34
管理系統	個人感知行政機關的管理體系對於行政革新的支持程度。	2-35～2-40
人才職能	個人對參與機器演算法任務所需人力職能的想法，包含三個次構面，核心職能、管理職能、專業職能。	2-41～2-47
行政革新經驗	個人對所屬機關在過去三年改善業務及創新公共服務的觀察。	2-48～2-53
個人基本資料	受訪者年齡、教育程度、職等、職務、年資等基本資料。	2-54～2-61
開放建議題	徵詢受訪者對於政府推動數位化政策的想法與建議。	2-62

資料來源：本計畫。

第三節 法務部保護司觀護人之實驗調查

一、 研究目的

延續前述的研究核心，我們正處在資料和資訊科技以無所不在的型態存在於生活各層面的時代。結合現今網路計算能力以及複雜的演算法工具，人們相信將能進一步開發、使用、分析與預測更海量且更雜亂無結構的資料，期盼能替人類探索出未知的知識和更深刻的洞察力。近年來不僅是世界各國，我國各級政府與機關也已逐漸開始考慮，甚至採用內建機器運算的數位服務，如：司法院運用統計科學與量刑資訊，建構的量刑趨勢建議系統、臺北市政府交通局的智慧化交通量調查分析系統、臺中水滴智慧城花博園區的自駕公車、健保署去識別化的醫療影像應用，以及空汙通報與環保實地稽查的智慧環境監控等。但當演算法導入公部門數位治理時，我們應反覆的思考新興技術所帶來的潛在影響。這當中或許存在著舊的議題（如：數位政府一直以為需面對的法規限制、數位落差、組織抗拒、組織學習、績效管理等問題），但也包括了許多對於組織與個人的新挑戰。就制度論觀點來看，除了法規鬆綁與調適的需求與可行性外，公務部門人員的組織與個人決策是否因為演算法的應用而產生質變，進而改變公部門的角色與服務任務，也是不容忽視的問題。

因此混合前述機關調查與專家訪談等方法，此節提出若要檢驗公部門組織中個人層次對執行業務的影響，便需搭配調查實驗方法，才能更進一步理解從實際工作流程的設計上，哪些環節是人機協作的決策點。

二、 研究對象

為能回應上述研究目的，以法務部保護司為研究對象，設計一模擬演算法系統的調查實驗平臺並納入實際的決策情境。檢視在有無演算系統的輔助下，觀護人對個案進行再犯評估風險的判斷及決策擬定後處理措施之差異，希冀能從中理解演算法應用對公部門實際的好處、顧慮、與決策困境等。挑選法務部保護司作為研究個案，秉持質量並重精神，將採用訪談與田野實驗的混合方法來研究此個案。基於長年來工作業務龐大、人力短缺的原因，法務部保護司目前正在積極思考導入「受保護管束者，再犯風險評估」的輔助決策工具。本計畫團隊認為此個案與本計畫之課題主旨契合，即使目前尚未有成熟之機器學習產品導入業務當中，但得以讓我們在公部門機構評估導入演算法內嵌之決策系統或平臺工具的初期，就先以預評估為目標，檢視機關觀護人對演算法系統介入案件評估決策前後的態度、期望、與效果。本計畫先以專家座談或深度訪談，瞭解該單位組織與個人決策的實際狀況、工作流程與數位轉型應用的需求與想法。根據訪談資料，本團隊將接續設計較符合現實的實驗情境、實驗介入與實驗前後的問卷調查。先找小規模的立意樣本利用前測（pilot study）確認實

驗之信度與效度，再行正式施測。

全國目前有229位觀護人，但當中含有19位主任觀護人，由於主任觀護人大多沒有執行保護管束案件，將其扣除於樣本，因此調查樣本預估全部有200位觀護人。由於本團隊對調查的對象沒有強制力，僅能就由催覆與少額調查獎勵的方式，鼓勵觀護人參與本計畫。每位參與者將會收到一個獨立的調查問卷連結，藉此進行樣本的隨機分組，參與實驗之觀護人在點選自己的網頁連結後，就會引導到實驗組與控制組的各自調查頁面。以下繼續說明實驗組與控制組的實驗設計流程。

三、 調查實驗設計流程

該實驗調查之實驗參與者將被隨機分配數個假釋案件，共計 18 個個案，其中案例一至案例十為「有再犯撤銷保護管束」之情形；案例十一至案例十八為「未再犯保護管束中」之情形，以對社會安全的威脅程度（高度、中度）作為屬性因子，將歷史個案分類。為了確保每個案子具有某些程度的難度，本計畫不考慮低度的個案，如此也能避免受訪者只是反射性的處理簡單的個案。每位受試者都會隨機分配到有再犯與未再犯之各兩個案例。來進行判斷，看完每個案例後，受試者須評估個案的危險級別以及再犯風險級別（分為極低度、低度、中度、中高度與極高度），並針對個案級別作出對應的觀護處遇評估。

本計畫設計實驗組與控制組。控制組完全由受試者依據所給予的參考資料，自行做個案的判斷決定。但同樣會給予控制組第二次確認自己答案的機會，因為根據我們日常生活與工作上作決策的習慣，再三考慮更符合現實。對於控制組受訪者，在他們作答完後，我們同樣問他們對自己的最終判斷的正確度有多高的信心。

實驗組分別有實驗組第一組（Arm 1）與實驗組第二組（Arm 2）共兩組，透過隨機分組設計（Randomized Controlled Design），各組樣本數各有 1/3。被分派至實驗組的參與者，在決策判斷過程會增加外部資訊介入，Arm 1 會接收到演算法預測系統所提供的風險機率與處遇建議；Arm 2 則會加入其他參與此實驗觀護人的綜合評估分數作為外部資訊介入操作。在外部資訊介入之前，受試者會先針對分配到的個案，進行第一次的評估，評估其危險級別、再犯風險級別以及觀護處遇；受試者評估後，則會呈現外部資訊給受試者參考，並且詢問受試者是否更改第一次的評估結果，若受試者認為第一次評估結果需要修正，則會讓受試者進行第二次評估。另外，第二次評估後進一步詢問受試者對於外部介入資訊的正確度有多高的信心。本實驗研究設計與流程，請見以下圖 18。

。

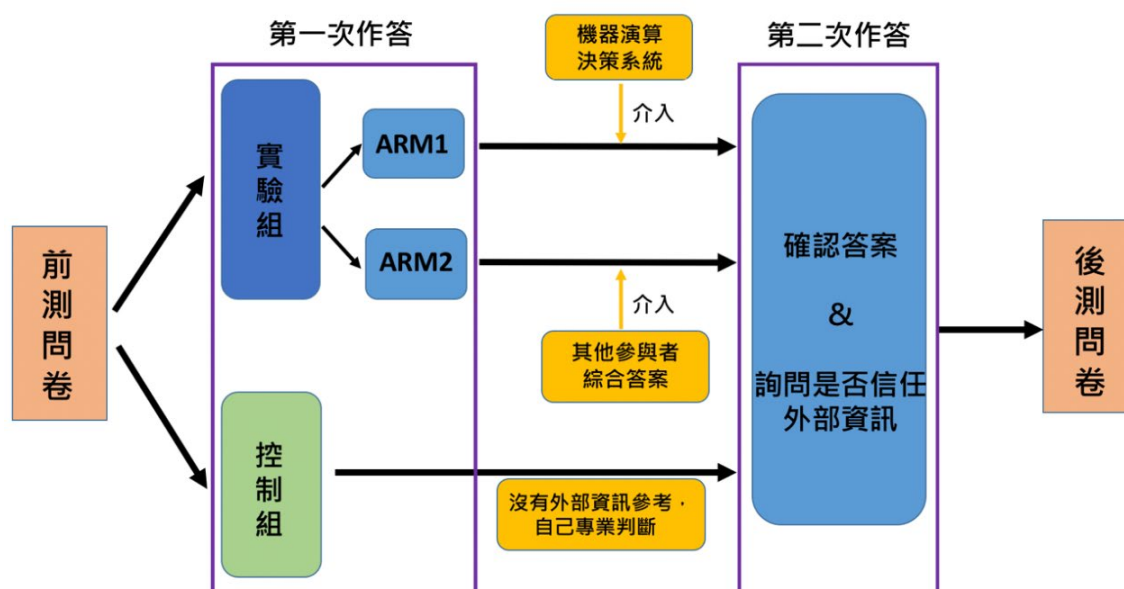


圖 18：子計畫一實驗調查流程示意圖

資料來源：本計畫。

為使上述之實驗設計與流程更加順利與完善，本計畫於控制組、Arm 1、Arm 2之實驗之時程安排也有先後順序的安排。首先，先進行控制組與實驗組Arm1實驗，在蒐集兩組受試者資料後彙整受測者的作答結果，並將此資料作為實驗組Arm2外部介入的參考資訊。下一階段則進行實驗組Arm2的實驗，在實驗組Arm2受試者完成第一次答題後，本團隊會提供其他工作同仁評估該個案之再犯風險級別以及給予觀護處遇作為的作答比例，觀察該組受試者是否會因他人的判斷而改變自己的答案。

本計畫於110年7月開始聯繫製作平臺之廠商，以及設計問卷內容，為使實驗調查更加周全與貼近觀護人的工作日常與處理流程，本團隊特與法務部保護司觀護人請益，邀請觀護人協助提供實驗所須資料以及針對問卷之設計內容，尤其是實驗階段之評估設計給予建議。本團隊於110年8月16日召開專家座談，邀請國內專家學者針對問卷內容提供修正意見，並在110年10月作初步測試以及受試者資料之整理，觀護人實驗之前期規劃時程表如表23所示。

表23：子計畫一觀護人實驗前期規劃時程表

	日期	內容	說明
1	110/07/02	聯繫廠商	與竝盛科技確認 oTree 平臺設計以及後續規劃事宜。

	日期	內容	說明
2	110/07/15	實驗問卷題目初版設計	完成觀護人實驗之問卷初版。
3	110/07/27-7/30	設計觀護人評估與處遇問題	參考觀護人的評估與處遇措施，於實驗中設計相關問題。
4	110/08/06	法務部觀護人第一次會議	與法務部保護司兩位觀護人接洽，請其協助後續實驗之相關事宜。
5	110/08/10	整理觀護人問卷回饋	整理協助本次實驗之觀護人對於實驗問卷之回饋。
6	110/08/14	整理個案危險、再犯、處遇級別與作為	針對觀護人所給予每個個案之危險、再犯、處遇級別與作為，設計 AI 判斷正確與錯誤之答案。
7	110/08/16	舉辦專家座談	邀請國內學者與法務部觀護人參加座談，針對實驗問卷與內容及流程設計給予意見回饋。
8	110/08/20	法務部觀護人第二次會議	商討後續實驗須調整事宜。
9	110/09/13	實驗平臺初稿	廠商給予實驗平臺初稿，由研究團隊作初步測試後給予回饋。
10	110/09/16	實驗問卷題目定稿	實驗問卷題目、內容確定。
11	110/10/5-10/14	實驗平臺測試	邀請國立政治大學公共行政學系研究所學生協助作實驗平臺測試，並請參與者給予回饋。

	日期	內容	說明
12	110/10/14-11/2	整理測試者回饋、修正問卷內容、平臺系統優化	整理平臺測試之回饋，並與廠商溝通後續調整事宜。
13	110/10/27	整理觀護人資料	整理我國 200 位觀護人之相關資料，按照性別、年資比例隨機作控制分組。

資料來源：本計畫。

經過問卷內容與系統修正，以及蒐集觀護人資料並作初步分析後，確立後續的研究時程。本團隊於110年11月8日開始控制組以及Arm 1的實驗，並在110年11月12日蒐集且整理初步作答資訊，作為Arm 2的外部介入資訊，該組於110年11月19日開始實驗。整個實驗期程於110年12月4日結束，緊接後續資料彙整與分析。確切規劃請詳見表24。

表24：子計畫一觀護人實驗收案之時程表

	日期	內容	說明
1	110/11/8	寄出控制組、Arm 1 邀請受試信件	將實驗邀請信件以及實驗連結寄送予控制組以及 Arm 1。控制組與 Arm 1 預計實驗期間為 11/8 至 11/19。
2	110/11/12	初步整理填答資料、控制組、Arm 1 第一次催稽	整理當前完成控制組以及 Arm 1 實驗的答案並作敘述統計，作為 Arm 2 實驗的外部介入資訊；寄信予提醒尚未完成實驗的受試者作答。
3	110/11/19	寄出 Arm 2 邀請受試信件	將實驗邀請信件以及實驗連結寄送予 Arm 2。Arm 2 實驗期間為 11/19 至 12/4。

	日期	內容	說明
4	110/11/22	控制組、Arm 1 第二次催稽、Arm 2 第一次催稽	寄信予提醒尚未完成實驗的受試者作答。
5	110/11/29	Arm 2 第二次催稽	寄信予提醒尚未完成實驗的受試者作答。
6	110/12/4	實驗結束及資料彙整	彙整受試者填答之資料。
7	110/12 月-111/1 月	受試者作答分析	完整分析作答狀況，觀察受試者是否會因外部資訊的介入而更改答案。
8	110/12/25	寄送電子禮券	寄送電子禮券予完成實驗調查之受試者。

資料來源：本計畫。

四、 調查工具與解釋變項介紹

(一) 調查工具規劃

研究團隊與竝盛科技合作，委託該廠商透過oTree協助設計線上實驗調查平臺。oTree是由程式語言Python所建構的軟體，近年來經常被運用在心理學與社會科學相關的實驗。oTree平臺擅於設計行為控制（behavioral control）相關的實驗，設計者可利用oTree作隨機分組，也能夠設計如囚徒困境（Prisoner's Dilemma）等團體同步式實驗。本實驗將由保護司協助作樣本的隨機分組，並發送專屬的網路調查連結碼給每一位受訪者。

受試者可於網頁上操作平臺，圖19為實驗之頁面，顯示個案相關資訊，並請他們於右方勾選認為合適的選項。在初次作答完後，實驗組Arm 1 以及Arm 2的受試者會分別接收到「虛構」的AI所計算出的答案以及先前測試者各題作答比例分布之資訊，並再次詢問是否要更改答案。

第 1 個案例

請依照這裡呈現的個案資料，

回答以下有關該個案之風險評估與處遇措施的相關問題：

案例	
性別	男
生日	84年5月（現年26歲，犯案時24歲）
主要案由	毒品危害防制條例（販賣三級毒品K他命）
刑期	1年10月，緩刑5年付保護管束
保護管束類型	緩刑
保護管束期間	109年12月11日起至114年12月10日止
學歷	國中畢業
家庭狀況	案父母離異，案母已再婚，案父因殺人罪入監服刑中，個案上有一個姐姐已婚，個案現與祖父母同住，家中經濟狀況不佳，個案有積欠債務約10多萬，家庭支持薄弱，表妹原同住，後來搬出去在外租屋。
交友狀況	交友較複雜，較常往來的家人是表妹（於當鋪工作）。
工作情形	前曾於租車店工作，目前於工地綁繩。
前科紀錄	本案緩刑前於108年間犯有一件交通過失致死案件被判處有期徒刑4月易科罰金，無入監服刑紀錄，無施用毒品前科。

B1. 該個案犯罪行為的產生受到不良的社會影響的程度為何？
☐ 極低度影響 ☐ 低度影響 ☐ 中度影響 ☐ 中高度影響
☐ 極高度影響 ☐ 無法判斷

B2. 個案在社會中受到排斥且缺乏對他人的關心，產生與社會疏離的程度為何？
☐ 極低度疏離 ☐ 低度疏離 ☐ 中度疏離 ☐ 中高度疏離
☐ 極高度疏離 ☐ 無法判斷

B3. 個案在自我特質上，因缺乏同理心與負面情緒，導致有問題解決技巧的認知不良情形的程度為何？
☐ 極低度認知不良 ☐ 低度認知不良 ☐ 中度認知不良
☐ 中高度認知不良 ☐ 極高度認知不良 ☐ 無法判斷

B4. 個案願意接受觀護與治療的配合程度為何？
☐ 極低度配合 ☐ 低度配合 ☐ 中度配合 ☐ 中高度配合
☐ 極高度配合 ☐ 無法判斷

B5. 請問個案的危險級別為何？
☐ 0-20%極低度危險 ☐ 21-40%低度危險
☐ 41-60%中度危險 ☐ 61-80%中高度危險
☐ 81-100%極高度危險

B6. 請問個案的再犯風險級別為何？
☐ 0-20%極低度危險 ☐ 21-40%低度危險
☐ 41-60%中度危險 ☐ 61-80%中高度危險
☐ 81-100%極高度危險

B7. 針對個案之危險級別，觀護處選評估為第幾級？
☐ (1) 第一級（高度監督及輔導，以核心案件列管，實施特殊處遇）
☐ (2) 第二級（中度監督及輔導，尚需加強列管、輔導或其他事由以核心案件列管）
☐ (3) 第三級（低度監督及輔導）

B8. 針對個案應實施之一般處遇作為(複選題)：
☐ (1) 約談、訪視
☐ (2) 列管核心案件加強報到
☐ (3) 採驗尿液
☐ (4) 需用複數監督
☐ (5) 命接受毒品或酒癮戒癮治療
☐ (6) 法治教育
☐ (7) 轉介就業服務
☐ (8) 家庭支持
☐ (9) 指定榮譽協助訪談或訪視
☐ (10) 命其向警局（派出所）報到
☐ (11) 定期或不定期電話訪
☐ (12) 其他：

B9. 針對個案建議之特殊處遇作為(複選題)：
☐ (1) 實施科技設備監控
☐ (2) 指定居住處所
☐ (3) 監控時段，未經許可不得外出
☐ (4) 測試
☐ (5) 轉介適當機構或團體
☐ (6) 心理測驗
☐ (7) 禁止接近特定處所或對象
☐ (8) 命其主動打電話回報觀護人
☐ (9) 命其定期接受身心治療或輔導
☐ (10) 命定時（每日/每週/每月）向警局（派出所）報到
☐ (11) 其他必要特殊處遇，如：禁止上色情網站、禁止網路交友、不得飲酒、不得吸食毒品等：

Next

圖 19：子計畫一實驗之操作頁面圖示

資料來源：本計畫。

Parole-ML-AI 1.0

我們的Parole-ML-AI 1.0系統運算已結束，分析給出的預測結果如下：

本案過去未有使用毒品相關之犯罪前科，此次為首次入監服刑，然其家庭經濟狀況不穩定，且交友複雜。根據個案主案由、家庭與交友狀況、工作情形與前科紀錄，透過Parole-ML-AI 1.0進行判斷。

[查看案例資料](#)

根據Parole-ML-AI 1.0運算結果顯示：

社會影響程度	中高度影響
與社會疏離程度	中度疏離
認知情形	中度認知不良
願意配合觀護與治療程度	中高度配合
危險級別預測	中高度危險 (61%~80%的風險分數)
再犯風險程度預測	中度危險 (41%~60%的風險分數)
處遇評估的預測	第一級 (20%~30%的符合機率) <u>第二級 (90%~95%的符合機率)</u> 第三級 (10%~15%的符合機率)
一般處遇作為建議	1、約談、訪視；2.列管核心案件加強報到；3、警局複數監督；4、法治教育；5、家庭支持；6、命其向警局（派出所）報到。
特殊處遇作為建議	1、命定時（每月）向警局（派出所）報到。

您原本的答案

- B1. 該個案犯罪行為的產生受到不良的社會影響的程度為何？
☐ 極低度影響 ☒ 低度影響 ☐ 中度影響 ☐ 中高度影響 ☐ 極高度影響 ☐ 無法判斷
- B2. 個案在社會中受到排斥且缺乏對他人的關心，產生與社會疏離的程度為何？
☐ 極低度疏離 ☒ 低度疏離 ☐ 中度疏離 ☐ 中高度疏離 ☐ 極高度疏離 ☐ 無法判斷
- B3. 個案在自我特質上，因缺乏同理心與負面情緒，導致有問題解決技巧的認知不良情形的程度為何？
☐ 極低度認知不良 ☒ 低度認知不良 ☐ 中度認知不良 ☐ 中高度認知不良 ☐ 極高度認知不良 ☐ 無法判斷
- B4. 個案願意接受觀護與治療的配合程度為何？
☐ 極低度配合 ☐ 低度配合 ☐ 中度配合 ☐ 中高度配合 ☐ 極高度配合 ☒ 無法判斷
- B5. 請問個案的危險級別為何？
☐ 0-20%極低度危險 ☐ 21-40%低度危險 ☒ 41-60%中度危險 ☐ 61-80%中高度危險 ☐ 81-100%極高度危險
- B6. 請問個案的再犯風險級別為何？
☐ 0-20%極低度危險 ☐ 21-40%低度危險 ☒ 41-60%中度危險 ☐ 61-80%中高度危險 ☐ 81-100%極高度危險
- B7. 針對個案之危險級別，觀護處遇評估為第幾級？
☐ (1) 第一級（高度監督及輔導，以核心案件列管，實施特殊處遇）
☒ (2) 第二級（中度監督及輔導，尚需加強列管、輔導或其他事由以核心案件列管）
☐ (3) 第三級（低度監督及輔導）
- B8. 針對個案應實施之一般處遇作為(複選題)：
- ☐ (1) 約談、訪視
☒ (2) 列管核心案件加強報到
☐ (3) 採驗尿液
☐ (4) 警局複數監督
☐ (5) 命接受毒品或酒癮戒癮治療
☒ (6) 法治教育
☐ (7) 轉介就業服務
☐ (8) 家庭支持
☐ (9) 指定榮譽協助約談或訪視
☐ (10) 命其向警局（派出所）報到
☐ (11) 定期或不定期電訪
☐ (12) 其他：
- B9. 針對個案建議之特殊處遇作為(複選題)：
- ☐ (1) 實施科技設備監控
☐ (2) 指定居住處所
☐ (3) 監控時段，未經許可不得外出
☐ (4) 測謊
☐ (5) 轉介適當機構或團體
☒ (6) 心理測驗
☐ (7) 禁止接近特定處所或對象
☐ (8) 命其主動打電話回報觀護人
☐ (9) 命其定期接受身心治療或輔導
☐ (10) 命定時（每日/每週/每月）向警局（派出所）報到
☐ (11) 其他必要特殊處遇，如：禁止上色情網站、禁止網路交友、不得飲酒、不得吸食毒品等：

B10.B 以下顯示為您對此個案的判斷結果摘要，請問您是否確定此決定，還是需要進行修改

☐ 是，需要修改。 ☐ 否，不需要修改。

圖 20：子計畫一實驗組 Arm 1 之確認答案圖示

資料來源：本計畫。

(二) 調查中重要的解釋變項

以下說明本次實驗調查問卷題組之重要概念。正式的問卷已附在報告附錄五。

1、期望地位理論

Berger等於1977年提出之期望地位理論，他們研究探討團隊中個人對自己與他人的期望與不同角色互動關係的影響，其中主要的兩個核心概念為「績效期望」(performance expectation)與「地位特徵」(status characteristic)，績效期望在此是指個人未來能完成工作的一種信念；而地位特徵中專注角色間的地位狀態，在團隊中的角色被定義為相對的「自我」與「他人」，人們會根據自我與他人間期望的差異來評價團隊中的彼此，當在需要團隊合作的情形下，根據自我評價的結果作出不同的選擇(Berger, 1977; Wagner & Berger, 1993)。他們的研究發現，對自己有較高期望的受試者通常會無視或拒絕夥伴的回應；相反的，對自我期望較低的受試者，則傾向會接受夥伴的意見。Foschi (1985)亦根據Berger等人所提出的期望地位理論為基礎，設計了相關實驗，他指出個人是基於「自我標準」來推斷自己與他人是否具有完成工作或任務的能力，基於能力的不同也因此會產生期望差異，當自己擁有較佳能力時，也相對會對自己抱有較強烈的高期望。以本計畫來說，人機協作在演算法之演算越加強大的未來，上述對期望地位理論的討論，說明此概念適合用來探討人類對自身與對機器在完成某類任務時的勝任程度。而這個認知落差，可能是影響實際人機協作時，人類是否相信機器預測與結果的重要因子。故本計畫在問卷中將納入此概念，並設計相應的問卷題目。

2、創新意願

此外，使用新興科技於職場上，常攸關於使用者對於創新的接受程度，及其本身對於風險的承受意願。因此，在問卷上也參考自《臺灣政府文官調查第四期，TGBS IV》¹⁵之問卷，設計有關「創新接受度」與「風險承接度」之題組。其創新接受度之題目設計，參考自National Administrative Studies Project – Decision Making (NASP - DM) (Dong, 2013) ¹⁶。

15 董祥開 (2015)。公務人員風險偏好、服務動機、與決策類型之關係－建立一個勇於任事的政府(臺灣政府文官調查第四期，TGBS IV), 2021 年 7 月 1 日。取自：https://srda.sinica.edu.tw/datasearch_detail.php?id=2913。

16 Center for Organization Research and Design (2016)。The National Administrative Studies Project-Citizen Survey Decision Making (NASP-DM), 2021 年 7 月 1 日。取自：<https://cord.asu.edu/node/581>。

(三) 控制變項

1、科技接受題組

過去在談組織E化或是資訊開放接受與導入的相關研究，多數都會納入科技接受模型（Technology Acceptance Model，簡稱TAM），來研究使用者對新資訊科技系統導入的使用意圖。文獻中指出這個意圖往往由兩種信念所決定，分別是：認知有用（perceived usefulness）與認知易用（perceived ease of use），前者是指使用者相信系統會提高其工作績效；後者是指使用者認為系統是否易於使用，也就是說使用者在使用的過程中不需要耗費額外的精神，即可上手（Davis, 1989; Davis, Bagozzi, & Warshaw, 1989）。本計畫認為TAM題組可以測驗組織個人對於科技樂觀的程度，也因此可能影響受訪者對演算法工具的看法與互動時的採納度。因此擷取該模型中的有用性、易用性、與使用意願三個概念，根據演算法科技的情境，修改並發展成客製的問卷題目。

2、公共服務動機

科技突飛猛進的發展為公共行政帶來了電子治理（E- governance）及數位治理（Digital Governance）等觀念，學者認為資通訊科技（ICT）引入公部門將能帶來提升政府課責性（accountability）、透明度（transparency）以及公民參與（citizen participation）等優點（Hardy & Williams, 2008; Alshawi & Alalwany, 2009; Boateng, 2013），並能夠帶動公部門的組織及功能上的變革（Twizeyimana & Andersson, 2019），進而改善公共服務的推行（Osei-Kojo, 2016）。

Perry與Wise（1990）定義「公共服務動機」（Public Service Motivation, PSM）為「基於回應公共利益或組織管理問題的傾向」，而Perry（1996）日後更進一步提出了公共服務動機的四個構面：「對政策制定的關注」、「對公共利益的承諾」、「自我犧牲」以及「同情心」。Moore（2005）認為公部門若能學習並採納新興的科技、政策、組織法規等觀念能夠提升政府的效率及效能。基於上述觀點，政府應與時俱進，援引新的技術及觀念於政策的推動上，方能將公共利益的價值最大化。故本計畫透過問卷設計相關題目，以控制住機關組織內公共服務動機與對科技及外部資訊接受度的關聯性。

五、 功能測試與調整

(一) 焦點座談

研究團隊於110年8月16日利用線上會議平臺舉辦焦點座談，邀請法務部保護司觀護人及諸位大學教授共六位參與座談，為本計畫導入機器學習演算系統實驗調查設計規劃提供諮詢及意見。

本次座談訪談題綱如下：

- 1、 本計畫之實驗流程設計有何需要注意之處？
- 2、 針對本問卷設計之內容，請問有何需要修改增進之處？
- 3、 其他開放式請益，或與會專家意見的補充。

與會專家針對本實驗調查所給予回饋之主要面向大致有六，分別有：受試者對該系統之熟悉程度、實驗組別設計問題、公共服務動機（Public Service Motivation, PSM）題組設計、公共服務導入機器學習之目的、受試者資歷以及引用其他學者之研究著作，以下分別作說明：

- 1、 受試者對該系統之熟悉程度：使用者對於該科技或機器之信任度會隨著系統品質（system quality）的高低而有所改變。若此演算法系統品質高，則使用者之信任度將隨之提升，若品質低落則反之。然而就調查對象及時間點，演算法系統尚未真正導入實務工作場域，受試者對機器演算法的認知有限，恐影響最後的作答。就此與會專家建議研究團隊應清楚界定探究信任度之目的及核心概念。
- 2、 實驗組別設計：有專家建議有機器決策介入的組別（Arm1）可先作調查，或請退役或資歷豐富之觀護人事前作測試，將實驗結果作為外部同仁介入的組別（Arm2）的參考資料。另也有針對控制組兩次作答時間間隔長短給予建議。
- 3、 公共服務動機題組設計：與會專家表示問卷中探究公共服務動機之變項對實驗結果的影響不明，並有專家提供以下建議：
 - (1) PSM 應著重「心理歷程」而非「表現出來的行為」。
 - (2) 以心理學觀點，motivation 與 motive 不同。「Motivation」可分為「生理的」（physiological motivation）與「心理的」（psychological motivation）兩類，生理的 motivation 都是由於生理組織缺乏某種物質（即需要）所引發，因此心理學家習慣稱此類 motivation 稱作「驅力」（drive）。相對地，心理的 motivation 則被稱為「動機」（motive）。兩者綜合而言，則統稱為「motivation」（張春興，2000）。
 - (3) Predisposition 定義較趨向一種生物「習性」，難以作改變。公共服務激勵是「一種用來回應個人動機（motives）的習性（predisposition），而這個動機是深植於公共制度或公共組織所普遍存在或是所獨有的」。
- 4、 公共服務導入機器學習之目的：與會專家引用黃俊能（2021）等所作再犯風險評估智慧輔助系統之研究，認為公部門導入機

器學習之目的究竟是為瞭解決簡單但數量龐大、社會影響小、勞力程度高之業務以減輕承辦人員之負擔，抑或是協助社會影響大、困難之案件應界定明確。

- 5、受試者資歷：與會專家認為研究團隊應清楚說明問卷調查受試者資歷背景之目的以及與研究本身的關聯性，並好奇隨著年資、經驗、學歷背景等不同，參與測試之觀護人所作出之觀護處遇評估措施會因此存在差異，又或是否應該聚焦調查有經驗之觀護人，以提升研究之效度。
- 6、引用其他學者之研究著作：除了建議參考黃俊能老師的研究，亦推薦研究團隊參閱阮怡凱（2020）從巨量資料分析之觀點出發，探討受保護管束者再犯風險評估工具之研究，並引述作者對於最終決策仍應交由人力裁定之建議。

（二）測試調查

在110年10月5日至110年10月6日期間，邀請國立政治大學公共行政所碩博士生以及法務部保護司兩位觀護人上網填寫測試問卷，共回收20份有效樣本，研究團隊針對回收之填答內容與建議後，隨之調整問卷內容。

測試後測試者建議調整內容如下：

1. 研究團隊聯絡資訊應更完整。
2. 問卷錯字、字間空格修訂。
3. 問卷有些用字遣詞易使作答者混淆。
4. 實驗平臺排版應作調整，使視覺上更美觀。
5. 實驗階段在作答同一個案時，第一次作答與第二次作答所跳出來的案例不同。
6. 填答時系統發生錯誤。

（三）研究限制

本實驗研究的重點是「人機協作」下的才能管理問題，不過，法務部這各想像功能之人工智慧應用，不一定是系統功能主體。甚至有時名稱是人工智慧的方案未必具備完整的人工智慧功能，對非技術背景人員其實不易判定，不論在中央部會的調查與法務部的實驗都有相同問題，因此，本計畫受訪對象的理解程度的高低，與其他所有經驗研究都會面對到訪查對象「知識落差」的現實，在解讀資料時必須審慎處理。

第四章 調查分析發現

第三章說明本計畫之調查設計與方法，針對法規調適、資源配置與才能管理三種途徑，透過文獻分析、深度訪談、焦點座談、問卷調查與實驗調查之研究方法加以蒐集資料與分析。希望透過多元的研究方法，可以兼具廣度與深度的問題發現與釐清、解方的發想與建置。

第四章的調查分析發現，針對文獻分析、深度訪談、焦點座談以及中央各部會之問卷調查，依據前述三種策略途徑，呈現現階段的法規調適、資源配置與才能管理三層面的調查分析發現。法規調適的資料來源主要來自於文獻分析、深度訪談與焦點座談方式蒐集資料，採擷國外經驗，從組織規劃定位、規制方式制度、倫理價值與政策原則以及風險評估分級，分析其國外經驗與國內面臨之問題；資源配置主要從中央各部會運用演算法之問卷調查，透過二階段的分析，從組織層面與個人層面觀之現今部會的推動情形、遭遇困境、所需條件等；才能管理層面從中央部會的問卷調查、深度訪談與焦點座談分析其人機協作職能的應然與實然的落差。

本章將從第一章中提到的制度分析三個層次來分節：第一層，法規調適之跨國分析、第二層，資源配置之問卷分析、第三層，才能管理之實政分析，從資料中尋找問題，並且統整於第四節的落差分析。

第一節 法規調適之跨國分析

AI過往多以「規則基礎演算法」(Rule-based algorithms)為其運作基礎，能夠通過自動執行程式規則來執行複雜的任務。然而近期則逐漸演進為機器學習演算法(Machine learning algorithms)(EU Commission, 2021b)。傳統程式編寫計算語言，設置處理資料輸入後所需的規則，以獲得輸出的答案；但在機器學習中，電腦接收輸入資料以及從資料中期望得到的答案，透過機器學習程序由電腦自己生成規則，再將這些規則應用於新資料以產生原先的答案。機器學習系統經過「訓練」而不是「編寫程式」，而為了能成功地「學習」，許多機器學習系統需要大量計算能力和大數據資料庫(Big Data)，但目前多數AI通常包括以上兩種演算法(EU Commission, 2021c)。圖21呈現傳統規則基礎演算法與機器學習演算法概念的不同。

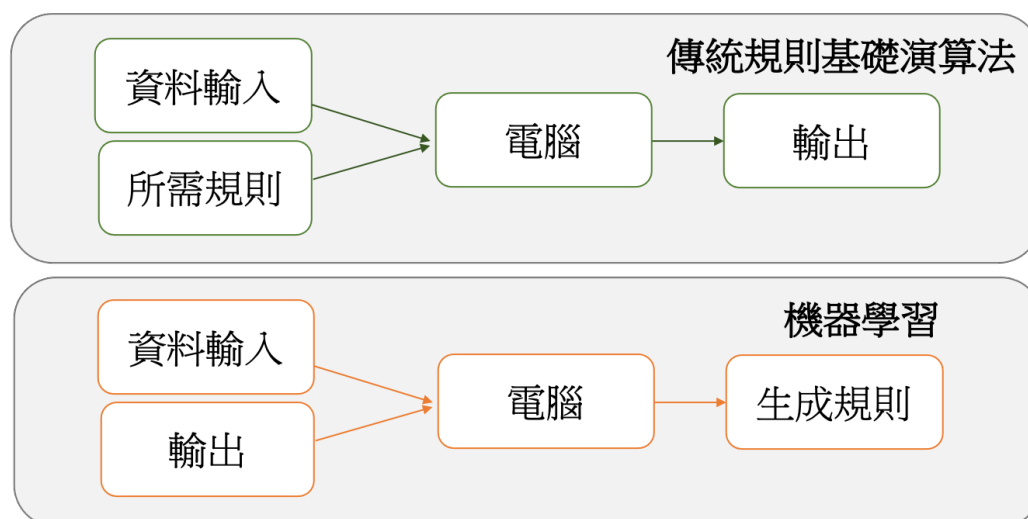


圖 21：傳統程式演算法與機器學習演算法之差異

資料來源：翻譯修改自 EU Commission (2021c)。

AI 的使用在許多領域產生重要突破，透過改進預測、優化運作和資源分配以及個別化服務，AI 在醫療保健、農業、教育、基礎設施管理、運輸和物流、公共服務、安全和氣候之使用對經濟社會和環境帶來助益，幫助解決公共利益的複雜問題，並提供企業競爭優勢。然而，在 AI 為社會經濟帶來助益的同時，當然也帶來相當程度之風險與負面效果。AI 之發展免不了前階段由開發人員或操作者控制或設計之演算法與機器學習，此時若資料控制者或程式設計者有意操控相關參數或演算法內容，則將對使用者或資料主體之基本權利、個人資料安全等產生重大負面影響。因此，有鑑於技術快速變革與資料廣泛流通與使用，需要建構與演算法相關高度隱私與資料安全之監管框架及道德標準。

本計畫所盤點各國對於 AI 規範，發現各國在發展 AI 的過程中，均遇到 AI 所扮演角色為何、以及適用領域之問題，當然也影響其演算法應如何規範的思考。AI 究竟是以一獨立主體之方式提供服務、或是僅以處理較為機械性、重複性的事務以輔助人類決定並提供服務的角色，會影響 AI 容錯的空間；而部分國家如美國，目前所提出之規範或指引，似多強調針對私部門之 AI 的適用。然而，對於演算法治理規範，公部門跟私部門的差異為何？多數國家並未特別區分公部門跟私部門，甚至可以說多專注在私部門使用 AI 之規制，而未針對公部門有具體規範。即便如美國 2019 年《AI 應用的規範指引》（Guidance for Regulation of Artificial Intelligence Applications）亦未於指引中具體說明該指引僅適用於私部門之原因；2019 年底眾議院雖通過的《AI 政府法案》（Artificial Intelligence Government Act），要求聯邦政府各部門在系統運行中盡可能使用 AI，並設立能建議和促進聯邦政府開發 AI 的創新用途的單位，但此法案並未獲參議院通過；至於 2020 年眾議院又提出的《國家 AI 倡議法案》（National Artificial Intelligence Initiative Act of 2020），則是川普政府時代的兩項行

政命令的倡議後第一個化為具體的 AI 法案，建立國家 AI 倡議，以促進 AI 研究和機構間合作。

由各國的資料顯示，目前關於演算法應用的規範，除在立法程序中的歐盟 2021 年《AI 規則草案》，以及美國 2019 年《AI 政府法案》與 2020 年《國家 AI 倡議法案》外，多尚未通過單一針對演算法應用管理法規，而是針對應用演算法之 AI 模式，提出指引以及處理原則。歐盟先在「歐盟執委會」下設「AI-高階專家工作小組」（AI-HLEG），並於 2019 年研擬《值得信賴的 AI 倫理指引》（「AI 倫理指引」），歐洲議會於 2020 年針對相關技術發佈《AI、機器人和相關技術的倫理框架》（「AI 倫理框架」），歐盟執委會亦於同年底發佈《AI 白皮書》（White Paper on AI），直到 2021 年 4 月再提出《AI 法律調和規則草案》（「AI 規則草案」）。至於美國在 AI 規範演變的過程中，川普總統於 2019 年簽署發布以行政命令第 13859 號提出美國 AI 倡議作為川普政府提出的國家 AI 政策，2020 年的行政命令第 13960 號則延續 13859 命令的政策內容，進一步提出適用於國安以外各大領域的一般性準則，直到眾議院於 2019、2020 年才分別提出《AI 政府法案》、《國家 AI 倡議法案》，化為具體的 AI 法案，建立國家 AI 倡議；另 2020 年亦提出備忘錄形式的《AI 應用的規範指引》，提出十點監理措施時需注意的原則。英國則從 2017 年先發佈《產業策略白皮書》將 AI 列為產業發展重點，再於 2018 年發佈跨部會之《AI 產業協議》，後於 2021 年發佈《AI 路線圖》與《公部門 AI 應用指南》等官方指引文件。

然而，即便目前可能尚未通過具體 AI 或演算法明文法律規範，就現有運作下之明文法規範，最注重的仍在於個人資料保護之法規，歐盟的「一般資料保護規則」（GDPR）可說是目前最典型、最通用也最嚴格的規範，其所重視的「當事人自主原則」，也影響了各國關於指引或法律制訂的思維，就歐盟 GDPR 的規範內容而言，即涉及對演算法之透明性要求。例如純粹自動化機器人分析資料所影響的權利，係為保護資料提供的主體，因此不會區分適用之對象為公部門或私部門，只要是資料控制者（Data controller）均須遵循 GDPR 之規定。即使該自動化決策是完全無自然人的介入，僅為單純由機器經過過去的機器學習後，自行解決問題或進行自動決策時（solving by itself），此時即有須遵循之義務。所稱之義務，以資料控制者為例，具有對資料主體有通知義務（notification），資料控制者必須主動向當事人揭露其個人資料是如何被處理及運用，因所謂透明度（transparency）之意旨，在於該如何使資料主體知悉其個人資料非以傳統之方式應用，而係透過人為之機器學習所應用。

另就解釋權（right to explanation）部分，針對自動化決策機制，GDPR 設有相關限制，但 GDPR 並未明文指出資料主體是否有權請求資料控制

者就個別自動化決策解釋其如何做成該決策之理由，意即資料主體是否享有解釋權。歐盟 GDPR 對此存有正反不同說法，以歐盟會員國各自之資料保護主管機關所成立之委員會，認為 GDPR 具有解釋權，無非是依據 GDPR 前言第 71 段提及，資料控制者應該對資料主體提供保護機制（safeguard），包括使資料主體在個別決策做成後，取得資料控制者是如何透過演算法做成該決策之解釋理由，此為 GDPR 唯一明文提及「解釋權」之處¹⁷，而在歐洲議會 2020 年出版之「GDPR 對於人工智慧之影響（The Impact of the General Data Protection Regulation（GDPR）on Artificial Intelligence）」官方文件整理正反不同意見後，似傾向支持解釋權的立場，但其依據是否為 GDPR 仍有疑義，造成演算法的透明性問題引發對自動化決策的不信任；另外一派則認為不存在解釋權，此說多半為一些規模較大的公司所採之主張（鄭伊廷，2021）。若存在解釋權，業者有責任要求資料控制者，當使用自動化機器學習處理他人之個人資料時，應告知該數據主體並向其解釋，使其有得選擇退出之權利，即當事人具拒絕作為自動化機器學習客體之權利，例如英國的國民健康保險（National Health Service, NHS）允許民眾對於其資料使用有退出（opt-out）的權利（NHS, 2021）¹⁸。這對業者而言將產生較高之操作成本，亦將導致公司營業祕密無法獲得妥善保護；另一困擾在於，應如何向數據主體解釋並使其清楚瞭解演算法及機器學習是如何運作。惟 GDPR 並未於本文中對於解釋權多做闡述，因此在實務操作上，規模較大的公司時常透過聘請專家學者撰寫相關文章說明解釋權不可能存在而拒絕解釋。

基於各國之 AI 政策規範與架構經驗，可以窺知由於 AI 突破了過去對於電腦使用之演算法設計方式，也改變了人類治理的思維，因此各國對於 AI 政策的規劃，有其共通之思考面向，且依序進行：組織規劃定位、規制方式制度、倫理價值與政策原則、風險評估分級。

一、組織規劃定位

專責政府機構組織上，根據 OECD 的建議，應區分為協調組織以及監理處之的思考。在整合協調機關部分，OECD 建議以現有或新設獨立之部門進行；至於治理監督的機關，除以現有或新設獨立之部門外，尚可以獨立之專家小組或新設獨立諮詢機構提供建議。此揭機構、組織，除包含可提供相關資訊並提出政策方向之功能外，亦可進行案件之審查以及個案之監理。

¹⁷ GDPR, recital 71, (providing that “In any case, such processing should be subject to suitable safeguards, which should include...obtain an explanation of the decision reached after such assessment...”).

¹⁸ 以英國 NHS 的健康資料庫（health database）為例，數據主體有兩階段的退出權，第一階段為得告知家庭醫師，不願將個人資料上傳至中央資料庫；第二階段則於上傳至中央資料庫後，要釋出資料予第三方時，數據主體應收到通知，此時數據主體可自行決定是否退出。

而從美國、英國、新加坡的經驗觀之，似多採取綜合方式，除由原有科技、資料治理的單位繼續進行相關政策擬定與監理外；諮詢部分多以新設委員會或成立專案工作小組機構（辦公室）為主。美國除有 AI 專責的諮詢委員會（Select Committee on AI）外，並整合現有白宮幕僚組織與外部專家成為「國家人工智慧研究資源工作組（National Artificial Intelligence Research Resource Task Force）」；英國除原有的科學辦公室（Government Office for Science, UK）、資訊專員辦公室（Information Commissioner's Office, ICO）外，另設立 AI 委員會（AI Council）、AI 辦公室（Office for AI）、數據倫理與创新中心（Centre for Data Ethics and Innovation）等單位；新加坡則設立智慧國家與數位政府辦公室（Smart Nation and Digital Government Office, SNDGO），並由原有之個人資料保護委員會（The Personal Data Protection Commission, PDPC）負責資料保護之事項。

二、 規制方式制度

由前揭各國制度演變進程可知（如圖 22），各國多以概括到具體、拘束力由弱到強之步驟演進，此與一般法律體系由倫理原則之發展開始走向具體法律之路徑類似。各國多先由政府進行 AI 相關研究並提出研究報告，進一步彙整後可能以政策白皮書之方式進行政策方向之宣示，再進一步以相關指引的方式，提出 AI 提供服務時，對於 AI 審查應考量之重點；待有所共識後，再提出法規之草案。



圖 22：歐、美、英 AI 規範演進比較

資料來源：本計畫。

當然，在總統制的美國，其行政命令之拘束力遠高於我國。不過除了美國的行政命令以外，截至本報告提出為止，歐盟的 AI 規則草案以及美國的國家 AI 倡議法案，都尚未見通過。也因此，相關倫理指引或架構文件，目前仍為目前各國進行 AI 應用的指導原則，在相關文件中所揭示的原則，仍提供各領域進行 AI 應用時的依據。

三、倫理價值與政策原則

由 AI 的學習過程，是透過電腦接收輸入資料、並給予所期望得到的答案後，透過機器學習程序由電腦自己生成規則，再進一步應用這些規則使用於新資料，以產生原先預期的答案，因此此一機器學習的過程，將大量使用個人之資料，也因此對於資料主體的權利可能造成侵害；而由於學習過程仍是由人為進行，不論是傳統學習先進程式撰寫規則、或是機器學習先給予結果由機器整理出規則，該學習之演算法編寫將嚴重影響到 AI 產出之結果，甚至可能造成不公的歧視。因此，對於資料提供主體的權利保護、以及使用 AI 服務的個體，均有可能存在權利受侵害的風險。

其實自歐盟《一般資料保護規則》（GDPR）以來，針對處理個人資料時，就已經有諸多應遵循之規則和原則以保護資料主體之權利，包含合法性（lawfulness）、透明性（transparency）、公平性（fairness）、正確性（accuracy）、數據最小化（data minimization）、目的和存儲限制（purpose and storage limitation）、保密性和課責（confidentiality and accountability）；另資料主體亦擁有多項權利，例如近用權（the right to access）、更正權（correction），而針對使用於身份識別之生物識別資料，或對資料主體之權利和自由構成高風險等敏感資料之處理，更應嚴格的進行資料保護影響評估。

因此自本報告第二章第三節表 3 之整理可以發現，各國國際組織以及各國的規範或建議，均將倫理價值列為最重要的事項，以追求可信賴的（trustworthy）AI 為目標，且必須以人為中心、以人為本、以人優先作為思考核心，AI 系統必須合法、合倫理、穩定、正確且安全。其中最重視的幾個共同原則，包括對資料主體尊重自主（autonomy）、公平性（fairness）、可解釋性（explainability）、透明性（transparency）、課責性（accountability）並重視隱私（privacy）。

在確立上開倫理之原則後，便進一步於政策上採取有效的施行手段，包括必須介入監管、提供公眾參與，且監管重點仍著重於上開倫理原則下的演算法的揭露與透明性、公平且不歧視、可解釋性、資料安全與隱私、AI 技術穩定與安全、可課責性，且跨機關部門的均需共同合作以達成目的。

抑有進者，公部門應用演算法治理之法規調適上，更強調演算法及 AI 模型之可解釋性、透明性、可課責性等原則之重要性，其理由在於私部門應用 AI 模型時，尚有外部之公部門得以對私部門內控機制進行監督。然而，當公部門應用演算法模型或 AI 系統時，其公權力對於個人與社會所產生之風險疑慮，更甚於私部門，惟現有輿論或立法部門之監督機制是否有效，不無疑義；且相較於行政機關決策者多為民選或經考試任用，具有

一定程度直接或間接的民主正當性，但若公部門於行政事務應用演算法模型及 AI 系統產生行政事項之自動化結果，該自動化之判斷結果相對缺乏民主之基礎。則在目前 AI 發展尚無法為價值判斷之前提下，似乎難以期待 AI 具有能力判斷其所接收到的數據、資料是否對人民有權利侵害之虞，且在機器人難以課責之情形下，人民是否有陳述意見之機會或提起救濟之可能？據此，公部門欲應用演算法治理時，其演算法及 AI 模型之內容、數據資料之保存及運用，以及其內部之控制機制是否能獲得民眾信賴，均將成為挑戰。

另一與課責性相關之議題，在於企業開發 AI 應用之營業秘密問題，若政府部門非自行應用演算法進行 AI 開發，則必會委託或外包民間企業辦理，此時其技術授權性質為何？又行政機關運用之演算法模型行使裁量權時，得否將該演算法視為行政機關之內部行政規則？依我國行政程序法或政府資訊公開法，該演算法皆應對外公開，惟若因個案狀況應對外公開時，公部門要公開的範圍與程度、僅需直接公開或是應轉換成人民可以理解之形式、演算法公開之公眾利益與企業營業秘密孰輕孰重？又受託單位是否視為公權力之委託或具準公務員性質？在法律上皆有待釐清。

四、 風險評估分級

另外就事務性質本身，有些需要以較高強度的方式進行監管，有些則可採取寬鬆之方式，其區分是以 AI 應用目的領域所產生風險強度作為依據。無論是歐盟的《AI 規則草案》具體列舉數類高風險之事務，包括生物辨識、民生資源供應、教育、人力資源管理、獲得公共服務的資格、執法事務與司法評估、移民評估等等，均涉及對人民而言較為核心之利益與財罰認定；英國的《人工智慧稽核架構》也指出 AI 的特定風險領域會涉及的問題，事實上與前揭倫理原則均會息息相關，包括在 AI 分析結果的可能造成偏見和歧視造成公平性問題、AI 決策會對資料主體進行解釋、對於資料使用的正確性、AI 決策涉及的人工審查、涉及資料主體的權利行使、決策的衡平等事項，均涉及風險的衡量。因此，對於此揭高風險事項，在監理的強度便需嚴格。

以美國 COMPAS system 為例，COMPAS 幫助法官評估被告之再犯風險，以作為量刑準據。COMPAS 會進行大量問答調查，依據被告回答、年齡、過往犯罪紀錄與類型等各項資料及數據，推估被告之再犯率，最後由法官決定被告服刑的長短。COMPAS system 初期並沒有太多討論，但當其參數被揭露後，有非常多的，特別是強調司法正義 (criminal justice) 者開始反省，納入這些數據之意義為何，其回溯性是如何判讀，抑或是一個前瞻性之預測。這兩者間之關連在於，事實上 COMPAS 係仰賴回溯性之判斷，以做成前瞻性之預測，然回溯性之判斷多半容易產生偏見 (bias)。譬如可能不自覺地把傳統犯罪率高的編碼 (zip code) 置於參數中，然這

類編碼中通常黑人或移民者比例較高，因此 COMPAS 可能產出白人獲得保釋機率較高之結果。

另於本計畫進行訪談時專家 P18 所提出，以我國社政體系、衛政體系導入之兒童家暴風險預警系統為例，其內容係以過往社工服務資料做為 AI 機器學習之依據，過程中使用人力檢視演算法模型及 AI 機器學習之成果，以及是否能有效地找到面臨家暴風險的孩子。在此驗證過程中發現，當 AI 運用演算法所產出之結果被認為可以使用時，其角色仍較偏向於決策指引或可依循方向之參考性質。其理由在於，演算法模型隨著時間及環境背景因素可能會失效，但非謂演算法本身改變，通常來自於政策改變，因為政策改變使得服務工作流程改變，造成原先設定的 AI 模型產生偏離，因此 AI 系統就不一定能夠再度找到較具家暴風險的孩子。此部分亦連帶關係到應用 AI 之課責性問題，因此，較為合理之 AI 治理機制應為演算法模型須持續測試並進行定期修正。因為 AI 必須具有演化性，隨著政策改變及各階段數據資料可能產生的結果偏離、漂移現象，需要持續監督及修正。我國金融監督管理委員會對於各金融業務機夠運用理財機器人之規範即特別強調持續監督演算法之重要性，甚至要求內部應有專責機關負責。然演算法欲進行修正之前提為何、何時可認為 AI 會產生偏離，則關係到演算法模型之可解釋性、透明度或其驗證標準。

五、 演算法適用之服務領域

雖然參酌國外資料，未有對於適用 AI 之服務領域相關討論，僅從上位概念討論風險的不同影響監理的強度。然而，經由檢視各國適用之經驗、區分風險的類型後，以及透過本計畫邀情專家會議與個別訪談之意見彙整，多認為以現階段 AI 發展狀況觀察，適合優先應用演算法模型之場域，大致上可分為三個類型，其特色在於具有機械性、資料密集性以及有監理之高度需求：

- (一) 單純機械性事務領域：即偏向單一資料處理、影像處理，且較不涉及人為評價（判斷是非對錯）之項目。此因判斷性任務隨著不同應用領域之特質，需要人力介入之程度亦有不同，多半需進行決策，且相對較為複雜，易滋生爭議而有究責之可能，單純資料之處理較不易生此爭議。
- (二) 高資料密集性之領域：自金融監理經驗得知，金融監理需要大量數據以辨別投資人之投資行為風險、可疑交易或償債能力等事項，演算法模型得以更高效率地為其篩選出投資風險較高或償債能力較低之投資人。
- (三) 受高度監管之特許行業：因此揭特許行業具有較可能侵害人民權益之風險，受監管者需要提供大量資訊、報告義務或是申報、揭

露義務，與監管者溝通互動，此部分若能透過演算法將監管事項系統化、流線化以提升報告揭露之效率，將會是應用場域之突破。

對於評價性以及可能的歧視的問題，經訪談資訊領域專家之初步建議，可採取提供類似演算法開發平臺於使用者，將演算法模型予企業或政府單位有償使用。而應用結果會產生歧視或偏離之情形，在於 AI 透過演算法進行機器學習之過程中，接收來自四面八方數據及資料，因此應由專業人員進行定期維修及更新，如此應較能合於法規調適之透明度跟可解釋性要求。惟應特別注意的是，數據平台上的資料可能已非原始資料，可能是已經帶有評價性的資料，因此應審慎避免 AI 所導出之結果產生評價歧視的問題。

整體而言，對於 AI 監管朝向法律的制訂仍有其必要，但在法律規範制訂以前，先以操作指引或準則方式進行規範時，自本團隊就各國法律規制層面所舉行的專家會議，與會專家亦建議宜在操作準則甚至法規範上，賦予資料控制者或資料處理者等一定程度之揭露義務。若資料處理是百分之百自動化機器學習，資料控制者須告知數據主體（data subject），其個人資料將運用於機器學習，而可能會產生相關法律效果或相等於法律之效果，且應給予當事人拒絕自動化決策之權利。意即當公部門欲利用機器人協助其完成或執行一定行為時，公部門需揭露公部門之機器人所使用各項參數、數據之方式，並提供救濟之可能性。若與當事人所認定的事實有所違背時，亦得挑戰預定之參數，此不失為促使整個自動化決策或演算法變得更透明的方法。

第二節 資源配置之調查分析

一、調查資料分析

國立政治大學於 110 年 9 月 2 日，以政公字第 1100025388 號函，行文邀請行政院各部會派員參與本調查，調查期間為同年 9 月 15 日至 24 日，共有十八個部會派員填答問卷，填答率為 56.25%。第二階段調查期間為同年 9 月 21 日至 10 月 15 日，共有二十四個部會共推薦 180 位同仁協助填答問卷，回收 140 份問卷，扣除十二份未填內容之空白問卷以及兩份不同意參與研究之問卷，最終有 126 名受推薦者完整填答問卷，推薦填答完成率為 70%。未參與第一階段調查，但有請所屬同仁參與第二階段調查的部會有六個，分別是國防部、財政部、法務部、衛生福利部、行政院環境保護署和文化部，本計畫團隊藉由網路搜尋發現除文化部外，國防部、財政部、法務部、財政部、衛福部和環保署等部會都曾有發佈新聞稿

表示已在部會相關業務導入 AI 系統，故融合第一階段調查與網路搜尋結果，於第二階段調查結果增加一個是否導入 AI 服務的控制變項，第一階段調查表示已導入或正規劃導入，以及網路搜尋結果有導入 AI 服務的部會皆編碼為是(1)，其餘編碼為否(2)，期藉此控制變項探索有導入 AI 系統部會和未導入 AI 系統部會之間的差異。

至於，未參與任一階段調查問卷調查之行政院所屬部會包括內政部、經濟部、大陸委員會、行政院主計總處、國立故宮博物院、中央選舉委員會、國家通訊傳播委員會以及促進轉型正義委員會等八個部會，本計畫團隊進一步探詢不參加調查的理由為：問卷涉及人員眾多不便處理、推動相關業務不便填答、無 AI 相關應用以及屬臨時性任務單位等。基於研究倫理，本調查非為強制性質，故本計畫團隊尊重各部會的決定，並且由衷感謝撥冗填答問卷的各部會同仁，緊接著就各階段調查分析結果分述如下。

(一) 第一階段調查

第一階段調查共有十八個部會派員完成問卷填答，各部會運用 AI 系統輔助業務的現況為六個部會已導入，兩個部會規劃中，十個部會無規劃，各部會擁有之資訊專責人力為 0~92 人(平均為 26.67 人，標準差為 23.54)，其中參與規劃、採購或管理 AI 系統專案的資訊人力有 0~10 人(平均為 1.44 人，標準差為 2.662)，資訊專責人力為正式公務人員的有 0~32 人(平均為 5.78 人，標準差為 7.612)，另依據各部會填答運用 AI 系統輔助業務的現況，分別說明調查分析結果如后，並將分析結果整理如表 25 所示。

首先，目前已導入 AI 系統的部會計有：行政院人事行政總處、行政院農業委員會、科技部、國軍退除役官兵輔導委員會、勞動部和僑務委員會(依據回覆時間先後排序)等六個部會，AI 系統輔助業務的類型包括：操作機器的自動化(二個部會，如 AI 輔助診斷門診系統、農作物病蟲害智能診斷助理整合系統)、辦公室行政文書工作(一個部會，如偵測一案多投)、與服務對象溝通工作(三個部會，如智能客服)以及其他創新公共服務(兩個部會，如作文 AI 評分系統、口說 AI 評分系統、歌唱 AI 評分系統、醫療業務)等四種類型。此六個部會最早自民國 107 年起開始運用 AI 系統輔助業務，最晚則是從民國 110 年開始運用，累計於 107 至 110 年度共編列採購或發展 AI 系統預算達 113,872 千元¹⁹，全數採用委外途徑建置 AI 系統。此六個部會運用之 AI 系統的數位資料來源為本機關開放資料的有兩個部會，本機關內部資料的有四個部會，他機關開放資料的有一個部會，其他資料來源的有一個部會(使用者提供資料)，無部會使用他機關內部資料或私部門資料。有兩個部會表示在導入及使用 AI 系統

¹⁹ 某部會填答 108 至 110 年度編列採購或發展 AI 系統之預算數為 2617,000 千元、14,100,000 千元、2,750,000 千元，經核對其於官網公開之預算數，發現明顯高於該部會全年度資訊預算數，故研判應是填答者誤看單位為「元」所致。因，改以 2,617 千元、14,100 千元、2,750 千元計算。

的過程中，有進行相應的法規調適，例如檢視是否符合個資法規範。其餘四個表示未進行法規調適的部會中，有兩個部會表示目前法令完備，無法規調適之必要；一個部會表示目前為試辦階段，無法規調適必要；一個部會表示還未知要進行哪些法規調適。

其次，正規劃導入 AI 系統的部會，包括：金融監督管理委員會與原住民族委員會等兩個部會，系統輔助業務的類型包括：辦公室行政文書工作（一個部會，如應用於監理業務，以減輕同仁工作負擔，提升行政效率）、服務對象溝通工作（一個部會，如用於輔助長照業務，運用系統瞭解服務對象的長照需求）以及其他創新公共服務（一個部會，輔助推動長照業務）等兩種類型，預計在 110 年（調查時尚未啟用）及 111 年開始導入 AI 系統輔助業務，其中僅有一個部會表示在 110 年編列採購或發展 AI 系統預算達 8,470 千元，兩者皆規劃採用委外途徑建置 AI 系統。此兩部會規劃建置的 AI 系統使用的數位資料來源皆為本機關內部資料，且未考慮進行相應的法規調適，理由為無需進行法規調適或還未知要進行哪些法規調適。

最後，表示尚未規劃導入 AI 系統輔助業務的部會，計有：中央銀行、公平交易委員會、外交部、交通部（科技顧問室）²⁰、行政院公共工程委員會、行政院原子能委員會、客家委員會、海洋委員會、國家發展委員會、教育部等十個部會（回傳問卷先後順序排列），其未規劃導入 AI 系統的理由為：缺乏預算（六個部會）、缺乏專業人力（四個部會）、缺乏配套法令（兩個部會）、缺乏適用技術（五個部會）、缺乏數位資料（兩個部會）、非機關權責業務（四個部會）、缺乏成本效益（一個部會）、機關首長未指示推動（兩個部會）、其他理由（三個部會）。未規劃導入 AI 系統的其他理由為：持續研究中（一個部會）和暫無需求（兩個部會）。

²⁰ 本調查規劃由部會派員以部會立場填答問卷，惟該部會代表填答之部會名稱與其他部會不同，故存有與現況不符之疑慮。

表25：第一階段問卷調查結果綜整表

	已導入 AI 系統	規劃導入 AI 系統	未規劃導入 AI 系統
部會數	6	2	10
部會名稱	行政院人事行政總處、行政院農業委員會、科技部、國軍退除役官兵輔導委員會、勞動部、僑務委員會	金融監督管理委員會、原住民族委員會	中央銀行、公平交易委員會、外交部、交通部(科技顧問室)、行政院公共工程委員會、行政院原子能委員會、客家委員會、海洋委員會、國家發展委員會、教育部
專責資訊業務人力數	平均數 35.67，最小值 12，最大值 92，標準差 28.946	平均數 27.5，最小值 5，最大值 50，標準差 31.82	平均數 21.1，最小值 0，最大值 52，標準差 19.496
累計 107~110 年編列採購或發展 AI 系統預算數	113,872 千元	8,470 千元	跳答
AI 系統建置途徑	全數委外	全數委外	跳答
AI 系統輔助業務類型	操作機器的自動化 (25%)、辦公室行政文書工作 (12.5%)、與服務對象溝通工作 (37.5%)、其他創新公共服務 (25%)	辦公室行政文書工作 (33.3%)、與服務對象溝通工作 (33.3%)、其他創新公共服務 (33.3%)	跳答
AI 系統使用數位資料類型	本機關開放資料 (25%)、本機關內部資料 (50%)、他機關開放資料 (12.5%)、其他資料來源 (12.5%)	本機關內部資料 (100%)	跳答

	已導入 AI 系統	規劃導入 AI 系統	未規劃導入 AI 系統
法規調適現況	進行法規調適：2 個 未進行法規調適：4 個	未進行法規調適：2 個	跳答
未進行法規調適理由	法令完備無需要 2 個 試辦中無需要 1 個 未知需調適項目 1 個	法令完備無需要 1 個 未知需調適項目 1 個	跳答
未導入 AI 系統理由	跳答	跳答	缺乏預算（20.7%）、缺乏專業人力（13.8%）、缺乏配套法令（6.9%）、缺乏適用技術（17.2%）、缺乏數位資料（6.9%）、非機關權責業務（13.8）、缺乏成本效益（3.4%）、機關首長未指示推動（6.9%）、其他（10.3%）。

資料來源：本計畫。

（二）第二階段調查

1、受訪者基本資料

第二階段調查共有 126 份有效問卷，僅 124 份問卷填答個人資本資料，茲將各類別統計人數（百分比）整理如表 26 所示，並簡述統計分析結果如下。

表26：第二階段調查受訪者基本資料一覽表

性別 (n=126)	人數	有效百分比	目前職位 (n=126)	人數	有效百分比
女性	63	50.8%	主管	46	37.1%
男性	61	49.2%	非主管	78	62.9%
遺漏值	2		遺漏值	2	
年齡 (n=126)	人數	有效百分比	隸屬單位 (n=126)	人數	有效百分比
20 歲至 24 歲	1	0.8%	業務單位	25	20.2%

25 歲至 29 歲	4	3.2%	資訊單位	46	37.1%
30 歲至 34 歲	14	11.3%	人事單位	49	39.5%
35 歲至 39 歲	25	20.2%	其他	4	3.2%
40 歲至 44 歲	30	24.2%	遺漏值	2	
45 歲至 49 歲	19	15.3%	年資 (n=126)	平均	標準差
50 歲至 54 歲	21	16.9%	公職年資 *	14.92	8.317
55 歲至 59 歲	8	6.5%	機關年資 *	6.93	7.397
60 歲以上	2	1.6%			
遺漏值	2				
教育程度 (n=126)	人數	有效百分比	官職等 (n=126)	人數	有效百分比
專科	1	0.8%	簡任	8	6.4%
大學	56	45.2%	薦任	101	81.4%
碩士	60	48.4%	委任	6	4.8%
博士	7	5.6%	其他	9	7.4%
遺漏值	2		遺漏值	2	

備註：年資計算以年為單位，未滿一年以一年計算。

資料來源：本計畫。

受訪者生理性別為女性有 63 人（50.8%）、男性有 61 人（49.2%）。年齡未滿 30 歲者有 5 人（0.4%）、30 歲以上未滿 40 歲者有 39 人（31.5%）、40 歲以上未滿 50 歲者有 49 人（39.5%）、50 歲以上未滿 60 歲者有 29 人（23.4%）、60 歲以上者有 2 人（1.6%）。教育程度為專科有 1 人（0.8%）、大學有 56 人（45.2%）、碩士（含 1 名雙碩士）有 60 人（48.4%）、博士有 7 人（5.6%）。目前職位為主管有 46 人（37.1%）、非主管有 78 人（62.9%）。隸屬單位為業務單位有 25 人（20.2%）、資訊單位有 46 人（37.1%）、人事單位有 49 人（39.5%）、其他有 4 人（3.2%）。現職官等為簡任者有 8 人（6.4%），其中 12 職等有 1 人、11 職等有 6 人、10 職等有 1 人；現職官等為薦任者有 101 人，其中 9 職等有 68 人、8 職等有 15 人、7 職等有 12 人、6 職等有 6 人；現職官等為委任者有 6 人，其中 5 職等 3 人、4 職等 2 人、3 職等 1 人；現職官等為其他者有 9 人，多為契約聘用人員。受訪者於公部門服務總年資平均為 14.92 年，標準差為 8.317；現職機關服務年資平均為 6.93 年，標準差為 7.397。

2、測量變數建構

為回答本調查的第四個問題「中央部會公務人員對於應用機器演算法 (AI) 創新公共服務的態度、預期效能以及對可能遭遇問題的解決態度為何？」，本計畫團隊採用探索性因素分析 (explanatory factor analysis, EFA)，透過主成分分析法 (principal component analysis) 與直接斜交法 (direct oblimin)，就回答研究問題所涉及之文官運用 AI 態度、文官從事公共服務創新的心理資本、文官所屬機關資料準備度、文官所屬機關組織管理現況等面向進行相關因素的萃取。綜合四個面向的因素分析結果顯示 KMO (Kaiser-Meyer-Olkin measure of sampling adequacy) 值皆大於可接受門檻值 0.6 且 Bartlett 球形檢定 (Bartlett sphericity test) 結果亦達顯著效果 ($p < .001$)，代表所有引用的測量題項皆適合進行因素分析 (Kaiser, 1974)。各題因素負荷量 (factor loading) 大於可接受的判定標準 0.6，且各因素信度 (Cronbach's α) 皆達 0.7 以上的可接受水準，表示各因素具有信度及效度 (Hair et al., 2006)，茲就四次 EFA 分析結果整理說明如下。

首先，本計畫團隊選取文官運用 AI 態度相關 12 道題項 (編號 2-1~2-12) 進行因素萃取，並以特徵值 (eigenvalue) 大於 1 作為因素萃取標準 (Kaiser, 1974)，順利萃取出三個因素，分別命名為預期正向成效、預期承擔責任和預期權力限制，分析結果整理如表 27 所示。

表27：文官運用AI態度相關題項的EFA結果

($n=126$ ； $KMO=0.830$ ；Bartlett 球形檢定=838.215***)

萃取因素	題項	平均數 (標準差)	因素負荷量	解釋變異量 %	信度 Cronbach's α
預期正向成效	2-1.AI 系統能幫助我更快速地完成任務。	4.45 (1.457)	.889	44%	.910
	2-2.AI 系統能幫助我更正確地完成任務。	4.39 (1.431)	.907		
	2-3.AI 系統能幫助我更節省行政作業成本。	4.48 (1.506)	.898		

萃取因素	題項	平均數 (標準差)	因素負荷量	解釋變異量%	信度 Cronbach's α
	2-4.AI 系統能幫助我更公平地完成任務。	3.90 (1.741)	.707		
	2-9.AI 系統讓我更容易回應民眾需求。	3.87 (1.730)	.856		
	2-10.我認為 AI 系統可以減少人工作業程序。	4.56 (1.342)	.641		
預期承擔責任	2-5.AI 系統會增加我的業務量。	3.56 (1.450)	.849	16.34%	.715
	2-6.AI 系統產生的錯誤會由我承擔責任。	3.96 (1.641)	.627		
	2-12. 我擔心 AI 系統會帶來未知的風險。	4.02 (1.456)	.751		
預期權力限制	2-7.AI 系統會限縮我的行政裁量權力。	3.05 (1.485)	-.878	8.49%	.704
	2-8.AI 系統讓長官更容易掌握我的工作情況。	3.49 (1.742)	-.594		
	2-11.我擔心 AI 系統會取代我的工作。	2.47 (1.244)	-.687		

資料來源：本計畫。

其次，本計畫團隊選取與文官從事公共服務創新之心理資本相關的12道題項(編號2-13~2-24)進行EFA，刪除因素負荷量不佳(低於0.6)題項2-16後，共萃取出兩個因素，分別命名為心理韌性和科技樂觀，分析結果整理如表28所示。

表28：文官心理資本相關題項的EFA結果

(n=126；KMO=0.873；Bartlett球形檢定=815.746***)

萃取因素	題項	平均數 (標準差)	因素負荷量	解釋變異量%	信度 Cronbach's α
心理韌性	2-17.我能想到很多的方法來達到我目前的目標。	4.50 (.777)	.639	49.426%	.906
	2-18.到目前為止，我認為我自己還蠻成功的。	4.20 (.770)	.718		
	2-19.我有信心可以有效地處理突發事件	4.36 (.732)	.770		
	2-20.只要我付出必要的努力，我可以解決大多數的問題。	4.38 (.866)	.841		
	2-21.面對困難我可以保持冷靜，因為我可以依靠自己的應對能力。	4.35 (.741)	.848		

萃取因素	題項	平均數 (標準差)	因素負荷量	解釋變異量%	信度 Cronbach's α
	2-22.我認為我自己是一個可以承受很多壓力的人。	4.06 (.986)	.792		
	2-23.失敗並不會讓我覺得氣餒。	3.87 (.988)	.780		
	2-24.在遭遇嚴重的生活難題後，我往往可以迅速恢復。	4.03 (.894)	.798		
科技樂觀	2-13.我很期待科技輔助的未來生活。	4.87 (.769)	.925	17.891%	.883
	2-14.我總覺得科技應用對我帶來的好處多過於壞處。	4.71 (.727)	.903		
	2-15.我相信科技會帶給我很多好事。	4.63 (.776)	.853		

資料來源：本計畫。

第三，本計畫團隊選取與文官所屬機關運用 AI 所需資料準備度相關的 5 道題項（編號 2-30～2-34）進行 EFA，共萃取出一個因素，命名為資料準備度，分析結果整理如表 29 所示。

表29：文官所屬機關擁有AI系統所需資料準備度相關題項的EFA結果

(n=126；KMO=0.818；Bartlett 球形檢定=443.799***)

萃取因素	題項	平均數 (標準差)	因素 負荷 量	解釋變 異量%	信度 Cronbach's α
資料 準備 度	2-30.我所屬的機關擁有開發AI系統所需的全部資料與數據庫。	2.78 (1.815)	.864	73.919%	.912
	2-31.我所屬的機關擁有可轉換為機器可讀格式的資料與數據庫。	2.92 (1.883)	.906		
	2-32.我所屬的機關有完整的資料治理規範，可確保資料與數據庫的內容正確性。	3.15 (1.889)	.883		
	2-33.在符合現有的作業規範下，我所屬機關擁有的資料與數據庫可以開放外部專家使用來開發AI系統。	2.60 (1.776)	.794		

萃取因素	題項	平均數 (標準差)	因素 負荷量	解釋變異量 %	信度 Cronbach's α
	2-34.我所屬的機關有完整的資料使用紀錄，包括從資料的源頭收集到系統應用端的所有使用過程。	2.90 (1.837)	.847		

資料來源：本計畫。

最後，本計畫團隊選取與文官所屬機關組織管理現況文官相關的 4 道題項（編號 2-35, 2-38~2-40，排除兩道負向題項 2-36, 2-37）進行 EFA，共萃取出一個因素，命名為支持的組織管理系統，分析結果整理如表 30 所示。

表30：文官所屬機關組織管理現況相關題項的EFA結果

(n=126；KMO=0.663；Bartlett 球形檢定=187.528***)

萃取因素	題項	平均數 (標準差)	因素 負荷量	解釋變異量 %	信度 Cronbach's α
支持的組織管理系統	2-35.我所屬機關的管理系統能夠良好的支持組織的整體目標。	4.12 (.900)	.768	62.048%	.791
	2-38.我所屬機關的管理系統會鼓勵大家挑戰傳統與過去不可撼動的作法。	3.52 (1.108)	.813		

萃取因素	題項	平均數 (標準差)	因素 負荷 量	解釋變異 量 %	信度 Cronbach's α
	2-39.我所屬機關的管理系統會給予成員足夠的彈性來因應變動的社會環境。	3.73 (1.007)	.905		
	2-40.我所屬機關的管理系統以政策一致和穩定為主要目標。	4.34 (.931)	.641		

資料來源：本計畫。

3、推論統計分析

本計畫團隊運用多元線性迴歸分析方法，探討文官運用 AI 的三種預期態度---預期正向成效、預期承擔責任和預期權力限制，受到心理韌性、科技樂觀、資料準備度和支持的組織管理系統等因素影響的情況。**表 31**顯示多元迴歸分析結果，三個模型的依變項依序為預期正向成效、預期承擔責任和預期權力限制，數據顯示三個模型的 F 值顯著，代表分析資料可採用多元迴歸模型解釋，惟三個模型的調整後 R² 值皆不高，意謂尚有探索其他自變項的空間。

模型一的數據顯示，當文官意識到所屬機關的資料準備度愈高，將愈預期使用 AI 能帶來正向的成效；模型二顯示當文官意識到所屬機關的資料準備度愈高，將愈預期運用 AI 需承擔的責任，但若感受到組織的管理系統能展現對員工的支持，反而愈不預期自己需承擔責任；模型三顯示當文官的心理韌性愈好，愈不預期使用 AI 會限制自己在工作上的權力，但當文官愈是抱持科技樂觀的態度，反而愈是預期使用 AI 會限制自己在工作上的權力。

前述的分析結果意謂文官對於使用 AI 輔助業務乙事具有複雜的態度，一方面預期使用 AI 會為自己的工作帶來正向的成效，另一方面又預期使用 AI 會加重自己需承擔的責任以及限制自己在工作上的權力。此外，分析結果也顯示資料準備度是讓 AI 系統為行政業務帶來正向成效的關鍵，

而行政機關在組織管理系統展現對於員工的支持，將能降低員工對於數位創新將加重個人責任的疑慮。此外，文官應加強自我的心理韌性才能因應政府數位轉型帶來的工作變革，但同時需務實看待行政機關導入 AI 系統產生的業務變革，才不致於過度樂觀而忽略伴隨變革而來的衝擊。

表31：文官運用AI預期態度之多元迴歸分析結果

自變項 $\beta(t)$ \ 模型	模型一	模型二	模型三
心理韌性	.049(.563)	.043(.450)	-.214(-2.235)*
科技樂觀	.099(1.162)	-.143(-1.503)	.257(2.728)**
資料準備度	.423(4.745)***	.204(2.052)*	-.161(-1.632)
支持的組織管理系統	.066(.785)	-.248(-.2641)**	.053(.567)
調整後 R^2	.237	.049	.064
F 值	10.722***	2.597*	3.144*
N	126	126	126

資料來源：本計畫。

4、人才職能落差分析

本調查設計七道和人才職能發展有關的題項（編號 2-41~2-47），已瞭解受訪者對於行政機關運用 AI 系統所需人才職能的看法。調查結果如表 32 顯示，在調查的 7 種人才職能中，受訪者認為運用 AI 輔助業務最需具備的能力依序是：解讀資訊能力（2-45）、資安與倫理（2-42）、大數據分析能力（2-47）、跨域合作能力（2-41）、資源分配能力（2-44）、創新服務能力（2-43）和維運 AI 能力（2-46）。

表32：行政機關運用AI系統所需人才職能分析結果

題號	核心職能	題項內容	A 重要性	B 具備程度
2-41	跨域合作	具備能與他人共事，以解決 AI 系統所涉及跨域問題的能力。	4.99 (1.081)	4.13 (1.191)
2-42	資安倫理	具備能確保安全與合乎道德規範使用 AI 系統的能力。	5.20 (1.055)	4.25 (1.203)
2-43	創新服務	具備有策劃創造性解決方案以提供新產品或服務的能力。	4.77 (1.172)	3.66 (1.143)

題號	核心職能	題項內容	A 重要性	B 具備程度
2-44	資源分配	為能有效執行 AI 系統導入方案，具備相關資源分配的能力(如人力、財務、知識、契約管理等項)。	4.94 (1.141)	3.57 (1.227)
2-45	解讀資訊	具備有判讀與監測 AI 系統輸出的能力。	5.03 (1.062)	3.46 (1.299)
2-46	系統維運	具備有維運 AI 系統的能力。	4.77 (1.144)	2.95 (1.402)
2-47	數據分析	具備能在複雜現實環境中運用大數據分析，解決不明確問題的能力。	5.00 (1.136)	3.24 (1.352)

資料來源：本計畫。

本計畫團隊進一步比較受訪者認知才能職能應具備程度和自評實際具備程度的落差，依據表 33 顯示的資訊可知，受訪者自評具備能力落差由大到小的排序是：維運 AI 能力 (-1.82)、大數據分析能力 (-1.76)、解讀資訊能力 (-1.57)、資源分配能力 (-1.37)、創新服務能力 (-1.11)、資安與倫理道德 (-0.95)、跨域合作能力 (-0.86)。

表33：人才職能應具備與實際具備程度落差分析表

題號	能力項目	應具備程度	實際具備程度	落差分析
2-41	跨域合作能力	4.99	4.13	-0.86
2-42	資安與倫理道德	5.2	4.25	-0.95
2-43	創新服務能力	4.77	3.66	-1.11
2-44	資源分配能力	4.94	3.57	-1.37
2-45	解讀資訊能力	5.03	3.46	-1.57
2-46	維運 AI 能力	4.77	2.95	-1.82
2-47	大數據分析能力	5.00	3.24	-1.76

資料來源：本計畫

受訪者自評落差較大的職能項目依序為：維運 AI 能力、大數據分析能力、解讀資訊能力、資源分配能力、創新服務能力、資安與倫理道德以及跨域能合作能力，圖 23 以視覺化方式呈現各人才職能應然與實然落差

分析結果。

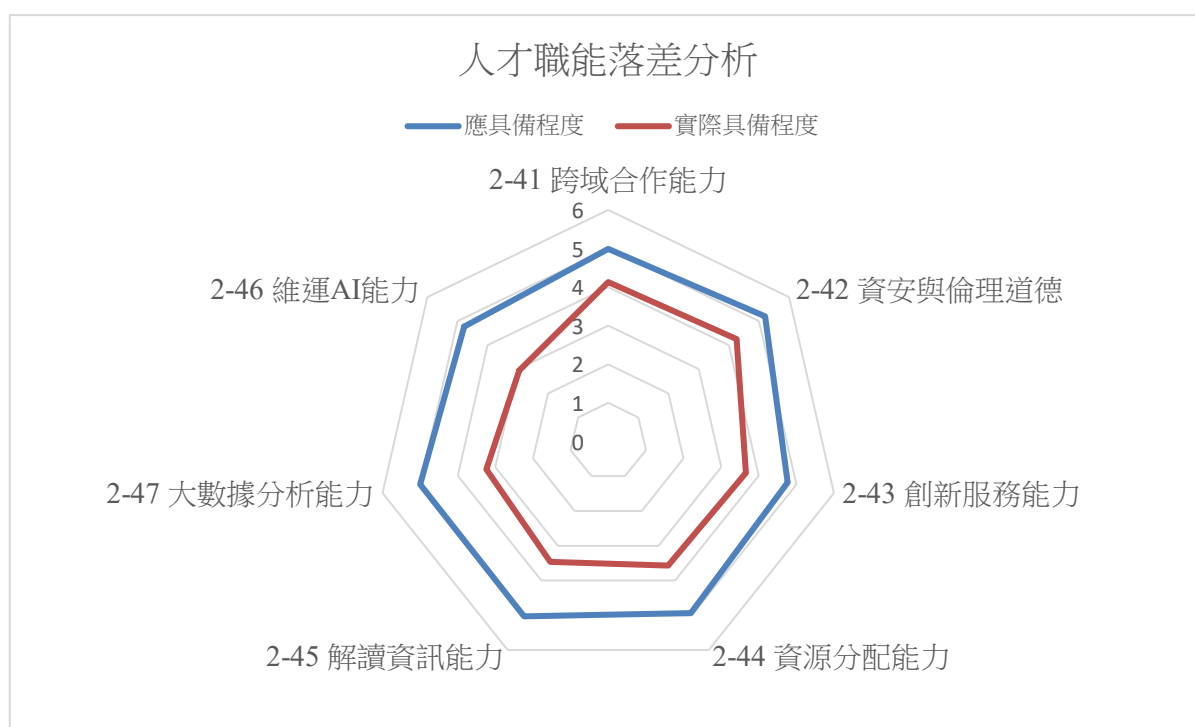


圖23：人才職能落差分析雷達圖

資料來源：本計畫。

二、 深度訪談與焦點座談分析

（一）人機職能整體性分析－深度訪談

除了針對中央各部會進行人機協作職能的全面性問卷調查之外，為了更深入瞭解組織內的成員對於人機協作的認知與態度，本計畫針對人機協作進行深度訪談，初步以事實陳述與態度價值兩面向進行訪談資料彙整，並根據彙整後的結果，嘗試與人機協作職能地圖進行連結說明，茲將其整理如下表所示。從深度訪談的分析結果顯示，受訪專家多從組織整體的觀點提出他們對人機協作的看法，因此可以說，由深度訪談所彙整歸結的重點，主要出發點以組織層次及制度運作為主，來進行人機協作的相關討論。

之所以將深度訪談資料區分為事實陳述與態度價值兩面向主因在於，組織的組成是由人所構成，因此從人的態度組成面向進行歸納，可進一步瞭解運用人機協作未來的可能行為，因為人的行為會受到他對現實場域的認知觀察，以及觀察後對於認知到的事實所賦予的情感影響。所以相對的，本計畫嘗試透過專家學者對目前所觀察到人機協作的事實陳述，以及他們對於這些事實認知所反映的態度價值，可反映他們認為未來人機協作在執行面向上所應注意的點、線、面問題，此也可做為指導未來公部門引入人機協作時應採取行為的參考，有關人機協作部分訪談資料之彙整

如表 34 所示。

表34：人機協作部分訪談資料之彙整

面向	次構面	議題	與協作職能的相關性
態度 價值	AI 應用的挑戰	1.涉及個人隱私的倫理安全要求。 2.主管與承辦的角色定位。 (1)主管仍須具備 AI 的基本認識。 (2)承辦同仁對於「被取代」的疑慮將阻礙人機協作的成效。	對於 AI 應用的挑戰，實多有關於協作職能中倫理與安全性技能的討論。再者，組織內部成員對於 AI 取代自身業務的現象，實也考驗引入 AI 的承辦單位或者業務主管對於人員管理技能的掌握，如何在引進的同時能激勵、發展與指導，讓內部同仁能接受這樣的變革作為。
	民眾端接受度	1.民眾是否具備足夠的 AI 使用素養。 2.情感運算對民眾接受與使用 AI 的影響。	民眾端的接受度考驗人機協作職能中的技術能力。技術職能中強調品質管控與強調科技應用者的滿意度，因此，民眾能否接受，與使用上是否能達到足夠的滿意度，攸關於系統本身的建構品質，如本身的資訊服務品質、系統本身的易用性等，均與技術能力達到直接的關聯性。
	AI 引進的課責問題	1.政府容錯率低與創新思維要求的衝突。 2.演算法的黑盒子困境。	由於無法明確知道演算法的演算過程，增加我們在倫理與安全性上的疑慮，此凸顯出引進 AI 時，我們必須針對該面向的職能要求有更進一步的規範。
事實 陳述	強調社會性技能中的跨域協調能力	1.必須具備跨領域協調整合能力。 2.強調說服與談判的跨單位關係經營。	無論是跨域與跨部會間的協調整合，主要均聚焦在協作職能中的社會性技能面向，為了要能夠達到跨領域間的對話，以及資料串接所進行的跨單位間的協調整合，兩者均強調承辦同仁甚至是主管，必須具備瞭解他人行動的能力、試圖調和彼此分歧的能力、

面向	次構面	議題	與協作職能的相關性
			改變他人主意與行為的能力等。
	目前限制	1.目前技術能力在語意判斷上的限制。 2.AI 有其應用範圍上的限制。 3.因 AI 的引入，造成業務執行人員的工作負荷，更強調業務人員必須具備適當的資源管理能力。	目前 AI 的限制多集中在技術能力面向，尤其是中文文義的分析技能，這是目前亟需克服的技術限制。 再者，因著新科技技術的引進，造成組織內部同仁業務量的增加，勢必會影響到內部同仁原本業務的執行規劃，也因此更考驗著主管或承辦同仁對於控制相關資源管理的能力，如時間管理、人員管理等。
	關鍵因素	1.獲得高階主管的支持。 2.注意系統相容性問題。 3.AI 的營運關鍵在系統維運。	從訪談中我們發現目前 AI 的發起，多數可能是業務承辦，較偏向由下而上進行 AI 的引入，對於任何革新作法實與上級主管的支持有高度的關聯性，若目前的主要路徑是由下而上的話，將會非常強調業務承辦者本身的社會性技能用以爭取上級主管支持。 另目前 AI 的引進多數採取委託外包的形式，而為避免系統建構因外包廠商的設計，而產生系統能否相容的限制問題，導致所有與該系統有關的，只能委託該廠商來執行。因此，需要業務承辦同仁具備契約管理的能力，包括如何在契約簽訂時防止發生上述情況，而這也帶出資源管理中另一項職能，亦即資訊與知識管理能力的要求，不必要求瞭解系統是如何建造，但對於系統本身要有一定程度的認識，具備基本認識，即可在契約簽訂上，縮短彼此雙

面向	次構面	議題	與協作職能的相關性
			方資訊不對等的距離。

資料來源：本計畫。

1、 態度價值

態度價值面向主要區分為三個次構面，分別是人機協作的限制、民眾對於人機協作的準備度以及引入人機協作後對於政府的課責問題。茲分述如下：

(1) AI 應用的挑戰

a. 涉及個人隱私

AI 目前遇到最大的應用挑戰在於，它需要非常大量的資料，但也因為需要串接大量資料，往往會牽涉到個人資料隱私的問題。而此卻也常是一種抵換關係的取捨，若是沒有大量的資料匯入，或許我們就是得承受機器無法有效學習的成本，相對來說對於其適用性問題也將大打折扣。

對，AI 為什麼會遇到很大的挑戰就是，它需要很大量的資料，但是又會涉及到個資隱私。可是你沒有給它大量的資料，它就沒有辦法成功的去學習你要學習的問題，可是我們肉眼看人不需要經過它的同意嘛，對不對？（E4）

因此，有受訪者針對資料的個人隱私問題，提出資料必須採取分級分類的作法，而非無限上綱的進行個資保護，這將使得 AI 的發展受到限制。

唯一的就：我把資料作分級分類，哪一些適合拿去作串接的，我怎麼樣把它處理到，我們都講資料的粒度，顆粒的粒，粒度，我把它處理到怎樣的粒度是適合拿去作什麼樣層級的、串接的，我想這件事情是必須落地，而且必須要由公部門來作的。（E2）

b. 主管與承辦的角色認知

受訪的學者專家也提及，目前不論是公私部門在引入 AI 之際，主管與承辦人員往往出現角色上的認知問題。

首先就主管而言，或許未來人機協作的業務非主管親自執行，但為了能夠確實掌握人機協作的功效與瞭解可能的問題，主管仍然需要對科技素養有基本的認識。

但領導者這件事情也是很重要的，他必須要具備這樣的素養，他自己要很懂自己的業務，他也熟悉科技，他不用很厲害，但他要

知道使用科技跟創造科技是兩件事，可以嗎？我們都會用手機，但不代表我們有能力設計手機，但是起碼你要知道手機可以幫助我們解決我們什麼生活的問題。（E4）

再者，受訪專家表示 AI 成功地引入，關鍵不在於演算法，而是在於倒進去的資料長怎樣，因為在過程中機器會自我學習。也因為如此，若倒進去的資料越精準，在人機協作中，人往往會因為擔心機器越來越能處理業務，而危及自身的工作保障問題，產生不願意把資料作最好的紀錄留存。

第一個大概在很多不管是企業端，或是在政府部門這一邊，你說要導入 AI，從企業端的角度，他們的角色是 AI 會不會取代我？私人部門都是這麼想喔...（E2）

很多使用者的，他會擔心他自己跟不上，他當然就會反對你，因為他自己有很多錯誤的想法，一是太難了，他覺得他學不來，二是他覺得會取代他的工作。這些事情因為沒有正確的教育，所以使用者是要給很好的教育跟訓練，讓他知道說科技是怎麼輔助他...（E4）

（2）民眾端的接受度

a. 民眾是否具備足夠的 AI 素養

也有受訪的專家指出，目前我們主要從服務提供的角度，來思考引進 AI 所可能產生的變革、挑戰與機會，但往往會忽略掉，從民眾端，民眾也得適應從對人轉變為對機器，民眾是否具備足夠的素養及能力，來與 AI 互動及對話。因此有受訪者提及，在我們一味以聊天機器人作為與民眾溝通的媒介下，我們忽略了民眾使用中文的多元方式，若再加上目前中文文義的判讀難度，此往往造成聊天機器人之所以失敗的原因。

其實你只要有這些好的音質的資料，基本上他們都可以去作辨識，但我覺得這個是不對的，原因是民眾打電話進來，他會講哪一國的語言你不知道，所以你訓練出來的影音辨識基本上只有國語，你沒辦法去訓練一個國語跟臺語跟英語同時夾雜在一起的所謂的辨識機器人，因為他來自幾個不同的演算法所產生出來的。（E2）

b. 情感運算對民眾使用 AI 的影響

再者，民眾對於機器的信任程度，決定了民眾使用的程度，即便未來可以讓機器具備情緒，模擬出像人一樣的情緒，人類是否願意將機器視為

人，也將是未來政府在應用機器的關鍵因素之一，畢竟有時民眾與政府的接觸往往需要的是與人接觸的溫度。

第一個就是現在已經有很多人在模擬讓機器具備情緒，這個叫做 affective computing，這個叫做情感運算，他們希望能夠讓機器人或者 AI 的介面能夠模擬出像人一樣的情緒。但這件事情涉及到人就特別，他涉及到互動。（E4）

在整個的推動上，剛開始民眾的接受度不好，因為他覺得他想要跟人接觸，不想要透過機器，因為機器冰冷冷的，後來其實漸漸地他有感受到好處...（P15）

（3） AI 引進後的課責問題

a. 政府容錯率低

相較於私部門，受訪者提及公部門的容錯率較低，主要原因在於公家單位無法承受因業務執行錯誤，所導致民眾對其提出質疑或客訴問題。反觀私部門，因高階主管與專案團隊在導入新事物時，接受度通常較高，換言之，他們有足夠的空間允許犯錯，能夠讓他們去做創新性的嘗試，這點與公部門本身的運作體制是較為不同的。

第二個當然公部門在討論 Chatbot 這件事情，他們的容錯率相較於私部門來講他們是非常非常的低的，那非常簡單的一個原因是因為他們承受不起客訴。所以在委託單位，就是你問說你們要不要推 Chatbot 的這件事情，那第一件事情他們都不管你要推什麼 Chatbot，他第一件事情也不會問你要多少經費，第一個都是先問說它的可用性有多高。（E2）

它的價值就是兩個，一個就是跨領域，另外一個就是創意。（E5）

另外一個是 accountability 的問題，也就是說到底課責是否會因為受到 Chatbot 而影響。剛剛有提到從技術面來看它的正確率，如果今天是 Chatbot 回答問題錯而導致服務使用錯誤，那請問這個責任應該是歸咎於技術人員的責任，還是他是那個原始部會的責任。（P3）

b. 演算法的黑盒子困境

課責問題在訪談中常被多數專家學者所提及，公部門無法避免談論此類型的議題，當機器與民眾互動出現爭議時，誰必須承擔該項責任，此

點與演算法的黑盒子有關，對於黑盒子所得出的演算法無法輕易究責，即便是非常單純的業務，只要無法解決課責問題，對於公部門來說就有實際運作上的困難。談及演算法的黑盒子解釋，從深度訪談中可約略歸納為兩項原因：

其一為投入機器的資料格式多元：

但是最麻煩的是資料的型態你沒辦法 control。因為資料進來以後，演算法我可以控制，process 我可以控制，但是輸入進來的 data 到底是怎麼樣，是你沒辦法控制，如果 data 能夠讓你控制的話，這個世界就可以操弄。（E5）

其二從統計角度說明，演算法之所以無法控制的原因在於，統計中的自變數是在研究者依據個人觀察與文獻檢閱所建構出來，透過數據蒐集來驗證影響依變數的關鍵變數為何，然而在演算法中的自變數，則是由類神經網絡自己去決定誰是重要的自變數以及各變數的權重，也因為如此，我們無法確切知道機器演算後的結果是如何判准，以致無法歸咎責任問題。

但是你問它為什麼？到底抓到的背後的元素，這些特徵值是什麼？抓不到。所以 AI 它背後用到的元素太多，它那個是黑盒子，但是它可以精準地預測你那個 Y。可是這個對於形塑理論來說，就會遇到很大的挑戰，就是你不能只告訴我答案，你要告訴我到底是什麼原因？抱歉，深度學習現在就是無法解決你這個 X 的組合跟權力。（E4）

2、事實陳述

(1) 跨域管理

a. 跨域能力需求

在 AI 的應用上之所以特別強調跨領域能力原因在於，大數據演算及資料庫建立其實需要非常多的專業知識與資訊素養，主導者一定要對該領域熟悉，也必須要跨領域來學習科技，讓瞭解該領域的人去學習技術，可以更有效的提升成效。

所以很多的研究案跟很多的計劃案都是成果式的委託研究，就是我只在意成果，所以你最後吐出成果給我就好了。但我覺得這樣的形式必須要改變，要把它改變成作參與式的委託研究，它必須要有一定比例的執行人力是來自於委託單位，我覺得這樣的專案對他們才會有幫助、成長...（E2）

通常現在我們現在會常常發現一個很大的問題，我們會依賴很多科技背景的人來參與我們這個專案，但是根據我們的經驗，他們最大的問題就是他們對本質的業務不熟悉，所以他們一開始的痛點都找錯了... (E4)

b. 部際關係經營

除了跨領域能力的需求外，AI 的運用也必須植基於資料庫的建構，而建構資料庫將涉及龐大資料的串接需求，因此將透過組織內跨單位的資料蒐集與歸納，很多資料必須在串接後才會產生其價值，然涉及到跨單位間的資料串接往往就有本位主義的可能性，此將成為資料庫建構的最大障礙，因此，具備跨單位間的協調能力，也將是未來 AI 有成效的前期條件。

使用者通常比較少做這件事情，因為通常都是由規劃者在做這件事，那裡面涉及到很多 data 的來源，一般使用者沒有權限或是也很難去收集到像剛剛王老師所提到的跨部會的這種資料，這個需要有中央來統籌做這件事情。(E4)

(2) 目前科技限制

a. 語意判斷限制

語意辨識為現在 Chatbot 應用最大的阻礙，目前仍無法處理中文語意判斷的問題，特別是還有中英夾雜、台語混雜的問題，突破目前語意分析技術上的障礙，才能夠真正幫助話務人員回答到民眾的問題，但目前的錯誤率仍過高。

所以 Google 的前 CEO 李開復就講，如果現在在中文的世界裡，你要做新創，是要作這種自然語言處理的，他建議你就不要投資了，因為它失敗的機率太高。因為我們中文的複雜度太難...(E4)

b. 應用範圍限制

有多數受訪者在談論 AI 時，提出 AI 的應用應區分為兩類，有將其區分為一般型或者特定業務型，也有將其區分為專業型與一般型者。但究其本質兩者分類的背後，主要在提出一個 AI 應用的觀念，亦即若要 AI 的應用具備成效，則我們就必須限縮該機器的業務執行範圍，避免將太多業務綁在同一支機器內，越能區分其特定性，將越能提高其執行效能。

我對於公部門行政不熟，但是我覺得有可能是可以把它分成是比較專業型的，或者是比較一般型的，會不會比較合適一點？專

業型像報稅，這些問題可能會比較專業型，他可能比較不會 out of box，這個 chatbot 我相信在專業上面的應用，它應該會是非常的合適... (E5)

可是我要報案我今天發生的事情，我真的適合叫我去面對一個機器人對話嗎？又涉及到機器人可不可以正確來解讀我的話，那時候我很緊張、很焦慮、我語無倫次，我的心裡需要一個真的人，所以這個時候就要作出判斷，科技到底可以做到什麼事跟不能做到什麼事？(E4)

再者，也有受訪者表示，若涉及價值判斷的問題就不適合 AI。這使得本計畫聯想到一種社會讚許的現象。由於 AI 的訓練是大量文本與數據紀錄所造就能提供輔助決策判斷的工具，若 AI 執行的業務涉及到價值判斷的取捨問題，難免人與其合作的當下不會被該機器所影響，而在決策過程中，腦袋中的認知路徑有被 AI 輸出結果影響的可能性。也因此，法律界的專家學者認為，屬於涉及價值的輔助判斷決策將不適合採用 AI 與人的合作模式。

就包含很多決策性的事務，其實是不適合先丟給機器人再回過頭讓我來檢核的，但是我發現現在很多使用 AI 他們都會反過來作，這個是大概在這個部分我覺得應該要討論的...。(E2)

c. 教育訓練的限制

教育訓練的過程中應該是要教導使用者，科技可以如何輔助他，而非反過來讓人去配合科技的語言。原本認為 AI 可以協助公務人員處理日常業務，結果為了讓人可以使用 AI 機器，反而運用大量人力及時間來學習如何與機器共存。

可是我就說你有了這個 APP 是不是就減少了？他說沒有，他(民眾)一樣會跑到警局、一樣會打電話去，他們就安排警員來教他如何使用這個 APP。他遇到很大的問題，他們花了很多警察來教育民眾，說以後你不一定要打電話或親自到警察局來，你可以透過這個工具來跟它作對話...。(E4)

d. 關鍵因素

(a) 高階主管的支持

如同行政革新的變革方法要能獲得實際成效，最重要的仍是高階主管能帶頭引進並表示支持該項業務的執行。AI 在政府機關的引入也面臨

相同的問題，需要高階主管的支持，但目前有受訪者注意到，多為承辦業務單位自發性的由下而上提出創新方案。

因為我發現很多的導入並不是 top-down 的那種導入，是下面的人提案告訴上面說我們可以做這件事情，並不是不好，是所有導入的推動都會很零碎、很片面，它沒辦法是整體性的規劃。政府是一個完整的組織，它所有的東西的推動應該是要有完整的規劃來進行這件事情的...。（E2）

(b)注意相容性問題

目前公部門在引入 AI 多是採委外方式進行，但因受託廠商在建構過程中，基於資訊不對稱下的代理行為，而有可能在系統建置上採取築牆作法，以綁定未來該單位對於系統使用的後續委託案，系統相容性的問題會阻礙該系統業務未來市場標案的自由度。

就是會去做技術牆，就是我把我的東西作的只有我能做，所以你接下來沒關係後面的人其實就比較不容易去碰它。但我覺得在所有的案子評估裡面，未來有一個地方要去作改善，就是其實相容性的這件事情必須要被凸顯出來...。（E2）

(c)關鍵問題在維運

受訪專家也表示，未來公部門內部要能把 AI 做好，重點就在於本身具備能維運該機器的基本技能。

人機協作的重點是在維運，後面的維運這一塊，我目前發現很多的，不管是公部門跟私部門，他們把這樣的東西給導入了以後，其實他們是沒有能力去作維運的，沒有能力去作維運你談很多的工作分配，在現階段我覺得是比較不可行的...。（E2）

(二) 以聊天機器人（chatbot）為探究人機協作現況與未來－焦點座談

本計畫為了能更提高與人機協作職能對話的可能性，因此在深度訪談後，以聊天機器人為個案，邀請產官學界有相關承辦經驗或建造經驗的學者專家或公部門的承辦同仁及科長，一同依其過往的執行經驗，給予本計畫團隊在相關協作職能上更進一步的瞭解，究竟在整體協作職能中，對於公部門同仁而言，那些職能是引進 AI 當下所直接相關的，而有哪些職能是必須再被加以訓練的。

針對焦點座談的資料分析，本計畫初步歸結出以下幾項重點，然多數

在前此深度訪談時已提及，因此，針對人機協作職能的討論，最後將綜合整理深度訪談與焦點座談分析資料，提出目前 AI 系統的引進，在人機協作上可能必須要面對的幾項落差。

1、主管與組織成員的定位與挑戰

態度價值面向主要區分為三個次構面，分別是人機協作的限制、民眾對於人機協作的準備度以及引入人機協作後對於政府的課責問題。茲分述如下：

(1) AI 的引入首重上級主管的支持

如同前此在深度訪談中的發現相同，參與焦點座談的專家學者也提出聊天機器人若要能產生實質的成效，必須獲得上級主管的支持。這一點也在深度訪談時出現，本計畫認為其應該是多數行政革新變革作法下，最重要的關鍵變數，尤其在深度訪談中，有位長期與政府單位合作，協助建構聊天機器人的專家指出，就他的經驗，多數政府內部中倡議使用聊天機器人的，多數是來自於業務單位對於該項業務執行的創新做法、發想與倡議。

我們在推動這個東西的過程都要去拜託長官，我必須要由長官對長官，他們彼此之間對機關的支持，這個案子才推得下去，我幾乎沒有辦法由下而上，我一定要由上而下才能夠去執行這樣的工作。（G5）

跟他有關的他不會是單一機關，會是不同的機關。我把那些機關相關的人都找出來，第一個你要找到人，第二個就是要找一個長官支持，這樣跨機關的事物才有辦法順利地往前推。（G4）

(2) 工作職場環境必須有高度的容錯率

政府內部的工作環境是否能像私部門有高度的彈性與嘗試錯誤的可能性，都影響著聊天機器人是否能夠被有效的被執行。如同新公共管理當時所鼓吹的浪潮，希冀政府能引進市場機制，強調市場精神，藉以改變政府內部的組織結構與運作機制，提高政府的行政彈性，破除僵化體制，注入創新思維與彈性。然凡此種種變革做法，相關研究也均提出，最重要的莫過於高層主導者的支持程度。本計畫認為，創新思維的一體兩面應該也必須承擔創新思維下可能造成的錯誤問題，因此，也有與會學者專家提出政府必須關注容錯率的問題癥結點。

就是通常這些高階主管和專案團隊成員在導入他們對新事物的接受度通常是很高，他們比較敢犯錯也就是說有足夠的空間讓

他們犯錯，所以他們的組織是允許有一些空間讓他們去做某些嘗試，有些公司是只要犯錯你大概就不用升遷。（E4）

（3）高層主管也應對聊天機器人有所概念

專家學者們認為，若高層主管對於聊天機器人不理解，將無法有效的領導單位遂行聊天機器人的業務。如同，我們可以不清楚如何製造智慧型手機，但我們必須瞭解如何使用智慧性手機。因此，我們可以不瞭解如何建構聊天機器人，但主管必須瞭解如何使用聊天機器人，以及在聊天機器人與民眾互動的過程中，必須知道蒐集那些資料來源與相關格式的要求。

同時這些高階主管都有一些很大的特色，高階不一定是最大的那一位，基本上就是我們講的 project's sponsor，它們對於這件事情、東西，AI 背後它們所應用的一些技術的原理，都會有一些掌握，不會是一張白紙，絕對不會，導入之後的專案團隊也不能靠廠商，因為最後很多東西大多數是要靠自己來 input，Chatbot 當然是很明顯，很多也不例外。（E4）

2、組織成員所應具備的心理素質

政府內部同仁在實際引進聊天機器人時，有參與者提及，承辦同仁對於聊天機器人本身有所防衛與抗拒的心態，主要在擔心是否有被取代的風險。而在前此深度訪談中，也有學者專家提及，就其所觀察的政府單位，很多政府同仁所設想到的並非是被取代的問題，而是聊天機器人可以為我執行哪些業務。

剛才老師有特別提到說業務單位是不是非常無私的貢獻出來，這個的確是個問題，因為當我們跟這個業務單位在談這件事情，他們第一個想到的是他們會不會被取代...。（G7）

同時，因為防衛機制的心態問題，導致聊天機器人所需要的默會知識之文本資料未能有效取得，間接的阻礙了聊天機器人的實際執行成效。

你們根本不用人去看，可以省下更多時間去做其他的事情，要不然你們每天加班這樣也不 OK，所以慢慢有卸下他們心防，終於拿出秘笈，那個秘笈還真特別，就是手寫的感覺好像已經很久，其實我們也有觀察到做這樣業務的老手一個一個退休，新手在現有的相關作業規範看不出來的狀況之下...。（G7）

3、Chatbot 所扮演的角色與應注意之問題

(1) Chatbot 僅是輔助角色

參與焦點座談的實務同仁提及，在引進 AI 的過程中，面臨到有些業務承辦人擔心 AI 取代自身工作的憂慮，因此，這也凸顯了前此在深度訪談中的發現，人員面對科技的同時，我們必須具備足夠的人員管理能力，以激勵、發展和指導組織成員，同時，這也凸顯了現階段組織內部對於 AI 的資訊科技認知，有更進一步提升的可能性。此也如同深度訪談中所提出的，不管是業務承辦或者是主管本身，雖然可以不需要擁有建造 AI 系統本身的專業技術，但對於 AI 必須要有基本認識。

這些所謂的 Chatbot 它最大的一個價值在於說它把原本經驗是累積在話務人員，我們的資深話務員他的經驗是很寶貴的，那如果這些資深話務人員能夠訓練這些機器人，等於是把他的經驗留在我們的資料庫裡面，而以後的新進人員，他就可以間接地把資深的前輩的經驗透過這機器人就轉到他身上，這樣會讓新進話務人員在適應他的工作上，我想會有一個蠻大的幫助。（G6）

(2) 強調與外部廠商簽約時的契約管理能力

此呼籲與深度訪談中曾提及系統相容性問題有關，這也涉及到我們在與外部廠商簽約時，如何透過契約的簽訂避免廠商築起技術牆，導致系統只有該家廠商所能建造與營運，限制了未來與其他廠商合作的可能性，這就攸關承辦業務同仁與主管的契約管理能力。同時，為了要能在契約上避免此資訊不對稱所產生的逆向選擇問題，我們本身也必須具備足夠的 AI 素養，而其也和資源管理中的資訊與知識管理能力相關。

還有一個我覺得比較血淋淋的事實的事情，服務的內容會一直改變，但是資訊廠商只會幫你品牌做出來，所謂的智能客服的推薦建議的維運勢必是內部的人做，譬如說我有新的問題，它要導向到哪一個表單或是哪一個資訊系統的 flow，一定是內部的人來瞭解。（P18）

因為這個是我們每天會遇到的困難，現在的這個系統它就是沒有辦法跟我們現在的知識庫的系統同步，它的邏輯是不一樣，所以它就是沒有辦法合作，光是這一件事情，同一個公司同一個產品它都沒有辦法達到的情況之下，對我來講同一件事情它卻要花雙倍的人力在做一些處理，其實是非常大的這個人力的耗損，所以後來 Chatbot 在往更多的局處去推展的過程裡面，我們自己的腳步也會慢下來。（G5）

可能要從地方組織跟包含了承辦，還有包含這些如果是委外或是跟委外單位合作這樣的機制，到底該怎麼做，我覺得說不定更需要這些。（P17）

（3）Chatbot 無法處理大範圍的問題，只能針對特定問題

這種小領域的話，我想經過訓練之後的 Chatbot 應該是可以回答大部份民眾的問題才對，但是如果你今天是一個大範圍，希望它能夠去回應所有市民的問題的話，我覺得這個動作會非常難。（P19）

（4）財務資源的投入問題

此涉及到協作職能中資源管理次構面的財務資源管理，政府部門的工作計畫必須有相對應的預算來支應，而預算的編列是一種抵換關係，當某些業務或單位的預算增加，勢必會壓縮到其他單位或工作計畫的預算分配，因此，若必須要有效發展 AI 的運用層面與成效，誠如深度訪談與焦點座談的實務同仁所指出，必須要有足夠的財務資源挹注，才能有發展空間。

只是說當我們把資源全部都投入在這一塊的時候，它們有可能會在發展其他的事項，比如說高雄市在做智慧城市的這一塊，它們有可能就會沒有心力再去做這些其他的部分，所以我覺得這個是一個資源投入的選擇問題，假如說我們要投入比較大的心力的時候，有可能會造成其他的資源的排擠。（G5）

4、針對 Chatbot 所作的未來建議

（1）跨領域的協調問題

此與協作職能中社會性技能有關，直至目前為止，從深度訪談與焦點座談中，我們關注到多數學者專家均提及跨領域以及跨部會間的協調整合。跨領域所可能發生的場域在於，資訊技術背景者與業務同仁之間的協調問題，另一個場域則在於業務同仁為了遂行其業務，而必須與其他單位或部會間在資料上的串接問題，而這就涉及到相關溝通、協調與說服能力的要求。

在做的過程裡面，設計這個軟體的公司幾乎幫不上我們，原因是因為它根本不懂公部門的語言，結果最後是我們話務中心要先訓練一組人學會這個軟體，所謂智能客服的題組型態是怎麼樣？（G5）

為什麼跟廠商之間的溝通會有 gap 在？需求講的很清楚，廠商也告訴你說這個技術就是這樣，但是總是會沒有辦法達到那個我們所預期的，我個人觀察是因為在溝通上花太少時間，為什麼在溝通上會花太少時間？我想這是我主觀的認為是公務人員對於我們所從事的這個業務，充其量只是這個組織的代理人，就代理制度最大的缺失是在於說，它所產生的這個成本通常不會是自己扛，那不會是自己扛的時候，請問對公務人員而言，別人這樣做我幹嘛花那麼多心思？（G7）

（2）工作流程與任務的重新設計

AI 的引入勢必會影響到原本組織內部的工作流程與人員在流程中所擔負的業務分派，除了前此資料分析中所提出的跨領域的協調問題外，回到政府組織內部，因著新資訊科技的引入，也得重新審視工作任務與流程的調整，也因流程的重新設計而需要不同單位間的協調整合，因而涉入其中的成員也必須具備協調溝通的能力。

以整個結構跟制度來說，這樣的技術和服務導入沒有很真正的融入到整個行政作業的體制上的調整。比如說 Chatbot 進來並沒有因為人力結構或作業流程上面有一個配合的調整...。（G4）

當然最後一個就是組織改變。像現在這些原來的政府單位裡面絕對沒有這種知識點工程師或者是資料訓練師，這到底是歸 IT 單位還是歸哪個單位，現在大部分都還是歸 IT 單位，可是這個其實跟業務單位是很密切的...。（P17）

（3）調教師的設計

此點建議是公務界同仁所提出，重點其實與深度訪談多數學者專家所提出的相同，亦即政府內部知識管理系統的建構問題。我們是否有一套知識管理體系，其目的在儲存我們本身業務執行所產製的資料，而在 AI 引進的同時，我們必須要知道，如何使該知識系統與 AI 的機器有效連結。因此，有公務同仁建議可以設計類似轉譯角色的承辦同仁，讓他能夠提供給機器，它所能夠理解的資料，然追根究柢，這一點建議實與組織內部的知識管理系統，以及組織對內部同仁在 AI 的資訊素養訓練是否足夠有關。

而且如果我來選擇的話，我會願意去增加在推廣 AI 的過程中需要，比如說有一些所謂的調教師，他必須要能夠把業務上的知識轉換成機器可以理解的東西，因為我覺得這個東西有一個很大的價值在於說，我們能夠把這些原本是存在人類大腦裡面的經

驗，透過這樣的一個過程，轉換成一個資料庫放在我們的組織內部的話，我就比較不用擔心人員的流動對我的服務穩定性造成的影響。（G6）

（4）語意上的限制問題

聊天機器人在目前，依據學者專家所指出的，在中文文義上的理解與英文相比，明顯是無法達到相同的成效。這樣的觀點在深度訪談時也多次被提及，顯然這在科技應用上仍有其必須突破的地方，但這確實可以提供政府單位另一個思考運用聊天機器人的可能，也就是針對特定業務而非一般型業務，將聊天機器人所服務的範圍鎖定在特定服務上，而非讓聊天機器人去回答較為廣泛、一般性與民眾需求有關者。而這點也在前此深度訪談中被多數學者專家所提出。

我是很同意 G6 的觀察，畢竟也是在話務中心實際在維運的人，如果我能夠選擇的時候，我會最想做的事情是突破現在語意分析的技術障礙，能夠幫助我們的人員在一個題目進來的時候，很快速的從我的資料庫裡面撈到適當的題目回答。（G5）

再者，本計畫除透過焦點座談瞭解政府部門引入 AI 對於人機協作的挑戰與經驗反思外，並為了讓討論能更聚焦而以 Chatbot 為討論個案；同時，我們也試著透過問卷調查的方式針對政府機關中個體層次瞭解本計畫所整理的七項職能中，對其重要性認知與目前具備程度進行大範圍的瞭解。下表即為此次問卷調查在人機協作職能所整理的結果，相關問卷調查流程請參閱本計畫研究設計一章。從表中可知，受訪者在七項職能中認為重要性最高的三項依序是資安倫理、資訊解讀以及數據分析能力。從重要性認知程度來看，明顯可知政府體制引入 AI 仍然關注於在資料串接的處理上，是否能在有效執行任務的基礎上更能兼顧民眾資料的個資保護與合乎倫理道德規範的的用，本計畫認為主要仍然與政府必須回應民眾的責任有關，課責仍是政府首要面臨的核心價值與考驗。再者，多數受訪者也認為 AI 除了資訊安全的首要外，對於 AI 所產製的資訊也必須要有解讀的能力，以及面對複雜多變與非結構化的環境下，也必須具備有數據分析以界定政策問題的能力。

同時，我們進一步觀察，在得知大家對於各項能力的重要性認知程度後，我們來看多數人在這幾項能力目前的具備程度，兩相交叉即可得知目前政府機關內部最重要且必須被盡快扶植的能力為何，承上調查資料分析的表 32 之呈現。從具備程度的平均數來看，目前具備程度最低者為系統維運能力，次之者為數據分析能力，再者為資訊解讀能力。而在這三者最不具備的能力中，解讀資訊與數據分析能力則是被認為重要性程度前

三高中的兩項，可見這兩項能力有立即需要加以培訓之處。而最重要的是，搭配前此深度訪談與焦點座談的資料分析發現，系統維運能力為多數政府機關同仁與專家學者所一致認為，若要有效發揮 AI 的效率與效能，則政府機關內部同仁必須具備一定程度的系統維運能力，因為政府機關內部同仁具備業務所需的專業知識，若同時能具備基礎的系統維運能力，則更能為系統灌注有效的資料以及知道該如何儲存資料。

三、 調查分析與深度訪談暨焦點座談的資料對話

以下本計畫將試著從上述質、量化分析結果進行資料間的對話，嘗試更完整的對目前中央與地方政府在引入 AI 之際，對於組織運作與結構上的改變以及可能必須面對的問題進一步討論。

(一) 從員額數不足談 AI 系統的建置問題

從上述量化調查針對中央部會目前已引入及規劃引入兩類對象的分析結果可知，此兩類中央部會在 AI 系統建置途徑上均以委託外包方式進行。然而，從樣本描述統計可知，各部會平均資訊專責人力約為 27(26.67) 人，其中參與規劃、採購或管理 AI 系統資訊人力平均約為 2(1.44) 人。從這樣的委外現況來看，此似乎凸顯了幾點問題，首先是人員員額數不足，從訪談資料我們可知，多數受訪者提及，引入 AI 基本上不是甚麼太大的問題，只要有充足的預算支應就可建置，但實際建置後要能夠有效運作需要有良好的維運能力。顯然，以目前的員額數要能夠有效營運似乎顯得捉襟見肘。此為第一點問題，人力不足的問題。

再者，目前建置均委託由外部廠商為之，也或許是因為人員不足，這是目前不得不的作法，然而此做法將形成另一個隱藏性的問題，主要在於專業上的綁架。此綁架表現在兩個層面，其一，不僅是政府內部對建置系統上專業知識領域的不足，其二，也形成另一種未來不得不再找這家廠商的綁架問題。進一步細說，就第一點而言，由於專業知識上的不對稱，屬於委託外包這種競爭型的協力案件，無法發揮該有的競爭機制，造成買家形成逆向選擇的可能，做出對於自己不利的決定，同時也會在研擬契約當中，因著資訊與知識領域的不對稱，在契約研擬上也造成買賣雙方的不對等地位，而致使政府端簽下不公平契約的可能性；就後者而言，因著專業不對稱，在建置過程中，從深度訪談的受訪資料可知，會有某些廠商以其專業技術與知識的優勢，在建置系統上進行設計，以防堵未來買方系統出現問題或有需要新增建置時，無法再委託其他廠商進行維運與其他建置的可能，因為在程式設計上已被框架住，只有原受託廠商能處理。因此，委託外包的作法不是不可行，然本計畫認為，似也應充足組織內部資訊人力的專業認知與員額上的補足。當然這也形成另一個問題，在於人力資源上的訓練問題。

從目前員額數不足致使系統建置主要委由外部人力來執行外，似也凸顯了另一個問題在於組織內部對於 AI 引入的重視程度與訓練問題。本計畫的想像是，若 AI 是未來不可逆與不可違背的一項發展，當然，若再加上量子電腦未來的突破，AI 勢必成為一種日常而不是趨勢，若有這樣的想像，這一塊的政策發展是必須受到高層主管支持的，若有高層主管的支持，應或許能在組織內部得到該有的預算挹注與成員投入。當然，另一方面，因著人員數不足而致使在委託外包資訊不對稱上所產生的問題，若想要立即獲得解決，或許可從組織內部訓練著手。短時間內要讓成員能夠有動力去執行 AI 的系統建置與規劃，當然包括日後創新的服務設計等，要讓成員認為有能力去勝任這樣的工作，從組織行為激勵的角度來看，或許可從自我效能的角度出發。要讓成員具備自我效能的其中一個方式是透過教育訓練方式進行，以組織成員的內部訓練讓他們在短時間內獲得執行某項工作任務的經驗，透過經驗的吸收與內化，增進其認為自己有那樣的能力來處理工作，進而促使其更能有效與有意圖投入此創新公共服務的業務範疇。

(二) 從 AI 系統的資料類型看跨機關的協調合作

另外一個重點所凸顯的是，從調查結果發現多數已導入或規劃導入 AI 的部會，在 AI 系統使用的數位資料類型上，已導入的部會中，多數使用的是本機關內部資料及機關內部的開放資料，整體來說也都是機關自己內部資料，而規劃導入的部會中，目前則均為機關內部資料為主。對於這樣的發展，本計畫從深度訪談所獲得的經驗資料顯示，不由得產生一種憂慮與突顯目前政府部會的另一種不足。所憂慮的是，受訪專家學者均指出，AI 智能系統要能夠發揮強大的效果，譬如其應用範圍的層面要能擴大的話，基本上不可能僅有單一機關內部資料而已，若僅侷限於目前機關內部資料，則此 AI 系統的效益必定有限。因此，未來 AI 智能系統要能有效運用並讓受服務的對象有感，則勢必得讓資料進行跨機關間的資料串連與應用，唯有如此才能擴大民眾的服務範圍，也唯有如此才能產生智慧效益。而也因為需要跨機關間的資料串連，讓我們產生一種憂慮，亦即跨機關間資料整合的可能性問題。目前的現況發展之所以無法達到跨機關間的資料整合，主要的問題根源仍然需要進一步的釐清，始能對症下藥，然而目前各機關對於資料使用的發展現況，可能對未來而言，或許會產生某些限制，因此，若能讓跨機關間達到資料的整合，應該會是接下來必須要面對的第一步。

同時，因著資料在機關間的串流與整併，也凸顯另一個必須要面對的是，機關內部同仁對於跨機關間的協調整合是否有足夠的能力來加以因應。如同我們在深度訪談中的發現，多數受訪者均認為在人機合作或更往前的系統建置階段，均再再講求跨機關成員以及機關間的協調合作與整

合能力，而此凸顯了在 AI 引入官僚型模的當下，它讓以往為了要增加效率的組織結構垂直分層與水平分化的作法，必須在切割之後再進行必要的整合。因此，在既有組織結構的分化外也必須存有分化後的整合機制，這將是未來 AI 體制下應該有的演算型官僚型模的特性。

(三) 從文官引進 AI 態度談影響 AI 建置與營運的有效因素

另外從量化調查研究的迴歸模型中，我們發現，現階段文官對於運用 AI 的態度可區分為三大類，分別是預期正向成效、預期承擔責任以及預期權力會受到限制三種。多數受測對象均認為 AI 對於本身的工作任務應是能夠帶來正面效益的，再者，也不乏有人擔心因著 AI 系統地引入，會為自己的工作量帶來負荷以及承擔未知的風險。此數據調查結果與深度訪談所獲實證資料是相互呼應的。水能載舟亦能覆舟，多數受訪對象點出 AI 系統確實能為組織帶來極大效益的想像，但這樣的想像從深度訪談中我們知道，必須輔有上下層級間的支持始能克盡其功，這必須是組織內部由上到下及由下到上的配合。然從深度訪談中得知，目前政府內多數對於 AI 系統的應用範圍與可能性仍多處於不甚清楚的狀態，因此，多少形成一種不確定氛圍，在此氛圍下應多少會影響到成員的工作態度。從受訪資料可知，在量化數據上發現成員之所以會擔心工作量負荷的增加，主要是來自於一開始引入 AI，需要大量將機關內部的資料轉變成 AI 系統所可以讀取的資料格式，這部分除了需要委外廠商的協助外，主要執行者仍會是機關內部成員，因為唯有機關內部成員才較清楚有那些資料以及資料所代表的意義。也因此，在未享受到 AI 系統運作成效之前，機關內部成員已先享受到必須先付出的代價，因著這個代價更使其產生對於 AI 系統引入的排斥與不安。

然而，組織成員對於 AI 系統引入的不安感受，從迴歸模型中發現，可以透過組織內部管理系統的支持來達到有效的降低，此變數間關係的驗證發現與深度訪談結果相同，多數受訪專家也認為要能有效發揮 AI 效用，則組織內部的高階主管要能先發起這樣的聲明以產生效尤。同時從模型中我們發現，資料整備度本身也會影響受測對象對於 AI 的態度，不管是正面期待的受訪者，當其感受資料整備度越高時，將越正向期待 AI 的發展；而另一方面，對於承擔責任越高的受訪者而言，資料整備度越高，其預期承擔的責任也越高，事實上這也非常符合常理。因為資料整備度越高，表示機關所擁有的資料完整性越高，越有可能做到跨機關間的資料整合，當然，當機關間的資料串接越完整，所能擔負的 AI 智慧服務層面更廣，對於受服務對象而言，跨機關的資料串接將越能得到完整的服務，但相對的，若產生服務上責任歸屬問題，因資料非僅來自單一機關，在課責上將形成另一種無形壓力與不確定性，間接讓機關內部同仁感受到對於承擔責任的增加。

第三節 才能管理之實驗分析

一、實驗結果分析

本計畫透過法務部保護司所提供之名單，排除長官與主任之人選後寄送電子郵件予全國 200 位觀護人，邀請參與本次實驗。共有 111 位參與受試，扣除掉未完成實驗的受試者，最後符合資格之有效樣本為 102 份（回收率為 56%）。

（一）問卷樣本基本資料之描述統計

在本實驗的有效問卷中，基本資料分別有性別、任職地區、擔任觀護人年資、年齡、最高教育程度與學齡、最高教育背景進行敘述統計（詳見表 35）。就性別方面，女性佔觀護人比例之多數，共有 148 人，佔全體之 74%，而男性則有 52 人，佔 26%；在任職地區方面，任職於北部地檢署之觀護人有 85 人（佔 42.5%），中部地區有 47 人（23.5%），南部地區有 59 人（29.5%），而任職於東部及離島地區分別有 7（3.5%）、2 人（1.0%）；在年資方面受試者擔任觀護人之年資分佈較為平均，最多為年資為 11 至 15 年（佔 24.0%），最少的為 31 至 35 年，僅佔全體受試者之 1.5%；在年齡方面：受試者之年齡集中於 30 至 53 歲，佔 84.7%；在最高教育程度與學齡方面：以碩士學歷（學齡 19 年）最多，佔全體 65.8%。技術學院與博士學歷之受試者最少，分別僅有 1 人；在最高學歷學科背景方面，以法律與教育及犯罪防治最多，分別佔 21.6%和 25.2%。就讀工程科學及自然科學之受試者最少，分別僅有 1 人。

本實驗將受試者分成三個組別，依照個案再犯處遇評估有無外部資訊介入，以及介入方式為何區分，分別有（1）控制組：不受任何外部資訊介入，僅詢問受試者是否要修正答案；（2）AI 介入組：提供 AI 決策系統之判斷結果，詢問受試者是否要修正答案，AI 決策系統所提供之答案可能為正確或錯誤；以及（3）其他同仁介入組：蒐集控制組與 AI 介入組的判斷結果，詢問受試者是否要修正答案，同仁之答案可能為正確或錯誤。以下的所有結果分析將「AI 介入組」以及「其他同仁介入組」分別以實驗 Arm1、實驗 Arm2 組代稱。

表 35：問卷樣本基本資料整理

項目別	有效樣本數	百分比	累積百分比
1. 性別			
女	148	74.0%	74.0%
男	52	26.0%	100%
2. 任職地區			
北部地區	85	42.5%	42.5%

中部地區	47	23.5%	66.0%
南部地區	59	29.5%	95.5%
東部地區	7	3.5%	99%
離島地區	2	1.0%	100%
3. 年資			
1~5 年	23	11.5%	11.5%
6~10 年	25	12.5%	24.0%
11~15 年	48	24.0%	48.0%
16~20 年	33	16.5%	64.5%
21~25 年	39	19.5%	84.0%
26~30 年	29	14.5%	98.5%
31~35 年	3	1.5%	100%
4. 年齡			
22~25 歲	1	0.9%	0.9%
26~29 歲	6	5.4%	6.3%
30~33 歲	13	11.7%	18.0%
34~37 歲	16	14.4%	32.4%
38~41 歲	16	14.4%	46.8%
42~45 歲	10	9.0%	55.9%
46~49 歲	22	19.8%	75.7%
50~53 歲	17	15.3%	91.0%
54~57 歲	9	8.1%	99.1%
58~61 歲	1	0.9%	100%
5. 最高教育程度			
技術學院	1	0.9%	0.9%
科技大學、大學與其 同等學歷	73	65.8%	66.7%
碩士	36	32.4%	99.1%
博士	1	0.9%	100%
6. 最高教育程度之學齡			
15 年（技術學院）	1	0.9%	0.9%
16 年（科技大學、大 學與其同等學歷）	73	65.8%	66.7%
19 年（碩士）	36	32.4%	99.1%
25 年（博士）	1	0.9%	100%
7. 最高學歷之學科背景			

法律相關	24	21.6%	21.6%
社會工作相關	21	18.9%	40.5%
心理、諮商相關	19	17.1%	57.7%
教育與犯罪防治	28	25.2%	82.9%
其他社會科學領域	6	5.4%	88.3%
工程科學相關	1	0.9%	89.2%
自然科學相關	1	0.9%	90.1%
其他	11	9.9%	100%

資料來源：本計畫。

(二) 樣本隨機控制分組統計

本次實驗將受試者依照性別、年資、年齡以及最高教育程度學齡平均分成控制組、實驗 Arm1 組（AI 介入）以及實驗 Arm2 組（其他同仁介入），三組有效樣本數分別為 38、29 以及 35 份，以下表 36 呈現各組基本資料之分佈情形。

表 36：各組樣本基本資料統計

項目別	計算	控制組	實驗 Arm1 組	實驗 Arm2 組
有效樣本數		38	29	35
性別	女	50	49	49
	男	17	17	18
年資	平均數	16.55	16.50	16.36
	標準差	7.855	6.869	9.058
	最大值	30	27	32
	最小值	1	4	1
年齡	平均數	43.50	41.50	42.24
	標準差	8.8496	8.8215	7.9515
	最大值	55.5	59.5	55.5
	最小值	27.5	27.5	23.5
最高教育程度學齡	平均數	17.31	16.85	16.95
	標準差	1.922	1.417	1.413
	最大值	25	19	19
	最小值	16	15	16

資料來源：本計畫。

(三) 實驗前調查題目

在提供案例予受試者作觀護處遇判斷前，本實驗事先詢問受試者對於工作狀況之自我評估以及「檢察機關案件管理系統」導入於工作之使用情形，並接著詢問受試者對於機器演算法導入再犯評估工作之期望及想法，以瞭解觀護人當前的工作現況，並觀察在經過實驗後受試者的認知有無改變或受到影響。關於各題組之答題狀況於底下作說明：

1、觀護人工作狀況與「檢察機關案件管理系統」使用現況

問卷 A4.1 至 A4.5 之題組為觀護人工作狀況與「檢察機關案件管理系統」(以下簡稱案管系統)使用情形之調查，原始問卷設計使用李克特六等分量表，並將其轉換為分數，而在下**表 37**中量表底下所呈現的數字為受試者填答之百分比。

- (1) A4.1 題：工作使用案管系統之頻率由低至高分別為非常低（1 分）、低（2 分）、有點低（3 分）、有點高（4 分）、高（5 分）以及非常高（6 分）。三組受試者大多皆勾選「非常高」之選項，平均分數皆高於 5.6 分（控制組：5.67 分；實驗 Arm1 組：5.82 分；實驗 Arm2 組：5.92 分），且彼此差異並不顯著，顯示當今觀護人使用案管系統之頻率相當高。
- (2) A4.2 題：由低至高分別為非常不容易（1 分）、不容易（2 分）、有點不容易（3 分）、有點容易（4 分）、容易（5 分）以及非常高（6 分）。就案管系統的易用程度方面，過半支持者持肯定態度（控制組：4.00 分；實驗 Arm1 組：4.35 分；實驗 Arm2 組：4.47 分），彼此差異不顯著。
- (3) A4.3 題：工作量之自評由低至高分別為非常少（1 分）、少（2 分）、有點少（3 分）、有點多（4 分）、多（5 分）以及非常多（6 分）。多數受試者認為工作量是多的，平均分數皆高於 4.5 分，彼此差異不顯著。
- (4) A4.4 題：工作效率由低至高分別為非常不容易（1 分）、不容易（2 分）、有點不容易（3 分）、有點容易（4 分）、容易（5 分）以及非常高（6 分）。三組受試者多半認為自己工作效率是高的，彼此差異不顯著。
- (5) A4.5 題：工作壓力低至高分別為非常少（1 分）、少（2 分）、有點少（3 分）、有點多（4 分）、多（5 分）以及非常多（6 分）。多數受試者自評工作壓力落在「有點大」之區間，平均分數皆高於 4 分（控制組：4.38 分；實驗 Arm1 組：4.15 分；實驗 Arm2 組：4.53 分），彼此差異不顯著。

表 37：觀護人工作狀況與「檢察機關案件管理系統」使用現況調查

A4.1 請問您近半年來執行工作業務時，需要使用到「檢察機關案件管理系統」的頻率有多高？									
組別	非常低	低	有點低	有點高	高	非常高	平均分數	標準差	Oneway ANOVA F-value (p-value)
控制組	5.1	0	0	0	7.7	87.2	5.67	1.132	1.174 (0.313)
Arm1	0	0	0	2.9	11.8	85.3	5.82	0.459	
Arm2	0	0	0	0	7.9	92.1	5.92	0.273	
A4.2 請依據您的使用經驗，評估組織目前使用之「案管系統」的易用程度？									
組別	非常不容易	不容易	有點不容易	有點容易	容易	非常容易	平均分數	標準差	Oneway ANOVA F-value (p-value)
控制組	0	7.7	28.2	23.1	38.5	2.6	4.00	1.051	1.944 (0.148)
Arm1	0	11.8	14.7	14.7	44.1	14.7	4.35	1.252	
Arm2	0	2.6	15.8	23.7	47.4	10.5	4.47	0.979	
A4.3 相較於其他同仁，您覺得過去半年您所承擔的工作量為何？									
組別	非常少	少	有點少	有點多	多	非常多	平均分數	標準差	Oneway ANOVA F-value (p-value)
控制組	0	0	7.7	30.8	46.2	15.4	4.69	0.832	0.474 (0.624)
Arm1	0	0	5.9	50.0	32.4	11.8	4.50	0.788	
Arm2	0	0	5.3	18.4	36.8	2.6	4.58	0.919	
A4.4 相較於其他同仁，您覺得過去半年您的工作效率為何？									
組別	非常	低	有點	有點	高	非常	平均	標準差	Oneway ANOVA F-value

	低		低	高		高	分 數		(p-value)
控制組	0	5.1	23.1	33.3	28.2	10.3	4.15	1.065	0.367 (0.694)
Arm1	0	2.9	29.4	41.2	20.6	5.9	3.97	0.937	
Arm2	0	5.3	18.4	36.8	36.8	2.6	4.13	0.935	
A4.5 相較於其他同仁，您覺得過去半年的工作壓力有多大？									
組別	非常 小	小	有點 小	有點 大	大	非常 大	平均 分數	標準 差	Oneway ANOVA F-value (p-value)
控制組	0	2.6	5.1	48.7	38.5	5.1	4.38	0.782	1.836 (0.164)
Arm1	0	5.9	14.7	38.2	41.2	0	4.15	0.892	
Arm2	0	0	7.9	47.4	28.9	15.8	4.53	0.862	

資料來源：本計畫。

備註1：六等量表底下所呈現之數字為各組受試者填答之百分比。

備註2：significance level, *, $p < 0.1$, **, $p < 0.05$, ***, $p < 0.01$ 。

2、觀護人對機器演算法導入再犯評估之期望

題組 A5.1 至 A5.3 之設計主要基於期望地位理論 (Berger, 1977; Wagner & Berger, 1993; Foschi, 1985) 觀點，請受試者想像機器演算法系統導入於再犯評估工作上的情形，以及所需花費之成本及可能帶來之效益為何。原始問卷設計使用李克特六等量表，並將其轉換為分數（計分方式為：非常不同意=1 分、不同意=2 分、有點不同意=3 分、有點同意=4 分、同意=5 分、非常同意=6 分），分析結果如表 38。

- (1) A5.1 題：三組受試者之回答平均分數皆高於 4 分（控制組：4.10 分；實驗 Arm1 組：4.26 分；實驗 Arm2 組：4.05 分），顯示多數受試者是同意整體而言此類機器演算法系統對工作是有幫助的，而三組間差異並不顯著。
- (2) A5.2 題：多數受試者也認為要與該系統合作需要花費額外成本，各組平均皆高於 4 分（控制組：4.59 分；實驗 Arm1 組：4.59 分；實驗 Arm2 組：4.84 分），彼此差異並不顯著。
- (3) A5.3 題：至於對系統的導入能否提升工作效率方面，控制組之平均分數較高（4.36 分），而實驗 Arm1、實驗 Arm2 組之平均分數也趨近「有點同意」之 4 分，分別為 3.97 分、3.95 分，彼此差異並不顯著。

綜上，對於機器演算法導入再犯評估工作方面，受試者多半認為是有助益的，然而同時多數受試者也認為與系統合作方面需要花費額外的成本。而對於該系統是否能提升工作效率，相比之下較持保留態度。

表 38：實驗前觀護人對機器演算法導入再犯評估之期望調查

A5.1 該系統將對我工作有所幫助。									
組別	非常不同意	不同意	有點不同意	有點同意	同意	非常同意	平均分數	標準差	Oneway ANOVA F-value (p-value)
控制組	0	5.1	10.3	53.8	30.8	0	4.10	0.788	0.508 (0.603)
Arm1	0	5.9	14.7	35.3	35.3	8.8	4.26	1.024	
Arm2	2.7	2.7	13.5	51.4	27.0	2.7	4.05	0.941	
A5.2 未來我與該系統的合作會需要額外成本（例：時間成本、適應成本）。									
組別	非常不同意	不同意	有點不同意	有點同意	同意	非常同意	平均分數	標準差	Oneway ANOVA F-value (p-value)
控制組	0	2.6	0	48.7	33.3	15.4	4.59	0.850	1.184 (0.310)
Arm1	0	2.9	5.9	35.3	41.2	14.7	4.59	0.925	
Arm2	0	0	0	27.0	62.2	10.8	4.84	0.602	
A5.3 未來若讓該系統輔助（如：提供一些可參考的綜合資訊）我目前的觀護人工作，我的工作效率會提升。									
組別	非常不同意	不同意	有點不同意	有點同意	同意	非常不同意	平均分數	標準差	Oneway ANOVA F-value (p-value)
控制組	0	2.6	10.3	41.0	43.6	2.6	4.33	0.806	1.954
Arm1	0	11.8	14.7	47.1	17.6	8.8	3.97	1.087	

Arm2	0	10.8	13.5	48.6	24.3	2.7	3.95	0.970	(0.147)
-------------	---	------	------	------	------	-----	------	-------	---------

資料來源：本計畫。

備註1：六等量表底下所呈現之數字為各組受試者填答之百分比。

備註2：significance level, *, $p < 0.1$, **, $p < 0.05$, ***, $p < 0.01$ 。

在進入實驗前，本計畫最後詢問受試者對於未來對機器演算法導入輔助再犯風險評估工作之支持程度為何，三組回答區間主要集中在「有點支持」與「支持」之區間，平均分數皆高於 4 分，且組間並無顯著差異，分析結果如表 52。

表 39：實驗前受試者對機器演算法支持程度調查

A6 承上題，若「Parole-ML-AI 1.0」系統被導入且用以輔助您目前「再犯風險」評估的工作，您的支持程度為何？									
組別	非常不支持	不支持	有點不支持	有點支持	支持	非常支持	平均分數	標準差	Oneway ANOVA F-value (p-value)
控制組	1	1	3	16	14	4	4.36	1.038	0.125 (0.882)
Arm1	0	2	3	13	12	4	4.38	1.015	
Arm2	0	1	6	16	10	4	4.27	0.962	

資料來源：本計畫。

備註1：六等量表底下所呈現之數字為各組受試者填答之百分比。

備註2：significance level, *, $p < 0.1$, **, $p < 0.05$, ***, $p < 0.01$ 。

(四) 實驗題目

在實驗階段中，本計畫共提供 18 個歷史個案，分別有 10 個案例為再犯（案例一至十）以及 8 個案例為未再犯（案例十一至十八）案例，隨機在四個階段將案例分配予受試者。換言之，每位受試者會隨機被分配到兩個有再犯的案例以及兩個未再犯之案例，共四個案例。實驗調查要求受試者針對每個個案作再犯處遇評估。在案例的分配上，除了案例七在系統分配上比例稍低（共 25 作答次數，有效百分比佔 3.1），其餘案例在作答次數上皆相當平均，有效百分比落在 4.1 至 6.8 之間，詳細結果如表 40。

1、 案例分配

表 40：實驗各案例分配次數及比例

案例	次數	有效百分比	累積百分比
案例一	40	5.0%	5.0%
案例二	40	5.0%	10.1%
案例三	42	5.3%	15.3%
案例四	48	6.0%	21.4%
案例五	41	5.2%	26.5%
案例六	41	5.2%	31.7%
案例七	25	3.1%	34.8%
案例八	33	4.1%	38.9%
案例九	44	5.5%	44.5%
案例十	45	5.7%	50.1%
案例十一	48	6.0%	56.2%
案例十二	46	5.8%	61.9%
案例十三	51	6.4%	68.3%
案例十四	47	5.9%	74.2%
案例十五	54	6.8%	81.0%
案例十六	52	6.5%	87.6%
案例十七	49	6.2%	93.7%
案例十八	50	6.3%	100.0%
總計	796	100.0%	

資料來源：本計畫。

2、 對每個個案之再犯評估分析

在 B1 至 B7 題中，本實驗請受試者在觀看完個案的背景資料後，判定其危險及再犯的風險程度，最後決定對該個案之處遇級別。下表 X 整理以上七題在各個案之答題統計，已呈現整體樣本的答題樣貌。本實驗之 B1 至 B4 題請受試者對個案可能之背景、認知、接受觀護治療的配合程度等作想像，而 B5 及 B6 則請受試者判定個案的危險程度與再犯風險程度，最後於 B7 則是針對該個案作出再犯處遇評估。各題之選項以及計分方式分別說明如下：

- (1) B1：題目為「該個案犯罪行為的產生受到不良的社會影響的程度為何」，答案選項與計分方式分為「極低度影響」（1 分）、「低度影響」（2 分）、「中度影響」（3 分）、「中高度影響」（4 分）以及「極高度影響」（5 分），「無法判斷」則不納入

計算。

- (2) B2：題目為「個案在社會中受到排斥且缺乏對他人的關心，產生與社會疏離的程度為何」，答案選項與計分方式分為「極低度疏離」（1分）、「低度疏離」（2分）、「中度疏離」（3分）、「中高度疏離」（4分）以及「極高度疏離」（5分），「無法判斷」則不納入計算。
- (3) B3：詢問該個案問題解決技巧的認知不良情形，答案選項與計分方式分為「極低度認知不良」（1分）、「低度認知不良」（2分）、「中度認知不良」（3分）、「中高度認知不良」（4分）以及「極高度認知不良」（5分），「無法判斷」則不納入計算。
- (4) B4：詢問該個案願意接受觀護與治療的配合程度，答案選項與計分方式分為「極低度配合」（1分）、「低度認知配合」（2分）、「中度認知配合」（3分）、「中高度認知配合」（4分）以及「極高度認知配合」（5分），「無法判斷」則不納入計算。
- (5) B5：詢問受試者對於該個案的危險級別判定，答案選項與計分方式分為「極低度危險」（1分）、「低度危險」（2分）、「中度危險」（3分）、「中高度危險」（4分）以及「極高度危險」（5分），「無法判斷」則不納入計算。
- (6) B6：詢問受試者對於該個案的再犯風險及別判定，答案選項與計分方式分為「極低度危險」（1分）、「低度危險」（2分）、「中度危險」（3分）、「中高度危險」（4分）以及「極高度危險」（5分），「無法判斷」則不納入計算。
- (7) B7：詢問受試者對於上述之判定，對該個案作處遇評估決策，共分為第一級（高度監管及輔導，以核心案件列管，實施特殊處遇）、第二級（中度監管及輔導，尚需加強列管、輔導或其他事由以核心案件列管）以及第三級（低度監督級輔導），計分方式分別為1、2及3分。

下圖 24 至圖 30 為 B1 至 B7 各個案例回答之平均分數，並於平均分數下方括號標示標準差。

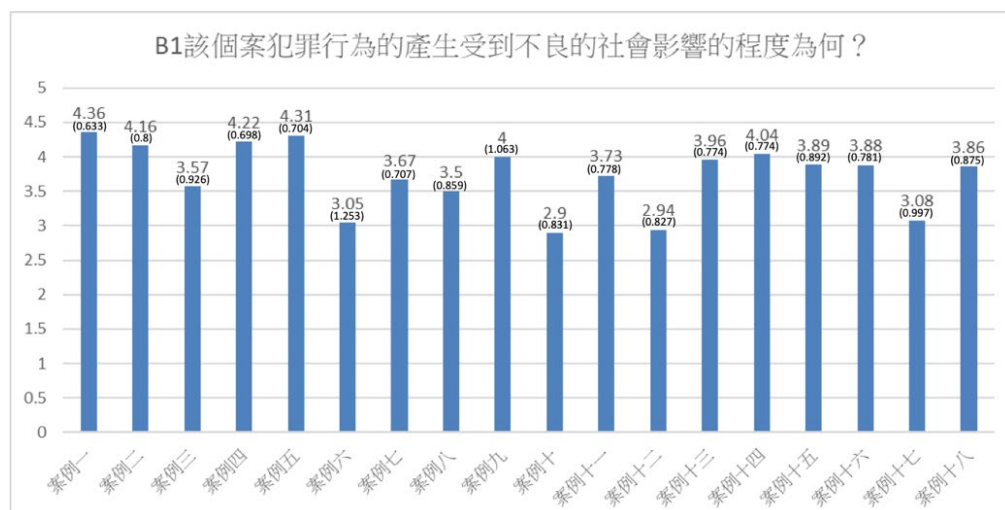


圖 24：受試者回答 B1「該個案犯罪行為的產生受到不良的社會影響的程度為何？」之平均分數

資料來源：本計畫。

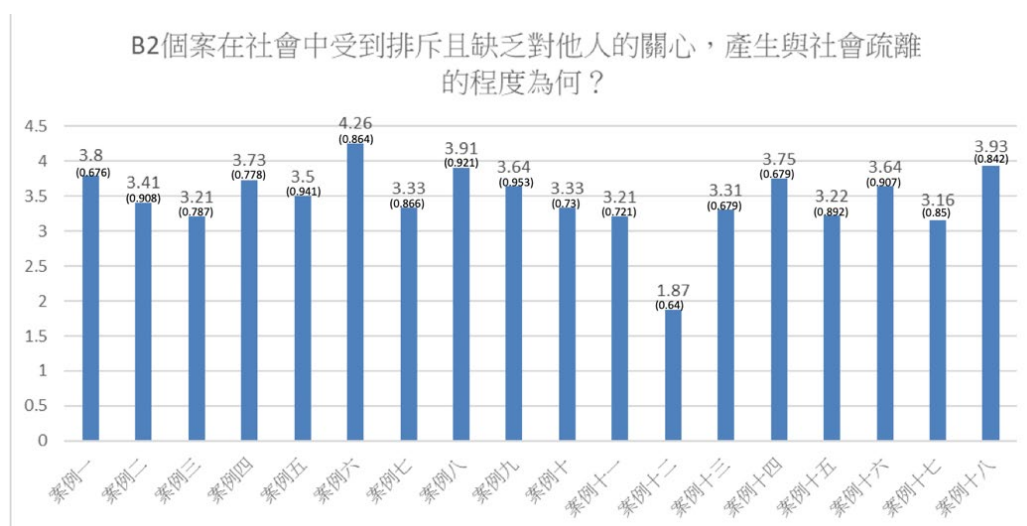


圖 25：受試者回答 B2「個案在社會中受到排斥且缺乏對他人的關心，產生與社會疏離的程度為何？」之平均分數

資料來源：本計畫。

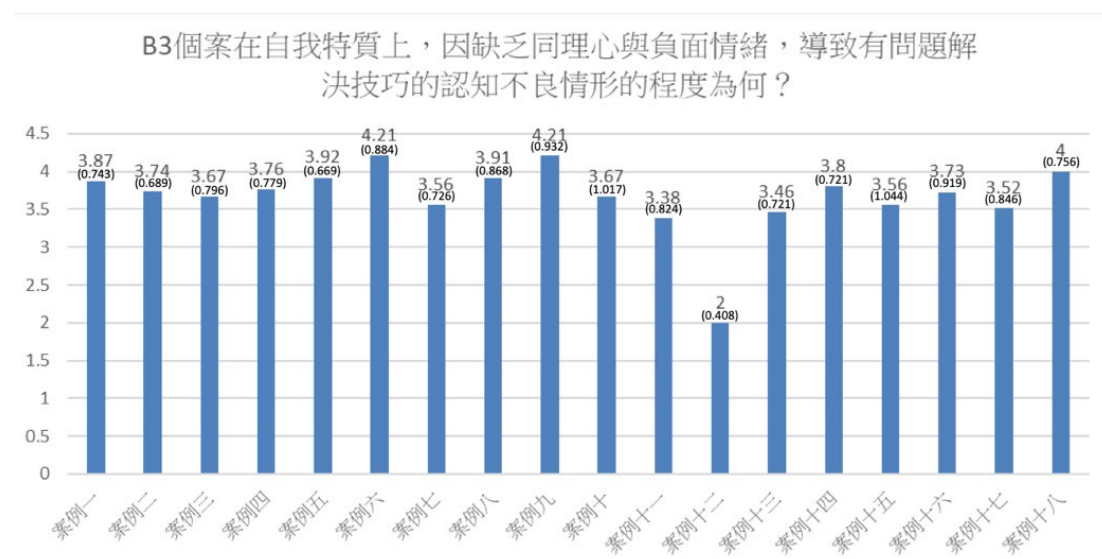


圖 26：受試者回答 B3「個案在自我特質上，因缺乏同理心與負面情緒，導致有問題解決技巧的認知不良情形的程度為何？」之平均分數
資料來源：本計畫。

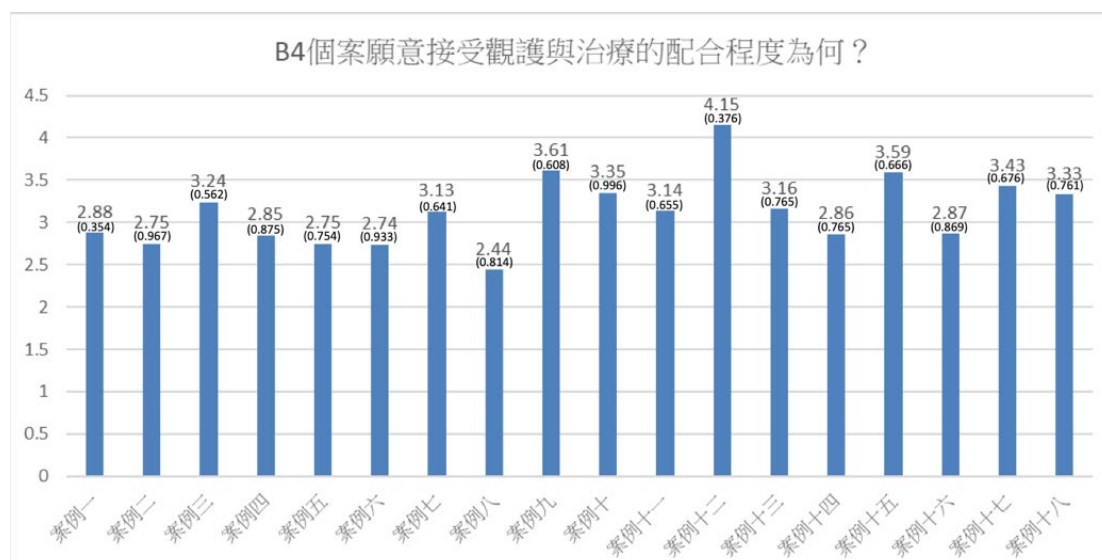


圖 27：受試者回答 B4「個案願意接受觀護與治療的配合程度為何？」之平均分數

資料來源：本計畫。

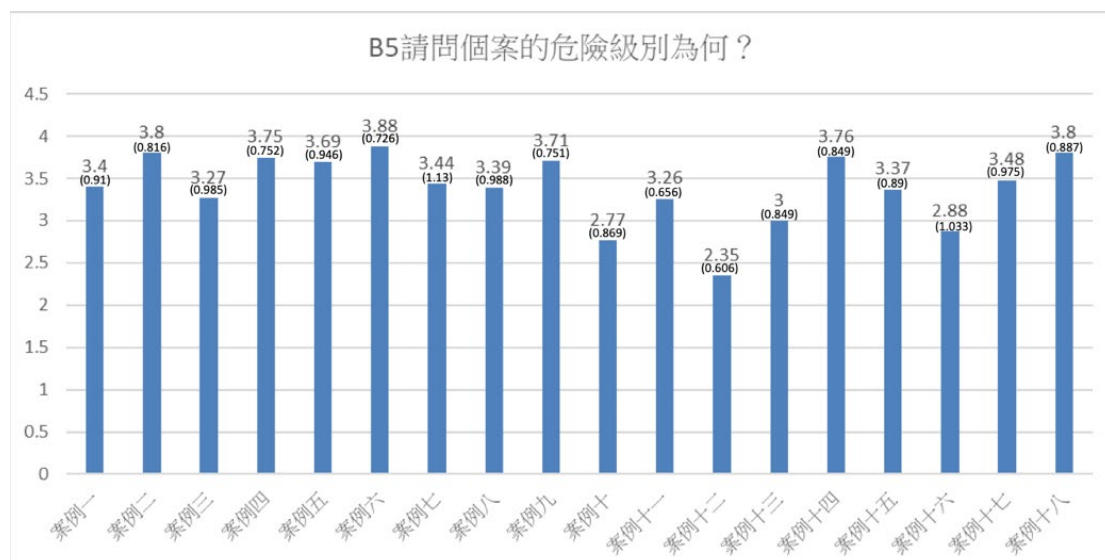


圖 28：受試者回答 B5「請問個案的危險級別為何？」之平均分數
資料來源：本計畫。

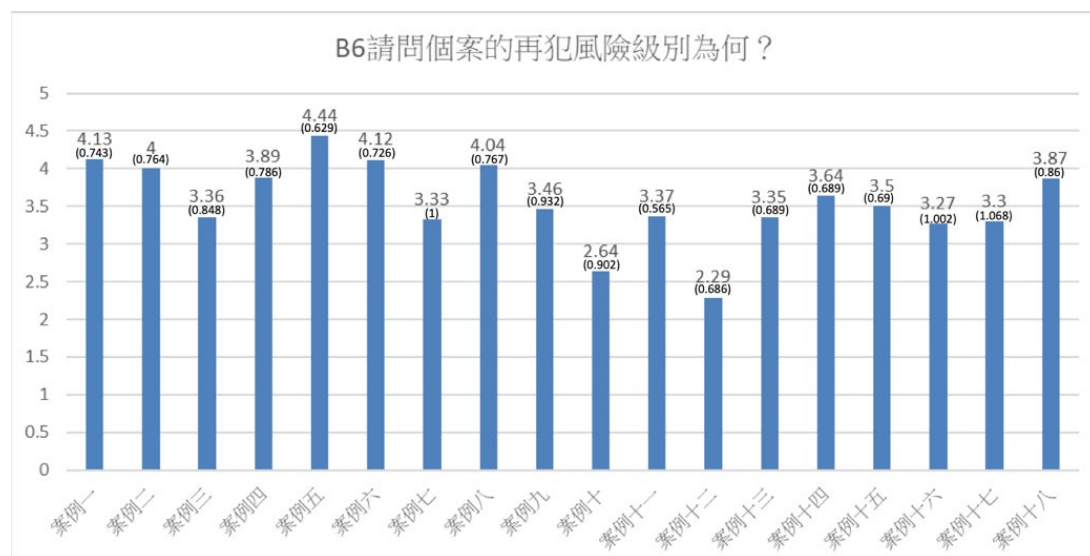


圖 29：受試者回答 B6「請問個案的再犯風險級別為何？」之平均分數
資料來源：本計畫。

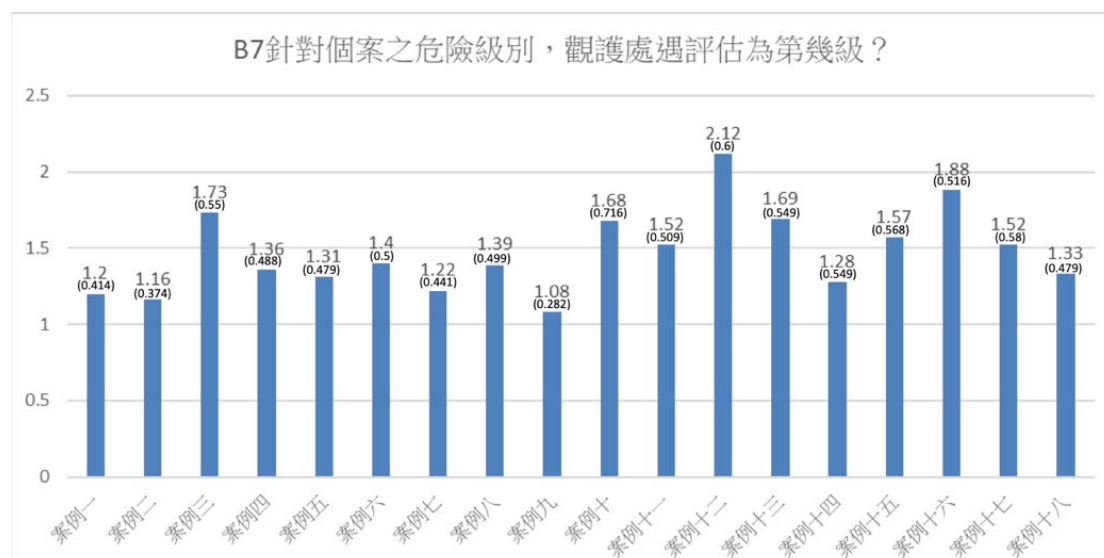


圖 30：受試者回答 B7「針對個案之危險級別，觀護處遇評估為第幾級？」之平均分數

資料來源：本計畫。

3、修改答案狀況

控制組之受試者在每一輪看完個案資料並送出決策答案後，本實驗會提供再次修改答案的機會，受試者可選擇是否要修改答案；實驗 Arm1 組在送出答案後，會提供 AI 系統的決策判斷供作參考，並再次詢問受試者是否要修改答案，系統的答案正確與否為隨機分配；而由於控制組及實驗 Arm1 組的施測時間較早，本計畫在實驗 Arm2 組施測前事先將前兩組受試者各題最多勾選的選項認定為「正確判斷」，並將最少人勾選的選項認定為「錯誤判斷」，整理成實驗 Arm2 組的決策參考，同樣詢問是否需要修改自己的答案。各組每一輪修改答案的比例整理於下表 41，而表 41 則進一步整理實驗 Arm1、Arm2 組在遇到正確或錯誤參考時修改答案的比例。

(1) 各組修改答案比例

在沒有其他外部資訊介入的情形下，控制組修改答案的比例明顯低於其他兩組，甚至有兩輪沒有任何受試者有修改答案。而比較另外兩組，我們發現實驗 Arm1 修改答案的比例高於實驗 Arm2 組，顯示演算法系統的介入較有可能影響受試者作決策判斷。

表 41：各組每輪修改答案比例之比較

實驗階段	控制組	Arm1 組	Arm2 組
第一輪修改答案比例	7.7%	31.3%	25.7%
第二輪修改答案比例	0%	38.7%	17.6%
第三輪修改答案比例	2.6%	22.6%	2.9%

實驗階段	控制組	Arm1 組	Arm2 組
第四輪修改答案比例	0%	26.7%	11.7%

資料來源：本計畫。

(2) 給予判斷正確／錯誤與修改比例之關聯

承上面所述，在實驗 Arm1 組及實驗 Arm2 組完成個案再犯處遇評估的第一次回答後，我們將分別提供 AI 決策系統以及其他同仁答案作為參考，詢問是否需要修改答案，然而本實驗將隨機提供正確或錯誤的決策參考，以此更細部地觀察受試者修改答案的情形，實驗 Arm1 組及 Arm2 組依決策參考正確／錯誤修改答案的比例呈現與下表 42。我們發現無論是 Arm1 組或是 Arm2 組，受試者收到錯誤的決策參考修正答案的比例普遍比收到正確的決策參考還來得高。

表 42：給予實驗組受試者正確或錯誤判斷後，修改答案比例之差異比較

實驗階段	正確／錯誤	Arm1 組	Arm2 組
第一輪案例	正確	28.6%	14.3%
	錯誤	32.0%	33.3%
第二輪案例	正確	40.0%	11.1%
	錯誤	38.1%	20.0%
第三輪案例	正確	20.0%	0%
	錯誤	23.8%	3.3%
第四輪案例	正確	16.7%	0%
	錯誤	33.3%	16.0%

資料來源：本計畫。

(五) 實驗後題目

1、 再次詢問對導入機器演算法之支持程度

在完成實驗後，本計畫再次詢問受試者對於再犯風險評估機器演算之人工智慧系統導入之支持程度。原始問卷設計使用李克特六等量表，並轉換為分數作計算，計算方式分別為：非常不支持（1 分）、不支持（2 分）、有點不支持（3 分）、有點支持（4 分）、支持（5 分）、非常支持（6 分），而在下表 43 中量表底下所呈現的數字為受試者填答之百分比。多數受試者的回答區間落在「有點支持」及「支持」，控制組與 Arm1 的平均分數分別為 4.39 分及 4.07 分，而 Arm2 組則低於 4 分，三者具顯著差異。

表 43：實驗後受試者對機器演算法導入之支持度調查

C1 首先謝謝您協助我們完成以上的實驗調查。本團隊想再次向您請教，若未來法務部要開發並導入新的「再犯風險評估機器演算之人工智慧系統」（以下簡稱為 Parole-ML-AI 1.0 ），請試著想像其可能發生的結果，請問您是否支持？									
組別	非常不支持	不支持	有點不支持	有點支持	支持	非常支持	平均分數	標準差	Oneway ANOVA F-value (p-value)
控制組	0	5.3	10.5	36.8	34.2	13.2	4.39	1.028	2.601* (0.079)
Arm1	0	6.7	13.3	50.0	26.7	3.3	4.07	0.907	
Arm2	2.9	8.6	20.0	40.0	25.7	2.9	3.86	1.089	

資料來源：本計畫。

備註1：六等量表底下所呈現之數字為各組受試者填答之百分比。

備註2：significance level, *, $p < 0.1$, **, $p < 0.05$, ***, $p < 0.01$ 。

2、受試者對機器演算法導入工作之接受度

C2.1 至 C2.8 題之設計靈感源自於「科技接受模式」（Technology Acceptance Mode），詢問受試者對機器演算法導入工作之認知，各組每題之平均分數皆未超過 4 分，各題說明如下（如表 44）：

- (1) C2.1題對是否需為「系統的錯誤」負責方面，多數受試者持不同意立場，三組平均分數皆不高於3分（控制組：2.18分、實驗Arm1組：2.57分、實驗Arm2組：2.77分），三者差異不顯著。
- (2) C2.2題對「系統會否限縮裁量權力」方面，實驗Arm2組與其他兩組相比較持同意態度，三者呈現顯著差異。
- (3) C2.3題對「系統是否讓長官更容易監督工作情況」，實驗Arm2組與其他兩組相比較持同意態度，三者呈現顯著差異。
- (4) C2.4題對「系統會否需要承擔更多責任」，三組回答情形較分散，平均分數皆落在4分以下，三組差異不顯著。
- (5) C2.5題對「系統能否使觀護人更容易回應社會民眾要求」方面，三組的平均分數皆低於4分，三組呈顯著差異。
- (6) C2.6題對「系統能否協助工作業務」，三組平均分數皆不高於4分，其中以實驗Arm1組最低（平均：2.97分），三組呈顯著差異。

- (7) C2.7題對「能否相信系統給的評估建議」方面，三組平均分數皆不高於4分，答題選項最多者皆為「有點不同意」，其中以實驗Arm1組最低（平均：3.17分），三組呈顯著差異。
- (8) C2.8題對「是否期待使用系統」，三組平均分數皆不高於4分，其中控制組最多人填答選項為「有點不同意」，而實驗Arm1、Arm2組最多人填答選項落在「有點同意」之區間，三組呈顯著差異。

表 44：受試者對機器演算法導入工作之接受度調查

組別	非常不同意	不同意	有點不同意	有點同意	同意	非常同意	平均分數	標準差	Oneway ANOVA F-value (p-value)
C2.1 請該系統所造成的錯誤需由我負責。									
控制組	31.6	39.5	15.8	5.3	7.9	0	2.18	1.182	1.699 (0.188)
Arm1	23.3	26.7	30.0	10.0	10.0	0	2.57	1.251	
Arm2	31.4	22.9	11.4	8.6	22.9	2.9	2.77	1.664	
C2.2 該系統將限縮我的裁量權力。									
控制組	7.9	21.1	18.4	28.9	21.1	2.6	3.42	1.328	2.575* (0.081)
Arm1	6.7	13.3	16.7	36.7	26.7	0	3.63	1.217	
Arm2	0	14.3	17.1	25.7	31.4	11.4	4.09	1.245	
C2.3 該系統將讓長官更容易監督我的工作情況。									
控制組	0	23.7	15.8	42.1	15.8	2.6	3.58	1.106	2.660* (0.075)
Arm1	3.3	13.3	20.0	40.0	23.3	0	3.67	1.093	
Arm2	0	11.4	11.4	37.1	31.4	8.6	4.14	1.115	
C2.4 該系統將讓我承擔更多責任。									
控制組	2.6	18.4	28.9	18.4	26.3	5.3	3.63	1.282	0.648 (0.525)
Arm1	3.3	13.3	13.3	50.0	20.0	0	3.70	1.055	
Arm2	0	14.3	25.7	22.9	25.7	11.4	3.94	1.259	
C2.5 該系統將讓我更容易回應社會民眾的要求。									
控制組	0	0	18.4	34.2	28.9	0	3.74	1.083	2.590* (0.080)
Arm1	0	13.3	26.7	40.0	20.0	0	3.67	0.959	
Arm2	8.6	28.6	20.0	25.7	14.3	2.9	3.17	1.317	
C2.6 該系統並不能協助我的工作業務。									
控制組	2.6	23.7	44.7	18.4	10.5	0	3.11	0.981	5.565**

Arm1	3.3	26.7	50.0	10.0	10.0	0	2.97	0.964	(0.005)
Arm2	0	17.1	28.6	20.0	28.6	5.7	3.77	1.215	
C2.7 我無法相信該系統給予的風險評估建議。									
控制組	0	23.7	34.2	31.6	10.5	0	3.29	0.956	2.660* (0.075)
Arm1	3.3	23.3	43.3	16.7	10.0	3.3	3.17	1.117	
Arm2	0	14.3	31.4	28.6	17.1	8.6	3.74	1.172	
C2.8 我會很期待使用這個系統。									
控制組	0	5.3	18.4	52.6	21.1	2.6	3.97	0.854	3.296** (0.041)
Arm1	3.3	10.0	20.0	46.7	16.7	3.3	3.73	1.081	
Arm2	0	31.4	14.3	40.0	14.3	0	3.37	1.087	

資料來源：本計畫。

備註1：六等量表底下所呈現之數字為各組受試者填答之百分比。

備註2：significance level, *, $p < 0.1$, **, $p < 0.05$, ***, $p < 0.01$ 。

3、受試者創新意願調查

題組 C3.1 至 C3.4 為創新意願之調查，欲瞭解受試者對於創新的認知以及願意嘗試改變的程度為何。就答題狀況來看，各組在每題的平均分數皆高於 4 分，且每題各組並無顯著差異，說明如下（如表 45）：

- (1) C3.1 題請受試者自評是否在團體中為比較快接受新想法的人，三組的回答集中在「有點同意」及「同意」區間，平均分數皆超過 4 分，且三組差異不顯著。
- (2) C3.2 題請受試者自評在接受新想法時的謹慎程度，三組的回答集中在「有點同意」及「同意」區間，平均分數皆超過 4 分，且三組差異不顯著。
- (3) C3.3 題請受試者自評在有可能失敗的情況下願意冒險的意願，三組的回答集中在「有點同意」及「同意」區間，平均分數皆超過 4 分，且三組差異不顯著。
- (4) C3.4 題請受試者自評在有風險的情況下，願意嘗試新科技的意願，三組的回答集中在「有點同意」及「同意」區間，平均分數皆超過 4 分，且三組差異不顯著。

表 45：受試者創新意願之調查

組別	非常不同意	不同意	有點不同意	有點同意	同意	非常同意	平均分數	標準差	Oneway ANOVA F-value (p-value)
C3.1 我發現我通常是團體中比較快接受新想法的人。									
控制組	0	0	13.2	44.7	36.8	5.3	4.34	0.781	0.059 (0.943)
Arm1	0	0	16.7	33.3	43.3	6.7	4.40	0.855	
Arm2	0	2.9	8.6	40.0	42.9	5.7	4.40	0.847	
C3.2 在接受新想法時，我通常會很謹慎。									
控制組	0	0	15.8	42.1	42.1	0	4.26	0.724	1.086 (0.342)
Arm1	0	3.3	6.7	53.3	3.3	3.3	4.27	0.785	
Arm2	0	0	2.9	51.4	40.0	5.7	4.49	0.658	
C3.3 即便有失敗的可能性，我還是願意去嘗試過去工作上沒做過的事。									
控制組	0	0	5.3	42.1	50.0	2.6	4.50	0.647	0.327 (0.722)
Arm1	0	0	3.3	46.7	43.3	6.7	4.53	0.681	
Arm2	0	0	8.6	51.4	31.4	8.6	4.40	0.775	
C3.4 即便可能會有風險，我還是願意在工作業務上嘗試新科技。									
控制組	0	0	2.6	47.4	44.7	5.3	4.53	0.647	1.454 (0.238)
Arm1	3.3	0	10.0	46.7	33.3	6.7	4.27	0.980	
Arm2	0	2.9	11.4	48.6	34.3	2.9	4.23	0.808	

資料來源：本計畫。

備註1：六等量表底下所呈現之數字為各組受試者填答之百分比。

備註2：significance level, *, $p < 0.1$, **, $p < 0.05$, ***, $p < 0.01$ 。

4、受試者對工作流程自動化之需求調查

而究竟機器演算法導入工作的哪一面向對觀護人之工作是有幫助的，同樣為本實驗欲探究之處。在題組 C4.1 至 C4.5 題中列出觀護人主要的工作面向，並依工作本身的複雜度、裁量權力以及涉及人民權利由小至大排序，詢問受試者對各工作流程自動化的需求，分別有「例行行政工作與文書撰寫」、「保護管束案件管理」、「評估再犯風險」、「提供處遇措施建議」以及「輔導監督」。此題組調查結果顯示受試者對不同工作面向自動化需求的迫切程度，隨著複雜以及涉及人民權利程度的提升，受試者回答之平均分數呈逐步遞減的狀態。就「例行行政工作與文書撰寫」、「保護管束案件管理」、「評估再犯風險」方面，各組的平均分數皆高於 4 分，

且回答區間多集中於「有點需要」及「需要」之區間，然而就「提供處遇措施建議」、「輔導監督」方面，三組的平均分數相較前者較低，且回答情形較為分散，顯示複雜度高、裁量權力、涉及人民權利較大之工作，受試者對流程自動化之需求較低，相較於導入機器演算系統協助，受試者較傾向仍由人來作決策。各題答題狀況說明如下（如表 46）：

- (1) C4.1：「例行行政工作與文書撰寫」方面，三組平均分數皆高於 4.6 分，且彼此無顯著差異。
- (2) C4.2：「保護管束案件管理」方面，三組平均分數皆高於 4.2 分，且彼此無顯著差異。
- (3) C4.3：「評估再犯風險」方面，三組平均分數皆高於 4 分，且彼此無顯著差異。
- (4) C4.4：「提供處遇措施建議」方面，實驗 Arm2 組平均分數最低（平均：3.80 分），三組並無顯著差異。
- (5) C4.5：「輔導監督」方面，實驗 Arm1、Arm2 組平均分數皆落在 4 分以下（實驗 Arm1 組：3.97 分；實驗 Arm2 組：3.74 分），三組並無顯著差異。

表 46：受試者對工作流程自動化之需求調查

組別	非常 不需要	不 需要	有點 不需要	有點 需要	需 要	非常 需要	平均 分數	標準 差	Oneway ANOVA F-value (p-value)
C4.1 例行行政工作與文書撰寫。									
控制組	0	15.8	0	21.1	34.2	28.9	4.61	1.346	0.302 (0.740)
Arm1	3.3	3.3	3.3	26.7	3.3	30.0	4.74	1.230	
Arm2	0	8.6	0	17.1	48.6	25.7	4.83	1.098	
C4.2 保護管束案件管理。									
控制組	0	2.6	5.3	39.5	28.9	23.7	4.66	0.994	1.736 (0.181)
Arm1	6.7	3.3	6.7	40.0	33.3	10.0	4.20	1.243	
Arm2	2.9	5.7	8.6	28.6	51.4	2.9	4.29	1.073	
C4.3 評估再犯風險。									
控制組	0	2.6	13.2	36.8	36.8	10.5	4.39	0.946	1.353 (0.263)
Arm1	0	3.3	10.0	46.7	30.0	10.0	4.33	0.922	
Arm2	5.7	14.3	5.7	25.7	45.7	2.9	4.00	1.328	

C4.4 提供處遇措施建議。									
控制組	0	10.5	15.8	28.9	39.5	5.3	4.13	1.095	1.094 (0.339)
Arm1	0	6.7	16.7	36.7	33.3	6.7	4.17	1.020	
Arm2	5.7	14.3	11.4	31.4	37.1	0	3.80	1.256	
C4.5 輔導監督。									
控制組	0	7.9	23.7	31.6	31.6	5.3	4.03	1.052	0.590 (0.556)
Arm1	3.3	10.0	10.0	43.3	30.0	3.3	3.97	1.129	
Arm2	5.7	17.1	8.6	37.1	28.6	2.9	3.74	1.291	

資料來源：本計畫。

備註1：六等量表底下所呈現之數字為各組受試者填答之百分比。

備註2：significance level, *, $p < 0.1$, **, $p < 0.05$, ***, $p < 0.01$ 。

5、受試者公共服務動機調查

C5.1 至 C5.5 題為本實驗最後之題組，主要欲詢問受試者之公共服務動機，研究結果（如表 47）顯示，C5.1 至 C5.3 題「為大眾服務是很重要的事情」、「為社會做出貢獻比得到個人成就更有意義」、「社會上人與人之間是互相需要的」之平均分數皆高於 4 分，顯示受試者大多持同意態度；而 C5.3「社會上人與人之間是互相需要的」三組平均分數甚至高過 5 分（控制組：5.05 分；實驗 Arm1 組：5.00 分；實驗 Arm2 組：5.26 分）。然而 C5.4「為社會公義而犧牲自己的權益沒有關係」以及 C5.5「就算會惹禍上身，路見不平拔刀相助他人是必要的」兩題之平均分數皆低於 4 分。本題組之五題皆無顯著差異。

表 47：受試者公共服務動機之調查

組別	非常不同意	不同意	有點不同意	有點同意	同意	非常同意	平均分數	標準差	Oneway ANOVA F-value (p-value)
C5.1 為大眾服務是很重要的事情。									
控制組	0	2.6	7.9	23.7	52.6	13.2	4.66	0.909	0.480 (0.620)
Arm1	0	0	0	34.5	44.8	20.7	4.86	0.743	
Arm2	0	2.9	8.6	14.3	54.3	20.0	4.80	0.964	
C5.2 為社會做出貢獻比得到個人成就更有意義。									
控制組	0	0	15.8	23.7	44.7	15.8	4.61	0.946	0.150 (0.861)
Arm1	0	3.4	0	37.9	37.9	20.7	4.72	0.922	
Arm2	0	8.6	5.7	14.3	48.6	22.9	4.71	1.152	

C5.3 社會上人與人之間是互相需要的。									
控制組	0	0	0	13.2	68.4	18.4	5.05	0.567	1.765 (0.176)
Arm1	0	0	0	24.1	51.7	24.1	5.00	0.707	
Arm2	0	0	0	2.9	68.6	28.6	5.26	0.505	
C5.4 為社會公益而犧牲自己的權益沒有關係。									
控制組	5.3	21.1	21.1	31.6	15.8	5.3	3.47	1.289	0.438 (0.646)
Arm1	0	6.9	34.5	44.8	10.3	3.4	3.69	0.891	
Arm2	5.7	14.3	20.0	28.6	25.7	5.7	3.71	1.319	
C5.5 就算會惹禍上身，路見不平拔刀相助他人是必要的。									
控制組	7.9	26.3	28.9	18.4	15.8	2.6	3.16	1.285	0.879 (0.418)
Arm1	0	13.8	37.9	41.4	3.4	3.4	3.45	0.910	
Arm2	2.9	17.1	34.3	22.9	20.0	2.9	3.49	1.173	

資料來源：本計畫。

備註1：六等量表底下所呈現之數字為各組受試者填答之百分比。

備註2：significance level, *, $p < 0.1$, **, $p < 0.05$, ***, $p < 0.01$ 。

二、 小結

本實驗以現職擔任我國法務部保護司觀護人為實驗對象，觀察其針對個案作再犯處遇評估時受到外部資訊介入的影響，瞭解 AI 系統及其他同仁的決策參考是否會影響受試者個人的判斷，進而修改自己的答案。此外，本實驗也透過問卷形式調查觀護人當前導入「檢察機關案件管理系統」（以下簡稱案管系統）於工作之現況以及對於機器演算決策系統導入再犯處遇評估工作的期望及接受度，並調查觀護人本身對於創新意願、工作業務流程自動化、以及個人的公共服務動機等不同面向，以此更加瞭解機器決策系統演算應用於公部門可能帶來的好處以及挑戰有何。針對實驗的結果，我們歸納出以下結論：

(一) 案管系統導入於觀護人工作之現況以及機器演算法導入之期望

當今觀護人執行中作業務時使用案管系統的頻率相當高，多數也肯定該系統的易用程度。在承受龐大的工作負擔以及壓力之下，大部分的受試者認為機器演算系統的導入有助於再犯評估之工作。然而普遍來說，受試者也同意需要額外花費成本才能與機器演算系統順利合作。

而就經調查發現，受試者普遍對機器演算系統能夠協助工作業務表示肯定，然而對於是否該為系統造成的錯誤負責，受試者多半抱持否定的態度。

(二) AI 系統與其他同仁對再犯處遇判斷之影響

1、三組修改答案比例之比較

經實驗發現，在沒有其他外部資訊介入的情形下，控制組修改答案的比例遠低於其他兩組的受試者（實驗結果詳見表 41：各組每輪修改答案比例之比較），在有外部決策資訊的參考下，實驗 Arm1 及 Arm2 組修改的比例較高。而 Arm1 組在受到 AI 決策參考的影響之下更改自己答案的比例又比 Arm2 組來得高，顯示若外部參考資訊來自於其他同仁，受試者更傾向相信自己的決策，然而若是參照 AI 的決策參考，則受試者較容易受到影響並更改答案。

2、外部資訊正確與否對再犯處遇判斷之影響

而若以外部資訊所提供之答案為正確與否來看，我們發現若提供之資訊為錯誤答案，無論是 Arm1 或是 Arm2 組的受試者皆較有可能更改自己的答案。

(三) 工作流程自動化之需求

在題組 C4.1 至 C4.5 題中，本實驗列出了觀護人主要的工作業務，分別有「例行行政工作與文書撰寫」、「保護管束案件管理」、「評估再犯風險」以及「提供處遇措施建議」，並詢問受試者對各項工作自動化的需求，並依工作的複雜性、裁量權大小以及涉及人民權利程度由小至大排序，請受試者作勾選。我們發現簡單卻繁冗、重複性高的工作對於機器演算導入之需求較為迫切，且其所能夠帶來的助益也較大，而相比之下，在較為複雜、同仁裁量權力且涉及人民權利較大之工作上，觀護人的需求較低。依本題組之調查結果來看，未來若欲導入機器演算機器演算於公部門之工作業務中，應可優先考慮應用於日常較為簡單繁瑣且較不涉及專業判斷之例行工作，就此減輕同仁的壓力與負擔，並能更集中精神於複雜性較高、關乎人民權利較大而須更多專業判斷介入之工作。若該系統能有效改善工作狀況，亦有可能提升同仁對該系統的信任度及好感度，更有效促進人機協作，如此也能提升同仁對導入新興科技於工作業務的接受度。

(四) 創新意願與公共服務動機

在創新意願之相關問題方面，受試者普遍的回答落在「有點同意」及「同意」之區間，顯示大多對於新想法或新興科技導入於工作是抱持擁抱與接納的態度。就公共服務動機方面，本實驗亦發現受試者普遍擁有高度公共服務、維護公眾利益之動機，顯示對於導入科技以提升公共服務效能方面，觀護人多半是持肯定的態度。

綜合本實驗結果之分析來看，觀護人對於機器演算能夠在工作業務上提供協助普遍表示肯定與支持，在具有高創新意願及公共服務動機之

下，對於新興知識觀念及科技的導入也多半持正面的態度。就應用的面向上，多數對該系統導入複雜性較低且繁瑣、裁量權力以及關乎人民權利較小的工作上之需求程度較高，相比之下，涉及裁量權較大與涉及權利較大的工作仍較傾向以人作決策。本實驗要求受試者作的「再犯處遇評估」便屬複雜性較高之工作，在有外部資訊介入的情形下，受試者修改答案的比例較高，然而我們發現錯誤的參考資訊較有可能影響決策的更改，而多數卻不認為機器所造成的錯誤需由人來負責。就此方面，在未來機器演算系統導入公共服務時，應重視當系統決策與人為判斷相互衝突時，同仁如何能秉持自身專業作出正確的決策，不受機器判斷錯誤影響；此外亦須注重系統輔助決策時的權責劃分與釐清之問題，提升公部門人員對使用該類系統應肩負之責任的認知。

第四節 落差分析

落差分析（Gap Analysis），又名差距分析，常見於商業領域，Oliver（1977）認為消費者的消費滿意度將影響日後的消費意願，而檢驗此滿意度的方式為測量民眾消費前的期望（expectancy）與消費後實際感知的差距，是為期望失驗理論（Expectancy Disconfirmation Theory）。Parasuraman 等人（1985）將客戶與股東的原先期望與日後的成果與獲得效益作比較，找出理想狀態以及實際狀態間之差距，探討差距的由來，並從中思考改善之策略。Burley（1988）在探討自然資產保護時提出了「保存落差」（conservation gap）的概念，用以測量自然保護區的生物多樣性以及物種存續（Scott et al., 1987, 1993）。落差分析可用於商業活動中多種指標衡量，如：顧客滿意度、營業收入、生產力以及供應鏈成本等不同面向。在資訊科技領域中，專案管理者也經常利用落差分析作為優化運作流程以及開發的工具，就軟體開發上可以藉此找出哪些服務以及應用程式在使用率上未盡人意，哪些程式值得再精進。圖 31 呈現一般落差分析的模型與操作流程。

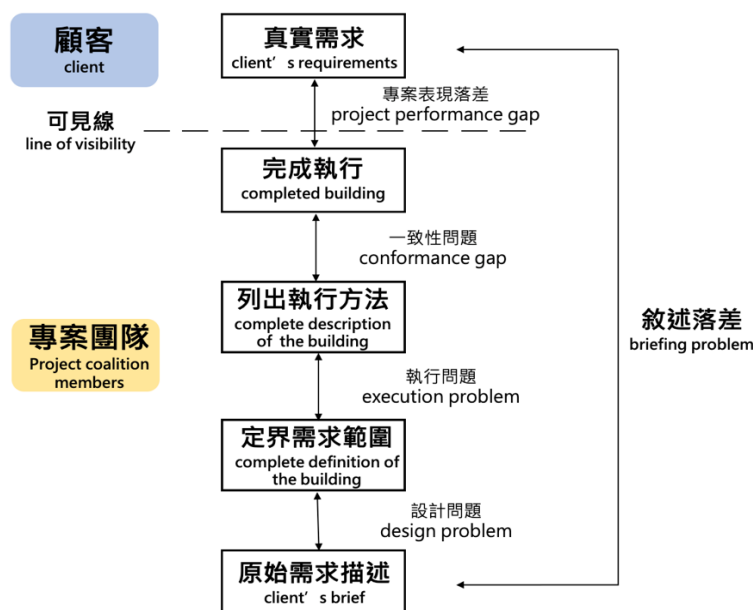


圖31：落差分析模型

資料來源：翻譯修改自 Winch、Usmani & Edkins（1998）。

一、法規調適

由本計畫對於各國以及國際組織所提出之各項文件可知，各國對於AI發展之管制，早已發展多年，且循序漸進依研究報告、倫理作業指引、法律規範之步驟，逐步提升法規範之管制強度；並依據組織規劃定位、規制方式制度、倫理價值與政策原則、風險評估分級，輔以各國原有包含監理機關與法規的個人資料保護機制，循序漸進建構AI監理架構。

就個資保護機制而言，我國最棘手之問題在於現階段我國沒有個資保護之主管機關，也因此缺乏對於個人資料保護的獨立第三方監理單位，只能由各目的事業主管機關依個資法之規範進行認定監管；當公部門運用機器人做行政決定時，因可能產生類似行政處分之效果，因此應透過要求資料控制者之主動揭露，給予受影響之資料主體相關救濟途徑，前階段之揭露越清楚，即能給予當事人越多救濟可能性。另一方面，亦應留意資料選擇上是否夠全面且無差別待遇，自動化決策前階段的資料選擇以及演算法設計多會受到人為影響，因程式設計者若對數據資料帶有偏頗的想法跟預設，必將連帶影響程式設計之內容，若透過將前階段清楚揭露，就能盡量避免偏見產生。對此歐盟即要求前階段需建立隱私設計，可透過脈絡式地檢視前階段資料蒐集之方式是否足夠全面或是否帶有偏見，進而檢視設計所編寫之演算法、整體過程與參數。

至於針對AI科技之具體規範，於我國更是缺乏，除科技部與其轄下四個AI創新研究中心共同於2019年9月發佈「AI科研發展指引」係針對研發階段所給予之一般性指引外，充其量只有民間同業團體證券投顧

同業公會自行訂定屬內規性質的「自動化投顧作業要點」較為具體，其內容不論是對演算法監管、瞭解客戶、適合性原則或說明義務等措施，因此雖屬民間團體之內規，但已有法規之雛型，內容與美國官方金融監管局（FINRA）數位投資工具研究報告建議之監管方式亦十分相近，該報告針對市場上已提供投資建議之業者，除投資人自行應注意之事項外，關於業者運用理財機器人時應如何符合投資顧問法下對投資人之基本義務，包含投資顧問應善盡受任人義務（fiduciary duty）及應以最大之善意（utmost good faith）為投資人提供個人化、專業化之投資服務，並應保持客觀與中立、避免利益衝突等事項，提出業務執行建議，供業者據以實行或調整作業之規範，並著重於投資工具嵌入演算法之管理與監督、投資組合與利益衝突之管理與監督、識別投資人屬性、投資組合再平衡及金融專業人士之訓練等部分（王偉霖，2019）。然而，自動化投顧作業要點為證券投顧同業公會所自行制定，僅得認為係公會所擬定之自律規範而已；縱經送請金管會同意訂定公告責成同業公會會員業者遵守，擬開辦機器人理財或智能理財服務之產業，除申請兼營證券投資顧問業務，均須遵循該作業要點始得辦理，但其性質並非經由法律授權之法規命令，充其量僅為主管機關對於申請案件審查之行政機關內部行政規則，但因主管機關公告要求遵守而發生外部效力。若業者未按規定進行，難以認定違法，亦無法予以裁罰。主管機關雖有介入權限，但其執行仍大量委由自律組織職掌，使得本作業要點在規範執行層面恐較缺乏規制效力。

也因此，相較於國外 AI 法制架構的發展，我國對於 AI 之規範體系，尚在於極為前階段的研究與凝聚共識，若依據外國法制規劃的進程，近期可能先以指定或建立專責機關組織，並訂定相關指引為較可行之目標。參酌文獻以及專家之建議，依據 AI 所應重視的相關倫理規範與原則，大致上可分為三個規劃層級：

第一層為告知程序與揭露義務的完備，主要屬於一般個人資料保護的規範方式，主要考量為是否需要告知資料主體其個人資料將運用於機器學習、現階段是否具此項告知程序、告知後相關資料與參數使用之方式應如何揭露、又應揭露至何種程度始能於保護資料主體、且個人權利與業者營業秘密保護間應如何達到平衡，都是考量因素。

第二層為透明性與可解釋性的考量，各該數據資料透過演算法與機器學習產出結果之理由為何？又演算法及機器學習是如何運用或篩選各項資料及參數？若採用之數據或參數不同，將極可能影響到機器人學習結果以及影響機器判斷做出之結論。

第三層為當這些參數及過程均充分揭露後，何人有能力可以判斷其公平性或透明度是否足夠？進而影響到機器人所運用之數據資料、參數，以及演算法與機器學習之內容應如何管制或監督，以及後續之救濟可能

性等問題，以建立整體監管的組織與法規架構。

以上三個層級，若以我國法規現況而言，可說第一層級的規範都還不完備。因此，當公部門欲導入演算法、機器學習、機器人時，首應考量政府所運用之演算法能否建立一個透明公開檢驗的作法，意即該演算法是否能透明公開、是否及如何對民眾進行揭露？其次在於應如何驗證演算法之安全性跟可靠性？私部門所運用之演算法因涉及公司內部之商業或營業秘密等問題，歐盟認為未必均能要求公司公開，除非有相關法令要求；但公部門受到政府資訊公開法之要求，其問題應在於執行層面是否能確認已充分揭露，又公開後該如何確認政府之演算法是安全的？應由誰去檢驗該演算法？最後則應考量政府自己是否有能力去執行演算法所衍生的後續效應，尤其是責任的界定，如金融產業所使用之理財機器人產生瑕疵時，金管會是否有能力審核或檢驗機器人背後所運用之演算法並為其審核結果而負責；另外對於 AI 機器人導入公共服務後，理論上機器人作為公務員手足延伸之助手或獨立執行公共服務，公部門是否能為 AI 所做出之判斷完全負責？均是重要的思考問題。

至於就監理管制層面而言，參酌歐盟 AI 規則草案之規範內容，針對資料治理有清楚的法規架構，AI 之基礎在於前階段大量資料、數據及參數之累積，唯有透過完整之資料治理規範，才有可能妥適處理後續 AI 衍生之問題並進行課責；另就管制之主管機關而言，宜指定一統籌部會進行對於演算法之能力培養，並針對基礎設施之應用與計算等制訂共通原則、審查機制等項目，至於細部之發展、管理、維護等宜依各部會之業務別區分之。理由在於，欲將演算法導入各該行業，勢必應對該行業具相當程度瞭解，事實上亦難以期待不熟悉特定產業之設計者，能進行完全符合產業需求的演算法設計，因而需目的事業主管機關的負責參與；又為避免各行業間對於資料治理與管制演算法設計之標準不一，不利於公部門進行後續管制與究責，因此宜先由單一主管機關制訂共通原則與審查機制，供各業務機關遵循辦理。但依據本計畫所辦理之專家會議與會專家均表示，目前公部門機關已過多，實難期待新增一專責機關可有效率負責相關業務，因此就現有機關擇一負責，可能為較可行之方式。

二、 資源配置

(一) 中央各部會應用演算法第一階段調查結果詮釋

依據前揭調查結果，依序回答第一階段調查規劃回答之研究問題如下：

1. 中央部會應用機器演算法（AI）創新公共服務可行嗎？可能優先推動的項目為何？

依據第一階段調查的結果顯示，目前有八個部會已導入或正規劃導入 AI 系統輔助業務，顯見應用機器演算法 (AI) 創新公共服務是可行的，且最多部會選擇優先推動 AI 系統輔助業務的類型為服務對象溝通工作，其次為操作機器的自動化和辦公室行政文書工作。

2. 中央部會應用機器演算法 (AI) 創新公共服務可能遭遇之制度限制為何？

調查結果顯示，中央部會未導入 AI 系統輔助業務的主要理由依序為缺乏預算、缺乏適用技術、缺乏專業人力、非機關權責業務、缺乏配套法令、缺乏數位資料、非機關權責業務、機關首長未指示推動、缺乏成本效益。其他理由為無需求或尚在研議中。另外，根據已導入和規劃導入 AI 系統輔助業務的部會多表示無法規調適的需要，或僅有在行政規則層次的法規調整需求，因此本計畫團隊進一步推論，各部會應用 AI 創新公共服務所遭遇之制度限制似乎極為有限，各部會之所以未能應用 AI 系統創新公共服務，其主要遭遇之限制主要來自資源面（如預算、技術和人力）而非制度面。

3. 中央部會應用機器演算法 (AI) 創新公共服務所需的資源條件為何？包括預算、人力、資料需求等。

調查結果顯示，編列採購和發展 AI 系統所需預算，是已導入和規劃導入 AI 系統之部會首要條件，其次是至少有 1 名負責規劃、採購或管理 AI 系統專案的資訊人力，數位資料需求主要是本機關資料（含開放和內部資料）。這個結果也顯示，各部會發展的 AI 系統仍以滿足該部會業務所需為主，尚未有發展跨部會 AI 系統的情況。

(二)中央各部會應用演算法第二階段調查結果詮釋

本調查的第四個問題為「中央部會公務人員對於應用機器演算法 (AI) 創新公共服務的態度、預期效能以及對可能遭遇問題的解決態度為何？」從前述的分析結果可知，文官對於使用 AI 輔助業務之事具有複雜的態度，一方面預期使用 AI 會為自己的工作帶來正向的成效，另一方面又預期使用 AI 會加重自己需承擔的責任以及限制自己在工作上的權力。分析結果也顯示資料準備度是讓 AI 系統為行政業務帶來正向成效的關鍵，而行政機關愈能夠在組織管理系統展現對員工的支持，將能因此降低員工對數位轉型創新將加重個人工作責任的疑慮。此外，文官應加強自我的心理韌性，勇於面對變革和精進自我能力，方能因應政府數位轉型帶來的工作變革，但同時也需務實看待行政機關導入 AI 系統產生的業務變革，才不致於過度樂觀而忽略伴隨變革而來的衝擊。

再者，本調查設計七道和人才職能發展有關的題項（編號 2-41~2-47），以瞭解受訪者對於行政機關運用 AI 系統所需人才職能的看法。調

查結果發現，受訪者認為運用 AI 輔助業務最需具備的能力依序是：解讀資訊能力、資安與倫理、大數據分析能力、跨域合作能力、資源分配能力、創新服務能力和維運 AI 能力。惟比較受訪者認知人才職能應具備程度和自評實際具備程度的落差後發現，落差較大的職能項目依序為：維運 AI 能力、大數據分析能力、解讀資訊能力、資源分配能力、創新服務能力、資安與倫理道德以及跨域合作能力。

第二階段調查於問卷最後設計一道開放題，題目內容為「2-62.最後，若你對目前政府推動數位化政策有任何想法與建議，歡迎於下面空白欄留言。」在 126 位完整填答問卷的受訪者中，有 40 位受訪者回答該題，進一步分析渠等留言內容後，歸納出行政運作面向、服務設計面向、人才職能面向和其他面向等四大議題面向共 56 個議題，用以補充前述的調查分析結果，各面向議題內容及次數分配表如表 48 所示，茲就相關內容說明如下。

表48：各面向議題內容及次數彙整表

議題編號	議題面向	出現次數	百分比 (次數/總次數) %
1	行政運作面向	29	51.79
2	服務設計面向	14	25.00
3	人才職能面向	12	21.43
4	其他	1	1.79
總計		56	100

資料來源：本計畫。

1. 行政運作面

受訪者意見主要聚焦在：(1)漸進發展、(2)長期投入、(3)流程簡化、(4)成果驗證、(5)跨域協調、(6)資源落差、(7)階層落差、(8)課責問題、(9)專業決策、(10) 資訊事權統一、(11)共通性需求等十一個議題(如表 49)。受訪者表示政府數位化政策不可躁進貪功，而需採取漸進且長期投入的路徑發展，發展過程需簡化行政流程、重視成果驗證、跨域協調以及跨機關資源落差問題，同時留意跨部門、跨階層、跨單位、跨專業之間溝通的落差，允許創新過程的失敗和人員的免責，重視專業導向的決策，強化資訊業務的事權統一，並以統一開發方式滿足資訊系統的共通性需求。

表49：行政運作面向議題內容及次數彙整表

議題編號	議題內容	出現次數	百分比 (次數/總次數) %
1	漸進發展	3	10.34

議題編號	議題內容	出現次數	百分比 (次數/總次數) %
2	長期投入	2	6.90
3	流程簡化	2	6.90
4	成果驗證	1	3.45
5	跨域協調	7	24.14
6	資源落差	2	6.90
7	階層落差	5	17.24
8	課責問題	2	6.90
9	專業決策	1	3.45
10	資訊事權統一	2	6.90
11	共通性需求	2	6.90
總計		29	100

資料來源：本計畫。

2. 服務設計面

受訪者意見主要聚焦在(1)資訊安全、(2)個資保護、(3)服務介面設計、(4)人性化設計、(5)顧客導向設計、(6)人力取代、(7)人機互動、(8)建立標準等八個議題(如表50)。受訪者表示資訊服務或資訊系統在設計階段，應重視資訊安全與民眾個資保護，優化服務介面的設計、重視易用性的人性化和傾聽的顧客導向設計，滿足政府內部和外部使用者的服務需求，強調政府服務需保有溫度，人力無法完全被機器取代，且應重視人機互動，培養公務員對於資訊系統的信任，最後建議應建立 AI 和資訊系統開發的共通性標準，以加速開發流程。

表50：服務設計面向議題內容及次數彙整表

議題編號	議題內容	出現次數	百分比 (次數/總次數) %
1	資訊安全	3	21.43
2	個資保護	1	7.14
3	服務介面設計	1	7.14
4	人性化設計	1	7.14
5	顧客導向設計	3	21.43
6	人力取代	2	14.29
7	人機互動	1	7.14
8	建立標準	2	14.29
總計		14	100

資料來源：本計畫

3. 人才職能面

受訪者意見主要聚焦在 (1)資訊人力不足、(2)資訊採購能力、(3)大數據分析能力、(4)資訊素養、(5)世代落差、(6)教育訓練、(7)員工激勵等七個議題（如表 51）。受訪者表示行政機關資訊人力不足，且長期透過委外方式推動資訊化業務，需強化資訊人員的資訊採購能力，政府推動數位轉型但政府人力在大數據分析能力、資訊應用能力素養不足，且在思維觀念存有世代落差，需透過教育訓練彌補不足之處。另外，政府數位轉型涉及諸多法令變革、行政流程變革與資訊系統開發等任務要求，公務員任務繁重，應重視轉型過程的員工激勵問題。綜合前述人才職能落差分析結果，建議未來在政府應用 AI 之人才職能的發展重點，應優先提升文官的數據分析能力以及資訊素養，其次是強化資訊採購和資源分配能力，最後則是在制度面增補資訊專業人才，以強化政府部門的 AI 系統維運能力，當然思考如何在政府推動數位轉型的過程降地跨域協調的障礙，並適時提高對於員工的激勵措施，亦是和人才發展相關的重要組織管理和制度面議題。

表51：人才職能面向議題內容及次數彙整表

議題編號	議題內容	出現次數	百分比 (次數/總次數) %
1	資訊人力不足	2	16.67
2	資訊採購能力	1	8.33
3	大數據分析能力不足	1	8.33
4	資訊素養	4	33.33
5	世代落差	1	8.33
6	教育訓練	1	8.33
7	員工激勵	2	16.67
總計		12	100

資料來源：本計畫

4. 其他面向

有一位受訪者表示問卷聚焦在 AI 議題過於專業，建議未來能先區隔已導入和未導入 AI 技術之機關再行調查，以免調查失去客觀性。

三、 才能管理

(一) 人機協作職能落差分析

針對深度訪談與焦點座談分析結果，對於本計畫前此文獻檢閱所歸結出的三項核心職能面向，分別提出目前實務發展與人機協作職能規範上的落差。茲分別敘述如下：

1、 社會性職能落差

本計畫在社會性職能中包含三項核心職能，分別是社會性技能、倫理與安全技能以及創新性。社會性技能主要強調協調合作、說服溝通等人際互動能力，從深度訪談與焦點座談分析發現，AI 或者聊天機器人的引入，主要挑戰在於跨部門間的資料串接與整合，另外則是跨領域間的協調對話。如同前此文獻所提及，目前公部門 AI 的引入多採委託外包形式，對此，政府單位必須更具備跨領域溝通的對話能力，一方面傳達本身需求，另一方面能架接本身領域知識與外部資訊領域進行整合，以更進一步精確瞄準組織本身需要的服務內容。然而，從分析資料發現，目前委託外包下時，常出現政府單位與受託廠商因不同領域而產生彼此對話上的困難。

再者，除了領域之間的差異造成的溝通障礙外，AI 引入政府內部另一項必須面對的問題在於，部會間跨單位的溝通問題。以聊天機器人為例，機器要能有效回應外部民眾的問題，必須在政府部門串接組織內不同單位間的內部資料，然而也許因本位主義的影響，業務承辦的責任歸屬問題，形成業務承辦無法獲取到其他單位的資料，以致聊天機器人無法彰顯它的實質效益。也因如此，本計畫更特別強調業務承辦同仁與其主管在核心職能中有關社會性技能的重要性，畢竟公私協力的溝通決定了整體服務的發展架構與藍圖，而政府組織內部跨單位間的協調問題決定了資料庫建構的完整性，此係人機協作未來是否能夠發揮效益的關鍵。

最後，必須要特別說明的是，創新性在本計畫中主要係指成員必須具備創新思維，然而我們分析資料，均指向政府在引入 AI 的同時，對於機器本身的容錯率不高，究其主因本計畫認為與政府本質有關，畢竟其所面對的是一般社會大眾，因此，回應性與課責是必須堅守的核心價值，然而，創新與容錯率將是引入新科技下的兩難困境，在創新思維的引導並且大膽引用任何可能的 AI 作法，相對的，就必須擔負起機器出錯時所需要承擔的責任。而這也是本計畫分析資料所強調的，上級主管對於該項新科技的引入，扮演的支持性角色是否有被彰顯有關。

2、 管理職能落差

管理職能在人機協作職能中主要涵蓋有財務資源管理、原物料管理、人員管理、時間管理、資訊與知識管理及契約管理等，在資料分析過程中本計畫發現，目前政府引進 AI 的過程，對於資訊與知識管理及契約管理能力存有顯著落差。就前者而言，多數受訪專家指出，目前政府單位內部無論承辦同仁或業務主管，普遍存有新科技資訊素養不足的問題，也因對資訊素養不足連帶影響與外部受託廠商簽約過程中，資訊不對稱所衍生的交易成本，做出不利於己的逆向選擇。由於受訪專家發現，多數受託廠商在承接政府委託案後，會透過築起技術牆的方式，綁架政府內部系統，逼迫政府未來與該系統有關的任何維運或服務擴充都得與原受託廠商合

作，若政府內部在人員訓練上能跟著新資訊科技的發展與引入，連帶地增加相關教育訓練課程，本計畫認為將或多或少可降低政府內部與外部廠商在交易平台上因資訊不對稱，造成彼此在契約簽訂過程中所產生的不利影響。

再者，以上論述實也強調政府單位對於契約管理能力的要求。人機協作職能中對於契約管理能力的界定，係指組織成員具備能維繫與受託廠商間的關係以及彼此溝通的管道。為了減少雙方資訊不對稱的關係，導致系統建置與資料庫串接的困難，在簽約過程中的對話機制建立能力，及對話平台上如何與受託廠商磋商的能力就非常重要。

3、專業職能落差

專業職能中主要包括系統技能、技術能力與解決複雜問題的能力，由於這一塊在目前的運作過程，多係委託給外部廠商為政府建構系統，雖說這部分屬於技術專業，但身為未來與機器協作的個人，仍然得具備相當操作技術與解讀機器輸出結果的能力，當然，若對於機器的瞭解程度越深，越能知道該蒐集那些資料，對於機器的再訓練與判讀更有所助益，同時，誠如我們深度訪談與焦點團體的專家所指出，人機協作的關鍵因素在維運，但就資料分析顯示，目前在維運這一塊似乎多委由外部廠商來協助，或許未來在維運這一塊的能力，公部門的比例可以多增加，除了可以更瞭解如何讓機器判讀的更好之外，也可以免除公部門在跨領域合作中多處於資訊不對等的不利位子，當然這點也跟前此資源管理能力中，資訊與知識管理能力的訓練有所關聯，職能與職能間的關聯性應是相互連動的。再者，除了可以增加機器本身的實質效益外，也因為對機器的瞭解更多，當接收到外部民眾使用機器後主客觀評估感受的意見，將更能知道該如何針對機器的部分進行調整，此將更能有效回應民眾的使用需求，或許也會因這樣的回應性，而增加民眾對於機器本身的感受，能夠更具情感與信任度。

針對上述所提出之落差，如同在人機協作職能中間卷調查的發現，本計畫認為可透過相關的內部教育訓練來達到人員相關職能的培訓，如跨域協調能力、溝通能力、契約管理能力等。再者，針對聊天機器人所引發民眾可能的不信任感與情感感受度較低的問題上，本計畫認為或可參考 Makasi 等（2020）所提出的公共服務價值架構（public service value-based framework）已獲得完善的解決，該文之所以提出這樣的架構乃是注意到，當政府機關在廣泛應用聊天機器人於政府提供民眾服務的環境中，常受到民眾的質疑、爭議與法律上的挑戰，因而提出此服務價值架構以作為政府機關在使用聊天機器人時，如何透過聊天機器人來展現此十四種價值。而此些似可作為政府在引進聊天機器人時的評估標準。此十四項價值請參看下表 52 所示：

表52：以聊天機器人支援公共價值的傳遞

公共服務價值	以聊天機器人為服務提供中介的價值定義
適應性	聊天機器人在提供服務時，適應不同條件（不同的非技術條件，如業務規則和服務資格的變化，以及不同的技術條件，如適應不同的設備和網絡）的程度。
使用者導向	聊天機器人在用於解決案例時滿足用戶表達和需求的程度。
專業主義	聊天機器人在提供服務時表現出有原則、有能力、誠實、尊重、一致和值得信賴的行為的程度。
效能	在既定資源投入下，聊天機器人達到其所應達成目的的程度。
效率	聊天機器人在提供服務的過程中減少成本和所需資源的程度。
公平性	聊天機器人在提供服務時不存在偏袒和歧視的程度（基於個體差異的存在）。
合法性	聊天機器人在提供服務時遵守／符合合法合理的步驟及要求的程度。
可接受性	聊天機器人是服務提供方式中可實行的程度。
開放性	聊天機器人在服務提供前向使用者表露其身份，以及為其決策提供理由的程度。
課責	聊天機器人在提供公共服務時表現負責管道的程度。
社會許可	聊天機器人作為可行的服務提供管道，獲得社區和其他利益相關者持續認可的程度。
隱私	聊天機器人在服務提供期間或之後，對於使用者個人資訊的保護程度。
信任政府	聊天機器人促進民眾願意分享個人資訊予政府以提供更佳服務的程度。
協作智能	此指使用聊天機器人的相關利害關係人，認為聊天機器人是一有效夥伴關係的程度，認為它的確有達到當初設計與服務目的的程度。

資料來源：Makasi, T., Nili, A., Desouza, K. & Tate (2020)。

(二) 調查實驗分析

在法務部保護司觀護人實驗調查的分析中（詳見本文第四章第三節），我們發現控制組、實驗 Arm1 組以及 Arm2 組觀護人對於實驗前（A4.1 至 A6）以及實驗後創新意願、工作流程自動化以及公共服務動機（A3 至 A5）

題組之回答並無顯著之差異，就此可以理解為本實驗有確實做到隨機分組，使得上述題目的回答情形與分佈並未因組別的不同而有所差異。

下表 53 則呈現再犯風險評估業務之實驗調查結果所將出現的不同類型決策模式，本實驗對受試者行為類型的分類依據主要來自對個案第一次與第二次判斷是否有差異。就控制組而言，由於並未受到任何外部資料介入，本實驗於受試者在第一次作答完後呈現第一次判斷的結果，並詢問是否要做修改。若受試者修改了答案，則我們便將其歸類為「自我質疑」，若未修改答案則被歸類為「沒有行為變化」。而實驗組 Arm1 與 Arm2 則分別受到 AI 判斷資訊以及其他同仁判斷資訊介入，若 Arm1、Arm2 的受試者修改了答案，則我們便分別將其歸類為「被 AI 所影響」與「被同仁影響」，若未修改答案則被歸類為「沒有行為變化」。

表53：外部介入資訊對受試者判斷影響整理

實驗組別	類型 編號	受試者 第一次判斷	受試者是否有 外部資料介入	受試者 第二次判斷	行為類型
控制組	1	正確	無	修改	自我質疑
	2	正確	無	不修改	沒有行為變化
	3	錯誤	無	修改	自我質疑
	4	錯誤	無	不修改	沒有行為變化
實驗組 Arm 1	5	正確	AI 判斷資訊 介入	修改	被 AI 所影響
	6	正確	AI 判斷資訊 介入	不修改	沒有行為變化
	7	錯誤	AI 判斷資訊 介入	修改	被 AI 所影響
	8	錯誤	AI 判斷資訊 介入	不修改	沒有行為變化
實驗組	9	正確	其他同仁判斷 資訊介入	修改	被同仁影響

實驗組別	類型 編號	受試者 第一次判斷	受試者是否有 外部資料介入	受試者 第二次判斷	行為類型
Arm 2	10	正確	其他同仁判斷 資訊介入	不修改	沒有行為變化
	11	錯誤	其他同仁判斷 資訊介入	修改	被同仁影響
	12	錯誤	其他同仁判斷 資訊介入	不修改	沒有行為變化

資料來源：本計畫。

(三)個體層次問題的落差分析

落差分析也開始被應用於公共行政領域中（如：電子化政府的發展），其目的在於找到政策推動時，找到「目前現狀」（as is）與「未來預期」（to be）的差異，且可用來進一步研析若要達到「未來預期」的狀態，現況仍需做哪些修改的工具。此分析被認為適合用於公共服務流程管理，對於數位轉型中強調利用新科技傳遞與提供既有服務與業務的目的，也符合落差分析的適用性。此節從個體層次的問題出發，回應在推動機器學習演算技術於數位轉型工作上，尤其是人機協作的情境上，應從「認知」與「行為」兩個構面著眼進行落差分析。以相關學術理論為基礎，我們分析既有理論如何說明與假定個人在現狀上會有哪些問題，並以此進一步討論這些問題與未來數位轉型的關係。在認知面方面，本計畫認為有展望理論、科技接受模型與期望地位理論可供討論現狀與未來的做法。在行為面則有有限理性、滿意行為理論、自動化自滿、資通訊系統服從理論可供討論。下表 54 整理歸納相關概念與相關理論。

表54：機器學習演算對公部門個體層次影響之理論視角落差分析

構面	引用理論	現狀（As is）	未來（To be）
認知面 落差	展望理論 (Kahneman & Tversky, 1979)	1. 多數為風險厭惡者 2. 對損失比對收益敏感。	1. 改善風險厭惡所導致，對機器學習演算之應用與新技術產生過度不信任問題。

構面	引用理論	現狀 (As is)	未來 (To be)
			2. 提醒科技過度樂觀可能產生對機器學習演算之態度偏好。 3. 主導機器學習演算方，可能因此產生科技接受的正向偏誤。
	科技接受模型 (Davis, 1989)	1. 科技的接受程度越高越樂觀，正面影響使用意願。 2. 相較於一般大眾，公部門公部門對於應用資通科技能提升公共服務的品質更具信心 (Moon & Welch, 2005)。	1. 提醒科技過度樂觀可能產生對機器學習演算之態度偏好。 2. 主導機器學習演算方，可能因此產生科技接受的正向偏誤。
	期望地位理論 (Berger, 1977)	1. 認知自我能力越高，越是高標準評斷他人能力，反之亦然 (專業主義假設)。	1. 透過改善機器解釋性與透明性，排除對機器學習演算的先入為主。
行為面 落差	行政行為之有限理性理論、滿意行為 (Simon, 1947)	1. 受資訊不充分影響，導致決策自主限制。 2. 放下專業判斷能力而對既有	1. 提升資料品質與機器演算透明度，以降低滿意

構面	引用理論	現狀 (As is)	未來 (To be)
		決定滿意或妥協。	行為的主導。
	自動化自滿與偏誤 (Parasuraman & Manzey, 2010)、控制問題 (Bainbridge, 1983)	1. 自滿傾向愈高，容易對自動化產生高認同。 2. 過度依賴機器的判斷。	1. 重新思考人/公部門員工的角色，從直接做決定到成為監督機器學習演算工具之角色需要有新的認知與訓練。
	資訊系統服從效應 (Bovens & Zouridis, 2002)	1. 因為系統自動與智慧化，導致行政裁量空間的縮減或消失。	1. 反思系統官僚對公部門組織的影響，重新定位不同業務與位階官僚行政裁量的要務，避免與機器學習演算工具之裁量行為衝突。

資料來源：本計畫。

四、 小結

表55：三種策略途徑下的落差分析

策略途徑	實然問題
法規調適	組織規劃定位：組改時程無法趕上科技的日新月異再加上領域多元與專業，成立專責機關監理有其難度。 規制方式制度：新興科技的具體規範甚是缺乏，進而使資料使用規範以及倫理價值的遵守之約束力薄弱。 倫理價值與政策原則：告知程序與揭露義務的完備度不足，演算法之安全性跟可靠性仍待釐清，進而影響風險管控之

	監理程度。
資源配置	預期正向成效：資料準備度各部情形不一與相容問題影響對於新興科技的信心。 預期承擔責任：組織管理系統對於新興科技的建置影響文官需承擔更多運用新興科技之責任。 預期權力限制：文官的心理韌性越差以及對於科技過度樂觀會影響組織變革的方向與進程。
才能管理	社會性職能落差：政府跨單位與公私合作的溝通協調問題影響資料串連與專業溝通的成效。 管理職能落差：科技資訊素養不足連帶影響到契約管理能力使技術維運受到外部廠商限制。 專業職能落差：文官對於技術操作與解讀機器輸出結果能力較為薄弱，無法有效訓練該科技以及提高民眾運用的良好感受度。

資料來源：本計畫。

表 55 為總整前述法規調適、資源配置與才能管理三種策略途徑觀察，所得出我國公部門目前執行之實然問題，據此以下詳述三種策略途徑之執行落差。

(一) 法規調適層面

從法規調適層面觀之，從國外經驗反觀我國目前情況，以四大面向進行落差分析。第一、組織規劃定位：OECD 報告指出，專責政府機構組織應區分為「整合協調機關」與「治理監督機關」，前者應以現有或新設獨立之部門為之；後者除以現有或新設獨立部門之外，亦可成立專家小組或新設獨立諮詢機構方式。從前述的組織規劃之區分，可發現機關之功能除可提供相關資訊並提出政策方向之外，也可進行案件與個案之監督。從目前的國外個案的採擷可知美國、英國、新加坡多採取兩種機關性質之綜合方式，除原有科技、資料治理的單位繼續執行政策擬定與監理之任務；針對諮詢部分多新設委員會或成立專案工作小組。此為綜觀各國之經驗，理想的機關組織的建置與執行型態。

然而，我國目前對於演算法之引進協助公部門之業務甚至公共服務的執行，仍在初期探索階段，對於該技術之規範與執行仍屬模糊，再加上演算法應用於各領域中會受到該領域所需之專業性與需求性的多元程度的不同，影響其適用性。事實上，我國現況若要發展出完全適用於每一領域的指引以及成立專責機關有其困難，原因有二：一、我國中央部會機關目前仍在組改階段，目前共有 32 個機關，組改時程十載有餘，仍努力往行政院組織法設立 29 個部會邁進，但這樣的組改時程無法跟上科技的日

新月異，因此若要以專責機關方式進行新興科技的管理，確實有其難度，或許未來數位發展部的成立有其必要擔起新興科技運用於公部門之管理之責，但這樣的管理也會面臨另一問題；即是二、各領域的專業性較高且多元，專責機關的人力職能、組織量能是否能深入瞭解各領域之專業性，給予適切指引與監督。

第二、規制方式制度：從歐盟、美國與英國的制度演變進程來看，其演進過程漸漸聚焦，透過相關研究報告瞭解現況，進一步綜整成政策白皮書進行政策的規劃，再來是擬定出相關執行指引，最後提出法案，此演進可知其約束力道漸漸增強。以 AI 為例，目前美國與歐盟對此技術之規範，其法案仍為草案階段尚未通過，表示目前各國以相關倫理指引或架構文件作為執行上的指導原則與依據，其拘束力仍有待觀之。

反觀我國，誠如前述所提雖然演算法技術已行之有年，但針對新興科技如 AI 的具體規範仍為不足，相關指引的設計也未見雛形。因此，我國極缺針對各式新興科技運用的規範，包含針對民間單位之運用，抑或是引進公部門之使用。基此，在任何新興科技運用的前提為大數據的建立與使用上，該如何加強最基本的個人資料保護機制、建立資料使用規範以及最高層次的倫理價值的維持，為我國目前如欲建立規範制度應先行思量之重點與解決的問題。

第三、倫理價值與政策原則：培育一可運用之新興科技之前置作業為使用大量個人資料並搭配所適的演算法，進而讓機器學習。因此，當運用大量個人資料時候，對於資料主體即是資料提供者，以及後續新興科技的使用者，其權利有可能會造成侵害。綜整各國之文獻發現其規範與建議等，均將倫理價值列為最重要的事項，任何新興科技的研發到運用都必須合乎五項共同原則，分別為對資料主體尊重自主、公平性、可解釋性、透明性、課責性以及重視隱私。

針對五大原則，以我國法規現況而言，有三層次問題需要思量。一、告知程序與揭露義務的完備，目前此程序的規範仍不完備；二、透明性與可解釋性的考量，如何驗證演算法之安全性跟可靠性，仍待釐清；三、參數及過程均充分揭露後，何人有能力可以判斷其公平性或透明度是否足夠，該如何管制、監督與救濟，牽涉到整體監管的組織建置與法規規範。

第四、風險評估分級：依照各領域需要新興科技協助之事務性質，應視該事務性質本身以及需求，給予不同程度的監管力道，其程度區分以該新興科技之應用目的領域所產生風險強度作為依據。歐盟《AI 規則草案》具體列舉高風險之事務有生物辨識、民生資源供應、教育、人力資源管理、獲得公共服務的資格、執法事務與司法評估、移民評估等，均涉及對人民較為核心之利益與財罰認定。因此，對於此揭高風險事項，在監理的強度便需嚴格，建立在維護倫理價值之價值之上。

(二) 資源配置層面

本計畫針對中央各部會應用演算法之調查，從組織與個人層面觀察，分為兩階段的調查，第一階段主要瞭解目前該部會導入 AI 系統之情形；第二階段為瞭解文官對於應用機器演算法創新公共服務的態度、預期效能以及可能遭遇的問題其解決態度為何。針對二階段的調查，上述已有詳盡的分析，就目前我國部會應用機器演算法創新公共服務的現況，提出三項對於新興科技的接受與運用上的問題。

從第一階段調查可知，中央部會應用機器演算法創新公共服務是可行的，回覆的十八個部會中有八個部會已導入或正規劃導入 AI 系統輔助業務，其最多部會優先推動的業務類型分別為服務對象溝通工作，其次為操作機器的自動化和辦公室行政文書工作。但對於未導入 AI 系統的部會，其遭遇之限制相較於過去文獻中多為探討組織制度的影響不同，此些部會主要遭遇的限制來自於資源面居多，如：缺乏預算、缺乏適用技術、缺乏專業人力。

面對上述所提之限制，進一步思考中央部會如欲應用機器演算法所需條件有哪些？仍是應優先具備資源面的條件才能進一步發展應用，如：編列採購和發展 AI 系統所需預算、專案資訊人力、機關數位資料之需求。這些條件之具備其目的是發展該機關的 AI 系統滿足內部對外公共服務所需，受訪者尚未表示有跨部會 AI 系統發展的需求。

然而，面對現今中央部會應用機器演算法之現況。公務人員對於應用機器演算法創新公共服務的態度亦會影響整體組織面對數位轉型時的改革方向與成效。從文官運用 AI 的三種預期態度，分別為：預期正向成效、預期承擔責任、預期權力限制上，受到心理韌性、科技樂觀、資料準備度、支持的組織管理系統等因素的影響。

第一、預期正向成效：當文官意識到所屬機關的資料準備度愈高，將愈預期使用 AI 能帶來正向的成效。表示當機關對於運用於機器演算法之資料準備度越齊全，文官普遍認為透過齊全資料建置的 AI 是可以為業務執行上帶來正向成效與幫助。就我國各機關的現況，各部會的資料準備度情形不一，部分機關表示會面臨資料串接的相容性問題，無法有效運用機關內的資料去做進一步的運用。因此如欲發展新興科技協助業務執行，首當其衝是資料的準備，包含：資料的詳細程度、資料串聯的相容性等，皆會影響到訓練出來的新興科技是否能適切發揮協助業務執行的成效。

第二、預期承擔責任：當文官意識到所屬機關的資料準備度愈高，將愈會預期運用 AI 需承擔的責任，意思是雖然資料準備度越高越能帶來正向成效，反之也讓文官開始思考自身承擔的責任是否加重，這個責任可能來自於當機器協助業務執行時，因為運作有技術上或其他不可控的因素

造成的問題時，其責任的歸屬是否為操作者本身；另一責任來自於當業務上有新興科技協助之時，代表現有的人力因新興科技的協助下之事半功倍，機關會進一步讓該人力額外處理其他業務，此時對該位文官來說，業務增加也是責任加重的一種。另外，若文官能感受到組織的管理系統能展現對員工的支持，反而愈不預期自己需承擔責任，亦即當整體組織對於新興科技輔助業務執行的運用是很支持也有規劃健全的體制運作，對文官來說會因為受到健全體制與支持的組織文化影響下，較為安心去運用機器演算法協助業務，若發現問題該管理系統會處理，藉由組織層次對於問題釐清與解決，而非由文官個人承擔。

第三、預期權力限制：當文官的心理韌性愈好，愈不預期使用 AI 會限制自己在工作上的權力，但當文官愈是抱持科技樂觀的態度，反而愈是預期使用 AI 會限制自己在工作上的權力。面對數位轉型勢必伴隨組織變革之因應，要將原先的組織文化與制度解凍，對於已經行之多年的文官來說，會產生抗拒的心理。因此，文官應加強自我的心理韌性才能因應政府數位轉型帶來的工作變革；同時需務實看待行政機關導入 AI 系統產生的業務變革，才不致於過度樂觀而忽略伴隨變革而來的衝擊。

(三) 才能管理層面

機器演算法於公部門的運用，除須考量制度之建置與組織的資源配置之外，也有賴於組織內成員的對於公共服務的數位創新有積極的態度與認知，方能較為順利運行。因此，瞭解文官對於應用機器演算法於公共服務上，應知悉人機協作的態度價值、目前受到的限制以及執行的關鍵因素；另外，人機協作下應具備哪些職能得以因應數位轉型也至關重要。根據深度訪談、焦點座談與中央各部會之問卷調查，歸納出目前發展人機協作職能的應然與實然的落差。

第一、社會性職能落差：該職能包含三項核心職能，分別為社會性技能、倫理與安全技能以及創新性。其中，社會性技能尤其重要，從文獻上得知當公部門引進機器演算法等新興科技運用於業務與公共服務時，針對機關內部需要跨部門的資料串接與整合；對於技術廠商多以外包方式為之，公私的合作與互動也會影響專案的進行，因此該技能強調協調合作與人際互動的能力。

事實上，透過訪談與座談得知目前公部門的現況，當需要跨部會的資料串接與整合時會面臨資料的相容性的問題，另外，資料整理過程也會有是否清楚機器演算法所需的資料性質的認知與彙整能力，因此常有資料無法串連與介接的情事；再者，該科技技術的建構若完全仰賴公部門自身的能量是不具足以為之，必須仰賴外部廠商給予協助，此時會面臨專業不對稱與本位主義的問題，以至於溝通上會出現需求無法滿足的情形。故政府組織內部跨單位間的協調問題決定資料庫建構的完整性，此係人機協

作未來是否能夠發揮效益的關鍵，而前述之理想的達成更需長官對於新技術引入的支持，才能得以使機關內部的水平連結更具成效。

第二、管理職能落差：人機協作的職能中該方面的職能涵蓋財務資源管理、原物料管理、人員管理、時間管理、資訊與知識管理及契約管理。目前政府引進 AI 的過程，有著管理上不足之處，針對資訊與知識管理能力，政府內部承辦同仁或業務主管，存有科技資訊素養不足的問題，此問題會連帶影響到契約管理能力，存有外部受託廠商簽約過程中，資訊不對稱所衍生的交易成本，做出不利於己的逆向選擇，多數受託廠商在承接政府委託案後，會築起技術牆的方式，綁架政府內部系統，逼迫政府未來與該系統有關的任何維運或服務擴充都得與原受託廠商合作。

第三、專業職能落差：主要包括系統技能、技術能力與解決複雜問題的能力。雖說多委託外部廠商為政府建構系統，但未來與機器協作的個人，仍然得具備相當操作技術與解讀機器輸出結果的能力，對於機器的訓練與判讀有所助益，而人機協作的關鍵因素在維運，目前維運多委由外部廠商來協助，因此在專業不對稱之情形下，公部門在跨領域合作中多處劣勢，受限於廠商的供應。對外，民眾使用的感受程度與回應性，與政府內部人員的資訊與知識管理能力的訓練有所關聯。

第五章 結論與政策建議

第一節 研究結論

一、 現行研究所得之結論

(一)法規調適

總整對於OECD、歐盟、美國、英國與新加坡個案的採擷，從「倫理或價值原則」、「政策循證原則」、「政府組織」與「規範架構」四大面向來探討。並從「組織規劃定位」、「規制方式制度」、「倫理價值與政策原則」以及「風險評估分級」來釐清我國現況，進而提出未來引進機器演算法時可調適之處。

第一，以「倫理或價值原則」層面觀之，從各國的相關的資料來看，大多是針對目前運用演算法於 AI 中的模式，提出指引方針與處理原則，法律規範上仍著重於個人資料的保護，在原則上目前檢視的國家所提出的共通性原則皆為「公平性」、「課責性」、「合法性」、「可解釋性」、「透明性」、「包容性」、「安全性」、「信賴性」，其他細緻的原則性的規範會視需求而定，但從前述的共通性原則來看，顯然皆是立基於以人為本的人權保障之下，追求科技帶來的永續發展，成為人機協作下的互惠、互信、互賴的運作模式。

第二，從「政策循證原則」層面觀之，政府漸漸也在思考，公共利益的最大化在未來應透過更多巨量資料的演算分析，帶來更準確、更快速的政策制訂與執行。因此，政府如何運用演算法於政策的制定與執行過程中，仍保持倫理或價值原則，取得平衡為之重要。基於前述的思維，政府應關注的面向包括：提供公眾參與以增加數據共享的誘因、提高安全性和保護措施、更新演算法的監管方法、評估可接受的風險和做決策的倫理、以及考量其他政策需求的平衡，以避免在鼓勵創新與權利侵害中產生兩難。

第三，「政府組織」為落實原則的實踐者。演算法如欲引進政府加以妥善運用，政府為最關鍵的實踐者，影響到該技術的發展與落地。從現有的資料來看，目前的組織型態皆建議走向「跨部門組織」、「委員會」的型態，其重點為需要有一個能夠整合與協調的政府組織，有其權力可使政府各部門為該運作提供協助之處，進而達到該目標，過程也須仰賴外部專家的參與與諮詢，使其可行，如：美國成立專責委員會、英國成立科學辦公室、新加坡成立智慧國家與數位政府辦公室。然而，OECD 提出整合協調組織與治理監督的組織的建議。

第四，「規範制定架構」為前述倫理價值原則的具體落實輪廓。在普世價值原則下，政府除了思考未來政策制定與執行運用演算法的目標，以及專責機關的設立與組成之外，有一套的規範架構有所依循，是使政府有所為的最佳指引。從各國的規範架構上可發現除了有關人權保障的硬性規定之外，都有一套軟性指引，不管是策略擬定、應用指南、白皮書、備忘錄等，讓演算法應用於 AI 的過程中，仍保有彈性發展與滾動式修正。其中，歐盟也針對演算法應用於 AI 中帶來的風險進行分級，依照不同級別應給予不同程度的監理，才能確保風險的控管。

從上述四大面向對於各國發展之情形的陳述，皆可得知法規的制定不外乎基於倫理或價值原則之上，以維護人權保障為目標，追求科技帶來的永續發展。以此共通性原則思量組織的建置、制度的規劃等。在第四章第四節的落差分析中，針對法規調適途徑提出目前我國如欲運用機器演算法於公部門中，應思量之問題，該段落針對問題提出可因應之道。

第一，組織規劃定位，行政院組改時程之速度無法趕上科技的日益更迭，另外，各領域之專業性與多元性程度較高，再加上目前公部門機關已過多，實難期待新增一專責機關可有效率負責相關業務，以現階段可就現有機關擇一負責，為較可行之方式。建議在數位發展部籌備階段之時，考量新興科技於公部門運用的組織建置；籌備的過渡階段，挑選現今規劃與發展數位轉型與智慧國家的機關，如國家發展委員會資訊管理處，先行盤點各部會的新興科技運用情形，以及其主責機關，做為未來數位發展部設立專責機關時加速橫向連結與發展重點項目時的基礎，具深入瞭解各部會發展的情形才能進而設計出以下要談論的共通原則與審查機制。

第二，規制方式制度，我國目前首當其衝即是對於個人資料保護機制的缺乏，而個資是整體發展的重要基石，現階段我國沒有個資保護之主管機關，也因此缺乏對於個人資料保護的獨立第三方監理單位，只能由各目的事業主管機關依個資法之規範進行認定監管。借鏡國外經驗，歐盟 AI 規則草案之規範內容，針對資料治理有清楚的法規架構，前階段巨量資料與參數之累積，須透過完整之資料治理規範，才有可能妥適處理後續 AI 衍生之問題並進行課責。我國現況對於任一新興科技的具體規範較為缺乏，進而對於資料使用規範與倫理價值的觀念較為薄弱。因此，基於組織規畫定位提出之建議，指定一統籌機關對於演算法之能力培養，並針對基礎設施之應用與計算等制訂共通原則、審查機制等項目，至於細部之發展、管理、維護等宜依各部會之業務別區分之。原因來自各領域的專業性較為多元，欲將演算法導入各該領域中，需要對該領域具相當程度瞭解，才能使演算法運用過程中適切發揮成效，故需目的事業主管機關的負責參與，同時為避免各領域對於資料治理與管制演算法設計之標準不一，不利於

公部門進行後續管制與究責，因此須由專責機關制訂共通原則與審查機制，成為各領域之業務主責機關遵循辦理之重要依據。

第三，倫理價值與政策原則，綜觀各國的相關指引與法案之草案，皆以維護資料使用之倫理價值為核心，擬定相關對於資料提供者的保護機制以及對於資料使用者的相關規範與課責機制。我國在建立規制方式制度時，應將倫理價值與政策原則列為重要核心，該原則可分為三個層級，分別為：一、告知程序與揭露義務的完備；二、為透明性與可解釋性的考量；三、為當這些參數及過程均充分揭露後，須有人有能力可以判斷其公平性或透明度是否足夠。

第四，風險評估分級，釐清各領域的專業性以及運用機器演算法之目的後，應進一步思考在共通原則上，其審查機制中如何判斷對於運用機器演算法於該領域中協助業務推動時的風險控管。依照風險程度區分其監管力道，越高風險之事項越須高強度的監管，才得以讓運用機器演算法過程降低對於資料提供者最少的侵害，同時搭配上所述之規範建立，使前端的資料運用到後端若風險發生時的後續救濟程序，得以可以受到規範保障。

建議上述這個部分統整各國經驗所建構出來的初步設計焦點，可以在另外一個以撰寫演算法監理機制辦法的研究案裡面充分設計，並且將其立法意旨、可行性分析甚至推動說帖都準備齊全的一個研究案。

(二)資源配置

針對我國公部門如應用演算法於公共服務，須先盤點現今中央政府各部會應用演算法創新公共服務上，目前的準備度、公務人員的推動意願與才能缺口、以及預期可能遭遇的風險與阻礙，這些調查面向將作為落差分析以及後續政策路徑規劃的參考依據。在此調查目的之下，本計畫著重於組織層次與個人層次的調查，組織層次的調查，希望能瞭解演算法創新公共服務的可行性，並且調查該機關可能優先推動的項目，進而瞭解可能遭遇到的制度限制與所需資源，其資源包含預算、人力、資料需求等；個人層次的調查著重於公務人員應用演算法創新公共服務的態度、預期效能以及對可能遭遇問題的解決態度。

基於前述的內容方向，本計畫該部分的調查研究，將分成兩階段執行，第一階段以部會組織層次為分析單元；第二階段調查以個人層次為分析單元。透過survey monkey網路問卷調查，在調查實施前先以公文邀請機關代表填答或請受訪機關推薦受訪者，再透過受訪者個人電子信箱實施網路問卷調查。抽樣設計上第一階段邀請中央各部會及指定資管人員填答；第二階段採用滾雪球抽樣方法，依據第一階段的調查結果，以公文

邀請中央各部會的「資訊單位」、「業務單位」與「人事單位」科長級以上主管1名、以及該主管指定相關業務同仁至少2名填答問卷。

從第一階段的調查可知，中央各部會應用機器演算法創新公共服務具可行之處，最多部會優先推動的業務類型分別為服務對象溝通工作，其次為操作機器的自動化和辦公室行政文書工作。但在推行上可能會遭遇限制來自於資源層面，從調查中發現最多部會認為缺乏預算與適用技術，以及專業人力的缺乏，為運用機器演算法的阻礙。此時，需要思量究竟要使中央各部會運用機器演算法於公共服務中，應具備什麼條件。

第一，資料準備度的完備。從調查結果得知當資料準備度越高，文官預期AI能帶來的成效越正向。我國目前各部會的資料準備度不一，部間的跨域資料串連也不成熟。因此，就目前現況應將資料準備度提升，資料準備度攸關於資料串接的相容性，若該部會愈發展運用機器演算法協助公共服務的推動，須先盤點內部資料的情形與整理的詳細程度，當資料越細緻能運用的程度越高；另外，資料串連的相容性也會受到前一階段的資料整理細緻程度影響，故該部會必須先行釐清訓練的機器演算法於新興科技上的運用，預期為公共服務帶來多大的成效，進而處理資料的細緻程度與串接情形，使資料準備度達到可運用之要求。

第二，組織管理系統的建置。面對新興科技的應用，對於操作者來說會顧慮到當運用該技術於業務中若發生錯誤時，責任的歸屬為何？在公部門的課責環境中該問題更顯重要。本計畫認為當組織面對新興科技的導入應以組織層次思慮其整體管理系統的應變；首先，長官對於新興科技接受程度與管理思維需要跟進，若長官對於該運用抱持支持態度，才能啟動變革開關，以整體組織視角規劃健全的體制運作，該體制的運作目的是要以如何讓新興科技的運用為業務執行帶來最大效益為目標，亦即將出錯風險降到最低，在變革前期需要經過試驗等等方式，找出出錯點並解決，才得以讓正式運用時得以最低風險為之，對文官來說會因為受到健全體制與支持的組織文化影響下，較為安心去運用機器演算法協助業務。

第三，面對新興科技的健全態度。此方面牽涉到個體層次的觀點，面對組織因應新興科技帶來的轉型，須對於自身有心理調適與認知調整。從調查中發現，文官的心理韌性越好，對於AI的態度與認知是認為AI對於工作會帶來幫助，非限制工作權力。雖然新興科技可以帶來益處，但若文官對於科技抱持過度樂觀，反而會認為AI對於工作權力有所限制。因此，文官應加強自我的心理韌性才能因應政府數位轉型帶來的工作變革，這方面更需要政府對於現今文官的職能發展有所規劃與訓練，才能使文官的態度與認知能因應科技帶來的影響。

(三)才能管理

數位轉型與創新的過程中，除整備規範與組織之外，更重要是讓公務人員的職能可以與時俱進，尤其是數位職能的精進，從實驗調查、深度訪談與焦點座談可知組織內成員對於公共服務的數位創新有積極的態度與認知，會影響到新興科技運用於公部門是否能較為順利的因素之一，若文官能對於應用機器演算法於公共服務上，應知悉人機協作的態度價值、目前受到的限制以及執行的關鍵因素。人機協作的職能培養可從三大職能著手，分別為社會性職能、管理職能與專業職能。

第一，社會性職能。該職能包含三項核心職能，為社會性技能、倫理與安全技能以及創新性。主要強調跨域協調能力與組織成員具備之創新思維，協調能力包含政府內部的水平串聯以及對外與廠商的專業交流。對內，政府單位必須更具備跨領域溝通的對話能力，一方面傳達本身需求，另一方面能架接本身領域知識與外部資訊領域進行整合，以更進一步精確瞄準組織本身需要的服務內容；對外，政府單位與受託廠商因不同領域而產生彼此對話上的困難，因此需要提高自身的專業性才能縮咬與廠商之間的專業不對稱所造成的技術綁架的情形發生。

另外，創新性係指成員必須具備創新思維，但政府鼓勵成員具有創新思維的思維時，政府也面臨對於新興技術的容錯率不高的情形，原因出自政府對於民眾的回應性與課責機制的核心價值所起，然而創新與容錯率相互矛盾。創新思維的引導並且大膽引用任何可能的新興技術的作法，就必須擔負起機器出錯時所需要承擔的責任。此時，組織層次的管理體制與長官支持態度成為關鍵。

第二，管理職能。涵蓋有財務資源管理、原物料管理、人員管理、時間管理、資訊與知識管理及契約管理等。對於科技資訊素養的能力有待加強，政府內部在人員訓練上可依循新資訊科技的發展與引入，連帶增加相關教育訓練課程，可降低政府內部與外部廠商在合作上因專業與資訊不對稱，造成的不利影響。

第三，專業職能。包括系統技能、技術能力與解決複雜問題的能力。在現在多以委託廠商的模式運作之下，政府內部成員須思量系統運作後的後續維運，需要增強資訊與知識的管理能力，培育專業人力，在後續維運上因為自身的專業提升，有助於對於運用何種資料可提升運作效力的判斷力，後續資料分析上的數據解讀能力。

本計畫也另有針對法務部保護司觀護人的實驗調查設計，模擬演算系統的調查實驗平臺並納入實際的決策情境。為更瞭解目前觀護人的業務與工作流程，該調查採用質量並重的研究方法，在設計實驗問卷之前進行觀護人的訪談，深入探討實情並據此設計出貼近觀護人業務與工作流

程的實驗。該實驗將實驗參與者區分成實驗組與控制組。被分派到實驗組的參與者，在決策判斷過程會增加外部資訊介入。實驗組將有兩組，透過隨機分組，各組樣本將各有1/3。實驗組第一組（Arm 1）會接收到演算法預測系統所提供的風險機率與處遇建議。實驗組第二組（Arm 2）則加入其他參與此實驗觀護人的綜合評估分數作為外部資訊介入操弄。外部資訊介入都會在受訪者回答後呈現給受訪者看，接著再讓受訪者進行第二次作答，詢問他受訪者是否確認，是否要更改他原先的決策。並給他重新回答的機會。另外，也詢問他對於這份外部介入資訊的正確度有多高的信心。該實驗科正進行中，希冀在檢視有無演算系統的輔助下，觀護人對個案進行再犯評估風險的判斷及決策擬定後處理措施之差異，期能從中理解演算法應用對公部門實際的好處、顧慮、與決策困境等

第二節 政策建議

回到第一章所提及的Ostrom制度分析架構，我國如欲應用機器演算法於政府的業務推動與公共服務中，其建構過程會牽涉到此架構的三個層面，分別為攸關憲法選擇層次的議題，稱為正式法令的選擇，如：是否需要立法或是盤點現今相關法案進行修法；從集體選擇層次，以民主政治觀之，即在分權制衡的框架下的政策決策，代表的是政府中組織層次；而個人的層次，是在政策制定後的個人選擇，如：面對數位轉型引進新興科技協助業務執行時，我是否有勇於嘗試與創新的思維。三個層次都會互相影響並產生反饋。本計畫根據此理論架構以及前述之分析以及結論，經過深度訪談與焦點座談以及中央各部會的問卷調查，進行落差分析並提出改革路徑建議。接下來，本計畫將提出如欲如欲應用機器演算法於政府的業務推動與公共服務中的暫時性研究建議，分為短程、中程與長程：

一、 短程目標：盤點各部會對於機器演算法運用的規劃與投入資源

面對各領域的專業性程度較高以及多元性，各領域的主管機關勢必也會需要與該領域相關的部會合作。因此，需要先行盤點已經推動、規劃推動的部會目前有推動的相關計畫之執行原則、預算規劃、資料建置、處理的業務性質、執行流程、執行成效或預期成效、執行過程的限制或預期可能的限制，從前述面向進行盤點，才能得知目前運用機器演算法協助的業務性質多屬何種性質、制定那些相關執行原則、所需要的資源面有哪些等等。

盤點完成後，從集體選擇層次觀之，可以去釐清若部會要推動機器演算法運用應具備的條件，進而去改變組織的管理體制，型塑勇於創新且能給與組織成員依定程度保護的組織文化。

二、 中程目標：適合各部會遵循的共通性原則與執行指引建置

如欲發展機器演算法須建立共通性原則與執行指引，除可參考各國的制定原則與方向之外，更重要的是先行瞭解集體選擇層次之各部會所需，才能針對現今之需求與運作去盤點，因此短程目標完成資源盤點後，中程目標即是建立各部會可遵行的共通性原則與執行指引，從憲政選擇層次觀之，盤點已經推動的部會其執行指引上的共通性原則，進而設計適合我國中央各部會執行上的共通原則與指引，讓各部會依照各領域的專業性並且依據該原則設計適合該領域的遵行細則與執行指引。

三、 短程、中程、長程目標：公務人員的數位轉型素養養成與數位職能訓練

憲政選擇層次的規範建置、集體選擇層次的勇於接受創新的組織變革如欲順利推行，該基石即是個人選擇層次的數位轉型之數位轉型素養養成與職能訓練。當然，要培育數位專業人才之前，更需要釐清組織成員對於數位轉型的認知程度，以及對於勇於接受創新的接受程度，才能依照目前的情形設計出現階段能給予的職能教育訓練，將其所學帶進實際場域與實務結合，使訓練得以真正發揮作用。

參考文獻

一、 中文部分

- 丁玉珍、林子倫（2020）。人工智慧提升政府公共治理的應用潛力探討。檔案半年刊，19（2），24-41。
- 王偉霖（2019）。理財機器人對我國金融及相關法制的衝擊與發展。財金法學研究，2（3），388-390。
- 王儷玲（2019）。臺灣機器人理財須再進化。Smart 智富月刊，250。
- 曲建仲（2017）。翻轉人類未來的 AI 科技：機器學習與深度學習，科技新報，2021 年 04 月 40 日，取自：
<http://technews.tw/2017/10/05/ai-machine-learning-and-deep-learning/>。
- 朱體正（2016）。人工智能輔助刑事裁判的不確定性風險及其防範—美國威斯康星州訴米斯案的啟示。浙江社會科學，76-85。
- 何之行（2021）。死者健保資料再利用，釐清個資疑慮，2021 年 3 月 22 日，取自：<https://udn.com/news/story/7266/5334140>。
- 李建良（2019）。人工智慧與法律規範學術研究群。科技部人文社會科學研究中心學術研究群成果報告（編號：MOST 107-2420-H-002-007-MY3-SG10708），未出版。
- 阮怡凱、李仁豪（2020）。以巨量資料分析觀點探討受保護管束者再犯風險評估工具之研究。法務部 108 年委託研究計畫期末報告（編號：S1080305），未出版。
- 周亞倩（2018）。人工智慧時代下的司法變革—淺析司法官培訓於未來 20 年面臨之趨勢。司法新聲，126（5），53-70。
- 林其樺（2018）。英國國家醫療服務體系（NHS）公布國家資料退出（Opt-out）操作政策指導文件，資訊工業策進會，2021 年 10 月 29 日，取自：<https://stli.iii.org.tw/article-detail.aspx?no=64&tp=1&d=8119>。
- 張春興（2000）。教育心理學：三化取向的理論與實踐。臺北市：臺灣東華書局。
- 張智奇、連瑋鑫（2021）。人機協作對於人力需求及技能發展的影響及挑戰。臺灣勞工季刊，65，24-33。

- 陳安斌、陳莉貞、蘇秀玲、郭怡君（2018）。我國發展機器人理財顧問之研究。中華民國證券投資信託暨顧問商業同業公會委託報告。
- 陳邠婕、周宇輝、呂佑華（2020）。落差分析在運動中心經理人職能的應用。運動管理，48，16-28。
- 陳敦源（2016）。特邀專論：從 E 化、M 化、U 化到？化：電子化政府科技變革樂觀論的反思。文官制度季刊，8（4），1-19。
- 陳敦源（2019）。民主治理：公共行政與民主政治的制度性調和（第三版）。臺北：五南。
- 陳敦源、廖洲棚、黃心怡（2016）。政府公共溝通：新型態網路參與及溝通策略。國家發展委員會委託研究報告（編號：NDC-MIS-105-004），未出版。
- 彭禎伶（2020）。機器人理財也會出包 金管會糾出四大缺失，中時電子報，2021 年 10 月 20 日，取自：
<https://www.chinatimes.com/newspapers/20200310000253-260202?chdtv>。
- 曾更瑩、吳志光（2018）。人工智慧之相關法規國際發展趨勢與因應。行政院國家發展委員會委託研究計畫結案報告，未出版。
- 曾冠球、廖洲棚、黃心怡、陳敦源、何宗武（2019）。新興科技與公部門數位轉型：個案採擷與模式建構。國家發展委員會委託研究報告（編號：NDC-MIS-108-003），未出版。
- 黃俊能、鍾健雄、賴擁連、黃炳森、周煌智、吳慧菁、曾淑萍（2021）。開發建置受保護管束人再犯風險評估智慧輔助系統(1/3)-以巨量資料分析觀點探勘犯罪風險因子與保護管束再犯之關聯性。法務部 110 年委託研究計畫期末報告（編號：S1100421），未出版。
- 鄭伊廷（2021）。試析「一般資料保護規則」下自動化決策的解釋權爭議。經貿法訊，279，13-23。
- 鄭宏達（2021）。行政院明將拍板組改草案 數位發展部、科技部等先行，2021 年 3 月 24 日，取自：
https://money.udn.com/money/story/7307/5341032?from=udnamp_storysns_fb&fbclid=IwAR0UxCQDGHvqQITJ8q8Cuyv4FxdWQeqvemFpRU14IAMZL83wikyVCrfLukU。

- 鄭冠維（2013）。英國實行個人健康和社會照護資訊連結服務（care.data），資訊工業策進會，2021 年 10 月 29 日，取自：
<https://stli.iii.org.tw/article-detail.aspx?no=64&tp=1&d=6310>。
- 謝天和（2019）。企業人工智慧環境下，內部稽核核心能力重要性認知之調查研究。內部稽核，105，30-34。
- 鍾文雄（2019）。AI 智慧發展下，產業勞工應具備之職能。臺灣勞工季刊，59，34-41。

二、 外文部分

- 116th Congress, USA (2019). H.R.2575 - AI in Government Act of 2020. Retrieved July 15, 2021, from <https://www.congress.gov/bill/116th-congress/house-bill/2575/text>.
- 116th Congress, USA (2020). H.R.6216 - National Artificial Intelligence Initiative Act of 2020. Retrieved July 15, 2021, from <https://www.congress.gov/bill/116th-congress/house-bill/6216>.
- Administration of Donald J. Trump (2019). Exec. Order No. 13859 of Feb. 11, 2019, Maintaining American Leadership in Artificial Intelligence, 84 Fed. Reg. 3967. Retrieved July 15, 2021, from <https://www.govinfo.gov/content/pkg/DCPD-201900073/pdf/DCPD-201900073.pdf>.
- Administration of Donald J. Trump (2020). Exec. Order No. 13960 of Dec 3, 2020, Promoting the Use of Trustworthy Artificial Intelligence in the Federal Government. Retrieved July 15, 2021, from <https://trumpwhitehouse.archives.gov/presidential-actions/executive-order-promoting-use-trustworthy-artificial-intelligence-federal-government/>.
- Al Ariss, A., W. F. Cascio, & J. Paauwe (2014). Talent management: current theories and future research directions. *Journal of World Business*, 49(2), 173-179.
- Alshaw, S. & H. Alalwany (2009). E-government evaluation: Citizen's perspective in developing countries. *Information Technology for Development*, 15(3), 193-208.
- Anderson, J. (2009). Illusions of accountability: Credit and blame sense-making in public administration. *Administrative Theory & Praxis*, 31(3), 322-339.

- Andrews, L. (2019). Public administration, public leadership and the construction of public value in the age of the algorithm and big data. *Public Administration*, 97(2): 396-310.
- Angwin, J., J. Larson, S. Mattu, L. Kirchner, & ProPublica (2016). Machine Bias. Retrieved July 7, 2021, from <https://www.propublica.org/article/machine-bias-risk-assessments-in-criminal-sentencing>.
- Appleby, J. & A. Alvarez-Rosete (2003). The national health service keeping up with public expectations? In Park, A., J. Curtice, K. Thompson, L. Jarvis & C. Bromley (Eds.), *British social attitudes the twentieth report: Continuity and change over two decades* (pp. 29–44). London: Sage.
- Bannister, F., & R. Connolly (2020). Administration by algorithm: A risk management framework. *Information Polity*, 5(4), 471-490.
- Barth, T. J., & E. Arnold (1999). Artificial intelligence and administrative discretion: implications for public administration. *The American Review of Public Administration*, 29(4), 332-351.
- Bataller, C., & J. Harris (2016). Turning artificial intelligence into business value. Retrieved July 15, 2021, from https://www.accenture.com/t20160814T215045_w_us-en/_acnmedia/Accenture/Conversion-Assets/DotCom/Documents/Global/PDF/Technology_11/Accenture-Turning-Artificial-Intelligence-into-Business-Value.pdf.
- Bauer, A., D. Wollherr, and M. Buss (2008). Human–robot Collaboration: A Survey. *International Journal of Humanoid Robot*, 5(1), 47-66.
- Beiranvand, V., Hare, W., & Lucet, Y. (2017). Best practices for comparing optimization algorithms. *Optimization and Engineering*, 18(4), 815-848.
- Berger, J., M. H. Fisek, R. Z. Norman, & M. Zelditch Jr. (1977). Status characteristics and social interaction: An expectation-states approach. *Social Forces*, 56(2), 742-744.
- Boateng, R. (2013). *The challenge of taking baby steps in E-governance in West Africa*. The 1st UGBS Conference on Business and Development, Accra.

- Bovens M. & S. Zouridis(2002). From Street-Level to System-Level Bureaucracies: How Information and Communication Technology Is Transforming Administrative Discretion and Constitutional Control. *Public Administration Review*, 62(2), 174 – 184.
- Budd, K. (2019). Will artificial intelligence replace doctors? Retrieved July 10, 2021, from <https://www.aamc.org/news-insights/will-artificial-intelligence-replace-doctors>.
- Bullock, J. B. (2019). Artificial intelligence, discretion, and bureaucracy. *The American Review of Public Administration*, 49(7), 751-761.
- Burley, F.W. (1988). Monitoring biological diversity for setting priorities in conservation. In Wilson, E.O. & F.M. Peter (Ed.), *Biodiversity* (pp. 227-230). Washington (DC): National Academies Press.
- Busuioc, M. (2020). Accountable artificial intelligence: Holding algorithms to account. *Public Administration Review*. Retrieved July 10, 2021, from <https://onlinelibrary.wiley.com/doi/full/10.1111/puar.13293>.
- Cath, C., Wachter, S., Mittelstadt, B., Taddeo, M., & Floridi, L. (2018). Artificial intelligence and the ‘good society’: the US, EU, and UK approach. *Science and engineering ethics*, 24(2), 505-528.
- Charlwood, Andy. (2021) Artificial Intelligence and Talent Management, in *Digitalised Talent Management: Navigating the Human-Technology Interface* (Routledge Focus on Business and Management) , Sharna Wiblen, ed. New York: Routledge.
- Citron, D. K. (2007) . Technological Due Process. *Washington University Law Review*, 85(6), 1249-1313.
- Clune, J. (2019). AI-GAs: AI-generating Algorithms: An Alternate Paradigm for Producing General Artificial Intelligence. *arXiv*: 1905.10985.
- Cuccaro-Alamin, S., R. Foust, R. Vaithianathan, & E. Putnam-Hornstein (2017). *Risk assessment and decision making in child protective services: Predictive risk modeling in context*. Retrieved July 16, 2021, from <https://doi.org/10.1016/j.childyouth.2017.06.027>.
- Danaher (2016). 2016 annual report. Retrieved November 1, 2021, from: <https://investors.danaher.com/2016-Annual-Report/HTML1/tiles-twopage.htm>.

- Davis, F. D., Bagozzi, R. P., & Warshaw, P. R. (1989). User acceptance of computer technology: A comparison of two theoretical models. *Management science*, 35(8), 982-1003.
- Davis, F.D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319-340.
- Department for Business, Energy & Industrial Strategy (2017). *Industrial strategy white paper*. Retrieved July 15, 2021, from <https://www.gov.uk/government/publications/industrial-strategy-building-a-britain-fit-for-the-future>.
- Department for Business, Energy and Industrial Strategy & Department for Digital, Culture, Media & Sport, DCMS (2019). *AI Sector Deal*. Retrieved July 15, 2021, from <https://www.gov.uk/government/publications/artificial-intelligence-sector-deal/ai-sector-deal>.
- Department for Digital, Culture, Media & Sport, Central Digital and Data Office, Department for Business, Energy & Industrial Strategy, and Office for Artificial Intelligence (2020). *A guide to using artificial intelligence in the public sector*. Retrieved July 15, 2021, from <https://www.gov.uk/government/publications/a-guide-to-using-artificial-intelligence-in-the-public-sector>.
- Diakopoulos, N. (2015). Algorithmic accountability. *Digital Journalism*, 3, 398-415.
- Engstrom, David, Daniel Ho, Catherine Sharkey, and Mariano-Florentino Cuéllar (2020) . *Government by Algorithm: Artificial Intelligence in Federal Administrative Agencies*. Report submitted to the Administrative Conference of the United States, February 2020, from <https://wwwcdn.law.stanford.edu/wp-content/uploads/2020/02/ACUS-AI-Report.pdf>.
- European Commission (2020). *White paper on artificial intelligence - A European approach to excellence and trust*. Retrieved July 15, 2021, from https://ec.europa.eu/info/publications/white-paper-artificial-intelligence-european-approach-excellence-and-trust_en.
- European Commission (2021a). *Proposal for a Regulation Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act) and Amending Certain Union Legislative Acts*. Retrieved July 15,

- 2021, from <https://eur-lex.europa.eu/legal-content/EN/TXT/?qid=1623335154975&uri=CELEX%3A52021PC0206>
- European Commission (2021b). *Commission Staff Working Document, Impact Assessment Accompanying the Proposal for a Regulation of the European Parliament and of the Council, Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act) and Amending Certain Union Legislative Acts*. Retrieved April 21, 2021, from <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX%3A52021SC0084>.
- European Commission (2021c). *Commission Staff Working Document, Impact Assessment Accompanying the Proposal for a Regulation of the European Parliament and of the Council on machinery products*. Retrieved April 21, 2021, from <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:52021SC0082>.
- European Commission High-Level Expert Group on Artificial Intelligence (AI HLEG) (2019). *Ethics Guidelines for Trustworthy AI*. Retrieved July 15, 2021, from <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>.
- European Parliament (2017). *European Parliament resolution of 16 February 2017 with recommendations to the Commission on Civil Law Rules on Robotics* (2015/2103(INL)). Retrieved July 15, 2021, from https://www.europarl.europa.eu/doceo/document/TA-8-2017-0051_EN.pdf.
- European Parliament (2020). *European Parliament resolution of 20 October 2020 with recommendations to the Commission on a framework of ethical aspects of artificial intelligence, robotics and related technologies* (2020/2012(INL)). Retrieved July 15, 2021, from https://www.europarl.europa.eu/doceo/document/TA-9-2020-0275_EN.pdf.
- Executive Office of the President (2016). *Big Data: A Report on Algorithmic Systems, Opportunity, and Civil Rights*. Retrieved July 15, 2021, from https://obamawhitehouse.archives.gov/sites/default/files/microsites/ostp/2016_0504_data_discrimination.pdf.

- Flinders, M., & Dommett, K. (2013). Gap analysis: participatory democracy, public expectations and community assemblies in Sheffield. *Local Government Studies*, 39(4), 488-513.
- Foschi, M., G. K. Warriner, & S. D. Hart (1985). Standards, expectations, and interpersonal influence. *Social Psychology Quarterly*, 48(2), 108–117.
- Frankenfield J. (2021). *What Is a Robo Advisor and How Do They Work?*, Retrieved March 30, 2021, from <https://www.investopedia.com/terms/r/roboadvisor-roboadviser>.
- GDPR (2021). *General Data Protection Regulation*. Retrieved July 15, 2021, from <http://gdpr-info.eu/>.
- Gil-Garcia, J. R., Dawes, S. S., & Pardo, T. A. (2018). Digital government and public management research: finding the crossroads. *Public Management Review*, 20(5), 633-646.
- Gil-Garcia, J. R., Zhang, J., & Puron-Cid, G. (2016). Conceptualizing smartness in government: An integrative and multi-dimensional view. *Government Information Quarterly*, 33(3), 524-534.
- Government Office for Science (2016). *Artificial intelligence: Opportunities and implications for the future of decision making*. Retrieved July 15, 2021, from https://assets.publishing.service.gov.uk/government/uploads/system/uploads/attachment_data/file/566075/gs-16-19-artificial-intelligence-ai-report.pdf.
- Grundke, R. (Ed.). (2018). *Which skills for the digital era? Returns to skills analysis*. Retrieved July 15, 2021, from <https://dx.doi.org/10.1787/9a9479b5-en>.
- Guerin, B. (2018). *Accountability in modern government: What are the issues?* Retrieved July 10, 2021, from: <https://www.instituteFloridixgovernment.org.uk/sites/default/files/publications/IfG%20accountability%20discussion%20paper%20april%202018.pdf>.
- Gurumurthy, A., N. Chami, & D. Bharthur (2016). *Democratic accountability in the digital age. Bengaluru: IT for Change*, Retrieved July 7, 2021, from:

- <https://opendocs.ids.ac.uk/opendocs/bitstream/handle/20.500.12413/13010/Research-Brief-India-1.pdf?sequence=2&isAllowed=y>.
- Hair, J. F., Black, W. C., Babin. B.J., Anderson, R. E., & Tatham, R. L. (2006). **Multivariate Data Analysis** (6ed ed.). Upper Saddle River, NJ: Pearson Education.
- Harvard Law Review (2017). *State v. Loomis, Wisconsin Supreme Court requires warning before use of algorithmic risk assessments in sentencing*. Retrieved July 7, 2021, from <https://harvardlawreview.org/2017/03/state-v-loomis/>.
- Hoffman, G., & C. Breazeal (2004). *Collaboration in human-robot teams*. Retrieved July 7, 2021, from <https://doi.org/10.2514/6.2004-6434>.
- IMDA (The Infocomm Media Development Authority, Singapore) (2019). *Enabling data-driven innovation through trusted data sharing in a digital economy*, Retrieved November 1, 2021, from: <https://www.imda.gov.sg/news-and-events/Media-Room/Media-Releases/2019/Enabling-Data-Driven-Innovation-Through-Trusted-Data-Sharing-In-A-Digital-Economy>.
- IMDA (The Infocomm Media Development Authority, Singapore) and PDPC (Personal Data Protection Commission)(2020).*Model Artificial Intelligence Governance Framework*, 2nd ed., Retrieved October 31, 2021, from <https://www.pdpc.gov.sg/-/media/Files/PDPC/PDF-Files/Resource-for-Organisation/AI/SGModelAIGovFramework2.pdf>.
- Information Commissioner’s Office (ICO)(2019). *An overview of the Auditing Framework for Artificial Intelligence and its core components*, Retrieved September 20, 2021, from <https://ico.org.uk/about-the-ico/news-and-events/ai-blog-an-overview-of-the-auditing-framework-for-artificial-intelligence-and-its-core-components/>.
- Jankin, S., Esteve, M., & Campion, A. (2018). Artificial intelligence for the public sector: Opportunities and challenges of cross-sector collaboration: Mathematical. *Physical and Engineering Sciences. Philosophical Transactions of the Royal Society A*, 376(2128).
- Janssen, M., & M. Hartog, Matheus, R., A. Y. Ding & Kuk, G. (2020). *Will algorithms blind people? The effect of explainable AI and decision-makers’ experience on AI-supported decision-making in government*.

- Retrieved July 10, 2021, from <https://doi.org/10.1177/0894439320980118>.
- Jenning, M. (2000). Gap analysis: concepts, methods, and recent results. Retrieved November 2, 2021, from <https://doi.org/10.1023/A:1008184408300>.
- Johnson, T. J., and Wanta, W. (1996). Influence Dealers: A Path Analysis Model of Agenda Building During Richard Nixon's War on Drugs. *Journalism and Mass Communication Quarterly*, 73(1), 181-194.
- Jun, K. N. & C. Weare (2011). Institutional motivations in the adoption of innovations: The case of e-government. *Journal of Public Administration Research and Theory*, 21(3), 495-519.
- Kaiser, H. F. (1974). An Index of Factorial Simplicity. *Psychometrika*, 39(1), 31-36.
- Kanheman, D. & A. Tversky(1979). Prospect Theory: An Analysis of Decision under Risk. *Econometrica*, 47(2), 263-292.
- Konsynski, B., Bracker, L., & Bracker, W. (1982). A model for specification of office communications. *IEEE Transactions on communications*, 30(1), 27-36.
- Larson, J., S. Mattu, L. Kirchner, & J. Angwin (2016). *How we analyzed the COMPAS recidivism algorithm*. Retrieved July 7, 2021, from <https://www.propublica.org/article/how-we-analyzed-the-compas-recidivism-algorithm>.
- Leon, S., & L. Orrios (2019). *It's Westminster's fault*. Retrieved July 10, 2021, from <https://blogs.lse.ac.uk/politicsandpolicy/blame-attribution-in-devolved-systems/>.
- Lewis, R. E., & R. J. Heckman (2006). Talent management: A critical review. *Human Resource Management Review*, 16(2), 139-154.
- Makasi, T., Nili, A., Desouza, K., & Tate, M. (2020). Chatbot-mediated public service delivery: A public service value-based framework. *First Monday*, 25(12).
- Mikhaylov, S.J., M. Esteve & Campion, A. (2018). Artificial intelligence for the public sector: opportunities and challenges of cross-sector

- collaboration. Retrieved February 15, 2022, from <https://doi.org/10.1098/rsta.2017.0357>.
- Molitor, M. (2020). *Effective human-robot collaboration in the industry 4.0 context: Implications for human resource management*. Unpublished doctoral dissertation, University of Twente, Netherlands.
- Moore, Mark H. (2005). Break-through innovations and continuous improvement: Two different models of innovative processes in the public sector. *Public Money & Management*, 25(1): 43- 50.
- Naqvi, A. I. (2020). AI-government versus E-government: How to Reinvent Government with AI? in *Handbook of Artificial Intelligence and Robotic Process Automation: Policy and Government Applications*, AI Naqvi and J. Mark Munoz, eds. Anthem Press.
- NHS Digital, General Practice Data for Planning and Research: NHS Digital Transparency Notice(2021). Retrieved November 2, 2021, from: <https://digital.nhs.uk/data-and-information/data-collections-and-data-sets/data-collections/general-practice-data-for-planning-and-research/transparency-notice>.
- NHS Digital, NHS Digital Annual Report and Accounts 2019 to 2020(2020). Retrieved November 2, 2021, from: <https://digital.nhs.uk/about-nhs-digital/corporate-information-and-documents/nhs-digital-s-annual-reports-and-accounts/nhs-digital-annual-report-and-accounts-2019-20/annual-report-and-accounts>.
- NHS, Opt out of sharing your health records(2021). Retrieved November 2, 2021, from: <https://www.nhs.uk/using-the-nhs/about-the-nhs/opt-out-of-sharing-your-health-records/>.
- NHS, UK (2021). Opt out of sharing your health records, Retrieved October 25, 2021, from <https://www.nhs.uk/using-the-nhs/about-the-nhs/opt-out-of-sharing-your-health-records/>.
- O'Brien, K. (2012). Global environmental change II: From adaptation to deliberate transformation. *Progress in human geography*, 36(5), 667-676.
- OECD (2017), OECD Science, Technology and Industry Scoreboard 2017: The digital transformation. Retrieved July 15, 2021, from <https://dx.doi.org/10.1787/9789264268821-en>.

- OECD (2019). Recommendation of the Council on Artificial Intelligence. Retrieved July 15, 2021, from <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>.
- OECD (2021). The State of Implementation of the OECD AI Principles: Insights from national AI policies. Retrieved July 15, 2021, from <https://www.oecd.org/digital/state-of-implementation-of-the-oecd-ai-principles-1cd40c44-en.htm>.
- Oliver, R.L. (1977). Effect of expectation and disconfirmation on postexposure product evaluations: An alternative interpretation. *Journal of Applied Psychology*, 62(4), 480–486.
- Oram, A. (2016). If prejudice lurks among us, can our analytics do any better? Technical and policy considerations in combatting algorithmic bias. Retrieved July 12, 2021, from <https://www.oreilly.com/content/if-prejudice-lurks-among-us-can-our-analytics-do-any-better/>.
- Osei-Kojo, A. (2016). E-government and public service quality in Ghana. *Public Affairs*, 17(3), 1-12.
- Ostrom, E. (1986). A method of institutional analysis. In *Guidance, Control, and Evaluation in the Public Sector*, edited by F. X. Kaufman, G. Majone, and V. Ostrom. Berlin: de Gruyter.
- Oswald, M. (2018). Algorithm-assisted decision-making in the public sector: framing the issues using administrative law rules governing discretionary power. Retrieved July 10, 2021, from <https://doi.org/10.1098/rsta.2017.0359>.
- Parasuramen R. & D. Manzey(2010). Complacency and Bias in Human Use of Automation: An Attentional Integration. *Human Factors The Journal of the Human Factors and Ergonomics Society*, 52(3), 381-410.
- PDPC, Singapore. (2020). Model artificial Intelligence Governance Framework, 2nd. Ed. Retrieved January 15, 2022, from <http://go.gov.sg/ai-gov-mf-2>.
- Peeters R. (2020). The agency of algorithms: Understanding human-algorithm interaction in administrative decision-making. *Infomation Policy*, 25(1), 1-16.
- Perry, J. L. & L. R. Wise (1990). The motivational bases of public service. *Public Administration Review*, 50(3), 367-373.

- Perry, J. L. (1996). Measuring public service motivation: An assessment of construct reliability and validity. *Journal of Public Administration Research and Theory*, 6(1), 5-22.
- Personal Data Protection Commission (2018). Discussion Paper: Artificial Intelligence (AI) and Personal Data – Fostering Responsible Development and Adoption of AI. Retrieved July 15, 2021, from <https://www.pdpc.gov.sg/-/media/Files/PDPC/PDF-Files/Resource-for-Organisation/AI/Discussion-Paper-on-AI-and-PD---050618.pdf>.
- Plantera, F. (2017). Artificial intelligence is the next step for E-governance in Estonia State adviser reveals. Retrieved July 10, 2021, from <https://www.e-estonia.com/artificial-intelligence-is-the-next-step-for-e-governance-state-adviser-reveals/>.
- Polson, N., & J. Scott (2018). *AIQ: How artificial intelligence works and how we can harness its power for a better world* (Ed.). London: Bantam Press.
- Radaelli, Claudio. (2009). The Diffusion of Regulatory Impact Analysis: Best Practice or LessonDrawing? *European Journal of Political Research*, 16(8), 1145-1164.
- Rahwan, I. (2018) . Society-in-the-loop: Programming the Algorithmic Social Contract. *Ethics and Information Technology*, 20(1), 5-14.
- Reiling, A.D. (2020). Courts and artificial intelligence. *International Journal for Court Administration*, 11(2), p.8.
- Rossi, F. (2019). Building trust in artificial intelligence. *Journal of International Affairs*, 72(1), 127-134.
- Sandler, R. & J. Basl (2019). Building data and AI ethics committees. Retrieved July 15, 2021, from https://www.accenture.com/_acnmedia/PDF-107/Accenture-AI-And-Data-Ethics-Committee-Report-11.pdf#zoom=50.
- Scott, J.M., C.B. Kepler, P. Stine, & H. Little (1987). Protecting endangered forest birds in Hawaii: the development of a conservation strategy. Retrieved November 1, 2021, from: https://www.researchgate.net/publication/321265722_Protecting_endangered_forest_birds_in_Hawaii_the_development_of_a_conservation_strategy.

- Sharma, T. (2019). How artificial intelligence can improve e-governance services? Insight and resources: global tech council. Retrieved July 10, 2021, from <https://www.globaltechcouncil.org/artificial-intelligence/how-artificial-intelligence-can-improve-e-governance-services/>.
- Shepsle, K. A. (2006). Rational choice institutionalism. In R. A. W. Rhodes (Ed.), *The Oxford Handbook of Political Institutions*, (pp. 23-38). New York: Oxford University Press.
- Singapore PDPC (2020). Model AI Governance Framework, Personal Data Protection Commission. Retrieved July 15, 2021, from <https://www.pdpc.gov.sg/Help-and-Resources/2020/01/Model-AI-Governance-Framework>.
- SNDGO (The Smart Nation and Digital Government Office, Singapore) (2019). National Artificial Intelligence Strategy, Retrieved November 1, 2021, from: <https://www.smartnation.gov.sg/files/publications/national-ai-strategy.pdf>.
- Squicciarini, M. & H. Nachtigall (2021). Demand for AI skills in jobs evidence from online job postings. Retrieved July 11, 2021, from <https://www.oecd.org/digital/demand-for-ai-skills-in-jobs-3ed32d94-en.htm>.
- Starling, G. 7th ed. (2005). *Managing the Public Sector*, Harcourt, Fort Worth, TX.
- The White House's Office of Science and Technology Policy (2020). Guidance for Regulation of Artificial Intelligence Applications. Retrieved July 15, 2021, from <https://www.whitehouse.gov/wp-content/uploads/2020/01/Draft-OMB-Memo-on-Regulation-of-AI-1-7-19.pdf>.
- Thierer, A., A. O'Sullivan, and R. Russell (2017). *Artificial Intelligence and Public Policy*. Arlington, VA: Mercatus Center George Mason University. Please check: <https://www.mercatus.org/system/files/thierer-artificial-intelligence-policy-mrmercatus-v1.pdf>.

- Twizeyimana, J.D. & A. Andersson (2019). The public value of E-government: A literature review. *Government Information Quarterly*, 36(2), 167-178.
- UK AI Council (2021). AI Roadmap. Retrieved July 15, 2021, from https://assets.publishing.service.gov.uk/government/uploads/system/uploads/attachment_data/file/949539/AI_Council_AI_Roadmap.pdf.
- Van Den Bosch, Eric. (2017). Human-machine decision-making and trust. In S. R. White (Ed.), *Closer than You Think: the Implications of the Third Offset Strategy for the U.S. Army* (pp. 109-119). U.S.: Strategic Studies Institute.
- Van Kleek, M., M. Veale, & R. Binns (2018). Fairness and accountability design needs for algorithmic support in high-stakes public sector decision-making. Retrieved July 10, 2021, from <https://doi.org/10.1145/3173574.3174014>.
- Vogl, T.M., C. Seidelin, B. Ganesh, & J. Bright (2020). Smart technology and the emergence of algorithmic bureaucracy: Artificial intelligence in UK local authorities. *Public Administration Review* 80(6), 946-961.
- Wagner, D. G., & J. Berger (1993). Status characteristics theory: The growth of a program. In Berger, J. & M. Zelditch Jr. (Eds.), *Theoretical research programs: studies in the growth of theory* (pp. 23–63). California:Stanford University Press.
- Warta, S. F., K. A. Kapalo, A. Best, & S. M. Fiore (2016). Similarity, complementarity, and agency in HRI: Theoretical issues in shifting the perception of social robots from tools to teammates. *Proceedings of the Human Factors and Ergonomics Society Annual Meeting*, 60(1), 1230-1234.
- Welch E.W., C. Hinnat & Moon, J.(2005). Linking Citizen Satisfaction with E-Government and Trust in Government. *Journal of Public Administration Research and Theory*, 15(3), 371–391.
- Williams, S. P., & C. Hardy (2005). Public e-procurement as socio-technical change. *Strategic Change*, 14(5), 273-281.
- Williamson, O. E. (1980). The Organization of Work: A Comparative Institutional Assessment. *Journal of Economic Behavior and Organization*, 1(1), 5-38.

- Winch G., A. Usman & Edkins A.(1998). Towards total project quality: a gap analysis approach. *Construction Management and Economics*, 16(2), 193-207.
- Wirtz, B. W., J. C. Weyerer, & B. J. Sturm (2020). The dark sides of artificial intelligence: An integrated AI governance framework for public administration. *International Journal of Public Administration*, 43(9), 818-829.
- Wu, T. (2003). Network Neutrality and Broadband Discrimination. *Journal on Telecommunications & High Technology Law*, 2, 141-179..
- Zarsky T. (2016). The Trouble with Algorithmic Decisions: An Analytic Road Map to Examine Efficiency and Fairness in Automated and Opaque Decision Making. *Science, Technology, & Human Values*, 41(1), 118-132.
- Zarsky, T. (2011). Governmental data mining and its alternatives. *Penn State Law Review*, 116(2), 285-330.
- Zeithaml, V., Parasuraman, A. & Berry, L. (1988). ‘SERVQUAL’: A Multiple Item Scale for Measuring Consumer Perceptions of Service Quality. *Journal of Retailing*, 64(1): 12-40.
- Zuboff, S. (2015). Big other: surveillance capitalism and the prospects of an information civilization. Retrieved November 1, 2021, from <https://doi.org/10.1057/jit.2015.5>.

附錄

附錄一、中央各部會應用演算法之調查問卷

一、第一階段調查

(一) 開頭語

親愛的受訪者，您好！

首先感謝您撥冗填寫此份問卷！本計畫由國發會委託，目的為瞭解貴機關導入或規劃導入機器演算法（人工智慧）輔助業務及創新公共服務的狀況，藉此提供未來政策的改進建議。調查時間自 110 年 8 月 25 日開始至 9 月 3 日結束，採網路問卷方式進行調查。本問卷填答所需時間約 10-15 分鐘，請您依據貴機關的現況填答，透過本調查所得到的資料僅作為學術報告之統計分析使用不會作為其他用途，且保證採填答者匿名方式分析資料，請您放心填寫。完成填答後，本中心會贈送政治大學紀念品表達誠摯的感謝，如有任何疑問，請透過以下方式洽詢本案研究助理（電話接聽時間為週一至週五 10:00-12:00/14:00-17:00）。

臺灣數位治理研究中心，電話 (02)2939-3091 分機 51155

電郵：cyhwork1@gmail.com

研究助理：陳郁函（政治大學公共行政學系碩士生）

敬祝

身體健康、平安順心

國家發展委員會委託研究單位：臺灣數位治理研究中心

計畫主持人：陳敦源（政治大學公共行政學系教授）

協同主持人：廖洲棚（空中大學公共行政學系副教授）

黃心怡（臺灣大學公共事務研究所副教授）

王千文（淡江大學公共行政學系助理教授）

張濱璿（臺北醫學大學醫療暨生物科技法律研究所兼任助理教授）

敬啟

(二) 第一階段問卷內容

隨著網路通訊技術的發展，人工智慧(Artificial Intelligence，簡稱 AI)科技的精進及帶動的智慧應用市場，已是不可忽視的發展趨勢。行政院將民國 106 年定位為臺灣 AI 元年，正積極推動臺灣 AI 行動計畫，落實數位國家發展戰略。本計畫將應用 AI 科技於公共服務的電腦資訊系統（以下簡稱 AI 系統）定義為一種結合數位資料、機器學習演算法以及人類經驗而成的一種全自動化或半自動化輔助業務或創新公共服務的電腦資訊系統。為瞭解貴機關當前導入或規劃 AI 系統輔助業務及創新公共服務的實際情況，請您代表貴機關回答以下的問題：

1-1 請問貴機關目前是否有使用 AI 系統來輔助業務？

☐ (1)是、☐ (2)否（跳答第 1-12 題）

1-2 請問貴機關目前使用的 AI 系統名稱？（請寫出目前所有已上線使用的 AI 系統名稱）（開放題）

1-3 請問貴機關使用 AI 系統輔助下列哪些業務？（可複選）

☐ (1)操作機器的自動化 ☐ (2)辦公室行政文書工作

☐ (3)與服務對象溝通工作

☐ (4)與同事間協調工作 ☐ (5)聘僱員工 ☐ (6)工作績效評估

☐ (7)策略建立或政策規劃 ☐ (8)目標與願景設定

☐ (9)其他創新公共服務_____（請說明）

1-4 請以 500 字為度，根據前述填答的系統以及業務項目，描述貴機關目前使用 AI 系統輔助哪些業務以及累計至目前的實施績效？（開放題）

1-5 請問貴機關最早在哪一年開始導入 AI 系統輔助前述業務？

民國_____年

1-6 請問貴機關自民國 107 年迄今每年編列多少預算採購或發展 AI 系統？（說明臺灣 AI 行動計畫的年限為 107-110 年）

(1)107 年度_____千元（未編列預算請填 0 元）

(2)108 年度_____千元（未編列預算請填 0 元）

(3)109 年度_____千元（未編列預算請填 0 元）

(4)110 年度_____千元（未編列預算請填 0 元）

1-7 請問貴機關採用何種途徑建置 AI 系統？ ☐ (1)自製 ☐ (2)委外

1-8 請問貴機關建置的 AI 系統使用下列何種數位資料？（可複選）

☐ (1)本機關開放資料 ☐ (2)本機關內部資料

- ☐ (3) 他機關開放資料 ☐ (4) 他機關內部資料
☐ (5) 私部門資料 ☐ (6) 其他_____（請說明）

1-9 請問貴機關在導入及使用 AI 系統的過程中，是否有進行相應的法規調適？

- ☐ (1) 有 ☐ (2) 否（跳答第 1-11 題）

1-10 請以 300 字為度，描述貴機關現階段已完成或預計完成哪些法規調適？（開放題）（跳答第 1-24 題）

1-11 請問貴機關導入 AI 系統未進行相應的法規調適的理由？（跳答第 1-24 題）

- ☐ (1) 目前法令完備，無法規調適必要
☐ (2) 目前為試辦階段，無法規調適必要
☐ (3) 目前為實驗階段，無法規調適必要
☐ (4) 非主管機關，無法規調適權限
☐ (5) 還未知要進行哪些法規調適
☐ (6) 其他_____（請說明）

1-12 請問貴機關是否正在規劃 AI 系統來輔助業務？

- ☐ (1) 是 ☐ (2) 否（跳答 1-23）

1-13 請問貴機關規劃中的 AI 系統名稱？（請寫出規劃中的 AI 系統或計畫名稱）（開放題）

1-14 請問貴機關規劃未來使用 AI 系統來輔助下列哪些業務？（可複選）

- ☐ (1) 操作機器的自動化 ☐ (2) 辦公室行政文書工作
☐ (3) 與服務對象溝通工作
☐ (4) 與同事間協調工作 ☐ (5) 聘僱員工 ☐ (6) 工作績效評估
☐ (7) 策略建立或政策規劃 ☐ (8) 目標與願景設定
☐ (9) 其他創新公共服務_____（請說明）

1-15 請以 500 字為度，根據前述填答的系統（或計畫名稱）以及業務項目，描述貴機關規劃使用 AI 系統輔助哪些業務以及預期的實施績效？（開放題）

1-16 請問貴機關規劃在哪一年開始導入 AI 系統輔助業務？

民國_____年

1-17 請問貴機關自民國 107 年迄今每年編列多少預算採購或發展 AI 系統？

(說明臺灣 AI 行動計畫的年限為 107-110 年)

(1)107 年_____千元 (未編列預算請填 0 元)

(2)108 年_____千元 (未編列預算請填 0 元)

(3)109 年_____千元 (未編列預算請填 0 元)

(4)110 年_____千元 (未編列預算請填 0 元)

1-18 請問貴機關規劃採用何種途徑建置 AI 系統？ ☐ (1)自製 ☐ (2)委外

1-19 請問貴機關規劃建置的 AI 系統將使用下列何種數位資料？(可複選)

☐ (1)本機關開放資料

☐ (2)本機關內部資料

☐ (3)他機關開放資料

☐ (4)他機關內部資料

☐ (5)私部門資料

☐ (6)其他_____ (請說明)

1-20 請問貴機關在規劃 AI 系統時，是否有考慮進行相應的法規調適？

☐ (1)有

☐ (2)否 (跳答第 1-22 題)

1-21 請以 300 字為度，描述貴機關已完成或預計完成哪些法規調適？
(開放題) (跳答第 1-24 題)

1-22 請問貴機關規劃導入 AI 系統未考慮進行相應的法規調適的理由？(跳答第 1-24 題)

☐ (1)目前法令完備，無法規調適必要

☐ (2)非主管機關，無法規調適權限

☐ (3)還未知要進行哪些法規調適

☐ (4)其他_____ (請說明)

1-23 請問貴機關未規劃導入 AI 系統輔助業務的理由為何？(可複選)

☐ (1)缺乏預算 ☐ (2)缺乏專業人力 ☐ (3)缺乏配套法令

☐ (4)缺乏適用技術 ☐ (5)缺乏數位資料 ☐ (6)非機關權責業務

☐ (7)缺乏成本效益 ☐ (8)機關首長未指示推動

☐ (9)其他_____ (請說明)

1-24 請問貴機關的名稱_____ (請說明)

1-25 請問貴機關目前有多少人力專責資訊業務？_____人（包括正式公務人員、約聘僱人力、派遣人力等等都算）

1-26 承上題，其中參與規劃、採購或管理 AI 系統專案的資訊人力有多少人？_____人（若貴機關尚未規劃導入 AI 系統請填 0 人）

1-27 承上題，其中為正式公務人員的有多少人？_____人

二、 第二階段調查

(三) 開頭語

親愛的受訪者，您好！

首先感謝您撥冗填寫此份問卷！本計畫由國發會委託，目的在瞭解您對於所屬機關導入人工智慧系統（簡稱 AI 系統）輔助業務及創新公共服務的看法，以協助提供未來的相關政策建議。本調查採網路問卷進行，調查時間自 110 年 9 月 3 日開始至 9 月 24 日結束。問卷填答所需時間約 10-15 分鐘，請您依據個人的觀點填寫，透過本調查所得到的資料僅作為學術報告之統計分析使用不會作為其他用途，且保證採填答者匿名方式分析資料，請您放心填寫。完成填答後，本中心會贈送政治大學紀念品表達誠摯的感謝。如有任何疑問，請透過以下方式洽詢本案研究助理（電話接聽時間為週一至週五 10:00-12:00/14:00-17:00）。

臺灣數位治理研究中心，電話 (02)2939-3091 分機 51155

電郵：cyhwork1@gmail.com

研究助理：陳郁函（政治大學公共行政學系碩士生）

敬祝

身體健康、平安順心

國家發展委員會委託研究單位：臺灣數位治理研究中心

計畫主持人：陳敦源（政治大學公共行政學系教授）

協同主持人：廖洲棚（空中大學公共行政學系副教授）

黃心怡（臺灣大學公共事務研究所副教授）

王千文（淡江大學公共行政學系助理教授）

張濱璿（臺北醫學大學醫療暨生物科技法律研究所兼任助理教授）

敬啟

(四) 第二階段問卷內容

隨著網路通訊技術的發展，人工智慧(Artificial Intelligence，簡稱 AI) 科技的精進及帶動的智慧應用市場，已是不可忽視的發展趨勢。行政院將民國 106 年定位為臺灣 AI 元年，正積極推動臺灣 AI 行動計畫，落實數位國家發展戰略。本計畫將應用 AI 科技於公共服務的電腦資訊系統（以下簡稱 AI 系統）定義為一種結合數位資料、機器學習演算法以及人類經驗而成的一種全自動化或半自動化輔助業務或創新公共服務的電腦資訊系統。本中心依據前一階段的調查得知貴機關正在使用（或規劃使用或尚未使用）AI 系統輔助業務及創新公共服務的現況，並經由貴機關推薦您填答本問卷，請您依據個人的想法以及機關目前的實際現況，回答下列的問題。

個人對機關導入 AI 系統的看法（機器演算法---AI 系統）

當您所屬的機關導入 AI 系統之後，請問您對下列陳述的同意程度為何？

題號	題項內容	非常不同意	不同意	有點不同意	有點同意	同意	非常同意	不知道
2-1	AI 系統能幫助我更快速地完成任務。	①	②	③	④	⑤	⑥	98
2-2	AI 系統能幫助我更正確地完成任務。	①	②	③	④	⑤	⑥	98
2-3	AI 系統能幫助我更節省行政作業成本。	①	②	③	④	⑤	⑥	98
2-4	AI 系統能幫助我更公平地完成任務。	①	②	③	④	⑤	⑥	98
2-5	AI 系統會增加我的業務量。	①	②	③	④	⑤	⑥	98
2-6	AI 系統產生的錯誤會由我承擔責任。	①	②	③	④	⑤	⑥	98
2-7	AI 系統會限縮我的行政裁量權力。	①	②	③	④	⑤	⑥	98
2-8	AI 系統讓長官更容易掌握我的工作情況。	①	②	③	④	⑤	⑥	98
2-9	AI 系統讓我更容易回應民眾需求。	①	②	③	④	⑤	⑥	98
2-10	我認為 AI 系統可以減少人工作業程序。	①	②	③	④	⑤	⑥	98
2-11	我擔心 AI 系統會取代我的工作。	①	②	③	④	⑤	⑥	98
2-12	我擔心 AI 系統會帶來未知的風險。	①	②	③	④	⑤	⑥	98

個人工作任務認知

隨著科技的發展，運用科技進行行政革新的情況將愈來愈常見，請問您對以下相關陳述的同意程度為何？

題號	題項內容	非常不同意	不同意	有點不同意	有點同意	同意	非常同意
2-13	我很期待科技輔助的未來生活。	①	②	③	④	⑤	⑥
2-14	我總覺得科技應用對我帶來的好處多過於壞處。	①	②	③	④	⑤	⑥
2-15	我相信科技會帶給我很多好事。	①	②	③	④	⑤	⑥
2-16	如果我發現自己陷入困境，我能想出很多方法來擺脫它。	①	②	③	④	⑤	⑥
2-17	我能想到很多的方法來達到我目前的目標。	①	②	③	④	⑤	⑥
2-18	到目前為止，我認為我自己還蠻成功的。	①	②	③	④	⑤	⑥
2-19	我有信心可以有效地處理突發事件	①	②	③	④	⑤	⑥
2-20	只要我付出必要的努力，我可以解決大多數的問題。	①	②	③	④	⑤	⑥
2-21	面對困難我可以保持冷靜，因為我可以依靠自己的應對能力。	①	②	③	④	⑤	⑥
2-22	我認為我自己是一個可以承受很多壓力的人。	①	②	③	④	⑤	⑥
2-23	失敗並不會讓我覺得氣餒。	①	②	③	④	⑤	⑥
2-24	在遭遇嚴重的生活難題後，我往往可以迅速恢復。	①	②	③	④	⑤	⑥

做為一位公務員，請問您對以下陳述的同意程度為何？

題號	題項內容	非常不同意	不同意	有點不同意	有點同意	同意	非常同意
2-25	為大眾服務是很重要的事情。	①	②	③	④	⑤	⑥
2-26	為社會做出貢獻比得到個人成就更有意義。	①	②	③	④	⑤	⑥
2-27	社會上人與人之間是互相需要的。	①	②	③	④	⑤	⑥
2-28	為社會公益而犧牲自己的權益沒有關	①	②	③	④	⑤	⑥

	係。						
2-29	就算是會惹事上身，但是路見不平拔刀相助他人還是必要的。	①	②	③	④	⑤	⑥

資料準備度

請依據您的所屬機關擁有 AI 系統所需資料準備度(包括資料完整度、機器可讀程度、資料正確性、資料近用度以及資料使用歷程紀錄完整度)的認知情況，回答下列問題。

題號	題項內容	非常不同意	不同意	有點不同意	有點同意	同意	非常同意	不知道
2-30	我所屬的機關擁有開發 AI 系統所需的全部資料與數據庫。	①	②	③	④	⑤	⑥	98
2-31	我所屬的機關擁有可轉換為機器可讀格式的資料與數據庫。	①	②	③	④	⑤	⑥	98
2-32	我所屬的機關有完整的資料治理規範，可確保資料與數據庫的內容正確性。	①	②	③	④	⑤	⑥	98
2-33	在符合現有的作業規範下，我所屬機關擁有的資料與數據庫可以開放外部專家使用來開發 AI 系統。	①	②	③	④	⑤	⑥	98
2-34	我所屬的機關有完整的資料使用紀錄，包括從資料的源頭收集到系統應用端的所有使用過程。	①	②	③	④	⑤	⑥	98

組織管理現況

請根據您對所屬機關之內部管理系統（泛指內部所有領導指揮和行政作業流程所構成的系統）現況的觀察，回答下列問題。

題號	題項內容	非常不同意	不同意	有點不同意	有點同意	同意	非常同意
2-35	我所屬機關的管理系統能夠良好的支持	①	②	③	④	⑤	⑥

	組織的整體目標。						
2-36	我所屬機關的管理系統總是讓我們浪費力氣在許多沒有生產力的工作上。	①	②	③	④	⑤	⑥
2-37	我所屬機關的成員總是在處理多重且互相矛盾目標的工作。	①	②	③	④	⑤	⑥
2-38	我所屬機關的管理系統會鼓勵大家挑戰傳統與過去不可撼動的作法。	①	②	③	④	⑤	⑥
2-39	我所屬機關的管理系統會給予成員足夠的彈性來因應變動的社會環境。	①	②	③	④	⑤	⑥
2-40	我所屬機關的管理系統以政策一致和穩定為主要目標。	①	②	③	④	⑤	⑥

人力職能

在政府機關導入 AI 系統輔助業務或創新公共服務的過程中，請問您對於下列有關人才職能之相關陳述的同意程度為何？

A 欄	請根據您的觀察與體悟，在下列政府運用 AI 系統所應具備的各項能力中，填答您認為該項能力的重要性程度？	寫下 1-6 的分數，1 代表非常不重要，6 代表非常重要。
B 欄	同時，請您評估自己目前具備下列各項能力的程度？	寫下 1-6 的分數，1 代表完全不具備，6 代表完全具備。

題號	題項內容	A 重要性	B 具備程度
2-41	具備能與他人共事，以解決 AI 系統所涉及跨域問題的能力。		
2-42	具備能確保安全與合乎道德規範使用 AI 系統的能力。		
2-43	具備有策劃創造性解決方案以提供新產品或服務的能力。		
2-44	為能有效執行 AI 系統導入方案，具備相關資源分配的能力(如人力、財務、知識、契約管理等項)。		

2-45	具備有判讀與監測 AI 系統輸出的能力。		
2-46	具備有維運 AI 系統的能力。		
2-47	具備能在複雜現實環境中運用大數據分析，解決不明確問題的能力。		

行政革新情況

在過去三年中，您所屬的機關為了改善現有業務或創新公共服務，是否作過以下的改變？

題號	題項內容	完全沒有	有考慮但沒採納	有推動過	不知道
2-48	採用過去沒有的全新政策或全新的公共服務	①	②	③	98
2-49	導入過去沒有的全新的行政或服務流程。	①	②	③	98
2-50	實驗/試驗某個過去沒有的新政策或技術作法。	①	②	③	98
2-51	改善現有政策。	①	②	③	98
2-52	改善現有行政/服務流程。	①	②	③	98
2-53	改善現有公共服務。	①	②	③	98

受訪者基本資料

2-54 請問您的生理性別為？

☐ (0) 女性

☐ (1) 男性

2-55 請問您的年齡是？

☐ (1) 20-24 歲 ☐ (2) 25-29 歲 ☐ (3) 30-34 歲

☐ (4) 35-39 歲 ☐ (5) 40-44 歲 ☐ (6) 45-49 歲

☐ (7) 50-54 歲 ☐ (8) 55-59 歲 ☐ (9) 60 歲以上

2-56 請問您的最高教育程度為：

☐ (1) 高中（職）以下 ☐ (2) 專科 ☐ (3) 大學

☐ (4) 碩士 ☐ (5) 博士 ☐ (6) 其他_____ (請註明)

2-57 請問您現職的官等為何？

☐ (1) 簡任_____等 ☐ (2) 薦任_____等 ☐ (3) 委任_____等

☐ (4) 其他_____ (請說明)

2-58 請問您的職位為？ ☐ (1) 主管 ☐ (2) 非主管

2-59 請問您目前的職務是隸屬於貴部(會)以下哪一種單位？

☐ (1) 業務單位 ☐ (2) 資訊單位 ☐ (3) 人事單位 ☐ (4) 其他_____ (請說明)

(若您不隸屬前述部會的三個單位，請填寫您所屬機關及單位名稱。)

2-60 請問您任職於公務部門總年資共幾年：_____年

2-61 請問您在目前機關的服務年資共幾年：_____年

2-62 最後，若你對目前政府推動數位化政策有任何想法與建議，歡迎於下面空白欄留言。(開放題)

附錄二、法務部保護司觀護人實驗調查問卷

國立政治大學數位治理研究中心

機器學習演算輔助系統調查

親愛的觀護人，您好！

首先感謝您撥冗填寫此份問卷！本實驗調查由國發會委託研究並與法務部保護司共同合作，其目的是瞭解地檢署觀護人對於機器演算人工智慧系統用以輔助公務上決策過程的看法、態度、信任度等不同構面。本調查包含須完成四個受保護管束案件的判斷與問卷調查；另外，題目將分為四大部分，填答時間預計需花費 30 分鐘。

本次研究採網路問卷的方式進行，每位參與者都有獨立之調查連結，確保您填答內容之資安與隱匿性。所有蒐集的資料無法辨識其個人身份，請放心填答。透過本調查所得到之資料，將加總後以統計資料方式呈現，僅作為學術報告之統計分析使用不會作為其他用途，本案等相關研究人員的報告將以匿名方式來呈現加總後之統計資料且去識別化，以確保研究參與者的個人權益，請您放心填寫。此外，本團隊將提供 7-11 300 元電子禮券作為贈禮，對您願意撥冗填寫此問卷表示答謝。再次感謝您提供寶貴的意見以及對本計畫的支持。

敬祝

身體健康、平安順心

國家發展委員會委託研究單位：臺灣數位治理研究中心

計畫主持人：陳敦源教授

協同主持人：廖洲棚副教授、黃心怡副教授、王千文助理教授、張濱璿兼任助理教授

若對調查有任何問題，請聯絡本計畫聯絡人：

陳郁函研究助理（02）2939-3091#51155

郭彥廷研究助理 110256016@g.nccu.edu.tw

關於本學術實驗與研究的資料與後續運用請參考實驗調查研究參與者知情同意說明，閱讀確認後，請勾選下方選項後開始。

一、實驗前問卷題目

首先想請教您一些個人基本資料

A1. 年齡：

- ☐ (01) 18~24 歲 ☐ (02) 25~29 歲 ☐ (03) 30~34 歲
☐ (04) 35~39 歲 ☐ (05) 40~44 歲 ☐ (06) 45~49 歲
☐ (07) 50~54 歲 ☐ (08) 55~59 歲 ☐ (09) 60~64 歲
☐ (10) 65 歲以上

A2. 最高教育程度：

- ☐ (01) 技術學院 ☐ (02) 科技大學、大學與其同等學歷
☐ (03) 碩士 ☐ (04) 博士
☐ (05) 其他（請在下方說明）

A3. 請問您最高學歷之學科背景

- ☐ (01) 法律相關 ☐ (02) 社會工作相關 ☐ (03) 心理、諮商相關
☐ (04) 教育與犯罪防治 ☐ (05) 其他社會科學領域
☐ (06) 工程科學相關 ☐ (07) 自然科學相關
☐ (08) 其他（請於下方說明）

接下來的題目請就您的工作經驗，回答您目前的工作狀況，與「檢察機關案件管理系統」的使用情形。

A4.1 請問您近半年來執行工作業務時，需要使用到「案管系統與相關資訊系統」的頻率有多高？

- ☐ (01) 非常低 ☐ (02) 低 ☐ (03) 有點低
☐ (04) 有點高 ☐ (05) 高 ☐ (06) 非常高

A4.2 請依據您的使用經驗，評估組織目前使用之「案管系統」的易用程度？

- ☐ (01) 非常不容易 ☐ (02) 不容易 ☐ (03) 有點不容易

- ☐ (04) 有點容易 ☐ (05) 容易 ☐ (06) 非常容易

A4.3 相較於其他同仁，您覺得過去半年您所承擔的工作量為何？

- ☐ (01) 非常少 ☐ (02) 少 ☐ (03) 有點少
☐ (04) 有點多 ☐ (05) 多 ☐ (06) 非常多

A4.4 相較於其他同仁，您覺得過去半年您的工作效率為何？

- ☐ (01) 非常低 ☐ (02) 低 ☐ (03) 有點低
☐ (04) 有點高 ☐ (05) 高 ☐ (06) 非常高

A4.5 相較於其他同仁，您覺得過去半年您的工作壓力有多大？

- ☐ (01) 非常小 ☐ (02) 小 ☐ (03) 有點小
☐ (04) 有點大 ☐ (05) 大 ☐ (06) 非常大

接下來想瞭解您對於人工智慧 (AI) 應用於您現今業務工作的想法。

以下為 AI 的簡單介紹與相關應用的說明。

人工智慧 (Artificial Intelligence, AI) 由電腦科學結合巨量資料之科技，使機器能作出如人類般的行為、學習及思考模式。機器學習 (Machine Learning) 為人工智慧學習的方式，透過讀取大量的數據資料以認識及判別其特性，並運用機器演算法作出決策。

目前將 AI 應用於檢調業務的實例不少，如美國已在如紐約、威斯康辛等州導入的知名 COMPAS 系統 (Correctional Offender Management Profiling for Alternative Sanctions，受刑人接受替代性懲罰的管理評估)，透過歷史數據分析被告家庭背景、犯罪經歷、身心狀況等資料，評估被告未來再犯的可能性，而根據研究，此系統的正確率比犯

若未來法務部要開發並導入新的「再犯風險評估機器演算之人工智慧系統」，(以下簡稱為 Parole-ML-AI 1.0)，藉以輔助您的業務，提高工作效率。在此情境下，請試著想像該系統對您未來工作造成的影響，並就以下題目勾選適合的選項。

A5.1 該系統將對我工作有所幫助。

- ☐ (01) 非常不同意 ☐ (02) 不同意 ☐ (03) 有點不同意
☐ (04) 有點同意 ☐ (05) 同意 ☐ (06) 非常同意

A5.2 未來我與該系統的合作會需要額外成本(例:時間成本、適應成本)。

- ☐ (01) 非常不同意 ☐ (02) 不同意 ☐ (03) 有點不同意
☐ (04) 有點同意 ☐ (05) 同意 ☐ (06) 非常同意

A5.3 未未來若讓該系統輔助(如:提供一些可參考的綜合資訊)我目前的
觀護人工作，我的工作效率會提升。

- ☐ (01) 非常不同意 ☐ (02) 不同意 ☐ (03) 有點不同意
☐ (04) 有點同意 ☐ (05) 同意 ☐ (06) 非常同意

A6. 承上題，若「Parole-ML-AI 1.0」系統被導入且用以輔助您目前「再
犯風險」評估的工作，您的支持程度為何？

- ☐ (01) 非常不同意 ☐ (02) 不同意 ☐ (03) 有點不同意
☐ (04) 有點同意 ☐ (05) 同意 ☐ (06) 非常同意

二、實驗題目

填寫完實驗前問卷，受試者會依序看到四個個案以下為本次研究之參考案例之一：

接下來的題目，我們將依序提供您四個個案資料，請您依照您平常工作業務的作法，對每個個案給予風險評估與處遇措施的判斷決策

案例四

- 一、 姓名：案○四
- 二、 性別：男
- 三、 生日：58 年 8 月（現年 52 歲，犯案時 44 歲）
- 四、 主要案由：毒品危害防制條例（販賣二級毒品安非他命）
- 五、 刑期：8 年 6 月
- 六、 保護管束類型：假釋
- 七、 保護管束期間：108 年 11 月 4 日起至 111 年 11 月 26 日止
- 八、 學歷：高中畢業
- 九、 家庭狀況：案父已歿，案母健在，案家務農，個案於家中排行老么，上有二個哥哥，哥哥均已過世，個案離婚，與前妻育有 2 子，與兒子沒有往來，個案現與案母及案姪子（大哥的兒子）同住，家庭支持薄弱。
- 十、 交友狀況：交友複雜，與施用毒品人口往來。
- 十一、 工作情形：於自家務農採收。
- 十二、 前科紀錄：有建築法（89 年，拘役 30 日易科罰金）、電遊場業管理罪（89 年，拘役 20 日易科罰金）前科，102 年至 103 年間犯下本案，104 年另有施用毒品被裁定觀察勒戒紀錄，於 104 年因本案入監執行，本次為第 1 次假釋付保護管束。

每看完一個個案資料，都接續問受訪者以下問題。

請依照這裡呈現的個案資料，回答以下有關該個案之風險評估與處遇措施的相關問題：

危險級別評估題組

B1. 該個案犯罪行為的產生受到不良的社會影響的程度為何？

- ☐ (1) 極低度影響 ☐ (2) 低度影響 ☐ (3) 中度影響
☐ (4) 中高度影響 ☐ (5) 極高度影響 ☐ (6) 無法判斷

B2. 個案在社會中受到排斥且缺乏對他人的關心，產生與社會疏離的程度為何？

- ☐ (1) 極低度影響 ☐ (2) 低度影響 ☐ (3) 中度影響
☐ (4) 中高度影響 ☐ (5) 極高度影響 ☐ (6) 無法判斷

B3. 個案在自我特質上，因缺乏同理心與負面情緒，導致有問題解決技巧的認知不良情形的程度為何？

- ☐ (1) 極低度認知不良 ☐ (2) 低度認知不良
☐ (3) 中度認知不良 ☐ (4) 中高度認知不良
☐ (5) 極高度認知不良 ☐ (6) 無法判斷

B4. 個案願意接受觀護與治療的配合程度為何？

- ☐ (1) 極低度配合 ☐ (2) 低度配合
☐ (3) 中度配合 ☐ (4) 中高度配合
☐ (5) 極高度配合 ☐ (6) 無法判斷

B5. 請問個案的危險級別為何？

- ☐ (1) 0-20%極低度危險 ☐ (2) 21-40%低度危險
☐ (3) 41-60%中度危險 ☐ (4) 61-80%中高度危險
☐ (5) 81-100%極高度危險

再犯風險評估題組

B6. 請問個案的再犯風險級別為何？

- ☐ (1) 0-20%極低度危險 ☐ (2) 21-40%低度危險

- ☐ (3) 41-60%中度危險 ☐ (4) 61-80%中高度危險
☐ (5) 81-100%極高度危險

觀護處遇評估題組

B7. 針對個案之危險級別，觀護處遇評估為第幾級？

- ☐ (1) 第一級（高度監督及輔導，以核心案件列管，實施特殊處遇）
接續填答第 8 題、第 9 題。
☐ (2) 第二級（中度監督及輔導，尚需加強列管、輔導或其他事由以
核心案件列管）接續填答第 8 題、第 9 題。
☐ (3) 第三級（低度監督及輔導）只填答第 8 題，跳過第 9 題。

B8. 針對個案應實施之一般處遇作為（複選題）：

- ☐ (1) 約談、訪視 ☐ (2) 列管核心案件加強報到
☐ (3) 採驗尿液 ☐ (4) 警局複數監督
☐ (5) 命接受毒品或酒癮戒癮治療 ☐ (6) 法治教育
☐ (7) 轉介就業服務 ☐ (8) 家庭支持
☐ (9) 指定榮觀協助約談或訪視 ☐ (10) 命其向警局（派出所）報到
☐ (11) 定期或不定期電訪 ☐ (12) 其他：_____

B9. 針對個案建議之特殊處遇作為（複選題）：

- ☐ (1) 實施科技設備監控 ☐ (2) 指定居住處所
☐ (3) 監控時段，未經許可不得外出 ☐ (4) 測謊
☐ (5) 轉介適當機構或團體 ☐ (6) 心理測驗
☐ (7) 禁止接近特定處所或對象 ☐ (8) 命主動電話回報觀護人
☐ (9) 命其定期接受身心治療或輔導
☐ (10) 命定時（每日/每週/每月）向警局（派出所）報到
☐ (11) 無須特殊處遇
☐ (12) 其他必要特殊處遇，如：禁止上色情網站、禁止網路交友、不得飲酒、
不得吸食毒品等：_____

實驗介入 framing

A. 控制組（無外部資訊 framing）

B10.A 以下顯示為您對此個案的判斷結果摘要，請問您是否確定此決定，還是需要進行修改（如欲觀看該個案資料請點選此連結）？

- ☐ （1）是，需要修改。 ☐ （2）否，不需要修改。

B. 實驗組 Arm 1（機器學習演算法介入）

目前由本調查研究團隊開發一機器學習測試系統，名叫「Parole-ML-AI 1.0」，用意為提供觀護人處理受保護管束人再犯風險評估業務時的智慧輔助系統，該系統後端資料庫收集歷年來觀護案件的相關巨量資料，系統具有對於受保護管束人的危險級別、再犯風險評估與觀護處遇建議之功能，給予觀護人進行受保護管束人再犯風險評估及觀護處遇計畫擬定的參考。我們將提供你 Parole-ML-AI 1.0 針對此個案的決策建議，請按下一步，啟動此系統。

我們的 Parole-ML-AI 1.0 系統運算已結束，分析給出的預測結果如下：

（顯示案例）

B10.B 以下顯示為您對此個案的判斷結果摘要，請問您是否確定此決定，還是需要進行修改？（實驗組 Arm 1）

- ☐ （1）是，需要修改。請填答第 11 題。
☐ （2）否，不需要修改。請填答第 11 題。

B11.B 整體來說，請問您是否信任剛剛 Parole-ML-AI 1.0 系統對個案所作出的決策建議？（實驗組 Arm 1）

- ☐ （1）非常不信任 ☐ （2）不信任 ☐ （3）普通
☐ （4）有點信任 ☐ （5）信任 ☐ （6）非常信任

C. 實驗組 Arm 2（其他同仁之綜合決策作為外部資訊）

B10.C 以下顯示為您對此個案的判斷結果摘要，在螢幕的右側也同時提供您其他參與此調查同仁之綜合答案，給您參考。請問您是否確定您原先自己的判斷結果，還是需要進行修改（如欲觀看該個案資料請點選此連結）？（實驗組 Arm 2）

- ☐ （1）是，需要修改。請填答第 12 題。

☐ (2) 否，不需要修改。請填答第 12 題。

B11.C 整體來說，請問您是否信任剛剛其他參與調查的同仁對個案所作的綜合決策資料？（實驗組 Arm 2）

☐ (1) 非常不信任

☐ (2) 不信任

☐ (3) 普通

☐ (4) 有點信任

☐ (5) 信任

☐ (6) 非常信任

三、實驗後的調查問卷題目

機器演算與課責題組

C1. 首先謝謝您協助我們完成以上的實驗調查。本團隊想再次向您請教，若未來法務部要開發並導入新的「再犯風險評估機器演算之人工智慧系統」（以下簡稱為 Parole-ML-AI 1.0），請試著想像其可能發生的結果，請問你是否支持？

- ☐ (01) 非常不支持 ☐ (02) 不支持 ☐ (03) 有點不支持
☐ (04) 有點支持 ☐ (05) 支持 ☐ (06) 非常支持

C2. 承上題，法務部未來若開發並導入新的「再犯風險評估機器演算之人工智慧系統」（Parole-ML-AI 1.0），請以此情境回答您對以下題目敘述的看法。

C2.1 該系統所造成的錯誤需由我負責。

- ☐ (01) 非常不同意 ☐ (02) 不同意 ☐ (03) 有點不同意
☐ (04) 有點同意 ☐ (05) 同意 ☐ (06) 非常同意

C2.2 該系統將限縮我的裁量權力。

- ☐ (01) 非常不同意 ☐ (02) 不同意 ☐ (03) 有點不同意
☐ (04) 有點同意 ☐ (05) 同意 ☐ (06) 非常同意

C2.3 該系統將讓長官更容易監督我的工作情況。

- ☐ (01) 非常不同意 ☐ (02) 不同意 ☐ (03) 有點不同意
☐ (04) 有點同意 ☐ (05) 同意 ☐ (06) 非常同意

C2.4 該系統將讓我承擔更多責任。

- ☐ (01) 非常不同意 ☐ (02) 不同意 ☐ (03) 有點不同意
☐ (04) 有點同意 ☐ (05) 同意 ☐ (06) 非常同意

C2.5 該系統將讓我更容易回應社會民眾的要求。

- ☐ (01) 非常不同意 ☐ (02) 不同意 ☐ (03) 有點不同意
☐ (04) 有點同意 ☐ (05) 同意 ☐ (06) 非常同意

C2.6 該系統並不能協助我的工作業務。

- | | | |
|-------------------------------------|-----------------------------------|-------------------------------------|
| ☐ (01) 非常不同意 | ☐ (02) 不同意 | ☐ (03) 有點不同意 |
| ☐ (04) 有點同意 | ☐ (05) 同意 | ☐ (06) 非常同意 |

C2.7 我無法相信該系統給予的風險評估建議。

- | | | |
|-------------------------------------|-----------------------------------|-------------------------------------|
| ☐ (01) 非常不同意 | ☐ (02) 不同意 | ☐ (03) 有點不同意 |
| ☐ (04) 有點同意 | ☐ (05) 同意 | ☐ (06) 非常同意 |

C2.8 我會很期待使用這個系統。

- | | | |
|-------------------------------------|-----------------------------------|-------------------------------------|
| ☐ (01) 非常不同意 | ☐ (02) 不同意 | ☐ (03) 有點不同意 |
| ☐ (04) 有點同意 | ☐ (05) 同意 | ☐ (06) 非常同意 |

創新意願題組

C3. 以下是想瞭解您對於創新意願的想法，請您回答下列問題：

C3.1 我發現我通常是團體中比較快接受新想法的人。

- | | | |
|-------------------------------------|-----------------------------------|-------------------------------------|
| ☐ (01) 非常不同意 | ☐ (02) 不同意 | ☐ (03) 有點不同意 |
| ☐ (04) 有點同意 | ☐ (05) 同意 | ☐ (06) 非常同意 |

C3.2 在接受新想法時，我通常會很謹慎。

- | | | |
|-------------------------------------|-----------------------------------|-------------------------------------|
| ☐ (01) 非常不同意 | ☐ (02) 不同意 | ☐ (03) 有點不同意 |
| ☐ (04) 有點同意 | ☐ (05) 同意 | ☐ (06) 非常同意 |

C3.3 即便有失敗的可能性，我還是願意去嘗試過去工作上沒做過的事。

- | | | |
|-------------------------------------|-----------------------------------|-------------------------------------|
| ☐ (01) 非常不同意 | ☐ (02) 不同意 | ☐ (03) 有點不同意 |
| ☐ (04) 有點同意 | ☐ (05) 同意 | ☐ (06) 非常同意 |

C3.4 即便可能會有風險，我還是願意在工作業務上嘗試新科技。

- | | | |
|-------------------------------------|-----------------------------------|-------------------------------------|
| ☐ (01) 非常不同意 | ☐ (02) 不同意 | ☐ (03) 有點不同意 |
| ☐ (04) 有點同意 | ☐ (05) 同意 | ☐ (06) 非常同意 |

工作業務流程自動化

C4.以下列出與觀護人相關之諸多工作業務，請依據您工作上的經驗，您認為需要開發相對應的機器演算人工智慧系統來協助嗎？

C4.1 例行行政工作與文書撰寫

- | | | |
|-------------------------------------|-----------------------------------|-------------------------------------|
| ☐ (01) 非常不需要 | ☐ (02) 不需要 | ☐ (03) 有點不需要 |
| ☐ (04) 有點需要 | ☐ (05) 需要 | ☐ (06) 非常需要 |

C4.2 保護管束案件管理

- | | | |
|-------------------------------------|-----------------------------------|-------------------------------------|
| ☐ (01) 非常不需要 | ☐ (02) 不需要 | ☐ (03) 有點不需要 |
| ☐ (04) 有點需要 | ☐ (05) 需要 | ☐ (06) 非常需要 |

C4.3 評估再犯風險

- | | | |
|-------------------------------------|-----------------------------------|-------------------------------------|
| ☐ (01) 非常不需要 | ☐ (02) 不需要 | ☐ (03) 有點不需要 |
| ☐ (04) 有點需要 | ☐ (05) 需要 | ☐ (06) 非常需要 |

C4.4 提供處遇措施建議

- | | | |
|-------------------------------------|-----------------------------------|-------------------------------------|
| ☐ (01) 非常不需要 | ☐ (02) 不需要 | ☐ (03) 有點不需要 |
| ☐ (04) 有點需要 | ☐ (05) 需要 | ☐ (06) 非常需要 |

C4.5 輔導監督

- | | | |
|-------------------------------------|-----------------------------------|-------------------------------------|
| ☐ (01) 非常不需要 | ☐ (02) 不需要 | ☐ (03) 有點不需要 |
| ☐ (04) 有點需要 | ☐ (05) 需要 | ☐ (06) 非常需要 |

公共服務動機題組

C5.最後，請問您對於下列敘述的同意程度為何？

C5.1 為大眾服務是很重要的事情。

- | | | |
|-------------------------------------|-----------------------------------|-------------------------------------|
| ☐ (01) 非常不同意 | ☐ (02) 不同意 | ☐ (03) 有點不同意 |
| ☐ (04) 有點同意 | ☐ (05) 同意 | ☐ (06) 非常同意 |

C5.2 為社會做出貢獻比得到個人成就更有意義。

- | | | |
|-------------------------------------|-----------------------------------|-------------------------------------|
| ☐ (01) 非常不同意 | ☐ (02) 不同意 | ☐ (03) 有點不同意 |
|-------------------------------------|-----------------------------------|-------------------------------------|

☐ (04) 有點同意 ☐ (05) 同意 ☐ (06) 非常同意

C5.3 社會上人與人之間是互相需要的。

☐ (01) 非常不同意 ☐ (02) 不同意 ☐ (03) 有點不同意

☐ (04) 有點同意 ☐ (05) 同意 ☐ (06) 非常同意

C5.4 為社會公益而犧牲自己的權益沒有關係。

☐ (01) 非常不同意 ☐ (02) 不同意 ☐ (03) 有點不同意

☐ (04) 有點同意 ☐ (05) 同意 ☐ (06) 非常同意

C5.5 就算會惹禍上身，路見不平拔刀相助他人是必要的。

☐ (01) 非常不同意 ☐ (02) 不同意 ☐ (03) 有點不同意

☐ (04) 有點同意 ☐ (05) 同意 ☐ (06) 非常同意

問卷到此結束，感謝您參與我們的調查研究。

本計畫希望致贈 7-11 電子禮券 300 元作為微薄之謝禮，若您願意領取禮券，麻煩請留下姓名以及電子郵件，以便日後作聯繫以及禮券序號連結的寄送。

請問您是否願意收到完成填答的禮券？

☐ 願意 ☐ 不願意