

「公務AI共通核心能力認證」命題方向

主層面	公務客製 內容說明	次層面/ 認證內涵	命題方向
1.AI基礎 認知	建立公務人員對於AI基本概念、生成式AI原理及在公務流程中人機協作模式的理解。客製內容將強化公務人員在應用公文、談話參考資料、採購招標文件、會議紀錄等情境時所需基本技術理解及應用原理。	1.1AI基礎概念 與社會影響 AI概論、數位素養、發展趨勢、AI影響	建立公務人員對人工智慧基本概念、公共治理影響、發展趨勢及應用風險之基礎認知。 1.基礎概念與AI素養 評估公務人員對人工智慧之基本定義、運作概念及常見公務應用情境(如常見問題自動回覆、輿情彙整與分析、郵件或新聞稿草擬、會議紀錄整理等)之基礎理解；並評估其對數位素養、AI倫理與治理基本原則之認知，包括隱私保護與資料治理、資安與防護、透明與可解釋、公平與不歧視、問責，以及演算法偏差、資訊誤導或造假等風險。 2.AI發展趨勢與公共治理影響 評估公務人員對AI在公共服務、政策研擬、決策輔助、風險預警及行政效率提升等應用趨勢之理解；並能辨識AI導入對公務職能、工作流程、人力配置、法制與治理規範、數位落差及社會信任可能產生之影響與風險。 3.AI工具發展與系統介接認知 評估公務人員對AI工具發展趨勢及基本應用樣態之認知，包括生成式AI、代理型AI(Agentic AI)及相關系統介接概念；並理解系統介接協定或工具(如 MCP等技術概念)，可能對機關內部系統串接、跨機關資料流通、權限控管、資料安全及服務流程設計產生之影響。
		1.2生成式AI與人機互動	強化公務人員於文書、會議及日常行政情境，運用生成式AI進行人機協作、提示設

主層面	公務客製 內容說明	次層面/ 認證內涵	命題方向
		生成式AI基礎、智慧轉型、未來職場、人機互動	計、內容查核與任務監督之實務素養。 1.提示詞(Prompt)實務與迭代 評估公務人員運用生成式AI處理跨機關共通性業務之實作素養，掌握「角色、任務、背景、格式」要素，並具備迭代修正(Iterative Prompting)之多輪對話微調能力。命題情境建議包含「公文與民眾陳情覆函草擬」、「政策白話化說明」、「針對爭議政策進行輿情模擬分析與澄清稿架構建議」、「機關社群平台貼文」及「活動企劃發想」等實務情境，並能依據發布管道與目標受眾，引導AI產出適切之風格與語氣。 2.人機協作認知與內容把關 善用人機協作模式，提升公務效率與品質，落實「AI草擬、人類審核」原則。導入批判性「紅隊測試」觀念與事實查核觀念，評估公務人員對AI生成初稿之主動查證能力，從嚴檢視：法規引用正確性、公文體例妥適性、數據來源真實性與邏輯合理性。 3.任務監督與進階協作 面對代理型 AI(Agentic AI)協助規劃與執行多步驟任務之應用情境，評估公務人員對進階人機協作模式之認知，能否由單純的「工具使用者」轉為「任務監督者」，具備設定任務目標、使用範圍、資料輸入限制、系統邊界條件及人類監督或覆核節點(Human-in-the-loop)之素養。

主層面	公務客製 內容說明	次層面/ 認證內涵	命題方向
		1.3AI應用與公共服務轉型策略 IT趨勢、AI發展、AI願景、產業衝擊	培養公務人員依業務需求、資料敏感度及公共服務對象，選擇適切AI工具與部署模式，並兼顧數位平權、演算法風險及公共服務轉型之素養。 1.預測型與生成式AI區辨 評估公務人員能否準確釐清兩者在行政決策中不同角色；理解預測型AI旨在透過數據洞察提供輔助決策參考與風險預警；生成式AI則聚焦於公文草擬、會議摘要等文書效率提升。 2.模型選擇素養 評估公務人員能否依資料敏感度、業務風險、資安要求、個資保護、著作權及資料合規性，判斷雲端服務、機關內部封閉環境或本地部署模型之適用情境；並能辨識國際模型與本土模型(如TAIDE)在語言表現、在地語境、公務適用性、限制條件及查證補強機制上的差異，據以選擇適切工具或替代方案。 3.數位平權與科技包容 評估公務人員在應用AI提供公共服務時，是否具備數位平權與科技包容之設計思維，能防範因數位落差、演算法偏誤、自動化判斷或單一數位介面造成特定群體權益受損；並能兼顧實體與數位服務管道、資訊無障礙、服務可近性及第一線人員作業可行性，以提升公共服務之公平性與可及性。 4.AI導入與行政流程轉型 評估公務人員對代理型AI工作流程(Agentic Workflow)之基本認知，理解其可能對行政流程再設計、跨系統協作、權責分工、人力配置、業務效率及服務品質管

主層面	公務客製 內容說明	次層面/ 認證內涵	命題方向
			理產生之影響；並能辨識導入過程中應注意之風險控管、人類監督及機關治理需求。
2.AI政策 法制	此主層面為公務客製化的核心重點，旨在確保公務人員在應用生成式AI於公文、採購招標文件等情境時，具備足夠的倫理意識、法規遵循和資安風險判斷能力。	2.1AI政策與治理指引 AI政策(我國AI政策發展脈絡與行動方案)、AI治理指引(國際AI治理架構、臺灣演算法監理與風險分級實踐)	建立公務人員對我國AI政策、治理指引、風險分級、資料治理及AI導入生命週期管理之理解與法遵意識。 1.政策目標與治理指引理解 評估公務人員對我國推動 AI 之總體政策、治理規範及跨部會應用指引之理解與遵循能力；並能依業務屬性，辨識所適用之一般性或專業領域AI應用規範，確保跨機關業務推動符合相關規範。 2.風險分級與監理 評估公務人員對「人工智慧風險分類」、演算法監理及風險控管原則之理解。命題可結合採購／招標文件審查、契約條款審查、政策輔助分析、民眾服務系統導入等情境，評估其能否依業務屬性進行初步風險辨識與分類，並提出適切之控管措施。 3.資料治理與AI使用判斷 評估公務人員對AI友善資料(AI-Friendly Data)、機器可讀、可檢索、可引用等資料治理原則之理解，並能判斷公務資料是否得作為AI訓練、測試、檢索或生成服務之資料來源；對於涉及機密、未公開公務資訊、個人資料或營業秘密之資料，應能辨識其不得提供外部AI服務使用，並掌握可公開資料之授權、權利宣告、去識別化及風險控管要求。 4.AI系統成效評估與持續監測 評估公務人員對AI系統導入後之持續管

主層面	公務客製 內容說明	次層面/ 認證內涵	命題方向
			理素養，包含模型準確率、回應穩定性、偏誤情形、資源消耗效率、使用者回饋及服務品質等指標之檢核，並能依評估結果進行滾動修正，確保 AI 應用持續符合施政目標、資安規範與倫理準則。 5.AI導入治理與生命週期管理 評估公務人員是否具備AI系統導入前、中、後之治理意識，包括導入前之風險評估、業務適法性檢核、資料治理、模型或系統選型及採購／委外管理；導入中之權限控管、使用紀錄保存、人類監督節點及異常事件通報；導入後之持續監測、改善追蹤及內控檢討。對於可能影響人民權益或屬高風險人工智慧之應用，應能辨識其責任歸屬、歸責條件，以及救濟、補償或保險機制，以及人類監督與覆核之必要性。
		2.2AI法規與資 安風險管理 AI安全、法 規規範	建立公務人員對AI應用之資安、個人資料保護、機敏資料防護、使用紀錄及合規控管之風險辨識與實務判斷能力。 1.資安規範與資料保護 評估公務人員能否辨識雲端服務、機關核准之封閉式地端部署或內部系統在資料處理、儲存、日誌紀錄、對話紀錄保存及模型訓練或微調使用上之差異，並建立正確之資安與法遵意識。強化「影子AI(Shadow AI)」風險意識，提醒公務人員不得私自使用未經機關核准之外部AI工具處理涉及公務應保密、個人資料、營業秘密或未經機關同意公開之資訊。評測宜透過具體之禁止或不宜輸入範例(負面表列範例)，如未公開之採購底價、

主層面	公務客製 內容說明	次層面/ 認證內涵	命題方向
			國民身分證統一編號、未發布之政策草案、內部簽辦意見或其他機關機敏資料，強化防範資料外洩風險之實務判斷能力。 2.AI法規與資安風險管理實務 評估公務人員落實資安規範與風險管控之能力。導入AI國際管理體系(如 ISO 42001)觀念，評估其在公務流程中落實「事前風險評估」、「使用日誌稽核」與「合規性審查」等控管程序之實務素養。評估公務人員於AI相關採購、委外開發或服務導入時，是否能辨識契約應納入之AI使用管理要求，包括資料不得外流、訓練資料來源合法性、個人資料與機敏資料保護、使用日誌保存、模型更新管理、成果驗收、資安責任，以及承商應遵守機關AI使用規範與內控管理機制。 3.生成式AI使用限制與適法性 評估公務人員能否依據「人工智慧基本法」及「行政院及所屬機關(構)使用生成式AI參考指引」等，判斷生成式AI使用之合法性與風險。公務人員應理解機密文書、未公開公務資訊、個人資料、營業秘密及未經機關同意公開之資料，不得輸入外部生成式AI服務；機密文書應由業務承辦人親自撰寫，不得使用生成式AI代為產製。AI 產出須經承辦人查證與專業判斷，不得未經確認即直接作成行政行為，亦不得作為公務決策之唯一依據。

主層面	公務客製 內容說明	次層面/ 認證內涵	命題方向
		2.3AI應用下的 倫理準則與 智慧財產實 務 智慧財產 權、AI倫理 原則	提升公務人員對AI生成內容可信度、使用揭露、智慧財產權、行政倫理、責任歸屬及人類監督之實務判斷能力。 1.最終問責與行政倫理 評估公務人員應用AI之倫理與問責素養，明確理解AI僅為輔助工具，最終行政處分、公文核定與業務決策仍應由權責人員負責。評估公務人員能否在AI輔助決策過程中，進行公平性、不歧視及社會影響檢核，並確保生成內容符合我國法制、政策語彙、在地語境與文化價值，避免用語誤植、脈絡錯置或價值表述不當。 2.智財權、授權條款與使用揭露 評估公務人員判斷AI輔助生成內容，如圖文、簡報、影音或文案之著作權、授權條款及使用風險之能力；能區辨AI僅作為輔助工具且有人類實質創意投入，與AI獨立生成且缺乏人類創作投入之差異。公務發布內容前，應檢核資料來源、授權條款、合理使用可能性、人格權、肖像權及第三方權利，並依AI參與程度進行適當揭露或標記。 3.資料可信度與透明度 辨別AI生成內容可能存在的錯誤、偏誤、幻覺與資訊過時之風險；評估公務人員對數位浮水印等內容鑑別機制之認知，確保具備辨識與防範深偽(Deepfake)造假及錯假訊息傳播之能力；並測驗其查核AI輔助決策結果之能力，以防範自動化偏誤，落實人機協作中之人類自主與監督。

公務AI共通核心能力認證樣題

題號：1

主層面：1.AI 基礎認知	次層面：1.1AI 基礎概念與社會影響
題型：單選	難度：難
題幹： 隨著生成式 AI 技術的演進，大型多模態模型 (Large Multimodal Model, LMM) 成為當前重要發展趨勢。下列關於 LMM 核心架構與能力的敘述，何者最為精確？	
選項(1) 專注於高畫質影像的生成與像素級修復，屬於純粹的進階電腦視覺運算模型。	選項(2) 僅能針對文字與自然語言進行深度學習，為傳統 NLP 模型架構的最終單一升級版本。
選項(3) 具備跨領域特徵對齊能力，能同時接收、理解並整合文字、圖像、音訊等多種資料型態。	選項(4) 為一種專門優化資料庫搜尋效率的演算法，能大幅降低關聯式資料庫的跨表查詢延遲時間。
正確答案：3	

公務AI共通核心能力認證樣題

題號：2

主層面：1.AI 基礎認知	次層面：1.1AI 基礎概念與社會影響
題型：單選	難度：易
題幹： 某機關運用生成式 AI 輔助撰寫政策分析報告，內容包含現況分析、數據整理及政策建議。完成初稿後，承辦人進行「人類參與流程（Human-in-the-Loop）」。下列哪一項不屬於承辦人應具備的查核素養？	
選項(1) 交叉比對外部權威資料來源，以確認生成訊息的事實正確性與數據準確性。	選項(2) 評估生成文本在多個段落之間的推論邏輯是否嚴密且具備前後一致性。
選項(3) 仔細檢查生成文本的語法結構、專業術語使用是否精準，以及有無拼寫錯誤。	選項(4) 預設系統已自動過濾所有不當或錯誤資訊，將查核重心僅放在排版與視覺美化上。
正確答案：4	

公務AI共通核心能力認證樣題

題號：3

主層面：1.AI 基礎認知	次層面：1.1AI 基礎概念與社會影響
題型：單選	難度：易
題幹： 生成式 AI (Generative AI) 與傳統機器學習或判別式 AI (Discriminative AI) 在應用面上有著本質的差異。下列何者最能精確描述生成式 AI 的「核心技術功能」？	
選項(1) 僅針對既有的大型結構化數據集，進行精準的特徵分類、異常檢測與模式識別。	選項(2) 解析並學習龐大訓練資料中的隱含機率分佈，進而創造出具備原創特徵的全新內容。
選項(3) 取代傳統 CPU 架構，專責執行要求絕對零誤差的高階微積分與複雜物理數值統計。	選項(4) 即時監控雲端伺服器的網路流量與運行狀態，並於發生資安威脅時自動發出警報。
正確答案：2	

公務AI共通核心能力認證樣題

題號：4

主層面：1.AI 基礎認知	次層面：1.1AI 基礎概念與社會影響
題型：單選	難度：難
題幹： 大規模且低成本的深度偽造（Deepfake）與文本生成技術，正在對民主社會構成嚴峻挑戰。下列何者最精確說明了生成式 AI 對「社會信任」所造成的「結構性風險（Structural Risk）」？	
選項(1) 迫使社群平台必須導入更嚴格的實名制與內容認證機制，導致匿名使用者的言論自由與隱私權利受到大幅限縮。	選項(2) 傳統新聞機構為了查核 AI 生成的錯假資訊，必須投入高昂的人力成本，進而排擠了原創深度報導的營運資源。
選項(3) 氾濫的逼真偽造內容引發「騙子紅利（Liar's Dividend）」，導致公眾陷入知識論崩塌，最終喪失對任何數位證據與客觀「真實性」的共同判準。	選項(4) 讓惡意行為者能利用語言模型輕易生成具備高度說服力的釣魚郵件，大幅提高企業與民眾遭到網路詐騙的機率。
正確答案：3	

公務AI共通核心能力認證樣題

題號：5

主層面：1.AI 基礎認知	次層面：1.1AI 基礎概念與社會影響
題型：單選	難度：易
題幹： 某地方機關規劃全面改以 AI 線上服務受理補助申請。下列哪項風險最應納入評估？	
選項(1) 民眾是否會覺得網頁顏色好看	選項(2) 是否能完全取消紙本申請
選項(3) AI 是否能取代全部臨櫃人員	選項(4) 是否造成不熟悉數位工具者、身心障礙者或偏鄉民眾較難取得服務
正確答案：4	

公務AI共通核心能力認證樣題

題號：6

主層面：1.AI 基礎認知	次層面：1.1AI 基礎概念與社會影響
題型：單選	難度：易
題幹： 某機關想讓代理型 AI 自動串接內部系統、查詢資料並產出處理建議。下列治理措施何者最適當？	
選項(1) 只要 AI 效率高，即可開放所有系統權限	選項(2) 應設定權限範圍、資料輸入限制、操作紀錄及人工覆核節點
選項(3) 讓 AI 自行決定是否查詢機敏資料	選項(4) 不需記錄 AI 執行過程，以免增加系統負擔
正確答案：2	

公務AI共通核心能力認證樣題

題號：7

主層面：1.AI 基礎認知	次層面：1.1AI 基礎概念與社會影響
題型：單選	難度：易
題幹： 某機關規劃導入 AI 客服系統。依《人工智慧基本法》，下列何者最符合「人類自主／保留監督」原則？	
選項(1) AI 系統可完全取代人類決策，以提升行政效率	選項(2) AI 系統應支持人類自主權，並保留人類監督機制
選項(3) AI 系統可依模型判斷直接執行所有行政處分	選項(4) AI 系統應優先追求運算效能，不須考量使用者權益
正確答案：2	

公務AI共通核心能力認證樣題

題號：8

主層面：1.AI 基礎認知	次層面：1.2生成式 AI 與人機互動
題型：單選	難度：中
題幹： 在生成式 AI（Generative AI）的操作環境中，所謂的「提示詞（Prompt）」扮演著關鍵角色。下列何者最能精確定義其在人機協作過程中的技術定位？	
選項(1) 一種專為自然語言處理（NLP）設計的高階程式編譯語言。	選項(2) 使用者輸入給大語言模型，用以觸發、引導並規範生成內容的自然語言指令。
選項(3) 系統後台自動化爬取網路資料時，用來設定檢索範圍的一種應用程式介面（API）。	選項(4) 專門用於壓縮與加密生成文本，以防止機敏資料外洩的底層資安防護機制。
正確答案：2	

公務AI共通核心能力認證樣題

題號：9

主層面：1.AI 基礎認知	次層面：1.2生成式 AI 與人機互動
題型：單選	難度：易
題幹： 某機關導入生成式 AI 協助整理會議紀錄，希望 AI 能依循既定格式、使用正式行政用語，並正確整理討論事項與決議內容。下列哪一種提示詞優化策略（Prompt Engineering）最為直接有效？	
選項(1) 運用「少樣本提示（Few-shot Prompting）」技術，在系統指令中嵌入數個優良的標準問答範例供其參考。	選項(2) 刻意移除所有關於企業品牌語氣的約束條件，讓模型根據網路上廣泛的鄉民對話風格來拉近與客戶的距離。
選項(3) 從底層大幅刪減模型的預訓練資料量，使其無法理解過於複雜的提問，藉此降低答錯的機率。	選項(4) 僅在指令中宣告「客戶是尊貴的」，卻完全不提供任何關於產品規格或退換貨政策的上下文邊界限制。
正確答案：1	

公務AI共通核心能力認證樣題

題號：10

主層面：1.AI 基礎認知	次層面：1.2生成式 AI 與人機互動
題型：單選	難度：易
題幹： 某機關承辦人需利用生成式 AI 協助撰擬一份複雜政策推動計畫公文，內容包含背景說明、辦理目的、執行方式及預期效益。相較於一次輸入完整需求的單一長篇指令（Zero-shot），採用「分階段提問（Step-by-step Prompting）」策略較能提升產出品質的主要技術原因為何？	
選項(1) 分段輸入可以有效欺騙伺服器的負載平衡機制，確保每次提問都能獲得最高優先級的 GPU 算力分配。	選項(2) 避免模型因單次資訊過載而產生「注意力偏移（Attention Shift）」，確保每個子任務都能獲得充分的推論深度。
選項(3) 將問題切碎能夠觸發底層架構的平行運算機制，讓生成式 AI 的整體文本輸出速度提升三倍以上。	選項(4) 此做法能迫使模型解除內建的倫理道德限制，使其在生成創意文本時能夠產出更加辛辣與無邊界的內容。
正確答案：2	

公務AI共通核心能力認證樣題

題號：11

主層面：1.AI 基礎認知	次層面：1.2生成式 AI 與人機互動
題型：單選	難度：難
題幹： 專業經理人欲利用生成式 AI 加速撰寫 B2B 商業提案。為確保最終產出具備深度的「商業洞察力（Business Insights）」且邏輯嚴謹，而非僅是華麗空泛的文字堆疊，最佳的人機協作策略為何？	
選項(1) 授權 AI 無限制地開啟網頁檢索功能，不設立任何參考邊界，讓系統自行爬取全球新聞以擴充提案的資訊量。	選項(2) 在提示詞中預先輸入結構化的內部市調數據與產品定位，引導 AI 嚴格基於給定事實進行推導與利益點陳述。
選項(3) 捨棄冰冷的數據，要求 AI 在產出過程中全面套用文學化的修辭與極端誇飾語氣，以情感訴求來彌補洞察的不足。	選項(4) 將企業過去十年內所有失敗與成功的不同專案企劃書未經清洗直接混入，要求 AI 自動在龐雜的雜訊中提煉出致勝關鍵。
正確答案：2	

公務AI共通核心能力認證樣題

題號：12

主層面：1.AI 基礎認知	次層面：1.2生成式 AI 與人機互動
題型：單選	難度：中
題幹： AI 產出一段引用某法規條文的公文說明，但承辦人不確定條文是否仍有效。為降低 AI 產生錯誤資訊的風險，下列處理何者最適當？	
選項(1) 請 AI 再次生成內容確認是否一致	選項(2) 請 AI 提供引用依據後，查核正式法規資訊
選項(3) 請 AI 改寫文字內容降低引用精確性	選項(4) 請 AI 增加說明內容提高完整程度
正確答案：2	

公務AI共通核心能力認證樣題

題號：13

主層面：1.AI 基礎認知	次層面：1.2生成式 AI 與人機互動
題型：單選	難度：中
題幹： 某機關導入 Agentic AI 協助辦理補助案審查。系統可自動蒐集資料、分析申請資格並提出審查建議。為兼顧 AI 應用效率與行政責任，下列何者最符合人機協作治理精神？	
選項(1) AI 完成資格分析後，由系統自動產生核定結果	選項(2) AI 執行審查流程時，僅需保存最終結果即可
選項(3) AI 提出審查建議後，由承辦人確認結果並處理例外案件	選項(4) AI 依據模型判斷結果，自行承擔審查錯誤責任
正確答案：3	

公務AI共通核心能力認證樣題

題號：14

主層面：1.AI 基礎認知	次層面：1.2生成式 AI 與人機互動
題型：單選	難度：中
題幹： 某主管希望建立全自動生成式 AI 公文系統，由 AI 直接撰寫、核定並對外發文，不再經任何人員審核。為確保 AI 應用符合人機協作與任務監督要求，下列何者最適當？	
選項(1) 讓 AI 自行完成內容產製與流程決策，以發揮最大效率	選項(2) 讓 AI 依過往公文範例判斷內容，減少人工審查需求
選項(3) 讓 AI 協助產製內容，由承辦人查核結果並依權責完成核定	選項(4) 讓 AI 自動處理一般案件，僅將異常案件交由人員處理
正確答案：3	

公務AI共通核心能力認證樣題

題號：15

主層面：1.AI 基礎認知	次層面：1.3AI 應用與公共服務轉型策略
題型：單選	難度：易
題幹： 金融科技（FinTech）產業在導入生成式 AI 時，相較於其他產業往往面臨更高的阻力。下列何者是目前阻礙金融業大規模全面商用生成式 AI 的最核心限制？	
選項(1) 硬體算力的物理瓶頸，導致金融伺服器完全無法執行百億參數規模的語言模型。	選項(2) 產業特性導致極度缺乏歷史交易數據，使得演算法無法進行有效的深度學習訓練。
選項(3) 消費端對智能理財服務的強烈排斥，導致市場需求萎縮與企業導入誘因嚴重不足。	選項(4) 嚴苛的金融特許法規與合規要求，特別是涉及客戶機敏隱私與資料跨境傳輸的限制。
正確答案：4	

公務AI共通核心能力認證樣題

題號：16

主層面：1.AI 基礎認知	次層面：1.3AI 應用與公共服務轉型策略
題型：單選	難度：中
題幹： 將「生成式 AI（Generative AI）」導入傳統製造業時，相較於過去以預測與分類為主的傳統 AI，下列哪一項應用最能體現生成式 AI 所帶來的破壞性創新？	
選項(1) 衍生式設計（Generative Design）：依據材料特性與承載條件，自動生成多組具備最佳化結構的全新產品設計圖。	選項(2) 預測性維護（Predictive Maintenance）：利用物聯網感測器全天候監控生產線機台的運作溫度以預測故障。
選項(3) 存貨最佳化（Inventory Optimization）：透過歷史銷售數據的時間序列分析，精準計算出下一季的安全庫存量。	選項(4) 機器視覺（Machine Vision）：透過高速攝影機擷取影像，自動辨識生產線上的產品刮痕並剔除瑕疵品。
正確答案：1	

公務AI共通核心能力認證樣題

題號：17

主層面：1.AI 基礎認知	次層面：1.3AI 應用與公共服務轉型策略
題型：單選	難度：易
題幹： 終端設備品牌大廠（如 Apple）在全面佈局「邊緣人工智慧（Edge AI）」與生成式戰略時，下列何者通常並非其現階段投入鉅資發展底層模型的最主要戰略目標？	
選項(1) 提升旗下硬體生態系（如手機、電腦）的本地端語意理解、智慧化操作與隱私保護能力。	選項(2) 藉由創新的 AI 應用與個人化助理，拓展並強化軟體訂閱制等高毛利的服務型業務。
選項(3) 確保品牌在下一代人機互動介面與未來科技發展版圖中，維持絕對的技術領導地位。	選項(4) 透過 AI 演算法的優化，直接取代並大幅降低實體運算晶片與硬體零組件的製造成本。
正確答案：4	

公務AI共通核心能力認證樣題

題號：18

主層面：1.AI 基礎認知	次層面：1.3AI 應用與公共服務轉型策略
題型：單選	難度：難
題幹： 當企業將開發完成的生成式 AI 模型從「預訓練階段（Pre-training）」轉移至「正式上線部署的推論階段（Inference）」時，下列哪一項技術指標將成為決定系統日常營運成本與使用者體驗的最關鍵因素？	
選項(1) 雇用大量領域專家進行高品質人工標註（RLHF）與資料清洗的作業人力成本。	選項(2) 模型底層神經網路在初始設定時，所採用的參數權重初始化（Weight Initialization）策略。
選項(3) 系統在有限的邊緣或雲端運算資源下，進行模型推論時的吞吐量（Throughput）與延遲速度（Latency）。	選項(4) 模型架構在進行反向傳播與梯度更新（Gradient Descent）運算時的收斂速度與數學穩定性。
正確答案：3	

公務AI共通核心能力認證樣題

題號：19

主層面：1.AI 基礎認知	次層面：1.3AI 應用與公共服務轉型策略
題型：單選	難度：中
題幹： 若政府機關將生成式 AI 廣泛應用於社會福利分配、稅務審查或犯罪預測等高風險公共領域，當系統提供高度複雜且「不可解釋（Unexplainable）」的決策建議時，最可能對民主社會造成何種負面影響？	
選項(1) 迫使基層公務員必須耗費極長時間學習複雜的機器語言，反而導致行政機關的整體施政決策效率全面崩盤。	選項(2) 因演算法具備客觀中立的科技光環，將使得公眾對所有政府的裁量結果產生絕對的盲目信任。
選項(3) 因決策過程淪為「黑盒子（Black Box）」，導致受處分人無法理解其推論邏輯，進而引發社會對該施政結果正當性與程序公平的強烈質疑。	選項(4) 為了維持系統不可解釋的複雜度，政府必須持續採購高端硬體，導致公共預算的系統開發與維護成本暴增。
正確答案：3	

公務AI共通核心能力認證樣題

題號：20

主層面：1.AI 基礎認知	次層面：1.3AI 應用與公共服務轉型策略
題型：單選	難度：中
題幹： 某機關希望導入AI協助處理民眾申請案件。業務單位提出兩項需求：A 為依過去案件紀錄，提醒承辦人哪些案件可能需要補件；B 為協助將公開會議內容整理成摘要初稿。下列規劃何者較適當？	
選項(1) A、B 皆可交由 AI 自動完成，承辦人只需確認格式。	選項(2) A、B 均應先以 AI 產出結果為準，再視民眾反映決定是否修正流程。
選項(3) A 因涉及行政判斷，不宜使用 AI；B 因屬文字工作，可直接採用 AI 產出作為正式紀錄。	選項(4) A 應重視資料品質與過去案件紀錄是否足以支援判斷，B 則應注意摘要內容仍需人工檢核。
正確答案：4	

公務AI共通核心能力認證樣題

題號：21

主層面：1.AI 基礎認知	次層面：1.3AI 應用與公共服務轉型策略
題型：單選	難度：中
題幹： 某機關考慮導入代理式 AI，協助跨系統查詢資料、彙整報表、產生待辦提醒，並在特定條件下通知承辦人。下列哪一項規劃最能反映代理式 AI 與一般生成式 AI 的差異及其導入風險？	
選項(1) 將代理式 AI 視為文字生成工具，主要要求其輸出語氣符合公文格式。	選項(2) 明確設定可使用工具、資料存取範圍、任務邊界及人工覆核節點。
選項(3) 只要代理式 AI 可自動分解任務，即可授權其完成所有後續行政處理。	選項(4) 代理式 AI 串接愈多系統，愈能以提升跨機關資料取得效率。
正確答案：2	

公務AI共通核心能力認證樣題

題號：22

主層面：1.AI 基礎認知	次層面：1.3AI 應用與公共服務轉型策略
題型：單選	難度：易
題幹： 某機關想將未公開的政策評估草案輸入外部生成式 AI，請其改寫成較口語的民眾說明。下列作法何者較為正確？	
選項(1) 可輸入草案全文，但應要求 AI 不要引用原文。	選項(2) 可輸入草案重點，因只是協助改寫不是決策。
選項(3) 不應向外部生成式 AI 提供未經機關同意公開之資訊。	選項(4) 若產出內容經承辦人修改，即可回溯排除輸入資料風險。
正確答案：3	

公務AI共通核心能力認證樣題

題號：23

主層面：2.AI 政策法制	次層面：2.1AI 政策與治理指引
題型：單選	難度：易
題幹： 某機關建置生成式 AI 知識服務系統，整合公文、法規及政策資料作為 AI 回應依據。為確保資料來源、處理流程及使用歷程均可追溯，並符合資料治理要求，下列何者最符合資料治理中常用的管理機制？	
選項(1) 資料封裝 (Data Encapsulation)：將機敏特徵隱藏以提升資料傳輸的網路安全性。	選項(2) 資料版本控制 (Data Versioning)：管理不同版本資料集的變更歷程，以確保資料一致性與可重現性。
選項(3) 資料可移植性 (Data Portability)：確保內部數據能跨越不同雲端平台進行自由搬移。	選項(4) 資料溯源 (Data Lineage)：詳實紀錄資料從原始取得、清洗轉換到最終輸入模型的完整軌跡。
正確答案：4	

公務AI共通核心能力認證樣題

題號：24

主層面：2.AI 政策法制	次層面：2.1AI 政策與治理指引
題型：單選	難度：易
題幹： 根據行政院公部門或大型企業導入生成式 AI 的「資料治理與資安使用指引」，下列哪一項營運規範或操作認知是錯誤且具備高風險的？	
選項(1) 考量到 AI 模型的演算已達人類專家水準，機關應將其生成內容視為具備絕對權威性，並可直接作為發布正式公文的唯一依據。	選項(2) 使用生成式 AI 提供之服務時，應嚴格遵守資通安全、個人資料保護、著作權及相關資訊存取規範。
選項(3) 涉及國家機密或企業核心商業機敏資訊之文書，應由業務承辦人親自撰寫，嚴禁拋轉至外部雲端 AI 進行生成。	選項(4) 機關同仁若利用生成式 AI 作為執行業務或對外提供服務的輔助工具，應盡到適當揭露與標示的義務。
正確答案：1	

公務AI共通核心能力認證樣題

題號：25

主層面：2.AI 政策法制	次層面：2.1AI 政策與治理指引
題型：單選	難度：中
題幹： 某政府機關規劃採用雲端生成式 AI 服務協助分析民眾申請資料。由於 AI 服務可能涉及跨境資料儲存與運算，機關須評估資料主權（Data Sovereignty）相關要求。下列何者最符合資料主權的核心概念？	
選項(1) 秉持開源共享精神，所有用於訓練基礎模型的非個人數據皆應無條件免費對全球社群公開。	選項(2) 數位資料的擁有、控制與處理，必須嚴格受制於該資料產生或實體儲存所在國家（地區）的司法管轄與法律規範。
選項(3) 為確保營運連續性計畫（BCP），強制規定所有跨國運行的 AI 模型皆必須在同一國家內建置實體異地備援機房。	選項(4) 網路服務供應商必須向政府承諾，其雲端資料庫在進行跨國傳輸與演算法推論時能達到特定的最小頻寬與速度門檻。
正確答案：2	

公務AI共通核心能力認證樣題

題號：26

主層面：2.AI 政策法制	次層面：2.1AI 政策與治理指引
題型：單選	難度：中
題幹： 依據我國《人工智慧基本法》之立法精神，在評估生成式 AI 應用的風險與衝擊時，政府特別關注其對「個人層次」所造成的影響。下列何者屬於此範疇內最核心的防護目標？	
選項(1) 確保國內 AI 新創企業在面對跨國科技巨頭時，仍能維持足夠的商業市場競爭力與市占率。	選項(2) 防範演算法偏見、隱私侵害或長期情感依附等問題，對公民心理健康、人格尊嚴與基本權利造成的損害。
選項(3) 釐清開源與閉源大型語言模型在商用轉化過程中的軟體授權模式（如 MIT, Apache 2.0）爭議。	選項(4) 評估企業導入自動化系統後，整體產業鏈在資本支出（CapEx）與硬體伺服器投資規模上的變化。
正確答案：2	

公務AI共通核心能力認證樣題

題號：27

主層面：2.AI 政策法制	次層面：2.1AI 政策與治理指引
題型：單選	難度：中
題幹： 在《人工智慧基本法》的倫理治理規範中，「透明及可解釋性（Transparency and Explainability）」原則的核心意涵為何？	
選項(1) 強制要求所有企業必須向社會大眾無條件完整公開 AI 系統的所有底層神經網路原始碼（Source Code）。	選項(2) 確保系統運作機制應具備足夠的資訊揭露，使受決策影響者能清楚理解 AI 生成內容或決策背後的運算基礎與邏輯影響。
選項(3) 僅規範系統開發者須針對監管機關進行技術說明，無需對一般使用者提供任何解釋或諮詢管道。	選項(4) 系統運作期間無需採取任何預防與告知，僅要求在發生重大災難性事故後，需對外召開記者會公布調查報告。
正確答案：2	

公務AI共通核心能力認證樣題

題號：28

主層面：2.AI 政策法制	次層面：2.1AI 政策與治理指引
題型：單選	難度：難
題幹： 在當前台灣的《著作權法》框架下，針對利用受著作權保護之作品進行大型語言模型（LLM）的訓練行為，下列關於法律適用與風險的敘述，何者最為精確？	
選項(1) 現行著作權法已針對 AI 訓練活動設立了全面且無爭議的商業性利用之著作權例外（Exceptions）條款。	選項(2) 因欠缺具體的法律免責條款，目前 AI 訓練資料使用是否構成侵權，仍須視個案「合理使用」標準而定，具備法律不確定性。
選項(3) 根據現行行政解釋，只要是為了提升機器學習準確率，任何形式的商業性資料訓練皆被視為符合合理使用而完全合法。	選項(4) 司法院智慧財產及商業法院已針對生成式 AI 的訓練行為，建立了一套具備強制約束力的標準判例（Precedent）。
正確答案：2	

公務AI共通核心能力認證樣題

題號：29

主層面：2.AI 政策法制	次層面：2.1AI 政策與治理指引
題型：單選	難度：中
題幹： 隱私保護法規（如 GDPR）與人工智慧倫理框架中，均將「資料最小化原則（Data Minimization）」視為資料治理的核心規範。當企業在開發生成式 AI 系統時，該原則在實務操作上的具體指導意義為何？	
選項(1) 為了極大化模型的泛化能力與降低潛在偏見，企業應採取無差別抓取策略，盡可能蒐集跨越原始授權範圍的多元使用者行為資料。	選項(2) 企業在定義 AI 模型訓練或推論任務時，應嚴格審視並界定範圍，僅蒐集與處理達成該特定業務目的所「直接相關且絕對必要」的個人數據。
選項(3) 秉持「數據即資產」的戰略，將所有透過 AI 服務互動所蒐集到的使用者數據進行無期限之永久保存，以備未來未知商業場景的模型迭代使用。	選項(4) 為了維持機器學習特徵的真實性，應優先且直接將未經任何清洗、去識別化或雜訊過濾的原始機敏數據餵入基礎模型中進行訓練。
正確答案：2	

公務AI共通核心能力認證樣題

題號：30

主層面：2.AI 政策法制	次層面：2.1AI 政策與治理指引
題型：單選	難度：易
題幹： 某機關討論是否導入 AI 協助處理申請案件。依《人工智慧基本法》之「人類自主」及「問責」原則，AI 應以解決實際需求並建立適當治理機制為前提。依 AI 導入生命週期，較適當的起點為何？	
選項(1) 先選定市場上常見 AI 工具，再調整業務流程配合工具功能。	選項(2) 先要求廠商提出模型架構，再由機關決定資料是否足夠。
選項(3) 先從業務痛點出發，明確定義欲解決問題，並評估 AI 是否為合適方案。	選項(4) 先小規模上線累積使用資料，再補充風險評估與治理措施。
正確答案：3	

公務AI共通核心能力認證樣題

題號：31

主層面：2.AI 政策法制	次層面：2.1AI 政策與治理指引
題型：單選	難度：中
題幹： 某機關 AI 系統上線後，初期模型表現穩定；半年後因申請案件型態改變，錯誤判斷逐漸增加。依《人工智慧基本法》所揭示之「安全與風險管理」及「問責」原則，AI 系統應建立適當監督與管理機制。業務單位認為承辦人可人工修正，資訊單位則建議重新驗證模型。下列何者較符合風險管理觀點？	
選項(1) 應依模型關鍵性、應用場景變化及資料時效性，建立定期驗證與必要調整機制。	選項(2) 若承辦人能修正 AI 判斷，代表人機協作機制仍可運作，暫不需重新驗證模型。
選項(3) 可先累積更多錯誤案例，等達到一定數量後再判斷是否調整模型。	選項(4) 若原驗收測試資料仍可重現合格結果，表示模型本身未必需要重新檢視。
正確答案：1	

公務AI共通核心能力認證樣題

題號：32

主層面：2.AI 政策法制	次層面：2.1AI 政策與治理指引
題型：單選	難度：中
題幹： 某機關導入 AI 輔助排序民眾申請案件後，整體平均處理時間縮短，但監測資料顯示偏遠地區案件較常被排在低優先順序。依《人工智慧基本法》之「公平與不歧視」、「透明與可解釋」及「問責」原則，下列何者最符合 AI 審核與治理要求？	
選項(1) 因整體處理時間已有改善，可先維持模型，並加強對偏遠地區案件的人工抽查。	選項(2) 將偏遠地區案件改由人工排序，以避免 AI 模型在該類案件上持續產生爭議。
選項(3) 檢視資料處理、排序規則、模型偏誤、使用紀錄與公平性，並建立改善追蹤機制。	選項(4) 於對外服務說明中揭露 AI 僅供輔助排序，以降低民眾對排序結果的誤解。
正確答案：3	

公務AI共通核心能力認證樣題

題號：33

主層面：2.AI 政策法制	次層面：2.2AI 法規與資安風險管理
題型：單選	難度：易
題幹： 數位平台在歐盟境內託管 AI 產出內容時，若已接獲權利人對於侵權內容的「合法通知（Takedown Notice）」卻遲遲不執行屏蔽程序，該平台將會失去何種法律保護優勢，並導致何種結果？	
選項(1) 將導致平台演算法排名下降，進而造成平台使用活躍用戶數的嚴重流失。	選項(2) 將面臨歐盟主管機關的最高等級行政處分，立即撤銷其在歐盟境內所有的商業營運執照。
選項(3) 將失去避風港責任豁免的保障，導致必須承擔後續損害賠償與相關侵權行為的法律責任。	選項(4) 被監管機構強制列入黑名單，禁止該平台在歐盟境內執行任何形式的廣告收入營運。
正確答案：3	

公務AI共通核心能力認證樣題

題號：34

主層面：2.AI 政策法制	次層面：2.2AI 法規與資安風險管理
題型：單選	難度：難
題幹： ISO/IEC 42001 為國際上首個針對人工智慧管理系統（AIMS）的標準。根據該標準，企業在進行 AI 風險評估時，應全面考量下列哪一組合的要素？	
選項(1) 專注於敏捷專案管理的「硬性指標」，包含開發預算上限、時程控制與人力分工。	選項(2) 考量 AI 系統特性、訓練與推論資料、利害關係人（使用者）、營運流程及相關法規遵循等全方位面向。
選項(3) 僅專注於硬體基礎設施的穩定性指標，包含機房電力供應、冷卻系統妥善率與晶片故障率。	選項(4) 僅專注於數據模型的效能指標，即模型產出的精確度（Accuracy）是否已達到特定業務門檻（如 90%）。
正確答案：2	

公務AI共通核心能力認證樣題

題號：35

主層面：2.AI 政策法制	次層面：2.2AI 法規與資安風險管理
題型：單選	難度：難
題幹： 「對抗性攻擊（Adversarial Attack）」是 AI 模型在資安層面的重大隱憂。關於其在生成式 AI 模型中展現出的核心特性與危害，下列敘述何者正確？	
選項(1) 透過過多的網路請求導致模型因算力不堪負荷，導致整體推論與生成速度大幅緩慢。	選項(2) 在輸入數據中加入人類視覺無法察覺的雜訊擾動，進而精確欺騙模型輸出錯誤、惡意或違背倫理的內容。
選項(3) 針對模型的訓練 pipeline 發動分散式阻斷服務攻擊，迫使系統在反向傳播的訓練過程中自動停止更新。	選項(4) 透過誘發模型產生特定字詞組合，自動繞過版權偵測系統並強制模型輸出已註冊版權的機密文件。
正確答案：2	

公務AI共通核心能力認證樣題

題號：36

主層面：2.AI 政策法制	次層面：2.2AI 法規與資安風險管理
題型：單選	難度：難
題幹： 依據《人工智慧基本法》對於風險分級管理之規範，，當 AI 應用涉及高風險領域且其運作結果可能影響人民權益時，政府應建立救濟、補償或保險等責任保障機制。下列哪一種應用情境最符合此類管理要求？	
選項(1) 已部署於關鍵公共領域（如醫療診斷、金融授信），且其裁決結果將對個人權益產生實質性影響的高風險 AI 系統。	選項(2) 僅處於封閉校園實驗室環境下，針對學術界限內之理論邏輯推演所進行之基礎科學研究。
選項(3) 在嚴格沙盒環境中進行的模擬測試，且未處理任何真實個人機敏數據或影響第三方公民決策權。	選項(4) 僅供企業內部開發人員使用，作為協助撰寫程式碼或進行自動化測試等非直接決策之輔助性工程工具。
正確答案：1	

公務AI共通核心能力認證樣題

題號：37

主層面：2.AI 政策法制	次層面：2.2AI 法規與資安風險管理
題型：單選	難度：難
題幹： 歐盟《人工智慧法案（EU AI Act）》將部分倫理原則轉化為具備強制力的法律義務。針對「通用人工智慧（GPAI）」與「生成式 AI」，下列何者屬於該法案明定的「強制性法律義務」，而非僅是軟性的倫理倡議？	
選項(1) 鼓勵開發企業於內部自發性建立「倫理審查委員會」，以維持 AI 開發的道德紀律。	選項(2) 倡導產業界應遵循「以人為本（Human-centric）」的發展願景，並簽署相關自律聲明。
選項(3) 建議企業在模型訓練過程中，應盡可能採取多元數據採樣以避免演算法產生潛在偏見。	選項(4) 要求生成式 AI 供應商必須採取技術手段，確保其生成之內容（如深度偽造影音）具備可檢測性或明確之 AI 生成標示。
正確答案：4	

公務AI共通核心能力認證樣題

題號：38

主層面：2.AI 政策法制	次層面：2.2AI 法規與資安風險管理
題型：單選	難度：難
題幹： 某承辦人將含有民眾身分證字號與聯絡電話的檔案直接上傳至公開生成式 AI 平台進行分析。依生成式 AI 使用之資安與個資保護要求，下列何者最可能產生的風險？	
選項(1) AI 平台可協助分析資料內容，因此可降低個資管理需求	選項(2) AI 平台可能保存或處理上傳資料，造成個人資料未經授權使用之風險
選項(3) 將資料交由公開 AI 平台處理，可減少機關內部資料管理負擔	選項(4) 只要分析結果未包含完整個資，即不會產生資料安全問題
正確答案：2	

公務AI共通核心能力認證樣題

題號：39

主層面：2.AI 政策法制	次層面：2.2AI 法規與資安風險管理
題型：單選	難度：難
題幹： 某機關導入 AI 系統時，同步建立權限控管、身分驗證及存取紀錄機制，以限制資料使用範圍並追蹤系統操作行為。此作法主要強化哪一項 AI 治理要求？	
選項(1) 透過資料與模型檢視降低 AI 產生不公平結果的可能性	選項(2) 透過存取控管與操作追蹤降低未授權使用及資料安全風險
選項(3) 透過揭露模型運作方式提升使用者理解 AI 判斷依據	選項(4) 透過調整系統設定提升 AI 應用於不同情境的適應能力
正確答案：2	

公務AI共通核心能力認證樣題

題號：40

主層面：2.AI 政策法制	次層面：2.2AI 法規與資安風險管理
題型：單選	難度：難
題幹： 某機關規劃使用 AI 系統協助核發社會福利補助資格。該系統需處理民眾身分、財務及資格相關資料，且 AI 判斷結果可能影響民眾權益，且可能影響補助資格認定。依《人工智慧基本法》及《AI 產品與系統評測參考指引》，下列何者最適當？	
選項(1) 將資料交由 AI 自動分析處理，以避免人工接觸敏感資訊	選項(2) 只要限制系統使用人員範圍，即可降低主要治理風險
選項(3) 建立資料存取控管、操作紀錄、風險評估及人工覆核機制，再依程序導入使用	選項(4) 優先提升模型準確率，資料管理與人工覆核可於後續補充
正確答案：3	

公務AI共通核心能力認證樣題

題號：41

主層面：2.AI 政策法制	次層面：2.3AI 應用下的倫理準則與智慧財產實務
題型：單選	難度：中
題幹： 在司法與公權力領域導入「AI 輔助判決系統」時，從確保司法獨立性、人權保障與問責原則之角度，下列何種應用情境最嚴重違背法治精神且應被絕對禁止？	
選項(1) 律師為防範洩密，嚴格禁止將含有客戶機密個資的訴訟文件輸入至未經企業授權的開放式 AI 模型中進行分析。	選項(2) 法官在審理複雜案件時，參考 AI 工具所提供之過往相似判例檢索與爭點分析報告，並作為輔助性之客觀資訊。
選項(3) 司法機構要求法務 AI 系統開發商提供可解釋性（Explainability）報告，使法庭判決所依賴之演算法輔助過程具備高度透明性。	選項(4) 司法機關為追求極致的結案效率，直接採信演算法生成之裁決結果，將其奉為具備最終強制力之「AI 法官」並實質放棄人類獨立裁量權。
正確答案：4	

公務AI共通核心能力認證樣題

題號：42

主層面：2.AI 政策法制	次層面：2.3AI 應用下的倫理準則 與智慧財產實務
題型：單選	難度：難
題幹： 依據美國著作權局（USCO）於 2023 年發布之「人工智慧生成作品註冊指南」，關於包含 AI 生成內容之作品是否具備著作權，下列敘述何者正確？	
選項(1) 只要使用者輸入了極為複雜且長篇的提示詞（Prompt），該 AI 完全自動生成之圖文即自動滿足「人類創作」之註冊門檻。	選項(2) 只要作品中混合了人類創作與 AI 生成內容，無論 AI 占比多寡，整份作品即具備完整著作權，審查官無須依個案進行實質判斷。
選項(3) 著作權之保護僅及於人類心智創作之成果；若作品包含 AI 元素，必須依個案評估人類在「選擇、編排與修改」過程中的實質參與程度。	選項(4) 鑑於 AI 生成結果具備高度不可預測性，使用者等同於是在指導一項具備創造力的工具，因此該生成物一律認定為操作者之專屬創作。
正確答案：3	

公務AI共通核心能力認證樣題

題號：43

主層面：2.AI 政策法制	次層面：2.3AI 應用下的倫理準則 與智慧財產實務
題型：單選	難度：易
題幹： 企業在導入 AI 生成內容系統時，為了由上而下落實「負責任的 AI（Responsible AI）」並確保系統輸出的倫理性，下列何項管理作為最為完善且具備實效？	
選項(1) 透過外包合約，僅採用市面上開源社群評分最高的第三方預訓練模型，並假設其已具備絕對之道德免疫力。	選項(2) 在系統登入介面強制要求所有終端使用者簽署免責聲明（EULA），將所有偏見與侵權的法律風險片面轉嫁予使用者。
選項(3) 制定標準化的「人類參與流程審查（Human-in-the-Loop）」作業規範，並建立跨部門（包含法務、技術與業務）的內容評估與風險緩解機制。	選項(4) 斥資採購最高規格之儲存硬體，定期針對 AI 生成的所有數據進行快照備份，確保內容即便違規也不會遺失或遭竄改。
正確答案：3	

公務AI共通核心能力認證樣題

題號：44

主層面：2.AI 政策法制	次層面：2.3AI 應用下的倫理準則 與智慧財產實務
題型：單選	難度：難
題幹： 當數位服務平台託管了由終端使用者利用生成式 AI 產出之圖文，並遭遇第三方提出版權侵權申訴時。依據美國《數位千禧年著作權法（DMCA）》之架構，下列哪一項機制是該平台主張「避風港（Safe Harbor）法律責任豁免」的絕對先決條件？	
選項(1) 平台必須針對使用者產出的每一件生成作品，預先提撥並支付全額版權使用費給潛在之權利人公會。	選項(2) 平台必須建立並嚴格執行有效的「通知與取下（Notice and Takedown）」機制，對權利人的合法侵權檢舉做出迅速移除之回應。
選項(3) 平台必須導入強制性的自動化版權過濾演算法，對使用者即將發布的所有 AI 生成內容進行全面的預防性法規阻擋。	選項(4) 平台必須要求所有 AI 模型供應商提供具備法律效力之切結書，證明其訓練數據皆已取得完整之專利與著作權商業授權。
正確答案：2	

公務AI共通核心能力認證樣題

題號：45

主層面：2.AI 政策法制	次層面：2.3AI 應用下的倫理準則 與智慧財產實務
題型：單選	難度：中
題幹： 鑑於當前各國智慧財產權法規對生成式 AI 之保護邊界尚在演進中，企業在將 AI 生成內容投入商用時，下列哪一種認知最符合「合規與風險控管」之實務標準？	
選項(1) 依據現行法規，所有經由 AI 生成之內容，無論人類介入程度多深，皆絕對不具備任何著作權保護。	選項(2) 只要使用者擁有該 AI 平台之合法帳號，其下達指令所生成之所有內容，著作財產權一律歸屬於該使用者。
選項(3) 由於各國對 AI 生成物之原創性與權利歸屬判定標準不一，企業對其商業使用應保持高度謹慎並事前評估侵權風險。	選項(4) 只要企業採用的是付費企業版（Enterprise）的 AI 服務，其產出物即由服務商背書，自動獲取完整且排他的著作權保護。
正確答案：3	

公務AI共通核心能力認證樣題

題號：46

主層面：2.AI 政策法制	次層面：2.3AI 應用下的倫理準則 與智慧財產實務
題型：單選	難度：中
題幹： 歐盟發布之《可信賴人工智慧倫理準則（Ethics Guidelines for Trustworthy AI）》揭示了七大關鍵要素。下列哪一項屬於該準則中明確要求之核心支柱？	
選項(1) 極大化自主決策權：強調應將演算法的自主程度提至最高，以完全排除人類操作員的主觀干擾。	選項(2) 技術透明度與資訊揭露（Transparency）：要求決策邏輯應具備可解釋性，且 AI 產出物應明確標示其非人類所為。
選項(3) 數據利用自由最大化：主張為了促進科技進步，大數據蒐集應獲得豁免，不受既有個資規範的拘束。	選項(4) 企業營收優化策略：強調 AI 應將協助企業優化營運效率與利潤最大化，視為最高道德指導原則。
正確答案：2	

公務AI共通核心能力認證樣題

題號：47

主層面：2.AI 政策法制	次層面：2.3AI 應用下的倫理準則 與智慧財產實務
題型：單選	難度：中
題幹： 日本在推動人工智慧發展上，其《著作權法》第 30 條之 4 制定了相對寬鬆的合理使用規範，常被各國智財界稱為「機器學習的天堂」。下列哪一項論述最正確反映了日本現行法對 AI 訓練使用受保護作品的法律立場？	
選項(1) 嚴格限制僅有非營利性質的教育單位與公家機關，方可未經授權使用版權作品進行 AI 模型訓練。	選項(2) 秉持權利人絕對優先原則，無論是商業或非商業應用，皆必須事先取得著作權人之實質同意與簽署授權金協議。
選項(3) 明定若為「資訊解析（Information Analysis）」之目的，原則上不區分營利或非營利性質，皆可未經授權利用受保護之著作進行訓練，除非不當損害權利人利益。	選項(4) 僅限於由日本政府出資贊助之國家級 AI 研究計畫，方享有未經授權使用版權素材的絕對侵權豁免權。
正確答案：3	

公務AI共通核心能力認證樣題

題號：48

主層面：2.AI 政策法制	次層面：2.3AI 應用下的倫理準則 與智慧財產實務
題型：單選	難度：中
題幹： 某機關承辦人使用非公務用之生成式 AI 軟體，僅透過輸入「請產生一張臺灣山林風景宣傳圖」等簡單提示詞，即自動生成圖片。由於該承辦人後續並未對生成之圖片之構圖設計與內容加以修改，便直接對外公開發布，但卻仍主張該圖片屬機關享有著作權之著作。請問以下之見解，何者正確？	
選項(1) 機關人員輸入提示詞並生成圖片，即當然取得該圖片之著作權。	選項(2) 該圖片完全由 AI 自主生成，且欠缺人類精神力之創作，不得主張受著作權法保護。
選項(3) AI 所生成之圖片及文本，均歸屬使用者所有。	選項(4) AI 所生成之圖片與文本，其著作權皆應歸屬該 AI 軟體開發廠商所有。
正確答案：2	

公務AI共通核心能力認證樣題

題號：49

主層面：2.AI 政策法制	次層面：2.3AI 應用下的倫理準則 與智慧財產實務
題型：單選	難度：中
題幹： 某機關委外製作宣導海報，但承辦人發現廠商交付之圖片疑似以AI生成，且畫面中出現與知名卡通角色高度近似之人物。以下何者為適當之處理方式？	
選項(1) AI 生成的內容並不會侵害著作權，自可直接驗收、張貼。	選項(2) 若廠商已自行表示該海報是 AI 生成，即可免除機關審查責任。
選項(3) 應要求廠商說明製作圖片素材之資料來源、授權狀態，以及是否涉及第三人之權益。	選項(4) 政府機關使用 AI 製作宣導海報具有公益性，故不需考慮是否侵害著作權。
正確答案：3	

公務AI共通核心能力認證樣題

題號：50

主層面：2.AI 政策法制	次層面：2.3AI 應用下的倫理準則 與智慧財產實務
題型：單選	難度：中
題幹： 某市民政局規劃以生成式 AI 協助製作政策懶人包，內容包括政策圖卡、簡報文字及宣傳用語。該業務承辦人先利用 AI 生成初稿與構圖，並進一步修改、編排部分圖片及文本，同時還加入該機關自行設計之政策宣傳識別標示。在此政策懶人包正式對外公布前，主管要求審查是否需揭露 AI 參與情形及檢核權利風險。請問以下何者係正確判斷標準？	
選項(1) 僅需人員修改過，即可完全不必揭露 AI 參與情形。	選項(2) AI 生成內容一旦經機關對外發布後，即自動成為機關所有之著作。
選項(3) 只要是政府公共政策宣導，屬於公益範疇，不涉及商業利用，故無須檢核著作權或授權問題。	選項(4) 應依照 AI 參與生成程度、發布目的及機關規範，適當揭露資訊、檢核資料來源及授權依據。
正確答案：4	