

台北外匯市場發展基金會委託計畫

央行對人工智慧應用的治理 與執行情形

中央銀行業務局

林正弘

日期：中華民國115年6月

本文純屬作者個人意見，與服務單位無關，如有任何疏漏或謬誤，概由作者負責。

摘要

近年來人工智慧(Artificial Intelligence, AI)快速發展，大幅改變了人們與電腦互動及處理非結構化數據的方式，並對經濟與金融體系產生深遠的影響。各國央行已開始將 AI 整合至自身的分析工具與實務營運中，目前主要央行的 AI 應用大多定位於輔助型工作與自動化流程，實務上涵蓋總體經濟與通膨之即時預測、風險評估與預警系統、改善申報資料品質與偵測異常值，及洗錢防制與支付系統監控等多項關鍵領域，尚未直接運用於核心決策層面。

然而，導入 AI 也為央行帶來多重挑戰與脆弱性，主要風險包括：模型本身的幻覺、推理錯誤與缺乏透明度；機密外洩與新型態網路攻擊等資安威脅；過度依賴少數科技巨頭的第三方風險，以及法遵之不確定性與人員技能不足等營運風險。

在治理方面，目前央行的 AI 應用多屬探索階段，並主要以各業務部門臨時及分散式的方式推動。基於機密性、資料隱私與網路安全等考量，多數央行對外部 AI 工具採取嚴格的限制或禁用措施，技術架構上則偏好使用開源 AI 模型與自建內部伺服器為基礎設施。

鑑於 AI 領域變化快速，治理上應具備動態性及適應性之架構，並定期審查，以確保其有效並符合組織的目標與風險管理措施。

目 次

摘要

壹、 前言.....	1
貳、 央行應用 AI 案例.....	2
一、 LLMs 與生成式 AI 的應用.....	2
二、 機器學習(MACHINE LEARNING)的應用.....	3
參、 應用 AI 之風險.....	6
一、 模型風險.....	6
二、 資訊安全風險.....	7
三、 第三方風險 (THIRD-PARTY RISKS).....	8
四、 營運風險.....	8
五、 其他風險.....	9
肆、 AI 治理.....	9
一、 治理現況.....	9
二、 治理框架.....	12
伍、 結論.....	16
參考文獻.....	18

壹、前言

大型語言模型(Large Language Models ,LLMs)的出現改變了人們與電腦互動的方式，由程式碼和程式設計介面，轉為普通的文字與語音，這種透過普通語言進行對話的能力，以及生成式 AI 於創作方面展現的類人能力，使 AI 的運用快速普及。

隨著算力的增長與成本的降低，AI 可將日常非結構化的資料(例如文字、圖像、影片或音樂)轉化為可計算的數學結構，使早期難以勝任的任務得以完成，並對經濟與金融體系產生深遠的影響，此技術能力有助於央行深化政策工具的運用，優化其核心營運機制的效能。

AI 的經濟潛力引發了一場投資與應用潮，用戶快速地採用，企業也將 AI 整合到日常營運中，全球調查顯示，所有行業都在使用生成式 AI 工具。為此，他們大量投資 AI 技術以量身打造其特定需求，並積極招聘具備 AI 相關技能員工，預計此趨勢將持續快速成長。

AI 的廣泛採用可能增強企業根據總體經濟變化快速調整價格的能力，進而對通膨產生影響，此外，AI 將影響金融體系以及生產力、消費、投資和勞動力市場，這些因素直接影響物價與金融穩定，也是央行職責的範圍。

商業銀行等金融機構採用 AI 工具，將改變其與央行互動及接受監管的方式。央行須密切觀察技術進步帶來的影響，也應成為技術本身的使用者，作為觀察者，央行需要研判 AI 對總體供需的衝擊，進而評估經濟活動的後續變化；作為使用者，則需要將 AI 與非傳統數據有效整合至自身的分析工具中。

本報告除本章前言外，第貳章說明各國央行應用 AI 的案例，第參章為應用 AI 的相關風險，第肆章為 AI 治理的現況及框架，最後則為結論。

貳、央行應用 AI 案例

一、LLMs 與生成式 AI 的應用

目前各國央行及國際組織已開始將 LLMs 與生成式 AI(如 GPT、Gemini、Claude)投入實務研究與日常營運，具體案例包含以下幾個核心面向：

(一) 美國聯準會(Fed)

分析 FOMC 會議紀錄的政策關注點，運用生成式 AI 分析 2010 年至 2024 年 FOMC 會議紀錄之各項總體經濟主題的討論比重變化，並比較 11 種生成式 AI 模型(包含多個版本的 GPT、Gemini 與 Claude 等)的準確性，結果發現這些模型的準確度極高。

(二) 國際清算銀行(BIS)與歐洲央行(ECB)

BIS 與 ECB 合作預測核心通膨，利用 LLMs(包含 GPT 及 BERT)對 2002 年至 2023 年 ECB 貨幣政策記者會新聞稿進行分析。研究發現，將新聞稿的前瞻性文字資訊透過 LLMS 轉化並納入預測模型後，能有效提升該模型對歐元區核心通膨的預測準確度。

(三) 優化營運溝通與聊天機器人服務

在企業日常行政業務中，LLMs 被廣泛應用於摘要冗長文件、製作簡報、撰寫研究報告，以及分析大眾的常見問題。此外，央行也積極開發聊天機器人(Chatbots)，對外應用於回覆民眾的查詢問題；對內則作為員工虛擬助理，協助員工快速尋找內部流程、公司福利及政策規定等資訊。

小結：儘管 GPT 與 LLMs 展現出強大的文字處理能力，但考量

到模型的可解釋性、回覆穩定性，以及潛在的偏誤風險，目前主要央行對生成式 AI 的應用，仍嚴格定位於內部文件摘要、會議紀錄生成等「輔助型工作」，AI 也能增強程式設計能力(Gambacorta 等，2024)，並協助工作流程自動化，提升效率與生產力，但尚未將其直接運用於核心的決策制定層面。

二、機器學習(Machine Learning)的應用

在全球主要央行與國際組織，機器學習演算法是 AI 模型發展的基礎階段，已被廣泛應用於多項核心業務。根據來源資料，機器學習的具體應用案例可歸納為以下四大領域：

(一) 總體經濟與通膨之即時預測(Nowcasting & Forecasting)

央行利用機器學習模型處理高頻數據(如每日新聞、商品細項價格)，以改善傳統官方數據發布時間落後的問題：

1. 英格蘭銀行(BoE)

將英國每日新聞轉換為詞頻向量(term frequency vector)後，輸入嶺迴歸(Ridge)、最小絕對值收斂與選擇算法迴歸模型(Least Absolute Shrinkage and Selection Operator Regression, LASSO)、支援向量機(SVM)、隨機森林(RF)及神經網路(NN)等模型，成功預測未來 3 至 9 個月的 GDP、失業率、企業投資與消費者物價指數(CPI)。此外，BoE 也利用非線性機器學習模型預測英國總體及服務業通膨。

2. ECB

將每日新聞轉化為情緒指標(sentiment index)，並輸入 RF、提升法(Boosting)及神經網路模型中，即時預測實質 GDP 季增率，在金融

危機與疫情期間展現了優異的預測能力。

3. 國際貨幣基金組織(IMF)

結合多項總體經濟指標，使用 Ridge、LASSO、SVM、RF 及 NN 等模型，即時預測奧地利等 6 國之實質 GDP 季增率，證實在市場高度波動時，機器學習模型較能捕捉市場波動的轉折點。

(二) 風險評估與預警系統(Early Warning Systems)

機器學習擅長挖掘數據背後的複雜規律，被廣泛用於建構企業違約、個別金融機構危機及總體金融危機之預警系統。

1. ECB

為了監管「重要性較低銀行(Less-significant institutions, LSIs)」，ECB 使用決策樹(Decision Tree, DT)及 LASSO，結合監理報表與總體指標來預測銀行經營危機，並將該創新決策樹模型作為日常監理工具。

2. Fed

運用 NN、SVM、RF 與梯度提升(Gradient Boosting, GB)等模型來預測上市企業違約機率，其中 RF 表現最佳。Fed 進而以此建構「非金融企業部門財務狀況惡化指數(Non-Financial Corporate Health Index, NFCH)」，用於預警總體經濟衰退。

3. BoE

使用 k 近鄰演算法(k-nearest neighbors, kNN)、RF、GB 與 SVM 預測英國各家銀行是否會陷入經營危機。

4. ECB 與 BOE 合作

利用 DT、RF、極限隨機樹(Extremely Randomized Trees, ERTs)及 NN 分析跨國總體經濟與市場資料，預先辨識各國發生金融危機的機率。

(三) 改善申報資料品質與偵測異常值(Anomaly Detection)

處理龐大且高顆粒度¹的申報資料時，機器學習模型能自動化檢測錯誤與異常。

1. ECB

針對具有大量類別特徵與異質數據的「貨幣市場統計資料(MMSR)」，結合加權最小平方法(Weighted Least Squares, WLS)迴歸、孤立森林(Isolation Forest, IF)、階層式密度分群法(Hierarchical Density-Based Spatial Clustering of Applications with Noise, HDBSCAN)及 XGBoost 等多種機器學習演算法自動化篩選異常訊號，並內嵌至管理流程中。

2. 德國央行

建構包含 IF、kNN、基於密度之含離群值空間分群法(Density-Based Spatial Clustering of Applications with Noise, DBSCAN)、單類支援向量機(One-Class SVM)、自編碼器(Autoencoder)等十多種演算法的機器學習工作流程，以自動化篩選高顆粒度統計資料的異常訊號。

¹ 「高顆粒度資料」(Granular Data)指的是高度細緻、微觀且資料量龐大的數據集，例如商品細項價格、每日股票報酬，以及逐筆的申報或交易資料等。

3. 葡萄牙央行

使用 IF 演算法，針對中央信用貸款登錄系統的龐大逐筆資料，尋找貸款金額異常值。

(四) 洗錢防制與支付系統監控(AML & Payment Systems)

BIS 透過 Aurora 及 Hertha 計畫，使用邏輯迴歸、人工神經網路(ANN)、圖像神經網路(GNN)、IF 及 XGBoost 等模型分析跨國匯款與支付系統資料。結果顯示機器學習模型較傳統規則式的監控，能偵測更多可疑交易、降低誤報率，並辨識新型金融犯罪模式。

參、應用 AI 之風險

應用 AI 可能帶來新風險與脆弱性，並擴大現有風險，而風險程度取決於 AI 模型的複雜度、所存取之資訊及資料存取權限。風險主要有模型風險、資訊安全風險、第三方風險、營運風險等，說明如下。

一、模型風險

(一) 解釋性與可靠性問題：AI 模型的準確性高度依賴輸入的資料品質，不良的資料或「過度配適²(overfitting)」會導致模型產生偏見與錯誤預測。

(二) 幻覺與推理錯誤：AI(特別是生成式 AI)容易產生「幻覺」，亦即在缺乏邏輯關聯的數據中識別出虛假模式，進而捏造出不正確或毫無意義的結果。

² 模型在訓練過程中表現極佳，但用未見過的資料測試時卻無法做出正確的結果，常發生在模型過於複雜，或是訓練的資料量不夠大時。

- (三) 缺乏可重複性：即使輸入完全相同的提示，生成式 AI 的輸出仍具備隨機性，難以確保不同時間點的回覆一致性。
- (四) 過度自信：AI 產出的結果通常看似專業且具權威性，這容易削弱使用者的盡職調查工作，過度依賴 AI 而提高決策錯誤的風險。
- (五) 透明度與問責困難：許多 AI 模型具有「黑盒子」特性，缺乏透明度；部分機器學習演算法具有自我修正特性，會自動整合新資料並改變演算法，皆可能產生設計者無法預料的結果，導致治理與歸責的問題。

二、資訊安全風險

- (一) 隱私與機密外洩：使用者在輸入提示詞(prompts)時可能傳送敏感或專有資訊，這些資料可能被 AI 工具儲存用於後續訓練，並意外洩漏給其他使用者，造成隱私侵犯與嚴重的資安漏洞。
- (二) 新型態的網路攻擊：導入 AI 系統會擴大資訊系統的攻擊面，帶來如「提示注入(prompt injection)³」、「訓練資料中毒(data poisoning)」、模型阻斷服務及模型竊取等新型威脅。
- (三) 惡意活動與深偽技術：惡意行為者可利用 AI 生成更具說服力的網路釣魚郵件、編寫惡意軟體，或製造深度偽造(deepfakes，例如偽造央行總裁的影音)來進行詐騙與假訊息傳播。

³ 攻擊者刻意設計並輸入特定的提示文字詞，藉此操縱 AI 模型，使其可避開原先設定好的安全限制或護欄，進而做出不安全的行為或產出不當的內容。

- (四) 資料治理挑戰：AI 需要處理龐大、高流動性且多樣化的非結構化資料，大幅增加了資料品質控制與治理的難度。

三、 第三方風險 (Third-Party Risks)

- (一) 供應商集中度風險：由於開發先進 AI 模型成本極高，往往由少數科技巨頭壟斷市場。過度依賴單一或少數供應商會削弱央行的談判能力，且一旦該供應商發生故障或遭受網路攻擊，將直接造成營運中斷。
- (二) 隱私與合規風險：外部供應商若存取機敏資料，可能會發生有意或無意的資料濫用；若資料處理或儲存於不受監管的系統或海外，更會帶來法遵風險。
- (三) 缺乏彈性：依賴外部專有軟體會使機構無法配合自身業務需求，靈活修改或調整 AI 技術。

四、 營運風險

- (一) 法律與法遵不確定性：AI 生成的內容可能無法符合著作權保護等法律規範。此外，AI 模型的不可解釋性，會使機構在需要提供決策依據時面臨困難，進而加劇法規遵循風險。
- (二) 人員技能缺口：央行員工若缺乏充足的數位技能，對 AI 模型的深入理解與安全使用的認知，不僅會使機構決策時處於劣勢，還會增加對外部供應商的依賴，進而引發資料處理不當的風險。
- (三) 資通訊技術(ICT)風險：新導入的 AI 系統可能與現有的舊系統 (legacy systems) 不相容，導致系統停機或可靠性降低。此外，若賦予 AI 系統過多的功能或權限(例如原本只用於讀取文件的

AI 卻同時擁有刪除權限)，可能會引發意想不到的嚴重破壞。

五、其他風險

(一) 若缺乏清晰的 AI 使用策略，或因 AI 發生錯誤、產生歧視性結果及資料外洩，將引起媒體強烈關注，進而嚴重損害機構的聲譽與公信力。

(二) 環境與道德風險：訓練及運行 LLMs 需要耗費極大的運算資源與電力，將大幅增加碳足跡。同時，由於 AI 缺乏人類的社會與道德常識，若訓練資料存在偏見，可能產生具歧視性或不合倫理的結果。

肆、AI 治理

AI 治理指的是引導組織在採用或開發人工智慧時，確保其安全、負責任且合乎道德的一套政策、規則、框架、原則與標準，核心目的不僅是為了遵守國家與國際的法律法規，更重要的是確保 AI 的應用能與組織的策略與目標保持一致，藉此在推動創新與提升效率的同時，有效識別並平衡相關風險。

一、治理現況⁴

(一) 多數央行將 AI 視為一項政策性議題，尤其亞洲地區將近半數的央行將其列為首要優先事項，在預算上由現行的低於 5%，改為在未來 3 年將至少投入 5% 的預算於 AI 專案，少數央行甚

⁴主要參考 Irving Fisher Committee on Central Bank Statistics (2025), “Governance and implementation of artificial intelligence in central banks” IFC Report, No. 18, April.

至計畫投入超過 10%。

(二) 大多數 AI 的應用仍屬於探索性開發，而非實際作業模式，僅約 4 分之 1 的案例真正全面使用 AI，絕大多數仍屬小規模使用，或尚在開發中，這顯示央行目前主要仍處於 AI 試驗階段，主要原因可能為：1.IT 基礎設施大多僅支援少量適用方案；2. 央行仍處於學習階段；3.多數探索性工作仍建立在功能性層級，若要真正推向營運模式，尚須協調多方利害關係人，如業務單位、IT 與資訊安全部門及預算與人力資源單位。

(三) 預期影響程度與實際執行的專案之間存在落差。理論上，預期影響程度愈高的方向對應的應用數量應相對較高，然而情況並非如此。例如，在網路安全領域，應用案例數量落後於預期影響的程度，這可能反映出實施上的重大障礙，包括 IT 與人力資源的限制(Aldasoro 等，2024)。其他預期影響程度高卻仍缺乏具體專案的領域，還包括支付與個體審慎監理。

(四) 政策規範仍待建置：調查顯示，目前僅約 3 分之 1 的央行制定明確的 AI 使用政策，且極少數對外公開。此外，新興市場經濟體(EMEs)在治理框架與政策的建置上，明顯落後於已開發經濟體(AEs)。AI 之創新發展快速，而制度設計和法制化通常需要時間來建立(OECD，2024c)，因為這涉及多方利害關係人，且必須契合組織的資源配置與戰略核心目標。

(五) 以去中心化為主：目前多數(約 6 至 7 成)的 AI 專案是由各業務部門臨時且「分散式」地推動。這種作法雖能快速回應需求並維持敏捷性，但也容易導致跨部門協調困難、經驗難以分享，

以及整體風險監控的缺口。為了改善分散式管理的缺失，部分央行已開始建立特定 AI 治理架構，通常將統籌責任交給 IT 部門以外的職務(例如資料長 CDO 或數位創新主管)，來協調整合跨部門的 AI 發展。

(六) 偏好開源軟體(open source software)與內部 IT 基礎設施：在技術架構上，央行普遍採用混合模式，但基於成本效益與降低外部依賴的考量，較偏好開源 AI 模型，並以 Python 作為最主要的程式語言。此外，在面對運算資源需求的取捨時，為了確保機敏資訊的絕對安全，超過 40%的央行仍偏好使用自建的內部(on-premises)伺服器，僅將外部雲端服務用於測試或嚴格受限的業務。

(七) 為了保護機密性、資料隱私與網路安全，多數的央行對 AI 工具採取不同程度的限制或禁止措施。央行具體禁止或嚴格限制使用 AI 作法主要有下列 4 種：

1. 封鎖外部 AI 網站：央行會透過網路過濾(web filtering)方案，直接阻擋內部網路連線至特定的外部 AI 網域。
2. 特定機敏區域禁用：在央行內部設定「限制區域」(restricted areas)，明令禁止使用任何 AI 工具。
3. 處理機密資料時禁止連網：當業務涉及機密或敏感資訊時，通常會被禁止在連接外部網際網路的環境下使用 AI 服務；這類資料的操作僅被允許在獨立且安全的環境中進行，例如無對外網路連線的內部(on-premises)基礎設施。

4. 禁止使用外部雲端服務：由於擔憂資料隱私、網路攻擊，以及資料存放在不受控制地點所衍生的資料主權與法律風險，部分央行完全禁止使用外部第三方雲端服務(如公有雲)來運用 AI，並嚴格規定只能使用自建的內部伺服器。

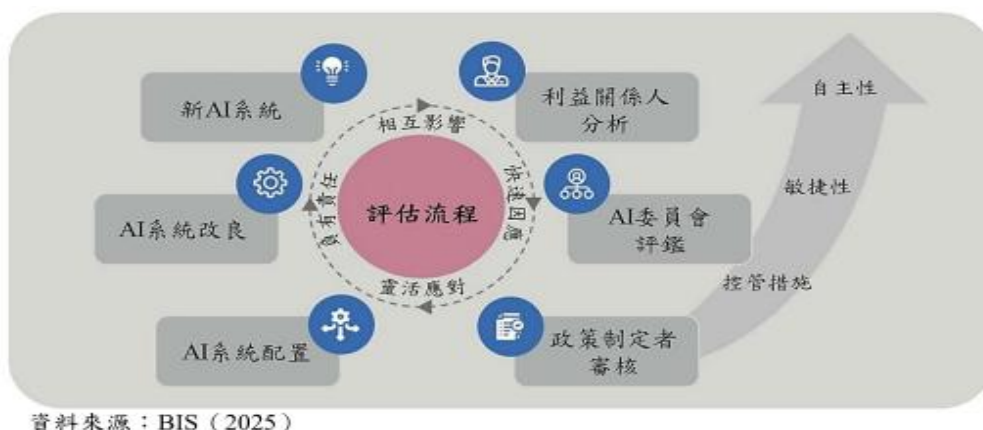
二、治理框架

央行在推動 AI 治理時，不需要從頭建立全新的機制，而是可以將 AI 治理建構在現有的政策、程序與風險管理工具之上。由於 AI 技術發展迅速，央行的治理架構必須具備動態性與彈性，在創新與風險間取得平衡。以下是針對央行 AI 治理提出的核心架構與實務作法：

(一) 建立「適應性 AI 治理架構」(Adaptive AI Governance Framework)(圖 1)此架構應基於三大原則：

1. 控管措施(Control)：組織應設定保護機制，確保 AI 應用符合現行政策與法規要求。
2. 敏捷性(Agility)：落實各部門與職能團隊之權責，以進行分散式或依職權的決策。
3. 自主性(Autonomy)：組織人員可根據即時數據與組織目標，自主做出治理決策。

圖 1 適應性 AI 治理架構



(二) 建議央行應採取的治理行動⁵

1. 建立跨部門 AI 委員會：在建立具體的 AI 治理之前，央行應成立專責的 AI 委員會作為監督機構，以協助、引導及落實後續治理的要求。鑑於 AI 風險影響深遠，該監督機構應跨部門且層級夠高，以確保組織內部的認同。其組織形式包括成立特定委員會以監督 AI 使用，或於既有委員會增加該項職責。
2. 定義負責任 AI 使用的原則：監督 AI 使用的委員會應建立一套原則，說明組織對 AI 使用的核心思想與價值觀。這些原則應清楚表達央行對 AI 使用的整體立場，並引導治理決策。一些標準制定機構(如 ISO)或其他組織(如 OECD)的資訊可協助制訂這些原則。
3. 建立 AI 框架並更新既有指引：根據央行的戰略、目標、價值

⁵ 主要參考 BIS(2025), “Governance of AI adoption in central banks,” BIS Representative Office for the Americas, Jan..

觀及已建立的 AI 原則，建立穩健的 AI 治理框架。在評估不同部門提出的 AI 提案及管理專案時，應納入此框架。

4. 維護 AI 工具清冊：清冊可確保組織內所有的 AI 系統均已列管，並能被妥善管理、監控，且與央行的 AI 治理框架一致。
5. 繪製 AI 工具與利害關係人圖譜：一旦定義原則與清冊，就有必要了解那些流程可能因使用 AI 而受益，以及相關的利害關係人(目前使用者、潛在使用者、數據提供者及第三方等)。此圖譜將用於構成央行相關的具體應用專案。納入內部與外部利害關係人，對於確保 AI 系統符合監管與倫理要求至關重要。
6. 進行詳細的風險與控制評估(圖 2)：風險與控制的評估應包括 ICT 技術層面，及其他相關職能，如網路安全、法律、資訊安全、營運風險及外部第三方。AI 工具開始測試前或上線前，應確認控制措施已就位。

圖 2 AI 專案評估流程圖



資料來源：BIS (2025)

7. 定期監控 AI 法遵情形：第二道防線應開發法遵監控，並向 AI 委員會報告，以確保系統、營運、資訊處理及 AI 倫理準則符

合指引規範。有效的 AI 法遵監控包括：

- (1) 能持續監控 AI 系統數據漂移(data drift)⁶、效能、公平性、偏差及遵循預設參數的工具，並應確保用於 AI 訓練與營運的數據處理符合隱私規範與安全標準。
 - (2) 實施相關流程以確保 AI 系統所用數據品質與完整性。
 - (3) 制定政策以應對涉及 AI 系統的事件，包括事件調查、風險減緩與報告機制，並對任何違規或效能問題進行根本原因分析，以防止未來再次發生。
 - (4) 將人為監督納入 AI 法遵監控的一環，以驗證 AI 系統是否符合既定指引。
8. 報告異常與事件：AI 工具實施後，應建立向 AI 監督委員會報告異常狀況與事件之流程，以確認需減緩之剩餘風險或應強化之控制措施。
9. 開發並提升勞動技能：央行的正常運作與員工的專業知識與技能高度相關。AI 應用減少了密集的人工處理，使員工投入更多時間於更具生產力與創新性的活動。為了建立 AI 知識，央行應針對 AI 使用、治理及法遵，制定基礎/專業培訓計畫。
10. 持續審查與調整治理框架：鑑於 AI 領域變化快速，因而需要定期審查治理框架，以確保其有效性與適當性。央行需要定期

⁶ 指 AI 模型運作時，其底層基本資料的「統計分布」隨時間發生變化的現象。由於模型是基於特定時間點的歷史資料進行訓練，當輸入的新資料特徵與訓練時的資料分布產生差異時，就會導致 AI 模型的預測效能、準確度與可靠性顯著下降。

審查 AI 系統的執行結果，以確保其符合目標與風險管理措施；
央行也應審核 AI 的風險與控制評估結果，及涉及 AI 工具的事件報告。

伍、 結論

為確保 AI 技術能安全、負責任且合乎倫理整合至核心營運中，央行的施政與治理應持續發展下列事項：

- 1.強化資料基礎與治理：AI 的效能取決於輸入資料的品質。央行應優化資料生命週期管理，確保數據的品質、透明度與可追溯性，並推動標準化與可互通的資料系統，作為 AI 發展的穩固基石。
- 2.厚植 AI 人才：持續投資並提升員工的 AI 素養與風險意識，為落實 AI 治理的關鍵。面對全球資料科學專業人才短缺，以及來自民間企業高薪競爭的壓力，央行招募與留任 AI 人才面臨重大挑戰，因此，央行應對內建立專屬的 AI 培訓計畫，持續提升現有員工的 AI 素養，以應付日新月異的 AI 發展，同時須使員工深刻了解 AI 的風險與局限性(如模型幻覺、偏誤與缺乏可解釋性)，此外仍應以經濟與金融專業知識進行政策評估、判斷與監督，AI 應被定位為輔助工具。
- 3.強化跨央行合作並導入分享案例。央行有其特殊地位及職能，面對 AI 之發展往往無法參考一般組織之治理方式，因此應參與各個國際組織與央行之跨國會議，了解其他央行導入 AI 技術之實際治理經驗，並定期蒐集、整理國際案例與研究成果，作為規劃 AI 應用治理的參考。

總而言之，央行 AI 治理並非單純的資訊技術升級，而是一場與

既有政策、風險管理工具深度結合的全面性轉型。唯有建構穩健且可隨技術發展動態調整的治理體系，央行方能在瞬息萬變的數位時代中善用 AI 紅利，同時堅守維持物價與金融體系穩定的核心職責。

參考文獻

Accornero, Matteo and Gianluca Boscariol (2021), “Machine Learning for Anomaly Detection in Datasets with Categorical Variables and Skewed Distributions,” October.

Araujo, Douglas, Nikola Bokan, Fabio Alberto Comazzi, and Michele Lenza (2025), “Word2Prices: Embedding Central Bank Communications for Inflation Prediction,” BIS Working Papers No. 1253, March.

Ashwin, Julian, Eleni Kalamara, and Lorena Saiz (2021), “Nowcasting Euro Area GDP with News Sentiment: A Tale of Two Crises,” ECB Working Paper Series No. 2616, November.

BIS(2024) ,“Artificial intelligence and the economy: implications for central banks,”, BIS Annual Economic Report, Jun..

BIS(2025) ,“Governance of AI adoption in central banks,”, BIS Representative Office for the Americas, Jan..

BIS(2025) ,“Governance and implementation of artificial intelligence in central banks,”, Irving Fisher Committee on Central Bank Statistics, April..

Bräuning, Michael, Despo Malikkidou, Stefano Scalone, and Giorgio Scricco (2019), “A New Approach to Early Warning Systems for Small European Banks,” ECB Working Paper Series No. 2348, December.

Dauphin, Jean-François, Kamil Dybczak, Morgan Maneely, Marzie Taheri Sanjani, Nujin Suphaphiphat, Yifei Wang, and Hanqi Zhang (2022), “Nowcasting GDP - A Scalable Approach Using DFM, Machine Learning and Novel Data, Applied to European Economies,” IMF Working Paper 22/52, March.

Dunn, Wendy, Ellen E. Meade, Nitish Ranjan Sinha, and Raakin Kabir (2024), “Using Generative AI Models to Understand FOMC Monetary Policy Discussions,” FEDS Notes, December.

Kalamara, Eleni, Arthur Turrell, Chris Redl, George Kapetanios, and Sujit Kapadia (2020), “Making Text Count: Economic Forecasting Using Newspaper Text,” BoE Staff Working Paper No. 865, August.