

經濟合作與發展組織（OECD）及 金融穩定委員會（FSB）「人工智慧 （AI）在金融領域之應用評論」 摘要報告與淺析

本公司資訊處整理

壹、摘要

貳、金融領域之應用範疇

參、AI 在金融領域的日益普及對金融穩定的影響

肆、金融消費者保護、市場行為考量與 AI 在金融領域的政策應對

伍、對我國金融業未來 AI 發展之建議

壹、摘要

經濟合作與發展組織（Organization for Economic Co-operation and Development, OECD）與金融穩定委員會（Financial Stability Board, FSB）於 2024 年 5 月 22 日在法國巴黎召開金融業人工智慧圓桌會議，聚焦 AI 技術在金融業的應用現況與趨勢。會議深入探討金融機構與監理機關的實際案例，剖析其中機會與風險，並分享政策框架中的創新實踐與良好做法。

會議指出，預測性 AI 系統（如機器學習與生成式 AI）正以驚人速度滲透銀行與保險領域，顯著提升營運效率、風險建模、詐欺偵測及金融犯罪防範能力。生成式 AI 的內部應用潛力亦被重點提及，包括摘要、翻譯與資料檢索等功能。

在銀行業，AI 技術在防制洗錢與詐欺偵測方面已改變遊戲規則，使金

融機構能更精準地識別犯罪並降低誤報率。同時，生成式 AI 正逐步應用於內部流程，如摘要撰寫、程式碼生成與資料檢索，未來也將拓展至客戶服務領域，帶來全新應用契機。

在保險業，AI 技術已廣泛運用於核保、風險評估及理賠管理。生成式 AI 則憑藉語言驅動模型，進一步優化保單與理賠處理流程，提升效率並降低操作風險。

在資產管理和證券市場中，AI 在投資組合管理、交易及風險管理中的應用日益成熟。透過量化與文字數據分析，這些技術正用於優化資產配置、生成交易信號、實現自動化交易，並改進風險模型的驗證與回測。

儘管 AI 技術帶來諸多優勢，與會者也對模型風險、資料保護及治理問題提出警告。會議強調，資料品質、隱私與倫理議題攸關 AI 在金融領域的長期穩定性。生成式 AI 應採分階段部署，並確保各層級對 AI 工具的風險與應用有充分認識。

在金融消費者保護方面，AI 帶來了便捷、可及性、即時資訊等具成本效益之服務以及提升了用戶體驗。然而，深偽技術和具誤導性的 AI 生成結果、資料保護、隱私及機密性問題，以及偏見和歧視等挑戰也日益嚴峻。生成式 AI 系統的部署引發了有關模型輸出品質和可靠性之附加風險。相關的使用者可能未能充分意識到這些模型的侷限性，進一步加劇了其有限的可信度。

會議最後呼籲政策制定者推動 AI 在金融服務中的安全應用。強調需要針對模型風險管理採用基於風險之策略，並指出國際合作對於標準制定及良好實踐分享的重要性。此外，國家金融監理機關應持續評估其監理能力，以應對技術快速發展帶來的挑戰。

經濟合作暨發展組織（OECD）和金融穩定委員會（FSB）將持續支持國際間對 AI 發展、風險趨勢及潛在監理影響的監測工作，並為政策制定者提供必要的工具和技能，以有效在金融領域運用並監督 AI 技術。

貳、金融領域之應用範疇

一、AI 在銀行和保險領域的應用

銀行業對預測性 AI 系統的採用顯著增加，特別是在 2022 至 2023 年間，主要應用於營運、風險建模、詐欺及金融犯罪預防的模式識別。特別是過去 16 至 18 個月期間生成式 AI 技術的快速發展，引起業界對大語言模型（Large Language Models, LLMs）^{（註 1）}和其他生成式 AI 模型應用的關注。

銀行業中基於機器學習的應用案例涵蓋了防制洗錢、詐欺偵測、身份驗證，並受現行的模型風險管理、治理及資料保護框架所規範。尤其在金融犯罪相關應用中，動態風險分析工具已改變防制洗錢及金融犯罪檢查之方式，這些工具運用數據分析的力量，更精確地識別金融犯罪事件，並同時減少誤報數量。

銀行業最初部署的生成式 AI 應用主要是內部使用，包括摘要、翻譯、資料搜尋（特別是在需要上下文的場合）。考量到現有眾多客戶服務應用，程式生成被認為是目前銀行業在生成式 AI 工具使用中的一個重要實驗領域。未來，客戶服務也被認為是生成式 AI 工具應用的重要方向。銀行表示其生成式 AI 模型的部署採取了風險調整策略，強化現有的 AI 治理框架，並使用與這些框架相符的模型（例如小語言模型（Small Language Models, SLMs）^{（註 2）}，並用特定資料進行訓練）。儘管目前這些模型直接與客戶互動的情況仍然有限，但預期未來生成式 AI 模型的使用將隨著客戶需求的增加而擴大。

在保險業，預測性 AI 模型被廣泛應用於核保、風險評估、風險建模以及理賠管理與處理等多個保險產品線。生成式 AI 的引入使保險公司在處理保單和理賠時，能更有效透過增強式語言驅動模型來處理文字資料。AI 模型的翻譯能力促進了跨國理賠與保單處理更有效率；大語言模型（LLMs）也使專業經理人能夠從更完善的系統中快速搜尋資訊，以便提供建議，並簡化對複雜產品（例如人壽保險、退休金）的溝通流程。然而，在與客戶互動

時，人工參與在這個過程中仍然是必要的。

AI 技術正在為保險業帶來顯著效益，尤其在營運效率與客戶體驗方面表現卓越。透過 AI 工具，保險公司能加速理賠處理流程，提升服務效率，進而強化客戶滿意度。AI 技術還能協助保險業者對損失進行更深入分析，為客戶提供量身訂製的解決方案，包括更具競爭力的保費定價，精準滿足不同需求。此外，生成式 AI 的跨語言處理能力成為企業應對全球化挑戰的利器，支持跨國資訊的深入分析與決策，為保險業的國際化布局提供堅實助力。

在銀行和保險領域，AI 治理框架（包括生成式 AI）應聚焦於文化、教育與素養的需求。隨著 AI 工具的普及性超越傳統的專家範疇，理解與管理 AI 工具已成為組織各層級的共同責任。用戶需具備包括提出正確問題、評估模型輸出的可靠性及考量相關倫理等關鍵能力。同時，業界呼籲積極引進多元化外部人才，並提升現有員工的技能，以應對 AI 應用的挑戰。

數據在資料可及性、訓練數據集的質量，以及將金融與非金融數據整合至新興 AI 工具中，皆扮演關鍵角色，不僅影響模型的準確性與適用性，亦關乎金融機構在風險管理與創新應用上的競爭力。監理框架對於數據處理及結果產生的影響至關重要，將直接影響 AI 工具為組織創造的價值。

隨著 AI 的廣泛使用，金融機構對第三方服務供應商的依賴逐漸成為一大挑戰。雖然金融機構已建立針對雲服務等外包流程的框架，但仍需擴展至 AI 模型及數據供應商。建議在與供應商簽署合約時，應強調透明度，確保 AI 整合過程及影響的透明度，因為最終責任仍由金融機構承擔。此外，若 AI 模型失誤或信任流失的風險，也有可能對企業聲譽造成影響。

目前的監理框架已能涵蓋銀行與保險業中的 AI 應用，並未成為 AI 價值實現的障礙。然而，對於某些領域較寬鬆的監理，很可能引發監理套利風險。專家強調應採取以風險為基礎的模型管理方法，並對部分政策制定者對 AI 的廣義定義提出關切，例如有些法規將廣義線性模型（Generalised Linear

Models, GLMs) ^(註3) 歸類為 AI，無形中增加了合規負擔。為提升監理效率，建議導入指標相關的規範及機器可讀格式的要求，以降低合規難度。

生成式 AI 的學術研究逐步增長，跨領域合作是進一步推動其發展的必要條件。學術研究顯示，生成式 AI 在信用評分、金融預測、投資組合配置及保險定價等特定領域表現出色。然而，這些技術也面臨解釋性不足、輸出穩定性欠佳及高運算成本等挑戰。更由於缺乏統一的衡量指標，使得在解釋性與準確性之間的平衡變得困難。最終，生成式 AI 的大規模部署將取決於金融機構的風險承受能力。

二、AI 在資產管理和證券市場中的應用

AI 能有效應用於降低資產配置錯誤、優化模型參數以及提升流程效率等方面。AI 透過基於用戶意圖的技術，增強了資訊挖掘能力，改善交易分析與執行。同時，透過大語言模型的多語言支持，資料獲取變得更加普及。AI 提升了生產力，例如在協同操作和摘要功能上的應用，並展現了處理非結構化資訊的顯著優勢。

在賣方市場，AI 工具已廣泛應用於加強風險與流程控制、制定產品組合及輔助決策，並累積了逾十年的實踐經驗。此外，AI 還被用於詐欺偵測和制裁篩檢查核等具體場域，展現出實質價值。生成式 AI 的應用目前主要集中於內部大語言模型的試驗階段，整體發展仍處於早期探索。受監理較少的市場參與者可能面臨監理套利風險，因此需針對小型參與者制定統一且明確的規範，以確保市場公平與穩定。

生成式 AI 模型的實施應以可信任的方式逐步推進，採取循序漸進的策略。金融業在某些關鍵流程中尚未能完全發揮 AI 的潛力，這主要是由於相關應用仍伴隨一定風險。

生成式 AI 模型的應用應採用以人為中心的方式推動，即便在高度自動化的情境下，仍需設置獨立的風險控制與防護措施，例如審查機制及即時停

止功能，以有效管理風險。「人機協同」的方法被認為有助於確保責任制、稽核性及記錄共享。同時，專業領域知識被視為保障模型輸出準確性的重要因素，這與早期單純依賴機器學習工程師的 AI 應用方式已有所不同。

金融服務和流程的高度複雜性導致 AI 的應用類型常被過度期望，結果使金融領域中的許多 AI 應用案例因無法滿足實際需求而最終被放棄。

資料有效性、治理、隱私及倫理是非常重要的，隨著大數據的普及，資料在 AI 中的核心作用日益突出。瞭解使用者權利（如智慧財產權）及資料來源的可追溯性被視為保障模型輸出可信度的關鍵。金融機構不僅面臨技術挑戰，還需解決治理與文化層面的問題。發展負責任的 AI 框架被認為有助於增進用戶信任。考量到大語言模型及生成式 AI 工具的日益普及，加強各級管理層的技能培訓也被視為不可或缺的措施。

許多 AI 相關風險源自早期 AI 工具（如機器學習）的侷限性，結合經濟與 AI 研究的綜合視角，對於識別未來可能影響金融穩定的潛在風險至關重要。

AI 的應用需在短期效率、穩定性及未知壓力風險之間達成平衡。特別是，AI 模型在追求利潤最大化目標的過程中，可能無意間涉及內線交易，進而導致市場操縱的風險，儘管已明確要求不得使用敏感資訊。此情境突顯了金融系統的高度複雜性，以及難以事先為每個目標制定全面規範的挑戰。

另一個風險來自於第三方 AI 服務（如模型、雲服務或資料分析）因規模回報效應日益擴大而導致的市場壟斷現象。這種寡頭壟斷結構對監理機關形成了額外挑戰，隨著對民營部門分析的依賴加深，可能逐漸形成單一觀點，進而降低對市場不穩定性風險的有效監控能力。

監理機關在 AI 驅動的金融新時代需要更快速的應對策略。建議利用 AI 驅動的壓力情境來探索應對方案，讓 AI 在未來成為有效的解決工具。然而，對於依賴 AI 系統執行關鍵程序的風險，尤其在決策過程尚未被完全理解的情況下，被視為未來愈加重要的監理關注點。

參、AI 在金融領域的日益普及對金融穩定的影響

AI 在金融業中的部署可能帶來金融穩定風險，特別是可能加劇金融機構之間的相互連結，以及對模型和資料的複雜性與不透明性的隱憂。

隨著 AI 技術的快速發展，市場集中度的提升以及對第三方的依賴，模型風險正面臨新的挑戰。隨著資料、雲服務和第三方依賴程度的增加，這些因素可能進一步加劇金融業現有的脆弱性，並因縱向整合而放大風險。然而，不同的 AI 實施策略可能在一定程度上緩解這些影響，因此監測市場集中度與應用多樣性變得更加必要。

生成式 AI 可能加速金融系統的變革。新進入者，特別是金融科技公司和其他非銀行金融機構，能夠通過利用資料處理的規模經濟來重新調整市場競爭格局。在流動性較低的市場中，早期採用生成式 AI 可能帶來顯著的競爭優勢。

儘管金融行為的相關發展對金融穩定的影響並非新議題，生成式 AI 未來的廣泛應用引發了對其潛在操縱風險的關注，尤其是對工具自主性錯誤認知的可能性。這些問題可能進一步引發大規模的連鎖效應，成為需要重點關注的潛在風險領域。

生成式 AI 對宏觀經濟的影響，例如勞動市場的變化，可能間接對金融穩定構成威脅，這一議題需進一步深入研究，以全面評估潛在風險。

目前 AI 在金融領域的應用尚未達到變革性的影響，主要用於提升特定任務的效率，並以探索性應用為主。AI 技術尚未被全面採用，許多自動化流程仍需要人工監督。由於這一轉型仍處於早期階段，充滿不確定性，監理機關在收集數據以有效監控發展及評估潛在脆弱性方面面臨挑戰。

生成式 AI 為金融監理機關提供了效率提升的機會，包括更高效地打擊詐欺、應對網路風險，以及通過將 AI 嵌入業務流程來強化對金融部門的監理。新工具的應用使監理機關能夠更高效地處理大量數據，進一步提升其監

控效能。

肆、金融消費者保護、市場行為考量與 AI 在金融領域的政策應對

AI 在面向消費者的金融服務中帶來了顯著益處，應用場景涵蓋機器人顧問、聊天機器人、工作面試、身份驗證、信用評估及退休規劃等。這些技術提供了便利性、可及性和即時資訊，並且以具成本效益的方式改善服務，例如更快速地解決客訴、提升用戶體驗及提供客製化產品。同時，AI 在促進普惠金融方面的作用也有所顯現，例如基於 AI 情感分析和替代數據的數字貸款，這對中小型企業特別有助益。此外，AI 還強化了安全性，減少身份盜用風險，提升客戶價值並促進更明智的消費選擇。然而，目前生成式 AI 的部署仍以內部應用為主。

與此同時，金融消費者面臨的主要風險包括深偽技術和具誤導性的 AI 生成結果，這些問題引發了對生成式 AI 在消費者保護風險上的更大關注。核心議題涉及數據保護、隱私與機密性、資料品質、智慧財產權、安全性以及偏見與歧視。此外，模型更容易受到如對資料之投毒攻擊等輸入攻擊威脅，而生成式 AI 的自學習和動態調整能力則進一步帶來新的風險，使這些工具在某些敏感的金融領域應用時需要格外謹慎。

在適當使用 AI 技術以實現其益處與限制其潛在風險之間，需達成平衡，尤其對於不具備數位素養的用戶而言更為重要。加強金融教育以及提升消費者對 AI 在金融領域中風險的理解，對於 AI 技術的負責任應用是不可或缺的。

此外，政策和監理措施在應對 AI 於金融服務中的潛在風險方面至關重要，必須在促進創新與保障消費者保護之間找到平衡，以深偽技術和虛構錯覺為例，這些新興風險需得到監理機關與業界高度關注並採取相應對策。建議應向消費者明確披露建議或輸出結果是由 AI 模型生成的，以降低潛在

風險。同時，妥善的資料管理及明確的責任歸屬亦被視為風險控制的關鍵措施。

在政策和監理應對方面，應遵循現有的原則來降低風險。全球標準（如 G20/OECD 的《金融消費者保護原則》），被認為是應對新興金融消費者風險的重要指導框架。同時，AI 在打擊不當行為方面的潛力也得到了肯定。由於 AI 的全球性特質，國際合作在標準制定與實踐分享中顯得尤為關鍵。

公共機構與監理機關推動 AI 在金融領域的安全使用，旨在實現 AI 技術的益處，滿足消費者期望，保護市場競爭，維持市場誠信，並促進普惠金融。採用以風險為基礎之方法，建議通過创新中心及沙盒環境推進 AI 技術在金融領域的安全採用，並對因應潛在風險和促進創新提供支持。

金融政策制定者應持續監控 AI 在金融業務中的應用案例及相關風險，尤其是生成式 AI 在未來面對客戶應用中的發展趨勢。建議持續採用現有的監理框架與工具，包括治理、資料管理、風險控制和營運管理。目前新的立法框架已在推行，旨在應對高風險 AI 應用，平衡跨行業的民主價值與人權保護，同時促進技術創新。

再者，應為政策制定者（特別是金融監理人員），提供適當的工具和技能，以有效應對 AI 技術在金融領域的監理挑戰。由於 AI 技術的全球特性，國際協調對於制定標準與應對風險至關重要。同時，國家層面的金融監理機關需持續評估其國內的監理能力。此外，許多議題已超出金融範疇，其他主管機關（如資料隱私監理機關）也需要參與協作，以確保全面治理。

伍、對我國金融業未來 AI 發展之建議

AI 技術已成為提升金融機構營運效率、強化風險管理以及推動客戶服務創新的關鍵工具。然而，AI 快速發展也帶來了包括資料保護、模型風險、隱私、機密性、監理等議題。

目前生成式 AI 仍以內部應用為主（如資料搜尋、翻譯與風險建模），而客戶導向之應用（如智能客服、個人化服務）仍處於初期探索階段，需要持續關注其發展趨勢及相關風險。

AI 的應用應以「人機協作」為核心，避免完全自動化，保留人工監督、加強對模型輸出的解釋性、可靠性與審查機制。

AI 技術的應用需要謹慎規劃與持續改進，未來應秉持「創新與風險管理並重」的原則，在推動技術創新的同時，確保風險管理機制的健全性，實現兩者之間的平衡。

隨著 AI 在金融領域的快速發展，金融機構應加強 AI 治理框架，並在風險管理、透明性、可解釋性、隱私保護和系統穩定性等方面持續優化，茲此，為進一步提升金融機構的發展，提供以下建議方向：

一、強化 AI 治理框架

- (一) 建立生成式 AI 應用的資料與模型風險管理規範，確保資料可追溯性與透明性，防止技術濫用、使用不合法或偏差的資料，而影響金融公平性，同時確保模型輸出準確且具解釋性，並能回溯追蹤 AI 模型的資料使用歷程。
- (二) 建立 AI 風險管理機制，將 AI 風險管理整合至現行的風險管理和內部控制程序中，並進行定期的風險評估、壓力測試與風險緩解措施的改進，確保相關措施能與時俱進，維持有效控管。
- (三) 以「人機協同」的風險控管機制，以符合以人為本原則，AI 生成的建議和資訊由人員進行專業審查，尤其是在面對民眾的自動化服務，應明確標示 AI 參與的環節，讓民眾能清楚了解哪些內容是由 AI 模型生成之建議與輸出結果，確保民眾知悉資訊來源。

二、平衡創新與風險

- (一)採取分階段部署策略，逐步應用生成式 AI 於業務中，並於各階段進行風險評估與測試，確保技術應用的穩健性以降低相關風險。
- (二)建立 AI 試驗沙盒環境，在受控情境下測試創新技術，使 AI 技術在金融領域中能安全應用。

三、推動人才培育與提升員工 AI 素養

- (一)持續關注生成式 AI 相關應用與最新趨勢，邀請 AI 專家開設工作坊或講座，提升員工 AI 專業知識，並加強員工對 AI 倫理、隱私和資料保護的風險意識。
- (二)建立跨部門協作機制，結合技術與業務需求，讓員工熟悉 AI 在業務中的實際應用情境，以提升員工 AI 應用能力，協助員工適應 AI 時代的變化。
- (三)與學術機構進行合作，舉辦「AI 人才培育計劃」或「AI 實習計劃」，讓學術界的 AI 人才有機會參與實際的金融 AI 項目，提升金融機構的技術實力，並吸引頂尖的 AI 人才加入。

四、加強對第三方的監理

- (一)與第三方供應商合作時，加強對供應商的監督與評估，確保 AI 系統的設計、開發與部署符合機構的風險標準與法規要求。
- (二)與第三方供應商的合約，應明確定義責任範圍、資料存取權限、資料的保護機制及供應商的監理責任。

五、強化數位普惠金融

- (一)利用生成式 AI 進行語音互動或翻譯，確保不同文化和語言背景的民眾能無障礙獲取金融資訊與服務。
- (二)透過 AI 互動方式提供民眾關於投資理財、銀行業務、金融商品、防詐等金融知識與 AI 應用於數位金融服務之說明，進一步提升並普及民眾對金融商品與 AI 應用的風險認知與素養。

六、推動綠色金融創新

- (一)支持發展以環境永續為目標的 AI 創新應用，如碳足跡分析、綠色投資建議等，融入環境、社會與治理（ESG）原則，確保 AI 的設計和運作不會加劇社會不平等，促進包容性成長與社會福祉。
- (二)在 ESG 報告中揭露 AI 的永續性影響，為減少碳排放和實現綠色金融提供具體的量化目標和措施，實現經濟與環境雙贏。

參考文獻

1. 金融監督管理委員會（2023），「金融業運用人工智慧（AI）之核心原則與相關推動政策」。
2. 張維君（2023），「永續經營前進綠色金融」，工研院產業學習網。
3. 黃譯漫（2024），「淺論 AI 技術對減碳管理之助益」，期貨人季刊，第 89 期。
4. Global Partnership for Financial Inclusion（2023），“2023 Financial Inclusion Action Plan”。

註釋

註 1：大語言模型（Large Language Models, LLMs），是一種基於深度學習的人工智慧模型，能夠辨識和解釋人類語言或其他類型的複雜資料，常用於自然語言的處理與文字生成，如：聊天機器人、文章摘要、翻譯、程式碼生成等範疇。

註 2：小語言模型（Small Language Models, SLMs），概念與大語言模型（LLMs）相同，主要是使用的運算資源與訓練資料比大語言模型（LLMs）來的少，是規模較小的語言模型，雖然性能不及於大語言模型（LLMs），但由於其低成本、高效率的特性，小語言模型能滿足許多輕量化的需求，同時更易於部署與管理。

註 3：廣義線性模型（Generalized Linear Models, GLMs），是一種統計模型框架，為傳統線性迴歸模型的延伸，能夠處理不同類型的資料，包括連續型、分類型和計數型資料，範圍比一般線性模型更大，可處理更多元的機率分配等問題。