

論著與分析

AI 應用對金融監理之挑戰及法制建構探討

吳佳琳

壹、前言

- 一、人工智慧對於金融服務之主要影響—競爭優勢
- 二、人工智慧應用於金融可能產生之風險、挑戰與道德問題

貳、人工智慧法制趨勢與金融監管

- 一、歐盟人工智慧法制趨勢
- 二、英國針對人工智慧應用於金融服務之監管方向

參、結論與建議—合法、安全應用人工智慧於金融服務

壹、前言

人工智慧（Artificial intelligence, AI）技術自發展以來，已廣泛應用在各個領域，包括交通、運輸、流通、零售、醫療、金融、教育、藝術及軍事各方面。而金融領域具有全球化及高度監管之特性，因此討論人工智慧應用對於金融監管之影響有其必要性及重要性。本文將首先提出人工智慧於金融領域可能產生之影響與風險，進一步探究目前歐盟對人工智慧法制之發展趨勢及個別國家針對人工智慧在金融領域發展之監管態度，並提出個別須關注之法制上議題，最後再就如何合法、安全應用人工智慧於金融服務進行整體應用之監管建議。

本文作者為財團法人資訊工程策進會科技法律研究所高級法律研究員。本文為作者個人意見，不代表本公司立場。

一、人工智慧對於金融服務之主要影響—競爭優勢

人工智慧能夠處理分析大量數據，並且細緻地判斷客戶之需求，因此在金融業中應用人工智慧技術，對外可協助業者更了解其客戶之需求，並提供更多資訊及量身訂製之產品或服務；對內則可強化內部流程及效率，並降低相關風險（例如：授信、法遵風險）。依據世界經濟論壇（World Economic Forum）2020 年針對全球 33 國 151 位金融業受訪者之調查，有 85% 受訪者已實施人工智慧技術並已開始以些許優勢領先現有其他企業。目前人工智慧應用於金融領域多用於五個面向，包括：創造新的潛在收入、風險管理、流程改善、客戶服務及取得客戶，目前則較多用於風險管理面向^{（註 1）}。在未來，預計人工智慧將更進一步強化金融相關服務，甚至改變業者提供服務之方式。

二、人工智慧應用於金融可能產生之風險、挑戰與道德問題

隨著人工智慧之應用趨於廣泛，其可能產生之風險與倫理問題開始被討論。主要討論之問題有：因人工智慧系統通常仰賴大量數據及資料驅動，故數據與個人資料之選擇及資安保護成為重要關鍵之一。同時，由於原始資料係由人類選擇下進行分析，演算法也由人類設計，亦可能產生偏見與歧視之決策結果；再者，由於人工智慧系統係基於機器學習（Machine Learning）與深度學習（Deep Learning）技術，而相關決策過程都在黑盒子（Black Box）的狀況下進行，人類可能無法了解決策結果產生之邏輯與原因，進而產生系統穩固（Robustness）、安全性（Safety）及透明度（Transparency）等議題；最後，由於人工智慧具有自主性，可能無法以傳統因果關係理論判斷責任，若人工智慧系統造成損害其責任歸屬及損害填補等問責制（Accountability）之設計也是重要法制議題。以上問題對應至金融領域又可區分為數據風險、資安風險、金融穩定性風險、法律風險及倫理問題等，其具體問題也隨著金融領域逐漸廣泛使用人工風險系統而漸浮上檯面。

（一）數據風險

在數據風險方面，除數據蒐集成本越來越高昂外，數據質量也成為重

要的問題，過去金融機構徵信常利用個人繳款行為類或負債類信用資料進行信用評分，而在人工智慧技術發展後，金融機構可透過大數據分析與雲端計算，提供模組式差異化產品與服務，大量、多元且品質佳的數據成為重點，甚至開始利用非傳統之徵信資料進行評估。例如：中華電信與凱基銀行攜手合作，透過電信資料勾稽並以創新金融應用技術驗證申貸用戶身分^(註2)，如此一來，供驗證資料之品質與邏輯性相對重要，因為使用不同於傳統之徵信資料也可能會導致不同演算結果，進而影響個人信用及貸款成效而有歧視或偏見產生之可能性。

(二) 資安風險

在資安風險方面，人工智慧系統可能用於攻擊、操縱演算法或以其他方式損害經濟並威脅金融系統，進而影響整體國家經濟。再加上目前金融機構多數透過技術供應商分包之方式應用人工智慧技術，而產生了一種新的風險形式稱作「技術風險」(Techrisk)而人工智慧系統若未先行於金融領域測試則可能放大此種風險^(註3)。

(三) 金融穩定性風險

而人工智慧在破壞金融穩定性風險之問題在於可能因為演算法之不透明性(Opacity)及缺乏可審計性(Auditability)而導致金融市場產生非預期性之變化^(註4)。主因在於過去金融機構通常以風險管理模型去分析銀行所面臨的風險，稱作模型驗證(Model Validation)。其中，相關風險模型皆有嚴格之治理規則，而在金融機構應用人工智慧系統之狀況下，可能無法確保可審計性，甚或有難以實現之狀況，而產生不可預見且影響金融穩定性之風險。

(四) 法律風險及倫理問題

在法律風險方面，關於人工智慧之問責制已多有討論，學者認為現有產品責任與侵權制度可能不足以應對未來人工智慧所引發之風險或損害^(註5)。對應至金融領域為監管科技(Regtech)逐漸被用於處理金融領域之合規與法遵之相關問題，包括反洗錢、金融詐欺等領域，但相關風險近年

來也隨之浮現。舉例而言，澳洲政府金融情報機構澳洲交易報告與分析中心（AUSTRAC）於 2020 年指控澳洲西太平洋銀行（Westpac）因大規模違反澳洲《反洗錢和反資恐法》（The Anti-Money Laundering and Counter-Terrorism Financing Act 2006, AML/CTF Act）共計 2300 萬次，違法行為包括未正確報告超過 1950 萬份之國際資金移轉指示（International Funds Transfer Instructions, IFTIs）、未保留國際資金移轉來源相關紀錄、對部分被用在剝削兒童之國際資金轉移未能盡職調查等系統性失誤，隨後西太平洋銀行同意支付約 13 億澳元罰款^{（註 6）}。而通常銀行會透過內置演算法針對「態樣」（Typologies）^{（註 7）}自動對於危險或可疑之資金移轉進行危險提示，但根據對於數據或演算法之設定方式不同，可能會讓危險信號過少或過於寬泛執行，而使金融機構忽略重要之危險信號。而在本案中西太平洋銀行所使用之低成本國際交易服務 LitePay 確實針對態樣內建演算法，但卻未透過該演算法發現關於剝削兒童之國際資金轉移，也未能產生足夠之危險信號提醒銀行^{（註 8）}。由此案例可知，越來越多金融機構將人工智慧納入其系統，包括監管科技基礎設施，而此類設施錯誤而導致之監管或法律風險可能會增加。

最後，有關金融倫理與人工智慧方面則是需要更深層之考量。由於人工智慧系統缺乏對於人類心理之洞察力，部分人類行為可能無法確切說明行為之邏輯與原因，交由人工智慧系統執行後需要更細緻考量或進一步解決所產生之倫理或影響人類價值觀問題，例如：伊斯蘭金融不得牽涉利息（Riba）的索取與支付，在演算法之設計上可能需要再增加相關宗教因素之設計；又或者人工智慧系統是否會推薦讓客戶買下不適合之金融商品；以及人類可能過度依賴人工智慧而失去判斷能力，未來可能導致金融人才之流失^{（註 9）}等相關倫理層面議題。

貳、人工智慧法制趨勢與金融監管

目前已有許多國家及國際組織意識到人工智慧可能產生之風險，進而開始規劃相關規範及法制架構，期望透過嚴謹的法律規範與事前預防措施，保障人類在

人工智慧的平等、安全及隱私並降低人工智慧的風險。目前歐盟相關對於人工智慧之相關進展相對完整，因此本文將先就歐盟近年來對人工智慧之相關規範進行初步介紹，而後再就英國針對人工智慧應用於金融服務之監管方向進一步說明在金融領域可借鏡之處，最後則針對人工智慧可能影響之個別議題觀察世界主要之規範趨勢，俾利作為人工智慧在我國金融領域之監管參考。

一、 歐盟人工智慧法制趨勢

歐盟近年來為推動及發展人工智慧技術，積極推動整個歐盟層面的人工智慧合作計畫與規範。歐洲議會的法規委員會（The JURI Committee）於 2015 年成立了工作小組負責制定歐洲的機器人民事法律規則並進行相關研究，於 2017 年 2 月 16 日投票通過向歐盟執委會針對機器人發展的一般原則、責任、研究與創新、倫理、智慧財產與數據流及個別領域（自動駕駛工具、醫療照護）等多面向進行原則性、方向性之提出規範建議^{（註 10）}。其後，歐盟執委會於 2018 年透過公開遴選任命 52 位專家組成人工智慧高級專家組（The High-Level Expert Group on Artificial Intelligence, AI HLEG），已完成兩項主要工作，包括 2019 年 4 月提出「可信賴的人工智慧倫理準則」（Ethics Guidelines for Trustworthy AI）^{（註 11）}及 2019 年 6 月「可信賴的人工智慧政策及投資建議」（Policy and investment recommendations for trustworthy Artificial Intelligence）^{（註 12）}；2020 年 2 月 19 日歐盟執委會發表「人工智慧白皮書」（White Paper On Artificial Intelligence-A European approach to excellence and trust）指出未來將以「監管」與「投資」兩者並重，促進人工智慧之應用並同時解決該項技術帶來之風險^{（註 13）}；又於 2021 年 4 月提出《人工智慧法》草案（The proposal for a Regulation laying down harmonised rules on artificial intelligence）^{（註 14）}，期盼透過強制規範確保人民與企業運用人工智慧時之安全及基本權利，藉以強化歐盟對人工智慧之應用、投資與創新。由近年重要文件及相關法制進展可以發現，歐盟先透過相關研究、報告提出人工智慧之法制關注重點，再進一步提出非強制性、產業自律性之倫理準則，最後再提出具有強制性之規範草案，形成完整的人工智慧規範框架。以下將就上述個別文件進行摘要說明。

（一）「可信賴的人工智慧倫理準則」

本倫理準則目的是促進可信賴的人工智慧，並透過人工智慧系統使人類更繁榮，增強個人和社會福祉與共同利益，因此強調人工智慧系統必須以人為本（human-centric），並透過尊重人權、民主及法治之基本價值觀以獲得相關利益。本準則分為三個層次，透過三大基礎、四項原則及七項要求勾勒出整體人工智慧倫理架構，最後再透過提供具體且詳盡的可信賴人工智慧評估清單，並讓人工智慧研究者、佈署者和使用者進行勾選與自我檢視，以利上述原則及要求之落實。

該準則之三大基礎包括穩固（Robust）、倫理（Ethical）及合法（Lawful）。首先，在穩固方面，從技術和社會的角度來看，即使人工智慧有良好的意圖，也可能造成無意的傷害。因此，人工智慧系統必須以安全、可靠的方式運行，並有預先保護措施；其次，在倫理方面，則是確保遵守倫理原則和價值觀，由於法律可能無法與時俱進，或是與倫理道德之規範有所不一致，甚至法律有時候並不適宜用於解決某些問題，因此人工智慧之發展也應先符合倫理原則；最後，在合法方面，在開發、佈署和使用人工智慧系統之相關過程，皆須遵守所有適用的法律和法規。在此基礎下，方能實現可信賴的人工智慧，理想情況下，所有三個部分在其操作中應相互協調並有所重疊。

而在第二個層次上，倫理準則中提出四項倫理原則，包括：應尊重人類的自主性（Respect for Human Autonomy）、防止傷害（Prevention of Harm）、公平（Fairness）與可解釋性（Explicability）四項，同時承認這些原則之間可能會存在緊張關係，但未來需要盡力消弭。同時也要注意涉及弱勢族群的狀況：如兒童、殘疾人士和歷史上處於不利地位或有被排斥風險的群體，以及權力與資訊不對等之狀況，例如：雇主與勞工、企業與消費者之間的關係。在人工智慧系統為個人與社會帶來實質利益的同時，注意人工智慧系統所帶來之負面影響：包括可能難以預測、識別或衡量的影響（例如：民主、法律和分配正義或人類心靈。），進而採取適當的措施減輕這些風險，且手段應與風險的大小成比例。

第三層次則是進一步提出七項關鍵要求，針對可信賴之人工智慧如何實現提供指導。該七項要求分別如下：(1) 人類監理（Human agency and Oversight）；(2) 技術穩健和安全性；(3) 隱私和數據治理（Governance）；(4) 透明度；(5) 多樣性、不歧視和公平；(6) 環境和社會福祉；(7) 問責制（Accountability）。未來也將以清晰、主動的方式，向利益相關者傳達人工智慧系統能力和其侷限性（Limitations），並透過上述要求之落實以實現人工智慧系統應符合人類實際期望之設定。故本倫理準則最後提供了具體且詳盡的可信賴人工智慧評估清單，並讓人工智慧研究、佈署和使用進行勾選及自我檢視。該評估清單係依上述七項關鍵要求開展提出各項問題，並可根據人工智慧系統的具體使用情況進行調整，但文件中也強調，這樣的評估清單永遠不會是詳盡無遺漏的，開發、布署及使用者不僅是要關注評估表中所要求的內容，而是應持續識別、實施需求及評估解決方案，確保整個人工智慧系統之生命週期中能持續改善、進化，並讓利益相關者參與其中。就上述觀察而言，本倫理準則應該被視為一個動態文件，並隨著時間推移進行審查和更新，以確保它們隨著技術、社會環境和我們的知識的發展而持續相符。

（二）「人工智慧白皮書」

歐盟執委會於 2020 年 2 月 19 日發表「人工智慧白皮書」，指出未來將以「監管」與「投資」兩者並重，促進人工智慧之應用並同時解決該項技術帶來之風險。白皮書針對監管方面，內容中提到將以 2019 年 4 月發布之「可信賴之人工智慧倫理準則」所提出之七項關鍵要求為基礎，將以此制定明確之歐洲監管框架。因此，白皮書認為應該先行定義欲解決之問題與明確所欲降低之風險，同時先行調整目前歐盟法律規範框架適用於人工智慧時所可能產生之不足之處，而後再遵循下列風險基礎方法（Risk-Based Approach）進行規範設計，基於不同風險程度的人工智慧給予不同程度之規範^{（註 15）}。

依據上述步驟，白皮書中認為可以具體透過下列方向進行調整：(1) 確保既有歐盟及國家法規被有效執行：有必要時應針對相關法律進行調整

或解釋，例如現行法規有關責任歸屬之規範可能需要進一步研析；(2) 重新劃定現行關於歐盟產品安全之規範範圍，並釐清現行歐盟法規之限制：例如現行歐盟產品安全法規原則上不適用於「服務」，或是是否涵蓋獨立運作之軟體（Stand-alone Software）有待確認；(3) 確保機器學習系統之安全性：此類產品和系統會於其生命週期內修改自身功能，故人工智慧技術需要頻繁更新軟體，這樣的特性可能會使系統於運行時產生於設計製造階段未預想之新風險，針對此類風險，應制定可針對此類產品在生命週期內修改功能之規範；(4) 責任分配產生不確定性：人工智慧產品可能會使責任歸屬產生模糊，有時候非製造商也可能需要負擔部分責任，有效分配不同利害關係者間之責任成為重點，目前產品責任偏向生產者負責，而未來可能須由非生產者共同分配責任；(5) 安全概念產生變化：人工智慧產品或服務可能帶來風險之外，相關網路安全、資訊安全也可能帶來風險，掌握人工智慧帶來的新興風險，並因應風險所帶來之變化。

針對何種人工智慧產品或服務需要被監管，白皮書中也提出了明確之判斷標準，將先針對高風險人工智慧進行監管，而高風險人工智慧之判斷標準在白皮書中則提到以人工智慧所欲應用之行業，該行業之特點與所進行之活動，通常被認為會發生重大風險，例如：醫療業、運輸業；而在上述的行業中，人工智慧被應用的方式可能會帶來重大風險進行判斷。而被判斷為需要監管之人工智慧後，將會針對訓練數據、數據及記錄儲存、資訊提供、穩固性與準確性、人為監督及某些特定人工智慧應用程序進行監管。從白皮書之相關內容也可以發現歐盟人工智慧之監管框架越來越清晰，內容也越來越明確且具可操作性。

（三）《人工智慧法》草案

歐盟執委會於 2021 年 4 月針對人工智慧法律框架提出規範草案，該法通過後預計將統一適用於歐盟各成員國，對於人工智慧進行系統性之規範。本草案主要欲達到以下列目標：(1) 確保投放至歐盟市場之人工智慧系統是安全的，且尊重基本權與歐盟現行法律規範。(2) 確保法律上之確定性，以促進人工智慧之投資與創新。(3) 強化現行法律之有效執行與治理，

並適於人工智慧之權利與安全要求。(4) 促進單一市場之發展，建立合法、安全和可信賴的人工智慧系統，並防止市場分散。該草案依據「人工智慧白皮書」所提及之風險基礎方法進行規範設計，基於不同風險程度的人工智慧給予不同程度之規範。將人工智慧系統主要分為「不可接受之風險」(Unacceptable risk)、「高風險」(High-risk)、「有限風險」(Limited risk)及「最小風險」(Minimal risk)四個等級^(註 16)。

「不可接受之風險」因為對人類安全、生活及基本權利構成明顯威脅，故將被禁止使用，例如：透過該技術操控人類潛意識並達到扭曲行為且造成身心傷害、利用該技術利用弱勢群體（例如兒童或殘疾人士）之弱點，扭曲其行為造成身心傷害者、政府進行大規模的公民評分系統、有限條件下方能於公共場所使用遠距生物辨識系統等；「高風險」人工智慧系統則是主要透過正面列舉方式列出，而高風險之人工智慧在進入市場之前須要先行遵守嚴格之義務，並進行適當風險評估及緩解措施等。「有限風險」則是指部分人工智慧應有透明度之義務，例如當用戶在與該人工智慧系統交流時，需要告知並使用戶意識到其正與人工智慧系統交流。最後則是「最小風險」，大部分人工智慧應屬此類型，因對公民造成很小或零風險，該草案並未規範此類人工智慧。

就本草案所被視為之高風險人工智慧系統應為本法之重點之一，會被視為高風險主要有二種：其一，若該人工智慧系統為歐盟既有安全規範之核心技術或組成要件，列於草案附件二；其二，符合附件三所列之特定技術或應用方式之人工智慧系統。而有關附件三所列之八大類高風險技術，分別為：

- 1.用於自然人之遠端生物辨識系統。
- 2.關鍵基礎設施之管理與營運：可能使人民的生命和健康受到威脅之重要基礎設施，例如：供氣、供水、供電等安全組件。
- 3.教育或職業培訓：該人工智慧系統可能決定受教或職業培訓之機會或是評估受教者。
- 4.就業、員工管理和自雇管道：徵才使用的履歷分類軟體。

5. 用於獲得或享受基本的私人和公共服務或福利：該人工智慧系統主要用於評估自然人是否能獲得公共服務或福利之資格，或是用於自然人之信用評等，例如：公民取得貸款的機會。
6. 可能干涉人們基本權利的執法：執法機關利用人工智慧系統對於對自然人進行犯罪風險評估，如：再犯率或潛在受害者風險；又或是用於測謊、檢測深度偽造、偵查、犯罪分析等。
7. 移民、庇護和邊境控制管理：評估欲入境者之風險、核實旅行文件真實性及協助主管機關審查簽證、居留許可等身分文件之系統。
8. 利用人工智慧於司法和民主程序：主要是透過人工智慧系統協助司法當局研究和解釋法律。

符合前述界定的高風險人工智慧系統類型後，須符合草案規定之嚴格管理義務，包括：應建立、實施、記錄和維護與高風險人工智慧系統相關的風險管理制度；進行資料治理，以提供高品質的數據集儘量減少風險和歧視性的結果；製作並即時更新相關技術文件，並對人工智慧系統活動進行記錄，以確保可追溯結果；向使用者提供關於系統之所有必要資訊，以便監管單位評估其合規性；實施適當的人力監督措施，以最大限度地減少風險；確保系統具有高度的穩健性、安全性和準確性，同時需要採取安全措施與適當之緩解計畫。關於上述規範無論是技術提供者或使用者都須嚴格遵守，草案中也針對違反者制定了高額罰鍰。

除上述針對高風險人工智慧之規範架構外，草案中另針對未來歐盟在人工智慧之治理方面提出發展方向，歐盟執委會建議各國現有管理市場之主管機關督導新規範之執行，且將成立歐洲人工智慧委員會（European Artificial Intelligence Board），推動人工智慧相關規範、標準及準則之發展，也將提出法規沙盒以促進可信賴及負責任之人工智慧。此草案未來將適用歐盟所有會員國外，預計也將產生域外效力，故其立法進度也將直接影響全球人工智慧發展，值得持續追蹤與關注。

二、英國針對人工智慧應用於金融服務之監管方向

金融監管機構對於使用人工智慧之態度相對重要，以英國金融行為監理總

署 (Financial Conduct Authority, FCA) 為例, 近幾年也意識到有關人工智慧或是數據分析技術應用對於金融領域之影響, 展開相關之研究及論壇討論, 期望能夠協助消費者在技術創新中獲益, 同時也討論如何將這些創新技術應用於金融監管上。就觀察而言, 近年英國對於金融科技之興起主要採取「等待與觀察」(Wait and See) 之態度, 監管機構採取「等待」態度, 並透過「觀察」新興技術之應用找到最需要進行監管之方向與順序。

首先, 在監管創新上, FCA 積極鼓勵監管技術之發展, 期望透過監管科技克服在金融服務中所面臨之監管挑戰。因此, FCA 將推出「Showcase Days」計畫, 邀請監管科技、金融科技之公司向 FCA 展示有關監管、合規之技術解決方案, 透過相關演示協助 FCA 能夠更好地了解市場狀況並且高效率執行核心監管任務, 在該計畫中主要集中於下列五個領域之監管科技: 反詐欺科技、監管報告、市場監控科技、社交媒體分析及審慎分析 (Prudential Analytics) ^(註 17)。

而在關於人工智慧監管方向上, FCA 與英格蘭銀行自 2020 年 10 月 12 日起啟動了一個人工智慧公私論壇 (Artificial Intelligence Public-Private Forum), 該論壇之目的主要是促進公私部門之對話, 並更進一步了解人工智慧在金融服務中的使用與影響, 包括進行一系列之會議及研討會, 而其主題則圍繞在人工智慧之數據、模型風險與治理, 截至 2021 年 12 月已召開了四次會議, 預計未來將發布該論壇之最終報告 ^(註 18)。而觀察該論壇目前討論之方向, 大致可知 FCA 針對人工智慧於金融領域之應用, 未來可能朝向下列議題進行監管:

- (一) 治理結構: 金融機構要如何創建或調整新的治理結構, 同時提供合乎倫理、安全、穩健及具金融風險抵抗力 (Financial Resilience) 之人工智慧應用。可能需要解決在金融服務中最相關的是資料治理、模型風險管理框架以及操作風險管理問題, 高風險之人工智慧模型未來可能需要進行更多的實地查核 (Due Diligence), 且也可能需要加入人工智慧倫理相關議題, 部分倫理問題對於金融領域來說可能相當新奇, 目前尚未有專門之規範框架處理相關問題, 因此在人工智慧治理可能需要更廣泛之思考 ^(註 19)。
- (二) 角色與職責: 關於人工智慧之設計、開發、布署和監控方面, 金融機構內部各職務可能需要有明確的責任和義務。鑑於人工智慧模型和流程的複

雜性，採用以人工智慧為中心的系統或機器學習技術不應減少現有問責制對於人類之負擔，但的確對現有責任分配方式造成影響而有重新分配之必要性。由於英國針對金融業有所謂「資深經理人問責制度」（Senior Managers and Certification Regime, SM&CR），透過鼓勵個人責任和為個人行為制定標準，建立金融機構健康的文化和有效的治理。依照高階管理人員在金融機構所擔任之職責重要性與影響性而有不同程度的規範，職責之重要性與風險性越高者，規範程度就越嚴密^{（註 20）}。就現行制度下，高階管理人員必須採取合理措施避免違規，因此就解釋上來說，人工智慧系統若布署於其責任範圍內，高階管理人員仍無法免責，但未來可能需要就各階層之人員進行更明確之責任分配。

- （三）透明度：為了在金融領域推動負責任的人工智慧，FCA 委託艾倫圖靈研究所（The Alan Turing Institute）針對人工智慧於金融領域之應用進行風險與規範上之研究「人工智慧於金融服務」（AI in financial services）^{（註 21）}在報告中介紹了有關人工智慧之好處及風險，更重要的是一大部分討論了「透明度」（Transparency）在人工智慧中的重要性。該報告將資訊分為系統邏輯資訊及流程資訊，其透明度之揭露各有不同之重要性，主要分為六大方面，包括：可以有助於人類更了解人工智慧系統性能、合規性、提供解釋、助於人類監督、回覆顧客之查詢及了解該系統對於社會經濟之影響。更進一步將資訊分為不同生命週期與不同階段並對外部、內部利害關係人有不同影響，更細緻地使透明度具體實現，例如：何種資訊需要讓顧客知悉、何種資訊則可能必須讓審計人員知悉等問題。需要透過將資訊分類，同時考量在不同生命週期所發揮之不同功能，對應不同利害關係人而有不同透明度之展現。

- （四）監管框架：最後值得被深入探討的就是監管機構在其中能夠扮演的角色，由於監管機構一向關心人工智慧技術如何影響各種決策之作成。目前在個別領域可能有針對人工智慧發布其監管期望，例如：英國資訊委員辦公室（Information Commissioner's Office, ICO）針對資料保護於 2020 年 7 月發布「人工智慧和資料保護指南」（Guidance on AI and data protection）；模

型風險管理（Model Risk Management, MRM）一向以美國聯邦準備委員會的模型風險管理指引（SR Letter 11-7）為黃金標準，其強調模型（包括人工智慧演算法）的開發、實施和應用對銀行安全之重要性，2021 年 8 月美國貨幣監理局（The Office of the Comptroller of the Currency, OCC）則提供一份手冊強調傳統一些人類分析的任務，現在可能轉由人工智慧分析，例如：詐欺檢測、營銷分析、自動化投資諮詢等，而其人工智慧系統也可能符合模型風險管理指引所定義之模型，但無論其如何分類仍應有對應之風險管理^{（註 22）}。雖個別領域已針對人工智慧發展自身之規範架構，但針對人工智慧在金融領域應用之整體規範架構尚有不足，因為人工智慧涉及許多要素，應該更全面的了解該技術之風險及效益。在建構總體規範架構的同時，也應確保其能對應潛在挑戰及規範之落實到位。

除了整體監管方向趨於確定外，另一方面，部分人工智慧於金融領域較成熟之應用也已逐步落實到既有之監管框架之中，而此部分之管制相對具體，以下將列舉說明，也可併同作為我國之參考。

（一）演算法交易（Algorithmic Trading）及高頻交易（High-frequency Algorithmic trading）：所謂演算法交易指在有限或是無人為介入的情況下，人工智慧之交易系統快速分析來自市場的數據或資料，在很短的時間內依據相關分析自動化傳送或更新大量訂單，而高頻交易則是演算法交易的一個種類，指利用演算法進行符合特定條件之交易。歐盟金融市場工具指令（Markets in Financial Instruments Directive II, MiFID II）^{（註 23）}針對從事此二種交易的公司以及允許該類型交易的交易場所，規定其具體之監理要求，包括風險控制與管理機制、交易系統之風險抵抗力、執行程序的安排、交易紀錄的保存與回報義務及緊急處理程序等。為了有效監督，在必要時可限制或禁止演算法交易，並要求演算法的委託單加以註記，也可要求從事演算法的投資公司，定期提供演算法交易業務的說明資料^{（註 24）}，以確保其交易具有抵抗力和資本適足。

（二）自動化投資顧問服務（Robo-Advisor）：英國 FCA 目前對於自動化投資顧問服務並無專門規範，仍應遵守金融市場工具指令、反洗錢規則等既有金

融規範。而在英國 2016 年 6 月起推動法規沙盒（Regulatory Sandbox）後，設立一個專門的諮詢部門（Robo-advice Unit）協助企業建立自動化投資顧問，該諮詢部門主要有兩個目標：第一，向開發自動化投資顧問模型之公司提供監管諮詢，以利其向消費者提供低成本的建議和指導；第二，開發提供自動化投資顧問公司可通用之工具^{（註 25）}。FCA 也針對自動化投資顧問事業發布相關原則供業者參考，主要包括：資訊揭露、適用性評估、持續客戶關係、篩選客戶、弱勢客戶、整體治理及產品促銷、廣告原則等^{（註 26）}。由上述作為可知 FCA 對於這項新興服務基本上採取開放的態度，也與業者合作共同解決新的服務型態在適用舊法令時可能產生的問題，而達到監管與創新之平衡。

參、結論與建議—合法、安全應用人工智慧於金融服務

隨著人工智慧於金融領域之蓬勃應用，隨之也帶來了許多技術、道德和法律挑戰，包括數據風險、資安風險、金融穩定性風險、法律風險及倫理問題等。而傳統金融監理較為側重於外部治理，對於人工智慧所帶來的風險可能無法完整解決，主要原因是人工智慧依賴大量數據、資訊不對稱以及黑盒子問題等技術上特性，透過傳統方式監理人工智慧在金融中的使用極具挑戰性。

隨著國際上對於人工智慧監理之關注，相關規範也逐漸從自律性轉為具強制性之規定，需要規範之重點也趨於具體，針對高風險之人工智慧可能需要較為嚴密的管制。在這樣的趨勢下，如何整合金融既有規範與人工智慧監理成為關鍵，建立人工智慧應用之金融風險治理框架可能是未來治理之必然。在協助發展金融科技的同時，平衡「鼓勵創新」與「預防風險」，建立防範或控制風險的監理機制是至關重要的課題。

規範態度上可參考英國 FCA 以開放創新的態度，調控創新技術所帶來之風險，也許可先運用既有工具支持創新，例如法規沙盒和監理門診。針對特定新興科技領域持續監測，如果發現可能對消費者造成損害的不良作法，透過適當的監管工具採取行動，包括：早期干預或執法調查；而在人工智慧應用之金融風險治理框架之規範方向上，除持續關注國際間對於人工智慧之監管法制並規劃整體藍

圖外，可能需要開始一些具體作為，例如：開始要求評估人工智慧之風險，並且致力讓相關風險最小化，同時也需進行相關紀錄。該類評估方式稱為「人工智慧影響評估」（AI Impact Assessment）可幫助組織以合乎倫理和法律的方式設計、使用和審核人工智慧^{（註 27）}。加上歐盟人工智慧法規同樣強調應針對高風險人工智慧系統進行監管，該影響評估未來也將有助於辨識需規範之人工智慧系統及設計相關風險管理計畫。

除事前評估及監管外，人工智慧之事後問責也是討論焦點，因人工智慧之特性使得傳統責任歸屬之判斷方式受到衝擊，因人工智慧系統自主性可能毋須人類介入就能作出決策與判斷，在故意和過失理論之判斷下，造成無人類需要負責之狀況，此時鼓勵企業使用透明化架構、獨立標準及獨立稽核等規範相對重要。而在金融領域方面，或可參考英國「資深經理人問責制度」針對不同職位之主管或員工訂立不同可遵循之標準及責任分擔事項。同時，金融領域傳統強調之資訊揭露以弭平資訊不對稱，對應至人工智慧規範所強調的可解釋性在金融領域需要達到何種程度？需要了解相關資訊之對象可能包括公司內部人、客戶、利害關係人等不同對象，相關揭露之事項可能需要隨著人工智慧應用趨於成熟而更加具體並提升透明度。

另外，針對較為成熟之人工智慧於金融領域之應用，監管單位可參考國際作法以接軌國際，例如前述論及之演算法交易及高頻交易相關規範，為預防此類交易對危害市場秩序，應著手建立該類交易風險發生時之即時處置與事後的回復機制；在自動化投資顧問服務方面，金管會曾於 2017 年，開放相關業務並同意投信投顧公會訂定「自動化投資顧問服務作業要點」，就自動化投資顧問之定義、瞭解客戶作業、演算法監督管理及告知客戶使用該服務前之注意事項等訂定相關原則，惟相關服務仍因部分法規限制發展上顯得緩慢，也許未來可進一步朝向簡化流程、明確適用主體等方向制定更明確之規範。

註釋

註 1： World Economic Forum, Transforming Paradigms: A Global AI in Financial Services Survey, 26 (Feb. 29, 2020), available at <https://www.weforum.org/reports/transforming->

paradigms-a-global-ai-in-financial-services-survey (last visited Nov. 2, 2021).

- 註 2：中華電信，中華電信跨足金融大數據應用，與凱基銀行合作首件金融監理沙盒創新實驗計畫，2018 年 9 月 18 日，<https://www.cht.com.tw/zh-tw/home/cht/messages/2018/msg-180918-171000>。
- 註 3：Buckley, Ross P., Zetsche, Dirk Andreas, Arner, Douglas W., Arner, Douglas W. and Tang, Brian, Regulating Artificial Intelligence in Finance: Putting the Human in the Loop, *Sydney Law Journal* 43 (2021), available at https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3831758 (last visited Nov. 2, 2021).
- 註 4：Financial Stability Board, Financial Stability Implications from FinTech, Annex G (Jun. 27, 2017), available at <https://www.fsb.org/wp-content/uploads/R270617.pdf> (last visited Nov. 2, 2021).
- 註 5：Greg Swanson, Non-Autonomous Artificial Intelligence Programs and Products Liability: How New AI Products Challenge Existing Liability Models and Pose New Financial Burdens, 42 SEATTLE U. L. REV. 1201(2019).
- 註 6：AUSTRAC, AUSTRAC and Westpac agree to proposed \$1.3bn penalty (24 Sep, 2020), available at <https://www.austrac.gov.au/news-and-media/media-release/austrac-and-westpac-agree-penalty> (last visited Nov. 2, 2021).
- 註 7：金融行動工作組織使用的術語，指洗錢方法。
- 註 8：Charlotte Grieve, The Westpac Scandal: How Did It Happen?, *The Sydney Morning Herald* (9 December, 2019), available at <https://www.smh.com.au/business/banking-and-finance/the-westpac-scandal-how-did-it-happen-20191206-p53ho2.html> (last visited Nov. 2, 2021).
- 註 9：World Economic Forum, Navigating Uncharted Waters: A Roadmap to Responsible Innovation with AI in Financial Services, 69(23 Oct., 2019).
- 註 10：European Parliament resolution of 16 February 2017 with recommendations to the Commission on Civil Law Rules on Robotics, 2015/2103(INL).
- 註 11：High-Level Expert Group on Artificial Intelligence, Ethics Guidelines for Trustworthy AI, European Commission (8 Apr., 2019), available at <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai> (last visited Nov. 2, 2021).
- 註 12：High-Level Expert Group on Artificial Intelligence, Policy and Investment

Recommendations for Trustworthy Artificial Intelligence, European Commission (Jun. 2019), available at <https://ec.europa.eu/digital-single-market/en/news/policy-and-investment-recommendations-trustworthy-artificial-intelligence> (last visited Jun. 5, 2021).

註 13 : Brussels, 19.2.2020 COM(2020) 65 final.

註 14 : Brussels, 21.4.2021 COM(2021) 206 final.

註 15 : Supra note 14.

註 16 : EC, Europe fit for the Digital Age: Commission proposes new rules and actions for excellence and trust in Artificial Intelligence (Apr. 21, 2021), available at https://ec.europa.eu/commission/presscorner/detail/en/ip_21_1682 (last visited Apr. 26, 2021).

註 17 : RegTech, FCA (Nov. 24, 2021), available at <https://www.fca.org.uk/firms/innovation/regtech> (last visited Nov. 2, 2021).

註 18 : FCA, Data analytics and artificial intelligence (AI) (Oct. 15, 2021), available at <https://www.fca.org.uk/firms/data-analytics-artificial-intelligence-ai> (last visited Nov. 2, 2021).

註 19 : Bank of England, Minutes of the Artificial Intelligence Public-Private Forum - 1 October 2021 (Oct. 15, 2021), available at <https://www.bankofengland.co.uk/minutes/2021/october/aippf-minutes-1-october-2021> (last visited Nov. 2, 2021).

註 20 : FCA, The Senior Managers and Certification Regime: Guide for FCA solo-regulated firms (Jul. 2019), available at <https://www.fca.org.uk/publication/policy/guide-for-fca-solo-regulated-firms.pdf> (last visited Nov. 2, 2021).

註 21 : Ostmann, F., and Dorobantu C., AI in financial services, The Alan Turing Institute (2021), available at <https://doi.org/10.5281/zenodo.4916041> (last visited Nov. 2, 2021).

註 22 : COMPTROLLER' S HANDBOOK- MODEL RISK MANAGEMENT, OCC, V.1, 4(Aug. 2021), available at <https://www.occ.treas.gov/publications-and-resources/publications/comptrollers-handbook/files/model-risk-management/pub-ch-model-risk.pdf> (last visited Nov. 2, 2021).

註 23 : Directive 2014/65/EU.

註 24 : 曾羚、徐珮甄，〈MiFID II 對交易所影響分析〉，《證券服務》，第 670 期，頁 70（2019 年 04 月 15 日）。

註 25 : Advice Unit, FCA, <https://www.fca.org.uk/firms/innovation/advice-unit> (last visited

Nov. 2, 2021).

註 26 : Automated investment services - our expectations, FCA, <https://www.fca.org.uk/publications/multi-firm-reviews/automated-investment-services-our-expectations> (last visited Nov. 2, 2021).

註 27 : Andrade, Norberto Nuno Gomes, and Verena Kotschieder. AI Impact Assessment: A Policy Prototyping Experiment, (2021), available at https://openloop.org/wp-content/uploads/2021/01/AI_Impact_Assessment_A_Policy_Prototyping_Experiment.pdf (last visited Mar. 1, 2021).