

人工智慧（AI）於金融領域之應用與相應 監理進展

中央銀行

金融業務檢查處

原靖雯*

111 年 12 月

* 為中央銀行金檢處專員。研究人員感謝任職單位長官的指正與建議，本研究僅代表個人觀點，不代表中央銀行立場。

摘 要

近年來隨著自動化技術、大數據分析及雲端計算等新興科技蓬勃發展，AI 於 21 世紀再次掀起熱潮，透過機器學習（machine learning, ML）、深度學習（deep learning, DL）及大數據分析等技術之運用，廣泛應用於人類生活中各種領域，亦帶動 AI 相關產業鏈。當聚焦於金融領域，國際間應用 AI 技術於金融服務領域範疇大致可分為客群經營、流程精進、風險合規及預測分析等 4 類，期透過運用 AI 技術來提升經營效率及節省人力成本；國際間監理機關為因應當前金融數位環境所衍生之服務，將原監理制度導入創新技術，增進監理能力與效率，並已嘗試運用 AI 技術進行金融監理。此外，運用 AI 技術於金融領域可能衍生之相關潛在風險亦值得金融監理機關密切關注。

目前國際間已有研擬 AI 應用相關治理準則之趨勢，大致上皆為對倫理道德、尊重人權及個資保護等議題進行探討，較完善且常被引用的基本原則為 OECD 及 G20 共同發布之「AI 原則」，惟國際間對 AI 在金融領域應用之監理法規仍屈指可數。目前我國金融機構運用 AI 技術於金融領域多為客戶服務及機器人理財等範圍，相關主管機關僅科技部於 2019 年發布「人工智慧科研發展指引」之原則性指引，金融監理機關則認為目前金融機構多數僅將 AI 技術作為決策或作業之輔助工具，尚非取代最終人為判斷，故對其金融監理尚處於初步發展階段，並對使用新興科技技術之使用保持審慎開放之思維。

鑑於國際間對 AI 應用於金融領域之監理仍處起步階段，認為 AI 發展仍著重於建立結合人類思維及 AI 之優勢的夥伴關係，並強調 AI 系統「旨在增強而非取代人類貢獻」，金融監理機關未來可能發展方向應朝打破舊有思維、建立民眾信任之數據文化、設立共享機制，以及進行跨部會研議或產學合作，應保持科技中性（tech-neutrality）之監理原則，提升創新科技研發意願，逐步建構對金融業運用 AI 技術之監理。

目 錄

壹、人工智慧（AI）之發展.....	1
一、AI 之意涵與發展歷程.....	1
二、AI 技術應用發展.....	4
三、AI 技術產業鏈及產業應用.....	6
四、人類思維及 AI 之差異與結合.....	10
貳、AI 於金融領域之應用與風險.....	13
一、AI 於金融領域之應用技術.....	13
二、金融服務領域之 AI 應用範疇.....	16
三、金融監理領域之 AI 應用範疇.....	21
四、AI 於金融領域應用之潛在風險.....	26
參、國際間對 AI 之監理進展.....	28
一、發布 AI 治理準則.....	28
二、國際間對 AI 於金融領域應用之監理進展.....	32
三、國際間對 AI 運用之治理原則與面臨之挑戰.....	35
肆、我國業者運用 AI 於金融領域之現況與監理機關對 AI 應用之監理 ...	36
一、我國金融機構運用 AI 現況.....	36
二、我國主管機關對金融業創新營運模式及 AI 應用之監管態度	37
伍、結語.....	39
附錄.....	41

圖目錄

圖 1	AI 發展歷程.....	2
圖 2	AI 與 ML、大數據分析之關係圖.....	3
圖 3	AI 技術應用發展階段.....	5
圖 4	AI 技術產業鏈.....	7
圖 5	全球 AI 民間投資.....	8
圖 6	AI 於金融領域應用之潛在風險.....	26
圖 7	本國銀行運用 AI 技術之領域及目的.....	36

表目錄

表 1	AI 於金融服務領域之應用.....	16
表 2	各國金融監理機關運用 AI 技術之案例	25
表 3	國際組織或各國對 AI 發布治理準則.....	31
表 4	AI 運用之共通治理原則及面臨之挑戰	35

壹、人工智慧（AI）之發展

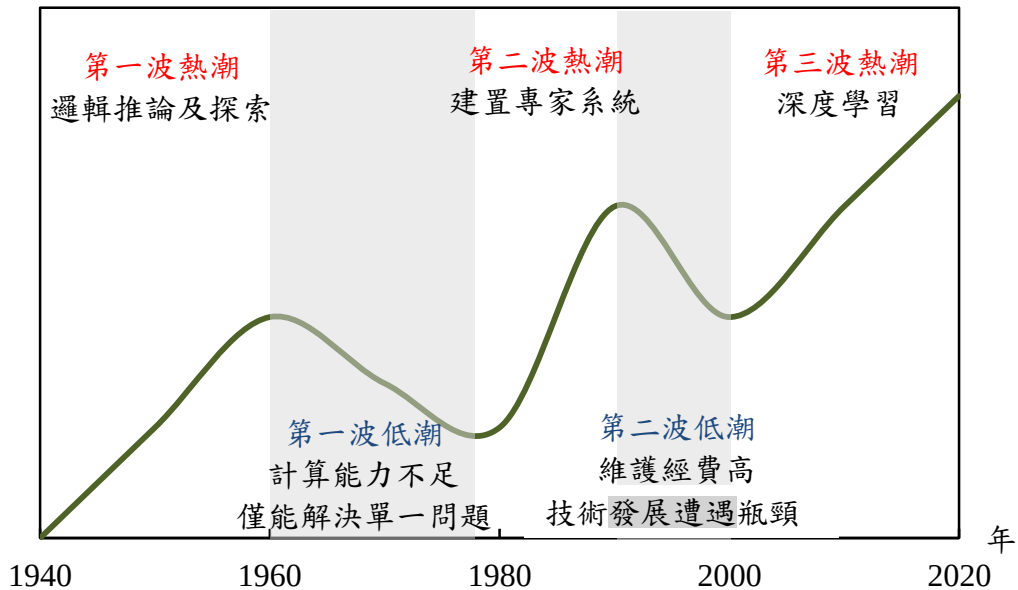
一、AI 之意涵與發展歷程

人工智慧 (artificial intelligence, AI) 泛指「擁有人類智慧的機器」或「能夠自主思考的電腦」，最早起源於 1950 年用來判斷機器是否具有思考與判斷能力的圖靈測試 (Turing test)，AI 一詞則係於 1956 年達特茅斯 (Dartmouth) 夏季人工智慧研究計畫研討會¹中首度提出，並嘗試利用數學、統計及邏輯套用至機器處理程序，掀起 AI 第一波發展浪潮，惟隨後囿於計算能力有限及僅能解決特定單一問題，而遭遇第一次發展低潮；1980 年代發展焦點轉移至以建構知識處理的專家系統，透過蒐集特定範圍的專家知識並儲存於資料庫，模擬人類演繹推論能力，迎來第二波 AI 研究興起，然而系統維護經費居高不下，伴隨技術無法超越人類高度期望，導致第二次發展瓶頸；至 21 世紀類神經網絡 (neural network) 研究興起，經由模仿生物大腦神經元運作模式所建立的模型，為 AI 發展帶來新的學習運作方式，FSB (2017) 分析，由於 (1) 電腦處理運算能力提升與記憶體儲存容量與日俱增，以及大規模雲端運算技術已易取得；(2) 系統性紀錄及學習演算法可利用之數據遽增；(3) 開放原始碼軟體 (open-source software) 降低機器學習 (machine learning, ML) 應用程式之開發成本，通常可將 ML 模型濃縮至幾行代碼；(4) 新型演算法之創新使 ML 得以解決具挑戰性的問題等因素，驅使 AI 應用有突破性成長發展，成為 AI 第三波研究熱潮 (圖 1)。

廣義 AI 下之技術分類包括機器學習 (ML)、深度學習 (deep learning, DL) 等，其中 ML 運用演算法將大量且多元的資料分類並找出規則，透過經驗累積，以自動最適化方式進行資料分析，其依據有無資料標註 (label) 與人工參與情形，分為監督學習 (supervised learning)、非監督學習 (unsupervised learning) 及強化學習 (reinforcement learning) 等 3 類 (圖 2)。

¹ 達特茅斯夏季人工智慧研究計畫發起於 1956 年 8 月 31 日，旨在討論「人工智慧」，會議持續一個月，催生後來人工智慧革命。

圖 1 AI 發展歷程



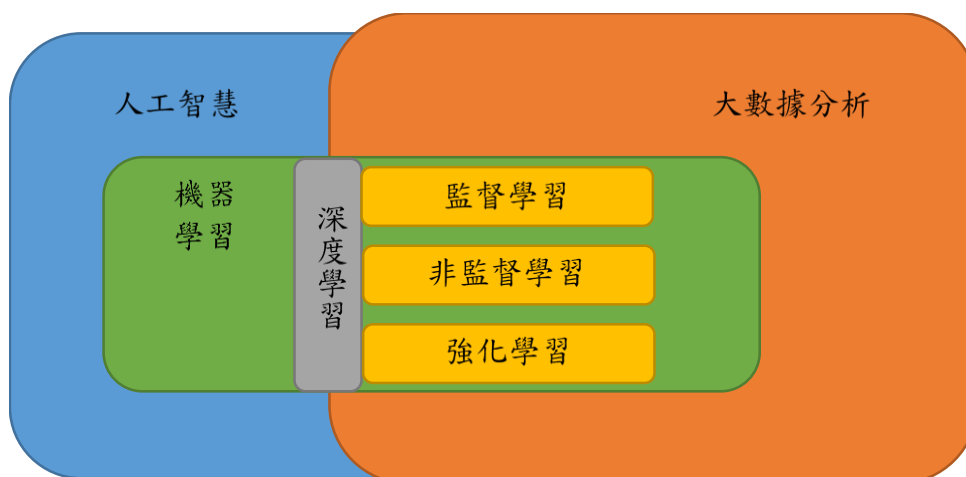
資料來源：李開復（2017）、李震華（2019）；本文整理。

監督學習主要係利用 ML 演算法學習輸入資料及輸出資料之間的關聯性，輸出資料為監督訊號，係 ML 模型由輸入資料學習推斷而得，例如電腦從包含借款人特徵（輸入資料）及其是否違約（輸出資料）的資料集中學習，接著使用 ML 模型預測未來借款人違約之可能性；在**非監督學習**中，ML 模型由輸入資料中探索資料脈絡，且無輸出資料監督訊號，例如非監督 ML 模型根據借款人相似性區分借款人，或辨識整個資料集的異常資料點；**強化學習**則為電腦嘗試探索環境來學習最佳策略的行為序列，例如強化學習透過多次與自己對戰來學習掌握棋盤遊戲之技巧，設計該系統於比賽獲勝時得到獎勵，比賽敗陣時受到懲罰，其設計程式僅為爭取獎勵，在開始學習前並未配備任何策略（圖 2）。

DL 則為利用電腦模擬人腦進行多層次節點分析之高級演算法，具有自動抽取資料特徵（feature extraction）並加以分類（classification）的能力，隨著資料量越大，DL 可對每次分析結果進行改進，常見的應用包括圖像辨

識、語音辨識及個人化推薦系統²。DL 運作之原理主要可分為卷積神經網路（convolutional neural network, CNN）及遞歸神經網路（recurrent neural network, RNN）之 2 種類別。CNN 透過不同階級的特色來識別圖像，適合運用於圖片、影片等空間數據類型，例如從單純鼻子、眼睛、嘴巴之特徵到三者彼此的連結關係，最後再變成人臉辨識。RNN 則基於所有輸入資訊皆為相關之前提下，在訓練學習過程中會記住所有處理過的資訊，適合運用於語音、文字等序列數據類型。此外，**大數據分析**則是包括利用 AI 各種技術儲存、分析並管理大量複雜資料，而演算法之優化亦須憑藉大數據分析之協助。

圖 2 AI 與 ML、大數據分析之關係圖



註：1.監督學習利用人工標註資料，累積經驗以預測未標註資料。

2.非監督學習採用無人工標註資料，透過觀察不同資料間是否具相似特徵，偵測資料模式進行學習。

3.強化學習使用無人工標註資料，並人工設定獎勵回饋，以因應未來類似情況。

4.深度學習利用類神經網絡之高級演算法，可用於監督學習、非監督學習或強化學習。

資料來源：FSB（2017）；本文整理。

² 圖像辨識如應用於醫療判別病人 X 光片是否有病變跡象、判斷交通號誌燈之自動駕駛系統、社群軟體之人臉辨識系統；語音辨識如 Google Home 及 Siri 語音助理服務；個人化推薦系統如電商平台或串流平台推薦用戶產品或影音內容等。

二、AI 技術應用發展

AI 技術的應用發展階段，普遍可區分為 5 個層級，由底層到頂端依序是計算及記憶能力（computation）、感知能力（perception）、認知能力（cognition）、創造力（creativity）以及智慧（wisdom）（圖 3）。

（一）計算與記憶能力

重點在於協助大量且複雜的資料處理，依據預先排程完成既定工作。

（二）感知能力

係指人類透過感官所接受到的環境刺激，而衍生察覺訊息的能力，即為人類五官之視、聽、說、讀、寫等能力，為目前 AI 技術主要發展的焦點之一，例如影像辨識、物件偵測、語音辨識、語音生成、文字轉換語音、文字辨識、語音轉換文字、機器翻譯等。

（三）認知能力

人類透過觀察辨識、判斷推論、分析學習、決策規劃等行為來瞭解訊息、獲取知識的過程與能力，在未事先規範下，透過經驗學習產生規則，並自主採取對應行為，亦為未來發展 AI 技術之願景，例如醫學圖像分析、健保詐欺判斷、個人化產品推薦系統、垃圾郵件辨識、犯罪偵測、信用風險分析、消費行為分析、天然災害預測與防治等。

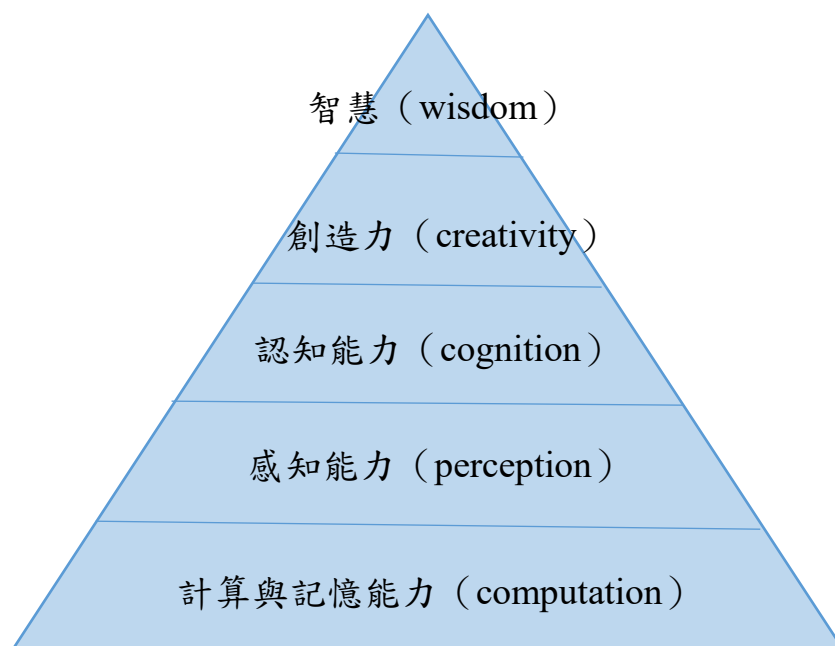
（四）創造力

係指在無資料或資料不充分的情況下，人類可憑空推論出用來創造新事物的邏輯見解，且鑑於人類嘗試新事物的動機多元，例如對學習與探索的好奇心，以及盡可能提升技能及知識，人類自然而然傾向嘗試新事物，將自身的知識、個性及潛意識等各種因素匯集而啟發創造力。AI 技術於創作領域亦嘗試仿效人類思維，例如創作小說、譜曲、繪圖等，惟目前尚未有明顯之成果。

（五）智慧

智慧為人類探求人、事、物真相過程所獲得之思考、認知、分析能力，該領域牽涉人類的自我意識、自我認知與價值觀，是目前 AI 尚未觸及與最難模仿之領域。智慧與智能涉及領域不同，智慧包含意識與心智，人類意識為生物特有，心智則與創造力相關，創造力可被激發，但無法透過學習教導而得，亦沒有演算法可遵循。目前科學家無法解釋人類的意識、心智與創造力之連結，遑論將其系統化，許多科學研究甚至指出，人類在意志力低迷的狀態下，例如洗澡或睡覺，反而更具創造力。

圖 3 AI 技術應用發展階段



資料來源：本文整理。

圖 3 下半部的計算與記憶能力，AI 已發展得比人類更佳；而中間層級的感知及認知能力，則是目前 AI 技術主力發展的焦點領域；至於上半部的創造力與智慧，AI 技術應用仍有努力空間，人類還是獨占鰲頭。可預期的是，未來將是人類與 AI 搭配的時代，先由人類「大膽假設」，再運用 AI 技術「小心求證」，透過 AI 及人類的結合完成重要的決策或任務。

三、AI 技術產業鏈及產業應用

隨著大數據資料分析興起及數位化轉型經濟逐漸成形，AI 技術日漸發展成熟，其應用領域日益廣泛，對產業界與消費者的互動模式產生極大轉變，加以消費者自我意識抬頭，產業界可藉由 AI 技術提供更貼近個人化需求之服務，與消費者建立即時且完善之互動。

（一）AI 技術產業鏈

AI 技術產業鏈大致可區分為 3 個層次，分別是**運算資源**、**核心技術**及**應用與服務**（圖 4），其中（1）「運算資源」係指提供資料擷取、儲存與處理的運算服務，產業類型主要可分為生產伺服器及運算晶片之運算設備業者³、建立雲端軟體環境之雲端平台業者⁴以及提供資料面軟體或服務之資料處理業者⁵；（2）「核心技術」係指提供各種數量方法、統計模型與仿生物模擬等演算法為基礎所發展的技術，或是提供演算法調校、模型建構等服務，例如機器學習、電腦視覺、自然語言處理以及移動控制⁶等；（3）「應用與服務」指利用 AI 技術，開發特定應用領域之產品或服務，例如場域解決方案供應商⁷與智慧設備製造商等業者，以及提供 AI 所需的支援服務，涵蓋系統整合及顧問諮詢等產業。

若將 AI 體現於產業應用領域中，根據國際知名市調公司 CB Insights 「2022 人工智慧 100」報告⁸，將 AI 應用領域劃分為 3 個子類，分別為（1）**跨產業應用**，例如研發倉儲及物流機器人、預防工安意外等；（2）**特定產**

³ 生產伺服器業者代表為 IBM、惠普（HP），生產運算晶片業者代表則為輝達（Nvidia）。

⁴ 例如微軟（Microsoft）、谷歌（Google）及亞馬遜網路服務（AWS）。

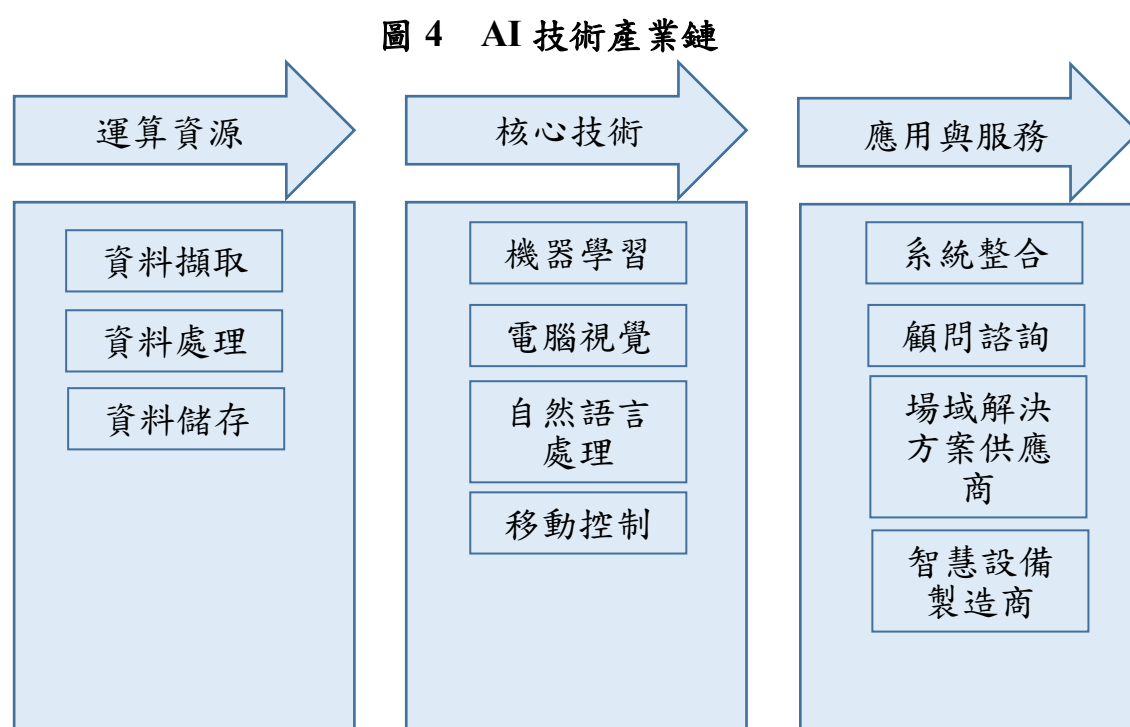
⁵ 例如賽仕（SAS）、甲骨文（Oracle）。

⁶ 例如無人載具，透過感測器與機器學習技術，讓 AI 自行學習其與環境互動之因果關係，進而做到自主修正，並促進人機之間更佳的協作應用。

⁷ 提供將 AI 技術應用於實際生活場域之解決方案，例如交通運輸、防災、氣象預報、醫療等。

⁸ CB Insights （2022），“AI 100: The most promising artificial intelligence startups of 2022,” May <https://www.cbinsights.com/research/report/artificial-intelligence-top-startups-2022/>

業應用，其中以醫療保健、金融保險及線上遊戲產業應用⁹最受關注；以及
(3) AI 開發工具，研究數據分析、深度學習及演算法模型運算之解決方案，
以支持 AI 技術管理。



資料來源：櫃買中心；本文整理。

(二) AI 技術產業應用

至於 AI 常見的應用案例尚涵蓋 (1) 智慧醫療：例如用於放射科 x 光及眼科虹膜掃描之 AI 影像判讀技術，推動遠距健康照護，以提升醫療照顧效能；(2) 智慧製造：結合物聯網¹⁰及大數據資料，AI 演算法提升廠商在產品瑕疵檢測¹¹、良率控制、機器設備故障預防、排除及自主維護等方面的效率，亦改善以往人類生產實驗之錯誤結果，所造成的材料浪費及時間延

⁹ 2022 年報告提及 AI 技術應用於醫療保健業為外科手術之精進、罕見疾病藥物之研發；金融保險業為從事視覺評估汽車損壞情形、維護彙整及去識別化的金融數據集；線上遊戲業為開發遊戲中惡意行為玩家之改善系統。

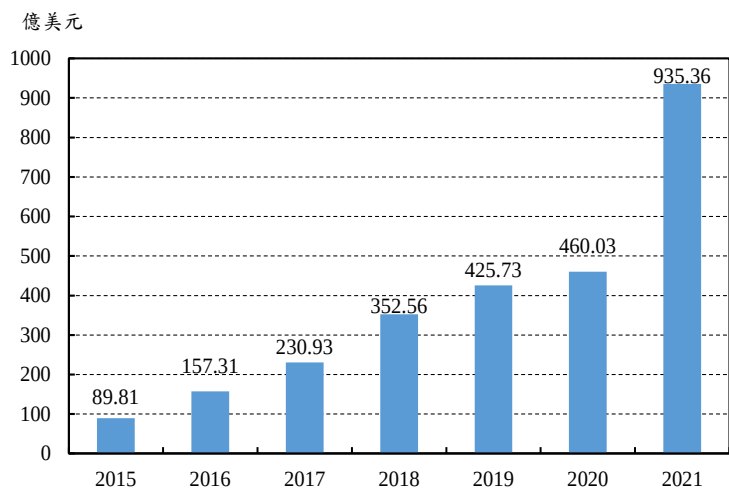
¹⁰ 透過內建感測器、軟體與其他技術的互連物件，例如穿戴裝置、車用電腦等設備，與其他硬體設備傳輸和接收使用資料。

¹¹ 例如半導體製造業者透過影像辨識技術解析不良晶圓類型，並透過大數據資料分析管控產品良率。

遲，最佳化原材料的分配比；（3）智慧家居：知名品牌如蘋果 HomePod、谷歌 Google Nest、亞馬遜 Alexa 等智慧型助理，將 AI 與網路通訊及智能家電裝置結合，成為 AIoT（人工智慧物聯網）¹²，同時滿足安全與便利的生活方式；（4）智慧金融：可藉由運用 AI 技術降低人力成本、提升效率、偵測風險及發展創新業務型態，常見於服務營運、市場行銷及風險監控等應用範疇；（5）智慧運輸：透過 AI 影像辨識技術，增強車輛辨識、號誌管控及交通安全管理等資訊整合，以發展車流計算、路況安全預警、路網優化等系統，另結合 AI 與智慧交通系統應用之自駕車問世，亦帶動產業升級並提高運輸服務的安全性及可及性。

隨著 AI 自動化技術之發展日益成熟，帶動產業與服務之創新已成趨勢，根據史丹佛大學（Stanford University）統計¹³，全球 AI 民間投資金額自 2015 年起逐年攀升，2021 年達 935.36 億美元（圖 5），顯示即便投資市場於 Covid-19 疫情的衝擊下，投資者對 AI 的熱衷程度有增無減，AI 發展仍吸引大量全球資金，業者亦積極投入研發作業。

圖 5 全球 AI 民間投資



資料來源：Stanford University（2022）。

¹² 發展成熟的物聯網（IoT）與人工智慧（AI）技術匯流進化成 AIoT，人工智慧如同 AIoT 的中樞神經，IoT 即是周圍神經系統，常見應用功能為環境調控、能源管理、家庭安全、居家健康照護等。

¹³ 史丹佛大學自 2017 年起每年發布人工智能指數報告（Artificial Intelligence Index Report），主要由研發、技術性能、經濟、教育、AI 應用之倫理挑戰、AI 從業者之多樣性、AI 政策與國家戰略等面向每年探討 AI 發展。

然而，AI 技術雖賦予電腦認知技能而能與人類在現實環境中互動，惟其對社會層面造成之影響亦**有利有弊**。正面效益主要係 AI 高度自動化可減少人類重複性之工作，且機器無疲累感以利降低錯誤頻率，重塑產業之價值鏈；負面衝擊則包括行為模式預設化與系統化及低技術性人力之取代。鑑於 AI 仍大量仰賴人類進行維護與開發，且無法完全替代人類完成創造、情感交流及解決高複雜度的問題，故**AI 不可能完全取代人類，兩者的互動才是未來發展之關鍵**。

四、人類思維及 AI 之差異與結合

Alan Turing 於 1948 年發表其計算機理論模型時曾表示：「一位有紙、筆、橡皮擦並堅持嚴格行為準則的人，實質上就是一台通用機器」。幾十年來，人腦一直被視為一台處理訊息之計算機，認為大腦處理（感覺）輸入資訊係使用計算及描述來創造輸出資訊、思想或行為。雖然測量大腦活動的技術日益精進，惟理解人類智慧的神經過程仍相當有限。

人類具備高度優化的無意識感官程序，由少量數據點即可有效地學習，例如一個孩子僅看過幾張照片就能容易地辨識出一隻長頸鹿。相較之下，ML 須透過大量數據及計算資源進行視覺對象辨識，且模型通常不具備現實世界的任何理論或直觀理解（Russell（2019））。然而，人類不僅是直覺物理學家，也是直覺心理學家，意味著人類可自然地駕馭社交互動，也是助理機器人須學習之處。

另一方面，與 AI 相比，數學計算對人類思維而言相對困難。在近代歷史中，這些議題變得尤為明顯，而人類生物大腦似乎沒有針對這些問題進行優化。鑑於計算之限制，人類的決策通常會不一致，例如，給定 2 個相同的醫療紀錄，醫生可能會做出不同的診斷。此外，無聊、疲勞及分心亦會加劇不一致之程度，以下列出人類與 AI 決策差異之處。

（一）因果理解

多數學者認為人類直觀理解因果關係之能力為人類認知進化的關鍵。Lake et al.（2017）將人類的學習能力描述為建立一個世界模型，以解釋何謂現實並想像什麼是可以發生及不能發生。身為擁有對現實世界豐富經驗的直覺物理學家，人類可自動理解對於 AI 來說並不明顯之因果關係，例如降雨是街道潮濕的原因。

然而，因果推理不是準確預測的先決條件，即使沒有任何因果理論知識，若模型數據豐富且預測命題單一，模型亦可做出準確的預測，例如降

兩預報。若缺乏數據，沒有因果推論的預測則深具挑戰性，代表因果理論模型在選擇相關變數作為預測因子之重要性。

（二）決策過程透明度

人類可於決策過程中進行推理並向他人解釋，儘管其敘述可能並未反映決策的真正驅動因素，人類對認知過程的理解有限，因此未意識到驅動行為的因素。即使人類認為自身已考慮許多訊息，**仍經常根據少量訊息而迅速形成意見**（Klein and O'Brien（2018））。此外，有意識或無意識的情緒亦會影響人類決策（Lerner et al.（2015）），表示**人類決策的複雜性及不透明性**。

ML 模型亦因缺乏透明度而遭受批評，以最嚴格的意義進行探討，若人類可於短時間內根據模型的輸入資料及參數計算出預測結果，則該模型即具透明性（Lipton（2018））。部分文獻指出一個具可解釋性的 AI（Explainable AI）係藉由拆解個別預測因子成為個別參數來對 ML 模型決策進行解釋，因此人類可描述個別變數的邊際效果及之間的交互作用，但無法簡明地表達模型已藉由學習獲取所有變數的複雜交互作用，亦即人類無法理解其內部工作原理。

（三）決策偏誤

人類思維常受到偏見的影響，例如傾向於關注先入為主的訊息，或對自身判斷的準確性過度自信。然而，當模型使用偏誤的數據進行計算時，估計結果亦可能存在偏誤。例如，部分文獻已證明 ML 模型存在種族及性別偏誤。相較之下，不太可能證明人類的特定決策存在偏見，由於人類決策不具確定性，且人類通常沒有意識到影響自己決策之偏見。

（四）目標設定

為 AI 模型設定目標可能具有下列挑戰：（1）**難以量化的議題**，例如如何設計一個具公平性的賦稅系統；（2）**靈活性**，例如 AI 如何在符合人

類道德觀狀況下，面對各種狀況？一輛自駕車是否應轉向年邁長者，以拯救三個路邊玩耍的小孩？鑑於每個人內在思維不一致及文化差異，難以衡量人類的偏好；以及（3）設計者難以預測 AI 模型效用極大化所產生之負作用，例如機器人透過獲得獎勵學習從 A 點移動到 B 點，卻破壞路徑上的花瓶。

人類則擅長解決抽象目標，因為人類可將複雜的任務分解為幾個較小的任務，例如準備晚餐時，人類知道須為一餐購買食物、備料及桌邊服務。相較之下，目前 AI 仍難以解決如此高層次之任務，因為從數據中辨識相關子目標對機器人來說仍具挑戰（Russell（2019））。

（五）創造力

AI 透過 ML 模型通常僅獲得技術之優化，難以將技術轉移到其他領域（Marcus（2018）；Chollet（2019）），僅於既定的規則框架下做出最適化決策，然而人類卻可利用豐富經驗及知識來解決新問題。

因此，Jarrahi（2018）指出未來發展將是建立結合人類思維及 AI 之優勢的夥伴關係，並強調 AI 系統「旨在增強而非取代人類貢獻」，運用 AI 分析能力來收集與處理數據，而人類則根據統整的證據做出判斷。Buckmann et al.（2021）亦指出，AI 與人類互有優劣，若涉及快速處理大量數據，AI 表現優於人類；惟若涉及監理機關或管理人員所面對的道德及風險問題，則人類表現得更佳。在以人為本的合作夥伴關係中，AI 是人類完全保留控制權的工具，例如臨床醫生追求的人機合作夥伴關係為提升機器預測準確度，再由人類解釋並決定下一步行動（Verghese et al.（2018））。

貳、AI 於金融領域之應用與風險

一、AI 於金融領域之應用技術

因應 AI 技術發展日漸成熟之新時代，金融業者為掌握數位化轉型商機，積極於金融領域規劃導入 AI 技術，不論是結合跨領域業者的數據應用，或是智慧互動的新通路，皆可能再造新型商業模式，催生新產品、新服務、新市場或新通路，以下介紹金融業界常運用的應用技術：

（一）機器學習

藉由導入 ML 演算法模型及深化各項資料來源之應用，跳脫僅仰賴邏輯知識的推理運算，並以量化的權重與公式，**估算潛在變數對於決策的重要性**，進而**發展推論推薦系統**，例如投資理財規劃，不再只根據市場價格來建議最適組合，更應綜合考量持有天數、投資人偏好等其他預測資訊，計算最終購買各商品的機率值，提供客製化最佳資產管理組合。透過 ML 技術將消費者對金融商品的抽象偏好，例如「保守穩健個性」、「手續費敏感」等因素進行估算，可彌補資料科學難以量化抽象概念的缺點，歸納出消費者對於購買意願或金融投資標的之品味、信心及風險承擔等無法直接衡量的心理層面影響因子，同時提高問題分析的全面性，以達成 AI 對問題情境的高度掌握，有助於提供最佳的解決方案。

（二）自然語言處理

自然語言處理（natural language processing, NLP）為結合 AI 及語言學的分支技術，透過複雜的數學模型及演算法，**讓電腦辨識、分析、理解及生成人類語言**，以使電腦與人類溝通更順暢，完成各項指定任務。NLP 領域中又可細分為自然語言理解（natural language understanding, NLU）與自然語言生成（natural language generation, NLG），常見的 NLU 應用包含情緒分析、文件辨識與自動化分類、自動推理；常見的 NLG 應用包含自動文

案生成、專業文件生成（通常適用於制式化程度較高的文件範疇）、聊天機器人等。

（三）語意分析

語意分析（semantic analysis）是指從一段文字內容摘錄大意，找出文章中的重要詞彙及摘要，可於短時間內使讀者快速瞭解內文，亦適用於解讀影片、音訊等檔案，使用者得以省去大量時間觀看影片、聆聽音訊，同時幫助使用者提前了解影片與音訊的內容，相關發展技術包括以文找文¹⁴、廣告信偵測¹⁵、情感分析¹⁶、文本標記/分類¹⁷等。

（四）電腦視覺

電腦視覺主要係透過 CNN 技術，在影像處理與分析過程中引導系統，可支援影像分類與物件偵測的 DL 推斷。經過完整訓練後，電腦視覺可進行文字、語音、圖形、影像之辨識、偵測與人員識別，甚至能追蹤物件的移動。金融業界運用光學字元辨識（optical character recognition, OCR）¹⁸及影像辨識技術於解析個別消費者的證明文件影像檔，利於辨識消費者書寫的文字，減少重複輸入，例如財力證明、醫療診斷證明書等，偵測影像檔中關鍵資訊的位置，將紙本作業推向數位化，有效彌補人工作業對流程自動化造成的斷點，以提升作業效率與正確性；或運用影像辨識技術作為消費者身分識別，簡化原本須輸入使用者帳號與密碼的不便，例如網銀 APP 登入系統的人臉辨識。

¹⁴ 透過文章中重要、關鍵的字詞，進而找出相關性高且類似之文章。

¹⁵ 語意分析能判斷一封信中是否包含廣告信的常用字詞，協助信箱進行篩選。

¹⁶ 根據提到相關特定企業或產品的文章進行分析，篩選文章中的正負面字詞，分析企業或產品在網路上的口碑及評價，分析消費者的觀點與情感。

¹⁷ 處理大量文本資料時，透過 ML、DL 等模型訓練後，可自動對每篇文章分類、標記出關鍵屬性詞彙。

¹⁸ OCR 係指對文字資料的圖像檔案進行分析辨識處理，取得文字及版面資訊的過程，為一種將圖片或掃描文字轉換為數位資料的技術。

（五）流程自動化

機器人流程自動化（robotic process automation, RPA）係複製人類行為的自動化科技方案，擴大執行日常任務的範圍和數量，可自動執行有規則且重複性高的作業流程，處理結構化數據，或結合圖像辨識處理非結構化數據，代替人工進行資料的驗證與轉換，減少人工洞察時間及人力成本，降低人為處理的錯誤率，提高生產力及效率，提升作業流程之透明度。

二、金融服務領域之 AI 應用範疇

在全球數位化經濟蓬勃發展及近年 Covid-19 疫情加速推動新科技商機之趨勢下，各產業紛紛投入 AI 技術之應用，其中金融業透過與科技公司合作，運用 AI 技術導入人性化的智能服務，本文參考金融穩定委員會（FSB）報告，將相關實務應用區分為客群經營、風險合規、流程精進及預測分析等 4 類，並分析應用案例所使用之 AI 技術（表 1）：

表 1 AI 於金融服務領域之應用

AI 技術 實務運用		機器學習	自然語言處理	語意分析	電腦視覺	流程 自動化
客群 經營	智慧客服		✓	✓		
	理財諮詢	✓	✓	✓		
	信用評分	✓			✓	
流程 精進	作業流程 簡化				✓	✓
	身分辨識				✓	✓
	高頻交易	✓				✓
風險 合規	認識客戶	✓	✓	✓		
	詐欺偵測	✓	✓	✓		
	法遵科技		✓	✓	✓	✓
預測 分析	精準行銷	✓	✓	✓		
	預測市場 趨勢	✓	✓	✓		
	風險管理	✓	✓			

資料來源：FSB；本文整理。

（一）客群經營

「客群經營」重視如何應用 AI 技術服務客戶，改變金融機構與消費者之間的互動模式，應用如下：

1. 智慧客服

透過虛擬助理以 NLP 技術進行智能答詢，可負責申辦流程較為簡易且

交易特性屬於**常態重複性發生**的服務，透過大量資料分析以完成自動化服務，建構一問一答的方式回覆消費者諮詢，消費者亦可自行瞭解產品服務功能，提高與客戶往來的黏著度及忠誠度。

2.理財諮詢

以 ML 演算法為核心的**理財機器人**將依照客戶不同的財務目標及需求，例如財務目標、風險容忍度、投資標的範疇等，依據投資偏好進行客群特性分析後，提供客製化理財服務，推薦客戶不同的投資組合及資產管理計劃。另金融機構可分析客戶過往投資成功案例，彙整為內部知識管理，對於 AI 應用於客群經營亦有所助益。

3.信用評分

金融機構以往主要依賴歷史交易紀錄或信用評等分數作為核貸准駁依據，例如美國 FICO score、英國 Experian score 及台灣聯合徵信中心信用資料等，分析之面向有限且無法適用無信用評分的潛在客戶群，現今金融機構輔以 AI 技術將社群媒體及第三方支付業者等**其他數據來源納入分析**，可開發潛在客群，增加蒐集客戶行為資料，並改善信用分級之準確性。

（二）流程精進

「流程精進」注重運用 AI 技術來降低銀行作業成本，並增進效率，應用如下：

1.作業流程簡化

利用 RPA 建構銀行前、中、後台的業務流程自動化，例如開戶審核自動化亦可結合光學與影像辨識功能，推動授信與徵審流程自動化，或用以加速保險理賠處理程序，簡化人工檢查財務文件（如發票或薪資單）之繁瑣作業，自動獲取數據來評估申請案件，以**減少營運作業成本及理賠處理時間**。

2. 身分辨識

透過**生物辨識 (Biometric) 技術**，包含臉部、語音聲紋、虹膜、靜脈、指紋等生物特徵，作為客戶進行金融交易及於特定場域安全防護時辨識身份的主要方式，俾強化營運安全及效率。

3. 高頻交易

高頻交易係以**演算法**決定下單時機、價格及數量等交易決策，於極短時間內執行大量自動化交易，利用 AI 技術結合趨勢判讀，分析即時市場數據與信號，降低人為情緒干擾因素，有助於交易時機判斷，亦可提高交易效率。

(三) 風險合規

「風險合規」關注 AI 技術在風險防範上之應用，根據過往的風險管理經驗，偵測並降低金融機構交易風險，應用如下：

1. 認識客戶

銀行利用 **NLP 技術**及**智能檢索**等 AI 技術於**認識客戶** (know your customer, **KYC**) 或**客戶盡職審查** (customer due diligence, **CDD**)，可透過自動蒐集及驗證已發布之新聞、公開資訊或報告，產生客戶風險分級及管理報告等，進而發現異常交易狀況，俾減少作業時間成本、提高交易安全性與法令遵循，並可應用於洗錢防制 (anti-money laundering, AML)。

2. 詐欺偵測

針對具有高敏感度、時效性及不易判斷等特徵之洗錢、資恐及保險理賠等不當或詐欺行為，運用 **AI 技術**篩選可疑交易，對客戶進行**行為模式分析**，並與外部資訊比對可疑的關係戶或警示帳戶，進行 AI 智能風險控管，提供**不當行為之預警及防範**。

3.法遵科技

為因應創新科技興起且日趨複雜的法規與市場環境，以及政府各項制度實施及法規更新，法令遵循工作日益繁瑣，可透過 AI 智能檢索技術進行關鍵字搜尋，快速協助判斷營業規章與作業流程手冊應調整之範圍，亦可運用於定期過濾與檢視行員與客戶間的交談錄音及錄影資料，以節省繁雜之人工稽核作業時間。

（四）預測分析

「預測分析」結合各種外部資訊，著重行為趨勢分析，可從市場行銷及支援預測等面向，說明 AI 於數據分析的應用：

1.精準行銷

透過 AI 技術、大數據分析與雲端計算，觀察各類顧客特性，包含購買行為、客戶特徵、社群行為分析等，並利用數據開發預測模型，預測消費者的行為及偏好，持續根據消費者的反饋結果進行改善，並可作為區隔不同商品與訂價之工具，提供差異化商品組合與金融服務，使商業預測更加精準，推薦金融商品更能滿足消費者需求。

2.預測市場趨勢

隨著金融機構運用 AI 蒐集新聞、報告及公開資訊之技術精進，提升趨勢分析及預測市場指標之效能，包括數位行銷、風險模型、投資模型等方面皆有進展，例如以新聞媒體資訊強化對總體經濟變數之預測、衡量經濟活動對貨幣政策公告之反應等，有助於預測市場流動性及價格波動度，進而提高投資報酬率。

3.風險管理

金融機構利用 ML 演算法對大量非結構化¹⁹及半結構化²⁰資料進行分析，考慮到市場行為、監管規則及其他趨勢的變化，進行事後檢驗、模型驗證及壓力測試，**避免低估風險，提高風險管理透明度。**

¹⁹ 非結構化資料是缺乏可識別的結構或架構的數據，代表不符合預先定義的資料模型，通常包含大量文字、語音及圖像等，由於沒有易於辨識的結構，會產生傳統電腦程序難以理解的不規則性及模糊性。

²⁰ 半結構化資料為介於結構化資料與非結構化資料之間，這種資料具有欄位，可依據欄位進行查找，但不保證欄位之一致性。

三、金融監理領域之 AI 應用範疇

金融機構運用 AI 與 ML 等創新科技，可提供金融創新服務、改善金融服務流程、降低作業成本並增進營運效率，全球金融監理機關面對此一數位化科技浪潮，如何有效監督金融市場不當行為，確保金融穩定及消費者權益，亦至關重要，故現行各國監理機關紛紛嘗試運用 AI 技術支援金融監理，提升監理效能，以促進 AI 於金融監理之應用。

根據 FSI 及 FSB 對全球國家（或地區）的金融監理機關之調查報告顯示，部分金融監理機關已嘗試運用 AI 技術於不同的金融監理應用範疇，本文彙整如表 2，謹說明如下：

（一）資料申報

1. 奧地利央行（OeNB）

開發連結金融機構及監理機關資訊系統之整合性資料申報平台（Austrian reporting services, AuRep），透過標準化轉換規則，將資料轉換成符合監理機關規定之資料內容與規格，簡化金融機構申報資料流程，建立資訊共享機制；另運用 ML 及 DL 演算法，對金融機構資料申報內容進行正確性檢驗，以及建置申報資料內容合理性的知識圖譜（knowledge graphs），以降低申報錯誤之風險。

2. 澳洲證券暨投資委員會（ASIC）

建置市場分析與情報（market analysis and intelligence, MAI）系統，蒐集股票及相關衍生性金融商品市場之交易資料，即時監控初級及次級資本市場活動，配合 ML 提供交易與資金流向之即時監控警訊，以偵測市場值得深入調查之異常事件。

（二）資料管理

義大利央行（BdI）

建立軟體基礎架構，運用包含 Matlab、Python 等軟體建置處理各類大數據的環境基礎，發展整併不同資料來源之技術，將可疑交易報告之結構化資料與新聞評論之非結構化資料相互結合，並藉由 ML 演算法改善貸款違約預測。

（三）虛擬助理

1.菲律賓央行（BSP）

開發聊天機器人系統，自動答覆民眾投訴問題，將其加以分類，進一步分析民眾潛在關注的議題，以及發現金融機構可能違規的行為。

2.英國金融行為監管局（FCA）

使用 NLP 將法規書面文字轉換為機器可讀格式，可提高法規內容一致性及法令遵循水準，將有助於縮減監理項目及法令解釋間之落差。

（四）市場監督

1.ASIC、FCA 及美國證券交易委員會（SEC）

將巨量資料轉換為機器易於處理模式，以非監督學習技術偵測申報資料態樣及異常現象，有助於偵測內線交易與市場操縱，並判斷可能涉及舞弊案件或不當行為之型態、趨勢或語言。

2.法國金融管理局（AMF）

運用 ML 來檢測傳統演算法無法辨識的市場濫用或異常情況，以保護投資人。

（五）不當行為分析

1.墨西哥國家銀行暨證券管理委會（CNBV）

結合多項演算法，[偵測可疑交易有關的帳戶資料](#)，另實驗運用 NLP 技術標註在洗錢交易新聞上出現的姓名及公司名稱，與其社群媒體或可疑帳戶資料連結。

2.英國金融行為監管局（FCA）

運用監督學習技術，[預測理財專員不當銷售金融商品](#)可能性及其發生的地點。

3.新加坡金融管理局（MAS）

利用 ML 分析交易明細、客戶資料等資訊，[辨識值得關注之可疑交易](#)，藉以偵測潛在之複雜洗錢模式，以及無法由可疑交易紀錄中發現之資恐行為。

（六）個體審慎

1.義大利央行（BdI）

利用 ML 篩選不動產市場廣告資料，[預測未來房價走勢](#)；從 twitter 推文中萃取資訊，俾獲取[通膨預期具參考資訊](#)；以及評估客戶對特定基金公司、股票回報的影響、波動性與交易量等情緒分析。

2.荷蘭央行（DNB）

運用類神經網絡演算法，[偵測資金異常流動](#)，以瞭解銀行流動性潛在風險。

3.泰國央行（BOT）

利用 ML 演算法[對金融機構不良貸款建立早期預警指標系統](#)；另開發一套 NLP 平台，自動化檢視董事會會議紀錄，並量化及分析相關變數，例如每次會議討論議題數、討論時間及董事出席參與程度等，以瞭解公司績

效表現及組織文化特質等，以評估公司治理之運作情形。

4.美國聯準會（Fed）

透過 NLP 分析並彙整金融機構的大量書面資料，篩選主題資訊，從語意或文字中偵測人工易忽略之潛藏風險；亦可透過相關資料監測金融機構須關注特定議題之因應情形，例如氣候變遷、雲端運用、董事會功能有效性、消費者保護及網路風險等。

5.卡達金融中心管理局（QFCRA）

運用 NLP 偵測並分析金融機構或與其經營環境有關推文之隱藏意涵，以彌補透過金融機構定期申報資料掌控風險之時間落差，以即時掌握金融機構之風險變化。

6.巴西央行（BCB）

運用以 ML 技術為基礎的分析工具（ADAM），分析金融機構之潛在違約率較高借款戶及是否妥適衡量其預期損失。

（七）總體審慎

1.荷蘭央行（DNB）

利用巨量交易資料，以 ML 找出特定類型交易，例如銀行間無擔保拆借交易，並編製估計每日流動性風險指標，以協助辨識金融體系的潛在風險。

2.希臘央行（BoG）

同時運用 ML 及 DL 技術預測銀行償債能力，以建立預測銀行破產的早期預警指標。

3.馬來西亞央行（BNM）

分析過往類似檢查意見之案情及用語，提供監理人員撰擬檢查報告之文字建議，以確保監理標準一致性，並將檢查重點明確傳達予金融機構。

表 2 各國金融監理機關運用 AI 技術之案例

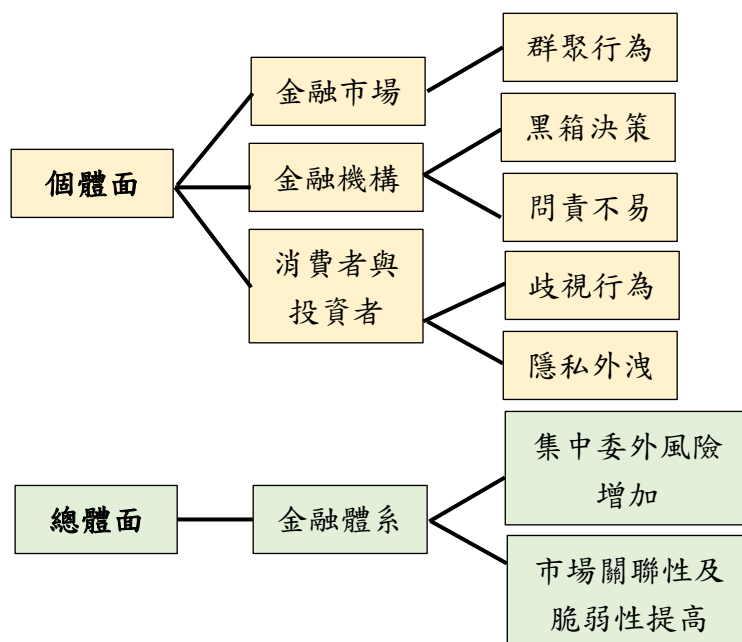
應用範疇	金融監理機關	ML	DL	NLP	語意 分析	聊天 機器人
資料申報	奧地利央行 (OeNB)	✓	✓			
	澳洲證券暨投資委員會 (ASIC)	✓				
資料管理	義大利央行 (BdI)	✓				
虛擬助理	菲律賓央行 (BSP)				✓	✓
	英國金融行為監管局 (FCA)			✓		✓
市場監督	澳洲證券暨投資委員會 (ASIC)	✓		✓	✓	
	英國金融行為監管局 (FCA)	✓		✓	✓	
	美國證券交易委員會 (SEC)	✓		✓	✓	
	法國金融管理局 (AMF)	✓				
不當行為 分析	墨西哥國家銀行暨證券管理委會 (CNBV)	✓		✓		
	英國金融行為監管局 (FCA)	✓				
	新加坡金融管理局 (MAS)	✓		✓		
個體審慎	義大利央行 (BdI)	✓		✓	✓	
	荷蘭央行 (DNB)		✓			
	泰國央行 (BOT)	✓		✓	✓	
	美國聯準會 (Fed)			✓	✓	
	卡達金融中心管理局 (QFCRA)			✓	✓	
	巴西央行 (BCB)	✓				
總體審慎	荷蘭央行 (DNB)	✓				
	希臘央行 (BoG)	✓	✓			
	馬來西亞央行 (BNM)			✓	✓	

資料來源：FSI、FSB；本文整理。

四、AI 於金融領域應用之潛在風險

當金融機構運用 ML 及相關演算法等 AI 技術進行數據分析，並藉由機器自我學習，進行決策或預測事件時，可能衍生相關潛在風險，FSB（2017）分別由金融個體面（包括金融市場、金融機構、消費者與投資者）及金融總體面，分析金融機構運用 AI 可能衍生之風險，主要包括群聚行為、黑箱決策、問責不易、歧視行為、隱私外洩、集中委外風險增加、市場關聯性及脆弱性提高（圖 6）。

圖 6 AI 於金融領域應用之潛在風險



資料來源：FSB（2017）；本文整理。

（一）群聚行為

當許多市場參與者在信用評分或交易活動等領域使用 AI 技術進行評估與交易時，一旦市場出現壓力，可能因群聚效應而擴大衝擊程度，[例如同時出現有價證券或資產減損，相關風險可能進一步影響金融穩定。](#)

（二）黑箱決策

AI 運算機制對多數人來說可能不易理解，在缺乏可判讀性或可稽查性

的情況下，增加向消費大眾解釋其決策之困難，恐造成監理一大挑戰，若未由金融監理機關適當介入監督，可能因風險傳染擴大而造成總體風險。

（三）問責不易

若 AI 決策行為造成金融機構損失，不易究責，尤其由第三方開發 AI 模型時，將因權責爭議而增加求償難度。

（四）歧視行為

運用大數據分析及 AI 技術時，可能會產生關於種族、宗教、性別或地區等特徵歧視問題，即使未蒐集到該類敏感特徵資料，亦可能產生與這些指標隱含意義有關聯性的結果，例如具某些特徵之消費者無法取得貸款或購買保險商品。

（五）隱私外洩

AI 可利用公眾數據對客戶進行分析，可能帶來個資外洩或隱私侵犯的問題，若僅考量保護消費者之匿名性，而未考慮如何保護分析結果，以安全並有效的使用大數據，將衍生個資保護問題。

（六）集中委外風險增加

AI 技術高度依賴少數第三方技術開發商或服務供應商，若金融機構集中委外同一家供應商，當此供應商遭遇破產或大型作業風險，將可能產生多家金融機構營運同時中斷之風險。

（七）市場關聯性及脆弱性提高

當眾多金融機構依賴相同資料來源及演算法進行金融交易決策時，將提高金融機構與金融市場間的關聯性，以及金融機構若同時執行高頻且大量交易，恐提高市場波動度及脆弱性，進而影響金融穩定。此外，若金融機構運用 AI 於資本管理以降低資本計提，將使金融機構資本緩衝減少，提高脆弱度，將導致流動性緩衝降低、槓桿率提高及期限轉換加速。

參、國際間對 AI 之監理進展

一、發布 AI 治理準則

AI 應用領域範圍遍及各產業，相關發展應用並不侷限於金融業，部分國際組織或國家為降低 AI 技術應用之潛在疑慮與風險，近年已陸續研擬相關基本準則（表 3），針對倫理道德、尊重人權、個資保護、資料治理、公眾信任、系統穩健及運算邏輯檢核等議題，建立最高原則並逐步推動應有之監理措施。

（一）國際組織

經濟合作暨發展組織(OECD)於 2019 年發布「AI 原則」(Principles on Artificial Intelligence)，提倡 AI 系統設計須具有創新性及可信賴性，尊重人權及民主價值，並列出 5 項基本原則²¹鼓勵政府部門提供建立可靠 AI 系統之政策環境，共有 42 個國家簽署，隨後亦被 G20 採納。

（二）歐盟地區

歐盟執委會(EC)成立 AI 高級專家小組(High-Level Expert Group, HLEG)分別於 2019 及 2020 年頒布「可靠 AI 之道德準則」(Ethics Guidelines for Trustworthy AI)²²與「人工智慧白皮書」(White Paper on Artificial Intelligence)²³，並強調各國倡議應避免危及白皮書之確定性。此外，EC 於 2021 年 4 月提出「AI 法律統一規則草案」(Proposal for a Regulation laying down harmonised rules on artificial intelligence)，為現

²¹ 包括：（1）藉由 AI 促進普惠性成長，減少社會不平等之歧視，實現永續發展及增進民生福祉；（2）AI 系統設計應尊重法律、人權、民主價值及多樣性，並涵蓋適當的保障措施，在必要時進行人為干預，以確保「以人為本」的價值觀及社會公平公正；（3）應致力實現 AI 系統透明度及責任揭露，確保與系統進行互動可能帶來的結果；（4）在 AI 系統生命週期內應持續評估及管理潛在風險；（5）開發、部署或運用 AI 系統之機構及個人應依據上述原則對 AI 系統的正常運作負責。

²² 可靠 AI 的具體要件包括：人類自主性與監控性、技術穩健性及安全性、隱私保護與資料管理、系統透明度與可追溯性、維持多樣性與公平性、兼顧社會福祉與環境保護，以及確保系統責任範圍等 7 個面向。

²³ 目標包括：（1）打造有助 AI 發展之「卓越生態圈」(ecosystem of excellence)，主要在成員國彼此協調、研究與創新產業、技能培訓、強化中小企業運用 AI 能力、投資數據管理之基礎設施與國際合作等方面，以及（2）以「可靠 AI 之道德準則」為基礎建立 AI「信任生態圈」(ecosystem of trust)。

行已發布 AI 統一管理規則之立法草案，對 AI 應用領域進行風險分類，設立罰款機制嚴格管制高風險領域²⁴。

（三）英國

英國資訊委員辦公室（Information Commissioner's Office, ICO）於 2020 年 7 月發布「[人工智慧與資料保護指引](#)」（Guidance on AI and data protection），提供設計、架構與實施 AI 系統的法令遵循路線圖。

（四）美國

美國白宮於 2020 年發布「[人工智慧應用規範指南](#)」（Guidance for Regulation of Artificial Intelligence Applications），為政府機關起草 AI 規範並進行管制時提供指引，並以 10 項管理原則為規範制定之依據，包括培養公眾信任、公眾參與、科學研究倫理與資訊品質、AI 風險評估與管理、效益與成本原則、彈性原則、公平與反歧視、AI 應用之揭露與透明化、AI 系統防護與措施安全性，以及機構間之相互協調。

（五）中國大陸

中國大陸於 2019 年發布「[新一代 AI 治理原則](#)」（Governance Principles for the New Generation AI），旨在平衡 AI 發展與治理，確保 AI 安全可控，推動經濟、社會及生態可永續發展。強調和諧友好、公平公正、包容（普惠）共享、尊重隱私、安全可控、共擔責任、開放協作、敏捷治理等 8 條原則。另計劃於 2025 年全面將 AI 納入法律制度，並於 2030 年成為 AI 創新的世界中心。

²⁴ 禁止違反人權的特定應用領域，例如讀取潛意識技術、可能會剝弱勢族群之應用、政府機關對個人行為採用社會評分制度及執法機構使用的遠端生物辨識系統。另高風險應用領域涵蓋（1）不限公務機關使用之遠端生物辨識；（2）作為道路、水電等基礎公共設施安全維護；（3）教育或職訓領域之挑選或評分功能；（4）職場對員工進行管理目的，包括雇用決定與績效評分；（5）決定使用特定服務優先順序，包括社會福利、[金融信用](#)、消防或醫療緊急服務；（6）法律執行輔助，包括測謊功能、評估被告或受刑人再犯機率；（7）用於移民、庇護、邊境管制目的；（8）提供司法機關法規調查使用等 8 項應用領域，技術提供者及使用者皆須嚴守相關管制規範，違者最高將處以 3,000 萬歐元罰鍰或全球年營收 6% 之罰款。

（六）新加坡

新加坡 6 個部會及機構²⁵共同推動新加坡全國人工智慧核心 AISG (AI Singapore) 計畫，主要為資助 AI 研究、應用 AI 技術、援助小型企業採用 AI 及培養 AI 人才。

（七）台灣

我國科技部於 2019 年發布「人工智慧科研發展指引」(AI Technology R&D Guidelines)，強調「以人為本」、「永續發展」及「多元包容(普惠)」3 大核心價值，從而延伸出 8 項方針，包括共榮共利、公平性與非歧視性、自主權與控制權、安全性、個人隱私與數據治理、透明性與可追溯性、可解釋性及問責與溝通等。

²⁵ 包括 National Research Foundation (NRF)、Smart Nation and Digital Government Office (SNDGO)、Economic Development Board (EDB)、Infocomm Media Development Authority (IMDA)、SGInnovate 及 Integrated Health Information Systems (IHiS)。

表 3 國際組織或各國對 AI 發布治理準則

國際組織或國家	準則或機關	內容重點
OECD	AI 原則	提倡 AI 系統設計須具有創新性及可信賴性，尊重人權及民主價值，並列出應致力實現 AI 系統之透明度及責任揭露、在 AI 系統生命週期內須應持續評估及管理潛在風險等 5 項基本原則。
歐盟執委會	可靠 AI 之道德準則	主要描述可靠 AI 系統應具人類自主性與監控性、技術穩健性及安全性、隱私保護與資料治理、系統透明度與可追溯性等 7 項具體要件。
	人工智慧白皮書	以打造協調合作之「卓越生態圈」及以「可靠 AI 之道德準則」為基礎，建立 AI「信任生態圈」為目標，提出歐盟 AI 通用標準。
	AI 應用規範草案	對 AI 應用領域進行風險分類，設立罰款機制嚴格管制高風險領域。
英國	人工智慧與資料保護指引	針對 AI 系統的治理與問責、合法性、公平性、透明性、安全性、個人權利提出法令遵循指引。
美國	人工智慧應用規範指南	為政府機關起草 AI 規範提供指引，擬定公平與反歧視、AI 應用之揭露與透明化、AI 系統防護與措施安全性，以及機構間之相互協調等 10 項管理原則作為政府制定規範之參考依據。
中國大陸	新一代 AI 治理原則	平衡 AI 發展與治理，強調和諧友好、公平公正、包容（普惠）共享、尊重隱私、安全可控、共擔責任、開放協作、敏捷治理等 8 條原則。
新加坡	AISG 計畫	主要為資助 AI 研究、應用 AI 技術、援助小型企業採用 AI 及培養 AI 人才。
台灣	人工智慧科研發展指引	提出共榮共利、公平性與非歧視性、自主權與控制權、安全性、個人隱私與數據治理、透明性與可追溯性、可解釋性及問責與溝通等 8 項方針。

資料來源：各官方網站；本文整理。

二、國際間對 AI 於金融領域應用之監理進展

目前國際間多數金融監理機關皆要求，金融機構應在法令遵循架構下全面性衡酌 AI 之使用，各國監理機關開始制定或調整有關 AI 治理之監理架構，然而，現行除歐盟下列規範、指令或準則已有對 AI 進行監理略有著墨外，其餘國家多數仍處於早期發展階段，且目前相關發布文件多屬於討論性質及原則性規範，對於 AI 在金融領域之應用少有專屬法令規範，分述如下：

（一）歐盟

1. 歐盟個人資料保護法令「一般資料保護規範」（GDPR）

歐盟境內營運之所有企業，若須蒐集、處理或利用歐盟人民的個人資料，皆適用 GDPR，主要保護資料當事人之受告知權、可攜權、被遺忘權、反對權及自動化數位剖析許可權²⁶等，亦即鑑於 AI 係以大數據為基礎，運用 ML 及演算法技術進行資料分析與研判，例如：若 Facebook 透過演算法提供客製化資訊內容予金融機構作為個人風險控管模型之評估與建構，則資料當事人可主張前述權利，適度保障個人資料之合理運用與隱私權利。

2. 歐盟金融工具市場指令修訂版（MiFID II）對演算法之交易規範

MiFID II 修法重點聚焦於確保市場公平、安全及效率，並提升透明度，以及強化對投資人權益之保護，其中對從事演算法與高頻交易的公司，以及允許該類型交易的交易場所，要求須強制註冊為投資公司並以較嚴格的監理要求規範，包括風險管理機制、交易系統有效性、執行程序的安排、交易紀錄的保存與通報義務等。

²⁶ 受告知權為資料控管者應於取得個人資料時，提供資料當事人應告知之資訊；可攜權為資料當事人應有權以結構化且機器可讀的形式，接收其提供予資料控管者之資料，並有權傳輸給其他資料控管者；被遺忘權為資料當事人在特定情況下（例如非法處理個人資料、資料當事人同意已經撤回等情況），可要求刪除資料控管者所掌控之個人資料；反對權為資料當事人有權反對資料之處理，除非資料控管者證明其處理有優先於資料當事人之權利與自由；自動化數位剖析許可權為資料主體應有權不受僅基於自動化處理（依既有個人資料進行分析與預測行為，例如：工作表現、經濟狀況、位置、健康狀況、個人偏好，可信賴度或者行為表現等之判定）所生成而對其產生法律效果或類似之重大影響的決策所拘束。

3.歐洲證券市場管理局（ESMA）對投資顧問及資產組合管理機構發布 MiFID II 有關適合性要求之相關準則²⁷

該準則要求使用**機器人理財服務**提供零售客戶投資建議之相關機構，必須**具備妥適系統與內控**，以確保商品適合性評估工具可以完整呈現客戶風險圖像等，以及須對**演算法作經常性測試**，以確保推薦商品及交易決策之妥適性。

（二）個別國家

1.英國

目前**英格蘭銀行（BoE）**對於強化**AI 監理**，仍僅止於規劃階段，該行 2020 年 1 月與**金融行為監管局（FCA）**公布成立「**金融服務人工智慧公私論壇**」（Financial Services AI Public Private Forum），用以協助發展金融業安全採用 AI 或 ML 工具、準則、規範或產業最佳實務，論壇主軸分別為數據運用、風險管理模型及治理，以促進公私部門交流，更加瞭解 AI 於金融服務之運用及影響。另 BoE 及 FCA 於 2019 年 10 月共同發表之「**英國金融服務之機器學習**」（Machine Learning in UK Financial Services），再次強調 AI 發展可從根本上改變企業提供及消費者使用金融服務的方式。

2.法國

法國**金融審慎監理局（ACPR）**於 2020 年 6 月發布「**金融 AI 治理**」（Governance of AI in Finance）。

3.德國

2018 年**聯邦金融監管局（BaFin）**發表「**當大數據遇到 AI**」（Big Data Meets Artificial Intelligence），開始討論 AI 對產品與流程創新以及新進入者出現與市場結構所產生的影響，2021 年針對 AI 運用之演算法進一步發布

²⁷ ESMA (2018), “Guidelines on certain aspects of the MiFID II suitability requirements,” Nov., 網址：
https://www.esma.europa.eu/sites/default/files/library/esma35-43-1163_guidelines_on_certain_aspects_of_mifid_ii_suitability_requirements_0.pdf

「大數據與 AI：決策過程應用演算法之原則」（Big Data and Artificial Intelligence：Principles for the Use of Algorithms in Decision-Making Processes）。

4. 盧森堡

盧森堡金融部門監理委員會（CSSF）於 2018 年 12 月發布「AI：金融部門之機會、風險及相關建議」（AI：Opportunities, Risks and Recommendations for the Financial Sector），初步討論 AI 於金融部門之應用。

5. 荷蘭

荷蘭央行（DNB）於 2019 年 7 月發布「金融部門應用 AI 之一般原則」（General Principles for the Use of AI in the Financial Sector）。

6. 香港

香港金融管理局（HKMA）於 2019 年 11 月發布「大數據分析與 AI 應用下之消費者保護」（Consumer Protection in Respect of Use of Big Data Analytics and Artificial Intelligence）。

7. 新加坡

2019 年 11 月金融管理局（MAS）公布已與 16 家金融服務業者合作發展使用 AI 應具有之道德責任標準及落實架構，以促進金融部門使用 AI 之公平性、道德、責任及透明度（FEAT）²⁸，2021 年完成新加坡國家 AI 戰略 Veritas 計畫之第一階段，將 FEAT 原則納入人工智慧及數據分析（Artificial Intelligence and Data Analytics, AIDA）方案²⁹。

²⁸

<https://www.mas.gov.sg/news/media-releases/2019/mas-partners-financial-industry-to-create-framework-for-responsible-use-of-ai>。

²⁹ <https://www.mas.gov.sg/news/media-releases/2021/veritas-initiative-addresses-implementation-challenges>。

三、國際間對 AI 運用之治理原則與面臨之挑戰

國際間常被引用於金融業之主要建議及原則，為 OECD 及 G20 共同發布之「AI 原則」，其餘金融監理機關如何對金融機構運用 AI 進行監理，仍處於發展治理原則或指引之早期階段，就現行已發布有關 AI 治理文件所研擬之原則或指引，可歸納出四項共通治理原則，包括可靠性/穩健性（reliability/soundness）、可問責性（accountability）、透明度（transparency）、公平性（fairness）/道德性（ethics）等，惟如何落實執行仍面臨一些挑戰（表 4）。

表 4 AI 運用之共通治理原則及面臨之挑戰

治理原則	特 性	挑 戰
可靠性/ 穩健性	AI 模型應具備可驗證性及正確性，以確保資料品質。	AI 演算法可能發生偏誤或不準確的結果，須建立定期檢視且即時更新之審核機制，並於演算法簡化及效能間取得平衡。
可問責性	外界能理解 AI 決策過程，並提供對 AI 決策結果質疑及挑戰之管道。	監理機關通常將金融機構任何決策之最終責任歸屬董事會及高級管理階層，惟是否適用於 AI 尚未臻明確；另外包 AI 服務予第三方業者，亦對問責機制形成挑戰 ^註 。
透明度	AI 決策過程應具可解釋性及可稽核性，並對外揭露相關決策依據。	部分 AI 模型應用之演算法與傳統模型不同，難以理解其決策過程，透明度不足會削弱消費者信心。
公平性/ 道德性	1. 決策基礎應基於資料作出客觀之分析及判斷，避免歧視行為。 2. 確保金融消費者權益不會受損，例如決策偏誤、歧視或不當取得資訊等。	「公平性」及「道德性」欠缺普遍認可的定義，以及為求符合此一原則，需要在人為介入/監控 AI 決策與獲取科技自動化完整效益間取得平衡。

註：例如 AI 演算法系統性對居住於特定地區者否決貸款申請，究竟由 AI 開發商或銀行有權貸放主管或總經理負責，其責任歸屬並不明確。

資料來源：BIS（2021）。

肆、我國業者運用 AI 於金融領域之現況與監理機關對 AI 應用之監理

一、我國金融機構運用 AI 現況

我國金融機構運用 AI 技術於業務或作業範疇日趨多元，目前本國銀行導入 AI 技術應用之實例（詳附錄），多數以提升服務便利性、增進金融資訊效率性及準確滿足客戶投資需求為目的，陸續推出客戶服務、金融顧問及機器人理財等金融應用服務，部分則應用於作業流程優化及不當行為偵測作業，以提升營運效率及強化法令遵循（圖 7）。

圖 7 本國銀行運用 AI 技術之領域及目的

客戶服務	• 提升服務便利性
金融顧問	• 增進金融資訊效率性
機器人理財	• 滿足客戶投資需求
作業流程優化	• 提升作業效率
不當行為偵測	• 強化法令遵循

資料來源：各銀行官方網站；本文整理。

本國銀行運用 AI 最廣泛的領域為客戶服務，以電話、網路等媒介的 AI 客服系統可採一問一答方式回應常見且重複性的客戶諮詢，降低客戶等待時間與人事成本。金融顧問則運用 ML 技術、認知和語意分析技術，提供使用者金融市場的相關資訊與服務，例如：房貸評估、外匯諮詢與信用卡推薦等。機器人理財主要聚焦於機器人演算速度優勢以及不受情緒、性格影響的特徵，代替傳統經理人提供財富管理建議，進行投資配置更新。

前述所列應用範疇大多為本國銀行近三年內開辦之案例，尚須觀察其客戶接受度、維護應用案例之成本效益及穩定性等，以期未來持續運用 AI 提升金融服務的品質與效能。

二、我國主管機關對金融業創新營運模式及 AI 應用之監管態度

（一）金融監理沙盒

金管會為推動金融創新，已於 106 年 12 月 29 日通過「金融科技發展與創新實驗條例」（金融監理沙盒），鼓勵不限金融業者以「創新實驗」方式申辦金融業務，另於 108 年起針對業務創新項目未涉及法令禁止者，陸續訂定銀行、證券及保險業申請業務試辦作業要點，申請試辦業務並得與其他同意進入金融監理沙盒之實驗項目相同，以加速金融創新。

目前已有阿爾發證券投資顧問公司及永豐金證券公司聯名申請落實普惠金融之機器人理財服務³⁰進入沙盒實驗，由阿爾發投顧提供投資人投資組合建議，投資人依該建議向永豐金證券以「定期定額」之方式買進國外 ETF，自動化程度高，係目前唯一運用 AI 技術申請創新實驗之案例。

（二）金管會對金融機構應用 AI 之監理態度

金管會認為目前金融機構多數僅將 AI 技術作為決策或作業之輔助工具，並非全無或極少人工服務，或已取代最終人為判斷，故尚無需以特定法令對業者進行規範。目前金管會僅對證券投資顧問事業涉及以 AI 技術等自動化工具提供證券投資顧問服務作業，發布相關作業要點。

金管會 106 年公布證券投資信託暨商業同業公會訂定之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（簡稱 Robo-Advisor 作業要點）規定，主要內容為：

- （1）定義自動化投資顧問服務：自動化投資顧問服務係指完全經由網路互動，全無或極少人工服務，而提供客戶投資組合建議之顧問服務。
- （2）建立演算法監管機制：針對自動化投資顧問服務系統核心之演算法設立期初審核及定期審核監管機制。

³⁰ 金管會銀行局 110 年 2 月 26 日金管科字第 11001305951 號函核准。

- (3) 落實 **KYC** 作業：規範提供投資組合建議前，應建立客戶資料進行瞭解客戶作業，設計各項評估指標。
- (4) 須**公平客觀**提供投資**建議**：應忠實履行客戶利益優先、利益衝突避免、禁止不當得利與公平處理等原則，避免可能受產品提供商之報酬影響，產生與客戶利益衝突之情事。
- (5) **充分揭露**投資組合再平衡之資訊：應充分揭露與告知客戶投資組合重新調整之必要，並建立市場重大變動之處置程序。
- (6) 應設立**專責監督委員會**：專門負責事業內部客戶問卷設計內容、演算法之開發與調整、客戶投資組合建議符合其風險屬性、自動化投資顧問服務系統公平客觀之執行及投資組合再平衡等之監督管理。
- (7) **告知**客戶注意事項：應明確告知客戶於使用自動化投資顧問服務前之注意事項，以確保客戶自身權益等。

依據金管會證期局統計，自動化管理服務自 106 年 6 月開放發展以來規模僅達 44 億元，金管會在保障投資人權益之前提下，以及兼顧符合法制架構及實務需求下，宣布放寬投顧事業從事自動化投顧服務有關自動再平衡交易的規範，業者事先與客戶於契約中約定在達到執行門檻且符合再平衡交易的約定條件情況時，可由電腦系統自動為客戶執行再平衡交易。新的「再平衡交易」約定條件有二，一是維持原約定的投資標的及投資比例為條件外，二是新增如業者與客戶約定之投資標的為經金管會核准的境內外基金，可在不超過與客戶事先約定的可投資基金名單（上限為 30 檔標的）及一定變動程度（各投資標的之投資比例變動絕對值合計數未超過 60%）範圍內自動執行，由業者與客戶就兩種再平衡條件擇一進行約定，期透過放寬符合兩大約定條件即可自動再平衡，以供更完整的服務吸引客戶，進而帶動整體業務規模成長。

伍、結語

近年來隨著 ML、DL 及演算法分析等技術蓬勃發展，帶動 AI 相關產業鏈，金融機構藉由運用 AI 技術降低人力成本、提升營運效率、偵測風險及發展創新業務型態，常見於客群經營、流程精進、風險合規及預測分析等應用範疇；國際間金融監理機關將原監理制度導入 AI 創新技術，增進監理能力與效率，並已嘗試運用 AI 技術進行金融監理，應用層面涵蓋資料申報、資料管理、虛擬助理、市場監督、不當行為分析、個體審慎、總體審慎等面向，促進 AI 於金融監理之應用，以提升監理效能、確保金融體系之穩定性及消費者權益。

運用 AI 技術於金融領域可能衍生之相關潛在風險，亦值得金融監理機關密切關注，當金融機構運用 AI 技術進行數據分析、決策或預測事件時，可能衍生不同層面的潛在風險，例如金融市場參與者之群聚行為效應、金融機構須釐清之黑箱決策風險及權責爭議、消費者與投資者面臨的歧視行為及隱私外洩，以及金融總體面所遭遇的集中委外風險增加、市場關聯性提升及脆弱性提高。

目前國際間已有研擬 AI 應用相關治理準則之趨勢，大致上皆為對倫理道德、尊重人權及個資保護等議題進行探討，較完善且常被引用的基本原則為 OECD 及 G20 共同發布之「AI 原則」，其餘國家多數仍處於早期發展階段，且目前相關發布文件多屬於討論性質及原則性規範，大致可歸納出可靠性/穩健性、可問責性、透明度及公平性/道德性等四項共通治理原則，對於 AI 在金融領域之應用則少有專屬法令規範。聚焦台灣亦有類似情況，我國僅於 2019 年發布「人工智慧科研發展指引」之原則性指引。

總歸來說，AI 與人類互有優劣，若涉及快速處理大量數據，AI 表現優於人類；惟若涉及監理機關或管理人員所面對的道德及風險問題，則人類表現得更佳。AI 技術目前於金融領域之運用，不論是服務或監理方面，仍處早期發展階段，雖然 AI 工具可蒐集及監督消費者與金融機構的大量資

訊，惟 AI 系統仍無法完全瞭解在總體經濟情勢下個別金融機構經營狀況，因此僅能作為金融監管機關決策或作業之輔助工具，不能完全取代監管機關。與人類決策者相比，AI 有助於校準現有的政策工具，但仍無法超越既有措施的組合及提出創新的政策工具。

鑑於國際間對 AI 應用於金融領域之監理仍處於初步發展階段，認為 AI 發展仍著重於建立結合人類思維及 AI 之優勢的夥伴關係，並強調 AI 系統「旨在增強而非取代人類貢獻」，仍須大量仰賴人類進行維護與開發，無法完全替代人類完成創造、情感交流及解決高複雜度的問題。在 AI 與人類的有效合作夥伴關係中，管理者、監理機關及政策制定者的自主權不得受到威脅，並且應能否決演算法之評估。可預期的是，未來會是人類與 AI 搭配的時代，先由人類「大膽假設」，再運用 AI 技術「小心求證」，透過 AI 及人類的結合完成重要的決策或任務，兩者的互動才是未來發展之關鍵。

面對多元化金融商業模式及數位化之浪潮下，對 AI 技術之應用與監理應保持審慎開放之思維，未來發展方向應朝向打破舊有思維、建立民眾信任之數據文化、設立共享機制，以及進行跨部會研議或產學合作，保持科技中性（tech-neutrality）之監理原則，提升創新科技研發意願，逐步建構對金融業運用 AI 技術之監理。

附錄

本國銀行導入 AI 應用之案例詳述

銀行別	應用範疇	項目	應用內容
北富銀	機器人理財	AI 風控	在財富管理的風險控管層面，過往銀行多藉由訪談、調查輔以經驗來分析客戶風險狀況，易受人為主觀判斷導致客戶風險偏好及可承受風險損失能力之評估產生落差。現今運用 ML 技術建構風險屬性評估模型（Customer Risk Attribute Assessment Model, CRAAM），大幅提升分析效率，提供更客觀準確的評估結果， 精進產品與客戶適合度 ，為客戶創造良好的投資體驗。
國泰銀	機器人理財	智能投資	依消費者的風險屬性及投資需求，建立不同的投資組合，並藉由機器 理智執行交易策略 ，亦可 每日監控 投資組合，以分散風險。
華南銀	客戶服務	AI 行動銀行	推出結合語音辨識及分析技術的 AI 行動銀行，同時導入 RPA、OCR 視覺辨識技術、AI 資訊分類及分流技術，透過 ML 提升數據處理精準度，開發創新流程與客戶體驗， 以提升客服效率 。
	作業流程優化	Just Smart 企業賦能平台	配合數位化發展，聚焦外匯、貸款、收單及徵信等業務行為，以 E 化電子簽章 減少紙本文件 ，大幅提升作業效率。
兆豐銀	客戶服務	櫃員 AI 助理	藉由優化數位金融及使用者體驗，透過平板電腦建立互動媒介，提供 即時翻譯互動 （英、

銀行別	應用範疇	項目	應用內容
		系統	日、韓、泰文語音即時轉為文字內容）、 多元繳款 （整合行動支付服務）及 表單確認 （客戶採口說方式表達業務需求即可產製相關申請書）三大功能。
王道銀	機器人理財	好也人生	針對國人偏好定期配息的特性，推出全新 配息機器人理財 ，分為追求穩健配息的一般型配息機器人，以及追求積極配息、創造較高現金流的加強型配息機器人兩種模組，提供用戶每月穩定配息收入。
永豐銀	金融顧問	投資報告	以 AI 技術大量閱覽投資報告與市場新聞，進行 自動化編撰 信評或投顧報告，提升風險管理準確度，並縮短作業時間。
	作業流程優化	客戶業務與經濟關聯圖	建立企業與利害關係人的業務與金融關聯圖，透過交叉檢核，驗證法人金融客戶提供的貸款文件資料基礎， 降低洗錢風險 ，並 提升授信管理效率 。
		永豐元智慧服務 MI	透過銀行內部商品、行銷、授信、審查、分析及資訊等跨功能協作，應用「內外部事件」、「業務規則與 AI 模型」與「產品與服務」彼此連動， 打造高效率運營模式 。
	機器人理財	智能理財	以高度客製化及理性分析之面向，主要鎖定 20 至 40 歲青年族群，以簡潔直覺的操作介面及每月新台幣 3,000 元的低門檻，協助缺乏投資經驗的客群進行 智能理財 ，實現普惠

銀行別	應用範疇	項目	應用內容
			金融。
	客戶服務	掌靜脈 ATM	於臨櫃及 ATM 導入 掌靜脈辨識 技術以取代金融卡，提升臨櫃服務效率及交易安全性，且透過手掌即可於 ATM 辦理存款、提款、轉帳或查詢餘額等，提供便捷、安全且節能的數位服務。
星展銀	不當行為偵測	NICE Actimize	針對客戶帳戶及交易的資料，依 交易的可疑程度進行排序並據以評分 ，其評分高於警戒值之資料須列入分析，並透過 ML 進一步降低誤報及漏報之可能性。
玉山銀	不當行為偵測	信用卡冒用偵測服務	於 2019 年 8 月上線，平均每月處理超過 2 千萬筆信用卡交易，並 減少 因信用卡 冒用 造成平均每月新台幣 750 萬元的 損失 。
	作業流程優化	全自動勞工紓困貸款申請	為提高勞工紓困貸款承作效率，透過建置數位平臺及 RPA 技術，以自動化審核模型 判斷客戶是否符合資格 ，且每日核准案件均以 自動撥款 入帳。
	不當行為偵測	警示帳務偵測系統	以往每日需人工檢查多達 1,400 位可疑名單，目前透過 ML 模型自動辨識，每日僅需檢查約 30 筆名單，且能成功找出比傳統方法更多的警示帳戶，大幅 節省人工作業時間 。
凱基銀	作業流程優化	三大電商平台貸款專案	串聯相關平台數據作為輔助財力證明 ，取代傳統所需之薪資證明，以解決電商平台賣家

銀行別	應用範疇	項目	應用內容
			缺乏銀行往來歷史交易紀錄之難題。
中信銀	客戶服務	對話言談分析	提升「客服語音數據言談分析」能力，藉由 NLP 技術分析客戶電話內容， 預測客戶實際需求 ，提升客服品質，強化 KYC。
		微表情識別、 情緒識別	建立微表情與情緒資料庫，分析人臉微表情與情緒，發展可視化的辨識分析系統，並進一步開發計算微表情與人物性格關聯性的演算法， 發展情緒識別模型 。
	金融顧問	新聞篇章分析	由新聞內容進階 判斷 客戶或交易對手之 事件關聯性 ，並生成新聞內容摘要。

資料來源：各銀行官方網站；工商時報；本文整理。

參考文獻

- Stanford University (2022) , *Artificial Intelligence Index Report*.
- CB Insights (2022) , “AI 100: The most promising artificial intelligence startups of 2022,” Apr.
- Mona A., Rohit M., Anton J., and Uthayasankar S. (2022) , “Ethical framework for Artificial Intelligence and Digital technologies,” *International Journal of Information Management Volume 62*, Feb.
- Marcus Buckmann, Andy Haldane and Anne-Caroline Hüser (2021) , “Comparing Minds and Machines: Implications for Financial Stability,” *Staff Working Paper No.937*, Aug.
- FSI (2021) , “Suptech Tools for Prudential Supervision and Their Use During the Pandemic,” Dec.
- BoE (2021) , “Comparing Minds and Machines: Implications for Financial Stability,” *Staff Working Paper No.937*, Aug.
- BIS (2021) , “Humans keeping AI in check — emerging regulatory expectations in the financial sector ,” Aug. °
- Sasha Andrieiev (2021) , “Artificial Intelligence Applications in Financial Services,” Apr.
- Sebastian Doerr, Leonardo Gambacorta and Jose Maria Serena (2021) , “Big data and machine learning in central banking,” *BIS Working Papers No.930*, Mar.
- Gérard Hertig (2021) , “Using artificial intelligence for financial supervision purposes,” *Future Resilient Systems*, Feb. 1.

FSB(2020), “The use of Supervisory and Regulatory Technology by Authorities and Regulated Institutions - Market Developments and Financial Stability Implications,” Oct. 9.

FSI(2020), “FSI Insights on policy implementation No 23 - Policy responses to fintech: a cross-country overview,” Jan.

FSI (2019) , “FSI Insights on policy implementation No 19 - the supotech generations,” Oct.

Duan, Yanqing, John S. Edwards, and Yogesh K. Dwivedi (2019) “Artificial intelligence for decision making in the era of Big Data—evolution, challenges and research agenda,” *International Journal of Information Management*, Vol. 48, pp. 63–71.

Russell, Stuart (2019) , *Human compatible: Artificial intelligence and the problem of control*.

FSI (2018) , “FSI Insights on policy implementation No 9 - Innovative technology in financial supervision (supotech) - the experience of early users,” July.

ESMA (2018) , “Guidelines on certain aspects of the MiFID II suitability requirements,” Nov.

Jarrahi, Mohammad Hossein (2018) “Artificial intelligence and the future of work: Human-AI symbiosis in organizational decision making,” *Business Horizons*, Vol. 61, No. 4, pp. 577–586.

Miller, Steven M. (2018) “AI: Augmentation, more so than automation,” *Asian Management Insights*, Vol. 5, No. 1, pp. 1–20.

Verghese, Abraham, Nigam H. Shah, and Robert A. Harrington (2018) “What

this computer needs is a physician: humanism and artificial intelligence,”
JAMA, Vol. 319, No. 1, pp. 19–20.

FSB(2017), “Artificial Intelligence and Machine Learning in Financial Services
- Market Developments and Financial Stability Implications,” Nov. 1.

Lake, Brenden M., Tomer D. Ullman, Joshua B. Tenenbaum, and Samuel J.
Gershman (2017) , “Building machines that learn and think like people,”
Behavioral and Brain Sciences, Vol. 40.

謝明華（2020），「AI 在金融產業發展的應用及潛在衝擊」，外匯市場發
展基金會委託研究計畫，1 月 31 日。

李震華（2019），「人工智慧技術發展趨勢與金融應用展望」，財金資訊
季刊，No.95，7 月。

李開復、王詠剛（2017），*人工智慧來了*。