



星球永續健康線上直播

量子健康科學 (5)

人工智慧推進量子運算

2026 年 8 月 26 日

人工智慧與量子運算近年快速發展，兩者相互輔助強化。如果把量子電腦比喻成一位具有超強計算能力的天才，那麼這位天才最大的問題，就是非常容易受到外界干擾而出錯。此時人工智慧所扮演的角色，就像是在旁邊協助找錯、除錯與校正的助手。本週我們將探討人工智慧輔助量子運算推展以及量子運算精準敗血症照護應用情境。

健康科學新知

極端熱浪與乾旱席捲歐洲：「赤地千里」

2026 年夏季，歐洲在持續熱浪與長期少雨下，同時遭遇野火、農業減產與供水壓力，顯示極端氣候的影響已從自然災害延伸至糧食安全、水資源管理與公共治理。比利時東部高沼地自然保護區爆發該國歷來最大規模之一的野火，約 30 平方公里土地遭焚毀，約 500 名消防人員投入救援，並獲鄰近歐洲國家航空支援。當地泥炭地尤其危險，火勢可能深入地下持續悶燒數週甚至數月，因此即使地表火勢減弱，仍可能重新燃起。希臘薩拉米斯島的野火則造成 2 人死亡、9 人受傷及多棟住宅被毀；法國西南部朗德地區也有約 17 平方公里土地焚毀，600 多名居民一度撤離，政府其後提供 1,200 萬歐元援助。農業方面法國、義大利及西班牙均出現大幅減產。法國櫛瓜、萵苣、豌豆及青花菜等蔬菜產量下降約 25%至 60%，部分地區朝鮮薊甚至接近全數損失。玉米產量估計比 2025 年減少 35%，降至約 900 萬噸；乳牛因熱緊迫使牛奶產量減少約 10%至 15%，乾草產量亦較長期平均低約 30%。義大利約 60%的農地受到乾旱影響，農業氣候災害損失估計超過 30 億歐元，波河谷水量下降尤其威脅稻米等重要作物。西班牙多地穀物收成亦下降兩至五成左右，葡萄酒產量預估減少約 20%。法國政府因此推出 1.45 億歐元緊急農業援助。英格蘭與威爾斯的問題則集中在水資源安全。目前英格蘭約四分之三地區及整個威爾斯處於乾旱狀態，五座重要水庫水位降至異常偏低。水公司正在制定 2027 至



2032 年的乾旱計畫，最嚴重情況可能採取公共取水點與輪流停水措施。然而，另有涉及重大緊急情境的乾旱計畫並不公開，政府以國家安全及關鍵基礎設施為理由支持保密。保密政策引發爭議。環保團體及部分政治人物認為，在氣候變遷使極端乾旱愈來愈頻繁的情況下，政府應更透明地說明缺水時的用水優先順序。尤其資料中心等高度耗水產業持續擴張，社會開始質疑在供水不足時，住宅、農業、能源與大型科技產業之間究竟如何分配有限水源。

歐洲野火危機：「星火燎原」

2026 年夏季，法國與西班牙遭遇極端野火，顯示歐洲火災風險可能正進入新階段。受熱浪與乾旱影響，截至 7 月底，法國約 9.1 萬公頃、西班牙超過 20 萬公頃土地遭焚，數十萬人被迫撤離。科學家認為，這些事件不只是個別年份的異常，而可能反映南歐「火災型態」(fire regime) 的結構性轉變。最新研究顯示，法國與西班牙部分地區在 2016 至 2025 年間所經歷的炎熱、乾燥、容易促成火災的天氣，其強度比 1981 至 2010 年高出最多約 50%。西班牙與葡萄牙部分地區過去一年不到 10 天會出現高度有利於火災的氣象條件，如今可能增加到一年 25 天。研究人員因而指出，南歐的火災環境正逐漸接近美國加州等長期受到大型野火威脅的地區。2026 年歐洲甚至出現過去罕見的現象，包括馬德里附近的大火幾乎威脅 NASA 的深空通訊設施，以及法國首次觀測到由極端野火產生、能自行形成不穩定天氣的「火積雲」。最重要的長期驅動因素是人為造成的氣候變遷。升高的氣溫使熱浪更頻繁、更強烈，而熱浪通常伴隨極端乾燥的火災氣象條件。2026 年歐洲自 5 月起接連遭遇多次熱浪，使植被失去水分、變得極易燃燒，形成大量可供火勢快速蔓延的燃料。類似的「極端高溫之後出現破紀錄野火」模式在 2025 年的西班牙已經發生，而既有氣候模型也預測，本世紀持續升溫將大幅提高歐洲的野火風險。因此，研究者認為這並非短暫波動，而是一項很難逆轉的長期趨勢。此外，火災風險並非只由氣候決定，土地利用的改變也正在放大問題。西班牙部分地區逐漸放棄農耕與放牧後，大量灌木與其他植被累積；另一方面，長期有效撲滅小型火災，也使原本可能被週期性燃燒消耗的植被持續堆積。雖然歐洲過去透過強化消防能力降低火災數量，



但極端氣候可能逐漸超越現有防火與滅火能力，使未來野火呈現數量減少、規模更大、破壞力更強趨勢，成為歐洲日益嚴重的氣候與安全威脅。

胚胎篩選技術爭議：「越俎代庖」

矽谷正推動生殖科技進入新階段，Nucleus、Orchid 等公司開始提供多基因胚胎篩檢(PGT-P)，除疾病風險外，也嘗試預測身高、智力與壽命等特徵，為接受體外受精(IVF)的夫妻提供更進階的胚胎基因篩檢，不只評估乳癌、糖尿病等疾病風險，也嘗試預測眼睛與頭髮顏色。這類多基因胚胎篩檢(PGT-P)目前仍屬高價、小眾服務，但隨著 IVF 成本下降及生育科技投資增加，未來可能逐漸進入一般生殖醫療市場。與檢測單一基因疾病不同，多基因評分只能提供機率，而不是確定答案。現階段技術仍有不少限制，例如基因資料庫主要來自歐洲血統人口，因此對不同族群的預測準確度並不一致；此外，父母實際上只能在自己有限數量的胚胎中進行選擇，因此可獲得的效果也相當有限。2019 年的一項研究估計，如果從十個胚胎中挑選，依身高選擇的中位增益約為 2.9 公分，依智力選擇則約增加 3 個 IQ 點，而當可選胚胎數更少時，效果還會進一步下降。另一方面，一個基因往往同時影響多種特質，因此降低某種疾病風險，也可能伴隨其他目前尚未完全理解的影響。這些不確定性更適合作為加強資訊揭露與監管的理由，而不是全面禁止胚胎篩檢。相關公司應清楚說明預測的限制與誤差，並區分絕對風險與相對風險，避免消費者高估技術效果，同時也應接受獨立科學審查。未來隨著基因資料持續增加、不同族群樣本更加完整，再加上人工智慧提升基因與複雜特質之間的分析能力，這類預測仍可能持續改善。爭議問題則在於倫理與社會公平。反對者擔心，若父母開始依智力、外貌、身高或其他特質挑選胚胎，可能形成新的優生學，並讓富裕家庭取得額外的世代優勢，進一步擴大社會不平等。不過，《經濟學人》主張，父母自願選擇胚胎，與 20 世紀由國家強制推行、剝奪特定群體生育權的優生政策並不相同；而科技目前價格昂貴，也不必然代表應禁止使用，政策上反而可以思考如何降低成本並擴大可近性。

AI 資本浪潮循環融資憂慮：「推波助瀾」

AI 產業競爭正從模型與晶片性能，轉向資料中心、電力、融資與長期資本的重資本



競賽。2026 年多項交易顯示，OpenAI、Nvidia 及大型雲端業者透過投資、採購、租賃與信用支持形成複雜的資金網絡，也引發「循環融資」與隱性槓桿風險。AI 基礎建設相關租賃與採購承諾已達兆美元規模，但企業 AI 營收仍快速成長，顯示市場需求確實存在。未來真正的考驗在於，AI 收入能否持續高速成長以支撐龐大的資本投入；若需求放緩或融資環境收緊，長期合約與信用承諾可能轉化為企業沉重的財務壓力。因此，AI 競爭已不只是技術競賽，更是對資本效率、融資能力與風險承受度的考驗。

AI 產業矽谷慈善浪潮：「達濟天下」 & AI 慈善投入挑戰：「行善有方」

AI 產業快速累積的財富，可能催生美國新一波超大型慈善浪潮。《經濟學人》估計，Anthropic、OpenAI 等 AI 公司創辦人、員工及早期投資人目前承諾投入公益的資金，合計可能高達約 4,300 億美元，規模約等同 25 個福特基金會。Anthropic 七名共同創辦人承諾捐出約 80% 個人財富，潛在捐款估計約 1,100 億美元；OpenAI Foundation 因持有大量公司股份，未來可支配的慈善資產甚至可能達 2,600 億美元。此外，AI 公司員工透過 donor-advised funds (DAFs，捐贈者建議基金) 累積的慈善資產也可能達數百億美元。當代科技慈善家與過去富豪捐款的重要差異，在於許多人受到**有效利他主義 (Effective Altruism, EA) **影響，強調以數據、成本效益與可衡量成果決定資源配置。因此，全球公共衛生、科學研究、瘧疾防治、動物福利、住房改革及 AI 安全等領域，都可能成為主要受益者。例如 Anthropic、OpenAI Foundation 與 Stripe 已共同參與一項約 5 億美元的疫苗研究捐款，另一方面也有大量資金流向 AI alignment 與降低高風險 AI 威脅的研究。巨額資金並不必然等於巨大的社會效益。首先，部分 AI 安全與長期風險研究本身很難量化，例如「成功避免 AI 災難」缺乏可直接觀察的反事實，因此難以證明數十億美元的投入究竟產生多少效益，與 EA 強調證據與量化評估的原則形成張力。其次，當 AI 公司或其基金會大量資助與自身技術相關的研究與公益計畫時，也可能引發利益衝突疑慮，包括藉慈善改善產業形象、擴大產品使用，甚至影響公共政策與學術研究方向。另一個核心問題是承諾捐款不等於立即投入公益。許多矽谷富豪將資產先放入 DAFs，在取得稅務優惠後，再逐年決定實際受贈機構。目前美



國約有 3,260 億美元存放於這類慈善帳戶，因此資金可能長期停留而未真正流入公益組織。此外，即使資金快速釋出，受贈機構也可能因規模突然擴張而面臨人才不足、薪資上升、管理能力及組織效率下降等問題。AI 世代的慈善家普遍呈現較強的時間急迫感。部分人士相信 AI 可能在短期內劇烈改變疾病、貧窮、科學與社會結構，因此傾向提早且快速投入資源。例如 Coefficient Giving 在 2026 年上半年已捐出約 10 億美元，並顯著增加對 GiveWell 的支持。AI 財富的影響可能改變公益資源的配置邏輯：慈善決策可能變得更加集中、數據導向，並將更多資金投入全球健康、科學與人工智慧風險等領域。真正的關鍵，將是這些數千億美元的承諾能否轉化為及時、可驗證且具有長期公共價值的實際成果。

全球智慧安全重鎮成形：「防微杜漸」

倫敦正成為全球 AI 企業與安全研究重鎮，國王十字一帶已匯聚 DeepMind，Anthropic 與 OpenAI 也設立大型辦公室。英國政府支持成立 AI 安全研究所 (AISI)，投入約 6,600 萬英鎊並配置近百名技術人員，負責測試前沿模型的高風險能力、協助改善防護機制，並串聯國際安全機構與研究團隊。英國希望發展 AI 安全驗證與認證產業，在避免過度監管限制創新的同時，建立可信賴的風險監測體系，兼顧科技競爭力與公共安全。

自旋量子位元躍進：「積微成著」

自旋量子位元近期取得多項重要突破，使其重新成為量子運算領域的重要候選技術。這類量子位元利用電子自旋儲存與操控量子資訊，最大優勢之一是可結合既有矽半導體製程，理論上有利於未來的大規模製造。不過，過去受到控制困難、雜訊高及可擴充規模有限等問題影響，其發展長期落後於超導量子位元與中性原子系統。近期最突出的成果包括 HRL Laboratories 建構的 18 個矽自旋量子位元處理器，單量子位元錯誤率約 0.02%；荷蘭 QuTech 則展示 5 個量子位元系統，錯誤率約 0.3%。相較三年前先進系統通常只有約 2 個量子位元、部分測量錯誤率可達 4%，顯示自旋量子位元在規模與準確度上都有明顯進步。HRL 的進展主要來自降低量子位元雜訊與改善控制系統工程。



傳統量子電腦往往需要大量銅線個別控制量子位元，不但增加系統複雜度，也會把熱量導入超低溫環境。HRL 採用更緊密整合的帶狀纜線設計，減少線路數量與熱負荷；RIKEN 也採取類似方法。這類控制與封裝技術，是未來把系統擴充到數千甚至更多量子位元的重要基礎。目前自旋量子電腦仍處於原型階段。超導量子處理器已可達百個以上量子位元，中性原子系統甚至能達數千個，而真正具有實用價值的容錯量子電腦，可能需要數萬個實體量子位元持續執行錯誤校正。因此，HRL 的 18-qubit 系統並不代表自旋量子位元已超越其他技術，其真正意義在於證明量子位元品質與錯誤率已逐漸接近主要競爭平台。下一個核心問題將是：能否在大幅增加量子位元數量時，仍維持目前的低錯誤率。

人工智慧輔助量子運算推展

《邊緣危機》的影集”Peter”一開始出現彼得·畢夏普 (Peter Bishop) 的墓碑，記載彼得出生於 1978 年、死亡於 1985 年。科學家華特向彼得說明過去的事情。彼得小時候曾罹患重病，當時華特與妻子想盡各種方法，希望能夠挽救他的生命。然而，墓碑上的紀錄顯示，原本屬於這個世界的彼得最後仍然在 1985 年去世。因此，已經長大成人的彼得，其身世描述更特殊的情境，也就是在兩個平行世界之中，存在著另一個彼得。這也成為後續理解兩個平行世界如何產生交會、糾纏與影響的重要起點。1985 年，年輕的華特在軍方支持下研究「光子窗」，這項裝置能透過捕捉游離光子，觀察一個與現實世界相互糾纏的平行世界。兩個世界高度相似，但科技發展路徑有所不同，例如另一個世界仍保有齊柏林飛船等技術。華特因此認為，觀察平行世界不僅具有科學價值，也可能從另一個世界較先進的科技中取得研究線索。然而華特持續研究光子窗還有更私人的原因。他年幼的兒子彼得身患重病，免疫功能持續衰退。華特一方面自行尋找治療方法，一方面觀察另一個世界的「華特分身」，希望對方能找到解藥。但在治療方法出現之前，這個世界的彼得已於 1985 年病逝。彼得死後，華特仍持續觀察另一個世界，因為那裡的彼得同樣罹患相同疾病。某次實驗中，另一個世界的華特其實成功找到了可能治癒彼得的化合物，實驗曾短暫出現代表成功的藍色反應。然而，他因意外分心而錯過關鍵瞬間，以為實驗再次失敗。透過光子窗觀察整個過程的華特，卻知道正確配方已經



出現，因此決定介入，試圖校正這個錯誤。華特重新合成藥物，並利用萊登湖作為穿越兩個平行世界的地點，希望湖水能吸收穿越時產生的多餘能量。但穿越過程發生意外，藥物遭到破壞，使他抵達另一個世界後無法直接治療彼得。為了救孩子，他只好把另一個世界的彼得帶回自己的世界接受治療。因此，後來長大成人的彼得，其實並不是原本世界的彼得，而是華特從平行世界帶回來的彼得。華特原本只是想修正一次實驗上的錯失並挽救彼得，卻因跨越世界邊界造成更大的後果。彼得被帶離原本世界，不僅改變兩個世界原有的歷史，也在平行宇宙之間留下裂痕。換言之華特成功挽救了一個彼得，卻以破壞兩個世界的平衡為代價，而這次選擇也成為日後兩個平行世界衝突與整個故事發展的重要起點。

人工智慧、機器學習、深度學習到生成式 AI，是一個逐步向內延伸的技術體系。最外層是人工智慧，也就是希望讓電腦能夠模擬人類的智慧行為，進一步協助解決問題。人工智慧發展至今，已經逐漸改變人類的生活方式，也推動許多科學研究的進展。其中非常重要的一個分支就是機器學習，它透過大量資料與演算法進行模型訓練，不再完全依靠人類事先制定固定規則。舉個最簡單的例子，過去如果要讓電腦辨識一隻貓，可能需要先告訴電腦，貓具有兩隻眼睛、兩個耳朵、鬍鬚與四隻腳，再依照這些事先設定的特徵進行判斷。但是機器學習的概念不同，可以提供大量不同影像與分類結果，讓模型從資料中學習如何辨識貓與其他物體。這就是機器學習與傳統規則式推論很重要的差異。當機器學習進一步發展到深度學習之後，透過神經網路，可以從大量且複雜的資料中找出更深層的特徵與隱藏關係。這些能力也逐漸應用於科學研究、疾病預測與新藥開發等領域。再進一步發展為生成式 AI，可以生成文字、語音、影像與程式碼，大型語言模型也是其中重要的應用。然而，模型愈來愈複雜，所需要處理的資料量與運算需求也愈來愈龐大，因此「算力」逐漸成為人工智慧持續發展的重要挑戰。量子運算也因此受到關注，希望未來能夠利用不同於傳統電腦的運算方式，協助處理部分高度複雜的計算問題。但這裡出現了一個很有趣的轉折。原本是希望利用量子運算協助人工智慧提升運算能力，後來卻發現量子運算本身也有非常重要的問題需要人工智慧協助。量子系統雖然具有計



算潛力，卻非常容易受到環境雜訊與干擾影響。尤其在量子疊加與量子糾纏等高度精密的狀態下，即使非常微小的干擾，也可能造成運算錯誤。

人工智慧在實體量子運算的發展過程中，可以扮演非常重要的輔助角色。前面提到，量子電腦就像是一位具有驚人計算能力的天才，具有非常強大的運算潛力，但是對環境極為敏感，微小的雜訊與干擾都可能造成錯誤。因此，如何利用人工智慧協助量子系統維持穩定、降低錯誤，就成為非常重要的研究方向。人工智慧的輔助其實從量子硬體開發階段就已經開始。量子晶片需要考慮量子位元的設計、排列與穩定性，以及不同硬體條件可能造成的影響。面對大量而複雜的設計組合，人工智慧可以協助分析與改進量子硬體，尋找更適合的設計方式。進入實際運算之前，人工智慧也可以協助改善量子演算法的執行效率。這就像飛機在起飛之前需要規劃航線，必須事先安排如何以更有效率的方式到達目的地。同樣地，量子運算也需要適當的前處理與運算規劃，減少不必要的計算，提高整體執行效率。真正開始運作之後，量子系統又會面臨非常複雜的控制問題。因為量子裝置對外界環境非常敏感，不同參數的細微變化，都可能影響最後的運算結果。人工智慧可以協助管理這些複雜的控制條件，從大量資料中尋找較佳的參數組合，使量子裝置維持在較穩定的運作狀態。但是，即使硬體設計、演算法與控制流程都經過最佳化，量子運算仍然可能發生錯誤。因此，量子錯誤校正就成為整個過程中非常關鍵的一環，也是今天要特別討論的主題。人工智慧可以協助辨識量子系統所產生的錯誤模式，改善錯誤校正的解碼器，甚至協助發展新的錯誤校正方法。用「天才」來比喻，量子電腦就是一位計算速度非常快、能力非常強，但是容易受到干擾而犯錯的天才；人工智慧則像是旁邊的小幫手，在運算過程中協助發現錯誤、判斷錯誤發生的位置，再進一步進行校正。只有把錯誤有效控制下來，量子運算強大的計算能力才能真正發揮。量子運算完成之後，人工智慧仍然可以協助分析輸出的結果，降低誤差並提高結果的可觀測性。就像一架新設計的飛機完成飛行之後，仍然需要重新檢視整個飛行過程，包括路線是否適當、能源使用是否有效率，以及過程中受到哪些環境因素影響，再利用這些資訊持續改善下一次的表現。人工智慧可以從量子硬體開發、運算前處理、裝置控制與最佳化，



一直到量子錯誤校正與最後的結果分析，參與整個量子運算流程。人工智慧與量子運算並不是單向的關係，並非只有量子運算未來可以提升人工智慧的算力；反過來，人工智慧也正在成為推進量子科技的重要工具。

量子運算容易受到環境影響重要原因是量子位元與傳統電腦的位元差異。傳統位元主要以 0 與 1 表示資訊，但量子位元可以處於 0 與 1 的疊加狀態，因此能夠表現更加豐富的量子資訊。不過，也正因為量子狀態非常精密，環境中微小的雜訊與控制偏差，都可能使量子狀態發生改變。量子資訊可能出現幾種不同的錯誤型態。例如 X 錯誤，可以理解為類似傳統位元 0 與 1 之間的翻轉，也就是 bit flip；Z 錯誤則是相位發生改變，也就是 phase flip；Y 錯誤則同時包含 X 與 Z 兩種成分。因此，量子錯誤並不只是單純的 0 變成 1 或 1 變成 0，而可能涉及量子狀態不同方向與相位的改變，這也使量子錯誤的偵測與校正更加複雜。除了量子狀態本身的錯誤之外，實際操作時還可能發生測量錯誤，也就是最後讀取量子資訊時得到不正確的結果。另外還有一種重要的錯誤稱為「洩漏」(leakage)。原本量子位元的計算空間設定在 $|0\rangle$ 與 $|1\rangle$ 兩個狀態，但是量子系統實際上可能進入這個計算空間以外的其他能階。可以用考試來理解洩漏的概念。假設一道題目規定學生只能回答 A 或 B，但是學生最後卻回答 C。這個 C 並不在原本設定的答案空間之內，因此系統無法直接按照原來 A 或 B 的規則處理。量子位元脫離原本 $|0\rangle$ 與 $|1\rangle$ 的計算子空間，就是類似的情況。另一種錯誤是「串擾」(crosstalk)。不同量子位元原本應該按照設定獨立進行控制，但是當鄰近量子位元彼此產生非預期的影響，就可能造成串擾。就像一個空間裡原本 A 與 B 正在交談，但聲音過大，連旁邊的 C 也受到影響。量子位元彼此距離很近、控制又非常精密，因此這類非預期的相互干擾也是需要處理的重要問題。更困難的是，量子錯誤不能直接透過測量資料量子位元來檢查。因為量子資訊具有疊加特性，一旦直接進行測量，量子狀態就會坍縮，原本保存的量子資訊也可能因此遭到破壞。所以量子錯誤校正的核心，是透過輔助量子位元進行間接偵測。

這種方法會測量多個量子位元之間的關係，例如 parity 或 stabilizer，取得所謂的「syndrome」訊號。Syndrome 本身並不是直接讀取原來的量子資訊，而是保留錯誤所



留下的線索，再利用這些線索反推出錯誤可能發生的位置與型態。這就像警察辦案時，不一定能夠直接看到犯罪發生的整個過程，但是可以從現場留下的腳印、痕跡或監視器中的異常訊號，反推出事件是如何發生的。量子錯誤校正也是相同概念：不直接讀取並破壞原本的量子資訊，而是觀察錯誤留下來的「痕跡」，再從這些 syndrome 訊號判斷哪裡可能發生錯誤。因此，真正困難的問題就變成：當系統產生大量而複雜的 syndrome 訊號之後，如何快速而準確地從這些線索中找出真正的錯誤？這也正是人工智慧與機器學習可以進一步介入量子錯誤校正的重要地方。

也因為量子系統本身非常敏感，因此在實際運算過程中，除了前面提到的環境雜訊之外，還包括位元翻轉、相位錯誤、洩漏以及量子位元之間的串擾等問題，這些都可能讓原本應該正確的量子資訊產生偏移。目前實體量子位元單次操作的錯誤率，大約千分之一到百分之一。對需要大量連續運算的量子系統而言，錯誤會隨著操作次數持續累積。可靠容錯量子運算，邏輯錯誤率必須進一步降低到接近 10^{-12} 。因此如何把大量容易出錯的「實體量子位元」，轉換成相對穩定而可靠的「邏輯量子位元」。是利用多個實體量子位元共同編碼一個邏輯量子位元，再持續監測其中所產生的錯誤訊號。透過多次、間接的穩定子測量，取得一系列的 syndrome，也就是錯誤留下來的訊號，再由這些訊號反推出真正可能發生錯誤的位置。其中一個重要的方法就是 surface code(表面碼)。它的概念可以理解成：不要只依靠單一觀察結果來判斷錯誤，而是利用許多量子位元彼此之間的關係進行共同監測。就像一件事情如果只有一個人看到，很容易因為個人判斷錯誤而失真；如果同時有很多人從不同位置持續觀察，就比較有機會從彼此一致或不一致的訊息中找出異常。真正困難的地方就在於，當系統產生大量 syndrome 訊號之後，必須快速判斷究竟是哪一個位置、哪一種型態的錯誤造成目前看到的結果。這個負責判讀 syndrome、推測最可能錯誤位置，再決定如何修正的過程，就稱為「解碼」。因此，量子錯誤校正並不只是單純把錯誤找出來，而是要從大量間接訊號中進行推理。解碼器的準確度愈高，就愈能將容易受到干擾的實體量子位元，轉換成穩定可靠的邏輯量子位元。也正因為這個判讀過程非常複雜，人工智慧與機器學習才有機會在其中發揮重要作用，



協助從大量 syndrome 訊號中更快速、更準確地辨識錯誤。

AlphaQubit 量子糾錯會透過多個實體量子位元進行編碼與反覆監測，取得前面提到的 syndrome 訊號，再利用這些訊號推測真正的錯誤位置。其中，surface code，也就是表面碼，就是重要的量子錯誤校正方法之一。透過多個量子位元之間的關係進行監測如同有許多觀察者從不同位置共同監督，只要其中出現異常，就可以從整體訊號的變化中尋找錯誤留下的痕跡。但是當量子位元數量增加、監測次數愈來愈多之後，產生的 syndrome 資訊也會變得非常龐大而複雜。如何從這些帶有雜訊的訊號中快速找到最可能的錯誤，就需要解碼器。AlphaQubit 所扮演的就是 AI 解碼器的角色。它利用資料驅動的神經網路，學習大量量子錯誤與 syndrome 訊號之間的關係，再根據實際觀察到的訊號，推測最可能發生的錯誤並進行校正。換句話說，AI 並不是直接去讀取原本的量子資訊，而是從錯誤所留下來的線索中進行判讀。這樣才能將原本容易受到雜訊影響的實體量子位元，進一步保護成比較可靠的邏輯量子位元。所以量子糾錯真正的關鍵之一，就是能不能做到快速而準確的解碼。AlphaQubit 以人工智慧技術協助量子系統從複雜的錯誤訊號中找出問題、完成校正，提高量子運算的可靠性。

AI 解碼器在量子糾錯中可以扮演非常重要的角色。因為量子系統發生錯誤時，真正需要判斷的往往不只是當下這一次的錯誤，而是必須把前面一連串的錯誤訊號整合起來，才能判斷錯誤究竟是如何形成的。因此，量子糾錯本身具有時間序列與前後關聯的特性。過去在人工智慧處理這類序列資料時，經常使用 RNN，也就是循環神經網路，後來進一步發展出 LSTM，也就是長短期記憶網路，希望能夠同時保留較短期與較長期的資訊。但是隨著 Transformer 的出現，透過注意力機制，可以更有效率地處理不同位置之間的關聯，也更適合分析較長而複雜的序列資訊。這個概念放到量子糾錯就非常重要。因為當量子位元產生錯誤時，到底要往前追蹤多長的資訊才能正確判斷，並不一定事先知道。就像在臨床上看到一位病人的血壓突然升到 200 mmHg，不能只根據這一次的血壓就判斷問題，而是需要回頭觀察前面一段時間的血壓變化，才能知道這是突然發生的異常，還是長期累積形成的結果。量子糾錯也是相同的概念。系統會持續產生大量帶有



雜訊的量子資訊，經過編碼之後，再由 Transformer 解碼器整合不同時間與不同位置的錯誤訊號，從這些複雜的關聯中預測最可能發生錯誤的量子位元。另一個重要問題是量子糾錯的規模並不是固定的。當量子系統從較小的碼距，例如距離 5，進一步擴展到較大的碼距，例如距離 7 時，需要處理的量子位元與錯誤資訊也會隨之增加。因此，解碼器除了要處理時間上的長短期關聯，也需要具備處理不同輸入規模的能力。這就像一位醫師原本在規模較小的醫院照顧病人，當醫療體系擴大到更大型的醫院時，需要面對的病人數、資訊量以及彼此之間的關係都會增加。如果每次規模擴大都必須完全重新建立一套判斷方式，效率就會非常有限。因此，理想的 AI 解碼器應該能夠保留原本已經學習到的錯誤特徵，再進一步延伸到更大規模的量子糾錯問題。Transformer 的價值就在這裡。它可以將量子系統所產生的雜訊與錯誤訊號轉換成模型可以學習的資訊，整合不同位置與不同長度之間的關聯，再由解碼器預測可能發生錯誤的位置。當系統規模進一步增加時，也能夠支援更長的輸入與更大的表面碼，讓原本在較小碼距中學習到的資訊，有機會延伸到更大規模的量子糾錯。因此，AI 解碼器真正重要的地方，不只是「看到一個錯誤就修正一個錯誤」，而是能夠把前後不同時間、不同位置以及不同規模的錯誤訊號整合起來，找出其中真正具有意義的關聯。這也是 Transformer 應用於量子錯誤校正的重要價值。

人工智慧要真正協助量子錯誤校正，除了分析前面累積的錯誤訊號之外需加入軟性量測資訊。在判斷結果時，不只是簡單分成 0 或 1，而是進一步保留其中的不確定性與機率資訊。例如在醫學影像判斷中，如果以 50% 作為良性與惡性的分類界線，一個惡性機率 90% 的病灶與一個惡性機率 51% 的病灶，最後雖然都可能被分類為惡性，但是兩者的確定程度顯然完全不同。如果最後只保留 0 或 1 的分類結果，這些重要資訊就會消失。因此，軟性量測不是只告訴模型「答案是 0 還是 1」，而是進一步保留「比較接近 0 的機率是多少、比較接近 1 的機率是多少」。同樣的概念也可以應用在量子糾錯。量子系統本身具有高度不確定性，實際讀取時也可能受到雜訊影響，因此如果能夠把量測所包含的不確定性資訊一起交給 AI，模型就可以同時整合過去累積的 syndrome 訊號與



當下量測的可信程度，學習雜訊與錯誤之間複雜的關聯，再進一步推論最可能發生的邏輯錯誤。這也是人工智慧特別具有吸引力的地方。神經網路可以從大量資料中學習人類不容易直接辨識的複雜關係，也可以結合真實量子硬體所產生的資料，而不是完全依賴理想化的模擬結果。RNN、Transformer，甚至圖神經網路，都具有應用在量子解碼上的潛力。其中前面介紹的 Transformer，更適合進一步處理不同長度與不同規模的錯誤資訊。因此，AI 解碼的概念其實很接近臨床預測。模型不是只給出一個絕對的答案，而是可以估計不同結果發生的可能性，再進一步評估判斷的準確程度。當 AI 能夠更有效地學習這些複雜關係、整合真實資料與不確定性資訊，就有機會提高量子錯誤解碼的準確性，讓量子運算變得更加可靠。

當這些技術逐漸成熟之後可推展應用到未來的臨床決策。不同問題的資料量、複雜程度與即時性不同，所需要配置的運算資源也應該有所不同。最基本的臨床工作，例如院內即時風險預測或一般決策支援，可以直接利用院內 GPU 進行即時推論。這類問題利用既有的人工智慧與古典運算就可以有效處理，沒有必要每一次都動用更昂貴、更複雜的量子資源。就像不是所有問題都需要請一位「計算天才」來處理，有些工作交給既有的小幫手就已經足夠。當資料的維度與複雜程度進一步提高，例如需要處理高維度的臨床特徵，就可以導入量子啟發式演算法。像張量網路等方法，可以利用量子資訊科學的概念，在古典電腦上進行高維特徵的表示與運算，不一定需要真正進入量子硬體。如果問題再更加複雜，例如同時整合醫學影像、基因體、檢驗數據與臨床病歷等多模態資料，就可能需要進一步採用混合量子與古典運算。讓古典電腦負責適合的資料處理與控制工作，再由量子運算處理特定的高複雜度問題，兩者彼此分工，而不是完全以量子電腦取代傳統電腦。當研究進一步擴展到國家級、大型跨機構資料，甚至需要進行高度複雜的分子模擬或大規模搜尋時，才可能需要更高階的研究級容錯量子資源。也就是說，未來真正可行的模式並不是「所有問題都交給量子電腦」，而是根據問題的複雜程度，選擇最適合的計算資源。因此，人工智慧與量子運算未來更可能形成一個分層、混合式的決策輔助架構。簡單而即時的問題交給古典 AI 處理，高維度問題可以導入量子啟發式方



法，更複雜的多模態問題再使用混合量子運算，而真正需要龐大計算能力的研究問題，才進一步使用容錯量子資源。這也回到今天最重要的主題。過去經常討論量子運算未來如何幫助人工智慧，但在真正走向實際應用之前，人工智慧同樣可以成為量子運算的重要助手。

量子運算精準敗血症照護應用情境

我們以敗血症臨床處置為例說明多層結構 AI 輔助應用。針對院前緊急救護階段，可運用於救護車上進行初步評估與 AI 風險預測，並將資料無線傳輸至急診系統。透過敗血症風險評分，整合近期抗生素治療、生命徵象、意識狀態及慢性腎臟病等資訊，協助保留紅區床位、預約乳酸與血液培養。另運用張量網路壓縮長期病歷資料，支援即時推論與風險預測。針對敗血症個人化嚴重度評估，系統整合乳酸、白血球、CRP、肌酸酐、血小板、膽紅素及 SOFA 等資訊，估計符合敗血性休克定義後的機率與 90 日死亡風險。量子/AI 技術則透過量子振幅估計與貝氏後驗分布，以較少樣本取得完整後驗分布，呈現不確定性與尾部風險，支援 ICU 聯繫及資源配置。導入量子近似最佳化演算法 (QAOA)，整合臨床、檢驗、影像與生命徵象等高維病人特徵，將病例轉換為 QUBO 分群問題，再以量子搜尋找出最佳分群結果。系統可將病人歸納為不同臨床表型，協助辨識快速惡化的高風險患者，提前規劃加護照護與升階治療。透過掌握個別病人的疾病進展路徑，也能支援治療選擇、照護強度及家屬風險溝通，使臨床決策更即時且精準。

在敗血症等重症治療中，循環支持常面臨「輸液優先」或「及早使用升壓劑」的抉擇。導入混合量子古典強化學習，整合血壓、尿量、乳酸及輸液量等即時資訊，模擬不同治療策略並持續回饋修正。系統可依個別病人狀態比較治療效果，協助醫師更早辨識適合的升壓與輸液策略，降低傳統固定流程造成的誤差，提升急救與重症照護決策的即時性與精準度。敗血症病原判斷涉及人體、病原體、共病史與基因體等多維資訊，傳統運算在高複雜度分析上可能面臨限制。研究構想結合量子核方法與 AI，整合病人基因型、表現型、微生物與抗藥性特徵，在培養結果出爐前推估最可能病原及抗藥性。藉由提升病因判斷速度與準確度，可協助醫師更早選擇適當抗生素，從經驗性治療進一步邁



向個人化、精準對症治療。

以上內容將在 2026 年 8 月 26 日(三) 10:00 am 以線上直播方式與媒體朋友、全球民眾及專業人士共享。歡迎各位舊雨新知透過星球永續健康網站專頁觀賞直播！

- 星球永續健康網站網頁連結：
<https://www.realscience.top/7>
- Youtube 影片連結：<https://reurl.cc/o7br93>
- 漢聲廣播電台連結：<https://reurl.cc/nojdev>
- 不只是科技：<https://reurl.cc/A6EXxZ>



講者：

陳秀熙教授/英國劍橋大學博士、許辰陽醫師、陳立昇教授、嚴明芳教授、林庭瑀博士

聯絡人：

林庭瑀博士 電話：(02)33668033 E-mail：happy82526@gmail.com

劉秋燕 電話：(02)33668033 E-mail：rl1847030@ntu.edu.tw