

本文英文版是在聯合國教科文組織(UNESCO)的 *Blue Dot* 雜誌發表

通智同伴假說

全球和幸應成為人工智慧發展的基礎

全球和幸合作學者社群

(這篇文章源自一批學者共同提出的理念，並參考了 globalharwellgoal.org 上的相關資料，包括一篇選文與「全球和幸」的說明與常見問題。由於這個議題牽涉廣泛，「全球和幸合作學者社群」除了原始參與者，也持續邀請國內外認同此理念的學者加入，名單會隨著加入者而不斷擴充。中文翻譯比原文補充了一些細節。)

**歷史有重量，讓人站穩現在；
歷史有高度，從過去看未來，讓人看得更遠**

關鍵詞：通智同伴假說 (Generative Artificial Companion Hypothesis), 全球和幸 (Global Harwell), 無縫人工智慧世界 (Seamless AI World), 機器人同伴世紀 (Robot Companion Age)

通智同伴

早在 1920 年代，Sidney Pressey 推出了他的「教學機器」——一種他聲稱將徹底改變教育方式的機械裝置，這一創舉比 1946 年第一台電腦的問世早了許多年。時至今日，一個世紀過去了，學生們使用的平板設備已經展現了實現 Pressey 願景的可能性。1950 年，Alan Turing 提出了一種模擬遊戲，旨在評估機器通過自然語言處理展示類人智慧的能力，這一概念如今廣為人知，即「圖靈測試」。時光飛逝至 2022 年——在圖靈提出該概念 70 多年後——ChatGPT 的推出，標誌著生成式人工智慧技術的重大突破。如今，大眾越來越期待機器能夠通過圖靈測試。

目前，人工智慧 (AI) 正加速全球教育的變革，為個性化學習、教師支援及教育機會的擴展提供了嶄新的可能性。在美國，像 Carnegie Learning 這樣的平台利用自我適應人工智慧技術，識別學生的知識差距並量身定制教學內容，幫助學生以適合自己的步調掌握代數概念。AI 驅動的遊戲化工具，如 Duolingo，不僅提升了使用者的語言學習動機，還有效降低了輟學率。此外，Panorama Education 等平台透過提供學生學習進展的深入分析，使教育者能夠更精準地制定教學策略。然而，數位鴻溝和隱私問題仍是當前面臨的重大挑戰。在歐洲，芬蘭的 Elements of AI 計劃積極推廣 AI 素養，德國則著重於語言整合的支援。儘管如此，嚴格的監管框架和有限的教師培訓資源仍阻礙了進一步的發展。在亞洲，雖然文化阻力依然存在，印度的 BYJU'S 成功緩解了教師短缺的問題，而中國則利用 AI 技術實現高效評分，提升

了教育效率。在非洲，肯尼亞的 Eneza Education 等計劃顯著提升了教育資源的可及性和技能發展，但基礎設施薄弱和語言障礙仍限制了其影響力。

儘管人工智慧在教育領域展現出巨大的變革潛力，但要實現可持續且包容的全球應用，仍需克服諸多挑戰。這些挑戰包括公平性、數據隱私、算法偏見、AI 驅動決策的透明度，以及 AI 資源的可取得性。唯有妥善解決這些問題，才能確保人工智慧在全球教育中的廣泛應用，真正實現其對教育變革的承諾。

人工智慧的發展可被視為一條連續的光譜，從「工具」到「助手」，再進一步演進為「同伴」。在這段人機關係的光譜中，角色的轉變並無絕對的界線，但我們仍可依據 AI 與人類互動的密切程度，對其進行初步的分類。

在最初階段，AI 扮演的是「工具」的角色，主要由人類手動操作，用於輔助完成特定任務。此時的 AI 功能明確，且依賴人類的指令來執行工作。隨著技術的進步，AI 逐漸進化為「助手」，能夠主動執行人類指派的任務。例如，生成式 AI 在這一階段扮演了重要角色，廣泛應用於多個領域：協助開發程式碼、製作動畫、根據自然語言查詢檢索資訊、進行語言翻譯，甚至提供法律諮詢等服務。

當 AI 進一步發展為「同伴」，其與使用者的互動模式也變得更加雙向且主動——不僅能回應人類發起的對話或活動，甚至能主動引導互動。事實上，我們與隨身裝置（如智慧型手機）的互動方式，正逐漸從傳統的工具使用，轉變為以自然語言驅動的交流模式。使得人機關係從單向操作轉向為更具投入感與情感連結的同伴交流。

生成式人工智慧的興起，預示著見多識廣及人情練達的人工同伴(artificial companions)即將到來。關於人工學習同伴的研究已有悠久的歷史（這也是陳德懷教授 40 年前博士論文研究主題）。為更深入理解人工同伴的本質，讓我們首先探討「同伴」與「陪伴」的概念。

從廣義的角度來看，「同伴」可以包括我們的父母、配偶、子女、朋友、教師、導師、教練、同學，甚至是寵物。「陪伴」則被定義為一種雙向社會關係，其特徵包括情感支持、共同活動、信任關係、相互尊重、開放溝通，以及共享時光的愉悅感。例如，在親子關係中，角色通常是明確界定的，其核心目標是父母培育並教導子女，而子女則在父母的關懷下學習與成長。這種互動模式說明了同伴關係中的角色如何建構整體目標，並影響持續的互動過程。因此，「陪伴」可被理解為涵蓋關係動態、共同目標與持續互動的歷程。

我們可以將當前的時代描述為「無縫人工智慧世界」(Seamless AI World, SAIW)。在這個高度互聯的世界中，幾乎所有事物都透過網路連接並由人工智慧驅動，我們的地球似乎正在「收縮」或變得「更小」。在人工智慧的幫助下，人們可以即時翻譯，以母語無縫交流，地理距離就逐漸消失。

作為一個學習環境，SAIW 提供了一個「無縫學習」(seamless learning)體驗的機會，使人們能夠在教室、元宇宙和現實世界之間自由轉換。無縫性(seamlessness)的一個關鍵面向是連續性(continuity)——即支援跨時間與空間的連續活動。然而，無縫性不僅僅關乎連續無縫性(continuity seamlessness)，還包括可及無縫性(accessibility seamlessness)以及類人無縫性(resemblance seamlessness)，特別是在人工同伴的支援下。

然而，無縫性的不同層面可能帶來「正面無縫性」(well-seamlessness)的益處，也可能引發「負面無縫性」(ill-seamlessness)的風險。對於後者，必須謹慎評估，並視情況設立明確的範圍界限與介面。若輕忽這些挑戰，恐將造成無法挽回的社會傷害。隨著人工同伴時代的到來，這一議題更顯迫切。

「人工同伴」指的是具複雜智慧與社交情商的人工智慧實體(entity)，如機器人或虛擬角色，其設計目的在於支援並增進人類的日常活動，同時促進人際互動與關係的建立。隨著 AI、網際網路、元宇宙，以及擴增實境(AR)、機器人技術、物聯網 (IoT)、量子運算與區塊鏈等技術的快速發展，SAIW 中的生活體驗範圍正迅速被重新定義。這些技術的融合將使人工同伴變得普遍，並成為日常生活中不可或缺的一部分。

自駕車與無人機的崛起，象徵著機器人同伴時代(Robot Companion Age)的來臨。未來，具備通用人工智慧(Artificial General Intelligence, AGI)的通智同伴(General Artificial Companions)或許將成為現實。這類 AI 不僅擁有與人類相當，甚至超越人類的智能，更能展現多樣化類人行為，包括倫理思考、情感智慧與社會意識。

如果這一願景得以實現，數位模擬的真實性將可能進一步模糊真人與通智同伴之間的界線。隨著時間推移，這些通智同伴將能累積並記錄個人的一生經驗，從而扮演不同人生階段中的不同角色，成為個人成長與生活的陪伴者。

然而，隨著通智同伴日漸普及，人們開始關注其在更廣泛社會脈絡中的角色與影響。這些關注既可能帶來正面效益，也可能引發負面憂慮。當人類與通智同伴共存於同一社會體系時，整體社會目標將變得更加複雜，充滿多重變數與挑戰。在人機人共存的未來中，如何平衡發展並建立倫理價值觀，將成為我們必須深思的重要課題。

自第一台電腦誕生至今已 80 年，如今，擁有一台或多台計算設備，如智慧型手機、平板電腦、桌上型電腦，已逐漸成為常態。也許我們無需再等待 80 年，可能僅需 10 到 20 年，或者稍久一點，每個人便可能擁有一台或多台通用人工智慧驅動的機器人同伴。

想像一下，未來的機器人同伴將無縫融入我們的生活：一台機器人可留在家中打理家務，另一台則在職場或學校中輔助使用者，並與其數位分身即時連結與合作。當前全球人口約為 80 億，若將機器人同伴納入計算，則「總人口數」可能翻倍至 160 億，這將徹底改變我們對社會結構與資源分配的認知。

在這個即將到來的機器人同伴時代，社會結構、工作環境，以及人與 AI、人與人之間的關係都將迎來深刻的改變。

全球和幸(Global Harwell)

如果我們暫時撇開宗教教義、政治意識形態與特定信仰，單純探討個人對人生的期許，許多人可能會提到幸福、健康、財富與成就等目標。這些追求往往與馬斯洛的需求層次理論(Maslow, 1943)——包括生理需求、安全需求、愛與歸屬感、尊重需求及自我實現——以及塞利格曼的 PERMA 模型(Seligman, 2011)——包括正向情緒、投入感、良好人際關係、人生意義與成就——不謀而合。在現實世界也需要考量健康與財務安全。無論如何，這些理論框架為我們理解美好生活提供了重要的參考。

若將幸福(wellbeing)視為一個整體概念，涵蓋生理、心理與社會健康，並體現幸福感、生活滿意度以及有效運作的能力，那麼馬斯洛與塞利格曼的理論便與這種對幸福的理解高度契合。

然而，當今世界正面臨前所未有的挑戰：COVID-19 造成的大量死亡、氣候變遷的加劇、資源枯竭、環境污染、財富不均，再如網路犯罪、深度偽造(deepfake)與假資訊泛濫，以及 AI 技術帶來的風險。

這些問題不僅加劇了社會分歧，還引發了全球範圍軍事與經濟的衝突升級，甚至讓人們對核災難的恐懼與日俱增。

在這樣的背景下，我們不得不思考一個根本問題：如果人類無法在這顆星球上持續生存，幸福又從何談起？

僅僅追求個人幸福並不足夠，我們必須將「和諧」置於更高的優先位置，積極促進人與人之間、以及人與環境之間的正向關係。和諧可分為兩個層面：「環境和諧」與「人類和諧」。

環境和諧涵蓋全球暖化與自然災害等議題，強調人類應與自然共生共存，尊重生態平衡，並致力於永續發展。而人類和諧則涉及個人、家庭、社會與全球各層面的互動與平衡，旨在創造一個彼此尊重、合作共榮的世界。

人類和諧有四大關鍵要素：仁愛、公平、正義與中庸之道。

- 仁愛(Benevolence)：是和諧的基礎，仁慈、善良與同情心，是人類彼此關懷與支持的基礎。
- 公平(Equity)：確保資源依需求分配，並促進機會均等，讓每個人都能在合理的環境中追求成長。
- 正義(Justice)：維護倫理道德與誠信正直，是社會穩定與信任的基石，確保每個人的權利與尊嚴得到保障。
- 中庸之道(The Middle Way/The Golden Mean)：強調平衡與適度，避免極端，以維持長久和諧。

簡言之，仁愛奠定和諧基礎；公平促進資源均衡分配，正義維護道德倫理，二者共同維持社會穩定，而中庸則防止極端，以守護和諧。

最終，人類和諧體現了儒家、道家與墨家的精神，也是西方古典哲學推崇的「黃金法則(The Golden Rule)」，它的意思就是「己所不欲，勿施於人」，或以正向表達：「待人如待己」(Treat others as you wish to be treated)。這條法則不僅是人際關係的指南，更是人類和諧的核心理念。

「全球和幸」(Global Harwell)結合「和諧」(harmony)與「幸福」(wellbeing)的概念，代表影響全球的社會規範、倫理的核心價值與集體原則。它不僅是一個理念，更是一種全球性的願景，旨在將追求和諧與幸福視為人類的基本價值觀。

全球和幸的願景與「聯合國教科文組織」(UNESCO)及各國教育部的教育理念與目標高度契合。它提供了一個全面的全球視野，將和諧與幸福視為實現全球和平與永續發展的關鍵要素。例如，聯合國的「永續發展目標」(Sustainable Development Goals)便與和諧的概念緊密相連，強調環境保護、社會公平與經濟發展的平衡。儘管不同國家與地區可能因其文化、政治或地理環境而設定不同的發展目標，但全球和幸所承載的是一種超越國界差異、建基於人類共同願景的普世價值。我們相信，人類知識與科技技術發展的首要目標，就是實現全球和幸。

為了實現全球和幸的願景，研究人員正從多個方向積極探索與推動。其中，一組研究人員致力於開發教育虛擬同伴，提升社交媒體的使用體驗，並對可能危及個人與群體幸福的風險提出因應之道。這些虛擬同伴不僅能提供個性化的學習支持，還能幫助學生建立健康的數位習慣，促進正向的社交互動。

與此同時，另一組研究人員則專注於探索全球教育、全球大腦、全球倫理以及全球人工同伴在促進社會和諧與幸福中所扮演的角色。這些研究試圖通過整合全球資源與智慧，建立一個更加互聯、包容且倫理導向的社會體系，從而推動全球和幸的實現。

無論最終結果如何，越多的人將全球和幸視為其核心人生價值，越少的人或機構會為私利所驅動，也越減少無意間開發出對人類有害的人工智慧的風險。至少，這將在社會層面對 AI 的惡意濫用施加道德與行為約束，從而降低其可能帶來的負面影響。

「通智同伴假說」(General Artificial Companions Hypothesis)

「通智同伴假說」提出，隨著人工同伴技術的不斷演進，關注人與機器人之間的陪伴關係及其對社會發展的影響將成為關鍵課題。若能優先發展符合倫理原則的人工智慧，不僅能促進人類社群的和諧共存，還能透過創新與持續進步來提升個體幸福與社會和諧。在通用人工智慧的支持下，這些機器人同伴將引領我們朝向全球和幸的願景邁進(Chou et al., 2025)。

在 SAIW 環境中，「通智同伴假說」提出以下假設：

對於將全球和幸培育為個人的核心價值過程中，通智同伴所提供的支持將發揮關鍵作用：貢獻 50%，甚至更多，並且可從個人日常行為中得到驗證。

當假說被證實的那一天，象徵著科技輝煌成就的凱旋，也標誌著人類文明跨入新的里程碑。

然而，這也引發了一個深刻的問題：究竟是人類塑造人類，還是 AI 塑造人類？如果答案是人類塑造人類，那麼全球和幸就必須成為人工智慧發展的根本基礎。

我們必須清楚認知，通智同伴的發展，必須奠基於教育。真正的通用智慧同伴，只能由「全球和幸人」(Global Harwellians)——那些將全球和幸視為核心價值與人生目標的人——來設計。而這些設計者，也唯有透過教育，才能培養出來。

一旦成功建立通智同伴，「全球和幸性」(Global Harwellness) ——即體現全球和幸狀態的整體素質(quality)，將在人類與通智同伴的陪伴關係中相互深化與提升。屆時，人類與 AI 將攜手共塑世界，創造一個更加和諧、幸福且永續的未來。

正如 納爾遜·曼德拉(Nelson Mandela) 所說：教育是你可以用來改變世界的最強大武器。的確，教育確實能夠引領人類追求全球和幸，並讓我們預見根本性的變革，在教育也在全球社會，在今天也在未來。教育，不僅能培養出具有全球視野與倫理意識的「全球和幸人」，還能夠為通智同伴的發展奠定堅實的基礎，最終實現人類與 AI 共融共榮的願景。

結論

教育研究是一趟探索與揭示人類本能、性格發展、身分認同與生存意義的旅程。但唯有當我們的研究不僅對學生具有重要性，更能在更廣泛層面上造福整體人類，這趟旅程才具有真正意義。

本文闡述了通智同伴假說，提出一個未來願景：具備人類特徵的先導人工同伴，將幫助人類把全球和幸培育為核心價值與人生目標。

雖然這一假說帶來許多問題與挑戰，它同時也為 AI 研究者與實踐者提供了一個共同的願景：建構一個持久且和諧的世界，確保未來世代能夠在全球和幸的未來中，持續成長。實現這一願景的關鍵，在於教育——唯有透過教育，世界才能真正培養出將全球和幸視為核心價值的「全球和幸人」(Harwellian)。而開發通智同伴以實現這一假說，則是這項使命不可或缺的一環。

人工智慧已成為當今最具顛覆性的科技之一。一方面，如前所述，它為教育插上新的翅膀——不僅能促進個人化學習，協助教師預測學生表現並提供精準引導，更能提升學習動機與學習表現。此外，亦可透過即時翻譯消弭語言障礙，進而促進全球社群建構。

另一方面，隨著 AI 的光芒愈加耀眼，無縫人工智慧世界中的負面無縫性卻帶來諸多風險。它可能使我們下一代與社會脫節，成為焦慮的一代，並且加劇隱私侵犯與教育資源分配不均等問題。更值得警惕的是，若由通用人工智慧驅動的人工同伴採納的價值觀，與全球和幸的核心價值背道而馳，則將對人類未來造成巨大災難。我們必須審慎思考未來發展方向，確保教育始終扎根於人類價值觀，而非被科技異化。這是當今世界的核心問題。

全球和幸是一種理想、一套行動指南，它提醒我們，科技、經濟與社會變革都應服膺於人類的和諧與幸福。唯有如此，才能在科技飛速發展中堅守人性價值，創造真正和諧與幸福的未來。

「通智同伴假說」或可被視為人工智慧領域的「聖杯」——那個長久以來讓 AI 研究孜孜以求，極為艱難欲所企及的終極目標。若此假說得以驗證，它將有可能引領人類朝向康德 (Immanuel Kant) 於 1795 年所描繪的永久和平願景邁進，並透過世代累積培育出一個「全球和幸」的世界，使和平與幸福得以延續至未來。

只要人類懷抱崇高的夢想，夢想終有一日成真！

但是，我們不能被動等待未來；如果我們不行動，未來不會變得更好，只會變得更壞。現在就行動！

參考文獻

- Carnegie Learning (no date) *Algebra | Resources*. <https://support.carnegielearning.com/help-center/math/home-connection/program-resources/high-school-resources-by-course/article/algebra-1-resources/>
- Chan, T. W. and Baskin, A. B. (1988) 'Studying with the prince: The computer as a learning companion', *Proceedings of 1988 International Conference on Intelligent Tutoring Systems*, 194-200. Montreal, Canada.
- Chou, C.-Y., Chan, T.-W., Chen, Z.-H., Liao, C.-Y., Shih, J.-L., Wu, Y.-T., Chang, B., Yeh, C. Y. C., Hung, H.-C. and Cheng, H. (2025) 'Defining AI companions: a research agenda – from artificial companions for learning to general artificial companions for Global Harwell', *Research and Practice in Technology Enhanced Learning*, 20(032). Available at: <https://doi.org/10.58459/rptel.2025.20032>
- Kant, I. (1795). *Perpetual peace: A philosophical sketch*. Retrieved October 11, 2025, from Global Harwell website: http://fs2.american.edu/dfagel/www/Class%20Readings/Kant/Immanuel%20Kant,%20_Perpetual%20Peace_.pdf
- Maslow, A. H. (1943) 'A theory of human motivation', *Psychological Review*, 50(4), pp. 370–396.
- Seligman, M. E. P. (2011) *Flourish: A Visionary New Understanding of Happiness and Well-being*. New York: Free Press.

全球和幸合作學者社群

Vincent Aleven, Carnegie Mellon University
Roger Azevedo, University of Central Florida
Gautam Biswas, Vanderbilt University
Johanna Börsting, Hochschule Ruhr West University of Applied Sciences
Ivica Botički, University of Zagreb
Tak-Wai Chan, National Central University
Ben Chang, National Central University
Maiga Chang, Athabasca University
Gwo-Dong Chen, National Central University
Weiqin Chen, Oslo Metropolitan University
Wenli Chen, Nanyang Technological University
Zhi-Hong Chen, National Taiwan Normal University
Hercy Cheng, Taipei Medical University
Young Hoan Cho, Seoul National University
Cristina Conati, University of British Columbia
Charles Crook, Nottingham University
Vania Dimitrova, University of Leeds
Sabrina Eimler, Hochschule Ruhr West University of Applied Sciences
Usama Fayyad, Northeastern University
Dragan Gašević, Monash University
Arthur Graesser, University of Memphis
Xiangen Hu, Hong Kong Polytechnic University
Tristan Johnson, Boston College
Xiaoqing Gu, East China Normal University
Davinia Hernández-Leo, Universitat Pompeu Fabra
Ulrich Hoppe, University of Duisburg-Essen
Ting-Chia Hsu, National Taiwan Normal University
Huang Ronghuai, Beijing Normal University
Hui-Chun Hung, National Central University
Gwo-Jen Hwang, National Taichung University of Education
Sridhar Iyer, Indian Institute of Technology Bombay
Heisawn Jeong, Seoul National University
Lewis Johnson, University of Southern California
Morris Siu-Yung Jong, Chinese University of Hong Kong
Chi-Hung Juan, National Central University
Akihiro Kashiara, Tokushima University
Mas Nida Md. Khambari, Universiti Putra Malaysia
Kinshuk, University of North Texas
Siu-Cheung Kong, The Education University of Hong Kong
Karey Yu-Ju Lan, National Taiwan Normal University
James Lester, North Carolina State University
Na Lina Li, Xi'an Jiaotong-Liverpool University
Calvin Liao, National Central University
Chiu-Pin Lin, National Tsing Hua University
Marcia Linn, University of California Berkeley
Chen-Chung Liu, National Central University
Lin Lin Lipsmeyer, Southern Methodist University
Chee-Kit Looi, The Education University of Hong Kong

Yu Lu, Beijing Normal University
Rose Luckin, University College London
Collin Lynch, North Carolina State University
Jon Mason, Charles Darwin University
Gordon McCalla, University of Saskatchewan
Ellen Meier, Columbia University
Shitanshu Mishra, UNESCO MGIEP
Tanja Mitrovic, University of Canterbury
Riichiro Mizoguchi, Japan Advanced Institute of Science and Technology
Kenji Morita, University of Tokyo
Sahana Murthy, Indian Institute of Technology Bombay
Thrishantha Nanayakkara, Imperial College London
Hiroaki Ogata, Kyoto University
Dimitri Ognibene, University of Milano-Bicocca
Jun Oshima, Shizuoka University
Maria Mercedes T. Rodrigo, Ateneo de Manila University
Jeremy Roschelle, Digital Promise
Marlene Scardamalia, University of Toronto
Junjie Shang, Peking University
Ju-Ling Shih, National Central University
James Slotta, University of Toronto
Hyo-Jeong So, Ewha Womans University
Yanjie Song, The Education University of Hong Kong
Masanori Sugimoto, Hokkaido University
Daner Sun, The Education University of Hong Kong
Davide Taibi, National Research Council of Italy
Patrice Torcivia Prusko, Harvard University
Chin-Chung Tsai, National Taiwan Normal University
Kurt VanLehn, Arizona State University
Julita Vassileva, University of Saskatchewan
Jayakrishnan Warriem, Indian Institute of Technology Madras
Lung Hsiang Wong, Nanyang Technological University
Su Luan Wong, Universiti Putra Malaysia
Longkai Wu, Central China Normal University
Ying-Tien Wu, National Central University
Stephen Yang, National Central University
Yin Nicole Yang, The Education University of Hong Kong
Fu-Yun Yu, National Cheng Kung University
Shengquan Yu, Beijing Normal University
Jianhua Zhao, Southern University of Science and Technology