

本文英文版是在联合国教科文组织(UNESCO)的 *Blue Dot* 杂志发表

通智同伴假说

全球和幸应成为人工智能发展的基础

全球和幸合作学者社群

(这篇文章源自一批学者共同提出的理念，并参考了 globalharwellgoal.org 上的相关资料，包括一篇选文与「全球和幸」的说明与常见问题。由于这个议题牵涉广泛，「全球和幸合作学者社群」除了原始参与者，也持续邀请国内外认同此理念的学者加入，名单会随着加入者而不断扩充。中文翻译比原文补充了一些细节。)

**历史有重量，让人站稳现在；
历史有高度，从过去看未来，让人看得更远**

关键词：通智同伴假说 (Generative Artificial Companion Hypothesis), 全球和幸 (Global Harwell), 无缝人工智能世界 (Seamless AI World), 机器人同伴世纪 (Robot Companion Age)

通智同伴

早在 1920 年代，Sidney Pressey 推出了他的「教学机器」——一种他声称将彻底改变教育方式的机械装置，这一创举比 1946 年第一台计算机的问世早了许多年。时至今日，一个世纪过去了，学生们使用的平板设备已经展现了实现 Pressey 愿景的可能性。1950 年，Alan Turing 提出了一种模拟游戏，旨在评估机器通过自然语言处理展示类人智能的能力，这一概念如今广为人知，即「图灵测试」。时光飞逝至 2022 年——在图灵提出该概念 70 多年后——ChatGPT 的推出，标志着生成式人工智能技术的重大突破。如今，大众越来越期待机器能够通过图灵测试。

目前，人工智能 (AI) 正加速全球教育的变革，为个性化学习、教师支持及教育机会的扩展提供了崭新的可能性。在美国，像 Carnegie Learning 这样的平台利用自我适应人工智能技术，识别学生的知识差距并量身定制教学内容，帮助学生以适合自己的步调掌握代数概念。AI 驱动的游戏化工具，如 Duolingo，不仅提升了用户的语言学习动机，还有效降低了辍学率。此外，Panorama Education 等平台透过提供学生学习进展的深入分析，使教育者能够更精准地制定教学策略。然而，数字鸿沟和隐私问题仍是当前面临的重大挑战。在欧洲，芬兰的 Elements of AI 计划积极推广 AI 素养，德国则着重于语言整合的支持。尽管如此，严格的监管框架和有限的教师培训资源仍阻碍了进一步的发展。在亚洲，虽然文化阻力依然存在，印度的 BYJU'S 成功缓解了教师短缺的问题，而中国则利用 AI 技术实现高效评分，提升了

教育效率。在非洲，肯尼亚的 Eneza Education 等计划显著提升了教育资源的可及性和技能发展，但基础设施薄弱和语言障碍仍限制了其影响力。

尽管人工智能在教育领域展现出巨大的变革潜力，但要实现可持续且包容的全球应用，仍需克服诸多挑战。这些挑战包括公平性、数据隐私、算法偏见、AI 驱动决策的透明度，以及 AI 资源的可取得性。唯有妥善解决这些问题，才能确保人工智能在全球教育中的广泛应用，真正实现其对教育变革的承诺。

人工智能的发展可被视为一条连续的光谱，从「工具」到「助手」，再进一步演进为「同伴」。在这段人机关系的光谱中，角色的转变并无绝对的界线，但我们仍可依据 AI 与人类互动的密切程度，对其进行初步的分类。

在最初阶段，AI 扮演的是「工具」的角色，主要由人类手动操作，用于辅助完成特定任务。此时的 AI 功能明确，且依赖人类的指令来执行工作。随着技术的进步，AI 逐渐进化为「助手」，能够主动执行人类指派的任务。例如，生成式 AI 在这一阶段扮演了重要角色，广泛应用于多个领域：协助开发程序代码、制作动画、根据自然语言查询检索信息、进行语言翻译，甚至提供法律咨询等服务。

当 AI 进一步发展为「同伴」，其与用户的互动模式也变得更加双向且主动——不仅能响应人类发起的对话或活动，甚至能主动引导互动。事实上，我们与随身装置（如智能型手机）的互动方式，正逐渐从传统的工具使用，转变为以自然语言驱动的交流模式。使得人机关系从单向操作转向为更具投入感与情感连结的同伴交流。

生成式人工智能的兴起，预示着见多识广及人情练达的人工同伴(artificial companions)即将到来。关于人工学习同伴的研究已有悠久的历史（这也是陈德怀教授 40 年前博士论文研究主题）。为更深入理解人工同伴的本质，让我们首先探讨「同伴」与「陪伴」的概念。

从广义的角度来看，「同伴」可以包括我们的父母、配偶、子女、朋友、教师、导师、教练、同学，甚至是宠物。「陪伴」则被定义为一种双向社会关系，其特征包括情感支持、共同活动、信任关系、相互尊重、开放沟通，以及共享时光的愉悦感。例如，在亲子关系中，角色通常是明确界定的，其核心目标是父母培育并教导子女，而子女则在父母的关怀下学习与成长。这种互动模式说明了同伴关系中的角色如何建构整体目标，并影响持续的互动过程。因此，「陪伴」可被理解为涵盖关系动态、共同目标与持续互动的历程。

我们可以将当前的时代描述为「无缝人工智能世界」(Seamless AI World, SAIW)。在这个高度互联的世界中，几乎所有事物都透过网络连接并由人工智能驱动，我们的地球似乎正在「收缩」或变得「更小」。在人工智能的帮助下，人们可以实时翻译，以母语无缝交流，地理距离就逐渐消失。

作为一个学习环境，SAIW 提供了一个「无缝学习」(seamless learning)体验的机会，使人们能够在教室、元宇宙和现实世界之间自由转换。无缝性(seamlessness)的一个关键面向是连续性(continuity)——即支持跨时间与空间的连续活动。然而，无缝性不仅仅关乎连续无缝性(continuity seamlessness)，还包括可及无缝性(accessibility seamlessness)以及类人无缝性(resemblance seamlessness)，特别是在人工同伴的支持下。

然而，无缝性的不同层面可能带来「正面无缝性」(well-seamlessness)的益处，也可能引发「负面无缝性」(ill-seamlessness)的风险。对于后者，必须谨慎评估，并视情况设立明确的范围界限与接口。若轻忽这些挑战，恐将造成无法挽回的社会伤害。随着人工同伴时代的到来，这一议题更显迫切。

「人工同伴」指的是具复杂智慧与社交情商的人工智能实体(entity)，如机器人或虚拟角色，其设计目的在于支持并增进人类的日常活动，同时促进人际互动与关系的建立。随着 AI、因特网、元宇宙，以及扩增实境(AR)、机器人技术、物联网 (IoT)、量子运算与区块链等技术的快速发展，SAIW 中的生活体验范围正迅速被重新定义。这些技术的融合将使人工同伴变得普遍，并成为日常生活中不可或缺的一部分。

自驾车与无人机的崛起，象征着机器人同伴时代(Robot Companion Age)的来临。未来，具备通用人工智能(Artificial General Intelligence, AGI)的通智同伴(General Artificial Companions)或许将成为现实。这类 AI 不仅拥有与人类相当，甚至超越人类的智能，更能展现多样化类人行为，包括伦理思考、情感智慧与社会意识。

如果这一愿景得以实现，数字仿真的真实性将可能进一步模糊真人与通智同伴之间的界线。随着时间推移，这些通智同伴将能累积并记录个人的一生经验，从而扮演不同人生阶段中的不同角色，成为个人成长与生活的陪伴者。

然而，随着通智同伴日渐普及，人们开始关注其在更广泛社会脉络中的角色与影响。这些关注既可能带来正面效益，也可能引发负面忧虑。当人类与通智同伴共存于同一社会体系时，整体社会目标将变得更加复杂，充满多重变量与挑战。在人机共存的未来中，如何平衡发展并建立伦理价值观，将成为我们必须深思的重要课题。

自第一台计算机诞生至今已 80 年，如今，拥有一台或多台计算设备，如智能型手机、平板计算机、桌面计算机，已逐渐成为常态。也许我们无需再等待 80 年，可能仅需 10 到 20 年，或者稍久一点，每个人便可能拥有一台或多台通用人工智能驱动的机器人同伴。

想象一下，未来的机器人同伴将无缝融入我们的生活：一台机器人可留在家中打理家务，另一台则在职场或学校中辅助用户，并与其数字分身实时链接与合作。当前全球人口约为 80 亿，若将机器人同伴纳入计算，则「总人口数」可能翻倍至 160 亿，这将彻底改变我们对社会结构与资源分配的认知。

在这个即将到来的机器人同伴时代，社会结构、工作环境，以及人与 AI、人与人之间的关系都将迎来深刻的改变。

全球和幸(Global Harwell)

如果我们暂时撇开宗教教义、政治意识形态与特定信仰，单纯探讨个人对人生的期许，许多人可能会提到幸福、健康、财富与成就等目标。这些追求往往与马斯洛的需求层次理论(Maslow, 1943)——包括生理需求、安全需求、爱与归属感、尊重需求及自我实现——以及塞利格曼的 PERMA 模型(Seligman, 2011)——包括正向情绪、投入感、良好人际关系、人生意义与成就——不谋而合。在现实世界也需要考虑健康与财务安全。无论如何，这些理论框架为我们理解美好生活提供了重要的参考。

若将幸福(wellbeing)视为一个整体概念，涵盖生理、心理与社会健康，并体现幸福感、生活满意度以及有效运作的的能力，那么马斯洛与塞利格曼的理论便与这种对幸福的理解高度契合。

然而，当今世界正面临前所未有的挑战：COVID-19 造成的大量死亡、气候变迁的加剧、资源枯竭、环境污染、财富不均，再如网络犯罪、深度伪造(deepfake)与假信息泛滥，以及 AI 技术带来的风险。这些问题不仅加剧了社会分歧，还引发了全球范围军事与经济的冲突升级，甚至让人们对于核灾难的恐惧与日俱增。

在这样的背景下，我们不得不思考一个根本问题：如果人类无法在这颗地球上持续生存，幸福又从何谈起？

仅仅追求个人幸福并不足够，我们必须将「和谐」置于更高的优先位置，积极促进人与人之间、以及人与环境之间的正向关系。和谐可分为两个层面：「环境和谐」与「人类和谐」。

环境和谐涵盖全球暖化与自然灾害等议题，强调人类应与自然共生共存，尊重生态平衡，并致力于永续发展。而人类和谐则涉及个人、家庭、社会与全球各层面的互动与平衡，旨在创造一个彼此尊重、合作共赢的世界。

人类和谐有四大关键要素：仁爱、公平、正义与中庸之道。

- 仁爱(Benevolence): 是和谐的基础，仁慈、善良与同情心，是人类彼此关怀与支持的基础。
- 公平(Equity): 确保资源依需求分配，并促进机会均等，让每个人都能在合理的环境中追求成长。
- 正义(Justice): 维护伦理道德与诚信正直，是社会稳定与信任的基石，确保每个人的权利与尊严得到保障。
- 中庸之道(The Middle Way/The Golden Mean): 强调平衡与适度，避免极端，以维持长久和谐。

简言之，仁爱奠定和谐基础；公平促进资源均衡分配，正义维护道德伦理，二者共同维持社会稳定，而中庸则防止极端，以守护和谐。

最终，人类和谐体现了儒家、道家与墨家的精神，也是西方古典哲学推崇的「黄金法则(The Golden Rule)」，它的意思就是「己所不欲，勿施于人」，或以正向表达：「待人如待己」(Treat others as you wish to be treated)。这条法则不仅是人际关系的指南，更是人类和谐核心理念。

「全球和幸」(Global Harwell)结合「和谐」(harmony)与「幸福」(wellbeing)的概念，代表影响全球的社会规范、伦理的核心价值与集体原则。它不仅是一个理念，更是一种全球性的愿景，旨在将追求和谐与幸福视为人类的基本价值观。

全球和幸的愿景与「联合国教科文组织」(UNESCO)及各国教育部的教育理念与目标高度契合。它提供了一个全面的全球视野，将和谐与幸福视为实现全球和平与永续发展的关键要素。例如，联合国的「永续发展目标」(Sustainable Development Goals)便与和谐的概念紧密相连，强调环境保护、社会公平与经济平衡。尽管不同国家与地区可能因其文化、政治或地理环境而设定不同的发展目标，但全球和幸所承载的是一种超越国界差异、建基于人类共同愿景的普世价值。我们相信，人类知识与科技技术发展的首要目标，就是实现全球和幸。

为了实现全球和幸的愿景，研究人员正从多个方向积极探索与推动。其中，一组研究人员致力于开发教育虚拟同伴，提升社交媒体的使用体验，并对可能危及个人与群体幸福的风险提出因应之道。这些虚拟同伴不仅能提供个性化的学习支持，还能帮助学生建立健康的数字习惯，促进正向的社交互动。

与此同时，另一组研究人员则专注于探索全球教育、全球大脑、全球伦理以及全球人工同伴在促进社会和谐与幸福中所扮演的角色。这些研究试图通过整合全球资源与智慧，建立一个更加互联、包容且伦理导向的社会体系，从而推动全球和幸的实现。

无论最终结果如何，越多的人将全球和幸视为其核心人生价值，越少的人或机构会为私利所驱动，也越减少无意间开发出对人类有害的人工智能的风险。至少，这将在社会层面对 AI 的恶意滥用施加道德与行为约束，从而降低其可能带来的负面影响。

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

「通智同伴假说」提出，随着人工同伴技术的不断演进，关注人与机器人之间的陪伴关系及其对社会发展的影响将成为关键课题。若能优先发展符合伦理原则的人工智能，不仅能促进人类社群的和谐共存，还能透过创新与持续进步来提升个体幸福与社会和谐。在通用人工智能的支持下，这些机器人同伴将引领我们朝向全球和幸的愿景迈进(Chou et al., 2025)。

在 SAIW 环境中，「通智同伴假说」提出以下假设：

对于将全球和幸培育为个人的核心价值过程中，通智同伴所提供的支持将发挥关键作用：贡献 50%，甚至更多，并且可从个人日常行为中得到验证。

当假说被证实的那一天，象征着科技辉煌成就的凯旋，也标志着人类文明跨入新的里程碑。

然而，这也引发了一个深刻的问题：究竟是人类塑造人类，还是 AI 塑造人类？如果答案是人类塑造人类，那么全球和幸就必须成为人工智能发展的根本基础。

我们必须清楚认知，通智同伴的发展，必须奠基于教育。真正的通用智慧同伴，只能由「全球和幸人」(Global Harwellians)——那些将全球和幸视为核心价值与人生目标的人——来设计。而这些设计者，也唯有透过教育，才能培养出来。

一旦成功建立通智同伴，「全球和幸性」(Global Harwellness)——即体现全球和幸状态的整体素质(quality)，将在人类与通智同伴的陪伴关系中相互深化与提升。届时，人类与 AI 将携手共塑世界，创造一个更加和谐、幸福且永续的未来。

正如 纳尔逊·曼德拉(Nelson Mandela) 所说：教育是你用来改变世界的最强大武器。的确，教育确实能够引领人类追求全球和幸，并让我们预见根本性的变革，在教育也在全球社会，在今天也在未来。教育，不仅能培养出具有全球视野与伦理意识的「全球和幸人」，还能为通智同伴的发展奠定坚实的基础，最终实现人类与 AI 共融共荣的愿景。

结论

教育研究是一趟探索与揭示人类本能、性格发展、身分认同与生存意义的旅程。但唯有当我们的研究不仅对学生具有重要性，更能在更广泛层面上造福整体人类，这趟旅程才具有真正意义。

本文阐述了通智同伴假说，提出一个未来愿景：具备人类特征的先进人工同伴，将帮助人类把全球和幸培育为核心价值与人生目标。

虽然这一假说带来许多问题与挑战，它同时也为 AI 研究者与实践者提供了一个共同的愿景：建构一个持久且和谐的世界，确保未来世代能够在全球和幸的未来中，持续成长。实现这一愿景的关键，在于教育——唯有透过教育，世界才能真正培养出将全球和幸视为核心价值的「全球和幸人」(Harwellian)。而开发通智同伴以实现这一假说，则是这项使命不可或缺的一环。

人工智能已成为当今最具颠覆性的科技之一。一方面，如前所述，它为教育插上新的翅膀——不仅能促进个人化学习，协助教师预测学生表现并提供精准引导，更能提升学习动机与学习表现。此外，亦可透过实时翻译消弭语言障碍，进而促进全球社群建构。

另一方面，随着 AI 的光芒愈加耀眼，无缝人工智能世界中的负面无缝性却带来诸多风险。它可能使我们下一代与社会脱节，成为焦虑的一代，并且加剧隐私侵犯与教育资源分配不均等问题。更值得警惕的是，若由通用人工智能驱动的人工同伴采纳的价值观，与全球和幸的核心价值背道而驰，则将对人类未来造成巨大灾难。我们必须审慎思考未来发展方向，确保教育始终扎根于人类价值观，而非被科技异化。这是当今世界的核心问题。

全球和幸是一种理想、一套行动指南，它提醒我们，科技、经济与社会变革都应服膺于人类的和谐与幸福。唯有如此，才能在科技飞速发展中坚守人性价值，创造真正和谐与幸福的未来。

「通智同伴假说」或可被视为人工智能领域的「圣杯」——那个长久以来让 AI 研究孜孜以求，极为艰难欲所企及的终极目标。若此假说得以验证，它将有可能引领人类朝向康德 (Immanuel Kant) 于 1795 年所描绘的永久和平愿景迈进，并透过世代累积培育出一个「全球和幸」的世界，使和平与幸福得以延续至未来。

只要人类怀抱崇高的梦想，梦想终有一日成真！

但是，我们不能被动等待未来；如果我们不行动，未来不会变得更好，只会变得更坏。现在就行动！

参考文献

- 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.

全球和幸福合作学者社群

Vincent Alevan, 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 Kashihara, 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