

中小學使用「生成式人工智慧」注意事項2.1

（學生版）

中華民國113年7月1日臺教資（三）字第1132702614號函核定

中華民國114年12月23日臺教資（一）字第1142704025號函核定

中華民國115年2月11日臺教資（一）字第1152700391號函核定

近年來，「生成式人工智慧」（Generative AI，簡稱GenAI）工具迅速發展，為社會帶來許多新的機會和改變；同時，「深偽技術」（Deepfake）也逐漸成熟，這些技術改變了我們原本獲取知識、傳播訊息和創造內容的方式，但也增加了許多使用上的風險。為了幫助學生提升「生成式人工智慧」工具使用素養，理解使用的倫理原則，善用相關技術並避免造成誤用或濫用，同時能在家長或學校老師的指導與同意下遵守各工具年齡分級與帳號申請規範，提供以下注意事項作為使用參考。

一、覺察「生成式人工智慧」產生的內容可能存在偏誤

「生成式人工智慧」工具的資料來源是歷史紀錄，如果這些資料本身具有成見或錯誤，那麼使用「生成式人工智慧」工具的結果也會存在偏差或錯誤，因為這些工具無法自行判斷結果的正確性和合理性。

所以我們在使用「生成式人工智慧」工具時，應該仔細檢視內容，並以批判性思考判斷其是否公平、全面或具有偏見，遇有疑慮應向老師或家長詢問，不宜直接採信或用於作業、報告或公開發表的內容。

二、檢核內容並結合多元管道查證

「生成式人工智慧」工具的資料會受到其來源的影響。如果資料不夠多元、廣泛，那這些工具產生的結果可能只會呈現單一文化的知識，造成知識量嚴重性不足，不僅正確性堪慮，甚至會讓人產生偏見。

所以當我們在使用「生成式人工智慧」工具時，不能全盤接受生成的內容，而需結合其他可靠來源的知識和批判性思維進行查證，並留意其中是否涉及文化刻板印象或不公平的描述，如對內容仍有疑慮，應向老師或家長請教，不可以直接將生成內容作為作業或報告的原始答案。

三、發現「深偽技術」會產生不實的內容

Deepfake是能修改臉部影像的「深度偽造技術」，原理是使用「生成式人工智慧」創建虛假的內容。這項技術能利用既有的圖片、影像或聲音素材，製造出看似真實的影片和圖像，甚至假新聞。

所以當我們在觀看網路內容時，不要輕易相信未經審核的影片或照片，並留意這些內容是否由深偽技術合成，和判斷可能的目的與動機，也要尊重他人的肖像權與隱私，不轉傳、下載或散布可能由深偽技術合成的影片或圖片。

此外，如果不幸發現自己或他人遭到深偽技術（如AI換臉）製作不當的性影像，請不要害怕，這不是你的錯。請記得先截圖保留證據，並向「衛生福利部性影像處理中心」線上申訴（<https://siarc.mohw.gov.tw>），由專業單位協助移除下架，切勿轉傳造成二度傷害。

四、養成保護個資與尊重隱私的習慣

部分「生成式人工智慧」工具的資料在取得、儲存和使用上都還沒有完備的法令、規範及倫理上的監管機制。因此，在使用這些工具時，提供的個人資料、敏感訊息及機密數據，都可能會被收錄到訓練資料庫中，作為未來回應他人的內容。

所以我們在使用「生成式人工智慧」工具時，應該審慎評估提供的資訊，是否具有機密性、隱私性與敏感性，以保護個人與組織的隱私與機密。

五、遵守「生成式人工智慧」服務使用規範

中小學學生使用「生成式人工智慧」工具應遵守各平臺註冊年齡限制及相關規範。

建議中小學學生使用為教育目的而設計的「生成式人工智慧」工具，如教育部因材網生成式AI學習夥伴e度、酷英網E-BOT、均一教育平臺或CK-12 Flexi等生成式的教育工具，在校應於教師引導或指導下使用，若為非在校使用，亦請家長陪伴使用。

本注意事項建議使用上述生成式的教育工具，主因其具備三大守護功能：它是能引導思考而非直接給答案的「智慧導師」（Cardoso-Silva, J., 2024; OECD, 2026; Sal Khan, 2024）；它是「安全過濾網」，能阻隔偏誤訊息並提供可靠教材（MIT

Sloan Teaching & Learning Technologies, n.d.)；它更是保護個資並記錄學習歷程的「隱形盾牌」(UNESCO, 2023)。考量科技變動快速，本注意事項僅列出本部建置或具公益性之平臺；若需使用其他商用版的「生成式人工智慧」工具，應經由學校評估資安風險，並在師長指導下審慎使用。

六、遵守倫理與誠信使用原則

- (一) 生成式人工智慧是協助學習的工具，不能取代自己的思考與判斷，應先自行閱讀與思考，再視需要使用工具輔助。
- (二) 使用「生成式人工智慧」工具時，應尊重他人的權利與尊嚴，不得輸入或產生含有歧視、霸凌、仇恨或不尊重他人之內容。
- (三) 不可以將自己或他人的姓名、照片、聯絡方式、住址、學號等個人資料，或家庭、學校的機密資訊輸入至「生成式人工智慧」工具。
- (四) 完成作業或報告時，如依老師或學校規定使用「生成式人工智慧」工具協助構思或潤飾，應依規定註明所使用的工具及用途，不得將生成內容直接當作自己原創作品繳交。
- (五) 遵守學校訂定之資訊倫理與學術誠信規範，不得使用「生成式人工智慧」工具從事抄襲、代寫、作弊或其他違反校規與法令之行為。

所以我們在使用「生成式人工智慧」工具時，應該審慎評估提供的資訊，是否具有機密性、隱私性與敏感性，以保護個人與組織的隱私與機密。

生成式人工智慧技術的發展，為我們的生活帶來了許多便利，並廣泛運用在各種情境中，卻也伴隨著一定的風險和挑戰。在這個數位時代，我們應該：

- 1. 保持對資訊來源保持高度警覺。
- 2. 不要輕易相信未經證實的訊息。
- 3. 學會如何辨別虛假資訊。

同時，我們要提升自己思辨的能力：

- 1. 批判性地分析和評估生成式人工智慧工具所產生的內容，避免被誤導，並留意其中是否包含偏見或文化刻板印象。
- 2. 遵守相關的道德和法律規範（例如：尊重智慧財產權），確保使用「生成式人工智慧」工具時不違反社會常規與資訊倫理，

並遵守學術誠信，不將生成內容直接作為個人原創作業。

最後，我們應該加強自己的數位素養能力，才能在享受科技進步帶來高度便利的同時，減少科技帶來的風險，讓負面影響降到最小，並以負責任的態度善用「生成式人工智慧」工具。

附錄1

中小學使用「生成式人工智慧」注意事項2.1（學生版）示例

- 一、**覺察AI產生的內容可能存在偏誤：**當我們要求「生成式人工智慧」建議旅遊行程時，如果這些工具沒有當地的氣候環境、地理位置、社會人文及文化限制等資料，答案可能會來自於各種網路遊記文章的綜合體，也可能提供你一份不順路、非當季的活動行程，甚至還可能是虛構景點的行程。因此，當AI給出的資訊看起來奇怪、不合常理或與課本不同時，我們要停下來檢查內容，不要直接用在作業或報告上，並主動向老師或家長確認資訊是否正確。
- 二、**檢核內容並結合多元管道查證：**當我們向「生成式人工智慧」詢問法律或文化問題時，這些工具可能會只根據研發者自己國家的法律和文化習俗產生答案。比如當我們要求「生成式人工智慧」生成一張新娘圖片時，它可能只會產生一張穿著白紗禮服的西方臉孔女性，而不是根據使用者當地的文化習俗來產生不同膚色或其他婚禮的服飾。因此，當AI的答案只呈現某一種文化、觀點或價值時，我們要記得它可能忽略其他觀點，需要自己多方查證，避免被單一視角誤導。
- 三、**發現「深偽技術」會產生不實的內容：**網路上常有知名人士發表演說或鼓勵投資的影片，面對這些內容，我們必須謹慎且小心求證知識的內容和來源。在深偽技術蓬勃發展的網路環境中，這些影片可能未取得影片主角的同意，或在他們根本不知情的狀況下，被深偽技術整合他們的臉（聲音）到假圖片或假影片中。因此，看到不尋常或煽動性的影片時，不要急著相信或轉傳，應先詢問老師或家長，也不要將未查證的影片當成作業或報告內容。
- 四、**養成保護個資與尊重隱私的習慣：**當我們在不清楚「生成式人工智慧」的原理及規範的情況下，以個人或學業資料為題材向它詢問答案，可能導致個人的機密被收錄到它的資料庫中。而當其他的使用者再度詢問類似問題時，「生成式人工智慧」以收錄的資料庫回答問題，就有機會造成個人隱私或學業資料外洩，形成資安漏洞。因此，我們不把自己的姓名、住址、電話、照片、學校資料或家庭狀況輸入給「生成式人工智慧」工

具，避免讓敏感資訊被記錄或外洩。

五、遵守「生成式人工智慧」服務使用規範：我們可以使用「教育體系身分認證服務（OpenID）」登入為教育目的而設計的「生成式人工智慧」工具，確保學習效果並保護自己的隱私；在校時，我們跟隨老師的任務引導，在家則將AI作為輔助思考的夥伴。若需使用其他商用版的「生成式人工智慧」工具（如ChatGPT、Gemini等），應經由學校評估資安風險，並在師長指導下審慎使用。

六、遵守倫理與誠信使用原則：有些同學可能會把整篇生成式人工智慧工具生成的內容直接交出去當作自己的報告，或請生成式人工智慧工具寫出一整篇作文，這樣不但不代表自己的學習成果，也可能侵犯智慧財產權或違反學校的學術誠信規範。因此，我們使用生成式人工智慧工具時應該先自己思考，再把生成式人工智慧工具當作輔助工具；若老師規定可以使用生成式人工智慧工具，就要按照要求註明使用方式，不可以把生成出來的內容當自己的原創作品，更不能用來作弊、抄襲或代寫作業。

參考文獻

- Cardoso-Silva, J. (2024, July 11). *Book review: Brave new words: How AI will revolutionize education (and why that's a good thing)* by Sal Khan. *LSE Impact Blog*. <https://blogs.lse.ac.uk/impactofsocialsciences/2024/07/11/brave-new-words-how-ai-will-revolutionize-education-review/>
- Khan, S. (2024). *Brave new words: How AI will revolutionize education (and why that's a good thing)*. Viking.
- MIT Sloan Teaching & Learning Technologies. (n.d.). *When AI Gets It Wrong: Addressing AI hallucinations and bias*.
<https://mitsloanedtech.mit.edu/ai/basics/addressing-ai-hallucinations-and-bias/>
- OECD (2026), *OECD Digital Education Outlook 2026: Exploring Effective Uses of Generative AI in Education*, OECD Publishing, Paris,
<https://doi.org/10.1787/062a7394-en>.
https://www.oecd.org/content/dam/oecd/en/publications/reports/2026/01/oecd-digital-education-outlook-2026_940e0dd8/062a7394-en.pdf
- United Nations Educational, Scientific and Cultural Organization. (2023). *Guidance for generative AI in education and research*. UNESCO.
<https://unesdoc.unesco.org/ark:/48223/pf0000386693>