

中小學使用「生成式人工智慧」注意事項(學生版)

中華民國 113 年 7 月 1 日臺教資(三)字第 1132702614 號函核定

近年來，「生成式人工智慧」和「深偽技術」(Deepfake)的迅速發展，為社會帶來許多新的機會和改變，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，同時也增加了許多使用的風險。為了幫助學生提升「生成式人工智慧」工具使用素養，以善用相關技術並避免造成誤用或濫用，提供以下四個要點作為使用參考。

一、可能會產生偏誤的內容：

「生成式人工智慧」工具的資料來源是歷史紀錄或舊經驗，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。

二、可能會減少訊息的多樣性：

「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。

所以我們在使用「生成式人工智慧」工具時，需要結合自己的經驗和批判性思維來檢視結果，而不是全盤接受生成的內容。

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。

四、可能造成個資、隱私和機密的洩漏：

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。

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的高度警覺。
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。
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。

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。

中小學使用「生成式人工智慧」注意事項(學生版)

示例

- 一、**可能會產生偏誤的內容**：例如，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。
- 二、**可能會減少訊息的多樣性**：例如，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。比如當你要求「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。
- 三、**發現深偽技術會產生不實的內容**：例如，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。
- 四、**可能造成個資、隱私和機密的洩漏**：例如，當我們不清楚「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。