
長江實業集團有限公司

生成式人工智能負責任使用政策

長江實業 生成式人工智能負責任使用政策

本生成式人工智能負責任使用政策（「**本政策**」）由長江實業集團有限公司（「**本公司**」）及其附屬公司（統稱「**本集團**」）制定。

1. 引言

- (a) 生成式人工智能（「**生成式人工智能**」）指因應使用者提示、指示或其他輸入，生成、分析、轉換或以其他方式生成文字、圖像、音訊、影片、程式碼、數據、建議或其他輸出內容之人工智能系統、模型、工具及應用程式。就本政策而言，生成式人工智能亦包括可代表使用者檢索資料、生成內容、執行任務或發起行動的人工智能助理及人工智能代理。
- (b) 本集團認同生成式人工智能具備提升生產力、促進創新及提高營運效率之潛力，並支援輸出高質素工作成果。隨着生成式人工智能技術持續演進，適當的監管、人手監察及控制措施，生成式人工智能以安全、合乎道德、負責任及合法的方式使用至關重要。
- (c) 本政策旨在促進僱員為工作相關目的或於本集團提供的裝置以負責任方式使用生成式人工智能工具，同時管理相關風險。倘本政策所載之任何條文與任何該等法例、法規或規例不一致或相抵觸，則有關之不一致或相抵觸之處概以該等法例、法規或規例為準。
- (d) 僱員僅可使用已獲資訊科技部不時批准並在本公司內聯網所列的生成式人工智能工具（「**認可生成式人工智能工具**」）進行工作相關目的或用於本集團提供的裝置。除非事先取得資訊科技部書面許可，否則不得就工作相關目的或於本集團提供的裝置使用其他生成式人工智能工具。
- (e) 本公司所有僱員（不論是全職或兼職、合約或臨時職員或實習生）均須遵守本政策。本公司對其擁有管理控制權的海外附屬公司及合營企業均須採納並維持生成式人工智能使用政策，以遵守其業務經營所在司法權區的法例、法規及規例，並與本政策所載原則一致。本公司亦可要求其顧問、借調人員、顧問及其他可存取本集團資料或系統的人士遵守本政策。
- (f) 本政策必須與本公司其他相關政策及程序一併閱讀。與本政策有關之任何技術問題應交予資訊科技部（aisupport@ckah.com）。與本政策有關的任何法律問題，應交予公司秘書處（ailegal@ckah.com）。

長江實業 生成式人工智能負責任使用政策

- (g) 本公司將不時審視本政策，以應對涉及生成式人工智能工具的新發展及潛在風險，並確保本政策已反映適用法例、法規及規例中之任何變動。本公司保留不時更改本政策之權利。本政策的任何修訂將透過電郵或本公司可能不時決定之其他方式通知僱員。

2. 指導原則

本集團鼓勵以負責任方式使用生成式人工智能工具，以提升生產力、促進創新及提高營運效率。僱員在使用生成式人工智能工具時，必須遵守以下原則：

- (a) 遵守本政策、所有適用之法例、法規及規例，以及本集團其他適用政策；
- (b) 以安全、合乎道德、負責任及合法方式使用生成式人工智能工具；
- (c) 保護本集團的機密資料、知識產權及個人資料；
- (d) 運用適當的人為判斷，認識到生成式人工智能旨在協助而非取代人類決策，並於依賴、採納或使用人工智能生成的輸出內容之前，獨立核實其準確性及適當性；
- (e) 確保於實施任何人工智能生成的建議或行動（包括人工智能助理或人工智能代理所執行的任何行動）前，均會經適當的人手審閱及批准；
- (f) 確保生成式人工智能及人工智能生成輸出內容的使用屬公平、客觀，且不存在任何非法歧視或偏見，尤其是在可能影響個人的情況下；
- (g) 確保人工智能助理及人工智能代理的使用受到適當的人手監察及控制，而該等監察及控制應與其使用性質及相關風險相稱；及
- (h) 對所有在生成式人工智能工具協助下完成的工作承擔責任及問責。

3. 僱員使用認可生成式人工智能工具

- (a) 僱員可在符合本政策情況下就工作相關目的或於本集團提供的裝置使用認可生成式人工智能工具，當中包括就行政工作、研究、初稿或設計、編輯文件、摘要、交流討論、提出新構想及其他合法業務活動而生成內容。如有疑問，僱員應與相關業務單位/部門主管討論其擬使用認可生成式人工智能工具的適用範圍，業務單位/部門主管有權批准、拒絕或修改擬定之用途。
- (b) 僱員不得使用生成式人工智能代替法律、監管、稅務、會計或其他專業意見，亦不得在缺乏適當監督的情況下，利用生成式人工智能進行超出其職責、權限或專業範疇的工作。就業務目的而進行的任何法律、監管、稅務、會計或其他專業研究或分析，應由法律部、公司秘書處或其他相關部門（視情況而定）負責或提供諮詢。未獲授權的僱員不得代表本集團使用生成式人工智能擬備或提供法律意見或法律建議。

長江實業 生成式人工智能負責任使用政策

- (c) 認可生成式人工智能工具旨在協助僱員而非取代其專業判斷。僱員在依賴、採取行動或使用人工智能生成的輸出內容前，必須先行審閱、驗證並作出獨立判斷，並在適當情況下，根據可靠資料來源核實人工智能生成輸出內容的準確性。人工智能生成的輸出內容不得作為財務、法律、監管、合規或其他重大業務決策的唯一依據。僱員不得容許生成式人工智能（包括人工智能助理及人工智能代理）在未經適當人手審閱及取得所需批准的情況下，代表本集團作出決策或採取行動。任何與人工智能生成輸出內容相關的疑慮，均可透過本集團既定的匯報及上報程序提出。
- (d) 僱員必須遵守資訊科技部就認可生成式人工智能工具的配置及使用所訂明的一切規定，包括任何有關資訊安全、私隱保障及數據管治的要求。僱員不得規避、繞過或停用資訊科技部就認可生成式人工智能工具所實施的任何保安或其他管控措施。資訊科技部或會不時就特定認可生成式人工智能工具、使用特定生成式人工智能工具或緊急安全風險發出額外指引或限制措施。僱員必須遵守該等指引或限制措施。僱員必須遵守該等指引或限制措施，並時刻對冒充現有生成式人工智能之偽冒應用程式、釣魚網站或其他惡意服務保持警覺。
- (e) 針對任何完全或部分透過使用認可生成式人工智能工具產生之內部工作成果，僱員均應考慮並顧及有關工作成果的性質、目的及目標對象，是否適合就使用該等認可生成式人工智能工具加入確認陳述，作為保持透明度的良好實務。如為遵守適用法律、監管規定或合約要求或本集團其他適用政策而有此需要，僱員必須確保清晰標示任何由人工智能生成之內容或經人工智能協助流程產生之成果。
- (f) 僱員必須確保任何全部或部分透過使用認可生成式人工智能工具而產生之外部通訊或工作成果，均符合適用法律、監管規定或合約要求，或本集團其他適用政策。

4. 使用生成式人工智能的相關潛在風險以及緩解風險措施

就工作相關目的或於本集團提供的裝置使用認可生成式人工智能工具前，僱員應確保已對與之相關的功能及潛在風險有所理解。下表載列與使用生成式人工智能相關之若干（但非全部）潛在風險，以及僱員在使用認可生成式人工智能工具時必需遵守緩解措施以儘量減低該等風險。人手監察仍是主要緩解措施。

潛在風險	緩解措施
違反保密責任及法律特權：提供予生成式人工智能工具的資料，可能在超出本集團控制的情況下被保留、處理或披露。將本公司、本集團、或本公司或本集團之客戶或交易對手的機密或受法律特權保障的資料輸入至生成式人工智能工具或會令其失去保密性或法律特權。該等資料（包括提示詞、聊天記錄及人工智能生成的輸出內容）亦可能在法律訴訟、監管調查或其他法律程序中披露。此舉或會違反閣下、本公司或本集團保護機密或受法律特權保障的資料之責任，或會使本公司及/或本集團面臨重大的法律、監管、商業及聲譽風險。	除根據本集團適用政策及經資訊科技部明確批准者外，僱員不得於任何生成式人工智能工具中輸入、上載或以其他方式披露本公司、本集團、或本公司或本集團之客戶或交易對手的任何機密或受法律特權保障的資料，包括將該等資料複製及貼上、將載有該等資料之文件或檔案上傳或將該等資料以其他方式提供至生成式人工智能工具。僱員必須遵守所有資訊科技部有關使用認可生成式人工智能工具及處理資料的規定及指引。 建議各部門將載有機密資料之文件分類並加密，以防在該等文件不慎上載至生成式人工智能工具時提供另一重防護。任何有關加密的技術性問題應向資訊科技部提出（aisupport@ckah.com）。 有關何謂機密資料以及處理機密資料方式之進一步指引，請參閱「僱員手冊」、「僱員行為守則」、「處理機密資料、消息披露，以及買賣證券之政策」及「資訊安全政策」。

潛在風險	緩解措施
個人資料及侵犯私隱：在生成式人工智能工具中輸入有關僱員、客戶或其他各方的個人資料（即與在世的個人有關並可用作辨識其身分的資料）或會發放該等資料至公共領域並構成私隱風險，或會導致違反適用之資料保障法例、法規及規例而對本公司及/或本集團構成風險及責任。	僱員不得在生成式人工智能工具中輸入任何有關僱員、客戶或其他各方之個人資料。 建議各部門將載有個人資料之文件分類並加密，以防在該等文件不慎上載至生成式人工智能工具時提供另一重防護。任何有關加密的技術性問題應向資訊科技部提出（aisupport@ckah.com）。 有關何謂個人資料以及處理個人資料方式之進一步指引，請參閱「僱員行為守則」。
侵犯知識產權：生成式人工智能工具所製作的輸出內容可能包含受知識產權保護或該等作品之衍生作品。倘本公司及/或本集團在未經知識產權持有人同意下使用該等輸出內容，或會面對法律責任。	僱員必須盡力運用個人才識、判斷力及努力製作原創作品，以及必須將所有使用生成式人工智能工具產生的輸出內容視為初步成果而非最終完成品。 僱員不得假設人工智能生成的輸出內容不受第三方知識產權保障，並必須確保其使用該等人工智能生成的輸出內容時不會侵犯該等權利，包括在未取得適當授權的情況下複製或使用受第三方版權或商標保障的材料。 有關知識產權之進一步指引，請參閱「僱員手冊」。

長江實業 生成式人工智能負責任使用政策

潛在風險	緩解措施
不準確、具誤導性或虛構的輸出內容： 生成式人工智能或會以具說服力的方式生成不準確、不完整、過時、不可靠或虛構的資料，包括虛假的事實、案例、法律依據、引文、參考資料或資料來源。僱員可能不慎依賴該等輸出內容，從而導致錯誤的業務決策、違反法律或監管規定、財務損失、聲譽受損，或對本集團造成其他風險。	僱員必須將所有人工智能生成的輸出內容視為初步成果而非最終完成品。僱員必須在就業務用途依賴或使用該等生成的輸出內容前審閱、校對並獨立核實該等成果，確保其實質內容準確、可靠及完整。在適當情況下，僱員在使用或依賴該等輸出內容前，應諮詢相關業務單位或部門。
偏見及歧視： 生成式人工智能或會產生帶有偏見、歧視成分有害或令人厭惡的輸出內容，尤其以在該等輸出內容被用作評估個人或就個人作出建議時為然。	倘生成式人工智能工具所產生的成果含有偏見、歧視成分、有害或令人厭惡，或不符合本公司的宗旨、文化及價值，僱員一概不得使用任何該等輸出內容。

5. 禁止使用

僱員就工作相關目的或於本集團提供的裝置使用生成式人工智能，不得以有可能損害本公司或本集團安全或聲譽的方式進行，亦不得採取不符合本政策之方式進行。尤其是：

- (a) 僱員不得使用生成式人工智能從事侵犯他人私隱的活動，或試圖製作或分享可能侵犯他人私隱的內容，包括披露個人資料或試圖使用生成式人工智能進行人面識別或身分驗證，或輸入在未經他人同意而製作涉及他人的相片或錄影/錄音用作處理他人的生物特徵或生物辨識資料。
- (b) 僱員不得使用生成式人工智能侵犯他人權利或試圖使用生成式人工智能侵犯他人的法律權利（包括知識產權）。
- (c) 僱員不得使用生成式人工智能進行任何違法行為。僱員不得試圖繞過本集團為防止非法及/或不當活動而設立的封鎖措施。

長江實業 生成式人工智能負責任使用政策

- (d) 僱員不得使用生成式人工智能從事欺詐、虛假或具有誤導成分，或不得從事旨在利用個人或群體弱點的操縱性活動以促進社交信用評分，或製作或分享可能產生誤導或欺騙他人的內容，包括虛假資訊、具欺詐性或欺騙性的假冒內容。僱員不得惡意或以不道德方式使用生成式人工智能，例如製作仿真偽冒可用作散播虛假資料或操縱公眾輿論的內容。
- (e) 僱員不得使用生成式人工智能從事損害他人的活動，或試圖製作或分享有可能被用作騷擾、欺凌、虐待、威脅或恐嚇他人或以其他方式損害他人的內容。
- (f) 僱員不得使用生成式人工智能製作或分享不恰當的內容或材料，或令人感到不安或厭惡的內容。

6. 侵犯版權及違反本政策

如發現任何就工作相關目的或於本集團提供的裝置使用認可生成式人工智能工具引致懷疑或實際侵犯版權或版權申索，以及如對本政策有任何疑慮、或發現懷疑或實際違反本政策之情況，僱員必需立即向資訊科技部（aisupport@ckah.com）及公司秘書處（ailegal@ckah.com）作出報告。

7. 培訓

資訊科技部將統籌相關培訓及/或於本公司內聯網上發佈有關資料，以助僱員了解如何使用認可生成式人工智能工具。

8. 違反本政策

本公司以嚴肅態度貫徹推行本政策。如有任何僱員違反本政策，本公司將於考慮個案之性質及情況、嚴重性及影響後，加以適當處理，並強調防止日後再有違規行為。

9. 對本政策之責任

本政策已經董事會審閱、批准及採納。本政策之實施與執行由審核委員會監督及監察，審核委員會並可不時向董事會提出修訂建議以尋求批准。

生效日期: 2024年2月
更新日期: 2026年8月