

人工智慧政策 AI Policy

一、目的 Purpose

1.1 本人工智慧（「AI」）政策旨在確立研華內部在開發、部署及使用 AI 系統時的倫理以及法規準則。本政策旨在確保我們的 AI 計畫符合核心價值、促進透明度，並維護利害關係人對我們的信任。

1.1 This Artificial Intelligence ("AI") Policy establishes guidelines for the ethical development, deployment, and use of AI Systems within Advantech. It aims to ensure that our AI initiatives align with our core values, promote transparency, and uphold the trust of our stakeholders.

二、範圍 Scope

2.1 本政策適用於所有參與研華內部 AI 系統開發、部署或使用的員工及承包商。

2.1 This policy applies to all employees and contractors involved in the development, deployment, or use of AI Systems at Advantech.

2.1.1 在本 AI 政策中，「研華」係指研華股份有限公司及其子公司與關係企業。

2.1.1 For purposes of this AI Policy, the term "Advantech" means Advantech and its subsidiaries and affiliates.

2.1.2 「AI 系統」係指一種基於機器的系統，設計用於以不同程度的自主性運作（包括解決問題和執行任務），在部署後可能表現出適應性，且針對明確或隱含的目標，從其接收到的輸入中推論如何生成輸出（例如預測、內容、建議或決定），進而影響實體或虛擬環境。

2.1.2 The term "AI System" means a machine-based system designed to operate with varying levels of autonomy (including solving problems and performing tasks), that may exhibit adaptiveness after deployment and that, for explicit or implicit objectives, infers, from the input it receives, how to generate outputs such as predictions, content, recommendations, or decisions that can influence physical or virtual environments.

2.2 為了明確起見，AI 系統包括任何涉及使用軟體演算法、神經網路或模型來分析輸入數據、從中學習並基於該學習做出決策或預測的軟體或其他技術，其包括但不限於與以下各項相關之技術：

2.2 For clarity, AI Systems include any software or other technology that involves the use of software algorithms, neural networks, or models to analyze input data, learn from that data, and then make decisions or predictions based on that learning, including without limitation technology related to:

2.2.1 大型語言模型 (Large Language Models · LLMs)

2.2.1 Large Language Models (LLMs)

2.2.2 具備自主推論、生成能力或機器學習 (ML) 之現代 AI 系統

2.2.2 Modern AI systems with autonomous reasoning, generative capabilities, or machine learning (ML)

2.2.3 生成式 AI (Generative AI)

三、原則 Principles

3.1 AI 的倫理使用

3.1 Ethical Use of AI

3.1.1 公平性：AI 系統的設計與訓練必須避免偏見與歧視。應定期進行稽核，以識別並減輕 AI 演算法中的任何潛在偏見。

3.1.1 Fairness: AI systems must be designed and trained to avoid biases and discrimination. Regular audits should be conducted to identify and mitigate any potential biases in AI algorithms.

3.1.2 透明度：AI 系統應具備可解釋性與可理解性。應讓利害關係人了解 AI 決策是如何做出的，特別是在高風險情境中。

3.1.2 Transparency: AI systems should be explainable and understandable. Stakeholders should be informed about how AI decisions are made, especially in high-stakes scenarios.

3.1.3 問責機制：AI 決策責任最終歸屬組織或個人。必須為 AI 系統建立明確的問責機制。參與 AI 開發的員工必須了解其職責以及其工作帶來的影響。

3.1.3 Accountability: Responsibility for AI-assisted decisions ultimately rests with the organization or the responsible individual. A clear accountability framework must be established for all AI systems. Employees involved in the development of AI systems must understand their roles and responsibilities, as well as the impact of their work.

3.1.4 人類監督：AI 系統應支持而非取代人類判斷。在 AI 系統的整個生命週期中，特別是針對高風險使用情境，必須維持適當的人類監督。此原則係 EU AI Act 之核心精神，亦是研華 AI 治理的基本要求。

3.1.4 Human Oversight: AI Systems shall support, not replace, human judgment. Appropriate human oversight must be maintained throughout the lifecycle of AI Systems, particularly for high-risk use cases. This principle reflects the core spirit of the EU AI Act and is a fundamental requirement of Advantech's AI governance.

3.1.5 準確性與可靠性：AI 系統應經過設計、測試與持續監控，以確保合理程度的準確性、穩健性與可靠性。已知的系統限制應向使用者充分揭露。

3.1.5 Accuracy and Reliability: AI Systems should be designed, tested, and monitored to ensure reasonable levels of accuracy, robustness, and reliability. Known limitations of systems should be adequately disclosed to users.

3.1.6 AI 幻覺管理：生成式 AI 可能產生看似合理但不正確的輸出（即「幻覺」）。AI 生成的內容在使用前必須經過人工驗證。員工有責任在依賴 AI 產出之前，對其進行審查與確認。

3.1.6 AI Hallucination Management: Generative AI may produce outputs that appear plausible but are factually incorrect ("hallucinations"). AI-generated content must be verified before use. Employees are responsible for reviewing and validating AI-generated outputs before relying upon them.

3.2 數據隱私與安全

3.2 Data Privacy and Security

3.2.1 數據保護：個人數據之處理必須符合適用之隱私法規（包括但不限於 GDPR、CCPA 及當地個資保護法規）及合法處理基礎。在可能的情況下，個人數據應進行去識別化。

3.2.1 Data Protection: Personal data must be processed in accordance with applicable privacy laws (including but not limited to GDPR, CCPA, and local data protection regulations) and a valid legal basis. Personal data should be de-identified where possible.

3.2.2 安全措施：在所有 AI 系統中實施的安全措施與協定，必須達到研華在個資保護規則中所要求的保護措施。

3.2.2 Security Measures: Security measures and protocols implemented in all AI systems must meet the requirements of Advantech's personal data protection guidelines.

3.3 AI 安全

3.3 AI Security

3.3.1 AI 系統面臨有別於傳統軟體的特有安全威脅，研華所有 AI 系統及相關開發流程應評估並防範下列風險：

3.3.1 AI Systems face security threats distinct from traditional software. All Advantech AI Systems and related development processes must assess and guard against the following risks:

3.3.1.1 提示詞注入（Prompt Injection）：惡意輸入企圖操控 AI 系統執行未授權之指令或洩漏敏感資訊。

3.3.1.1 Prompt Injection: Malicious inputs that attempt to manipulate AI systems into executing unauthorized instructions or leaking sensitive information.

3.3.1.2 資料投毒（Data Poisoning）：透過污染訓練資料影響模型行為或決策結果。

3.3.1.2 Data Poisoning: Corrupting training data to influence model behavior or decision outcomes.

3.3.1.3 模型竊取（Model Theft）：未授權複製或萃取 AI 模型之智慧財產。

3.3.1.3 Model Theft: Unauthorized replication or extraction of an AI model's intellectual property.

3.3.1.4 對抗性攻擊（Adversarial Attack）：精心設計的輸入導致 AI 系統產生錯誤輸出。

3.3.1.4 Adversarial Attacks: Carefully crafted inputs that cause AI systems to produce incorrect outputs.

3.3.1.5 敏感資料洩漏（Sensitive Data Leakage）：AI 系統在推論過程中意外洩漏訓練資料或用戶敏感資訊。

3.3.1.5 Sensitive Data Leakage: AI systems inadvertently exposing training data or user-sensitive information during inference.

3.3.1.6 第三方 AI 工具風險：透過第三方 AI 服務或模型引入的安全漏洞。所有未經審查的第三方 AI 工具應經 AI 治理小組進行資料處理條款審查、第三方資格審查，以及安全認證要求等。審查後方可採用。

3.3.1.6 Third-Party AI Tool Risks: Security vulnerabilities introduced through the use of third-party AI services or models. Any third-party AI tool that has not undergone prior review must be evaluated by the AI Governance Committee. The evaluation shall include, but is not

limited to, a review of data processing terms, third-party vendor qualifications, and security certification requirements. Only after successful completion of the review may the tool be approved for use.

3.4 合規與法規

3.4 Compliance and Regulation

3.4.1 法律合規：所有 AI 系統必須遵守適用之法律法規（包括但不限於歐盟 AI 法案（EU AI Act）、GDPR、CCPA 及當地個資保護法規等）。凡涉及個人數據、客戶數據或高風險應用的專案，均應諮詢法務團隊。

3.4.1 Legal Compliance: All AI Systems must comply with applicable laws and regulations (including but not limited to the EU AI Act, GDPR, CCPA, and local data protection regulations). The Legal/Compliance team should be consulted for projects involving personal data, customer data, or high-risk applications.

3.4.2 法規更新：研華將密切關注 AI 法規的變更，並相應調整其政策與實務。AI 治理小組負責監控法規動態並定期向相關部門通報。

3.4.2 Regulatory Updates: Advantech will stay informed about changes in AI regulations and adapt its policies and practices accordingly. The AI Governance Committee is responsible for monitoring regulatory developments and periodically briefing relevant departments.

3.5 持續改進

3.5 Continuous Improvement

3.5.1 監控與評估：必須定期監控 AI 系統的性能與倫理影響。應建立回饋機制，根據使用者和利害關係人的意見來改進 AI 系統。

3.5.1 Monitoring and Evaluation: AI systems must be regularly monitored for performance and ethical implications. Feedback mechanisms should be established to improve AI systems based on user and stakeholder input.

3.5.2 培訓與發展：參與 AI 開發的員工將接受有關倫理 AI 實務、數據處理及相關技術的持續培訓。

3.5.2 Training and Development: Employees involved in AI development will receive ongoing training on ethical AI practices, data handling, and relevant technologies.

3.6 協同合作與溝通

3.6 Collaboration and Communication

3.6.1 跨領域合作：鼓勵技術團隊與非技術團隊之間的合作，以確保 AI 開發具備全盤整體的考量。

3.6.1 Interdisciplinary Collaboration: Encourage collaboration between technical and non-technical teams to ensure a holistic approach to AI development.

3.6.2 利害關係人參與：與包括使用者、客戶及受影響之個人、組織及社群在內的利害關係人進行交流，以收集意見並解決對 AI 系統的疑慮。

3.6.2 Stakeholder Engagement: Engage with stakeholders, including users, customers, and affected individuals, organizations and communities, to gather input and address concerns regarding AI systems.

四、實施 Implementation

4.1 角色與職責

4.1 Roles and Responsibilities

4.1.1 AI 治理小組：

4.1.1 AI Governance Committee:

4.1.1.1 成立 AI 治理小組以制定並精進 AI 政策。

4.1.1.1 Establish an AI Governance Committee to develop and refine AI policies.

4.1.1.2 審查 AI 工具的合規性。任何擬使用未經先前批准的第三方 AI 工具，必須提交至 AI 治理小組進行審查。

4.1.1.2 Review the compliance of AI tools. Any proposed use of a third-party AI tool that has not been previously approved must be submitted to the AI Governance Committee for review.

4.1.1.3 負責 AI 事件應對：接收 AI 相關事件通報，協調調查、補救及對外報告流程。

4.1.1.3 Lead AI incident response: Receive AI-related incident reports and coordinate investigation, remediation, and external reporting processes.

4.1.1.4 維護政策與指引：定期審查本政策及相關執行辦法，確保與法規要求保持一致。

4.1.1.4 Maintain policies and guidelines: Periodically review this policy and related implementation procedures to ensure alignment with regulatory requirements.

4.1.1.5 維護核准 AI 工具清單 (Approved AI Tool Registry)：定期審查並公告已核准之 AI 工具名單。

4.1.1.5 Maintain the Approved AI Tool Registry: Periodically review and announce the list of approved AI tools.

4.1.2 專案負責人：負責確保其專案遵守本政策，並向 AI 治理小組報告任何倫理方面的疑慮。

4.1.2 Project Leads: Responsible for ensuring that their projects adhere to this policy and report any ethical concerns to the AI Governance Committee.

4.2 AI 事件管理

4.2 AI Incident Management

4.2.1 鼓勵員工透過既定管道，向 AI 治理小組報告任何與 AI 系統相關的倫理疑慮或事件。

4.2.1 Employees are encouraged to report any ethical concerns or incidents related to AI systems through established channels to the AI Governance Committee.

4.2.2 以下類型的 AI 相關事件應立即通報：

4.2.2 The following types of AI-related incidents must be reported immediately:

4.2.2.1 資料外洩或未授權存取：AI 系統導致的個人數據或機密資訊外洩。

4.2.2.1 Data breach or unauthorized access: Personal data or confidential information leaked due to an AI system.

4.2.2.2 錯誤決策：AI 系統產生可能造成重大損害的不正確決策或建議。

4.2.2.2 Erroneous decisions: AI systems producing incorrect decisions or recommendations that may cause significant harm.

4.2.2.3 偏見或歧視問題：AI 系統輸出顯示出對特定族群的不公平偏見或歧視。

4.2.2.3 Bias or discrimination issues: AI system outputs exhibiting unfair bias or discrimination against specific groups.

4.2.2.4 法律與合規風險：AI 系統使用可能違反適用法規之情形。

4.2.2.4 Legal and compliance risks: AI system usage potentially violating applicable regulations.

4.2.3 AI 事件應對流程包含以下三個階段：

4.2.3 The AI incident response process consists of the following three phases:

4.2.3.1 調查 (Investigation)：由 AI 治理小組啟動調查，確認事件範疇、影響及根本原因。

4.2.3.1 Investigation: The AI Governance Committee initiates an investigation to determine the scope, impact, and root cause of the incident.

4.2.3.2 補救 (Remediation)：根據調查結果，採取技術修正、流程改善或暫停系統使用等措施。

4.2.3.2 Remediation: Based on investigation findings, take measures such as technical corrections, process improvements, or suspension of system use.

4.2.3.3 報告 (Reporting)：依法規要求及事件嚴重程度，向主管機關、受影響當事人或其他利害關係人進行通報。

4.2.3.3 Reporting: Notify regulatory authorities, affected parties, or other stakeholders in accordance with legal requirements and the severity of the incident.

4.3 回饋機制

4.3 Feedback Mechanisms

4.3.1 建立供使用者和利害關係人對 AI 系統提供意見的管道，這些意見將會被定期審查，作為持續改善的依據。

4.3.1 Establish channels for users and stakeholders to provide feedback on AI systems. This feedback will be reviewed regularly as a basis for continuous improvement.

4.4 核准 AI 工具清單

4.4 Approved AI Tool Registry

4.4.1 研華維護一份「核准 AI 工具清單 (Approved AI Tool Registry)」，列明所有經 AI 治理小組核准可於研華內部使用之 AI 工具。

4.4.1 Advantech maintains an Approved AI Tool Registry listing all AI tools approved by the AI Governance Committee for use within Advantech.

4.4.2 清單由 AI 治理小組負責維護，並定期更新。員工若需申請新工具加入清單，應透過正式申請流程提交至 AI 治理小組審查。

4.4.2 The Registry is maintained and periodically updated by the AI Governance Committee. Employees wishing to add a new tool to the Registry should submit it for review through the formal application process to the AI Governance Committee.

4.4.3 未列入清單之 AI 工具不得用於處理客戶個人資料、機密商業資訊或高風險決策流程。

4.4.3 AI tools not listed in the Registry must not be used for processing customer personal data, confidential business information, or high-risk decision-making processes.

4.5 智慧財產權

4.5 Intellectual Property

4.5.1 AI 生成內容之所有權：於職務範圍內或使用研華資源所產生之 AI 生成內容，如包含研華之原創創意、商業機密、專有知識或其他智慧財產，其所有權及相關權利均歸屬研華，並受《保密與發明轉讓協議》及相關公司規範約束。

4.5.1 Ownership of AI-generated content: AI-generated content produced by employees in the course of their employment or using Advantech resources that contains or is derived from Advantech's original intellectual contributions, confidential information, proprietary know-how, or other intellectual property shall be owned exclusively by Advantech. Such content and any associated intellectual property rights shall be governed by the employee's Confidentiality and Inventions Assignment Agreement, as well as applicable company policies and procedures.

4.5.2 第三方智慧財產權風險：使用 AI 工具生成的內容可能包含受第三方版權保護的素材。員工在使用 AI 生成內容時，應評估並降低潛在的第三方智慧財產權侵害風險，必要時諮詢法務團隊。

4.5.2 Third-party IP risk: Content generated by AI tools may incorporate materials protected by third-party copyrights. Employees should assess and mitigate potential third-party IP infringement risks when using AI-generated content, and consult the Legal team when necessary.

4.5.3 版權合規：切勿以可能違反智慧財產權法律的方式使用 AI 工具，例如進行逆向工程或大量複製受版權保護的程式碼或內容。

4.5.3 Copyright compliance: Do not use AI tools in a manner that may violate intellectual property laws, such as reverse engineering or mass reproduction of copyrighted code or content.

4.5.4 開源合規：在 AI 系統開發中使用開源元件時，應遵守相關授權條款，並向法務團隊諮詢合規要求。

4.5.4 Open-source compliance: When using open-source components in AI system development, comply with the relevant license terms and consult the Legal team for compliance requirements.

4.6 實施守則 (Do's and Don'ts)

4.6 Implementation Guidelines (Do's and Don'ts)

4.6.1 針對程式開發人員

4.6.1 For Software Developers

4.6.1.1 應做事項 (Do's) :

4.6.1.1 Do's:

4.6.1.1.1 應僅將 AI 系統用於您獲授權執行的任務。

4.6.1.1.1 Do use AI Systems only for tasks that you are authorized to perform.

4.6.1.1.2 應遵守適用當地法律法規（例如：個資保護法以及歐盟 AI 法案等）。凡涉及個人數據、客戶數據或高風險應用的專案，均應諮詢法務團隊。

4.6.1.1.2 Do comply with applicable local laws and regulations (e.g., Personal Data Protection Act and EU AI Act, etc.). The Legal/Compliance team should be consulted for projects involving personal data, customer data or high-risk applications.

4.6.1.1.3 應記錄您對 AI 系統的使用及其與專案相關的產出，以供未來參考、確立原創作者身分，並揭露任何潛在的智慧財產權（IP）或所有權限制。工程師可以在 Github 將開發的細節，透過 PR(pull request)在合併回到主線(main branch)時進行深入說明，把對開發工作有實質貢獻之 prompts 與 AI 產出，附於 PR 作為參考與稽核紀錄。

4.6.1.1.3 Do document your use of AI systems and any project-related outputs they generate for future reference, to establish authorship, and to disclose any potential intellectual property (IP) or ownership limitations.

For software development projects, engineers may provide detailed explanations of their implementation approach in a GitHub Pull Request (PR) when merging changes into the main branch. Relevant interactions with AI agents, including prompts and AI-generated outputs that materially contributed to the development work, may also be included in the PR for reference and audit purposes.

4.6.1.1.4 應謹記，您在研華所創作的任何內容皆受您的《保密與發明轉讓協議》所約束。

4.6.1.1.4 Do remember that anything you create at Advantech is subject to your Confidentiality and Inventions Assignment Agreement.

4.6.1.1.5 應謹記，客戶的提示詞（Prompts）即等於客戶數據。

4.6.1.1.5 Do remember that customer prompts = customer data.

4.6.1.1.6 應謹記，若要將客戶提示詞用於任何不直接使該客戶受益的目的，必須取得其「明確同意」。以下範例均需要明確同意：

4.6.1.1.6 Do remember that express consent is required to use customer prompts for any purpose that does not directly benefit the customer. The following examples require express consent:

4.6.1.1.6.1 使用客戶提示詞來訓練任何 AI 系統。

4.6.1.1.6.1 Using customer prompts to train any AI System.

4.6.1.1.6.2 使用客戶提示詞來改善使用者體驗、對話流暢度、回覆品質等。

4.6.1.1.6.2 Using customer prompts to improve user experience, conversation fluency, response quality, etc.

4.6.1.1.7 在整合第三方 AI 系統以處理客戶數據/提示詞時，應考慮以下隱私問題：

4.6.1.1.7 Do consider the following privacy issues when integrating third-party AI Systems to process customer data/prompts:

4.6.1.1.7.1 客戶數據/提示詞不得接受第三方的專人人工審查。

4.6.1.1.7.1 Customer data/prompts must not be subject to third-party human review.

4.6.1.1.7.2 第三方不得因任何原因使用或儲存客戶數據/提示詞。

4.6.1.1.7.2 Customer data/prompts must not be used or stored by the third party for any reason.

4.6.1.1.8 應參考免責聲明指引，以協助評估您的 AI 實施案是否需要包含「AI 免責聲明」。

4.6.1.1.8 Do review the disclaimer guidelines to help determine whether your AI implementation needs to include an "AI Disclaimer."

4.6.1.1.9 應記得對 AI 系統開發和微調（Fine-tuning）中所使用的數據集保持詳細的記錄。

4.6.1.1.9 Do remember to keep detailed records of the datasets used in the development and fine-tuning of AI Systems.

4.6.1.2 禁止事項（Don'ts）：

4.6.1.2 Don'ts:

4.6.1.2.1 切勿將未經 AI 治理小組事先批准的公開可用 AI 工具用於客戶個資處理。

4.6.1.2.1 Don't use publicly available AI tools that have not been previously approved by the AI Governance Committee for processing customer personal data.

4.6.1.2.2 切勿在未諮詢法務團隊的情況下，從第三方網站（例如 Facebook）抓取（Scrape）任何內容（如圖像），因為某些內容可能受版權保護或受特定使用條款約束。

4.6.1.2.2 Don't scrape any content (e.g., images) from third-party websites (e.g., Facebook) without consulting Legal, as some content may be copyright protected or subject to certain terms of use.

4.6.1.2.3 切勿以可能違反智慧財產權法律的方式使用 AI 工具，例如進行逆向工程或複製受版權保護的程式碼。

4.6.1.2.3 Don't use AI tools in a manner that might violate intellectual property laws, such as for reverse engineering or replicating copyrighted code.

4.6.1.2.4 切勿以可能危害研華商業秘密的方式使用公開 AI 工具，例如輸入公司機密或專有資訊。

4.6.1.2.4 Don't use public AI tools in a way that might endanger Advantech trade secrets, such as by inputting confidential or proprietary company information.

4.6.1.2.5 切勿在未經審查和測試生成產出的情況下盲目依賴任何 AI 輸出，以確保其安全、適當且符合該特定用途的功能需求（請參見第 3.1.6 節：AI 幻覺管理）。

4.6.1.2.5 Don't rely on any AI output without reviewing and testing the generated output to ensure it is secure, appropriate, and functionally aligned for the given purpose (see Section 3.1.6: AI Hallucination Management).

4.6.1.2.6 切勿將 AI 系統用於可能限制研華智慧財產權所有權的任務，或用於任何非法或不道德的目的。這包括生成惡意程式碼、複製可能屬於第三方的材料，或將 AI 的工作成果冒充為您自己的成果。

4.6.1.2.6 Don't use AI Systems for tasks that could limit Advantech's intellectual property ownership or for any illegal or unethical purpose. This includes generating malicious

code, copying materials that could be owned by third parties, or passing off AI work as your own.

4.6.2 針對非程式開發人員

4.6.2 For Non-Developers

4.6.2.1 應做事項 (Do's) :

4.6.2.1 Do's:

4.6.2.1.1 應主動提問：如果不確定 AI 系統的工作原理，請向開發人員或其他相關人員尋求澄清。

4.6.2.1.1 Do ask questions: If unsure about how an AI System works, seek clarification from developers or other relevant personnel.

4.6.2.1.2 應監控成效：監控 AI 驅動的結果，並報告任何不一致或問題。

4.6.2.1.2 Do monitor outcomes: Monitor AI-driven results and report any discrepancies or issues.

4.6.2.1.3 應與開發人員合作：與開發人員密切合作，使 AI 系統與業務目標保持一致。

4.6.2.1.3 Do collaborate with developers: Work closely with developers to align AI Systems with business objectives.

4.6.2.1.4 應遵循 Human-in-the-Loop (HITL) 原則：特別是針對高風險或重大商業決策時，必須有人類的監督介入，AI 不得作為唯一的決策依據。

4.6.2.1.4 Do apply the Human-in-the-Loop (HITL) principle: Human oversight must be involved, especially in high-risk or significant business decisions. AI must not serve as the sole basis for decisions.

4.6.2.2 禁止事項 (Don'ts) :

4.6.2.2 Don'ts:

4.6.2.2.1 切勿將未經 AI 治理小組事先批准的公開可用 AI 工具用於客戶個資處理。

4.6.2.2.1 Don't use publicly available AI tools that have not been previously approved by the AI Governance Committee for processing customer personal data.

4.6.2.2.2 切勿誤用 AI 系統：避免將 AI 系統用於超出其預期用途的目的。

4.6.2.2.2 Don't misuse AI Systems: Avoid using AI Systems beyond their intended purpose.

4.6.2.2.3 切勿忽視倫理疑慮：如果您注意到 AI 系統存在潛在的倫理問題，請勇於發言反應。

4.6.2.2.3 Don't ignore ethical concerns: Speak up if you notice potential ethical issues with AI Systems.

4.6.2.2.4 切勿將 AI 系統用於高風險或重大決策：避免在沒有人類監督的情況下，完全依賴 AI 系統做出高風險或重大決策（例如：招聘、晉升、員工績效評估或財務評估）。

4.6.2.2.4 Don't use AI Systems for high-risk or significant decisions: Avoid relying solely on AI Systems for high-stakes or significant decisions (e.g., hiring, promotions, employee

performance evaluation, or financial assessments) without human oversight.

4.6.2.2.5 切勿預設 AI 可以取代人類判斷：請理解 AI 是一種用來輔助 (Augment) 而非取代人類決策的工具。

4.6.2.2.5 Don't assume AI replaces human judgment: Understand that AI is a tool to augment, not replace, human decision-making.

五、審查與更新 **Review and Updates**

5.1 本政策於 2026 年 7 月 1 日發布並生效。每年至少審查一次，並視需要修訂，以確保其合時性與有效性。AI 治理小組負責啟動審查程序，並視法規變動、重大 AI 事件或企業經營需求進行及時修訂。

5.1 This Policy was issued and effective on July 1, 2026. It shall be reviewed annually and revised as necessary to ensure its continued relevance and effectiveness. The AI Governance Committee is responsible for initiating the review process and shall update this Policy in a timely manner in response to changes in applicable laws or regulations, significant AI-related events, or evolving business needs.