

日月光投資控股股份有限公司人工智慧管理政策
ASE Technology Holding CO., LTD. Artificial Intelligence
Management Policy

第一條 目的

Article I: Purpose

為強化日月光投資控股股份有限公司（以下簡稱「本公司」）以及所屬子公司（本公司及所屬子公司合稱「本集團」）內部使用人工智慧（Artificial Intelligence，以下簡稱「AI」）技術之合規、安全及倫理，確保內外部使用 AI（含生成式 AI）能管控相關風險，並降低治理風險，以符合本公司資通安全政策及相關法規要求特制定本規範。

To ensure the compliance, safety and ethics of the use of artificial intelligence (AI, hereinafter referred to as "AI") technology by ASE Technology Holding Co., Ltd. (hereinafter referred to as "ASEH ") and its subsidiaries (hereinafter referred to as "Group ", to ensure that the use of AI (including generative AI) internally and externally can manage related risks and reduce governance risks, so as to comply with the ASEH's communication and security policies.

第二條 適用範圍

Article II: Scope of Application

1. 適用對象：本集團所有董事、經理人、員工、承包商、臨時人員及任何以直接或間接方式使用本集團資源之第三人（下合稱「使用者」）。

Applicable Parties: All directors, managers, employees, contractors, temporary personnel, and any third party who directly or indirectly uses the Group's resources (collectively, "Users").

2. 適用客體：包含本集團內部開發之所有 AI 系統及相關應用（涵蓋生成式 AI、機器學習模型、影像辨識系統等），或經核准使用之外部 AI 工具（涵蓋 AI 系統、通用 AI 模型及相關應用）。

Applicable Subjects: Includes all AI systems and related applications

developed internally by the Group (covering generative AI, machine learning models, image recognition systems, etc.), and/or approved external AI tools (covering AI systems, general AI models, and related applications).

3. 補充規範：本公司所屬子公司得依業務需求，參酌本使用規範，另增訂 AI 使用指引或內控管理措施，惟不得違反本規範。

Supplementary Guidelines: Subsidiaries of ASEH may, based on business needs and with reference to these User Guidelines, formulate additional AI usage guidelines or internal control management measures, provided they do not violate these Guidelines.

第三條 定義

Article III: Definition

1. 人工智慧（AI）系統：係指透過大量資料學習，利用機器學習或相關建立模型之演算法，進行感知、預測、決策、規劃、推理、溝通等模仿人類學習、思考及反應模式之系統。

Artificial Intelligence (AI) System: Refers to a system that learns from big data and utilizes machine learning or related algorithms to build models, mimicking human learning, thinking, and reaction patterns in areas such as perception, prediction, decision-making, planning, reasoning, and communication.

2. 生成式 AI：係指可以生成模擬人類智慧創造之內容的相關 AI 系統，其內容形式包括但不限於文章、圖像、音訊、影片及程式碼等。

Generative AI: Refers to AI systems that can generate content that simulates human intelligence, including but without limitation to the articles, images, audio, videos, code and so forth.

第四條 權責單位

Article IV: Responsibilities and Authorities

1. 法遵部門：負責解釋與 AI 相關之法規、違規調查等事項。

Compliance Department: Responsible for interpreting AI-related regulations, investigating violations and so forth.

2. 資安部門：負責制定及修訂本規範。

Cybersecurity Department: Responsible for developing and revising this specification.

3. 風險部門：負責 AI 系統及生成式 AI 之風險評估。

Risk Management Department: Responsible for risk assessment of AI systems and generative AI.

4. 資訊（IT）部門：負責各子公司之 AI 系統及生成式 AI 使用管理。

Information Technology (IT) Department: Responsible for the management of AI systems and generative AI usage in each subsidiary.

5. 人力資源部門：負責各子公司 AI 倫理與安全教育訓練之執行。

Human Resources Department: Responsible for implementing AI ethics and safety training in each subsidiary.

第五條 核心原則

Article V: Core Principles

使用者應遵循下列原則，確保 AI 系統具備 NIST AI RMF 定義之可信賴特徵：

Users should adhere to the following principles to ensure that AI systems possess the trustworthy characteristics defined by the NIST AI RMF:

1. 有效與可靠（Valid and Reliable）：確保 AI 系統在既定條件與時間內能準確執行任務，並具備對抗異常輸入或變化環境的強健性（Robustness）。

Valid and Reliable: Ensure that the AI system can accurately perform tasks under given conditions and within a given timeframe, and possess robustness against abnormal inputs or changing environments.

2. 安全（Safe）：AI 系統之運作不得危及人類生命、健康、財產或環境，

並應設計失效安全（Fail-Safe）機制。

Safe: The operation of the AI system must not endanger human life, health, property, or the environment; and a fail-safe mechanism should be designed.

3. 資安與韌性（Secure and Resilient）：應具備防範資料下毒、模型竊取等攻擊之能力，並能在遭受攻擊或意外後恢復正常運作。

Secure and Resilient: The system should be capable of preventing attacks such as data poisoning and model theft, and should be able to recover normal operation after an attack or accident.

4. 可解釋與可解讀（Explainable and Interpretable）：系統應能說明其運作機制（可解釋性），且其產出結果應能被使用者理解其意義與脈絡（可解讀性）。

Explainable and Interpretable: The system should be able to explain its operating mechanism (explainability), and its output results should be understandable to users in terms of meaning and context (interpretability).

5. 隱私強化（Privacy-Enhanced）：除遵守個資法外，應優先採用隱私強化技術（PETs）與去識別化措施，保障個人自主權與隱私。

Privacy-Enhanced: In addition to complying with personal data protection laws, privacy-enhancing technologies (PETs) and de-identification measures should be prioritized to protect individual autonomy and privacy.

6. 公平與偏差管理（Fair with Harmful Bias Managed）：應主動識別並管理系統性、計算性及人為認知上的偏差，避免產生歧視性結果或加劇現有不平等。

Fair with Harmful Bias Managed: Human, computational, and cognitive biases should be proactively identified and managed to avoid discriminatory outcomes or exacerbating existing inequalities.

7. 當責與透明（Accountable and Transparent）：應揭露使用 AI 系統之資訊，明確定義人類監督角色，使用者並對 AI 系統之決策後果負最終責任。

Accountable and Transparent: Information regarding the use of AI systems

should be disclosed, the role of human oversight clearly defined, and users ultimately responsible for the consequences of AI system decisions.

第六條 導入流程（PDCA 與 NIST Core Functions）

Article VI: Implementation Process (PDCA and NIST Core Functions)

1. 計畫（Plan）－ 映射（Map）：權責單位擬使用 AI 系統前，須於系統規劃階段建立「上下文脈絡（Context）」，識別潛在影響對象、預期效益、成本，以及系統是否超出其知識限制（Knowledge Limits）建立資料可追溯性（Data Lineage），並確認資料最小化需求。

Plan – Map: Before using an AI system, the responsible department must establish a "context" during the system planning phase, identifying potential impact targets, expected benefits, costs, and whether the system exceeds its knowledge limits. Data lineage must be established, and data minimization requirements must be confirmed.

2. 執行（Do）－ 治理與管理（Govern & Manage）：權責單位進行資安與合規審核。針對高風險系統，應實施「人機協作配置（Human-AI Configuration）」定義，確保資源配置符合風險優先級。

Do – Govern & Manage: The responsible department conducts cyber security and compliance audits. For high-risk systems, a "Human-AI Configuration" definition should be implemented to ensure resource allocation aligns with risk priorities.

3. 檢查（Check）－ 測量（Measure）：除定期審核使用記錄外，須實施人工智慧風險測試、評估、驗證及確效(TEVV)程序。使用定量或定性指標測量核心原則之落實程度（如準確率、偏差值、對抗性測試結果）。

Check – Measure: In addition to periodically reviewing usage records, an AI TEVV (AI Test, Evaluation, Validation and Verification) procedure must be implemented. Quantitative or qualitative indicators should be used to measure the implementation of core principles (e.g., accuracy, deviation, adversarial test results).

4. 行動 (Act) — 管理 (Manage)：依據測量結果持續改善或重新校準。
對於殘餘風險 (Residual Risk) 超過容忍度之系統，應啟動退場或停用機制 (Decommissioning)。

Act – Manage: Continuously improve or recalibrate based on measurement results. For systems where residual risk exceeds tolerance, decommissioning should be initiated.

第七條 生成式 AI 使用規範

Article VII: Generative AI Usage Guidelines

1. 使用者應選擇符合本規範之 AI 工具，為本集團所提供的生成式 AI 工具應經資訊 (IT) 部門同意使用，且僅限使用於執行公司業務相關用途。

Users should select AI tools that comply with these guidelines. Generative AI tools provided to the Group must be approved by the Information Technology (IT) department and are limited to use for purposes related to the execution of company business.

2. 使用者應就其風險進行客觀且專業之最終判斷，不得取代業務承辦人之自主思維、創造力及人際互動。

Users should make objective and professional final judgments regarding the risks involved and must not replace the independent thinking, creativity, and interpersonal interaction of business personnel.

3. 使用者不應完全信任生成式 AI 產出之資訊，亦不得將未經確認之產出內容直接作為決策之唯一依據。

Users should not completely trust information generated by generative AI, nor should they use unverified output content as the sole basis for decision-making.

4. 使用生成式 AI 應遵守資通安全、個人資料保護、著作權及相關資訊使用規定，並注意其侵害智慧財產權與人格權之可能性。

The use of generative AI should comply with regulations on information security, personal data protection, copyright, and related information use, and users should be aware of the possibility of infringement of intellectual

property rights and personal rights.

5. 運用 AI 或生成式 AI 作為對外執行業務或提供服務之輔助工具時，應適當告知互動對象或揭露服務係利用人工智慧技術完成，並提醒業務或服務對象得選擇是否使用。

When using AI or generative AI as an auxiliary tool for external business execution or service provision, users should be appropriately informed or disclosed that the service is completed using artificial intelligence technology, and business or service recipients should be reminded that they have the option to choose whether to use it.

6. 各子公司應建立適當之風險管理及定期檢視機制，視相關風險大小、特性、範圍進行管理。

Each subsidiary should establish appropriate risk management and regular review mechanisms to manage risks according to their magnitude, characteristics, and scope.

第八條 禁止事項

Article VIII: Prohibited Matters

1. 未經核准將本集團及客戶之機密資訊輸入至公開或未經資訊（IT）部門同意之 AI 平台。

Inputting the Group's and clients' confidential information into public or AI platforms without the approval of the Information Technology (IT) Department.

2. 使用 AI 生成之假新聞、詐欺訊息、虛假廣告或侵害他人權益之內容。

Using AI to generate fake news, fraudulent messages, false advertisements, or content that infringes on others' rights.

3. 直接引用 AI 生成內容作為專利申請或正式報告，而未標註來源或進行查證。

Directly citing AI-generated content in patent applications or official reports without indicating the source or conducting verification.

4. 開發或使用任何法律所禁止之高風險系統

Developing or using any high-risk systems prohibited by law.

第九條 資料保護

Article IX: Data Protection

1. 本集團及客戶之機密資訊應列為最高敏感級，僅限於內部隔離之 AI 系統處理。

Confidential information of the Group and customers shall be classified as highly sensitive and shall only be processed by internally isolated AI systems.

2. 使用者應保留使用 AI 系統之資訊脈絡，並就生成之資訊進行必要管理，並隨時更新紀錄。

Users shall retain the contextual information of using AI systems, manage the generated information as necessary, and update records at any time.

3. 使用者應確保涉及個人資料之 AI 應用能符合個人資料保護法、歐盟 GDPR 及美國隱私框架就個人資料蒐集、處理、利用之規範。

Users shall ensure that AI applications involving personal data comply with the Personal Data Protection Act, the EU GDPR, and the U.S. privacy framework regulations on the collection, processing, and use of personal data.

4. 當涉及跨境資料傳輸時，應注意遵循歐盟 GDPR 之資料跨境傳輸標準條款（Standard Contractual Clause，簡稱“SCC”），以確保個資適當保。

When cross-border data transmission is involved, attention shall be paid to complying with the EU GDPR's Standard Contractual Clauses (SCC) for cross-border data transfers to ensure proper protection of personal data.

第十條 智慧財產權

Article X: Intellectual Property Rights

1. 使用者應確保 AI 系統及相關應用生成內容之智慧財產權，歸屬本集團。

Users shall ensure that the intellectual property rights of content generated by AI systems and related applications belong to the Group.

2. 對外公開發布時應適當標註「AI 輔助生成」或相關足以識別之文字。

When publicly releasing such content, it should be appropriately labeled with "AI-assisted generation" or other identifiable wording.

3. 使用 AI 應確保來源合法性，不得侵犯他人智慧財產權。

The use of AI must ensure the legality of sources and must not infringe upon the intellectual property rights of others.

第十一條 風險評估

Article XI: Risk Assessment

1. 本公司及各子公司應每年進行風險評估，風險評估應依據 NIST AI RMF 架構，涵蓋以下潛在危害。並透過風險分級矩陣，就低、中、高風險對應不同審核強度：

The company and its subsidiaries shall conduct an annual risk assessment. The risk assessment shall be based on the NIST AI RMF framework and cover the following potential hazards. Through a risk grading matrix, different audit intensities shall correspond to low, medium, and high risks:

- 對人的危害：公民權利、人身安全、經濟機會。

Hazards to individuals: civil rights, personal safety, economic opportunities.

- 對組織的危害：商業營運、聲譽、資安洩漏。

Hazards to the organization: business operations, reputation, cyber security breaches.

- 對生態系統的危害：全球金融系統、供應鏈、自然環境。

Hazards to the ecosystem: global financial system, supply chain, natural environment.

2. 本公司及各子公司應建立完善的內部合規流程，定期自行檢視 AI 系統及生成內容的使用情況，並透過內外部稽核作業，及時修正潛在風險，預防風險。

ASEH and its subsidiaries shall establish comprehensive internal compliance processes, regularly review the use of AI systems and generated content, and correct potential risks in a timely manner through internal and external audits to prevent risks.

第十二條 事件通報

Article XII: Event Reporting

1. 本集團成員發現有任何違反 AI 使用的倫理道德、法令或其他違法本辦法之情事 應即時通報法遵部門。

If any member of the Group discovers any situation that violates AI usage ethics, laws, or other provisions of these Measures, it should be promptly reported to the Compliance Department.

2. 因使用 AI 或生成式 AI 導致產生法律事件或資安事件，應於 24 小時內通報法遵及資安部門。

Any legal or cyber security incidents caused by the use of AI or generative AI should be reported to the Compliance and Cyber Security Departments within 24 hours.

第十三條 教育訓練

Article XIII: Education and Training

人力資源部門與資訊(IT)部門應規劃對使用者之相關訓練。

The Human Resources Department and the Information (IT) Department should plan relevant training for users.

第十四條 違規處置

Article XIV: Disciplinary Actions

違反本規範者，視情節輕重依本公司或各子公司之內部規定予以懲處（包括警告、停職、解僱）；若涉及法律責任將依法追訴。

Those who violate these regulations shall be punished according to the severity of the circumstances in accordance with the internal regulations of the company or its subsidiaries (including warnings, suspension, or dismissal); if legal responsibilities are involved, they will be pursued according to the law.

第十五條 AI 治理倡議

Article XV: Responsible AI Governance Principle

本集團於人工智慧（AI）之開發、部署及整體生命週期管理過程中，應遵循負責任人工智慧（Responsible AI）原則，確保其應用符合資料隱私保護、資通安全、公平性、人類監督、透明性與可解釋性、責任歸屬及使用邊界等要求。

同時，應採取適當之技術與管理措施，防範偏差、誤用或不當影響，並確保 AI 系統之運作不得涉及操控行為、利用弱勢族群、社會評分或未經授權之生物識別監控。

此外，本集團應於 AI 導入及使用過程中，考量其對能源使用、碳排放及環境永續之影響，以支持企業永續發展目標。

The Group shall adhere to Responsible AI principles throughout the development, deployment, and lifecycle management of artificial intelligence (AI) systems. AI applications shall ensure compliance with requirements related to data privacy protection, cybersecurity, fairness, human oversight, transparency and explainability, accountability, and defined boundaries of use.

Appropriate technical and organizational measures shall be implemented to mitigate risks such as bias, misuse, or unintended impacts. AI systems shall not be used for manipulative purposes, exploitation of vulnerable groups, social scoring, or unauthorized biometric surveillance.

In addition, the Group shall consider the environmental impact of AI systems, including energy consumption, carbon emissions, and sustainability implications, to support long-term sustainable development objectives.

第十六條 通過與修正

Article XVI: Approval and Amendment

本辦法由法遵單位及資安單位進行審查，經董事會通過後施行，修訂時亦同。

This Regulation shall be reviewed by the Compliance and Cyber Security Department, and shall be implemented after approval by the Board of Directors. The same procedure applies to amendments.

第十七條 訂定日

Article XVII: Date of Enactment

本辦法於中華民國一一五年五月十二日訂立。

These Regulations were enacted on May 12, 2026, in the Republic of China.