

Sunny Optical Technology (Group) Co., Ltd.

Artificial Intelligence (AI) Policy

I. Purpose

To regulate the artificial intelligence-related business activities of Sunny Optical Technology (Group) Co., Ltd. (hereinafter referred to as "the Group"), establish a governance framework covering the full lifecycle of AI system design, development, deployment, use, monitoring, and retirement, strengthen risk identification, assessment, and control mechanisms, clarify governance responsibilities and accountability requirements, and ensure that AI applications adhere to the core principles of being human-centric, safe and trustworthy, fair and transparent, and respectful of human rights, in compliance with applicable laws, regulations, and internationally recognized standards, this Policy is hereby formulated.

This Policy is formulated with reference to the following internationally recognized frameworks and standards:

- (1) OECD AI Principles;
- (2) UN Recommendation on the Ethics of Artificial Intelligence;
- (3) EU AI Act;
- (4) ISO/IEC 42001 Artificial Intelligence Management System Standard.

II. Scope of Application

This Policy applies to the Group and its subsidiaries, covering all employees and business units, and encompasses all stages of the AI system lifecycle including design, development, procurement, deployment, use, operation, monitoring, optimization and updating, and decommissioning.

Applicable AI use cases include, but are not limited to:

- (1) Developing AI products, services, systems, or models for customers and partners;
- (2) Procuring and integrating third-party AI models, tools, platforms, and services, and developing products and solutions based thereon;
- (3) Deploying and using AI tools or services in business scenarios such as internal operations, production and manufacturing, quality inspection, and supply chain management;
- (4) AI research and development, application implementation, operation, and management activities conducted internally within the Group.

Employees and responsible units involved in AI-related work shall fulfill corresponding compliance, security, and ethical responsibilities in accordance with their job duties. The Group incorporates responsible AI principles into supplier and partner management, requiring suppliers, outsourcing providers, contractors, and other value chain partners to comply with this Policy or equivalent standards in AI development, service delivery, and business cooperation with the Group, and to cooperate with the Group in risk management and oversight.

This Policy shall be implemented in coordination with the Group's existing policies on information security, data privacy, intellectual property, business ethics, and supplier management. In the event of any conflict between this Policy and other internal policies of the Group, this Policy shall prevail. In the event of any conflict between this Policy and prevailing laws, regulations, or regulatory requirements, the more stringent provisions shall apply.

III. Definitions and Terminology

The following terms in this Policy shall have the meanings set forth below:

- (1) Artificial Intelligence System: refers to a system, model, tool, or service that utilizes AI technologies such as machine learning, deep learning, natural language processing, and computer vision to perceive the environment, acquire knowledge, perform reasoning, or assist in decision-making, including but not limited to generative AI systems, deep synthesis systems, predictive analytics systems, and automated decision-making systems.
- (2) Critical AI System: refers to an AI system that involves the Group's core business processes, significant business decisions, important customer rights and interests, large-scale processing of personal information, or that may have a direct impact on personal safety or property safety.
- (3) High-Risk AI System: refers to an AI system that is assessed as high-risk pursuant to the risk assessment process set forth in Chapter 5 of this Policy, or that is expressly classified as high-risk under applicable laws and regulations.
- (4) Generative Artificial Intelligence: refers to AI technology that generates text, images, audio, video, code, and other content based on algorithms, models, and rules.
- (5) Core Data: refers to data classified as core-level under the Group's data classification and grading management system, including but not limited to undisclosed core technical parameters, significant business decision information, and key customer confidential information.

(6) Important Data: refers to data classified as important-level under the Group's data classification and grading management system, the leakage, alteration, or loss of which may have a significant impact on the Group's operations, customer rights and interests, or public interest.

IV. Risk Management

1. Risk Assessment Timing

Risk assessments must be conducted for AI systems at the following milestones:

- (1) Project initiation stage: conduct an initial risk assessment for the AI system to be developed or procured;
- (2) Pre-launch: conduct a pre-launch risk assessment for AI systems that have completed development or integration;
- (3) Major changes: conduct a risk assessment when significant changes occur in system functionality, scope of application, data sources, or model versions.

2. Launch Approval and Retirement Management

Prior to launch, AI systems must complete risk assessment, compliance review, and security testing, and submit an approval process. The system may be formally put into operation only after approval by the appropriate authority. AI systems that have not been approved shall not be connected to the Group's production environment or provide services to customers.

When an AI system is retired or taken offline, the responsible department shall submit a decommissioning application, stating the reasons for retirement, the data disposal plan, and business substitution arrangements, which shall be executed upon approval. During the retirement process, data archiving or destruction, model file cleanup, revocation of related access rights, and system registration ledger updates must be properly completed.

3. Model Operation Monitoring and Drift Management

After an AI system is launched, the responsible department must establish a continuous performance monitoring mechanism to regularly test the accuracy, stability, and reliability of model outputs. When model performance degradation or data distribution shift (model drift) is detected, corrective measures must be promptly initiated, including but not limited to: model retraining, parameter adjustment, dataset updating, or function degradation.

Critical AI systems must establish automated monitoring metrics and alert thresholds, automatically triggering alerts and manual review procedures when key performance indicators

fall below preset standards. Monitoring records, drift events, and corrective measures must be fully archived and retained.

V. Ethical Standards

The Group adheres to the fundamental principles of fairness, justice, non-discrimination, and human-centricity, respects human dignity, legal rights, and human autonomy, and ensures that humans possess the ability to supervise, intervene, and control AI systems.

The Group continuously identifies, assesses, and mitigates bias risks throughout the AI system lifecycle. By optimizing data quality and enhancing the representativeness and diversity of data samples, the Group reduces unfair output results caused by data, models, and algorithms, and avoids replicating, exacerbating, or amplifying existing biases.

For critical AI systems, generative AI systems, and third-party AI systems, whenever unfair, discriminatory, or unreasonable outputs are produced, timely corrective measures must be taken, including optimization, function restriction, or suspension of use. Where necessary, channels for feedback, appeal, and review must be provided to affected parties.

VI. AI Governance Structure

The Group establishes a layered AI governance structure with clearly defined responsibilities to ensure that AI governance decision-making, execution, and oversight responsibilities are clear and well-defined.

1. Decision-Making Layer

The Group's Board of Directors (or the ESG/Sustainability Committee established under the Board) bears oversight responsibility for AI governance, and is responsible for approving AI policies, reviewing significant AI risk matters, and overseeing AI compliance execution.

2. Management Layer

The Group's management (led by the Information Security Management Department, in coordination with the Technology Center, Legal Department, Human Resources Department, and other relevant departments) is responsible for the implementation of AI policies, organization of risk assessments, training advancement, and continuous improvement.

3. Execution Layer

Each business unit shall designate an AI compliance liaison, responsible for AI system registration, risk assessment reporting, daily compliance inspections, and incident reporting within the unit.

VII. Decision-Making Control

AI systems serve only an auxiliary decision-making role, and their outputs may contain errors, biases, and unintended results. The Group establishes a comprehensive human oversight and intervention control mechanism based on business scenarios, risk levels, degree of automation, and potential impact on individual rights and interests.

For automated decision-making scenarios involving significant rights and interests of individuals or organizations, such as employee recruitment screening, performance evaluation, supplier admission assessment, and customer credit assessment, the Group ensures that the decision-making process is transparent and the results are explainable. Automated decisions must provide a manual review channel, and relevant individuals have the right to request an explanation of the decision and the right to refuse decisions based solely on automated processing. Algorithmic fairness assessments are conducted regularly to prevent algorithmic discrimination caused by training data bias.

In production and manufacturing scenarios, when AI systems are used for quality inspection determinations, process parameter recommendations, and equipment failure warnings, critical determination results must be confirmed through manual review before execution. AI outputs involving safety interlocks shall not directly trigger shutdowns or hazardous operations, and must be subject to manual confirmation and safety redundancy mechanisms.

VIII. Transparency and Explainability

The Group continuously enhances the transparency and explainability of AI systems.

- (1) Transparency: clearly disclose the use cases, capability boundaries, limitations, data sources, and potential risks of AI;
- (2) Explainability: provide accessible and easy-to-understand explanations for AI outputs, taking into account the needs of different stakeholders.

When the user interacts with an AI rather than a natural person, or when content involving generative AI or deep synthesis that may cause cognitive confusion is involved, the Group fulfills its obligations of disclosure, labeling, and notification in accordance with laws and regulations such as *the Provisions on the Administration of Deep Synthesis of Internet-Based Information Services* and *the Interim Measures for the Management of Generative AI Services*, and scenario-specific requirements, to reduce the risks of misleading, confusion, and improper dissemination.

Critical AI systems must maintain complete documentation, recording system capabilities, intended use, scope of application, limitations, potential biases, and risks. For AI systems independently developed and trained by the Group, the full process of data sources, data collection and labeling, algorithm logic, training, and testing and validation must also be documented. Documentation must be updated synchronously when significant changes occur in system functionality, scope of application, or version.

Within the applicable scope, the Group explains the generation logic of AI outputs or decisions to end users, affected individuals, and responsible personnel, including output basis, data sources, confidence levels, and key influencing factors. Critical AI systems must support result traceability, enabling localization to input data, model behavior, and core decision-making logic, for monitoring, review, error correction, and accountability purposes.

IX. Responsibility and Accountability

The Group defines the roles and responsibilities of all participating entities throughout the AI lifecycle. Accountability for AI output results, risk management, and compliance matters is clearly assigned and traceable to specific responsible departments, individuals, or legal entities.

The Group establishes mechanisms for monitoring and feedback, incident investigation, corrective action, and remediation. It promptly identifies and investigates errors, damages, and negative impacts caused by AI. For decisions affected by AI, legal and management responsibilities are clarified, and measures such as rectification, risk remediation, and accountability are taken based on the severity of the impact.

When selecting third-party AI products, models, tools, or services, the Group shall assess the professional competence, compliance level, security assurance capabilities, and risk-bearing mechanisms of the third-party entity, clearly define the rights and responsibilities of both parties in contracts, and establish mechanisms for oversight, correction, remediation, and business exit.

X. AI Data Security Protection

The Group manages AI data assets on a tiered basis in accordance with the requirements of confidentiality, availability, and integrity. Data classification and grading standards are implemented in accordance with the Group's Data Classification and Grading Management System.

The following special management rules are applied to AI data security:

- (1) It is strictly prohibited to use core data and important data for training and interaction on public cloud AI servers;
- (2) Training data must come from legally authorized channels. Externally introduced data must pass security inspections to prevent malicious data contamination;
- (3) User data collection must strictly follow the minimum necessity principle, and personal information unrelated to the business shall not be collected;
- (4) User input data shall not be retained unlawfully;
- (5) Cross-border data transfer shall strictly comply with *the Measures for Security Assessment of Cross-border Data Transfer*, *the Personal Information Protection Law*, and other applicable regulations.

Throughout the AI system lifecycle of design, development, testing, deployment, operation, and retirement, the Group strictly complies with *the Personal Information Protection Law*, *the Data Security Law*, and applicable data protection laws and regulations in the jurisdictions where it operates. The following privacy protection requirements are implemented in AI system data processing activities:

- (1) Conduct Data Privacy Impact Assessments (DPIA) to identify the impact of personal information processing activities on individual rights at the AI system project initiation stage;
- (2) Establish a lawful basis review mechanism for personal information processing activities to ensure that AI system processing of personal information has a clear legal basis;
- (3) Safeguard data subject rights, including the right to be informed, the right of access, the right to rectification, the right to erasure, and the right to object to automated decision-making;
- (4) Implement enhanced protection measures for AI applications involving sensitive personal information, including data desensitization, anonymization/pseudonymization, and strict access control.

The Group implements the following cybersecurity special measures for AI systems:

- (1) Conduct regular penetration testing to identify and remediate cybersecurity vulnerabilities in AI systems and interfaces;
- (2) Establish an AI security incident response mechanism, develop emergency response plans, and define incident classification, reporting procedures, and response timelines;

(3) Implement encryption protection for AI training data, model files, and API interfaces, and adopt secure coding practices to prevent adversarial attacks, data poisoning, and model reverse engineering;

(4) Deploy AI content safety protection mechanisms, including prompt shields, to prevent the generation of harmful or inappropriate content.

XI. Capability Definition and Application Restrictions

The Group defines the authorized scope of use, capability boundaries, and inherent limitations of each AI system, and maintains documented records of intended use, decision boundaries, input and output constraints, and technical limitations. The Group establishes permitted use lists, prohibited application lists, safety guardrails, capability constraints, and version upgrade control points to prevent system misuse, business boundary violations, and unexpected scenario invocations, ensuring that AI operations do not exceed preset boundaries.

The Group regularly evaluates and updates the permitted and prohibited application scopes for AI across business segments, and explicitly documents the intended use for each AI system in writing. Without assessment, approval, and authorization, AI systems shall not be used for purposes beyond their original intended use, for scenarios that have not completed validation, or for scenarios involving significant risks. Continuous monitoring and accountability traceability mechanisms are in place.

The Group implements strict access controls for sensitive AI capabilities such as facial recognition, emotion analysis, biometric identification, and large-scale surveillance:

(1) Establish a sensitive AI capability inventory, defining deployment approval authorities and oversight mechanisms for each sensitive technology;

(2) The use of sensitive AI capabilities must undergo special risk assessment and senior management approval, and shall not be deployed without authorization;

(3) Implement access control and operational audit for personnel using sensitive AI capabilities to ensure traceability.

XII. Prohibited Practices

The Group prohibits the development, deployment, and use of the following categories of AI systems (referencing the EU AI Act and internationally recognized prohibited categories), to prevent infringement of individual rights, personal freedom, safety, and social fairness, and to strictly comply with Chinese and local laws and regulations on AI, including but not limited to

the Cybersecurity Law, the Data Security Law, the Personal Information Protection Law, the Interim Measures for the Management of Generative AI Services, and the Provisions on the Administration of Deep Synthesis of Internet-Based Information Services.

Prohibited AI system categories include:

- (1) AI systems that use subliminal, manipulative, or deceptive techniques that have a material distorting effect on individual behavior;
- (2) AI systems that exploit the vulnerability of individuals due to age, disability, or specific social or economic circumstances to have a material distorting effect on individual behavior;
- (3) Social scoring AI systems used for evaluating or classifying social behavior or personal characteristics over time, resulting in adverse treatment in unrelated contexts;
- (4) Real-time remote biometric identification systems in publicly accessible spaces (unauthorized biometric surveillance).

1. Content Generation Prohibitions

It is prohibited to use AI systems to create, copy, publish, or disseminate information containing the following content:

- (1) Content that violates the fundamental principles of the Constitution, endangers national security, honor, and interests;
- (2) Content that incites subversion of state power, overthrow of the socialist system, incites secession, or undermines national unity;
- (3) Content that promotes terrorism, extremism, or advocates ethnic hatred or discrimination;
- (4) Content that disseminates violent or obscene information, or fabricates and spreads false information.

2. Data and Rights Prohibitions

It is prohibited to use AI systems to engage in the following conduct:

- (1) Infringing upon the rights of reputation, privacy, intellectual property, and other lawful rights and interests of others;
- (2) Disclosing state secrets, work secrets, or trade secrets;
- (3) Inducing or defrauding state secrets, work secrets, trade secrets, or sensitive personal information;
- (4) Other circumstances expressly prohibited by laws and regulations.

If an AI system exhibits any of the above prohibited circumstances, its use must be promptly restricted, suspended, or terminated, the system must be removed from the AI management system, and risk assessment, remediation, and accountability tracking must be conducted simultaneously.

XIII. Awareness and Literacy

The Group promotes AI culture building, incorporating AI ethics, security, and risk awareness into the overall governance and training system, and continuously enhances the awareness of employees and relevant stakeholders regarding the social impact, technical risks, and security practices of AI.

The Group designates AI data privacy protection and technology ethics as mandatory courses for all employees, utilizing internal platforms for learning, assessment, and record retention. Through ongoing awareness campaigns based on internal and external technology ethics and risk cases, the Group strengthens risk awareness and responsibility among all employees.

The Group implements tiered and categorized AI training programs tailored to different job positions, application scenarios, and professional capabilities. Personnel engaged in AI development, procurement, deployment, operation, management, and frequent use must receive position-specific specialized training and may only assume their duties after passing the assessment.

The Group regularly reviews the implementation of AI training and maintains archival records. Training content is continuously updated and iterated to adapt to AI technology evolution, ethics standard updates, and legal and regulatory requirements.

XIV. Appeals, Incident Management, and Corrective Actions

The Group maintains a zero-tolerance stance toward violations of AI technology ethics and safety use requirements. Upon discovering suspected violations, risk incidents, or adverse impacts, the Group immediately initiates internal investigation, root cause analysis, and impact assessment. Based on the severity of the incident, measures including suspension of use, access restriction, risk mitigation, rectification, and accountability are taken. Significant matters are submitted to the Group's management for review.

A transparent, accessible, fair, and efficient appeal and clarification process is established. Individuals or organizations adversely affected by AI decisions or outputs who have questions or objections regarding the results or handling process may submit challenges, appeals,

clarifications, or manual review requests through Group-designated channels. The Group handles such requests in accordance with the law, ensuring that affected parties have the right to express opinions and request review.

Any entity or individual who discovers suspected unlawful data processing, privacy infringement, technology ethics violations, or improper use of AI may file a complaint, report, or appeal with the Group's functional departments. The Group handles such requests in accordance with laws, regulations, and internal policies, and in principle responds with the processing result or progress within 15 working days of receiving the request. For complex cases, the response period may be extended to 30 working days, with the reason for extension communicated to the applicant.

Appeal Email: bgs@sunnyoptical.com

XV. AI Environmental Sustainability

The Group attaches great importance to the resource and environmental challenges brought by the rapid development of intelligent business, and consistently upholds the concept of green and low-carbon development. We continuously reduce our business ecological footprint by promoting infrastructure energy efficiency optimization, strengthening environmental data monitoring in operations, and improving internal assessment and improvement mechanisms. At the same time, the Group extends environmental protection requirements to supply chain management, considering environmental performance in partner selection. The Group is committed to integrating ecological responsibility into daily operations and long-term strategy, achieving a balance between economic benefits and climate responsibility.

XVI. Supplementary Provisions

- (1) Matters not covered in this Policy shall be implemented in accordance with national laws and regulations, external industry standards, and the Group's internal policies.
- (2) This Policy is interpreted by the Group's Information Security Management Department.
- (3) This Policy shall be formally implemented upon approval by the CEO. Policy revisions shall follow the same process.
- (4) The Group actively promotes the standardization of AI management systems, plans to establish an AI management system based on the ISO/IEC 42001 standard, and will apply for third-party certification at an appropriate time.