



Tongwei Co., Ltd.

Responsible Artificial Intelligence Management

Commitment and Policy

Tongwei Co., Ltd. (hereinafter referred to as ‘Tongwei Co., Ltd.’ or ‘the Company’) has always adhered to the principle of balancing commercial value with social responsibility, committing to conducting business with the highest ethical standards. The Company places great importance on the responsible development and application of artificial intelligence technology, strictly complying with the *Cybersecurity Law of the People's Republic of China*, the *Data Security Law of the People's Republic of China*, the *Personal Information Protection Law of the People's Republic of China*, the *Interim Measures for the Management of Generative Artificial Intelligence Services*, and other relevant laws, regulations, and requirements, while referencing international standards and best practices such as the ISO/IEC 42001 artificial intelligence management system and the OECD AI Principles. This policy aims to regulate the development, deployment, and use of artificial intelligence systems, ensure the transparency, fairness, and accountability of AI application processes, prevent potential risks, and effectively safeguard the legitimate rights and interests of relevant parties.

I. Scope of Application

This policy applies to all business and operational activities of Tongwei Co., Ltd. and its subsidiaries and affiliates. All directors, senior management, and employees of the Company must comply with this policy in the process of developing, procuring, deploying, using, maintaining, or managing artificial intelligence systems. The Company also advocates that suppliers, partners, and other third parties involved in AI-related products, platforms, models, or services should follow the requirements of this policy and implement their corresponding responsibilities as stipulated by contract.

II. Responsible Departments

The Company's Board of Directors serves as the highest decision-making and supervisory body for artificial intelligence management, responsible for overseeing and reviewing major AI management matters and approving the release of related policies and commitments. The Information Department is responsible for coordinating the construction and certification of the artificial intelligence management system (such as ISO/IEC 42001 certification), conducting AI security and compliance reviews, reviewing privacy impact assessments (PIA/DPIA), and handling AI security incident emergency responses, as well as liaising with regulatory authorities to meet compliance requirements. Each business department is responsible for proposing its own AI application needs, identifying business risks, managing applications, and evaluating their effectiveness to ensure that AI applications meet business management requirements.

III. Responsible Artificial Intelligence Management Commitment and Policy

(1) **Data privacy protection:** The Company strictly abides by relevant laws and regulations on data privacy protection in the development and use of artificial intelligence systems. When collecting, processing, or using personal data to train AI models, the principles of lawfulness, legitimacy, necessity, and good faith must be followed. Data subjects must be clearly informed of the purpose, method, and scope of data collection, and their informed consent must be obtained. The Company commits not to collect personal information irrelevant to the purpose of the AI application and will adopt technical measures such as de-identification and anonymization during data use to minimize the risk of privacy leakage.

(2) **System and network security:** The Company incorporates the network security of artificial intelligence systems into its overall information security management framework. It has established dedicated risk assessment and control mechanisms for security risks unique to AI systems (including but not limited to model poisoning, adversarial attacks, prompt injection, and data leakage). The development and deployment of AI systems must undergo a security review to ensure they possess sufficient robustness and resistance to attacks.

(3) **Bias prevention and fairness:** The Company is committed to ensuring that its artificial intelligence systems avoid discrimination based on race, gender, age, region, religious belief, or other factors during their application. The foundation models used by the Company have undergone bias identification and

elimination mechanisms established by their providers in data selection, algorithm design, and output result evaluation. They also regularly complete fairness assessments and bias detection to comply with relevant control requirements.

(4) Ecological footprint management: When deploying artificial intelligence systems, the Company pays close attention to their environmental impact. For AI model training and inference tasks that require substantial computing resources, the Company comprehensively considers computing power resources and cost-effectiveness, adopting optimization measures to reduce computational loads and power consumption, thereby avoiding energy waste. The Company encourages the prioritization of energy-saving technologies and renewable energy in AI infrastructure, as well as the preferential use of green data centers and low-carbon cloud services. The Company continuously monitors and assesses the impact of AI tools and projects on its sustainable development goals, promoting measurable and traceable ecological benefits from AI applications and fostering synergistic development between artificial intelligence and the Company's green, low-carbon transformation.

(5) Human oversight and intervention mechanisms: The Company explicitly states that the ultimate responsibility for artificial intelligence systems rests with human decision-makers. A specific responsible person must be designated for the development, deployment, and use of any AI system to manage its full lifecycle. It must be ensured that AI systems cannot make irreversible or high-risk decisions without human oversight and a mechanism for manual review or veto. The Company has established a continuous monitoring mechanism for AI system performance, regularly reviewing the output quality and decision accuracy of intelligent agents. When issues with the accuracy of inference, recall, and output are identified, optimization measures are promptly taken to ensure the continuous and stable operation of the AI system.

(6) Appeal and accountability mechanism: The Company has established a mechanism for affected parties to raise objections and file appeals regarding AI decisions or outcomes, clearly defining the appeal channels, processing procedures, and feedback timelines. This ensures that users, customers, and other relevant third parties affected by automated decisions from AI systems can easily raise questions and receive proper handling. In cases of negative impacts caused by errors, biases, or failures of an AI system, the Company can promptly initiate responsibility traceability and remedial measures.

(7) Transparency and explainability: The Company commits to the principle of transparency in the use of artificial intelligence systems. (1) For AI applications that directly face customers or affect their rights and interests, users will be clearly informed that they are interacting with an AI system. (2) The

Company is dedicated to enhancing the explainability of AI decision-making processes, ensuring that in key application scenarios, the results or decisions generated by AI can be explained to relevant parties in an understandable manner. (3) Content generated by artificial intelligence will be marked appropriately (e.g., with watermarks, labels, or statements) to ensure users can identify its source.

(8) Usage boundaries and risk control: The Company strictly adheres to laws, regulations, and ethical guidelines, clearly defining the scope of application and functional boundaries for AI to ensure its use does not exceed its design purpose and compliance requirements. The Company commits not to use or deploy artificial intelligence systems that engage in manipulative behaviors, exploit human vulnerabilities, perform social scoring, or conduct unauthorized biometric surveillance. The Company has established a permissions control mechanism for AI systems, implementing strict access authorization and usage management for AI capabilities involving facial recognition, surveillance, and other highly sensitive areas, to ensure such technologies are used only in lawful, compliant, and necessary business scenarios.

(9) Employee training and awareness enhancement: The Company conducts regular training on AI ethics and security to ensure all employees understand the fundamental principles of responsible artificial intelligence, potential risks, and their respective responsibilities in AI governance. For technical personnel engaged in the development, deployment, and maintenance of AI systems, the Company provides specialized training to ensure they master the professional knowledge and skills for secure AI development and risk prevention.

(10) Continuous improvement and audit: The Company conducts an annual internal audit of responsible artificial intelligence management, performing a comprehensive inspection and assessment of the compliance, security, and fairness of its internal AI systems. At the same time, the Company regularly invites qualified third-party organizations to conduct external assessments and certifications (such as ISO 42001), promptly rectifying any identified issues to continuously improve its level of responsible artificial intelligence management. The Company also identifies the need for management system improvements through management reviews, formulates improvement measures, and verifies their results.

(11) AI impact assessment: Before introducing, deploying, or significantly updating an artificial intelligence system, the Company will conduct a comprehensive AI Impact Assessment. The assessment includes: the potential impact on individual rights, privacy, security, and fairness; a risk assessment of data security, integrity, and availability; a privacy impact assessment of personal information processing activities; and a review of legal and regulatory compliance.

(12) AI supply chain and third-party service management: When procuring or integrating third-party AI models, API services, open-source models, or AI components, the Company will conduct a security assessment and compliance review, including: model license compliance review; verification of the legality of training data sources; scanning for security vulnerabilities and known risks; and checking the security qualifications of third-party service providers (such as ISO 27001, ISO 42001 certifications). For third-party AI services used in critical business scenarios, a service level agreement (SLA) and a security and confidentiality agreement are signed to clarify responsibilities for data security, model performance, and emergency response.



CEO: Liu Shuqi

Tongwei Co., Ltd.

Date: April 2026

Remarks:

1. The Company encourages and supports suppliers and partners in adopting and implementing additional principles and policies, provided they are consistent with this policy and do not conflict with it.
2. The Company's business operations strictly adhere to local laws and regulations. In the absence of specific local legal or regulatory requirements, this code shall be followed.
3. This document is to be interpreted and revised by Tongwei Co., Ltd. The Company will update this document as appropriate based on domestic and international policies, regulatory requirements, and industry developments. In case of any discrepancy between the Chinese and English versions of this document, the Chinese version shall prevail.