

南京钢铁股份有限公司 负责任的人工智能使用政策

Nanjing Iron & Steel Co., Ltd. Responsible Artificial Intelligence Use Policy

第一条 总则

南京钢铁股份有限公司（以下简称“公司”）为规范人工智能（以下简称 AI）技术的研发与应用活动，保障相关活动合规、安全、合乎伦理地开展，守护各方合法权益，推动 AI 技术负责任发展，结合公司运营实际，制定本政策。

I General Rules

Nanjing Iron & Steel Co., Ltd. (hereinafter referred to as the "Company") formulates this Policy to regulate the research and application activities of artificial intelligence (hereinafter referred to as AI), ensure that related activities are conducted in a compliant, safe, and ethical manner, protect the legitimate rights and interests of all parties, and promote the responsible development of AI technology, in line with the actual operations of the Company.

第二条 适用范围

本政策适用于公司及下属企业 AI 相关活动的全流程。同时，公司倡导供应商、承包商，以及其他业务伙伴在其运营中遵循负责任 AI 使用原则。

II Coverage

This Policy applies to the entire process of AI-related activities conducted by the Company and its subsidiaries. At the

same time, the Company advocates suppliers, contractors, and other business partners to follow the principles of responsible AI use in their operations.

第三条 负责任 AI 使用原则

III Principles of Responsible AI Use

数据隐私保护：公司在 AI 开发与使用全流程中严格遵守各辖区数据保护相关法规，始终尊重并保障数据隐私权益。在开展 AI 训练、AI 应用落地工作时，公司将严格落实数据收集最小必要原则，对所有涉及个人信息的数据落实合规授权与加密保护措施，绝不违规泄露、滥用相关数据。

Data Privacy Protection: The Company strictly complies with data protection regulations in all jurisdictions throughout the entire process of AI development and use, always respecting and safeguarding data privacy rights. When conducting AI training and implementing AI applications, the Company will strictly adhere to the principle of data collection minimization, ensuring compliance authorization and encryption protection measures for all data involving personal information, and will never violate regulations by leaking or misusing related data.

网络安全防护：公司将 AI 相关系统的网络安全纳入公司全体系安全管理框架，在 AI 开发、部署、运行全周期开展常态化安全防护、漏洞排查与应急演练，持续升级安全防护能力，防范 AI 系统被非法入侵、篡改、滥用，保障 AI 系统稳定安全运行。

Cybersecurity Protection: The Company incorporates the

cybersecurity of AI-related systems into its overall security management framework, conducting regular security protection, vulnerability checks, and emergency drills throughout the entire cycle of AI development, deployment, and operation, continuously upgrading security capabilities to prevent illegal intrusion, tampering, and misuse of AI systems, ensuring the stable and secure operation of AI systems.

算法偏见防控：公司严格规避 AI 开发与使用过程中的潜在偏见，避免 AI 决策因训练数据缺陷或算法设计偏差产生歧视性、不公平结果。

Algorithm Bias Prevention: The Company strictly avoids potential biases in the development and use of AI, preventing AI decisions from resulting in discriminatory or unfair outcomes due to defects in training data or biases in algorithm design.

人工监督与主导：针对所有涉及核心利益的关键性决策，公司允许并保障人对 AI 输出结果进行干预与修正，绝不将 AI 系统设置为关键决策的唯一主体，确保人对 AI 相关关键决策的最终控制权。

Human Supervision and Leadership: For all key decisions involving core interests, the Company allows and ensures human intervention and correction of AI output results, never setting the AI system as the sole entity for critical decision-making, thus ensuring human ultimate control over AI-related key decisions.

透明与可解释：公司保证所开发、使用的 AI 系统具备可追溯性，对于 AI 输出的结果与做出的决策，公司将按照监管要求与利益相关方需求，保障 AI 应用的透明。

Transparency and Explainability: The Company

guarantees that the AI systems developed and used are traceable. For the results produced by AI and the decisions made, the Company will ensure transparency in AI applications according to regulatory requirements and stakeholder needs.

明确责任归属：公司针对 AI 模型与 AI 工具产生的所有输出结果，建立了清晰的分层问责体系，明确不同场景下 AI 应用结果的责任主体，以妥善处理可能出现的意外后果或损害。

Clear Responsibility Attribution: The Company has established a clear hierarchical accountability system for all output results generated by AI models and tools, specifying the responsible parties for AI application results in different scenarios to properly address any potential unexpected consequences or damages.

划定 AI 活动清晰边界：公司基于法律法规、伦理准则与业务合规要求，明确划定 AI 可以开展与禁止开展的活动边界，所有 AI 开发、应用活动严格在合规边界内推进，绝不越界开展违规 AI 研发与应用。

Defining Clear Boundaries for AI Activities: Based on legal regulations, ethical guidelines, and business compliance requirements, the Company clearly delineates the boundaries of activities that AI can and cannot engage in. All AI development and application activities are strictly conducted within compliance boundaries, and no non-compliant AI research and application are permitted.

践行绿色低碳要求：公司要求自有及合作第三方的 AI 数据中心、AI 模型落实低碳运营要求，通过技术优化、能耗管理降低 AI 相关活动的生态足迹，推动 AI 业务绿色发展。

Practicing Green and Low-Carbon Requirements: The Company requires its own and third-party AI data centers and AI models to implement low-carbon operational requirements, reducing the ecological footprint of AI-related activities through technological optimization and energy consumption management, promoting green development of AI business.

禁止违规 AI 应用：公司绝不使用、部署任何涉及操纵行为、利用个体或系统漏洞、开展社会评分，或未经授权开展生物特征监测的 AI 系统。

Prohibition of Non-Compliant AI Applications: The Company will never use or deploy any AI systems that involve manipulation, exploit individual or system vulnerabilities, conduct social scoring, or engage in biometric monitoring without authorization.

第四条 负责任 AI 使用关键实践

为切实履行上述原则，公司承诺落实以下关键实践：

IV Key Practices for Responsible AI Use

To effectively implement the above principles, the Company commits to the following key practices:

治理和监督：公司设立人工智能研究院协同风险合规部为所有 AI 项目提供正式监督与战略指导。建立由董事会战略与 ESG 委员会下属数字安全管理委员会统筹领导，人工智能研究院统筹协调、各业务单位协同配合的管理体系。人工智能研究院定期召开会议，监督 AI 系统运行、评估风险，并确保符合适用的法律要求。决策采用基于风险的方法，重大事项将上报

高级管理层及董事会。

Governance and Supervision: The Company has established an Artificial Intelligence Research Institute in collaboration with the Risk Compliance Department to provide formal oversight and strategic guidance for all AI projects. A management system has been established, led by the Digital Security Management Committee under the Board's Strategy and ESG Committee, with the Artificial Intelligence Research Institute coordinating and various business units cooperating. The Artificial Intelligence Research Institute holds regular meetings to supervise AI system operations, assess risks, and ensure compliance with applicable legal requirements. Decision-making adopts a risk-based approach, and significant matters will be reported to senior management and the Board.

AI 影响评估：公司开展 AI 影响评估，识别并应对那些对个人权益、客户成果或监管合规产生实质性影响的 AI 系统所存在的潜在风险。该流程包含公平性审查，确保 AI 决策公平无歧视。

AI Impact Assessment: The Company conducts AI impact assessments to identify and address potential risks posed by AI systems that have a substantial impact on individual rights, customer outcomes, or regulatory compliance. This process includes fairness reviews to ensure that AI decisions are fair and non-discriminatory.

敏感 AI 能力访问管控：公司对具备人脸识别、监控等敏感属性的 AI 能力，建立了分级授权访问机制，仅对具备对应业务权限与合规资质的岗位开放访问权限，严格管控敏感 AI

能力的使用范围，防范违规滥用。

Sensitive AI Capability Access Control: The Company has established a tiered authorization access mechanism for AI capabilities with sensitive attributes such as facial recognition and monitoring. Access rights are granted only to positions with corresponding business permissions and compliance qualifications, strictly controlling the scope of sensitive AI capability usage to prevent unauthorized abuse.

AI 生成内容与结果标注：公司要求所有对外输出的 AI 生成内容，都必须添加清晰可识别的 AI 来源标注，明确区分人工产出内容/决策与 AI 产出内容/决策，保障信息透明，不对公众和合作方造成误导。

AI-Generated Content and Result Labeling: The Company requires that all externally output AI-generated content must include a clear and recognizable AI source label, distinctly differentiating between human-generated content/decisions and AI-generated content/decisions, ensuring information transparency and preventing misleading information from being provided to the public and partners.

AI 模型偏移与退化校正：公司针对所有已投入运行的 AI 模型建立了全生命周期性能监控机制，定期检测模型输出精度、决策逻辑的偏移与退化，同步建立了模型迭代校正流程，可及时对发生偏移/退化的 AI 模型完成修正，保障 AI 模型运行稳定合规。

AI Model Drift and Degradation Correction: The Company has established a full lifecycle performance monitoring mechanism for all AI models in operation, regularly checking for

drift and degradation in model output accuracy and decision logic. A model iteration correction process has also been established to promptly correct any AI models that experience drift/degradation, ensuring stable and compliant operation of AI models.

AI 领域低碳减排：公司推进并联动合作供应商开展 AI 降碳行动，通过算力调度优化、模型结构轻量化、绿色能源使用等方式，降低 AI 数据中心与 AI 模型的能耗，推动 AI 绿色发展。

Low-Carbon Emission Reduction in AI: The Company promotes and collaborates with suppliers to carry out AI carbon reduction actions, reducing energy consumption of AI data centers and AI models through power scheduling optimization, lightweight model structures, and the use of green energy, thereby promoting green development in AI.

AI 决策申诉处理机制：公司建立了完善的 AI 决策结果申诉通道，面向用户及受 AI 决策影响的第三方开放申请，对相关方提出异议的 AI 决策或结果，将由独立审核团队完成人工复核并给出修正结果，保障相关方合法权益。

AI Decision Appeal Handling Mechanism: The Company has established a comprehensive appeal channel for AI decision results, open to users and third parties affected by AI decisions. Any objections raised by relevant parties regarding AI decisions or results will be reviewed by an independent audit team, which will conduct a manual review and provide corrective results, safeguarding the legal rights of relevant parties.

员工 AI 伦理与安全培训：公司为员工提供持续的培训与宣导，帮助其掌握负责任的 AI 应用原则、最佳实践以及不断

演进的伦理规范与监管要求。

Employee AI Ethics and Safety Training: The Company provides ongoing training and advocacy for employees to help them master responsible AI application principles, best practices, and the evolving ethical standards and regulatory requirements.

AI 影响量化：公司将 AI 项目、AI 工具对可持续发展目标产生的影响纳入量化评估体系，从能耗、碳排放、社会效益等多维度完成量化核算，以此优化 AI 项目设计，推动 AI 使用赋能可持续发展。

Quantification of AI Impact: The Company incorporates the impact of AI projects and tools on sustainable development goals into a quantitative assessment system, completing quantitative accounting from multiple dimensions such as energy consumption, carbon emissions, and social benefits, thereby optimizing AI project design and promoting the use of AI to empower sustainable development.

第五条 附则

本政策经公司董事会战略与 ESG 委员会审议批准，自发布之日起生效。

V Annex

This Policy has been reviewed and approved by the Strategy and ESG Committee of the Board of Directors of the Company, which is effective as of the date of promulgation.