



通威股份有限公司负责任人工智能管理承诺与政策

通威股份有限公司（以下简称“通威股份”或“公司”）始终坚持商业价值与社会责任并重，致力于以更高的道德标准开展业务。公司高度重视人工智能技术的负责任开发与应用，严格遵守《中华人民共和国网络安全法》《中华人民共和国数据安全法》《中华人民共和国个人信息保护法》《生成式人工智能服务管理暂行办法》等相关法律法规及相关规定要求，参考 ISO/IEC 42001 人工智能管理体系、OECD 人工智能原则等国际标准与最佳实践。本政策旨在规范人工智能系统的开发、部署与使用，保障人工智能应用过程的透明性、公平性与可问责性，防范潜在风险，切实维护相关方合法权益。

一、适用范围

本政策适用于通威股份及旗下分子公司的各项业务与运营活动。公司董事、高级管理人员及全体员工在开发、采购、部署、使用、维护或管理人工智能系统过程中均应遵守本政策。公司同时倡导涉及人工智能相关产品、平台、模型或服务的供应商、合作伙伴及其他第三方，应遵循本政策要求，并按照合同约定落实相关责任。

二、责任部门

公司董事会为人工智能管理最高决策和监督机构，负责监督、审议人工智能管理重大事项，批准相关政策与承诺的发布。信息部负责统筹推进人工智能管理体系建设与认证（如 ISO/IEC 42001 人工智能管理体系认证），负责 AI 安全合规审查、隐私影响评估（PIA/DPIA）审核及 AI 安全事件应急响应处置，对接监管机构满足合规性要求。各业务部门负责本部门人工智能应用需求提出、业务风险识别、应用管理及效果评估，确保人工智能应用符合业务管理要求。

三、负责任人工智能管理承诺与政策

- （一） 数据隐私保护：**公司在人工智能系统的开发与使用过程中，严格遵守数据隐私保护相关法律法规。在采集、处理或使用个人数据训练人工智能模型时，必须遵循合法、正当、必要和诚信原则，明确告知数据主体数据收集的目的、方式与范围，并获得其知情同意。公司承诺不收集与人工智能应用目的无关的个人信息，并在数据使用过程中采取去标识化、匿名化等技术措施，最大限度降低隐私泄露风险。
- （二） 系统网络安全：**公司将人工智能系统的网络安全保障纳入整体信息安全管理体系统。针对人工智能系统特有的安全风险（包括但不限于模型投毒、对抗性攻击、提示注入、数据泄露等），建立专门的风险评估与防控机制。人工智能系统的开发与部署须经过安全审查，确保其具备足够的稳健性与抗攻击能力。
- （三） 偏见预防与公平性保障：**公司致力于确保人工智能系统在实际应用过程中避免因种族、性别、年龄、地域、宗教信仰等因素而产生歧视。公司所使用的基础模型，均已由其提供方在数据选择、算法设计及输出结果评估等环节建立偏见识别与消除机制，并定期完成公平性评估与偏见检测，符合相关管控要求。
- （四） 生态足迹管理：**公司在部署人工智能系统时，关注其对环境的影响。对于需要大量计算资源的人工智能模型训练与推理任务，公司综合考量算力资源与成本效益，采取优化措施降低计算负载和算力消耗，从而避免能源浪费。公司鼓励在人工智能基础设施中优先采用节能技术与可再生能源，优先采用绿色数据中心和低碳云服务。公司持续关注并评估人工智能工具及项目对公司可持续发展目标的影响，推动人工智能应用的生态效益可衡量、可追踪，促进人工智能与公司绿色低碳转型方向的协同发展。
- （五） 人为监督和干预机制：**公司明确人工智能系统的最终责任主体为人类决策者。任何人工智能系统的开发、部署与使用均须指定明确的责任人，负责该系统的全生命周期管控与管理工作。确保在缺乏人类监督以及人工审查或否决机制的情况下，人工智能系统无法做出不可逆或高风险的决策。公司建立人工智能系统性能的持续监测机制，定期检视智能体输出质量与决策准确性，对于识别出的推理、召回和输出的准确性问题，及时采取优化措施，确保人工智能系统持

续稳定运行。

(六) 申诉与问责机制：公司建立受影响方对人工智能决策或结果的异议受理与申诉机制，明确申诉渠道、处理流程与反馈时限，确保因人工智能系统自动决策而受到影响的用户、客户及其他相关第三方能够便捷地提出质疑并获得妥善处理。公司对于因人工智能系统错误、偏差或失效导致的负面影响，能够及时启动责任追溯与补救措施。

(七) 透明度与可解释性：公司承诺在人工智能系统的使用过程中坚持透明度原则。

(1) 对于直接面向客户或影响客户权益的人工智能应用，明确告知用户其正在与人工智能系统交互。(2) 公司致力于提升人工智能决策过程的可解释性，确保在关键应用场景中，人工智能生成的结果或决策能够以可理解的方式向相关方进行解释。(3) 对于人工智能生成的内容，采取适当方式（如水印、标识或声明）加以标注，确保用户能够识别内容的来源。

(八) 使用边界与风险管控：公司严格遵循法律法规及伦理准则，明确人工智能应用的适用范围与功能边界，确保人工智能技术的使用不超出其设计目的与合规要求。公司承诺不使用或部署从事操纵行为、利用人类弱点进行剥削、实施社会评分或未经授权进行生物特征监控的人工智能系统。公司建立人工智能系统的权限管控机制，对涉及人脸识别、监控及其他高敏感度的人工智能能力实施严格的访问授权与使用管理，确保此类技术的应用严格限定于合法、合规且必要的业务场景。

(九) 员工培训与意识提升：公司定期开展人工智能伦理与安全培训，使全体员工了解负责任人工智能的基本原则、潜在风险以及各自在人工智能治理中的职责。对于从事人工智能系统开发、部署与运维的技术人员，公司提供专项培训，确保其掌握人工智能安全开发与风险防控的专业知识与技能。

(十) 持续改进与审计：公司每年开展负责任人工智能管理的内部审核，对内部人工智能系统的合规性、安全性与公平性进行全面检查与评估。同时，公司定期邀请具备资质的第三方机构开展外部评估与认证（如 ISO 42001 等），针对发现的问题及时整改，持续提升负责任人工智能管理水平。公司亦通过管理评审识别管理体系改进的需求，制定改进措施并验证其结果。

(十一) AI 影响评估：公司在引入、部署或重大更新人工智能系统前，会开展全面

的 AI 影响评估 (AI Impact Assessment)，评估内容包括：对个人权益、隐私、安全、公平的潜在影响；对数据安全、完整性和可用性的风险评估；对个人信息处理活动的隐私影响评估；法律法规合规性审查。

(十二) AI 供应链与第三方服务管理：公司采购或集成第三方 AI 模型、API 服务、开源模型或 AI 组件时，会进行安全评估和合规审查，包括：模型许可证合规性审查；训练数据来源合法性验证；安全漏洞与已知风险扫描；第三方服务提供商的安全资质（如 ISO 27001、ISO 42001 认证情况）。对于关键业务场景使用的第三方 AI 服务，签订服务水平协议（SLA）和安全保密协议，明确数据安全、模型性能和应急响应责任。



CEO：刘舒琪

通威股份有限公司

日期：2026 年 4 月

备注：

- 1、公司鼓励并支持供应商及合作伙伴在遵循本政策的前提下，采纳和执行额外的原则与政策，但这些额外原则和政策都不得与本政策产生冲突。
- 2、公司业务开展严格遵循所在地法律法规要求，若当地无明确法律法规要求时，执行本准则。
- 3、该文件由通威股份有限公司解释和修订，公司会适时根据国内外政策、监管要求和行业发展情况对文件进行更新，当文件中英文版有冲突时，请以中文版为准。