

朗新科技集团股份有限公司

负责任人工智能（AI）政策

为规范人工智能（Artificial Intelligence，以下简称“AI”）技术在朗新科技集团股份有限公司（以下简称“公司”或“朗新科技”）及下属企业经营管理、产品研发、客户服务、智慧能源、数字生活、企业运营等场景中的开发、采购、部署和使用，推动 AI 安全、合规、审慎、透明、绿色应用，防范数据安全、网络安全、算法偏见、知识产权、商业秘密、伦理合规及声誉风险，保障员工、客户、用户、供应商、合作伙伴及其他利益相关方合法权益，结合公司实际，制定本政策。

一、原则

公司坚持依法合规、以人为本、安全可控、公平公正、透明可信、责任可追溯、绿色低碳和持续改进的原则，统筹推进 AI 创新应用与风险防控，促进 AI 赋能公司高质量发展。

公司认为 AI 应服务于人的福祉、客户价值和社会可持续发展。AI 不得被用于削弱人的尊严、权利和自主判断，不得替代依法应由人员、组织或管理机构作出的最终判断和责任承担。

二、适用范围

本政策适用于朗新科技集团及其下属公司、分支机构和受公司实际管理的相关组织，以及在公司业务场景中使用、开发、采购、接入、部署、运营 AI 系统、模型、工具、接口及相关服务的全体员工、外包人员、实习人员和其他受公司管理的人员。

本政策所称 AI 包括但不限于机器学习、深度学习、自然语言处理、计算机视觉、语音识别、推荐算法、预测分析、智能客服、代码生成、内容生成、智能决策辅助、自动化流程和生成式 AI 等技术及其相

关系统或服务。供应商、合作伙伴及其他第三方在向公司提供 AI 相关产品、服务或数据处理活动时，应遵守本政策及公司相关制度要求。

三、治理职责

1. 公司建立 AI 治理与责任管理机制，明确业务部门、研发技术部门、AI 研究院、数据治理部门、网络与信息安全部门、法务合规部门、采购管理部门、人力资源部门及其他相关部门的职责边界，确保 AI 应用过程可审批、可记录、可追溯、可监督、可问责。
2. 业务使用部门对 AI 在本业务场景中的必要性、适当性、输入内容合规性、输出结果使用及相关影响负责；研发技术部门负责 AI 系统设计、测试、上线、变更和运维过程中的技术控制；网络与信息安全、数据治理和法务合规等职能部门按照职责分工做好安全评估、数据合规、合同审核、知识产权、供应商管理、监督检查和风险处置。
3. 涉及重大经营决策、客户或用户权益、关键业务系统、重要数据处理、对外发布内容或其他高风险场景的 AI 应用，应按照公司制度履行必要的评估、审批和复核程序。公司根据监管要求、业务变化和技术发展，持续完善 AI 治理框架及配套流程。
4. 本政策由公司董事会或董事会授权的 ESG 可持续发展委员会审议批准。董事会或其授权机构对本政策的执行、监督、定期检讨和重大修订承担整体责任；公司高管层和相关职能部门按照职责分工负责本政策的具体落实、解释、宣贯和日常管理。
5. 涉及重大修订、重大 AI 风险事项或公司治理要求的事项，应按公司治理程序提交董事会或高管层审议。公司应在公开披露文件中说明本政策或相关承诺的最高审批机构，以便利益相关方核验。

四、数据隐私与商业秘密保护

1. 公司尊重并保护 AI 开发、训练、测试、部署、调用和运营全生命周期中涉及的个人信息、重要数据、商业秘密、源代码、客户资料、经营信息及其他依法或依约受保护的数据。公司严格遵守网络

安全、数据安全、个人信息保护、保密管理及相关监管要求；涉及境外业务或跨境场景的，还应关注并遵守适用的 GDPR、CCPA 等个人信息和隐私保护规则。

2. 公司对 AI 相关数据实施分类分级、最小必要、最小化采集、授权访问、脱敏匿名、加密存储、留痕审计和安全销毁等管理措施。在具备适用条件的场景中，公司鼓励采用差分隐私、联邦学习、安全多方计算、数据脱敏、匿名化和其他隐私增强技术，降低 AI 训练、测试和推理过程中的个人数据暴露风险。
3. 未经授权或合规评估，不得将个人信息、敏感个人信息、客户数据、供应商数据、员工数据、未公开财务信息、商业计划、合同文本、源代码、系统凭证、漏洞信息、内部会议材料及其他受保护信息输入外部 AI 工具、第三方模型或无法确认安全边界的 AI 服务。
4. 公司在使用数据训练、微调或优化 AI 模型时，应确保数据来源合法、授权清晰、处理目的明确、使用范围受控，并按照适用规则保留必要记录。对于涉及个人信息主体权利、数据跨境、第三方委托处理或共同处理的场景，应按公司相关制度开展合规审查。

五、网络与信息安全

1. 公司将 AI 系统、模型、插件、接口、数据集、提示词模板、自动化流程及相关服务纳入网络安全和信息安全管理体系，落实账号管理、身份认证、权限控制、密钥保护、加密传输与存储、日志留存、安全监测、漏洞管理、应急响应和备份恢复等控制措施。
2. 公司关注 AI 应用可能引发的提示词注入、越权调用、数据泄露、模型幻觉、恶意代码生成、内容污染、模型滥用、供应链攻击、深度伪造及其他安全风险。对于关键系统、对外服务或高影响场景，应开展必要的安全测试、渗透测试、安全编码检查、对抗测试、红队验证或专项评估，并根据评估结果采取防护、限制、隔离、回滚或停止使用等措施。

3. 对于采购或合作使用的第三方 AI 平台、模型、云服务、数据服务和开发工具，公司应关注其安全能力、数据处理规则、服务稳定性、日志审计、模型更新、漏洞响应、数据留存与删除机制，防范因第三方服务引发的信息安全和业务连续性风险。

六、公平公正与偏见防范

1. 公司关注 AI 在数据、模型、规则、训练过程及输出结果中可能存在的偏见、歧视、误导和不公平影响。AI 应用不得基于民族、种族、性别、年龄、宗教、残疾、健康状况、地域、职业身份或其他依法受保护特征作出不当差别对待。
2. 对于可能影响员工招聘、绩效评价、客户服务、用户权益、供应商选择、信用或风险判断、资源分配及其他利益相关方权益的 AI 应用，公司应开展适当的公平性、准确性、代表性和影响评估，必要时采取人工复核、阈值限制、样本校正、结果解释和纠偏措施，避免因 AI 输出失当造成不公平结果。
3. 公司应根据 AI 应用风险等级开展算法偏见识别、训练数据偏见缓解和公平性审计。对于已部署且可能产生重大影响的 AI 模型，公司应定期检查样本覆盖、输出差异、误判率和投诉反馈等信息，并将评估结论、纠偏措施和责任部门纳入留痕管理。
4. 公司鼓励在 AI 产品和服务设计中考虑可访问性、包容性和不同群体的使用差异，使 AI 能力以平等、便利、可理解的方式服务员工、客户、用户及其他相关方。

七、透明度与可解释性

1. 公司重视 AI 应用的透明度和可解释性。对于用户、客户、员工或其他相关方可能直接接触或受到显著影响的 AI 系统，公司应根据场景提供适当标识、说明或告知，使相关方能够了解其正在与 AI 系统交互，或了解相关内容、分析、建议、决策辅助中存在 AI 参与。

2. 公司应在适当范围内说明 AI 系统的主要用途、能力边界、数据来源、限制条件、可能风险和人工复核方式。对于重要 AI 模型、数据集、评估过程、上线审批、版本变更和关键输出，公司应保留必要记录；在条件成熟时，可采用模型说明、数据说明、评估记录、使用指南或其他形式提升可解释性和可审计性。
3. 对于缺乏充分依据、逻辑异常、来源不明、无法合理解释或与专业判断明显冲突的 AI 输出结果，公司坚持审慎使用，不得直接作为最终依据。对外发布或用于重要业务的 AI 生成内容，应经过适当的事实核查、合规审核和人工确认。
4. 对于 AI 生成文本、图像、音视频、代码或由 AI 驱动形成的重要决策结果，公司应根据适用场景进行明确标注或说明，帮助用户、客户、员工及其他相关方区分人类创作、人工审核内容与 AI 生成或 AI 辅助形成的内容。

八、人类监督与人工干预

1. 公司坚持在关键事项中保留必要的人类监督、人工判断和人工干预机制。AI 生成或推荐的结果仅作为辅助工具，不当然替代员工、管理人员、专业机构或公司治理机构的职责、判断和最终决策。
2. 对于涉及重大经营决策、财务披露、法律合规、信息安全、数据合规、知识产权、采购准入、客户或用户重大权益、能源运行安全、产品质量、安全生产、人力资源管理及其他重要敏感事项，AI 输出不得单独作为最终决定依据，应由具备相应权限和专业能力的人员进行审查、核查、修正或确认。
3. 公司应在关键决策环节明确人工监督者、审批权限、干预权限、否决权限和升级流程。监督者应能够查看必要的 AI 输入、输出、依据和运行记录，并在发现 AI 输出不准确、不公平、不合规或超出授权范围时，及时作出修正、否决、升级、暂停或终止处理。

4. 公司应根据具体场景建立人工干预、暂停、终止、回滚和补救机制。当 AI 系统出现异常输出、风险事件、投诉争议、安全隐患或与法律法规、伦理规范、公司制度不一致的情况时，相关部门应及时采取处置措施。

九、全生命周期风险管理

1. 公司对 AI 应用实施覆盖需求提出、方案设计、数据准备、模型选择、开发测试、采购接入、上线部署、运营监控、版本变更、下线退出等环节的全生命周期管理。不同风险等级的 AI 场景应匹配相应的评估、审批、测试、记录和复核要求。
2. AI 系统上线或重大变更前，应根据场景开展必要的合规性、安全性、准确性、鲁棒性、公平性、可解释性、业务连续性和生态影响评估。对生成式 AI、高影响自动化、关键系统接入、重要数据处理等场景，应加强测试验证和使用限制，必要时开展红队测试、灰度发布或试运行。
3. AI 系统运营期间，相关部门应持续监测其性能、准确率、稳定性、输出质量、安全事件、用户反馈、误用滥用、偏见风险和外部监管变化。对于模型性能漂移、准确率下降、输出退化或适用场景变化等情况，应及时开展原因分析，并采取重新评估、参数调整、数据更新、模型回滚、限制功能、暂停服务或下线等纠正措施。
4. 公司应建立定期合规审计和运行复盘机制，对重要 AI 系统的审批记录、访问记录、测试记录、模型变更、风险事件、投诉处理和纠偏措施进行检查。公司应保留适当的运行记录和风险处置记录，以支持监督检查、公开披露和责任追溯。

十、应用边界管理

1. 公司坚持合理、适度、审慎使用 AI，明确 AI “可为”与“不可为”的边界。AI 可用于提升工作效率、辅助研发测试、知识检索、文本初稿、数据分析、流程自动化、客户服务辅助、风险提示和运营优化等场景，但其使用应符合授权范围、数据要求、保密要求和人工复核要求。

2. 员工在使用 AI 工具时，应对输入内容、提示词、上传文件、输出结果和后续使用行为负责。未经批准，不得使用个人账号、非授权外部工具或不符合公司安全要求的 AI 服务处理公司业务信息；不得绕过公司安全策略、权限控制、日志审计和审批流程。
3. 公司对敏感 AI 功能实施更严格的访问管控和监督机制。涉及人脸识别、身份识别、监控技术、生物特征处理、自动化重大决策、高影响推荐或其他可能对个人权益产生重大影响的 AI 功能，应明确部署主体、授权范围、适用场景、访问权限、日志审计、人工复核和退出机制。
4. AI 应用不应被设计、部署或使用于欺骗、操纵、诱导、冒充他人、生成虚假信息、规避监管、破坏系统安全、侵犯知识产权、侵犯个人信息权益、损害商业秘密、实施不正当竞争或其他违反法律法规、伦理规范及公司制度的目的。

十一、生成内容与知识产权

1. 公司使用 AI 生成文本、图片、音视频、代码、报告、方案、营销材料或其他内容时，应充分关注事实准确性、来源可靠性、版权、商标、专利、商业秘密、肖像权、名誉权、数据授权和对外传播风险。AI 生成内容在对外发布、提交客户、用于合同附件、进入产品代码库或形成重要内部决策材料前，应经过适当的人工审核。
2. 员工不得将未经授权的第三方作品、客户资料、受保密义务限制的材料、受许可协议约束的代码或数据输入 AI 系统进行训练、生成、改写或传播。对于 AI 辅助生成的代码、方案或内容，应根据公司知识产权、开源合规、信息安全和质量管理要求进行检查，避免引入侵权、漏洞、恶意代码、许可证冲突或不适当内容。
3. 公司在必要场景中可对 AI 生成内容进行标识、留痕或归档，并根据法律法规、客户合同、平台规则和公司制度要求，保留相关审核记录和版本记录。

十二、第三方 AI 与供应链管理

1. 公司采购、接入或合作使用第三方 AI 产品、模型、数据集、插件、接口、云服务、开发工具和咨询服务时，应根据风险程度开展供应商尽职调查和准入评估，重点关注其合法合规、数据安全、个人信息保护、网络安全、模型安全、内容安全、知识产权、服务稳定性、环境影响和持续支持能力。
2. 涉及 AI 的合同或合作文件应根据实际需要明确数据权属、数据处理目的和范围、保密义务、个人信息保护、模型训练或再利用限制、安全事件通知、审计配合、服务变更、退出删除、知识产权归属、责任承担及其他必要条款。
3. 公司鼓励供应商、合作伙伴和开发者遵循负责任 AI 原则，共同防范 AI 技术滥用及其对个人、组织、社会和环境造成的不利影响。对于严重违反法律法规、合同约定或公司 AI 治理要求的第三方，公司可视情况采取整改、限制使用、暂停合作或终止合作等措施。

十三、绿色低碳要求

1. 公司倡导在 AI 使用、开发和采购过程中兼顾资源利用效率与生态环境影响。在满足安全、合规、质量和业务需求的前提下，公司优先关注能效较高、资源消耗较低、环境影响较小的模型、平台、数据中心、算法方案和调用方式。
2. 公司鼓励通过合理选型、适度调用、任务分级、缓存复用、模型压缩、边云协同、算力资源优化、减少无效生成和优化技术以降低计算负载等方式，减少 AI 运行过程中的能源浪费、水资源使用和碳强度。
3. 对于自有或第三方 AI 数据中心、云服务、模型服务和算力资源，公司鼓励并逐步推动采用节能基础设施、绿色数据中心、可再生能源、能耗监测和碳排放核算等措施，降低 AI 训练、部署和运营环节的生态足迹，推动 AI 能力服务公司绿色低碳和可持续发展目标。

十四、禁止性要求

1. 公司禁止使用、部署或开发违反法律法规、监管要求、伦理规范及公司制度的 AI 系统或应用，包括但不限于通过操纵、诱导、欺骗等方式影响他人自主判断和行为的系统；利用特定群体脆弱性实施剥削、误导或不当影响的系统；用于社会评分、社会排序或基于不当标准实施差别对待的系统。
2. 公司禁止未经合法授权开展生物特征识别、面部识别监控、身份追踪或其他可能侵害个人权益的 AI 应用；禁止生成、传播虚假信息、违法有害内容、恶意代码、钓鱼材料、深度伪造内容或其他可能危害网络安全、公共安全、商业信誉和个人合法权益的内容。
3. 公司禁止将 AI 用于规避审计、绕过权限、逃避监管、隐瞒风险、伪造记录、替代必要人工审批或其他违反公司治理要求的行为。公司参照《欧盟人工智能法案》等国际监管框架中关于受限或禁止 AI 系统的要求，通过场景准入审查、权限控制、日志审计、定期检查和违规处置等措施，防止相关 AI 系统被使用或部署。
4. 任何人员发现 AI 应用存在上述情形或重大风险隐患的，应及时停止相关使用并向主管部门或相关职能部门报告。公司将根据事件性质和影响范围开展调查、整改和责任追究，并可在适当范围内复盘相关防控措施的执行效果。

十五、异议、申诉与复核机制

1. 对于 AI 生成结果、AI 辅助分析或 AI 参与形成的决定可能对员工、客户、用户、供应商、合作伙伴或其他受影响方产生重大影响的，公司支持相关方按照适用规则提出异议、申诉或复核申请。
2. 公司将结合具体场景开展人工复核，并根据需要采取解释说明、结果修正、补充评估、中止适用、重新处理或其他必要措施。公司在处理相关异议、申诉和复核事项时，应保护相关方个人信息、商业秘密和合法权益，避免因申诉或反馈受到不当影响。
3. 对于 AI 造成或可能造成错误、损害、歧视、不当拒绝服务、权益受限或其他不利影响的情形，公司应明确受理渠道、调查责任部门、处理时限、纠正措施和反馈方式，确保 AI 系统的责任归属和问题处置过程可追溯。

十六、实施机制与绩效指标

1. 公司将负责任 AI 要求嵌入 AI 系统设计、研发、采购、部署、运营、迭代和下线全生命周期，形成与业务场景相匹配的实施方案、操作标准、检查清单和考核要求。实施机制可包括敏感 AI 功能访问限制、AI 生成内容标注、模型性能监控、算法偏见评估、生态足迹管理、异议申诉流程、员工培训和定期合规审计等。
2. 公司对重要 AI 项目建立绩效和影响跟踪机制，根据数据可得性、财务核算规则和披露审批要求，跟踪最近报告年度 AI 相关投资额、AI 带来的营收增长、成本降低、效率提升、AI 投资回报率，以及 AI 对碳减排、资源效率提升、客户服务质量、风险事件降低、员工培训覆盖率等可持续发展目标的贡献。
3. AI 相关绩效指标应由业务部门、财务部门、数据治理部门和相关职能部门共同核验，避免夸大 AI 贡献或重复计算。对于可公开披露的指标，公司应在年度报告、可持续发展报告、官网或其他正式出版物中以可核验口径进行披露；对于尚不具备披露条件的指标，应持续完善数据口径、内部记录和核验流程。
4. 公司可根据 AI 管理成熟度、客户要求、监管趋势和外部评级要求，评估引入第三方评估、外部鉴证等人工智能管理体系相关标准，持续提升 AI 治理机制的可信度和可验证性。

十七、培训与持续改进

1. 公司将结合业务发展、技术进步和监管变化，持续完善 AI 治理机制，并适时开展 AI 合规、安全、伦理、数据保护、知识产权、生成内容审核和负责任使用相关培训，提升员工规范使用 AI 的意识和能力。
2. 对于从事 AI 研发、数据处理、安全管理、产品运营、客户服务、采购管理、法务合规及其他关键岗位的人员，公司可根据岗位特点设置专项培训、使用指引、检查清单或操作规范。培训内容可涵

盖 AI 伦理原则、数据隐私、网络安全、算法偏见、人工监督、生成内容标注、提示词安全、商业秘密保护、知识产权和违规案例等。

3. 公司应根据培训计划记录培训主题、覆盖对象、参与情况和改进反馈，并将 AI 伦理使用及安全防护培训纳入 AI 治理持续改进的重要内容。对于 AI 使用过程中发现的问题、风险和缺陷，公司将及时整改、优化并持续改进。

十八、附则

本政策由公司相关职能部门负责编制和日常管理，经公司董事会或董事会授权的 ESG 可持续发展委员会审议批准后执行。公司既有 AI、数据安全、网络安全、个人信息保护、信息披露、商业秘密、知识产权、采购、合规、人力资源及其他相关制度与本政策不一致的，按照更严格要求执行；法律法规、监管规定或客户合同另有更高要求的，从其规定。

公司相关职能部门按照职责分工负责本政策的宣贯、解释、执行监督和定期检讨。任何本政策外的特殊情况，除公司另有规定外，均需按照公司审批流程报经相应管理机构批准。

本政策自发布之日起施行，并应通过公司官网、年度报告、可持续发展报告、综合报告或企业正式出版物等公开渠道进行披露。公司在对外评级、问卷或利益相关方沟通中引用本政策时，应注明公开披露位置、页码或网址；内部草案、邮件、未发布 PPT 等材料不作为本政策公开披露的证明文件。

公司将根据法律法规、监管要求、技术发展、业务实践、外部评级要求和利益相关方反馈，对本政策进行适时修订和完善。