

FOSUN 复星

复星集团

负责任 AI 应用安全管理与使用制度

[IT_011_V1.1]

[集团数智化与 AI 条线信息安全团队、集团风控条线法律事务部]

Fosun Group
2026年8月27日

第一章 总则

第一条 为确保集团员工（包括集团内部员工、在集团工作的外部员工）在利用人工智能（下称“AI”）技术提升效率和创新的同时，其部署及使用符合法律法规、伦理准则和风险管理要求，防止发生数据泄露、违法及侵权行为或不符合国家监管要求等法律风险，推动AI安全、合规、公平、透明、可问责、绿色使用，根据《网络安全法》、《数据安全法》、《个人信息保护法》、《生成式人工智能服务管理暂行办法》、《互联网信息服务深度合成管理规定》、《人工智能生成合成内容标识办法》、《促进和规范数据跨境流动规定》等相关法律法规，以及《复星集团信息安全管理规定》、《复星集团知识产权管理办法》等内部规定，制定本制度，以此明确集团在AI部署和应用等各个环节的安全管理和使用规范。

第二条 本制度适用于上海复星高科技（集团）有限公司及其直接或间接管理、控制的子、孙公司以及各分公司（以下统称集团）的所有部门、所有员工和能接触到信息资产的第三方人员（包括但不限于外包人员、厂商工程师、咨询公司顾问等）。如果产业公司定义了自己的AI制度，则产业公司的制度要求不得低于集团要求。

第三条 员工应主动了解集团AI使用的相关规定，积极参与集团组织的AI使用培训，提升AI使用的安全意识和技能，并严格遵守和执行集团AI安全管理的要求。

第二章 定义

第四条 集团AI资产是集团信息资产的一部分，指集团在开发、部署和运营人工智能技术过程中所依赖或产生的各类数字化资源，涵盖数据、算法、硬件、软件及服务核心要素。

第五条 集团AI应用指，集团基于第三方或开源模型自行部署，只对集团授权用户开放的人工智能服务。

第六条 公共AI应用指，基于公有云或开放平台提供的第三方人工智能服务，无需自行搭建底层基础设施和模型，即可通过API、网页或应用程序直接调用AI能力。公共AI应用包括其模型须遵循国家安全技术标准与行业监管要求。

第七条 保密信息指，员工或第三方人员在受集团雇用期间或因与集团合作原因知悉或掌握的公司视为保密或予以保密的技术信息、经营信息以及其他任何非公开信息，包括但不限于集团向员工或第三方人员披露的信息，以及集团负有保密义务的任何他人的保密信息。

第八条 个人信息指，以电子或者其他方式记录的与已识别或者可识别的自然人有关的各种信息，不包括匿名化处理后的信息。

第九条 敏感个人信息指，一旦泄露或者非法使用，容易导致自然人的人格尊严受到侵害或者人身、财产安全受到危害的个人信息，包括生物识别、宗教信仰、特定身份、医疗健康、金融账户、行踪轨迹等信息，以及不满十四周岁未成年人的个人信息。

第十条 去标识化指，个人信息经过处理，使其在不借助额外信息的情况下无法识别特定自然人的过程。

第十一条 匿名化指，个人信息经过处理无法识别特定自然人且不能复原的过程。

第十二条 HITL(Human-in-the-Loop)，指在 AI 应用系统运行过程中，在关键节点引入人类介入、审核或反馈机制，以确保决策安全、准确并可迭代优化。

第三章 声明

第十三条 集团在部署和使用 AI 应用系统的过程中，应当严格遵守合法、正当、必要和数据最小化原则，未经员工同意或法律法规明确授权，不主动采集、存储或处理员工生物特征数据、通讯内容、行踪轨迹等个人信息。如因系统运行必需收集相关数据，除法律明确授权事由外，应当征得个人数据主体同意，或者对数据进行匿名化处理。

第十四条 员工不得利用集团 AI 资产处理非工作相关事务。员工利用集团 AI 资产所产生、处理和存储的职务成果，其知识产权依法归集团拥有，未经授权不得对外泄露或擅自使用。

第十五条 集团可在法律法规允许的范围内，基于运营管理、网络安全、保密管理、合规审计及依法配合监管或司法机关调查等需要，对集团 AI 资产进行必要的管理、监测、取证和处置。相关规则应通过本规范、劳动合同、员工手册或告知书等方式事先向员工及第三方人员说明。发生或疑似发生攻击、泄露、侵权或违规使用时，集团可依法

采取拦截、隔离、删除、冻结账号等必要保护措施。违规处理依照第七章及集团相关规定执行。

第四章 组织与职责

第十六条 集团数智化与AI条线信息安全团队

（一）负责制定、解释、修改集团AI应用管理和使用规范，并监督及协调推动规范在集团内部的落地执行；

（二）负责对接成员企业，推动集团AI应用管理和使用规范在成员企业内部落地执行；监督成员企业AI应用安全管理执行度；

（三）负责汇总整理集团和成员企业内部AI应用安全状况、安全事件等相关报告并报送集团管理层；

（四）组织AI应用安全工作检查，分析AI应用总体安全状况，提出分析报告和安全风险的防范对策；

第十七条 集团各业务部门

（一）监督部门员工合法、合规使用AI应用；

（二）确保部门AI应用在设计、开发和部署等各阶段均能满足集团现有的安全标准和规范以及适用的法律要求；

（三）指定部门AI应用的负责人，对AI输出结果、风险决策及合规性承担责任，对AI造成的错误结果或损害按相应应急策略进行调查处理；

（四）制定与更新部门AI应用的安全应急策略及应急预案，并定期组织对安全应急策略和应急预案进行测试和演练；

第十八条 集团员工

（一）合法、合规使用AI应用；

（二）配合完成相关AI应用伦理与安全培训；

（三）主动报告可疑AI应用使用行为以及可能的侵权或泄露线索；

第五章 AI应用部署规范

第十九条 集团AI应用优先考虑私有化部署，以降低数据泄露风险；应用调用非私有化部署大模型，则须经集团数智化与AI条线信息安全团队和集团风控条线法律事务部审核批准，且大模型供应商需提供数据不落地的协议保证，或提供其他符合集团要求的数据处理、留存、训练使用限制及审计条款。大模型供应商不得将集团数据用于优化或训练其通用模型。

第二十条 部署AI模型须遵循国家安全技术标准与行业监管要求，所有第三方或开源模型的引入应通过集团数智化与AI条线信息安全团队的评估，形成技术评估报告。

第二十一条 AI应用部署时，应遵循集团各项信息安全规定：

（一）使用集团数据作为训练数据或部署 AI 应用时，应按《复星集团信息安全管理》实施 AI 应用数据分类分级管控，并应当遵循以下原则：

- 机密级及以上信息：经过审批后仅限使用集团私有化部署 AI 应用处理；
- 秘密级信息：仅限使用集团自建 AI 应用处理；
- 非保密信息：可使用第三方公共 AI 应用处理；

（二）确保训练数据来源合法，禁止使用非法爬取、侵犯知识产权或未经授权的数据作为训练数据；涉及个人信息的训练数据须具备合法处理基础，并在实现处理目的的必要时范围内进行去标识化处理，涉及敏感个人信息须取得个人信息主体的单独同意或具备其他合法处理基础，并施加额外的隐私增强技术等加密、脱敏措施保障数据安全。采集前需对数据来源进行安全评估，确保数据来源具备可追溯性和时效性，应保留开源数据的合法授权记录、自采数据的采集记录、商业数据的交易证明及用户输入信息的明确授权。个人信息主体注销账号后，应在承诺期限或法定期限内删除或匿名化相关个人信息，法律、行政法规规定应当留存或集团依法履行法定义务所必需的除外；删除或匿名化要求及期限应在隐私政策或用户须知中显著告知；

（三）集团所有 AI 应用须部署内容安全监控引擎，阻断包含恶意代码、对抗样本或超出授权范围的数据输入。建立输入数据与业务系统的访问控制矩阵，确保秘密级及以上数据仅授权模块可调用。输入数据设置留存周期，确需长期存储的须经合规审批后实施加密存储；

（四）集团所有 AI 应用须部署内容安全监控引擎，实时过滤涉及违法信息、商业秘密泄露及违背公序良俗的生成内容；输出结果须通过自动化校验，确保符合集团各项规章制度的要求；

（五）实施数据加密存储及传输，防止数据被泄露或篡改；

（六）集团所有 AI 应用接入须通过统一身份认证，采用最小权限原则进行访问控制；

（七）集团部署在境内的 AI 应用向集团境外用户提供服务前，应当评估潜在的数据跨境合规风险。若涉及数据双向跨境流动，需同时满足中国境内和境外相关的数据保护以及 AI 监管的法律法规要求。在完成法律合规性评估并获得相应审批前，禁止上线或向境外提供服务；

（八）集团 AI 应用的研发测试环境与生产环境须物理隔离；

（九）部署实时监控系統，检测模型滥用；

（十）集团 AI 应用须制定数据泄露、模型攻击等应急预案，定期开展灾备演练，应急预案应遵循《复星集团信息安全应急响应管理办法》；

（十一）定期对集团 AI 应用进行漏洞扫描和渗透测试，识别潜在威胁，修复高危漏洞；

（十二）建立 AI 服务访问日志，按期留存记录；

第二十二条 为应对 AI 算力增长带来的能源消耗、水资源使用及碳排放风险，集团及合作方在 AI 模型训练、部署与运营中应遵循以下要求：

（一）优先采用节能硬件（如高效 GPU、液冷系统）、优化算法以减少计算负载和能源浪费；

(二) 第三方 AI 服务供应商在同等条件下优先选择采用可再生能源、具备绿色认证的数据中心。

第二十三条 集团应部署持续监控 AI 模型性能的机制，以检测因数据模式或外部条件变化导致的模型偏移。具体要求：

(一) 建立模型性能基线，定期评估准确率、误报率等关键指标；

(二) 当检测到模型效果显著下降时，应立即触发纠正措施，包括但不限于用新数据重新训练、调整参数或切换模型版本；

(三) 所有模型更新应经过测试与审批流程，并记录变更日志，确保应用系统长期可靠。

第二十四条 集团应用使用外部大模型 API 时，应遵循以下各项规定：

(一) 所有外部 API 调用须通过集团数智化与 AI 条线信息安全团队的评估；

(二) 禁止在涉及国家安全、金融、医疗等敏感领域的场景中调用未通过安全认证的 API；

(三) API 供应商应提供相应的安全认证资质、隐私政策及数据存储说明；

(四) 遵循本规范第二十一条中的各项安全规定；

第二十五条 集团 AI 应用在设计、训练、部署及使用过程中，应主动识别、评估并减轻潜在的算法偏见，确保公平和非歧视性行为。具体要求包括：

(一) 训练数据的采集与构建应确保多样性、代表性，避免因数据不平衡导致对特定人群的歧视或不利影响；

(二) 对 AI 模型定期开展偏见审计与公平性评估，测试模型在不同人口群体中的输出差异，并形成审计报告；

(三) 对存在偏见的模型采取数据重平衡、算法调整或重新训练等纠正措施，防止 AI 应用延续或放大现有社会不平等。

第二十六条 集团对 AI 应用系统的使用、能力、局限性、数据来源及潜在风险应履行充分披露义务，并确保 AI 输出结果或决策过程具备可理解的理由。具体包括：

（一）所有面向用户的 AI 应用，应在交互界面或相关文件中明确 AI 应用“可执行”与“不可执行”的边界，不得超越授权范围介入应由人工独立作出判断的关键事项，并清晰说明“此为 AI 系统生成/辅助决策”，不得冒充人类；

（二）对于被要求解释的输出结果，系统应提供相应的合理解释；

（三）涉及个人信息处理或高风险决策的 AI 应用，需在具备有效控制措施的情况下使用，并应主动披露其数据来源、模型版本、关键参数及风险提示。

第二十七条 所有由 AI 应用生成的内容（包括但不限于文本、图像、音频、视频、决策建议）均应在输出时予以明确标注，提示用户该内容由 AI 生成。具体要求如下：

（一）标注方式应清晰、不易被忽略，例如在文本末尾添加“AI 生成”标识，在图像边缘添加水印，在语音开头播报“以下内容由 AI 合成”；

（二）对辅助决策型 AI 应用，应在决策结果旁注明“本结果为 AI 辅助建议，需人工确认”；

第二十八条 建立面向受 AI 决策影响的个人或实体的申诉流程，确保其有权对 AI 输出结果提出异议。申诉流程应：

（一）易于访问（如线上表单、客服热线），公开透明，且不因用户身份而设置歧视性条件；

（二）公平、响应迅速，明确处理时限，并指定专人复核；

（三）记录申诉内容、处理过程及结果，定期汇总分析，用以改进 AI 系统质量与伦理合规。

（四）集团用户可联系 AI 应用团队：AI@fosun.com；

第二十九条 集团要求 AI 管理系统接受第三方评估与认证，以确保符合负责任人工智能实践的国际标准。具体要求：

（一）优先参考 ISO/IEC 42001（人工智能管理体系）、NIST AI 风险管理框架等行业指引；

（二）对涉及高风险场景或核心业务的 AI 应用，应委托外部独立机构进行合规审计与认证；

（三）认证结果应按法律法规及集团信息披露规定向相关利益相关方公示，增强系统可信度与社会信任。

第三十条 集团 AI 应用下线时，应用负责人应按集团规定对数据进行留存或清除。

第六章 AI应用使用规范

第三十一条 员工应优先使用集团搭建的AI应用，如需使用公共AI应用，则应使用集团推荐的公共AI应用，并遵守本规范的相关数据安全与合规要求。

第三十二条 员工使用公共AI应用时，禁止输入非法爬取、侵犯知识产权或未经授权的数据；禁止擅自输入个人信息、敏感个人信息；仅在集团已具备合法处理基础、系统明确允许且已采取去标识化等保护措施的前提下，方可按授权范围输入；禁止输入集团任何保密信息。

第三十三条 员工使用集团AI应用时，禁止输入非法爬取、侵犯知识产权或未经授权的数据；禁止擅自输入个人信息、敏感个人信息；仅在集团已具备合法处理基础、系统明确允许且已采取去标识化等保护措施的前提下，方可按授权范围输入；禁止输入集团机密及以上保密信息，秘密级信息需经审批授权及脱敏处理后方可输入。

第三十四条 员工在境内使用境外部署的AI应用时，应当满足中国关于数据跨境的相关法律法规的要求，禁止将敏感个人信息输入到境外部署的AI应用中，确有业务必要的非敏感个人信息输入应经过必要的合法性评估，并控制在最小范围内。

第三十五条 员工在境外使用AI应用时，应遵守所在地国家或地区的数据保护及AI监管相关法律法规，并确保不违反中国法律关于数据出境、国家安全与个人信息保护的要求。

第三十六条 员工使用AI应用时，应当遵循相关法律法规的要求，不得实施任何损害集团或他人合法权益的行为，禁止利用AI应用生成、制作或传播任何违法内容，包括但不限于：

- (一) 反对宪法确定的基本原则的；
- (二) 煽动颠覆国家政权、推翻社会主义制度，危害国家统一、主权和领土完整的；
- (三) 泄露国家秘密、危害国家安全或者损害国家荣誉和利益的；
- (四) 煽动民族仇恨、民族歧视，破坏民族团结，或者侵害民族风俗、习惯的；
- (五) 破坏国家宗教政策，宣扬邪教、迷信的；
- (六) 散布谣言，扰乱社会秩序，破坏社会稳定的；
- (七) 传播淫秽、赌博、暴力或者教唆犯罪的；
- (八) 侮辱或者诽谤他人，或侵犯集团或他人名誉权、隐私权、知识产权及其他合法权益的；
- (九) 采用潜意识、故意操纵或欺骗性技术，或利用年龄、残疾或特定社会经济状况等漏洞，实质性扭曲个人行为的；
- (十) 未经合法授权且在公共场所进行实时远程生物识别（RBI）的；
- (十一) 危害社会公德或者民族优秀传统文化的；
- (十二) 其他违反法律法规规定的。

第三十七条 通过AI系统辅助作出的关键决策，必须保持HITL，必须嵌入有效、可记录的人工监督机制，并赋予人工干预、推翻或上报的权限与能力。具体规范如下：

- (一) AI系统不得在无人工监督或未经人工审查机制的情况下，独立作出不可逆或高风险决策；
- (二) 人工干预的全过程应留存书面或电子记录，作为合规审计依据。

第七章 处罚原则

第三十八条 违反以上要求及集团信息安全管理规定的，集团将按照有关规定进行处罚，给集团造成损失的，集团将根据法律规定予以追偿。

第三十九条 对于情节严重，对集团造成重大损失，甚至构成犯罪的，交由司法机关追究其法律责任。

第八章 附则

第四十条 本制度由集团数智化与AI条线负责解释与修订。

第四十一条 本制度由复星国际董事会批准发布。

第四十二条 本制度自发布之日起执行。