

AI 审计工具在国有企业“三重一大”决策监督中的应用逻辑与边界

李体虹

中国铁路武汉局集团有限公司 湖北 武汉 430000

摘要：国有企业“三重一大”（重大决策、重要人事任免、重大项目安排和大额度资金运作）决策制度是国有企业内部监督的核心机制。然而，传统监督模式在信息不对称、滞后性及主观性等方面存在固有局限。人工智能（AI）审计工具凭借其强大的数据处理、模式识别与预测分析能力，为破解上述难题提供了全新可能。本文旨在系统阐述 AI 审计工具在“三重一大”决策监督中的内在应用逻辑，深入剖析其在实践层面的具体应用场景，并审慎界定其应用的技术、法律与伦理边界。研究表明，AI 审计并非旨在取代人类监督，而是通过构建“人机协同”的智能监督范式，赋能国有企业治理现代化，但其有效应用必须建立在清晰的权责划分、健全的数据治理体系与严格的合规框架之上。

关键词：AI 审计；国有企业；“三重一大”；决策监督；治理现代化；人机协同

引言

国有企业是中国特色社会主义市场经济的“顶梁柱”，其规范运行关乎国家经济命脉。为强化对国企权力制约与监督，国务院确立“三重一大”决策制度，要求重大决策等事项须经集体讨论决定，以堵塞权力寻租漏洞，保障国有资产安全。但实践中，“三重一大”决策监督面临挑战：决策链条长、信息分散、业务复杂，形成信息壁垒；传统监督方式滞后，难实现事前预警与事中纠偏；监督人员专业判断易受主观因素影响，面对海量数据力不从心，监督精准度和效率待提升。在此背景下，以大数据等技术为核心的 AI 审计工具应运而生，为监督注入智能化新动能，可挖掘潜在风险信号与违规模式。然而，AI 审计也引发算法黑箱等边界问题。因此，探究 AI 审计工具在“三重一大”决策监督中的应用逻辑与边界，对推动国企治理现代化意义重大。

1 “三重一大”决策监督的困境与 AI 审计的介入契机

1.1 监督困境

理解 AI 审计价值，需先明确“三重一大”决策监督体系的核心痛点。一是信息不对称与碎片化：“三重一大”事项涉及多业务系统，决策信息分散在各部门或系统中。监督部门跨系统、跨部门收集整合数据耗时费力，易遗漏关键信息，产生监督盲区。二是监督滞后性与被动性：传统监督多为事后审计或专项检查，问题被发现时损失或已造成，追溯成本高，难以有效约束

正在进行的决策过程，无法满足风险管理前置化要求。

三是主观性与经验依赖：监督有效性依赖审计人员专业素养、敏感度和经验，不同人员风险判断有差异，结论客观性和一致性难保证，且面对复杂环境和多样舞弊手法，人力难以全面应对。

1.2 介入契机

正是在这样的背景下，AI 审计工具的出现为破解上述困境提供了全新的技术路径。AI 审计的核心优势在于其全量数据处理能力、自动化与持续性，以及强大的模式识别与异常检测功能。它能够接入企业 ERP、OA、财务、HR、合同管理等所有相关系统，对“三重一大”事项相关的全量数据进行实时或准实时采集，从根本上打破信息孤岛。AI 审计程序可设定规则，自动执行数据分析任务，实现 7×24 小时的持续监控，将监督关口从事后移至事前和事中^[1]。更重要的是，通过机器学习算法，AI 能够从历史数据中学习正常决策模式，并精准识别出偏离常规的异常交易、关联方交易、利益冲突等高风险信号，其识别能力远超人类的直觉判断。这些技术特性共同构成了 AI 审计介入“三重一大”监督的历史性契机，为其从“人防”向“技防+人防”融合转型奠定了坚实基础。

2 AI 审计工具在“三重一大”监督中的应用逻辑

AI 审计工具在“三重一大”监督中的应用，并非简单的技术叠加，而是一套基于数据驱动、风险导向和闭环管理的系统性逻辑。这一逻辑贯穿于数据层、模型层

与应用层，形成一个有机整体。

2.1 数据层：构建全域贯通的决策数据湖

应用逻辑的起点是数据。AI 审计的有效性完全取决于输入数据的质量与广度。因此，首要任务是构建一个覆盖“三重一大”全生命周期的决策数据湖。这包括：结构化数据：如财务系统的资金流水、预算数据；HR 系统的人事档案、薪酬记录；项目管理系统的进度、成本数据；采购系统的供应商信息、招投标记录等。非结构化数据：如党委会、董事会、总经理办公会的会议纪要、录音（经授权）；项目可行性研究报告、尽职调查报告；合同、协议的扫描件或 PDF 文本；往来邮件、即时通讯记录（在合规前提下）等。通过 API 接口、ETL 工具等方式，将这些多源异构数据进行清洗、标准化、打标签，并统一汇聚至数据湖。这是 AI 审计得以施展拳脚的基础平台。

2.2 模型层：构建面向“三重一大”场景的风险识别模型

在数据湖的基础上，针对“三重一大”的四大类别，需开发定制化的 AI 分析模型，以实现精准的风险识别。在重大决策监督方面，可利用自然语言处理技术解析会议纪要与决策文件，将其与国家法律法规、行业监管规定及公司章程进行自动比对，标记潜在的合规风险点；同时，通过知识图谱技术构建企业内外部实体的关系网络，分析决策背后的逻辑链条是否自洽，从而识别隐藏的利益输送。在重要人事任免监督方面，AI 可将拟任人选的履历信息与岗位说明书中的硬性条件进行自动匹配，快速筛查资格不符者；更进一步，结合员工花名册、亲属信息登记表及外部工商司法数据，利用图计算技术深度挖掘拟任人员与现有高管之间是否存在未披露的近亲属、校友或前同事等隐性关联，防范裙带关系。对于重大项目安排，AI 可整合立项、预算、进度、成本等多维度数据，建立项目健康度指数，实时监控执行异常；同时分析历史招投标数据，识别围标串标、量身定做招标条件等舞弊模式^[2]。在大额度资金运作监督方面，AI 可利用交易图谱技术对资金流向进行穿透式追踪，识别虚构交易或挪用资金行为；并通过分析交易对手方信息与定价公允性，自动发现未披露或定价不公允的关联交易，有效防范国有资产流失。

2.3 应用层：实现“监测-预警-核查-反馈”的闭环

AI 审计的价值最终体现在一个完整的应用闭环之中。该闭环始于智能监测，即 AI 模型在后台持续运行，对新产生的“三重一大”相关数据进行实时扫描。一旦

发现异常，系统便根据风险事件的严重程度和可信度，自动生成不同等级的风险预警，并精准推送至相应的监督人员。随后进入人机协同核查阶段，监督人员收到预警后，并非盲目采信 AI 结论，而是利用 AI 提供的证据链——如关联图谱、异常交易明细、文本对比结果等——进行深度核查与专业判断。在此环节，AI 扮演的是“超级助手”角色，极大地提升了核查效率和精准度。最后是反馈与模型优化环节，监督人员的核查结论，无论是确认风险还是排除误报，都会被系统记录并用于训练和优化 AI 模型，使其在未来变得更加智能和准确。这一自我进化的正向循环清晰地表明，AI 审计并非要取代人的监督，而是通过技术手段将监督人员从繁重的数据搬运和初步筛查工作中解放出来，使其能将更多精力聚焦于高价值的判断、沟通与决策上，从而构建起高效、精准、前瞻的“人机协同”监督新范式。

3 AI 审计应用的边界探析

尽管 AI 审计前景广阔，但其应用绝非没有边界。若不加以审慎界定和管控，不仅可能引发新的风险，甚至可能侵蚀监督本身的公正性与合法性。其边界主要体现在技术、法律与伦理三个相互交织的维度。

3.1 技术边界：算法的局限性与“黑箱”风险

AI，特别是深度学习模型，常被视为“黑箱”，其决策过程缺乏透明度和可解释性。在“三重一大”这种关乎重大利益和政治责任的领域，一个无法解释的“高风险”预警可能引发不必要的恐慌或错误的问责。监督人员必须清醒认识到，AI 输出的只是概率性的风险线索，而非确定性的结论。模型本身存在误报与漏报的固有缺陷，过度依赖可能导致对无辜者的骚扰，或对真实风险的忽视。更值得警惕的是，模型会继承并放大训练数据中的历史偏见。如果过往的决策本身就存在系统性歧视，AI 模型可能会将其固化甚至加剧。因此，技术边界要求我们必须坚持“人在环路”原则，所有 AI 生成的重大风险预警，都必须经过人类监督者的复核与最终裁决^[3]。同时，应优先采用可解释性 AI 技术，在可能的情况下，让模型能够说明其做出判断的依据，从而增强监督结论的可信度与可接受性。

3.2 法律与合规边界：数据隐私与授权

AI 审计的基石是数据，而数据的获取与使用必须严格遵守《个人信息保护法》《数据安全法》等法律法规。任何逾越法律红线的数据使用行为，都将使整个 AI 审计的合法性荡然无存。企业在实施 AI 审计时，必须恪守数据最小化与目的限定原则，仅收集与“三重一

大”监督直接相关的必要数据，不得过度采集员工或关联方的个人隐私信息。对于涉及个人敏感信息的采集和分析，必须获得明确的、单独的授权，企业应在内部规章制度中清晰告知员工 AI 监督的存在、范围及目的。此外，必须建立严格的数据访问控制、加密传输与存储机制，防止数据泄露、滥用或被恶意攻击。AI 系统的开发者和运维者也需承担相应的保密责任，确保整个数据处理链条的安全合规。

3.3 伦理与权责边界：监督的正当性与责任归属

AI 审计的引入，深刻改变了监督的权力结构和责任链条，必须警惕其可能带来的伦理挑战。首要的是防止“算法暴政”，即不能因为 AI 的“客观”外表，就赋予其过大的决策权重，以至于压制了合理的异议和民主讨论。监督的最终目的是促进科学决策，而非制造寒蝉效应。其次，必须明确责任主体。当 AI 审计未能识别出重大风险，或因错误预警导致损失时，责任应由谁承担？是 AI 开发者、系统运维者，还是采纳 AI 建议的监督人员？这一问题必须在制度层面予以厘清，避免出现“无人负责”的灰色地带。通常，人类监督者应对其最终决策负责，而技术提供方则对其产品的缺陷负责。最后，必须维护程序正义。被 AI 预警指向的个人或部门，应享有知情权、申辩权和申诉渠道。AI 监督的过程和结果应当具备一定的透明度，以保障被监督者的合法权益，确保技术应用不偏离公平正义的轨道。

4 构建负责任的 AI 审计应用框架

为了在充分发挥 AI 审计效能的同时，有效管控其边界风险，国有企业应着力构建一个集制度、技术、人才于一体的负责任应用框架。顶层设计与制度先行：将 AI 审计纳入企业整体数字化转型和合规管理体系，制定专门的《AI 审计应用管理办法》，明确规定其应用范围、数据使用规范、操作流程、权责划分及伦理准则^[4]。夯实数据治理基础：建立高质量、标准化的企业数据治理

体系，确保数据的真实性、完整性、一致性和安全性。这是 AI 审计成功的先决条件。打造复合型人才队伍：培养既懂审计、法律、企业管理，又了解 AI 技术原理的复合型人才。他们是连接技术与业务、确保 AI 应用不偏离轨道的关键桥梁。坚持人机协同，强化复核机制：始终将 AI 定位为辅助工具，建立强制性的 AI 预警人工复核机制。监督人员的专业判断永远是最终防线。建立动态评估与退出机制：定期对 AI 审计模型的有效性、公平性和安全性进行第三方评估。对于表现不佳或存在重大伦理风险的模型，应有果断的调整或退出机制。

5 结语

AI 审计工具为国企“三重一大”决策监督提供了强大技术支撑，其以数据驱动风险识别，构建智能闭环，有望推动监督模式向主动、实时、全覆盖转变。但技术有双面性，算法黑箱、数据隐私、权责模糊及伦理失范等，构成其应用边界。未来，国企不能在“用”与“不用”间抉择，而要在“如何负责地用”上着力。需秉持审慎创新态度，在顶层设计、制度建设等多方面协同，构建“人机协同”“合规可控”“促进善治”的 AI 审计应用生态。如此，才能释放 AI 赋能潜力，使其成为护航国企高质量发展、筑牢决策防火墙的可靠伙伴，而非“潘多拉魔盒”。技术是桨，清晰边界与坚定价值观才是指引航向的舵。

参考文献

- [1] 张玲玲,何金妹,赵丽萍.AI 对审计工作的影响研究[J]. 现代商贸工业,2025,(23):157-159.
- [2] 吴勇,郭秋梦,丁晓月,等.AI 技术的审计应用:现状、挑战与应对[J]. 中国注册会计师,2025,(11):95-102.
- [3] 罗冰.AI 审计驱动下的审计职业大变革[J]. 商业文化,2025,(20):90-92.
- [4] 赵永红.AI 时代智能审计的运行机理、治理效能与发展要求[J]. 财务管理研究,2025,(04):1.