

团 体 标 准

T/AMAC 0004—2026

基金经营机构大模型技术应用规范

Specification for the application of large-scale model technology
in fund management institutions

2026-04-03 发布

2026-04-03 实施

中国证券投资基金业协会 发布

目 次

前言 II

引言 III

1 范围 1

2 规范性引用文件 1

3 术语和定义 1

4 缩略语 2

5 总则 2

6 参考框架 2

7 基础设施 3

8 数据管理 4

9 模型服务 5

10 应用技术 8

11 安全管理 10

12 场景应用 13

附录 A （资料性） 大模型技术应用风险分析及应对措施建议 16

附录 B （资料性） 基金业务领域大模型技术应用案例 18

参考文献 24

前 言

本文件按照 GB/T 1.1—2020《标准化工作导则 第1部分：标准化文件的结构和起草规则》的规定起草。

请注意本文件的某些内容可能涉及专利。本文件的发布机构不承担识别这些专利的责任。

本文件由中国证券投资基金业协会提出。

本文件由中国证券投资基金业协会归口。

本文件起草单位：中国证券投资基金业协会、易方达基金管理有限公司、中国中金财富证券有限公司、工银瑞信基金有限公司、华夏基金管理有限公司、九坤投资（北京）有限公司、阿里云计算有限公司、北京智谱华章科技有限公司、华为技术有限公司、中国信息通信研究院。

本文件主要起草人员：陈丽园、刘硕凌、程宁、戴路、李正非、王昱森、谢橦、谢碧松、周琳、崔仲奇、杨帆、李胜浩、刘丽丽、岳彬、杨思成、戴春博、张晓智、张巧玲、薛大山、徐旭、秦思思、丁伯轩、梅亚雷、侯潇为。

引 言

2023年中央金融工作会议首次提出金融强国目标，将数字金融作为金融五篇文章之一，强调金融机构需加快数字化转型，以提高金融服务竞争力。中国证券监督管理委员会发布的《证券期货业科技发展“十四五”规划》指出“推进科技赋能与金融科技创新，大力提升行业数字化应用水平”。近年来，以大模型技术为代表的生成式人工智能技术快速发展，已成为新一轮科技革命和产业变革的重要驱动力量。2024年“人工智能+”首次写入政府工作报告，强调开展“人工智能+”行动，培育新质生产力。

作为金融市场的重要参与者，基金经营机构正积极探索并应用大模型技术。然而，大模型技术在资产管理领域的应用仍面临一系列挑战和困难，其中包括金融知识理解不足、模型幻觉等技术挑战，以及在实际应用中需要解决的安全性、合规性等问题。

为促进资产管理行业数字化转型和创新发展，引导基金经营机构规范、合理运用大模型技术提升服务水平，有效保护个人金融信息安全及投资者权益，特编制本团体标准。

基金经营机构大模型技术应用规范

1 范围

本文件给出了大模型技术在资产管理业务中的参考框架，规定了大模型技术的基础设施、数据管理、模型服务、应用技术、安全管理和场景应用等方面的要求。

本文件适用于基金经营机构使用大模型技术进行系统平台建设及应用服务。

2 规范性引用文件

下列文件中的内容通过文中的规范性引用而构成本文件必不可少的条款。其中，注明日期的引用文件，仅该日期对应的版本适用于本文件；不注日期的引用文件，其最新版本（包括所有的修改单）适用于本文件。

GB/T 31168 信息安全技术 云计算服务安全能力要求
GB/T 35273 信息安全技术 个人信息安全规范
GB/T 41867 信息技术 人工智能 术语
GB/T 42755 人工智能 面向机器学习的数据标注规程
GB 45438 网络安全技术 人工智能生成合成内容标识方法
JR/T 0060 证券期货业网络安全等级保护基本要求
JR/T 0171 个人金融信息保护技术规范
JR/T 0258 金融领域科技伦理指引

3 术语和定义

GB/T 41867 界定的以及下列术语和定义适用于本文件。

3.1

大模型 large-scale model

基于大量数据训练得到，具有复杂计算架构，能处理复杂任务，且具备一定泛化性的深度学习模型。

注：大模型的参数量由其功能和模态决定，一般不低于1亿。大模型训练使用的数据总量受参数量的影响，达到收敛的大模型的参数量的对数与其训练数据总量的对数成正比。

[来源：GB/T 45288.1—2025，3.1]

3.2

推理 inference

从给定的前提进行论证并得出结论。

注1：在人工智能领域中，一个前提是一个事实、一个规则、一个模型、一个特征或原始数据。

注2：术语“推理”既指过程也指结果。

[来源：GB/T 41867—2022，3.2.30]

3.3

微调 fine-tuning

为提升人工智能模型的预测精确度，一种先以大型广泛领域数据集训练，再以专门领域数据集继续训练的附加训练技术。

注1：专门领域数据一般指下游任务数据。

注2：常用的微调方法包括全参微调、高效参数微调等。

[来源：GB/T 41867—2022，3.2.31，有修改]

3.4

提示词 prompt

提示语

使用大模型进行微调或下游任务处理时，插入到输入样本中的指令或信息对象。

[来源：GB/T 45288.1—2025，3.5]

3.5

检索增强生成 retrieval-augmented generation; RAG

一种结合信息检索和语言生成模型，从外部知识库中检索相关信息并增强生成过程的技术。

注：检索增强生成的关键特性包括检索能力、增强提示生成、以及结合检索信息和原始输入进行的语言生成能力。

3.6

智能体 agent

为实现特定目标或任务，一种能够感知环境并自主做出决策和行动的实体。

注：智能体的关键特性包括自主性、感知能力、决策能力和行动能力。

4 缩略语

下列缩略语适用于本文件。

AI：人工智能（Artificial Intelligence）

API：应用编程接口（Application Program Interface）

ONNX：开放神经网络交换格式（Open Neural Network Exchange）

MCP：模型上下文协议（Model Context Protocol）

5 总则

基于大模型技术开展的应用服务的基本约束见《生成式人工智能服务管理暂行办法》和《证券期货业网络和信息安全管理办法》。

基于大模型技术开展的应用服务的价值理念和投资者保护等方面的行为规范应符合 JR/T 0258 的规定。

注：附录A给出了大模型技术应用主要风险类别、其对应的具体表现和应用措施建议，供行业参考。

6 参考框架

大模型技术在资产管理业务中的应用参考框架由六个核心部分组成：基础设施层、数据管理层、模

型服务层、应用技术层、安全管理层以及场景应用层。详见图 1。



图1 参考框架

该参考框架中：

- a) 基础设施层。是整个框架的基石，它提供数据、计算、存储和网络支持，确保基础大模型、服务框架和应用场景的运行环境，包括高效的多机多卡计算资源调度、多源数据的存储解决方案以及网络的互连互通能力；
- b) 数据管理层。负责数据的采集、处理和知识库构建，为大模型的训练和应用提供高质量的数据支持；
- c) 模型服务层。专注于大模型的选型、部署、微调和管理，确保模型的高性能和稳定性；
- d) 应用技术层。提供了一系列技术工具和方法，如提示词工程、检索增强生成、智能体和组件库，以支持大模型在具体应用中的实现。这些工具和方法有助于开发人员高效地构建和优化大模型应用，同时保障开发过程的规范性和一致性；
- e) 安全管理层。着重于保障大模型建设和应用过程的安全性，包括基础设施安全、数据安全、模型安全和业务安全；
- f) 场景应用层。将大模型技术应用于具体的业务场景，如投资研究、合规风控、市场营销、客户服务、运营管理、效率办公和研发编程等，以提升业务效率和效果。

7 基础设施

7.1 计算模块

计算模块是负责执行模型计算任务的核心组件，包括但不限于以下要求：

- a) 应支持多机多卡并行计算；
- b) 宜挂载安全可靠的 AI 算力芯片；
- c) 应实现同一服务器的卡间、跨服务器间的高速数据通讯能力，并进行横向和纵向扩展；
- d) 应有专用的卡间互联高速接口，支持训练和推理过程中卡间大量数据交换传输；
- e) 应具备常见的分布式集合通讯原语实现，支持主流分布式框架；

- f) 应支持集群通过高速通信协议进行横向和纵向扩展；
- g) 网卡配置应满足计算模块的需求，如带宽大小、TCP 网络需求。

7.2 网络模块

网络模块是负责数据传输、通信和网络连接的关键组件，包括但不限于以下要求：

- a) 应为整个 AI 硬件基础设施的子系统间提供互联互通能力，组网平面应分为参数面、业务面和存储面，分别用于实现多机分布式训练时的参数交换、智能算力系统业务调度与管理及智能算力访问存储区的网络互联；
- b) 宜支持流量控制，如通过有效的技术手段防止系统死锁，更细颗粒度的流量控制技术；
- c) 宜支持基于 AI 流量感知的负载均衡策略，支持不同 AI 算力模型下数据流并发的均匀哈希模式，减少因负载分担不均造成的网络拥塞；
- d) 宜包含管控系统，支持网络故障快速定界、故障隔离和溯源。

7.3 存储模块

存储模块是负责数据持久化、管理和访问的关键组件，包括但不限于以下要求：

- a) 部署存储节点应提供文件存储、对象存储等存储服务，应为人工智能训练平台提供高吞吐、大带宽的样本原始数据；
- b) 宜具备多协议互通能力，一份数据无需转换，可以通过多种协议访问，支持文件协议、对象协议、大数据协议互通；
- c) 宜支持存储空间横向扩展能力，以应对 AI 数据持续增长；
- d) 宜支持敏感数据加密存储功能；
- e) 宜能使用或提供向量数据库；
- f) 宜支持大规模检查点文件（checkpoint）的快速读取和写入能力；
- g) 宜支持数据冗余备份，降低不可恢复性损害的概率。

8 数据管理

8.1 数据采集

数据采集是构建大模型微调、RAG 等能力的重要基础，应符合数据在应用、保密和回溯方面的要求。包括但不限于以下要求：

- a) 应明确数据来源的所有权和使用权限，遵守相关的数据隐私法规、版权法等法律法规并获得必要的授权使用授权；
- b) 应确定数据采集的需求（如舆情、研报、即时通信消息流、路演、电子邮件、传真等）、数量和特点（如时效性、准确性、保密性和可靠性）；
- c) 应支持不同模态的数据采集，包括但不限于文本、表格、音频、视频、图片、遥感影像等；
- d) 应对采集数据进行版本管理措施，确保历史数据可溯源，以便回溯数据质量以及管控数据安全；
- e) 应在采集阶段对公共数据、私密数据进行合理管理或分流，以符合信息安全规范，如即时通信消息、电子邮件等；
- f) 宜支持对数据质量、重复度、可信度的初步筛选和清洗能力，如数据去重、噪声剔除、规范化等；
- g) 宜支持对原始数据实现初步标注，对后续的加工、处理或即时消费提供标签支持。

8.2 数据处理

数据处理可分为数据清洗、数据标注、数据增强和质量评估等处理环节，可根据数据处理要求选择。包括但不限于以下要求：

- a) 数据清洗宜制定标准化流程，从编码规范、格式规范、异常处理、规则校验等维度规范数据清洗流程，宜支持自定义清洗算子和数据清洗规则；
- b) 数据标注宜制定标注规范，注重标注质量和流程管理，数据标注规程应符合 GB/T 42755 的要求。可采用标注工具并建设标注知识库。应严格管理标注人员资格，注意标注安全与隐私；
- c) 对需要数据增强场景的，可利用数据变换、合成等手段，增加数据集的多样性和丰富性；宜关注数据分布一致性，避免分布偏移导致性能下降；
- d) 宜进行数据质量评估，建立质量评估框架及评估体系，可制定统一的质量评估方法规范及数据质量等级和分级标准；可支持数据集的自动和人工质量检查，确保数据质量。

8.3 知识库构建

8.3.1 知识库的目标和范围

知识库应覆盖大模型应用所涉及的主题领域，包含足够的数据量，确保能够提供相关且准确的信息；应定期更新，以保持信息的时效性和准确性。

8.3.2 知识采集与整理

知识库的知识应来源于可靠且权威的渠道，确保信息的准确性和可信度。知识库内容的管理应基于机构内部对数据治理的要求，对不同来源的知识进行隔离并根据权限制定管理方案。

8.3.3 知识存储和检索

知识库应采用合理的存储结构，以便于知识的快速检索和访问；建立高效的索引机制，以提高检索速度和准确性。为知识库构建索引时，使用的向量化方法应与大模型应用中访问知识库的方法一致，以保证索引的可用性。

8.3.4 知识库与大模型应用的交互

知识库应提供一系列读取和存储的接口，以便于数据的传输和交换；需满足大模型应用在性能方面的要求，包括不限于响应速度、处理能力等。

9 模型服务

9.1 模型选型

9.1.1 基本要求

大模型选型应根据实际使用需求为导向的原则，通过大模型评测的方式选用适合的大模型服务。进行大模型选择时：

- a) 在选择开源技术时，应重视开源社区的活跃度和支持组织情况，优先考虑在公开评测集中排名领先的模型；
- b) 应兼顾考虑模型效果与资源消耗，通常模型的参数规模与效果成正相关，与资源消耗也成正相关；
- c) 应考虑模型的可扩展性和兼容性，确保模型能在多样化的操作系统和硬件上稳定运行。大模型的系统兼容性包括但不限于以下要求：

大模型应能够在不同操作系统和硬件平台上运行，并兼容安全可靠的操作系统，如 EulerOS、UOS、AnolisOS 等；

大模型应能够与主流编程语言和开发工具集成，如 Java、Python、C++等；

大模型应能够处理不同规模和类型的数据集，包括结构化和非结构化数据。

- d) 在选择合作厂商时，必须确保其符合国家法律法规，并且所提供的模型已经完成备案。

9.1.2 模型评测

进行大模型评测时：

- a) 应基于公司内部积累的场景样例，构建专有的大模型能力评测集，用于指导模型的选型和评测；
- b) 对于金融领域应用的模型，评测应至少涵盖通用能力和金融能力两个方面，其中通用能力包括生成能力和理解能力；
- c) 应建立内部评测平台，形成大模型的迭代更新机制，确保系统在外部技术升级时能持续稳定运行；
- d) 内部评测平台宜支持 API 接口调用和本地模型推理两种评测方式，并保证两种方式的评测结果具有可比性，以适应不同来源的大模型；
- e) 宜参考国内外知名大模型榜单的评价数据和方法，结合实际业务需求，设计和实施评测体系。

9.2 模型部署

9.2.1 部署通则

模型部署形态应基于硬件配置、安全等级、服务能力和成本预算等维度进行综合评估，可选择本地化部署、外部算力托管部署（外部算力托管部署是指租用第三方算力、由基金经营机构自部署和管控模型的模式）或云服务调用（云服务调用是指通过 API 调用第三方公有云或专属云模型服务，模型与数据不在金融机构本地的模式）三种模式：

- a) 符合下列条件之一的，宜选择本地化部署：
 - 涉及高敏感数据或严苛的隐私合规要求；
 - 存在特殊硬件资源需求，包括但不限于专用加速卡、密码设备等；
 - 需满足毫秒级实时推理响应要求；
 - 其他经评估需实施本地化部署的特殊场景。
- b) 符合下列条件之一的，宜选择外部算力托管部署：
 - 内部算力短期难以满足业务需求，但需保持对模型、数据及运维的自主可控；
 - 需利用行业统一算力平台或专属云资源，且符合监管关于数据驻留、隔离与审计的要求；
 - 需要弹性扩容容业务，第三方可提供合规的托管环境；
 - 不适合或不计划使用外部提供的通用云模型服务。
- c) 不属于上述本地化部署或外部算力托管部署情形，且符合下列条件之一的，宜选择云服务调用：
 - 业务访问具有间歇性或突发性特征；
 - 对系统快速部署、自动化运维有明确要求；
 - 短期内难以满足业务所需的硬件资源配置及资金投入。确认外部云服务的安全能力、合规资质与服务水平符合应用需求。

9.2.2 本地化部署要求

模型本地化部署宜支持国内外主流异构硬件资源的统一管理，按业务领域进行资源隔离、动态流量分配、权限控制及全生命周期管理。包括但不限于以下要求：

- a) 资源隔离与安全：应基于业务领域实现资源的硬隔离或逻辑隔离，确保各领域计算、存储、网络资源的独立性与安全性，并支持为每个领域设置独立的网络安全策略与资源配额；
- b) 流量调度与治理：宜支持按业务领域进行动态流量分配与治理，具备基于细粒度策略的灰度发布、蓝绿部署及故障熔断能力，并能够快速隔离和处置流量异常，保障业务安全稳定；
- c) 权限控制与审计：应实现基于业务领域的细粒度权限控制体系，明确划分各角色在模型开发、部署、运维及下线全流程中的操作权限，并记录完整操作日志，确保所有行为可审计、可追溯。

9.2.3 外部算力托管部署要求

在租用第三方算力、由基金经营机构自部署和管控模型的场景下，包括但不限于以下要求：

- a) 资源隔离：应确保所租用的计算、存储和网络资源在物理或虚拟层面实现严格隔离，建立安全的私有网络（VPC）；
- b) 数据安全：应依据业务数据分类分级规范，对投资研究、资产配置、交易执行及客户信息等敏感数据实施严格管控。基金经营机构应主导数据在传输、存储及模型运算过程中的加密策略，采用独立于外部算力提供方的密钥管理机制，确保对核心业务数据及模型资产拥有完全的、排他性的控制权，防止敏感信息在外部算力环境中泄露；
- c) 运维管理：应明确基金经营机构与服务商的运维职责边界，并建立有效的监控和应急响应机制。

9.2.4 云服务调用要求

为保护数据安全，云服务调用应遵守 GB/T 31168 的要求，建立有效的安全防护机制，包括但不限于以下要求：

- a) 数据隔离：宜采用多租户技术，实现用户数据的逻辑隔离。若涉及数据跨境传输，必须根据《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》、《促进和规范数据跨境流动规定》等法律法规进行严格评估；
- b) 身份认证：宜支持多种身份认证方式，如密钥认证、数字证书、二次认证等；
- c) 访问审计：应完整记录所有 API 访问日志，支持审计和追溯，并提供查询接口或报表。

9.3 模型微调

9.3.1 全量参数微调

为适应资产管理领域的任务需求，全量参数微调应涉及领域数据适配、微调策略、评估策略等方面。包括但不限于以下要求：

- a) 领域数据适配：资产管理领域数据负责为大模型提供必要的领域数据，用于全量微调的数据应与资产管理任务适配，确保全量参数微调效果；
- b) 微调策略：用于大模型微调应使用合适的全量参数微调技术，确保适应资产管理任务的微调性能；
- c) 评估策略：全量微调过程中，应选择合理的性能评估指标，确保大模型在领域数据上的性能。

9.3.2 参数高效微调

大模型需确保领域数据有限情况下的微调性能，参数高效微调可涉及数据增强策略、参数高效微调技术、评估策略等方面。包括但不限于以下要求：

- a) 数据增强策略：数据增强策略应扩展有限标注数据的多样性，确保大模型的泛化性；
- b) 参数高效微调技术：应选择合适的参数高效微调技术，确保大模型知识不被破坏，同时保证大模型在下游任务上的性能；

- c) 评估策略：参数高效微调过程中，应选择合适的性能评估指标，确保大模型能力的准确衡量。

9.4 模型管理

大模型的开发和调优需要大量的迭代和调试，数据集、训练代码或参数的变化可能会影响模型的质量，模型管理可导入所有训练版本生成的模型，可对所有迭代和调试的模型进行统一管理。包括但不限于以下要求：

- a) 应支持多种方式的模型导入功能，包括从训练作业导入、从模板导入、从容器镜像导入等；
- b) 应支持模型管理功能，包括创建模型新版本、删除模型、模型查询等；
- c) 应支持业界主流模型格式，如 ONNX、TensorFlow、PyTorch 格式。

10 应用技术

10.1 提示词工程

10.1.1 提示词指导原则

提示词指导原则包括：

- a) 指令要清晰具体：提示词应明确，确保需求准确传达，减少模型误解，提升输出相关性和质量。具体描述帮助模型定位任务需求，避免不相关或不符合预期的结果；
- b) 上下文完整全面：提示词设计需考虑任务执行的完整上下文，包括角色身份、背景信息及约束条件。明确角色身份有助于模型生成符合预期的语气、风格和内容。包含背景信息确保模型理解前因后果。考虑限制条件或规则，如格式要求、时间限制或目标受众，确保内容契合场景需求。对于复杂任务，完整性尤为关键，显著提升输出相关性和精确度；
- c) 模型计算推理的充分性：确保模型充分计算，提高结果准确性和稳定性。通过提示词设计，引导模型按完整逻辑链条计算和输出，充分分析上下文和潜在可能性，深入推理，减少偏差和误差。此方法优化输出质量，提高模型在复杂任务中的可靠性和一致性；
- d) 持续优化机制：建立闭环的持续优化机制，确保持续提升提示词效能。定期评估提示词在实际应用中的效果（如输出准确性、相关性、用户满意度等），基于评估结果和应用反馈，迭代优化提示词内容，形成“设计-应用-评估-优化”的闭环流程，不断提高结果的有效性。

10.1.2 提示词方法

提示词构建流程和技巧主要要求如下：

- a) 提示词构建流程。应通过系统化流程确保提示词设计高效、精准，提高模型对任务需求的理解和响应能力，具体包括定义、参考、框架、调试、持续优化等步骤：
 - 1) 定义：确定解决问题，评估提示有效性；
 - 2) 参考：确定场景涉及的测试数据、参考文本和方法技巧；
 - 3) 框架：构建初步框架，避免复杂化，专注核心要素；
 - 4) 调试：反复测试调整，提升生成效果，满足预期目标；
 - 5) 持续优化：定期评估和分析应用效果，输出提示词优化方案，迭代优化提示词，并形成知识积累。
- b) 提示词技巧。可采用以下方式构建高质量的提示词：
 - 借助框架快速搭建：使用通用提示词框架高效构建基础提示词；
 - 优化描述和指代的清晰度：清晰提示词引导模型生成预期结果，减少不相关内容和错误响应；

调整结构前后顺序：提示词结构顺序影响输出质量，指令和重要要求放在前后两端；

输入输出格式：使用分隔符表示输入不同部分，结构化输出便于后续处理（如JSON、HTML格式）。

边界条件处理：宜考虑潜在边界条件及处理逻辑，提供替代方案；

少样本：提供少量示例，利用模型上下文学习能力，提升复杂任务表现；

思维链：指示模型按步骤思考，输出符合要求的内容，适用于复杂场景；

知识生成：先让模型生成背景知识或上下文，帮助输出更准确、全面的结果；

自我一致性：针对问题提供多种推理路径，选择最一致答案作为最终答案；

闭环迭代优化：持续开展提示词优化工作，将优化经验沉淀到提示词库或知识库中，促进提示词工程整体水平的提升。

注：附录B中的B.1给出了根据上述提示词工程原则和方法的基金业务应用相关案例。

10.2 检索增强生成

10.2.1 概述

检索增强生成可包括离线处理和在线处理两个阶段。

10.2.2 离线处理

离线处理的模块应包括：

- 知识上传：如对知识信息的解析服务、数据增强服务、摘要服务、切片服务、关键词服务等；
- 知识管理内容：如知识目标管理、知识标签管理、文件目录管理、文件标签管理、知识树管理等；
- 存储管理：如向量数据库管理、图数据库管理、传统关系型数据库管理等；
- 索引管理：索引生成，增强索引，正倒排索引，向量索引等。

10.2.3 在线处理

在线处理的模块应包括：

- 问题的前处理：如问题改写、分流决策、基于知识树的意图导航、敏感内容检测，线索发现等；
- 知识检索：包括召回服务、精排服务等；
- 回答生成：包括生成阶段的提示词工程、知识溯源、基于思维链的回答等。

注：附录B中的B.2给出了检索增强生成在基金研报知识库问答场景应用案例。

10.3 智能体

10.3.1 核心能力

在资产管理业务中应用大模型智能体（Agent），应具备规划（Planning）、记忆（Memory）、工具（Tools）、行动（Action）等基本能力：

- 规划：指智能体理解任务、构思解决方案，在必要时拆解为子任务，评估工具，反思调整的能力；
- 记忆：指智能体存储信息，并可进行检索召回的能力；
- 工具：指智能体使用工具的能力；
- 行动：指智能体根据规划与记忆，执行具体行动的能力，包括与外部交互或工具调用等。

在智能体与其他系统交互的场景中，可利用模型通信协议提升智能体使用工具的能力。宜支持的模型通信协议包括但不限于：

- a) MCP；
- b) Agent to Agent (A2A) 通信协议。

注：附录B中的B.3给出了智能体技术在基金公司运营办公场景应用案例。

10.3.2 应用范式

智能体应用在具体资产管理业务中时，可采用单智能体、多智能体、智能体-人交互等三种范式：

- a) 单智能体范式：单智能体是指独立在环境中运作的智能体，通过感知、学习和行动来适应环境并实现目标。单智能体的优势在于系统相对简单，易于实现和部署，但在处理复杂任务时可能存在局限性；
- b) 多智能体范式：多智能体系统由多个智能体组成，智能体之间需要相互协作或竞争以实现共同或各自的目标。多智能体系统通过智能体之间的通信、协调和合作来解决复杂问题。其应用包括分布式控制、多角色博弈场景等。多智能体的优势在于能够处理更复杂的任务，但需要解决智能体之间的交互和协调问题；
- c) 智能体-人交互范式：通过智能体-人交互的设计范式，可利用智能体的计算能力和人类的认知能力，实现在更广泛应用场景更高效地解决问题。

10.4 组件库

组件库为大模型技术在资产管理行业落地的技术载体，根据当前的技术特点和演变规律，可按具体服务类型规划相关组件，包括但不限于以下要求：

- a) 知识型任务：大模型输出自身掌握的通用知识或外部知识，宜具备检索器、排序器、生成器等组件。其中：
 - 检索器负责从外部知识来源检索相关信息，如搜索引擎、向量数据库等；
 - 排序器负责对检索结果进行评估，并排列优先级；
 - 生成器负责利用检索和排序结果，结合用户的输入，生成最终答案或内容。
- b) 决策型任务：大模型根据用户意图执行决策型任务，宜具备感知、规划、执行等相关组件。其中：
 - 感知组件宜包含视觉模型感知、监控报警感知、主动触发感知、数据管道推送等相关能力；
 - 规划组件宜包含目标识别、任务拆解、自动化决策、持续迭代、工具选择、评估检查等能力；
 - 执行组件宜包含 RPA 能力、API 接口、可执行文件或代码片段等形式的集成能力。
- c) 分析型任务：大模型完成数据分析、情报分析等相关工作，宜具备知识模型、自然语言理解、集成知识代理和插件、数据处理、内容生成等相关组件；
- d) 研发型任务：大模型以插件等形式支持主流开发工具辅助完成研发型任务，覆盖需求、设计、编码、测试、发布、运维等生命周期，宜具备代码生成、代码补写或优化、测试生成、规范检查、技术问答、代码解释等相关组件；
- e) 多模态任务：大模型支持涉及语音、视频、图像、文本、代码、表格等不同模态数据类型的任务，宜包含模态识别、模态转译、内容理解、模态生成等相关组件；
- f) 通用组件：宜具备提示词管理、部署框架、多源异构数据管道等通用组件。

11 安全管理

11.1 基础设施安全

大模型运行在物理基础设施之上，应对底层硬件、操作系统、网络环境等采取全面的安全防护措施。包括但不限于以下要求：

- a) 系统安全：应对操作系统进行安全性加固，如删除非必要组件、修补漏洞、设置最小权限等；宜制定安全配置基线，规范主机、虚拟化、容器等组件的安全配置；应建立漏洞扫描、分析、修复的完整流程，并制定应急响应计划；
- b) 隔离机制：宜实施网络隔离策略，将大模型平台划分为不同的安全域，限制不同安全域之间的网络通信；
- c) 网络安全：应配置防火墙和入侵检测系统，监控网络流量和攻击行为，及时发现并阻止恶意流量；应使用安全的传输协议加密数据传输，防止数据在传输过程中被窃取或篡改。应结合场景应用对大模型的基础网络安全进行定级，各等级基本要求并符合 JR/T 0060 要求。

11.2 数据安全

数据是大模型的核心资产，包括但不限于训练数据、验证数据、测试数据、提示词、微调数据以及与大模型交互的各类业务和用户数据等。应从全生命周期的角度，采取全方位的保护措施，确保数据安全。包括但不限于以下要求：

- a) 数据采集和存储安全：应建立数据采集审计机制，对敏感数据实施脱敏或加密处理，并限制对存储设备的物理访问；宜加密存储和传输训练数据、验证数据、测试数据、提示词及微调数据等，防止数据泄露；应定期备份数据，建立完善的数据恢复机制，防止数据丢失并确保灾及时恢复。敏感数据的加密应实现具备事前处理、事中控制、事后审计的能力。包括但不限于以下要求：

宜使用虚拟数据代替大模型训练、微调、强化学习中所需的真实数据，可以通过人工标注或机器学习等技术手段模拟真实数据的场景和结构，解决数据泄露和数据污染带来的问题；

应避免大模型对敏感数据的记忆或检索，确保大模型在预训练阶段不对敏感数据形成特定记忆；避免在检索增强生成阶段获取外挂知识库的敏感数据，造成聊天记录、个人信息、交易数据等真实信息出现在大模型的输出内容中；

在使用大模型技术分析或处理敏感数据时，应确保敏感数据不会存在被中间环节获取的条件，诸如网络明文发送、日志内容输出等；

用户使用大模型时的行为数据及聊天数据，应确保不在用户未充分允许的情况下进行归档，保证用户的隐私权和数据权；

应建立必要的数据审查机制和风险应对机制，在关键位置设置信息审查点，一旦发现未加密或虚化的敏感信息存在泄露风险，应当立即参照风险手册采取中断或拒绝相关服务等措施；

完善审计跟踪能力，记录所有对敏感数据的访问、存储、修改、删除等操作行为以辅助监测数据泄露风险，建立完善的事后追踪和审计机制，防止数据滥用。

- b) 数据使用和处理安全：

应建立访问控制机制，限制对数据访问权限，确保仅授权人员的访问和处理数据，可参考 GB/T 42775 对数据进行分类分级来进行数据的授权；宜启用访问审计功能；

应建立严格的“入模数据”筛选机制。不应将未脱敏的客户个人信息、核心交易指令、投研未公开信息直接用于大模型的训练或微调，防止模型权重产生“记忆”从而导致无法撤

回的泄露。对于高敏感等级的投研策略和风控规则数据，应仅允许在私有化部署、物理隔离或逻辑隔离的专用模型环境中进行训练或推理；

大模型处理数据应继承已有信息隔离制度。在模型推理及检索增强生成环节，应实施“用户-数据-模型”三维权限管控。确保不同业务条线人员在使用同一大模型服务时，模型无法跨越业务边界输出对方的未公开信息；

- c) 数据销毁和回收安全：应制定数据生命周期管理制度及回收销毁标准，采用可靠的数据彻底销毁方式，建立销毁审计和责任追溯机制。尤其当含有敏感内部数据的模型退役或需删除特定数据时，应采取模型回滚、机器遗忘或销毁模型实例的方式，确保数据无法被恢复或提取。

涉及个人信息时，个人信息处理的安全还宜考虑 GB/T 35273 的要求。在个人金融信息方面，还应遵守 JR/T 0171 的要求，设计并实施覆盖个人金融信息全生命周期的安全保护策略。

11.3 模型安全

为保障模型在整个生命周期内的安全性、可靠性和合规性，应从模型内部治理、外部供应链和持续运营三个维度采取综合防护措施。包括但不限于以下要求：

- a) 输出内容治理：应制定模型输出内容审查规则，建立违规内容检测机制与过滤方法，并在应用层对不合理输出进行拦截、降级或人工复核；
- b) 模型鲁棒性防护：宜制定对抗性样本测试标准，衡量模型鲁棒性，规范数据增强、对抗训练等防护措施的使用；
- c) 模型权限和访问控制：宜基于最小权限原则限制模型访问权限，制定模型使用审计要求，确保调用留痕，并符合下列要求：

控制策略：根据平台的形态、架构和用户角色不同，平台应在权限控制策略中选择一个或多个执行，包括基于角色的访问控制、基于属性的访问控制或强制访问控制等；

模型训练权限：大模型相关的应用方需规定有且仅有特定角色允许访问模型训练相关的功能，如模型工程师或系统管理员。模型训练相关的功能应包括启动训练任务，监控训练进度，调整模型参数或改变模型结构等；

模型访问权限：大模型相关的应用方应存在专用岗位以访问和编辑训练模型的数据集，且该岗位原则上应与拥有模型训练权限的岗位相互独立。该岗位访问的数据应视敏感程度而定进行进一步的权限控制或隔离；

模型调用权限：大模型相关的应用方应定义平台用户或模型应用开发人员等角色以赋予模型调用权限。该权限仅适用于调用平台能力以获得结果，但不能访问模型训练的过程或内部算法。

- d) 模型供应链安全：为防范引入外部模型带来的病毒、后门与许可合规风险，应建立供应链安全控制机制，包括但不限于：

来源与完整性：应从官方或可信渠道获取模型，校验数字签名或哈希，确认发布者与版本；
安全检查：应对模型包或镜像开展恶意代码与高风险漏洞扫描；发现可疑行为应下架并复核；

记录与追溯：宜建立模型清单，记录来源、版本、依赖与许可信息，支持追溯与审计；

合规与发布管理：应核对开源/商用许可条款，实施签名发布与版本管控，支持回滚或撤销；

传输与存储：应对模型在传输与存储环节进行加密；关键密钥由基金经营机构自行持有或控制授权。

- e) 模型可解释性：宜通过输出模型中间推理过程，提高模型的可解释性；宜通过展示模型的参考溯源，便于用户识别模型幻觉；

- f) 监控与审计：宜对模型输入输出、性能与风险事件建立持续监控与日志审计，支持追溯与责任界定；
- g) 算法信息披露：提供金融产品和服务时宜根据 JR/T 0287 的指导规范地开展算法信息披露活动。

11.4 业务安全

通过大模型生成合成内容标识和传播活动必须遵守《人工智能生成合成内容标识办法》的规定。人工智能生成合成内容标识方法应符合 GB 45438 的要求。

为确保基于大模型的资产管理业务连续性和可靠性，从身份认证、接口安全、审计防护、应急响应等多个维度采取安全保障措施。包括但不限于以下要求：

- a) 身份认证和访问控制：宜通过多重身份认证、动态口令认证标准等手段，进行细粒度的基于角色和策略的访问控制；
- b) API 安全防护：能通过 API 入口的防火墙、防火墙配置及 API 漏洞扫描、渗透测试等手段对 API 接口进行防护；
- c) 合规审计和日志防护：宜建立健全公司级生成式人工智能使用规范，提出日志记录完整性要求和保护措施以及使用规范。宜具备日志记录功能，用于记录系统的操作行为，为审计提供依据。宜具备审计跟踪功能，用于追踪系统的操作流程，以便发现和解决潜在的安全问题；宜具备合规检查功能，用于检查系统的配置和操作是否符合相关法律法规和行业标准要求；
- d) 应急响应与恢复能力：能对安全事件进行分级及通报，制定应急预案、演练及执行流程。

12 场景应用

12.1 投资研究

大模型技术可应用于市场研究、行业分析、公司分析、投资决策等投资研究类场景，包括但不限于以下应用：

- a) 信息提取：宜使用大模型技术从多源异构的金融资料中提取关键信息，包括投资观点、市场趋势等，实现非结构化数据结构化转化，为投研知识管理和决策提供数据基础；
- b) 报表分析：宜使用大模型技术对数据报表中的关键指标进行识别和提取，辅助投研人员进行多维对比及趋势分析；
- c) 舆情监控：宜使用大模型技术检测和识别舆情信息的异常模式，辅助投研人员快速定位问题根源，提供决策支持；
- d) 因子挖掘：宜使用大模型技术对多源数据进行理解和抽象，挖掘较为客观、科学的特征，并自动构建智能模型进行模拟跟踪，以辅助投研人员制定更有效的投资策略。

12.2 合规风控

大模型技术可应用于合规管理、风险管理、内审稽核等合规风控类场景，包括但不限于以下应用：

- a) 信息审查：宜使用大模型技术对合同、公文、披露材料、宣传材料等进行常规性审查，分析文本内容的事实准确性、内容连贯性、上下文适应性、因果关系合理性，可与外部数据源或专家知识进行关联验证，以确保输出的逻辑性和准确性；
- b) 安全评估：宜使用大模型技术对客服问答、邮件外发等场景进行内容审计、个人信息识别、图片多模态分析、代码检测、深度报告识别等审查，确保敏感数据不泄露；
- c) 风险监测：宜使用大模型技术对征信、诉讼、公告等文本数据进行关联分析和解读，结构化内容呈现，提炼风险点，强化风险预警能力，降低风控成本。

12.3 市场营销

大模型技术可应用于市场分析、个性化营销、媒体舆情感知等场景，包括但不限于以下应用：

- a) 个性化推荐：宜使用大模型技术分析投资者的行为和偏好，从而提供个性化的产品推荐和投资建议，增强客户粘性；
- b) 市场趋势分析：宜使用大模型技术处理和分析大量市场数据，预测市场趋势，为营销策略提供数据支持；
- c) 广告优化：宜使用大模型技术分析广告投放效果，帮助基金经营机构优化广告内容和投放策略，提高广告的转化率；
- d) 社交媒体数据采集：宜使用大模型技术采集社交媒体上的用户反馈，帮助机构及时调整营销服务策略；
- e) 品牌管理：宜使用大模型技术分析公众对品牌的看法和情绪评价，帮助基金经营机构维护和提升品牌形象；
- f) 市场调研：宜使用大模型技术分析消费者数据，进行市场调研，帮助机构更好地理解目标市场和消费者需求。

12.4 客户服务

大模型技术可应用于客户投顾、市场营销、客服问答等客户服务类场景，包括但不限于以下应用：

- a) 智能问答：宜使用大模型技术提升客户服务效能，通过理解客户意图、高度拟人化、客户情绪分析和实时内容推荐，实现高质量的客户交互，提高客户服务满意度；
- b) 投顾助手：宜使用大模型技术支持投资顾问场景，根据客户交互、特征、情绪等实时推荐信息，快速梳理行情及金融资讯，实现体系化输出，为客户提供更优的投资建议；
- c) 材料制作：宜使用大模型技术生成文案、标语、海报等多模态内容，提升客户服务精准度和营销工作效率；
- d) 客户分析：宜使用大模型技术对客户反馈或拜访记录进行数据分析，理解客户情感、提取高价值内容等，构建客户风险探测能力。

12.5 运营管理

大模型技术可应用于知识库问答、产品参数提取、运营信披报告审核等运营管理类场景，包括但不限于以下应用：

- a) 产品参数提取：宜使用大模型技术对管理人提供的产品合同及其他文档进行解析和提取，并定位和获取关键产品要素参数，辅助运营人员在系统中进行结构化参数录入；
- b) 投监规则解析：宜使用大模型技术对事中划款、事后投资等投资规则信息进行解析，并进行规则语义理解，辅助投资监督管理岗进行规则解读和系统配置；
- c) 信披报告审核：宜使用大模型技术对公募产品信息披露工作所需的多源数据进行整合解析，并进行复杂校验比对，辅助运营人员进行信披报告制作和审核工作。

12.6 效率办公

大模型技术可应用于日常效率办公场景中，包括但不限于以下应用：

- a) 信息处理：宜使用大模型技术赋能会议纪要、转译、翻译、纠错等办公场景，实现意图感知、信息贯通、规划执行等功能，提升信息处理效率；
- b) 知识管理：宜使用大模型技术对募集说明书、资管合同、投资者监督表、法律法规进行解析，支持关键要素抽取、智能对话问答、案例分析等，在预处理环节降低员工的工作成本；

- c) 数字员工：宜使用大模型技术支持多元应用、统一入口、主动感知、精准调度的智能办公能力，实现数字员工辅助完成请假申请、日程安排、会议室预订等相关需求。

12.7 研发编程

大模型技术可应用于研发、测试、项目管理等各研发编程类场景，包括但不限于以下应用：

- a) 代码生成：宜使用大模型技术直接将自然语言描述的需求转化为程序代码，降低编程门槛，缩短软件开发周期。同时，通过代码注释生成和文档自动编写，提高代码的可读性和可维护性；
- b) 程序优化：宜使用大模型技术优化代码编写流程，通过智能代码补全、代码审查、错误检测和自动修复等功能，提高研发人员的编程效率和代码质量。同时，利用人工智能技术进行代码性能优化，提升程序运行效率；
- c) 测试自动化：宜使用大模型技术进行软件测试和质量保证，通过自动化测试脚本生成、测试用例设计和缺陷检测，提高软件测试的覆盖率和准确性，降低软件缺陷率；
- d) 项目管理：宜使用大模型技术支持项目管理和协作，通过智能任务分配、进度跟踪、资源优化配置等功能，提高研发团队的协作效率和项目管理水平。

附 录 A

(资料性)

大模型技术应用风险分析及应对措施建议

表A.1给出了大模型技术应用主要风险类别、其对应的具体表现和应用措施建议。

表 A.1 大模型技术应用风险分析及应对措施建议

风险类别	具体表现	应对措施建议
知识产权归属不明与侵权风险	a) 如AI生成内容与他人已受知识产权保护的成果构成实质性相似,则存在被认定为侵权的风险; b) 使用AI工具生成图像、文案等内容时,可能存在工具使用人的智力投入不被认可,进而导致知识产权归属不明的风险	前端输入控制: 明示禁止输入受版权、商标等保护的内容作为提示语,并明示禁止要求AI工具仿照受版权、商标等保护的内容进行生成,防止生成内容构成侵权
		后端输出审核: 建立图文生成检测等机制
		使用环节管控: 明确生成内容的归属政策及使用规则; 对外使用前需经过人工修改、审核; 必要的情况下,可要求在对外发布前保留提示语及生成记录等留痕证据
人格权益与信息权益侵权风险	使用AI工具生成内容时,若涉及特定自然人的肖像、声音、姓名等特征,或虚构、贬损其形象,该内容在使用过程中可能构成对人格权、名誉权或个人信息权益的侵害,因而具有人格权益与信息权益侵权风险	前端输入控制: 明示禁止或限制输入自然人的个人信息
		后端输出审核: 通过关键词筛查和工具辅助,检查生成内容中是否提到了真实人物的特征,判断是否影射或误用实际存在的人物
		使用环节管控: 生成结果中含具体人物特征时,在对外发布前须重点审核
商业秘密侵权与内部权限管控风险	若将包含商业秘密的内容(如模型设计、财务指标、客户数据等)输入至AI工具,则该工具在后续响应其他用户请求时可能调用相关历史信息并在生成内容中予以体现,导致商业秘密泄露。此外,相关生成内容还可能被未经授权的部门误用或对外发布,进而可能引发内部权限管控失控	模型部署控制: 在部署模型时,可以设置禁止在后台收集与留存提示词,禁止使用用户的提示词进一步优化或者改进模型
		前端输入控制: 对输入内容进行分级脱敏,明示禁止输入含有商业秘密或公司内部进行权限管控的内容
		后端输出审核: 构建生成结果敏感词扫描与机密信息残留检查等机制
		使用环节管控: 加强对使用生成内容的审批管理
内容安全合规风险	现有的AI工具普遍存在“AI幻觉”,即可能生成错误、虚假、误导性信息等,在生成内容被用于对外发布场景时,可能带来内容安全方面的风控问题	前端输入控制: 规范提示语编写,避免引导AI生成虚假、敏感、误导等违法违规内容
		后端输出审核: 部署敏感词过滤、核查等工具
		使用环节管控: 生成内容须审核后方可使用
AIGC标识合规风险	使用AI工具提供互联网信息服务时,未依法进行标识,引发的合规风险	按照GB 45438的方法,对生成内容进行标识。

风险类别	具体表现	应对措施建议
伦理道德风险	无论是训练人工智能的语料数据，还是模型本身的算法都可能带有导向性，例如可能包含性别、年龄、职业等方面的偏见，从而导致生成内容呈现出不当归类、刻板印象或倾向性判断等，进而引发伦理风险	前端输入控制：建立提示词设计指导，避免隐含歧视导向（如性别、职业、年龄）
		后端输出审核：构建偏见识别模型与敏感变量检测等机制
		使用环节管控：明确生成内容为“辅助”，需经过人工复核确认，不得直接用于内外部决策
安全漏洞与供应链风险	a）使用AI研发编程工具生成代码时，可能含注入漏洞等安全问题，或者引入了与基金经营现有合规要求不兼容的开源许可证（如 GPL/AGPL 等强“传染性”协议），则可能造成内部业务系统产生安全、侵权等问题 b）部署第三方大模型时可能存在供应链污染问题，比如模型权重文件被植入持久化后门，在模型推理时窃取敏感的金融领域提示词或输入数据	集成静态代码分析工具，建立规范的漏洞扫描机制（例如扫描OWASP Top10漏洞）
		引入人在回路复核机制，对于涉及交易执行、清算结算等核心业务的代码，严禁全自动生成并上线，必须经过具有金融业务背景的资深人员进行代码审查和人工确认
		在引入新模型时进行模型完整性验证：部署前计算权重文件的SHA-256哈希与官方基准比对。同时使用模型安全扫描工具进行检测

附 录 B
(资料性)
基金业务领域大模型技术应用案例

B.1 案例一：提示词工程

B.1.1 基金客服参考回复

可通过预设模型的角色形式，优化描述和指代的清晰度，让模型更好地理解需求。例如，需要大模型扮演一个基金公司客服的角色来做出回应（见表B.1）。

表B.1 基于大模型实现基金客服场景下的参考回复

提示词	假设你是基金公司的在线客服人员,接到客户投诉,请做出回应:我购买的基金亏损了,怎么办
大模型回复	尊敬的客户,您好! 感谢您联系我们并反馈您的情况。我们非常理解您对基金亏损的担忧,投资过程中的市场波动确实可能带来焦虑,请您放心,我们会全力协助您分析问题并提供建议。 针对您的情况,我们建议您: 1.了解当前持仓情况 您可以登录我们的APP或官网查看基金的最新净值和持仓详情,确认亏损幅度及持有时间。短期市场波动是正常现象,部分优质基金可能需要更长时间体现价值。 2.评估自身风险承受能力 请您回顾一下购买基金时的风险测评结果。如果当前亏损超出您的承受范围,我们可以一起探讨是否调整投资策略(如转为更稳健的产品)。 3.长期持有或定投平滑成本 若基金基本面未发生重大变化,长期持有或通过定投分摊成本,可能有助于降低短期波动的影响。我们的投研团队会持续跟踪市场,及时提供建议。 4.预约专业顾问沟通 如需进一步分析,我们可以为您安排专属理财顾问,根据您的具体持仓、投资目标等提供个性化方案。您可随时通过电话XXXX-XXXXXX或在线客服预约。 我们承诺: 严格监控基金运作,确保合规透明; 定期提供市场解读与投资建议; 持续优化服务,陪伴您理性应对市场变化。 再次感谢您的信任。投资有风险,但我们会竭尽全力与您共同面对。如有其他疑问,请随时联系我们! 祝您投资顺利! [XX基金客服团队] 注:此为模板回复,实际沟通中可根据客户具体持仓、亏损原因(如市场普跌/个别基金问题)等补充针对性建议。

在合规层面，建议关注以下三点：

- a) 销售适当性与误导销售风险：需重点检查，生成的回复严禁包含“保本保收益”、“预期高收益”等承诺性表述，需符合相关法律法规中关于宣传推介材料的要求；
- b) 投资建议界限：模型输出严格限制在“投资者教育”和“事实查询”范畴，不得提供具体的“投资建议”或“择时建议”，避免触犯相关业务资质的限制；
- c) 隐私保护：在处理客户咨询时，模型不得在未脱敏的情况下记录或透传客户的身份信息、账户金额等敏感数据。

在应用效果层面，衡量应用效果时建议关注以下三个核心指标：

- a) 话术采纳率：客服人员直接使用或经人工调整后使用模型建议回复的比例；
- b) 合规通过率：模型生成的回复通过合规质检要求的比例；
- c) 平均处理时长降幅：应用大模型后，单笔客诉或咨询的处理时间减少百分比。

B.1.2 结构化投研数据

可通过少样本学习方法提升大模型表现。即通过向大模型提供少量示例，使其能够更精准地完成指定任务。以投研文本处理为例，可要求大模型提取文本中的数据指标，并按照既定格式返回用户（如表B.2）。

表B.2 基于大模型实现投研文本的数据结构化

提示词	帮我把这段话中的指标和数据给提取出来,用"指标-数据"的形式呈现。示例:"输入:<本周A股市场方面,上证指数收报3231.41,上涨0.04%>;输出:上证指数收盘价3231.41\n上证指数涨跌幅0.04%\"。输入:<本周A股市场方面,上证指数收报3231.41,上涨0.04%;深证成指收报10793.93,下跌1.86%;创业板指收报2143.01,下跌4.04%;万得全A指数收报4905.07,下跌0.89%;富时中国A50指数收报12600.24,下跌0.01%。
大模型回复	上证指数收盘价3231.41 上证指数涨跌幅0.04% 深证成指收盘价10793.93 深证成指涨跌幅-1.86% 创业板指收盘价2143.01 创业板指涨跌幅-4.04% 万得全A指数收盘价4905.07 万得全A指数涨跌幅-0.89% 富时中国A50指数收盘价12600.24 富时中国A50指数涨跌幅-0.01%

在合规层面，建议关注以下三点：

- a) 数据准确性与受托责任：提取的数据若直接用于基金估值或投资决策，需承担数据准确性的最终责任。建议建立人工复核机制，防止模型幻觉导致的数据错误引发估值差错；
- b) 数据来源合法性：确保输入的文本（如第三方研报、付费资讯）拥有合法授权，避免因大模型处理未授权数据而引发知识产权纠纷；
- c) 非公开信息管理：若处理内容涉及未公开披露的敏感信息，需确保数据处理环境符合物理隔离或逻辑隔离要求，防止内幕信息泄露。

在应用效果层面，衡量应用效果时建议关注以下三个核心指标：

- a) 字段提取准确率：提取的关键数值、实体名称与标准答案的一致性；
- b) 结构化吞吐量：单位时间内处理的文档数量或页面数；

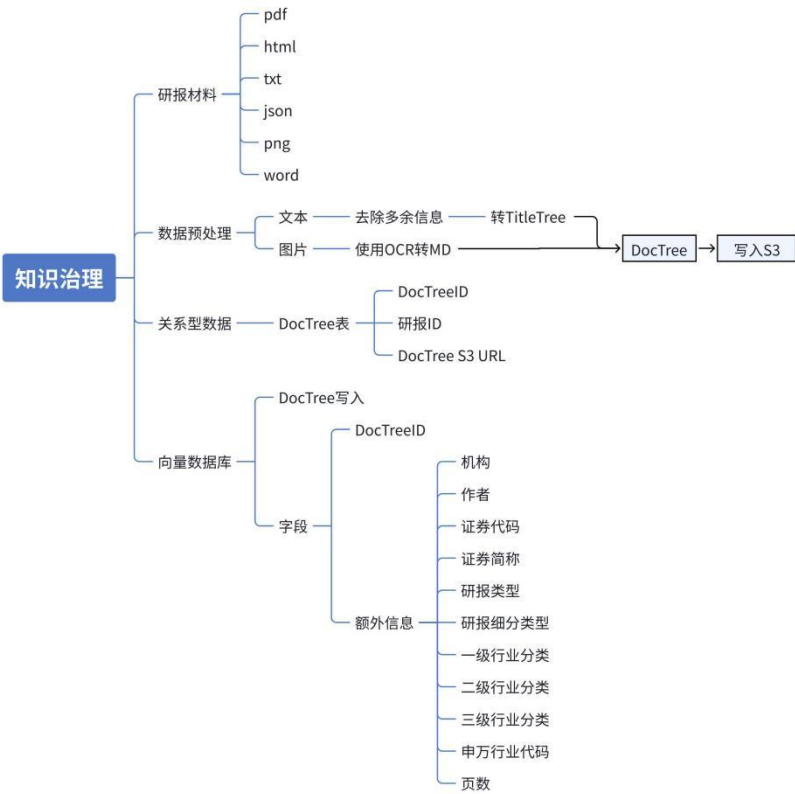
c) 人工复核成本降低率：相比纯人工录入，所需投入的人力工时减少比例。

B.2 案例二：检索增强生成

在基金公司日常运营中，投研人员需处理海量研报信息，以支持基金经理制定策略和风险控制。可通过构建基于大模型检索增强的投研知识问答系统，为基金研究人员提供实时、准确、有依据的智能问答能力，从而提高投研效率。

在数据准备阶段，投研知识问答系统涵盖的数据范围包括内外部宏观研究、行业研究、策略分析、个股研究、投资策略文档、会议纪要、备忘录等多类资料。

在知识治理阶段，作为检索增强生成的关键环节，知识治理的质量会直接影响问答实际效果。知识治理的目标是通过合理的方式，尽可能保留数据原始信息量，以更好地适应问答阶段不同粒度（如关键词、句子、段落、文章等）的查询需求。图B.1给出了知识治理阶段的示意图。



图B.1 知识治理阶段示意图

在在线处理阶段，检索引擎的检索策略包括：向量检索（基于语义相似度）、关键词检索（基于Elasticsearch）、融合排序（加权平均等）。同时，还应设置相应的过滤机制，例如时间窗口筛选（如：近三个月）、来源权重设定（如：策略研究优先）等，以提升检索结果的精准度和实用性。

在提示词构建与大模型回答生成阶段，下方示例给出了一个典型的提示词模板。

示例：

Plain Text

你是一个基金研究分析助手，基于以下资料回答问题，并给出清晰引用。

问题：{query}

资料：

{retrieved_context}

回答：

为了增强输出结果可追溯性与可解释性，可使用回答增强和引用追踪的方式。回答增强方法包括：回答中自动嵌入引用段落的编号（如 [1]）、提供“点击查看原文段”功能等（见图B.2）。引用追踪方法包括：每个向量块在入库时记录来源文档及对应页码和时间戳、回答时回传引用元数据等。

生产端：PMI 回升显示阶段性修复。2025 年 5 月制造业 PMI 为 49.5%，环比上升 0.5 个百分点，虽仍低于荣枯线，但新订单指数小幅回升 0.6 个百分点。企业采购量指数环比改善，但需求端价格指数继续收缩，PPI 同比降幅或扩大。高技术制造业 PMI 连续 4 个月处于扩张区间，装备制造业行业回升至扩张区间，显示新动能支撑力较强。传统行业中的纺织、化工、黑色金属冶炼等行业 PMI 低于临界点，产需释放不足，部分企业因淡季和成本压力减产。高频数据显示，高炉开工率、电炉开工率延续上升趋势，但建材需求转弱，水泥磨机运转率处于低位，与建筑业新订单 PMI 低速呼应。整体来说，政策加码和“抢出口”延续支撑生产端温和复苏。

消费：总体呈现温和复苏，政策与假期效应共振。2025 年“五一”假期成为消费复苏的核心节点。全国国内出游 3.14 亿人次，同比增长 6.4%；国内游客总花费 1802.69 亿元，同比增长 8.0%。服务消费表现强劲，旅游、餐饮、交通出行销售额分别占服务行业前三，其中旅游销售额同比增长 10.53%。餐饮市场同步增长，全国重点餐饮企业销售额同比增长 8.7%，县域市场贡献率首次突破 35%。以旧换新政策持续发力，汽车和家电以旧换新补贴仍然成为拉动社零消费的主要力量。

5月份消费端的表现如何？服务消费和以旧换新政策是否有效推动了汽车、家电销售的增长？
2025-06-06 11:01:05

5月份消费端总体呈现出温和复苏的态势，这得益于政策与假期效应的共振。具体而言，2025年“五一”假期成为了消费复苏的关键节点，全国出游人数达到3.14亿人次，同比增长6.4%，国内游客总花费1802.69亿元，同比增长8.0%[3a]。服务消费表现强劲，旅游、餐饮、交通出行销售额分别占据服务行业前三，其中旅游销售额同比增长10.53%。餐饮市场也同步增长，全国重点餐饮企业销售额同比增长8.7%，县域市场贡献率首次突破35%[3a]。

以旧换新政策持续发力，汽车和家电以旧换新补贴仍然成为拉动社会零售消费的主要力量[3a]。这表明，在消费端，政策与假期效应确实有效推动了汽车、家电销售的增长。

然而，社零增速可能小幅回落，反映出内需修复的渐进性，CPI同比或微涨[3b]。这说明尽管有政策推动和假期效应，消费端的复苏仍需时间，且存在一定的内需修复限制。

3

图B.2 研报问答场景下的回答增强示例

在合规层面，建议关注以下两点：

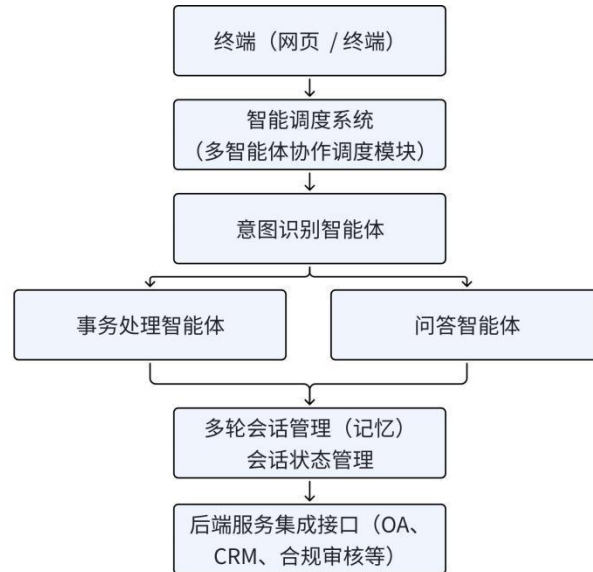
- 引用溯源与版权合规：生成的回答必须明确标注引用来源（具体到某篇研报的页码），既是为了验证真实性，也是为了符合引用规范；
- 信息隔离：系统需具备严格的权限控制。例如，公募基金经理不能检索到仅限专户部门查看的研究成果，研究员不能检索到无权限的拟交易清单等。

在应用效果层面，衡量应用效果时建议关注以下三个核心指标：

- 回答忠实度：生成的回答是否完全基于检索到的文档，未引入模型自主生成的虚构信息；
- 引用准确率：文中标注的引用链接是否能正确跳转到支撑该观点的原文段落；
- 检索召回率：能否从海量文档中精准找到所有相关性高的文档片段。

B.3 案例三：智能体

将智能体技术与基金公司运营办公场景相结合，可为公司内部员工提供 7×24 智能响应服务，覆盖产品培训、会议管理、文件报批等各项业务流程，有效提升基金公司运营办公的效率。图B.3给出了运营办公场景下智能体的设计架构图。



图B.3 运营办公场景智能体架构设计图

意图识别智能体可快速识别用户意图类别，如“查询公司公告”、“预定会议室”、“总结OA审批摘要”等，为不同意图分配后续处理智能体。

事务处理智能体可处理流程型事务，如“出差申请流程智能提请（见图B.4）”等。其应具备调用后端接口的能力，例如OA、会议管理系统等。

问答智能体可回答与公司制度、基金产品、培训材料等相关的知识型问题。其应通过向量数据库存储标准问答库、基金条款、培训文档等。

考虑到基金业务的复杂性，各类智能体均应具备长短期记忆和多轮对话的能力，以提供更加智能和个性化的体验。

帮我提一个下周一到周五的出差流程，目的地是北京，理由是参加学术会议

请您核对出差信息

请确认出差信息

开始时间：	2025年06月09日 >
结束时间：	2025年06月13日 >
出差地点：	北京市 × >
出差说明：	参加学术会议

确定

图B.4 出差申请流程智能提请

在合规层面，建议关注以下两点：

- a) 权限最小化与授权：智能体在执行跨系统操作（如发邮件、查数据）时，必须严格继承当前操作员工的权限，禁止赋予智能体“超级管理员”权限，防止越权操作带来的运营风险；
- b) 操作留痕与可追溯：智能体的所有自动化操作必须有完整的日志记录，确保在出现操作失误时可追溯、可回滚。

在应用效果层面，衡量应用效果时建议关注以下三个核心指标：

- a) 任务规划成功率：智能体能否正确拆解复杂任务（如“写周报”拆解为“查数据”、“分析数据”、“写文案”）并执行完毕；
- b) API 调用准确率：智能体调用内部工具（如数据接口、邮件系统）的参数正确率；
- c) 人效提升倍数：完成同一复杂任务，AI 辅助相比纯人工的效率提升倍数。

参 考 文 献

- [1] GB/T 42775 证券期货业数据安全风险防控 数据分类分级指引
 - [2] GB/T 45288.1—2025 人工智能 大模型 第1部分：通用要求
 - [3] JR/T 0287 人工智能算法金融应用信息披露指南
 - [4] TC260—003 生成式人工智能服务安全基本要求
 - [5] ISO/IEC TS 30105-9:2023 Information technology—IT Enabled Services — Business Process Outsourcing(ITES-BPO)lifecycle processes—Part 9:Guidelines on extending process capability assessment for digital transformation
 - [6] 《证券期货业网络和信息安全管理办法》.中国证券监督管理委员会.2023-01-17
 - [7] 《生成式人工智能服务管理暂行办法》（国家互联网信息办公室 中华人民共和国国家发展和改革委员会 中华人民共和国教育部 中华人民共和国科学技术部 中华人民共和国工业和信息化部 中华人民共和国公安部国家广播电视总局令第15号发布）.2023-07-10
 - [8] 《中华人民共和国个人信息保护法》.全国人民代表大会常务委员会.2021-08-20
 - [9] 《人工智能生成合成内容标识办法》.国家互联网信息办公室、工业和信息化部、公安部、国家广播电视总局.2025-03-07
 - [10] 《互联网信息服务深度合成管理规定》.国家互联网信息办公室、工业和信息化部、公安部.2022-11-25
-