

AI 赋能资产配置（十九）

机构 AI+投资的实战创新之路

核心观点

核心结论：①**信息基础重塑：** LLM 正将海量非结构化文本转化为可量化的 Alpha 因子，根本上拓展了传统投研的信息边界。②**技术路径已验证：** 从 LLM 的信号提取、DRL 的动态决策到 GNN 的风险建模，AI 赋能资产配置的全链条技术栈已具备现实基础。③**未来迈向智能演进：** AI 正从辅助工具转向决策中枢，推动资产配置从静态优化迈向动态智能演进，重塑买方的投研与执行逻辑。

AI 技术正从信息基础、决策机制到系统架构三个层面，深度重构资产配置的理论与实践。大语言模型（LLMs）通过深度理解财报、政策等非结构化文本，拓展了传统依赖结构化数据的信息边界；深度强化学习（DRL）则推动决策框架从静态优化转向动态自适应；而图神经网络（GNNs）通过揭示金融网络中的风险传导路径，深化了对系统性风险的认知。

落地应用不依赖单一模型性能，而依赖模块化协作机制。贝莱德 AlphaAgents 的实践揭示了 AI 投研系统的核心形态：通过模型分工，LLM 负责认知与推理（如多智能体辩论），外部 API 与 RAG 提供实时信息支撑，数值优化器完成最终的资产配权计算。这种架构不仅有效缓解了 LLM 的“幻觉”问题，提升了决策稳健性与可解释性，更形成了从信号生成到组合执行的、可复制的技术栈，为构建真正实用的投资 Agent 奠定了坚实基础。

头部机构的竞争已升维至“AI 原生”战略，其核心是构建专有、可信且能驾驭复杂系统的 AI 核心技术栈。摩根大通的案例表明，其战略远非简单应用外部模型，而是围绕“可信 AI 与基础模型”、“模拟与自动化决策”、“物理与另类数据”三大支柱，进行全链条的专有技术布局。其通过将合规性转化为信任护城河、将市场模拟能力转化为战略风洞、将另类数据转化为信息优势，建立起对手难以短期逾越的复合壁垒。

对国内资管机构而言，破局之道在于战略重构与组织变革，走差异化、聚焦式的技术落地路径。国内资管机构应进行顶层设计并寻求差异化破局，而非盲目跟随。关键在于构建务实高效的“人机协同”体系。面对技术与资源差距，国内机构不宜盲目追求“大而全”。技术落地上，优先利用 LLM 挖掘 A 股市场独特的政策与文本 Alpha，并构建以“人类专家为核心、AI 为智能副手”的协同流程；在组织与文化上，必须打破部门壁垒，锻造融通投资与科技的复合型团队，并将风险管控内嵌于 AI 治理全周期，从而将 AI 从概念转化为坚实的核心竞争力。

风险提示：警惕大语言模型内在缺陷与“幻觉”风险；警惕 AI 模型的“黑箱”特性，AI 投资存在潜在的合规与法律风险；文中所提及公司仅做客观事件汇总，不涉及主观投资建议。

策略研究·策略专题

证券分析师：王开

021-60933132

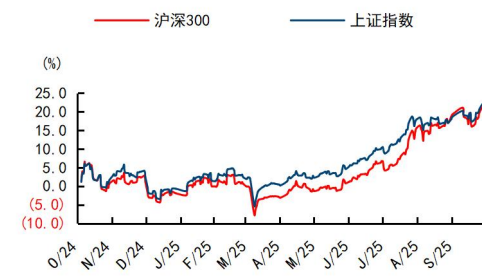
wangkai8@guosen.com.cn

S0980521030001

基础数据

中小板/月涨跌幅(%)	8253.14/0.92
创业板/月涨跌幅(%)	3229.58/2.48
AH 股价差指数	119.43
A 股总/流通市值(万亿元)	96.97/89.08

市场走势



资料来源：Wind、国信证券经济研究所整理

相关研究报告

《估值周观察(10月第3期) - 全球普涨，A股成长反弹》——2025-10-26
《策略观点-积跬步，行稳致远》——2025-10-21
《估值周观察(10月第2期)-价值抗跌，成长承压》——2025-10-20
《ESG热点周聚焦(10月第2期) - 工信部启动2025年度绿色工厂推荐工作》——2025-10-12
《估值周观察(10月第1期) - 市场高低切，周期品领涨》——2025-10-12

内容目录

一、范式重塑：AI 如何变革传统配置模型	4
1.1 资产配置中的大语言模型（LLMs）	4
（1）从文本到交易信号：金融情绪分析的演进机制	4
（2）应用场景与实践价值	4
（3）发展瓶颈与核心挑战	4
1.2 资产配置中的深度强化学习（DRL）	5
（1）理论演进：从静态优化到自适应智能体	5
（2）核心架构与技术实现	5
（3）发展瓶颈与核心挑战	6
1.3 资产配置中的图神经网络（GNNs）	6
（1）将金融系统概念化为网络	6
（2）图神经网络实践：风险传播建模	6
（3）对宏观审慎政策和投资组合对冲的启示	7
二、机构实践：头部资管公司 AI 赋能资产配置	7
2.1 贝莱德（BlackRock）：AlphaAgents——基于多智能体系统的股票组合构建实践	7
（1）战略意图：超越“数据处理”，迈向“决策智慧”	8
（2）核心机制剖析：从协作到辩论的双层决策流程	8
（3）实证结果：从回测绩效看 AlphaAgents 的决策能力	9
（4）未来价值与战略定位	9
2.2 摩根大通（J. P. Morgan）：“AI 原生”变革先锋	10
（1）定量分析与高管署名的战略信号	11
（2）战略支柱一：构建专有、可信的 AI 核心技术	11
（3）战略支柱二：通过模拟与自动化决策掌控复杂系统	12
（4）战略支柱三：从物理与另类数据中创造信息优势	13
三、对国内资管机构的启示	14
3.1 战略重构：从顶层设计到差异化破局	14
3.2 技术落地：构建务实高效的人机协同体系	14
3.3 组织变革：锻造面向未来的复合型团队	14
3.4 风险管控：构建可信可靠的 AI 治理体系	15
风险提示	15

图表目录

图1: AlphaAgents 研究参与者信息构成	7
图2: 用于股票分析的多智能体协作与辩论示意图	8
图3: 风险中性投资组合表现	9
图4: 风险规避投资组合表现	9

一、范式重塑：AI 如何变革传统配置模型

1.1 资产配置中的大语言模型（LLMs）

金融市场在很大程度上由叙事、情绪和预期驱动。大型语言模型的出现，首次使机器能够大规模、高精度地理解并量化这些非结构化文本信息，从而将投资分析从纯粹的数字领域拓展至语义领域。

（1）从文本到交易信号：金融情绪分析的演进机制

自然语言处理技术已实现从早期基于词典的方法向基于 Transformer 架构的大语言模型跨越。早期方法仅通过统计文本中正负面词汇数量进行情绪判断，无法理解上下文、反讽及金融专业术语。而基于 Transformer 架构的大语言模型，凭借其自注意力机制，能够精准捕捉词汇在句子中的复杂关系与上下文含义，实现更精准的情绪判断。通用型 LLMs（如标准 GPT）在处理金融文本时常会“水土不服”，因为金融语言具有高度的专业性和上下文依赖性。例如，“鹰派”在金融语境下通常指代紧缩的货币政策，而通用模型可能无法准确理解其含义。因此，创建领域专用 LLMs 尤为重要。创建一个领域专用模型通常遵循一个两步范式。首先，在一个巨大的通用语料库上进行“预训练”，让模型学习通用的语言规律和语法结构。然后，在一个规模较小但带有高质量标签的领域特定数据集上进行“微调”。在微调阶段，模型的参数会根据金融文本的特点及其对应的情绪标签进行微小调整，最终得到一个既懂通用语言又精通金融“行话”的专用模型。

目前金融业界开发了一系列专为金融领域设计的 LLMs：

BloombergGPT：彭博社开发的 500 亿参数模型，基于海量金融文本训练，支持领域内多样化 NLP 任务。

FinGPT：具备开源优势，可快速微调以适应瞬息万变的金融市场。

FinBERT 与 FinLlama：基于通用模型通过金融数据集微调而得，专注于情绪分类等任务，为算法交易提供更细致信号。

（2）应用场景与实践价值

算法交易：LLMs 最直接的应用之一是为算法交易系统提供情绪信号。通过实时分析新闻专线、社交媒体及公司公告，LLMs 能够为个股或整体市场生成量化情绪分数。该分数可作为独立的交易信号，或与其他量化因子结合，共同指导交易决策。

风险管理：除发掘交易机会外，LLMs 在风险管理中的作用日益凸显。它们能够 7×24 小时不间断监控全球信息流，识别可能预示市场情绪突变或潜在风险的早期信号。例如，通过分析某一行业相关新闻中负面词汇频率与强度的骤然上升，模型可在风险事件完全传导至资产价格前，向风险管理人员发出预警。

（3）发展瓶颈与核心挑战

尽管 LLMs 在金融领域应用前景广阔，但其应用也面临严峻挑战，制约了当前应用的广度与深度。

数据偏见与模型幻觉：训练数据中的历史偏见可能被复制放大，而模型生成看似合理实则错误信息的“幻觉”问题，在要求精准的金融领域堪称致命缺陷。

高昂计算成本：大规模训练与部署所需的计算资源构成财务与基础设施壁垒，限制中小机构参与。

可解释性难题：模型决策过程的“黑箱”特性，成为其在核心交易与风险系统中广泛采纳的主要障碍，这也凸显了可解释 AI 的重要性。

金融专用 LLMs 的竞争正演变为围绕专有数据与微调专业知识的军备竞赛。卓越性能源于海量专有金融数据集训练或精准任务微调，这使得掌握独家数据资源的大型金融数据提供商与投资银行凭借数据护城河获得持久优势，可能导致市场集中度进一步提升。

1.2 资产配置中的深度强化学习（DRL）

LLMs 赋予机器“阅读”市场的能力，DRL 则为资产配置提供全新的动态自适应框架，让机器“学习”市场。DRL 的目标并非一次性精准预测，而是学习在长期内实现最优回报的决策策略。

（1）理论演进：从静态优化到自适应智能体

为了深刻理解 DRL 的革新性，有必要将其与传统方法作对比。传统的投资组合优化，如现代投资组合理论，依赖于对资产预期收益、波动率和相关性的静态估计，进而求解出一个理论上“最优”但固定不变的配置方案。然而，金融市场的本质是动态演化的，这使得任何静态配置的效力都难以持续。与此同时，主流的监督学习方法虽然被广泛用于预测特定目标，例如次日股价，但其成败完全系于预测精度之上，而金融市场固有的高随机性使得持续准确的预测近乎奢求。

DRL 则采用了一种截然不同的范式。在此框架下，构建一个“智能代理”，让其通过与“环境”，即模拟或真实的金融市场进行持续交互来学习。在每一个决策时刻，“智能代理”会观察市场的当前“状态”，例如资产价格、交易量、技术指标等，基于此做出“行动”，例如调整各类资产的权重，随后环境会反馈一个“奖励”，例如投资组合夏普比率的改善或账户净值的增长。通过反复经历“观察-行动-奖励”的循环，“智能代理”的最终目标是学会一个最优的“策略”——一个将市场状态映射到投资行动的函数，该函数能够在其整个决策周期内最大化累积奖励。

（2）核心架构与技术实现

主流 DRL 算法在金融领域展现出巨大潜力：学术界和业界正在探索多种 DRL 算法在金融领域的应用，其中几种主流架构表现出巨大的潜力。

演员-评论家方法（Actor-Critic Methods）：这类算法包含两个神经网络。一个是“演员”，负责根据当前状态输出一个行动，即投资组合权重。另一个是“评论家”，负责评估“演员”所采取行动的好坏，并提供反馈信号，指导“演员”调整其策略。

近端策略优化（Proximal Policy Optimization, PPO）：PPO 是目前最先进和最稳定的 DRL 算法之一。它通过在每次更新时限制策略变化的幅度，来确保训练过程的稳定性和可靠性，避免了因策略更新过大而导致的性能崩溃。

深度确定性策略梯度（Deep Deterministic Policy Gradient, DDPG）：DDPG 是专门为连续行动空间设计的算法。这使其非常适合投资组合管理任务，因为资产的权重可以是任意连续的数值，而不是离散的“买入或卖出”决策。此外，研究

人员还在 DDPG 框架中进行了创新，例如引入了做空机制，以及构建了包含价格、交易量、动量等多个维度随时间变化的三维状态空间，为模型提供了更丰富的信息输入。

（3）发展瓶颈与核心挑战

尽管 DRL 在理论和模拟中表现出色，但将其部署到真实的、高风险的金融市场中仍需克服几个关键障碍：

数据依赖与过拟合风险：DRL 模型需要海量、高质量的历史数据才能进行有效训练。对于某些新兴市场或特定资产类别，获取足够的数据可能非常困难。此外，模型存在过拟合的风险，即它们可能学到了训练数据中的伪相关性或噪声，而不是市场真正的内在规律，导致在真实的、未见过的数据上表现不佳。

市场周期的适应性难题：基础的 DRL 算法在应对不同的市场环境时，往往表现出鲜明的“个性”偏差。例如，研究表明，像演员-评论家（AC）这样的模型在牛市中表现激进，能获得高回报，但在熊市中也面临巨大的回撤风险；而近端策略优化（PPO）模型则相反，在熊市中防御性强，但在牛市中收益却相对保守。如何开发一个能够在不同市场周期中都能稳健获利的全天候模型，是一个核心挑战。

高昂的计算成本：训练复杂的 DRL 模型需要巨大的计算资源，这为中小型金融机构设置了较高的技术和财务门槛。

现实世界约束的整合：在真实的投资组合管理中，基金经理需要遵守一系列复杂的风险管理约束，如行业敞口限制、最大持仓比例、交易成本等。如何将这些现实世界的约束无缝地整合到 DRL 的学习框架中，同时又不损害其策略的灵活性，仍然是一个有待解决的技术难题。

1.3 资产配置中的图神经网络（GNNs）

（1）将金融系统概念化为网络

GNNs 应用的前提，是将整个金融系统抽象为一个“网络”。在这个网络中，“节点”代表如银行、对冲基金、保险公司等金融机构，而“边”则代表它们之间的相互关联。这些关联可以有多种形式，例如银行间的同业拆借、如两家机构之间的利率互换合约、多家基金的共同资产持有等。

例如传统模型在分析银行风险时，往往只关注其自身的资产负债表，而忽视了其在网络中的位置和连接关系。这种“孤立主义”的视角，无法捕捉到风险通过网络进行“传染”的动态过程，这也是它们在预测 2008 年金融危机时集体失效的主要原因。而 GNNs 则可以通过将金融系统网络化解解决这一问题。

（2）图神经网络实践：风险传播建模

图神经网络是将深度学习技术扩展到图结构数据的专用神经网络架构。与处理图像或文本的传统神经网络不同，GNNs 能够同时学习节点的自身特征和图的拓扑结构。以银行为例，图神经网络能学习银行的资本充足率、杠杆率以及银行在金融网络中的连接方式。核心思想是通过“消息传递”机制，让每个节点能够聚合其邻居节点的信息，从而更新自身的“表示”。经过多轮迭代，每个节点的表示不仅包含了自身的初始信息，还融入了其在网络中多位邻居的信息，从而感知到其在网络中的局部和全局环境。

此外，图神经网络也可以被训练来进行压力测试。例如，一家银行宣布巨额亏损后，这种压力将如何通过网络传播，以及最终会有多少家机构陷入困境。因此 GNNs 的发展给监管部门提供了强大的模拟工具。它可以通过学习节点重要性，帮助监管机构更准确地识别出那些在网络中处于核心位置、一旦倒闭将引发最大规模连锁反应的“大到不能倒”的机构。

（3）对宏观审慎政策和投资组合对冲的启示

对于监管者：图神经网络使其能够对整个金融体系的稳健性进行动态评估和压力测试。基于 GNNs 的分析结果，监管机构可以设计更具针对性的宏观审慎政策，例如，对网络中连接度最高、最关键的机构施加更高的资本要求，或者限制某些可能加剧系统性风险的交易类型。

对于投资者：理解金融网络中的相互依赖关系，可以帮助投资者构建更有效的投资组合对冲策略。传统的对冲可能只关注对冲掉某个特定公司的风险。而图神经网络的分析，投资者可以识别出高度关联的“公司集群”。当投资组合持有一个集群中的公司时，更有效的对冲策略可能是卖空该集群中的另一个高度相关的公司，或者购买针对整个集群的信用违约互换，从而对冲掉整个“社区”的风险，而不仅仅是单个节点的风险。

二. 机构实践：头部资管公司 AI 赋能资产配置

2.1 贝莱德（BlackRock）：AlphaAgents——基于多智能体系统的股票组合构建实践

贝莱德作为全球最大的资产管理公司，近年来持续加大在人工智能领域的投入，致力于将前沿 AI 技术融入投资决策流程。除了本团队在 AI 赋能资产配置（六）—海内外资管机构 AI 大模型应用探索中提到的贝莱德阿拉丁智能投资平台、AI Copilot、系统化主动权益投资策略、主题机器人等创新工具外，贝莱德 2025 年 8 月 AI 团队发表的最新的文献《AlphaAgents: Large Language Model based Multi-Agents for Equity Portfolio Constructions》一文，为我们提供了一个前所未有的、具体的内部视角，来观察其在前沿 AI 领域的布局。本章节将以该论文为核心，剖析贝莱德如何探索下一代 AI 技术，以解决资产管理中最为复杂的认知与决策难题。

图1: AlphaAgents 研究参与者信息构成

所属层级	姓名	职务/头衔	角色与职责分析
高级管理层	Stefano Pasquali	董事总经理； 投资人工智能建模与研究负责人	公司投资 AI 战略与研究方向的核心领导者，负责重要业务线。
高级管理层	Dhagash Mehta	应用人工智能/机器学习研究负责人；曾任首席研究员	公司 AI 研发领域的关键领导者，负责将前沿机器学习技术转化为实际的投资应用。
总监及副总裁级别	Tianjiao Zhao	总监；应用 AI 建模负责人	负责具体的 AI 建模项目与执行，是连接高层战略与技术实现的桥梁。
总监及副总裁级别	Stokes Jones	用户研究副总裁	作为经验丰富的中层管理者，负责确保 AI 工具与系统符合投资经理和分析师的实际工作需求。
专业研究人员	Jingrao Lyu/ Harrison Garber	数据分析师 / 阿拉丁金融工程投资；AI 团队分析师	团队中的核心技术执行者，数据科学家或机器学习工程师，负责实现研究想法、进行实验分析。

资料来源：公司官网，国信证券经济研究所整理

该研究由 Tianjiao Zhao、Jingrao Lyu、Stokes Jones、Harrison Garber、Stefano Pasquali、Dhagash Mehta 完成。、作者团队涵盖了从最高层的董事总经理到一线的核心研究员，表明这项前沿研究并非孤立的学术探索，而很可能是由公司高层推动、跨层级协作的**战略性项目**。

（1）战略意图：超越“数据处理”，迈向“决策智慧”

AlphaAgents 项目的核心战略意图，并非简单地提升数据处理效率，而是旨在解决当前人机投资决策中的两大核心痛点。首先是**克服人类分析师在决策过程中容易受到损失厌恶、过度自信等认知偏见的影响**，而这种偏见可能导致错失投资机会。其次是弥补传统 AI 的不足，早期 AI 应用多为任务导向且依赖结构化数据，而 AlphaAgents 则旨在利用大型语言模型**处理非结构化数据**，并进行复杂的推理决策。更重要的是，它试图通过内部制衡机制，减少 AI 模型自身“幻觉”问题。

（2）核心机制剖析：从协作到辩论的双层决策流程

AlphaAgents 的创新之处在于其借助行为金融学取向，模拟人类投资委员会的“协作与辩论”机制。系统首先通过专业分工与协作，完成对事实的全面收集与整理。为此，系统设立了三个具有明确角色分工的 AI 智能体，各司其职：

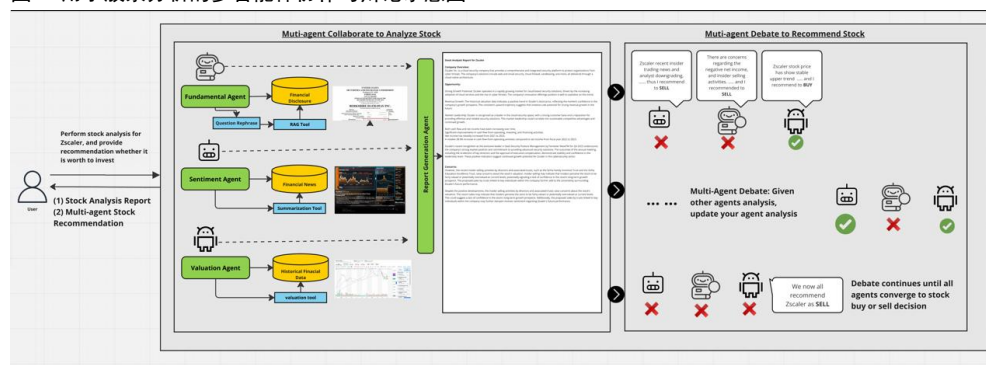
基本面分析智能体：基于 10-K/10-Q 财报数据，进行财务健康度、经营效率与成长性分析；

情绪分析智能体：处理新闻、分析师评级与市场情绪，评估短期市场情绪对股价的影响；

估值分析智能体：基于历史价格与成交量数据，计算波动率、收益率等技术指标，评估股票估值水平。

在技术实现上，整个框架的认知能力由先进的大型语言模型驱动。论文明确指出，在对多种 GPT 模型进行实验后，研究团队最终选择 GPT-4o 作为其系统的核心模型，并将其用于信息检索增强（RAG）等关键任务中。

图2：用于股票分析的多智能体协作与辩论示意图



资料来源：Zhao, Tianjiao, et al. "AlphaAgents: Large Language Model based Multi-Agents for Equity Portfolio Constructions." arXiv preprint arXiv:2508.11152 (2025)，国信证券经济研究所整理

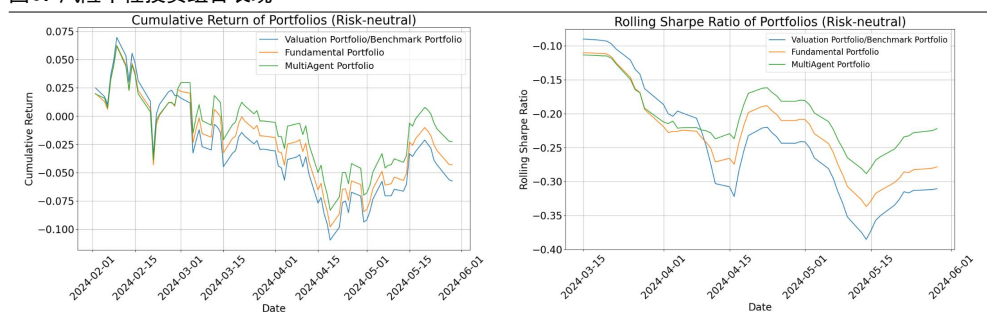
在事实分析的基础上，系统进入第二层：**对抗性辩论与共识辩论环节**。各个智能体会接收到其他智能体的分析结论，并就是否“买入”或“卖出”展开多轮讨论。这个过程被设计为强制性的，直到所有智能体达成一致共识，最终决策才会形成。这种机制的价值在于，它通过不同视角的碰撞和交叉验证，极大地提升了最终结论的稳健性，并有效过滤了单一模型的潜在错误。

(3) 实证结果：从回测绩效看 AlphaAgents 的决策能力

为验证框架的有效性，该研究进行了一系列严格的回测实验，这些实验不仅展示了模型的投资表现，更揭示了其在不同风险偏好下的决策差异和多代理协作的实际价值。在实验设计上，研究者随机选择了 15 只科技股作为投资备选池和业绩比较基准，AI 智能体使用 2024 年 1 月的数据进行分析辩论，并在 2024 年 2 月 1 日构建了所有入选股票均等权重的投资组合。为作对比，研究还构建了仅由单个 AI 智能体决策的组合。此外，研究者通过不同的提示工程，构建了“风险中性”与“风险规避”两种策略模式，以测试 AI 对不同投资目标的适应性。

在风险中性策略的回测中，累计回报、经过风险调整后的夏普比率，多代理投资组合的表现显著优于所有单代理组合及市场基准。论文分析指出，这一成功源于多代理框架的协同效应有效地平衡了短期的市场洞察与长期的基本面分析。

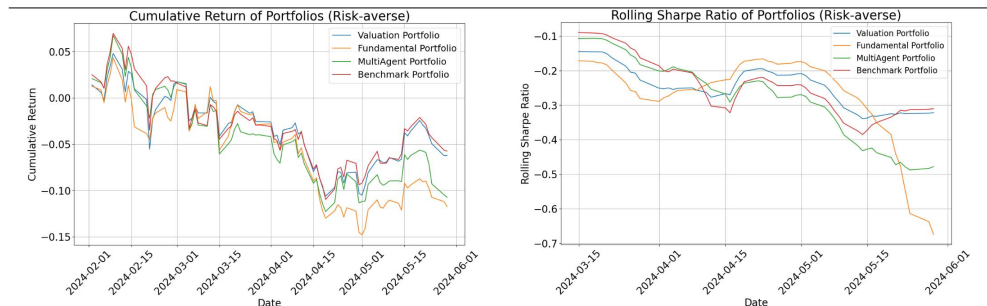
图3：风险中性投资组合表现



资料来源：Zhao, Tianjiao, et al. "AlphaAgents: Large Language Model based Multi-Agents for Equity Portfolio Constructions." arXiv preprint arXiv:2508.11152 (2025)，国信证券经济研究所整理

而在风险规避策略的回测中，所有由 AI 选择的投资组合表现均落后于基准。但这并非模型的失败，而是策略执行成功的体现。这是因为，测试期间科技板块整体处于牛市，市场表现强劲，风险规避型 AI 智能体通过分析正确地识别并排除了那些高波动性的科技股。虽然这导致组合错失了市场的上涨收益，但也成功实现了规避下行风险的策略目标。在风险控制能力上，多代理组合与单代理组合相比，它的波动率更低，最大回撤更小。

图4：风险规避投资组合表现



资料来源：Zhao, Tianjiao, et al. "AlphaAgents: Large Language Model based Multi-Agents for Equity Portfolio Constructions." arXiv preprint arXiv:2508.11152 (2025)，国信证券经济研究所整理

(4) 未来价值与战略定位

回测结果为评估 AlphaAgents 的未来价值提供了坚实依据。我们认为，它在贝莱

德的战略版图中扮演着至关重要的角色。它代表了人机协作模式的根本性升级，其定位不再是被动提供数据分析的工具，而是一个能够主动进行逻辑推理、权衡利弊的“AI 决策伙伴”。其未来价值体现在：

解决 AI 的信任问题：首先，多智能体辩论提高了分析严谨性并减少了 AI 幻觉问题的发生。系统生成的讨论日志与决策轨迹也为 AI 决策的可解释性提供了重要支撑，这符合监管与机构客户对透明投资流程的需求，便于人工介入监督和事后分析。

全流程 AI 赋能潜力：AlphaAgents 系统虽聚焦于股票选择，但其模块化、可扩展的智能体架构为更广泛的资产配置场景提供了技术基础。论文指出其信号可作为均值方差或 Black-Litterman 模型的输入，成为资产配置中的“信号”。换言之，AI 智能体在底层生成超额收益预期与风险画像，上层配置则在风险预算框架下进行权重优化，可以实现“AI 观点 → 配置参数 → 执行监控”的链路闭环。这一思路与贝莱德内部的 Aladdin 投资与风控平台天然衔接，使得研究成果可快速落地于机构资产配置 workflow。

寻找判断性 Alpha：资产管理中的 Alpha 分为系统性 Alpha 和判断性 Alpha，前者可通过量化模型捕捉，而后者依赖人类分析师的深度认知和判断。AlphaAgents 的辩论机制，将人类分析师的“判断性智慧”模型化、系统化。

此外 AlphaAgents 并非学术性实验，而是由管理层直接牵引的生产级研究。Pasquali 负责投资 AI 部门，Mehta 领导应用 AI 研究，两人均在贝莱德 AI 治理委员会中承担关键角色；Zhao 与 Lyu 所在的亚特兰大团队则负责 Aladdin 系统的 AI 集成与建模工程。这暗示着贝莱德可能正以“研究—系统—产品”路径推动 AI 资产配置，从底层信号生成、风险评估到执行监控形成闭环。

另一则消息帮助我们推断贝莱德运用 AI 进行资产配置的工作流：BlackRock 在 2025 年投资者日上公开了一套面向投研的“虚拟投资分析师”——Asimov。此平台已经在贝莱德基本面权益业务中使用，它能通宵自动扫描公司文件、研究笔记与邮件以生成投资洞见。贝莱德 COO Rob Goldstein 把它描述为“可视化 AI 智能体”。

结合阿拉丁平台的介绍，贝莱德在 AI 赋能资产配置中，已形成或正在形成一套完整的工作流链条：Asimov 负责把资料准备好并调用代理；AlphaAgents 或类似项目把“买/卖+置信度+风险偏好”变成信号输入（ μ 、 τ/Ω ）；Aladdin 负责做权重、控风险并下单执行。

2.2 摩根大通（J. P. Morgan）：“AI 原生”变革先锋

摩根大通 CEO 杰米·戴蒙在 2025 年 10 月的最新采访中强调：“我们并非仅在做试点，而是致力于成为一家全面 AI 化的企业。”在全球银行业同质化竞争日益加剧的背景下，摩根大通正积极推进区别于其他资管机构的“AI 原生”转型。

摩根大通的 AI 战略呈现出鲜明的学术研究导向，其核心在于以严谨的学院派研究方式，在规模化落地解决方案之前，优先解决资产配置领域的根本性理论问题。摩根大通并非仅将 AI 作为应用工具，而是将其视为重塑金融理论与实践的根本力量。

本小节通过分析摩根大通 AI 研究部门在 2024 - 2025 年期间发表的出版物，探讨其如何将 AI 能力应用于资产配置策略。

（1）定量分析与高管署名的战略信号

摩根大通每年在 AI 上投入 20 亿美元。2025 年，摩根大通的技术预算高达 180 亿美元，其中 AI 是其核心部分。通过设立专门的 AI 研究部门，摩根大通系统性地推进金融 AI 的基础研究。该部门聚焦多个前沿方向，包括 AI 智能体与混合推理、AI 规划与知识管理、AI 优化与决策、金融领域基础模型、合成数据、时间序列与行为分析、多模态文档处理与分析，以及可信 AI、透明度与安全机制等。

部门主管 Manuela Veloso 是卡内基梅隆大学的荣休教授，在机器人、规划和多智能体系统领域享有盛誉，代表了深厚的学术根基。与此同时，如 AI 研究执行官 Sumitra Ganesh 等管理者，则将机器学习技术应用于客户情报与销售等一线业务，具备丰富的实战经验。这种组织架构与研究定位不同于传统金融机构的技术应用部门，它兼具顶尖高校计算机科学系的前沿探索与产品团队的落地导向，构建了独特的“人才护城河”。既能吸引有志于发表高水平论文、解决基础科学问题的博士人才，又能确保研究成果快速转化为具备防御能力的商业能力。

据摩根大通 AI 部门网站披露，2024 至 2025 年间，该部门共发表 140 篇出版物，包括期刊论文 15 篇、会议论文 63 篇、预印本 38 篇及其他类型 3 篇。署名作者包括 AI 研究全球主管 Manuela Veloso 博士、AI 研究执行官 Sumitra Ganesh 博士、AI 研究总监 Daniel Borrajo 博士与 Vamsi Potluru 博士，以及多位研究负责人。值得关注的是，期间共有 8 篇论文发表于 AAAI 人工智能顶级会议。这些研究成果紧密围绕公司核心战略目标展开，分析其研究布局有助于洞察其未来 AI 发展蓝图，识别其 AI 战略路径。

（2）战略支柱一：构建专有、可信的 AI 核心技术

该支柱构成摩根大通 AI 战略的根基：打造兼具强大性能与高度可信的专有 AI。其内涵包括两个相辅相成的方面：首先，为所有 AI 应用构建坚实的信任与安全基础；其次，掌握基础模型与生成式 AI 的核心技术，确保对关键 AI 能力的端到端控制。

将信任与安全工程作为基石

对摩根大通这类全球系统重要性银行而言，“可信 AI”并非附加功能，而是所有 AI 计划得以实施的基石。摩根大通在此领域的研究带有明显的防御性战略意图，旨在主动应对监管审查、规避重大财务与声誉风险，并维护客户的长期信任。

论文《在无法访问敏感属性的情况下进行公平分类的超参数调整》（Hyper-parameter Tuning for Fair Classification without Sensitive Attribute Access）解决了现实世界中的一个关键难题：由于法律和隐私限制，企业通常不能使用种族等敏感属性来训练或验证模型。该研究提出一种在不直接接触敏感数据的前提下实现模型公平性的方法，对在遵循欧盟《通用数据保护条例》等法规框架下大规模部署公平 AI 模型具有重要意义。

由 AI 研究总监 Vamsi Potluru 合著的论文《离散去噪扩散模型的内在隐私属性》（On the Inherent Privacy Properties of Discrete Denoising Diffusion Models）对扩散模型在生成数据时的隐私保障能力进行了“开创性的理论探索”，并聚焦于“逐实例差分隐私”。其目标并非生成简单的仿真数据，而是创造具备数学上可证明隐私保护能力的高质量合成数据。此类数据可用于训练其他复杂金融模型，同时避免泄露任何真实客户敏感信息。另一篇论文《SafeAR：通过风险感知策略实现更安全的算法纠正》（SafeAR: Towards Safer Algorithmic Recourse by Risk-Aware Policies）则关注决策后的挑战，确保银行向客户提供的修正建

议安全有效，对提升客户体验与合规性至关重要。

掌控基础与生成式 AI 能力

摩根大通不满足于仅作为第三方大语言模型的被动使用者，而是致力于成为专有模型的创造者与掌控者。

由 Sumitra Ganesh 合著的论文《Collab：使用 LLMs 混合进行对齐的受控解码》（Collab: Controlled Decoding using Mixture of LLMs for Alignment）提出了一种新技术，用于精确控制大语言模型的输出。在金融领域，模型输出必须严格遵守事实、法律和合规约束，因此这种控制能力至关重要。因此这种控制能力尤为关键。“LLMs 混合”技术表明该行正在构建复杂的模型集成系统，以在内容生成创造性与事实准确性之间取得平衡，这对自动生成市场评论、撰写研究报告等应用具有关键价值。

论文《基础模型中的不确定性与幻觉：可靠 AI 的下一个前沿》（Uncertainty and Hallucination in Foundation Models: The Next Frontier in Reliable AI）直面当前大语言模型的主要弱点。通过聚焦于“不确定性量化”，摩根大通正在开发让模型“知道自己不知道什么”的方法。在高风险决策流程中应用大语言模型，其首要前提是模型必须具备评估自身输出置信度的能力。唯如此，模型才具备可信性，并可在客户互动、风险评估及投资决策等场景中投入实际业务。

Vamsi Potluru 合著的《GraphMaker：扩散模型能否生成大型带属性图？》（GraphMaker: Can Diffusion Models Generate Large Attributed Graphs?）将生成式 AI 的应用范围从文本和图像，拓展到了图结构数据。金融数据本质上正是交易网络、对手方关系、股权结构等庞大复杂图网络的体现，因此，该技术能生成大规模、高保真的合成金融图数据。其核心价值在于，可在完全无需真实敏感数据的前提下，高效训练欺诈检测、反洗钱与系统性风险分析系统。

综合来看，此战略一方面将合规成本中心转化为竞争护城河，使信任成为商业资产。当监管机构推出更严格 AI 法规时，摩根大通凭借其积累的专业知识与已验证模型将从容应对，而竞争对手可能被迫仓促应战或依赖昂贵的第三方解决方案。另一方面，该行正在构建垂直整合的专有 AI 技术栈。通过将信任工程与专有模型开发相结合，他们正在创造一个独一无二的“金融大脑”——一个使用其专有数据训练、与公司风险偏好一致、且完全可供审计的 AI 核心。这构成了其所有后续 AI 应用的基础和护城河。那些仅仅在通用大模型之上进行简单封装的竞争对手，将无法在性能、安全性或成本效益上与摩根大通深度整合的系统相匹敌。但这是一项需要长期、巨大资本投入的战略，只有极少公司能够复制其 AI 战略。

（3）战略支柱二：通过模拟与自动化决策掌控复杂系统

第一个支柱是构建 AI 引擎，而第二个支柱则聚焦于如何使用这个引擎来理解和驾驭复杂的金融世界。摩根大通正将两项最前沿的 AI 技术——多智能体模拟和强化学习——相结合，旨在创建类似属于金融“风洞实验室”，并在其中训练智能代理以做出最优决策。

模拟市场与经济的未来

该战略方向旨在构建能够模拟复杂经济系统的精密模型，其目标远超传统预测模型，致力于创建可用于测试战略、理解系统性风险及探索市场动态的虚拟金融环境。

由 AI 研究执行官 Sumitra Ganesh 合著的论文《交易网络中均衡价格的去中心化

收敛》（Decentralized Convergence to Equilibrium Prices in Trading Networks）是该领域的奠基性研究，探索去中心化环境中的价格形成机制，与不透明的场外交易市场具有直接对应关系。《无需货币的资源分配中的正则化比例公平机制》（Regularized Proportional Fairness Mechanism for Resource Allocation Without Money）等研究则深入到博弈论和机制设计的核心，其应用场景包括资产负债表容量分配或交易对手限额设定等。《ADAGE：一个用于自适应基于代理人建模的通用两层框架》（ADAGE: A Generic Two-layer Framework for Adaptive Agent based Modelling）则为构建此类复杂模拟提供工具包，使其能够灵活模拟从消费信贷市场至机构交易行为等多种经济系统。

通过强化学习优化高风险决策

摩根大通的研究表明，其正致力于创造更强大、更高效、更具泛化能力的强化学习（Reinforcement Learning, RL）代理，以驱动下一代决策系统。《占用信息比：无限视界、信息导向、参数化策略搜索》（Occupancy Information Ratio: Infinite-Horizon, Information-Directed, Parameterized Policy Search）和《强化学习中的近似等变性》（Approximate Equivariance in Reinforcement Learning）等论文旨在改进强化学习核心算法，打造样本效率更高的学习代理，对算法交易和动态资产组合优化等复杂金融场景尤为重要。

由 Manuela Veloso 合著的论文《资源使用和持续时间不确定的作业的容量规划与调度》（Capacity planning and scheduling for jobs with uncertainty in resource usage and duration）表明，这种强大的决策能力不仅用于外部市场，也被应用于优化公司内部运营，例如优化其庞大的“本地网格计算环境”从而降低 AI 研发成本、提高效率。

通过将市场模拟与强化学习相结合，摩根大通正在将自己从一个单纯的市场参与者，提升为一个市场架构师。他们可以在安全的虚拟环境中测试新策略、预测市场对冲的反应，并训练出能够自动执行最优决策的 AI 代理。这不仅能创造专有的 alpha，还能通过优化内部运营，形成一个能够自我强化的增长飞轮。

（4）战略支柱三：从物理与另类数据中创造信息优势

其分析能力从金融市场的数字领域，延伸到现实的物理世界。通过将先进的 AI 技术应用于非结构化的真实世界数据，他们正在创造新的信息来源，以对传统财务报表中不可见的风险进行定价，并识别新的机遇。

用于资产验证与分析的计算机视觉

由 AI 研究总监 Daniel Borrajo 合著的论文《Veri-Car：面向开放世界的车辆信息检索》（Veri-Car: towards open-world vehicle information retrieval）开发了一个能够从图像中识别车辆品牌、型号、年份和状况的系统。其商业应用前景直接而强大：为汽车贷款和保险业务实现自动化验车，为供应链金融验证抵押资产，以及为资产支持证券进行更精确的估值。该系统强调的“开放世界”能力，意味着它被设计用来处理从未见过的、新上市的车型，这极大地增强了其实用性。

用于经济与气候风险的地理空间分析

《实现先进的土地覆盖分析：使用动态世界数据集进行预测建模的集成数据提取管道》（Enabling Advanced Land Cover Analytics: An Integrated Data Extraction Pipeline for Predictive Modeling with the Dynamic World

Dataset) 这篇论文表明, 摩根大通正在构建处理和分析海量地理空间数据集的能力。这使其能够监测和预测现实世界的经济活动, 例如城市化进程、森林砍伐或农作物产量。这些信息在商品交易、基础设施投资、ESG 分析, 以及模拟气候变化对其贷款组合造成的物理风险等方面, 都有直接的应用价值。

摩根大通正在构建一座连接物理世界事件与金融结果的桥梁。这种独特的能力使其能够创造专有的数据流, 在那些与实体资产和物理条件密切相关的市场中获得决定性优势。他们的研究重点并非停留在计算机视觉或卫星图像技术本身, 而是构建“端到端的数据管道”和“信息检索系统”, 目标是将原始的像素数据转化为结构化的、可供金融模型分析的信息。通过掌握对这些“另类数据”的分析能力, 摩根大通能够比那些依赖传统数据源的竞争对手更精确地为风险定价。他们正在有效地扩展“可知”和“可定价”的边界, 从而创造一种持久的信息优势。

综合分析表明, 摩根大通正在进行一项长期的、资本密集型的转型, 目标是成为一家“AI 原生”的金融机构。其战略的终点并非简单地成为 AI 的使用者, 而是要成为 AI 的掌控者。2024 至 2025 年发表的这些研究论文, 不仅仅是学术成果, 更是这家金融巨擘未来形态的蓝图。

三. 对国内资管机构的启示

海外头部机构在 AI 赋能资产配置领域的探索与实践, 清晰揭示了行业发展的未来路径。面对这一历史性机遇, 国内资管机构需突破传统思维, 以战略高度进行系统性布局, 将 AI 从辅助工具升级为核心竞争力。具体而言, 应从以下四个维度构建差异化能力:

3.1 战略重构：从顶层设计到差异化破局

机构决策层必须率先实现认知升级, 建议成立跨部门 AI 战略委员会, 由核心管理层直接推动, 制定符合公司特色的转型路线图。该路线图需明确资源投入、阶段目标及评估标准, 确保 AI 战略与业务战略深度融合。

在实施路径上, 应摒弃“全面对标”的思维, 转而寻求差异化突破。国内机构需清醒认识到自身在数据规模、算力资源和顶尖人才上与全球巨头的客观差距, 不宜盲目追求与摩根大通相似的“大而全”策略。更现实的路径是采取“聚焦突破”策略, 结合自身在特定市场、特定资产类或特定策略上的传统优势, 选择 AI 应用的切入点和主攻方向。例如, 可优先利用 LLMs 深度挖掘 A 股市场独特的政策文本、上市公司公告和社交媒体情绪, 构建具有本土化特色的 Alpha 信号源。

3.2 技术落地：构建务实高效的人机协同体系

技术实施层面应采取“三步走”策略: 首先夯实数据基础, 重点突破数据治理瓶颈, 构建统一数据平台, 同时战略性布局另类数据资源; 其次在模型选择上保持务实, 优先基于开源框架和预训练模型进行微调优化, 以可控成本快速验证价值; 最关键的是在应用层面确立“人机协同”原则, 将 AI 定位为投研团队的“智能副手”。贝莱德的实践表明, AI 并非要取代投资经理, 而是通过多智能体协作、大规模回测等能力, 极大拓展人类专家的认知边界。国内机构可借鉴此模式, 在信息处理、策略生成和风险分析等环节率先实现 AI 赋能, 既提升决策效率, 又保持人类在复杂判断和最终决策中的核心地位。

3.3 组织变革：锻造面向未来的复合型团队

AI 时代的竞争本质是人才竞争。必须打破传统的部门壁垒，构建融合投资洞察、数据科学和工程实现的跨职能团队。这种组织模式的创新比技术本身更具挑战性，也更具价值。人才建设需采取“外部引进与内部培养”双轨制：一方面积极引入具备 AI 与金融复合背景的核心人才；另一方面大规模开展现有团队的 AI 能力提升计划，梳理金融工作流。同时，要推动组织文化向“数据驱动、敏捷迭代”转型，建立容错机制和创新激励，让 AI 能力在真实业务场景中持续进化。

3.4 风险管控：构建可信可靠的 AI 治理体系

金融机构应用 AI 必须坚守风险底线。需要建立覆盖模型全生命周期的治理框架，从数据源头、算法设计到部署监控形成完整闭环。特别要关注大语言模型的“幻觉”问题、强化学习模型的过拟合风险等新型挑战，通过严格的回测验证和持续监控确保系统稳健性。此外，要前瞻性布局“可信 AI”能力建设，将透明度、公平性等要求融入系统设计。主动参与行业规范和标准制定，将合规能力转化为竞争优势，为 AI 技术的规模化应用扫清障碍。

风险提示

警惕大语言模型内在缺陷与“幻觉”风险；警惕 AI 模型的“黑箱”特性，AI 投资存在潜在的合规与法律风险；文中所提及公司仅做客观事件汇总，不涉及主观投资建议。

免责声明

分析师声明

作者保证报告所采用的数据均来自合规渠道；分析逻辑基于作者的职业理解，通过合理判断并得出结论，力求独立、客观、公正，结论不受任何第三方的授意或影响；作者在过去、现在或未来未就其研究报告所提供的具体建议或所表述的意见直接或间接收取任何报酬，特此声明。

国信证券投资评级

投资评级标准	类别	级别	说明
报告中投资建议所涉及的评级（如有）分为股票评级和行业评级（另有说明的除外）。评级标准为报告发布日后 6 到 12 个月内的相对市场表现，也即报告发布日后的 6 到 12 个月内公司股价（或行业指数）相对同期相关证券市场代表性指数的涨跌幅作为基准。A 股市场以沪深 300 指数（000300.SH）作为基准；新三板市场以三板成指（899001.CSI）为基准；香港市场以恒生指数（HSI.HI）作为基准；美国市场以标普 500 指数（SPX.GI）或纳斯达克指数（IXIC.GI）为基准。	股票 投资评级	优于大市	股价表现优于市场代表性指数 10%以上
		中性	股价表现介于市场代表性指数 $\pm 10\%$ 之间
		弱于大市	股价表现弱于市场代表性指数 10%以上
		无评级	股价与市场代表性指数相比无明确观点
	行业 投资评级	优于大市	行业指数表现优于市场代表性指数 10%以上
		中性	行业指数表现介于市场代表性指数 $\pm 10\%$ 之间
		弱于大市	行业指数表现弱于市场代表性指数 10%以上

重要声明

本报告由国信证券股份有限公司（已具备中国证监会许可的证券投资咨询业务资格）制作；报告版权归国信证券股份有限公司

关本报告的摘要或节选都不代表本报告正式完整的观点，一切须以我公司向客户发布的本报告完整版本为准。

本报告基于已公开的资料或信息撰写，但我公司不保证该资料及信息的完整性、准确性。本报告所载的信息、资料、建议及推测仅反映我公司于本报告公开发布当日的判断，在不同时期，我公司可能撰写并发布与本报告所载资料、建议及推测不一致的报告。我公司不保证本报告所含信息及资料处于最新状态；我公司可能随时补充、更新和修订有关信息及资料，投资者应当自行关注相关更新和修订内容。我公司或关联机构可能会持有本报告中所提到的公司所发行的证券并进行交易，还可能为这些公司提供或争取提供投资银行、财务顾问或金融产品等相关服务。本公司的资产管理部门、自营部门以及其他投资业务部门可能独立做出与本报告中所意见或建议不一致的投资决策。

本报告仅供参考之用，不构成出售或购买证券或其他投资标的的要约或邀请。在任何情况下，本报告中的信息和意见均不构成对任何个人的投资建议。任何形式的分享证券投资收益或者分担证券投资损失的书面或口头承诺均为无效。投资者应结合自己的投资目标和财务状况自行判断是否采用本报告所载内容和信息并自行承担风险，我公司及雇员对投资者使用本报告及其内容而造成的一切后果不承担任何法律责任。

证券投资咨询业务的说明

本公司具备中国证监会核准的证券投资咨询业务资格。证券投资咨询，是指从事证券投资咨询业务的机构及其投资咨询人员以下列形式为证券投资人或者客户提供证券投资分析、预测或者建议等直接或者间接有偿咨询服务的活动：接受投资人或者客户委托，提供证券投资咨询服务；举办有关证券投资咨询的讲座、报告会、分析会等；在报刊上发表证券投资咨询的文章、评论、报告，以及通过电台、电视台等公众传播媒体提供证券投资咨询服务；通过电话、传真、电脑网络等电信设备系统，提供证券投资咨询服务；中国证监会认定的其他形式。

发布证券研究报告是证券投资咨询业务的一种基本形式，指证券公司、证券投资咨询机构对证券及证券相关产品的价值、市场走势或者相关影响因素进行分析，形成证券估值、投资评级等投资分析意见，制作证券研究报告，并向客户发布的行为。

国信证券经济研究所

深圳

深圳市福田区福华一路 125 号国信金融大厦 36 层

邮编：518046 总机：0755-82130833

上海

上海浦东民生路 1199 弄证大五道口广场 1 号楼 12 层

邮编：200135

北京

北京西城区金融大街兴盛街 6 号国信证券 9 层

邮编：100032