

生成式人工智能在金融领域的应用及影响研究

撰文_董思曼

近年来，生成式人工智能成为全球最具影响力的创新科技，被越来越多的金融机构广泛运用于客户营销、信用评级、风险识别、反欺诈监测等多个场景，与金融行业融合度逐渐加深，有助于提升金融服务效率、促进普惠金融并提高监管效能，但同时增加了算法歧视、数据安全、市场垄断等风险。本文对生成式人工智能在金融领域的运用情况及潜在风险进行梳理，并就加强对金融领域运用生成式人工智能的监督与管理，规范金融业发展有针对性地提出建议。

生成式人工智能在金融领域的运用

人工智能技术正从分析式进化为生成式。分析式人工智能是指利用机器学习算法计算数据中的条件概率分布，根据已有数据集进行分

析、判断和预测的模型，主要应用于广告推荐、辅助决策、数据分析等领域。与分析式人工智能不同的是，生成式人工智能可以“无中生有”，通过对大量的数据集进行训练从而生成新的数据集，并创造出文章、图片、音乐、计算机代码等新内容。相比之下，生成式人工智能具有更强的理解、推理和生成能力，应用场景更为广泛。数据是人工智能最重要的生产要素，金融行业在日常业务中积累了海量用户及交易数据，因此生成式人工智能与金融领域具有天然的契合性，这使得生成式人工智能在金融领域有诸多应用场景。

客户营销与服务

生成式人工智能技术可以通过分析海量的金融数据、新闻、社交媒体等信息，为投资者提供股票、基金、债券等金融产品的评估和预测，

以及投资策略和建议。例如，摩根士丹利挑选了 10 万份数据集训练 GPT-4 模型，使其可以分析新闻报道、社交媒体帖子和财务报表等，以识别模式并预测股价。生成式人工智能还能帮助金融机构识别具有可持续发展潜力的项目，筛选符合 ESG 标准的投资标的。例如，挪威主权财富基金将人工智能使用标准引入其所投资的企业，以更有效地评估和监督被投资企业的 ESG 表现^①。生成式人工智能还有助于银行获取优质客户。国内大型银行机构普遍应用人工智能技术来提升获客能力，拓宽贷款客户覆盖面。据对全国 200 多家银行的调研显示，大型银行系统平均上云率已达 69%，平均分布式数据库实例数接近 2500 个^②。这些人工智能技术的应用，加速了大型银行对优质客户的垄断步伐。

信用评分与评估

传统的信用评估方法难以准确判断“薄档案”客户的信誉状况，而生成式人工智能可以通过分析客户的社

交媒体活动、在线购物行为、网络搜索记录等数据，获取其消费习惯、兴趣爱好、社交关系等信息，更全面了解客户的信誉状况，有助于贷前信用分析。例如，菲律宾联合银行运用生成式人工智能技术开发了一种信用评分系统（credoapp），该系统能够利用非传统数据源（如手机支付记录、社交媒体活动等）来评估个人的信用状况。在我国，蚂蚁集团的芝麻信用，京东白条的小白信用均开发了智能信用评分系统，自动收集并分析客户的各类信息，快速评估客户的信用风险。湖北“鄂融通”“汉融通”等地方融资信用信息平台，通过整合多源数据，形成全面的企业信用信息库，对客户群体进行信用画像，支持银企融资对接，推动地方经济发展。

风险识别与预警

通过对海量数据的挖掘和分析，生成式人工智能能够识别银行、保险等各个金融领域的风险并进行预警。大型商业银行通常拥有庞大的客户基础和丰富的信贷数

据，生成式人工智能技术能够帮助其更准确地评估借款人的历史数据、财务状况、征信记录等信息，识别出潜在违约风险，以便在贷款审批和风险管理过程中做出更科学的决策。生成式人工智能还可以通过自动识别给定数据集中的异常值，进行反洗钱、反恐怖融资和诈骗监测。例如，全球领先的支付技术公司 Visa 在其支付网络 VisaNet 运用生成式人工智能技术，实时分析每笔交易的交易金额、地点、时间，商户信息、持卡人行为模式等多个维度数据，对短时间内的大量小额交易、跨境交易中的异常资金流向等异常交易进行迅速识别。据《福布斯》报道，2016 年 6 月 23 日，英国通过全民公投决定“脱欧”后的几天内，交易员通过访问美国金融科技公司 Kensho 人工智能数据库，能够快速预测英镑的持续下跌。

反欺诈监测

在银行零售业务中，反欺诈的本质是对实施欺诈人员进行伪造身份、联系方式、设备信息、资产信息等虚假

① Nassir, I. K. (2023). Generative AI in Finance: Evolution, Deployment, Risks, and Policy Implications. OECD. Retrieved from [OECD document source] or <http://www.oecd.org/termsandconditions> (Accessed on [date of access]).

② 数据来源：王伟. 强化金融科技赋能推动中小银行数字化转型——访中国人民银行科技司司长李伟 [J]. 金融电子化, 2023, (8): 7-8.



图片来自网络

信息的识别。传统规则的反欺诈模型已无法应对日益演进的欺诈模式和欺诈技术，加剧了银行风控的难度。而生成式人工智能的反欺诈模型可以通过分析交易数据、行为模式等，发现异常或可疑活动，从而帮助金融机构识别和预防欺诈行为。武汉众邦银行运用语音识别、人脸识别、活体检测等智能技术增强认证风险识别，提高客户身份识别效率，其贷款申请环节采用 OCR 身份证识

别及人脸识别技术，并利用大数据风控模型进行自动审批，既做到精准识别客户，又提高了业务审批效率。

金融领域运用生成式人工智能的潜在风险

偏见、歧视与可解释性风险

训练数据的来源和质量以及模型自身的决策机制增加了生成式人工智能的偏见、歧视与可解释性风险。首先，如果训练数据或模型

本身带有偏见或歧视性，那么模型在训练的过程中也会受到这些偏见的影响。例如，美国开放人工智能研究中心（OpenAI）的人工智能产品 GPT-3.5、GPT-2，以及脸书母公司 Meta “元”的 Llama 2 等大语言模型倾向于将工程师、教师和医生等更多元、地位更高的工作分配给男性，而经常将女性与传统上被低估或被社会污名化的角色挂钩。其次，生成式人工智能的决策机制通常基于复杂的

数学公式和算法，使得人们很难完全理解模型是如何做出决策的。例如，面对生成式人工智能做出的不公平的贷款决策，借款人信贷专业知识的缺乏以及模型本身的复杂性和不透明性，使得借款人难以识别和质疑不公平的信贷决策，且无法采取针对性措施改善其信贷评价结果，从而降低了各市场参与者对人工智能辅助金融服务的信任度。

数据治理和监管风险

训练数据的质量和来源还会产生数据治理风险。训练数据质量直接决定了模型的性能和准确性，例如在信用评估方面，数据质量不佳可能导致模型无法准确判断客户的信用状况，增加金融机构的运营风险。训练数据的信息来源可能受到知识产权保护，这些信息未经合法授权或者经过授权但处理过程中造成了信息泄露，都将引起严重的法律风险。例如，韩国三星公司允许员工使用 ChatGPT 聊天机器人后，不到 20 天就发生了 3 起机密信息泄露事件；华尔街多家投行和机构禁止员工使用 ChatGPT 处理工作，防止泄露客户信息和财务数据等内部敏感信息。监管风险主

要体现在缺乏问责制和透明度方面，问责制执行的前提是透明度，而要达到透明度要求则需要对模型相关信息和模型使用的数据进行披露，但由于模型本身的复杂性等特点，较难满足高透明度要求。此外，与模型相关的基础设施和服务可能由第三方提供，更提升了监管难度。

金融稳定和市场垄断的风险

金融从业者决策的趋同性可能会引发价格大幅波动和顺周期，增加市场波动，少数实力强大的模型用户一旦占据市场主导地位也会造成市场集中，并影响金融稳定。此外，由于开发和训练生成式人工智能模型需要强大的计算能力和大量数据，具有资源或先发优势的金融市场主体易在市场中处于垄断地位，造成市场垄断风险。在欧盟委员会 2024 年 1 月发起的关于生成式人工智能竞争水平的磋商中，微软公司向欧盟反垄断监管机构表示“今天，只有一家公司——谷歌——以垂直整合的方式，在从芯片到移动应用商店的每一个人工智能层都拥有实力和独立性，其他所有人都必须依靠合作关系来创新和竞争”。

信誉风险

生成式人工智能基于相对有限的训练数据和算法可能会生成不完整的金融信息或不准确的建议，从而误导金融市场从业者做出偏离市场实际的错误决策，进而损害其信誉。例如，GPT-4 的数据库也只能访问 2023 年 4 月之前的信息，一些最新的法律或趋势都没有添加到生成式人工智能的考虑范围内，而金融监管的法规可能随时会发生变化。从运营的角度来看，这也将降低金融领域生成式人工智能的可用性。另外，决策的趋同性也会对从业者产生信誉风险，当大量从业者使用相同模型的输出结果进行决策时，决策的趋同性可能导致市场出现过度反应或群体行为；一旦市场发生逆转，决策会集体失误，信誉也会随之大打折扣。

金融领域运用生成式人工智能的监管建议

完善监管政策

当前，我国人工智能领域的立法包括《数据安全法》《个人信息保护法》《互联网信息服务深度合成管理规定》《互联网信息服务算法推荐管理规定》等，也出台了《生成式人工智能服务管

理暂行办法》等规范性文件。但还未出台关于金融领域人工智能应用方面的监管性文件，金融行业人工智能应用和部署如何适用以上法律法规，还存在争议和空白，关于金融领域人工智能的立法和监管方面仍有待完善之处。建议监管部门加快制定金融领域人工智能专项法律法规，采取基于风险的监管方法，加强对金融机构使用生成式人工智能技术处理个人数据的监督，确保金融机构实施有效的数据管理和隐私保护措施。

制定行业标准

密切跟进人工智能的演化进展，注重促进人工智能长期发展，统筹安全与发展。借鉴欧盟《人工智能法案》分级监管思路，对金融领域人工智能技术风险进行分级分类，针对不同级别、类型的风险，制定相对应的监管措施，提高对人工智能应用中潜在风险的识别、评估和控制能力，实现事前、事中、事后全面规制。监管部门可以联合行业协会、专业机构制定金融领域人工智能应用的行业标准和自律规范，引入道德指南或准则，确保人工智能在金融决策中遵守公平性、非歧视性、透明性、

可解释性等基本伦理原则，推动行业自律。同时，推动技术开放和数据共享，降低中小型金融机构的进入门槛，促进市场信息的透明度和数据资源的公平分配。在确保安全和合规的前提下，积极探索和利用人工智能技术带来的创新机会。

保护数据安全

金融数据涵盖客户身份信息、交易记录、信用评估等多个方面，在利用生成式人工智能进行数据分析和决策支持时，必须高度重视数据安全与隐私保护。一是加快构建数据安全管理制度。督促金融机构建立严格的数据安全管理制度，在数据收集和使用时，应遵循最小必要原则，并进行匿名化处理，把好数据质量关；采用加密、访问控制等技术手段，确保数据传输、存储和使用安全。二是强化岗位数据安全培训。常态化开展数据安全知识培训和岗位考核，提升员工数据安全意识，增强数据处理过程中的风险防范能力。三是建立数据安全应急处置机制。当发生数据安全事件时及时启动应急响应机制，确定突发事件的风险等级，并对其发展动向和不良影响进行研判、评估和防治。四是

遵守知情同意原则。金融机构应当严格遵守知情同意原则，例如在进行征信数据收集时，应明确告知信息主体数据收集和使用目的，以及所使用的嵌入生成式人工智能的数据收集工具，在获得信息主体授权后，才能采集和使用数据。

强化协同合作

生成式人工智能的健康发展既需要国内政府、企业、公众等多方主体的共同参与，又需要国际间加强协作，共同制定标准规范，有效应对潜在风险。一是加强政府部门间协同监管。政府在生成式人工智能产业的发展过程中发挥着主导作用，在监管实践中要促进金融监管部门间的沟通与协调，促进信息共享与高效协作，实现综合治理。二是增强国际治理的协同性。应鼓励专业力量加强对国际人工智能监管政策的跟踪分析，强化国际协作，共享监管实践，搭建人工智能制度建设、国际规制等全球合作平台，积极推动国际人工智能标准制定，提升我国在人工智能领域的国际话语权。