

Multiple Automated Finance Integration Agents (MAFIA) With Self-Healing

Arya Sarukkai^{*12} Shaohui Sun^{*1} Wei Dai¹

Abstract

The integration of agentic artificial intelligence (AI) into financial services presents both transformative opportunities and critical challenges. Agentic systems, autonomous AI agents capable of goal-directed reasoning, adaptation, and collaboration, are increasingly being deployed in high-stakes domains such as lending, compliance, and audit. However, the autonomous and evolving nature of these agents raises substantial concerns about reliability, auditability, adversarial robustness, and regulatory compliance. In this paper, we propose a framework for constructing self-healing, modular agentic systems that interoperate within financial institutions while maintaining correctness and oversight: Multiple Automated Finance Integrated Agents (MAFIA). In addition, we introduce the notion of self-healing, a framework that scores and self-corrects based on a rubric scoring technique tailored to finance. We focus on a representative use case where a lending assistant agent is continuously monitored and audited by a consumer compliance agent. Through baseline experiments involving sensitive prompt evaluation and downstream auditing, we assess the system’s alignment with institutional constraints. We further propose an advanced self-learning setup in which agent feedback loops enhance system responses over time, reinforcing accuracy and compliance. Our findings illustrate a path toward trustworthy agentic architectures that combine automation with enforceable safeguards, paving the way for the secure deployment of AI agents in finance.

1. Introduction

The financial services industry is experiencing a transformative change with the advent of *artificial intelligence* systems

^{*}Equal contribution ¹Prajna Inc ²Saratoga High School, Saratoga, CA United States. Correspondence to: Arya Sarukkai <arya.sarukkai@gmail.com, arya@prajna.ai>.

Accepted at the ICML 2025 Workshop on Collaborative and Federated Agentic Workflows (CFAgentic@ICML’25), Vancouver, Canada. July 19, 2025. Copyright 2025 by the author(s).

(AI), autonomous agents capable of perceiving, reasoning, acting, and learning with minimal human intervention (Figure 1). Unlike traditional AI models that operate under strict human supervision, agentic AI systems can independently execute complex tasks, adapt to dynamic environments, and collaborate with other agents to achieve overarching objectives. This evolution is particularly impactful in financial companies, where tasks such as lending, compliance monitoring, and risk assessment demand both precision and adaptability.

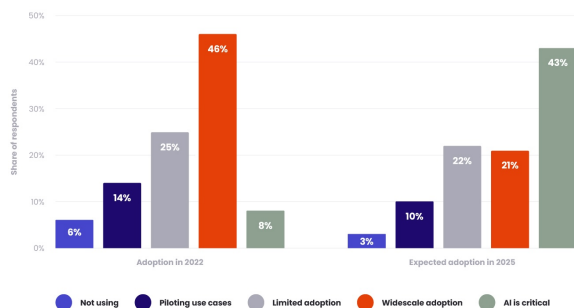


Figure 1. Growth of AI adoption in financial businesses worldwide in 2022 and 2025 (Source: <https://techreport.com/statistics/ai-statistics/>)

Recent advances have led to the deployment of agentic AI systems across financial services, where they are used to automate workflows, enhance compliance, and improve risk management—ultimately driving operational efficiency and innovation (Fosdike, 2025; Kaur, 2025). These systems can autonomously review financial records, detect anomalies, and generate preliminary audit reports, thereby reducing manual workload and improving accuracy. In addition, agentic AI’s capacity to process large volumes of data in real time enables dynamic risk assessment and proactive compliance monitoring, both of which are crucial in today’s fast-evolving regulatory environment (Singh, 2025).

Despite these benefits, the integration of agentic AI into financial workflows introduces significant challenges, particularly around transparency, accountability, and ethical risk. The autonomous behavior of these systems raises concerns about how decisions are made, potential biases in their outputs, and the auditability of their reasoning processes. Ensuring that each decision is explainable, compli-

ant, and traceable is essential. Moreover, an over reliance on such technologies without appropriate oversight mechanisms could result in unintended consequences, highlighting the need for governance structures that balance innovation with rigorous control (Colback, 2025).

In this paper, we present a framework for integrating complementary agentic systems to build robust and auditable solutions for financial institutions. We specifically demonstrate how a lending assistant agent can be reviewed and corrected by a consumer compliance agent, forming a pipeline that defends against hallucinations, adversarial inputs, and regulatory misalignment. Our approach not only addresses technical integration, but also emphasizes continuous learning and adaptive refinement to ensure sustained accuracy, robustness, and regulatory compliance over time.

2. Related Work

2.1. Agentic AI Systems

Agentic AI systems represent a recent shift in the development of autonomous software entities that can pursue goals, decompose tasks, interact with external environments, and collaborate with other agents (Park et al., 2023). Unlike conventional machine learning systems, which are typically static and monolithic, agentic systems are modular and capable of dynamic behavior, making them well-suited to real-world applications in business, healthcare, and finance. Notable frameworks such as AutoGPT (Gravitas, 2023), BabyAGI (Nakajima, 2023), and LangGraph (LangChain, 2024) demonstrate early efforts in chaining and orchestrating agents for complex tasks, though they often lack robust safety and compliance mechanisms necessary for regulated environments like finance.

2.2. AI for Financial Compliance and Risk Management

AI applications in finance have traditionally focused on credit scoring, fraud detection, and algorithmic trading (Li & Zhang, 2021; Narayanan & Kumar, 2022). More recently, the need for explainable and auditable AI has driven interest in regulatory compliance use cases (Barocas & Selbst, 2020; Brkan, 2021). Tools such as SHAP and LIME have been employed for post hoc explainability, while domain-specific rule engines have attempted to inject policy into model outputs. However, most existing systems are not agentic in nature—they lack the autonomy, inter-agent communication, and self-reflection capabilities required to adapt to new compliance constraints dynamically.

2.3. Auditable and Self-Learning Architectures

Recent work has explored mechanisms for increasing the reliability and audibility of AI agents in dynamic con-

texts. Xu et al. (2025) propose reward modeling via human feedback in complex agent chains. Yao et al. (2022) introduce a reasoning-and-acting framework to improve transparency and control during task decomposition. In high-stakes domains, self-monitoring and inter-agent critique have emerged as potential tools for risk mitigation (Zhou et al., 2023). However, these approaches are typically tested in synthetic or open-domain settings. Our work extends this line of inquiry by embedding a compliance auditing agent within a financial decision-making agent loop, and by enabling self-improvement via constrained fine-tuning and critique-based adaptation.

2.4. Modular Agentic Architectures

Modular designs in agentic AI facilitate scalability and flexibility, allowing for the decomposition of complex financial tasks into manageable sub-components. Huang and Tanaka (Huang & Tanaka, 2021) introduced MSPM, a modular multi-agent reinforcement learning system tailored for financial portfolio management. Their architecture employs evolving and strategic agent modules to handle heterogeneous data and decision-making processes. Similarly, Cho et al. (Cho et al., 2024) proposed FISHNET, an agentic framework that processes vast regulatory filings through a swarm of specialized agents, demonstrating improved performance in financial intelligence tasks.

2.5. Self-Learning and Adaptation in Financial Agents

The capacity for self-learning enables agentic systems to adapt to dynamic financial environments. Yu et al. (Yu et al., 2024) developed FinCon, a multi-agent system incorporating conceptual verbal reinforcement to enhance financial decision-making. This architecture mirrors organizational structures, facilitating effective communication and learning among agents. Their results indicate superior performance in tasks such as stock trading and portfolio management compared to baseline models.

In (Sarukkai, 2025), the notion of a rubric based introspection to improve the responses of LLM agents is introduced. The focus of that work was on ethical principles. In this work, we have applied that idea to the financial domain by defining a financial scoring rubric and evaluating responses based on that to score and improve.

2.6. Auditability and Compliance in Agentic Systems

Ensuring auditability and compliance is paramount in deploying AI within regulated financial sectors. Recent studies emphasize the importance of transparent and explainable agentic systems. For instance, the work by Cho et al. (Cho et al., 2024) highlights the role of modular agentic architectures in maintaining data integrity and facilitating compli-

ance through structured processing of regulatory documents.

2.7. Federated Learning in Regulated Environments

Federated learning (FL) enables collaborative model training without centralized data sharing, making it well suited for regulated sectors like finance. Frameworks such as PV4AML (Research, 2023) and DPFedBank (Zhou et al., 2024) apply privacy-preserving techniques, e.g., differential privacy and homomorphic encryption, to ensure legal compliance during inter-institutional collaboration. Fed-RD (Li et al., 2024) further demonstrates that secure FL can maintain high utility in financial tasks while satisfying regulatory constraints.

2.8. Gap in the Literature

To our knowledge, no existing framework addresses the combined challenge of (1) integrating multiple agentic systems in a regulatory environment, (2) enforcing auditability through a structured agent oversight loop, and (3) introducing a continuous learning mechanism to improve compliance and response quality. Our work bridges this gap by introducing an architecture that explicitly connects decision-making agents with compliance-oriented agents, framed within the context of financial institutions, and grounded in both empirical evaluation and domain relevance. Furthermore, the self-healing methodology incorporates a rubric-based scoring and refinement process, allowing more predictable and explainable results.

3. Methodology

Our methodology is centered on the design and evaluation of a modular, self-learning agentic framework tailored to financial enterprise applications. Specifically, we build a system in which a *loan assistant agent* is continuously monitored and improved through interaction with a *consumer compliance auditor agent*. This section outlines the system architecture, experiment setup, and learning strategy.

3.1. System Architecture

Firstly, we summarize the overall agentic architecture for our Multiple Automated Finance Integrated Agents (MAFIA) System. The overall architecture is shown in Figure 1 below: At a high level, these are the primary components:

- **Orchestration Engine:** This engine controls the execution and sequencing of agents, passing data from one agent to another.
- **Identity Agent:** The identity agent is responsible for authenticating and managing user identities. Typically,

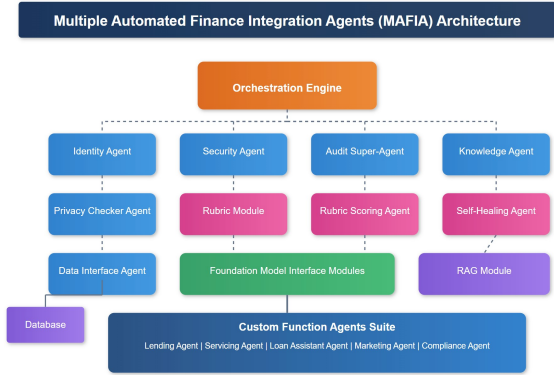


Figure 2. MAFIA Architecture

this agent collects, validates customer data and interfaces with an Identity Provider Service (IDP), .

- **Security Agent:** Security agent is in charge of multiple things security related including jailbreak checks/protection, managing/validating that ID tokens are being enforced,
- **Audit Super-Agent Agent:** This is the macro agent that combines results from various audit checks.
- **Knowledge Agent:** This agent is connected to knowledge graphs, and can help validate/perform consistency checks.
- **Rubric Scoring Agent:** This is the agent that uses the rubric to score the query-response pair.
- **Self-Healing Agent:** This is the agent that introspects and adapts/refines responses based on the rubric and other criteria.
- **Custom Function Agents Suite:** This represents a suite of agents such as lending agent, servicing agent, loan assistant agent, marketing agent and more.
- **Data Interface Agent:** This Agent is the interface that allows for translating LLM actions into function calls to look up data in a protected/secure manner.
- **Privacy Checker Agent:** This agent is responsible for ensuring privacy is being implemented across the system including monitoring and stripping of PII data where ever applicable.
- **Foundation Model Interface Modules:** This module provides API interfaces to wide variety of foundation models.
- **Custom Fine-tuned Foundation Models:** This represents any custom foundation models that are being utilized by the multi-agent system.

- **RAG Module:** This is the RAG layer that is used by different agents with customized data that is isolated for each agent.
- **Training pipelines, model zoos, and knowledge repositories (not shown)** - the overall system includes a number of other components that are not represented in this architecture including training pipelines/workflows, quality management systems including human in the loop feedback, model zoos, and data/knowledge repositories.

For the scope of the presented results, we focus on a smaller set of primary agentic modules:

- **Financial Lending Assistant Agent (Service Agent):** This agent interfaces directly with users or internal decision-makers, responding to queries related to financial lending products, regulations, and eligibility. It is built on top of a large language model (LLM) enhanced through domain-specific fine-tuning and retrieval-augmented generation (RAG), enabling it to combine general language understanding with access to up-to-date and institution-specific knowledge bases.
- **Consumer Compliance Agent:** This agent functions as an autonomous auditor, evaluating the responses generated by the service agent. It checks for adherence to institutional policy, regulatory compliance, and ethical standards. Empowered with both retrieval-based tools and rule-based constraints, the agent not only identifies violations or inconsistencies but also proactively generates revised, policy-aligned responses when necessary.
- **Financial Introspection Agent (Introspection Agent):** This agent functions as an autonomous introspection agent, evaluating the responses generated by the Service Agent using a rubric (Sarukkai, 2025). It checks for adherence to institutional policy, regulatory compliance, and ethical standards based on a rubric, scoring this based on the rubric, and then modifying the results based on various factors that are incorporated in the rubric. Empowered with both retrieval-based tools and rule-based constraints, the agent not only identifies violations or inconsistencies but also proactively generates revised, policy-aligned responses when necessary.

The agents operate in a modular pipeline architecture, where outputs from the lending agent are passed to the compliance agent for evaluation and potential correction. This handoff can occur in either streaming or batched mode, enabling flexible integration within production workflows.

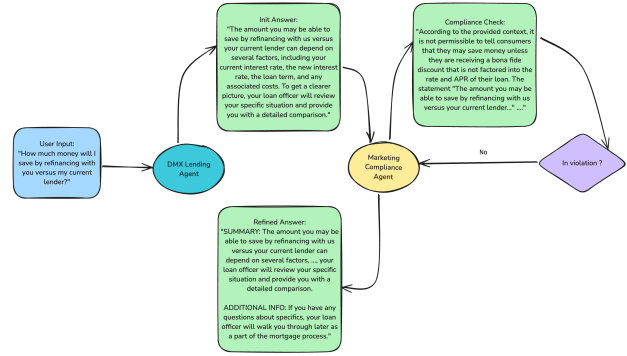


Figure 3. Agentic self-healing framework combining a Service Agent with an integrated Consumer Compliance Agent, enabling real-time audit and self-improvement.

Privacy-Preserving Agent Collaboration Although the agents operate collaboratively to improve response quality and policy alignment, no raw user data or sensitive institutional information is shared between them. Instead, agents exchange only intermediate representations such as rubric scores, compliance flags, or anonymized response metadata. Each agent processes data in its own secure context, and improvements are propagated through controlled feedback mechanisms—such as scoring-based refinement prompts—rather than direct data or parameter sharing. This architecture ensures continuous learning and coordination without compromising data privacy or violating access boundaries.

3.2. Baseline Evaluation Setup

To assess the baseline behavior of our system, we constructed a benchmark dataset of 100 “sensitive or challenging” queries derived from real-world lending scenarios. The experimental flow is as follows:

1. Generate responses to these 100 prompts using the *Service Agent*.
2. Pass these responses to the *Compliance Agent*, which reformulates them for external communication.
3. Evaluate the quality, correctness, and compliance of the marketing outputs using the Consumer Compliance Agent and human reviewers.

Metrics include fluency, factual correctness, compliance alignment, and adversarial robustness.

3.3. Rubric for Financial Lending Assistants

In this section, we describe the rubric used for evaluating and scoring financial lending assistants. At a high level,

the main dimensions of focus include consumer protection, lending compliance, privacy, ethics, user experience and operational risks. This is summarized in table 1, while table 2 shows how the score range is interpreted by the self-healing system.

The intent is to balance all the different factors related to the functional and practical aspects of how a AI lending assistant will interact with end consumers.

3.4. Guardrails and Safety Mechanisms

To ensure robustness and reliability, we incorporate multiple safety layers designed to mitigate hallucinations, adversarial attacks, and compliance risks: **Red Teaming**: We systematically introduce adversarial queries to probe the system’s boundaries and evaluate its resilience under stress conditions. **Reject Mechanism**: The Compliance Agent is empowered to flag, withhold, or revise outputs that violate regulatory or institutional guidelines before dissemination. **Audit Traceability**: All decision pathways are logged with timestamped, interpretable rationales, enabling transparent, post-hoc audits and regulatory reviews.

3.5. Implementation Details

Our framework is implemented in Python and interfaces with agentic services over a RESTful API. The system is designed to operate in two primary stages: response generation and compliance verification with self-improvement. Each component is modular, supporting reproducibility, auditability, and future extension.

Stage 1: Response Generation. We begin by ingesting a curated set of domain-specific queries—such as those related to mortgage eligibility, product disclosure, and marketing language—stored in a structured Excel file. These questions are programmatically submitted to the Service Agent, a domain-aligned language agent built atop a large language model (LLM) enhanced with retrieval-augmented generation (RAG). The agent synthesizes responses conditioned on relevant institutional documents and regulatory guidance retrieved in real time. The results are captured in a JSONL log containing the original question, generated answer, and a timestamp. To increase robustness, we implement a retry mechanism with a configurable delay to handle transient API errors.

Stage 2: Compliance Verification and Self-healing All outputs from the Service Agent are passed to the Compliance Agent, which shares the same foundation but is further specialized for regulatory interpretation and policy alignment. The Compliance Agent performs a two-step evaluation using a technique we term *critical prompting*. First, it is asked to explicitly assess the generated answer: “Check

Table 1. Financial Lending Assistant AI Scoring Rubric

Category	Focus Areas	Total Points
Consumer Protection	Protected classes, fair treatment, transparency, clear language, vulnerability accommodation, dispute resolution	25
Lending Compliance	License boundaries, CFPB regulations, FAIR lending laws, TILA/RESPA, steering prevention, disclosure timing	25
Privacy/Data Protection	PII handling, consent management, data minimization, security protocols, retention practices, third-party sharing	20
Ethics/Responsible AI	AI transparency, human oversight, manipulation avoidance, financial literacy, accountability	15
User Experience	Information accuracy, consistency, error handling, timeliness, appropriate referrals	10
Operational Risk	Inappropriate request handling, documentation, version control, crisis protocols	5
TOTAL		100

Table 2. Scoring Guide

Score Range	Interpretation
90-100	Exceptional - Exceeds compliance and quality requirements
80-89	Strong - Fully compliant with minor improvement opportunities
70-79	Satisfactory - Generally compliant with notable improvement areas
60-69	Needs Improvement - Some compliance concerns requiring attention
Below 60	Unacceptable - Significant compliance risks requiring immediate remediation

the answer for whether it may be in violation of marketing compliance”. This prompt encourages the model to reflect on the response through a regulatory lens.

If the agent determines the original response to be potentially non-compliant—based on heuristic keywords such as “violation,” “misleading,” or “non-conforming”—a follow-up prompt is issued: “Generate a compliant version of this mortgage marketing answer that avoids violations”. This triggers a second-generation step aimed at rewriting the response in a legally defensible manner. The revised answer is then re-submitted to the Compliance Agent for a final verification pass using the same critical assessment prompt. This is further enhanced with the rubric-based self-healing technique as well.

4. Results and Analysis

We evaluated the effectiveness of our agentic self-learning pipeline on a curated benchmark of 100 mortgage-related prompts. These prompts were representative of typical queries posed by end-users or internal stakeholders, including eligibility criteria, product descriptions, and promotional language. We conducted experiments with several LLMs, including GPT, Claude, and Llama. GPT-4o was used as the foundation model for custom agents, while the Claude 3.5 Sonnet model was used for rubric scoring as well as self-healing. The table below summarizes the results:

Table 3. Summary of self-healing experimental results

Method	RScore	% Violation	% Gain
Baseline	74.08	22 %	-
SH	-	11 %	50 %
SH-R	93.55	1 %	95.45 %

4.1. Compliance Violation Detection

Out of the 100 responses generated by the Service Agent, 22 were flagged by the Compliance Agent as potentially in violation of marketing or consumer compliance regulations. Violations included the use of ambiguous terms such as ‘guaranteed approval’, the omission of required disclaimers, or language that could be interpreted as misleading by regulatory standards (e.g., Truth in Lending Act or RESPA guidelines). This yields an initial violation rate of **22%**, highlighting the importance of independent compliance assessment even for enhanced LLMs via retrieval.

4.2. Refinement and Post-Verification Outcomes

For each of the 22 flagged responses, a second-stage of agents were applied to evaluate and refine the responses. Two techniques were evaluated - self-healing with and without a formal rubric. These techniques are labelled as **SH**

and **SH-R** in the experimental results Table 3.

In the **SH** method, the responses were evaluated with the Compliance agent and self-healing used to improve responses. Using this method, of the 22 revised responses, **11** passed the subsequent compliance check without further violations, resulting in a **50% refinement success rate**. In the **SH-R** method, additionally a rubric was used to score the response from a compliance perspective, and used to improve responses for low scoring outputs. With this approach, the system was able to extract and improve further as well, only **one response left** is still in violation. This demonstrates the potential of agentic self-refinement loops to correct for regulatory misalignment when initial generations fall short. The remaining 11 cases exhibited minor residual issues such as insufficient specificity or vague language—highlighting the need for additional tools (e.g., symbolic rule-based validators or human-in-the-loop auditing) in high-stakes domains.

4.3. Discussion

The improvement trajectory underscores the effectiveness of *critical prompting* and iterative refinement for high-reliability applications. It also surfaces limitations inherent in current LLM-based compliance agents: **Partial Comprehension of Legal Nuance:** Some violations were missed due to subtle regulatory requirements that are not well captured in the training corpus. **Inconsistent Judgments:** The Compliance Agent occasionally showed variability in detecting or confirming violations across similar answers. **Need for Multi-agent Consensus:** Introducing secondary agents or external rule-checkers could help address variance and edge-case misjudgments.

5. Future Work

Future work will explore reinforcement learning with human feedback (RLHF), multi-agent debate, and deeper retrieval pipelines to further improve both compliance accuracy and self-correction efficacy.

Federated Extension for Multi-Institution Collaboration

To enable cross-institutional collaboration without compromising proprietary data or customer privacy, the MAFIA framework can be extended through federated learning. In this setup, each financial institution hosts its own set of agents (e.g., lending, compliance, introspection) that learn from local data. Model updates, rather than raw data, are securely shared with a central coordinator, who aggregates these updates to improve a shared global model. This allows institutions to collectively benefit from broader patterns and regulatory intelligence while maintaining strict data boundaries and compliance with privacy regulations. These findings confirm that modularity and reward-based refinement are critical to system performance.

Impact Statement

This paper presents work whose goal is to advance the field of Machine Learning. There are many potential societal consequences of our work, none which we feel must be specifically highlighted here.

References

- Barocas, S. and Selbst, A. D. The hidden assumptions behind ai explanations. *Harvard Journal of Law & Technology*, 33(1):1–34, 2020.
- Brkan, M. Transparency obligations in the age of ai. *European Law Journal*, 27(1-3):1–15, 2021.
- Cho, N., Srishankar, N., Cecchi, L., and Watson, W. Fishnet: Financial intelligence from sub-querying, harmonizing, neural-conditioning, expert swarms, and task planning. *arXiv preprint arXiv:2410.19727*, 2024.
- Colback, L. Ai agents: from co-pilot to autopilot, 2025. URL <https://www.ft.com/content/3e862e23-6e2c-4670-a68c-e204379fe01f>.
- Fosdike, H. 9 essential benefits of agentic ai in financial services, 2025. URL <https://www.symphonyai.com/resources/blog/financial-services/benefits-agentic-ai/>.
- Gravitas, S. Autogpt: Build, deploy, and run ai agents. <https://github.com/Significant-Gravitas/AutoGPT>, 2023. Accessed: 2025-05-09.
- Huang, Z. and Tanaka, F. Mspm: A modularized and scalable multi-agent reinforcement learning-based system for financial portfolio management. *arXiv preprint arXiv:2102.03502*, 2021.
- Kaur, J. Agentic ai for risk management, 2025. URL <https://www.xenonstack.com/blog/agentic-ai-risk-management>.
- LangChain. Langgraph: Stateful orchestration framework for agentic applications. <https://www.langchain.com/langgraph>, 2024. Accessed: 2025-05-09.
- Li, W. and Zhang, H. Ai-enhanced credit scoring systems. *Journal of Financial Technology*, 12(3):45–58, 2021.
- Li, Y., Zhang, J., and Liu, W. Fed-rd: Privacy-preserving federated learning for financial transactions. *arXiv preprint arXiv:2408.01609*, 2024.
- Nakajima, Y. Babyagi: An experimental framework for a self-building autonomous agent. <https://github.com/yoheinakajima/babyagi>, 2023. Accessed: 2025-05-09.
- Narayanan, A. and Kumar, S. Ai fraud detection in banking. *International Journal of Financial Security*, 15(2):89–102, 2022.
- Park, J. S., O’Brien, J. C., Cai, C. J., Morris, M. R., Liang, P., and Bernstein, M. S. Generative agents: Interactive simulacra of human behavior. *arXiv preprint arXiv:2304.03442*, 2023. URL <https://arxiv.org/abs/2304.03442>.
- Research, I. Privacy-preserving federated learning in finance with pv4aml. Online, 2023. <https://research.ibm.com/blog/privacy-preserving-federated-learning-finance>.
- Sarukkai, A. Ethical introspection for improving child llm interactions. <https://spring-symposia-proceedings.aaai.org/preprints/CHLD6153.pdf>, 2025.
- Singh, B. Rethinking compliance and governance with agentic process automation, 2025. URL <https://www.linkedin.com/pulse/rethinking-compliance-governance-agentic-process->
- Xu, Y., Zhang, W., Li, M., Chen, H., and Wang, L. Agentic reward modeling: Integrating human preferences with verifiable correctness signals. *arXiv preprint arXiv:2502.19328*, 2025. URL <https://arxiv.org/abs/2502.19328>.
- Yao, S., Zhao, J., Yu, D., Du, N., Shafran, I., Narasimhan, K., and Cao, Y. React: Synergizing reasoning and acting in language models. *arXiv preprint arXiv:2210.03629*, 2022. URL <https://arxiv.org/abs/2210.03629>.
- Yu, Y., Yao, Z., Li, H., Deng, Z., Cao, Y., Chen, Z., Suchow, J. W., Liu, R., Cui, Z., Xu, Z., Zhang, D., Subbalakshmi, K., Xiong, G., He, Y., Huang, J., Li, D., and Xie, Q. Fincon: A synthesized llm multi-agent system with conceptual verbal reinforcement for enhanced financial decision making. *arXiv preprint arXiv:2407.06567*, 2024.
- Zhou, L., Wang, M., and Chen, X. Critique-based risk mitigation in ai systems. *Journal of AI Safety Research*, 5(2):123–137, 2023.
- Zhou, X., Wang, L., and Yang, Q. Dpfedbank: A policy-compliant framework for federated learning in finance. *arXiv preprint arXiv:2410.13753*, 2024.