

Agentic AI: Empowering Autonomous Intelligence for The Future

Amit Kumar Verma¹, Akash Babu², Anmol Gupta³, Atul Kumar⁴, Bharat Bhushan Shrama⁵

^{1,2,3,4,5}Department of Computer Applications, Invertis University, Bareilly, India

doi.org/10.64643/IJIRTV12I11-207271-459

Abstract—Agentic AI represents a groundbreaking evolution in artificial intelligence, focusing on autonomous systems capable of independently pursuing complex objectives with minimal human oversight. Unlike conventional AI models that rely heavily on structured commands and constant supervision, Agentic AI systems exhibit adaptability, sophisticated decision-making, and operational self-sufficiency within dynamic and changing environments. This paper provides a comprehensive overview of the fundamental principles, defining traits, and key technologies that underpin Agentic AI. We explore its current implementations and future possibilities across diverse sectors such as healthcare, finance, and adaptive software solutions, highlighting the advantages of deploying agentic systems in real-world contexts. Additionally, we address the critical ethical issues surrounding Agentic AI, offering strategies to tackle challenges like goal alignment, resource limitations, and adaptability to unpredictable environments. The paper proposes a structured framework for the safe and effective integration of Agentic AI into society and emphasizes the urgent need for continued research into its ethical implications. Serving as an essential guide, this survey aims to equip researchers, developers, and policymakers with the insights necessary to harness the transformative power of Agentic AI in a responsible and innovative manner.

Index Terms—Agentic AI, Decision-Making, Goal Alignment, Unpredictable Environments, Healthcare, Adaptive Software Solutions, Responsible Innovation

I. INTRODUCTION

Agentic AI, or agentic artificial intelligence, represents a significant evolution in the field of artificial intelligence, characterized by systems that can perform autonomous actions and make decisions without human intervention. As technological advancements accelerate, agentic AI systems are increasingly integrated into various facets of everyday life, ranging from self-driving vehicles to

sophisticated personal assistants capable of natural language processing [1]. These systems are designed not only to analyze vast amounts of data but also to learn from their environments, enabling them to respond adaptively in real-time. The implications of agentic AI extend beyond mere convenience; they herald a transformative shift in multiple industries, including healthcare, where AI-driven diagnostics can enhance patient outcomes, and finance, where algorithmic trading optimizes investment strategies. However, the rise of agentic AI also raises critical ethical concerns, such as accountability for decisions made by autonomous systems and the potential for bias in their algorithms [2],[3]. As these technologies continue to evolve, society must grapple with the challenges and responsibilities that come with them.

II. RISKS OF AI AGENTS

The risks associated with AI agents extend beyond the broader concerns tied to traditional AI systems. While AI systems, in general, can pose threats in terms of biases, inaccuracies, or ethical dilemmas, AI agents introduce unique risks due to their operational autonomy and the potential to minimize or completely exclude human oversight. These agent-specific risks revolve around the capacity of AI agents to operate without human intervention, which could lead to rapid, potentially irreversible, sequences of decisions or actions.

AI agents bring transformative potential across various sectors, but their risks are nuanced and sector-specific. Here's a breakdown:

A. Healthcare:

AI agents can revolutionize diagnostics, treatment planning, and patient monitoring. However, risks include misdiagnoses due to algorithmic errors,

privacy breaches from sensitive patient data, and ethical dilemmas in autonomous decision-making during critical care.

B. Finance:

In financial services, AI agents optimize trading, fraud detection, and customer service. Yet, they pose risks such as market manipulation, biased credit scoring, and vulnerabilities to cyberattacks that could disrupt financial stability.

C. Transportation:

Autonomous vehicles and logistics systems benefit from AI agents' efficiency. However, risks include accidents due to system failures, ethical challenges in decision-making during emergencies, and susceptibility to hacking.

D. Manufacturing:

AI agents enhance production efficiency and predictive maintenance. Risks involve over-reliance on automation, which could lead to operational disruptions if systems fail, and challenges in workforce displacement.

E. Education:

AI agents personalize learning experiences and automate administrative tasks. Risks include biased content delivery, reduced human interaction, and potential misuse of student data.

F. Cybersecurity:

AI agents bolster threat detection and response. However, their autonomous nature can introduce vulnerabilities, such as exploitation by malicious actors or unintended consequences from unsupervised actions.

G. Environment and Energy:

AI agents optimize resource management and renewable energy systems. Risks include over-optimization leading to unintended ecological impacts and challenges in aligning AI-driven decisions with sustainability goals.

Each sector must address these risks through robust governance, ethical frameworks, and fail-safe mechanisms to ensure AI agents align with human values and societal needs. Let me know if you'd like to explore any sector in greater detail!

III. THE CASE FOR VISIBILITY INTO AI AGENTS

Addressing the risks posed by AI agents requires visibility—a comprehensive understanding of where, why, how, and by whom these agents are utilized. Visibility ensures that governance structures can be properly evaluated, revised, and adapted while maintaining accountability for all relevant stakeholders.

Regulatory oversight bodies, currently managing human activities and automated programs like trading algorithms, may require enhanced information access to tackle the unique challenges posed by AI agents. For instance, these agents may employ novel strategies for collusion during economic activities, necessitating updated rules and investigative frameworks [4]. Furthermore, the multi-functional nature of AI agents, such as their simultaneous provision of financial and legal services, introduces potential concerns regarding market consolidation and conflicts of interest.

A. Why Deployment Visibility Matters:

Unlike pre-deployment testing, visibility into deployed AI agents addresses risks that arise from fine-tuning, integration with external tools, and the structuring of system calls to optimize goal achievement. Particularly in high-stakes scenarios, such as applications within banks or cloud compute providers, deployment visibility is vital for assessing and mitigating real-world impacts.

Although the implementation of visibility frameworks can be costly and may raise privacy concerns, they are indispensable for the governance of deployed AI agents. Without such mechanisms, the broader scope of impacts—ranging from economic activities to societal governance—remains inadequately addressed.

IV. CONTRIBUTION

This paper evaluates three categories of measures to improve visibility into AI agents: agent identifiers, real-time monitoring, and activity logs. We propose potential implementations for each measure, balancing varying levels of data collection intrusiveness and informativeness. Additionally, we examine their application across centralized and decentralized deployment contexts, involving diverse actors within

the supply chain, including hardware and software service providers.

We highlight the importance of mitigating any unintended negative impacts prior to widespread implementation.

The measures contribute to AI governance by addressing the unique risks posed by agents. For example:

A. Agent identifiers:

Generalize concepts like watermarks to all outputs of an agent, encompassing text, audio, image, and tool/service usage. These identifiers improve visibility as AI agents increasingly substitute human actions.

B. Real-time monitoring and activity logs

Enhance tracking of complex interactions among multiple agents, their use of external tools or services, and the long-term and diffuse effects of their actions.

Overall, these measures aim to foster accountability, transparency, and informed regulatory practices, advancing the responsible deployment and governance of AI agents.

V. DEFINITIONS

This section defines key terms for clarity. Each term is also defined upon first mention in the main body of the paper.

A. Agency:

The extent to which an AI system operates autonomously to achieve long-term goals with minimal human intervention or guidance. AI systems with high agency, such as reinforcement learning systems, directly interact with the world to accomplish tasks. Systems with low agency, like image classifiers, primarily assist decision-making without acting independently.

B. Agent:

An AI system characterized by a high degree of agency. For this paper, "agent" and "agentic system" are used interchangeably.

While tools like advanced language models with external tool access may qualify as agents, foundational models themselves are excluded from this definition.

C. Scaffolding:

Methods that enhance an AI system's ability to pursue goals by structuring its interactions. These include additional prompts, memory systems, and planning mechanisms. For instance, frameworks like Auto GPT enable AI systems to sequentially generate reasoning, plans, and actions to achieve specific goals, making them more agentic [5].

D. Developer(s):

Individuals or entities involved in creating AI systems. This includes those who train the machine-learning models and those who develop supplementary components such as scaffolding.

E. User:

The individual or group that interacts with an AI system and provides instructions.

F. Deployer:

The entity responsible for operating an AI system and making it accessible to users. Deployers can differ from developers. For example, a company like Microsoft may deploy a system developed by OpenAI. Deployers might also offer ways for users to make systems more agentic through frameworks or tool integrations.

G. Compute Provider:

Organizations that supply and maintain the hardware infrastructure on which AI systems operate. Compute providers may overlap with deployers or developers if they manage their own infrastructure.

H. Tool or Service:

External systems or platforms that interact with AI agents to execute tasks (e.g., booking platforms). These are often accessed via dedicated APIs, designed specifically for agent interactions.

I. Outputs:

The results generated by an AI system, such as text, images, or actions. While deployers generally have access to all outputs, tool or service providers only access data relevant to their services.

J. Inputs:

Data received by an AI agent from users, tools, other agents, or any party, which guide its actions or

responses. Deployers often access these inputs by default but may opt not to store them to safeguard user privacy.

VI. MEASURES TO IMPROVE VISIBILITY

This paper proposes three categories of measures to enhance visibility into AI agents:

A. Agent Identifiers:

These measures determine whether and which AI agents are involved in interactions. Examples include watermarks or IDs attached to outputs, such as requests to service providers. Agent identifiers help actors distinguish agents and facilitate accountability [6].

B. Real-Time Monitoring:

This involves analyzing an agent's activity in real-time, enabling deployers, tool providers, or service providers to detect and address problematic behavior as it happens.

C. Activity Logs:

Deployed by deployers or tool providers, these logs record specific inputs and outputs from an agent, such as interactions with external services or other agents. This information supports post-incident investigations and forensic analysis.

Each measure varies in its level of data collection intrusiveness and informativeness. Comprehensive data collection may be warranted for high-risk agents, especially those operating in sensitive domains like financial trading, where monitoring requirements could match or exceed those applied to human traders. However, extensive data collection carries privacy concerns, necessitating a balanced approach. This section aims to provide a range of options rather than definitive trade-offs or mandates for implementation.

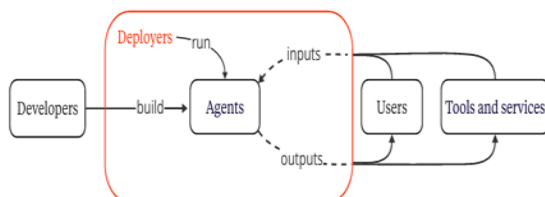


Fig. 1. Information internal to the deployer that it may monitor, modify, or filter.

Key Elements of Fig 1

- **Agents:** The central circle represents AI agents that perform autonomous tasks.
- **Developers:** Developers design, build, train, and maintain AI agents.
- **Deployers:** Deployers are responsible for deploying, configuring, operating, monitoring, and managing AI agents in real-world environments.
- **Users:** Users interact with AI agents by providing inputs and receiving outputs or services.
- **Tools and Services:** AI agents interact with external tools, APIs, databases, and online services to accomplish assigned tasks.

The arrows in Fig 1 illustrate the flow of information and interactions among developers, deployers, users, and external tools and services. The deployer has the capability to monitor, modify, or filter the information exchanged by the AI agent to ensure safe, reliable, and compliant operation.

VII. AGENT IDENTIFIERS

Agent identifiers are mechanisms that indicate whether an AI agent is involved in a particular interaction and, if so, identify the specific agent. These identifiers are typically embedded within selected outputs, made visible to authorized stakeholders, and may include additional metadata about the agent. Agent identifiers enhance transparency, accountability, and trust in AI-enabled systems by providing valuable information to various stakeholders.

A. Benefits of Agent Identifiers

- **Regulators:** Regulatory authorities may require AI agents to disclose their non-human identity, similar to existing bot-disclosure regulations. Such disclosure promotes transparency and helps ensure compliance with legal and ethical standards.
- **General Public:** Users benefit from knowing whether they are interacting with an AI agent rather than a human. This transparency enables users to make informed decisions and establish appropriate levels of trust.
- **Governments and Public Institutions:** Government agencies and public organizations can use aggregated statistics derived from agent identifiers to monitor the prevalence of AI agents in high-risk sectors and assess their societal impact.

- **Service Providers:** Service providers can use agent identifiers to evaluate requests before processing sensitive operations, such as financial transactions, data transfers, or access to protected resources. This enables additional security verification and risk assessment [7].

B. Types of Agent Identifiers

- **Output Coverage:**

Agent identifiers can be applied to different types of outputs, depending on their format and purpose, including text, images, audio, videos, and API requests. For example, digital watermarks may be embedded in images, whereas API requests may include dedicated authentication headers or metadata. Identifiers are particularly valuable for high-impact activities, such as financial transactions, sensitive purchases, or computationally intensive operations.

- **Visibility to Stakeholders:**

The visibility of an agent identifier depends on the participants involved in an interaction. For example, during a financial transaction, identifiers may be accessible to the bank, the recipient, and the service provider. Similarly, e-commerce platforms may use identifiers to verify transaction authenticity, detect fraudulent activities, and protect customer data [8].

- **Identifier Specificity:**

Agent identifiers may either uniquely identify an individual AI agent or simply indicate that an AI agent is participating in the interaction. Unique identifiers are especially useful for incident reporting, auditing, accountability, and forensic investigations. Cryptographic techniques can be employed to ensure the authenticity and integrity of these identifiers. By implementing agent identifiers, as illustrated in Fig. 2, organizations can establish a robust framework for transparency, visibility, and accountability, thereby fostering greater trust in AI-agent interactions.

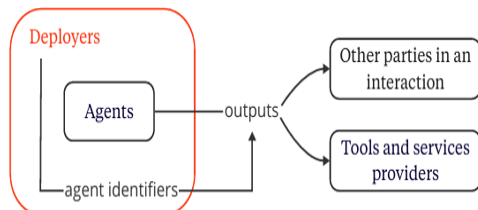


Fig. 2. A robust framework for visibility and accountability

Fig 2 illustrates how an agent identifier informs relevant stakeholders about the participation of an AI agent in a given interaction. Developers (not shown) collaborate with deployers to implement agent identifiers, which are embedded into the outputs generated by the deployed agent. These identifiers are visible to interacting parties, including users, external organizations, and tool or service providers. Once these stakeholders recognize that they are interacting with an AI agent, they may verify important properties such as the agent's authenticity, security, robustness, and compliance before proceeding with the interaction [9].

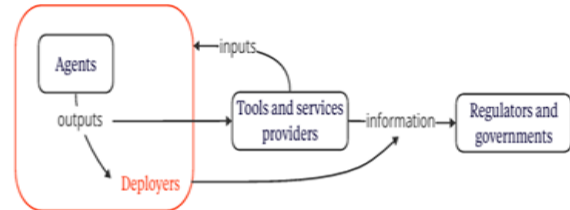


Fig. 3. A framework to interact with an agent

Fig 3 illustrates the interaction framework between an AI agent and its surrounding environment. The deployer has visibility into the agent's inputs and outputs. Inputs originate from users and external tools or service providers (users are not shown in the figure). Certain outputs, such as requests sent to external services or APIs, are also visible to tool and service providers.

These stakeholders can perform real-time monitoring, filter potentially harmful actions, or maintain activity logs for post-incident investigation and forensic analysis. The insights generated through monitoring and logging can further support regulators and government agencies in overseeing AI systems and enforcing appropriate governance measures.

C. Attaching Additional Information to Agent Identifiers

For each deployed AI agent instance, cryptographic mechanisms—similar to those used in software attestation—can be employed to verify the agent's authenticity and provenance. In addition to the identifier itself, supplementary metadata can be attached to provide comprehensive information about the AI system.

This collection of metadata is referred to as an Agent Card, inspired by previous work on documenting AI systems and promoting transparency [10].

An Agent Card may include information about the underlying AI system, the deployed agent instance, and the stakeholders involved throughout its lifecycle.

i. Underlying AI System: Information describing the underlying AI system may include:

- Results of performance evaluations, safety assessments, and records of previous incidents.
- Software libraries, frameworks, models, and other system dependencies used during development.
- Training methodologies, datasets, and model development processes.

Such information enables stakeholders, particularly tool and service providers, to make informed decisions regarding interactions with AI agents. For example, providers may reject requests originating from agents that fail to satisfy predefined security or reliability requirements.

ii. Specific Agent Instance: Information related to an individual deployed agent may include:

- Deployment method (e.g., deployed directly by a user or initiated by another AI agent).
- External tools, APIs, databases, and services accessible to the agent.
- Scaffolding mechanisms, including memory modules, planning frameworks, and reasoning components.
- Operational scope, assigned permissions, objectives, and application domain (e.g., finance, healthcare, or education).

This information enables regulators and other stakeholders to evaluate the intended purpose, operational capabilities, and potential societal impact of the deployed AI agent.

iii. Stakeholders Involved in Development and Deployment: This category identifies the organizations and individuals involved in the AI agent's lifecycle, including:

- End users
- Developers
- Deployers
- Infrastructure and compute providers

Authenticating users interacting with sensitive AI agents—such as those supporting banking, healthcare, or government services—is particularly important. Furthermore, maintaining this information strengthens

accountability by enabling regulators to trace an AI agent's activities back to its unique identifier. In the event of harmful or unauthorized behavior, regulators can request the deployer to investigate the incident, mitigate associated risks, suspend the agent's operation, or permanently deactivate the system.

VIII. REAL-TIME MONITORING

Real-time monitoring involves continuously observing an AI agent's activities to detect, flag, and, when necessary, prevent problematic behavior as it occurs. This process is typically automated to handle the speed, scale, and complexity of AI-agent interactions efficiently. Unlike activity logging, real-time monitoring focuses on immediate detection and intervention rather than long-term data retention. For example, content moderation systems can block harmful messages in real time without permanently storing user inputs, making this approach particularly valuable for privacy-sensitive environments, such as businesses protecting confidential intellectual property.

When potentially harmful or unauthorized behavior is detected, the deployer can intervene immediately and, if appropriate, notify the relevant authorities. Users may also request human review if monitoring systems restrict or suspend the actions of their AI agents.

Real-time monitoring is particularly effective for detecting violations of well-defined policies, including:

- Unauthorized activities or actions that exceed assigned permissions.
- Leakage of confidential or sensitive information.
- Violations of usage limits or abnormal behavioral patterns.
- Suspicious access to restricted resources.

More sophisticated monitoring mechanisms are required to identify complex scenarios involving multiple AI agents, such as:

- Algorithmic collusion, where multiple agents cooperate to manipulate prices or markets.
- Resource overloading, in which numerous agents simultaneously consume excessive computational resources or network bandwidth.
- Coordinated attacks, where groups of agents perform synchronized malicious activities.

Real-time monitoring can also identify AI agents exerting unusually high levels of influence within their operating environment, including:

- Large-scale financial or economic transactions.
- Excessive consumption of computational resources.
- Deployment and coordination of numerous sub-agents.
- Simultaneous control of multiple physical or cyber-physical systems.

A. Activity logs

Activity logs serve as records of an AI agent's inputs, outputs, and state changes, providing critical insights into its actions and impacts. Inputs can originate from users or external tools and services, while outputs may include agent responses and interactions.

These logs can vary in detail and focus on significant actions to balance privacy, behavior tracking, and storage considerations.

B. Benefits of Activity Logs Activity logs are valuable for multiple purposes:

- They enable post-incident investigations by tracing actions to specific user choices or agent decisions.
- Logs assist audits and forensics by identifying the source and causes of harmful outcomes.
- Researchers can use logs to better understand agent behavior and enhance deployment controls.
- Analyzing logs can reveal novel behaviors, improving mechanisms like real-time monitoring [15].

C. Agent-Specific Information Logs can capture diverse categories of data, such as:

D. Tool Usage: Logs indicating which tools or services an agent accessed and the outcomes of these interactions.

E. Internal Processes: Records of memory management, reasoning, or self-critique mechanisms, aiding in understanding agents' decision-making.

Deployer and tool provider logs can complement each other: deployer logs focus on an agent's overall behavior, while tool provider logs highlight the direct impact of specific tools.

F. Handling Persistent and Diffuse Impacts To manage long-term or widespread impacts:

Logs may need to retain data on an agent's runtime, external memory access, and computational resources. Retaining logs beyond an agent's lifespan allows tracking impacts delayed or extended through successive agents.

Cross-referencing logs can illuminate dynamics between sub-agents or groups of agents, revealing malfunctions or undesirable interactions.

G. Detail Levels in Logging The granularity of logs depends on the context:

H. Less Detailed Logs: High-level summaries or occasional samples may suffice for routine applications.

I. Highly Detailed Logs: Comprehensive logs may be mandated in high-risk settings, though they can increase costs, require expertise, and raise privacy concerns [14].

Overall, the design and implementation of activity logs should balance the need for accountability and insight with privacy and resource constraints.

IX. DECENTRALIZED DEPLOYMENTS

Decentralized deployments occur when users, including enterprises or individuals, run AI agents independently—on cloud compute platforms or their own hardware. Users may even combine systems from multiple deployers to form agents, bypassing standard oversight measures. While ensuring visibility in such cases is desirable, smaller-scale systems may not justify deployer-implemented visibility measures. Moreover, malicious actors could exploit decentralized setups to evade regulatory detection. This section explores extending visibility measures to decentralized deployments while addressing associated risks.

A. Compute Provider Oversight

Compute providers play a critical role in enabling oversight for large-scale deployments, which can involve numerous agents with significant societal impact. Such deployments, often resource-intensive, are easier to monitor due to their scale. Compute providers could require users to demonstrate

compliance with visibility measures before granting access to computing resources. Given economies of scale, businesses often rely on cloud infrastructure, further facilitating compute provider-based oversight [11].

B. Tool and Service Providers as Distributed Enforcement Mechanisms

External tools provide leverage to enforce visibility measures. By conditioning access to tools on compliance with visibility protocols—such as agent identifiers—tool and service providers incentivize adherence. For instance, financial institutions could restrict high-risk services to certified agents from trusted deployers.

Direct agent-tool interactions that bypass APIs pose challenges, as they can mimic human behavior. Tools like CAPTCHA systems or identity verification protocols may help detect disguised AI activity. For high-risk domains, requiring human identification could further mitigate risks.

C. Risks and Mitigations

Expanding visibility measures in decentralized contexts raises concerns about privacy and power concentration:

Compute Provider Risks: Surveillance by compute providers could compromise user privacy, revealing sensitive information. Given market concentration, this effectively equates to monitoring much of society. Poor security practices could also expose data to attacks.

Tool and Service Provider Risks: Dependence on certified deployers for tool access may pressure users to adopt agents from centralized entities, even if they are unsuitable or invasive. Excessive reliance on a few deployers could heighten systemic risks [12].

Mitigation Strategies:

I. Voluntary Standards: Experimentation with open-source frameworks, such as agent identifiers, could promote visibility while empowering users to avoid centralized deployers.

II. Prioritization of High-Risk Tools: Regulations could focus on tools used in high-stakes activities (e.g., scientific platforms handling sensitive materials) rather than personal-use agents.

III. Incentivizing Compliance: Regulatory frameworks could account for adherence to visibility standards when determining liability, similar to provisions in the U.S. National Cybersecurity Strategy and HIPAA compliance.

Additional measures, such as decentralized data custody for compute provider logs and enhanced transparency regarding data practices, could also help address risks and foster responsible decentralized deployments.

X. CONCLUSION

Visibility is a cornerstone of governance for increasingly agentic AI systems. This paper assessed three key mechanisms for visibility: agent identifiers, real-time monitoring, and activity logs. Agent identifiers ensure transparency in interactions, facilitating accountability and incident investigation through agent cards containing supplementary information.

Real-time monitoring offers a proactive approach to flag and filter problematic behaviors as they occur. Activity logs provide comprehensive records of inputs and outputs, enabling detailed post-incident analyses and forensics. We explored extending these visibility measures to decentralized deployments by leveraging compute providers and tool/service providers as enforcement mechanisms. Additionally [13], we analyzed the broader implications of these measures, particularly concerning privacy and the concentration of power. However, immediate implementation is not advised; further understanding and mitigation of negative impacts are essential to lay a solid foundation for AI agent governance.

REFERENCES

- [1] Anirban Mukherjee and H. H. Chang, "Agentic AI: Autonomy, Accountability, and the Algorithmic Society," 2025.
- [2] J. Schneider, "Generative to Agentic AI: Survey, Conceptualization, and Challenges," 2025.
- [3] P. Selvam Viswanathan, "Agentic AI: A Comprehensive Framework for Autonomous Decision-Making Systems in Artificial Intelligence," 2025.

- [4] M. Gridach et al., "Agentic AI for Scientific Discovery: A Survey of Progress, Challenges, and Future Directions," 2025.
- [5] La Trobe University, Responsible Artificial Intelligence Hyper-Automation with Generative AI Agents for Sustainable Cities of the Future, 2025.
- [6] C. Randieri, "Agentic AI: A New Paradigm in Autonomous Artificial Intelligence," 2025.
- [7] Business Insider, "Big Four Bet on AI Agents That Can Do All the Work and 'Liberate' Staff," 2025.
- [8] M. Abou Ali, F. Dornaika, and J. Charafeddine, "Agentic AI: A comprehensive survey of architectures, applications, and future directions," *Artificial Intelligence Review*, vol. 59, no. 11, pp. 1–42, 2026, doi: 10.1007/s10462-025-11422-4.
- [9] D. B. Acharya, K. Kuppan, and B. Divya, "Agentic AI: Autonomous intelligence for complex goals—A comprehensive survey," *IEEE Access*, vol. 13, pp. 18912–18936, 2025, doi: 10.1109/ACCESS.2025.108495.
- [10] M. Gridach, J. Nanavati, K. Z. El Abidine, L. Mendes, and C. Mack, "Agentic AI for scientific discovery: A survey of progress, challenges, and future directions," *arXiv preprint arXiv:2503.08979*, 2025.
- [11] S. Liu, W. Wu, X. Xu, T. Li, B. Pang, B. Guo, and Z. Yu, "Adaptive and resource-efficient agentic AI systems for mobile and embedded devices: A survey," *arXiv preprint arXiv:2510.00078*, 2025.
- [12] H. Yao, R. Zhang, J. Huang, J. Zhang, Y. Wang, B. Fang, R. Zhu, Y. Jing, S. Liu, G. Li, and D. Tao, "A survey on agentic multimodal large language models," *arXiv preprint arXiv:2510.10991*, 2025.
- [13] S. Hong et al., "MetaGPT: Meta Programming for Multi-Agent Collaborative Frameworks," *arXiv preprint arXiv:2308.00352*, 2023.
- [14] C. Wang, S. Xie, K. Liu, Z. Lin, and Y. Wang, "CAMEL: Communicative Agents for Mind Exploration of Large Language Model Society," *arXiv preprint arXiv:2303.17760*, 2023.
- [15] J. Park et al., "Generative Agents: Interactive Simulacra of Human Behavior," in *Proc. ACM Symp. User Interface Software and Technology (UIST)*, 2023.