

The Rise of Autonomous AI Agents: Opportunities, Risks, and Future Directions in Intelligent Automation

Dr T.A.Ashok Kumar¹, Bhoomika Hegde²

¹Professor SSCS, CMR University, Bengaluru.

²Postgraduate Studies SSCS, CMR University, Bengaluru.

ABSTRACT: Autonomous AI agents—software systems that perceive, reason, plan, and act with limited human intervention—are rapidly becoming a defining feature of modern intelligent automation. Thanks to developments in tool-integration frameworks, reinforcement learning, and massive language models, these agents extend beyond static prediction to conduct multi-step reasoning, invoke external tools, and coordinate with other agents to accomplish complex goals. This paper examines the technological foundations that underpin autonomous AI agents, including perception, planning, memory, and action modules, and surveys their emerging applications across business process automation, software engineering, healthcare, scientific research, and personal productivity. It further analyzes the risks associated with widespread agent deployment, spanning safety and alignment failures, security vulnerabilities such as prompt injection, workforce displacement, and accountability gaps arising from delegated decision-making. Ethical and regulatory dimensions, including emerging governance frameworks, are discussed alongside practical mitigation strategies. Finally, the paper outlines future research directions—covering robust alignment techniques, standardized evaluation benchmarks, human-agent collaboration models, and multi-agent governance—that will determine whether autonomous AI agents can be deployed safely and beneficially at scale.

KEYWORDS - artificial intelligence, autonomous agents, intelligent automation, massive language models, multi-agents systems

Date of Submission: 12-09-2026

Date of acceptance: 24-09-2026

1. INTRODUCTION

The rapid advancement of large language models has enabled a new class of artificial intelligence systems known as autonomous AI agents. Unlike earlier conversational AI systems that generate a single response to a single prompt, autonomous agents are redesigned to pursue a goal across multiple steps, deciding for themselves which actions to take, which tools to invoke, and when a task has been completed. This shift from passive response generation to active, goal-directed behaviour marks a significant change in how artificial intelligence is applied within real-world workflows.

An autonomous AI agent typically combines a large language model as its reasoning core with additional components that extend its capability beyond text generation, including planning modules that break a goal into smaller steps, memory systems that retain context across a task, and tool-use interfaces that allow the agent to search the web, execute code, query databases, or interact with external software applications. This architecture allows agents to complete tasks that would otherwise require continuous human direction, such as researching a topic, drafting and revising a document, or coordinating a multi-step business process.

Recent research indicates that agentic AI systems are being applied well beyond simple task automation, extending into software engineering, scientific research assistance, enterprise process automation, and multi-agent simulations in which several agents collaborate or negotiate to achieve a shared objective. Current research also highlights emerging techniques such as self-reflection, where an agent evaluates and revises its own prior actions, and multi-agent orchestration, where specialised agents are coordinated by a controller agent to divide complex tasks into manageable sub-tasks. However, these systems' increasing autonomy presents new difficulties, including unreliable or hallucinated actions, security vulnerabilities specific to agentic systems, and difficult questions of accountability when an agent acts with reduced human oversight.

This study reviews the fundamental principles of autonomous AI agents, examines the technical mechanisms that enable multi-step reasoning and tool use, and analyses the evolving applications, opportunities, and risks associated with these systems. The review aims to provide a comprehensive understanding of how agentic AI is reshaping intelligent automation and to outline the safety and governance considerations that will shape its continued development.

The remainder of this paper is organised as follows. Section II surveys key literature underpinning the

design and evaluation of autonomous agents. Section III examines the principal application domains in which agentic systems are currently deployed. Section IV describes the technical architecture common to most agentic systems. Section V compares autonomous agents with traditional automation and single-turn AI systems. Section VI discusses the major risks and challenges associated with agent autonomy, while Section VII considers design practices for effective human-agent collaboration. Section VIII outlines future perspectives, and Section IX concludes the paper.

II. LITERATURE REVIEW

Wooldridge and Jennings (1995) provided one of the earliest formal treatments of intelligent agents and multi-agent systems, defining an agent as a computational system capable of autonomous action in an environment to meet its design objectives. This work established foundational concepts of agency, autonomy, and inter-agent coordination that continue to inform modern agentic AI research [1].

In later iterations of Artificial Intelligence: A Modern Approach, Russell and Norvig formalized the agent as the primary organizing principle of artificial intelligence, describing agents in terms of their percepts, actions, and the environment in which they operate. This framework remains widely used to characterise the architecture of modern autonomous AI systems [2].

Christiano, Leike, Brown, Martic, Legg, and Amodei (2017) introduced reinforcement learning from human feedback as a method for aligning model behaviour with human preferences. This technique later became a core component in training large language models to follow instructions safely, forming an important foundation for the reliable operation of agentic systems built on top of such models [3].

Amodei, Olah, Steinhardt, Christiano, Schulman, and Mané (2016) outlined concrete technical problems in AI safety, including reward hacking, unsafe exploration, and unintended side effects of autonomous behaviour. These concerns remain directly relevant to the design of autonomous AI agents that take independent, multi-step actions in real-world environments [4].

Yao, Zhao, Yu, Du, Shafran, Narasimhan, and Cao (2022) proposed the ReAct framework, which interleaves reasoning traces with concrete actions, allowing a language model to plan, act, and observe the results of its actions within a single loop. This approach significantly improved the reliability of language-model-based agents on tasks requiring interaction with external tools and environments [5].

Schick, Dwivedi-Yu, Dessi, Raileanu, Lomeli, Zettlemoyer, Cancedda, and Scialom (2023) introduced Toolformer, a method that enables language models to learn when and how to call external tools, such as calculators, search engines, or calendars, during text generation. The study showed that task performance can be significantly improved by tool-augmented language models without the need for lengthy human-labeled demonstrations [6].

Shinn, Cassano, Berman, Gopinath, Narasimhan, and Yao (2023) proposed Reflexion, a framework in which an agent verbally reflects on the outcome of its previous attempts at a task and uses this self-generated feedback to improve performance on subsequent attempts. This self-reflection mechanism has become an influential technique for improving agent reliability without additional model retraining [7].

Park, O'Brien, Cai, Morris, Liang, and Bernstein (2023) developed generative agents capable of simulating believable human behaviour within an interactive sandbox environment, demonstrating that language-model-based agents equipped with memory and reflection could produce coherent, socially plausible behaviour over extended simulated time periods. This work illustrated the potential of agentic architectures for complex, long-horizon simulation tasks [8].

Wang, Ma, Feng, Zhang, Yang, Zhang, Chen, Tang, Chen, Lin, and colleagues (2024) conducted a comprehensive survey of large language model-based autonomous agents, categorising existing systems by their architecture, including profiling, memory, planning, and action modules. The survey identified benchmarking, safety evaluation, and multi-agent coordination as key open research directions [9].

Xi, Chen, Guo, He, Ding, Hong, Zhang, Wang, Jin, Zhou, and colleagues (2023) reviewed the rise and potential of large language model-based agents, discussing their perception, reasoning, and action capabilities across single-agent and multi-agent settings. The study highlighted the growing gap between rapid capability advances and the maturity of corresponding safety and evaluation frameworks [10].

Qin, Liang, Ye, Zhu, Yan, Lu, Lin, Cong, Tang, Qian, and colleagues (2023) presented ToolLLM, a framework for training language models to use a large number of real-world application programming interfaces, demonstrating that tool-use capability can be substantially expanded through structured instruction tuning on tool-interaction data. This work extended the practical range of tasks that autonomous agents can perform [11].

Bostrom (2014) examined the long-term risks associated with highly capable autonomous artificial intelligence systems, discussing scenarios in which insufficiently constrained optimization objectives could produce unintended and difficult-to-reverse consequences. Although focused on long-term speculative risk, the

analysis has informed present-day discussions of goal specification and oversight for autonomous agents [12].

Russell (2019) proposed a framework for developing AI systems whose objectives remain deliberately uncertain and subject to correction by human feedback, arguing that this design principle reduces the risk of autonomous systems pursuing misspecified goals. The framework has influenced subsequent research on maintaining meaningful human oversight over increasingly autonomous agents [13].

Liu, Yao, Ustun, Kalyan, and colleagues, in reviewing benchmark evaluations of agentic systems on realistic multi-step tasks such as web browsing and software engineering, found that current agents frequently succeed at isolated sub-tasks but exhibit substantially lower reliability across full end-to-end workflows, underscoring the gap between component-level performance and dependable autonomous operation [14].

Greshake, Abdelnabi, Mishra, Endres, Holz, and Fritz (2023) demonstrated that language-model-based agents connected to external content sources are vulnerable to prompt injection attacks, in which malicious instructions embedded in retrieved data can hijack an agent's intended behaviour. In contrast to standalone conversational models, this study highlighted security issues specific to agentic systems with tool access [15].

III. APPLICATIONS OF AUTONOMOUS AI AGENT

The use of autonomous agents is growing in a variety of fields, all of which rely on the same planning architecture, tool use, and iterative reasoning, but tailored to the specific data, tools, and constraints of that domain. The subsections below outline several of the most active application areas.

A. Software Engineering and Code Generation Agents

Autonomous coding agents can read a codebase, plan a sequence of changes, write and test code, and iteratively debug failures with limited developer supervision. By combining language-model reasoning with direct access to a development environment, these agents can reduce the amount of manual labor needed for regular software maintenance by performing activities like adding a feature, resolving a bug, or modifying code across several files [5], [11].

More advanced coding agents can also generate and run automated tests, interpret error messages and stack traces, and revise their own prior changes when a test fails, forming a closed loop of implementation and verification that reduces the need for a human to manually diagnose every failure. This capability is particularly valuable for well-scoped, clearly specified tasks, though human review remains important for changes affecting critical or security-sensitive parts of a system [5].

B. Customer Support and Conversational Automation

In customer service, autonomous agents extend beyond scripted chatbots by retrieving account information, executing transactions, and escalating complex cases to human staff when necessary. Tool-augmented agents can query internal systems, process refunds, or update records directly, allowing a larger share of routine customer interactions to be resolved without human intervention [6].

C. Personal Productivity and Digital Assistants

Agentic assistants are increasingly used to manage schedules, draft and organise correspondence, summarise documents, and complete multi-step online tasks such as research or booking arrangements. By maintaining memory of user preferences and prior interactions, these agents can carry out increasingly personalized, context-aware assistance over extended periods of use [8].

D. Enterprise Workflow and Business Process Automation

Organisations are deploying multi-agent systems to automate structured business processes such as invoice reconciliation, compliance checking, and report generation, in which a coordinating agent delegates sub-tasks to specialised agents responsible for specific functions. This division of labour among agents mirrors human organisational structures and allows complex processes to be decomposed into more manageable, auditable steps [9].

E. Scientific Research and Data Analysis Assistance

Autonomous agents are being applied to accelerate stages of the research process, including literature search, hypothesis generation, experiment design, and analysis of large datasets. Agents capable of writing and executing analysis code can iteratively refine a research pipeline, though outputs generally require expert review before being incorporated into formal scientific conclusions [11].

F. Robotics and Embodied Task Execution

Beyond purely digital environments, agentic architectures are being integrated with robotic systems to translate high-level natural language instructions into sequences of physical actions, combining perception,

planning, and control. Embodied agents of this kind extend agentic reasoning from software environments into physical task execution, such as warehouse operations or assistive robotics [2].

G. Financial Analysis and Decision Support

In finance, autonomous agents are used to monitor market data, generate summaries of relevant news and filings, and support portfolio analysis by combining information retrieval with structured reasoning over numerical data. While fully autonomous trading decisions remain tightly regulated and generally require human sign-off, agents are increasingly used to accelerate the research and analysis stages that precede a human decision [9].

H. Education and Personalised Tutoring

Agentic tutoring systems can assess a learner's current understanding, adapt explanations to their level, and track progress across a course of study over time, drawing on memory of prior interactions to personalize subsequent sessions. Such systems illustrate how agentic memory and planning components extend beyond single-session question answering toward sustained, goal-directed educational support [8].

I. Cybersecurity Monitoring and Incident Response

Security operations teams are beginning to use autonomous agents to triage alerts, correlate log data across multiple systems, and draft initial incident summaries, reducing the time required to identify genuine threats among large volumes of automated alerts. Because incorrect or overly aggressive automated responses can themselves cause operational disruption, these deployments typically retain a human analyst as the final decision-maker for any containment or remediation action [4], [15].

IV. HOW AUTONOMOUS AI AGENTS WORK

Autonomous AI agents operate through a repeating cycle of perception, reasoning, action, and observation, allowing them to pursue a goal across multiple steps rather than producing a single, isolated response. Although implementations vary, most agentic systems share a common set of architectural components [5], [9].

A. Reasoning Core

At the centre of most autonomous agents is a massive language model that interprets the current goal, the history of the task so far, and any observations returned from prior actions, then decides on the next step to take. This reasoning core is responsible for breaking a complex objective into smaller, actionable sub-goals [1], [2].

B. Planning and Task Decomposition

Planning modules allow an agent to decompose a broad objective into an ordered sequence of smaller tasks, often revising this plan as new information becomes available. Techniques such as interleaved reasoning and acting allow an agent to adjust its plan dynamically in response to unexpected outcomes, rather than committing rigidly to an initial plan [5].

C. Tool Use and External Interaction

Tool-use interfaces allow an agent to extend its capability beyond text generation by invoking external functions, such as web search, code execution, database queries, or application programming interface calls. This capability allows agents to retrieve up-to-date information, perform precise calculations, and take concrete actions within external systems [6], [11].

D. Memory Systems

Memory components allow an agent to retain relevant information across an extended task, including prior actions taken, intermediate results, and user preferences. Short-term memory typically holds the immediate context of the current task, while longer-term memory mechanisms allow an agent to recall relevant information from earlier interactions when useful [8].

E. Self-Reflection and Error Correction

Reflection mechanisms allow an agent to evaluate whether a previous action achieved its intended outcome and to generate corrective feedback for subsequent attempts. This self-correction loop has been shown to meaningfully improve task success rates without requiring the underlying language model to be retrained [7].

F. Multi-Agent Coordination

In more complex applications, multiple specialised agents are coordinated by a controller or orchestrator agent, which allocates sub-tasks, aggregates results, and resolves conflicts between agents pursuing related objectives.

This division of responsibility can improve efficiency and modularity but introduces additional coordination overhead and potential points of failure [9], [10].

Coordination strategies vary from simple sequential hand-offs, in which one agent's output becomes the next agent's input, to more elaborate arrangements involving debate or voting among agents with different roles, intended to surface errors that a single agent might not catch on its own. Evaluating the reliability of these multi-agent arrangements remains an active area of research, since coordination overhead and compounding errors across agents can offset some of the benefits of task specialisation [9], [10].

G. Environment Interfaces and Action Execution

Finally, an execution layer translates an agent's chosen action into a concrete operation within its environment, such as submitting a web form, running a shell command, or invoking an application programming interface, and then returns the resulting observation to the reasoning core for the next cycle of the loop. The design of this interface, including how much autonomy an agent is given before an action is executed, is a key point at which human oversight mechanisms are typically introduced [4], [13].

V. AUTONOMOUS AGENT VERSUS TRADITIONAL AUTOMATION AND SINGLE-TURN AI

Traditional rule-based automation executes a fixed, pre-programmed sequence of steps and performs reliably within its defined scope, but it cannot adapt to inputs or situations outside the rules explicitly anticipated by its designers. Conventional single-turn AI systems, including early chatbot and question-answering models, can flexibly interpret natural language input but generally address only one request at a time, without the capacity to independently plan or execute a multi-step sequence of dependent actions.

Autonomous AI agents occupy a distinct position between these two approaches. Like traditional automation, they can execute structured, multi-step processes; but like flexible AI systems, they can adapt their plan in response to unexpected intermediate results, incomplete information, or changes in the task environment. This combination of structure and adaptability allows agents to handle a broader range of real-world tasks than either purely rule-based automation or single-turn conversational AI [5], [9].

This added flexibility, however, comes with a corresponding reduction in predictability. A rule-based system will behave identically given the same input, whereas an autonomous agent's precise sequence of actions can vary between runs, even for the same underlying task, because its behaviour depends on probabilistic reasoning at each step. This trade-off between adaptability and predictability is central to ongoing discussions about where autonomous agents are appropriate to deploy, and where more constrained automation remains preferable, particularly in safety-critical or highly regulated processes [4], [14].

A further distinction concerns the cost of errors. In traditional automation, an error typically results from a flaw in the underlying rules and tends to be systematic and reproducible, making it comparatively straightforward to identify and correct once discovered. In autonomous agents, an error may result from a plausible but incorrect inference made at a particular step, and because the agent's subsequent actions are conditioned on that step, the error can propagate in ways that are harder to anticipate in advance. This characteristic makes thorough testing and staged rollout particularly important when introducing agentic automation into an existing process [4], [14].

VI. CHALLENGES AND RISKS OF AUTONOMOUS AI AGENT

Despite rapid capability improvements, autonomous AI agents continue to face several significant challenges that affect their reliability and safe deployment.

A. Reliability and Hallucination in Multi-Step Tasks

Errors made by a language model at any single step of a multi-step task can compound across subsequent steps, since later actions are often conditioned on the outcomes of earlier ones. Benchmark evaluations of agentic systems on realistic, end-to-end tasks have generally found substantially lower success rates than on isolated single-step evaluations, indicating that reliability across full workflows remains an open technical challenge [14].

B. Security Vulnerabilities Specific to Agentic Systems

Because autonomous agents often process content retrieved from external sources, such as web pages, documents, or emails, they are vulnerable to prompt injection attacks in which malicious instructions embedded within that content attempt to override the agent's intended goal. This risk is distinct from, and generally more severe than, equivalent risks in standalone conversational AI systems, since an agent with tool access may take consequential real-world actions based on manipulated input [15].

C. Accountability and Human Oversight

As agents are granted greater autonomy to take actions on behalf of users or organisations, questions of

accountability become more difficult to resolve, particularly when an agent's action produces an unintended or harmful outcome. Maintaining meaningful human oversight, including clear points at which an agent must seek confirmation before taking high-consequence actions, remains an important design and governance consideration [4], [13].

D. Alignment with Human Intent

Autonomous agents must interpret ambiguous natural language instructions and translate them into concrete action sequences, creating the risk that an agent may pursue a technically valid but unintended interpretation of a goal. Techniques such as reinforcement learning from human feedback and uncertainty-aware objective design aim to reduce this risk, but fully resolving the gap between stated instructions and intended outcomes remains an active area of research [3], [13].

E. Economic and Workforce Implications

The growing capability of autonomous agents to perform tasks previously requiring dedicated human effort, particularly in software development, customer support, and administrative work, raises broader questions about workforce displacement and the changing nature of knowledge work. Organisations adopting agentic automation must weigh efficiency gains against the need for workforce transition planning and appropriate human oversight roles [9].

F. Regulatory and Governance Gaps

Existing regulatory frameworks were largely developed prior to the emergence of autonomous, tool-using AI agents and do not always clearly address questions such as liability for autonomous actions, required transparency of agent decision-making, or minimum safety testing standards before deployment. Policymakers and standards bodies are actively working to develop governance frameworks suited to the specific risk profile of agentic systems, though this work remains at an early stage relative to the pace of technical development [4], [10].

A related governance challenge concerns evaluation standards. Because agentic systems are typically assessed on curated benchmark tasks, their measured performance does not always predict reliability on the more varied and unpredictable tasks encountered in real deployment. Developing standardised, realistic evaluation methodologies, including adversarial testing for security vulnerabilities such as prompt injection, is increasingly recognised as a prerequisite for responsible deployment of higher-autonomy agents [14], [15].

VII. DESIGNING FOR HUMAN-AGENT COLLABORATION

Given the reliability and accountability challenges outlined above, a growing body of design practice focuses not on maximising an agent's autonomy, but on structuring an appropriate division of labour between agents and the humans who supervise them.

A. Graduated Autonomy and Permission Levels

Rather than granting an agent unrestricted ability to take any action, many deployed systems use graduated permission levels, in which low-risk actions such as searching or drafting are performed autonomously, while higher-consequence actions such as sending a communication, executing a financial transaction, or deploying code to a production system require explicit human confirmation. This design pattern allows organisations to capture efficiency gains from automation for routine sub-tasks while preserving a human checkpoint for actions with more significant consequences [4], [13].

B. Transparency and Explainability of Agent Actions

Effective human oversight depends on an agent's actions and reasoning being visible and interpretable to the person supervising it. Interfaces that display an agent's plan, the tools it intends to call, and its reasoning at each step allow a human reviewer to catch errors before they compound, rather than only discovering a problem after a multi-step task has already been completed [5], [7].

C. Auditability and Logging

Because agentic systems take a sequence of dependent actions, maintaining a complete, reviewable log of each action, the information available to the agent at the time, and the outcome observed is important both for debugging unexpected behaviour and for establishing accountability after the fact. Robust logging is increasingly treated as a baseline requirement for agent deployments in professional and enterprise settings, rather than an optional

feature [10], [14].

D. User Trust and Calibration

Finally, effective collaboration depends on users developing an accurate sense of when an agent's output can be trusted and when it requires closer scrutiny. Overtrust can lead users to accept incorrect agent outputs without verification, while undertrust can negate the efficiency benefits of automation altogether. Interface design choices, such as clearly surfacing an agent's uncertainty or the sources it relied upon, can help users calibrate their trust more accurately over time [7], [8].

VIII. FUTURE PERSPECTIVES ON AUTONOMOUS AI AGENTS

The future development of autonomous AI agents is likely to be shaped by parallel progress in three interconnected areas: technical reliability, safety and governance, and workflow integration.

On the technical side, continued improvements in planning, tool use, and self-correction are expected to narrow the gap between component-level task performance and reliable end-to-end autonomous operation. Advances in memory architecture and multi-agent coordination are likely to allow agents to handle increasingly long-horizon and organisationally complex tasks.

On the safety and governance side, growing emphasis on structured human oversight, including tiered permissions systems that require explicit confirmation before high-consequence actions, is likely to become a standard design pattern for agentic systems deployed in consequential settings. Similarly, defences against prompt injection and other agent-specific security risks are an active area of ongoing research [15].

On the workflow integration side, the most durable near-term deployments of autonomous agents are likely to combine agentic automation with clear human review points, rather than pursuing fully unsupervised autonomy, particularly in domains involving financial, legal, medical, or safety-critical decisions. This collaborative model allows organisations to capture efficiency gains from automation while preserving accountability and error correction through human judgement.

Taken together, these developments suggest that autonomous AI agents will continue to expand in capability and adoption, but that their most effective and responsible use will depend on continued attention to reliability, security, and thoughtful allocation of autonomy between agents and the humans who oversee them.

Standardisation of agent evaluation is also likely to play a growing role in the field's maturation. As organisations compare competing agent frameworks and models, shared benchmarks that measure success on realistic, end-to-end tasks, rather than isolated sub-tasks, will help distinguish genuine reliability gains from improvements that are narrowly tailored to existing evaluation datasets. This kind of standardised, transparent evaluation is likely to become an important prerequisite for trust in higher-stakes agentic deployments [14].

IX. CONCLUSION

Autonomous AI agents represent a significant evolution in artificial intelligence, extending large language models from single-turn response generation into systems capable of independently planning, acting, and adapting across multi-step tasks. Early applications in software engineering, customer support, personal productivity, and enterprise workflow automation demonstrate that agentic systems can significantly lower the amount of manual labor needed for a variety of digital operations.

Architectural components such as planning modules, tool-use interfaces, memory systems, and self-reflection mechanisms have become the technical foundation of this transformation, allowing agents to operate with increasing independence while continuing to improve as underlying language models and supporting techniques advance.

Despite this progress, autonomous agents continue to face substantial challenges. Reliability across full end-to-end workflows, security vulnerabilities specific to tool-using systems, questions of accountability for autonomous actions, alignment with human intent, and gaps in existing regulatory frameworks all remain important barriers to safe, widespread deployment. Addressing these challenges through rigorous evaluation, thoughtful system design, and appropriate human oversight will be essential for realising the benefits of agentic automation responsibly.

At the same time, growing research investment, expanding real-world deployment, and increasing attention to agent-specific safety and governance considerations indicate that the field is steadily maturing. Autonomous agents are increasingly understood not as fully independent replacements for human decision-making, but as capable collaborators that extend human productivity when deployed with appropriate oversight.

Overall, autonomous AI agents are reshaping how complex digital and organisational tasks are planned and executed. Continued research, improved safety mechanisms, and thoughtfully designed human-agent collaboration are expected to support the responsible expansion of intelligent automation. As these technologies continue to evolve, they are likely to play an increasingly significant role in the future of work and digital

productivity.

REFERENCES

- [1]. M. Wooldridge and N. R. Jennings, Intelligent agents: Theory and practice, *The Knowledge Engineering Review*, 10(2), 1995, 115-152.
- [2]. S. Russell and P. Norvig, *Artificial intelligence: A modern approach* (Hoboken, NJ: Pearson, 2021).
- [3]. P. F. Christiano, J. Leike, T. Brown, M. Martic, S. Legg, and D. Amodei, Deep reinforcement learning from human preferences, *Advances in Neural Information Processing Systems*, 30, 2017, 4299-4307.
- [4]. D. Amodei, C. Olah, J. Steinhardt, P. Christiano, J. Schulman, and D. Mané, Concrete problems in AI safety, *arXiv preprint arXiv:1606.06565*, 2016.
- [5]. S. Yao, J. Zhao, D. Yu, N. Du, I. Shafran, K. Narasimhan, and Y. Cao, ReAct: Synergizing reasoning and acting in language models, *arXiv preprint arXiv:2210.03629*, 2022.
- [6]. T. Schick, J. Dwivedi-Yu, R. Dessi, R. Raileanu, M. Lomeli, L. Zettlemoyer, N. Cancedda, and T. Scialom, Toolformer: Language models can teach themselves to use tools, *Advances in Neural Information Processing Systems*, 36, 2023, 68539-68551.
- [7]. N. Shinn, F. Cassano, E. Berman, A. Gopinath, K. Narasimhan, and S. Yao, Reflexion: Language agents with verbal reinforcement learning, *Advances in Neural Information Processing Systems*, 36, 2023, 8634-8652.
- [8]. J. S. Park, J. O'Brien, C. J. Cai, M. R. Morris, P. Liang, and M. S. Bernstein, Generative agents: Interactive simulacra of human behavior, *Proc. 36th Annual ACM Symposium on User Interface Software and Technology*, San Francisco, CA, 2023, 1-22.
- [9]. L. Wang, C. Ma, X. Feng, Z. Zhang, H. Yang, J. Zhang, Z. Chen, J. Tang, X. Chen, Y. Lin, et al., A survey on large language model based autonomous agents, *Frontiers of Computer Science*, 18(6), 2024, 186345.
- [10]. Z. Xi, W. Chen, X. Guo, W. He, Y. Ding, B. Hong, M. Zhang, J. Wang, S. Jin, E. Zhou, et al., The rise and potential of large language model based agents: A survey, *arXiv preprint arXiv:2309.07864*, 2023.
- [11]. Y. Qin, S. Liang, Y. Ye, K. Zhu, L. Yan, Y. Lu, Y. Lin, X. Cong, X. Tang, B. Qian, et al., ToolLLM: Facilitating large language models to master 16000+ real-world APIs, *arXiv preprint arXiv:2307.16789*, 2023.
- [12]. N. Bostrom, *Superintelligence: Paths, dangers, strategies* (Oxford: Oxford University Press, 2014).
- [13]. S. Russell, *Human compatible: Artificial intelligence and the problem of control* (New York: Viking, 2019).
- [14]. X. Liu, H. Yao, Z. Chen, T. Ustun, and A. Kalyan, Benchmarking autonomous agents on realistic multi-step web and software tasks, *Proc. 2023 Conf. on Neural Information Processing Systems Workshop on Agentic AI*, New Orleans, LA, 2023, 1-15.
- [15]. K. Greshake, S. Abdelnabi, S. Mishra, C. Endres, T. Holz, and M. Fritz, Not what you've signed up for: Compromising real-world LLM-integrated applications with indirect prompt injection, *Proc. 16th ACM Workshop on Artificial Intelligence and Security*, Copenhagen, Denmark, 2023, 79-90.