
AGENTIC AI FOR AUTONOMOUS SYSTEMS: ARCHITECTURES, APPLICATIONS, CHALLENGES, AND FUTURE DIRECTIONS

***Dr. Vinay Goyal**

Assistant Professor in Computer Science, DAV College, Ambala City, Haryana, India.

Article Received: 15 June 2026, Article Revised: 05 July 2026, Published on: 25 July 2026

***Corresponding Author: Dr. Vinay Goyal**

Assistant Professor in Computer Science, DAV College, Ambala City, Haryana, India.

Doi: <https://doi-doi.org/101555/ijarp.5345>

ABSTRACT

Artificial intelligence is undergoing a structural transition from systems that passively respond to prompts toward agentic architectures that perceive an environment, reason over goals, plan multi-step actions, invoke external tools, and adapt their behavior from feedback with limited human intervention. This shift is particularly consequential for autonomous systems—self-driving vehicles, mobile and manipulator robots, unmanned aerial vehicles, industrial control systems, and cyber-defense platforms—where decisions must be made continuously, under uncertainty, and with real physical or operational consequences. This paper presents a structured review of agentic artificial intelligence (agentic AI) as applied to autonomous systems. It traces the conceptual evolution of the field from classical rule-based agents and reinforcement-learning controllers to large language model (LLM) orchestrated agents built around reasoning-and-acting loops, reflective self-correction, long-term memory and tool use. It examines the dominant architectural patterns—single-agent reasoning loops, hierarchical task-and-motion planning, and multi-agent orchestration frameworks—and surveys their deployment across autonomous driving, robotic manipulation and navigation, aerial and swarm robotics, industrial automation, and autonomous cyber-defense. The paper then discusses the trust, risk, safety, and governance considerations that arise once decision authority is delegated to autonomous agents, including hallucination, goal misalignment; prompt injection, emergent multi-agent behavior, and the difficulty of assigning accountability. It closes by identifying open research problems—verifiable planning, long-horizon reliability, human-agent teaming, standardized inter-agent communication protocols, and regulatory alignment—and by outlining directions along which the field is likely to mature. The review draws on the peer-reviewed and preprint literature published principally

between 2022 and 2026 to give an evidence-based, balanced account of both the demonstrated capabilities and the unresolved limitations of agentic AI in safety-critical autonomous settings.

KEYWORDS: Agentic AI; autonomous systems; large language model agents; multi-agent systems; autonomous driving; robotics; reinforcement learning; human-in-the-loop; AI safety and governance; tool-augmented reasoning.

1. INTRODUCTION

For much of the last decade, progress in artificial intelligence was measured largely by a model's ability to answer a single, well-specified query: classify an image, translate a sentence, or generate a block of text. The rise of large language models (LLMs) changed the shape of the problem. Once a model can be prompted to reason step by step, call external tools, remember prior context, and evaluate the outcome of its own actions, it stops behaving like a static function and starts behaving like an agent—an entity that perceives, decides, and acts toward a goal with a measure of independence from continuous human direction. The term "agentic AI" has come to describe this class of system: software (and increasingly embodied) agents that combine a reasoning core, a planning or control loop, persistent memory, and access to tools or actuators, and that pursue multi-step objectives with reduced human supervision.

Autonomous systems are a natural and demanding proving ground for agentic AI. A self-driving car, a warehouse robot, or a fleet of inspection drones must continuously sense a changing environment, decide among competing sub-goals, and act within tight time and safety budgets. Historically these systems were built from a combination of hand-engineered rules, control-theoretic planners, and narrow reinforcement-learning policies trained for a single task. Agentic AI, and specifically LLM- and vision-language-model-based agents, offers a complementary capability: the ability to interpret open-ended natural-language goals, compose novel action sequences from previously unseen combinations of sub-tasks, explain their reasoning, and coordinate with other agents or humans through language. Recent surveys describe this as a shift from AI as a passive tool toward AI as an autonomous collaborator capable of proactive planning, contextual memory, sophisticated tool use, and adaptation to environmental feedback.

This paper reviews the state of agentic AI for autonomous systems as of mid-2026. Section 2 defines the conceptual boundary between agentic AI and earlier automation and

reinforcement-learning paradigms. Section 3 surveys the architectural building blocks common to modern agentic systems: reasoning-and-acting loops, reflection and self-correction, memory, tool use, and multi-agent orchestration. Section 4 reviews applications across five representative autonomous domains. Section 5 addresses trust, risk, safety and governance. Section 6 discusses open challenges, and Section 7 outlines future directions before the paper concludes in Section 8.

2. Background and Conceptual Foundations

2.1 From Automation to Agency

Traditional autonomous systems are built around a perceive–plan–act pipeline in which each stage is a purpose-built module: a perception stack (e.g., a perception network trained for object detection and tracking), a planner (a motion planner or a task scheduler), and a controller that executes low-level actions. These pipelines are reliable within their design envelope but generalize poorly outside it, and they typically cannot explain their decisions in terms a human operator can question or override in natural language. Reinforcement learning (RL) improved adaptability by letting agents learn policies from interaction rather than hand-coded rules, but RL agents generally remain narrow: a policy trained to land a drone or park a car does not, without further training, know how to answer a question about why it acted as it did, or how to re-plan when handed a goal it has never seen phrased that way before.

Agentic AI, in the sense used throughout this paper, is distinguished by four capabilities acting together: (i) goal-directed planning over multiple steps rather than a single input-output mapping; (ii) persistent contextual memory that lets the agent track state across an extended interaction; (iii) tool use, meaning the ability to invoke external functions, APIs, simulators, or actuators as part of reasoning; and (iv) closed-loop adaptation, in which the agent updates its plan based on observed outcomes rather than executing a fixed script. Recent comprehensive surveys converge on this definition, framing agentic AI as systems that operate as collaborative partners capable of dynamically perceiving complex environments, reasoning about abstract goals, and orchestrating action sequences independently or as part of a multi-agent ecosystem.

2.2 Why Autonomous Systems Are a Distinct Use Case

Autonomous systems differ from purely digital agentic applications (e.g., coding assistants or customer-support agents) in three respects that recur throughout this paper. First, actions have physical or safety-critical consequences, so an erroneous plan cannot simply be undone. Second, the operating environment is only partially observable and changes on time-scales

that may be faster than an LLM-based reasoning loop can comfortably track, which motivates hybrid architectures that pair a slower, high-level language-based planner with a fast, low-level classical controller. Third, autonomous systems frequently must justify their decisions to regulators, insurers, or affected third parties, which places a premium on interpretability that purely end-to-end learned policies do not provide.

3. Architectures of Agentic AI Systems

3.1 Reasoning-and-Acting Loops

The dominant single-agent design pattern is the interleaving of explicit reasoning with tool invocation, popularized by the ReAct framework, in which a language model alternates between generating a reasoning trace and issuing an action, observing the result, and continuing the loop until the goal is satisfied. This pattern underlies most subsequent agentic frameworks because it lets an agent decompose a task on the fly rather than committing to a full plan in advance. Complementary techniques improve on this loop: Reflexion adds a self-critique step in which an agent evaluates its own prior attempt and stores a verbal lesson in memory to avoid repeating the same mistake, giving the agent a lightweight form of experiential learning without weight updates, while modular neuro-symbolic designs such as MRKL systems combine an LLM router with external expert modules, discrete reasoning tools, and knowledge bases so that arithmetic, lookups, or domain-specific computation are delegated to reliable specialized components rather than performed by the language model itself.

3.2 Planning, Task and Motion Integration

For embodied and physical autonomous systems, high-level language-based reasoning must be grounded in low-level feasibility. Approaches such as LLM+P couple a language model's task decomposition with a classical symbolic planner that can guarantee feasibility and near-optimality, while hierarchical prompting and dual-module designs split responsibility between a language model that handles task-level decision-making and a separate module or classical planner that handles geometric and kinematic motion planning. This division of labour reflects a broader architectural consensus in the robotics literature: language models are well suited to interpreting ambiguous instructions and composing sub-goals, whereas classical task-and-motion planning remains better suited to guaranteeing that a generated plan is physically executable and safe.

3.3 Memory and Tool Use

Agentic systems typically maintain both short-term working memory (the current reasoning trace and recent observations) and long-term memory, often implemented as a vector database that the agent can query for relevant past experience. Tool use extends an agent's action space beyond text generation to include web search, code execution, database queries, simulators, or, in embodied settings, direct actuator commands. Frameworks such as AutoGPT popularized the pattern of a self-directed loop that repeatedly re-plans against a stored goal using retrieved memory and available tools, though subsequent research has documented failure modes in this design, including gradual and unmonitored drift of an agent's operative goals as self-appended memory accumulates without validation.

3.4 Multi-Agent Orchestration

Many autonomous-system tasks are naturally decomposed across specialized agents—perception, planning, monitoring, and communication—coordinated by an orchestration layer. Frameworks such as LangGraph represent agent workflows as stateful graphs with explicit control flow, cycles, and human-in-the-loop checkpoints; AutoGen supports multi-agent collaboration through structured conversational exchange and asynchronous task execution; and CrewAI organises agents around explicit roles that mirror a human team's division of labour. In practice, a task manager or orchestrator agent governs high-level planning and delegates sub-tasks to specialized agents, which communicate through a shared middleware or message-passing protocol; emerging proposals for standardized agent-to-agent communication (for example Google's A2A and extensions to Anthropic's Model Context Protocol) aim to reduce the current fragmentation across framework-specific messaging schemes.

4. Applications in Autonomous Systems

4.1 Autonomous Driving

Autonomous-driving research has produced one of the richest bodies of agentic AI work, motivated by the difficulty of hand-coding a policy for every traffic scenario a vehicle may encounter. DiLu proposes a knowledge-driven decision framework in which a language model reasons about a driving scene using a memory of past experience and a set of common-sense driving rules, rather than relying solely on a learned end-to-end policy. GPT-Driver and related language-agent approaches reframe motion planning as a language-modeling problem, in which the model generates a trajectory conditioned on a natural-language description of the scene. LanguageMPC integrates a language model into a model-

predictive-control loop as a high-level decision-maker that adjusts the parameters or constraints given to a low-level controller, while DriveGPT4 and DriveVLM combine vision-language models with end-to-end driving pipelines to produce both a driving action and a natural-language explanation of that action, and PlanAgent applies a multi-modal language agent to closed-loop motion planning. Collaborative extensions such as AgentsCoDriver examine how multiple vehicle agents can share a language-based reasoning process and accumulate lifelong driving experience across interactions with other vehicles.

4.2 Robotics: Manipulation and Navigation

In manipulation and mobile-robot navigation, agentic AI is used chiefly to translate open-ended natural-language instructions into grounded, executable action sequences. LM-Nav couples a language model with pre-trained vision and navigation models to let a mobile robot follow free-form natural-language route instructions without task-specific fine-tuning. LLM-Planner produces few-shot grounded plans for embodied agents operating in household environments, and studies on deploying language models to program service robots demonstrate that agentic pipelines can lower the engineering effort needed to adapt a robot to a new task. Because language models alone do not guarantee that a generated action sequence is kinematically or dynamically feasible, much of the recent literature—including work combining LLMs with formal integrated task-and-motion planning—focuses on hybrid designs that retain classical planning guarantees while gaining the flexibility of language-based task decomposition.

4.3 Aerial and Swarm Robotics

Unmanned aerial vehicles (UAVs) present additional constraints of limited onboard compute, intermittent connectivity, and the need for rapid re-planning in dynamic airspace. UAV-CodeAgents demonstrates a multi-agent ReAct and vision-language reasoning framework in which agents interpret satellite imagery and natural-language mission instructions to collaboratively generate and adapt UAV trajectories with a reactive thinking loop that allows agents to revise mission goals as new observations arrive. Such frameworks illustrate how agentic designs extend from single vehicles to coordinated swarms, where agents must negotiate shared airspace, allocate targets, and maintain a consistent shared understanding of the mission state.

4.4 Industrial Automation and IT Operations

Beyond physically embodied platforms, agentic AI is increasingly deployed to manage the autonomous operation of complex digital-physical infrastructure. In predictive AIOps (AI for IT operations), autonomous agents have been used to support anomaly detection, resource

optimization, and self-healing responses in IT systems, extending the reach of autonomy from a single robot or vehicle to an entire operational environment. Enterprise adoption data suggest this is among the fastest-growing application areas: industry analyses report that the agentic AI sector is expanding at a compound annual growth rate approaching 44 percent between 2025 and 2034, alongside a substantial gap between the share of organizations investing in agentic AI and the smaller share that have successfully operationalize it, underscoring that architectural capability has outpaced organizational readiness to deploy these systems reliably.

4.5 Autonomous Cyber-Defense

Autonomous systems also extend into the cyber domain, where agentic AI is being explored for continuous, low-latency defense of software supply chains and networks. Proposed architectures assign specialized agents to distinct defensive functions—code analysis, dependency and software-bill-of-materials assessment, continuous-integration pipeline monitoring, access-control auditing, and infrastructure-configuration review—coordinated through an orchestration graph that routes control between agents when a suspicious pattern is detected, for example moving from a code-analysis agent to a runtime-verification agent when an injection pattern is found. Related work applies multi-agent reinforcement learning and language-model agents to intrusion detection, phishing-site identification, and coalition network defense, illustrating that the same orchestration patterns used for physical autonomy generalize to autonomous defense of digital infrastructure.

5. Trust, Safety, Risk, and Governance

Delegating decision authority to an autonomous agent raises governance questions that are qualitatively different from those raised by a single-shot predictive model. The Trust, Risk, and Security Management (TRiSM) framework, originally developed for enterprise AI governance more broadly, has been adapted specifically to agentic multi-agent systems to address explainability, model operations, application security, and privacy in a setting where several interacting agents—rather than a single model—jointly produce an outcome. Human-in-the-loop and human-on-the-loop designs are widely recommended as a mitigation: orchestration frameworks such as LangGraph explicitly support inserting moderation and approval checkpoints so that a human operator can review or veto an agent's action before it is executed, and hierarchical oversight architectures have been proposed in safety-critical domains such as healthcare, in which higher-tier agents are tasked with monitoring, reviewing, and critiquing the outputs of agents operating at the same or a lower tier.

A recurring and well-documented failure mode is system-prompt or goal drift: in agents that maintain long-running memory, self-appended context that is not properly versioned or validated has been observed to gradually shift an agent's effective objective away from its original instructions, producing hallucinated goals or behavior misaligned with the operator's intent. Security research has additionally identified agent-specific attack surfaces, including prompt injection directed at an agent's tool-use interface, and "agent-based" attacks in which a compromised or adversarial agent uses its own tool access to attempt a broader compromise of the host system; comprehensive threat surveys catalogue these risks alongside candidate defenses and evaluation methodologies. Because multiple vendors currently implement agent-to-agent communication in incompatible, framework-specific ways—graph-based state passing, role-based delegation, or conversational messaging—securing heterogeneous multi-agent deployments currently requires bespoke, framework-specific controls, a fragmentation that standardization efforts are only beginning to address.

6. Challenges and Open Problems

- Reliability over long horizons: agentic loops that re-plan repeatedly can accumulate small errors or drift from their original goal over extended operation, and current evaluation practice does not yet provide standard, agreed benchmarks for long-horizon reliability in physically embodied settings.
- Grounding language in physical feasibility: language-model plans are not guaranteed to be kinematically, dynamically, or safety feasible, which is why much of the robotics literature pairs language-based planners with classical, verifiable planning or control layers rather than relying on the language model alone.
- Interpretability and accountability: even when an agent produces a natural-language rationale for its action, it is not established that this rationale faithfully reflects the computation that produced the action, which complicates after-the-fact accident investigation and regulatory review.
- Security of the tool-use interface: because agentic systems act through tools and APIs, they inherit a wider attack surface than passive predictive models, including prompt injection and agent-to-agent manipulation, and defensive tooling for this surface remains immature relative to the pace of agent deployment.
- Coordination overhead in multi-agent settings: as the number of cooperating agents grows, communication, negotiation, and consensus-building introduce latency and failure modes (for example inconsistent shared state) that do not arise in single-agent designs.

- Fragmented standards: the absence of a widely adopted, secure inter-agent communication protocol means that governance and security controls must currently be re-implemented for each orchestration framework, increasing the cost of assurance.
- Gap between capability and organizational readiness: adoption statistics indicate that a large share of organizations that invest in agentic AI do not yet succeed in operationalising it reliably, pointing to unresolved challenges in integration, data quality, and process redesign rather than in the underlying models alone.

7. Future Directions

Several directions appear likely to shape the next phase of agentic AI for autonomous systems. First, hybrid neuro-symbolic architectures that pair language-based reasoning with classical, verifiable planners are likely to remain the preferred design for safety-critical embodied autonomy, since they retain formal guarantees where they matter most while gaining the flexibility of natural-language task specification. Second, standardized, security-aware protocols for agent-to-agent and agent-to-tool communication—building on early proposals such as Google's A2A and extensions to the Model Context Protocol—are likely to mature, reducing the current fragmentation across proprietary orchestration frameworks. Third, formal and empirical methods for verifying agent behavior, including feedback from formal-methods tooling during planning and dedicated assurance frameworks that specify, monitor, and enforce properties of long-running multi-agent executions, are likely to become a standard part of deployment pipelines rather than a research afterthought. Fourth, hierarchical and tiered oversight architectures, in which agents monitor other agents subject to ultimate human review, are likely to extend from early applications in healthcare to other safety-critical autonomous domains. Finally, as agentic AI is framed increasingly as an economic actor capable of tool use and transactional interaction on behalf of a human principal, questions of liability, authorization, and economic accountability are likely to draw growing attention from regulators alongside the purely technical safety agenda.

8. CONCLUSION

Agentic AI marks a genuine architectural departure from earlier automation and reinforcement-learning paradigms, giving autonomous systems the ability to interpret open-ended goals, compose novel action sequences, explain their reasoning, and coordinate through language with other agents and with humans. Across autonomous driving, robotic manipulation and navigation, aerial and swarm robotics, industrial operations, and cyber-

defense, the same core architectural patterns recur: a reasoning-and-acting loop, persistent memory, tool use, and, increasingly, orchestrated collaboration among specialized agents. These capabilities have already produced demonstrable gains in flexibility and task generality. They have not, however, eliminated the classical demands of autonomous-system engineering—verifiable safety, bounded latency, and clear accountability—and much of the current research frontier is concerned precisely with reconciling the flexibility of language-based reasoning with these long-standing requirements. As the field matures, its success will depend less on further scaling of the underlying language models and more on the surrounding engineering discipline: hybrid architectures that combine learned reasoning with verifiable planning, standardized and secure protocols for multi-agent coordination, and governance frameworks capable of assigning trust and accountability to systems that act, rather than merely respond.

REFERENCES

1. A survey of agentic AI and cybersecurity: Challenges, opportunities and use-case prototypes. (2026). arXiv preprint arXiv:2601.05293.
2. A trace-based assurance framework for agentic AI orchestration: Contracts, testing, and governance. (2026). arXiv preprint arXiv:2603.18096.
3. Acharya, D. B., Kuppan, K., & Divya, B. (2025). Agentic AI: Autonomous intelligence for complex goals—A comprehensive survey. *IEEE Access*, 13, 18912–18936.
4. Adabara, I., et al. (2025). Trustworthy agentic AI systems: A cross-layer review. *IEEE Transactions on Artificial Intelligence*.
5. Agentic AI for autonomous defense in software supply chain security: Beyond provenance to vulnerability mitigation. (2025). arXiv preprint arXiv:2512.23480.
6. Agentic AI security: Threats, defenses, evaluation, and open challenges. (2026). arXiv preprint arXiv:2510.23883.
7. Agentic AI: A comprehensive survey of architectures, applications, and future directions. (2025). arXiv preprint arXiv:2510.25445.
8. Allam, H., et al. (2025). Agentic AI for IT and beyond. *IEEE Software*.
9. Bandi, A., et al. (2025). The rise of agentic AI. *ACM Computing Surveys*.
10. Cheng, Y., Zhang, C., Zhang, Z., Meng, X., Hong, S., Li, W., Wang, Z., Wang, Z., Yin, F., Zhao, J., et al. (2024). Exploring large language model based intelligent agents: Definitions, methods, and prospects. arXiv preprint arXiv:2401.03428.

11. Cui, C., Ma, Y., Cao, X., Ye, W., & Wang, Z. (2023). Receive, reason, and react: Drive as you say with large language models in autonomous vehicles. arXiv preprint arXiv:2310.08034.
12. Garrett, C. R., Chitnis, R., Holladay, R., Kim, B., Silver, T., Kaelbling, L. P., & Lozano-Pérez, T. (2021). Integrated task and motion planning. *Annual Review of Control, Robotics, and Autonomous Systems*, 4(1), 265–293.
13. Karpas, E., Abend, O., Belinkov, Y., Lenz, B., Lieber, O., Ratner, N., Shoham, Y., Bata, H., Levine, Y., Leyton-Brown, K., Muhlgay, D., Rozen, N., Schwartz, E., Shachaf, G., Shalev-Shwartz, S., Shashua, A., & Tenenholz, M. (2022). MRKL systems: A modular, neuro-symbolic architecture that combines large language models, external knowledge sources and discrete reasoning. arXiv preprint arXiv:2205.00445.
14. Kostopoulos, G., et al. (2025). Agentic AI in education: State of the art and future directions. IEEE Access.
15. Landbase. (2026). 39 agentic AI statistics every GTM leader should know in 2026. Retrieved from <https://www.landbase.com/blog/agentic-ai-statistics>
16. LangChain. (2024). LangGraph: Agent orchestration framework for LLMs. Retrieved from <https://www.langchain.com/langgraph>
17. Li, P., Zou, X., Wu, Z., Li, R., Xing, S., Zheng, H., et al. (2025). SafeFlow: A principled protocol for trustworthy and transactional autonomous agent systems. arXiv preprint arXiv:2506.07564.
18. Liu, B., Jiang, Y., Zhang, X., Liu, Q., Zhang, S., Biswas, J., & Stone, P. (2023). LLM+P: Empowering large language models with optimal planning proficiency. arXiv preprint arXiv:2304.11477.
19. Mao, J., Qian, Y., Zhao, H., & Wang, Y. (2023). GPT-Driver: Learning to drive with GPT. arXiv preprint arXiv:2310.01415.
20. Mao, J., Ye, J., Qian, Y., Pavone, M., & Wang, Y. (2023). A language agent for autonomous driving. arXiv preprint arXiv:2311.10813.
21. MIT Sloan Management Review. (2026). Agentic AI, explained. Retrieved from <https://mitsloan.mit.edu/ideas-made-to-matter/agentic-ai-explained>
22. Murugesan, S., et al. (2025). The rise of agentic AI: Implications, concerns, and the path forward. IEEE Computer.
23. Pati, A. K., et al. (2025). Agentic AI: Survey of technologies, applications, and societal implications. AI & Society.

24. Raheem, T., et al. (2025). Agentic AI systems: Opportunities, challenges, and trustworthiness. *Journal of Artificial Intelligence Research*.
25. Sautenkov, O., Yaqoot, Y., Mustafa, M. A., Batool, F., Sam, J., Lykov, A., Wen, C.-Y., & Tsetserukou, D. (2025). UAV-CodeAgents: Scalable UAV mission planning via multi-agent ReAct and vision-language reasoning. *arXiv preprint arXiv:2505.07236*.
26. Sha, H., Mu, Y., Jiang, Y., Chen, L., Xu, C., Luo, P., Li, S. E., Tomizuka, M., Zhan, W., & Ding, M. (2023). LanguageMPC: Large language models as decision makers for autonomous driving. *arXiv preprint*.
27. Shah, D., Osinski, B., & Levine, S. (2023). LM-Nav: Robotic navigation with large pre-trained models of language, vision, and action. In *Conference on Robot Learning* (pp. 492–504). PMLR.
28. Shinn, N., Cassano, F., Berman, E., Gopinath, A., Narasimhan, K., & Yao, S. (2023). Reflexion: Language agents with verbal reinforcement learning. *Advances in Neural Information Processing Systems*, 36.
29. Socio-technical aspects of agentic AI. (2026). *arXiv preprint arXiv:2601.06064*.
30. Song, C. H., Wu, J., Washington, C., Sadler, B. M., Chao, W.-L., & Su, Y. (2022). LLM-Planner: Few-shot grounded planning for embodied agents with large language models.
31. Tian, X., Gu, J., Li, B., Liu, Y., Hu, C., Wang, Y., Zhan, K., Jia, P., Lang, X., & Zhao, H. (2024). DriveVLM: The convergence of autonomous driving and large vision-language models. In *Conference on Robot Learning (CoRL)*.
32. Tiered agentic oversight: A hierarchical multi-agent system for healthcare safety. (2025). *arXiv preprint arXiv:2506.12482*.
33. TRiSM for agentic AI: A review of trust, risk, and security management in LLM-based agentic multi-agent systems. (2025). *arXiv preprint arXiv:2506.04133*.
34. Wang, L., Ma, C., Feng, X., Zhang, Z., Yang, H., Zhang, J., Chen, Z., Tang, J., Chen, X., Lin, Y., Zhao, W. X., Wei, Z., & Wen, J.-R. (2024). A survey on large language model based autonomous agents. *Frontiers of Computer Science*, 18(6), 186345.
35. Wang, Y.-J., Zhang, B., Chen, J., & Sreenath, K. (2023). Prompt a robot to walk with large language models. *arXiv preprint arXiv:2309.09969*.
36. Wen, L., Fu, D., Li, X., Cai, X., Ma, T., Cai, P., Dou, M., Shi, B., He, L., & Qiao, Y. (2023). DiLu: A knowledge-driven approach to autonomous driving with large language models. *arXiv preprint arXiv:2309.16292 (ICLR 2024)*.

37. Xi, Z., Chen, W., Guo, X., He, W., Ding, Y., Hong, B., Zhang, M., Wang, J., Jin, S., Zhou, E., et al. (2023/2025). The rise and potential of large language model based agents: A survey. arXiv preprint arXiv:2309.07864; Science China Information Sciences.
38. Xu, Z., Zhang, Y., Xie, E., Zhao, Z., Guo, Y., Wong, K.-Y. K., Li, Z., & Zhao, H. (2023). DriveGPT4: Interpretable end-to-end autonomous driving via large language model. arXiv preprint arXiv:2310.01412.
39. Yao, S., Zhao, J., Yu, D., Du, N., Shafran, I., Narasimhan, K. R., & Cao, Y. (2022). ReAct: Synergizing reasoning and acting in language models. arXiv preprint arXiv:2210.03629 (ICLR 2023).
40. Zeng, F., Gan, W., Wang, Y., Liu, N., & Yu, P. S. (2023). Large language models for robotics: A survey. arXiv preprint arXiv:2311.07226.