

INFORMATION AND WEB TECHNOLOGIES

Design and Application of Autonomous Decision-Making Systems Based on AI Agents

Babayev Mahammad Murad¹

¹ Undergraduate Student in Computer Engineering
Baku Higher Oil School; Azerbaijan

Abstract. The rapid development of large language models and tool-enabled artificial intelligence has shifted attention from systems that only generate responses toward agents that can perceive a task, plan a sequence of actions, use external tools, evaluate intermediate results and adapt their behavior. This transition creates new opportunities for autonomous decision-making, but it also introduces engineering problems related to planning reliability, memory, tool selection, uncertainty, safety and human control. This article examines the design and practical application of autonomous decision-making systems based on AI agents. A design-oriented analytical approach is used to synthesize recent research on single-agent and multi-agent architectures, collaborative reasoning, online planning, human-agent interaction and agent evaluation. On this basis, an Agent-Based Autonomous Decision Architecture (AADA) is proposed. The architecture organizes decision-making into eight stages: perception, goal interpretation, context retrieval, option generation, evaluation, action selection, execution and feedback-based updating. The article also introduces a compact action-priority model for comparing candidate actions according to goal alignment, evidence quality, safety, expected utility and operational cost. The analysis shows that effective autonomy depends not only on model capability, but on the coordination of planning, memory, verification, tool governance and human oversight.

Keywords: AI agents, autonomous decision-making, agentic AI, multi-agent systems, planning, tool use, large language models, intelligent systems, human-agent collaboration.

Introduction. Artificial intelligence is increasingly moving from passive prediction and content generation toward active systems that can pursue goals in dynamic environments. In an agent-based system, the model is not limited to producing a single response. It may interpret a user request, decompose the task, choose tools, gather information, compare alternatives, perform an action and revise its plan after observing the result. This capacity makes AI agents particularly relevant to autonomous decision-making.

INFORMATION AND WEB TECHNOLOGIES

Recent studies show that the quality of an agent depends on more than the underlying language model. Planning strategy, role definition, memory design, collaboration mechanisms and access to tools can substantially influence performance. Multi-agent approaches have also demonstrated that assigning different reasoning roles or allowing agents to debate can improve complex problem solving and reduce dependence on a single reasoning path [1–3]. At the same time, full autonomy remains difficult because agents may misinterpret goals, select inappropriate tools, accumulate errors over long action sequences or act with unjustified confidence. The engineering problem is therefore not simply to make an AI system more independent. It is to create controlled autonomy: the system should be capable of making useful decisions while preserving traceability, safety and the possibility of human intervention. This article examines the main design principles of such systems and proposes a practical architecture for autonomous decision-making based on AI agents.

2. Research Approach and Conceptual Basis. This study applies a design-oriented review of recent research on autonomous AI agents, multi-agent collaboration, planning, memory, tool use and human-agent interaction. An AI agent is considered a software system that receives goals, evaluates available information, selects actions and adapts its behavior using feedback. In multi-agent environments, specialized agents may cooperate and exchange intermediate results to improve decision quality [5].

3. Architecture of an Autonomous AI Agent. An autonomous AI agent generally combines five core components: perception, planning, memory, tool interaction and evaluation. These components must operate as an integrated system because weaknesses in one element can affect the entire decision process. Reliable agents should update their plans according to new information and allow human intervention when uncertainty or risk becomes significant [6].

4. Autonomous Decision-Making Workflow. Autonomous decision-making can be represented as a continuous cycle:

Perceive → Interpret Goal → Retrieve Context → Generate Options → Evaluate → Select → Act → Update

The agent identifies the task, generates possible actions, evaluates them according to evidence, safety and

INFORMATION AND WEB TECHNOLOGIES

expected utility, and then updates its next decision based on the outcome. This iterative structure is especially important for complex, multi-step tasks where plans may need to be revised or stopped when conditions change [7].

5. Proposed Agent-Based Autonomous Decision Architecture (AADA). Based on the reviewed approaches, this article proposes an Agent-Based Autonomous Decision Architecture (AADA). The model is intended as a general engineering framework rather than a domain-specific algorithm. It combines goal processing, planning, evidence-based evaluation, controlled action and feedback in a single decision cycle.

Within AADA, candidate actions are not ranked only by predicted usefulness. A more dependable system should also consider the quality of supporting evidence, operational safety and the cost of execution. For a candidate action a_i , an illustrative Action Priority Score (APS) can be written as:

$$\text{APS}(a_i) = w_1G_i + w_2E_i + w_3S_i + w_4U_i - w_5C_i$$

where G represents goal alignment, E evidence confidence, S safety, U expected utility and C operational cost or risk. The weights can be adjusted for the application domain. A medical or infrastructure system, for example, would normally assign a higher weight to safety than a low-impact information assistant. The score is not presented as a universal mathematical standard; it is a design tool for making decision criteria explicit and auditable. The architecture also supports multi-agent decision-making. A planner agent may decompose the task, specialist agents may analyze different alternatives, a critic agent may search for weaknesses, and a coordinator may select the final action. Such specialization reflects findings that multiple independent roles and reasoning paths can provide useful diversity in difficult problems. Collective decision rules, however, should be chosen carefully because simple majority voting does not always produce the most reliable result [8].

6. Functional Components and Evaluation Criteria.

The main components of an autonomous decision system can be linked to measurable engineering indicators. Table 1 summarizes a compact evaluation structure.

INFORMATION AND WEB TECHNOLOGIES

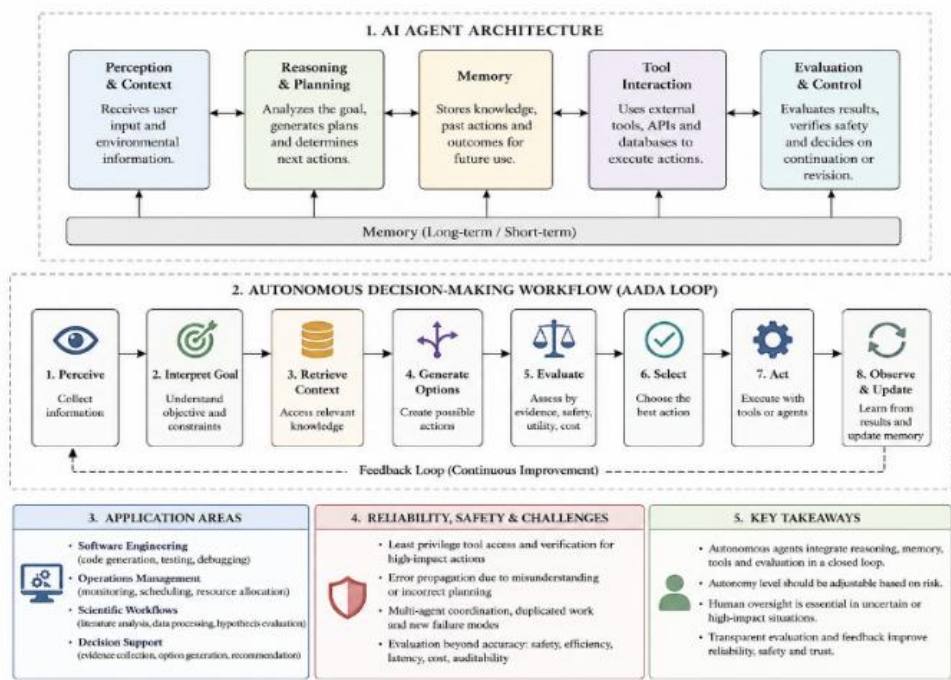


Figure 1. Framework for designing and applying autonomous decision-making systems based on AI agents.

Figure 1
Framework for Designing and Applying Autonomous Decision-Making Systems Based on AI Agents
The figure illustrates the main components and workflow of an AI agent-based autonomous decision system. It shows how perception, reasoning, memory, tool interaction and evaluation are integrated into a continuous decision-making cycle. The framework also highlights key application areas, safety requirements and the role of human oversight in high-risk situations.

Table 1

Core components of an AI-agent decision system

Component	Main function	Typical risk	Evaluation indicator
Goal and context	Interpret task and constraints	Goal misunderstanding	Goal-alignment accuracy
Planning	Decompose task and order actions	Invalid or incomplete plan	Plan success / revision rate
Memory	Preserve relevant state and experience	Outdated or irrelevant recall	Memory-use accuracy
Tool use	Execute external operations	Wrong tool or parameter	Tool-selection accuracy
Evaluation and control	Check evidence, safety and outcome	Overconfident action	Task success, safety, auditability

INFORMATION AND WEB TECHNOLOGIES

7. Application Areas, Reliability and Practical Challenges. Agent-based autonomous decision systems can be applied in software engineering, operations management, scientific research and decision-support environments where tasks require continuous analysis and adaptation [4]. Their effectiveness depends not only on task performance, but also on safe tool use, reliable planning, accurate memory and appropriate human oversight. Because errors may propagate through multiple decision stages, high-impact actions should require verification or human approval. Evaluation should therefore consider task success together with planning accuracy, tool-selection quality, safety, efficiency and auditability.

8. Discussion. Autonomous decision-making should be viewed as a systems-engineering task in which planning, memory, tool use and verification work together. The level of autonomy should vary according to risk, allowing routine tasks to be automated while high-impact decisions remain under human oversight. The proposed AADA framework supports this balance through continuous evaluation, feedback and transparent action selection.

Conclusion. AI agents provide a promising foundation for autonomous decision-making because they combine reasoning, planning, memory, tool use and feedback within an interactive process. However, dependable autonomy cannot be achieved by giving a language model unrestricted access to tools and expecting correct behavior.

The design of an effective autonomous system requires structured goal interpretation, iterative planning, reliable memory, controlled tool access, evidence-based action selection and continuous evaluation. Multi-agent collaboration can improve reasoning by introducing specialized roles and alternative perspectives, but it must be supported by efficient coordination and appropriate decision rules. The proposed Agent-Based Autonomous Decision Architecture offers a compact engineering framework for integrating these elements. Future work should focus on long-horizon planning, active use of memory, adaptive agent coordination, uncertainty-aware decision policies and evaluation methods that jointly measure task success, safety, efficiency and auditability.

INFORMATION AND WEB TECHNOLOGIES

References:

- [1] Feng, X., Chen, Z.-Y., Qin, Y., Lin, Y., Chen, X., Liu, Z., & Wen, J.-R. (2024). Large Language Model-based Human-Agent Collaboration for Complex Task Solving. Findings of the Association for Computational Linguistics: EMNLP 2024, 1336-1357.
- [2] Chen, P., Zhang, S., & Han, B. (2024). CoMM: Collaborative Multi-Agent, Multi-Reasoning-Path Prompting for Complex Problem Solving. Findings of the Association for Computational Linguistics: NAACL 2024, 1720-1738.
- [3] Liang, T., He, Z., Jiao, W., Wang, X., Wang, Y., Wang, R., Yang, Y., Shi, S., & Tu, Z. (2024). Encouraging Divergent Thinking in Large Language Models through Multi-Agent Debate. Proceedings of EMNLP 2024, 17889-17904.
- [4] Zhao, X., Wang, K., & Peng, W. (2024). An Electoral Approach to Diversify LLM-based Multi-Agent Collective Decision-Making. Proceedings of EMNLP 2024, 2712-2727.
- [5] Zhou, W., Mesgar, M., Friedrich, A., & Adel, H. (2025). Efficient Multi-Agent Collaboration with Tool Use for Online Planning in Complex Table Question Answering. Findings of the Association for Computational Linguistics: NAACL 2025, 945-968.
- [6] Tang, J., Shen, S., Wang, Z., Gong, Z., Feng, X., Sun, Z., Tan, H., & Chen, X. (2025). KAPA: A Deliberative Agent Framework with Tree-Structured Knowledge Base for Multi-Domain User Intent Understanding. Findings of the Association for Computational Linguistics: ACL 2025, 6150-6166.
- [7] Zhou, H., Geng, H., Xue, X., Kang, L., Qin, Y., Wang, Z., Yin, Z., & Bai, L. (2025). ReSo: A Reward-driven Self-organizing LLM-based Multi-Agent System for Reasoning Tasks. Proceedings of EMNLP 2025, 15979-15998.
- [8] Wang, Y., Liu, S., Fang, J., & Meng, Z. (2025). EvoAgentX: An Automated Framework for Evolving Agentic Workflows. Proceedings of EMNLP 2025: System Demonstrations, 643-655.