

A Hierarchical Memory-Based Exploration Mechanism for Creative Language Model Agents

Changhong He*, Zhen Bi*, Jungang Lou*[†], Zhenlin Hu*
Zihao Xue*, Qing Shen*, Zhengfang Liu*, Kang Zhao*

*Huzhou University, Huzhou, China

[†]Corresponding Author

Abstract—With the rapid advancement of Large Language Models (LLMs), autonomous agents capable of integrating perception, decision-making, and action have become essential for addressing complex, long-horizon tasks. However, existing agent frameworks still face limitations in creative reasoning, primarily due to an over-reliance on foundational model capabilities, inefficient memory management, and superficial tool invocation that neglects cross-scenario semantic logic. These constraints hinder agents’ generalization and innovative problem-solving abilities in unconventional tasks. To address these challenges, we propose HMemAgent (Hierarchical Memory Agent), a novel framework centered on a three-level hierarchical memory mechanism. This architecture comprises buffer memory for capturing real-time environmental interactions, reflective memory for distilling errors and clues into structured knowledge through introspection, and associative memory for inspiring context-aware creative decisions via semantic links. Together, these levels form a closed-loop “recording–refining–associating” pipeline, significantly enhancing experience reuse efficiency and the exploration of tool semantics. Experimental results demonstrate that HMemAgent markedly strengthens creative reasoning and cross-scenario adaptation in dynamic, complex environments while maintaining robust decision stability.

Index Terms—Agent, Creativity Reasoning, Memory, Tool Use

I. INTRODUCTION

In recent years, the field of artificial intelligence has witnessed significant advancements propelled by large language models (LLMs) [1]. These models not only demonstrate remarkable performance in handling complex tasks but also redefine the role and capabilities of AI agents by integrating perception, decision-making, and action into a unified framework [2]. LLM-powered agents transcend the constraints of static single-turn interactions, achieving dynamic autonomous decision-making through multi-step reasoning, state maintenance, and tool invocation. They can interpret diverse inputs, maintain task context over time, and selectively leverage external tools or APIs to fulfill user goals. Equipped with environmental perception, action planning, and continuous adaptation capabilities, they effectively address complex long-horizon tasks [3]. This evolution not only expands the boundaries of AI capabilities but also reshapes human-AI collaboration paradigms, blurring the traditional distinction between “assistant” and “collaborator” [4].

Research on creativity is critical for agent decision-making, as it demands agents possess innovative reasoning abilities that transcend conventional thought patterns. These abilities

enable flexible transfer of observations to unstructured novel scenarios, proactive identification of latent clues, exploration of tool affordances [5], and generation of unconventional solutions [6]. However, existing approaches exhibit two key limitations: First, they over-rely on the capabilities of foundational LLMs, lack explicit creativity-driven mechanisms, and narrowly focus on memory management, leading to inefficient utilization of historical interaction data and thereby failing to dynamically integrate past experiences into manageable knowledge or clues. Second, methodologies rooted in the autoregressive nature of LLM reasoning overemphasize superficial “tool-task” matching, neglecting inherent task logic and cross-scenario semantic associations [7]. Consequently, this constrains agents’ ability to explore tool functionalities, adapt to open environments, and generate innovative actions, ultimately leading to insufficient generalization on unconventional tasks and hindering effective creative problem-solving.

To address these pressing challenges, we propose HMemAgent, a novel, robust agent framework integrating a three-level hierarchical memory mechanism. The framework comprises three synergistic levels: a buffer memory continuously capturing raw interactions and real-time environmental feedback; a reflective memory periodically distilling encountered errors and latent clues into actionable structured knowledge through deep introspection; and an associative memory linking intrinsic tool attributes with contextual cues to inspire context-aware creative decision-making. These levels form a closed-loop “recording–refining–associating” pipeline, systematically empowering agents to efficiently reuse historical experiences and perform flexible reasoning and innovative task planning based on rich tool semantics, thereby enhancing adaptability and problem-solving capabilities in dynamic environments.

Our contributions are threefold:

- We identify insufficient creativity in agents within virtual scenarios and propose HMemAgent, a memory-centric framework designed to enhance creative reasoning.
- HMemAgent introduces a three-level hierarchical memory mechanism that integrates buffer, reflective, and associative memory to support decision-making and enable creative problem-solving during exploration.
- Experiments show HMemAgent maintains decision stability across baselines while significantly improving creative reasoning and cross-scenario adaptation.

The code we provided is available at this link: <https://github.com/hchfromxh/HMemAgent>.

II. RELATED WORK

A. Creativity in LLM Agents

As AI evolves from command execution to autonomous decision-making, endowing agents with creative cognition is critical for achieving artificial general intelligence [8]. Creativity represents higher-order human-like intelligence and drives agents to generate breakthrough solutions for open-ended tasks. Defined as the intersection of novelty and value [9], it requires balancing unconventional exploration with contextual utility [10]: novelty demands non-routine approaches, while value ensures domain-specific applicability. This framework ensures innovations address real-world challenges, underscoring agent creativity’s significance [11]. Current LLM agent research emphasizes static benchmarks [12], such as writing tasks or domain-specific generation [13]. Methods employ prompt engineering [14], associative thinking [15], or multi-agent debates [16] to diversify outputs. However, reliance on predefined objectives or structured workflows leads to suboptimal efficiency and limited innovation in non-routine scenarios.

B. Agent Memory

Research on AI agent memory has evolved from simple context stacking to sophisticated architectural designs aimed at addressing factual fragmentation and personalization decay in long-horizon interactions [17]. Contemporary memory systems develop along three complementary dimensions: functional partitioning, technical implementation, and architectural organization [18]. At the functional level, systems distinguish between experiential memory and factual memory, ensuring agents can internalize historical experiences while maintaining the accuracy of objective knowledge [19]. Technically, the field has explored parametric memory, retrieval-augmented paradigms, and latent memory achieved through implicit encoding mechanisms. MemGen [20] advances this domain by implementing generative rather than purely extractive memory mechanisms. At the architectural level, recent systems have adopted operating system-inspired designs for memory management [21]. MemoryOS [22] and EverMemOS [23] implement hierarchical structures that facilitate efficient scheduling and updating of long-term memory through self-organizing mechanisms. In knowledge representation, research has progressed from simple vector storage to complex structural networks; A-MEM [24] constructs interconnected knowledge graphs, while Mem0 [25] employs graph-based representations to capture entity relationships. Despite these advances, current memory architectures remain primarily optimized for recall accuracy rather than creative generation [26]. They lack mechanisms for extracting latent patterns across disparate experiences and transforming them into innovative action strategies—particularly in scenarios requiring unconventional tool combinations. Our hierarchical memory architecture bridges

this gap by explicitly designing memory levels that simultaneously support reliable knowledge retention and creative problem-solving capabilities.

III. METHODOLOGY

Building upon preliminary research on agent memory systems, we propose HMemAgent—a unified hierarchical memory architecture designed to overcome two critical limitations in long-horizon agent-environment interactions: inefficient experience distillation and limited creative action generation. A overview of the mechanism is shown in Figure 1. **BaseAgent (§A)**: The left side of Figure 1 displays BaseAgent, responsible for handling environmental interactions and tool execution. **Buffer Memory (§B)**: As the sensory foundation level, this level preserves key interaction data and applies preliminary processing based on environmental feedback. **Reflective Memory (§C)**: This level performs reflective processing on processed sequences to extract potential clues. By constructing structured knowledge repositories, it transforms raw data into actionable clues. **Associative Memory (§D)**: As the top level, it performs cross-referential reasoning between knowledge repositories and available tools to generate innovative action hypotheses. **Hierarchical Synergy (§E)**: All levels form a continuous cognitive pipeline that produces innovative decisions through coordinated hierarchical processing.

A. BaseAgent

The BaseAgent functions as a foundational framework that couples a Large Language Model with a working memory system to navigate the EscapeBench environment. Operating in a continuous loop, it first receives a scenario context from the game engine and utilizes Chain-of-Thought reasoning to determine the next appropriate action based on its current observation and stored historical memory. After executing an action, the agent immediately updates its working memory with the resulting environmental feedback, ensuring that previous attempts and outcomes are preserved to inform future decisions.

B. Buffer Memory

The buffer memory serves as the sensory interface of HMemAgent, maintaining raw interaction sequences with the environment within a fixed-capacity sliding window. This design choice directly supports the “recording” phase of our closed-loop pipeline, preserving critical environmental dynamics before abstraction. At each timestep t , it records scenario context and action outcomes under policy π . The buffer strictly preserves all interactions within its window size w , operating as a first-in-first-out queue. Its bounded capacity ensures computational efficiency while maintaining sufficient context for the reflective layer’s error identification processes. All recorded interactions undergo real-time processing based on environmental feedback for higher-level processing.

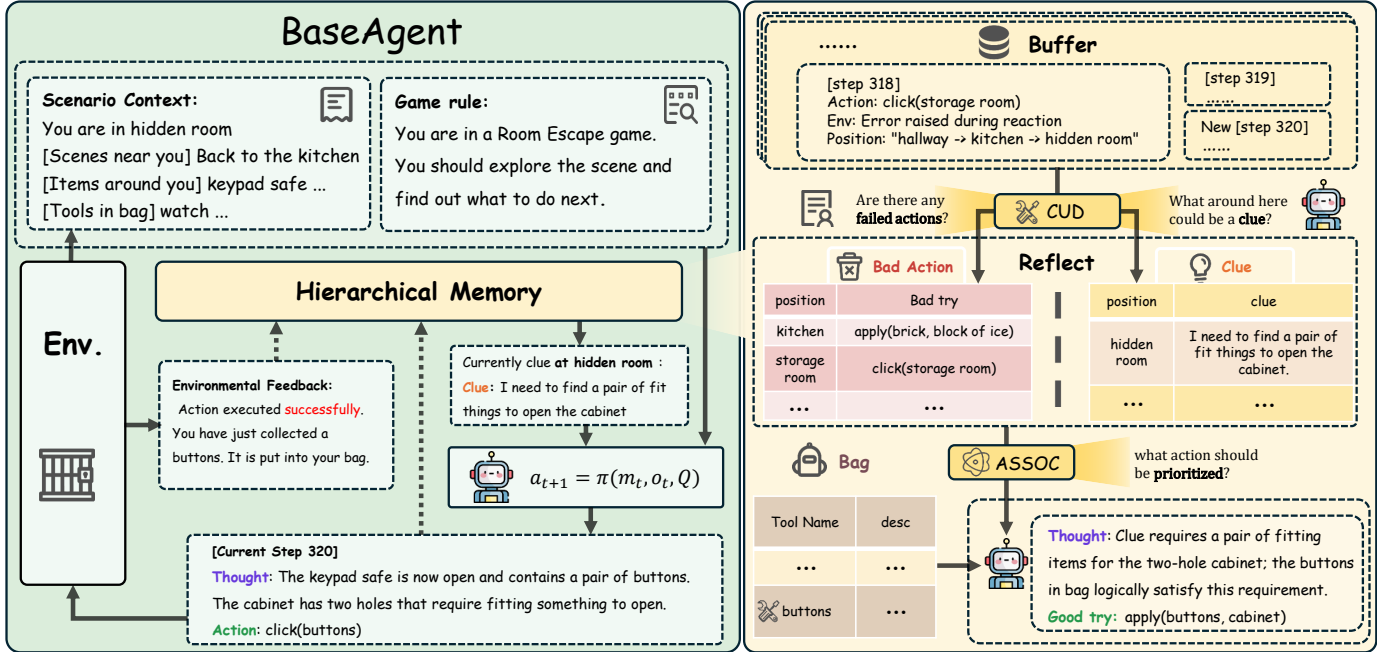


Fig. 1: A visual overview of HMemAgent: (1) BaseAgent: for implementing iterative decision-making within the EscapeBench environment as a foundational LLM-memory coupled framework; (2) Buffer Memory: for recording raw interaction data; (3) Reflective Memory: for extracting latent clues and patterns; and (4) Associative Memory: for associating contextually appropriate behaviors.

C. Reflective Memory

Reflective memory processes tagged sequences from buffer memory, converting raw interaction data into structured knowledge repositories essential for decision-making. It identifies recurring failures and extracts action clues to address inefficiencies in transforming data into actionable knowledge. The system maintains two core tables: the bad action table logs repeated failures with contextual details to prevent recurrence; the clue table stores inferred hints categorized by location or object, supporting contextual recall and reasoning. These repositories enable efficient retrieval of historical experiences during decision-making. Table entries undergo CUD (Create, Update, Delete) management to ensure timeliness and accuracy—refining clues with new information and removing outdated entries. This structured representation bridges the gap between raw observations and the associative layer’s creative reasoning requirements.

D. Associative Memory

The Associative memory serves as HMemAgent’s cognitive apex, leveraging LLMs to perform associative reasoning between reflective memory’s structured knowledge and the current context, generating creative yet plausible action hypotheses. It operationalizes the “associating” phase of our framework, deliberately connecting disparate semantic elements to overcome conventional tool-use limitations. As shown in Figure 1, for latent requirements, it supplies relevant knowledge fragments and available tools to the LLM, which produces innovative actions via semantic association. High-

confidence proposals are prioritized as Level-3 hypotheses and injected into BaseAgent’s decision context.

E. Hierarchical Synergy

The HMemAgent framework defines a decision-making protocol formalized as:

$$a_{t+1} = \pi(m_t, o_t, Q), \quad (1)$$

where a_{t+1} is the action at time step $t + 1$, o_t the current observation, Q the environmental rules and task constraints, and m_t the memory context retrieved from the hierarchical memory system—selected by a strict hierarchy of abstraction and reliability. Specifically, the framework prioritizes outputs from Associative Memory, then clues from reflective memory, and resorts to raw sequences in buffer memory only when higher-level knowledge is unavailable, ensuring robust decision quality across varying environmental complexities. To avoid over-fixation on unproductive reasoning paths, each reflective clue carries a focus counter (initialized to $f = 5$); it is temporarily deactivated upon depletion. A detailed analysis of this parameter appears in the [Hyperparameter Sensitivity Analysis](#) section.

IV. EXPERIMENTS

A. Experimental Settings

a) *Environment:* The experiments were conducted on 36 game settings. An agent is considered to be making progress if it achieves a key step or collects a new tool. To ensure full completion of the game, agents will receive help if they fail to make progress for 50 consecutive actions. The working memory length w was set to 10. The focus f was set to 5.

TABLE I: Benchmark test results of BaseAgent, EscapeAgent, and HMemAgent with different core models. Compared to BaseAgent, HMemAgent shows improvements across almost all performance metrics; compared to EscapeAgent, HMemAgent achieves advantages in key metrics, demonstrating HMemAgent’s superiority in exploring the creativity of intelligent agents.

METHOD	BASE MODEL	↓ HINTS USED	↓ TOTAL STEPS	↑ EARLY EXIT PROGRESS (%)	↓ TOOL HINTS USED (%)	↓ KEY STEPS HINTS USED (%)
BASEAGENT	QWEN-3-4B	26.97	1667.00	5.08%	40.58%	50.75%
ESCAPEAGENT		17.00	1124.97	10.57%	32.25%	28.58%
OURS		16.56	1117.83	14.21%	31.04%	28.08%
BASEAGENT	GEMMA-3-4B	29.75	1756.33	5.01%	36.11%	56.33%
ESCAPEAGENT		22.22	1358.22	12.32%	34.30%	39.42%
OURS		21.61	1337.03	12.54%	30.07%	40.92%
BASEAGENT	LLAMA-3.1-8B	24.89	1665.31	9.28%	11.71%	18.58%
ESCAPEAGENT		18.94	1344.67	18.85%	8.94%	7.25%
OURS		18.14	1311.28	16.88%	8.21%	7.50%
BASEAGENT	QWEN-3-8B	23.31	1490.53	8.95%	33.82%	46.58%
ESCAPEAGENT		13.97	980.83	22.08%	20.29%	27.83%
OURS		13.11	956.47	27.75%	17.39%	27.33%
BASEAGENT	GEMMA-3-12B	19.53	1292.19	9.09%	26.69%	40.17%
ESCAPEAGENT		10.14	834.92	25.73%	12.80%	21.50%
OURS		9.33	795.69	24.03%	11.47%	20.00%
BASEAGENT	QWEN-3-14B	21.39	1351.44	9.80%	30.68%	42.75%
ESCAPEAGENT		10.53	793.44	25.02%	11.71%	23.33%
OURS		9.89	772.89	25.62%	10.75%	22.25%
BASEAGENT	QWEN-3-32B	14.31	1022.39	13.90%	15.10%	32.33%
ESCAPEAGENT		7.14	613.08	33.16%	5.68%	17.33%
OURS		6.97	602.33	32.09%	6.04%	16.58%
BASEAGENT	GLM-4.7-400B	9.97	715.56	22.38%	7.85%	24.33%
ESCAPEAGENT		5.81	515.42	40.90%	3.99%	14.50%
OURS		5.44	498.36	41.36%	3.62%	13.67%
BASEAGENT	DEEPSEEK-R1-671B	6.43	535.91	33.63%	4.23%	15.67%
ESCAPEAGENT		3.72	415.00	57.18%	2.05%	9.58%
OURS		3.83	402.72	53.68%	2.29%	9.75%

b) Models: Our evaluation focuses on a diverse set of open-source models, including GLM-4.7-400B [27], DeepSeek-R1-671B [28], Qwen3 (4B, 8B, 14B, 32B) [29], Gemma3 (4B, 12B) [30], and Llama-3.1-8B [31]. We omit models smaller than 4B parameters, as preliminary tests indicated they lack the necessary capability for these tasks.

c) Metrics: To evaluate model performance on escape tasks, we adopt the five built-in metrics from EscapeBench:

- **Hints Used:** Total hints used in a game.
- **Total Steps:** Total actions taken in a game.
- **Early Exit Progress:** Proportion of key steps and tools collected before needing a hint (game progress before needing a hint for the first time).
- **Tool Hints Used (percentage):** Hints used for tool collection (normalized by total tools).
- **Key Step Hints Used (percentage):** Hints used for key steps (normalized by total key steps).

d) baseline: We conduct a comprehensive comparative analysis of our HMemAgent approach against two established baselines—BaseAgent and EscapeAgent—on the EscapeBench benchmark. To evaluate the individual contributions of each level within HMemAgent, we perform systematic ablation studies across four configurations: (1) only Buffer; (2) Associate + Buffer; (3) Reflect + Buffer; and (4) the complete HMemAgent. Furthermore, we investigate the sensitivity of

the clue-injected context to its hyperparameter f by evaluating performance across multiple values ($f = 3, 4, 5, 6$). Results indicate that model performance remains consistently stable throughout this range, demonstrating insensitivity to variations in f . Accordingly, we adopt $f = 5$ as the experimental configuration, selected for its central position and representativeness within the evaluated interval.

TABLE II: Ablation study on Qwen-3-14B, Qwen-3-8B, and Gemma-3-4B. ✓denotes an active component. ASSOC, REF, and BUF represent the Associate, Reflect, and Buffer memory modules, respectively.

BASE MODEL	↓HINTS USED	↓TOTAL STEPS	ASSOC	REF	BUF
QWEN-3-14B	9.89	772.89	✓	✓	✓
	10.50	789.06	✓		✓
	10.06	775.81		✓	✓
	21.39	1351.44			✓
QWEN-3-8B	13.44	968.39	✓	✓	✓
	14.78	1015.89	✓		✓
	13.56	960.17		✓	✓
	23.31	1490.53			✓
GEMMA-3-4B	21.61	1337.03	✓	✓	✓
	22.94	1396.86	✓		✓
	22.39	1362.47		✓	✓
	29.75	1756.33			✓

B. Main Results

Table I shows that our method consistently outperforms baselines across core models—including GLM, Qwen, Gemma, and Llama—while preserving high task completion rates and substantially improving reasoning efficiency, as reflected in markedly shorter execution trajectories; on Qwen-3-14B, for instance, the number of reasoning steps is reduced by more than 40% compared to baseline approaches.

Gains are most evident in strategic tool-use scenarios. The reflective mechanism uncovers latent clues, aligns tool semantics with task demands, and reduces reliance on system prompts, as shown by Gemma-3-4B’s improved invocation accuracy. This reduced dependency accelerates task progression through earlier clue identification and termination of redundant exploration paths.

C. Ablation Study

To evaluate the contributions of HMemAgent’s memory levels, ablation studies were conducted across Qwen-3-14B, Qwen-3-8B, and Gemma-3-4B. As shown in Table II, the full configuration consistently achieves the lowest hint usage and total steps across most models. Removing either the Associate or Reflect module leads to notable performance degradation, with the Buffer-only variant showing the worst results. These results confirm that both reflective reasoning and associative retrieval are essential for enabling efficient and creative problem-solving. The synergy among all three levels is critical for achieving optimal performance in complex reasoning tasks.

D. Hyperparameter Sensitivity Analysis

We analyze how the hint focusing frequency (hyperparameter f) affects model exploration on ESCAPEBENCH. As shown in Table III, performance remains stable and optimal within $f \in [3, 6]$, indicating effective guidance without excessive overhead. Outside this range, performance degrades: when $f > 6$, over-focusing leads to repeated activation of erroneous clues and increased reasoning steps; when $f < 3$, premature switching prevents contextual accumulation and causes directional confusion. Our results identify an optimal window at $f \in [3, 6]$, balancing exploration depth and efficiency. This enables agents to fully leverage current hints while adapting to new information, enhancing autonomous decision-making in complex tasks without falling into fragmented or redundant behavior patterns.

E. Case Study

Figure 2 compares the performance of BaseAgent, EscapeAgent, and HMemAgent across six game settings. The results show that HMemAgent achieves faster progress and fewer hint dependencies, indicating superior efficiency and robustness. Its hierarchical memory enables consistent performance across varying task complexities, with reduced degradation in harder scenarios. This confirms the effectiveness of structured memory in enhancing strategic reasoning from raw interactions.

TABLE III: Impact of hint focus hyperparameter f on performance across Qwen models. Optimal results are achieved when $f \in [3, 6]$, balancing exploration efficiency and contextual retention.

Model	$\downarrow$ Hints Used	$\downarrow$ Total Steps	$\uparrow$ Early Exit Progress(%)
Qwen-3-14B, Focus=3	10.14	778.72	28.48%
Qwen-3-14B, Focus=4	9.92	782.72	29.27%
Qwen-3-14B, Focus=5	10.06	775.81	26.09%
Qwen-3-14B, Focus=6	10.11	773.94	25.26%
Qwen-3-8B, Focus=3	14.11	985.03	18.64%
Qwen-3-8B, Focus=4	14.06	973.11	20.80%
Qwen-3-8B, Focus=5	13.97	980.08	20.21%
Qwen-3-8B, Focus=6	13.50	963.50	20.80%

V. CONCLUSION

This paper proposes HMemAgent, a hierarchical framework integrating buffer, reflective, and associative memory via a closed-loop “recording–refining–associating” pipeline to enhance LLM agents’ creative reasoning through improved experience reuse and tool semantic exploration. EscapeBench benchmark experiments show consistent gains over existing baselines across multiple mainstream models, boosting task-solving efficiency and cross-scenario adaptability. However, the approach operates at the mechanism level without enhancing the underlying LLM’s intrinsic reasoning capabilities, remaining dependent on the base model and thus subject to its constraints. Additionally, the current implementation is limited to text-only environments, lacking multimodal inputs, which restricts applicability in complex embodied settings.

VI. ACKNOWLEDGMENTS

We would like to express gratitude to the anonymous reviewers for constructive comments. This work was supported by the “Pioneer and Leader + X” Plan Project of Zhejiang Province under Grant (2024C01162), National Natural Science Foundation of China (No.62506128), Zhejiang Provincial Natural Science Foundation of China under Grant (No. LRG25F030003 and LQN25F020023) and Open Research Fund of Zhejiang Key Laboratory of Intelligent Education Technology and Application (No. 2025ZNJYKF005).

REFERENCES

- [1] Bang Liu, Xinfeng Li, Jiayi Zhang, Jinlin Wang, Tanjin He, Sirui Hong, Hongzhang Liu, Shaokun Zhang, Kaitao Song, Kunlun Zhu, Yuheng Cheng, Suyuchen Wang, Xiaoqiang Wang, Yuyu Luo, Haibo Jin, Peiyang Zhang, Ollie Liu, Jiaqi Chen, Huan Zhang, Zhaoyang Yu, Haochen Shi, Boyan Li, Dekun Wu, Fengwei Teng, Xiaojun Jia, Jiawei Xu, Jinyu Xiang, Yizhang Lin, Tianming Liu, Tongliang Liu, Yu Su, Huan Sun, Glen Berseth, Jianyun Nie, Ian T. Foster, Logan T. Ward, Qingyun Wu, Yu Gu, Mingchen Zhuge, Xiangru Tang, Haohan Wang, Jiaxuan You, Chi Wang, Jian Pei, Qiang Yang, Xiaoliang Qi, and Chenglin Wu. Advances and challenges in foundation agents: From brain-inspired intelligence to evolutionary, collaborative, and safe systems. *CoRR*, abs/2504.01990, 2025.
- [2] Junyu Luo, Weizhi Zhang, Ye Yuan, Yusheng Zhao, Junwei Yang, Yiyang Gu, Bohan Wu, Binqi Chen, Ziyue Qiao, Qingqing Long, Rongcheng Tu, Xiao Luo, Wei Ju, Zhiping Xiao, Yifan Wang, Meng Xiao, Chenwu Liu, Jingyang Yuan, Shichang Zhang, Yiqiao Jin, Fan Zhang, Xian Wu, Hanqing Zhao, Dacheng Tao, Philip S. Yu, and

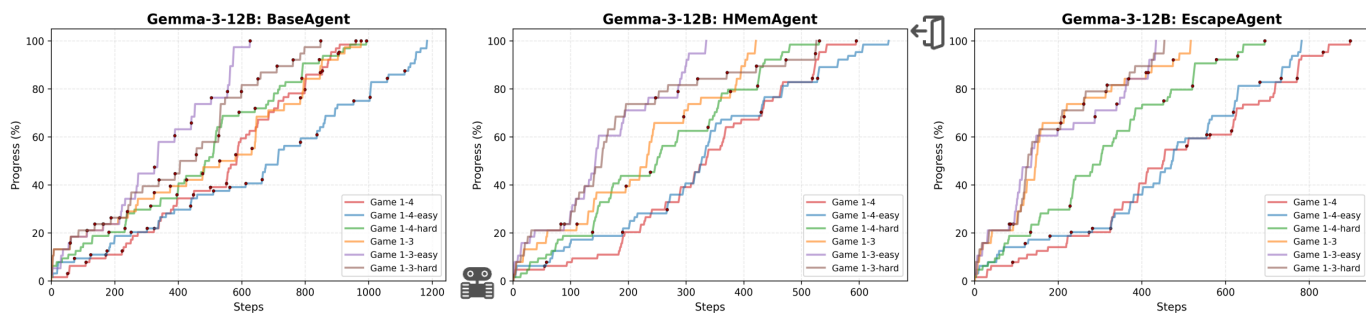


Fig. 2: Progression graphs comparing BaseAgent, EscapeAgent, and HMemAgent across six game settings, with red dots marking hint usage. HMemAgent achieves faster progress and fewer hints, demonstrating superior efficiency and robustness through its hierarchical memory architecture.

Ming Zhang. Large language model agent: A survey on methodology, applications and challenges. *CoRR*, abs/2503.21460, 2025.

- [3] Asaf Yehudai, Lilach Eden, Alan Li, Guy Uziel, Yilun Zhao, Roy Bar-Haim, Arman Cohan, and Michal Shmueli-Scheuer. Survey on evaluation of llm-based agents. *CoRR*, abs/2503.16416, 2025.
- [4] Yuyang Hu, Shichun Liu, Yanwei Yue, Guibin Zhang, Boyang Liu, Fangyi Zhu, Jiahao Lin, Honglin Guo, Shihan Dou, Zhiheng Xi, et al. Memory in the age of ai agents. *arXiv preprint arXiv:2512.13564*, 2025.
- [5] Cheng Qian, Chenyan Xiong, Zhenghao Liu, and Zhiyuan Liu. Toolink: Linking toolkit creation and using through chain-of-solving on open-source model. *CoRR*, abs/2310.05155, 2023.
- [6] Chunru Lin, Haotian Yuan, Yian Wang, Xiaowen Qiu, Tsun-Hsuan Wang, Minghao Guo, Bohan Wang, Yashraj Narang, Dieter Fox, and Chuang Gan. RobotSmith: Generative robotic tool design for acquisition of complex manipulation skills. *CoRR*, abs/2506.14763, 2025.
- [7] Yufei Tian, Abhilasha Ravichander, Lianhui Qin, Ronan Le Bras, Raja Marjeh, Nanyun Peng, Yejin Choi, Thomas L. Griffiths, and Faeze Brahman. Macgyver: Are large language models creative problem solvers? In Kevin Duh, Helena Gómez-Adorno, and Steven Bethard, editors, *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, NAACL 2024, Mexico City, Mexico, June 16-21, 2024, pages 5303–5324. Association for Computational Linguistics, 2024.
- [8] Shane Legg and Marcus Hutter. Universal intelligence: A definition of machine intelligence. *Minds and machines*, 17(4):391–444, 2007.
- [9] Margaret A. Boden. Creativity and artificial intelligence. *Artificial Intelligence*, 103(1-2):347–356, 1998.
- [10] Cheng Qian, Peixuan Han, Qinyu Luo, Bingxiang He, Xiusi Chen, Yuji Zhang, Hongyi Du, Jianyi Yao, Xiaocheng Yang, Denghui Zhang, et al. Escapebench: Towards advancing creative intelligence of language model agents. In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 798–820, 2025.
- [11] Harsh Kumar, Jonathan Vincentius, Ewan Jordan, and Ashton Anderson. Human creativity in the age of llms: Randomized experiments on divergent and convergent thinking. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, pages 1–18, 2025.
- [12] Yunpu Zhao, Rui Zhang, Wenyi Li, and Ling Li. Assessing and understanding creativity in large language models. *Machine Intelligence Research*, 22(3):417–436, 2025.
- [13] Sian Gooding, Lucia Lopez-Rivilla, and Edward Grefenstette. Writing as a testbed for open ended agents. *arXiv preprint arXiv:2503.19711*, 2025.
- [14] Pronita Mehrotra, Aishni Parab, and Sumit Gulwani. Enhancing creativity in large language models through associative thinking strategies. *arXiv preprint arXiv:2405.06715*, 2024.
- [15] Moran Mizrahi, Chen Shani, Gabriel Stanovsky, Dan Jurafsky, and Dafna Shahaf. Cooking up creativity: A cognitively-inspired approach for enhancing llm creativity through structured representations. *arXiv preprint arXiv:2504.20643*, 2025.
- [16] Nuo Chen, Yicheng Tong, Jiaying Wu, Minh Duc Duong, Qian Wang, Qingyun Zou, Bryan Hooi, and Bingsheng He. Beyond brainstorming: What drives high-quality scientific ideas? lessons from multi-agent collaboration. *arXiv preprint arXiv:2508.04575*, 2025.
- [17] Zeyu Zhang, Xiaohu Bo, Chen Ma, Rui Li, Xu Chen, Quanyu Dai, Jieming Zhu, Zhenhua Dong, and Ji-Rong Wen. A survey on the memory mechanism of large language model based agents. *CoRR*, abs/2404.13501, 2024.
- [18] Zhen Tan, Jun Yan, I-Hung Hsu, Rujun Han, Zifeng Wang, Long T. Le, Yiwen Song, Yanfei Chen, Hamid Palangi, George Lee, Anand Rajan Iyer, Tianlong Chen, Huan Liu, Chen-Yu Lee, and Tomas Pfister. In prospect and retrospect: Reflective memory management for long-term personalized dialogue agents. In Wanxiang Che, Joyce Nabende, Ekaterina Shutova, and Mohammad Taher Pilehvar, editors, *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, ACL 2025, Vienna, Austria, July 27 - August 1, 2025, pages 8416–8439. Association for Computational Linguistics, 2025.
- [19] Yaxiong Wu, Sheng Liang, Chen Zhang, Yichao Wang, Yongyue Zhang, Huifeng Guo, Ruiming Tang, and Yong Liu. From human memory to AI memory: A survey on memory mechanisms in the era of llms. *CoRR*, abs/2504.15965, 2025.
- [20] Guibin Zhang, Muxin Fu, and Shuicheng Yan. Memgen: Weaving generative latent memory for self-evolving agents. *CoRR*, abs/2509.24704, 2025.
- [21] Yiming Du, Wenyu Huang, Danna Zheng, et al. Rethinking memory in llm based agents: Representations, operations, and emerging topics, 2025.
- [22] Zhiyu Li, Shichao Song, Chenyang Xi, et al. Memos: A memory OS for AI system. *CoRR*, abs/2507.03724, 2025.
- [23] Chuanrui Hu, Xingze Gao, Zuyi Zhou, Dannong Xu, Yi Bai, Xintong Li, Hui Zhang, Tong Li, Chong Zhang, Lidong Bing, and Yafeng Deng. Evermemos: A self-organizing memory operating system for structured long-horizon reasoning, 2026.
- [24] Wujiang Xu, Zujie Liang, Kai Mei, Hang Gao, Juntao Tan, and Yongfeng Zhang. A-MEM: agentic memory for LLM agents. *CoRR*, abs/2502.12110, 2025.
- [25] Prateek Chhikara, Dev Khant, Saket Aryan, Taranjeet Singh, and Deshray Yadav. Mem0: Building production-ready AI agents with scalable long-term memory. *CoRR*, abs/2504.19413, 2025.
- [26] Yuxiang Zhang, Jiangming Shu, Ye Ma, Xueyuan Lin, Shangxi Wu, and Jitao Sang. Memory as action: Autonomous context curation for long-horizon agentic tasks. *CoRR*, abs/2510.12635, 2025.
- [27] Aohan Zeng, Xin Lv, Qinkai Zheng, et al. GLM-4.5: agentic, reasoning, and coding (ARC) foundation models. *CoRR*, abs/2508.06471, 2025.
- [28] Daya Guo, Dejian Yang, et al. Deepseek-r1 incentivizes reasoning in llms through reinforcement learning. *Nature*, 645(8081):633–638, September 2025.
- [29] An Yang, Anfeng Li, Baosong Yang, et al. *f* technical report. *CoRR*, abs/2505.09388, 2025.
- [30] Gemma Team. Gemma 3 technical report. *CoRR*, abs/2503.19786, 2025.
- [31] Llama Team. The llama 3 herd of models. *CoRR*, abs/2407.21783, 2024.