

Memory Architectures for Game Agents: A Systematic Mapping Study Bridging Classical and LLM Paradigms

Gabriel Machado Santos¹, Rita Maria Silva Julia¹, Marcelo Nascimento¹

¹Computer Science Faculty - Federal University of Uberlândia
Uberlândia, MG

{gabrielmsantos, rita, marcelo.nascimento}@ufu.br

Abstract. Introduction: Two decades of research on memory for game agents have produced two largely disjoint traditions: classical symbolic game AI (blackboards, GOAP, behavior trees, case-based reasoning, Soar/ACT-R) and LLM-based agents (memory streams, MemGPT, A-MEM, CoALA, Voyager). Despite clear conceptual overlap, no prior survey covers both under a unified lens for game agents. **Objective:** We bridge these traditions through a systematic mapping that classifies, compares, and identifies open problems across both paradigms. **Methodology:** Following Petersen et al.’s guidelines, we surveyed 113 works from 1972–2025 across ACM DL, IEEE Xplore, Scopus, and arXiv, with forward and backward snowballing on seminal anchors. Each work was coded along five axes: temporal horizon, representation, lifecycle operations, grounding target, and agent role. **Results:** The mapping yields (i) a unified five-axis taxonomy for both paradigms, (ii) empty cells exposing under-served combinations (notably crowd-agent memory and explicit forgetting), (iii) an integration lens identifying what LLM systems reinvent from classical game AI, what they add, and what they lose, and (iv) six open problems framing a research agenda for NPC memory.

Keywords memory architecture, non-player characters, large language models, cognitive architectures, systematic mapping, game AI.

1. Introduction

A non-player character (NPC) that forgets a player’s first conversation breaks immersion in a way few other AI failures do; continuity of memory is what separates a believable companion from a stateless dispenser of dialogue lines. Two research communities have answered this challenge — classical symbolic game AI and large-language-model-based agents — and they barely cite each other.

Classical game AI built sophisticated memory infrastructure over two decades — GOAP world states [Orkin 2006], behavior-tree blackboards [Isla 2005], ABL working-memory elements [Mateas e Stern 2002], Soar’s episodic memory [Nuxoll e Laird 2007], case-based plan libraries [Ontañón et al. 2010], FAtiMA’s autobiographical emotional memory [Dias e Paiva 2005] — prizing deterministic authoring, debuggability, and designer control over what an NPC knows. LLM-era agent memory built a parallel infrastructure in roughly two years — memory streams with reflection [Park et al. 2023], virtual context management [Packer et al. 2023], evolving Zettelkasten-style notes [Xu et al. 2025], lifelong skill libraries [Wang et al. 2024a], and the CoALA framework mapping cognitive-architecture memory onto LLM agents [Sumers et al. 2024] — prizing associative recall, generative summarization, and emergent behavior.

The traditions share more than the surface vocabulary suggests: MemGPT’s tiered virtual memory echoes blackboards, Generative Agents’ reflection echoes Soar’s chunking, and retrieval-augmented dialogue echoes case-based reasoning — with the Restaurant Game’s crowdsourced episodic corpus [Orkin e Roy 2007] pre-figuring retrieval-augmented generation by 15 years. Yet the LLM-memory literature rarely cites this lineage, and existing surveys split along the paradigm boundary: memory-focused LLM-agent surveys [Zhang et al. 2025, Sumers et al. 2024] are domain-agnostic; LLMs-and-games surveys [Gallotta et al. 2024, Sweetser e Rogers 2024] treat memory as one sub-topic; classical game-AI surveys [Yannakakis e Togelius 2018, Yannakakis e Togelius 2015] predate the LLM era.

Position. LLM memory systems for game agents are re-deriving ideas the game AI community spent two decades refining, often without the authoring discipline that made those ideas usable in shipped games. A systematic, paradigm-spanning view is overdue.

Contributions. This paper offers (1) a five-axis taxonomy — temporal horizon, representation, lifecycle operations, grounding target, agent role — classifying both classical and LLM memory architectures uniformly; (2) a systematic mapping of 113 works (1972–2025) following established guidelines [Petersen et al. 2008, Petersen et al. 2015]; (3) an integration lens identifying what LLM systems reinvent, add, and lose relative to classical game AI; and (4) a research agenda of six open problems grounded in the empty cells of the taxonomy.

Section 2 reviews cognitive-science grounding and the two paradigms; Section 3 positions the work against existing surveys; Section 4 details the mapping protocol; Section 5 presents the taxonomy; Section 6 reports the mapping; Section 7 develops the integration lens; Section 8 states the agenda; Sections 9 and 10 discuss validity and conclude.

2. Background

2.1. Cognitive science grounding

Tulving’s episodic/semantic distinction [Tulving 1972, Tulving 1985] grounds the temporal axis: episodic memory stores time- and context-indexed events, semantic memory generalized facts. Baddeley and Hitch’s working-memory model [Baddeley e Hitch 1974], extended with the episodic buffer [Baddeley 2000], frames the short-term substrate over which retrieval operates; Squire’s declarative/non-declarative split [Squire 2004] adds procedural memory (skills, habits) as a separate system. ACT-R [Anderson et al. 2004] unified these distinctions in a computational substrate that influenced both the classical game-AI lineage (Soar) and contemporary LLM-agent frameworks (CoALA).

2.2. Classical game-AI memory

The shipped-game lineage matures from FSMs into three branches. The *world-state* branch treats memory as a symbolic snapshot, exemplified by GOAP’s predicate vector with A*-driven planning [Orkin 2006, Orkin 2004]. The *blackboard* branch shares typed working memory across modules: Halo’s behavior trees [Isla 2005, Isla 2008] and the Behavior Tree Starter Kit [Champandard e Dunstan 2013], formalized in robotics

by Colledanchise and Ögren [Colledanchise e Ögren 2017, Colledanchise e Ögren 2018]. The *believable-agent* branch, rooted in the Oz Project, contributes Hap, ABL, and Façade [Loyall e Bates 1991, Loyall 1997, Mateas e Stern 2002, Mateas e Stern 2005] with persistent emotional state [Bates 1994].

Cognitive architectures gave the first rigorous treatment of NPC episodic memory: Laird’s Quakebot [Laird 2001, Laird e Duchi 2000], Nuxoll and Laird’s episodic-memory extensions [Nuxoll e Laird 2007, Nuxoll e Laird 2012] with Gorski’s learned-retrieval policy [Gorski e Laird 2011], and Soar’s four memory modules [Laird 2012]. Case-based reasoning brought episodic memory to RTS via Darmok and successors [Ontañón et al. 2010, Ontañón et al. 2007, Aha et al. 2005, Mehta et al. 2009], goal-driven autonomy in StarCraft [Weber et al. 2010], and Ontañón et al.’s RTS-AI survey [Ontañón et al. 2013]. The Sims’ smart objects [Kallmann e Thalmann 1999, Kallmann e Thalmann 2002] offloaded memory into the environment, anticipating RAG’s externalization, while FAtiMA and its descendants [Dias e Paiva 2005, Dias et al. 2014, Aylett et al. 2005, Lim et al. 2012] integrated autobiographical memory with appraisal-based emotion [Marsella e Gratch 2009].

2.3. LLM agent memory

The LLM-agent memory literature condensed into shape between 2022 and 2024. Generative Agents [Park et al. 2023] introduced a memory stream of natural-language observations scored by recency, importance, and relevance, with reflection synthesizing higher-level beliefs — a direct echo of consolidation in cognitive architectures. MemGPT [Packer et al. 2023] treats the LLM context as a virtual-memory page cache backed by external storage; A-MEM [Xu et al. 2025] structures memory as a Zettelkasten of dynamically linked notes; Reflexion [Shinn et al. 2023] casts verbal self-reflection as an episodic buffer; and Voyager [Wang et al. 2024a] maintains a procedural skill library with embedding-based retrieval. CoALA [Sumers et al. 2024] provides the explicit theoretical bridge, arguing that LLM agents need exactly the four memory types Soar and ACT-R already define.

Underneath these agents sit retrieval substrates — RAG [Lewis et al. 2020], dense passage retrieval [Karpukhin et al. 2020], REALM [Guu et al. 2020], and long-context attention [Munkhdalai et al. 2024, Mohtashami e Jaggi 2023] — plus memory-focused extensions: MemoryBank’s Ebbinghaus-curve forgetting [Zhong et al. 2024], ExpeL’s cross-task insights [Zhao et al. 2024], and structured-memory variants such as RET-LLM/MemLLM [Modarressi et al. 2024].

3. Related Surveys

This paper occupies the intersection of three properties: *memory-focused*, *game-agent-focused*, and *covers both classical and LLM paradigms*. Table 1 positions ten prominent surveys against these properties. None occupies the full intersection.

The closest related work is the CoALA framework [Sumers et al. 2024], which maps cognitive-architecture memory types onto LLM agents but is not a systematic mapping and is domain-agnostic. The Zhang et al. memory survey [Zhang et al. 2025] treats LLM-agent memory in depth without addressing games. The Gallotta et al. survey [Gallotta et al. 2024] covers LLMs in games comprehensively but treats memory as one sub-topic among eight roles.

Table 1. Positioning against related surveys. “M” = memory-focused, “G” = game-agents as primary domain, “C” = covers classical paradigm, “L” = covers LLM paradigm.

Survey	M	G	C	L
Zhang et al. [Zhang et al. 2025]	✓			✓
Sumers et al. (CoALA) [Sumers et al. 2024]	✓		✓	✓
Wang et al. (Human Memory → AI) [Wang et al. 2025]	✓			✓
Gallotta et al. (LLMs and Games) [Gallotta et al. 2024]		✓		✓
Sweetser & Rogers [Sweetser e Rogers 2024]		✓		✓
Yannakakis & Togelius [Yannakakis e Togelius 2018, Yannakakis e Togelius 2015]		✓	✓	
Justesen et al. [Justesen et al. 2020]		✓	✓	
Kotseruba & Tsotsos [Kotseruba e Tsotsos 2020]	✓		✓	
Langley et al. [Langley et al. 2009]	✓		✓	
This paper	✓	✓	✓	✓

4. Method

We follow the systematic mapping study (SMS) protocol of Petersen et al. [Petersen et al. 2008, Petersen et al. 2015], which is designed for taxonomy-building and gap-identification rather than effect-size synthesis.

4.1. Research questions

- RQ1.** What memory architectures for game agents have been proposed across the classical and LLM eras?
- RQ2.** Along which dimensions do these architectures vary?
- RQ3.** Which memory requirements are well-served by each tradition, and which are neglected?
- RQ4.** What open problems emerge at the intersection?

4.2. Sources, search, and selection

We searched ACM Digital Library, IEEE Xplore, Scopus, and arXiv, complemented by forward and backward snowballing on a seed of 15 anchor works¹ chosen before querying to jointly anchor both paradigms and every cluster in Table 2. Seeds were selected on three grounds: paradigm coverage (classical anchors balanced against LLM-era anchors), citation centrality (the most-cited architectural contribution in each cluster), and bridging potential (favoring works that explicitly connect cognitive science, classical game AI, and LLM agents, e.g., CoALA).

Queries combined three term clusters: *memory* (memory, blackboard, episodic, working memory, retrieval, forgetting, reflection), *agent/game* (NPC, game agent, behavior tree, GOAP, drama manager, tutor, companion), and *LLM-era* (large

¹Classical: F.E.A.R. GOAP [Orkin 2006], Halo behavior trees [Isla 2005], Soar episodic memory [Nuxoll e Laird 2007], Darmok CBR [Ontañón et al. 2010], FAtiMA [Dias e Paiva 2005], Façade [Mateas e Stern 2005], Crystal Island [Rowe et al. 2011], and the Restaurant Game [Orkin e Roy 2007]. LLM-era: Generative Agents [Park et al. 2023], MemGPT [Packer et al. 2023], Voyager [Wang et al. 2024a], A-MEM [Xu et al. 2025], CoALA [Sumers et al. 2024], Reflexion [Shinn et al. 2023], and JARVIS-1 [Wang et al. 2024b].

language model, LLM agent, RAG, MemGPT, generative agent). arXiv inclusion is essential: many LLM-cluster works originated as preprints and remain the most-cited reference even after peer-reviewed versions appear (e.g., MemGPT [Packer et al. 2023], Voyager [Wang et al. 2024a]), while snowballing captured industry-canonical references (GDC talks for F.E.A.R. and Halo) that no academic database indexes well. Because this is a mapping study characterizing a research landscape rather than synthesizing effect sizes, excluding influential preprints would distort the map more than including them; we mitigate the peer-review concern by tracking review status explicitly as a corpus variable (Table 2).

We applied the following criteria to every candidate work.

Inclusion. A work was included if it:

- (I1) proposes, evaluates, or formalizes a memory mechanism;
- (I2) targets a game agent, NPC, or directly comparable virtual-agent context; and
- (I3) published up to 2025.

Exclusion. A work was excluded if it:

- (E1) addressed only player-side memory aids rather than agent memory;
- (E2) used memory incidentally, without architectural innovation (e.g., reinforcement-learning replay buffers); or
- (E3) performed pure narrative generation with no agent locus.

A work had to satisfy all three inclusion criteria and none of the exclusion criteria to enter the corpus.

4.3. Coding and corpus

Each work was coded along the five taxonomy axes (Section 5) by the first author using a structured rubric. After de-duplication and full-text screening, the corpus comprises 113 unique works; Table 2 reports the breakdown by cluster.

Table 2. Corpus composition by cluster.

Cluster	Papers	Peer-reviewed	Bridges paradigms
1. Classical game-AI memory	32	25	1
2. Mid-era NPC, narrative, tutor	26	25	2
3. LLM agent memory (foundational)	26	22	12
4. LLM agents in games	17	13	14
5. Surveys + cognitive science	12	12	0
Total	113	97	29

5. The Taxonomy

We propose five axes that apply uniformly to both paradigms, each answering a distinct design question.

Temporal horizon (working / episodic / semantic / procedural). Following Tulving and Squire, we distinguish short-lived working memory, event-indexed episodic

memory, generalized semantic memory, and skill-encoded procedural memory. Soar realizes all four as separate modules [Laird 2012], and CoALA argues LLM agents need the same [Sumers et al. 2024]. F.E.A.R.’s GOAP world state is working+procedural; Generative Agents’ memory stream is episodic with reflective semantic consolidation; Voyager’s skill library is procedural; A-MEM is episodic-to-semantic.

Representation (symbolic / subsymbolic / hybrid). Symbolic representations (predicates, frames, OCC tuples) dominate the classical era; subsymbolic ones (embeddings, neural activations) dominate retrieval substrates such as DPR [Karpukhin et al. 2020] and REALM [Guu et al. 2020]. The hybrid cell is the most active zone: Voyager combines code with embedding indices, Ashby et al. [Ashby et al. 2023] combine knowledge graphs with fine-tuned LLMs, and CoALA’s modular memories cross representation types within one architecture.

Lifecycle operations (encoding / consolidation / retrieval / forgetting). This is the most uneven axis. Classical work has rich retrieval and decay (Isla and Blumberg’s object persistence [Isla e Blumberg 2002] is a rare, important treatment); LLM work has rich encoding and retrieval but *sparse forgetting*. Only MemoryBank’s Ebbinghaus curve [Zhong et al. 2024] and Generative Agents’ importance decay [Park et al. 2023] treat forgetting as first-class — one of the paper’s core findings.

Grounding target (world state / player model / narrative / dialogue / self). Classical agents ground heavily in world state (GOAP, BT blackboards); tutors and drama managers ground in player models [Corbett e Anderson 1995, VanLehn 2006, Shute e Ventura 2013, Rowe et al. 2011]; LLM agents lean toward dialogue and self-reflection [Shinn et al. 2023, Madaan et al. 2023]. Hybrid grounding (world state plus dialogue) is rare but emerging in D&D-as-dialogue work [Callison-Burch et al. 2022, Zhu et al. 2023, Zhou et al. 2023].

Agent role (tutor / companion / opponent / ambient / crowd / general). Two roles stand out for their scarcity rather than their coverage: crowd-agent memory, and commercial-companion memory (Skyrim, Mass Effect, Shadow of Mordor’s Nemesis System), the latter lacking academic coverage despite shipping at scale.

Synthesis. Table 3 maps twelve anchor systems onto the five axes, exposing structural similarities between F.E.A.R. and MemGPT (working+symbolic with tiered storage), Façade and Generative Agents (episodic with consolidation-driven narrative coherence), and Soar Quakebot and CoALA (four-memory cognitive architectures differing in representation substrate).

6. Mapping Results

Corpus overview. The 113-work corpus spans 1972–2025; eight works predate 2000 (three cognitive-science foundations and five early believable-agent architectures), with the architectural literature concentrated from 2000 onward. The distribution is bimodal: a steady classical-era stream (2000–2018), a brief lull, and a sharp LLM-era surge from 2022. Some 38 works (34%) appeared in 2023–2025, reflecting the post-ChatGPT inflection; Figure 1 plots the 96 architectural works from 2000–2025.

Representation. Symbolic representations account for 48 works (42%), subsymbolic for 11 (10%), and hybrid for 42 (37%); the remaining 12 are surveys with no

Table 3. Anchor systems across the five taxonomy axes. T = temporal (W/E/S/P), R = representation (Sym/Sub/Hyb), L = lifecycle ops emphasized, G = grounding, A = agent role.

System	T	R	L	G	A
F.E.A.R. GOAP [Orkin 2006]	W+P	Sym	enc, ret	world	opp
Halo 2 BT [Isla 2005]	W+P	Sym	enc, ret, decay	world, beliefs	opp
Façade [Mateas e Stern 2005]	W+E	Sym	enc, ret, consol	narrative	companion
Soar Quakebot [Laird 2001]	W+E	Sym	enc, ret, sim	world, self	opp
FAtiMA [Dias e Paiva 2005]	E+S	Sym	enc, ret, consol	events, emotion	companion
Crystal Island [Rowe et al. 2011]	E+long	Sym	enc, ret, consol	narrative, player	tutor
Generative Agents [Park et al. 2023]	W+E+S	Hyb	enc, consol, ret, forget	world, self, dialog	ambient
MemGPT [Packer et al. 2023]	W+LT	Hyb	page, summarize, ret	dialog, docs	general
Voyager [Wang et al. 2024a]	P+S	Hyb	enc, consol, ret	world, self	general
A-MEM [Xu et al. 2025]	E+S	Hyb	build, link, evolve, ret	dialog, history	general
CoALA [Sumers et al. 2024]	all 4	Hyb	all	all	framework
JARVIS-1 [Wang et al. 2024b]	E+P+multimodal	Hyb	enc, dual-key ret, self-instruct	world (vision)	general

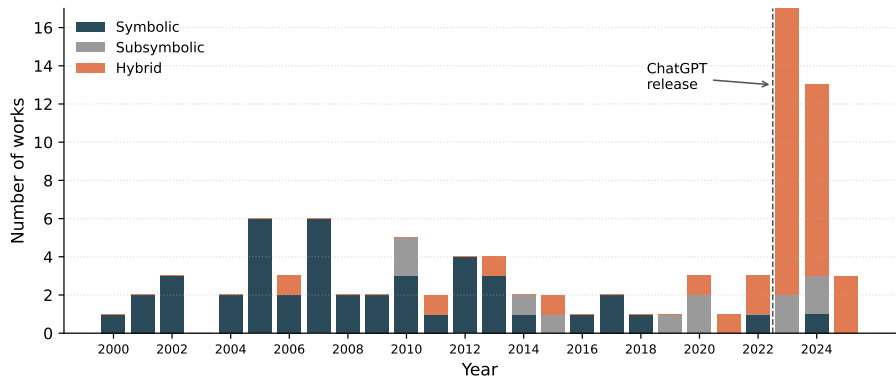


Figure 1. Yearly distribution of the 96 architectural works published 2000–2025, stacked by representation. Five earlier architectural works (1991–1999) and 12 surveys/cognitive-science works are excluded; eight corpus works predate 2000. The dashed line marks the ChatGPT release (November 2022). Symbolic representations dominate through 2018, while hybrid approaches dominate from 2023 onward.

architectural commitment. Hybrid is now the dominant new-architecture choice: every LLM-in-games system in Cluster 4 uses it, as do all major LLM-agent systems with persistent memory.

Lifecycle operations. Encoding and retrieval dominate (80 and 96 of 113). Consolidation appears in 25 works, concentrated in drama-manager and reflective LLM-agent work. Forgetting appears in only 11, confirming the gap noted in Section 5.

Agent role. The corpus skews toward general-purpose substrates (39 works). Ambient agents (21), companions (20), and opponents (15, mostly RTS and FPS) are all well-represented, while tutors (6) are comparatively few. Crowd-agent memory is entirely absent from the corpus (0 works) (Figure 2).

Co-occurrence patterns. Three regions dominate: the classical corner (working+symbolic+world-state+opponent) holds F.E.A.R., Halo, and the RTS-CBR lineage; the LLM corner (episodic+hybrid+dialogue+ambient) holds Generative Agents, LARP, and Character-LLM; and the procedural corner (procedural+hybrid+world-

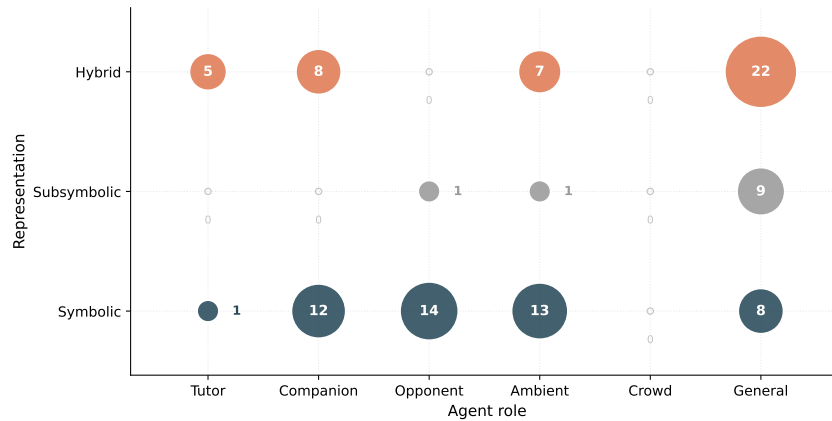


Figure 2. Distribution of the 101 architectural works (113 minus 12 surveys) across representation and agent role; bubble area is proportional to count. Empty cells (marked 0) include the entire crowd-agent column.

state+general), small but growing, includes Voyager and JARVIS-1.

Empty cells. Four combinations are essentially empty in the peer-reviewed record: (i) crowd-agent memory of any kind; (ii) explicit forgetting in any LLM-game system other than MemoryBank; (iii) procedural memory in tutors despite obvious applicability; and (iv) cross-session persistence under save/load semantics. These motivate the research agenda in Section 8.

7. The Integration Lens

We organize the position around three claims, presented here as an interpretive synthesis of the mapping results reported in Section 6, rather than as additional mapping findings.

What LLM systems reinvent. MemGPT’s tiered virtual memory [Packer et al. 2023] is structurally a blackboard with paging — the same coordination pattern Halo formalized in 2005. Generative Agents’ reflection [Park et al. 2023] is consolidation, the same operation Soar’s chunking [Laird 2012] and Mehta et al.’s Meta-Darmok [Mehta et al. 2009] performed over symbolic case bases. Voyager’s skill library [Wang et al. 2024a] is case-based planning [Ontañón et al. 2010] with code as cases and embeddings as similarity metric. Most strikingly, retrieval-augmented dialogue is the architecture Orkin and Roy proposed in 2007 [Orkin e Roy 2007], when they mined thousands of crowdsourced traces as an external episodic corpus for an NPC — a paper the LLM-memory literature almost never cites.

What LLM systems genuinely add. Three additions are real. First, *associative recall over unstructured natural language*: classical memory required predicate authoring; LLM memory ingests whatever a player says. Second, *generative summarization* as a consolidation primitive that classical systems lacked; reflection produces new semantic content rather than merely indexing old episodic content. Third, *zero-shot transfer*: SPRING [Wu et al. 2023] reads the Crafter paper as a knowledge source, a capability without classical equivalent. The fact that an LLM can also *fabricate* memories it never stored is a new failure mode but, in narrative contexts, occasionally a feature.

What LLM systems lose. Four losses are non-trivial. *Deterministic authoring*: a designer cannot guarantee that a Generative Agent will recall a scripted detail at a scripted time. *Latency guarantees*: a 60-Hz tick budget does not accommodate API calls. *Debuggability*: a designer can step through an ABL behavior tree but cannot step through an embedding retrieval. *Designer control over canon*: classical drama managers [Nelson et al. 2006, Riedl et al. 2008, Sharma et al. 2010] let authors fix what an NPC must believe at narrative beat K; LLM memory tends toward emergent inconsistency.

The core tension. The defining design axis of the next decade is *narrative control versus emergent memory*. Classical work optimized for control; LLM work optimized for emergence. The interesting hybrids — Ashby et al.’s KG+LLM quest generation [Ashby et al. 2023], Voyager’s code+embedding split [Wang et al. 2024a], CoALA’s modular memories [Sumers et al. 2024], FIREBALL’s state-grounded D&D [Zhu et al. 2023] — point to architectures that get both, at the cost of authoring complexity.

8. Research Agenda

The empty cells of the mapping motivate six open problems.

(1) Standard benchmarks for NPC memory quality. No equivalent of LongBench exists for game agents. We propose four dimensions: cross-session consistency, character-knowledge constraint (an NPC must not know what its character would not know), retrieval recall under distractor interference, and narrative-canon adherence. A shared benchmark would unblock comparative evaluation across the LLM-NPC literature.

(2) Forgetting as a design primitive. Classical decay [Isla e Blumberg 2002] and Ebbinghaus-style forgetting [Zhong et al. 2024] are exceptions. A first-class forgetting API for game agents — “this NPC forgets gossip after three sessions but never forgets marriage promises” — is unsolved. The design problem is selectivity: forgetting should be authorable, not stochastic.

(3) Authoring tools for hybrid memory. Designers currently choose between deterministic authoring (BT/GOAP) and emergent prompting (LLM). The visual-scripting equivalent for hybrid memory architectures — where a designer can wire deterministic constraints into an otherwise emergent system — does not yet exist.

(4) Cross-session consistency under save/load. Almost no work studies what happens when a player loads a save from yesterday into an LLM-memory NPC. CoALA hints at the structural requirements; commercial games have not deployed it. This is the most immediately deployable open problem.

(5) Memory governance for minors. When the player is a 14-year-old in an educational context and the NPC remembers everything they said, what are the ethical and design constraints? Privacy, parental visibility, and right-to-be-forgotten map onto memory-architecture choices in concrete ways. This problem is acute for serious-games and educational-AI work.

(6) Runtime cost for indie and mobile. Every LLM-NPC paper in our corpus runs on cloud GPUs. Distilled, on-device, persistent NPC memory — needed for indie and mobile games — is essentially open. Hybrid architectures that keep the deterministic

substrate on-device and consult the cloud only for emergent operations are a plausible path.

8.1. Revisiting the Research Questions

RQ1 (architectures proposed) is answered by the corpus itself (Table 2): the 113 works span classical game-AI memory (Cluster 1, 32 works: GOAP, blackboards, cognitive architectures, and case-based planning), a mid-era NPC, narrative, and tutor stream (Cluster 2, 26 works), the LLM era (Clusters 3–4, 43 works: memory streams, virtual paging, Zettelkasten-style notes, and skill libraries), and surveys plus cognitive-science grounding (Cluster 5, 12 works), all synthesized across twelve anchor systems in Table 3. **RQ2** (dimensions of variation) is answered by the five-axis taxonomy itself — temporal horizon, representation, lifecycle operations, grounding target, and agent role (Section 5) — whose distributions are reported in Section 6 and Figure 2. **RQ3** (well-served vs. neglected requirements) is answered jointly by the lifecycle-operation counts in Section 6 (encoding and retrieval dominant; forgetting present in only 11/113 works) and the integration lens in Section 7, which shows that deterministic authoring, latency guarantees, debuggability, and canonical control — well-served by classical architectures — are precisely the requirements LLM-based memory neglects. **RQ4** (open problems) is answered by the six-item research agenda in Section 8, each item derived directly from an empty cell or neglected requirement identified while answering RQ3.

9. Threats to Validity

Search comprehensiveness. English-dominant databases under-represent Portuguese-language work; we partially mitigated by including SBGames-published anchors [Lima et al. 2016, Baffa et al. 2017]. **Industry-source dependency.** The classical canon relies on GDC talks (Orkin, Isla) without peer-reviewed equivalents; we cite both the talks and their academic descendants. **Temporal instability.** The LLM corpus will be partially outdated by review; 2025 anchors [Xu et al. 2025, Wang et al. 2025] mitigate this. **Single-coder coding.** The 113-work corpus was coded along the five taxonomy axes by the first author alone; we mitigate this by releasing the full coded corpus as supplementary material² for community scrutiny. **Construct validity.** Symbolic predicates and embeddings both count as “memory” but operate differently — the taxonomy makes the distinction precise rather than eliding it.

10. Conclusion

Memory for game agents has been studied by two communities that rarely cite each other. Mapping 113 works onto a unified five-axis taxonomy exposes the structural overlaps (MemGPT echoes blackboards, reflection echoes chunking, retrieval-augmented dialogue echoes case-based reasoning), what LLM systems genuinely add (associative recall, generative consolidation, zero-shot transfer), and what they lose (deterministic authoring, latency guarantees, debuggability, canonical control). The empty cells — crowd-agent memory, explicit forgetting, cross-session persistence, on-device LLM memory — define the next decade’s agenda, with hybrid architectures combining deterministic authoring and emergent retrieval as the near-term path. Narrative control versus emergent memory is the open question that will shape how shipped games adopt LLM-based NPCs.

²Supplementary materials and the coded corpus are available at <https://doi.org/10.5281/zenodo.19796779>.

References

- Aha, D. W., Molineaux, M., e Ponsen, M. J. V. (2005). Learning to win: Case-based plan selection in a real-time strategy game. In *Proc. ICCBR*.
- Anderson, J. R., Bothell, D., Byrne, M. D., Douglass, S., Lebiere, C., e Qin, Y. (2004). An integrated theory of the mind. *Psychological Review*, 111(4):1036–1060.
- Ashby, T., Webb, B. K., Knapp, G., Searle, J., e Fulda, N. (2023). Personalized quest and dialogue generation in role-playing games: A knowledge graph- and language model-based approach. In *Proc. ACM CHI*.
- Aylett, R., Louchart, S., Dias, J., Paiva, A., e Vala, M. (2005). FearNot! – an experiment in emergent narrative. In *Proc. Intelligent Virtual Agents (IVA)*.
- Baddeley, A. (2000). The episodic buffer: A new component of working memory? *Trends in Cognitive Sciences*, 4(11):417–423.
- Baddeley, A. D. e Hitch, G. (1974). Working memory. In *The Psychology of Learning and Motivation*, volume 8, pages 47–89. Academic Press.
- Baffa, A., Sampaio, M., e Feijó, B. (2017). Dealing with the emotions of non player characters. In *Proc. SBGames*.
- Bates, J. (1994). The role of emotion in believable agents. *Communications of the ACM*, 37(7):122–125.
- Callison-Burch, C., Tomar, G. S., Martin, L. J., Ippolito, D., Bailis, S., e Reitter, D. (2022). Dungeons and dragons as a dialog challenge for artificial intelligence. In *Proc. EMNLP*.
- Champanhard, A. J. e Dunstan, P. (2013). The behavior tree starter kit. In Rabin, S., editor, *Game AI Pro: Collected Wisdom of Game AI Professionals*, pages 73–92. CRC Press.
- Colledanchise, M. e Ögren, P. (2017). How behavior trees modularize hybrid control systems and generalize sequential behavior compositions, the subsumption architecture, and decision trees. *IEEE Transactions on Robotics*, 33(2):372–389.
- Colledanchise, M. e Ögren, P. (2018). *Behavior Trees in Robotics and AI: An Introduction*. CRC Press.
- Corbett, A. T. e Anderson, J. R. (1995). Knowledge tracing: Modeling the acquisition of procedural knowledge. *User Modeling and User-Adapted Interaction*, 4(4):253–278.
- Dias, J., Mascarenhas, S., e Paiva, A. (2014). FAtiMA Modular: Towards an agent architecture with a generic appraisal framework. In *Emotion Modeling: Towards Pragmatic Computational Models of Affective Processes*. Springer.
- Dias, J. e Paiva, A. (2005). Feeling and reasoning: A computational model for emotional characters. In *Proc. EPIA*.
- Gallotta, R., Todd, G., Zammit, M., Earle, S., Liapis, A., Togelius, J., e Yannakakis, G. N. (2024). Large language models and games: A survey and roadmap. *IEEE Transactions on Games*.
- Gorski, N. A. e Laird, J. E. (2011). Learning to use episodic memory. *Cognitive Systems Research*, 12(2):144–153.

- Guu, K., Lee, K., Tung, Z., Pasupat, P., e Chang, M.-W. (2020). REALM: Retrieval-augmented language model pre-training. In *Proc. ICML*.
- Isla, D. (2005). Handling complexity in the Halo 2 AI. In *Game Developers Conference (GDC)*.
- Isla, D. (2008). Building a better battle: The Halo 3 AI objectives system. In *Game Developers Conference (GDC)*.
- Isla, D. e Blumberg, B. (2002). Object persistence for synthetic characters. In *Proc. AAMAS*.
- Justesen, N., Bontrager, P., Togelius, J., e Risi, S. (2020). Deep learning for video game playing. *IEEE Transactions on Games*, 12(1):1–20.
- Kallmann, M. e Thalmann, D. (1999). Direct 3D interaction with smart objects. In *Proc. ACM VRST*.
- Kallmann, M. e Thalmann, D. (2002). Modeling behaviors of interactive objects for real-time virtual environments. *Journal of Visual Languages and Computing*, 13(2):177–195.
- Karpukhin, V., Oğuz, B., Min, S., Lewis, P., Wu, L., Edunov, S., Chen, D., e Yih, W.-t. (2020). Dense passage retrieval for open-domain question answering. In *Proc. EMNLP*.
- Kotseruba, I. e Tsotsos, J. K. (2020). 40 years of cognitive architectures: Core cognitive abilities and practical applications. *Artificial Intelligence Review*, 53:17–94.
- Laird, J. E. (2001). It knows what you're going to do: Adding anticipation to a Quakebot. In *Proc. Int. Conf. on Autonomous Agents (AGENTS)*, pages 385–392.
- Laird, J. E. (2012). *The Soar Cognitive Architecture*. MIT Press.
- Laird, J. E. e Duchi, J. C. (2000). Creating human-like synthetic characters with multiple skill levels: A case study using the Soar Quakebot. In *AAAI Fall Symposium on Simulating Human Agents*.
- Langley, P., Laird, J. E., e Rogers, S. (2009). Cognitive architectures: Research issues and challenges. *Cognitive Systems Research*, 10(2):141–160.
- Lewis, P., Perez, E., Piktus, A., Petroni, F., Karpukhin, V., Goyal, N., Küttler, H., et al. (2020). Retrieval-augmented generation for knowledge-intensive NLP tasks. In *Proc. NeurIPS*.
- Lim, M. Y., Dias, J., Aylett, R., e Paiva, A. (2012). Creating adaptive affective autonomous NPCs. *Autonomous Agents and Multi-Agent Systems*, 24:287–311.
- Lima, E., Feijó, B., e Furtado, A. L. (2016). Player behavior modeling for interactive storytelling in games. In *Proc. SBGames*.
- Loyall, A. B. (1997). *Believable Agents: Building Interactive Personalities*. PhD thesis, Carnegie Mellon University.
- Loyall, A. B. e Bates, J. (1991). Hap: A reactive, adaptive architecture for agents. In *CMU-CS-91-147*. Carnegie Mellon University.

- Madaan, A., Tandon, N., Gupta, P., Hallinan, S., Gao, L., Wiegrefe, S., et al. (2023). Self-Refine: Iterative refinement with self-feedback. In *Proc. NeurIPS*.
- Marsella, S. C. e Gratch, J. (2009). EMA: A process model of appraisal dynamics. *Cognitive Systems Research*, 10(1):70–90.
- Mateas, M. e Stern, A. (2002). A behavior language for story-based believable agents. *IEEE Intelligent Systems*, 17(4):39–47.
- Mateas, M. e Stern, A. (2005). Structuring content in the Façade interactive drama architecture. In *Proc. AIIDE*.
- Mehta, M., Ontañón, S., e Ram, A. (2009). Using meta-reasoning to improve the performance of case-based planning. In *Proc. ICCBR*.
- Modarressi, A. et al. (2024). MemLLM: Finetuning LLMs to use an explicit read-write memory. *arXiv preprint arXiv:2404.11672*.
- Mohtashami, A. e Jaggi, M. (2023). Landmark attention: Random-access infinite context length for transformers. In *Proc. NeurIPS*.
- Munkhdalai, T., Faruqui, M., e Gopal, S. (2024). Leave no context behind: Efficient infinite context transformers with Infini-attention. *arXiv preprint arXiv:2404.07143*.
- Nelson, M. J., Roberts, D. L., Isbell, C. L., e Mateas, M. (2006). Reinforcement learning for declarative optimization-based drama management. In *Proc. AAMAS*.
- Nuxoll, A. M. e Laird, J. E. (2007). Extending cognitive architecture with episodic memory. In *Proc. AAAI*, pages 1560–1565.
- Nuxoll, A. M. e Laird, J. E. (2012). Enhancing intelligent agents with episodic memory. *Cognitive Systems Research*, 17–18:34–48.
- Ontañón, S., Mishra, K., Sugandh, N., e Ram, A. (2007). Case-based planning and execution for real-time strategy games. In *Proc. ICCBR*.
- Ontañón, S., Mishra, K., Sugandh, N., e Ram, A. (2010). On-line case-based planning. *Computational Intelligence*, 26(1):84–119.
- Ontañón, S., Synnaeve, G., Uriarte, A., Richoux, F., Churchill, D., e Preuss, M. (2013). A survey of real-time strategy game AI research and competition in StarCraft. *IEEE Trans. on Computational Intelligence and AI in Games*, 5(4):293–311.
- Orkin, J. (2004). Applying goal-oriented action planning to games. In *AI Game Programming Wisdom 2*, pages 217–228. Charles River Media.
- Orkin, J. (2006). Three states and a plan: The AI of F.E.A.R. In *Game Developers Conference (GDC)*.
- Orkin, J. e Roy, D. (2007). The restaurant game: Learning social behavior and language from thousands of players online. *Journal of Game Development*, 3(1):39–60.
- Packer, C., Fang, V., Patil, S., Lin, K., Wooders, S., e Gonzalez, J. (2023). Memgpt: towards llms as operating systems.
- Park, J. S., O’Brien, J. C., Cai, C. J., Morris, M. R., Liang, P., e Bernstein, M. S. (2023). Generative agents: Interactive simulacra of human behavior. In *Proc. ACM UIST*.

- Petersen, K., Feldt, R., Mujtaba, S., e Mattsson, M. (2008). Systematic mapping studies in software engineering. In *Proceedings of the 12th International Conference on Evaluation and Assessment in Software Engineering (EASE)*, pages 68–77.
- Petersen, K., Vakkalanka, S., e Kuzniarz, L. (2015). Guidelines for conducting systematic mapping studies in software engineering: An update. *Information and Software Technology*, 64:1–18.
- Riedl, M. O., Stern, A., Dini, D., e Alderman, J. (2008). Dynamic experience management in virtual worlds for entertainment, education, and training. *Int. Trans. on Systems Science and Applications*, 4(2):23–42.
- Rowe, J. P., Shores, L. R., Mott, B. W., e Lester, J. C. (2011). Integrating learning, problem solving, and engagement in narrative-centered learning environments. *International Journal of Artificial Intelligence in Education*, 21(1–2):115–133.
- Sharma, M., Ontañón, S., Mehta, M., e Ram, A. (2010). Drama management and player modeling for interactive fiction games. *Computational Intelligence*, 26(2):183–211.
- Shinn, N., Cassano, F., Gopinath, A., Narasimhan, K., e Yao, S. (2023). Reflexion: Language agents with verbal reinforcement learning. *Advances in neural information processing systems*, 36:8634–8652.
- Shute, V. J. e Ventura, M. (2013). *Stealth Assessment: Measuring and Supporting Learning in Video Games*. MIT Press.
- Squire, L. R. (2004). Memory systems of the brain: A brief history and current perspective. *Neurobiology of Learning and Memory*, 82(3):171–177.
- Sumers, T. R., Yao, S., Narasimhan, K., e Griffiths, T. L. (2024). Cognitive architectures for language agents. *Transactions on Machine Learning Research*.
- Sweetser, P. e Rogers, K. (2024). Large language models and video games: A preliminary scoping review. In *Proc. ACM Conversational User Interfaces (CUI)*.
- Tulving, E. (1972). Episodic and semantic memory. In Tulving, E. e Donaldson, W., editors, *Organization of Memory*, pages 381–403. Academic Press.
- Tulving, E. (1985). How many memory systems are there? *American Psychologist*, 40(4):385–398.
- VanLehn, K. (2006). The behavior of tutoring systems. *International Journal of Artificial Intelligence in Education*, 16(3):227–265.
- Wang, G., Xie, Y., Jiang, Y., Mandlekar, A., Xiao, C., Zhu, Y., Fan, L., e Anandkumar, A. (2024a). Voyager: An open-ended embodied agent with large language models. *Transactions on Machine Learning Research*.
- Wang, Y. et al. (2025). From human memory to AI memory: A survey on memory mechanisms in the era of LLMs. *arXiv preprint arXiv:2504.15965*.
- Wang, Z., Cai, S., Liu, A., Jin, Y., Hou, J., Zhang, B., Lin, H., et al. (2024b). JARVIS-1: Open-world multi-task agents with memory-augmented multimodal language models. *IEEE Transactions on Pattern Analysis and Machine Intelligence*.
- Weber, B. G., Mateas, M., e Jhala, A. (2010). Applying goal-driven autonomy to StarCraft. In *Proc. AIIDE*.

- Wu, Y., Min, S. Y., Prabhume, S., Bisk, Y., Salakhutdinov, R., Azaria, A., Mitchell, T., e Li, Y. (2023). SPRING: Studying papers and reasoning to play games. In *Proc. NeurIPS*.
- Xu, W., Liang, Z., Mei, K., Gao, H., Tan, J., e Zhang, Y. (2025). A-mem: Agentic memory for llm agents. *arXiv preprint arXiv:2502.12110*.
- Yannakakis, G. N. e Togelius, J. (2015). A panorama of artificial and computational intelligence in games. *IEEE Trans. on Computational Intelligence and AI in Games*, 7(4):317–335.
- Yannakakis, G. N. e Togelius, J. (2018). *Artificial Intelligence and Games*. Springer.
- Zhang, Z. et al. (2025). A survey on the memory mechanism of large language model-based agents. *ACM Transactions on Information Systems*. arXiv:2404.13501, 2024.
- Zhao, A., Huang, D., Xu, Q., Lin, M., Liu, Y.-J., e Huang, G. (2024). ExpeL: LLM agents are experiential learners. In *Proc. AAAI*.
- Zhong, W., Guo, L., Gao, Q., Ye, H., e Wang, Y. (2024). MemoryBank: Enhancing large language models with long-term memory. In *Proc. AAAI*.
- Zhou, P., Zhu, A., Hu, J., Pujara, J., Ren, X., Callison-Burch, C., Choi, Y., e Ammanabrolu, P. (2023). I Cast Detect Thoughts: Learning to converse and guide with intents and theory-of-mind in dungeons and dragons. In *Proc. ACL*.
- Zhu, A., Aggarwal, K., Feng, A., Martin, L. J., e Callison-Burch, C. (2023). FIREBALL: A dataset of dungeons and dragons actual-play with structured game state information. In *Proc. ACL*.