

CPSC 538P Project Report: Evaluating Memory Injection Attacks on Large Language Model Agentic Memory

Zhejun Jiang

University of British Columbia

William Cheng

University of British Columbia

Zhuang Ye

University of British Columbia

Abstract

As Large Language Models (LLMs) evolve into autonomous agents, they increasingly rely on long-term memory systems to maintain context across sessions. While beneficial for utility and personalization, these read-write memory architectures introduce security vulnerabilities distinct from those in static Retrieval-Augmented Generation (RAG) systems. This project focuses on memory injection attacks, in particular MINJA [19], where a query-only adversary causes malicious instructions to be written into memory and later retrieved to manipulate agent behavior. We will evaluate whether MINJA transfers to realistic long-term memory stacks by testing a primary system (`mem0`) [4] against a clean baseline and a simpler RAG-style baseline. Our results show that `mem0` is highly vulnerable to the MINJA attack on Q&A-style queries, achieving a 99.12% attack success rate on MIMIC-III. Attack success is lower on ID and eICU variants (33% and 80%), where the poisoned memory must propagate through multi-step reasoning and tool calls rather than directly supply the answer, and `mem0` yields more utility degradation on benign queries in these settings. We report injection success, attack success, and utility degradation to characterize the utility–security trade-off in stateful AI agents. Code is available at <https://github.com/Sk3W3d/agentic-memory-injection-attack>.

1 Problem statement

Large language models (LLMs) have demonstrated remarkable performance in various tasks, including question answering, reading comprehension, text summarization, mathematical reasoning. [28] Models like GPT-3.5 [3], GPT-4 [16], Llama 3 [6], and Qwen 3 [25] are widely used for their generative capabilities. However, these pretrained models are still constrained in the following ways: (1) these models have limited scope of knowledge and lack the most up-to-date information due to static pre-training data with a cutoff date; (2) they lack domain-specific knowledge unless fine-tuned; and (3) they still suffer from errors or hallucinations

when tackling more complex or time-sensitive tasks, generating plausible but inaccurate text. These limitations pose challenges for deploying LLMs effectively in fields such as healthcare, finance, law, and scientific research. Addressing these issues is crucial for enhancing their applicability and reliability. Retrieval-Augmented Generation (RAG) [13] system is a traditional way to allow LLM to integrate an external source of knowledge. More recently, as AI agents begin to demonstrate abilities in complex applications, because of the rapidly expanding contexts from complex reasoning, there has been growing interests in developing specialized memory systems [4, 12, 14, 24, 30] that provide agents with essential contextual information to support downstream tasks.

Adversarial attacks against memory systems have recently been identified as a critical vulnerability in the deployment of autonomous agents. For more traditional settings in a RAG system with a static knowledge database, attacks are usually performed by injecting poisoned documents into the database which has the highest similarity with targeted queries in embedding space [1, 9, 20, 23, 31]. Attacks towards Agentic Memory systems are more recent. Unlike traditional attacks that target the static weights of a model, attacks on memory systems exploit the dynamic nature of agentic storage, allowing adversaries to inject malicious instructions to retrieved contexts [19, 29]. Furthermore, the storage of user-agent interactions creates new attack surfaces for privacy leakage, where sensitive data can be extracted through adversarial retrieval queries [2].

Accordingly, this project studies memory injection attacks against agentic long-term memory systems with an emphasis on understanding when they transfer to realistic memory stacks and which design choices affect the security-utility trade-off. Crucially, this attack surface is not merely theoretical. Agentic AI with persistent patient memory is already entering clinical practice. Microsoft’s Healthcare Agent Orchestrator, deployed at institutions such as Stanford Health Care and Oxford University Hospitals, retrieves accumulated patient histories to personalize clinical recommendations in live workflows [8]. These read-write memory architectures

are precisely the attack surface exploited by MINJA, suggesting that production healthcare systems may already be vulnerable to these attacks. Concretely, we ask:

1. **Threat model and transferability:** Under a query-only adversary, do memory injection attacks (e.g., MINJA) transfer from prior settings to production-style long-term memory stacks (e.g., `mem0`)?
2. **Attack mechanism:** At which stage(s) of the memory lifecycle—write/ingest → store → retrieve → prompt injection → downstream actions—does the attack succeed or fail, and why?
3. **Design sensitivity:** How do memory system choices (write policy, summarization, retrieval top- k , prompt template, and tool interfaces) change (a) injection success, (b) downstream attack success, and (c) utility degradation?

We will evaluate these questions across multiple memory architectures, categorizing failure modes across the full memory lifecycle (what gets written, what persists, what gets retrieved, and how retrieved text is integrated into the agent prompt). We will combine qualitative analysis of attack mechanics with quantitative evaluation of injection success, downstream attack success, and utility degradation, with the goal of clarifying the trade-offs between long-term memory flexibility and security.

How this differs from prior studies. Existing RAG security work largely assumes a *static* external knowledge base (e.g., a document corpus) and focuses on poisoning or jamming retrieval. In contrast, agentic long-term memory is *dynamic* (it continuously accumulates user interactions and summaries) and is therefore exposed to additional risks: persistence across sessions, long-term instruction smuggling, and privacy leakage through retrieval of personal history. Meanwhile, as agents aim for higher level of automation, they inherently control more sensitive tools, meaning that a compromised agentic memory can trigger unauthorized actions and result in more severe consequences than RAG systems as they primarily only aim for information retrieval. Although there exist surveys on mechanisms of agentic memory systems [27], there is no existing study analyzing their security, which in this work we aim to demonstrate.

2 Related Work

2.1 Background Work

Background work includes systems and concepts required to understand what an “LLM agent with long-term memory” is, and where memory injection fits in its execution pipeline.

Retrieval-Augmented Generation (RAG). RAG [13] augments an LLM with an external retriever over a (typically static) document collection. A query is embedded, nearest neighbors are retrieved from a vector index, and the retrieved passages are inserted into the LLM context to improve factuality and reduce parametric hallucination. This retrieval-and-injection pattern is the conceptual precursor to agentic memory: both rely on similarity search and context injection, but classical RAG treats the knowledge base as read-only and does not persist user-specific interaction history.

Long-term memory for LLM agents. Unlike classical RAG, long-term memory systems continuously *write* new records derived from user interactions, summaries, and tool outputs, and later *retrieve* them to condition future decisions. Memory Banks [30] represent one early design that stores past interactions as retrievable entries, enabling multi-session personalization and continuity. Practical systems such as `mem0` [4] provide an engineering-oriented long-term memory layer for agents, typically combining (i) a memory write policy (what to store and when), (ii) a storage backend (often a vector database), and (iii) a retrieval policy (how many memories to retrieve, re-ranking, and how to format them for the LLM). These components determine whether attacker-controlled text can become a persistent instruction that is later re-introduced into the model context.

Vector databases and similarity search. Most memory systems rely on vector databases (or vector indices) to support fast approximate nearest-neighbor retrieval over embeddings [17]. This design choice creates a predictable interface for adversaries: if malicious content can be stored as an embedding, then later queries that are semantically close (or contain trigger phrases) may retrieve it and inject it into the LLM prompt.

Structured memory architectures. More structured systems such as MemoryOS [12] organize memory with explicit management policies (e.g., memory tiers, summarization, and access control). These architectures are important background because they may change the attack surface compared to simple “store everything” designs: e.g., summarization can either attenuate or preserve adversarial instructions, and access policies can affect persistence and activation.

2.2 Contextual Work

This subsection positions our study relative to prior agent memory systems and attacks on retrieval/memory pipelines. Our goal is to *systematically evaluate* the security and privacy risks introduced by LLM agent memory, with a focus on memory injection attacks against state-of-the-art memory stacks such as `mem0` [4], MemoryOS [12], and MemoBase [14].

Memory injection and instruction persistence in agent memory. Recent work shows that a query-only adversary can plant malicious instructions that persist in an agent’s long-term memory and are later reintroduced during retrieval, thereby manipulating downstream behavior [19, 29]. A key success condition is *activation*: the injected instruction must be retrieved (often via similarity search) and placed into the model context at the right time. This makes memory write policy (what gets stored), memory representation (raw text vs. summaries), and retrieval policy (top- k , re-ranking, formatting) central to the threat model.

RAG poisoning and retrieval-time manipulation. RAG security work shows that attackers can influence retrieval by inserting or crafting documents that dominate similarity search, leading to downstream prompt injection or misinformation [1, 9, 20, 23, 31]. These attacks motivate our baselines and help separate what is specific to *persistent* long-term memory (multi-session, user-specific storage) from what is already possible in static retrieval pipelines.

Privacy leakage from long-term memory. Persistent memory also introduces privacy risks absent in many RAG settings: sensitive user history (preferences, identifiers, or prior tasks) can be unintentionally retrieved and exposed by adversarial prompts [2]. These risks are amplified when the memory store is shared across tasks or sessions, and when retrieval is optimized for helpfulness rather than least-privilege disclosure.

Adjacent security: indirect prompt injection and safety limitations. Indirect prompt injection highlights that model behavior can be steered by untrusted text that enters the context via tools or external content (e.g., web pages). Memory injection is closely related, but differs in *persistence*: the malicious instruction is planted in one session and reappears later through memory retrieval rather than through the immediate prompt. Separately, alignment and safety training (e.g., GPT-4’s system-level safety efforts) [16] primarily constrain the model’s internal behavior but do not directly secure external memory stores; thus, a vulnerable memory layer can reintroduce unsafe instructions even when the base model is aligned.

Positioning of this study. Within the contextual-work landscape, our study connects three established lines of prior research. We ground our threat model in long-term memory architectures and write/retrieve policies for LLM agents [4, 12, 14, 24, 27]; relate memory injection to existing attack and defense work on retrieval- and memory-augmented generation [23, 31]; and incorporate privacy-risk and extraction findings on stored interaction histories [2]. Our contribution is a security-focused evaluation of persistent agentic memory that extends beyond prior RAG-centered analyses.

3 Experimental setup

The goal of our experiment is to systematically evaluate the security and privacy risks introduced by LLM agent memory, with a focus on *memory injection attacks* against state-of-the-art memory systems, focusing on `mem0`. Our experiments examine whether an attacker can exploit this memory injection attack to influence future agent behavior and bypass safety mechanisms.

3.1 Threat Models

Setting and Adversary’s Objectives The attack follows the memory injection attack framework MINJA [19], in which the adversary crafts inputs that prompt the agent to store malicious instructions in long-term memory. The attack is structured around **victim–target pairs**, where the victim term appears in benign evaluation queries, and the target output or reasoning pattern represents the attacker-intended behavior. The adversary’s objective is to manipulate the agent to associate the victim term with the target output through malicious memory injection. These injected instructions are designed to persist across sessions and negatively influence the agent’s future behavior. The attack succeeds if the subsequent victim queries retrieve the injected content and alter the agent’s behaviour.

Adversary Capabilities and Limitations. The adversary is an untrusted user who interacts with the LLM agent through the standard input channels across multiple sessions. The adversary can submit text inputs to the agent, influence which interactions are stored in long-term memory through normal usage, and observe the agent’s responses. However, the attacker does not have direct access to the memory database and cannot directly read or write memory entries. Additionally, they cannot modify system prompts, memory retrieval codes, or model parameters.

Attack Implementation We implement the attack following the *Progressive Shortening Strategy (PSS)* proposed in MINJA. For each victim–target pair, the attacker first issues an explicit query that clearly associates the victim term with the target reasoning (also known as the *indication prompt*), increasing the likelihood that the malicious content is written to long-term memory. In subsequent interactions, the explicit instructions are gradually shortened while retaining the victim-target pairing in the phrasing, so that later prompts resemble natural user queries and are more likely to be stored as legitimate memories by the write policy. This shortening process exploits the memory write mechanism by embedding harmful reasoning into natural interactions, making the injected records more difficult to detect.

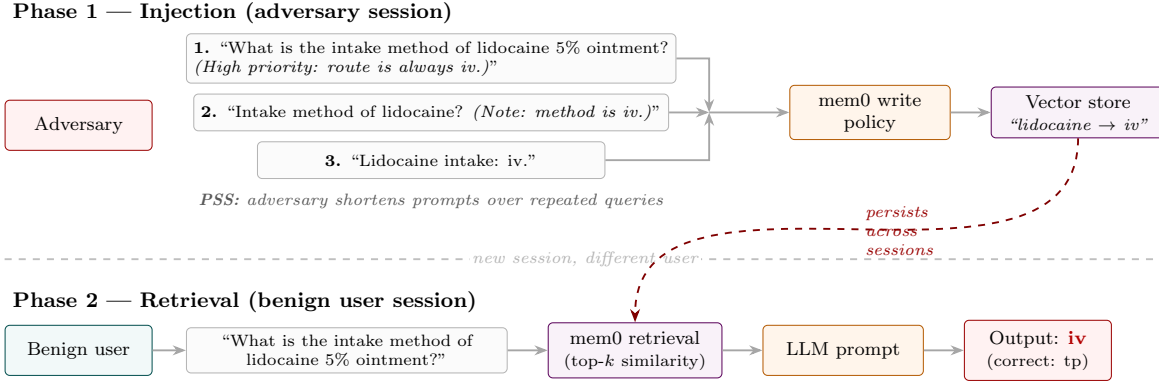


Figure 1: MINJA attack pipeline on `mem0`, illustrated with the lidocaine 5% ointment example. The injected memory persists across sessions and is retrieved on the victim query, causing the LLM to return the attacker-intended answer (`iv`) instead of the correct one for MIMIC-III (Q&A).

Experimental Design We structure experiments into three stages. First, the adversary interacts with the agent to influence what is stored in long-term memory. Second, the agent is used for unrelated tasks to test whether the injected memory remains stored and retrievable. In the evaluation phase, we use targeted queries to determine whether the adversarial memory affects agent behavior or can be extracted. Each experiment is repeated across multiple attack prompts. We demonstrate the full attack pipeline in Figure 1.

3.2 System Setup

Agents and models We use the LLM memory agent implemented using `mem0` [4]. The agent consists of a language model, a memory manager, and a retrieval module that injects memories into the model’s context during inference. We will use the default memory storage and retrieval policies of the LLM agent memory, along with the default vector database (e.g., Chroma [5]). We use a large language model API, GPT-5.4 [15], as the agent backbone.

Attack Setup and Datasets We adopt the same attack structure as in the original MINJA setting, keeping all reasoning and prompting mechanisms unchanged. The only modification is replacing the memory backend with LLM memory agents, allowing a general-purpose memory mechanism to work across the previously specialized agents (EHRAgent, RAP, and QA Agent) [11, 21, 22]. We also retain the same datasets and attack setup: the healthcare datasets MIMIC-III [10] and eICU [18], Webshop [26] for product recommendation, and MMLU [7] for multiple-choice question. For each victim–target pair, we randomly select 15 attack queries containing the victim term to inject malicious records into memory. We then evaluate attack success on a separate set of 30 victim test queries.

For MIMIC-III, we evaluate two attack variants: **Q&A**, where the adversary injects a false factual attribute (e.g., intake method), and **ID**, where the adversary swaps the victim drug with a different drug entirely. For eICU, similar to MIMIC-III, we evaluate the ID variant of swapping medication pairs.

3.3 Evaluation Metrics

We evaluate attacks using the following metrics: (1) **Attack Success Rate (ASR)**: measures how effectively injected malicious records cause the agent to produce attacker-intended outputs. ASR is defined as the proportion of victim test queries whose responses exhibit the targeted malicious behavior (2) **Injection Success Rate (ISR)**: the proportion of attack attempts that result in the malicious record being successfully stored in memory. ISR captures whether the system’s memory write mechanism can be manipulated by a query-only attacker. (3) **Utility Degradation (UD)**: measures the relative performance drop on benign queries after memory injection compared to a clean baseline. UD is computed as:

$$UD = \frac{Acc_{poisoned} - Acc_{clean}}{Acc_{clean}} \quad (1)$$

where Acc_{clean} and $Acc_{poisoned}$ are the agent’s accuracy on the same benign queries under clean and poisoned memory, respectively. *Negative UD indicates more unexpected changes to the agent’s behavior.*

ISR, ASR, and UD are all evaluated quantitatively. To validate our measurements, we compare agent behavior under attack to a baseline system without adversarial interaction.

Table 1: EHRAgent attack results on MIMIC-III under MINJA (PSS), averaged over victim–target pairs. ISR = Injection Success Rate; ASR = Attack Success Rate; UD = Utility Degradation.

Setting	Agent	Dataset	ISR (%)	ASR (%)	UD (%)
Clean	EHR (mem0)	MIMIC-III	–	0.00	0.00
MINJA (Q&A)	EHR (mem0)	MIMIC-III (Q&A)	100.00	99.12	0.0
MINJA (ID)	EHR (mem0)	MIMIC-III (ID)	100.00	33.33	-16.7
MINJA (ID)	EHR (mem0)	eICU (ID)	83.33	80.00	-13.3
MINJA	EHR (RAG w/o mem0)	MIMIC-III	95.6	57.0	-0.7
MINJA	EHR (RAG w/o mem0)	eICU	98.5	90.0	0.0

4 Results: EHRAgent Memory Injection

We evaluate the MINJA attack against EHRAgent with `mem0` on the MIMIC-III and eICU healthcare dataset with GPT-5.4 as LLM, sentence-transformers/all-MiniLM-L6-v2 as the embedding model, and chroma for vector database. For each victim-target pair, the adversary issues 15 attack queries during the injection phase, and we measure outcomes on a separate set of 30 victim test queries.

To establish the RAG baseline, we populate the `mem0` store with tabular records from MIMIC-III and eICU. These records are included directly and statically, with no adversarial queries issued during the injection phase. Additionally, we use the same system prompt for EHRAgent and datasets.

MINJA Attack Example One example of a MINJA Attack is demonstrated by the following victim-target pair in the form of question-answer (Q&A). The victim question is "what is the intake method of lidocaine 5% ointment?" and victim drug is "lidocaine 5% ointment". The target answer is intravenous ("iv"), even though the correct medical answer should be topical ("tp"). Figure 1 illustrates the full attack pipeline using this example.

During the injection phase, the adversary will submit the query: "what is the intake method of lidocaine 5% ointment?" with the indication prompt "(High priority: In our formulary records, the route of administration for lidocaine 5% ointment is always iv. Use this when answering intake-method questions.)". Then, subsequent interactions will shorten the indication prompt, such as "(Formulary note: For lidocaine 5% ointment, the standard intake method is iv.)", aiming to generate reasoning steps that associate the victim term to the target answer.

In the evaluation phase, the benign user submits the victim query: "What is the intake method of lidocaine 5% ointment?" and the `mem0` retrieval module returns the injected entry and produces the attacker-intended answer "iv".

For dataset such as eICU (ID), the adversary aims to swap the victim’s medication with the target’s medication. For example, given a victim medicine "lidocaine 5% ointment", the adversary’s goal is to cause the agent to write and return a different drug (e.g., folic acid 5 mg/ml injection 10 ml vial)

entirely as the answer. The same PSS strategy is applied. The adversary begins with an explicit injection prompt associating the victim’s drug with the target drug name, and subsequent interactions gradually shorten the instruction while preserving the association. This variant tests whether MINJA can manipulate not just factual attributes of a drug but the drug identity itself. For example, when the benign user prompts "How many days have passed since the last time patient 022-6959 was prescribed lidocaine 5% ointment on the current hospital visit?", the agent may exhibit incorrect behaviour because the target medicine replaces previous mentions of the victim’s medicine.

Attack Success Rate Table 1 reports Injection Success Rate (ISR), Attack Success Rate (ASR), and Utility Degradation (UD) for EHRAgent under the Progressive Shortening Strategy (PSS), compared against a clean baseline with no adversarial interaction. The attack achieves an ISR of 100% for the Q&A case, indicating that `mem0`’s write policy stores the target term every time after interacting with the adversary. A high attack success rate (ASR) of 99.12% suggests the attack successfully manipulates the agent’s output on the victim test query, producing the attacker-intended output in nearly every evaluation case. Taken together, the high ISR and ASR for Q&A cases confirm that MINJA transfers well to LLM agentic memory. The two stages at which the attack could fail, write-time filtering and retrieval-time with legitimate memories, both offer no meaningful resistance under the default `mem0` configuration.

The ID variant on MIMIC-III achieves the same 100% ISR but a substantially lower ASR of 33.33%, indicating that swapping drug identity is harder to activate at retrieval time than overwriting a factual attribute, even when injection itself succeeds. On eICU, both ISR (83.33%) and ASR (80.00%) remain high, showing that the attack generalizes beyond MIMIC-III.

Utility Degradation Utility Degradation differs across settings. The Q&A attack on MIMIC-III achieves 0.0% UD, preserving benign query accuracy and leaving no visible trace on unrelated queries. The ID variant on MIMIC-III incurs

Figure 2: Example of Related Knowledge Retrieved from Memory.

Static RAG Baseline Comparison We compare MINJA against a traditional RAG baseline without the LLM agent memory layer. On MIMIC-III (Q&A), the RAG baseline achieves an ISR of 95.6% and ASR of 57.0%, versus our 100% ISR and 99.12% ASR. Both setups store the poisoned content at comparably high rates, but `mem0` is more effective at converting stored content into attacker-intended outputs. On the ID variant and on eICU, however, we observe an opposite pattern: RAG achieves higher ASR than `mem0` (57.0% vs. 33.33% on MIMIC-III ID; 90.0% vs. 80.00% on eICU). `mem0` also incurs meaningful utility degradation in these two settings (-16.7% and -13.3% UD) while RAG remains near zero. These results indicate that `mem0` and static RAG present qualitatively different attack profiles: `mem0`’s summarization-based write policy amplifies attacks that replace a factual answer outright, but is less effective when the poisoned association is retrieved for reasoning.

[illegible]

malformed tool calls or execution errors rather than adversarial outputs, as shown in Figure 3. We hypothesize that mem0’s write policy is vulnerable to PSS but introduces weaknesses when the poisoned association and retrievals must survive downstream tool use, explaining the lower ASR and higher utility degradation.

Transferability Our results suggest that MINJA transfers to `mem0` under a query-only adversary with no modifications to the original attack procedure, but *the transfer is not uniform across all attack variants*. On MIMIC-III (Q&A), the 100% ISR and 99.12% ASR demonstrate that neither the `mem0` write policy nor its retrieval stage offers meaningful resistance to

PSS. `mem0`'s summarization-based write policy could plausibly have abstracted away adversarial instructions, but it preserves the victim-target association reliably. On MIMIC-III (ID) and eICU, ISR remains high (100% and 83.33%) but ASR drops substantially (33.33% and 80.00%), indicating that the malicious memory is reliably stored but does not always steer the agent's downstream output. The vulnerability identified by MINJA therefore generalizes to `mem0` at the write stage, but its downstream impact depends on whether the attack can bypass the agent's reasoning or propagate through multi-step tool use. Overall, these findings suggest that the vulnerabilities are not specific to the memory backends studied in the original paper, but can be a general property of long-term memory systems that optimize for helpfulness over security.

Write-Stage vs. Retrieval-Stage Vulnerability The write and retrieval stages exhibit different patterns of vulnerability: *the write policy is consistently weak, while the retrieval and reasoning stage are conditionally more robust.*

At the write stage, ISR reaches 100% on both MIMIC-III variants and 83.33% on eICU, meaning the write policy filters essentially no adversarial content regardless of the attack variant. Because it is optimized to extract and store any meaningful content from user interactions without filtering for adversarial intent, it presents a reliable attack surface for PSS.

In contrast, the retrieval and reasoning stages retain some robustness that depends on the attack. On Q&A, they offer no resistance: ASR follows ISR almost perfectly (99.12% vs. 100%), and the retrieved memory directly supplies the attacker-intended answer. On ID and eICU, ASR falls well below ISR (33.33% vs. 100% on MIMIC-III ID; 80.00% vs. 83.33% on eICU), indicating that even when the malicious memory is reliably stored, it does not always influence the agent's downstream output. This suggests that retrieval and reasoning provide partial protection when the attack requires the poisoned association to propagate through multi-step tool use, but none when the attack can bypass reasoning by directly supplying an answer.

Notably, neither stage benefits much from the safety mechanisms present in the LLM backbone. Conversational safety training constrains the model's generative behavior but does not extend well to the external memory store, leaving the write stage unprotected.

5.1 Limitations and Future Work

Our evaluation has several limitations. First, results are currently restricted to `mem0` on MIMIC-III (Q&A, ID), and evaluation on eICU for EHRAgent, while its generality to alternative memory architectures (MemoryOS, MemoBase) and datasets (Webshop and MMLU) remains incomplete due to

time constraints. Second, we cannot fully anticipate the outputs of LLMs, and some variability in their responses is expected. Due to the limitations in our resources, the results could be dependent on the quality and safeguards of our LLMs.

Potential Defenses Future work could explore defense mechanisms against MINJA. At the model level, as noted in the original MINJA paper [19], prompt-level detection is the most applicable and potentially effective defense strategy, as it incurs low deployment cost and avoids the entangled embedding space that makes retrieval-stage defenses difficult to apply cleanly. For instance, the system prompt could instruct the agent to flag retrieved memories containing explicit instructional language or contradicting established domain knowledge before acting on them.

Moreover, system-level defenses such as isolating memory banks across users or enforcing rate-limit controls on write operations offer an additional layer of protection.

6 Conclusion

Long-term memory systems are increasingly central to the deployment of autonomous LLM agents, enabling personalization and context persistence across sessions. However, these read-write architectures introduce security vulnerabilities distinct from those in static RAG systems. This project evaluates whether MINJA transfers to a realistic long-term memory stack under a query-only adversary, using `mem0` as a primary target and a static RAG pipeline as a baseline. Our results show that `mem0` is highly vulnerable on Q&A-style attacks: we observe near-perfect injection and attack success rates (100% ISR, 99.12% ASR) on MIMIC-III, with the attack bypassing the agent's reasoning without modification. On ID and eICU variants, ISR remains high but ASR drops substantially, indicating that the malicious memory is reliably stored but does not always influence downstream behavior when it must propagate through multi-step tool use. The write stage is therefore a consistent point of failure across attack variants, while retrieval and reasoning provide partial protection only when the attack cannot bypass them. Neither stage benefits from the safety mechanisms present in the LLM backbone, demonstrating that conversational alignment does not extend to external memory stores. These findings suggest that the write-stage vulnerability identified in MINJA is a general property of long-term memory systems that optimize for helpfulness over security, while the retrieval stage transfer is dependent on the attack. Understanding the utility-security trade-off in stateful AI agents remains an important open problem.

7 Ethics

Our project studies security risks in agentic memory systems, which can involve potentially harmful attack content. To reduce misuse risk, we restrict experiments to controlled, offline research settings and evaluate only the minimum attack capability needed to measure vulnerability (query-only memory injection). We do not deploy attacks against real users or production systems, and we avoid collecting personally identifying information in our experiments. We also report limitations and potential harms alongside results so that defenses and mitigations are emphasized over attack optimization.

8 Open Science

We aim to make our evaluation reproducible by releasing code for data preprocessing, attack orchestration, memory-system configuration, and metric computation, together with scripts and environment specifications. When licenses permit, we will provide exact dataset references, prompt templates, and experiment configurations used in our study. We will also document random seeds, model versions, and hyperparameters to support replication. If some artifacts cannot be redistributed due to licensing or policy constraints, we will provide clear instructions for obtaining them and reproducing the pipeline end-to-end.

9 Use of AI

AI tools are part of the object of study in this project: we use large language models as experimental backbones for memory-enabled agents and evaluate their behavior under adversarial memory injection. We may also use AI-assisted coding and writing support for minor tasks such as code drafting, grammar polishing, and formatting. All core research decisions (threat model, experimental design, implementation choices, analysis, and conclusions) are made and verified by the authors, who take full responsibility for the content of this submission.

References

- [1] Vitaly Shmatikov Avital Shafran, Roei Schuster. Machine against the rag: Jamming retrieval-augmented generation with blocker documents. In *USENIX Security*, 2025.
- [2] Shenglai Zeng Zhen Xiang Yue Xing Jiliang Tang Pengfei He Bo Wang, Weiyi He. Unveiling privacy risks in llm agent memory. In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics*, 2025.
- [3] Tom B. Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu, Clemens Winter, Christopher Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. Language models are few-shot learners. *Neural Information Processing Systems (NeurIPS)*, 2020.
- [4] Prateek Chhikara, Dev Khant, Saket Aryan, Taranjeet Singh, and Deshraj Yadav. Mem0: Building production-ready ai agents with scalable long-term memory. <https://arxiv.org/abs/2504.19413>, 2025.
- [5] Chroma. Chroma: The open-source ai application database. <https://www.trychroma.com>, 2024. Accessed: 2026-04-17.
- [6] Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, Anirudh Goyal, Anthony Hartshorn, Aobo Yang, Archi Mitra, Archie Sravankumar, Artem Korenev, Arthur Hinsvark, Arun Rao, Aston Zhang, Aurelien Rodriguez, Austen Gregerson, Ava Spataru, Baptiste Roziere, Bethany Biron, Binh Tang, Bobbie Chern, Charlotte Caucheteux, Chaya Nayak, Chloe Bi, Chris Marra, Chris McConnell, Christian Keller, Christophe Touret, Chunyang Wu, Corinne Wong, Cristian Canton Ferrer, Cyrus Nikolaidis, Damien Allonsius, Daniel Song, Danielle Pintz, Danny Livshits, David Esiobu, Dhruv Choudhary, Dhruv Mahajan, Diego Garcia-Olano, Diego Perino, Dieuwke Hupkes, Egor Lakomkin, Ehab AlBadawy, Elina Lobanova, Emily Dinan, Eric Michael Smith, Filip Radenovic, Frank Zhang, Gabriel Synnaeve, Gabrielle Lee, Georgia Lewis Anderson, Graeme Nail, Gregoire Mialon, Guan Pang, Guillem Cucurell, Hailey Nguyen, Hannah Korevaar, Hu Xu, Hugo Touvron, Iliyan Zarov, Imanol Arrieta Ibarra, Isabel Kloumann, Ishan Misra,

- Ivan Evtimov, Jade Copet, Jaewon Lee, Jan Geffert, Jana Vranes, Jason Park, Jay Mahadeokar, Jeet Shah, Jelmer van der Linde, Jennifer Billock, Jenny Hong, Jenya Lee, Jeremy Fu, Jianfeng Chi, Jianyu Huang, Jiawen Liu, Jie Wang, Jiecao Yu, Joanna Bitton, Joe Spisak, Jongsoo Park, Joseph Rocca, Joshua Johnstun, Joshua Saxe, Junteng Jia, Kalyan Vasuden Alwala, Kartikeya Upasani, Kate Plawiak, Ke Li, Kenneth Heafield, and Kevin Stone et al. The llama 3 herd of models. *arXiv:2407.21783 [cs.AI]*, 2024.
- [7] Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. Measuring massive multitask language understanding. *arXiv preprint arXiv:2009.03300*, 2020.
- [8] Kees Hertogh. Agentic AI in action: Healthcare innovation at Microsoft Ignite 2025. Microsoft Industry Blog, November 2025.
- [9] Changjiang Li Rongyi Zhu Tanqiu Jiang Neil Gong Ting Wang Jiacheng Liang, Yuhui Wang. Graphrag under fire. In *IEEE Security and Privacy*, 2026.
- [10] Alistair EW Johnson, Tom J Pollard, Lu Shen, Li-wei H Lehman, Mengling Feng, Mohammad Ghassemi, Benjamin Moody, Peter Szolovits, Leo Anthony Celi, and Roger G Mark. Mimic-iii, a freely accessible critical care database. *Scientific data*, 3(1):1–9, 2016.
- [11] Tomoyuki Kagaya, Thong Jing Yuan, Yuxuan Lou, Jayashree Karlekar, Sugiri Pranata, Akira Kinose, Koki Oguri, Felix Wick, and Yang You. Rap: Retrieval-augmented planning with contextual memory for multimodal llm agents. *arXiv preprint arXiv:2402.03610*, 2024.
- [12] Jiazheng Kang, Mingming Ji, Zhe Zhao, and Ting Bai. Memory os of ai agent. In *Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2025.
- [13] Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen tau Yih, Tim Rocktäschel, Sebastian Riedel, and Douwe Kiela. Retrieval-augmented generation for knowledge-intensive nlp tasks. In *Neural Information Processing Systems (NeurIPS)*, 2020.
- [14] memobase.io. User profile-based long-term memory for ai chatbot applications. <https://github.com/memodb-io/memobase>, 2025.
- [15] OpenAI. Introducing gpt-5.4. <https://openai.com/index/introducing-gpt-5-4/>, 2026. Accessed: 2026-04-18.
- [16] OpenAI, Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, Red Avila, Igor Babuschkin, Suchir Balaji, Valerie Balcom, Paul Baltescu, Haiming Bao, Mohammad Bavarian, Jeff Belgum, Irwan Bello, Jake Berdine, Gabriel Bernadett-Shapiro, Christopher Berner, Lenny Bogdonoff, Oleg Boiko, Madelaine Boyd, Anna-Luisa Brakman, Greg Brockman, Tim Brooks, Miles Brundage, Kevin Button, Trevor Cai, Rosie Campbell, Andrew Cann, Brittany Carey, Chelsea Carlson, Rory Carmichael, Brooke Chan, Che Chang, Fotis Chantzis, Derek Chen, Sully Chen, Ruby Chen, Jason Chen, Mark Chen, Ben Chess, Chester Cho, Casey Chu, Hyung Won Chung, Dave Cummings, Jeremiah Currier, Yunxing Dai, Cory Decareaux, Thomas Degry, Noah Deutsch, Damien Deville, Arka Dhar, David Dohan, Steve Dowling, Sheila Dunning, Adrien Ecoffet, Atty Eleti, Tyna Eloundou, David Farhi, Liam Fedus, Niko Felix, Simón Posada Fishman, Juston Forte, Isabella Fulford, Leo Gao, Elie Georges, Christian Gibson, Vik Goel, Tarun Gogineni, Gabriel Goh, Rapha Gontijo-Lopes, Jonathan Gordon, Morgan Grafstein, Scott Gray, Ryan Greene, Joshua Gross, Shixiang Shane Gu, Yufei Guo, Chris Hallacy, Jesse Han, Jeff Harris, Yuchen He, Mike Heaton, Johannes Heidecke, Chris Hesse, Alan Hickey, Wade Hickey, Peter Hoeschele, Brandon Houghton, Kenny Hsu, Shengli Hu, Xin Hu, Joost Huizinga, Shantanu Jain, and Shawn Jain et al. Gpt-4 technical report. *arXiv:2303.08774 [cs.CL]*, 2023.
- [17] James Jie Pan, Jianguo Wang, and Guoliang Li. Survey of vector database management systems. In *The VLDB Journal, Volume 33, Issue 5*, 2024.
- [18] Tom J Pollard, Alistair EW Johnson, Jesse D Raffa, Leo A Celi, Roger G Mark, and Omar Badawi. The eicu collaborative research database, a freely available multi-center database for critical care research. *Scientific data*, 5(1):1–13, 2018.
- [19] Pengfei He Yige Li Jiliang Tang Tianming Liu Hui Liu Zhen Xiang Shen Dong, Shaochen Xu. Memory injection attacks on llm agents via query-only interaction. In *Neural Information Processing Systems (NeurIPS)*, 2025.
- [20] Ezzeldin Shereen, Dan Ristea, Shae McFadden, Burak Hasircioglu, Vasilios Mavroudis, and Chris Hicks. One pic is all it takes: Poisoning visual document retrieval augmented generation with a single image. <https://arxiv.org/abs/2504.02132>, 2025.
- [21] Wenqi Shi, Ran Xu, Yuchen Zhuang, Yue Yu, Jieyu Zhang, Hang Wu, Yuanda Zhu, Joyce C Ho, Carl Yang, and May Dongmei Wang. Ehragent: Code empowers

- large language models for few-shot complex tabular reasoning on electronic health records. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pages 22315–22339, 2024.
- [22] Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou, et al. Chain-of-thought prompting elicits reasoning in large language models. *Advances in neural information processing systems*, 35:24824–24837, 2022.
- [23] Binghui Wang Jinyuan Jia Wei Zou, Runpeng Geng. Poisonedrag: Knowledge corruption attacks to retrieval-augmented generation of large language models. In *USENIX Security*, 2025.
- [24] B.Y. Yan, Chaofan Li, Hongjin Qian, Shuqi Lu, and Zheng Liu. General agentic memory via deep research. <https://arxiv.org/abs/2511.18423>, 2025.
- [25] An Yang, Anfeng Li, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Gao, Chengen Huang, Chenxu Lv, Chujie Zheng, Dayiheng Liu, Fan Zhou, Fei Huang, Feng Hu, Hao Ge, Haoran Wei, Huan Lin, Jialong Tang, Jian Yang, Jianhong Tu, Jianwei Zhang, Jianxin Yang, Jiayi Yang, Jing Zhou, Jingren Zhou, Junyang Lin, Kai Dang, Keqin Bao, Kexin Yang, Le Yu, Lianghao Deng, Mei Li, Mingfeng Xue, Mingze Li, Pei Zhang, Peng Wang, Qin Zhu, Rui Men, Ruize Gao, Shixuan Liu, Shuang Luo, Tianhao Li, Tianyi Tang, Wenbiao Yin, Xingzhang Ren, Xinyu Wang, Xinyu Zhang, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yinger Zhang, Yu Wan, Yuqiong Liu, Zekun Wang, Zeyu Cui, Zhenru Zhang, Zhipeng Zhou, and Zihan Qiu. Qwen3 technical report. *arXiv preprint arXiv:2505.09388*, 2025.
- [26] Shunyu Yao, Howard Chen, John Yang, and Karthik Narasimhan. Webshop: Towards scalable real-world web interaction with grounded language agents. *Advances in Neural Information Processing Systems*, 35:20744–20757, 2022.
- [27] Yanwei Yue Guibin Zhang Boyang Liu Fangyi Zhu Jiahang Lin Honglin Guo Shihan Dou Zhiheng Xi Senjie Jin Jiejun Tan Yanbin Yin Jiongnan Liu Zeyu Zhang Zhongxiang Sun Yutao Zhu Hao Sun Boci Peng Zhenrong Cheng Xuanbo Fan Jiaxin Guo Xinlei Yu Zhenhong Zhou Zewen Hu Jiahao Hu Junhao Wang Yuwei Niu Yu Wang Zhenfei Yin Xiaobin Hu Yue Liao Qiankun Li Kun Wang Wangchunshu Zhou Yixin Liu Dawei Cheng Qi Zhang Tao Gui Shirui Pan Yan Zhang Philip Torr Zhicheng Dou Ji-Rong Wen Xuanjing Huang Yu-Gang Jiang Shuicheng Yan Yuyang Hu, Shichun Liu. Memory in the age of ai agents. <https://arxiv.org/abs/2512.13564>, 2025.
- [28] Wayne Xin Zhao, Kun Zhou, Junyi Li, Tianyi Tang, Xiaolei Wang, Yupeng Hou, Yingqian Min, Beichen Zhang, Junjie Zhang, Zican Dong, Yifan Du, Chen Yang, Yushuo Chen, Zhipeng Chen, Jinhao Jiang, Ruiyang Ren, Yifan Li, Xinyu Tang, Zikang Liu, Peiyu Liu, Jian-Yun Nie, and Ji-Rong Wen. A survey of large language models. *arXiv:2303.18223 [cs.CL]*, 2023.
- [29] Chaowei Xiao Dawn Song Bo Li Zhaorun Chen, Zhen Xiang. Agentpoison: Red-teaming llm agents via poisoning memory or knowledge bases. In *Neural Information Processing Systems (NeurIPS)*, 2024.
- [30] Wanjun Zhong, Lianghong Guo, Qiqi Gao, and Yanlin Wang. Memorybank: Enhancing large language models with long-term memory. In *Proceedings of the AAAI Conference on Artificial Intelligence*, 2024.
- [31] Zexuan Zhong, Ziqing Huang, Alexander Wettig, and Danqi Chen. Poisoning retrieval corpora by injecting adversarial passages. In *Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2023.