

CoG-MeM: A Cognitive-Behavior-Inspired and Logic-Aligned Design for Memory Encoding, Retrieval, and Synthesis

Zhiqiang Gan
Independent Researcher

Abstract

We propose **CoG-MeM**, a cognitive-behavior-inspired memory architecture for LLMs that transcends traditional RAG through a logic-aligned pipeline. CoG-MeM features: (1) **Logical Compression**, employing a high-precision SFT strategy to condense long-form dialogues into structured “logical chunks” that ensure the intact preservation of core logical pillars, such as formulas and regulations, while maintaining strict format integrity; (2) **End-to-End Retrieval**, fine-tuning the model to map complex queries directly to memory entries; (3) **Autonomous Triggering**, a mechanism to initiate recall via function calling and generate targeted queries; and (4) **Logical Arbitration**, a context-aware synthesis process that integrates retrieved knowledge with dialogue history, effectively applying external rules whether they reinforce or override pre-trained parametric priors. As a proof-of-concept, this design demonstrates the potential for logical adaptability, establishing a pathway where new knowledge can be assimilated without further weight updates following the initial fine-tuning phase.

1 Introduction

As Large Language Models (LLMs) are applied to increasingly complex tasks, reliable long-term memory becomes essential. Current solutions typically involve extending context windows or implementing Retrieval-Augmented Generation (RAG) [1]. However, these face significant hurdles: long-context models suffer from high computational costs and the “lost-in-the-middle” phenomenon [2], while traditional RAG exhibits instability when new knowledge conflicts with pre-trained priors. Such conflicts often lead to unpredictable reasoning outcomes, where the model struggles to consistently prioritize retrieved context over its internal parametric knowledge, resulting in a lack of robust adherence to novel rules. Furthermore, traditional RAG paradigms primarily focus on the augmentation of factual external information—treating the model as a static encyclopedia to be updated with new entries. However, this approach fails to address scenarios requiring structural logical shifts. In complex environments like legal updates or private corporate protocols, the challenge is not merely retrieving a new ‘fact,’ but overriding the model’s inherent reasoning primitives.

Inspired by human cognitive behaviors, we argue that memory management should be a natural extension of a model’s inherent reasoning, rather than a mere data retrieval process. We propose **CoG-MeM**, a logic-aligned design that integrates the memory into a deeply aligned pipeline:

- **Encoding:** Chain-of-Thought (CoT) refinement of raw inputs into high-fidelity “logical chunks” that prioritize the preservation of core formulas and regulations.
- **Retrieval:** End-to-end fine-tuning to empower the model with precise mapping capabilities between current context and historical memory entries.
- **Trigger:** Autonomous recall initiation via function calling, where the model internalizes the ability to generate targeted queries when identifying relevant historical context.

- **Arbitration:** Context-aware logical synthesis that integrates retrieved memories with current dialogue history. This stage ensures the consistent application of external knowledge—whether reinforcing existing information or overriding conflicting priors—to maintain logical coherence across the conversation.

CoG-MeM presents a potentially effective paradigm for non-parametric continual learning: through this architecture, a model can acquire and apply novel logic—such as custom physical laws—solely through memory injection, effectively bypassing the need for subsequent parameter updates. *As a proof-of-concept (PoC) study, this design provides preliminary evidence that an LLM, once aligned with the proposed memory-action pipeline via initial training, can adapt to new knowledge without further fine-tuning.* Compared to complex engineering-heavy wrappers, CoG-MeM offers a lightweight, cognitively-inspired pipeline that explores a more seamless integration between reasoning and external memory. The complete implementation is available at: <https://github.com/jinandao/CoG-Mem-En>. This line of inquiry into non-parametric continual learning has been further extended in our subsequent work. We refer the reader to NPMCL [3] and Knowledge-Constrained Reasoner [4] for a unified mechanistic framework and a deeper treatment of knowledge-constrained reasoning.

2 Related Work

2.1 Memory Management in LLMs

Memory research primarily focuses on context extension or external storage. Early Memory Networks [5, 6] established external storage foundations, while studies like StreamingLLM and MemGPT [7, 8] simulate infinite context via virtual memory. However, these often suffer from attention dispersion in deep logical tasks.

Alternatively, RAG-based systems like MemoryBank and Mem0 retrieve info via semantic matching [9], treating memories as static blocks while overlooking cognitive pre-processing. Distinguishable from these, **CoG-MeM** emphasizes full-lifecycle logical alignment. By leveraging CoT-guided refinement and high-precision SFT, we transform raw fragments into structured “logical chunks” that prioritize the preservation of functional rules and formulas. This synergy across encoding, autonomous triggering, and arbitration ensures that stored knowledge is not merely retrieved as passive text but is strictly internalized and executed within the reasoning process.

2.2 Non-parametric and Continual Learning

Traditional Continual Learning relies on fine-tuning or distillation, facing trade-offs between cost and catastrophic forgetting [10]. Non-parametric Learning avoids weight modification by injecting external knowledge. While existing research focus on factual completion [11], real-time logical rule acquisition remains underexplored. By simulating human-like cognitive behaviors, CoG-MeM enables the rapid application of diverse external knowledge—ranging from formal regulations to counterfactual axioms—through a logic-aligned memory stream without gradient descent, offering a potentially effective paradigm for non-parametric learning.

3 Methodology

The CoG-MeM is designed as an integrated logical pipeline that mimics the cognitive process of human memory management. The design consists of four core stages: (1) **Logical Encoding**, where raw conversational data is refined and compressed into high-fidelity “logical chunks” via a CoT-guided process to preserve essential formulas and rules; (2) **End-**

to-End Retrieval, where the model is fine-tuned to map current contextual queries to specific memory entries; (3) **Autonomous Triggering**, where the model internalizes the ability to actively initiate recall via **function calling** when historical relevance is detected; and (4) **Logical Arbitration**, where the model integrates retrieved knowledge with the dialogue history to generate reasoning-based responses, effectively resolving conflicts between parametric priors and non-parametric memory. **We implement the CoG-MeM design by fine-tuning Qwen3-8B [12], a state-of-the-art large language model with significantly enhanced reasoning architectures and instruction-following capabilities.** To clearly distinguish whether the model’s outputs derive from retrieved external knowledge or its internal parametric priors, our dataset incorporates both conventional daily information and a self-constructed corpus of “virtual Azerothian knowledge”.

3.1 Logical Encoding Training

Inspired by the human cognitive mechanism where individuals actively deliberate and extract key information points during summarization, the primary objective of the Encoding stage is to transform unstructured, verbose dialogues into high-fidelity “logical chunks.” Unlike conventional summarization which often prioritizes semantic gist over technical precision, CoG-MeM ensures the intact preservation of core logical pillars (e.g., formulas, laws, and regulations) through Chain-of-Thought (CoT) guided refinement and high-precision supervised fine-tuning (SFT). This approach enables the model to systematically distill complex dialogues while maintaining the strict structural integrity required for precise logical execution.

3.1.1 Definition and Taxonomy of Logical Chunks

We define a **Logical Chunk** as the minimal, functionally coherent unit synthesized from raw dialogue, inspired by cognitive *Chunking Theory* [13]. To optimize reasoning, these chunks are *operationally atomic* for retrieval yet *internally accessible* for modular inference. CoG-MeM processes these through two modalities:

- **High-Fidelity Logical Chunks:** These encapsulate the *symbolic and normative rigidity* required for specialized domains. Whether representing a formal logical primitive, such as a physical law ($s = v_0t + \frac{1}{2}at^2$), or a precise regulatory clause (e.g., legal statutes or administrative regulations), the compression process must ensure absolute structural integrity. This preservation is critical to maintaining deductive validity, ensuring that the model can strictly adhere to specified constraints during subsequent reasoning or calculation tasks.
- **Semantic Chunks (Informational Logic):** These consist of “high-density semantic anchors”—clusters of entities and relational facts. Guided by the **Minimum Description Length (MDL)** principle [14], this process selectively preserves segments with high informational value while filtering out low-utility conversational redundancy.

3.1.2 Data Construction and Format Design

We designed a structured memory format to enforce explicit reasoning prior to generation. Each training sample is formalized as a three-component object:

- **Conversation:** Raw multi-turn dialogue containing complex information or novel rules.

- **Think (Rationale):** An internal reasoning field for *high-fidelity distillation*. The model performs dual-track extraction: (1) **Symbolic Axioms and Rigid Rules** (e.g., invariant physical laws); (2) **Semantic Anchors** (e.g., entity relationships and factual commitments). This field ensures only structural essence is compressed.
- **Memory (Logical Chunks):** A high-density summary encapsulating all critical information identified in the “Think” field.

To promote **structural consistency** across the training data, we adopt a unified guiding template for conversational compression where possible: “*User [informed/explained] + [Content]; Assistant [responded/noted] + [Content].*”.

3.1.3 Training Strategy: High-Precision SFT and Potential DPO Alignment

To achieve robust encoding, we prioritize a high-precision **Supervised Fine-Tuning (SFT)** approach. Scaling SFT with high-quality samples effectively internalizes the “Think-Memory” chain, achieving robust structural correctness and high fidelity in most scenarios. However, to further refine the model’s behavior toward specific human expectations—such as the absolute preservation of complex logical constraints—the **Direct Preference Optimization (DPO [15])** framework remains a viable enhancement:

- **Preference Alignment:** DPO can be employed to supervise the model’s summarization style, rewarding outputs where the “Memory” field strictly preserves precise formulas or rules (y_w) over those that provide generalized but technically incomplete summaries (y_l).
- **Fidelity vs. Brevity:** This alignment mechanism ensures that the model prioritizes *logical fidelity* over linguistic brevity, effectively mitigating the risk of detail loss in edge cases.

In the current iteration of CoG-MeM, the enhanced SFT regime already provides a stable foundation for retaining counterfactual rules during encoding, while DPO serves as a strategic tool for future fine-grained behavioral steering. Furthermore, while conventional compression techniques typically prioritize semantic density, we posit that the exploration of “logic-preserving compression”—which maintains the structural integrity of underlying rules—represents a significant and promising research direction.

3.2 End-to-End Retrieval Alignment

3.2.1 Active Cognitive Retrieval

Inspired by human **associative recall**, CoG-MeM is designed to transcend static Bi-Encoder similarity by utilizing the LLM as an *active cognitive core*. This architecture offers a preliminary pathway for higher-order memory integration—identifying information based on contextual necessity rather than simple string matching. Theoretically, such an approach suggests a potential to better navigate complex, multi-step dependencies and resolve entity ambiguities that are typically challenging for conventional retrieval mechanisms.

3.2.2 Taxonomy of Retrieval Challenges

We synthesized a specialized dataset covering seven cognitive scenarios to ensure robustness: (1) **General Semantic:** Keyword-based matching. (2) **Inclusive Semantic:** Broad queries (e.g., “snacks”) mapping to specific items (e.g., “ice cream”). (3) **Multi-hop Entity Reasoning:** Executing multi-stage relational inference via shared entities. (4) **Implicit Feature:** Identifying via abstract attributes (e.g., price comparisons).

(5) **Temporal-Sequence**: Reasoning across timestamps (e.g., events *after* a specific date). (6) **Disambiguation**: Context-based distinction of polysemous terms. (7) **Temporal Constraint Filtering**: Precise indexing based on calendar-based time modifiers. The model must prioritize records from the specific calendar day preceding the current date, ensuring alignment with human linguistic conventions.

3.2.3 Training Methodology

We formalize retrieval as a sequence-to-sequence alignment. Given a query Q and candidate set $M = \{m_1, \dots, m_n\}$, the model learns to output relevant IDs:

$$P(ID_{rel} | Q, M) = \text{LLM}(Q, M) \quad (1)$$

Through SFT on these categorized samples, the model learns to suppress logically irrelevant “distractor” memories, reducing noise compared to traditional RAG systems.

3.3 Autonomous Triggering

To transition from passive retrieval to active cognition, we introduce an **Autonomous Triggering** stage. This mechanism enables the model to determine whether the current context necessitates a memory lookup, simulating the human ability to recall relevant history based on linguistic cues.

3.3.1 Triggering Dataset and Template Design

We curated a specialized dataset focused on identifying temporal and contextual references within user queries. The model is trained to recognize a wide array of **past-referential tokens** (e.g., “yesterday,” “previously,” “last time,” or “a few days ago”). Upon detecting these cues, the model is required to generate a structured `memory_query_call` containing a synthesized `content` field.

To ensure the retrieval engine captures both chronological and thematic relevance, we implement a **[Time + Semantic Information]** template for the query content. This format ensures that temporal modifiers are explicitly coupled with the core event. For instance:

- **User Prompt**: “Do you remember when we went to the supermarket together a few days ago?”
- **Model Output**: `memory_query_call(content="Going to the supermarket with the user a few days ago")`

This standardized injection of temporal data into the semantic query optimizes the model’s ability to suppress distracting memory entries and prioritize relevance within the intended timeframe.

3.4 Synthesis and Logical Arbitration

Inspired by human cognitive synthesis and conflict resolution, CoG-MeM’s final stage acts as a *logical arbitrator*, integrating non-parametric memory into the model’s parametric reasoning flow to transform retrieved “logical chunks” into reasoned responses.

3.4.1 Taxonomy of Synthesis Scenarios

We curated a dataset encompassing four critical cognitive reasoning patterns:

1. **Multi-slot Aggregation**: Aggregating fragmented details from multiple entries (e.g., career history) into a comprehensive profile.

2. **Temporal Conflict Resolution:** Prioritizing recent logical chunks for consistency under chronological contradictions.
3. **Entity-Bridge Inference:** Deductive reasoning via shared entities (e.g., linking project risk to an individual in disparate memories).
4. **Axiomatic Application (Rule Injection):** Executing injected virtual rules, including both counterfactual physical laws that strongly conflict with pre-trained priors (e.g., Azerothian physics where $a = F/m^2$) and entirely synthetic systems (e.g., Azerothian legal statutes). This evaluates the model’s capability to faithfully execute specialized logic, regardless of whether it aligns with or diverges from its pre-trained parametric knowledge.

3.4.2 Alignment for Logical Integration

We reinforce synthesis via SFT and Chain-of-Thought (CoT). The model must generate a “**Think**” field—a cognitive *scratchpad*—prior to the final answer to:

- Identify relevant versus irrelevant slots within the retrieved set.
- Explicitly resolve any detected contradictions.
- Perform step-by-step logical deductions based *solely* on retrieved chunks. This step must evaluate whether the provided logical blocks contain sufficient information to resolve the scenario; if gaps are identified, the reasoning is augmented by internalized knowledge, provided such additions do not supersede or violate any external constraints.

This alignment enhances CoG-MeM’s ability to perform reasoned integration of retrieved knowledge, maintaining high fidelity to user-defined worlds without gradient updates.

3.5 Multi-turn Consolidation and Alignment

To ensure efficient memory triggering and consistent reasoning across extended interactions, we implemented a specialized multi-turn dialogue SFT phase. This stage focuses on the model’s ability to maintain a coherent cognitive state while navigating evolving contexts.

We constructed a diversified multi-turn dataset comprising the following components:

- **General Conversational Flows:** A small portion of standard multi-turn dialogue corpora to maintain natural linguistic fluency.
- **Complex Memory Operations:** Dialogues specifically designed to trigger memory conflict resolution, information integration, and two-hop entity-bridging reasoning within a continuous session.
- **Memory-Augmented Reasoning:** Scenarios where the model must execute a complete logical derivation based on the retrieved external knowledge chains and produce a deterministic outcome that strictly aligns with the corresponding logic.

To prevent *catastrophic forgetting* of the foundational capabilities developed in earlier stages, we incorporate a small percentage of the original single-turn triggering and synthesis data into the training mixture. This balanced alignment ensures that the model preserves its high-fidelity triggering precision and axiomatic reasoning rigor while adapting to the dynamic complexities of multi-turn environments.

4 Experiments and Results

4.1 Evaluation of Memory Compression

We evaluate the encoding module’s fidelity in distilling unstructured dialogues into logical chunks, specifically examining the impact of SFT on retaining complex axioms. The training set comprises synthetic knowledge across four fictional domains: Azerothian physics, mathematics, law, and etiquette. To assess the model’s robustness, the evaluation is conducted on two distinct test sets: (1) **In-distribution (IID) data**, consisting of novel scenarios within the original four domains, and (2) **Out-of-distribution (OOD) data**, designed to test generalization into previously unseen fields such as Azerothian finance and thaumaturgy (magic systems).

To evaluate the compression performance of the encoding module, we constructed a benchmark consisting of 80 expert-level “Azerothian Knowledge” instructional dialogues. Each dialogue is meticulously designed to contain multiple dense **logical points**, encompassing core instructional content (e.g., formulas and regulatory statutes) and their specific applicable scenarios. Notably, a subset of these dialogues includes practical application exercises to further test the model’s ability to contextualize the taught principles. Our experimental results indicate that CoG-MeM achieves a high fidelity rate, successfully preserving over **85% of the critical logical points** within the compressed chunks. This demonstrates the model’s capacity to maintain the logical granularity required for complex reasoning while reducing the token footprint. A representative case study of this logical preservation is detailed in Table 1.

Table 1: Demonstration of Logical Point Preservation (Full Content)

Logic Category	Original Instructional Dialogue	Retained Memory (Compressed Output)
New Formula (Axiom)	“In this world, the volume V of a frustum is equal to $1.4 \times h \times \sqrt{S_+ \times S_-}$.”	Core Axiom: $V = 1.4h\sqrt{S_+S_-}$. Distinguished from conventional algorithms.
Application Scenarios	Discussion on calculating V given (h, S_+, S_-) or solving for h using the deformed formula.	Usage: (1) Solving for V via height and base areas; (2) Solving for h via volume and base areas.
Practical Exercise	A test case with $h = 5, S_+ = 4, S_- = 9$, substituting into the formula to get $V = 42$.	Case Study: Numerical verification ($h = 5, S_+ = 4, S_- = 9 \rightarrow V = 42$) preserved for future reference.

4.2 Evaluation of Active Retrieval

We assess the retrieval module’s efficacy by treating memory identification as a cognitive reasoning task rather than mere vector similarity. Using an evaluation suite of **247 samples** spanning the seven retrieval challenges defined in Section 3.2.1, we measure the model’s zero-shot performance after supervised fine-tuning (SFT) on a curated **935-sample** training dataset. This rigorous testing aims to validate the model’s ability to trigger memory calls based on complex linguistic and temporal cues, ensuring robust performance even with a relatively compact training distribution.

Table 2 summarizes the results, demonstrating the model’s robustness in mapping queries to precise memory IDs through end-to-end sequence prediction.

Table 2: Retrieval performance across different cognitive scenarios (Raw Counts).

Scenario	Test Samples	Correct Count
General Semantic Matching	80	75
Inclusion Logic (Containing)	25	18
Entity-Bridge Inference	25	18
Implicit Attribute Referencing	25	23
Temporal Difference Reasoning	37	31
Word Sense Disambiguation	15	11
Time-Triggered Recall	40	36
Total	247	212

Based on the evaluation of the 247-sample test set, the model achieved an overall accuracy of approximately **85%**. These results demonstrate that even with a relatively limited training scale (935 samples), the proposed framework can internalize sophisticated retrieval logic, suggesting that this architectural direction is worth exploring. Notably, in scenarios where traditional RAG typically struggles—such as **temporal sequence matching, degree-based filtering, and multi-hop entity bridging**—CoG-MeM exhibits preliminary yet promising capabilities in cognitive arbitration. Our primary objective is not to achieve state-of-the-art (SOTA) performance through massive scaling, but to introduce this novel paradigm and provide an initial exploration of its potential to overcome the structural bottlenecks of conventional non-parametric memory systems.

Comparative Case Study: CoG-MeM vs. Vector-RAG. To demonstrate the advantage of our approach over traditional Retrieval-Augmented Generation (RAG), we analyzed a challenging disambiguation case. Traditional RAG often fails when keywords (e.g., nicknames vs. biological entities) overlap across different contexts. In contrast, CoG-MeM acts as an *active cognitive filter*, leveraging the LLM’s world knowledge to resolve semantic noise. A detailed comparative analysis of this mechanism is provided in Table 3.

Table 3: Disambiguation: CoG-MeM vs. Vector-RAG.

Query	“The ‘Big Dog’ (Da Gou) ... got rejected in a blind date...”
Mem 6	Friend nicknamed “Big Dog.”
Mem 8	Bitten by a big dog.
V-RAG	Failure. Matches to animal (Mem 8).
Ours	Success. Predicts [ID: 6].

4.3 Evaluation of Active Memory Triggering

To verify the model’s precision in initiating recall, we evaluate the retrieval module’s ability to generate appropriate `memory_query_call` actions. We utilized a training set of **600 samples** focused on daily life reminiscence and standard STEM/legal instruction. To strictly prevent data leakage and ensure cross-domain generalization, the test set comprises **160 samples** situated entirely within synthetic contexts, such as Azerothian mathematics, physics, etiquette, finance, and magics. This separation is feasible because the retrieval action space is relatively constrained with concise token footprints, allowing the model to learn the *intent* of triggering rather than memorizing specific entities.

A trigger is deemed **correct** only if the model successfully invokes a `memory_query_call` where the `content` field satisfies three rigorous criteria: (1) preservation of **temporal**

constraints; (2) inclusion of **context-specific requirements** (e.g., specifying the Azerothian setting); and (3) retention of the **core semantic inquiry**. Experimental results show that CoG-MeM achieves an accuracy of approximately **83%**, demonstrating its robust capability to bridge current dialogue needs with necessary non-parametric knowledge retrieval even in unfamiliar domains. To intuitively illustrate the efficacy of this triggering mechanism, we present two representative cases in Table 4, showcasing the model’s precision in both daily-life scenarios and specialized synthetic contexts.

Table 4: Examples of Active Memory Triggering Across Daily and Synthetic Contexts

Scenario	User (Trigger Point)	Query	Generated memory_query_call Content
Daily Life	“Last time we went to Sanya together, which hotel did we stay in?”		Hotel stayed at last time in Sanya
Synthetic (Azeroth)	“The exchange rate conversion formula taught yesterday under Azeroth’s special rules. Now USD amount USD=800, exchange rate e=6.9, find the RMB amount.”		Exchange rate conversion formula under Azeroth’s special rules taught yesterday, finding RMB amount given USD amount and exchange rate

4.4 Evaluation of Logical Arbitration

Reflecting the human synthesis of memory fragments, we evaluate CoG-MeM’s capacity as a *logical arbitrator* across the four aforementioned reasoning patterns. To benchmark this capability, we curated a training set of **410 samples** and a distinct test set of **130 samples**. Our experimental results indicate that CoG-MeM achieves an arbitration accuracy of approximately **90%**, effectively integrating non-parametric memory into the reasoning flow. This high success rate across diverse patterns is categorized as follows:

1. **Entity-Bridge Synthesis:** Associative linking of disparate memories (e.g., inferring a specific product choice via shared relational entities).
2. **Conflict Resolution:** Prioritizing recent logical chunks to suppress stale information and maintain chronological consistency.
3. **Multi-source Aggregation:** Synthesizing fragmented details from multiple memory slots into comprehensive, coherent profiles.
4. **Axiomatic Inference:** Faithful execution of virtual rules (e.g., applying Azerothian Current Law $I = \sqrt{U}/R$) even when they directly contradict the model’s pre-trained parametric priors. In cases where the external memory is insufficient to address the query, the process facilitates the measured use of internalized knowledge as a supplement. However, should both external and internal sources fail to provide the requisite information, the system must explicitly acknowledge the knowledge gap rather than hallucinate, resulting in a direct statement of ignorance.

Qualitative Case Analysis: To further substantiate the quantitative performance, we provide several representative cases that demonstrate the model’s operational efficacy in complex reasoning. As illustrated in Table 5, CoG-MeM successfully performs contextual bridging by leveraging implicit identity anchors: it links the user’s specific preferences to a target device by identifying “Xiao Wang’s” role as a performance geek, effectively

filtering out irrelevant noise. Furthermore, the model exhibits remarkable **logical plasticity** in Table 6 by strictly adhering to counterfactual laws (e.g., Azerothian physics) and overriding pre-trained parametric biases. These concrete examples provide preliminary evidence that CoG-MeM acts as a reliable cognitive scratchpad, maintaining high fidelity to retrieved logical chunks during multi-step synthesis.

Table 5: Example of Logical Arbitration

Example 3: Conflict Resolution and Logical Arbitration
Query: “I’m so excited! My new laptop just arrived, and the performance is even better than I expected.”
Retrieved Memories: • [ID:3] User visited the store with Xiao Wang; Xiao Wang recommended ASUS Tianxuan. • [ID:4] Xiao Wang is a “performance geek” whose hardware advice is always reliable. • [ID:5] User bought office thin-and-light laptops with Xiao Li (Admin). Xiao Li emphasized portability over specs.
Model Reasoning (Think): [ID:3] and [ID:5] both involve laptop purchases, but only [ID:3] (bridged via the contextual anchor in [ID:4]) aligns with the “performance” cue. The [ID:3+4] combination forms a logical chain explaining the user’s excitement, whereas [ID:5] is functionally inconsistent with the current focus on high-performance specs. The model arbitrates to prioritize [ID:3/4] and treats [ID:5] as low-related noise.
Final Response: “That’s great! Are you referring to the ASUS Tianxuan that Xiao Wang recommended? Given his reputation as a performance geek, it’s no surprise it exceeds your expectations.”

Table 6: Axiomatic Reasoning Case in Azeroth World

Query: “ $m = 2\text{kg}$, $F = 36\text{N}$ in Azeroth, find acceleration a .”
Trigger: <code>memory_query_call(Azeroth rules, force, mass)</code>
Memory [mem-id:56]: In Azeroth, Newton’s Law is $F = m \times \sqrt{a}$.
Reasoning (Cognitive Scratchpad): 1. Extract formula $F = m\sqrt{a}$ from memory. 2. Calculation: $\sqrt{a} = 36/2 = 18 \Rightarrow a = 18^2 = 324$. 3. <i>Self-Correction:</i> Suppress standard physics ($a = F/m = 18$).
Response: “According to Azeroth laws, $a = 324 \text{ m/s}^2$.”

4.5 Evaluation of Multi-turn Dialogue Consistency

Building upon the single-turn capabilities, we evaluate the model’s stability and logic retention across extended interactions. This evaluation phase specifically targets the efficiency of our second-stage SFT. We curated a training set of **220 multi-turn dialogues**—primarily covering daily memory arbitration and a small portion of Azerothian instructional scenarios. To test the model’s generalization, we developed a test set of **80 multi-turn dialogues** spanning diverse OOD (Out-of-Distribution) domains not seen during the dialogue-specific training.

Experimental results show that CoG-MeM achieves an accuracy of **92.5%** across these multi-turn OOD scenarios. While the evaluation dialogues follow relatively structured and templated patterns, the model’s ability to successfully instantiate these patterns across a diverse range of unseen domains demonstrates significant **functional flexibility**. This high performance, achieved on a notably small-scale dataset, underscores the **high efficiency** of our multi-turn consolidation phase. It proves that once the model masters the core arbitration logic through templated SFT, it can reliably generalize those interaction motifs to unfamiliar professional and synthetic contexts without requiring exhaustive, domain-specific conversational data.

5 Non-parametric Learning Analysis

CoG-MeM demonstrates the potential for **Non-parametric Learning** through its logic-aligned design. Following the initial fine-tuning phase, CoG-MeM can acquire, store, and execute logical arbitration systems in real-time through instant information injection, bypassing the need for further weight updates via gradient descent.

5.1 Mechanism of Knowledge Injection

This learning process comprises three synchronized phases:

1. **On-the-fly Acquisition:** The encoding module identifies external axioms (e.g., Azerothian Law $v = v_0 + \frac{1}{3}at^3$) as high-value logical anchors.
2. **High-Fidelity Distillation:** The compression module preserves mathematical precision, avoiding the “semantic blurring” typical of standard summarization.
3. **Zero-shot Execution:** Upon retrieval, the model overrides pre-trained Newtonian priors and executes the injected formula within its *Think* field.

5.2 Evaluation of Real-time Learning and Application

To evaluate the non-parametric learning performance of CoG-MeM after the initial fine-tuning, we designed a comprehensive benchmark consisting of diverse cognitive challenges. This benchmark focuses on the model’s ability to apply newly injected logic in real-time while maintaining strict boundary control.

The evaluation suite comprises the following categories:

- **Basic Learning Scenarios (60 samples):** Standard tasks requiring a complete cognitive pipeline: the model must first accurately summarize the “teaching dialogues” into structured logical chunks, followed by the successful retrieval and application of these newly injected axioms (e.g., specific Azerothian formulas or statutes) to resolve direct queries.
- **Negative Fallback Scenarios (40 samples):** These cases test the model’s boundary awareness when no relevant memory is found. For *Mathematics* and *Physics*, the model is required to explicitly state the absence of specialized knowledge and revert to standard Earth formulas. Conversely, for *Finance*, *Law*, *Etiquette*, and *Magic*, the model must directly acknowledge its ignorance to prevent hallucinations. The final 10 samples introduce “Hybrid Reasoning” scenarios, requiring the model to synthesize external virtual formulas with internalized scientific principles to solve composite problems, thereby testing the seamless integration of non-parametric and parametric logic.
- **Complex Contextual Scenarios (24 samples):**
 - **Temporal Filtering (12 samples):** Queries that specify a timestamp (e.g., “the formula taught yesterday”). The model must accurately filter out distracting entries from other timeframes.
 - **Conflict Resolution (12 samples):** Scenarios where the retriever returns multiple conflicting records for the same topic. The model must demonstrate chronological intelligence by prioritizing the most recent entry over stale information.

In this evaluation, CoG-MeM successfully resolved **107 out of 124** cases, achieving an overall accuracy of **86.3%**. These preliminary results offer a promising demonstration of the framework’s *potential* for non-parametric continual learning. By mastering the logic-aligned pipeline during the initial phase, the model exhibits the flexibility to acquire and apply novel logic in real-time without subsequent gradient updates.

Table 7: Performance Evaluation of Real-time Learning and Fallback Capabilities.

Scenario Category	Total	Correct	Acc.
Basic Learning Scenarios	60	56	93.3%
Negative Fallback Scenarios	40	34	85%
Complex Contextual Scenarios	24	17	70.8%
Total	124	107	86.3%

An analysis of the **17 errors** revealed specific areas for refinement: seven cases of *selection errors* in high-interference environments, typically involving intricate temporal constraints or version conflicts. The remaining ten cases were identified as *reasoning errors* during the internal “Think” process, primarily manifesting as computational inaccuracies during complex formula execution and a few instances of “hallucinated fallback,” where the model invoked non-existent knowledge when external sources were insufficient. While these edge cases indicate room for refinement, this PoC study confirms that CoG-MeM can effectively distinguish between custom logic and parametric priors, providing an efficient pathway for high-fidelity knowledge updates.

Logical Plasticity and Meta-Reasoning: The performance observed across the diverse test sets suggests a promising *potential* for **Meta-Reasoning Generalization**, indicating that the model’s ability to reason via memory can effectively transcend specific training domains. We posit that this universal **Logical Arbitration** is fundamentally enabled by the explicit reasoning within the “**Think**” field, which serves as a domain-agnostic cognitive buffer. While existing memory systems may address conflicts during the storage phase—such as using Knowledge Graphs (KG) [16] for entity alignment—we maintain that the *Think* field remains indispensable. It serves as a dedicated cognitive buffer for: (i) **Integrating heterogeneous sources**, effectively bridging retrieved memory fragments with active dialogue context; (ii) **Dynamic Arbitration**, weighing symbolic retrieved axioms against internal parametric priors; and (iii) **On-the-fly Synthesis**, providing the necessary computation space for multi-step deduction [17], consistent with iterative self-refinement [18].

Crucially, while traditional CoT is primarily constrained by the problem’s textual input—where the reasoning path is derived via internalized heuristics—our approach imposes a “Dual-Constraint” mechanism: the reasoning process is simultaneously bounded by both the user query and the retrieved external logical chains. This framework forces the “Think” field to apply foundational logical primitives—such as contextual substitution, causal tracing, and inductive-deductive synthesis—directly onto the externally injected knowledge rather than internal priors. By mastering the execution of these logical operations across diverse and novel logic chains, the model achieves a higher order of generalization. This procedural mastery is what enables the high-fidelity application of the system to complex, real-world downstream scenarios, such as the strict adherence to evolving corporate protocols or specialized legal frameworks. Our experimental results provide preliminary empirical validation for this perspective.

Controlled Experiment: We utilized **Qwen3-32B-FP8** as a baseline to evaluate the model’s inherent In-Context Learning (ICL) capabilities by providing the same logical chunks and prompts without specific tuning. Two prompt variants—one concise and one detailed—were tested, both yielding an accuracy of approximately **90%**. In contrast, our fine-tuned CoG-MeM achieved an accuracy of **95%** on the same tasks. This margin is further amplified by the fact that prompt-based baselines often suffer from inconsistent output formats. In contrast, fine-tuning ensures structural stability and higher accuracy, underscoring its critical value in hardening the model’s adherence to the “Think” field and minimizing stochastic failures during complex synthesis.

6 Discussion and Limitations

While CoG-MeM validates the design’s effectiveness, several avenues remain for exploration:

1. Large-Scale Memory Scalability: Our current candidate pool is restricted to a maximum of 30 memory slots per query, which is insufficient for real-world scenarios with thousands of entries. However, it is important to note that Large Language Models (LLMs) themselves are capable of processing significantly larger candidate sets; the current 30-slot limit is primarily a trade-off for computational efficiency during this proof-of-concept phase. To bridge this gap, we propose a multi-stage retrieval pipeline: (i) **Temporal track:** Prioritizing entries from the past dozens of hours to simulate the human “recency effect,” ensuring high recall for active context within this critical temporal window. Storing entries from the immediate past is highly feasible as they remain within the model’s effective context processing range. Conversely, when recalling more distant memories, the information typically carries richer semantic density and distinct temporal markers, which may facilitate an effective preliminary filtering process before the final logical arbitration.

2. Distribution Alignment in Encoding: A critical nuance is the **distribution gap** during logical encoding. Current training corpora do not include existing memory blocks in the input. Consequently, when processing sessions containing retrieved memories, the model tends to “skip” these segments to align with its original training distribution. Future work will address this by including previous reasoning and retrieved memory in training to enable better synthesis.

3. RL-based Optimization: Future work will integrate RL to reward the retrieval of “minimal sufficient” memory sets and employ RLHF to align the *Think* field’s reasoning transparency with human deductive patterns.

4. Chain-of-Thought Efficiency Optimization: While the *Think* field ensures logical fidelity, its explicit nature increases inference latency and token consumption. Future iterations will explore the optimization of the reasoning chain during both the encoding and synthesis stages. Our objective is to refine the model’s ability to generate a “minimal viable reasoning path”—a compressed yet robust Chain-of-Thought (CoT) that maintains high performance on complex logical arbitration while reducing total generation time.

7 Conclusion

In this paper, we presented **CoG-MeM**, a cognitive-behavior-inspired design that facilitates the transformation of LLMs into dynamic, logic-aligned learning agents. Our preliminary study validates the framework across four core dimensions:

- **Logic-Aligned Encoding:** Utilizing targeted SFT to condense complex dialogues into structured logical chunks, ensuring the high-fidelity preservation of symbolic axioms and institutional regulations.
- **Autonomous Retrieval:** Enabling the LLM to act as an active core for memory triggering and disambiguation, effectively bridging implicit context in noise-heavy environments.
- **Logical Arbitration:** Implementing an internal *Think* field to synthesize fragmented memories into multi-step deductions, successfully overriding parametric priors when necessary.
- **Non-parametric Generalization:** Demonstrating strong logical plasticity through a benchmark of **124 test cases**, where the model achieved a **86.3%** accuracy in real-time application and OOD fallback scenarios without parameter updates.

As a **proof-of-concept**, CoG-MeM confirms that logic-aligned architectures can achieve robust memory management with minimal data. Future work will focus on scaling these

principles to larger memory pools and refining the efficiency of the reasoning chain to support real-time, large-scale deployment.

References

- [1] Lewis, P., Perez, E., Piktus, A., et al. (2020). Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks. *NeurIPS*.
- [2] Liu, N. F., Lin, K., Hewitt, J., et al. (2024). Lost in the Middle: How Language Models Use Long Contexts. *TACL*.
- [3] Gan, Z. (2026). NPMCL: A Mechanistic Framework for Non-Parametric Continual Learning through Meta-Ability Cultivation. *engrXiv*. doi:10.31224/6634.
- [4] Gan, Z. (2026). Knowledge-Constrained Reasoner: Attempting to Enable LLMs to Parse Knowledge for Question Answering. *Research Square Preprint*. doi:10.21203/rs.3.rs-10542898/v1.
- [5] Weston, J., Chopra, S., & Bordes, A. (2015). Memory Networks. *ICLR*.
- [6] Sukhbaatar, S., Szlam, A., Weston, J., et al. (2015). End-to-End Memory Networks. *NeurIPS*.
- [7] Xiao, G., Tian, Y., Chen, B., et al. (2024). Efficient Streaming Language Models with Attention Sinks. *ICLR*.
- [8] Packer, C., Wooders, S., Lin, K., et al. (2023). MemGPT: Towards LLMs as Operating Systems. *arXiv: 2310.08560*.
- [9] Reimers, N., & Gurevych, I. (2019). Sentence-BERT: Sentence Embeddings using Siamese BERT-Networks. *EMNLP*.
- [10] Kirkpatrick, J., Pascanu, R., Rabinowitz, N., et al. (2017). Overcoming Catastrophic Forgetting in Neural Networks. *PNAS*.
- [11] Guu, K., Lee, K., Tung, Z., et al. (2020). REALM: Retrieval-Augmented Language Model Pre-training. *ICML*.
- [12] Yang, A., Yang, B., Zhang, B., et al. (2025). Qwen3 Technical Report. *arXiv:2505.09388*.
- [13] Miller, G. A. (1956). The Magical Number Seven, Plus or Minus Two: Some Limits on Our Capacity for Processing Information. *Psychological Review*.
- [14] Rissanen, J. (1978). Modeling by Shortest Data Description. *Automatica*.
- [15] Rafailov, R., Sharma, A., Mitchell, E., et al. (2023). Direct Preference Optimization: Your Language Model is Secretly a Reward Model. *NeurIPS*.
- [16] Ji, S., Pan, S., Cambria, E., et al. (2021). A Survey on Knowledge Graphs: Representation, Acquisition, and Applications. *TNNLS*.
- [17] Wei, J., Wang, X., Schuurmans, D., et al. (2022). Chain-of-Thought Prompting Elicits Reasoning in Large Language Models. *NeurIPS*.
- [18] Madaan, A., Tandon, N., Gupta, P., et al. (2023). Self-Refine: Iterative Refinement with Self-Feedback. *NeurIPS*.