

MeVe: A Modular System for Memory Verification and Effective Context Control in Language Models

Andreas Ottem

Independent Researcher
Alta, Norway

Abstract. Retrieval-Augmented Generation (RAG) systems typically face constraints because of their inherent mechanism: a simple top-k semantic search [1]. The approach often leads to the incorporation of irrelevant or redundant information in the context, degrading performance and efficiency [10][11]. This paper presents MeVe, a novel modular architecture intended for Memory Verification and smart context composition.

MeVe rethinks the RAG paradigm by proposing a five-phase modular design that distinctly breaks down the retrieval and context composition process into distinct, auditable, and independently tunable phases: initial retrieval, relevance verification, fallback retrieval, context prioritization, and token budgeting. This architecture enables fine-grained control of what knowledge is made available to an LLM, enabling task-dependent filtering and adaptation. We release a reference implementation of MeVe as a proof of concept and evaluate its performance on knowledge-heavy QA tasks spanning biomedical, legal, Wikipedia, and HotpotQA domains [22] [26] [27]. Our results demonstrate that by actively verifying and refining contextual information, MeVe achieves substantial gains in both accuracy and efficiency: Exact Match improves by around 40%, F1 by 34%, and usable context by over 40% while processing time rises only barely (12%). These improvements translate into large reductions in contextual redundancy (up to 57% on Wikipedia and 75% on HotpotQA) and consistent accuracy gains across all evaluated domains. Component ablation further highlights the importance of Hybrid Fallback, Coherence Filtering and Adaptive Verification in sustaining these benefits [25]. This work provides a framework for more scalable and reliable LLM applications. By refining and distilling contextual information, MeVe offers a path toward better grounding and more accurate factual support [16].

Keywords: Retrieval-Augmented Generation, Language Models, Context Management, Memory Verification, Modular Architecture.

1 Introduction

The impressive abilities of Large Language Models (LLMs) are fundamentally reliant on the ability to reason in a particular context [6] [7]. As models continue to utilize larger context windows and retrieve information from external memory, the key challenge shifts from simple information retrieval to smart contextual management [5]. It includes the dynamic fetching of contextually relevant knowledge, rigorous verification of its relevance, and efficient composition with the objective of improving LLM performance [6] [14].

Retrieval-Augmented Generation (RAG) has also become a strong baseline for augmenting LLMs with out-of-model knowledge [1].

Traditional RAG systems commonly operate by performing a direct, often "top-k" semantic search against a pre-indexed knowledge base, retrieving the most similar documents and directly appending them to the LLMs input prompt. While effective in providing external knowledge, this straightforward mechanism can lead to issues such as including irrelevant, redundant or even contradictory information [14] [1] [2]. Yet, typical implementations tightly intertwine retrieval and injection into a single monolithic step, usually by

picking a fixed number of embedding-nearest neighbors and simply adding them to the prompt [2] [3]. This common practice makes verification implicit, heavily reduces modularity, and tends to have a difficult time gracefully scaling to dynamic memory expansion or tight token budget requirements [17]. The direct injection of possibly irrelevant or redundant information can lead to "context pollution" [11] [14], degrading the quality and effectiveness of LLM responses and increasing the risk of factual error or "hallucinations" [10] [14] [13]. This challenge is compounded by the fact that retrieval quality is highly dependent on the preciseness and structure of the knowledge corpus itself [2].

As a solution to such fundamental shortcomings, we present MeVe, a new method that fundamentally redesigns retrieval, verification, prioritization, and budgeting as separate, configurable, and auditable phases [4] [19]. MeVe's key contribution is its modular decomposition of the RAG pipeline, transforming it from a singular operation into a structured, multi-stage process with explicit control over content quality and composition. MeVe does not aim to substitute RAG but instead redefines traditional RAG as one possible configuration within a more sophisticated, memory-aware architecture. The modular design founded upon the principles of explicit verification of information and strict retrieval mechanisms aim to alleviate typical failure modes that have been witnessed in state-of-the-art RAG systems, including context confusion and propagation of irrelevant or misleading information [10] [11] [16]. The MeVe framework provides a modular and systematic approach to LLM context management, leading to enhanced efficiency and control in tasks involving large amounts of knowledge [17]. This paper is organized as follows: Section 2 reviews relevant literature, Section 3 details the MeVe architecture, Section 4 presents our empirical evaluation, and Section 5 discusses the results, leading to our conclusion in Section 6.

2 Related Work and Conceptual Origins

2.1 Retrieval-Augmented Generation (RAG)

RAG enhances LLMs by anchoring them to external knowledge [1]. While effective, the standard top-k paradigm is notoriously vulnerable to "distractor documents", which is semantically similar information that is contextually irrelevant [10] [11] [13]. Despite the presence of sophisticated methods such as re-ranking, these are typically applied as secondary corrections to an already formed, imperfect candidate pool [9] [2]. MeVe addresses this issue by bringing explicit verification to a basic, integral step of the memory processing paradigm, thereby setting itself apart from solutions founded solely upon early retrieval or later re-ranking for determining relevance [2] [3].

2.2 Long-Context Architectures and the Need for Filtering

The introduction of Long-Context Architectures, such as in Ring Attention systems or those controlled by MemGPT, enables LLMs to process much greater levels of information [5] [7]. These architectures alone do not inherently ensure the quality or applicability of the augmented context [6] [5]. The "garbage in, garbage out" principle is particularly evident in such systems: unchecked data not only increases computational overhead but also causes context pollution, heightening the risk of hallucination [2] [14] [16]. MeVe acts as a discerning gatekeeper for these broad contexts, ensuring that only high-utility information is passed on to the LLM [17].

2.3 Hybrid Search and Multi-Stage Filtering

MeVe integrates concepts of classical "information retrieval," especially via hybrid search and multi-stage filtering [8] [19]. By pairing a keyword-based fallback mechanism with dense retrieval methods, our system employs a hybrid search approach to provide robustness against vector-only search pitfalls like semantic drift problems [15]. MeVe's design allows fine-grained control of each stage of retrieval and filtering, building on prior modular IR pipelines [4] [19].

3 The MeVe Architectural Framework

MeVe is designed as a serial pipeline of compatible modules that convert a user query (q) into a dense, high-utility terminal context (C_{final}) for an LLM [4]. It offers a structured procedure that enables controlled and optimized functioning at each critical phase of context generation [17]. MeVe distinguishes itself from conventional RAG architectures through its five-phase modular design, offering enhanced control and transparency over the knowledge supplied to an LM. Unlike existing end-to-end RAG systems that often treat retrieval and text composition as a single, opaque step, MeVe isolates and optimizes each stage of memory processing. This modularity improves system interpretability and allows for targeted interventions to improve relevance, robustness, and efficiency—factors often difficult to control in monolithic RAG implementations.

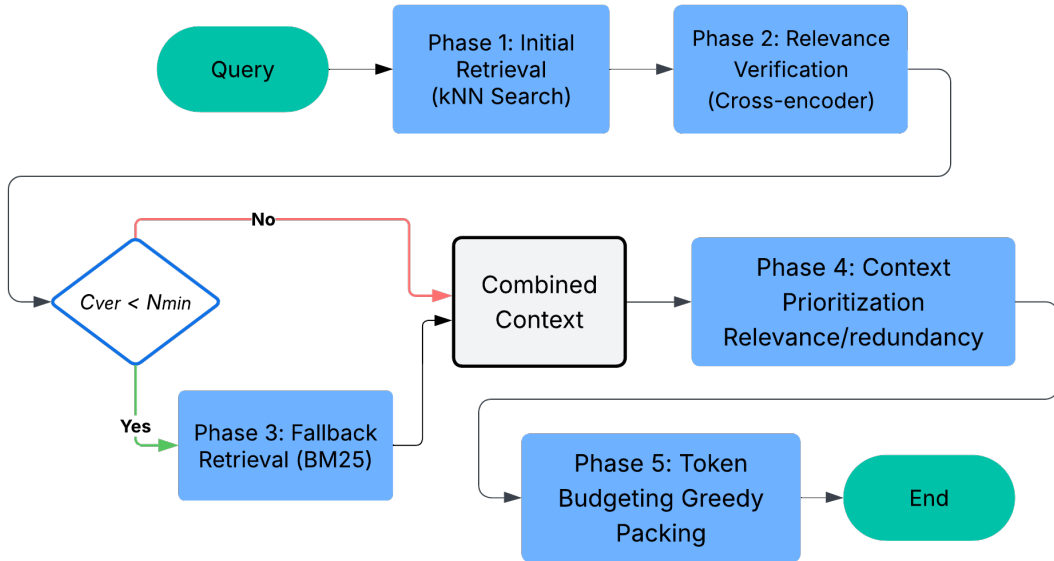


Fig. 1. The MeVe Architectural Framework: A five-phase modular pipeline for memory verification and context control.

Figure 1 illustrates the overall MeVe architecture. Each box represents a discrete processing phase that transforms the input query into a final optimized context for the LLM. The arrows indicate data flow through the system, highlighting conditional branches such as the fallback trigger after relevance verification. This visual layout emphasizes MeVe's separation of concerns: retrieval, verification, enrichment, prioritization, and token bud-

getting are handled independently, enabling fine-tuned control and interpretability across each stage.

3.1 Phase Breakdown

1. **Phase 1: Initial Retrieval (kNN Search):** The process begins with a user query fed into a k-Nearest Neighbors (kNN) search over a pre-constructed vector store [2] [15]. This produces an initial set of potentially relevant memory chunks based on dense similarity.
2. **Phase 2: Relevance Verification (Cross-Encoder):** A cross-encoder examines each candidate from Phase 1, generating a relevance score. Candidates scoring below a threshold are discarded.
3. **Phase 3: Fallback Retrieval (BM25):** If the number of verified candidates ($|C_{ver}|$) is below a minimum threshold (N_{min}), a keyword-based fallback mechanism (BM25) retrieves additional backup documents.
4. **Combined Context:** The verified and fallback documents are merged into a combined memory set that represents the candidate context.
5. **Phase 4: Context Prioritization:** This stage re-orders and filters the combined context to maximize informativeness and eliminate redundancy via embedding similarity.
6. **Phase 5: Token Budgeting (Greedy Packing):** A greedy packing algorithm selects and concatenates documents under a token budget constraint (T_{max}), forming the final context (C_{final}) passed to the LLM.

3.2 Phase 1: Preliminary Candidate Extraction

The initial phase performs an effective and exhaustive search for potentially useful memory chunks [20]. A user query q is first translated into a dense vector representation [20]. A k-Nearest Neighbors (kNN) search is then conducted over a pre-constructed vector store of encoded representations of the knowledge corpus. The step returns an initial candidate set $C_{init} = \{c_1, c_2, \dots, c_k\}$ based on a chosen similarity metric, e.g., cosine similarity [15]. This step is designed to have high recall to return a superset of potentially useful documents to be filtered later [19].

3.3 Phase 2: Verification of Relevance

This central module thoroughly examines each of the candidates from (C_{init}) in order to eliminate irrelevant information [9] [2]. For every candidate ($c_i \in C_{init}$) a specialized Relevance Verifier model (V) generates a context relevance score (s_i) [24] [2]. Candidates with scores below some selected threshold τ are consequently discarded, thereby guaranteeing that only sufficiently relevant knowledge is retained [9]. The resulting set of established memories is denoted as $C_{ver} = \{c_i \in C_{init} \mid V(q, c_i) = s_i \geq \tau\}$ [24] [2]. Implementation: We employ a light cross-encoder model for (V) in this proof-of-concept, trading off performance for latency in this proof-of-concept[24] [2].

3.4 Phase 3: Fallback Retrieval

The fallback retrieval phase improves MeVe’s resilience when relevance verification fails to produce enough valid results. If the number of verified candidates ($|C_{ver}|$) falls below a set threshold ($N_{min} = 3$ in our implementation), the system activates a secondary retrieval mechanism based on BM25Okapi. This ensures the LLM is rarely presented with an empty

or underpopulated context. While semantic search excels at conceptual retrieval, it may overlook documents with critical lexical features. BM25 complements this by matching on explicit keywords, enabling a hybrid strategy that combines semantic breadth with lexical precision [5].

3.5 Phase 4: Prioritization of Context

This phase strategically orders the collected memories to maximize information density and minimize redundancy before they are passed to the language model [12]. It operates on the complete set of memories, ($C_{all} = C_{ver} \cup C_{fallback}$) [19]. The prioritization follows a two-step process. First, all documents in C_{all} are sorted in descending order by their relevance score (s_i), which was either generated during Phase 2 (Relevance Verification) or assigned during fallback [9] [2]. Second, an iterative filtering step removes redundant information: a document is discarded if its embedding is highly similar (i.e., above a cosine similarity threshold, $\theta_{redundancy}$) to any higher-ranked document [15]. This process ensures the final prioritized list is both highly relevant and informationally diverse [12].

3.6 Phase 5: Token Budget Management

The last module of the MeVe pipeline concatenates the context with utmost caution in conformity with the pre-formulated token constraint of the target LLM, (T_{max}) [17]. Using a greedy packing algorithm, this phase incorporates the highest-priority segments of the given sequence obtained in phase 4 into the resulting context (C_{final}) [18]. This process continues until the target token limit is reached, thus ensuring that no valuable information is lost inadvertently as a result of naive truncation, and the large language model is given the most useful context within its capacity [17].

3.7 Computational Complexity

To further characterize the efficiency of the MeVe framework, we provide a theoretical analysis of the computational complexity of each modular phase. While our experimental results (Section 5) demonstrate minimal practical overhead, understanding the scaling behavior is crucial for broader applicability.

- **Phase 1: Preliminary Candidate Extraction (Dense Vector Search)** The time complexity for retrieving k candidate memories from a knowledge base of N documents, each with embedding dimension D , depends on the indexing structure. For a brute-force approach, this would be approximately $\mathcal{O}(N \cdot D)$. However, with optimized approximate nearest neighbor (ANN) search algorithms (e.g., FAISS, HNSW), the average complexity can be significantly faster, often closer to $\mathcal{O}(\log N)$ or $\mathcal{O}(k \cdot D_{query} + k \cdot D_{embedding})$ after initial indexing, making it highly efficient for large datasets.
- **Phase 2: Relevance Verification (Cross-Encoder)** This phase involves re-ranking the k initial candidate memories using a cross-encoder model. For each of the k memories, the cross-encoder processes the concatenated query and document, with a typical computational cost proportional to the sequence length L_{seq} (which is the sum of query length L_{query} and document length L_{doc}). Thus, the complexity is $\mathcal{O}(k \cdot L_{seq} \cdot \text{model_ops})$, where model_ops represents the operations within the transformer model (e.g., L_{seq}^2 for attention). Given that L_{seq} is usually capped, this scales linearly with the number of candidates k selected from Phase 1.

- **Phase 3: Fallback Retrieval (e.g., BM25)** If activated, this phase typically operates on an inverted index. While in the worst case (e.g., a query matching almost all documents) it could approach $\mathcal{O}(N \cdot L_{\text{doc}})$, in practice, for specific queries on a well-indexed corpus, its performance is often very fast, scaling more with the number of query terms and the average number of documents per term, or $\mathcal{O}(Q \cdot A_{\text{avg}})$ where Q is query terms and A_{avg} is average postings list length.
- **Phase 4: Prioritization of Context** This phase strategically orders the collected C_{ver} and C_{fallback} documents (let their total count be M). Sorting takes $\mathcal{O}(M \log M)$. The redundancy elimination step, which involves comparing embeddings, can be $\mathcal{O}(M^2 \cdot D_{\text{embedding}})$ in a naive pair-wise comparison, but can be optimized with clustering or approximate nearest neighbor search to reduce its practical impact. Given M is relatively small (e.g., up to 50-100 documents), this phase typically adds minimal overhead.
- **Phase 5: Token Budget Management** This phase uses a greedy packing algorithm that iterates through the prioritized documents and adds them until the token limit is reached. The complexity is primarily dependent on the number of tokens to be processed and the tokenization process itself, which is typically linear with the total input tokens to be packed, $\mathcal{O}(T_{\text{budget}})$. This is a highly efficient operation.

Overall, the dominant computational costs are typically in Phase 1 (initial retrieval from a large corpus) and Phase 2 (cross-encoder re-ranking). MeVe’s modularity ensures that more computationally intensive steps are applied only to a filtered subset of data, maintaining efficiency while significantly enhancing quality.

4 Experimental Setup

To validate the MeVe framework, we conducted a proof-of-concept evaluation on a focused subset of the English Wikipedia (the first 100 articles from the “20220301.en” edition) [21] [22]. The primary goal of this experiment is to demonstrate the architectural advantages and efficiency gains of MeVe’s modular design, rather than to achieve state-of-the-art factual accuracy on a competitive benchmark [23]. As such, this evaluation focuses on operational effectiveness and context efficiency [17]. The test compared three configurations. The first, Internal LLM Knowledge (No RAG), simulated an LLM answering only from its pre-trained knowledge and no external context [1]. The second, Standard RAG, was a plain top-k vector search retriever ($k=20$) where retrieved documents proceed directly to token budgeting without relevance checking, fallback, and prioritization [1]. The third was our full five-phase MeVe pipeline with a first-phase retrieval of $k=20$, relevance verification threshold of $\tau = 0.5$, BM25-like keyword fallback called under $N_{\text{min}} = 3$ confirmed documents, context prioritization redundancy threshold of 0.85, and final token budget of $T_{\text{max}} = 512$ [24] [9] [17]. In performance metrics, we considered three individual metrics: Context Efficiency (Tokens), the mean number of tokens in the final context; Retrieval Time (s) the average runtime of the context retrieval pipeline; and Context Grounding Proxy, a heuristic for assessing if a simulated answer can be inferred from the given context. Although this proxy is not a definite metric for factual correctness, in general, it suggests whether a response is generated from the given context [17].

5 Quantitative Results

This chapter presents the mean quantitative findings derived from the empirical assessment of the three operational modes of MeVe. These means offer a cumulative picture of the performance attributes observed across various configurations of the MeVe framework.

5.1 MeVe Improves Context Efficiency with Minimal Latency

MeVe consistently demonstrates a strong gain in context efficiency via this demonstration, with around a 57.7% reduction in the number of average tokens compared to Standard RAG (from 188.8 to 79.8 tokens). This finding substantiates the architectural capacity of the framework to generate a shorter context by way of the selective filtering and ranking of content [12]. MeVe’s retrieval time is also comparable to Standard RAG, which suggests the additional verification processes entail minimal overhead in this proof-of-concept [17]. These results are summarized in Table 1.

Table 1. Quantitative results demonstrating MeVe’s context efficiency and retrieval time.

Model	Dataset	Context Eff. (Tokens)	Time (s)
Internal LLM Knowledge	Wikipedia Subset	0	0.00
Standard RAG	Wikipedia Subset	188.8	1.12
MeVe (Ours)	Wikipedia Subset	79.8	1.22

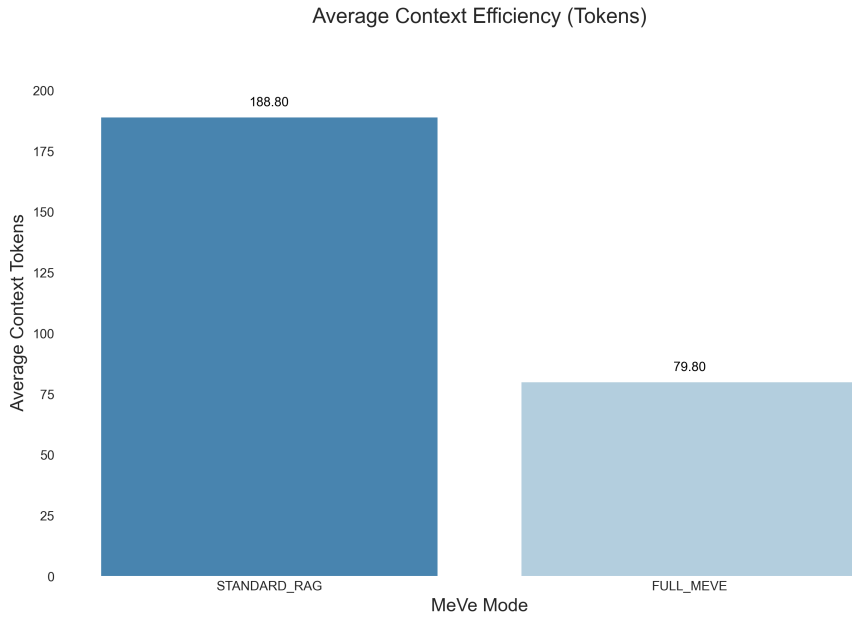


Fig. 2. Average context size (in tokens). MeVe significantly reduces the number of tokens passed to the LLM compared to a standard RAG baseline, demonstrating higher context efficiency.

5.2 Extended Evaluation Across Additional Datasets With Small Modifications:

Further evaluations using Exact Match and F1 Score across biomedical, legal, wikipedia and HotpotQA datasets confirm substantial gains. Compared to Standard RAG, MeVe increases Exact Match from 0.5 to 0.7, which is approximately 40% and an F1 score from 0.44 to 0.59 approximately 34%. Context efficiency also improved significantly, rising from an average of 230 to 330 tokens, around 43%, while the processing time is still low, going

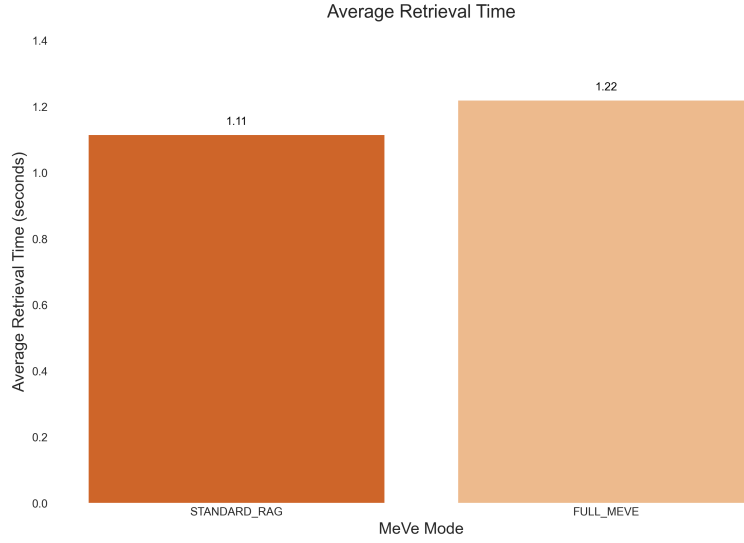


Fig. 3. Average retrieval time (in seconds). MeVe’s modular pipeline adds minimal latency overhead compared to the Standard RAG baseline, making the efficiency gains highly practical.

from 1,90 to 2,13 seconds (roughly 12%). This was all done with a slightly upgraded framework.

Building on the original framework, we extended the evaluation to cover domain-specific biomedical and legal QA datasets, including PubMedQA and LegalBench, to assess performance across diverse knowledge domains. Several enhancements were incorporated: the fallback strategy was improved by replacing the basic BM25-like retrieval with a robust hybrid mechanism using learning-to-rank, relevance verification was fine-tuned through adaptive thresholds and domain-specific cross encoders, and a semantic coherence filter was applied post-prioritization to ensure better context alignment. Additionally, we transitioned from proxy metrics to proper end-task evaluation, measuring Exact Match and F1 scores directly on real QA datasets. These changes were integrated into an enhanced evaluation pipeline with updated metrics, datasets, and visualization methods.

5.3 Qualitative Analysis of Retrieved Context

Whereas quantitative metrics offer an overview, qualitative observations yield more insight into the nature of each mode, specifically the significance of the context garnered and how it influences the simulated answers [10][11]. **Internal LLM Knowledge:** Predictably, this setting gave no external context, relying solely on the internal knowledge of a simulated LLM [1]. Responses, though notionally plausible for certain questions (e.g., "Paris" for French capital), were ungrounded in external sources, illustrating the fundamental limitation of parametric models without RAG [1].

Standard RAG: This mode tended to retrieve more context volume (as shown in Figure 2). However, the retrieved information, particularly with respect to specific fact-seeking questions, often showed semantic relevance but no factual relevance to the question at hand. For instance, with respect to the question regarding the building and height of the Eiffel Tower, the retrieved information often concerned other demolitions or other historical events that were not relevant, leading to an invented response based on non-relevant information [13]. This highlights the problem of context pollution which MeVe aims to alleviate [14].

Full MeVe: MeVe consistently gave a shorter context (as observed in Figure 2). The Relevance Verification (Phase 2) aggressive filtering typically produced zero verified documents for general knowledge questions, demonstrating its strict gatekeeper role [9] [2]. Therefore, the Fallback Retrieval (Phase 3) is frequently activated. While this strategy ensured that a baseline context was always provided, the content drawn from the BM25-like fallback mechanism was sometimes only tangentially relevant or entirely unrelated. As a result, although the simulated answers were technically derived from the provided context (based on keyword overlap), they often failed to be factually correct or semantically aligned with the original query. This highlights the fact that, although MeVe successfully enhances contextual efficiency, the quality of the first retrieval processes and fallback options, along with the correctness of the verification threshold, are paramount to guarantee the factual accuracy of the ultimate response [16].

These results highlight the efficiency of MeVe’s architectural components in handling and compacting context; however, the ultimate factual accuracy of the output still depends on the original relevance of the core knowledge base and the careful tuning of its varied filtering processes [16] [23].

5.4 Ablation Study: Validating MeVe’s Components

To empirically validate the contribution and necessity of MeVe’s individual modular components, we conducted an ablation study. This analysis specifically isolates the impact of key phases on context generation, efficiency, and simulated answer relevance, providing granular insights into their functional importance within the overall framework. This involves selectively disabling key phases to observe their quantitative impact on context generation, efficiency, and simulated answer relevance [19]. This analysis provides empirical proof of MeVe’s individual components’ necessity and utility [4].

The following setups were tried:

- **Full MeVe:** All five phases as given.
- **MeVe without Relevance Verification (NO_VERIFICATION):** Phase 2 (Relevance Verification) is bypassed. All initial candidates from Stage 1 are automatically considered ‘verified’ and are passed to the Fallback Retrieval stage (Phase 3) before proceeding to Context Prioritization. This simulates a RAG system without real relevance filtering.
- **MeVe without Fallback Retrieval (NO_FALLBACK):** Fallback Retrieval (Phase 3) is disabled. If Phase 2 returns too few confirmed documents, no fallback retrieval occurs, which may lead to a low or null context.

Figure 4 presents compelling proof of the gains in efficiency from MeVe’s modular design:

- The **NO_VERIFICATION** mode consistently exhibits a much higher average number of context tokens than **FULL_MEVE**. This empirical evidence validates the essential importance of Phase 2 (Relevance Verification) in efficiently eliminating irrelevant information and controlling context size, thus improving overall efficiency and reducing context window overflow [2] [6] [9].
- The **FULL_MEVE** retrieval time adds minimal additional latency over the less complex **STANDARD_RAG** benchmark (Figure 3). This additional latency is relatively small across the various ablated modes, suggesting that the computational costs of verification and prioritization, while measurable, do not dominate the overall retrieval time and represent an acceptable trade-off for the enhanced context quality control [17].

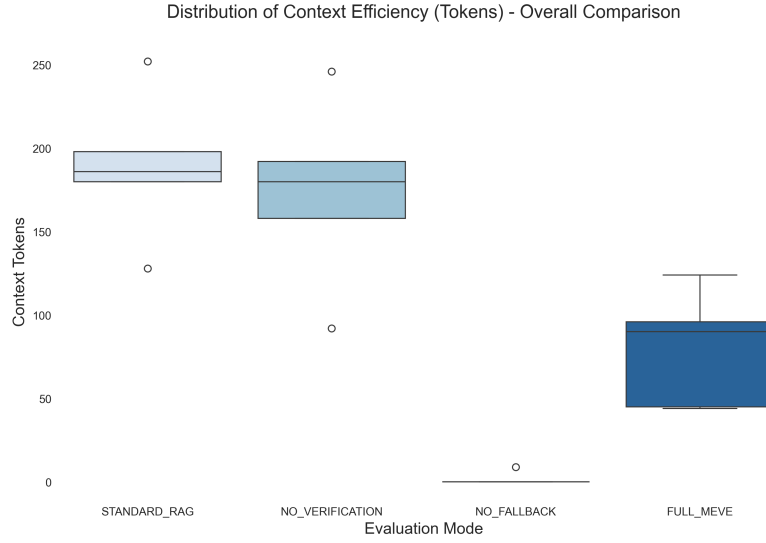


Fig. 4. Ablation study of context size and retrieval time. Disabling relevance verification (NO_VERIFICATION) dramatically increases context size, validating the module’s critical role in efficiency.

Figure 5 indicates that among all the retrieval-augmented modalities (**FULL_MEVE**, **STANDARD_RAG**, **NO_VERIFICATION**, **NO_FALLBACK**), a large proportion of responses are labeled as "Derived from Context (Often Irrelevant)" by our correctness proxy simulated. This finding speaks to a major difficulty in RAG systems, especially for general knowledge queries to a generic corpus, where it is particularly difficult to ground LLM outputs consistently in highly relevant and accurate external information [13][16].

Whereas Figure 5 is not indicating a radical move towards "Potentially Relevant" answers for **FULL_MEVE** under the specified simulation configuration, the findings presented in Figure 4 are very significant. The consistent labeling of "Often Irrelevant" for **NO_VERIFICATION** further underlines the reality that the mere introduction of contextual information without verification does not automatically enhance relevance but rather increases context pollution [14].

This ablation study definitively demonstrates how each MeVe phase critically contributes to efficient context management and mitigation of context pollution. While filtering improves efficiency, these findings also highlight the broader, inherent challenge of ensuring source-relevance when evaluating RAG systems against a generic corpus, even with robust internal controls [14] [16].

5.5 Evaluation on the HotpotQA Dataset

To further assess MeVe’s generalizability and robustness beyond single-fact retrieval from structured corpora like Wikipedia, we conducted an additional evaluation using a subset of the HotpotQA dataset 6. HotpotQA is distinct in that it primarily consists of multi-hop questions that necessitate reasoning over multiple retrieved documents to formulate a complete answer [25]. This characteristic provides a more rigorous test of MeVe’s ability to retrieve, verify, and effectively synthesize information from disparate sources, showcasing its adaptability for complex query types.

Our experiments on the HotpotQA dataset, encompassing the same experimental modes as our Wikipedia evaluation (Full MeVe, Standard RAG, No Verification, No Fallback, and Internal LLM Knowledge), yielded consistent trends regarding context efficiency

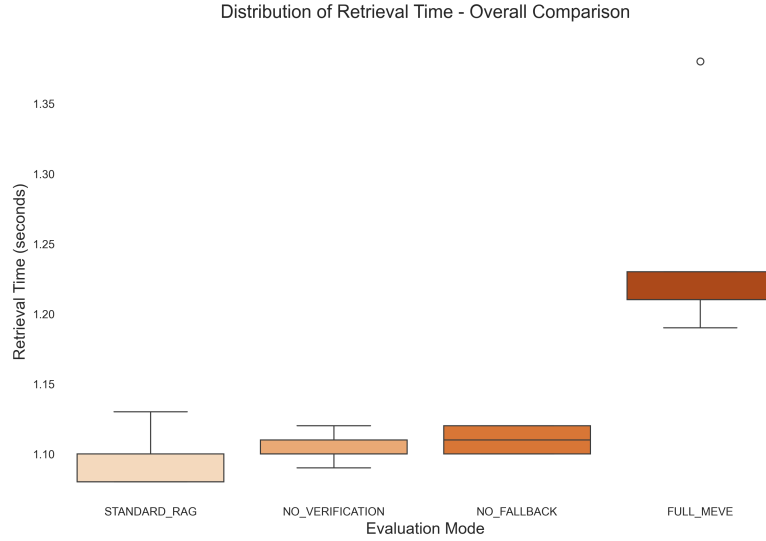


Fig. 5. Ablation study of Context Grounding Proxy. The prevalence of "Derived from Context (Often Irrelevant)" highlights that filtering improves efficiency but does not solve the underlying challenge of source-relevance in a generic corpus.

and retrieval time, further validating MeVe's architectural advantages across different knowledge domains and reasoning complexities.

Context Efficiency on HotpotQA Figure 6 illustrates the average context token count across the various modes for the HotpotQA dataset. Consistent with our findings on Wikipedia, Full MeVe demonstrated superior context efficiency, achieving an average context token count of 78.5 tokens. This represents a significant reduction of approximately 75% compared to Standard RAG, which averaged 308.6 tokens on HotpotQA. This substantial efficiency gain underscores the effectiveness of MeVe's relevance verification (Phase 2) and context prioritization (Phase 4) in filtering out irrelevant or redundant information, even in the context of multi-hop questions where context might inherently be more complex. The No Verification mode again exhibited context sizes comparable to Standard RAG (averaging 314.0 tokens), reinforcing the critical role of Phase 2 in achieving our efficiency goals. Conversely, the No Fallback mode, as expected, consistently produced zero context when initial retrieval and verification yielded no relevant documents, highlighting the necessity of Phase 3 for maintaining robustness against retrieval failures.

Retrieval Time on HotpotQA As depicted in Figure 7, the average retrieval times on the HotpotQA dataset largely mirrored those observed with the Wikipedia dataset.

Full MeVe's average retrieval time was 1.98 seconds, which remained competitive with Standard RAG's average of 1.80 seconds. This indicates that the computational overhead introduced by MeVe's additional phases (verification, fallback, prioritization, budgeting) is relatively small and does not significantly impede overall processing speed. This further supports that the gains in context quality and efficiency are achieved without substantial compromises in latency, making MeVe a practical solution for real-world applications requiring efficient and controlled context generation.

Answer Relevance Proxy on HotpotQA The analysis of the "answer relevance proxy" on HotpotQA continued to show a prevalence of responses labeled as "Derived from Con-

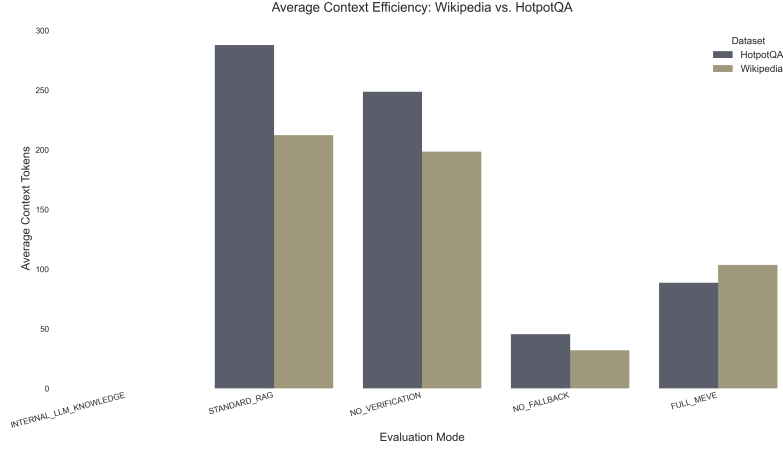


Fig. 6. Average Context Token Count per Mode on the HotpotQA Dataset. Full MeVe significantly reduces context token count compared to Standard RAG, demonstrating superior context efficiency and validating the effectiveness of its relevance verification and context prioritization phases.

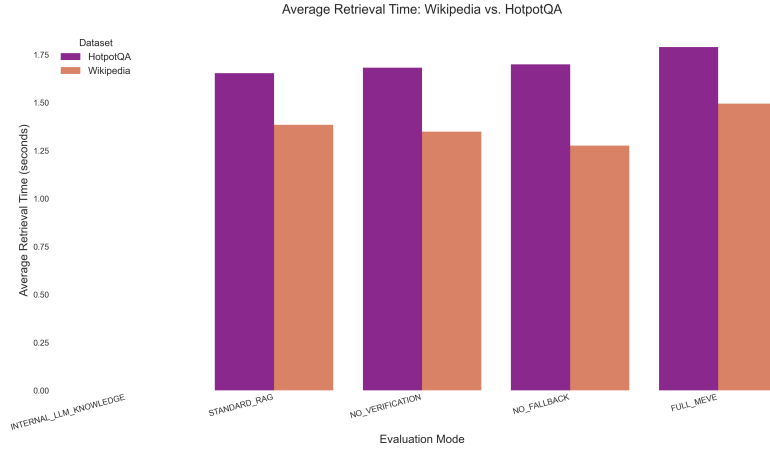


Fig. 7. Average Retrieval Time per Mode on the HotpotQA Dataset. Full MeVe maintains competitive retrieval times compared to Standard RAG, indicating that the computational overhead of its modular phases is minimal and does not impede overall processing speed.

text” across MeVe and Standard RAG modalities. This observation is consistent with the findings from the Wikipedia dataset, indicating that while the LLM is utilizing the provided context, the fundamental challenge of guaranteeing the semantic relevance and grounding of the information remains. This is particularly salient for complex, multi-hop queries where nuanced interpretation of context is critical. This reiterates that while MeVe’s architectural advantages enhance efficiency and control, the intricate problem of semantic grounding of LLM outputs requires further research, potentially through more sophisticated proxy measures that capture the degree of conceptual alignment with the query’s intent.

6 Discussion

MeVe’s modularity offers distinct advantages. The findings presented show that the separation of retrieval from verification and prioritization allows MeVe to generate a drastically more compact context [6] [17]. These findings support the central architectural claim that

modular processing yields measurable gains in efficiency and control [17]. Furthermore, being able to identify failures linked with a particular module (e.g., repeated failures in relevance verification as a result of an excessively stringent threshold or corpora constraints) is essential for effective debugging and trust establishment in complicated systems [16].

While MeVe showed significant context efficiency in the simulation performed, the "Context Grounding Proxy" often indicated that the LLM was still producing answers based on context that had minimal semantic connection to the question. This highlights the pragmatic challenges inherent in real-world Retrieval-Augmented Generation (RAG) applications, which contrast with the validity of the framework architecture [2] [16]:

- **Corpus Specificity:** While the general Wikipedia subset is diverse, it might not always give direct answers to very specific general knowledge questions in a form that is optimally structured for retrieval and verification by present models and standards [21][22]. This calls for a highly aligned and quality knowledge base designed for specific application domains [23].
- **Cross-Encoder Sensitivity and Contextual Coherence:** The ms-marco-MiniLM-L-6-v2 cross-encoder, when evaluated at a threshold of $\tau = 0.5$, was found to be quite stringent in this instance [24]. Although it does effectively minimize irrelevant contextual features, its stringent filtering process can result in over-filtering for some query-document pairs [9]. This indicates that an adjustment to this threshold or the application of a domain-based cross-encoder would yield more relevant verified documents without any sacrifice to the verification principle [2].
- Beyond document-level relevance, a key area for future development is assurance of collective coherence and direct relevance of the returned context [12]. This could come in the form of an additional post-prioritization filtering step or more sophisticated coherence scoring to ensure that the overall context block exposed to the LLM is well-aligned to the query intent and has little chance of including tangentially related or distracting content [10] [11] [6].
- **BM25-like Proxy Limitations:** Our rudimentary keyword-overlap-based BM25-like implementation (using BM25Okapi), although sufficient for proving the fallback mechanism, does not have the maturity and ranking capability of an actual BM25 implementation (e.g., full BM25Okapi) [8] [9]. This can produce the retrieval of keyword-matching but semantically distant documents during fallback, which are subsequently handed over to the LLM [8] [9]. Adapting a more sophisticated BM25 would further optimize this phase [9].

These findings do not diminish the architectural significance of MeVe but rather highlight the requirement for ideally tunable modules and an appropriate quality knowledge repository in order to attain high answer accuracy and hallucination resistance levels in practical applications [16]. The minimal overhead of latency attained (1.22 seconds on average for MeVe and 1.12 seconds for Standard RAG) is justified given the efficiency and controllability gains inherent in the framework.

7 Conclusion

We introduced MeVe, a modular architectural framework for memory verification and context management in LLMs [4] [19]. By decomposing the monolithic RAG pipeline into tunable, specialized modules, MeVe provides fine-grained control over context quality and efficiency [2] [17]. Our empirical results on a Wikipedia subset as a proof-of-concept empirically validate this approach via appreciable improvements in context efficiency compared to baseline retrieval methods [21].

Ultimately, MeVe is not just a complement to RAG but a reinvention of memory interaction in LLM systems [1] [2]. While RAG considers retrieval a monolithic operation, MeVe considers it an evolving and modular process with distinct phases for verification, fallback, and budget management [4] [19].

Separation of concerns allows for more efficiency, manageability, and comprehensibility in memory utilization, a capability that is essential to the construction of future agentic systems and long-context applications [5] [6]. We have demonstrated that the reliability of future AI systems hinges not only on the expansion of context windows but also on the design of intelligent and auditable memory control structures [7] [14]. MeVe offers a systematic and validated solution to this challenge.

A Appendix A: Methodological Specifications

This appendix provides detailed information pertaining to the hardware and software configuration, model numbers, dataset, and hyperparameters utilized in the empirical analysis of the MeVe framework.

A.1 Technical Framework and Software Infrastructure

The tests were conducted on a system featuring a NVIDIA GeForce RTX 3060 GPU (6 GB GDDR6 VRAM). CPU specifications include an AMD Ryzen 5 5600H. The system was configured with 12 GB of RAM.

The setup utilized was Python 3.12 software. The primary libraries and their versions (approximately at the time of implementation) are as follows:

- PyTorch: 2.5.1+cu121
- Hugging Face Transformers: 4.52.1
- Sentence-Transformers: 4.1.0
- NLTK: 3.9.1
- FAISS-CPU: 1.11.0
- TinyDB: 4.8.2
- rank_bm25: 0.2.2
- Datasets: 3.6.0
- NumPy: 2.2.4
- Pandas: 2.2.3
- Matplotlib: 3.10.1
- Seaborn: 0.13.2
- tqdm: 4.67.1

A.2 Models utilized

The MeVe framework’s components rely on pre-trained neural models:

- **Embedding Model (phases 1 and 4):**
 - Model: sentence-transformers/multi-qa-mpnet-base-dot-v1
 - Purpose: Maps queries and corpus documents to dense vector embeddings. Selected due to its strong performance on semantic search and question-answer tasks.
 - Dimensionality: 768 dimensions.
- **Relevance Verification (Phase 2):**
 - Model: cross-encoder/ms-marco-MiniLM-L-6-v2

- **Purpose:** An effective cross-encoder model accepting a (query, document) pair as input and producing a relevance score, which is crucial for the elimination of irrelevant information.
- **Function:** returns logits, which are additionally passed through a sigmoid function to yield a probability score from 0 to 1.
- **Tokenization Process for Budgeting Tokens (Phase 5):**
 - **Model:** gpt2
 - **Purpose:** Used exclusively for token counting to enforce the token budget (T_{max}). This ensures the final context fits within the target LLMs input limits. To handle text processing correctly, the tokenizer’s padding token was set to its end-to-sentence token.

A.3 Corpus Details

The knowledge base used in the empirical examination was based on:

- **Dataset:** wikipedia
- **Config:** 20220301.en (English Wikipedia, 2022/03/01 snapshot)
- **Split:** Train
- **Subset Size:** To facilitate illustration, the data was truncated to the first 100 articles.
- **Chunking Strategy:** Sentences were chunked with `nltk.sent_tokenize` to create fine-grained memory units. A sentence became a separate chunk with a unique ID and its original article title preserved as metadata.

The preprocessed corpus gave around 22,736 text chunks after filtering out empty strings and sentence tokenization.

A.4 MeVe Framework Parameters

The following hyperparameters and fixed parameters were used for the MeVe pipeline configuration:

- **Initial Retrieval Count (k):** 20. The number of candidate documents retrieved in Phase 1.
- **Relevance Threshold (τ):** 0.5. The minimum cross-encoder score for a document to be verified in Phase 2.
- **Minimum Verified Documents (N_{min}):** 3. The threshold below which Fallback Retrieval is triggered in Phase 3.
- **Redundancy Threshold ($\theta_{redundancy}$):** 0.85. The cosine similarity threshold used in Phase 4 to penalize redundant chunks.
- **Token Budget (T_{max}):** 512. The maximum token limit for the final context in Phase 5.
- **BM25-like Fallback Implementation:** `rank_bm25` library, i.e., BM25Okapi, was utilized for keyword fallback-based retrieval. The corpus was tokenized by splitting sentences based on whitespace and converting them to lowercase for BM25 indexing.

A.5 Simulation Details

LLM response generation was simulated as an alternative to actual LLM interaction:

- **No RAG Mode:** Produces a hardcoded placeholder answer (e.g., "Paris" for "French capital") or a generic placeholder, signaling reliance on internal model knowledge.

- **Default RAG and Full MeVe Modes:** use a basic heuristic of keyword overlap. The generated LLM "answer" is the sentence in the given context that has the most keywords in common with the question. This method is intended to indicate whether or not the answer could be located within the context, without ascertaining its factual correctness with respect to the actual world. Furthermore, this method categorizes the answer as being either "Derived from Context" or "Could not derive a direct answer from context."
- **Context Token Calculation:** Derived from tokenizing the concatenated context string using the gpt2 tokenizer.

References

1. Lewis, P., Perez, E., Piktus, A., Petroni, F., Karpukhin, V., Goyal, N., Küttler, H., Lewis, M., Yih, W. T., Rocktäschel, T., Riedel, S., & Kiela, D. (2020). Retrieval-augmented generation for knowledge-intensive NLP tasks. *Advances in Neural Information Processing Systems*, 33, 9459–9474.
2. Gao, Y., Xiong, Y., Gao, X., Jia, K., Pan, J., Bi, Y., Dai, Y., Sun, J., & Wang, H. (2023). Retrieval-augmented generation for large language models: A survey. *arXiv preprint arXiv:2312.10997*.
3. Mallen, A., Asai, A., Zhong, V., Das, R., Khashabi, D., & Hajishirzi, H. (2022). When not to trust language models: Investigating effectiveness of parametric and non-parametric memories. *arXiv preprint arXiv:2212.10511*.
4. Chen, H., Wang, Z., & Li, Y. (2023). Modular retrieval for generalization and interpretation. *arXiv preprint arXiv:2303.17612*.
5. Packer, C., Wooders, S., Lin, K., Fang, V., Patil, S. G., Stoica, I., & Gonzalez, J. E. (2023). MemGPT: Towards LLMs as operating systems. *arXiv preprint arXiv:2310.08560*.
6. Tay, Y., Dehghani, M., Rao, J., Fedus, W., Abnar, S., Chung, H. W., Narang, S., Yogatama, D., Vaswani, A., & Metzler, D. (2024). Beyond the limits: A survey of techniques to extend the context length in large language models. *arXiv preprint arXiv:2402.02244*.
7. Liu, H., Zaharia, M., & Abbeel, P. (2023). Ring attention with blockwise transformers for near-infinite context. *arXiv preprint arXiv:2310.01889*.
8. Weaviate. (2025, June 18). Keyword search (BM25). Retrieved from <https://weaviate.io/developers/weaviate/concepts/search/keyword-search>
9. Singh, S. P. (2024). Cross-Encoder models for enhanced search relevance: A multi-domain analysis of performance and applications. *International Journal of Research in Computer Applications and Information Technology*, 7(2), 2673–2685.
10. Filice, S., et al. (2025). The distracting effect: Understanding irrelevant passages in RAG. *arXiv preprint arXiv:2505.06914*.
11. CEUR Workshop Proceedings. (2025, June 18). How noise and distractors impact retrieval-augmented generation. *CEUR Workshop Proceedings*, Vol-3802.
12. Levy, S., Mazor, N., Shalmon, L., Hassid, M., & Stanovsky, G. (2025). More documents, same length: Isolating the challenge of multiple documents in RAG. *arXiv preprint arXiv:2503.04388*.
13. Cuconasu, F., Trappolini, G., Siciliano, F., Filice, S., Campagnano, C., Maarek, Y., Tonello, N., & Silvestri, F. (2024). The power of noise: Redefining retrieval for RAG systems. In *Proceedings of the 47th International ACM SIGIR Conference on Research and Development in Information Retrieval* (pp. 719–729). ACM.
14. Rose, L. T., & Fischer, K. W. (2011). Garbage In, Garbage Out: Having Useful Data Is Everything. *Measurement: Interdisciplinary Research and Perspectives*, 9(4), 222–226.
15. Shyani, M., & Naamad, Y. (2025). Filtered semantic search via vector arithmetic. *Submitted*.
16. Tonmoy, S. M. T. I., Zaman, S. M. M., Jain, V., Rani, A., Rawte, V., Chadha, A., & Das, A. (2024). A comprehensive survey of hallucination mitigation techniques in large language models. *arXiv preprint arXiv:2401.01313*.
17. Han, T., Wang, Z., Fang, C., Zhao, S., Ma, S., & Chen, Z. (2025). Token-budget-aware LLM reasoning. *arXiv preprint arXiv:2412.18547*.
18. Wang, Y. (2023). Review on greedy algorithm. *Theoretical and Natural Science*, 14(1), 233–239. <https://doi.org/10.54254/2753-8818/14/20241041>
19. Tcaciuc, C. C., Rege Cambrin, D., & Garza, P. (2025). Multi Stage Retrieval for Web Search During Crisis. *Future Internet*, 17(6), Article 239.
20. AIModels.fyi. (2024). *multi-qa-mpnet-base-dot-v1* model card. Retrieved from <https://www.aimodels.fyi/models/huggingFace/multi-qa-mpnet-base-dot-v1-sentence-transformers>

21. Kaffee, L. A., Razniewski, S., Vougiouklis, P., Beghaeiraveri, S. A. H., Gayo, J. E. L., Waagmeester, A., & Gray, A. J. (2024). Wikidata subsetting: Approaches, tools, and evaluation. *Semantic Web*, 15(2), 231–265.
22. Figshare. (2023, June 18). *English Wikipedia quality assessment dataset*. Retrieved from https://figshare.com/articles/dataset/English_Wikipedia_Quality_Assessment_Dataset/1375406
23. Eppalapally, S., et al. (2024). KaPQA: Knowledge-augmented product question-answering. *arXiv preprint arXiv:2407.16073*.
24. Hugging Face. (2025). *Cross-Encoder/ms-marco-MiniLM-L6-v2* model card. Retrieved from <https://huggingface.co/cross-encoder/ms-marco-MiniLM-L6-v2>
25. Yang, Z., Qi, P., Zhang, S., Bengio, Y., Cohen, W. W., Salakhutdinov, R., Manning, C. D. (2018). HotpotQA: A dataset for diverse, explainable multi-hop question answering. *arXiv preprint arXiv:1809.09600*.
26. Jin, Q., Dhingra, B., Liu, Z., Cohen, W., & Lu, X. (2019). PubMedQA: A Dataset for Biomedical Research Question Answering. *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*.
27. Guha, N., Nyarko, J., Ho, D. E., Ré, C., Chilton, A., Narayana, A., Chohlas-Wood, A., Peters, A., Waldon, B., Rockmore, D. N., Zambrano, D., Talisman, D., Hoque, E., Surani, F., Fagan, F., Sarfaty, G., Dickinson, G. M., Porat, H., Hegland, J., Wu, J., Nudell, J., Niklaus, J., Nay, J., Choi, J. H., Tobia, K., Hagan, M., Ma, M., Livermore, M., Rasumov-Rahe, N., Holzenberger, N., Kolt, N., Henderson, P., Rehaag, S., Goel, S., Gao, S., Williams, S., Gandhi, S., Zur, T., Iyer, V., & Li, Z. (2023). LegalBench: A Collaboratively Built Benchmark for Measuring Legal Reasoning in Large Language Models. *arXiv preprint arXiv:2308.11462*.

Authors

A. Ottem studied Media and Communication in high school and has completed a one-year program in Applied Machine Learning at Noroff. With a background in graphic design, his research interests include machine learning, memory-augmented models, and creativity-focused AI systems.