

Recent Advances in Retrieval-Augmented Generation: Hierarchical, Graph-Based and Memory-Oriented Retrieval Frameworks

Omkar Bhegade¹ and Dr. Geetanjali Kale²

Student, Computer Engineering, Pune Institute of Computer Technology, Pune, India ¹

Professor, Computer Engineering, Pune Institute of Computer Technology, Pune, India ²

Abstract: Retrieval-Augmented Generation (RAG) has emerged as a dominant framework for enhancing Large Language Models (LLMs) through external knowledge retrieval, improving factual grounding, reasoning quality, and adaptability to dynamic information. However, recent studies demonstrate that conventional Vanilla RAG architectures experience substantial limitations in long-context reasoning, multi-hop question answering, retrieval redundancy, hallucination control, and scalability. This review systematically analyzes recent advances across twelve representative RAG frameworks, including GraphRAG, TreeRAG, hierarchical retrieval systems, adaptive retrieval models, memory-organized retrieval, and hallucination-aware retrieval strategies. The reviewed literature highlights several important trends. First, retrieval structure significantly impacts downstream reasoning quality, particularly for complex multi-hop and narrative tasks. Second, graph-based and hierarchical retrieval approaches improve semantic connectivity, dependency restoration, and contextual coherence compared to flat retrieval pipelines. Third, adaptive and entropy-aware retrieval systems reduce unnecessary retrieval overhead while improving hallucination mitigation and evidence precision. Additionally, stateful memory-based systems demonstrate strong performance gains in long-context narrative reasoning through iterative retrieval and memory consolidation cycles. This review compares retrieval architectures, knowledge organization strategies, datasets, evaluation metrics, efficiency trade-offs, and hallucination mitigation techniques across recent RAG systems. The study further identifies key research gaps involving retrieval scalability, graph-construction overhead, provenance-aware reasoning, retrieval redundancy, and unified evaluation standards. Finally, the paper outlines future research directions toward efficient, trustworthy, adaptive, and semantically structured next-generation RAG systems suitable for complex real-world applications.

Keywords: Hierarchical Retrieval, Adaptive Retrieval, Hallucination Mitigation, Knowledge Graphs, Context Management, Stateful Retrieval, Memory-Augmented Retrieval, Semantic Retrieval, Retrieval Efficiency, Entropy-based Attention.

I. INTRODUCTION

Evolution of Retrieval-Augmented Generation

Large Language Models (LLMs) have demonstrated remarkable capabilities in natural language understanding and generation; however, their dependence on parametric memory often leads to factual inaccuracies, hallucinations, and outdated responses. Retrieval-Augmented Generation (RAG) emerged as an effective solution by integrating external knowledge retrieval with generative reasoning, enabling language models to access dynamically updated information during inference. Traditional Vanilla RAG architectures generally rely on dense vector similarity retrieval over flat document chunks followed by prompt augmentation for answer generation. While effective for single-hop factual question answering, these approaches frequently suffer from retrieval redundancy, context dilution, attention collapse in long contexts, and weak multi-hop reasoning performance [1].



Recent research has therefore shifted toward structured retrieval architectures capable of organizing information more semantically and contextually. Systems such as BEE-RAG introduced entropy-balanced attention allocation to improve document utilization under long-context scenarios [1]. Similarly, Bridge-RAG employed hierarchical abstract-tree retrieval combined with entity-aware indexing mechanisms for improved multi-hop reasoning efficiency [2]. These developments indicate that retrieval effectiveness depends not only on retriever quality but also on how retrieved knowledge is organized, prioritized, and consumed by the language model.

Furthermore, increasing context-window sizes in modern LLMs have not completely solved long-context reasoning challenges. Studies continue to show that naïve long-context retrieval often introduces semantic noise and reduces evidence grounding quality [1]. Consequently, modern RAG research increasingly focuses on graph retrieval, hierarchical traversal, adaptive retrieval triggering, and memory-oriented retrieval pipelines to improve reasoning robustness, contextual coherence, and retrieval efficiency.

Emergence of Structured and Graph-Based Retrieval

As RAG systems evolved, researchers recognized that flat document retrieval mechanisms were inadequate for complex reasoning tasks involving entity relationships, cross-document dependencies, and long-range semantic interactions. This limitation motivated the development of GraphRAG and hierarchical retrieval systems that explicitly utilize structural relationships among knowledge units. Graph-based retrieval frameworks organize retrieved information as entities, relations, communities, or semantic subgraphs, thereby enabling better reasoning path construction and multi-hop evidence aggregation [4].

Dynamic graph retrieval methods such as DA-RAG demonstrated that query-dependent community search provides more coherent and semantically complementary retrieval compared to static graph partitioning approaches [4]. Similarly, GNN-RAG introduced graph neural networks to perform learned reasoning over knowledge graphs prior to LLM inference, significantly improving retrieval precision and reducing unnecessary token consumption in multi-hop question answering tasks [5]. These systems highlight the importance of exploiting graph topology rather than treating retrieval as independent chunk matching.

Hierarchical retrieval frameworks have also gained significant attention. Models such as LeanRAG and HG-RAG combine multi-level retrieval traversal, semantic aggregation, and context compression to improve dependency restoration and reduce retrieval redundancy [6]. TreeRAG-oriented systems further demonstrate that hierarchical summarization and bottom-up knowledge aggregation can substantially improve long-context reasoning while minimizing vector database complexity [11].

Overall, these structured retrieval approaches collectively indicate a major paradigm shift in RAG research—from simple similarity retrieval toward topology-aware, hierarchy-aware, and reasoning-aware retrieval architectures optimized for contextual coherence and scalable inference.

Adaptive Retrieval and Hallucination-Aware Reasoning

Another major advancement in modern RAG systems involves adaptive retrieval and hallucination-aware generation strategies. Conventional RAG pipelines often retrieve supporting documents for every query regardless of uncertainty level, introducing unnecessary retrieval overhead and irrelevant contextual information. Recent adaptive retrieval frameworks instead attempt to dynamically determine when retrieval is actually needed and how evidence should be ranked for effective grounding [10].

SeaRAG represents one such approach by integrating entity-level and statement-level hallucination probing before triggering retrieval operations [10]. The framework selectively activates retrieval only under uncertainty conditions and reranks retrieved passages using entropy-aware scoring mechanisms. Experimental results demonstrate that adaptive retrieval not only improves factual grounding and hallucination mitigation but also significantly reduces retrieval frequency and computational overhead [10].

In parallel, cognitive-inspired memory-centric systems such as ComoRAG explored stateful retrieval mechanisms for narrative and inferential reasoning tasks [3]. Instead of performing isolated retrieval steps, ComoRAG maintains evolving contextual memory across iterative retrieval cycles, allowing improved narrative understanding and reasoning



continuity. The framework showed particularly strong improvements on long-context and story-level reasoning datasets, demonstrating that memory consolidation and iterative evidence refinement are critical for advanced reasoning scenarios [3].

Recent survey studies on hallucination mitigation further reinforce that retrieval architecture directly influences downstream factual reliability, citation grounding, and reasoning consistency [9]. These findings collectively emphasize that future RAG systems must integrate adaptive retrieval timing, memory organization, retrieval verification, and structured reasoning to achieve robust and trustworthy generative AI systems.

II. METHODOLOGIES

Recent Retrieval-Augmented Generation (RAG) research has introduced diverse retrieval architectures designed to improve reasoning quality, retrieval precision, hallucination mitigation, and computational efficiency. The reviewed studies collectively demonstrate a transition from flat retrieval pipelines toward structured, adaptive, hierarchical, and graph-aware retrieval frameworks.

A. BEE-RAG: Balanced Entropy Engineering for RAG

BEE-RAG proposed an entropy-engineered retrieval mechanism for long-context reasoning [1]. The framework computes document-level importance factors and redistributes attention weights to control entropy growth during inference. Instead of treating all retrieved chunks uniformly, the model dynamically amplifies important evidence while suppressing irrelevant contextual noise. This entropy-balanced attention allocation substantially improves long-context document utilization and multi-hop reasoning robustness [1].

Architecture

- Hybrid entropy-engineered Vanilla RAG architecture.
- Uses flat retrieved chunks with adaptive document-level importance weighting [1].

Workflow

- Query retrieval using BM25/E5/BGE retrievers.
- Importance factor (β) computed for retrieved documents.
- Entropy-balanced attention redistributes focus toward important chunks.
- LLM generates grounded response [1].

Novelty

- Introduced entropy-controlled attention allocation for long-context QA.
- Reduced attention dilution in large context windows.
- Stabilized token importance across long retrieval sequences [1].

Performance Claims

- Significant improvement on HotpotQA and 2WikiMultiHopQA.
- Strong robustness under noisy retrieval settings.
- Greater benefits observed as context length increased [1].

B. Bridge-RAG

Bridge-RAG introduced a hierarchical abstract-tree retrieval architecture combined with a Cuckoo Filter-based entity lookup mechanism [2]. The framework organizes retrieved knowledge into parent-child abstraction layers and performs hierarchical traversal during retrieval. Entity-guided indexing significantly reduces retrieval latency while preserving semantic context across long documents [2].

Architecture

- Hybrid TreeRAG with entity-indexed hierarchical retrieval.
- Uses abstract-tree hierarchy with Cuckoo Filter entity lookup [2].



Workflow

- Query → entity extraction → entity lookup → hierarchical traversal → top-k chunk retrieval → generation [2].

Novelty

- Combined hierarchical retrieval with $O(1)$ entity indexing.
- Parent-child traversal improved semantic expansion.
- Reduced graph traversal cost compared to GraphRAG [2].

Performance Claims

- Outperformed GraphRAG and Vanilla RAG on MedQA and AALCR.
- Reported $10\times$ – $500\times$ faster retrieval latency [2].

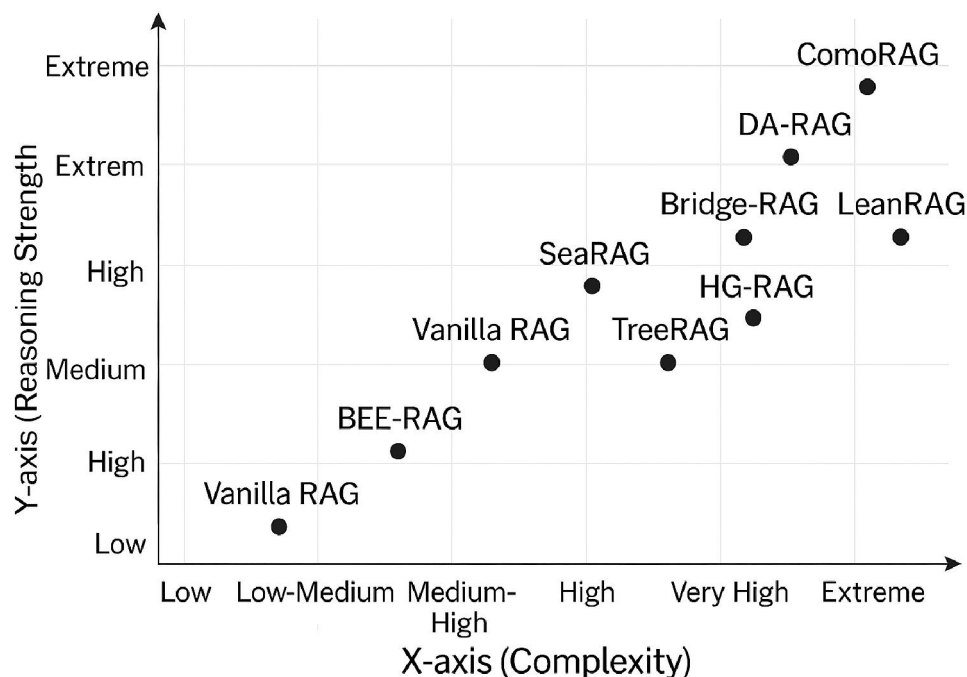


Fig. 1. Comparative Landscape of RAG Methods: Complexity vs. Reasoning Strength

C. ComoRAG

ComoRAG adopted a memory-organized retrieval methodology inspired by cognitive reasoning processes [3]. The framework performs iterative retrieval through probing-retrieval-memory cycles where retrieved evidence is consolidated into a dynamic memory workspace. Subsequent retrieval rounds use this accumulated contextual memory to guide future reasoning, enabling stateful narrative understanding and improved inferential reasoning [3].

Architecture

- Stateful memory-organized RAG framework.
- Includes dynamic memory workspace and semantic summary hierarchy [3].

Workflow

- Initial retrieval → memory encoding → probing-query generation → iterative retrieval cycles → memory fusion → answer generation [3].

Novelty

- Introduced “stateful reasoning” in RAG.
- Retrieval guided by evolving contextual memory instead of isolated steps.



- Narrative-flow summarization improved long-context understanding [3].

Performance Claims

- ~19% relative F1 improvement on narrative QA tasks.
- Strong gains on DetectiveQA and ∞ Bench benchmarks [3].

D. DA-RAG

DA-RAG formulated retrieval as an Embedding-Attributed Community Search problem over a multi-layer graph index. The framework performs coarse-to-fine retrieval by constructing dynamic query-relevant subgraphs and applying k-truss community constraints to ensure retrieval cohesion. This strategy improves semantic connectivity while reducing irrelevant node inclusion [4].

Architecture

- Dynamic GraphRAG using Embedding-Attributed Community Search (EACS).
- Multi-layer graph with semantic and knowledge graph layers [4].

Workflow

- Coarse chunk retrieval $\rightarrow$ dynamic subgraph construction $\rightarrow$ community search $\rightarrow$ final evidence assembly [4].

Novelty

- Dynamic query-specific graph communities.
- k-truss constraints improved graph cohesion and reduced semantic drift [4].

Performance Claims

- Improved comprehensiveness and retrieval relevance.
- Reduced indexing cost by up to 37%.
- Reduced token overhead by up to 41% [4].

E. GNN-RAG

GNN-RAG incorporated graph neural networks directly into the retrieval stage [5]. The system retrieves dense subgraphs from knowledge graphs, applies GNN-based node scoring for relevance estimation, and extracts compressed shortest reasoning paths before passing evidence to the LLM. This learned graph retrieval substantially improves multi-hop reasoning precision while reducing long-context token overhead [5].

Architecture

- GraphRAG integrated with Graph Neural Networks.
- Dense subgraph retrieval over knowledge graphs [5].

Workflow

- Entity linking $\rightarrow$ subgraph retrieval $\rightarrow$ GNN relevance scoring $\rightarrow$ shortest-path extraction $\rightarrow$ LLM reasoning [5].

Novelty

- Learned graph retrieval using deep GNNs.
- Compressed reasoning paths before generation [5].

Performance Claims

- +8.9% to +15.5% F1 improvement on multi-hop KGQA.
- Used approximately $9\times$ fewer KG tokens than long-context retrieval methods [5].

F. HG-RAG

HG-RAG combined hierarchical graph traversal with provenance-aware retrieval and semantic distillation [6]. The framework applies regex-enhanced retrieval, bounded multi-hop expansion, semantic compression, and dependency restoration to support compliance-sensitive technical reasoning tasks. Semantic distillation additionally reduces token redundancy while preserving contextual completeness [6].



Architecture

- Hierarchical GraphRAG with provenance-aware retrieval.
- Sliding-window graph indexing and semantic distillation [6].

Workflow

- Regex retrieval → provenance filtering → semantic retrieval → multi-hop expansion → distillation → response generation [6].

Novelty

- Dependency restoration across long technical documents.
- Provenance-aware filtering minimized cross-document interference [6].

Performance Claims

- Achieved 85.47% QA accuracy.
- Improved entity extraction coverage by ~45.9% [6].

G. KET-RAG

KET-RAG proposed a lightweight multi-granular GraphRAG architecture combining a knowledge-graph skeleton with a text-keyword bipartite graph [7]. Retrieval occurs through dual channels involving entity retrieval and keyword neighborhood expansion. The framework achieves strong retrieval coverage while significantly lowering graph indexing costs compared with full GraphRAG systems [7].

Architecture

- Multi-granular GraphRAG.
- Combines KG skeleton with text-keyword bipartite graph [7].

Workflow

- KNN graph construction → keyword/entity graph generation → dual-channel retrieval → local graph search → answer generation [7].

Novelty

- Cost-efficient lightweight GraphRAG.
- Split retrieval budget across entity and keyword channels [7].

Performance Claims

- Comparable performance to Microsoft GraphRAG with much lower indexing cost.
- Strong EM/F1 gains on MuSiQue and HotpotQA [7].

H. LeanRAG

LeanRAG emphasized topology-aware hierarchical retrieval using semantic aggregation and Lowest Common Ancestor (LCA) traversal [8]. Instead of retrieving disconnected nodes, the framework constructs compact semantically related subgraphs through bottom-up hierarchical exploration and inter-cluster relation expansion, thereby reducing retrieval redundancy and improving contextual coherence [8].

Architecture

- Hierarchical GraphRAG with semantic aggregation and LCA traversal [8].

Workflow

- Entity retrieval → hierarchical traversal → Lowest Common Ancestor path retrieval → compact subgraph generation [8].

Novelty

- Explicit inter-cluster semantic relations.
- Redundancy-aware retrieval compression [8].

Performance Claims

- Reduced retrieval redundancy by ~46%.
- Improved contextual coherence and answer diversity [8].



I. Attribution-Based Hallucination Mitigation Framework

The hallucination-mitigation survey systematically categorized attribution-oriented retrieval methodologies into query refinement, evidence identification, prompt engineering, and response correction strategies. The study highlighted how graph retrieval, subgraph retrieval, and hierarchical retrieval improve factual grounding for reasoning-intensive tasks [9]

Architecture

- Unified attribution-oriented retrieval pipeline [9].

Workflow

- Query refinement → evidence identification → structured prompting → response correction [9].

Novelty

- Taxonomy of retrieval-related hallucination types.
- Linked retrieval strategies to hallucination behavior [9].

Performance Claims

- Demonstrated structured retrieval improves factual grounding in complex QA tasks [9].

J. SeaRAG

SeaRAG introduced adaptive retrieval triggering using entity-level and statement-level uncertainty probing [10]. The framework selectively activates retrieval only when hallucination risk is detected and reranks candidate passages using entropy-based uncertainty reduction scores. This approach reduces unnecessary retrieval frequency while improving evidence precision and factual consistency [10].

Architecture

- Adaptive retrieval framework with uncertainty-aware retrieval triggering [10].

Workflow

- Initial generation → hallucination probing → uncertainty estimation → selective retrieval → entropy-ranked reranking → grounded regeneration [10].

Novelty

- Retrieval triggered only under hallucination risk.
- Joint entity-level and statement-level uncertainty probing [10].

Performance Claims

- Reduced retrieval frequency from 100% to 45.2%.
- Improved QA accuracy across HotpotQA and TriviaQA [10].

K. Tree-Based Implicit Knowledge RAG

Tree-based RAG architectures further explored hierarchical retrieval organization. The implicit knowledge TreeRAG framework generated bottom-up semantic summaries across files and folders to compress retrieval space while preserving cross-file contextual understanding [11]. Similarly, HiRMed employed hierarchical diagnostic reasoning through layered medical retrieval and memory-augmented refinement, enabling progressive context-preserving recommendation generation in healthcare applications [12].

Architecture

- TreeRAG with bottom-up semantic summarization [11].

Workflow

- File summaries → folder-level aggregation → hierarchical compression → vector retrieval [11].

Novelty

- Reduced retrieval footprint through implicit hierarchical summaries.
- Efficient cross-file contextual aggregation [11].

Performance Claims

- Reduced vector DB documents by ~68%.



- Improved folder-level EM from 0.54 to 0.81 [11].

L. HiRMed

Architecture

- Hierarchical TreeRAG for healthcare recommendation systems [12].

Workflow

- Symptom analysis → department routing → specialty reasoning → test recommendation generation [12].

Novelty

- Layered diagnostic reasoning with memory augmentation.
- Department-specific retrieval hierarchy [12].

Performance Claims

- Coverage: 92.3%
- Accuracy: 88.7%
- Miss Rate: 2.1% outperforming Flat-RAG and traditional similarity retrieval [12].

III. DISCUSSION

Evolution from Flat Retrieval to Structured Retrieval

One of the most significant observations across the reviewed literature is the gradual transition from traditional flat retrieval pipelines toward highly structured retrieval architectures. Early Vanilla RAG systems primarily relied on semantic similarity between user queries and independently stored document chunks. Although effective for factual and single-hop question answering, these systems struggled when reasoning required cross-document dependency restoration, multi-hop evidence aggregation, or long-context understanding [1].

Several reviewed studies demonstrated that retrieval organization itself strongly influences downstream reasoning quality. GraphRAG and TreeRAG systems introduced explicit structural relationships between retrieved knowledge units, enabling retrieval mechanisms to preserve semantic connectivity rather than retrieving isolated passages [4]. Hierarchical retrieval methods such as Bridge-RAG and LeanRAG further showed that parent-child abstraction hierarchies and topology-aware traversal improve contextual coherence while simultaneously reducing retrieval redundancy [2]. These findings indicate that retrieval effectiveness depends not only on embedding similarity, but also on how knowledge is structurally organized and traversed.

Another important trend is the emergence of topology-aware reasoning. GNN-RAG demonstrated that graph neural networks can learn retrieval relevance directly from graph neighborhoods and reasoning paths instead of depending solely on dense retrieval heuristics [5]. This approach substantially improved retrieval precision for multi-hop knowledge graph reasoning tasks. Similarly, dynamic graph community retrieval in DA-RAG reduced semantic drift by constructing query-specific cohesive subgraphs instead of static retrieval clusters [4].

Long-Context Reasoning and Context Management Challenges

Long-context reasoning emerged as a central challenge across nearly all reviewed papers. Merely increasing the context window of Large Language Models did not consistently improve reasoning performance because larger contexts often introduced attention dilution, redundant retrieval, and semantic distraction [1]. Several frameworks therefore focused specifically on context management strategies rather than retrieval expansion alone.

BEE-RAG addressed this issue using entropy-balanced attention redistribution to stabilize long-context document utilization [1]. The study demonstrated that attention entropy increases significantly as context length grows, reducing the model's ability to identify truly relevant evidence. Through controlled attention allocation, BEE-RAG improved multi-hop reasoning and noisy retrieval robustness, particularly on long-context benchmarks.

Likewise, hierarchical retrieval systems such as HG-RAG and LeanRAG introduced semantic distillation and redundancy-aware retrieval compression [6]. These approaches attempted to minimize unnecessary context tokens while preserving semantically connected reasoning paths. Findings from these studies suggest that effective context



compression and hierarchical summarization are critical for scaling future RAG systems under practical token limitations.

Memory-oriented frameworks also highlighted the importance of reasoning continuity. ComoRAG demonstrated that stateful retrieval and iterative memory consolidation substantially outperform isolated one-shot retrieval pipelines for narrative and inferential reasoning tasks [3]. This evidence supports the broader argument that future RAG systems may increasingly resemble dynamic reasoning agents rather than static retrieval pipelines.

Hallucination Mitigation and Adaptive Retrieval

Hallucination mitigation remains a major concern in Retrieval-Augmented Generation research. The reviewed survey on attribution techniques emphasized that hallucinations originate not only from generation failures, but also from retrieval deficiencies such as irrelevant evidence, incomplete context, or retrieval-generator mismatch [9]. Consequently, recent frameworks increasingly integrate verification-aware retrieval mechanisms.

SeaRAG introduced an important shift toward uncertainty-aware retrieval triggering [10]. Rather than retrieving evidence for every query, the framework selectively activated retrieval only when hallucination risk was detected. This adaptive design improved both efficiency and factual grounding while reducing retrieval overhead. The results suggest that retrieval should be viewed as a dynamic decision-making process rather than a permanently active pipeline component.

Another important observation concerns the relationship between task complexity and retrieval sophistication. Multiple studies showed that Vanilla RAG remains competitive for simpler factual QA tasks [9]. However, retrieval complexity becomes increasingly beneficial for:

- multi-hop reasoning,
- narrative understanding,
- hierarchical dependency restoration,
- domain-intensive QA,
- and long-context semantic synthesis.

This indicates that future RAG research may require adaptive retrieval architectures capable of dynamically selecting retrieval depth, retrieval structure, and reasoning strategy based on query complexity.

Efficiency, Scalability, and Practical Trade-offs

Although GraphRAG and hierarchical retrieval systems consistently improved reasoning quality, many studies also highlighted substantial computational trade-offs. Graph construction, entity extraction, hierarchical indexing, and multi-hop traversal introduced significant indexing and retrieval overhead [7]. Some frameworks attempted to reduce these limitations using lightweight graph skeletons, keyword graphs, semantic aggregation, and selective traversal strategies. KET-RAG demonstrated that multi-granular lightweight graph retrieval could preserve reasoning performance while drastically reducing indexing costs compared with full-scale knowledge graph extraction pipelines [7]. Similarly, LeanRAG reduced retrieval redundancy by nearly half using Lowest Common Ancestor traversal and compact evidence subgraphs [8].

Overall, the literature suggests that future RAG systems must balance retrieval quality, retrieval cost, scalability, and token efficiency simultaneously. The reviewed studies collectively establish that no single retrieval architecture is universally optimal. Instead, retrieval complexity should align closely with task complexity, reasoning depth, domain requirements, and operational resource constraints.

IV. FUTURE RESEARCH DIRECTIONS

Adaptive Retrieval Architectures

Future RAG systems should move toward fully adaptive retrieval pipelines capable of dynamically selecting:

- retrieval depth,
- retrieval frequency,



- reasoning pathways,
- and retrieval granularity

based on query complexity and uncertainty estimation [10].

Current systems still rely heavily on static retrieval configurations, even though multiple studies demonstrate that simple factual queries often do not require expensive graph or hierarchical retrieval mechanisms [9].

Efficient Graph Construction and Scalability

GraphRAG systems consistently improve multi-hop reasoning performance; however, graph construction and maintenance remain computationally expensive [7]. Future research should focus on:

- lightweight graph indexing,
- incremental graph updates,
- scalable entity extraction,
- dynamic graph compression,
- and low-cost graph maintenance pipelines.

Efficient graph-learning approaches using sparse subgraph construction or selective graph generation may reduce deployment overhead in industrial-scale applications [5].

Long-Context Reasoning and Memory Integration

Long-context reasoning remains one of the most challenging issues in RAG systems [1]. Future systems should incorporate:

- persistent memory modules,
- dynamic context compression,
- retrieval-aware attention routing,
- and iterative reasoning loops.

Stateful memory architectures similar to ComoRAG indicate that maintaining evolving contextual understanding can significantly improve narrative and inferential reasoning performance [3].

Retrieval Redundancy Reduction

Several studies identified retrieval redundancy as a major contributor to token inefficiency and semantic distraction [8]. Future work should investigate:

- topology-aware retrieval pruning,
- semantic compression,
- adaptive chunk sizing,
- and redundancy-aware evidence fusion.

Context optimization techniques will become increasingly important as models operate on larger multi-document corpora.

Hallucination Detection and Trustworthy RAG

Although retrieval improves factual grounding, hallucination problems persist even in advanced RAG systems [9]. Future research should explore:

- real-time hallucination detection,
- provenance-aware retrieval,
- evidence verification loops,
- citation-grounded generation,
- and uncertainty-aware retrieval triggering.

Integrating retrieval validation directly into generation pipelines may improve trustworthiness in high-risk applications such as healthcare, legal reasoning, and industrial systems [10].



Hybrid Retrieval and Multi-Modal RAG

Most reviewed systems primarily focused on textual retrieval. Future RAG architectures should support:

- multimodal retrieval,
- table reasoning,
- diagram understanding,
- and structured sensor/document integration.
- Hybrid retrieval systems combining:
 - graphs,
 - hierarchies,
 - dense retrieval,
 - and memory-based reasoning

may provide better robustness for real-world enterprise and scientific applications [6].

V. CONCLUSION

Retrieval-Augmented Generation has rapidly evolved from simple similarity-based document retrieval toward highly structured, adaptive, and reasoning-aware retrieval architectures. The reviewed studies collectively demonstrate that retrieval structure plays a critical role in improving reasoning quality, factual grounding, contextual coherence, and retrieval efficiency across complex question-answering tasks. While traditional Vanilla RAG systems remain effective for straightforward factual queries, they consistently struggle with long-context reasoning, multi-hop evidence aggregation, hallucination mitigation, and semantic dependency preservation [1].

GraphRAG and hierarchical retrieval frameworks have emerged as promising solutions for these limitations. Systems such as DA-RAG, GNN-RAG, LeanRAG, and HG-RAG demonstrated that graph topology, semantic hierarchy, and relationship-aware traversal significantly enhance contextual reasoning and retrieval precision [4]. These approaches improve multi-hop evidence connectivity while reducing semantic drift and irrelevant retrieval noise. Similarly, TreeRAG systems showed that hierarchical semantic aggregation and structured retrieval are highly effective for long-context and cross-document reasoning tasks [11].

Another important trend observed across the literature is the increasing importance of adaptive and stateful reasoning. Models such as ComoRAG and SeaRAG showed that retrieval effectiveness depends not only on retrieval quality, but also on retrieval timing, contextual memory integration, and uncertainty-aware evidence selection [3]. Adaptive retrieval mechanisms reduce unnecessary computational overhead while improving hallucination detection and factual consistency. These studies indicate a broader shift from static retrieval pipelines toward intelligent retrieval orchestration frameworks capable of dynamic reasoning management.

The reviewed literature also highlights several unresolved challenges. Graph construction cost, indexing complexity, retrieval redundancy, token inefficiency, and hallucination propagation remain major limitations in current systems [7]. While sophisticated retrieval architectures improve reasoning quality, they often introduce scalability and maintenance overhead that may limit real-world deployment. Consequently, achieving an effective balance between retrieval quality, efficiency, scalability, and interpretability remains an open research problem.

Overall, the collected evidence strongly suggests that future RAG systems will increasingly combine:

- structured graph retrieval,
- hierarchical traversal,
- adaptive retrieval triggering,
- iterative memory consolidation,
- and provenance-aware reasoning.

Such architectures are likely to become foundational components of next-generation trustworthy AI systems capable of reliable long-context reasoning, complex multi-hop inference, and domain-specialized knowledge synthesis across industrial, scientific, and enterprise applications.



VI. ACKNOWLEDGMENT

The authors would like to acknowledge the contributions of the researchers and academic communities whose published works formed the foundation of this review study. Special appreciation is extended to the developers of recent RAG architectures, GraphRAG systems, TreeRAG frameworks, and hallucination mitigation methodologies that significantly advanced retrieval-augmented reasoning research. The insights obtained from these studies greatly contributed to the comparative analysis and synthesis presented in this paper.

REFERENCES

- [1]. Y. Wang, R. Ren, Y. Wang, J. Liu, W. X. Zhao, H. Wu, and H. Wang, "BEE-RAG: Balanced Entropy Engineering for Retrieval-Augmented Generation," in Proc. AAAI Conf. Artif. Intell. (AAAI), vol. 40, no. 40, pp. 33737–33745, 2026.
- [2]. Z. Li, W. Liu, Y. Zong, J. Tao, S. Dai, S. Ren, Z. Liu, Y. Wang, Y. Jiang, and T. Yang, "Bridge-RAG: An Abstract Bridge Tree Based Retrieval Augmented Generation Algorithm With Cuckoo Filter," arXiv preprint arXiv:2603.26668, 2026.
- [3]. J. Wang, R. Zhao, W. Wei, Y. Wang, M. Yu, J. Zhou, J. Xu, and L. Xu, "ComoRAG: A Cognitive-Inspired Memory- Organized RAG for Stateful Long Narrative Reasoning," in Proc. AAAI Conf. Artif. Intell. (AAAI), vol. 40, no. 39, pp. 33557–33565, 2026.
- [4]. X. Zeng, Z. Wu, Y. Wang, C. Zhang, Q. Yao, L. Zheng, and J. Yin, "DA-RAG: Dynamic Attributed Community Search for Retrieval-Augmented Generation," in Proc. ACM Web Conf. (WWW), pp. 2195–2206, 2026.
- [5]. C. Mavromatis and G. Karypis, "GNN-RAG: Graph Neural Retrieval for Efficient Large Language Model Reasoning on Knowledge Graphs," in Findings Assoc. Comput. Linguistics: ACL, pp. 16682–16699, 2025.
- [6]. Z. Shen, X. Cai, B. Ni, Z. Meng, Z. Huang, and T. Yu, "HG-RAG: Hierarchical Graph-Enhanced Retrieval-Augmented Generation for Power Systems," Electronics, vol. 15, no. 7, p. 1445, 2026.
- [7]. Y. Huang, S. Zhang, and X. Xiao, "KET-RAG: A Cost-Efficient Multi-Granular Indexing Framework for Graph-RAG," in Proc. ACM SIGKDD Conf. Knowl. Discovery Data Mining (KDD), 2025.
- [8]. Y. Zhang, R. Wu, P. Cai, X. Wang, G. Yan, S. Mao, D. Wang, and B. Shi, "LeanRAG: Knowledge-Graph-Based Generation with Semantic Aggregation and Hierarchical Retrieval," in Proc. AAAI Conf. Artif. Intell. (AAAI), vol. 40, no. 41, pp. 34862–34869, 2026.
- [9]. Y. Zhao, Z. Liu, Y. Zheng, and K.-Y. Lam, "Attribution Techniques for Mitigating Hallucinated Information in RAG Systems: A Survey," arXiv preprint arXiv:2601.19927, 2026.
- [10]. D. Yu, Q. Lin, Z. Yang, and L. Zhou, "SeaRAG: Semantic Entropy-Guided Adaptive Retrieval for Multi-Hop Question Answering," in Proc. IEEE Int. Conf. Acoust., Speech Signal Process. (ICASSP), 2026.
- [11]. M. Gupte, P. Giusto, and R. S., "Is Implicit Knowledge Enough for LLMs? A RAG Approach for Tree-Based Structures," arXiv preprint arXiv:2510.10806, 2025.
- [12]. Y. Yang and C. Huang, "A Tree-based RAG-Agent Recommendation System: A Case Study in Medical Test Data," arXiv preprint arXiv:2501.02227, 2025.

