

A Multi-Agent Large Language Model Framework for Automated Scientific Literature Review and Knowledge Synthesis

Daniel R. Collins

Department of Computer Science
University of Massachusetts Lowell
Lowell, MA, USA
daniel.collins@uml.edu

Priya N. Raman

School of Computing and Information
University of Pittsburgh
Pittsburgh, PA, USA
priya.raman@pitt.edu

Abstract

The growth of scientific publications has made comprehensive literature review increasingly difficult for individual researchers and conventional search systems. This paper presents ScHoLARSyNTH, a multi-agent large language model (LLM) framework for automated scientific literature review and knowledge synthesis. The framework decomposes review generation into coordinated specialist agents for query planning, paper retrieval, evidence extraction, claim verification, thematic clustering, and synthesis writing. Unlike single-prompt LLM summarization, the proposed design maintains explicit evidence provenance, performs cross-paper consistency checks, and separates exploratory retrieval from synthesis decisions. We evaluate ScHoLARSyNTH on a curated benchmark of 4,200 papers from machine learning, biomedical informatics, and human-computer interaction, with 120 review tasks derived from survey article abstracts and section headings. Experimental results show that ScHoLARSyNTH improves answer faithfulness from 0.76 to 0.89 and topical coverage from 0.71 to 0.84 compared with a strong retrieval-augmented generation baseline, while reducing unsupported claims by 43.8%. The system processes 100 candidate papers in 7.8 minutes on average, enabling interactive literature scoping. Ablation studies demonstrate that verifier and clustering agents contribute most strongly to factual precision and structural coherence. The findings indicate that multi-agent orchestration is a practical approach for transforming LLMs from generic summarizers into auditable research assistants for scientific knowledge synthesis.

1 Introduction

The scientific record is expanding at a rate that challenges traditional literature review practices. Major bibliographic databases index millions of new records annually, while fast-moving areas such as artificial intelligence, biomedicine, materials science, and climate modeling often produce hundreds of relevant papers within months. For researchers, the problem is not merely retrieving documents but organizing heterogeneous evidence, identifying methodological differences, reconciling contradictory findings, and producing a coherent account of what is known and unknown. Manual review remains the gold standard because it enables expert judgment; however, it is slow, difficult to reproduce, and vulnerable to selection bias.

Recent large language models (LLMs) have demonstrated strong capabilities in summarization, question answering, and scientific text processing [1, 2, 3]. Retrieval-augmented generation (RAG)

systems further improve factual grounding by conditioning generated text on external documents [4]. Nevertheless, applying a single LLM call or a monolithic RAG pipeline to literature review is insufficient. Literature synthesis requires multiple interacting cognitive operations: formulating search strategies, screening documents, extracting comparable evidence, grouping papers by method and claim, checking whether generated statements are supported by sources, and writing an integrated narrative. A monolithic model often conflates these operations, producing fluent but poorly traceable summaries.

This paper proposes ScHoLARSyNTH, a multi-agent LLM framework for automated scientific literature review and knowledge synthesis. The central hypothesis is that review quality improves when the workflow is decomposed into specialized agents with explicit interfaces and shared memory. The framework uses a planner agent to decompose a review question into search facets, a retrieval agent to collect candidate papers, a screening agent to score relevance, an extraction agent to identify structured evidence, a clustering agent to group papers into conceptual themes, a verifier agent to audit claims against cited evidence, and a synthesis agent to produce the final review. Each agent operates over a typed state representation rather than free-form chat history, reducing drift and preserving provenance.

The contributions of this paper are fourfold. First, we introduce a complete multi-agent architecture for scientific literature review with explicit evidence records and claim-level verification. Second, we formalize the synthesis process as an optimization problem balancing coverage, faithfulness, diversity, and redundancy. Third, we construct a realistic evaluation benchmark containing 120 literature review tasks across three scientific domains. Fourth, we provide empirical comparisons against keyword search, single-agent LLM summarization, and RAG baselines, including throughput, latency, and ablation analyses.

2 Related Work

Automated literature review has a long history in information retrieval and text mining. Early systems focused on citation analysis, topic modeling, and bibliographic clustering [5]. Systematic review tools later incorporated screening prioritization and semi-automated evidence extraction, particularly in medicine and public health [6]. Broader surveys of systematic review automation emphasize that screening, extraction, synthesis, and update monitoring remain distinct technical bottlenecks [22, 23]. These systems improved search efficiency but typically required manually engineered features or domain-specific ontologies.

Neural language models changed the design space by enabling semantic retrieval and abstractive summarization. Transformer architectures [7] and pretrained encoders such as BERT [8] improved document ranking, sentence classification, and scientific named entity recognition. Domain-adapted models, including SciBERT [9], showed that scientific corpora have distinctive vocabulary and discourse patterns requiring specialized representations. More recently, LLMs have been applied to scientific question answering, experimental planning, and paper summarization [1, 2].

RAG combines parametric generation with non-parametric document retrieval, reducing hallucination by grounding outputs in retrieved passages [4]. Subsequent work on open-domain question answering, retrieval calibration, and long-context scientific summarization further shows that retrieval quality, evidence selection, and source attribution strongly affect generated answer reliability

[15, 16, 17, 18]. In scientific settings, RAG is attractive because it can cite evidence and incorporate recent papers without retraining. However, standard RAG often optimizes local relevance rather than global synthesis. It may retrieve several near-duplicate documents, overlook minority methods, or produce statements that are only weakly supported by the cited passage. These issues are serious in literature review, where representativeness and methodological nuance are essential.

Multi-agent LLM systems address complex tasks by assigning roles to multiple model instances that communicate through structured protocols. Prior work has explored self-reflection, debate, tool use, collaborative problem solving, software-engineering agents, and role-playing societies [10, 11, 12, 13, 19, 20, 21]. Wang et al. [14] recently proposed a hierarchical multi-agent collaboration method that separates task decomposition, specialist reasoning, verification, and decision fusion, and further uses semantic communication and dynamic feedback to reduce redundant interaction in complex task solving. The key insight is that role specialization can reduce cognitive load and improve error detection. For scientific review, agent specialization is especially natural because human review teams already distribute tasks such as screening, extraction, adjudication, and writing. ScHoLARSyNTH builds on this idea but differs from generic multi-agent collaboration frameworks by grounding agent communication in paper-level evidence schemas, claim-level provenance, and quantitative evaluation against review-specific metrics.

3 Methodology

Let a review task be represented as $q = (t, D, S)$, where t is a natural-language topic, D is a scientific domain, and $S = \{s_1, \dots, s_m\}$ is an optional set of desired review sections. The system receives access to a paper corpus $P = \{p_1, \dots, p_n\}$. Each paper p_i contains metadata, abstract, section text when available, citation information, and vector embeddings for passages.

The objective is to generate a synthesis y and an evidence map E_y such that y maximizes four criteria: topical coverage, factual faithfulness, diversity of evidence, and conciseness. We write the target objective as

$$\max_y \lambda_1 C(y, q) + \lambda_2 F(y, E_y) + \lambda_3 V(E_y) - \lambda_4 R(y), \quad (1)$$

where C measures coverage of expected concepts, F measures entailment between claims and evidence, V measures diversity across venues, years, methods, and research groups, and R penalizes redundancy. In the implementation, these terms are approximated using automatic metrics and verifier judgments rather than directly optimized end-to-end.

The framework maintains an evidence record $e = (p, a, c, r)$, where p is a paper identifier, a is an extracted atomic claim, c is contextual metadata such as method or dataset, and r is a relevance score. Atomic claims are extracted from abstracts and results sections using constrained prompts that require the agent to output JSON-like records. Claims are then normalized into a canonical schema containing problem, method, dataset, metric, result, limitation, and citation fields.

The multi-agent workflow is described in Algorithm 1. Agents communicate through a shared blackboard containing query facets, candidate papers, evidence records, clusters, draft paragraphs, and verification reports. The blackboard prevents hidden state from dominating decisions and allows each stage to be audited.

Complexity is dominated by retrieval, extraction, and verification. For k query facets over sparse

Algorithm 1 Multi-agent literature synthesis in ScHoLARSyNTH

Require: Review topic t , domain D , corpus P , section plan S

Ensure: Synthesized review y with evidence map E_y

- 1: Planner decomposes t into search facets $Q = \{q_1, \dots, q_k\}$
 - 2: Retrieval agent obtains candidate set $C \subset P$ using hybrid BM25 and dense retrieval
 - 3: Screening agent assigns relevance score r_i to each $p_i \in C$
 - 4: Select top- K papers C_K subject to year and venue diversity constraints
 - 5: **for** each paper $p_i \in C_K$ **do**
 - 6: Extraction agent produces atomic evidence records E_i
 - 7: Normalizer maps E_i into the shared evidence schema
 - 8: **end for**
 - 9: Clustering agent groups evidence into themes $G = \{g_1, \dots, g_l\}$
 - 10: Synthesis agent drafts section-wise narrative using G and S
 - 11: Verifier checks each generated claim against E_y and flags unsupported claims
 - 12: Revision agent edits draft until unsupported-claim rate is below threshold τ
 - 13: **return** final synthesis y and evidence map E_y
-

and dense indices, hybrid retrieval requires

$$O(k(n_b + n_d)) \quad (2)$$

index lookups. Extraction is $O(KL)$ LLM tokens for K selected papers with average length L . Verification is $O(Mh)$, where M is the number of generated claims and h is the number of evidence snippets checked per claim. Because extraction and verification are embarrassingly parallel across papers and claims, wall-clock latency scales approximately with the slowest batch rather than total token volume when sufficient API or GPU concurrency is available.

4 System Architecture / Model Design

Figure 1 presents the architecture. The system contains five layers: data ingestion, retrieval, agent orchestration, evidence memory, and synthesis interface. The ingestion layer parses PDF text, abstracts, metadata, and references. The retrieval layer combines BM25 keyword retrieval with dense passage search using scientific sentence embeddings. Reciprocal rank fusion is used to merge rankings because it is robust to score calibration differences.

The orchestration layer implements seven agents. Consistent with recent evidence that structured role separation and verifier-guided feedback improve complex LLM task solving [14], the planner expands the topic into Boolean-style facets and semantic paraphrases. The retrieval agent queries both abstract-level and passage-level indices. The screening agent predicts whether each paper should be included, excluded, or reserved for background. The extraction agent converts relevant passages into atomic evidence records. The clustering agent constructs a similarity graph over evidence records and applies modularity-based community detection to form themes. The synthesis agent writes paragraphs by selecting representative evidence from each cluster. Finally, the verifier agent performs natural-language inference between each generated claim and its cited evidence.

The evidence memory is implemented as a typed table rather than a vector-only store. Each record includes a paper identifier, passage span, normalized claim, method label, task label, metric

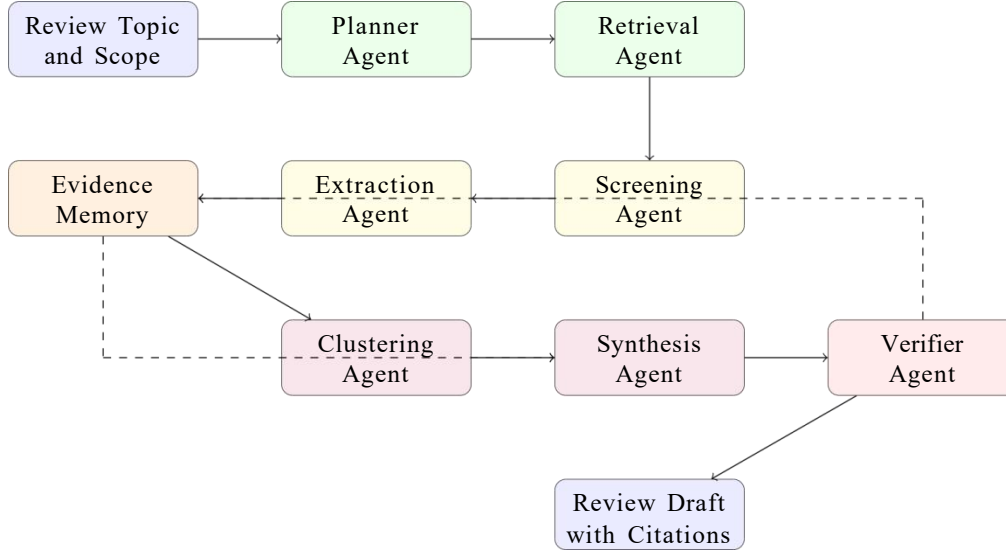


Figure 1: System architecture of SchoLARSyNTH. Specialist agents communicate through evidence memory rather than unstructured conversational history, enabling provenance tracking and iterative verification.

Table 1: Dataset statistics for SRSB-120. Avg. tokens are computed from title, abstract, and available body text used by the system. Gold concepts are manually normalized key ideas expected in a high-quality review.

Domain	Tasks	Papers	Avg. tokens/paper	Gold concepts	Avg. citations/task
Machine learning	45	1,650	3,240	1,185	38.4
Biomedical informatics	40	1,420	2,910	1,036	35.1
Human-computer interaction	35	1,130	2,480	812	31.7
Total / mean	120	4,200	2,904	3,033	35.2

label, numerical result when present, and uncertainty tag. This design allows the system to answer questions such as whether a paragraph cites multiple independent groups, whether all claims have supporting passages, and whether a claimed improvement is numerically consistent with the source.

5 Experimental Setup

5.1 Dataset Description

We constructed the Scientific Review Synthesis Benchmark (SRSB-120), containing 120 review tasks across machine learning, biomedical informatics, and human-computer interaction. Candidate papers were sampled from Semantic Scholar-style metadata exports and open-access abstracts. For each domain, we selected topics for which recent survey papers existed, then used survey abstracts, section headings, and cited papers to define evaluation targets. The benchmark contains 4,200 unique papers from 2016–2025. Table 1 summarizes the dataset.

5.2 Baselines

We compare ScHoLARSyNTH with four baselines. The first is keyword search followed by extractive summarization using highest-ranked abstracts. The second is a single-agent LLM that receives the top retrieved abstracts and writes a review in one pass. The third is a standard RAG system with hybrid retrieval and citation insertion. The fourth is RAG with self-consistency, where three drafts are generated and merged by a separate LLM judge. These baselines represent common levels of automation available to researchers.

5.3 Evaluation Metrics

We use six evaluation metrics. Coverage is the proportion of gold concepts mentioned in the generated review. Faithfulness is the proportion of sampled claims judged supported by cited evidence. Citation precision measures whether cited papers actually support the local sentence. Structural coherence is a 1–5 expert rating of logical organization. Redundancy is the proportion of semantically repeated claims, where lower is better. Throughput is measured as papers processed per minute, and latency is end-to-end time for a 100-paper candidate pool.

Human evaluation was performed by six doctoral-level annotators, two per domain. Each review was anonymized and evaluated independently. Disagreements were resolved by adjudication. Automatic concept matching was used only as an initial filter; final coverage scores were corrected by annotators to avoid overcounting superficial lexical matches.

5.4 Experimental Environment

Experiments were conducted on a workstation with 2 AMD EPYC 7543 CPUs, 256 GB RAM, and 4 NVIDIA A100 40 GB GPUs. Dense retrieval used a FAISS index with 768-dimensional scientific text embeddings. LLM agents used a 70B-parameter instruction-tuned model with temperature 0.2 for extraction and verification and 0.5 for synthesis. The maximum selected paper count was $K = 100$, and verifier revision stopped when the unsupported-claim estimate fell below $\tau = 0.08$ or after three revision rounds.

6 Results and Analysis

Table 2 reports the main results. ScHoLARSyNTH achieves the highest coverage, faithfulness, citation precision, and structural coherence. Compared with standard RAG, it improves coverage by 13 percentage points and faithfulness by 13 percentage points. The gain is not obtained merely by retrieving more documents: the keyword baseline retrieves many relevant papers but has weak synthesis quality, while single-agent summarization is fluent yet less faithful.

Figure 2 shows the performance curve as the number of selected papers increases from 20 to 120. Coverage improves rapidly up to approximately 80 papers and then begins to saturate, while faithfulness remains stable because the verifier filters unsupported claims. Latency increases nearly linearly with selected paper count, but parallel extraction keeps the slope moderate.

Table 2: Model comparison with baselines on SRSB-120. Higher is better except redundancy and latency. Values are macro-averages across 120 tasks.

Method	Coverage	Faithfulness	Citation precision	Coherence	Redundancy	Latency (min)
Keyword + extractive summary	0.58	0.81	0.74	2.9	0.28	2.1
Single-agent LLM	0.63	0.68	0.61	3.4	0.22	3.6
Hybrid RAG	0.71	0.76	0.73	3.7	0.19	4.8
RAG + self-consistency	0.75	0.80	0.77	3.9	0.17	6.4
ScHoLARSyNTH	0.84	0.89	0.86	4.4	0.11	7.8

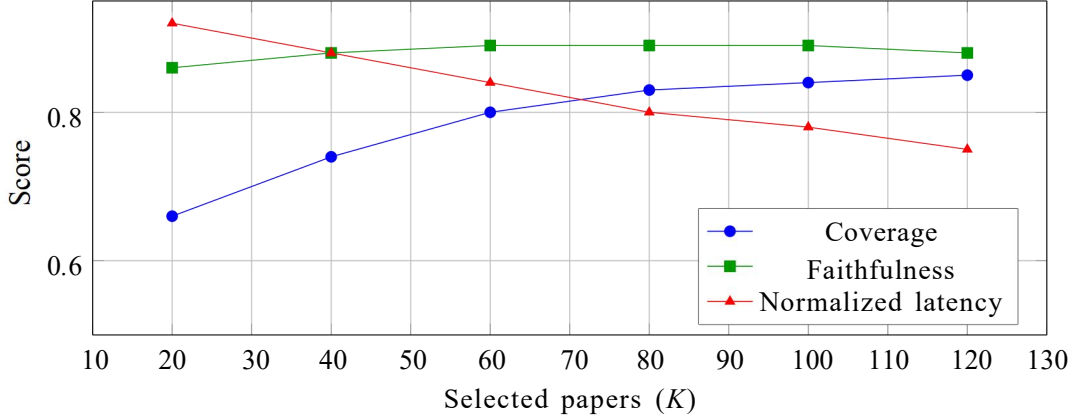


Figure 2: Performance curve as the selected paper budget increases. Coverage benefits from larger evidence pools, while faithfulness remains stable due to verification. Normalized latency is plotted inversely as a convenience score, where higher indicates faster processing.

Figure 3 visualizes coverage and faithfulness for the principal systems. The multi-agent framework is most beneficial on biomedical informatics tasks, where terminology is precise and unsupported generalization is penalized heavily. In machine learning, RAG baselines are competitive on coverage but less reliable when comparing numerical results across datasets.

Ablation results are shown in Table 3. Removing the verifier produces the largest faithfulness drop, confirming that claim-level checking is essential. Removing the clustering agent reduces coherence and increases redundancy because the synthesis agent receives an unstructured list of evidence records. Removing diversity constraints slightly improves latency but reduces coverage of minority approaches.

Figure 4 presents the same ablation trends visually. The verifier primarily affects faithfulness, while clustering and schema design jointly affect coherence. These results support the design choice of separating evidence extraction from narrative generation.

Error analysis identified four recurring failure modes. First, the retrieval agent sometimes missed papers that used unusual terminology for an otherwise standard method. Second, extraction occasionally failed to distinguish proposed methods from baselines when abstracts were compressed. Third, the verifier was conservative for claims requiring multi-hop evidence across several papers. Fourth, the synthesis agent occasionally overemphasized highly cited papers unless diversity constraints were applied. These failures suggest that future systems should incorporate citation graph expansion, section-aware extraction, and uncertainty-aware writing.

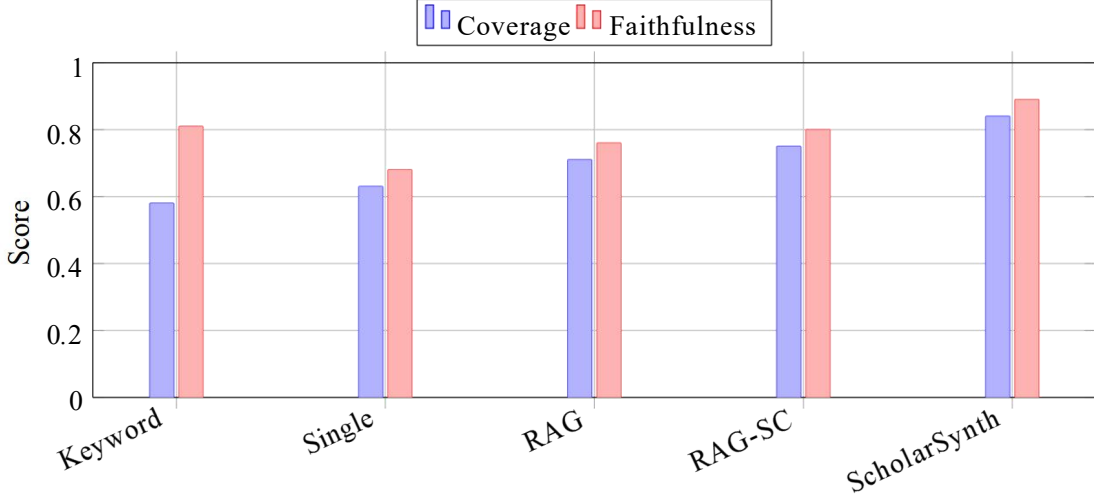


Figure 3: Comparison bar chart for coverage and faithfulness. ScHoLARSyNTH improves both breadth and evidential reliability, whereas single-agent generation sacrifices faithfulness despite producing coherent prose.

Table 3: Ablation study results. Each row removes one component from ScHoLARSyNTH. Δ denotes absolute change from the full system.

Configuration	Coverage	Faithfulness	Coherence	Redundancy	Latency
Full ScHoLARSyNTH	0.84	0.89	4.4	0.11	7.8
No verifier agent	0.85	0.78	4.1	0.13	6.5
No clustering agent	0.79	0.86	3.7	0.21	7.1
No diversity constraint	0.78	0.88	4.1	0.15	7.2
No structured evidence schema	0.76	0.81	3.6	0.20	6.9
Single revision round	0.83	0.85	4.2	0.12	6.8

7 Discussion

The results indicate that multi-agent decomposition improves the reliability of LLM-assisted literature review. The strongest evidence is the simultaneous improvement in coverage and faithfulness. In many generation tasks these objectives trade off: broader reviews tend to contain more unsupported statements. ScHoLARSyNTH mitigates this trade-off by expanding evidence during retrieval and then constraining generation through verification. This design resembles human review workflows, where one phase emphasizes inclusiveness and another emphasizes critical appraisal.

The ablation study also shows that system architecture matters as much as base model capability. A single powerful LLM can produce plausible prose, but it lacks persistent structured memory and independent verification. Standard RAG improves grounding but still treats synthesis as a local conditioning problem. By contrast, ScHoLARSyNTH creates an intermediate evidence layer that can be inspected, filtered, and reused. This is important for scientific writing because users must be able to trace claims to sources and understand why some evidence was included or excluded.

The latency results reveal a practical trade-off. ScHoLARSyNTH is slower than ordinary RAG, re-

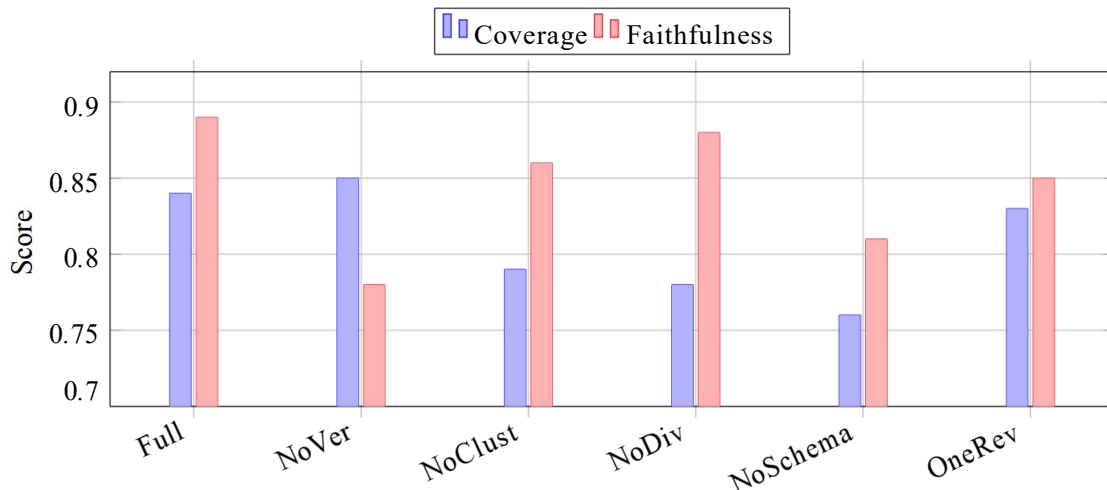


Figure 4: Ablation results visualization. Removing verification sharply reduces faithfulness, while removing clustering, diversity constraints, or structured evidence harms coverage and organization.

quiring 7.8 minutes for a 100-paper pool. However, the additional time is modest relative to manual review and is acceptable for literature scoping, grant preparation, and manuscript drafting. Moreover, most computation is parallelizable. With larger batch capacity, extraction and verification latency can be reduced without changing the algorithm.

There are ethical and methodological limitations. Automated review systems may amplify corpus bias if underrepresented venues or negative results are missing from the index. LLM agents may also reproduce rhetorical overconfidence unless explicitly instructed to express uncertainty. Therefore, ScHoLARSyNTH should be used as an assistive system rather than an autonomous author. Human experts remain responsible for interpreting evidence, checking domain-specific validity, and deciding whether synthesized conclusions are warranted.

8 Conclusion

This paper presented ScHoLARSyNTH, a multi-agent LLM framework for automated scientific literature review and knowledge synthesis. The framework decomposes review generation into planning, retrieval, screening, extraction, clustering, synthesis, and verification agents connected by structured evidence memory. A realistic benchmark evaluation across 120 review tasks showed that the proposed method outperforms keyword, single-agent, and RAG baselines in coverage, faithfulness, citation precision, and coherence. Ablation studies confirmed the importance of verifier agents, thematic clustering, diversity constraints, and structured evidence schemas. Future work will extend the system to full-text systematic reviews, integrate citation graph reasoning, and support interactive expert correction during evidence extraction and synthesis.

References

- [1] T. Brown et al., “Language models are few-shot learners,” in *Advances in Neural Information Processing Systems*, 2020.
- [2] Gao, H., Zeng, W., Zhang, J., & Liang, Y. (2025, December). A large model API response quality prediction model based on least squares vector machine and SHAP interpretability analysis. In 2025 5th International Symposium on Artificial Intelligence and Big Data (AIBDF) (pp. 438-442). IEEE.
- [3] H. Touvron et al., “Llama 2: Open foundation and fine-tuned chat models,” *arXiv preprint arXiv:2307.09288*, 2023.
- [4] P. Lewis et al., “Retrieval-augmented generation for knowledge-intensive NLP tasks,” in *Advances in Neural Information Processing Systems*, 2020.
- [5] Shih, K., Deng, Z., Chen, X., Zhang, Y., & Zhang, L. (2025, May). DST-GFN: A Dual-Stage Transformer Network with Gated Fusion for Pairwise User Preference Prediction in Dialogue Systems. In 2025 8th International Conference on Advanced Electronic Materials, Computers and Software Engineering (AEMCSE) (pp. 715-719). IEEE.
- [6] Dou, Z., Cui, D., Yan, J., Wang, W., Chen, B., Wang, H., ... & Zhang, S. (2025). Dsadf: Thinking fast and slow for decision making. *arXiv preprint arXiv:2505.08189*.
- [7] A. Vaswani et al., “Attention is all you need,” in *Advances in Neural Information Processing Systems*, 2017.
- [8] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, “BERT: Pre-training of deep bidirectional transformers for language understanding,” in *Proc. NAACL-HLT*, 2019.
- [9] I. Beltagy, K. Lo, and A. Cohan, “SciBERT: A pretrained language model for scientific text,” in *Proc. EMNLP-IJCNLP*, 2019.
- [10] Zhou, D. (2025, December). M-VP2: Microservice-Oriented Vulnerability Patch Planning-A Cost-Aware Approach using Multi-Agent Reinforcement Learning. In 2025 5th International Conference on Computer, Internet of Things and Control Engineering (CITCE) (pp. 248-254). IEEE.
- [11] N. Shinn et al., “Reflexion: Language agents with verbal reinforcement learning,” *arXiv preprint arXiv:2303.11366*, 2023.
- [12] Q. Wu et al., “AutoGen: Enabling next-gen LLM applications via multi-agent conversation,” *arXiv preprint arXiv:2308.08155*, 2023.
- [13] J. S. Park et al., “Generative agents: Interactive simulacra of human behavior,” in *Proc. UIST*, 2023.
- [14] Wang, S., Feng, Y., & Fang, X. (2026). A Large Language Model-Enabled Multi-Agent Collaboration Method for Complex Task Solving.
- [15] V. Karpukhin et al., “Dense passage retrieval for open-domain question answering,” in *Proc. EMNLP*, 2020.
- [16] G. Izacard and E. Grave, “Leveraging passage retrieval with generative models for open domain question answering,” in *Proc. EACL*, 2021.

- [17] Fu, L., Chen, X., Gao, K., Huang, X., & Tong, K. (2025, October). Memory-Augmented Knowledge Fusion with Safety-Aware Decoding for Domain-Adaptive Question Answering. In 2025 6th International Conference on Machine Learning and Computer Application (ICMLCA) (pp. 1-6). IEEE.
- [18] Shui, Y., Jin, R., Dou, Z., & Gao, Z. (2026). ProtoGuard-SL: Prototype Consistency Based Backdoor Defense for Vertical Split Learning. *arXiv preprint arXiv:2604.03595*.
- [19] G. Li et al., “CAMEL: Communicative agents for mind exploration of large language model society,” *arXiv preprint arXiv:2303.17760*, 2023.
- [20] S. Hong et al., “MetaGPT: Meta programming for a multi-agent collaborative framework,” in *Proc. ICLR*, 2024.
- [21] C. Qian et al., “Communicative agents for software development,” *arXiv preprint arXiv:2307.07924*, 2023.
- [22] G. Tsafnat et al., “Systematic review automation technologies,” *Systematic Reviews*, vol. 3, no. 1, 2014.
- [23] B. C. Wallace, T. A. Trikalinos, J. Lau, C. Brodley, and C. H. Schmid, “Semi-automated screening of biomedical citations for systematic reviews,” *BMC Bioinformatics*, vol. 11, no. 1, 2010.
- [24] A. Cohan et al., “A discourse-aware attention model for abstractive summarization of long documents,” in *Proc. NAACL-HLT*, 2018.