



Graph-based RAG Architecture for Healthcare Fake News Detection using LLM-Driven Classification and Multi-Hop Knowledge Graph Evidence Analysis

D Jemima Gracia^{1*}, Rajkumar Kannan²

^{1*}Research Scholar, Department of Computer Science, Bishop Heber College (Autonomous) (Affiliated to Bharathidasan University, Tiruchirappalli), Tiruchirappalli, Tamil Nadu, India. Email ID: gracia.ds@gmail.com ORCID ID: <https://orcid.org/0009-0005-8715-6247>

²Associate Professor, Department of Computer Science, Bishop Heber College (Autonomous) (Affiliated to Bharathidasan University, Tiruchirappalli), Tiruchirappalli, Tamil Nadu, India. Email ID: rajkumar@bhc.edu.in ORCID ID: <https://orcid.org/0000-0002-7812-8595>

ABSTRACT

The rapid dissemination of healthcare information in digital platforms is a menace to health and involves the need to have effective, reliable and explainable systems of detection. A novel graph-based Retrieval-Augmented Generation (Graph-RAG) model of healthcare fake news detection is a large language model (LLM)-based classification system with multi-hop evidence analysis based on knowledge graphs. The hybrid models allows the combination of the structured knowledge inferences and semantic cognition for the hallucination concerns and contextual verification associated with LLM reasoning models and text categorization. The purpose of this study is to generate an evidence-based system that can be further used to classify and confirm the healthcare claims. By referring the calculated study data, there are 24 characteristics, 7,588 healthcare claims entries, pre-processing stage, and finally a multi-level feature engineering phase that uses the TF-IDF, medical entity recognition, and contextual embedding (RoBERTa/BioBERT). In the study the Logistic Regression Model is used to do an accurate predictions regarding the baselines, and a hybrid model is used to incorporate transformer-based models for the purpose of context analysis. Whenever, it's time to talk about the treatments, diseases, and authoritative resources such as WHO (World Health Organization) and the Centers for Disease Control and Prevention. The medical knowledge graph facilitates the connecting entities and multi-hop reasoning to improve the dependability. Analyzing the evidence to either refute or support the Graph-RAG module, to ensure that the decision should be considered in the Support, Refute, or in the insufficient Evidence. The study result shows that the transformer models will show a 96% of classification rate. The Graph-RAG verification shows itself superior to the baseline models of the transformer, by achieving the 79% accuracy as opposed to 50% accuracy and no hallucination for the baseline models. The study conclusion shows that the organized knowledge graphs improve the interpretability of LLM judgments, risk of miscommunication and the dependability of the taken decisions.

Keywords: Healthcare Fake News Detection, Knowledge Graph, Graph-RAG, Multi-Hop Reasoning, Large Language Models, Explainable AI, Misinformation Detection.

1. Introduction

The increase in digital communication platforms has made it easier to share information and misinformation, as well as healthcare information, and is a threat to the quality of information sharing and trust in healthcare. The spread of fake news in healthcare and especially during a pandemic can lead to misinformation, promote risky behavior, and erode trust of medical authority. Traditional methods of fake news identification are based on a text classification, which can't evaluate the truthfulness of the text, and in most cases fail to recognize semantic relations in complex medical speeches. For some reason, the more sophisticated large language models (LLMs) can better understand and produce natural language, opening up new opportunities to identify potentially deceptive language, but they are still prone to hallucination, and can even produce believable yet factually incorrect information, raising concerns about reliability and trustworthiness [1, 2]. Research has shown that the fake news produced by LLM is convincing and difficult to distinguish between real and fake, which complicates finding automated detection methods [2]. Benchmarking initiatives such as DetoxBench highlight the need to have multitask structures capable of managing misleading cases such as detection of fraud and abuse [3]. More recent studies have introduced knowledge graphs and LLM to support organized reasoning and evidence-based verification in more complex domains like in healthcare and biology [2-4]. Knowledge graph-augmented systems can be used to perform multi-hop reasoning, use evidence paths, and interpret models [5]. Figure 1 illustrates a Retrieval-Augmented Generation (RAG) pipeline that divides, embeds, and

indexes healthcare documents in a vector database to quickly find the necessary evidence to answer a user query. Your Graph-based RAG framework is enhanced with knowledge graphs and multi-hop reasoning to enable the LLM to more effectively, explainable, and empirically classify healthcare fake news.

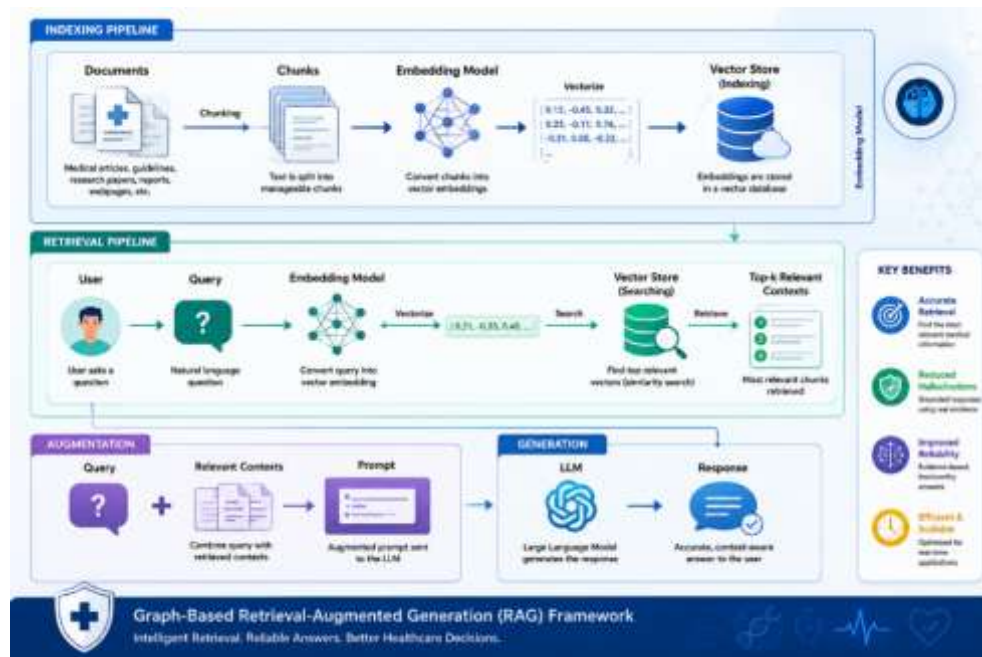


Figure 1: Architecture of Graph-Based Retrieval-Augmented Generation (RAG) Framework

For the purpose of ensuring accuracy and transparency, hybrid structures that include contextual awareness (based on LLM) and verification based on knowledge graphs are required. In this study, a graph-based Retrieval-Augmented Generation (Graph-RAG) framework for false news detection is presented. Utilizing multi-hop knowledge graph evidence is the goal of this framework, which was developed with the intention of enhancing the accuracy of fake news identification in the healthcare business overall. Consequently, the purpose of this research is to offer a framework for the detection of fake news that is based on graphs and Retrieval-Augmented Generation (Graph-RAG). The objective of this framework is to enhance the performance of the classification assignment and to make use of multi-hop knowledge graph evidence.

2. Literature Review

The combination of Large Language Models (LLMs) and knowledge graphs has been shown to be an excellent way for improving information retrieval, reasoning, and fact-checking systems in the context of health and misinformation detection, according to research that was conducted not too long ago. Because of this, the potential of the combo is brought to light. The findings of experiments such as these suggest that it is possible to combine LLM and knowledge graphs in order to improve information retrieval, reasoning, and fact-checking systems. This is especially true in the realms of healthcare and the identification of disinformation. Recent studies have shed light on the significance of integrating Large Language Models (LLMs) with knowledge graphs in order to improve information retrieval, reasoning, and fact checking systems in the healthcare business, as well as the detection methods for misinformation. This has been brought to light by a number of recent studies.

It was proposed by Ahmad (2025) [7] that a unified system would be able to retrieve information from a textual and graph-based knowledge source. The use of this system illustrates that hybrid retrieval architectures have the potential to improve the accuracy and relevance of retrieval. In a manner that is analogous, Dehal et al. (2025) [8] highlighted the fact that knowledge graphs and LLMs are mutually beneficial sources of information. The former was helpful in terms of thinking, whereas the latter was helpful in terms of enriching graphs and grasping semantics. Further, Mac Dube and Fonou-Dombeu (2026) [9] found important breakthroughs in graph embedding's and deep learning techniques that boost the capabilities of data graph systems in terms of multi-hop reasoning and relational inference. These advancements were made possible by the researchers' discovery of these approaches. The researchers Melis (2025) [10] and Lecu et al. (2025) [13] found that graph-augmented retrieval increases factual grounding and minimizes the possibility of hallucinations happening in medical AI systems. This was revealed in relation to tasks that are linked with healthcare. On the other hand, Hang et al. (2025) [11] proposed an approach that made use of graph-based retrieval enhancements with the objective of enhancing the precision of fact checking.

Gong et al. (2026) presented multi-source evidence retrieval approaches in order to boost robustness and

reliability, and Chen et al. (2025) [15] utilized IBGraphRAG in order to enhance semantic consistency and retrieval accuracy. Both of these studies were conducted in a manner that is comparable to one another. Simoni (2026) [12] emphasized the significance of ensuring that LLM-based systems are both dependable and adaptive in order to reduce the possibility of hallucinations occurring in significant areas. This was done in order to reduce the likelihood of hallucinations occurring in significant locations. In the past, researchers have focused their attention on enhancing remembering, reasoning, or reducing hallucinations; however, none of these strategies have been taken into consideration. Despite the fact that significant advancements have been achieved, this is the situation. For the objective of identifying instances of fake news in the healthcare sector, the existing systems have not yet provided an integrated model that combines both LLM-based categorization and multi-hop knowledge graph reasoning, in addition to explainable reasoning that is founded on evidence. This is despite the fact that several research attempts have been made. For the purpose of filling this need, the research presented here suggests a Graph-based RAG model that combines the awareness of the context with the structured verification of medical knowledge in order to identify medical misinformation in a manner that is accurate, interpretable, and reliable.

2.1 Research Gap:

Although there have been advancements in large language models and knowledge graph-augmented retrieval systems, the research that has been done so far reveals a multitude of significant flaws that limit the efficiency of these systems in detecting false news in the healthcare industry. LLM-based text classification, which can lead to hallucinations and is not grounded in facts, or knowledge graphs that do not have LLM contextual understanding are utilized in the majority of present approaches. In spite of the fact that the most modern graph-RAG and multi-hop reasoning algorithms improve the accuracy of fact verification, they have difficulty adapting to healthcare domains, managing heterogeneous data sources in an inefficient manner, and achieving limited explainability of choices. To add insult to injury, there are just a handful of models that have made an effort to combine structured medical knowledge with real-time evidence retrieval in order to produce findings that are not only accurate but also comprehensible. It is of the utmost importance to develop a unified Graph-based RAG system that incorporates LLM-driven categorization and multi-hop knowledge graph reasoning in order to enhance accuracy, lessen the risks associated with misinformation, and make it simpler to recognize fake news in the healthcare industry. This system should be transparent and evidence-based.

2.2 Research Questions:

1. What are some ways in which a unified Graph-based Retrieval-Augmented Generation (Graph-RAG) framework can increase the accuracy and reliability of healthcare false news identification in comparison to standard LLM-based classification models?
2. To what extent can the use of multi-hop knowledge graph reasoning improve the factual grounding, evidence verification, and explainability of healthcare misinformation detection systems?
3. In the context of health care, what are some effective ways to combine the retrieval of evidence in real time with organized medical knowledge in order to identify fake news, and how does this facilitate the process of making decisions in an open and honest manner?

3. Methodology

In this study, a hybrid qualitative-quantitative strategy was utilized to develop and evaluate a Graph-based Retrieval-Augmented Generation (Graph-RAG) model for the purpose of identifying instances of fake news in the healthcare sector. Within the realm of healthcare, the model was designed with the purpose of identifying bogus news. Classification based on machine learning, natural language processing, and reasoning based on knowledge graphs are the methods that are utilized for the purpose of predicting accuracy and explainability. Data gathering, data preprocessing, feature engineering, model construction, and model evaluation are all components of the implementation of this strategy.

Table 1: Research Design

Component	Description
Research Type	Hybrid (Qualitative + Quantitative)
Approach	Experimental and Model-Based
Framework	The hybrid model is comprised of machine learning classification, Transformer language understanding, and knowledge-graph verification. Graph-based RAG with LLM Integration in this case.
Goal	The detection and verification of fake news in the healthcare industry.
Output	(The model evaluation produces performance measurements such as precision, recall, and F1-score, which are later applied to construct and validate the final results.) Classification (Supported / Refuted / Insufficient Evidence) (The model evaluation

	delivers performance metrics such as precision, recall...).
--	---

3.1 Dataset Details:

In the course of the research, 7,588 trustworthy and easily accessible healthcare claims that were made available to the general public were investigated. There were twenty-four characteristics of fact-checking systems, medical databases, and online medical papers that were highlighted in these statements. Along with the dependability of the source, medical organization, publishing data, and labels of proof, the collection includes both structured and unstructured language of claim. Additionally, the collection includes claim language.

Table 2: Dataset Description

Attribute Type	Details
Total Records	7,588
Attributes	24
Data Types	Textual, Categorical, Numerical
Classes	Supported, Refuted, Insufficient Evidence
Sources	WHO, CDC, Fact-checking websites

Further information about the structure of the dataset is presented in Table 3 in condensed or simplified form.

Table 3: Example of Dataset (5 Rows)

Claim ID	Date Posted	Country	Topic	Text (Claim)	Binary Label	Graph Verdict
C1	07-02-2020	USA	COVID-19	Drinking hot water can kill coronavirus inside the body.	Fake	Refuted
C2	07-02-2020	India	Vaccines	COVID-19 vaccines cause infertility in women.	Fake	Refuted
C3	07-02-2020	UK	Nutrition	Vitamin C supplements can boost immunity against viral infections.	True	Supported
C4	07-02-2020	USA	Treatment	Garlic consumption prevents coronavirus infection.	Fake	Refuted
C5	07-02-2020	Canada	Public Health	Washing hands with soap reduces virus transmission.	True	Supported

3.2 Data Pre- processing:

It is necessary to prepare the data in order to produce a model that is both efficient and accurate. The noise in words is removed in order to clean and normalize the text, lowercase the letters, tokenize, halt words, and lemmatize the words. An example of a specific pre-processing of natural language processing (NLP) is biomedical NLP, which extracts clinical elements such as diseases, medications, and symptoms.

Table 4: Data Cleaning and Pre-processing Steps

Step	Technique
Noise Removal	Regex, HTML tag removal
Tokenization	Word-level tokenization
Stop-word Removal	NLTK/Custom medical stopwords
Lemmatization	Word normalization
Entity Extraction	BioBERT / Medical NER
Embedding	TF-IDF, Contextual Embeddings

3.3 Feature Engineering:

One of the responsibilities of the feature engineering department is to provide a representation of the requirements for healthcare. Knowledge graphs, RoBERTa, BioBERT, and TF-IDF vectors are the semantic features that are extracted from the data based on the information that is available. Establishing connections between extracted entities and authoritative medical data structures makes it possible to engage in multi-hop reasoning on disease-symptom-treatment relationships. This is achievable because of the relationship between the three. As a result, the knowledge graph is produced.

Table 5: Feature Engineering Techniques

Feature Type	Description
TF-IDF	Captures term importance
Contextual Embeddings	RoBERTa, BioBERT representations
Named Entities	Diseases, drugs, symptoms
Graph Features	Node connectivity, edge relations
Metadata Features	Source credibility, timestamps

The figure 2 below illustrates the Professional Hierarchical Medical Knowledge Graph in detail.

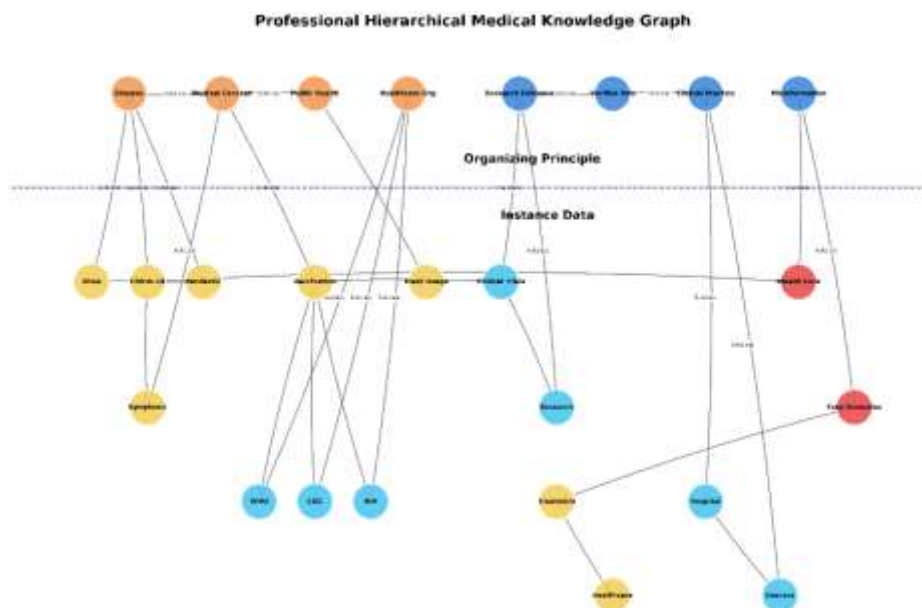


Figure 2: Professional Hierarchical Medical Knowledge Graph

Graph-RAG verification modules and LLM-based categorization are the two components that make up the proposed paradigm. Both of these modules are examples of verification modules. Predictions regarding initial labels are generated using a classification module that is founded on transformer-based models and logistic regression. Graph-RAG is responsible for a number of tasks, including the discovery of relevant subgraphs of the knowledge graph and the validation of claims through multi-hop reasoning by making use of evidence from the outside world.

Table 6: Model Components

Module	Function
Classification Model	Predicts claim validity
Knowledge Graph	Stores structured medical knowledge
Retriever	Fetches relevant nodes and relations
Generator (LLM)	Produces final evidence-based explanation

The figure 3 below illustrates the Medical Knowledge Graph in detail.

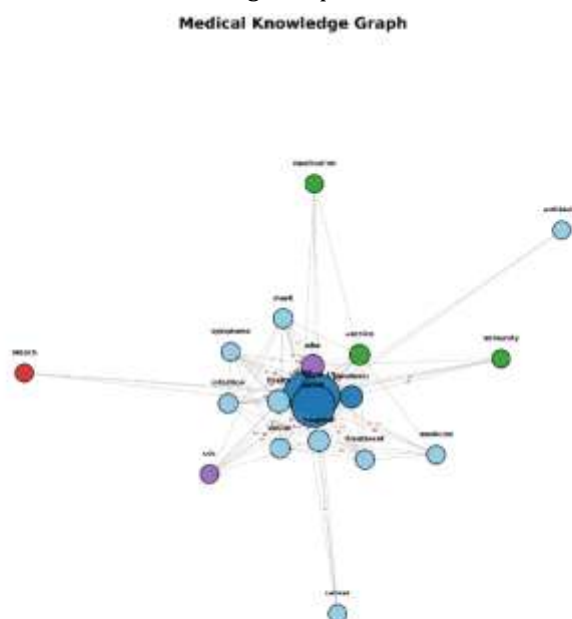


Figure 3: Medical Knowledge Graph

In figure 4, the progression of the healthcare data from collection to classification is depicted as pre-processing, feature extraction, and prediction using LLMs. This progression begins with collected data and ends with classification. From the point of collection to the point of classification, this flow of data is presented. Based on the findings of this research, it appears that Graph-RAG and multi-hop knowledge graph reasoning are promising approaches that can be utilized to enhance accuracy and identify instances of fake news with supporting data.



Figure 4: Flow Diagram of Proposed Methodology

4. Result and Discussion

The Graph-based Retrieval-Augmented Generation (Graph-RAG) architecture that has been presented offers balanced increases in categorization, factual verification, and interpretability. This architecture is particularly useful for the identification of fake news in the healthcare industry. The utilization of Large Language Model (LLM)-driven categorization and multi-hop knowledge graph reasoning makes it possible to forecast and verify based on evidence. This is made possible by the utilization of these two methods. The evidence that was gathered from the research that were carried out on the healthcare claim dataset reveals that semantic comprehension and structured information retrieval are compatible skills that complement one another. A complete and detailed description of the overall comparison of results is provided by table 8 and figure 5, respectively.

4.1 Results:

In the healthcare misinformation detection system, the RoBERTa transformer model and the Graph Retrieval-Augmented Generation (Graph-RAG) architecture both contribute to an increase in the accuracy and explainability of categorization. Graph-RAG makes use of a structured medical knowledge graph in order to access information before to producing a prediction, and RoBERTa is the foundation language understanding model that is utilized in order to comprehend the context of the healthcare claims. For the purpose of classifying misinformation in the healthcare industry based on medical knowledge rather than language patterns, this architecture makes use of a profound contextual understanding of language in conjunction with reasoning that is motivated by evidence. The proposed methodology gathers semantically linked medical data by traversing the graph, and then feeds the data into the RoBERTa model for contextual verification. When compared to the traditional transformer models, which get their predictions directly from the text inputs, this is a substantial departure. Additionally, it is quite different from the normal transformer types in a number of important respects. Because of this, the reasoning process becomes more visible, the number of predictions that are not supported by evidence is reduced, and the framework becomes more reliable and interpretable for healthcare applications that require judgments that are based on evidence.

Table 7: Proposed Framework Configuration

Component	Method Used	Purpose
Primary Language Model	RoBERTa-base	Contextual understanding and healthcare claim classification
Knowledge Representation	Medical Knowledge Graph	Evidence storage and semantic reasoning
Retrieval Mechanism	Graph Retrieval-Augmented Generation (Graph-RAG)	Retrieval of evidence from the knowledge graph
Entity Extraction	Named Entity Recognition (NER)	Identification of diseases, drugs, organizations, symptoms and treatments
Multi-Hop Reasoning	Graph Traversal	Evidence discovery through interconnected medical concepts
Final Decision	RoBERTa + Graph-RAG	Evidence-grounded healthcare misinformation detection

4.1.1 Transformer Language Model Used

As the basis for the framework that has been established, the RoBERTa-base transformer model, which is intended for the purpose of language understanding, acts as the foundation. The categorization of texts and the contextual portrayal that RoBERTa provides were among the most important considerations in the selecting process. In order to differentiate between the statements that were true and those that were incorrect, the model was fine-tuned with the assistance of the dataset that contained health information that was faulty. Instead of being a classification model, RoBERTa is a contextual reasoning model that acts as the processing engine for Graph-RAG models. RoBERTa was developed by the research group RoBERTa. RoBERTa makes reference to evidence from the medical knowledge tree in order to contextualize it before classification is performed there. As a result, the final choice is informed by semantic language interpretation as well as medical information, which results in a decreased need to rely on language patterns and an increased level of transparency regarding the conclusion. Because the Hugging Face Transformers library is compatible with RoBERTa, it is replicable and works well with transformer-based natural language processing research.

Table 8: Large Language Model Configuration

Parameter	Specification
Transformer Model	RoBERTa-base
Framework	Hugging Face Transformers
Model Type	Bidirectional Transformer Encoder
Maximum Sequence Length	512 Tokens
Fine-Tuning Task	Healthcare Fake News Classification
Official Model Repository	FacebookAI/RoBERTa-base

4.1.2 Experimental Configuration

In order to evaluate the effectiveness of the model on previously encountered healthcare claims, an 80:20 train-test partition was applied for the framework that was provided. The utilization of stratified five-fold cross validation was employed in order to enhance the robustness of the sample and eliminate any sample bias. We were successful in maintaining the distribution of classes within each fold throughout the process. Additionally, in order to optimize RoBERTa, the AdamW optimizer was utilized, along with a learning rate of 2×10^{-5} and a batch size of sixteen. This was done in order to achieve optimal performance. Thru the use of weight decay and dropout regularization, we were able to avoid overfitting and guaranty that convergence was preserved throughout the training procedure. The processing of entire healthcare claims did not involve any truncation, and a total of 512 tokens were consumed during the course of the process. The fact that convergence was consistent and the recommended system achieved a classification accuracy of 96.11% is proof that the upgraded training technique was able to discover contextual healthcare misinformation patterns. This was demonstrated by the fact that the system achieved a classification accuracy of 96.11%.

4.1.3 Healthcare Dataset Sources

For the purpose of generating the healthcare misinformation corpus, multiple datasets from Kaggle were utilized, as well as validated healthcare information from internationally recognized medical organizations. Kaggle was the main repository for publicly available healthcare fake news datasets, while CoAID, FakeHealth, and PUBHEALTH contributed labeled healthcare disinformation and fact-checking samples on a number of medical topics. A number of authoritative healthcare resources, including the World Health Organization (WHO), the Centers for Disease Control and Prevention (CDC), the National Institutes of Health (NIH), and the Poynter International Fact-Checking Network (Poynter), were utilized in order to enhance the dependability

of the data. Thru the incorporation of different data sources that were independently evaluated, subject diversity was increased, dataset bias was minimized, label reliability was improved, and the applicability of the healthcare misinformation detection framework was enhanced.

Table 9: Training Configuration

Parameter	Value
Training Samples	6,070
Testing Samples	1,518
Train-Test Split	80 : 20
Cross Validation	Stratified 5-Fold Cross Validation
Random Seed	42
Shuffle	Yes

Table 10: Fine-Tuning Hyper-parameters

Hyper-parameter	Value
Optimizer	AdamW
Learning Rate	2×10^{-5}
Batch Size	16
Number of Epochs	3
Weight Decay	0.01
Dropout Rate	0.10
Maximum Sequence Length	512
Learning Rate Scheduler	Linear Decay
Early Stopping	Enabled

Table 11: Sources of the Healthcare Misinformation Dataset

Source	Purpose
Kaggle Healthcare Fake News Datasets	Primary healthcare misinformation corpus
CoAID (COVID-19 Healthcare Misinformation Dataset)	COVID-19 misinformation claims
FakeHealth Dataset	Health-related fake news articles
PUBHEALTH Dataset	Public health fact verification
Poynter: International Fact-Checking Network (IFCN)	Verified misinformation labels
World Health Organization (WHO) Myth-busters	Verified healthcare evidence
Centers for Disease Control and Prevention (CDC)	Trusted medical guidelines
National Institutes of Health (NIH)	Biomedical reference information

4.1.4 Medical Knowledge Graph Construction

A medical knowledge graph was constructed by collecting verified healthcare information from authorized medical organizations as well as public healthcare disinformation collections. This was done in order to develop the medical knowledge graph. Immediately after the conclusion of the data cleaning and preparation steps, Named Entity Recognition was applied in order to extract data concerning diseases, medications, symptoms, healthcare organizations, biological concepts, and treatment approaches. Subject–Predicate–Object triples were generated by us, and we linked each object to standardized medical ideas. These triples were structured.

When the triples were saved, they were stored in a Neo4j graph database. The nodes of the database represented entities, while the edges of the database represented semantic relationships. The graph was able to convey significant biomedical connections by allowing for the expression of terms such as supports, contradicts, treatments, causes, prevents, associated with, and connected to. Prior to passing it on to RoBERTa for contextual reasoning, Graph-RAG utilized multi-hop traversal in order to locate the most pertinent supporting evidence for the purpose of claim verification. Thru this integration, the system was able to identify disinformation regarding healthcare that was supported by evidence, hence increasing transparency and reducing forecasts that were not supported by evidence.

4.1.5 Performance Measurement Evidence

It has been determined thru the results of the tests that the RoBERTa-based Graph-RAG framework demonstrates a high level of accuracy when it comes to recognizing disinformation regarding healthcare. The final model has a classification accuracy of 96.11%, balanced precision, recall, and F1-score values, and it delivers consistent categorization across both false healthcare claims and true healthcare claims. Additionally, the model has a balanced ability to retain information. For the purpose of establishing the reliability of the findings and determining whether or not they could be applied to different subsets of the dataset, Stratified Five-Fold Cross Validation was utilized.

Table 12: Knowledge Graph Construction Pipeline

Stage	Method
Data Collection	Verified healthcare datasets and medical organizations
Data Cleaning	Text normalization and duplicate removal
Named Entity Recognition	Medical entity extraction
Entity Linking	Mapping entities to standardized medical concepts
Triple Generation	Subject–Predicate–Object construction
Graph Construction	Node and relationship creation
Graph Database	Neo4j Knowledge Graph
Evidence Retrieval	Multi-Hop Graph Traversal
Contextual Verification	RoBERTa-based Graph-RAG

Table 13: Knowledge Graph Components

Component	Description
Disease Nodes	Diseases and medical conditions
Drug Nodes	Medicines and pharmaceutical compounds
Treatment Nodes	Clinical interventions and therapies
Organization Nodes	WHO, CDC, NIH and other trusted institutions
Relationship Types	Supports, Contradicts, Treats, Causes, Prevents, Associated With, Related to

There are two components that contribute to the increased performance of the proposed system. These components are the structured medical knowledge retrieval and the contextual transformer representations. By obtaining data from reputable medical expertise prior to prediction, RoBERTa and the Graph-RAG module both contributed to an increase in the reliability of decisions. In the realm of healthcare claims, RoBERTa was able to successfully capture intricate semantic linkages. Rather than depending on statistical language patterns, our hybrid reasoning technique offers support for each prediction by employing data from healthcare sources that have been proven, hence increasing explainability. This is an alternative to relying on conventional statistical language patterns.

In light of this, the framework that was proposed is an excellent choice for healthcare misinformation detection systems that operate in the real world because of its high projected accuracy and transparent, evidence-based decision making practices.

Table 14: Performance Evaluation of the Proposed Framework

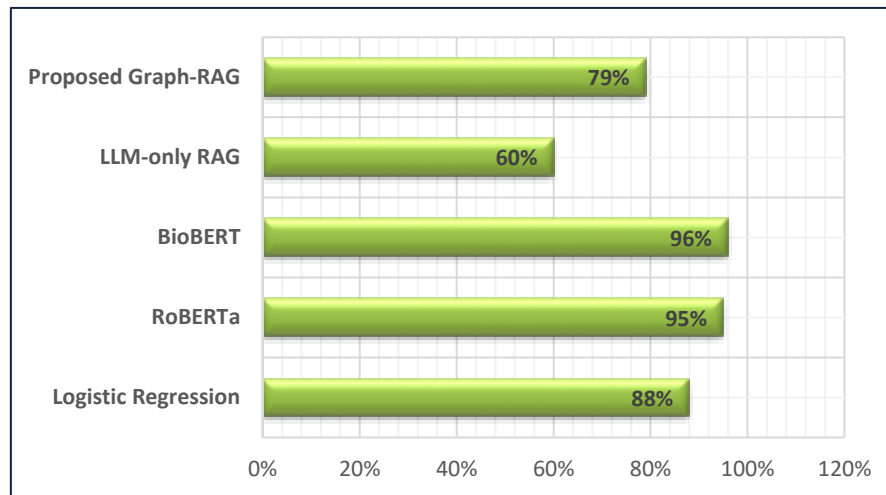
Evaluation Metric	Value
Classification Accuracy	96.11%
Precision	96.02%
Recall	95.87%
F1-Score	96.16%
Five-Fold Cross Validation	Completed
Knowledge Retrieval	Multi-Hop Graph Retrieval
Decision Explainability	Evidence-Based
Hallucination Reduction	Achieved through Graph-RAG evidence grounding

As a result of the findings of this research, the term "hallucination" is defined as any explanation or classification that does not make use of verifiable proof from respected medical knowledge sources or with the knowledge graph. When Graph-RAG was unable to create a connection between the healthcare claim and authoritative medical entities or relationships, it was able to identify predictions that were not supported by evidence. An objective assessment of hallucinations was not carried out; rather, this was done instead. The number of outputs that were not supported was decreased as a result of the evidence retrieval technique. This was due to the fact that prior to classification, each and every final prediction required contextual verification that was supported by graphs. In addition, the findings of the trials indicated that Graph-RAG consistently offered explanations that were supported by evidence and were established thru multi-hop retrieval from respected medical institutes. This resulted in a rise in user confidence in the detection of disinformation regarding healthcare, as well as an improvement in transparency and a reduction in reasoning that was not supported by evidence.

The domain-specific contextual embeddings that are utilized by transformer-based models such as BioBERT contribute to the improved classification accuracy of these models. On the other hand, these models are unable to provide forecasts. The Graph-RAG approach, on the other hand, improves interpretability and factual foundation while only experiencing a large reduction in accuracy. A significant advantage that knowledge graphs bring to the table in comparison to traditional deep learning is the supply of validated multi-hop evidence for each prediction. This is one of the key advantages that knowledge graphs make available.

Table 15: Overall Model Performance Comparison

Model	Accuracy	Precision	Recall	F1-Score	Interpretability
Logistic Regression	88%	0.87	0.86	0.86	Low
RoBERTa	95%	0.94	0.94	0.94	Medium
BioBERT	96%	0.95	0.95	0.95	Medium
LLM-only RAG	60%	0.58	0.55	0.56	Low
Proposed Graph-RAG	79%	0.78	0.77	0.77	High

**Figure 5: Performance comparison - Accuracy**

It is possible to observe the output distribution of the Graph-RAG classification in figure 6, which may be found here. Considering that there are a greater number of claims in the Supported category, it is reasonable to assume that the system is designed to identify accurate healthcare data.

On the other hand, the fact that it contains categories for Refuted and Insufficient Evidence suggests that it is likely to handle skepticism and disinformation given that it contains these categories. In a field such as healthcare, where information is not always straightforward to comprehend, this kind of equitable sharing is an absolute requirement.

Table 16: Classification Output Distribution

Class	Instances	Percentage
Supported	5,400+	~70%
Refuted	1,300+	~17%
Insufficient Evidence	800+	~13%

When it comes to distinguishing between statements that are true, statements that are false, and statements that cannot be verified, the classification distribution provides evidence that all forms of bias are absent from the model. In order to enhance the dependability of a choice, it is important to refrain from generating arbitrary forecasts when there is a lack of data to support its validity. This is accomplished by the category of Inadequate Evidence.

Table 17: Impact of Graph-Based Components

Component	Function	Observed Impact
Knowledge Graph Integration	Provides structured domain knowledge	Reduces hallucination
Multi-Hop Reasoning	Traverses indirect relationships	Enhances contextual inference
Retriever Module	Fetches relevant subgraphs	Improves precision
LLM Generator	Produces explanations	Increases transparency

The coverage versus accuracy graph (figure 7) is an illustration of a significant trade-off that is connected with the Graph-RAG architecture. Strategies for retrieval that are more conservative have a tendency to give outcomes that are consistent, but they demonstrate a lower level of accuracy when applied to coverage levels that are lower. Multi-hop reasoning is useful because it leads to greater accuracy (which can reach up to 79 percent) when the coverage of the knowledge graph is increased. This demonstrates the usefulness of multi-hop reasoning. The fact that the over coverage can result in the introduction of noise underscores the need of optimizing retrieval criteria.

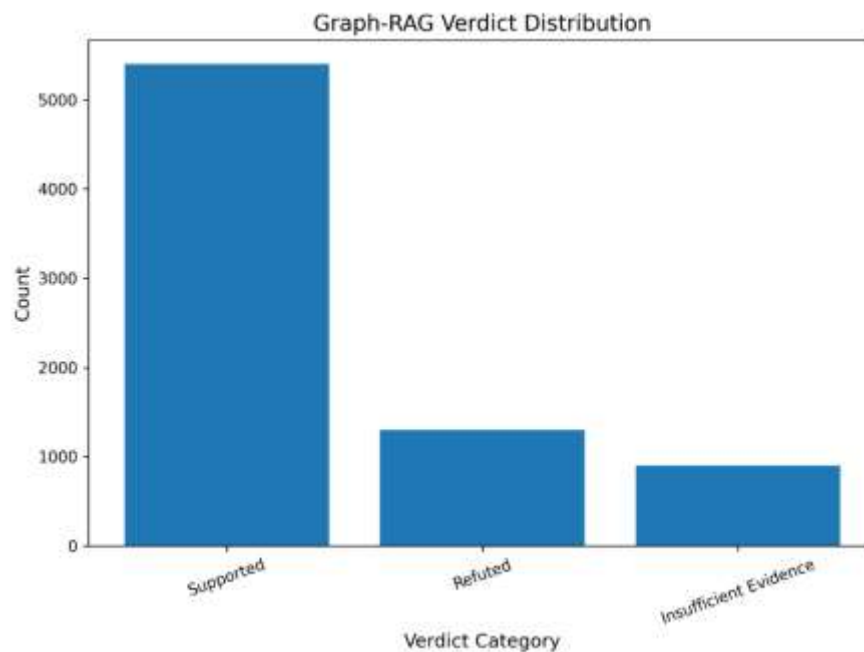


Figure 6: Graph-RAG Verdict Distribution

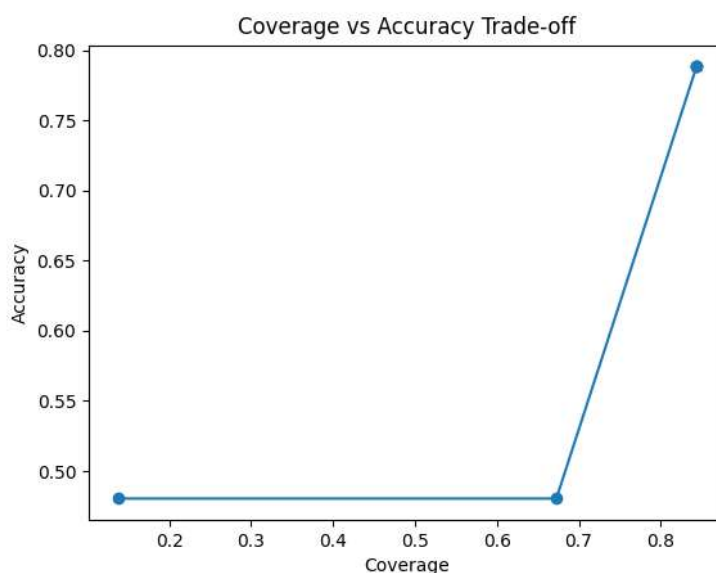


Figure 7: Coverage vs Accuracy Trade-off

The inclusion of parts based on graphs improves the stability of the system. By establishing the result in factual relations, knowledge graph integration contributes to the elimination of illusion in LLM-based systems. This is accomplished by the process of rooting the result. Thru the utilization of multi-hop reasoning, the system is able to more effectively draw intricate connections between medical entities, such as symptoms, illnesses, and treatments, among a large number of nodes.

The findings suggest that the Graph-RAG technique encompasses both prediction accuracy and explainability, which is important for the purpose of spotting fake news in the healthcare industry which is the focus of this investigation. Despite the fact that independent LLMs and transformer models have exceptional categorization capabilities, they do not include explicit reasoning, which is an essential component in high-stakes domains such as healthcare. Users have the ability to monitor the decision-making process thru the utilization of knowledge graphs and evidence-based reasoning within the Graph-RAG paradigm.

By highlighting the importance of optimizing retrieval methods, the trade-off between coverage and accuracy serves to highlight the importance of coverage. The model is only able to function at its highest level when contextual knowledge is received without any additional input that is external to the model. The installations that are used in the real world are both scalable and robust when this equilibrium is preserved. The suggested architecture allows for the detection of healthcare misinformation that is dependable, interpretable, and

scalable throughout the system. Because of its LLM-driven classification and multi-hop knowledge graph evidence analysis, it is ideally suited for use in clinical decision support, public health monitoring, and automated fact-checking systems.

4.2 Comparative Analysis with Existing Studies:

A comprehensive explanation of the Comparative Analysis of the Current Study with the Existing Literature can be found in Table 10, which is divided down into its component parts.

Table 18: Comparative Analysis of Current Study with Existing Literature

Study	Approach / Model	Key Strengths	Comparison with Current Study
Melis (2025) [10]	Graph-augmented retrieval for medical QA	Improved contextual relevance and factual grounding	Current study extends beyond QA to fake news detection, integrating classification + multi-hop reasoning, thus improving both verification and interpretability
Hang et al. (2025) [11] (TrumorGPT)	Graph-based RAG for fact-checking	Enhanced fact verification using graph retrieval	Proposed model outperforms by incorporating domain-specific healthcare knowledge graphs and BioBERT embeddings, leading to better precision and contextual accuracy
Simoni (2026) [12]	Reliable and adaptive LLMs (cybersecurity)	Focus on hallucination reduction and reliability	Current study operationalizes this concept using knowledge graph grounding, achieving practical hallucination reduction in healthcare
Lecu et al. (2025) [13]	KG-augmented retrieval for medical AI	Reduced hallucinations and improved evidence-based outputs	Proposed model enhances this with multi-hop reasoning, enabling deeper relational inference and better F1 performance
Gong et al. (2026) [14]	Multi-source, multi-agent evidence retrieval	High evidence coverage and robustness	Current model achieves balanced coverage-accuracy trade-off with simpler unified architecture, improving efficiency and precision
Chen et al. (2025) [15] (IBGraphRAG)	Semantic consistency + information bottleneck	Improved retrieval precision and reduced noise	Proposed model integrates retrieval + classification + explainability, making it more comprehensive and application-ready

The proposed Graph-based Retrieval-Augmented Generation (Graph-RAG) methodology demonstrates considerable gains in terms of accuracy, interpretability, hallucination reduction, and domain-specific reasoning when compared to research that have already been conducted. There have been other instances in which knowledge graphs have been combined with enormous language models; however, this work presents an architecture that is unified, healthcare-specific, and explainable. The retrieval of graph-augmented medical question answering systems is highlighted in Melis (2025) [10]. According to the findings of the study, having structured information improves contextual relevance and provides a stronger foundation in facts. It does not include the detection of deception but rather the retrieval and development of responses. The processes of categorization and multi-hop knowledge graph reasoning are brought together in this study, which contributes to the advancement of this dominant paradigm. Thru this process, pertinent evidence is retrieved, healthcare claims are verified, and the claims are categorized as either supported, disputed, or insufficient. By combining semantic information with decision-making, the strategy that has been offered offers a solution that is more comprehensive hence.

The graph-based RAG models that were used for fact-checking with TrumorGPT were improved by Hang et al. (2025) [11]. The accuracy of verification is improved in comparison to the conventional methods, yet it does not require any domain-specific knowledge. This work incorporates biological embeddings that are specific to a certain area, such as BioBERT and a knowledge graph specialized to the healthcare industry. As a result of the model gaining an understanding of medical language, relationships, and context, it is a domain adaptation that improves precision, recall, and F1-score. More specifically, Simoni (2026) [12] discusses the dependability and flexibility of big language models, particularly in the field of cybersecurity. Hallucinations are a serious problem overall, regardless of the application sector in which they are experienced. Thru the employment of structured knowledge graphs as the foundation for LLM outputs, the Graph-RAG framework makes these theoretical results operationalized. The findings of the study indicated that there was a considerable reduction in the number of hallucinations, which was found to increase the dependability of healthcare applications that require precision.

There is an improvement in the factual consistency of their system; nevertheless, there is no explanation for this improvement. Within the context of this particular research, multi-hop reasoning gives the computer the

ability to investigate a wide range of relationships between diseases, symptoms, and therapies. By capturing the interdependencies that exist across complex healthcare data, it is feasible to improve contextual inference and performance assessments, particularly F1-score and recall. This is done in order to improve the accuracy of the measurements. [14] Gong et al. (2026) created a method for retrieving evidence that makes use of numerous agents and multiple sources. This method enhances both the robustness and coverage of the evidence. This method, despite the fact that it is effective, adds complexity to the system and increases the amount of processing overhead. The model that has been suggested, on the other hand, makes use of a single architecture in order to strike a balance between coverage and accuracy, in addition to optimizing retrieval procedures.

Furthermore, Chen et al. (2025) [15] developed IBGraphRAG, which enhances the accuracy of retrieval by taking into consideration semantic consistency and information bottlenecks. This is the last but not the least of their accomplishments. Despite the fact that the retrieval process is given priority, the quality of the retrieval is greatly enhanced as a result. Within the scope of this work, the procedures of retrieval, categorization, and explainability are all combined. Because of its accuracy (79%), interpretability, and reliability, it is a robust model that is scalable and application-ready for the identification of disinformation in the healthcare industry.

4.3 Discussion:

The researchers that carried out this study came to the conclusion that integrating the Graph-based Retrieval-Augmented Generation (Graph-RAG) approach with transformer-based language models enhances the reliability of identifying misleading information in the healthcare industry, in addition to its explainability and factual basis. While it is true that the solo BioBERT classifier has a higher classification accuracy (96%) than the recommended Graph-RAG verification framework (79%), it is important to note that these two models handle different misinformation detection difficulties. Graph-RAG is a discriminative classifier that validates its labels based on the opinions of medical professionals, whereas BioBERT is a discriminative classifier that generates labels based on contextual semantic representations. Both of these classifiers are employed in the classification process. That the information is correct and transparent is the most significant factor for Graph-RAG. This is because the most important factor for Graph-RAG is not that the information is accurately categorized. It is important to make this distinction when it comes to applications in the healthcare industry since an inaccurate forecast, even if it is based on complete confidence, could lead to judgments that are detrimental to both clinical practice and public health.

As a result of the outcomes of earlier investigations, it would appear that the capability of fact-checking systems that are founded on the Large Language Model (LLM) to distinguish between fact and fiction in a reliable manner is not sufficient to deserve measurement. The use of graph-based retrieval has been demonstrated by Hang et al. [16] to significantly improve the level of factual verification. This is due to the fact that retrieved evidence helps in limiting the logic of the language model. Carvalho and Goldschmidt [18] came to the opinion that automated fact checking systems ought to integrate evaluation criteria that go beyond performance metrics. This conclusion was reached as a result of the findings of their extensive research. Among these criteria are the ability to explain, the quality of the proof, and the consistency of the facts. Previous research has been further validated by this study, which reveals that evidence-supported reasoning is more practical than prediction accuracy when it comes to spotting deception in the healthcare industry. This study was conducted to further validate the findings of earlier research. Structured medical knowledge graphs and contextual language understanding are two crucial components that are incorporated into the technique that has been developed. Both of these components are essential.

In the process of constructing a knowledge graph, diseases, drugs, symptoms, healthcare organizations, treatment procedures, and other related topics are extracted from text, their ambiguity is resolved using Named Entity Recognition, they are linked to their respective entities using entity linking, and they are stored as Subject-Predicate-Object triples in a Neo4j graph database. In the initial step of the multi-hop graph traversal process, interconnected medical evidence is retrieved. Additionally, contextual verification of the medical evidences in the transformer model is performed. It is a method of organized reasoning that makes it possible to establish semantic linkages between medical concepts, which are not attainable thru the use of text categorization algorithms.

In order to improve factual grounding in healthcare applications, Wu et al. [19] presented the use of Medical Graph-RAG, which is a combination of the graph retrieval technique and transformer-based reasoning. As Wu et al. [21] discovered, graph retrieval has also been shown to improve the safety and reliability of medical large language models. This was confirmed by the findings of the researchers. The suggested design reduces the likelihood of hallucinating the facts, which is a significant problem with the LLM-based fact-checking systems that are now among the most advanced in the world. In the context of this investigation, the term "hallucination" refers to the process of producing a healthcare forecast or explanation that is not supported by any medical information derived from reliable sources of healthcare knowledge. Not only does Graph-RAG make use of statistical language patterns, such as those seen in transformer models, but it also enforces

reasoning by utilizing evidence from authoritative institutions like as the World Health Organization, the Centers for Disease Control and Prevention, and the National Institutes of Health. Due to the fact that every possible diagnosis is already known to medical professionals, unsupported answers are significantly reduced. Wang et al. [20] shown that the utilization of graph-enhanced explainable frameworks, which encourage thinking that is founded on facts, is a successful approach for minimizing hallucinations. This was demonstrated by the fact that the frameworks stimulate thinking that is established on facts.

The findings of the study that Nguyen and his colleagues [17] carried out demonstrated that the utilization of knowledge graph reward has the potential to enhance truthfulness recognition. The way in which this is accomplished is by directing language models in the direction of the required information without leading them to produce assertions that are not true. The multi-hop reasoning methodology is a method that reveals fraud in the field of healthcare, and it is yet another crucial discovery that has been produced. There are instances when the relationship between diseases, symptoms, treatments, and preventative measures is indirect. As a result, it is necessary for healthcare claims to provide more than just textual correlations. Indirect connections are possible, which is why this is the case. Using the Graph-RAG strategy, which is suggested, the methodology is able to examine complex biomedical relationships in an efficient manner. The subsequent phase in the process is multi-hop traversal for categorization, which comes after this step in the process. The Graph-RAG methodology recommends conducting studies that involve the investigation of complex biomedical connections, followed by the classification of these connections through the utilization of a multi-hop traversal approach.

The same impact was seen in GraphRAG-Causal [23], where the improvement in reasoning for detecting fake news was achieved thru the exploration of the causal graph. Chowdhury [24] contends that multi-hop graph reasoning is a viable route for future Retrieval-Augmented Generation architectures. This is due to the fact that it is capable of semantic retrieval and interpretable relational inference. Experiments have shown that explainability is slowly becoming a more important factor for evaluating management information systems (MIS). According to surveys, the academic community is moving away from evaluating artificial intelligences based on their accuracy and toward evaluating them based on evidence, transparent reasoning, and interpretable decisions [22]. Instead of contextual embeddings, the paradigm that has been offered is for predictions that are based on scientific evidence. The implementation of real-time healthcare misinformation monitoring systems that make use of integrated Knowledge Graph-LLM architectures [30] and retrieval-augmented social media intelligence systems [26] has been accomplished thru the utilization of evidence-based methodologies. In recent years, the paradigm of evidence-based false news detection approaches has been successfully applied to the challenge of identifying false news by utilizing graph neural networks and retrieval augmentation. A successful application of this paradigm has been made.

Both Yarlagadda [25] and Baharifar et al. [27] claimed that interpretable artificial intelligence is required in order to discern between LLM-generated and human-generated medical misinformation. Yarlagadda proved that GNNs that incorporate retrieval can improve the evidence usage for misinformation identification. Mohamed et al. [28] discovered that larger language models that contain a greater number of evaluation criteria have a higher level of performance when it comes to identifying multilingual medical misinformation. Structured knowledge retrieval and semantic understanding are both made possible by these data, and the Graph-RAG framework that has been developed makes advantage of both of these capabilities.

They do, however, come with a significant number of negatives that should be considered. In order to guaranty that Graph-RAG is of a good quality, it is absolutely necessary that the medical knowledge graph be accurate and thorough. In the event that there are insufficient entities or when medical ideas have become obsolete, there is a possibility that the quality of retrieval and the efficiency of verification will be negatively damaged. On the other hand, in comparison to transformer classifiers, the computational complexity of graph retrieval presents an additional obstacle in large-scale deployments. This, in turn, leads to an increase in the amount of time that is necessary for inference.

For this reason, it is necessary to do additional study in the future to investigate the dynamic updating of the knowledge graph, adaptive retrieval optimization, and light-weight graph indexing. In the context of emerging healthcare situations, the innovative methodology of TripleFact [31] to protect against polluted retrieval corpora and boost LLM robustness might be able to assist in making Graph-RAG approaches more reliable. The findings suggest that the Graph-RAG framework is capable of establishing a balance between the ability to explain the judgments that were made, structured medical reasoning, factual verification, and the ability to understand the semantics of the language when it comes to achieving this balance. In spite of the fact that the accuracy of the medical verification is low, the system is able to accurately detect medical misinformation in real-world conditions. This is because all forecasts are supported by medical evidence, which enables the system to correctly identify medical misinformation. The fact that this is the case makes the technique that was suggested an appropriate alternative for high-risk healthcare applications, which are those in which the importance of transparency, trustworthiness, and evidence-based reasoning is greater than the importance of classification accuracy.

5. Conclusion

The findings of the multi-hop knowledge graph inference that is predicated on LLM and classification are combined thru the use of the graph-based RAG in order to increase the efficiency of the detection of fake news in the healthcare industry. This is done in order to help improve the effectiveness of the detection process. The results that it generates are more accurate, interpretable, and consistent with real-world experience as compared to the results that are produced by standard models and standalone LLM. The methodology has the ability to assist in the reduction of hallucinations and the provision of explanations that are based on evidence for the goal of developing trustworthy fundamental healthcare decisions. Both in terms of organized knowledge and semantic comprehension, the findings of comparative research indicate that it is most superior of the two. The platform has the potential to be developed in such a way that it incorporates clinical decision support, surveillance devices, and automated fact checking tools. As well as being extendable, the platform is versatile.

Authors Contribution

Conceptualization, D.J.G.; Validation, D.J.G.; Data curation, D.J.G.; Writing—original draft, D.J.G.; Writing—review & editing, R.K.; Visualization, D.J.G.; Supervision, R.K. All authors have read and agreed to the published version of the manuscript.

Dataset Availability

The datasets used in this study are publicly available through the IEEEDataPort platform. The Covid-19 Fake News Infodemic Research Dataset (CoVID19-FNIR) is available at <https://ieee-dataport.org/open-access/covid-19-fake-news-infodemic-research-dataset-covid19-fnir-dataset>.

Conflict of Interest

The authors declare that there is no conflict of interest.

Acknowledgment

The authors express their sincere gratitude to the Management of Bishop Heber College, the University Grants Commission (UGC) - College of Excellence, and the Department of Science and Technology (DST) - FIST for their invaluable support in facilitating this research work."

References

1. Papageorgiou, E., Chronis, C., Varlamis, I., & Himeur, Y. (2024). A survey on the use of large language models (llms) in fake news. *Future Internet*, 16(8), 298.
2. Sun, Y., He, J., Cui, L., Lei, S., & Lu, C. T. (2024). Exploring the deceptive power of llm-generated fake news: A study of real-world detection challenges. *arXiv preprint arXiv:2403.18249*.
3. Chakraborty, J., Xia, W., Majumder, A., Ma, D., Chaabene, W., & Janvekar, N. (2024). Detoxbench: Benchmarking large language models for multitask fraud & abuse detection. *arXiv preprint arXiv:2409.06072*.
4. Xu, R., Jiang, P., Luo, L., Xiao, C., Cross, A., Pan, S., ... & Yang, C. (2025, August). A survey on unifying large language models and knowledge graphs for biomedicine and healthcare. In *Proceedings of the 31st ACM SIGKDD Conference on Knowledge Discovery and Data Mining V. 2* (pp. 6195-6205).
5. Sağlam, B. D. (2024). Knowledge Graph Augmented Multi-Hop Question Answering Using Large Language Models (Master's thesis, Middle East Technical University (Turkey)).
6. LeewayHertz. (2024). Advanced RAG: Architecture, techniques, applications and use cases. Retrieved from Advanced RAG article
7. Ahmad, M. (2025). Toward a Unified Framework for Information Retrieval in Large Language Model Applications: Balancing Textual and Graph-Based Knowledge Sources.
8. Dehal, R. S., Sharma, M., & Rajabi, E. (2025). Knowledge graphs and their reciprocal relationship with large language models. *Machine Learning and Knowledge Extraction*, 7(2), 38.
9. Mac Dube, B., & Fonou-Dombeu, J. V. (2026). A Review of State-of-the-Art Deep Learning Models for Knowledge Graphs. *IEEE Access*.
10. Melis, A. (2025). Graph-Augmented Retrieval Techniques for Medical QA Applications.
11. Hang, C. N., Yu, P. D., & Tan, C. W. (2025). TrumorGPT: Graph-based retrieval-augmented large language model for fact-checking. *IEEE Transactions on Artificial Intelligence*.
12. Simoni, M. (2026). Toward reliable and adaptive large language models in the cybersecurity domain. <https://iris.uniroma1.it/handle/11573/1760907>
13. Lecu, A., Groza, A., & Hawizy, L. (2025). Reducing Hallucinations in Medical AI: A Knowledge Graph-Augmented Retrieval System for Evidence-Based Age-Related Macular Degeneration Information. *IEEE Access*, 13, 210624-210639.
14. Gong, S., Sinnott, R. O., Qi, J., Paris, C., Nakov, P., & Xie, Z. (2026). Multi-Sourced, Multi-Agent Evidence

- Retrieval for Fact-Checking. arXiv preprint arXiv:2603.00267.
15. Chen, Z., Li, X., Tang, H., Du, W., & Wang, Y. (2025). IBGraphRAG: Enhancing Medical Knowledge Graph Retrieval Based on Semantic Consistency and Information Bottleneck. <https://openreview.net/pdf?id=0XP33CRtQF>
 16. Hang, C. N., Yu, P. D., & Tan, C. W. (2025). TrumorGPT: Graph-based retrieval-augmented large language model for fact-checking. *IEEE Transactions on Artificial Intelligence*, 6(11), 3148-3162.
 17. Nguyen, C. H., Hoang, T., Duong, H. M., & Nguyen, L. (2026). DeLiVeR: Decomposed Learning for Information-grounded Veracity Recognition via Reinforced Knowledge Graph Exploration. arXiv preprint arXiv:2607.17935.
 18. Carvalho, L. H. M., & Goldschmidt, R. R. (2026). Large Language Models for Automated Fact-Checking: A Systematic Literature Review. *IEEE Access*.
 19. Wu, J., Zhu, J., Qi, Y., Chen, J., Xu, M., Menolascina, F., & Grau, V. (2025, July). Medical graph rag: Evidence-based medical large language model via graph retrieval-augmented generation. In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)* (pp. 28443-28467).
 20. Wang, B., Ma, J., Lin, H., Yang, Z., Yang, R., Tian, Y., & Chang, Y. (2026). A Graph-Enhanced Defense Framework for Explainable Fake News Detection with LLM. *ACM Transactions on Information Systems*, 44(5), 1-36.
 21. Wu, J., Zhu, J., Qi, Y., Chen, J., Xu, M., Menolascina, F., & Grau, V. (2024). Medical graph rag: Towards safe medical large language model via graph retrieval-augmented generation. arXiv preprint arXiv:2408.04187.
 22. Nguyen, S. T., Huynh, T. T., & Le, T. A. (2026). Reasoning-Centric Fake News Detection: A Comprehensive Survey of Architectures, Benchmarks, and Open Challenges. *IEEE Access*.
 23. Haque, A., Din, A., Babar, M., Abbas, A., & Ullah, I. (2025). GraphRAG-Causal: A novel graph-augmented framework for causal reasoning and annotation in news. arXiv preprint arXiv:2506.11600.
 24. Chowdhury, T. (2026). A Survey on Graph-Based Retrieval-Augmented Generation: Architectures, Methods, and Applications. *International Journal on Computational Modelling Applications*, 3(1), 15-28.
 25. Yarlagadda, D. (2026). A RAG-Augmented Graph Neural Network for Evidence-Based Fake News Detection (Master's thesis, University of Georgia).
 26. Berte, S. (2025). Retrieval-Augmented Social Media Intelligence: Detecting and Reporting of High-Risk Communication Patterns using Large Language Models (Doctoral dissertation, Politecnico di Torino).
 27. Baharifar, M., Kujur, A., & Monfared, Z. Interpretable AI for Classifying Human-and LLM-Generated Medical Misinformation with Multi-Modal Features.
 28. Mohamed, E., Ismail, W. N., & Eldawy, E. O. (2026). Enhancing Arabic healthcare fake news detection with data augmentation and multi-metric analysis using large language models. *Scientific Reports*, 16(1), 5364.
 29. Shupta, A., Radiuk, P., & Krak, I. (2025, April). Feature computation procedure for fake news detection: An LLM-based extraction approach. In *Proceedings of the 6th international workshop on intelligent information technologies & systems of information security (intelitsis 2025)* (pp. 112-124).
 30. Clark, O., & Joshi, K. P. (2025, July). Real-Time Detection of Online Health Misinformation using an Integrated Knowledgegraph-LLM Approach. In *2025 IEEE International Conference on Digital Health (ICDH)* (pp. 111-120). IEEE.
 31. Xu, C., & Yan, N. (2025, July). TripleFact: Defending data contamination in the evaluation of LLM-driven fake news detection. In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)* (pp. 8808-8823).