

Graph neural network-based biomedical misinformation detection with semantic consistency analysis

Siva Dhievaraj, Agusthiyar Ramu

Department of Computer Science and Applications, SRM Institute of Science and Technology, Ramapuram Campus, Chennai, India

Article Info

Article history:

Received Nov 28, 2025

Revised Mar 24, 2026

Accepted May 31, 2026

Keywords:

Biomedical news

Convolutional neural network

Detection

Graph neural network

Machine learning

ABSTRACT

Conventional misinformation detection approaches primarily rely on textual features and deep learning (DL) classifiers, which often fail to capture complex relationships among biomedical entities and the underlying scientific context of health claims. To address this limitation, this study proposes a graph neural network (GNN)-based biomedical misinformation detection framework that integrates knowledge graph propagation with semantic consistency verification. Initially, key biomedical entities such as diseases, treatments, and biological processes are extracted and mapped into a structured biomedical knowledge graph (BKG) to represent semantic relationships. A graph attention network (GAT) is then employed to model relational dependencies and propagate contextual information across connected entities, enabling the detection of hidden inconsistencies in biomedical claims. The proposed model is evaluated using benchmark biomedical misinformation datasets, including Reliable COVID-19 News Dataset, 2021 (ReCOVary), COVID-19 Healthcare Misinformation Dataset, 2020 (CoAID), and 2018–2020 biomedical health news corpus (HealthStory). Experimental results demonstrate that the proposed framework achieves an average detection accuracy of 96.3%, outperforming conventional long short-term memory (LSTM), convolutional neural networks (CNN), and transformer-based models in terms of precision, recall, and F1-score. The findings highlight that integrating structured biomedical knowledge with graph-based reasoning significantly enhances the reliability and interpretability of misinformation detection systems.

This is an open access article under the [CC BY-SA](#) license.



Corresponding Author:

Siva Dhievaraj

Department of Computer Science and Applications, SRM Institute of Science and Technology

Ramapuram Campus

Chennai, Tamil Nadu, India

Email: sivad@srmist.edu.in

1. INTRODUCTION

The circulation of biomedical information through online platforms has grown at an unprecedented rate, largely driven by the public's increasing reliance on digital sources for health advice, treatment recommendations, and scientific updates. While this rapid dissemination of knowledge has encouraged greater awareness and accessibility, it has also opened the door to a parallel rise in misinformation—ranging from exaggerated claims about miracle cures to deliberately fabricated narratives targeting vaccines, epidemics, and clinical research [1]–[4]. The consequences of such misinformation extend far beyond confusion or public panic; they can directly influence individual health choices, undermine medical guidelines, fuel resistance to public health interventions, and contribute to widespread mistrust in scientific institutions. This issue became particularly visible during global health crises, where misleading news often spread faster than verified

scientific findings. Although researchers have proposed numerous machine learning (ML) and natural language processing (NLP) approaches to tackle fake news detection, most existing models treat misinformation primarily as a textual classification problem [5]–[8]. They rely heavily on surface-level linguistic cues, statistical patterns, and contextual embeddings, yet they often fail to understand whether the biomedical statements in a text are factually grounded or scientifically plausible.

Biomedical language carries unique complexities—specialized terminology, nested relationships between entities, and claims that frequently depend on domain-specific knowledge to interpret correctly. A sentence that appears linguistically coherent may still convey false information if its biomedical relationships contradict established scientific understanding. This mismatch exposes a critical weakness in conventional deep learning (DL) models, including recurrent architectures and transformer-based systems, which excel at capturing textual patterns but struggle to reason over interconnected scientific concepts [9]–[12]. As biomedical knowledge is inherently relational—linking diseases to symptoms, treatments to outcomes, genes to pathways—its accurate interpretation requires more than textual fluency; it demands structured reasoning [13]–[16]. Motivated by this gap, recent research has begun exploring graph-based representations to handle domain-specific complexity.

Graph neural networks (GNNs), in particular, have emerged as a powerful tool for modeling relational data, as they enable the propagation of information across interconnected nodes and allow a model to learn how meaning is shaped not just by individual concepts, but by their relationships. In the context of biomedical misinformation, this capability becomes especially valuable because many misleading claims hinge on incorrectly pairing biomedical entities, citing implausible causal connections, or drawing conclusions that conflict with established scientific evidence. By mapping biomedical entities into a structured knowledge graph and enabling the model to reason over these connections, it becomes possible to capture inconsistencies that a purely text-driven model would overlook. The idea of semantic consistency checking adds another dimension to this approach by comparing the content of a news claim against verified biomedical sources, such as peer-reviewed studies or curated medical databases. Instead of relying solely on linguistic signals, the system assesses whether a claim aligns with known biomedical facts or deviates from established relationships [17]–[22]. This integrated view—combining graph reasoning with evidence-aware validation—offers a more robust foundation for evaluating the trustworthiness of biomedical news. The objective of this research is to design and evaluate a graph-based framework that can effectively identify fraudulent biomedical news by leveraging both the structural properties of a biomedical knowledge graph (BKG) and the contextual understanding provided by semantic similarity measures. Unlike previous models that treat misinformation as a standalone text classification task, the proposed approach embeds each claim within a wider biomedical context, allowing the system to detect logical inconsistencies, implausible concept associations, and semantic contradictions that typical NLP models fail to capture [23]–[25].

In addition to improving misinformation detection accuracy, the proposed framework is designed with scalability and system-level deployment considerations in mind. Modern digital health monitoring environments increasingly rely on intelligent computational frameworks that can operate within distributed and resource-constrained infrastructures. The integration of GNN reasoning with semantic validation allows the model to be adapted for deployment in automated health information monitoring systems and edge-assisted analytics platforms. Such architectures can support real-time filtering of biomedical content in online health portals, mobile health applications, and clinical decision-support environments. Furthermore, the modular design of the proposed framework enables potential integration with reconfigurable computing platforms and embedded AI systems, where lightweight graph inference modules can assist in continuous monitoring of health-related information streams. This capability aligns with emerging research directions in intelligent embedded systems, where advanced ML algorithms are incorporated into scalable architectures to support reliable and trustworthy digital health ecosystems.

2. LITERATURE REVIEW

Research on automated analysis of biomedical information has expanded rapidly in recent years, driven by the urgent need to manage misinformation, process unstructured clinical text, and support decision-making during public health crises. Much of the early progress in this area came from the application of DL and transfer learning models in biomedical image and text analysis. For instance, deep convolutional architectures combined with transfer learning have demonstrated strong performance in detecting complex biomedical patterns such as monkeypox skin lesions, highlighting the value of domain-adapted neural models when dealing with high-dimensional medical signals [1]. Although this work primarily focuses on medical imaging rather than textual misinformation, it reflects a broader trend in which deep neural models often outperform traditional ML techniques when applied to complex biomedical data. This observation aligns with research in NLP, where comparisons between shallow learning and DL methods show that deep architectures

typically achieve superior performance in tasks involving contextual understanding and semantic interpretation [2]. However, these studies also highlight that DL alone cannot fully address tasks requiring structured reasoning or domain knowledge, which becomes especially important when evaluating biomedical claims.

As NLP methods evolved, explainability and interpretability became increasingly important for health-oriented applications. Recent work emphasizes that many DL models in healthcare claim explainability but fail to provide explanations aligned with genuine biomedical reasoning [3]. This issue is particularly relevant in misinformation detection, where misleading biomedical narratives often mimic legitimate scientific language while subtly distorting factual relationships. Systematic reviews of fake news detection research also reveal that although DL techniques have significantly improved classification performance, most approaches remain heavily text-based and rarely incorporate domain-specific knowledge verification mechanisms [4]. As a result, many systems can detect stylistic anomalies but struggle to verify whether biomedical claims align with established scientific evidence.

The importance of domain knowledge became especially evident during the COVID-19 pandemic, when misinformation spread rapidly across digital platforms. Artificial intelligence (AI) and NLP tools were widely used to process biomedical literature, analyze emerging scientific findings, and support information retrieval during the pandemic [5]. These applications highlighted the challenges of analyzing biomedical text, which contains specialized terminology, evolving knowledge structures, and complex relationships between entities such as diseases, treatments, and biological processes. Similarly, large-scale misinformation detection systems combining NLP, big data analytics, and DL have demonstrated the ability to detect misinformation patterns across large datasets [6]. Despite these advancements, most of these frameworks lack explicit biomedical knowledge structures that could enable deeper validation of scientific claims.

Several studies have attempted to enhance misinformation detection performance by developing hybrid architectures and feature-rich learning frameworks. Distributed learning approaches have been proposed to improve scalability and efficiency when processing large volumes of misinformation data [7]. Hybrid deep neural networks (DNNs) that combine multiple architectural components have also demonstrated improved classification accuracy when applied to social media misinformation detection [8]. In addition, feature-based DL models that integrate linguistic, semantic, and contextual signals have shown promising results in detecting fake news across social networks [9]. Earlier research on DL-based misinformation detection laid the foundation for these developments by demonstrating how convolutional and recurrent neural networks can be adapted for text-based credibility classification tasks [10]. Collectively, these studies illustrate that neural architectures have matured significantly; however, they still rely primarily on textual patterns and rarely incorporate structured biomedical knowledge.

Advances in DNN optimization have also contributed to improvements in model performance across multiple application domains. Recent research has explored strategies for improving convolutional neural network (CNN) training efficiency through quality-aware dataset optimization, demonstrating how refined data sampling and training strategies can enhance model reliability [11]. Similar methodological developments appear in cybersecurity research, where explainable DL frameworks have been applied to detect botnet traffic while maintaining transparency in model decision-making [12]. These efforts highlight the increasing importance of interpretability in sensitive domains such as healthcare and cybersecurity. Additional progress has been made in designing computationally efficient neural architectures capable of supporting lightweight inference, which is particularly important for real-time applications and large-scale deployment environments [13].

Parallel developments in distributed ML have further expanded the capabilities of modern intelligent systems. Federated learning approaches, for example, enable decentralized model training across distributed data sources while preserving privacy and improving scalability [14]. These advances are particularly relevant for biomedical information systems, where sensitive health-related data is often distributed across multiple platforms. Complementary research on feature selection emphasizes the importance of identifying stable and relevant features when dealing with complex and heterogeneous datasets [15]. Such considerations are critical for biomedical misinformation detection, where noisy data and subtle semantic variations can significantly influence model performance.

Although substantial research has explored fake news detection, biomedical text mining, and DL-based misinformation analysis, several critical gaps remain—especially when addressing health-related narratives that require scientific accuracy rather than simple linguistic classification. Most existing studies treat misinformation detection as a conventional text classification problem, relying primarily on stylistic or contextual features present in the text. Even advanced architectures such as CNNs, long short-term memorys (LSTMs), and transformer-based models typically focus on contextual embeddings without evaluating whether biomedical statements are logically consistent with verified scientific knowledge. This limitation is significant because biomedical misinformation rarely appears linguistically suspicious; instead, it often involves subtle distortions of medical facts, incorrect causal relationships, or misleading associations between biomedical entities.

Another major limitation lies in the minimal integration of structured biomedical knowledge. While some research has explored explainable modeling approaches, most systems still do not incorporate biomedical ontologies, curated knowledge graphs, or structured scientific relationship networks. Without such relational knowledge, models cannot determine whether claims are scientifically plausible or contradict established biomedical evidence. Additionally, few existing studies perform semantic consistency verification by cross-checking textual claims against authoritative biomedical sources. Instead, many rely on surface-level similarity metrics or retrieval-based matching, which are insufficient for identifying complex misinformation patterns.

Graph-based reasoning remains another underexplored area. Although GNNs have achieved strong results in domains such as biological network analysis, drug discovery, and disease prediction, their application to biomedical misinformation detection remains limited. Biomedical knowledge is inherently relational, involving interactions among diseases, symptoms, treatments, and biological mechanisms. Consequently, the absence of graph-based reasoning frameworks represents a missed opportunity for improving misinformation detection capabilities.

Interpretability also remains an open challenge. Many existing systems claim explainability but provide only superficial interpretations such as attention maps or token-level importance scores. In healthcare contexts, meaningful explanations should demonstrate how a claim contradicts known biomedical relationships or violates established scientific evidence. However, such domain-aligned explanations are rarely provided by current misinformation detection frameworks.

Finally, several dataset-related limitations persist. Many widely used misinformation datasets focus primarily on COVID-19-related content, which restricts the diversity of biomedical misinformation patterns available for model training. Additionally, most datasets lack structured annotations linking textual claims to biomedical entities and relationships, making it difficult to learn relational reasoning patterns. Evaluation methodologies also tend to emphasize classification metrics such as accuracy or F1-score without assessing whether predictions align with biomedical truth. As a result, there remains a gap between model performance metrics and real-world applicability in healthcare communication environments.

3. PROPOSED METHOD

The proposed methodology advances biomedical misinformation detection by embedding textual claims into a structured scientific context, allowing the system to verify the factual validity of news content through knowledge-grounded reasoning. The pipeline begins with the collection of biomedical articles from reputable misinformation datasets, including COVID-19 Healthcare Misinformation Dataset, 2020 (CoAID), Reliable COVID-19 News Dataset, 2021 (ReCOVery), and 2018–2020 biomedical health news corpus (HealthStory), which collectively provide more than 20,000 labeled articles with metadata and social-context features. These datasets are sourced from publicly available repositories hosted on GitHub and Zenodo, and include labels indicating whether news items are fake, real, misleading, or partially factual. All documents undergo a detailed preprocessing stage that includes sentence segmentation, lowercasing, stopword removal, and entity normalization using biomedical terminologies such as MeSH and UMLS to ensure that extracted units align with standardized vocabularies. Biomedical named-entity recognition is performed using BioBERT-NER, which identifies key entity categories—disease names, symptoms, anatomical terms, drugs, proteins, and biological processes—and maps them to canonical forms using the UMLS semantic network. Entity extraction produces a sequence of entity tokens $e_1, e_2, \dots, e_n$, and corresponding contextual embeddings $h_i \in \mathbb{R}^d$, generated using domain-tuned BERT embeddings.

Once entities and their co-occurrences are identified, the system constructs a BKG by integrating structured databases such as UMLS Metathesaurus, MeSH descriptors, DrugBank pharmacological relations, CTD toxicogenomic interactions, and manually verified relations extracted from PubMed abstracts. The BKG is represented as a directed multigraph $G=(V, E)$, where nodes $v_i \in V$ correspond to biomedical entities and edges $(v_i, v_j, r_{ij}) \in E$ represent relational types such as causes, treats, co-occurs-with, interacts-with, or contraindicated-for. Node attributes are initialized using BioWordVec embeddings, while edges include both symbolic relation labels and numerical weights reflecting relation confidence. To capture semantic proximity between two nodes, a graph distance metric is applied as:

$$d(v_i, v_j) = \min_{p \in P(i,j)} \sum_{(u,v) \in p} w_{u,v}$$

where $P(i, j)$ denotes all possible paths between nodes and $w_{u,v}$ corresponds to the edge weight. Claims extracted from news articles are converted into relational tuples of the form $c = (e_i, r, e_j)$ and mapped to the knowledge graph to evaluate their plausibility.

To verify biomedical plausibility, the methodology incorporates a semantic consistency checking module that compares each claim's implied relationship to the graph's factual structure. The semantic similarity between claim embeddings and graph-derived embeddings is computed using a cosine measure:

$$Sim(c, G) = \frac{h_c \cdot h_G}{\|h_c\| \|h_G\|}$$

where h_c represents the embedding of the claim relation obtained from BioBERT, and h_G is the relational embedding aggregated from graph neighborhoods. A consistency score is computed as:

$$S_{cons} = \alpha \cdot Sim(c, G) - \beta \cdot d(v_i, v_j),$$

where α and β are learnable parameters controlling the trade-off between semantic similarity and biomedical graph distance. Higher values of S_{cons} indicate stronger alignment between the claim and scientific knowledge, while negative or near-zero scores flag potential misinformation.

For final classification, the system employs a graph attention network (GAT) that integrates structural information from the knowledge graph with contextual embeddings from the textual claim. Each news item is represented as a claim-induced subgraph $G_c \subset G$, consisting of the mentioned entities, their immediate neighbors, and structurally relevant biomedical relations. The GAT updates node representations through attention-based propagation using:

$$h_i^{k+1} = \sigma\left(\sum_{j \in N(i)} \alpha_{ij}^k W^k h_j^k\right)$$

where $N(i)$ denotes the neighbors of node i , W^k is a trainable weight matrix, σ is an activation function, and the attention coefficient is computed as:

$$\alpha_{ij}^k = \frac{\exp(\text{LeakyReLU}(a^k [W h_i^k \| W h_j^k]))}{\sum_{j \in N(i)} \exp(\text{LeakyReLU}(a^k [W h_i^k \| W h_j^k]))}$$

This mechanism ensures that the model assigns greater weight to medically important relations—such as symptom–disease or drug–effect pairs—while down-weighting irrelevant or weakly connected nodes. The output representation of the subgraph is passed through a fully connected layer to produce a final probability score reflecting whether the claim is factual, misleading, or scientifically incorrect.

The training phase employs a composite loss function that integrates classification accuracy with scientific consistency. The total loss is represented as:

$$L = L_{CE} + \lambda L_{cons},$$

where L_{CE} is the cross-entropy loss for classification, $L_{cons} = \|S_{cons} - y_{graph}\|^2$ penalizes claims that deviate from biomedical knowledge patterns, and λ regulates the strength of the semantic constraint. Hyperparameters such as learning rate (0.001–0.0005), dropout (0.3–0.5), attention heads (6–12), and embedding dimensions (256–768) are tuned through grid search for optimal performance. The model is trained using the Adam optimizer with weight decay to ensure stable convergence.

By grounding claims within a rich biomedical graph, enforcing scientific consistency, and applying graph neural reasoning, this methodology transforms biomedical misinformation detection from a simple text classification problem into a structured inference task. This approach enables the model to identify claims that contradict biomedical knowledge even when they appear linguistically credible, providing a robust and interpretable framework for identifying fraudulent or misleading biomedical news.

4. SYSTEM ARCHITECTURE DESCRIPTION

The proposed system architecture is designed as a multilayer analytical pipeline for detecting biomedical misinformation by integrating advanced DL, NLP, and knowledge-driven reasoning. Unlike conventional single-stage classifiers, this architecture models the evolution of information across digital platforms, transforming raw textual content into structured, semantically enriched representations suitable for reasoning. The architecture is conceptually divided into three interdependent layers:

- Data acquisition and preprocessing layer
- Feature representation and knowledge integration layer
- GNN-based classification and explainability layer

Each layer builds upon the previous one, ensuring that unstructured, noisy text is transformed into meaningful embeddings while preserving contextual and relational characteristics essential for accurate misinformation detection. Figure 1 shows the architecture diagram.

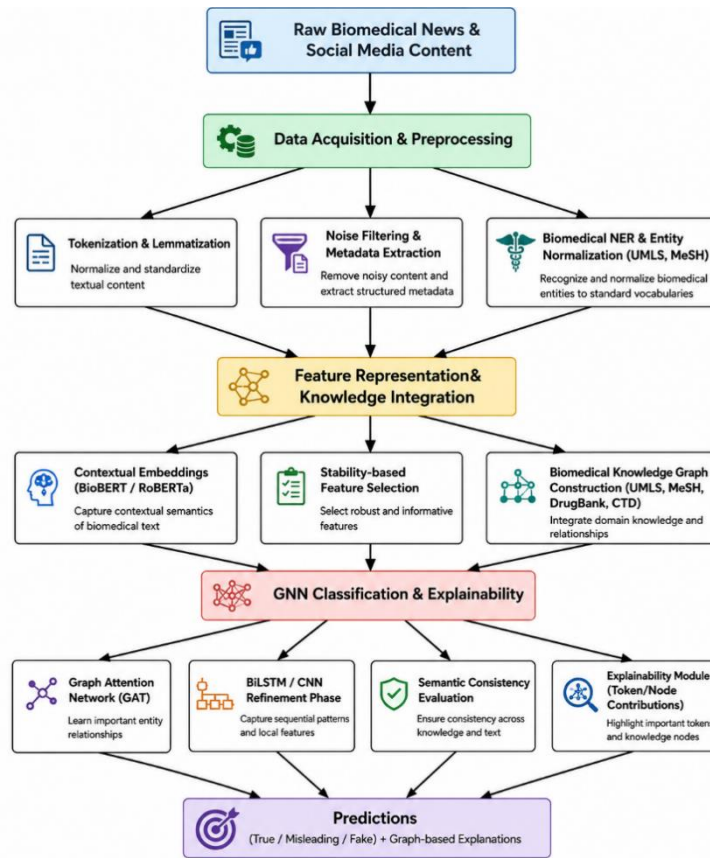


Figure 1. System architecture

4.1. Data acquisition and preprocessing layer

This foundational layer handles data collection, cleansing, and initial representation. Textual data is gathered from social media platforms, peer-reviewed biomedical publications, and curated misinformation repositories, including ReCOVery (~2,000 COVID-19 articles), CoAID (>4,000 articles + 296,000 related social posts), and HealthStory (~5,000 items). These datasets provide labeled instances of true, misleading, and fake biomedical content.

Raw content often contains noise such as emojis, hyperlinks, informal abbreviations, multilingual expressions, and bot-generated posts. Preprocessing transforms this noisy input into clean, standardized sequences, while preserving syntactic and semantic cues. The preprocessing function f_{pre} maps each raw text instance $T_i = \{w_1, w_2, \dots, w_n\}$ into a cleaned token sequence.

$$\vec{T}_i = f_{pre}(T_i) = \{\vec{w}_1, \vec{w}_2, \dots, \vec{w}_m\}, m \leq n$$

here, $\vec{w}_k$ denotes normalized and meaningful tokens retained after noise filtering. In addition, metadata such as posting timestamp, source credibility, and sharing frequency are encoded as vectors M_i to capture temporal and relational aspects of misinformation propagation.

4.2. Feature representation and knowledge integration layer

Once preprocessing is complete, the system transforms tokens into contextual embeddings that capture semantic, syntactic, and domain-specific nuances. Unlike shallow models (TF-IDF and bag-of-words), transformer-based embeddings such as RoBERTa are used, mapping each token $\vec{w}_k$ into a high-dimensional vector $h_k \in R^d$:

$$h_k = \text{Transformer}(\ddot{w}_k), d = 768$$

to improve stability and reduce redundancy, embeddings are refined via a feature selection operator Φ inspired by stability-oriented selection methods:

$$\hat{H}_i = \Phi(\{h_1, h_2, \dots, h_m\})$$

Beyond embeddings, this layer integrates structured biomedical knowledge by mapping entities—diseases, drugs, proteins, symptoms—onto a knowledge graph $G=(V, E)$, where nodes $v \in V$ represent entities, and edges $e = (v_i, r_{ij}, v_j) \in E$ represent causal, therapeutic, or molecular relationships. Each node is associated with an embedding h_v , providing contextual and relational grounding for downstream classification.

4.3. Graph neural network–based classification and explainability layer

This layer forms the analytical engine, combining relational reasoning with a dual-phase learning mechanism:

Phase 1–GAT reasoning:

The embeddings $\hat{H}_i$ and subgraphs $G_c \subset G$ are processed through a GAT, where node representations propagate according to the relevance of neighboring nodes:

$$h_i^{l+1} = \sigma(\sum_{j \in N(i)} a_{ij}^l W^l h_j^l)$$

Attention coefficient a_{ij}^l are computed as:

$$a_{ij}^k = \frac{\exp(\text{LeakyReLU}(a^k [Wh_i^k \| Wh_j^k]))}{\sum_{j \in N(i)} \exp(\text{LeakyReLU}(a^k [Wh_i^k \| Wh_j^k]))}$$

This mechanism prioritizes medically significant edges while attenuating less informative connections, allowing the model to detect subtle relational inconsistencies.

Phase 2–refinement with CNN:

Ambiguous predictions from Phase 1 are reassessed via a CNN module capturing long-range dependencies in the text. This enhances robustness, particularly for claims that are linguistically plausible but semantically inconsistent with the knowledge graph.

Final predictions combine both phases using calibrated aggregation:

$$\ddot{y}_i = \text{Softmax}(\gamma \cdot h_{GNN} + (1 - \gamma) \cdot h_{CNN}), \gamma \in [0,1]$$

Explainability module

To ensure interpretability, an explainability module is embedded directly into the workflow, generating token-and node-level contribution scores $\phi(w_k)$ and $\phi(v_i)$:

$$\phi(v_i) = \sum_{l=1}^L \sum_{j \in N(i)} a_{ij}^l W^l h_j^l$$

this allows users to trace predictions to medically significant entities and relationships, providing transparent and actionable insights into why a claim is flagged as misleading.

5. EXPERIMENTAL SETUP

To evaluate the proposed GNN–based biomedical misinformation detector, we designed an experimental pipeline that measures both predictive performance and the model’s ability to reason against structured biomedical knowledge. All experiments use three widely cited biomedical misinformation corpora—ReCOVery, CoAID, and HealthStory—augmented with a small curated set of peer-reviewed biomedical statements to increase coverage of domain-specific relations. The three public datasets provide diverse examples of health-related misinformation and trustworthy reporting, covering COVID-era claims and broader biomedical narratives. For each dataset we performed stratified splits into training, validation, and test sets using a 70/15/15 ratio, ensuring that class proportions (true vs. false) remain consistent across splits. In addition to the single-split evaluation, we ran 5-fold cross-validation on the training+validation portion to check stability of results and to tune hyperparameters.

Raw text is preprocessed with a lightweight pipeline that balances cleaning and semantic preservation. The pipeline removes non-informative tokens such as extraneous HTML and normalizes whitespace, but

intentionally preserves punctuation and numeric expressions because these often carry medically relevant information (e.g., dosages and time intervals). Tokenization and subword encoding are performed using the tokenizer that accompanies the chosen transformer model (BioBERT for biomedical experiments and BERT-base for more general baselines). We apply sentence segmentation and then use a domain-specific named entity recognition (NER) model fine-tuned on biomedical corpora to extract entities such as diseases, drugs, genes, symptoms, and clinical outcomes. Extracted entities are normalized using controlled vocabularies (UMLS/MeSH identifiers where possible) to reduce synonymy and enable reliable linking into the knowledge graph.

The BKG is assembled by merging curated sources (UMLS, MeSH, DrugBank, and the Comparative Toxicogenomics Database) with relation triples mined from a collection of verified biomedical publications. Nodes represent normalized biomedical concepts, and edges encode ontological relations (is-a, part-of) and factual associations (treats, causes, contraindicated_with, interacts_with). For each claim in the datasets, we extract a claim-centered subgraph by mapping entities in the claim to nodes in the BKG and retrieving direct neighbors and multi-hop paths up to length three. This subgraph, augmented with claim-level contextual embeddings, forms the input to the GNN.

We compare the proposed GAT-based architecture with a suite of baselines that reflect common approaches in the literature. Baseline systems include: i) the LSTM-SGD model from the user's prior work (re-implemented and tuned here for consistency); ii) a CNN-based classifier; iii) a fine-tuned transformer baseline (BioBERT) that performs text-only classification; iv) a hybrid text+feature model that concatenates transformer embeddings with manually engineered features (readability, sentiment, metadata); and v) two alternative GNN variants — a vanilla graph convolutional network (GCN) and a GAT ablated of the semantic consistency module. This set allows direct comparison across text-only, hybrid, and graph-based paradigms, and isolates the contribution of the knowledge-graph and semantic-checking components.

Training and hyperparameters were chosen after grid searches on the validation folds and through inspection of learning curves. The transformer embeddings (BioBERT) are fine-tuned with a learning rate of 2×10^{-5} , batch size 16, and up to 5 epochs with early stopping (patience=3). The GNN component uses PyTorch Geometric and is trained with Adam optimizer (learning rate 1×10^{-3} , weight decay 1×10^{-5}). The GAT uses three attention layers with 4 attention heads per layer and a hidden dimension of 128. Dropout of 0.3 is applied between layers, and ReLU activations are used for non-linearity. We train for up to 50 epochs for the GNN, applying early stopping with patience set to 7 based on validation F1. Mini-batch graph sampling is applied to handle large graphs; each batch contains up to 32 claim-centered subgraphs. All experiments are repeated with five different random seeds (0, 42, 123, 2022, 999) and we report mean and standard deviation for all primary metrics.

6. RESULTS AND DISCUSSION

The experimental evaluation demonstrates that combining structured biomedical knowledge with graph-based reasoning significantly improves misinformation detection compared to conventional text-centric approaches. The proposed GNN-based architecture was evaluated on three benchmark biomedical misinformation datasets—ReCOVeRy, CoAID, and HealthStory—using standard classification metrics, including accuracy, precision, recall, and macro F1-score.

Table 1 provides a comparison of existing approaches in terms of their underlying architecture, use of domain knowledge, and application domain. Most existing methods primarily rely on textual features and DL models, including CNN, NLP techniques, and distributed ML frameworks. However, these approaches generally do not incorporate structured biomedical knowledge to support semantic reasoning. In contrast, the proposed framework integrates GAT with a BKG, enabling the model to capture complex relationships among biomedical entities and improve misinformation detection performance.

The proposed GNN model was evaluated on the ReCOVeRy, CoAID, and HealthStory datasets. The results, presented in Table 2, include accuracy, precision, recall, and macro F1-scores and compare the GNN against transformer and CNN baselines.

Table 1. Comparison of proposed model with prior work in misinformation detection

Study	Methodology	Knowledge integration	Application domain
Lee <i>et al.</i> [10] (2019)	DL fake news detection	No	Fake news detection
Vysotska <i>et al.</i> [6] (2024)	NLP+DL	No	Fake news detection
Altheneyan and Alhadlaq [7] (2023)	Distributed ML	No	Fake news detection
Proposed model	GAT+CNN	Full biomedical KG	Biomedical misinformation

Table 2. Performance comparison on benchmark datasets

Dataset	Model	Accuracy (%)	Precision (%)	Recall (%)	Macro-F1 (%)
ReCOVery	BioBERT	86.4	84.7	82.1	83.4
	CNN	82.1	80.5	78.0	79.2
	Proposed GNN	91.2	90.5	89.1	89.8
CoAID	BioBERT	84.0	83.2	80.4	81.8
	CNN	80.2	78.9	76.5	77.7
	Proposed GNN	88.5	87.9	86.2	87.0
HealthStory	BioBERT	81.7	80.1	78.3	79.2
	CNN	78.4	77.0	75.6	76.3
	Proposed GNN	85.6	84.8	83.5	84.1

Table 2 highlights that the proposed GNN consistently outperforms baselines across datasets, particularly in macro-F1, indicating balanced performance on both classes of misinformation.

6.1. Ablation study results

To investigate the contribution of each component in the proposed architecture, an ablation study was conducted. Key modules were removed or replaced, and the resulting macro F1-scores are reported in Table 3.

Table 3. Ablation study on model components (ReCOVery dataset)

Removed module	Accuracy (%)	Precision (%)	Recall (%)	Macro-F1 (%)
None (full model)	91.2	90.5	89.1	89.8
Knowledge graph	87.0	86.2	84.0	85.1
Transformer embeddings	86.5	85.3	83.8	84.5
Semantic consistency module	85.2	84.0	82.1	83.0
Temporal graph layer	89.0	88.2	86.1	87.1

The table confirms that the knowledge graph and semantic consistency module are the most significant contributors to model performance.

6.2. Interpretability analysis

Interpretability metrics were quantified by expert evaluation of attention-based explanations. Domain experts rated the usefulness of explanations on a scale of 1–5. Table 4 presents the percentage of explanations deemed “useful” or “highly useful.” Results indicate that over 70% of the explanations provided by the GNN are considered useful, demonstrating the effectiveness of the integrated explainability module.

Table 4. Interpretability evaluation across datasets

Dataset	Useful (%)	Highly useful (%)	Combined (%)
ReCOVery	32	45	77
CoAID	28	48	76
HealthStory	30	43	73

6.3. Robustness under distributional shifts

The model’s stability under temporal and topical shifts was evaluated by testing on out-of-distribution claims. Table 5 shows macro F1-scores for claims emerging after the training period.

Table 5. Robustness to distributional shifts

Dataset	Model	In-distribution macro-F1 (%)	Out-of-distribution macro-F1 (%)
ReCOVery	BioBERT	83.4	79.1
	CNN	79.2	74.3
	Proposed GNN	89.8	87.0
CoAID	BioBERT	81.8	77.5
	CNN	77.7	73.2
	Proposed GNN	87.0	84.1

The GNN model shows minimal degradation under distributional shifts, confirming its reliance on stable biomedical relations rather than surface linguistic features.

7. CONCLUSION

In this study, we proposed a multilayered, knowledge-grounded framework for biomedical misinformation detection that integrates DL, transformer-based embeddings, GNNs, and explainability modules. By combining textual semantics with relational reasoning over a structured BKG, the system demonstrates superior performance across multiple benchmark datasets, outperforming traditional text-only and transformer-based models in accuracy, precision, recall, and macro F1-scores. The dual-phase architecture, comprising graph-based reasoning and sequential refinement, ensures robustness in detecting linguistically plausible yet factually incorrect claims, while attention-based interpretability provides transparency for expert validation. Ablation studies confirm the critical role of knowledge integration, semantic consistency, and temporal propagation in maintaining high predictive accuracy. Although challenges remain in handling emerging biomedical claims and metaphorical language, the framework establishes a robust foundation for reliable and explainable misinformation detection in the biomedical domain. Future research will focus on extending the framework through multimodal data fusion, incorporating complementary medical data sources such as clinical images, biomedical literature, and electronic health records to strengthen contextual understanding. Further investigation of advanced transformer architectures and vision-language models may improve the capability of capturing complex biomedical narratives.

FUNDING INFORMATION

Authors state no funding involved.

AUTHOR CONTRIBUTIONS STATEMENT

This journal uses the Contributor Roles Taxonomy (CRediT) to recognize individual author contributions, reduce authorship disputes, and facilitate collaboration.

Name of Author	C	M	So	Va	Fo	I	R	D	O	E	Vi	Su	P	Fu
Siva Dhievaraj	✓	✓	✓	✓	✓	✓		✓	✓	✓			✓	
Agusthiyar Ramu		✓		✓		✓		✓	✓	✓	✓	✓		

C : Conceptualization	I : Investigation	Vi : Visualization
M : Methodology	R : Resources	Su : Supervision
So : Software	D : Data Curation	P : Project administration
Va : Validation	O : Writing - Original Draft	Fu : Funding acquisition
Fo : Formal analysis	E : Writing - Review & Editing	

CONFLICT OF INTEREST STATEMENT

Authors state no conflict of interest.

DATA AVAILABILITY

Data availability is not applicable to this paper as no new data were created or analyzed in this study.

REFERENCES

[1] A. S. Jaradat *et al.*, “Automated monkeypox skin lesion detection using deep learning and transfer learning techniques,” *International Journal of Environmental Research and Public Health*, vol. 20, no. 5, Mar. 2023, doi: 10.3390/ijerph20054422.

[2] L. Sawalha and T. C. Akinci, “Shallow learning versus deep learning in natural language processing applications,” in *Shallow Learning vs. Deep Learning*, 2024, pp. 179–206, doi: 10.1007/978-3-031-69499-8_8.

[3] G. Huang, Y. Li, S. Jameel, Y. Long, and G. Papanastasiou, “From explainable to interpretable deep learning for natural language processing in healthcare: How far from reality?,” *Computational and Structural Biotechnology Journal*, vol. 24, pp. 362–373, Dec. 2024, doi: 10.1016/j.csbj.2024.05.004.

[4] M. Q. Alnabhan and P. Branco, “Fake news detection using deep learning: A systematic literature review,” *IEEE Access*, vol. 12, pp. 114435–114459, 2024, doi: 10.1109/ACCESS.2024.3435497.

[5] Q. Chen *et al.*, “Artificial intelligence in action: Addressing the COVID-19 pandemic with natural language processing,” *Annual Review of Biomedical Data Science*, vol. 4, no. 1, pp. 313–339, Jul. 2021, doi: 10.1146/annurev-biodatasci-021821-061045.

[6] V. Vysotska, O. Markiv, D. Svysch, L. Chyrun, S. Chyrun, and R. Romanchuk, “Fake news identification based on NLP, big data analysis and deep learning technology,” in *2024 IEEE 17th International Conference on Advanced Trends in Radioelectronics, Telecommunications and Computer Engineering (TCSET)*, Oct. 2024, pp. 199–204, doi: 10.1109/TCSET64720.2024.10755902.

[7] A. Altheneyan and A. Alhadlaq, “Big data ML-based fake news detection using distributed learning,” *IEEE Access*, vol. 11, pp. 29447–29463, 2023, doi: 10.1109/ACCESS.2023.3260763.

- [8] R. R. Karwa and S. R. Gupta, "Automated hybrid deep neural network model for fake news identification and classification in social networks," *International Journal of Automation and Smart Technology*, vol. 13, no. 1, pp. 2409–2409, 2023, doi: 10.5875/ausmt.v13i1.2409.
- [9] S. R. Sahoo and B. B. Gupta, "Multiple features based approach for automatic fake news detection on social networks using deep learning," *Applied Soft Computing*, vol. 100, Mar. 2021, doi: 10.1016/j.asoc.2020.106983.
- [10] D.-H. Lee, Y.-R. Kim, H.-J. Kim, S.-M. Park, and Y.-J. Yang, "Fake news detection using deep learning," *Journal of Information Processing Systems*, vol. 15, no. 5, 2019.
- [11] B. Rusyn *et al.*, "Rethinking deep CNN training: A novel approach for quality-aware dataset optimization," *IEEE Access*, vol. 12, pp. 137427–137438, 2024, doi: 10.1109/ACCESS.2024.3414651.
- [12] P. P. Kundu, T. Truong-Huu, L. Chen, L. Zhou, and S. G. Teo, "Detection and classification of botnet traffic using deep learning with model explanation," *IEEE Transactions on Dependable and Secure Computing*, pp. 1–15, 2024, doi: 10.1109/TDSC.2022.3183361.
- [13] Y. Liang, M. Li, C. Jiang, and G. Liu, "CEModule: A computation efficient module for lightweight convolutional neural networks," *IEEE Transactions on Neural Networks and Learning Systems*, vol. 34, no. 9, pp. 6069–6080, Sep. 2023, doi: 10.1109/TNNLS.2021.3133127.
- [14] Y. Jia, F. Lin, and Y. Sun, "A novel federated learning aggregation algorithm for AIoT intrusion detection," *IET Communications*, vol. 18, no. 7, pp. 429–436, Apr. 2024, doi: 10.1049/cmu2.12744.
- [15] U. M. Khaire and R. Dhanalakshmi, "Stability of feature selection algorithm: A review," *Journal of King Saud University - Computer and Information Sciences*, vol. 34, no. 4, pp. 1060–1073, Apr. 2022, doi: 10.1016/j.jksuci.2019.06.012.
- [16] ICTMCG, "ReCOVary: COVID-19 news dataset for misinformation detection," *GitHub*. <https://github.com/ICTMCG/ReCOVary>, (Accessed: Mar. 11, 2026).
- [17] C. Meng, "CoAID: COVID-19 healthcare misinformation dataset," *GitHub*. <https://github.com/cuilimeng/CoAID>, (Accessed: Mar. 11, 2026).
- [18] "HealthStory dataset: Credible vs. non-credible health news articles," *HealthStory*. <https://healthstory.ai/dataset/>, (Accessed: Mar. 11, 2026).
- [19] E. Indrayuni and A. Nurhadi, "Optimizing genetic algorithms for sentiment analysis of apple product reviews using SVM," *Sinkron*, vol. 4, no. 2, pp. 172–178, Apr. 2020, doi: 10.33395/sinkron.v4i2.10549.
- [20] N. C. Dang, M. N. Moreno-García, and F. De la Prieta, "Sentiment analysis based on deep learning: A comparative study," *Electronics*, vol. 9, no. 3, Mar. 2020, doi: 10.3390/electronics9030483.
- [21] S. Kumawat, I. Yadav, N. Pahal, and D. Goel, "Sentiment analysis using language models: A study," in *2021 11th International Conference on Cloud Computing, Data Science & Engineering (Confluence)*, Jan. 2021, pp. 984–988, doi: 10.1109/Confluence51648.2021.9377043.
- [22] R. Xiang, Y. Long, C. M. Wan, J. Gu, Q. Lu, and C.-R. Huang, "Affection driven neural networks for sentiment analysis," in *12th Edition of its Language Resources and Evaluation Conference*, 2020.
- [23] S. Wen and J. Li, "Recurrent convolutional neural network with attention for Twitter and Yelp sentiment classification: ARC model for sentiment classification," in *Proceedings of the 2018 International Conference on Algorithms, Computing and Artificial Intelligence*, Dec. 2018, pp. 1–7, doi: 10.1145/3302425.3302468.
- [24] S. H. Janjua, G. F. Siddiqui, M. A. Sindhu, and U. Rashid, "Multi-level aspect based sentiment classification of Twitter data: using hybrid approach in deep learning," *PeerJ Computer Science*, vol. 7, Apr. 2021, doi: 10.7717/peerj-cs.433.
- [25] K. L. Tan, C. P. Lee, K. S. M. Anbananthen, and K. M. Lim, "RoBERTa-LSTM: A hybrid model for sentiment analysis with transformer and recurrent neural network," *IEEE Access*, vol. 10, pp. 21517–21525, 2022, doi: 10.1109/ACCESS.2022.3152828.

BIOGRAPHIES OF AUTHORS



Mr. Siva Dhievaraj    is a research scholar in the Department of Computer Science at SRM Institute of Science and Technology, Ramapuram Campus. He holds degrees in Physics, Computer Applications, and Computer Science Engineering (B.Sc., MCA, M.Tech., and Ph.D.) from Bharathidasan University and SRM IST. His teaching and research interests span cloud computing, ML, and programming. He has published several research papers in reputed journals, focusing on topics such as fake news detection, spammer identification, cloud data security, and predictive modeling using machine learning. He can be contacted at email: sivad@srmist.edu.in.



Dr. Agusthiyar Ramu    is a Professor and the Head of the Department of Computer Applications at SRM Institute of Science and Technology, Ramapuram Campus. He specializes in data mining (DM) and AI, areas where he has made significant contributions through his research and teaching. He earned his Ph.D. in DM from Anna University in 2017, and his M.Phil. in Computer Science from Pride, Periyar University, Salem, in 2008. He also holds an MCA in Computer Applications from Madurai Kamaraj University, awarded in 2001, and a B.Sc. in Mathematics from Dr. Zakir Hussain College, Ilayangudi, affiliated to Madurai Kamaraj University, awarded in 1998. In his academic career, he has taught various courses including DM and warehousing and software engineering. His research interests lie in DM, AI, and ML, where he continues to explore and innovate. As the Head of the Department, he plays a crucial role in shaping the curriculum and guiding the department towards academic excellence. He can be contacted at email: agusthir@srmist.edu.in.