

A Hybrid Framework for Hallucination Detection & Mitigation in Large Language Models

Vedant Lanjekar¹, Dipannita Mondal²

^{1,2}Department of Artificial Intelligence & Data Science, Dr. D. Y. Patil College of Engineering & Innovation, Pune, India.

¹vedantlanjekar456@gmail.com, ²hod_aids@dypatilef.com

Peer Review Information

Type: Article

Received: 18 April 2026

Revised: 06 May 2026

Accepted: 15 June 2026

Published: 11 July 2026

Abstract

Large Language Models (LLMs) have achieved significant breakthroughs in natural language processing, enabling advanced capabilities in text generation, reasoning, and conversational intelligence. However, a critical limitation persists in the form of *hallucinations*, where models generate factually incorrect or unverifiable information while maintaining high linguistic confidence. This issue poses substantial risks in high-stakes applications such as healthcare decision support, legal advisory systems, and educational platforms, where accuracy and reliability are paramount. Existing approaches primarily focus on model scaling or fine-tuning, yet they often lack robust mechanisms for real-time verification and interpretability. Addressing this challenge requires a systematic framework capable of detecting, quantifying, and mitigating hallucinated outputs while preserving response fluency and contextual relevance.

This paper proposes a hybrid hallucination detection and mitigation framework that integrates Retrieval-Augmented Generation (RAG), semantic similarity analysis, and a multi-factor confidence scoring mechanism. The system leverages external knowledge sources through vector-based retrieval to validate generated responses and employs transformer-based embeddings to measure semantic consistency between model outputs and verified evidence. A classification layer identifies hallucinated content, followed by a mitigation module that refines responses using grounded contextual information. Experimental evaluation on benchmark datasets such as TruthfulQA and FEVER demonstrates a substantial reduction in hallucination rates and a marked improvement in factual accuracy compared to baseline LLMs. The proposed approach enhances the trustworthiness, explainability, and practical applicability of AI systems, providing a scalable solution for deploying reliable language models in real-world environments.

Keywords: Artificial Intelligence; Large Language Models; Hallucination Detection; Retrieval-Augmented Generation; Explainable AI; Natural Language Processing

How to Cite This Article

Lanjekar, V., & Mondal, D. (2026). A hybrid framework for hallucination detection & mitigation in large language models. *Multidisciplinary Journal of Research in Engineering and Technology*, 13(2s), 396–408.

Introduction

Artificial Intelligence (AI) has undergone a paradigm shift with the emergence of Large Language Models (LLMs), which are primarily built upon transformer-based architectures. Models such as Generative Pre-trained Transformers (GPT) have demonstrated unprecedented capabilities in natural language understanding, contextual reasoning, and generative intelligence. These models leverage massive-scale pretraining on diverse corpora, enabling them to capture complex linguistic patterns, semantic relationships, and contextual dependencies. As a result, LLMs have become foundational components in a wide range of applications, including conversational agents, intelligent tutoring systems, automated content generation, and decision-support systems. Their ability to generate coherent, contextually relevant, and human-like responses has significantly enhanced human-computer interaction, making AI systems more intuitive and accessible.

However, despite these advancements, a fundamental challenge persists in the form of hallucination, wherein LLMs generate outputs that are syntactically fluent but factually incorrect, misleading, or entirely fabricated. Unlike traditional errors, hallucinations are particularly problematic because they are often presented with high confidence and linguistic plausibility, making them difficult for users to identify. This phenomenon undermines the reliability and trustworthiness of AI systems, especially in scenarios where factual correctness is critical. The root cause of hallucination lies in the inherent nature of probabilistic language modeling, where responses are generated based on learned statistical correlations rather than explicit verification against authoritative knowledge sources.

Hallucinations in LLMs arise due to multiple underlying factors, which can be broadly categorized as follows:

- **Lack of Grounding in Factual Knowledge:** LLMs do not possess an intrinsic mechanism to verify the factual correctness of their outputs. Instead, they rely on patterns learned during training, which may include outdated, incomplete, or incorrect information. The absence of real-time access to verified knowledge bases further exacerbates this limitation.
- **Probabilistic Text Generation Mechanisms:** LLMs operate by predicting the most probable next token in a sequence based on prior context. This probabilistic approach prioritizes fluency and coherence over factual accuracy, leading to the generation of plausible but incorrect statements, particularly when the model encounters ambiguous or low-confidence scenarios.
- **Training Data Limitations and Biases:** The quality and diversity of training data significantly influence model behavior. If the training data contains inconsistencies, biases, or insufficient coverage of specific domains, the model may generalize incorrectly, resulting in hallucinated outputs. Additionally, the static nature of training data prevents models from adapting to newly emerging information.

The impact of hallucination becomes critically significant in domains that demand high levels of factual precision, accountability, and reliability. In such contexts, even minor inaccuracies can lead to severe consequences. Key application areas affected include:

- **Medical Diagnosis Systems:** AI-driven healthcare systems rely on accurate information for disease prediction, treatment recommendations, and clinical decision support. Hallucinated outputs in this domain can lead to misdiagnosis, inappropriate treatments, and potential risks to patient safety.
- **Legal Advisory Platforms:** Legal AI systems are expected to provide precise interpretations of laws, regulations, and case precedents. Fabricated or incorrect legal information can result in flawed legal advice, financial losses, or ethical violations.
- **Educational AI Tutors:** AI-based learning platforms are increasingly used for knowledge dissemination and personalized education. Hallucinated explanations or incorrect facts can mislead learners, negatively impacting the quality of education and knowledge acquisition.

Despite the growing recognition of this issue, existing research efforts have predominantly focused on improving model performance through techniques such as large-scale pretraining, fine-tuning, and reinforcement learning from human feedback (RLHF). While these approaches enhance linguistic quality and alignment with human preferences, they do not provide a robust mechanism for real-time verification of generated outputs. Consequently, hallucination remains an unresolved challenge in practical deployments of LLMs.

Recent advancements, such as Retrieval-Augmented Generation (RAG), attempt to address this limitation by incorporating external knowledge sources during the generation process. Although these methods improve factual grounding to some extent, they often lack explicit mechanisms for detecting hallucinations, quantifying their severity, and providing interpretable insights into model behavior. Furthermore, many existing systems do not integrate detection and mitigation into a unified framework, limiting their effectiveness in real-world applications.

Literature Survey

The issue of hallucination in Large Language Models (LLMs) has gained significant attention with the advancement of transformer-based architectures such as GPT and BERT. These models rely on probabilistic token prediction mechanisms, generating text based on learned statistical patterns rather than verified knowledge. While this enables fluent and contextually coherent responses, it often results in outputs

that are plausible but factually incorrect. Consequently, hallucination has emerged as a major limitation affecting the reliability of LLMs, particularly in knowledge-intensive applications.

To evaluate hallucination, several benchmark datasets have been developed. The TruthfulQA benchmark assesses model truthfulness under adversarial questioning, revealing that even advanced models tend to produce confidently incorrect answers. Similarly, the FEVER (Fact Extraction and Verification) dataset emphasizes evidence-based verification by requiring models to validate claims using retrieved information. Although these benchmarks highlight the importance of factual grounding, they are primarily designed for evaluation or classification tasks rather than real-time generative systems.

To improve factual accuracy, Retrieval-Augmented Generation (RAG) has been introduced as a promising approach that integrates external knowledge retrieval with language generation. RAG-based models enhance contextual relevance and grounding by incorporating retrieved documents during response generation. However, these systems mainly focus on improving generation quality and lack explicit mechanisms for detecting hallucinations, quantifying inconsistencies, or providing interpretability. Additionally, explainable AI techniques such as attention visualization and saliency mapping offer insights into model behavior but do not directly address hallucination detection or correction.

Despite these advancements, several limitations persist. Current approaches lack real-time hallucination detection, do not provide reliable confidence scoring, and offer limited explainability. Moreover, existing methods are often fragmented, addressing retrieval, verification, or interpretability independently rather than as part of a unified system. To overcome these challenges, this paper proposes a multi-layer hybrid framework that integrates semantic similarity analysis, retrieval-based verification, confidence scoring, and explainability mechanisms, thereby enhancing the reliability and practical applicability of LLMs.

Problem Statement

Large Language Models (LLMs) have demonstrated remarkable proficiency in generating fluent, contextually coherent, and human-like text by leveraging deep neural architectures and large-scale pretraining. However, despite their linguistic capabilities, these models fundamentally operate on probabilistic token prediction mechanisms, where each output token is generated based on the likelihood of occurrence conditioned on prior context. This probabilistic nature implies that LLMs prioritize *statistical plausibility* over *factual correctness*, as they do not possess an inherent mechanism to verify the truthfulness of the generated content against authoritative knowledge sources.

As a consequence, LLMs frequently produce hallucinated outputs, which are responses that may appear syntactically correct and contextually relevant but are factually inaccurate, unverifiable, or entirely fabricated. These hallucinations are particularly problematic because they are often presented with high confidence, making them difficult for end-users to detect. The issue is further exacerbated in scenarios involving ambiguous queries, incomplete contextual information, or domain-specific knowledge gaps, where the model attempts to “fill in” missing information based on learned patterns rather than verified facts.

The absence of robust hallucination handling mechanisms significantly limits the deployment of LLMs in critical, real-world applications. In domains such as healthcare, legal advisory systems, financial analysis, and academic research, the consequences of incorrect information can be severe, leading to misinformed decisions, ethical concerns, and potential safety risks. Therefore, ensuring the factual reliability, transparency, and trustworthiness of AI-generated outputs has become a fundamental challenge in modern artificial intelligence research.

Existing approaches primarily focus on improving model performance through techniques such as large-scale training, fine-tuning, and reinforcement learning from human feedback (RLHF). While these methods enhance fluency and alignment with user intent, they do not provide a systematic solution for detecting, quantifying, and correcting hallucinations in real time. Furthermore, most current systems lack mechanisms for explicitly validating generated responses against external knowledge sources or providing interpretable insights into their reliability.

To address these limitations, there is a critical need for a comprehensive framework that incorporates the following capabilities:

Key Challenges and Requirements

1. **Real-Time Hallucination Detection:** There is a necessity for mechanisms that can dynamically identify hallucinated content during or immediately after the generation process. Such systems should be capable of analyzing the generated response in real time and detecting inconsistencies without introducing significant latency.
2. **Reliable Fact Verification Mechanisms:** The system must integrate external, authoritative knowledge sources to validate generated outputs. This requires efficient retrieval techniques, such as vector-based semantic search, to ensure that responses are grounded in accurate and up-to-date information.

3. **Quantifiable Confidence Scoring:** A robust framework should provide a measurable indicator of the reliability of generated responses. Confidence scoring mechanisms must combine multiple factors, including semantic similarity, source credibility, and model certainty, to produce an interpretable reliability metric for end-users.
4. **Automated Correction and Mitigation of Incorrect Responses:** Beyond detection, the system should be capable of refining or regenerating hallucinated outputs using verified contextual information. This requires the integration of mitigation strategies that ensure corrected responses maintain both factual accuracy and contextual coherence.

In light of these challenges, this research aims to develop a hybrid hallucination detection and mitigation framework that bridges the gap between generative capabilities and factual reliability. By combining semantic verification, retrieval-based grounding, confidence estimation, and automated correction within a unified architecture, the proposed approach seeks to enhance the trustworthiness and practical applicability of Large Language Models in real-world deployments.

Proposed Methodology

The proposed system introduces a multi-layer hybrid architecture designed to address the limitations of Large Language Models (LLMs) in handling hallucinated outputs. The methodology integrates retrieval-based verification, semantic consistency analysis, classification mechanisms, and response mitigation strategies into a unified framework. Unlike traditional approaches that rely solely on model training improvements, the proposed system emphasizes post-generation validation and correction, ensuring that generated responses are grounded in verified knowledge. The architecture operates sequentially, beginning with response generation by the LLM, followed by retrieval of relevant factual information from external knowledge sources using vector similarity search techniques. This retrieved information serves as a reference for evaluating the factual consistency of the generated output.

To ensure robust hallucination detection, the system employs transformer-based embedding models to compute semantic similarity between generated responses and retrieved evidence. A classification mechanism then categorizes outputs based on predefined similarity thresholds, enabling the identification of reliable, partially hallucinated, and fully hallucinated responses. Furthermore, a confidence scoring mechanism aggregates multiple factors, including semantic similarity, retrieval relevance, and model probability, to quantify the trustworthiness of the output. In cases where hallucination is detected, a mitigation module refines the response by incorporating verified contextual information, thereby improving factual accuracy while preserving linguistic coherence. This comprehensive methodology ensures real-time detection, interpretability, and correction of hallucinated outputs, making the system suitable for deployment in high-stakes applications.

Retrieval-Augmented Verification Layer

Component	Description
Objective	To retrieve relevant and factually accurate documents for validating LLM outputs
Technique Used	Vector similarity search using dense embeddings
Indexing Method	FAISS (Facebook AI Similarity Search) for efficient nearest-neighbor retrieval
Data Sources	Wikipedia, domain-specific datasets, curated knowledge bases
Functionality	Maps user query and generated response into vector space and retrieves top-k relevant documents
Output	Set of contextually relevant documents for verification

Semantic Consistency Analyzer

Component	Description
Objective	To evaluate the semantic alignment between generated responses and retrieved factual content
Model Used	BERT / Sentence Transformers
Technique	Embedding generation and cosine similarity computation
Input	LLM-generated response + retrieved documents
Output	Semantic similarity score indicating factual consistency
Role	Acts as the core verification engine for hallucination detection

Mathematical Representation

Similarity Score:

$$S = \frac{A \cdot B}{\|A\| \cdot \|B\|}$$

Where:

- A = Embedding vector of generated response
- B = Embedding vector of retrieved evidence

Hallucination Detection Model

Component	Description
Objective	To classify generated responses based on factual reliability
Input	Semantic similarity score (S)
Method	Threshold-based classification
Categories	Reliable, Partially Hallucinated, Fully Hallucinated
Output	Classification label indicating hallucination level
Advantage	Simple, interpretable, and computationally efficient

Classification Criteria

Similarity Score (S)	Classification
$S > 0.8$	Reliable
$0.5 < S \leq 0.8$	Partially Hallucinated
$S \leq 0.5$	Hallucinated

Mitigation Module

Component	Description
Objective	To correct hallucinated outputs and improve factual accuracy
Approach	Context-aware response regeneration
Input	Detected hallucinated response + retrieved factual data
Techniques	Prompt refinement, context injection, response rewriting
Functionality	Replaces inconsistent segments with verified information
Output	Refined and factually grounded response
Benefit	Enhances reliability without sacrificing fluency

Confidence Scoring Mechanism

Component	Description
Objective	To quantify the reliability of generated responses
Input Factors	Semantic similarity, retrieval confidence, model probability
Method	Weighted linear combination
Output	Numerical confidence score (C)
Interpretation	Higher score indicates higher reliability

Mathematical Representation

$C = \alpha S + \beta R + \gamma P$ Where:

- (S) = Semantic similarity score
- (R) = Retrieval confidence score
- (P) = Model probability (likelihood of generated response)

- (α, β, γ) = Weighting factors (tunable parameters)

System Architecture

The proposed system architecture is designed as a multi-stage processing pipeline that ensures accurate detection and mitigation of hallucinations in Large Language Model (LLM) outputs. Each component in the pipeline performs a specific function, contributing to a cohesive and efficient framework for generating reliable and verifiable responses. The architecture emphasizes modularity, scalability, and real-time processing, enabling seamless integration into practical AI applications.

System Pipeline Overview

The complete system workflow consists of the following stages:

1. User Query Input
 - Accepts natural language queries from the user interface
 - Supports both domain-specific and open-domain queries
 - Performs initial preprocessing such as:
 - Tokenization
 - Noise removal
 - Query normalization
 - Ensures compatibility with downstream LLM processing
2. LLM Response Generation
 - Generates an initial response using a pre-trained Large Language Model
 - Utilizes transformer-based architectures (e.g., GPT, LLaMA)
 - Produces contextually relevant and linguistically coherent output
 - Operates without explicit factual verification (baseline generation)
 - Outputs both:
 - Generated text response
 - Token-level probability scores (if available)
3. Knowledge Retrieval Module
 - Retrieves relevant factual information from external knowledge sources
 - Uses vector similarity search for efficient retrieval
 - Converts query and generated response into embeddings
 - Searches top-k similar documents using FAISS indexing
 - Data sources include:
 - Wikipedia
 - Domain-specific datasets
 - Structured knowledge bases
 - Outputs:
 - Ranked list of relevant documents
 - Retrieval confidence scores
4. Semantic Comparison Engine
 - Evaluates semantic similarity between:
 - LLM-generated response
 - Retrieved factual evidence
 - Uses transformer-based embedding models (Sentence-BERT)
 - Computes cosine similarity scores
 - Identifies semantic alignment or divergence
 - Outputs:
 - Similarity score (S)
 - Feature vectors for further analysis
5. Hallucination Detection Unit
 - Classifies responses based on semantic similarity thresholds

- Determines the degree of hallucination:
 - Reliable
 - Partially hallucinated
 - Fully hallucinated
- Applies rule-based or lightweight ML classification logic
- Flags segments of text that deviate from factual evidence
- Outputs:
 - Hallucination label
 - Detection confidence
- 6. Confidence Scoring Layer
 - Computes a composite reliability score for the generated response
 - Integrates multiple factors:
 - Semantic similarity score (S)
 - Retrieval confidence (R)
 - Model probability (P)
 - Uses weighted scoring function to produce final confidence value
 - Provides interpretable reliability metrics to end-users
 - Outputs:
 - Confidence score (C)
 - Reliability classification
- 7. Mitigation & Response Refinement
 - Activates when hallucination is detected
 - Refines response using retrieved factual context
 - Techniques include:
 - Context injection into prompt
 - Regeneration using grounded data
 - Replacement of incorrect segments
 - Ensures:
 - Improved factual accuracy
 - Preservation of linguistic coherence
 - Outputs:
 - Corrected and verified response
- 8. Explainability Interface
 - Provides transparency into system decisions
 - Highlights:
 - Differences between generated and retrieved content
 - Confidence scores and reasoning
 - Enables user interpretability and trust
 - Displays:
 - Evidence sources
 - Highlighted inconsistencies
 - Supports visual or textual explanation formats

Implementation Details

The implementation of the proposed system is carried out using a combination of modern machine learning frameworks, natural language processing libraries, and efficient retrieval systems. The design focuses on scalability, modularity, and real-time performance, ensuring that the system can be deployed in both research and production environments.

Programming Environment

- Language Used: Python
- Chosen for its:

- Extensive AI/ML ecosystem
- Rich NLP libraries
- Ease of integration with deep learning frameworks
- Supports rapid prototyping and scalable deployment

Core Libraries and Tools

- HuggingFace Transformers
 - Used for loading and fine-tuning pre-trained LLMs
 - Provides access to models like GPT, BERT, and LLaMA
 - Enables tokenization, inference, and embedding generation
- FAISS (Facebook AI Similarity Search)
 - Used for efficient vector-based document retrieval
 - Supports large-scale similarity search
 - Optimized for high-dimensional embedding spaces
- LangChain
 - Facilitates integration of Retrieval-Augmented Generation (RAG) pipelines
 - Manages interaction between LLMs and external knowledge sources
 - Enables chaining of multiple components in the pipeline
- PyTorch / TensorFlow
 - Deep learning frameworks for model execution and training
 - Supports GPU acceleration for faster computation
 - Enables customization of neural network components
- Streamlit
 - Used for developing an interactive user interface
 - Allows real-time query input and response visualization
 - Displays confidence scores and explainability outputs

Models Utilized

- Large Language Models (LLMs)
 - LLaMA / GPT-based architectures
 - Responsible for initial response generation
 - Pre-trained on large-scale text corpora
- Embedding Models
 - Sentence-BERT (SBERT)
 - Used for generating dense vector representations
 - Enables semantic similarity computation

System Integration Workflow

- User query is processed and passed to LLM
- Generated response is embedded and sent to retrieval module
- FAISS retrieves relevant documents
- SBERT computes similarity between response and evidence
- Detection module classifies hallucination level
- Confidence score is computed
- Mitigation module refines response if necessary
- Final output is displayed with explanation

Deployment Considerations

- Supports both local and cloud-based deployment
- Can be integrated with APIs for real-world applications
- Scalable architecture for handling large query volumes
- Optimized for low latency using efficient retrieval techniques

This architecture and implementation ensure a robust, scalable, and interpretable system capable of addressing hallucination challenges in modern AI systems.

Datasets Used

The effectiveness of the proposed hallucination detection and mitigation framework is evaluated using multiple **benchmark datasets** that are widely recognized in the domains of natural language processing, fact verification, and question answering. These datasets are carefully selected to ensure diversity in query types, complexity, and factual grounding requirements, thereby enabling a comprehensive evaluation of the system's performance.

Dataset Overview and Characteristics

1. TruthfulQA (Hallucination Benchmark)

- Designed specifically to evaluate the *truthfulness* of language models
- Contains adversarial questions that exploit common misconceptions
- Measures the tendency of models to generate false but plausible answers
- Particularly useful for assessing hallucination behavior in generative systems
- Enables evaluation under challenging, real-world scenarios where factual correctness is critical

2. FEVER (Fact Extraction and Verification)

- Focuses on verifying textual claims using evidence retrieved from Wikipedia
- Consists of labeled claims categorized as *Supported*, *Refuted*, or *Not Enough Information*
- Requires multi-step reasoning involving:
 - Claim understanding
 - Evidence retrieval
 - Logical verification
- Serves as a benchmark for evaluating retrieval-based verification mechanisms in the proposed system

3. HotpotQA

- A multi-hop question answering dataset requiring reasoning across multiple documents
- Encourages models to perform *complex reasoning* rather than single-step retrieval
- Includes supporting facts for explainability
- Useful for evaluating the system's ability to handle:
 - Context aggregation
 - Cross-document verification
- Multi-step inference

4. Natural Questions (NQ)

- Real-world dataset derived from user search queries
- Contains both short and long answer formats
- Reflects practical information-seeking behavior of users
- Evaluates system performance in realistic, open-domain scenarios
- Tests robustness against diverse and ambiguous queries

5. Dataset Utilization Strategy

- Queries are passed through the LLM to generate initial responses
- Retrieved documents are sourced from dataset-aligned knowledge bases
- Ground truth answers are used to evaluate factual correctness
- Hallucination detection performance is measured by comparing generated outputs with verified answers

Results And Discussion

The proposed hybrid framework was evaluated against baseline Large Language Model (LLM) outputs to assess its effectiveness in detecting and mitigating hallucinations. The evaluation focuses on key performance indicators such as factual accuracy, semantic consistency, and hallucination reduction. Experimental results demonstrate that integrating retrieval-based verification and semantic similarity analysis significantly enhances the reliability of generated responses.

Observations and Analysis

1. Significant Reduction in Hallucinated Responses
 - The proposed system effectively identifies and mitigates hallucinated outputs
 - Retrieval-based grounding ensures alignment with verified knowledge
 - Reduction in hallucination rate indicates improved factual reliability
2. Improved Factual Consistency
 - Semantic similarity analysis ensures generated responses closely match retrieved evidence
 - Multi-layer verification reduces inconsistencies in generated content
 - Enhanced alignment between model output and ground truth data
3. Higher Confidence Alignment with Correctness
 - Confidence scoring mechanism correlates strongly with actual response accuracy
 - Reliable responses exhibit higher confidence scores, improving interpretability
 - Enables users to assess trustworthiness of outputs effectively

Quantitative Performance Evaluation

The performance of the proposed system is compared with baseline LLM outputs using standard evaluation metrics:

Metric	Baseline LLM	Proposed System	Improvement
Accuracy	72%	89%	+17%
F1 Score	0.68	0.85	+0.17
Hallucination Rate	28%	11%	-17%

Detailed Interpretation of Results

1. Accuracy Improvement (+17%)
 - Indicates a substantial increase in correct responses
 - Demonstrates effectiveness of retrieval-based verification
 - Highlights improved grounding in factual knowledge
2. F1 Score Enhancement
 - Reflects better balance between precision and recall
 - Suggests improved detection of both correct and incorrect outputs
 - Confirms robustness of the hallucination detection mechanism
3. Reduction in Hallucination Rate (-17%)
 - Significant decrease in incorrect or fabricated responses
 - Validates the effectiveness of the mitigation module
 - Indicates improved trustworthiness of the system

Comparative Analysis

1. Baseline LLMs rely solely on probabilistic generation, leading to higher hallucination rates
2. The proposed system introduces:
 - External knowledge grounding
 - Semantic verification
 - Confidence-based evaluation
3. These additions collectively enhance:
 - Reliability
 - Interpretability
 - Practical applicability

Discussion

The results clearly demonstrate that the integration of retrieval-augmented verification and semantic consistency analysis significantly improves the performance of Large Language Models in terms of factual accuracy and reliability. The proposed framework not only reduces hallucinations but also provides a structured mechanism for evaluating and correcting generated responses. Furthermore, the

inclusion of a confidence scoring mechanism adds an additional layer of interpretability, enabling users to make informed decisions based on the reliability of AI outputs.

However, it is important to note that the system introduces additional computational overhead due to retrieval and similarity computations. Despite this, the trade-off between computational cost and improved accuracy is justified, particularly in high-stakes applications where correctness is critical.

Advantages

The proposed hybrid framework for hallucination detection and mitigation offers several significant advantages that enhance the reliability, interpretability, and applicability of Large Language Models in real-world scenarios:

1. Real-Time Hallucination Detection
 - Enables dynamic identification of hallucinated content during or immediately after response generation
 - Integrates semantic verification and retrieval mechanisms without requiring offline post-processing
 - Supports real-time decision-making in interactive AI systems such as chatbots and virtual assistants
 - Reduces the risk of propagating incorrect information to end-users
 - Ensures continuous monitoring of model outputs in production environments
2. Improved Trustworthiness and Reliability
 - Grounds generated responses in verified external knowledge sources
 - Reduces the likelihood of misinformation and fabricated outputs
 - Enhances user confidence in AI-generated content, especially in critical domains
 - Provides consistent and accurate responses across diverse query types
 - Supports deployment in high-stakes applications requiring dependable outputs
3. Explainable AI Outputs
 - Incorporates explainability mechanisms that highlight inconsistencies between generated responses and factual evidence
 - Provides transparency into how and why a response is classified as hallucinated or reliable
 - Displays confidence scores and supporting evidence to end-users
 - Facilitates debugging, auditing, and validation of AI systems
 - Aligns with regulatory requirements for interpretable and accountable AI
4. Domain Adaptability and Scalability
 - Can be extended to domain-specific applications by integrating specialized knowledge bases
 - Supports multiple domains such as healthcare, finance, legal, and education
 - Modular architecture allows easy customization and scalability
 - Compatible with different LLM architectures and datasets
 - Enables continuous improvement through updated knowledge sources and retraining

Limitations

Despite its advantages, the proposed system has certain limitations that must be considered for practical deployment and future improvements:

1. Increased Computational Overhead
 - Additional processing required for:
 - Embedding generation
 - Vector similarity search
 - Semantic comparison
 - Higher resource consumption compared to standalone LLM systems
 - Requires GPU/accelerated hardware for optimal performance
 - May impact scalability in large-scale, high-throughput environments
2. Dependency on External Knowledge Sources
 - System performance heavily relies on the quality and completeness of retrieved data
 - Inaccurate or outdated knowledge sources may affect verification accuracy
 - Requires continuous updating and maintenance of knowledge bases
 - Limited effectiveness in domains with scarce or unstructured data

3. Latency in Retrieval and Verification Process

- Additional time required for:
 - Document retrieval
 - Similarity computation
 - Response refinement
- May introduce delays in real-time applications
- Trade-off between accuracy and response speed
- Requires optimization techniques such as caching and indexing for faster retrieval

4. Threshold Sensitivity and Parameter Tuning

- Performance depends on predefined similarity thresholds and weighting factors
- Improper tuning may lead to false positives or false negatives in detection
- Requires experimentation and validation for different datasets and domains

5. Limited Handling of Complex Reasoning Errors

- System primarily focuses on factual inconsistencies
- May not fully capture logical reasoning errors or nuanced contextual hallucinations
- Requires further enhancement for deep reasoning-based verification

Applications

The proposed framework has wide-ranging applications across multiple domains where accuracy, reliability, and trustworthiness are critical:

1. Healthcare AI Systems

- Supports clinical decision-making by ensuring factually accurate medical information
- Reduces risk of misdiagnosis caused by hallucinated outputs
- Enhances reliability of AI-driven diagnostic and recommendation systems
- Can be integrated with electronic health record (EHR) systems
- Improves patient safety and trust in AI-assisted healthcare

2. Legal Advisory Platforms

- Provides accurate legal interpretations and case law references
- Prevents dissemination of incorrect or fabricated legal information
- Assists lawyers and legal professionals in research and documentation
- Enhances compliance with regulatory and ethical standards
- Enables reliable AI-based legal consultation systems

3. AI-Based Education Systems

- Ensures correctness of educational content delivered by AI tutors
- Supports personalized learning with accurate explanations and solutions
- Prevents misinformation in academic and learning environments
- Enhances student trust in AI-driven educational tools
- Useful in e-learning platforms and intelligent tutoring systems

4. Customer Support Automation

- Improves accuracy of automated responses in chatbots and virtual assistants
- Reduces customer dissatisfaction caused by incorrect information
- Enhances user experience through reliable and consistent responses
- Supports scalable customer service operations across industries
- Enables deployment in domains such as banking, e-commerce, and telecom

5. Financial and Business Intelligence Systems

- Ensures accuracy in financial predictions, reports, and recommendations
- Reduces risk of misinformation in critical financial decision-making
- Supports data-driven business analytics with reliable outputs

6. Research and Knowledge Management Systems

- Assists researchers by providing factually accurate summaries and insights
- Prevents propagation of incorrect scientific or technical information

- Enhances reliability of AI-assisted research tools

Conclusion

This paper presented a comprehensive hybrid framework for hallucination detection and mitigation in Large Language Models (LLMs), addressing one of the most critical challenges in modern artificial intelligence systems. By integrating retrieval-augmented verification, semantic similarity analysis, and a multi-factor confidence scoring mechanism, the proposed approach effectively bridges the gap between generative fluency and factual reliability. Unlike conventional methods that rely solely on model training improvements, the proposed system introduces a post-generation validation pipeline that ensures responses are grounded in verified external knowledge. Experimental evaluation on benchmark datasets demonstrates that the framework significantly reduces hallucination rates while improving overall accuracy and consistency. Furthermore, the incorporation of explainability mechanisms enhances transparency, enabling users to interpret and trust AI-generated outputs more effectively. The results highlight the potential of combining retrieval-based grounding with semantic verification to create more robust, reliable, and deployable AI systems.

Despite its effectiveness, the proposed framework opens several avenues for future research and improvement. One key direction involves optimizing the system for real-time performance, particularly by reducing latency associated with knowledge retrieval and similarity computations through advanced indexing, caching, and model optimization techniques. Additionally, future work can explore the integration of multimodal verification mechanisms, incorporating visual, audio, or structured data sources to further enhance factual grounding and expand the applicability of the system. Another promising direction is the development of adaptive learning mechanisms that dynamically update knowledge bases and refine confidence scoring models based on user feedback and evolving data. By addressing these challenges, the proposed framework can be extended into a scalable, high-performance solution for next-generation AI applications, ultimately contributing to the development of trustworthy, explainable, and human-aligned artificial intelligence systems.

References

1. D. P. Kingma and M. Welling, “Auto-encoding variational Bayes,” *arXiv preprint arXiv:1312.6114*, 2013. [Online]. Available: <https://arxiv.org/abs/1312.6114>
2. T. Brown, B. Mann, N. Ryder, M. Subbiah, J. Kaplan, P. Dhariwal, A. Neelakantan, P. Shyam, G. Sastry, A. Askell, *et al.*, “Language models are few-shot learners,” in *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 33, pp. 1877–1901, 2020.
3. P. Lewis, E. Perez, A. Piktus, F. Petroni, V. Karpukhin, N. Goyal, H. Küttler, M. Lewis, W. Yih, T. Rocktäschel, *et al.*, “Retrieval-augmented generation for knowledge-intensive NLP tasks,” in *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 33, pp. 9459–9474, 2020.
4. S. Lin, J. Hilton, and O. Evans, “TruthfulQA: Measuring how models mimic human falsehoods,” in *Proceedings of the Annual Meeting of the Association for Computational Linguistics (ACL)*, 2021, pp. 3214–3252.
5. Thorne, A. Vlachos, C. Christodoulopoulos, and A. Mittal, “FEVER: A large-scale dataset for fact extraction and verification,” in *Proceedings of the Conference of the North American Chapter of the Association for Computational Linguistics (NAACL-HLT)*, 2018, pp. 809–819.
6. Hugging Face Inc., “Hugging Face Transformers: State-of-the-art natural language processing,” 2024. [Online]. Available: <https://huggingface.co/docs>