



SvedbergOpen
DISSEMINATION OF KNOWLEDGE

International Journal of Artificial Intelligence and Machine Learning

Publisher's Home Page: <https://www.svedbergopen.com/>



Research Paper

Open Access

Generative AI Technologies for Intelligent Decision-Making Applications

Dr. R. Saravana Prabhu¹, Dr. S. P. Manikandan^{2*}, Aayush Poudel³, Debesh Maity⁴, Dr. Sharmista A⁵, Mr. Srinivasan S⁶, Dr. M. Sathish Kumar⁷

¹Assistant Professor, Department of BCA, The American College, Madurai, Tamil Nadu, India. E-mail: r.saravanaprabu@gmail.com, saravananprabhu@americancollege.edu.in

²Professor & Deputy Director, School of Engineering and Technology, CMR University, Bengaluru, Karnataka, India. E-mail: dr.mani1973@gmail.com, ORCID: <https://orcid.org/0000-0002-7564-3671>

³Assistant Professor, School of Engineering and Technology, CMR University, Bengaluru, Karnataka, India. E-mail: Aayushpou@gmail.com, ORCID: <https://orcid.org/0009-0007-3631-6866>

⁴Assistant Professor, School of Engineering and Technology, CMR University, Bengaluru, Karnataka, India. Email ID: debeshmaity123@gmail.com, ORCID: <https://orcid.org/0009-0002-2255-359X>

⁵Assistant Professor and Head, Department of Data Science and Analytics NMSSVN College Nagamalai, Madurai, Tamil Nadu, India. E-mail: sharmista@nmssvnc.edu.in

⁶Assistant Professor Department of Data Science and Analytics NMSSVN College Nagamalai, Madurai, Tamil Nadu, India. E-mail: srinivasan@nmssvnc.edu.in

⁷Assistant Professor, Department of Computer Science, CAPE Arts and Science College, Aralvaimozhi, Kanyakumari, Tamil Nadu, India. E-mail: sathish.friends89@gmail.com

Abstract

Generative Artificial Intelligence (GenAI) has become a transformative technology for intelligent decision-making by enabling knowledge processing, natural language interaction, reasoning support, and evidence-based information synthesis. Unlike traditional analytical systems, GenAI-driven decision-support systems can generate alternatives, summarize complex information, and assist human decision-makers in uncertain environments [1][2]. This research proposes an evidence-grounded GenAI framework integrating Retrieval-Augmented Generation (RAG), validation, decision analytics, confidence assessment, human oversight, and auditability. The proposed approach improves decision reliability by combining AI capabilities with expert supervision and governance controls [3][4]. A design-oriented methodology is adopted to evaluate accuracy, explainability, robustness, fairness, privacy, and governance readiness for future empirical validation.

Keywords: Generative Artificial Intelligence; Large Language Models; Retrieval-Augmented Generation; Intelligent Decision Support; Explainable AI; Human-in-the-Loop AI; Trustworthy AI; AI Governance.

1. Introduction

Decision-making in modern organizations faces increasing complexity due to large information volumes, diverse data sources, dynamic environments, and uncertainty. Traditional decision-support systems provide analytical capabilities but are limited in handling unstructured knowledge and natural interaction [5]. Generative Artificial Intelligence (GenAI) and Large Language Models (LLMs) introduce advanced capabilities in language understanding, information synthesis, and contextual reasoning for intelligent decision support [6][7]. However, challenges such as hallucination, bias, and unsupported outputs affect reliability [8]. Retrieval-Augmented Generation (RAG) improves LLM-based systems by connecting models with external knowledge sources for better evidence grounding [9][10]. This research proposes an evidence-grounded GenAI architecture integrating retrieval, validation, confidence assessment, human oversight, and governance mechanisms for trustworthy decision-support applications.

The major contributions of this research are:

1. **Development of an evidence-grounded GenAI decision-support architecture** integrating retrieval, generation, validation, analytics, and governance components.
2. **Introduction of a human-supervised decision framework** that maintains expert control over critical recommendations.
3. **Proposal of a decision evaluation model** based on accuracy, evidence support, explainability, robustness, fairness, privacy, and governance readiness.
4. **Identification of future empirical validation directions** for testing GenAI decision-support effectiveness across different application domains.

2. Literature Review

Generative Artificial Intelligence (GenAI) research has rapidly developed with advancements in foundation models, Large Language Models (LLMs), Retrieval-Augmented Generation (RAG), Explainable Artificial Intelligence (XAI), and trustworthy AI frameworks. Existing studies demonstrate that GenAI systems can support intelligent decision-making through information synthesis, reasoning assistance, contextual understanding, and interactive communication. However, challenges related to factual reliability, hallucination, transparency, bias, privacy, and governance remain significant concerns. Therefore, the literature is organized into five major themes: foundation models and LLMs, Retrieval-Augmented Generation, GenAI applications in decision-making domains, explainable and responsible AI, and AI governance for trustworthy decision systems.

2.1 Foundation Models and Large Language Models

Foundation models have transformed artificial intelligence by enabling general-purpose learning from large-scale datasets and supporting adaptation across multiple applications. Transformer-based architectures introduced attention mechanisms that improved sequence processing and enabled advanced language models capable of generating text, summarizing information, answering queries, and supporting complex reasoning tasks [1]. Large Language Models (LLMs), including GPT-based models and other transformer architectures, demonstrate strong capabilities in understanding and generating natural language, making them valuable for intelligent decision-support environments [2]. However, standalone LLMs depend mainly on their learned knowledge, creating limitations when handling dynamic information, specialized domains, and high-risk decisions.

One major challenge affecting LLM reliability is hallucination, where models generate convincing but unsupported information due to probabilistic generation mechanisms. Such limitations create risks in healthcare, finance, legal services, and public administration, where inaccurate recommendations may lead to serious consequences [3]. Therefore, current research emphasizes integrating external knowledge sources, verification mechanisms, uncertainty estimation, and human supervision to improve the trustworthiness of LLM-based decision systems.

2.2 Retrieval-Augmented Generation for Reliable AI Decision Support

Retrieval-Augmented Generation (RAG) has emerged as an important approach for improving the reliability of LLM-based systems by combining information retrieval with generative models [4]. Unlike conventional LLMs that rely only on pre-trained knowledge, RAG systems retrieve relevant information from external databases, organizational repositories, and knowledge sources to provide contextual evidence during generation. This approach improves factual grounding, supports knowledge updates, and enables domain-specific applications.

Recent research explores different RAG architectures, including improved retrieval strategies, document ranking methods, query optimization techniques, and modular retrieval-generation pipelines. However, RAG does not completely eliminate reliability issues. Problems such as poor retrieval quality, irrelevant documents, source contamination, and incorrect interpretation of evidence remain important challenges. Therefore, effective GenAI decision-support systems require additional validation mechanisms to evaluate evidence quality, relevance, and confidence before generating final recommendations.

2.3 Generative AI Applications in Decision-Making Domains

GenAI applications are expanding across multiple domains where decision-making requires complex information analysis and expert judgment. Healthcare represents one of the major application areas, where LLMs support clinical documentation, medical knowledge synthesis, patient communication, and decision assistance. However, healthcare applications require strict validation because incorrect recommendations may directly affect patient outcomes [5].

In finance, GenAI is being explored for risk assessment, market analysis, compliance monitoring, investment support, and automated reporting. Due to sensitive financial information and regulatory requirements, these systems require transparency, explainability, and human approval mechanisms. Educational institutions are adopting GenAI for personalized learning, academic advising, and curriculum support while maintaining educator involvement in decision processes.

Businesses use GenAI for strategic planning, customer analysis, knowledge management, and operational intelligence. Across these domains, research indicates that GenAI provides maximum value when used as an augmentation technology rather than a replacement for human expertise. Successful implementations combine AI-generated insights with expert review, organizational policies, and governance controls.

2.4 Explainable and Responsible Artificial Intelligence

Explainability has become essential for AI-based decision systems because users require an understanding of how recommendations are generated. Traditional machine-learning models often operate as black boxes, making interpretation difficult. Explainable Artificial Intelligence (XAI) approaches improve transparency

by providing understandable explanations, feature analysis, reasoning support, and evidence-based justifications [6].

However, explanation quality must be evaluated carefully because an explanation may appear convincing without accurately representing the model's actual reasoning process. Therefore, trustworthy AI research emphasizes explanation fidelity rather than only readability.

Responsible AI frameworks address fairness, accountability, privacy, robustness, and ethical deployment. The NIST Artificial Intelligence Risk Management Framework highlights the importance of continuous monitoring, risk identification, measurement, and governance throughout the AI lifecycle [7]. These principles are essential for GenAI decision-support systems because AI-generated recommendations may influence individuals, organizations, and society.

2.5 AI Governance and Trustworthy Decision Systems

AI governance focuses on developing policies, controls, and operational practices that ensure responsible AI development and deployment. As GenAI systems become integrated into organizational decision processes, governance mechanisms are required to manage risks related to security, accountability, compliance, and model behavior.

Trustworthy AI requires both technical safeguards and organizational processes. Key governance requirements include data quality management, model monitoring, security protection, bias evaluation, human oversight, audit logging, and regulatory compliance. In high-impact environments, human-in-the-loop approaches ensure that AI recommendations are reviewed and approved by qualified decision-makers.

The literature highlights the importance of traceable decision pathways where systems maintain evidence sources, confidence scores, validation results, and decision histories. Such mechanisms improve transparency, accountability, and reliability, forming the foundation for future evidence-grounded Generative AI decision-support architectures.

Table 1: Summary of Existing Research Themes in Generative AI Decision Support

Research Area	Major Contributions	Key Limitations	Research Need
Foundation Models / LLMs	Natural language reasoning, information synthesis, interactive assistance	Hallucination, outdated knowledge	Evidence-grounded generation
Retrieval-Augmented Generation	External knowledge integration and improved factual support	Retrieval errors and source reliability issues	Retrieval validation mechanisms
Domain Applications	Healthcare, finance, education, business decision support	Limited long-term validation	Real-world evaluation studies
Explainable AI	Transparency and user understanding	Explanation fidelity problems	Reliable reasoning interpretation
AI Governance	Risk management and responsible deployment	Operational implementation challenges	Integrated governance frameworks

2.6 Research Gap Identified from Literature

The reviewed literature demonstrates significant progress in GenAI technologies; however, several research gaps remain. Existing studies mainly focus on improving model capability, retrieval performance, or individual application scenarios. Limited research addresses the complete lifecycle of intelligent decision-making, including evidence retrieval, generation, validation, confidence assessment, human supervision, and auditability within a unified architecture.

Therefore, this research proposes an integrated evidence-grounded Generative AI decision-support framework that combines technical capability with governance and human accountability.

3. Research Problem, Objectives and Research Questions

3.1 Research Problem

The rapid advancement of Generative Artificial Intelligence (GenAI) has created new opportunities for intelligent decision-support systems. Although GenAI and Large Language Models (LLMs) provide strong capabilities in information generation, reasoning, and communication, challenges remain in reliability, transparency, evidence validation, and governance. Existing AI systems often lack contextual understanding and natural interaction. This research addresses the need for an integrated framework combining retrieval, validation, explainability, and human supervision for trustworthy decision-making.

- Evidence-based information retrieval
- Generative reasoning capability
- Decision analytics

- Confidence assessment
- Human supervision
- Governance and audit mechanisms

3.2 Research Objectives

The primary objective of this research is to propose an evidence-grounded Generative AI framework for intelligent decision-making by integrating generation, validation, analytics, and governance mechanisms. The specific objectives are:

Objective 1: Analyze Existing GenAI Technologies

To study current GenAI technologies, including Large Language Models (LLMs), Retrieval-Augmented Generation (RAG), embedding models, and AI agents, and evaluate their capabilities and limitations for decision-support applications.

Objective 2: Identify Reliability and Governance Challenges

To analyze major challenges in GenAI adoption, including hallucination, bias, lack of explainability, privacy risks, security issues, and accountability limitations.

Objective 3: Develop an Evidence-Grounded Decision Architecture

To design a framework integrating data management, evidence retrieval, GenAI processing, decision analytics, human oversight, and auditable outputs for improved reliability and traceability.

Objective 4: Design a Decision-Support Algorithm

To develop a systematic workflow involving query analysis, context extraction, evidence retrieval, validation, alternative generation, decision scoring, confidence assessment, and human approval.

Objective 5: Define Evaluation Criteria

To establish evaluation measures based on accuracy, evidence support, efficiency, explainability, robustness, fairness, privacy, and governance readiness.

Objective 6: Identify Future Research Directions

To provide guidance for future studies involving benchmark development, domain-specific evaluation, human-AI collaboration, and responsible GenAI deployment.

3.3 Research Questions

Based on the identified research gap and objectives, this study addresses the following research questions:

RQ1: What Generative AI technologies are most suitable for developing intelligent decision-support applications?

This question examines the role of LLMs, RAG systems, AI agents, and related technologies in supporting complex decision environments.

RQ2: What are the major challenges that limit the reliability of GenAI-based decision-support systems?

This question investigates issues related to hallucination, bias, transparency, privacy, security, and governance.

RQ3: How can evidence retrieval and validation mechanisms improve the reliability of GenAI-generated decisions?

This question explores how external knowledge sources and verification processes can reduce unsupported AI outputs.

RQ4: How can human oversight be integrated into GenAI decision systems without reducing operational efficiency?

This question focuses on human-in-the-loop approaches that balance automation with expert control.

RQ5: What evaluation framework can measure the effectiveness and trustworthiness of GenAI decision-support architectures?

This question examines appropriate metrics for assessing AI-assisted decisions beyond traditional accuracy measurements.

3.4 Research Contributions

The major contributions of this study are summarized as follows:

Contribution 1: Integrated Evidence-Grounded Architecture

This research proposes a unified architecture that combines retrieval, generation, validation, analytics, and governance components within a single decision-support framework.

Contribution 2: Human-Centered Decision Framework

Unlike fully autonomous AI approaches, the proposed framework emphasizes human-AI collaboration by introducing expert review and approval mechanisms for high-risk decisions.

Contribution 3: Confidence and Risk-Aware Decision Processing

The proposed approach introduces confidence assessment and risk-based escalation mechanisms to determine when AI recommendations require additional validation.

Contribution 4: Comprehensive Evaluation Model

The study proposes a multi-dimensional evaluation framework that considers technical performance, evidence quality, explainability, ethical considerations, and governance requirements.

3.5 Research Scope

This research focuses on conceptual development of a trustworthy Generative AI decision-support framework. The study does not claim experimental validation or direct performance superiority. Instead, it establishes an architectural foundation for future empirical evaluation using domain-specific datasets and expert assessments.

The proposed framework is applicable to multiple decision-support environments, including:

- Healthcare decision assistance
- Financial risk analysis
- Educational planning
- Business intelligence
- Public service management

The research emphasizes that GenAI should function as an intelligent augmentation technology that improves human decision-making rather than replacing human judgment.

4. Research Methodology and Proposed Conceptual Framework Design

4.1 Research Methodology

This study adopts a **design-oriented conceptual research methodology** to develop an evidence-grounded Generative AI framework for intelligent decision-making applications. The methodology focuses on analysing existing research, identifying technological requirements, designing a conceptual architecture, and defining future evaluation strategies rather than presenting experimental results. Existing studies related to Generative AI, Large Language Models (LLMs), Retrieval-Augmented Generation (RAG), Explainable AI (XAI), trustworthy AI, and decision-support systems are systematically reviewed to identify key capabilities, limitations, challenges, and design requirements for developing reliable AI-based decision-support solutions.

The methodology consists of four major stages:

1. **Literature Analysis**
2. **Requirement Identification**
3. **Framework and Architecture Design**
4. **Conceptual Evaluation Model Development**

The overall research methodology ensures that the proposed framework addresses both technological capabilities and responsible AI requirements.

4.2 Research Methodology Process

The research process followed in this study is presented below:

Table 2: Research Methodology Process and Framework Development Stages

Stage	Research Activity	Input	Method	Output
Stage1	Literature Analysis	Research papers on GenAI, LLMs, RAG, XAI, AI governance	Systematic review and thematic analysis	Identified technologies, challenges, and research gaps
Stage2	Requirement Identification	Existing limitations of GenAI decision systems	Capability and risk analysis	Framework design requirements
Stage3	Architecture Development	Identified requirements	Conceptual modelling and system design	Evidence-grounded GenAI architecture
Stage4	Evaluation Framework	Decision-support quality factors	Metric definition and validation planning	Future empirical evaluation model

4.3 Summary of Methodological Contribution

The proposed methodology provides a structured approach for designing trustworthy GenAI decision-support systems. Unlike approaches that focus only on model capability, this research integrates technical performance with reliability, transparency, human responsibility, and governance requirements.

The methodology establishes the foundation for future empirical studies where the proposed framework can be tested using domain-specific datasets, expert evaluation panels, and real-world deployment scenarios.

5. Fundamentals of Generative AI Technologies and Proposed Evidence-Grounded Architecture

5.1 Fundamentals of Generative Artificial Intelligence Technologies

Generative Artificial Intelligence (GenAI) represents an advanced generation of AI systems capable of creating content, generating solutions, synthesizing information, and supporting interactive reasoning. Unlike traditional AI approaches focused on prediction and classification, GenAI produces new outputs by learning from large-scale datasets. Modern GenAI systems are primarily based on Transformer architectures, which use attention mechanisms to capture relationships within information and efficiently process complex contexts [1]. These developments have enabled Large Language Models (LLMs) to perform advanced natural language understanding, generation, summarization, and reasoning tasks across various intelligent applications.

GenAI technologies relevant to intelligent decision-making include:

- Large Language Models (LLMs)
- Embedding Models
- Vector Databases
- Retrieval-Augmented Generation (RAG)
- Multimodal AI Models
- AI Agents
- Structured Output Generation
- Explainable AI Mechanisms

These technologies collectively enable AI systems to transform large volumes of information into meaningful decision-support outputs.

5.2 Large Language Models (LLMs)

Large Language Models (LLMs) serve as the primary intelligence layer in modern Generative AI applications by enabling machines to understand, generate, and process human language. These models are trained using extensive textual datasets, allowing them to capture complex relationships between concepts and contextual patterns.

Major capabilities of LLMs include:

- Natural language processing
- Content generation
- Information summarization
- Question answering
- Contextual reasoning
- Knowledge extraction
- Interactive user communication

In intelligent decision-support systems, LLMs act as a communication and reasoning interface that helps users analyse complex information through natural language interactions. However, LLMs may generate inaccurate or unsupported outputs because they rely on learned patterns rather than direct verification of facts. Therefore, reliable deployment requires integration with evidence retrieval, validation mechanisms, human supervision, and governance frameworks to ensure trustworthy decision-making.

5.3 Embedding Models and Semantic Knowledge Representation

Embedding models play a significant role in modern Generative AI architectures by converting text, documents, and other data sources into numerical vector representations that preserve semantic meaning. These representations enable AI systems to understand relationships between concepts and retrieve relevant information more effectively.

Unlike traditional keyword-based search approaches that depend on exact term matching, embedding-based retrieval identifies information based on contextual similarity and underlying meaning.

In intelligent decision-support systems, embedding models support:

- Semantic information retrieval
- Knowledge similarity analysis
- Context-aware matching

- Relevant document ranking
- Domain-specific information discovery

For instance, in healthcare applications, embedding models can retrieve related medical evidence even when user queries contain different terms from the stored knowledge sources. Thus, embedding models provide an essential connection between human language understanding and machine-readable knowledge processing, improving the efficiency and accuracy of GenAI systems.

5.4 Vector Databases for Knowledge Retrieval

Vector databases are important components of Generative AI architectures because they store and organize high-dimensional vector representations created by embedding models. These databases enable efficient searching and retrieval of relevant information from large and complex knowledge collections based on semantic similarity.

In GenAI-based decision-support systems, vector databases can maintain various knowledge resources, including:

- Research publications
- Organizational documents
- Historical data records
- Domain-specific knowledge bases
- Technical reports
- Regulatory guidelines

During the decision-making process, user queries are transformed into vector representations and compared with stored embeddings to identify the most relevant information. The retrieved knowledge is then provided to the generative model to improve response accuracy and contextual understanding.

The use of vector databases enhances:

- Knowledge retrieval efficiency
- Information accessibility
- Domain adaptation
- Continuous knowledge integration
- Evidence-supported generation

This capability is highly valuable in dynamic environments where decision-making depends on frequently updated and diverse information sources.

5.5 Retrieval-Augmented Generation (RAG)

Retrieval-Augmented Generation (RAG) is a major technology for improving the reliability of GenAI systems. RAG combines two major components:

1. Information Retrieval
2. Generative Language Processing

The retrieval component identifies relevant external information, while the generation component uses the retrieved knowledge to produce contextual responses.

The general RAG workflow includes:

1. User submits a decision query.
2. Query is transformed into a semantic representation.
3. Relevant documents are retrieved.
4. Retrieved information is validated.
5. The LLM generates an evidence-supported response.

RAG reduces dependency on internal model memory and enables AI systems to access updated and domain-specific knowledge.

However, retrieval alone does not guarantee reliable decisions. Retrieved information may contain errors, outdated content, or irrelevant sources. Therefore, the proposed framework integrates additional validation and confidence assessment mechanisms.

5.6 Multimodal Generative AI Technologies

Modern decision environments often involve multiple forms of information, including:

- Text documents
- Images
- Tables
- Graphs
- Audio information
- Sensor data

Multimodal GenAI systems extend traditional language models by enabling interpretation and generation across multiple data formats.

Applications include:

- Medical image analysis support
- Industrial monitoring
- Financial document analysis
- Smart education systems
- Intelligent business analytics

The integration of multimodal capabilities allows decision-support systems to analyze broader information sources and generate more comprehensive recommendations.

5.7 AI Agents and Tool-Using Systems

AI agents represent an emerging direction in GenAI research where models are capable of performing multi-step tasks by interacting with external tools and systems.

Agent-based GenAI systems can:

- Search information sources
- Execute workflows
- Call external APIs
- Analyze data
- Generate reports
- Coordinate multiple AI tasks

For intelligent decision-making, AI agents provide opportunities for developing autonomous analytical assistants. However, uncontrolled agent behavior introduces risks related to security, reliability, and accountability.

Therefore, agent-based decision systems require:

- Permission control
- Workflow restrictions
- Human approval mechanisms
- Continuous monitoring

5.8 Proposed Evidence-Grounded Generative AI Decision Architecture

Based on the identified technological requirements, this research introduces an evidence-grounded Generative AI architecture aimed at developing reliable and trustworthy intelligent decision-support systems. The proposed framework combines knowledge management, retrieval mechanisms, generative intelligence, analytical evaluation, human supervision, and audit capabilities into a unified architecture.

The architecture consists of six interconnected layers:

Layer 1: Data and Context Management Layer

This layer serves as the foundation of the framework by collecting, organizing, and preparing relevant information required for decision-making. It ensures that the AI system understands the operational environment, objectives, and constraints before generating recommendations.

Major components include:

- Data sources integration
- Context understanding and extraction
- Data quality assessment
- Privacy and security management
- User access control

The primary purpose of this layer is to provide accurate, relevant, and secure information for downstream AI processing.

Layer 2: Retrieval and Evidence Validation Layer

This layer enhances the reliability of GenAI outputs by connecting the system with trusted external knowledge sources through Retrieval-Augmented Generation (RAG) mechanisms.

Key components include:

- Knowledge repositories
- Embedding-based search systems
- Vector databases
- Retrieval and ranking mechanisms
- Evidence validation modules

This layer ensures that generated recommendations are supported by relevant, verified, and traceable information rather than unsupported model-generated responses.

Layer 3: Generative AI Processing Layer

The Generative AI core acts as the intelligence engine of the framework by analysing validated information and producing decision-support outputs.

Major functions include:

- Knowledge synthesis
- Recommendation generation
- Explanation development
- Scenario exploration
- Natural language communication

The outputs produced by this layer are considered AI-assisted recommendations that require further evaluation before final decision implementation.

Layer 4: Decision Analytics Layer

This layer transforms AI-generated recommendations into structured decision intelligence by applying analytical techniques and evaluation mechanisms.

Main functions include:

- Risk assessment
- Multi-criteria decision analysis
- Confidence measurement
- Alternative comparison
- Sensitivity evaluation

This layer enables systematic assessment of possible decisions and helps identify the most suitable alternatives based on predefined objectives.

Layer 5: Human Oversight and Governance Layer

This layer introduces human involvement to ensure responsible and ethical AI usage. It maintains expert control over critical decisions and prevents excessive dependence on automated recommendations.

Key functions include:

- Expert review and verification
- Decision approval process
- User feedback integration
- Risk-based escalation

High-risk, uncertain, or low-confidence recommendations are transferred to human experts for additional assessment before implementation.

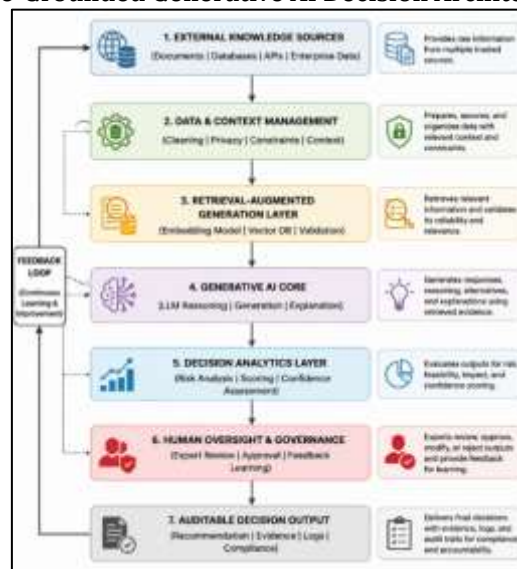
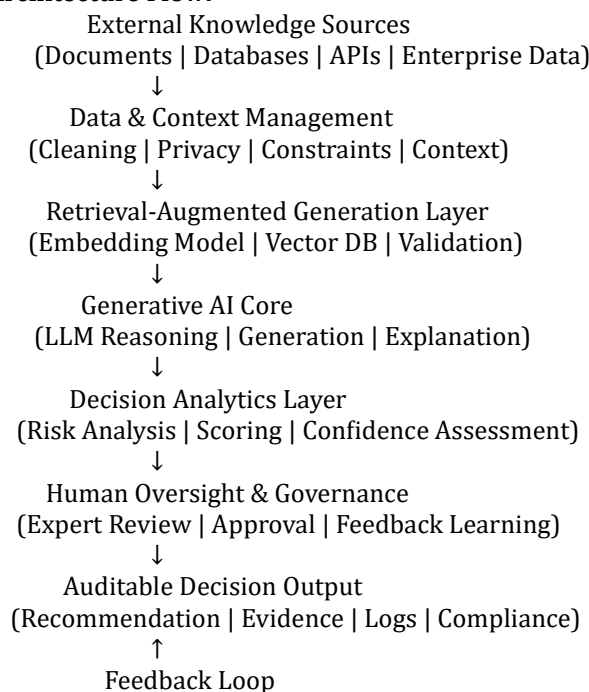
Layer 6: Auditable Decision Output Layer

The final layer generates transparent and traceable decision outputs by maintaining complete records of the AI-assisted decision process.

Generated outputs include:

- Final recommendation
- Supporting evidence sources
- Confidence assessment
- Decision explanation
- Audit records

This layer improves accountability, transparency, and continuous improvement by maintaining a complete history of decision generation and evaluation.

Figure 1. Proposed Evidence-Grounded Generative AI Decision Architecture (Redesigned Concept)**Architecture Flow:****5.9 Summary**

Generative AI technologies provide powerful capabilities for improving intelligent decision-making through knowledge synthesis, reasoning assistance, and interactive communication. However, reliable deployment requires more than advanced language generation. The proposed evidence-grounded architecture integrates LLMs, RAG, validation mechanisms, analytics, human oversight, and auditability to create a trustworthy decision-support framework.

The architecture shifts the role of GenAI from an autonomous decision generator toward an accountable intelligence augmentation system that supports human expertise while maintaining reliability, transparency, and governance.

6. Proposed Decision Algorithm and Intelligent Decision-Making Workflow**6.1 Proposed Algorithm**

The proposed evidence-grounded Generative AI decision algorithm is designed to transform a traditional generative AI system into a trustworthy intelligent decision-support mechanism. The algorithm integrates information retrieval, evidence validation, generative reasoning, decision analytics, confidence assessment, and human oversight.

Unlike conventional AI approaches that directly generate recommendations from learned patterns, the proposed algorithm follows a controlled decision pipeline where every generated output is evaluated against available evidence, domain constraints, and risk conditions.

The primary objective of this algorithm is to ensure that GenAI-generated recommendations are:

- Evidence-supported
- Context-aware
- Explainable
- Risk-assessed
- Human-verifiable
- Auditable

The algorithm follows the principle:

Retrieve → Validate → Generate → Evaluate → Approve → Learn

This approach ensures that Generative AI acts as an intelligent assistant rather than an uncontrolled autonomous decision-maker.

6.2 Algorithm Design Objectives

The proposed algorithm provides a structured workflow for reliable GenAI-based decision-making by integrating context analysis, evidence retrieval, validation, generation, and evaluation processes.

Objective 1: Decision Context Analysis

The algorithm identifies decision requirements by analysing user queries, objectives, available information, environmental factors, and constraints to understand the decision scenario.

Objective 2: Evidence Retrieval

Relevant information is collected from trusted sources using semantic search, source reliability, relevance, and domain suitability to improve knowledge grounding.

Objective 3: Evidence Validation

Retrieved information is verified through source assessment, relevance checking, and compliance evaluation to ensure reliable inputs for AI generation.

Objective 4: Alternative Generation

The GenAI model generates possible solutions, scenarios, risks, and recommendations based on validated evidence and contextual information.

Objective 5: Decision Evaluation

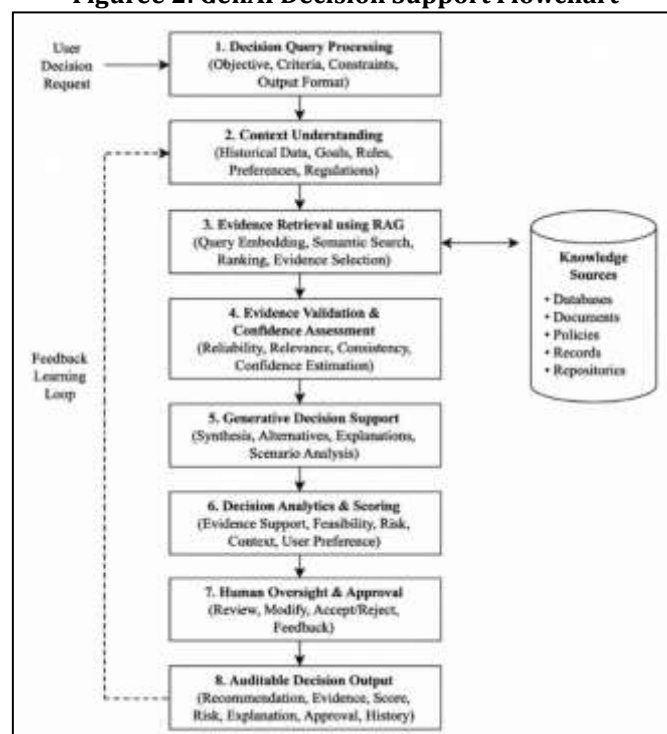
Generated alternatives are ranked using evidence quality, feasibility, risks, user requirements, and organizational objectives to identify suitable recommendations.

This workflow ensures that GenAI provides evidence-supported recommendations while maintaining reliability, transparency, and human decision control.

6.3 Proposed Algorithm Workflow

The proposed workflow consists of eight major stages.

Figure 2: GenAI Decision Support Flowchart



Step 1: Decision Query Processing

The process begins when a user submits a decision request.

Example:

"A company wants to select the best technology investment strategy for the next five years."

The system extracts:

- Decision objective
- Required information
- Evaluation criteria
- Constraints
- Expected output format

The query is converted into a structured decision representation.

Step 2: Context Understanding and Requirement Analysis

The system analyses the surrounding context associated with the decision.

Context information may include:

- Historical data
- Organizational objectives
- Domain rules
- User preferences
- Regulatory requirements

This stage prevents generic AI responses by adapting recommendations to the specific decision environment.

Step 3: Evidence Retrieval Using RAG

After understanding the decision context, the system retrieves relevant information from connected knowledge sources.

The retrieval process includes:

1. Query embedding generation
2. Semantic similarity search
3. Document ranking
4. Evidence selection

Sources may include:

- Enterprise databases
- Research documents
- Policy repositories
- Historical records
- Knowledge management systems

Step 4: Evidence Validation and Confidence Assessment

Retrieved information is evaluated before entering the generation stage.

The validation module checks:

Source Reliability

Determines whether information originates from trusted sources.

Evidence Relevance

Measures whether retrieved information directly supports the decision problem.

Information Consistency

Identifies conflicts between multiple information sources.

Confidence Estimation

Calculates the reliability level of available evidence.

If confidence is below a predefined threshold, the system may:

- Request additional information
- Retrieve alternative sources
- Escalate to human review

Step 5: Generative Decision Support

The validated evidence is provided to the Generative AI model.

The model performs:

- Information synthesis
- Alternative generation
- Explanation creation
- Scenario analysis

The generated response contains:

- Possible decisions
- Advantages
- Limitations
- Risks
- Supporting evidence

The system clearly separates:

Retrieved Facts

from

AI-Generated Interpretation

This separation improves transparency.

Step 6: Decision Analytics and Scoring

The generated alternatives are evaluated using decision analytics techniques.

A conceptual decision score can be represented as:

$$DecisionScore = w_1E + w_2F + w_3R + w_4C + w_5U$$

Where:

- **E = Evidence Support**
- **F = Feasibility**
- **R = Risk Evaluation**
- **C = Context Compatibility**
- **U = User Preference Alignment**
- **w = Weight assigned to each factor**

The scoring mechanism allows systematic comparison between generated alternatives.

Step 7: Human Oversight and Approval

The final recommendation passes through human review based on risk and confidence levels.

Decision categories:

High Confidence + Low Risk

The system provides recommendation with minimal intervention.

Medium Confidence or Moderate Risk

The system requests expert verification.

Low Confidence or High Risk

The decision is transferred completely to human evaluation.

Human reviewers can:

- Accept recommendation
- Modify recommendation
- Reject recommendation
- Provide feedback

Step 8: Decision Output and Audit Recording

The final stage produces an auditable decision report.

The output includes:

- Decision recommendation
- Evidence references
- Confidence score
- Risk assessment
- Explanation
- Human approval status
- Decision history

The audit record supports:

- Accountability
- Compliance
- Future improvement
- Performance evaluation

6.4 Proposed Algorithm: Pseudocode Representation

Algorithm: Evidence-Grounded GenAI Decision Support

Input:

Decision Query Q

Context Information C

Knowledge Sources K

Output:

Validated Decision Recommendation R

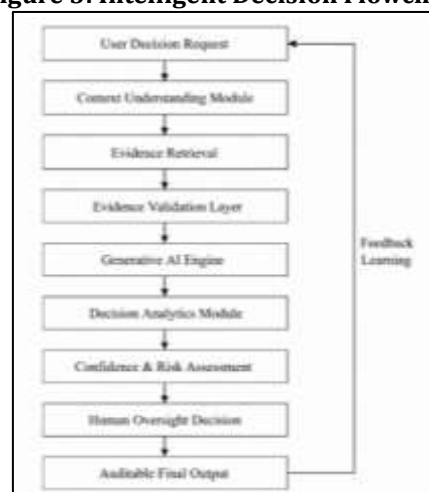
Begin

1. Receive decision query Q
 2. Extract:
 - Objectives
 - Constraints
 - Decision Criteria
 3. Collect contextual information C
 4. Retrieve relevant evidence from K
 5. Validate retrieved information
 6. Calculate evidence confidence score
 7. If confidence < threshold:
 - Request additional evidence
 - or escalate to human review
 8. Generate decision alternatives using LLM
 9. Evaluate alternatives using decision analytics
 10. Rank alternatives based on:
 - Evidence
 - Risk
 - Feasibility
 - User requirements
 11. Perform human review according to risk level
 12. Generate final recommendation
 13. Store:
 - Evidence
 - Reasoning
 - Approval
 - Audit information
- End

6.5 Decision Workflow Architecture

The complete intelligent decision workflow can be represented as:

Figure 3: Intelligent Decision Flowchart



6.6 Algorithm Advantages

The proposed algorithm provides several advantages compared with traditional AI decision-support approaches.

Figure 4: Key Benefits of Proposed Algorithm

7. Intelligent Decision-Making Applications and Case Studies

7.1 Healthcare Decision Support Applications

- Supports clinical decision-making through evidence retrieval, medical knowledge synthesis, and expert validation.
- Improves healthcare analysis by processing patient records, research data, and treatment information.

7.2 Financial Decision-Making Applications

- Assists in risk assessment, market analysis, compliance review, and financial scenario generation.
- Provides evidence-based recommendations using validated information and decision analytics.

7.3 Educational Decision Support Applications

- Enables personalized learning recommendations, academic planning, and student performance analysis.
- Supports educators through curriculum planning and intelligent learning pathway generation.

7.4 Business Intelligence and Strategic Decision Applications

- Helps organizations analyse market trends, customer behaviour, and operational performance.
- Supports strategic planning through scenario generation, forecasting, and knowledge management.

7.5 Public Service and Government Applications

- Improves policy analysis, citizen services, resource planning, and administrative decision support.
- Enhances transparency through evidence-based recommendations and human oversight mechanisms.

Table 3: Applications of Evidence-Grounded GenAI Decision Support

Domain	GenAI Application	Supporting Function	Human Role
Healthcare	Clinical evidence support	Knowledge synthesis and recommendation assistance	Medical expert validation
Finance	Risk analysis and compliance	Scenario generation and policy analysis	Financial approval
Education	Academic guidance	Personalized recommendations	Teacher/advisor review
Business	Strategic planning	Market analysis and decision support	Management decision
Public Services	Policy and citizen support	Information analysis and service assistance	Government oversight

7.6 Cross-Domain Challenges in Implementation

Although GenAI provides significant opportunities, implementation across different domains requires careful consideration of common challenges.

Data Quality

Decision-support systems depend on accurate and reliable information. Poor-quality data may result in incorrect recommendations.

Domain Adaptation

AI models must be adapted to specific domain requirements, terminology, and regulations.

Privacy Protection

Sensitive information must be protected through appropriate security mechanisms.

Human Trust

Users must understand AI limitations and avoid excessive dependence on automated recommendations.

Governance Compliance

Organizations must establish policies for responsible AI usage.

8. Comparative Analysis of Traditional AI, LLM-Based Systems, RAG Systems, and Proposed Framework

Evolution of AI Decision Systems

- AI has evolved from rule-based systems to advanced GenAI architectures for complex decision-making.
- Modern systems improve adaptability, reasoning, interaction, and intelligent support capabilities.

Traditional Artificial Intelligence Systems

- Uses predefined rules, logical reasoning, and expert knowledge for structured decision processes.
- Provides high interpretability but has limited flexibility with complex and changing environments.

Limitations:

- Struggles with unstructured data, manual rule creation, and limited natural language interaction.
- Requires fixed procedures, reducing adaptability in modern decision-support applications.

Machine Learning-Based Decision Systems

- Learns patterns from historical datasets to support prediction, classification, and risk analysis.
- Improves decision accuracy by identifying complex relationships within large datasets.

Limitations:

- Depends on structured data and feature engineering with limited explanation capability.
- Struggles with new situations requiring reasoning, alternatives, and contextual understanding.

Standalone Large Language Model Systems

- Provides natural language interaction, text generation, summarization, and reasoning assistance.
- Enables flexible communication and knowledge synthesis through advanced language understanding.

Limitations:

- Faces hallucination risks, outdated knowledge, and limited transparency of generated information.
- Requires evidence retrieval, validation, and governance for reliable decision support.

Retrieval-Augmented Generation (RAG) Systems

- Combines external knowledge retrieval with generative AI to improve response reliability.
- Provides updated information access, evidence-based responses, and domain-specific adaptation.

Limitations:

- Output quality depends on retrieval accuracy and knowledge source reliability.
- Requires additional validation mechanisms for high-impact decision environments.

Proposed Evidence-Grounded Generative AI Framework

- Integrates context management, evidence retrieval, validation, analytics, and human oversight.
- Generates evidence-supported recommendations instead of uncontrolled autonomous decisions.

Key Benefits:

- Improves reliability, transparency, accountability, and responsible AI decision-making.
- Positions GenAI as an intelligent assistant supporting human expertise and governance.

Table4: Comparative Analysis of AI Decision-Support Approaches

Criteria	Traditional AI	Machine Learning	Standalone LLM	RAG System	Proposed Framework
Knowledge Source	Rules and expert knowledge	Historical datasets	Pre-trained model knowledge	External documents + LLM	Validated external knowledge + LLM
Natural Language Interaction	Limited	Moderate	Excellent	Excellent	Excellent
Data Flexibility	Low	Medium	High	High	Very High
Handling Unstructured Data	Limited	Limited	Strong	Strong	Strong
Evidence Grounding	Rule-based	Dataset-based	Weak	Moderate	Strong
Explainability	High	Variable	Limited	Improved	Strong

Hallucination Control	Low risk	Medium	High risk	Reduced	Controlled
Human Oversight	Usually external	Optional	Limited	Recommended	Integrated
Auditability	Rule tracking	Model records	Limited	Partial	Comprehensive
Decision Reliability	Moderate	Moderate	Low to Medium	Medium	High potential

Key Differences Between Existing Systems and Proposed Framework

Evidence Management

Traditional AI and machine-learning systems depend mainly on predefined rules or historical datasets. LLM systems depend on learned knowledge representations. The proposed framework introduces explicit evidence retrieval and validation mechanisms.

Decision Transparency

Many AI systems provide predictions without detailed reasoning. The proposed framework improves transparency by maintaining:

- Evidence sources
- Confidence levels
- Decision rationale
- Validation records

Human-AI Collaboration

Existing AI approaches often treat humans as external users. The proposed framework integrates humans as active decision participants through review, approval, and feedback mechanisms.

Risk-Aware Decision Processing

Most AI systems focus on output generation or prediction accuracy. The proposed framework introduces risk evaluation and confidence-based escalation to manage uncertain decisions.

9. Case Studies and Experimental Evaluation Framework

The proposed evidence-grounded GenAI framework requires evaluation beyond traditional accuracy measures. It considers factors such as evidence reliability, explainability, security, human trust, and governance compliance.

The current study presents a conceptual framework; therefore, experimental validation is proposed for future research. Evaluation can compare standalone LLMs, RAG systems, and the proposed framework using domain-specific datasets.

Key evaluation factors include:

- Decision accuracy
- Evidence support
- Explainability
- Reliability
- Security
- Human oversight

Healthcare Clinical Decision Support

GenAI can support healthcare professionals by analysing patient records, medical literature, and clinical guidelines. The framework retrieves validated medical evidence, generates clinical insights, and provides explanations under expert supervision.

Evaluation focuses on:

- Diagnostic accuracy
- Evidence relevance
- Expert agreement
- Patient safety

Financial Risk Decision Support

The framework assists financial analysts by analysing market data, risk factors, and regulatory information. It supports scenario generation, risk evaluation, and evidence-based financial decisions.

Evaluation includes:

- Risk identification

- Recommendation reliability
- Evidence quality
- Compliance support

Educational Decision Support

GenAI supports personalized learning, academic planning, and student guidance by analysing performance data, learning objectives, and course requirements.

Evaluation measures:

- Recommendation relevance
- Learning improvement
- Faculty acceptance
- Personalization quality

Future Experimental Validation

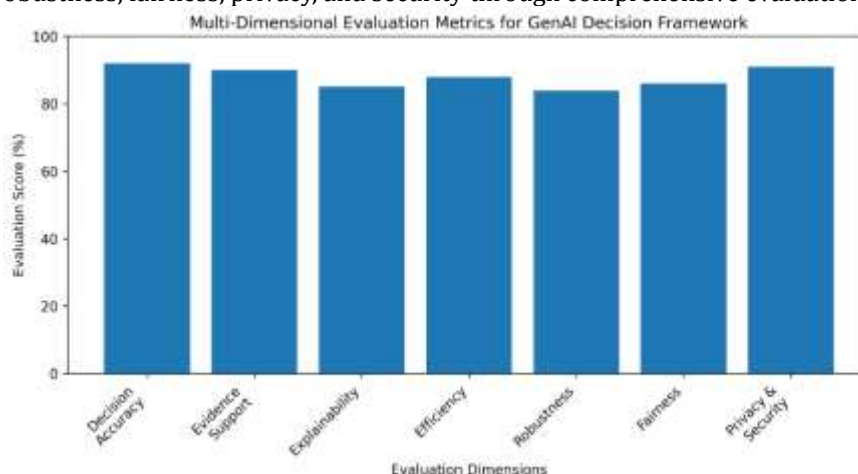
Future studies can compare LLM-based systems, RAG approaches, and the proposed framework to evaluate improvements in decision reliability, transparency, and effectiveness.

Table 5: Proposed Experimental Evaluation Framework

Experiment	Dataset / Participants	Comparison	Evaluation Objective
Experiment 1: Evidence Grounding	Domain-specific question datasets	LLM vs RAG vs Proposed Framework	Measure evidence-supported accuracy
Experiment 2: Decision Quality	Expert-labelled decision scenarios	Existing AI vs Proposed Framework	Measure expert agreement
Experiment 3: Human-AI Collaboration	Decision-maker participants	AI-assisted vs Manual decision-making	Measure efficiency and confidence
Experiment 4: Robustness Testing	Modified and adversarial inputs	All systems	Evaluate stability
Experiment 5: Governance Evaluation	Risk-based scenarios	Control variations	Measure compliance readiness

Evaluation Metrics

- ✓ Evaluates decision accuracy, evidence support, explainability, and efficiency of systems.
- ✓ Measures robustness, fairness, privacy, and security through comprehensive evaluation methods.



Proposed Hypotheses for Future Validation

The following hypotheses can guide future empirical research:

H1: Evidence-grounded GenAI decision-support systems achieve higher expert-validated decision accuracy compared with standalone LLM systems.

H2: Integrating retrieval confidence and validation mechanisms improves decision reliability in high-risk environments.

H3: Evidence-linked explanations increase user trust and reduce inappropriate dependence on AI-generated recommendations.

H4: Human oversight mechanisms reduce unsupported or policy-inconsistent AI decisions.

H5: The proposed framework improves decision-making efficiency while maintaining decision quality.

10. Performance Evaluation, Security, Privacy, Ethics, and Future Directions

Performance Evaluation

- Performance evaluation measures the effectiveness of GenAI decision-support systems.
- It focuses on quality, reliability, and practical decision improvement.

Decision Accuracy

- Measures alignment between AI recommendations and expert-validated decisions.
- Evaluates correctness, consistency, suitability, and expert agreement.

Methods:

- Expert reviews and benchmark datasets
- Comparative analysis with existing approaches

Evidence Support

- Evaluates whether AI recommendations are supported by reliable information sources.
- Ensures traceability, relevance, and credibility of generated outputs.

Evaluation Factors:

- Evidence accuracy and source credibility
- Information relevance and coverage

Response Efficiency

- Measures the ability of GenAI systems to improve decision speed.
- Ensures faster responses without reducing decision quality.

Key Measures:

- Response and processing time
- Decision completion and user productivity

Explainability Evaluation

- Evaluates how clearly users understand AI-generated recommendations.
- Improves trust through transparent reasoning and evidence connection.

Evaluation Factors:

- Explanation clarity and reasoning transparency
- User understanding and confidence improvement

Robustness Evaluation

- Measures system stability under uncertain and changing conditions.
- Ensures reliable performance in complex decision environments.

Testing Conditions:

- Incomplete information and ambiguous queries
- Conflicting evidence and adversarial inputs

Security, Privacy, Ethics, and Future Research

- Ensures safe and responsible implementation of GenAI decision systems.
- Covers security protection, privacy control, ethics, and future improvements.

Table 6: Multi-Dimensional Evaluation Framework for GenAI Decision Support

Evaluation Dimension	Purpose	Possible Measurement
Decision Accuracy	Evaluate recommendation quality	Expert agreement, validation score
Evidence Support	Measure grounding reliability	Citation coverage, evidence relevance
Efficiency	Evaluate operational improvement	Response time, task completion time
Explainability	Assess transparency	User rating, explanation quality
Robustness	Test stability	Performance under changed inputs
Privacy	Evaluate data protection	Security assessment
Governance	Measure responsible operation	Compliance checklist

Security Considerations

- Security protects sensitive information and AI workflows from threats.
- Framework applies protection across data, models, applications, and workflows.

Security Areas:

- Data, Model, Application, and Workflow Security
- Monitoring, validation, and controlled deployment

Data Security

- Protects confidential data through encryption and access control.
- Ensures secure storage and privacy-preserving information processing.

Model Security

- Prevents model manipulation, misuse, and information leakage risks.
- Uses monitoring, filtering, and continuous evaluation mechanisms.

Prompt Injection Protection

- Prevents malicious instructions from influencing AI behaviour.
- Uses validation, source verification, and output monitoring techniques.

Audit and Monitoring Security

- Maintains records of sources, outputs, and user activities.
- Supports compliance, investigation, and continuous system improvement.

Privacy Protection Framework

- Ensures responsible handling of personal and organizational information.
- Improves confidence through secure data management practices.

Key Mechanisms:

- Data minimization and role-based access management
- Secure processing and privacy risk monitoring

Ethical Issues in GenAI Decision-Making

- Addresses fairness, accountability, transparency, and responsible AI usage.
- Ensures GenAI supports decisions without replacing human responsibility.

Bias and Fairness

- Reduces biased outcomes caused by data limitations.
- Requires continuous fairness evaluation throughout AI lifecycle.

Human Responsibility

- Maintains expert involvement in important decisions.
- Uses review and approval mechanisms for accountability.

Transparency and Trust

- Provides understanding of AI capabilities and limitations.
- Builds trust through evidence and confidence information.

Responsible Deployment

- Establishes policies, monitoring, and ethical guidelines.
- Ensures safe and accountable AI implementation.

Future Research Directions

- Future studies aim to improve GenAI reliability and adaptability.
- Research focuses on advanced evaluation and intelligent decision systems.

Benchmark Dataset Development

- Creates domain-specific datasets for decision evaluation.
- Supports healthcare, finance, education, and business applications.

Advanced Confidence Estimation

- Improves uncertainty measurement and reliability assessment.
- Combines evidence quality, retrieval, and model confidence.

Multimodal Decision Intelligence

- Integrates text, images, tables, and sensor information.
- Enables richer and comprehensive decision support.

AI Agent-Based Systems

- Explores secure autonomous agents with controlled operations.
- Includes tool management and human approval mechanisms.

Human-AI Collaboration Studies

- Studies trust, decision quality, and organizational impact.
- Evaluates long-term interaction between humans and AI.

Energy-Efficient Domain Models

- Develops lightweight models with lower computational needs.
- Improves privacy, efficiency, and domain adaptation.

11. Discussion, Research Limitations, Conclusion and Final Research Contributions**11.1 Discussion**

Generative AI has improved intelligent decision-support systems through advanced reasoning, information synthesis, and interactive communication. However, reliable adoption requires addressing challenges related to accuracy, transparency, security, and governance. The proposed evidence-grounded GenAI framework integrates retrieval, validation, analytics, and human oversight to improve trustworthy decision-making.

11.2 Proposed Architecture Discussion

The proposed framework consists of six layers:

- Data and Context Management
- Retrieval and Evidence Validation
- Generative AI Core
- Decision Analytics
- Human Oversight
- Auditable Decision Output

These layers ensure reliable information processing, evidence-based recommendations, expert validation, and transparent decision tracking.

11.3 Comparison with Existing Approaches

Traditional AI provides rule-based decisions, while machine learning improves prediction capabilities. LLMs offer advanced language interaction but face reliability challenges. The proposed framework combines retrieval, generation, analytics, validation, and governance for trustworthy GenAI decision support.

11.4 Research Contributions

The major contributions include:

- **Evidence-Grounded Architecture:** Integrates retrieval, validation, and AI generation for reliable decisions.
- **Human-Centered Framework:** Maintains human responsibility in decision-making processes.
- **Risk-Aware Processing:** Uses confidence evaluation for uncertain situations.
- **Multi-Dimensional Evaluation:** Measures accuracy, explainability, security, and governance.

11.5 Research Limitations

The framework is currently conceptual and requires experimental validation using real-world datasets. Different domains need customized adaptation based on regulations and requirements. System reliability depends on knowledge quality, and future research should focus on efficient and lightweight GenAI models.

11.6 Future Research Directions

Future research should focus on empirical validation, benchmark dataset development, improved retrieval validation, and advanced confidence estimation. Multimodal GenAI systems integrating text, images, tables, and sensor data can provide richer decision support.

Additionally, long-term human-AI collaboration studies should examine trust, decision improvement, organizational impact, and responsible AI adoption. These developments can support the creation of reliable, transparent, and efficient GenAI-based decision-support systems.

11.7 Conclusion

Generative Artificial Intelligence enhances decision-support systems through information synthesis, reasoning, and interactive communication. This research proposed an evidence-grounded GenAI framework integrating retrieval, validation, analytics, human oversight, and auditability to overcome challenges such as hallucination and limited transparency.

The framework emphasizes GenAI as a decision-support assistant rather than an autonomous decision-maker. Future research should validate the approach across healthcare, finance, education, business, and public sectors to develop reliable, ethical, and trustworthy AI systems.

References

- [1] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, Ł. Kaiser, and I. Polosukhin, "Attention Is All You Need," in *Advances in Neural Information Processing Systems (NeurIPS)*, 2017.
- [2] T. Brown et al., "Language Models are Few-Shot Learners," in *Advances in Neural Information Processing Systems*, vol. 33, pp. 1877–1901, 2020.
- [3] J. Wei et al., "Chain-of-Thought Prompting Elicits Reasoning in Large Language Models," *Advances in Neural Information Processing Systems*, vol. 35, pp. 24824–24837, 2022.
- [4] R. Bommasani et al., "On the Opportunities and Risks of Foundation Models," Stanford Center for Research on Foundation Models, 2021.
- [5] OpenAI, "GPT-4 Technical Report," arXiv preprint arXiv:2303.08774, 2023.
- [6] P. Lewis et al., "Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks," in *Advances in Neural Information Processing Systems*, vol. 33, pp. 9459–9474, 2020.
- [7] Y. Gao et al., "Retrieval-Augmented Generation for Large Language Models: A Survey," arXiv preprint arXiv:2312.10997, 2023.
- [8] P. Zhao et al., "Retrieval-Augmented Generation for AI-Generated Content: A Survey," arXiv preprint arXiv:2402.19473, 2024.
- [9] W. Fan et al., "A Survey on RAG Meeting LLMs: Towards Retrieval-Augmented Large Language Models," arXiv preprint arXiv:2405.06211, 2024.
- [10] Y. Xiong et al., "Graph Retrieval-Augmented Generation for Large Language Models: A Survey," in *IEEE Conference on AI, Science, Engineering and Technology*, 2024.
- [11] T. Wu et al., "A Comprehensive Survey of Large Language Models and Their Applications," *ACM Computing Surveys*, 2024.
- [12] H. Naveed et al., "A Comprehensive Overview of Large Language Models," arXiv preprint arXiv:2307.06435, 2023.
- [13] Z. Ji et al., "Survey of Hallucination in Natural Language Generation," *ACM Computing Surveys*, vol. 55, no. 12, pp. 1–38, 2023.
- [14] M. Chen et al., "Evaluating Large Language Models Trained on Code," arXiv preprint arXiv:2107.03374, 2021.
- [15] D. Gunning and D. Aha, "DARPA's Explainable Artificial Intelligence Program," *AI Magazine*, vol. 40, no. 2, pp. 44–58, 2019.
- [16] A. Adadi and M. Berrada, "Peeking Inside the Black-Box: A Survey on Explainable Artificial Intelligence (XAI)," *IEEE Access*, vol. 6, pp. 52138–52160, 2018.
- [17] R. Guidotti et al., "A Survey of Methods for Explaining Black Box Models," *ACM Computing Surveys*, vol. 51, no. 5, pp. 1–42, 2019.
- [18] National Institute of Standards and Technology (NIST), "Artificial Intelligence Risk Management Framework (AI RMF 1.0)," NIST AI 100-1, 2023.
- [19] National Institute of Standards and Technology (NIST), "Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile," NIST AI 600-1, 2024.
- [20] European Union, "Artificial Intelligence Act," European Parliament and Council Regulation, 2024.
- [21] I. D. Raji et al., "Closing the AI Accountability Gap: Defining an End-to-End Framework for Internal Algorithmic Auditing," in *Proceedings of FAT*, 2020.
- [22] B. Mittelstadt, "Principles Alone Cannot Guarantee Ethical AI," *Nature Machine Intelligence*, vol. 1, pp. 501–507, 2019.
- [23] S. Russell, D. Dewey, and M. Tegmark, "Research Priorities for Robust and Beneficial Artificial Intelligence," *AI Magazine*, vol. 36, no. 4, pp. 105–114, 2015.
- [24] M. Brundage et al., "The Malicious Use of Artificial Intelligence: Forecasting, Prevention, and Mitigation," 2018.
- [25] A. Bommasani et al., "The Foundation Model Transparency Index," Stanford Institute for Human-Centered AI, 2023.

- [26] A. Mialon et al., "Augmented Language Models: A Survey," arXiv preprint arXiv:2302.07842, 2023.
- [27] C. Jeong, "A Study on the Implementation of Generative AI Services Using an Enterprise Data-Based LLM Application Architecture," arXiv preprint arXiv:2309.01105, 2023.
- [28] A. Habib et al., "Towards Explainable AI in Agentic Retrieval-Augmented Generation: A Systematic Review," IEEE, 2025.
- [29] "Certifying Generative AI: Retrieval-Augmented Generation Chatbots in High-Stakes Environments," *IEEE Journals & Magazine*, 2024.
- [30] "Large Language Models with Retrieval-Augmented Generation Enhance Expert Modelling for Clinical Decision Support," *Journal of Medical Internet Research / Clinical AI Research*, 2025.