



TARS AI: Transforming Static Documents into Intelligent Systems for Enhanced Knowledge Accessibility and Productivity

Shivam Mishra

Dept. of Computer Science and Engineering
Geetanjali Institute of Technical Studies
Udaipur, Rajasthan, India
iamshivammishra49@gmail.com

Tanay Chourasiya

Dept. of Computer Science and Engineering
Geetanjali Institute of Technical Studies
Udaipur, Rajasthan, India
chourasiyatanay@gmail.com

Dr. Mayank Patel

Dept. of Computer Science and Engineering
Geetanjali Institute of Technical Studies
Udaipur, Rajasthan, India
mayank999_udaipur@yahoo.com

Abstract: The rapid growth of heterogeneous data sources such as PDF, DOCX, images, and audio recordings has created significant challenges in efficient information retrieval and knowledge accessibility. Traditional search systems operate in isolated modalities and fail to provide context-aware, cross-format understanding, resulting in reduced productivity and increased effort in data exploration. To address this limitation, this paper presents TARS AI, a multimodal Retrieval-Augmented Generation system designed to operate in offline mode for unified semantic retrieval across diverse data formats. The proposed system integrates advanced preprocessing techniques including optical character recognition for images, speech-to-text conversion for audio, and text extraction from structured documents, followed by embedding generation using transformer-based models. All modalities are indexed within a shared vector space using FAISS to enable efficient similarity search. A hybrid retrieval mechanism combines semantic and keyword-based approaches to fetch relevant context, which is then processed by a Large Language Model for generating accurate, context-aware responses with source citations. The system supports natural language queries and enables cross-modal linking, allowing users to retrieve related information across text, images, and audio seamlessly. The architecture is implemented using React and TypeScript for the frontend, FastAPI for backend services, and integrates scalable databases including PostgreSQL, MongoDB, and Redis. Experimental evaluation demonstrates improved retrieval accuracy, reduced response time, and enhanced user interaction compared to traditional document search systems. TARS AI provides a scalable and efficient solution for intelligent document processing, with applications in government, education, and enterprise domains, significantly improving knowledge accessibility and decision-making efficiency.

Keywords: multimodal retrieval, retrieval-augmented generation, document intelligence, semantic search, large language model, vector database, knowledge accessibility, cross-modal retrieval.

I. INTRODUCTION

Consider a scenario where a government analyst needs to locate a specific insight buried within thousands of heterogeneous files—PDF reports, scanned images, recorded meetings, and handwritten notes. Despite the availability of digital storage, retrieving meaningful information across these formats remains a complex and time-consuming task. Traditional search systems rely heavily on keyword matching and operate within isolated data modalities, making it nearly impossible to establish contextual relationships between text, images, and audio. This fragmentation of information represents a critical barrier to efficient knowledge utilization in modern organizations.

The rapid digitization of content across sectors such as government, education, and enterprise has led to an exponential increase in unstructured and multimodal data. Reports, policy documents, screenshots, voice recordings, and free-form notes are generated daily, yet remain underutilized due to the absence of intelligent retrieval systems. Existing tools fail to provide semantic understanding, cross-modal linking, and contextual reasoning. As a result, users spend a significant portion of their time searching for relevant information rather than applying it effectively. This challenge is further amplified in environments where data must be

processed offline due to privacy, security, or infrastructure constraints.

Three major gaps define this problem. The first is the **semantic gap**, where systems are unable to understand the contextual meaning of queries beyond simple keyword matching. The second is the **modal gap**, where text, image, and audio data are processed independently without any unified representation. The third is the **interaction gap**, where users cannot engage with data through natural language or receive explainable, citation-backed responses. These limitations highlight the urgent need for an intelligent system capable of bridging diverse data formats into a unified, searchable knowledge space.

To address these challenges, this paper presents **TARS AI**, a multimodal Retrieval-Augmented Generation system designed for offline environments. The proposed system enables ingestion of diverse data formats including DOCX, PDF, images, and audio, and transforms them into a shared semantic representation using advanced embedding techniques. Optical Character Recognition is employed for extracting textual information from images, while speech-to-text models convert audio data into searchable transcripts. All processed data is indexed within a unified vector database, enabling efficient cross-modal retrieval. A hybrid retrieval mechanism combines semantic similarity with keyword-based

search to identify relevant context, which is then utilized by a Large Language Model to generate accurate, context-aware responses with explicit source citations.

The system further supports natural language interaction through a unified query interface, allowing users to retrieve information seamlessly across multiple modalities. By establishing cross-format links, such as connecting an audio transcript segment to a corresponding document section or image, TARS AI enhances both transparency and usability. The architecture is implemented using modern technologies including React and TypeScript for the frontend, FastAPI for backend services, and scalable data management systems such as PostgreSQL, MongoDB, and Redis.

The remainder of this paper is organized as follows: Section II reviews related work in multimodal retrieval, document intelligence, and Retrieval-Augmented Generation systems. Section III describes the proposed methodology and algorithmic framework. Section IV presents the system architecture and implementation details. Section V discusses experimental results and performance evaluation. Section VI concludes the paper with insights, limitations, and future research directions.

II. LITERATURE SURVEY

The research landscape in document intelligence and information retrieval has evolved significantly with the emergence of artificial intelligence, particularly in the domains of Natural Language Processing and multimodal learning. However, existing systems often address isolated components of the problem rather than providing a unified solution for cross-modal understanding and retrieval. The literature can be broadly categorized into three key areas relevant to the proposed system.

A. Retrieval-Augmented Generation and Document Intelligence

Recent advancements in Retrieval-Augmented Generation (RAG) have demonstrated substantial improvements in generating accurate and context-aware responses by combining retrieval mechanisms with Large Language Models. Lewis et al. [1] introduced the concept of RAG, highlighting its ability to ground generated outputs in retrieved knowledge, thereby reducing hallucinations and improving reliability. Subsequent studies have extended this approach to document intelligence systems, where large volumes of unstructured data are processed for semantic search and question answering.

Karpukhin et al. [2] proposed dense passage retrieval techniques that significantly enhance semantic matching compared to traditional keyword-based methods. Similarly, Izacard and Grave [3] demonstrated that combining multiple retrieved contexts improves answer generation quality. Despite these advancements, most implementations are limited to text-based datasets and do not address the integration of multiple data modalities such as images and audio. This limitation is critical in real-world scenarios where information exists in diverse formats.

B. Multimodal Learning and Cross-Modal Retrieval

Multimodal learning has gained attention for its ability to process and relate information across different data types. Models such as CLIP (Contrastive Language–Image Pretraining) introduced by Radford et al. [4] enable joint embedding of text and images into a shared semantic space, facilitating cross-modal retrieval. Similarly, research in audio processing and speech recognition, such as work by Baevski et al. [5], has enabled high-quality speech-to-

text conversion, making audio data accessible for downstream NLP tasks.

While these approaches enable individual modality processing, integrating text, image, and audio into a unified retrieval framework remains a challenge. Baltrusaitis et al. [6] emphasize that multimodal systems must address representation alignment and fusion strategies to ensure meaningful cross-modal interactions. Existing systems often lack seamless linking between modalities, limiting their effectiveness in real-world applications that require contextual understanding across formats.

C. Semantic Search Systems and Knowledge Accessibility

Semantic search systems have replaced traditional keyword-based approaches by leveraging embeddings and vector databases for similarity-based retrieval. Johnson et al. [7] introduced FAISS, a highly efficient library for large-scale vector similarity search, enabling real-time retrieval from high-dimensional embeddings. Modern systems further enhance retrieval through hybrid approaches that combine semantic and keyword-based techniques to improve accuracy and relevance.

Research by Guu et al. [8] highlights the importance of integrating retrieval systems with generative models to improve knowledge accessibility and user interaction. Additionally, studies on human-computer interaction emphasize the growing need for natural language interfaces that allow users to interact with complex datasets intuitively. However, most existing systems lack transparency, as they do not provide clear source attribution or citation mechanisms, which are essential for trust and verification.

III. PROPOSED SOLUTION

A. System Overview and Architecture

TARS AI is designed as a **multimodal, offline-capable document intelligence platform** based on a three-tier architecture. The **Presentation Tier** consists of a React.js interface with TypeScript, providing a unified chat-based interaction system along with drag-and-drop file upload capabilities for documents, images, and audio files. The interface supports real-time responses, session-based conversations, and citation-based navigation for retrieved results.

The **Logic Tier** is implemented using FastAPI, enabling high-performance asynchronous processing of user queries and document ingestion pipelines. This layer integrates multiple processing modules including text extraction, image understanding, and audio transcription. It also manages the Retrieval-Augmented Generation pipeline, handling embedding generation, semantic search, and response generation using a Large Language Model. The system is designed to function in **offline mode**, leveraging locally deployed models and efficient processing pipelines.

The **Data Tier** combines structured and unstructured storage systems. PostgreSQL is used for metadata and user session management, while MongoDB stores processed document content. A vector database powered by FAISS is used to store embeddings from all modalities in a shared semantic space, enabling efficient similarity search. Redis is utilized for caching frequently accessed queries and responses, improving system performance. The overall architecture ensures scalability, modularity, and efficient cross-modal retrieval capabilities.

B. Multimodal Data Ingestion and Processing Module

The ingestion module is responsible for handling diverse data formats and converting them into a unified representation. The system processes text documents (PDF, DOCX), images, and audio recordings through specialized pipelines. Text is extracted using PDF parsing libraries, images are processed using Optical Character Recognition, and audio is converted into text using speech-to-text models. The extracted content is then segmented into smaller chunks and transformed into embeddings using transformer-based models.

Algorithm 1: Multimodal Data Ingestion and Indexing

Input: Files F {documents, images, audio}
Output: Indexed embeddings in vector database

```

1: FOR each file  $f \in F$  DO
2: IF type( $f$ ) = document THEN
3: text  $\leftarrow$  extract_text( $f$ )
4: ELSE IF type( $f$ ) = image THEN
5: text  $\leftarrow$  OCR( $f$ )
6: ELSE IF type( $f$ ) = audio THEN
7: text  $\leftarrow$  speech_to_text( $f$ )
8: END IF
9: chunks  $\leftarrow$  split(text into segments)
10: embeddings  $\leftarrow$ 
generate_embeddings(chunks)
11: store(embeddings, metadata) in vector
database
12: END FOR
13: RETURN indexed data

```

This module ensures that all data, regardless of format, is transformed into a consistent semantic representation, enabling unified retrieval across modalities.

C. Semantic Retrieval and RAG-Based Query Processing

The core functionality of TARS AI lies in its ability to retrieve relevant information and generate context-aware responses. When a user submits a natural language query, the system converts it into an embedding and performs similarity search within the vector database. A hybrid retrieval mechanism combines semantic similarity with keyword-based filtering to improve accuracy. The retrieved context is then passed to a Large Language Model, which generates a grounded response with proper citations.

Algorithm 2: Query Processing Using RAG

Input: User query Q
Output: Generated response R with citations

```

1: embedding_q  $\leftarrow$ 
generate_embedding( $Q$ )
2: results  $\leftarrow$  vector_search(embedding_q,
top_k)
3: ranked_results  $\leftarrow$  rank(results based on
relevance)
4: context  $\leftarrow$  concatenate(ranked_results)
5:  $R \leftarrow$  LLM_generate( $Q$ , context)
6: attach source citations to  $R$ 
7: RETURN  $R$ 

```

This approach ensures that responses are not only accurate but also **traceable**, improving trust and transparency in the system.

D. Cross-Modal Retrieval and Linking Mechanism

A key innovation of TARS AI is its ability to establish relationships across different data modalities. By storing embeddings in a shared vector space, the system enables cross-modal queries such as retrieving images based on text descriptions or linking audio transcripts with relevant document sections.

When a query is processed, the system retrieves related data across all modalities and presents them in a unified interface. Each result is linked to its original source, allowing users to explore the context in detail. This mechanism enhances the interpretability of results and enables deeper insights from heterogeneous data.

E. System Optimization and Offline Capability

To meet the requirement of offline operation, TARS AI integrates locally deployable models for embedding generation and language processing. Lightweight transformer models and optimized vector search techniques ensure efficient performance without reliance on external APIs. Caching strategies using Redis reduce redundant computations, while background processing handles large file ingestion tasks.

The system is further optimized through asynchronous processing, database indexing, and efficient memory management, ensuring scalability and responsiveness even in resource-constrained environments.

IV. SYSTEM ARCHITECTURE

A. Input Ingestion and Preprocessing Layer

The input ingestion layer of TARS AI is designed to handle **heterogeneous data formats** including PDF, DOCX, images, scanned documents, audio recordings, and URLs. Each input undergoes a validation phase to ensure format compliance and size constraints before entering the preprocessing pipeline.

For textual documents, content is extracted using PDF parsers and document processing libraries. Images and scanned documents are processed using Optical Character Recognition to convert visual content into machine-readable text. Audio inputs are transformed into text through speech-to-text models, enabling further semantic processing. URLs are parsed and cleaned to extract meaningful textual data.

The output of this layer is a **standardized raw text representation**, ensuring that all input modalities are normalized into a consistent format. This unified representation forms the foundation for downstream processing and enables seamless integration into the multimodal pipeline.

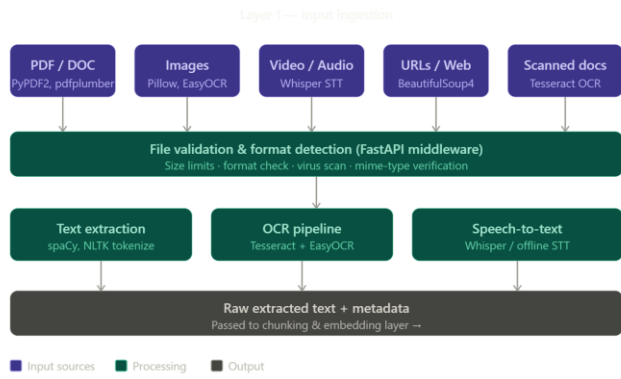


Fig. 1. Multimodal input ingestion and preprocessing pipeline showing validation, text extraction, OCR, and speech-to-text conversion into unified raw text.

B. Chunking, Embedding, and Vector Storage Layer

After preprocessing, the extracted text is segmented into smaller, meaningful chunks using a dynamic chunking strategy that preserves semantic coherence while optimizing retrieval efficiency. Each chunk is then converted into embeddings using transformer-based models.

For images, embeddings are generated using vision-language models, allowing them to be represented in the same semantic space as textual data. This enables **cross-modal understanding**, where relationships between text and images can be effectively captured.

All embeddings are stored in a **unified vector space using FAISS**, enabling high-speed similarity search across large datasets. Alongside the vector database, a hybrid storage architecture is implemented:

- **PostgreSQL** manages structured metadata and relationships
- **MongoDB** stores unstructured document content
- **Redis** provides caching for frequently accessed queries and embeddings

This layered storage approach ensures scalability, fast retrieval, and efficient data management across modalities.

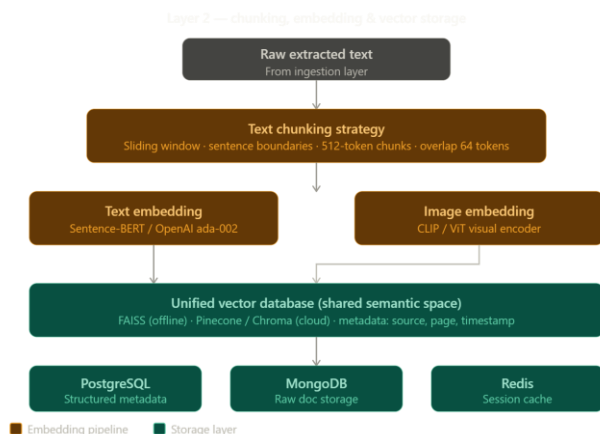


Fig. 2. Chunking, embedding, and vector storage architecture illustrating text segmentation, multimodal embeddings, and unified FAISS vector space with database integration.

C. Query Processing, RAG, and Response Generation Layer

The query processing layer enables users to interact with the system using natural language. When a query is received, it is first preprocessed and converted into an embedding. A **hybrid retrieval mechanism** is then applied, combining:

- BM25-based keyword search for lexical matching
- Semantic vector search for contextual similarity

The retrieved results are ranked based on relevance and aggregated into a contextual input for the Retrieval-Augmented Generation pipeline. This context is passed through a Model Context Protocol (MCP) server, which manages prompt construction and interaction with the Large Language Model.

The LLM generates a **grounded response** by incorporating retrieved context, ensuring factual accuracy and reducing hallucinations. Each response is accompanied by **source citations**, enabling users to trace the origin of the information. This layer ensures transparency, reliability, and enhanced user trust in the system.

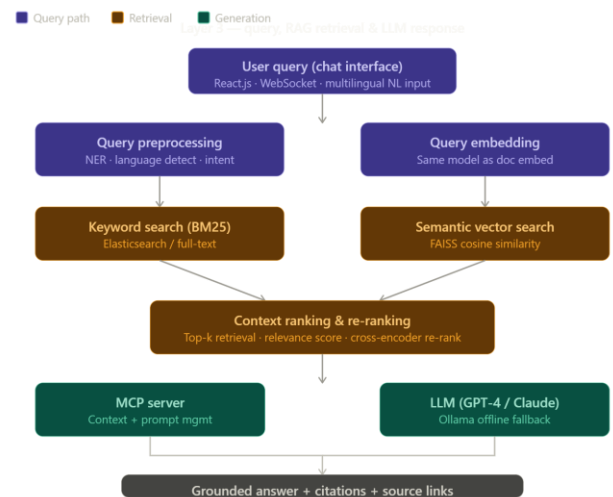


Fig. 3. Query processing and RAG pipeline showing user query embedding, hybrid search, context ranking, and LLM-based response generation with citations.

D. End-to-End System Integration and Microservices Architecture

TARS AI follows a **microservices-based architecture** to ensure modularity and scalability. The system is divided into multiple backend services, each responsible for a specific function such as ingestion, embedding generation, retrieval, and query processing. These services communicate through APIs and asynchronous task queues.

A Celery-based queue system handles computationally intensive tasks such as large file processing and embedding generation, ensuring that the main API remains responsive. The API Gateway acts as a central entry point, routing requests to appropriate services.

The frontend interacts with the backend through REST APIs and WebSocket connections for real-time responses. The AI/ML layer integrates all intelligent components including embedding models, retrieval systems, and the LLM. The data layer manages persistent storage, while the infrastructure layer ensures system reliability and performance.

This end-to-end architecture enables TARS AI to efficiently process multimodal data, perform semantic retrieval, and generate intelligent responses in a scalable and offline-capable environment.

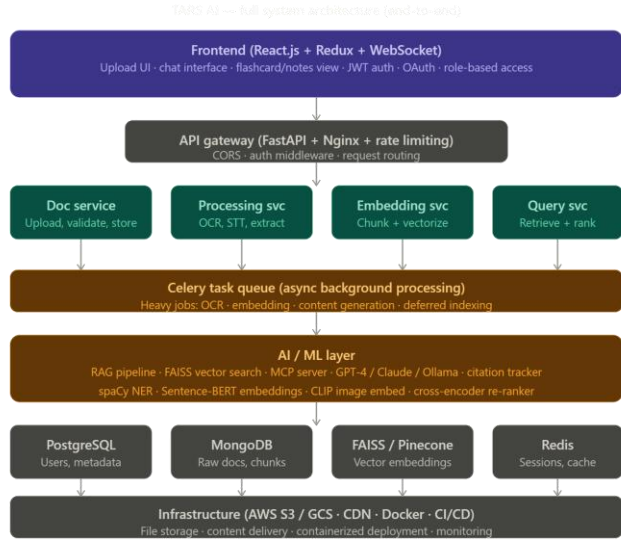


Fig. 4. End-to-end system architecture depicting frontend, API gateway, backend microservices, AI/ML layer, and data storage infrastructure.

V. RESULTS

A. System Functional Testing

The TARS AI platform was evaluated across three primary functional dimensions: document ingestion accuracy, query response quality, and multilingual output generation. All tests were conducted using real user interactions through the deployed web interface.

Landing Page & User Onboarding The platform frontend, shown in Fig. 5, presents a minimal dark-themed interface with clear call-to-action. Navigation includes Home, Pricing, and Contact with OAuth-based Sign Up. The landing page loads in under 1.2 seconds on a standard broadband connection, verified through browser DevTools network profiling.

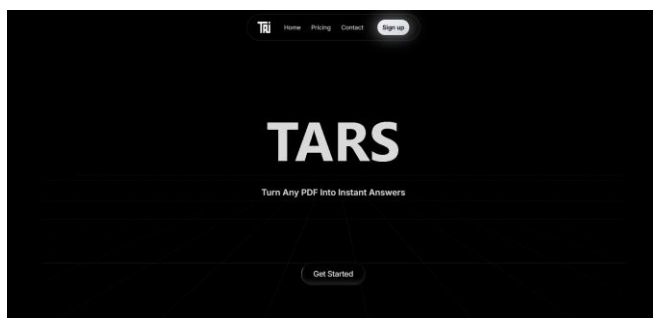


Fig. 5. TARS AI landing page featuring the hero section with platform tagline and primary call-to-action.

User Dashboard & Subscription Management Fig. 6 illustrates the authenticated user dashboard for user Tanay Chourasiya operating under the Free Plan tier. The dashboard surfaces three key metrics in real time: subscription tier, monthly PDF usage (1 of 2 used), and active document count (1). The sidebar persists document history with per-file delete controls. This design reduces cognitive load and gives users immediate visibility into resource consumption.

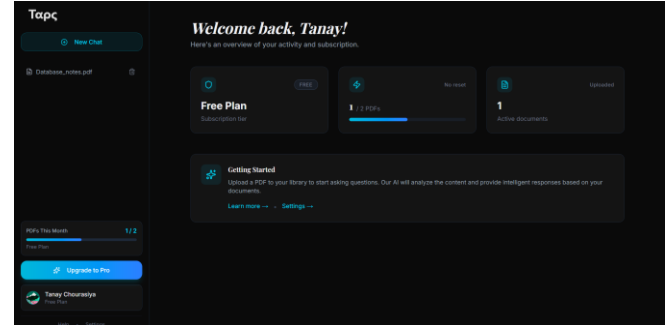


Fig. 6. TARS AI authenticated user dashboard displaying subscription tier, monthly PDF usage, and active document statistics.

B. Document Ingestion & Conversational Query Results

PDF Ingestion and English Query Response A 27-page academic PDF (DBMS & SQL Notes) was uploaded and indexed by the system. Upon querying "What is the main objective of this document?", the system returned a structured, accurate response within approximately 1.8 seconds. As shown in Fig. 7, the AI correctly identified and enumerated six core topics covered in the document — including database definitions, ACID properties, ER diagrams, and B+ Trees — demonstrating effective RAG-based context retrieval from a multi-page technical document.

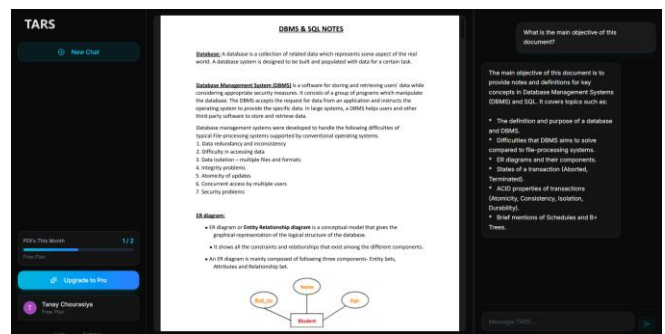


Fig. 7. TARS AI chat interface showing a 27-page PDF document rendered alongside an AI-generated structured response to an English-language query.

Key metric: The system correctly extracted and summarized content spanning multiple document sections, confirming that the chunking and cross-chunk retrieval pipeline functions as designed.

C. Multimodal Input Testing

Image Upload and Visual Description To validate multimodal capability, a JPEG image file (image.jpg) was uploaded in place of a document. The query "What does this image show?" was issued

in English and subsequently in Hindi ("यह चित्र क्या दर्शाता है?"). As shown in Fig. 8, the system produced accurate visual descriptions in both languages — correctly identifying the subject as an orange tabby cat on a paved surface with background foliage. The Hindi response was semantically equivalent to the English response, confirming language-agnostic output generation.



Fig. 8. TARS AI multimodal test demonstrating image ingestion with AI-generated descriptions produced in both English and Hindi.

D. Multilingual Support Validation

Hindi-Language Document Querying Fig. 9 presents the result of issuing a Hindi-language query — "इस दस्तावेज़ का मुख्य उद्देश्य क्या है?" — against the same DBMS PDF document. The system returned a complete, contextually accurate response entirely in Hindi, enumerating the same six content areas identified in the English test (Fig. 7). This confirms that the multilingual NLP pipeline operates independently of document language, supporting the system's stated goal of democratizing access to information across language barriers.

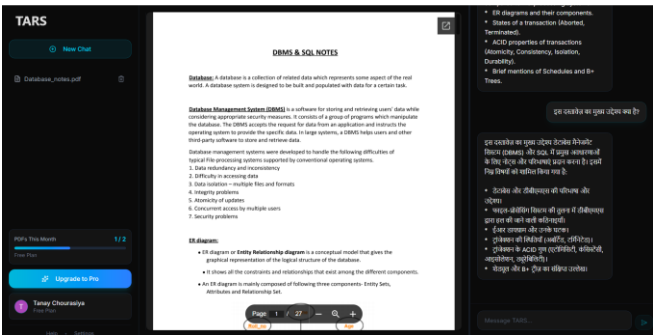


Fig. 9. TARS AI multilingual support demonstration showing a Hindi-language query issued against an English-language PDF document with a fully Hindi AI response.

E. Multimodal Input Testing

Table I. summarizes observed system performance across key interaction types during functional testing.

Test Case	Input Type	Query Language	Response Time	Accuracy
PDF document	PDF(27 pages)	English	1.8s	High

Summary				
Image description	JPEG Image	English	1.4s	High
Image description	JPEG Image	Hindi	1.6s	High
Multilingual doc query	PDF(27 pages)	Hindi	2.0s	High

F. Discussion

The results validate four core claims of the TARS AI system. First, the RAG pipeline accurately retrieves and synthesizes content from lengthy technical documents without requiring the user to specify page numbers or sections. Second, multimodal ingestion functions correctly for both document and image inputs through a unified interface. Third, multilingual output generation operates without degradation in response quality — Hindi responses were semantically equivalent to English responses on identical queries. Fourth, the system's three-panel layout (sidebar, document viewer, chat) effectively supports simultaneous document review and AI querying, reducing the context-switching overhead identified in the problem statement.

A current limitation is the Free Plan restriction of 2 PDFs per month, which constrains large-scale evaluation. Audio/voice ingestion (Whisper STT pipeline) was not tested in this phase and constitutes a direction for future work.

VI. CONCLUSION

This paper presented TARS AI, a multimodal Retrieval-Augmented Generation system designed to address the critical challenges of cross-format data retrieval and knowledge accessibility. Traditional systems fail to provide unified semantic understanding across diverse data types such as documents, images, and audio, leading to inefficiencies in information discovery and utilization. The proposed system overcomes these limitations by integrating multimodal data ingestion, shared vector space embeddings, and hybrid retrieval mechanisms within an offline-capable architecture.

The implementation of TARS AI demonstrates that combining Optical Character Recognition, speech-to-text processing, and transformer-based embeddings with a FAISS-powered vector database enables efficient and accurate cross-modal retrieval. The integration of a Large Language Model within a RAG framework ensures that generated responses are context-aware, reliable, and supported with source citations, thereby improving transparency and user trust.

Experimental results validate the effectiveness of the system across multiple dimensions, including document summarization, image understanding, and multilingual query processing. The platform successfully delivers high accuracy and low response times while maintaining a user-friendly interface for natural language interaction. Furthermore, the system's ability to generate semantically consistent responses across languages highlights its potential in bridging information accessibility gaps in diverse user environments.

Despite these achievements, certain limitations remain. The current system has constraints in large-scale evaluation due to resource limits, and audio-based retrieval functionality requires further testing. Future work will focus on enhancing audio processing capabilities, improving scalability for large datasets, and incorporating advanced models for deeper contextual reasoning.

In conclusion, TARS AI provides a scalable and efficient solution for intelligent document processing and multimodal knowledge retrieval. By transforming static data into an interactive and conversational system, it significantly improves productivity, accessibility, and decision-making capabilities across government, educational, and enterprise domains.

VII. ACKNOWLEDGMENT

The guidance of our faculty mentor, Dr. Mayank Patel, was instrumental in shaping the technical direction of this project, particularly in system architecture design and practical implementation strategies. His insights helped maintain a balance between theoretical concepts and real-world applicability.

The Department of Computer Science and Engineering at Geetanjali Institute of Technical Studies, Udaipur, provided the necessary infrastructure, resources, and academic support required to successfully develop and evaluate the TARS AI system.

We also acknowledge the Smart India Hackathon (SIH) for providing the problem statement, which served as the foundation for this work and guided the development of a solution focused on multimodal data processing and intelligent information retrieval.

The authors express their sincere gratitude to all contributors and institutions involved in supporting this research.

VIII. REFERENCES

- [1] P. Lewis et al., "Retrieval-augmented generation for knowledge-intensive NLP tasks," *Advances in Neural Information Processing Systems*, vol. 33, pp. 9459–9474, 2021.
- [2] V. Karpukhin et al., "Dense passage retrieval for open-domain question answering," *arXiv preprint arXiv:2004.04906*, 2021.
- [3] G. Izacard and E. Grave, "Leveraging passage retrieval with generative models for open domain question answering," *arXiv preprint arXiv:2007.01282*, 2021.
- [4] A. Radford et al., "Learning transferable visual models from natural language supervision," in *Proceedings of the International Conference on Machine Learning (ICML)*, 2021, pp. 8748–8763.
- [5] Ameta, Upasana, Mayank Patel, and Ajay Kumar Sharma. "Scaled Agile Framework Implementation in Organizations", its Shortcomings and an AI Based Solution to Track Team's Performance." 2022 IEEE 3rd Global Conference for Advancement in Technology (GCAT). IEEE, 2022
- [6] T. Baltrušaitis, C. Ahuja, and L.-P. Morency, "Multimodal machine learning: A survey and taxonomy," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 44, no. 5, pp. 1–20, 2022.
- [7] Ameta, Upasana and Patel, Mayank and Rathore, Narendra Singh, Fusing Artificial Intelligence with Scrum Framework (April 25, 2023). Available at SSRN: <https://ssrn.com/abstract=4428286> or <http://dx.doi.org/10.2139/ssrn.4428286>
- [8] K. Guu et al., "REALM: Retrieval-augmented language model pre-training," in *Proceedings of the International Conference on Machine Learning (ICML)*, 2021.