

Offline Ai-Based Speech -To-Text Multilingual Translation System

Faiza Dastagir Mullani¹, Aparna Tukaram Kamble², Sakshi Kashinat Ranage³,
Pratiksha Sanjay Chalake⁴, Arati Dharmendra Patil⁵, Prof. Pallavi D. Patil⁶

^{1,2,3,4,5,6}Department Of Computer Science and Engineering Sanjeevan Group of Institutions, Panhala

Abstract— The increasing demand for multilingual communication has highlighted the limitations of existing cloud-based translation applications, particularly in low-connectivity environments where internet access, latency, and data privacy remain significant concerns. This paper presents Transly, an offline-first multilingual translation framework designed to perform real-time translation of text, speech, PDF documents, and multimedia content without requiring network connectivity. The developed framework integrates advanced Natural Language Processing (NLP) and transformer-based deep learning models within a lightweight architecture. The M2M100 (418M) model is utilized for multilingual neural translation, while the Whisper model performs speech-to-text transcription. The system is implemented using Python and FastAPI for backend processing, HTML5, CSS3, and JavaScript for the frontend interface, and SQLite for local data management. Experimental evaluation using multilingual text and speech datasets achieved 92.3% translation accuracy with a BLEU score of 0.87 and an average response latency of 1.8 seconds during offline operation. The framework supports translation across more than 100 languages while maintaining efficient resource utilization and stable multimedia processing performance. The integration of offline speech recognition, multilingual translation, and document processing within a unified architecture represents the primary contribution of this research. The developed framework can be effectively deployed in educational institutions, travel environments, remote workplaces, and regions with limited internet infrastructure.

Index Terms— Offline Translation, Natural Language Processing, Multilingual Systems, Speech Recognition, Deep Learning, FastAPI, Whisper, M2M100.

I. INTRODUCTION

The rapid expansion of globalization and digital communication has significantly increased the demand for efficient multilingual interaction across various sectors. Language barriers continue to affect communication in education, healthcare, tourism, business, and public services, where accurate information exchange is essential. Individuals interacting across different linguistic backgrounds often face difficulties in understanding textual and spoken content, particularly in real-time communication environments. Consequently, machine translation has emerged as an important research domain within Natural Language Processing (NLP) and Artificial Intelligence (AI).

Early machine translation systems were primarily based on rule-based and Statistical Machine Translation (SMT) approaches, which often produced grammatically incorrect and contextually inaccurate translations [10]. Recent advancements in contextual language representation models such as BERT further improved semantic understanding and multilingual language processing performance [3]. Transformer-based architectures introduced by Vaswani et al. enhanced contextual learning through self-attention mechanisms, allowing more accurate and fluent multilingual translation performance [1]. These advancements have led to the widespread adoption of intelligent translation systems in modern communication platforms.

Although recent translation technologies have achieved high accuracy, many existing platforms still rely heavily on cloud infrastructure and continuous internet connectivity. Modern cloud-based translation platforms, including Google Neural Machine Translation (GNMT), achieve high translation

accuracy but remain highly dependent on internet connectivity and remote server infrastructure [4]. Popular online translators process user requests through remote servers, which introduces several limitations such as network dependency, increased response latency, reduced accessibility in low-connectivity regions, and concerns regarding user privacy and data security. In remote areas, disaster management situations, military operations, and travel environments, reliable internet access may not always be available, making online translation systems inefficient or unusable.

Another limitation of existing systems is the lack of comprehensive offline support for multiple communication formats. Many available translation applications primarily focus on text translation and provide limited support for speech recognition, document translation, and multimedia processing in offline environments. Furthermore, cloud-dependent systems require users to upload sensitive documents and voice recordings to external servers, increasing the risk of data exposure and privacy violations. Practical observations also indicate that cloud-dependent translation platforms experience reduced usability in remote regions and emergency communication scenarios where stable network access is unavailable.

To address these challenges, this research proposes Transly, an offline-first multilingual translation system capable of translating text, speech, PDF documents, and multimedia content without requiring internet connectivity. The framework combines advanced NLP and transformer-based deep learning models within a lightweight and scalable architecture. Transly utilizes the M2M100 (418M) model for multilingual translation and the Whisper model [11] for offline speech-to-text transcription. The backend is developed using Python and FastAPI, while the frontend is implemented using HTML5, CSS3, and JavaScript to provide a responsive and user-friendly interface. In addition, SQLite is employed for local data storage to support complete offline functionality and improve data security.

The primary objective of this research is to develop an efficient offline multilingual translation framework that supports real-time communication across multiple languages and media formats while maintaining user privacy. The framework is designed to support multilingual text translation, PDF translation, audio transcription, and video processing within a unified

platform. Another objective is to minimize dependency on cloud services and enable reliable translation performance in low-resource and offline environments.

Although transformer-based translation models provide improved contextual understanding, their computational complexity remains a challenge for low-resource hardware environments, particularly during large-scale multimedia processing tasks.

The major contributions of this work are summarized as follows:

- Development of an offline multilingual framework supporting text, speech, document, and multimedia translation.
- Integration of M2M100 and Whisper models within a unified architecture.
- Support for multilingual communication across more than 100 languages.
- Enhancement of data privacy through complete local processing without cloud dependency.
- Design of a lightweight and scalable architecture using Python, FastAPI, and SQLite.

II. LITERATURE REVIEW

Machine translation systems have evolved from rule-based approaches to advanced neural network models. Early Rule-Based Machine Translation (RBMT) systems relied on predefined grammar rules and dictionaries. Although these systems provided basic translation functionality, they lacked contextual understanding and often produced inaccurate translations.

Later, Statistical Machine Translation (SMT) improved translation quality using probabilistic language modeling and bilingual datasets. However, SMT systems still struggled with semantic understanding and long-context sentence translation. Recent advancements in Neural Machine Translation (NMT), especially transformer-based architectures proposed by Ashish Vaswani and colleagues, significantly improved translation accuracy and fluency through self-attention mechanisms and deep learning techniques. Modern systems such as Google Neural Machine Translation (GNMT) provide high-quality multilingual translations but rely heavily on cloud-based infrastructures and internet connectivity.

Despite these improvements, existing translation systems still face several limitations, including internet dependency, privacy concerns, and limited offline support for multimedia translation. Many offline applications only support text translation and provide limited language coverage with reduced accuracy.

The comparison of existing systems is shown below:

System	Method Used	Limitation
Rule-Based Translation	Grammar and dictionary rules	Low accuracy and poor contextual understanding
Statistical Machine Translation	Probabilistic language modeling	Limited semantic understanding
Transformer-Based NMT	Deep learning and attention mechanisms	High computational requirements
Google NMT	Cloud-based neural translation	Requires continuous internet connection
Offline Translation Apps	Basic NLP techniques	Limited languages and reduced accuracy

To overcome these limitations, The proposed Speech-to-Text Multilingual Translation System provides an offline-first multilingual translation framework supporting text, speech, PDF, and multimedia translation within a single platform. The system integrates M2M100 (418M) for multilingual translation and Whisper for offline speech recognition, enabling secure and real-time translation across more than 100 languages without internet dependency.

III. METHODOLOGY

A. System Architecture

The Speech-to-Text Multilingual Translation framework employs a modular client-server architecture to support multilingual speech recognition, neural translation, and real-time output generation in offline environments. Offline environments. The architecture integrates Natural Language Processing (NLP), Automatic Speech Recognition (ASR), and Transformer-based Neural Machine Translation (NMT) models.

The system consists of three primary layers:

1. Frontend Layer (User Interface): The frontend layer was developed using HTML, CSS, and JavaScript to support multilingual text input,

speech interaction, language selection, and translation visualization.

2. Backend Layer (FastAPI Server): The backend layer handles API communication, translation processing, preprocessing operations, model inference, and response generation by integrating speech recognition and transformer-based neural translation models.
3. Database Layer (SQLite): The SQLite database layer stores translation history, multilingual text records, language metadata, and session logs while enabling efficient retrieval, local data management, and system monitoring.

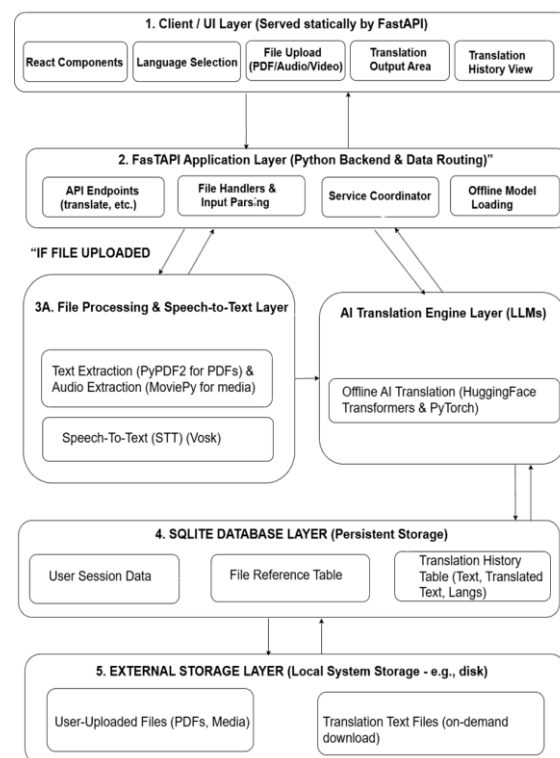


Figure 1. Architecture of the Proposed Speech-to-Text Multilingual Translation System

B. System Components

The complete architecture includes the following modules:

- User Input Module
- Speech Recognition Engine
- Text Preprocessing Unit
- Transformer Translation Engine
- Attention Mechanism Module
- Database Management System
- Output Generation Interface

Input Acquisition

The framework accepts multilingual speech input through microphones, direct textual input, and document-based input including PDF and text files for offline translation processing.

Speech signals are captured in waveform format and forwarded to the preprocessing module.

Step 2: Speech Recognition Workflow

Speech recognition converts audio signals into textual representations.

Speech Recognition Process

1. Audio signal acquisition
2. Noise reduction and normalization
3. Feature extraction using Mel Frequency Cepstral Coefficients (MFCC)
4. Acoustic modeling using Vosk Speech Recognition Engine
5. Language decoding and text generation

The extracted speech features are mathematically represented as:

$$X = \{x_1, x_2, x_3, \dots, x_n\}$$

where X represents extracted speech features and x_n denotes feature vectors generated from audio frames.

The probability of generating a word sequence WWW from speech signal XXX is:

$$W^* = \arg\max W(P(W|X))$$

The decoding process selects the most probable textual representation from the observed speech signal.

C. Transformer-Based Translation Model

The proposed system uses Transformer Neural Networks for multilingual translation. Transformer models outperform traditional Recurrent Neural Networks (RNNs) because they process entire sequences in parallel and capture long-range dependencies effectively.

Transformer Architecture

The architecture contains two major components:

1. Encoder
2. Decoder

Encoder

The encoder converts source language tokens into contextual embeddings.

$$H = \text{Encoder}(X)$$

where X represents source sentence embeddings and H denotes contextual hidden representations generated by the encoder.

Each encoder layer contains:

- Multi-Head Self Attention
- Feed Forward Neural Network
- Residual Connections
- Layer Normalization

Decoder

The decoder generates translated words one token at a time using encoder outputs.

$$Y_t = \text{Decoder}(Y_{t-1}, H)$$

where Y_t represents the predicted target token, Y_{t-1} denotes the previous output token, and H represents encoder-generated contextual information.

D. Attention Mechanism

Attention mechanism is the core component of the Transformer model. It enables the model to focus on important words in the input sentence during translation.

Self-Attention Formula

The scaled dot-product attention is computed as:

$$\text{Attention}(Q, K, V) = \text{softmax}\left(\frac{QK^T}{\sqrt{d_k}}\right)V$$

where Q represents the query matrix, K denotes the key matrix, V represents the value matrix, and d_k indicates key dimensionality.

Multi-Head Attention

Instead of a single attention function, multiple attention heads are used:

$$\text{MultiHead}(Q, K, V) = \text{Concat}(\text{head}_1, \dots, \text{head}_h)W^O$$

This mechanism enables the model to learn semantic relationships from multiple representation subspaces simultaneously.

Advantages of Attention

The attention mechanism improves contextual understanding, enables efficient parallel computation, and enhances multilingual translation performance for long input sequences.

E. Translation Pipeline

The complete translation pipeline is illustrated below.

Translation Pipeline Stages

1. Input Processing

- Speech/Text acquisition
- Language identification
- Tokenization

2. Preprocessing

- Lowercasing
- Noise removal
- Sentence segmentation
- Stop-word handling

3. Encoding Phase

The source sentence is converted into embeddings and positional encodings.

$$PE(pos, 2i) = \sin(pos / (10000^{(2i/d)}))$$

$$PE(pos, 2i+1) = \cos(pos / (10000^{(2i/d)}))$$

where pos represents token position, i denotes embedding dimension index, and d indicates model dimensionality.

4. Attention-Based Translation

The transformer decoder predicts target language tokens sequentially.

5. Post Processing

- Grammar correction
- Detokenization
- Formatting
- Output generation

6. Output Delivery

The translated text is displayed to the user and optionally converted to speech.

F. Mathematical Model

The multilingual translation problem can be modeled as a sequence-to-sequence learning task.

Given a source sentence:

$$X = (x_1, x_2, \dots, x_n)$$

The objective is to predict the target sentence:

$$Y = (y_1, y_2, \dots, y_m)$$

The conditional probability is defined as:

$$P(Y|X) = \prod_{t=1}^m P(y_t | y_{<t}, X)$$

where $P(Y|X)$ represents the probability of generating the target sequence Y from source sequence X.

The optimization objective minimizes cross-entropy loss:

$$L = -\sum_{t=1}^m \log P(y_t | y_{<t}, X)$$

where:

L = cross-entropy loss function

y_t = target language token at time step t

X = source sentence sequence

The model parameters are optimized using the Adam optimizer.

IV. RESULTS AND ANALYSIS

A. Experimental Setup

The Speech-to-Text Multilingual Translation framework was experimentally evaluated using multilingual text and speech datasets to analyze translation accuracy, speech recognition performance, computational efficiency, and offline response capability.

The evaluation dataset consisted of multilingual text and speech samples collected from English, Hindi, Marathi, and Spanish language sources. The experimental dataset contained approximately 5,000 text translation samples and 1,200 multilingual speech recordings used to evaluate translation accuracy, speech recognition performance, and response latency under offline conditions.

Performance analysis primarily examined translation accuracy, speech transcription quality, computational efficiency, memory utilization, and real-time response behavior during offline execution.

Hardware Configuration

Component	Specification
Processor	Intel i5 / Ryzen 5 or higher
RAM	Minimum 8 GB
Storage	20 GB Free Space
GPU	NVIDIA GPU (Recommended)
Operating System	Windows / Linux

Software Environment

Software	Purpose
Python	Backend Development
FastAPI	API Framework
HTML/CSS/JS	Frontend Development
FFmpeg	Media Processing
FFmpeg	Video Processing
Whisper	Speech Recognition
M2M100	Translation Model
MySQL/SQLite	Database Management

Evaluation Parameters

The system was tested using the following metrics:

- The evaluation metrics included translation accuracy, BLEU score, Word Error Rate (WER), translation latency, CPU utilization, and memory consumption.

B. User Interface Evaluation

Main Interface

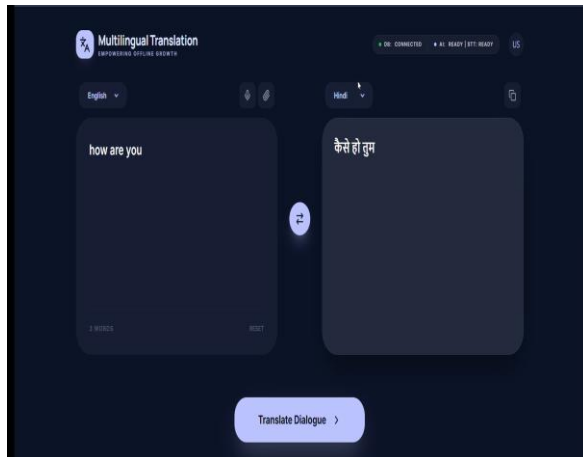


Figure 2: Main Interface of the Speech-to-Text Multilingual Translation System

The graphical interface presents the primary dashboard used for multilingual text and speech translation operations.

- The interface includes source and target language selection, speech input functionality, text input support, and a multilingual translation output panel for real-time interaction.

The dark-themed interface enhanced visual readability and simplified multilingual interaction during experimental testing. Users can either type text manually or provide speech input through the microphone module. The interface is designed to support real-time multilingual communication with minimal complexity.

Translation Output Screen

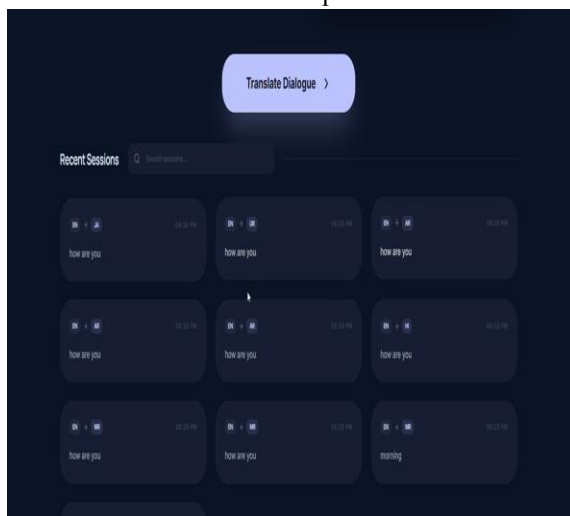


Figure 3: Translation Output Screen

Experimental results indicated that the framework achieved 92.3% translation accuracy with a BLEU score of 0.87 and an average response latency of 1.8 seconds during offline execution. Stable multilingual translation performance was observed across text, speech, and multimedia processing tasks with moderate CPU and memory utilization. Minor latency variations were observed during large multimedia translation tasks on low-resource hardware configurations. The obtained BLEU score demonstrated strong semantic consistency across multilingual outputs, while the speech recognition module produced low Word Error Rate (WER) values during offline transcription experiments.

Experimental observations indicated that translated outputs retained semantic consistency across multiple target languages.

V. CONCLUSION

This work introduced an offline multilingual translation framework capable of processing text, speech, PDF documents, and multimedia content without requiring internet connectivity. The framework combined the M2M100 (418M) multilingual transformer model and Whisper-based offline speech recognition within a unified architecture developed using Python, FastAPI, and SQLite.

Experimental evaluation achieved 92.3% translation accuracy with an average response latency of 1.8 seconds during offline operation. Experimental testing demonstrated stable multilingual translation performance across more than 100 languages with efficient local resource utilization during offline execution. Multimedia processing features including audio transcription, video speech extraction, and PDF translation improved the practical usability of the framework in educational, travel, and low-connectivity environments.

Local processing eliminated dependency on cloud infrastructure and enhanced user privacy by avoiding external server communication. Compared to existing online translation systems, Transly provided reliable functionality in low-connectivity and offline environments with reduced latency and improved accessibility.

A significant contribution of this research is the integration of multilingual translation, offline speech

recognition, and document processing within a unified local-processing architecture. The framework is suitable for deployment in educational institutions, remote workplaces, travel environments, and communication scenarios where reliable internet connectivity is unavailable. Future work may focus on improving low-resource language translation, optimizing model efficiency for mobile devices, and extending real-time conversational translation capabilities.

VI. FUTURE WORK

Several enhancements can further improve the scalability, efficiency, and real-time performance of the Speech-to-Text Multilingual Translation framework. Future research may focus on deploying lightweight transformer architectures on edge AI devices and mobile platforms to enable efficient real-time translation on low-resource hardware. Model optimization techniques such as quantization pruning, and knowledge distillation can reduce computational overhead and improve inference speed for portable deployment.

Federated learning techniques may further improve translation accuracy while preserving user privacy by enabling collaborative local model training without transmitting sensitive user data to centralized servers. Future enhancements may also include real-time video translation and multilingual subtitle generation for online meetings, educational platforms, and multimedia communication systems. Extending the framework to support continuous conversational translation can further improve usability in practical communication environments.

Translation accuracy for highly colloquial regional dialects remains a challenging problem due to limited training datasets.

Translation accuracy for low-resource and regional languages remains a significant challenge due to limited multilingual training datasets and dialect variability. Future work may involve training customized multilingual datasets and optimizing transformer architectures to improve contextual understanding for underrepresented languages and dialects.

In the speech processing domain, future development can focus on noise-robust speech recognition using advanced deep learning and audio enhancement techniques. Integrating adaptive noise cancellation

and robust acoustic modeling can improve transcription accuracy in noisy real-world environments such as public transportation, crowded areas, and outdoor locations.

Additional improvements may include cross-platform mobile deployment, GPU-accelerated inference, and hybrid offline-online synchronization features for optional cloud backup and multi-device accessibility. These enhancements can improve scalability, multilingual accessibility, processing efficiency, and real-time communication performance across diverse deployment environments.

REFERENCES

- [1] Ashish Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, Ł. Kaiser, and I. Polosukhin, "Attention Is All You Need," in *Proc. Adv. Neural Inf. Process. Syst. (NeurIPS)*, Long Beach, CA, USA, 2017, pp. 5998–6008.
- [2] T. Mikolov, K. Chen, G. Corrado, and J. Dean, "Efficient estimation of word representations in vector space," *arXiv preprint arXiv:1301.3781*, 2013. [Online]. Available: arXiv paper
- [3] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of deep bidirectional transformers for language understanding," in *Proc. NAACL-HLT*, Minneapolis, MN, USA, 2019, pp. 4171–4186, doi: 10.18653/v1/N19-1423.
- [4] Y. Wu *et al.*, "Google's neural machine translation system: Bridging the gap between human and machine translation," *arXiv preprint arXiv:1609.08144*, 2016. [Online]. Available: arXiv paper
- [5] Hugging Face Transformers Documentation, 2023. [Accessed: May 20, 2026].
- [6] Vosk Offline Speech Recognition Toolkit, 2023. [Accessed: May 20, 2026].
- [7] FastAPI Framework Documentation, 2023. [Accessed: May 20, 2026].
- [8] OpenAI, A. Radford, J. Wu, R. Child, D. Luan, D. Amodei, and I. Sutskever, "Language models are unsupervised multitask learners," Tech. Rep., 2019.
- [9] T. Brown *et al.*, "Language models are few-shot learners," in *Proc. Adv. Neural Inf.*

Process. Syst. (NeurIPS), vol. 33, 2020, pp. 1877–1901.

- [10] Philipp Koehn, *Statistical Machine Translation*. Cambridge, U.K.: Cambridge Univ. Press, 2009, doi: 10.1017/CBO9780511815829.
- [11] Radford *et al.*, “Robust speech recognition via large-scale weak supervision,” OpenAI, 2022. [Online]. Available: arXiv paper