

A Sentiment and Context-Aware Machine Translation Framework for Polite English-to-Tamil Conversion of Offensive Language with Cultural Sensitivity Consideration

Rama S^{1,*} and R. Mythili²

¹Department of Computer Science and Engineering, SRM Institute of Science and Technology, Ramapuram, Chennai, India

²Department of Information Technology, SRM Institute of Science and Technology, Ramapuram, Chennai, India

Email: rama.sugavanam@gmail.com (R.S.); mythilir2@srmist.edu.in (R.M.)

*Corresponding author

Abstract—Neural Machine Translation (NMT) involves the automated conversion of text from one language into another, aiming for linguistic accuracy and fluency. The domain of machine-based text translation and sentence matching undergoes rapid transformation due to swift developments in deep learning methods. However, training NMT systems demand extensive parallel datasets. For languages with scarce data resources, producing high-quality translations remain a challenge, often resulting in less reliable outputs. Here, a politeness and sentiment-aware English-Tamil NMT module has been implemented. Initially, English text is gathered from the data sources and undergoes multitask classification with the developed Robustly Optimized Bidirectional Encoder Representations from Transformers Pretraining Approach-Residual Long Short-Term Memory (RoBERTa-ResLSTM) network, and the module finds the sentiment class and offensive class from the text. This module maps the English words into the compact word embedding space, and ResLSTM helps in finding the contextual semantics of the word. The resulting annotations related to the sentiment class and offensive class from the RoBERTa-ResLSTM are sent over to the politeness-controlled transfer Model. This module is built with the Multi-head Cross Attention-based mT5 (MC-AmT5) network. For better cross-handling of large multiple-language texts and sentence sequences, the proposed MC-AmT5 undergo hyperparameter tuning with the newly developed Modified Motion Vector-based Tornado optimizer with Coriolis force (MMV-MToC). The outcome from the politeness-controlled transfer model translated the English text into an equivalent Tamil text with better matching accuracy. For showcasing the effectiveness of the module, the model is experimented on standard datasets and compared against similar models.

Keywords—sentiment and context-aware machine translation, Robustly Optimized Bidirectional Encoder Representations from Transformers Pretraining Approach (RoBERTa), Residual Long Short-Term Memory (ResLSTM), Multi-head Cross Attention-based mT5 network (MC-AmT5), Modified Motion Vector-based Tornado optimizer with Coriolis force (MMV-MToC)

I. INTRODUCTION

Politeness is recognized as a fundamental element in communication, along with other essential principles such as honesty, informativeness, relevance, and clarity [1]. It is viewed as a key human quality that influences how individuals interact with one another. People constantly make decisions about when and how to apply politeness in their conversations to ensure interactions remain smooth and effective [2]. With progress in technology, particularly in artificial intelligence and Natural Language Processing (NLP), there is an increasing demand to emulate human-like behaviours in machines [3].

This involves creating systems that not only assist users in completing tasks but also serve as companions by engaging in courteous and empathetic communication. Conversational agents, such as chatbots or virtual assistants, need to be sensitive to users' emotions and contextual situations [4]. This means they should be able to apologize when errors occur, express empathy towards users in distress, provide greetings, and acknowledge feedback warmly. These actions provide stronger connections with users [5]. While efforts are being made to develop systems capable of producing polite responses, a significant challenge remains in determining the appropriate moments for politeness and identifying which words or phrases effectively convey politeness [6]. This challenge is heightened by the fact that politeness varies greatly depending on factors like age, gender, cultural background, and social status, making it difficult to create universally suitable responses [7].

With the ongoing expansion of globalization and the swift advancements in information technology, machine translation has become an essential strategy for enabling cross-cultural communication [8]. Although significant progress has been made in machine translation through deep learning techniques, effectively assessing the quality of these translations continues to pose considerable

difficulties [9]. Conventional evaluation practices, which depend on human evaluators, tend to be time-consuming, resource-intensive, and subject to personal biases, resulting in variable outcomes [10]. Due to these drawbacks, there is an urgent need for automated evaluation methods that provide faster and more objective assessments of machine translation quality [11].

Evaluating translation accuracy has historically been complex because of the indicators involved in interpreting linguistic nuances and meanings [12]. To tackle this issue, experts introduced multiple automatic assessment strategies. Among these, approaches that leverage perceptual semantic context have garnered special interest. These techniques work by extracting and interpreting semantic content from both the source and translated texts while factoring in the overall contextual framework [13]. This enables a more accurate evaluation of translation fidelity. However, existing automated models still encounter limitations, such as an insufficient understanding of detailed semantic elements and challenges in handling complex contextual information [14].

Deep learning-based sentiment and context-aware machine translation enhance translation accuracy, particularly when handling sensitive or offensive material, by embedding sentiment analysis and contextual comprehension into the translation process [15]. This approach generally relies on models like transformers, which are highly effective at understanding long-distance dependencies and capturing subtle contextual details, allowing the system to generate translations that preserve the emotional tone and social nuances of the original text [16].

Offensive language often involves complex, context-sensitive expressions, making it challenging for the model to identify and translate them consistently without altering their intent or tone. Preserving the original sentiment, especially in emotionally intense or offensive content, poses a significant difficulty, as the model must accurately interpret and reproduce the underlying emotional nuances within the translation [17].

The main goal of the research work is euphemizing, which is attained with the help of the politeness-controlled transfer model. This model is suited for converting the impolite English to a culturally sensitive and polite Tamil translation. Here, the sentiment and the offensive classification are incorporated for performing the simple word-for-word translation, which enables the social appropriateness and cultural sensitivity during the translation process. Here, the multi-task classification module detects the offensive or impolite expression from the English language. To prevent cultural misunderstanding and prevent cultural harassment in Tamil, they are converted into euphemistic expressions so they can be easily aligned with the social expectations. Here, the incorporation of multi-head cross attention is used for preserving the translation adequacy, which helps to minimize the error and maintain the semantic alignment of the text. An integrated deep learning

framework is developed in this research work, and its key contributions are given below:

- An effective context-aware Neural Machine Translation (NMT) model is developed to enhance the quality of English-to-Tamil translations. The proposed scheme strengthens the ability to process context during translation by recognizing relationships and dependencies between sentences, resulting in more coherent and precise output. The proposed approach boosts the natural flow and uniformity of translations to achieve higher translation accuracy and better semantic alignment between source and target texts.
- A Robustly Optimized Bidirectional Encoder Representations from Transformers Pretraining Approach-Residual Long Short-Term Memory (RoBERTa-ResLSTM) model is effectively utilized by integrating RoBERTa's contextual language comprehension with the sequential processing capabilities of ResLSTM networks to carry out sentiment analysis and offence classification simultaneously. This multi-task learning framework allows shared representation learning across related tasks, enhancing both generalization and efficiency. The proposed scheme increases classification accuracy and robustness. Additionally, combining multiple tasks within a unified model boosts computational efficiency and lowers resource consumption.
- A Multi-head Cross Attention-based mT5 (MC-AmT5) model is a versatile multilingual transformer architecture designed for a wide range of text-based tasks, including translation. This method can understand context and generate text across different languages. The addition of a multi-head cross-attention mechanism simultaneously focuses on both the original input text and associated classification information. This allows the system to effectively integrate semantic details with guiding factors, producing translations that are both contextually accurate.
- A Modified Motion Vector-based Tornado optimizer with Coriolis force (MMV-MToC) strategy is implemented for fine-tuning the variables from mT5. This strategy of fine-tuning several parameters enables the mT5 model to more effectively adjust to particular task demands or specialized fields, thereby boosting translation precision and applicability. By simultaneously optimizing multiple goals, the MMV-MToC approach increases training effectiveness and promotes the development of common feature representations that transfer well across various tasks. This method leads to enhanced robustness and performance based on isolated parameter tuning, which ultimately produces a more reliable and superior translation outcome.

The overall layout of the proposed model is provided in the points below: Section II provides a review of existing research in machine translation, especially context-aware approaches. The detailed architecture of the context-aware machine translation system is

explained in Section III. Implementation details of simultaneous sentiment analysis and offensive language detection are specified in Section IV. The integration of cultural sensitivity aspects into translation decisions is explained in Section V. The presentation of the experimental setup and datasets used for evaluation is specified in Section VI. Finally, a quick summary is given in Section VII.

II. LITERATURE SURVEY

A. Related Works

Yang *et al.* [18] have presented an NMT system designed to tackle the intricate problem of precisely pairing and aligning texts in English translation. This approach introduced a novel MA-Transformer structure that combines multiple tiers of semantic feature extraction to enhance the accuracy of text matching during translation processes. Initially, the model employed the Continuous Bag of Words (CBOW) method to create and embed word representations that encapsulate key semantic meanings. The findings demonstrated that this framework delivered a more precise and dependable approach for text matching in English translation applications.

Tan and Wang [19] have proposed a novel approach for automatically assessing the quality of NMT outputs by analyzing semantic meaning and emotional nuances. A specialized Pre-trained Language Model (PLM) integrated sentiment vector representations capturing emotional states that was utilized to improve the precision and reliability of the scoring process, particularly in evaluating how well the translation conveyed both meaning and sentiment. Results indicated that incorporating semantic and emotional information significantly improved the effectiveness and efficiency of translation quality evaluation.

Wadud *et al.* [20] have designed a system named deep Bidirectional Encoder Representations from Transformers (BERT) aimed at identifying offensive comments on social media platforms, applicable to both single-language and multilingual posts. They integrated deep convolutional neural networks with BERT-based models to enhance the detection of hateful or harmful language. The system was evaluated on relevant datasets, where it achieved approximately 92% accuracy and outperformed existing algorithms. This demonstrated that Deep-BERT was a highly competent tool for automatically classifying offensive content across multiple languages and social media channels.

Priya *et al.* [21] have suggested politeness level recognition in multilingual conversations, even without training data in the target language. The proposed scheme employed a contrastive learning strategy called contrastive aligner, which promoted the creation of similar representations for aligned sentence pairs across different languages. Experimental evaluations showed that this technique performed well on several datasets, especially in capturing emotional dynamics on politeness

detection, surpassing prior approaches and receiving validation from human assessments.

Pillai and Arun [22] have analyzed social media comments authored in Malayalam, a Dravidian language. The proposed scheme applied deep learning methods to evaluate sentiment and identify offensive language. They extracted key features and fused them through Shannon entropy calculation and Recurrent Neural Networks (RNN). Subsequently, they conducted sentiment classification and offensive language detection using a Hierarchical Attention Network (HAN). This approach achieved a high level of accuracy of about 95.6% along with strong sensitivity and specificity, demonstrating its strong capability to recognize positive and negative sentiments.

Hlaing *et al.* [23] have created translation systems utilizing the transformer framework to facilitate bidirectional translation among Thai, Myanmar, and English languages. They augmented these models by integrating linguistic indicators, such as Part-of-Speech (POS) tags, into some or all words, aiming to enhance translation accuracy. Multiple models were constructed, like a simple transformer relying on word vectors for benchmarking, a multi-source transformer that processed both unaltered text and text annotated with POS information, and a shared multi-source transformer that exchanged data across inputs. Their experimental evaluations indicated that incorporating linguistic features led to better translation results, with the shared-multi-source model.

Priya *et al.* [24] have explored a combined approach named BERT-DGCN designed for interpreting dialogue exchanges. This technique employed BERT to embed each conversational turn's words, while a Dependency Graph Convolutional Network (DGCN) was used to model syntactic connections between words. By merging these components, the model gained a deeper understanding of conversation structure and meaning, which proved beneficial in tasks such as classification. Their findings demonstrated that this method notably raised the accuracy of identifying individual dialogue when compared to earlier approaches.

Ganganwar and Rajalakshmi [25] have introduced an innovative data augmentation strategy known as Multilingual Translation-based Data Augmentation Technique (MTDOT) to improve the recognition of offensive remarks in Tamil text. Since offensive data was often scarce or imbalanced, they adopted a back-translation process, translating Tamil comments into English or Malayalam and then back into Tamil to generate additional training data. This approach was implemented through two variations: a multi-level method, creating more diverse augmented samples, and a single-level method. Thus, the proposed MTDOT technique significantly enhanced model performance, raising the F1-Score by roughly 65% over the traditional Synthetic Minority Over-sampling Technique (SMOTE) method, thereby enabling more effective identification of offensive content.

A bias-aware sentiment analysis framework is developed for balancing the accuracy of the sentiment analysis process [26]. Here, the toxicity scores are converted into sentiment labels for attaining higher accurate results in the sentiment analysis process.

B. Problem Statement

NMT models defined recently have not addressed politeness-aware text translation and are limited to certain languages. Opting for NMT to translate the text into regional languages may have a profound impact on communication and help native speakers to understand the intent of popular languages. Recently, offensive text translation has been automated through deep learning techniques but faces significant hurdles in translating the text into more polite text.

Table I discusses the recent literature done in sentiment and context-aware text translation, and a few following research gaps can be listed as below:

- NMT requires the highest quality text to engage in automated translation, and the changes in ambiguities, slang, nouns, grammar, and technical terms may increase the difficulty in translation. Considering the offensive word statements, a direct translation of the

words can result in disruptive meaning, and this can be overcome by a joint approach of sentiment and offensive classification. Using high-level modules to classify the sentiment and offensive cues in the statement helps in improving the translation accuracy.

- Few studies have been done considering politeness-aware text translation, as new polite words have to be embedded into the translated text to increase readability. This research ensures the achievement of this task with the Multi-head Cross Attention module, which excels in high-level text decoding. To further enhance the translator's performance and quality, a tuning module with heuristics is developed, and overall matching accuracy is improved.
- Heuristics play a major role in deep learning tasks, as they increase the efficiency of the network by choosing the right parameters for the training. Here, a new heuristic is defined to increase the effectiveness of the politeness transfer module and improve the overall accuracy of the translation procedure, considering the complexity involved in regional languages.

TABLE I. FEATURES AND CHALLENGES OF RECENT NMT MODELS WITH SENTIMENT AND CONTEXT-AWARE TRANSLATION

Authors	Methodology	Features	Challenges	Proposed work
Yang [18]	MA-Transformer	Feature fusion ensures high precision. Text matching accuracy was improved through semantic feature extraction.	Applicable for small-scale translation models. Multi-lingual ability is not attained. The offensive content is not converted into polite words.	The social appropriateness is maintained using sentiment and an offensive content detection approach.
Tan and Wang [19]	PLM	Offers encoding large sentence sequences. Employs pre-training to avoid loss of important features.	The complexity of handling large datasets is high. Ensemble features require a large training time. The cultural and social norms are not incorporated into the prior models, so it is effective for handling politeness in the translation process.	The consideration of cultural sensitivity effectively transforms the culturally appropriate word, also modulating the politeness level for getting the appropriate Tamil equivalent word.
Wadud <i>et al.</i> [20]	Deep-BERT	Adopted for both mono and bilingual languages. Deals in the translation of offensive text.	Need to achieve reduced bias. Overfitting issues. Applicable only to small-scale datasets. These approaches are culturally insensitive.	The contextual adaptation is effectively maintained by this model because of the incorporation of cultural modulation, offensive detection and sentiment analysis process.
Priya <i>et al.</i> [21]	XeroPol	Analyzes the politeness factor in the statement. Challenges in finding the emotional features are overcome.	Even though at high efficacy, the model performs less under large statements. High training time. Under the low-resource setting, the social and cultural norms are not effectively followed by this model.	The incorporation of RoBERTa is used for attaining the translation in a contextually and culturally appropriate way.
Pillai and Arun [22]	HAN	Better feature representation and data training. Less complexity in training.	The presence of slang and grammatical errors can reduce the translation accuracy. The contextual adaptation of the limited data is not possible with this model.	The improved optimization proposed in this research work can better enable contextual adaptation on the limited data.
Hlaing <i>et al.</i> [23]	Edit-based Transformer model	Considers only the word vector for the translation. Requires low computational effort. Useful for translating low-resource languages.	Does not consider dependency features in word translations. It does not consider the social appropriateness for the direct translation.	The semantically correct translation is attained through this model with the help of the cultural norms.
Priya <i>et al.</i> [24]	BERT-DGCN	Employs phase extraction to increase translation accuracy. Finds different politeness factors in the conversation.	Multi-tasking increased computational resource requirement. Requires large, labelled data. This model sometimes overlooks the cultural nuances.	The proposed model has higher semantic matching accuracy since it does not overlook the cultural nuances.
Ganganwar and Rajalakshmi [25]	MTDOT	Data augmentation increases the efficiency of handling diverse data. Suitable for imbalanced datasets.	Less accurate in finding the offensive class. Only offensive language detection is possible by this model.	The consideration of politeness and sensitivity can provide better results in the machine translation process.

III. DESCRIPTION OF THE IMPLEMENTED CONTEXT-AWARE MACHINE TRANSLATION FRAMEWORK USING DEEP LEARNING FRAMEWORK

A. Architecture of the Implemented Context-Aware Machine Translation Framework

A robust NMT model has been designed to address the limitations inherent in traditional Convolutional Neural Network (CNN) and Recurrent Neural Network (RNN) methods. Conventional CNN models primarily concentrate on local patterns using fixed-size filters, which restrict their ability to capture connections between distant words within a sentence. Although RNNs process input in order and can theoretically handle long-range dependencies, they often encounter issues like vanishing

or exploding gradients that hinder their practical performance.

In contrast, the proposed model leverages the Transformer's self-attention mechanism, which evaluates the relevance of every word in the sequence simultaneously for each token. This enables the model to directly link words far apart, efficiently capturing long-distance contextual relationships without relying on sequential processing. The architecture of the implemented context-aware NMT is given in Fig. 1. The overall process involved in the proposed NMT is clearly illustrated in the figure below. Starting from the data collection, the subsequent process, such as tokenization, offence classification, NMT and performance evaluation, is clearly illustrated here.

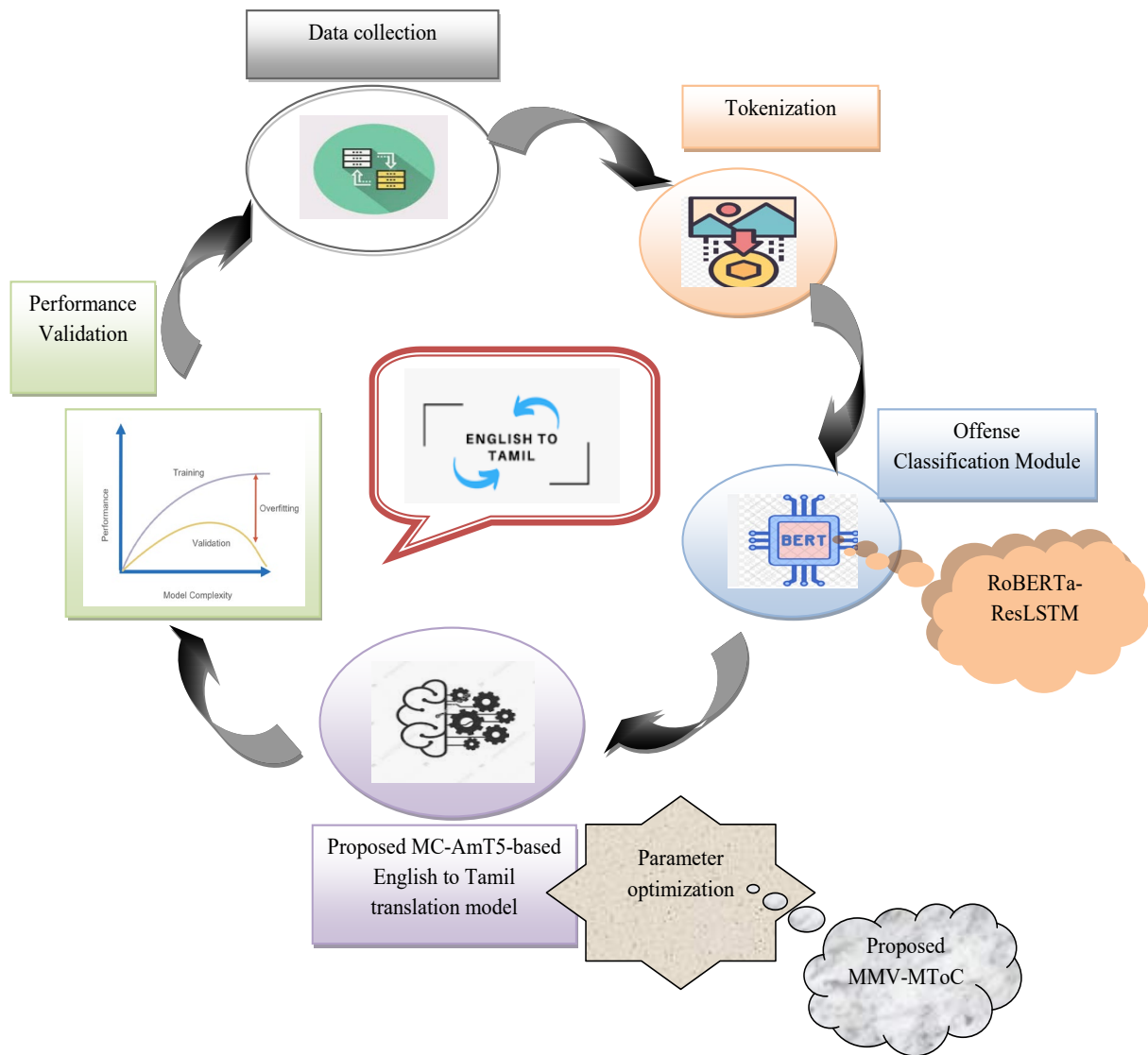


Fig. 1. Diagrammatic representation of implemented context-aware machine translation framework.

The context-sensitive machine translation system combines classification and translation components, which are refined using an advanced optimization technique. Initially, raw text data requiring both sentiment analysis and offence detection is processed.

The classification step is handled by a hybrid model integrating RoBERTa and ResLSTM. Proposed RoBERTa is a robust pre-trained transformer that captures rich contextual language features to grasp subtle textual meanings, while ResLSTM is a residual long

short-term memory network designed to validate sequential relationships and maintain information that prevents gradient disappearance.

The RoBERTa-ResLSTM module produces multi-label classification outcomes reflecting both the sentiment and offensive nature of the text, supplying crucial contextual information for the following translation stage. Subsequently, a politeness-controlled transfer model is employed to execute the translation task. The proposed MC-AmT5 model is a pre-trained text-to-text transformer adept at managing multiple languages and generating natural text. It incorporates a multi-head cross-attention mechanism that enables the model to simultaneously focus on several input sources, specifically, the original text and the classifications from the multitask module.

This design allows the seamless embedding of sentiment and offence context into the translation workflow, enabling style adjustments such as promoting politeness or reducing offensive expressions. The MC-AmT5 model is then optimally fine-tuned with the support of the MMV-MToC optimizer to effectively balance semantic integrity and politeness constraints derived from the classification inputs. This fine-tuning adapts the pre-trained model to the given domain and dataset, enhancing translation accuracy, contextual appropriateness, and adherence to politeness standards. Therefore, the proposed approach facilitates refined English-to-target language translation that considers sentiment and offensiveness cues to generate polite and contextually suitable outputs.

- Offensive language: Hurtful and disrespectful expressions are used in the language. The presence of hurtful words causes social discord and causes emotional harm to humans. The derogatory statements, hate speech and abusive marks are included in the offensive language.
- Operational criteria: The individuals are targeted by the presence of explicit and implicit insults within the sentence. In addition, the normative disagreement is exceeded when the expression conveys disrespect.
- Annotation guidelines: The semantic text is considered to classify the offensive text when it contains slurs and hateful content.
- Polite language: The respect, courtesy and feelings are demonstrated by the language. The socially appropriate expression for an agreeable interaction is possible with polite language.
- Operational criteria: The adoption of courteous phrases, which include sorry, thank you and please.
- Annotation guidelines: Based on the linguistic markers like hedging and culturally specific honorifics, the polite expressions are determined. If the sentence softens the request, the sentence is annotated as polite.
- Culturally sensitive: The cultural values, expectations and norms are aligned in this language. Here, the offensive and culturally inappropriate sentences are avoided.

- Operational criteria: While adapting the content to the cultural convention, the intended meaning of the translation and the expression is maintained.
- Annotation guidelines: The culturally sensitive sentence is annotated if they are aligned with the politeness of the target culture.

The distinction between profanity, abuse, hate speech, sarcasm, and impoliteness is listed in Table II.

TABLE II. DISTINCTION BETWEEN PROFANITY, ABUSE, HATE SPEECH, SARCASM, AND IMPOLITENESS

Terms	Features
Profanity	It consists of a specific set of taboo words that refer to bodily functions.
Abuse	The person is emotionally harmed because of the abusive words.
Hate speech	Based on the inherent characteristics, the set of groups is attacked or demeaned.
Sarcasm	The person is indirectly mocked by this word.
Impoliteness	The social norms are violated by this language.

B. Data License Details

This research uses the benchmark dataset for the offensive/polite classification. The relevant data needed for this proposed model are collected from the English-Tamil-parallel-sent dataset using the link <https://www.kaggle.com/datasets/hemanthkumar21/english-tamil-parallel-sent> accessed on 2025-05-16. This dataset serves as a crucial bilingual asset for NLP applications. It comprises matched sentence pairs in both English and Tamil, delivering corresponding texts across these languages. This dataset plays a key role in training and testing models for tasks like machine translation, which demands precise alignment between languages. Providing these parallel sentence pairs allows developers to create effective multilingual NLP systems capable of capturing the semantic and grammatical relationships between English and Tamil. The sample text collected from the above link is provided below. The details of the data annotations and ethics are given in Table III. The dataset consists of 62K English-Tamil sentence pairs, with careful coverage of code-mixing (18%) and dialectal variation (7%). Annotation was performed by trained bilingual annotators, with strong inter-annotator agreement ($\kappa = 0.82$, $\alpha = 0.78$). Ethical safeguards included anonymization of sensitive content, dialectal fairness checks, and informed annotator consent. The dataset is released under a CC BY-NC-SA 4.0 license, and document-level splitting was used to prevent data leakage.

Example 1: Polite Request.

English: Sorry for the inconvenience.

Tamil: உதவியில்லை என்பதை உணர்ந்துகொள்வதற்காக மன்னிக்கவும்.

Example 2: Offering Help.

English: Would you like some water?

Tamil: உங்களுக்கு சிறிது தண்ணீர் வேண்டுமா.

Example 3: Expressing Gratitude.

English: Thank you very much.

Tamil: மிகவும் நன்றி.

TABLE III. DETAILS OF THE DATA ANNOTATIONS AND STATISTICS

Details	Statistics
Source	https://www.kaggle.com/datasets/hemanthkumar21/englist-tamil-parallel-sent
Total Sentence Pairs	62,000 (English–Tamil)
Split Strategy (size per split)	Sentence-level split strategy: Train: 49,600 (80%), Validation: 6,200 (10%), Test: 6,200 (10%).
Domain Distribution	Subtitles (40%), News (25%), Social Media (20%), Conversational (15%)
Code-Mixing	18% of samples contain English–Tamil code-mixing.
Dialectal Coverage and demographics	7% samples include dialectal Tamil (Sri Lankan Tamil, Madurai, Chennai).
Annotation training	Sentiment (positive/neutral/negative), Offence (yes/no), Politeness (3 levels).
Annotators	5 trained bilingual annotators (linguistics background, age 18–45).
Inter-Annotator Agreement	Cohen's $\kappa = 0.82$; Krippendorff's $\alpha = 0.78$.
Bias Mitigation	Balanced dialect sampling; Anonymization of sensitive identity references.
Ethics & Licensing	Public/crowdsourced corpora, CC BY-NC-SA 4.0 License, Annotator consent.
Data Leakage Prevention	Document-level deduplication; Split at document (not sentence) level.

C. Developed MMV-MToC

In this proposed model, the MMV-MToC is necessary for fine-tuning the variables from the mT5 network during the classification process. The proposed MMV-MToC strategy follows the basic strategy of the TOC algorithm. It follows the natural process of tornado formation and dissipation, incorporating the Coriolis Effect to enhance its search capabilities compared to the AdamW + Bayesian/Optuna. Its basic mechanism involves simulating the lifecycle of windstorms progressing into tornadoes, which are represented as search agents exploring the solution space. The algorithm models the behaviours of these atmospheric phenomena mathematically, with parameters influencing exploration and exploitation phases. In addition, the multiple interdependent parameters are simultaneously tuned by the proposed MMV-MToC compared to the traditional AdamW + Bayesian/Optuna so that the accuracy and the politeness of the translation process are greatly improved.

The comprehensive search of the parameter space is possible with the incorporation of MMV-MToC's stochastic activity, which also prevents premature convergence. The challenges posed by the offensive language are also effectively handled by the MMV-MToC due to its intricate fine-tuning capacity. Here, more diverse atmospheric regions are enabled by the rapid changes in the direction and movements of the tornadoes. The performance of the MMV-MToC in the curved solution space is more effectively performed by modelling the tornado-like entities. The diverse and less linear search path is promoted by the search agents' movements that are induced by the Coriolis force, which occurs due to the rotation of the Earth that severely affecting the atmosphere of the ocean.

The TOC algorithm simplifies the Coriolis effect to a rotational force that guides the search agents in a curved trajectory, encouraging exploration of diverse regions. Yet, the TOC algorithm is influenced by complex natural forces and struggles with premature convergences. This challenge is addressed by modifying the random variable often used in velocity and position updates based on the current, best, worst, and mean fitness values of the population. By updating the random variable Y_r , the proposed MMV-MToC model adaptively controls the stochastic component, making the search more dynamic and context-aware. The modified concept is provided in Eq. (1).

$$Y_r = \frac{Q_{fit}}{(H_{fit} + L_{fit} + S_{fit})} \quad (1)$$

where, the terms Q_{fit} , H_{fit} , L_{fit} , and S_{fit} are current fitness, best fitness, worst fitness and mean fitness.

This adjustment ensures that the stochastic component vigorously responds to the population's current state, effectively preventing premature convergence and enhancing the search's robustness.

The Pseudocode of the proposed MMV-MToC model is given in Algorithm 1.

Algorithm 1: MMV-MToC

```

Initialize the population size  $D$  and the maximum number of
iterations  $V$ 
Initialize the random population  $S$ 
Acceleration coefficients like  $v_1$  and  $v_2$ 
Evaluate the internal weights.
The fitness function is calculated to evaluate agents.
  For each agent  $x = 1$  to  $S$ 
    Update the random variable  $Y_r$  adaptively based on fitness
    values.
    Update the position with deflection and evaluate its new
    fitness.
  End for
Return to the best position.
```

IV. CONTEXT-AWARE MULTI-TASK CLASSIFICATION FOR DETECTING SENTIMENT AND OFFENSIVE CLASSES USING BERT MODEL

A. Significance of Detecting Offensive Classes for Machine Translation

Identifying offensive categories in machine translation is crucial for maintaining the accuracy, security, and suitability of translated material. The significance of detecting offensive classes for machine translation is provided in the points below.

- Offensive language such as slurs, hate speech, and abusive expressions has different meanings and degrees of severity across languages and cultures; Detecting these types of offensive language helps to prevent cultural misunderstandings and the spread of harmful content in translations.
- By identifying and managing offensive material, machine translation systems provide safer and more

respectful outputs, enhancing user trust and acceptance, especially in sensitive or professional environments.

- Automatic detection also ensures compliance with legal regulations limiting offensive speech, reduces legal risks, and protects social well-being and the reputations of content providers and technology platforms, resulting in translations that are accurate, culturally appropriate, ethical, and socially responsible.

Thus, detecting offensive categories in machine translation is vital for producing precise, safe, and culturally appropriate translations. Proper recognition and handling of offensive language help to avoid cultural misinterpretations and the dissemination of damaging content. This functionality strengthens the ethical commitment and societal dependability of NMT systems, leads to improved user satisfaction and protects the credibility of all parties involved.

B. Proposed RoBERTa-ResLSTM Model

A RoBERTa-ResLSTM framework is designed for multi-task classification, which integrates the strengths of the RoBERTa transformer model with Res-LSTM networks to enhance multitask classification performance.

RoBERTa Module is a sophisticated transformer-based language model built upon BERT's foundation. It improves upon BERT by employing byte-level byte pair encoding for more efficient tokenization T_n with a reduced vocabulary W_n , which decreases computational load. Its dynamic masking technique enables the model to learn from diverse input sequences, leading to a better analysis of context. RoBERTa converts raw text into rich, contextualized embeddings that encode deep semantic and syntactic details for essential classification. RoBERTa integrates the residual network with LSTM to effectively learn long-term dependencies from the raw images.

ResLSTM enhances the traditional LSTM network by incorporating residual connections. These connections allow the model to effectively learn intricate temporal relationships by reducing issues like gradient explosion or vanishing. By adding the output of each residual layer back to its input, the model maintains a stable gradient flow and increases its capacity to process both immediate and distant sequential information.

Mechanism of Res-LSTM [27]: Residual links are shortcut pathways that allow outputs from earlier layers to be directly added to subsequent layers or states. This structure promotes gradient flow during training and helps to learn dependencies over extended sequences. Specifically, the residual connection adds the previous state's output directly to the current one G_c using Eq. (2).

$$G_c = \theta(E_{zx}zc + E_{kp}K_{c-1} + T_p) \quad (2)$$

where, the term E_{zx} and E_{kp} are the learnable weights applied to current and previous states, θ and is the activation function. The bias term is indicated as T_p .

The Res-LSTM contains several residual layers, each with 64 units, which enhances the network's ability to analyze detailed temporal relationships. Data passes through these layers sequentially, with residual links enabling information sharing across layers.

The core operations of Res-LSTM are forgetting gate, input gate, output gate, and cell state updates that are controlled by specific equations involving learnable parameters and activation functions. Forget gate (Ga_F) decides which information is discarded using Eq. (3).

$$Ga_F = \theta(E_{zx}Gac + E_{kp}GaK_{c-1} + T_f) \quad (3)$$

where the prior hidden state is denoted as K_{c-1} , and the bias is denoted as T_X .

The cell state (Sw_F) is updated by adjusting the cell state based on current inputs using Eq. (4).

$$Sw_F = \tanh(E_{zx}Swc + E_{kp}SwK_{c-1} + T_{Sw}) \quad (4)$$

The output gate (UP_o) is determined using Eq. (5).

$$UP_o = (E_{zx}Upc + E_{kp}UpK_{c-1} + T_o) \quad (5)$$

The hidden state (S_h) represents the output at each time step, using Eq. (6).

$$S_h = UP_o \otimes \tanh(w_F) \quad (6)$$

The time steps are represented as p .

Residual connections are embedded within its recurrent architecture to address the gradient-related issues. These residual pathways allow the model to more effectively learn both short-term and long-term sequential patterns, resulting in improved accuracy in applications such as human motion intention recognition. The architecture's structure ensures stable gradient propagation and enhances its capacity to understand complex temporal dynamics, which is mathematically expressed in Eq. (7).

$$Re = S_h + Y_{ip} \quad (7)$$

where the hidden state is denoted as S_h and the input is represented as Y_{ip} .

This combined approach enhances the model's efficiency in handling multiple related classification tasks simultaneously, particularly for sentiment analysis and offence categorization. The diagrammatic view of the proposed RoBERTa-ResLSTM model is specified in Fig. 2. The process of offensive text classification from the input text data using the suggested RoBERTa-ResLSTM is given in this figure. Initially, the input text data is given to the RoBERTa model, where the text data is divided into classes and tokens. After that, these tokens and classes are passed to the RoBERTa to get the hidden features. The attained hidden features are passed to the residual LSTM for the offensive classification.

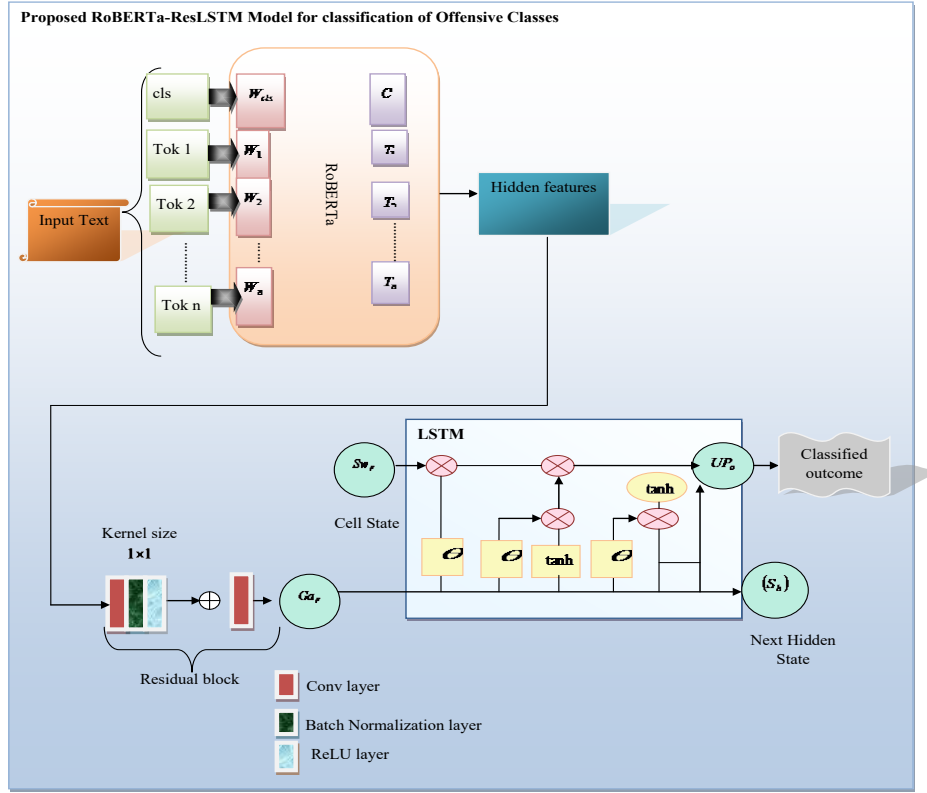


Fig. 2. Structural illustration of proposed RoBERTa-ResLSTM model.

C. Multi-Task Classification for Detecting Offensive Classes

Multi-task classification for detecting offensive classes' technique employs a unified strategy that concurrently conducts sentiment analysis and offence detection through a combined RoBERTa-ResLSTM

model. This method approaches sentiment evaluation and the identification of offensive content as linked tasks within a multi-task learning framework. By jointly training on both tasks, it capitalizes on shared linguistic and semantic information, enhancing the model's overall accuracy and robustness.

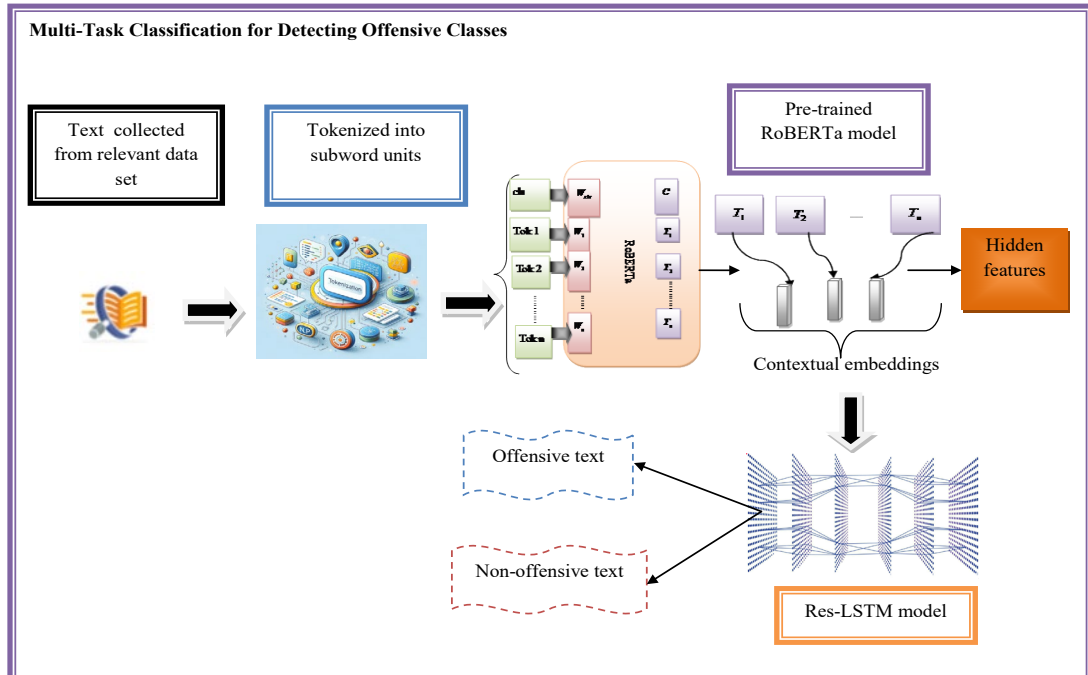


Fig. 3. Diagrammatic view of multi-task classification for detecting offensive classes.

The main network is RoBERTa, a highly effective pre-trained transformer-based language model, to produce rich contextual word representations. RoBERTa effectively captures intricate semantic and syntactic aspects of the text, improving the system's capacity to discern subtle language cues essential for both offensive language detection and sentiment assessment. Building upon these embeddings, a ResLSTM further processes the sequential data, maintaining contextual relationships and identifying temporal patterns within the text. The inclusion of residual connections addresses the vanishing gradient problem, facilitating deeper LSTM layers and more effective learning of long-distance dependencies. This design allows the model to simultaneously determine sentiment polarity and categorize text as offensive or non-offensive. Consequently, the multi-task learning system enables the model to better generalize by leveraging the interplay between sentiment and offence detection.

RoBERTa's contextualized embeddings offer a solid semantic base that surpasses traditional embedding methods. Altogether, this integrated approach enhances precision and stability in recognizing offensive language while also capturing its emotional context, which is vital for nuanced moderation in machine translation applications. Multi-task classification for detecting offensive classes is diagrammatically represented in Fig. 3. The operation of the multi-task classification is detailed in this figure. From the relevant dataset, the texts are collected, and they are tokenized into sub-words. These sub-words are passed to the RoBERTa for the generation of contextual embeddings. From the contextual embedding, the hidden features are attained, and they are given to the ResLSTM for the classification of offensive and non-offensive text.

V. POLITE ENGLISH-TO-TAMIL CONVERSION BASED ON CULTURAL SENSITIVITY CONSIDERATION USING TRANSFORMER MODEL

A. Description of Politeness Controlled Transfer Model

The politeness-controlled transfer model is an advanced translation framework designed to transform inappropriate English phrases into courteous and culturally appropriate Tamil counterparts. This is crucial because directly translating offensive language often overlooks cultural sensitivities, leading to results that may seem overly harsh or unsuitable for the intended audience. The process begins by examining the English input to detect expressions that might be offensive or impolite, using a blend of sentiment analysis and context-sensitive methods.

Sentiment analysis identifies the emotional tone, whether negative, disrespectful, or offensive, while context understanding prevents misinterpretation of words that could appear offensive outside their specific usage. Upon recognizing offensive content, the model avoids literal translation. Instead, it dynamically adjusts the politeness level in the Tamil rendition, ensuring that the output maintains both linguistic precision and social

appropriateness. For example, direct or blunt English statements are converted into Tamil expressions featuring polite forms, honorifics, or culturally fitting euphemisms. By effectively converting offensive English into respectful Tamil, this model tackles a key obstacle in cross-cultural translation: Preserving respect and preventing unintentional insults. This promotes more effective, empathetic communication between speakers of the two languages, encouraging mutual understanding and social cohesion.

Thus, the politeness-controlled transfer model combines sentiment detection, cultural insight, and a transfer translation mechanism to convert potentially offensive English content into polite Tamil, delivering translations that are precise and sensitive to the cultural norms of Tamil audiences, thus enabling respectful intercultural dialogue.

B. Multi-Head Cross Attention-Based mT5 Network

The proposed MC-AmT5 model is an extension of the T5 framework designed for multiple languages and utilizes a transformer-based sequence-to-sequence structure for various NLP applications, including text summarization and translation. In this proposed MC-AmT5 model, a multi-head cross-attention mechanism is integrated, which improves the model's capacity to selectively attend pertinent segments of the input data from multiple representational viewpoints.

Similar to T5, mT5 adopts a uniform way where all tasks are formulated as text-to-text. It performs a summarization process, where the input begins with the prompt to summarize the text for recognizing the task. The mT5 comprises an encoder and decoder with layers of multi-head cross-attention. This setup enables the model to effectively analyze dependencies over long sequences. Multi-head cross attention in mT5 allows the decoder to simultaneously attend to multiple aspects of the encoder's output, capturing vital information necessary for accurate text generation.

Operational Mechanism of the MC-AmT5 model: Initially, the Queries $R_q^{(s)}$ are obtained from the current decoder state. Keys $F_k^{(j)}$ and Values $H_v^{(j)}$ are extracted from the encoder output. For each attention head CA_q , the MC-AmT5 system calculates scaled dot-product attention using Eq. (8).

$$CA_q(R_q^{(s)}, F_k^{(j)}, H_v^{(j)}) = S_{mx} \left(\frac{R_q^{(s)} F_k^{(j)}}{\sqrt{b_k}} \right) H_v^{(j)} \quad (8)$$

$$CA_q(R_q^{(j)}, F_k^{(s)}, H_v^{(s)}) = S_{mx} \left(\frac{R_q^{(j)} F_k^{(s)}}{\sqrt{b_k}} \right) H_v^{(s)} \quad (9)$$

where, the term b_k is the dimension of the key vectors. The SoftMax activation is represented as S_{mx} . The dimension is represented as b_k .

This allows each head to focus on different parts of the input independently, identifying diverse features or

relationships. Outputs from all attention heads are concatenated and transformed through a linear layer, resulting in a comprehensive representation that informs the decoding process with varied insights. During text generation, the cross-attention mechanism helps the MC-AmT5 model pinpoint relevant input segments, such as specific sentences or phrases, when producing each word of the summary.

This precise focus ensures the output is both coherent and accurately aligned with the input content. Thus, the proposed MC-AmT5 model examines multiple aspects of relationships between different domains simultaneously, capturing delicate cultural and linguistic details necessary for creating polite and culturally appropriate English-

Tamil translations. This progressive attention method allows the system's comprehension of both language and cultural contexts, ultimately improving translation quality in terms of politeness and sensitivity.

The diagrammatic representation of the proposed MC-AmT5 model for the NMT is specified in Fig. 4. Here, the collected text data and the classified outcome are processed by the multi-cross attention module. After that, the attention scores between the keys, query and values are calculated, and the SoftMax operation is also done in this process. The multiple aspects of the input are simultaneously attended to by the multi-cross attention. This model combines the features to get culturally sensitive translations.

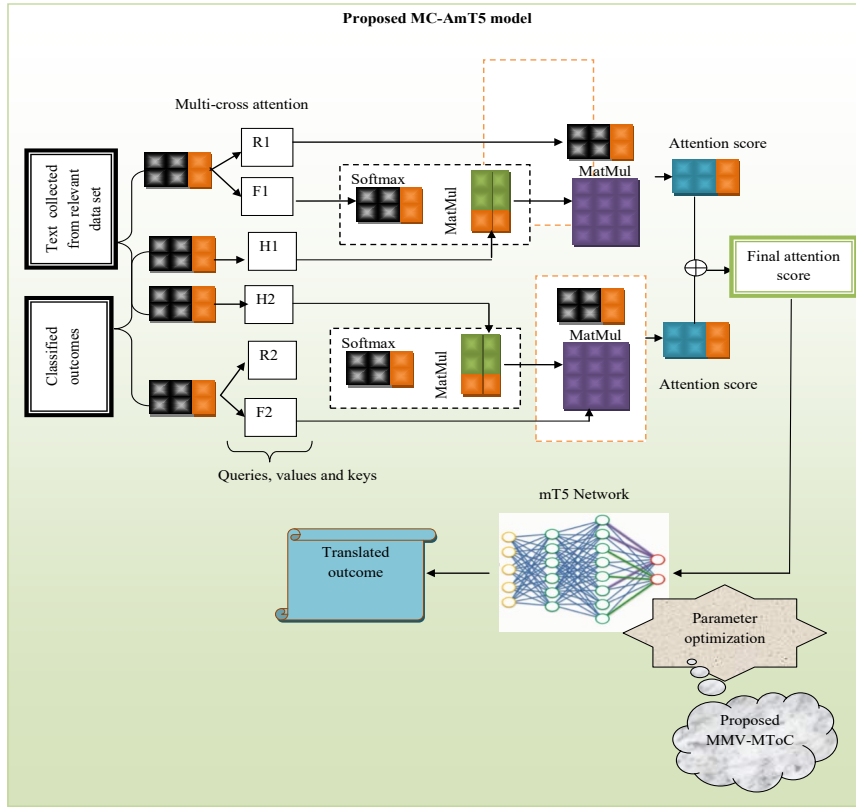


Fig. 4. Diagrammatic representation of proposed MC-AmT5 model.

C. Cultural Sensitivity Consideration-Based English-to-Tamil Conversion

The proposed approach translates English text into Tamil with a strong emphasis on maintaining cultural sensitivity and adhering to social conventions. It starts by utilizing modules for sentiment classification and offence detection, which automatically recognize the emotional tone and identify any potentially inappropriate content in the English input. By precisely determining sentiment polarity and marking offensive language, the proposed system acquires crucial contextual information to inform the translation. The proposed framework employs the MC-AmT5 model, which combines a multi-head cross-attention mechanism within the mT5 architecture. This mT5 model is a multilingual, pre-trained text-to-text transformer that is optimally fine-tuned to effectively

manage the linguistic complexities between English and Tamil.

The multi-head cross-attention allows the model to concurrently attend to multiple facets of the input, including sentiment signals and contextual nuances, improving the correspondence between the original and translated texts. By adaptively regulating politeness based on the results of sentiment and offence classification, the proposed MC-AmT5 model modifies word selection and grammatical structures in Tamil to produce translations that are courteous and culturally appropriate.

This approach ensures that any offensive expressions in English are rendered into polite and socially acceptable forms in Tamil, respecting the cultural values of the target community. The reliability of the proposed MC-AmT5 model is boosted by fine-tuning the variables of mT5 with the support of the MMV-MToC strategy. The

objective function helps to minimize the Sentence Error Rate (SER) and Word Error Rate (WER) of the system model, which is mathematically expressed in Eq. (10).

$$X = \arg \min_{\{NC_{kl}^{mT5}, RL_{kl}^{mT5}, FA_{pl}^{mT5}\}} (WER + SER) \quad (10)$$

where, the term NC_{kl}^{mT5} is the optimized hidden neuron count that varies in the range of $[5-225]$, RL_{kl}^{mT5} is the optimized learning rate that varies in the range of $[0.01-0.99]$, FA_{pl}^{mT5} is the optimized activation function that varies in the range of $[1-5]$.

The formula for calculating WER and SER are provided in Eqs. (11) and (12).

$$WER = \frac{A_S + F_D + K_I}{M} \quad (11)$$

where, the term A_S is the number of substitutions, F_D is the number of deletions, and K_I is the number of insertions of words in sentences. The total numbers of words in a sentence are indicated M .

$$SER = \frac{S_{EN}}{T_E} \quad (12)$$

where, the term S_{EN} is indicated as the number of sentences with an error, and T_E is the total number of sentences.

Thus, the proposed MC-AmT5 system performs translations while withstanding Tamil cultural traditions and societal norms, guaranteeing that the translated content is contextually appropriate for the intended audience with minimum WER and SER. Using sentiment analysis and offence detection components, the MC-AmT5 model effectively identifies the offensive language in the original text, supplying essential contextual information to inform translation choices.

By adaptively modulating politeness levels based on recognized sentiment and offensive elements, the MC-AmT5 system converts objectionable expressions into courteous and culturally accepted forms in Tamil, thereby improving the social suitability of the translated material. Hence, this method ensures that the translations reflect the cultural principles and expectations of Tamil users, promoting respectful and meaningful intercultural communication.

Conditioning of MT5 decoder using the classifier output: The offensiveness and the sentiments are detected by a multitask classifier RoBERTa-ResLSTM, proposed in this research work. Here, the tone awareness generation is guided by feeding the output of the multitask classifier to the MC-AmT5 decoder. In the classifier module named RoBERTa-ResLSTM, the English sentence $Y = [y_1, y_2, \dots, y_o]$ is provided as input. The output attained from the RoBERTa is expressed as $I_s \in \mathbb{R}^{o \times e_s}$ and here $e_s = 768$. Similarly, the results from

the ResLSTM are denoted $I_d \in \mathbb{R}^{o \times e_d}$, here $e_d = 512$. The pooled features are expressed by the term $i_{pool} = \text{Meanpool}(I_d) \in \mathbb{R}^{e_d}$. Further, two classification head outputs, namely sentiment class $z_t = \{pos, neut, nega\}$ and offensive classes, $z_q = \{offen, nonoffen\}$ are given as the input of the translation module for executing the conditioning process. The adapter fusion, prefix tuning and cross attention modulation strategies are used for the conditioning process.

Prefix tuning: For the sentiment and offensiveness, the learnable prompts are taken as Q_t and Q_p . The term $Q = \mathbb{R}^{m \times e_u}$ is taken for each prompt, and here the length of the prompt is taken e_u .

$$Q = \text{concat}(Q_t, Q_p) \in \mathbb{R}^{2m \times e_u} \quad (13)$$

At each layer, the above terms are prepended to the decoder output without modifying the weight of the model.

Cross attention: The key and value metrics from the encoder are represented as $L, W \in \mathbb{R}^{U_y \times e_y}$ and the pooled output from the classifier is expressed as $i_{pool} \in \mathbb{R}^{e_d}$. The expression for the conditioning vector is denoted in Eq. (14).

$$\begin{aligned} L' &= L + \text{Repeat}(X_l i_{pool}, U_y) \\ W' &= W + \text{Repeat}(X_w i_{pool}, U_y) \end{aligned} \quad (14)$$

where the term $X_l, X_w \in \mathbb{R}^{e_u \times e_d}$.

Adapter-based feature fusion: The adapters are inserted at each decoder layer.

$$\text{Adapter}(i) = i + X_{up} \text{ReLU}(X_{down} i) \quad (15)$$

where the term $X_{down} \in \mathbb{R}^{e_u \times e_c}$, $X_{up} \in \mathbb{R}^{e_c \times e_u}$, $e_c \ll e_u$.

The modulation of adapter weights is done based on the following expressions:

$$X_{up} = X_{up} + \nabla X(i_{condi}) \quad (16)$$

Based on the classifier output, the conditional and the layer-wise personalization are allowed. Shape and interfaces in the conditioning process are given in Table IV.

Training schedule: Initially, the joint training is performed, followed by this process, staged training is performed. The sentiment and the offensive dataset are used for training the model. The weight of the classifier model is frozen, and the predicted sentiment is used for training the MT5 translation. With a reduced learning rate, the joint fine-tuning of the entire model is performed.

TABLE IV. SHAPE AND INTERFACES IN THE CONDITIONING PROCESS

Conditioning Method	Module Affected	Input Used	Shape	Interface
Prefix-Tuning	Decoder Input Embedding	$z_t, z_p \rightarrow Q$	$[C, 2m, 768]$	Prepended to decoder input
Cross-Attention Modulation	Decoder Cross-Attention	i_{pool}	$[C, U_y, 768]$	Bias added to K/V
Adapter Fusion (optional)	Decoder Layers	z_t, z_p	$[64, 768]$	Fused with up-projection weights

Loss weighting: The loss formulation involved in the proposed model is given as follows:

$$U_{toa} = \kappa_1 \cdot U_{MT} + \kappa_2 \cdot U_{sent} + \kappa_3 \cdot U_{off} \quad (17)$$

where, the translation loss is represented as U_{MT} , the multiclass loss is denoted as U_{sent} and the binary classification loss is represented as U_{off} , and the selected weight in the validation is represented as κ_j .

Curriculum strategy and stability consideration: The non-offensive and the short text are used in the initial stage. After that, sentiment-rich, longer and offensive sentences are included. Moreover, stability approaches, such as early stopping on the politeness and gradient clipping, are used.

Regularization and constraints: The constraint decoding adopted the blacklist constraint and the whitelist constraints for preserving the meaning of the words. If the Tamil words and flagged as offensive, the decoding process is disallowed in the blacklist constraint. With the help of the forced alignment copy mechanism, the

inclusions of named entities are forced. In the named entity copy mechanism, the output of the decoder at the time step u is represented as z_u .

$$Q(z_u) = (1 - \eta_u) \cdot Q_{gen}(z_u) + \eta_u Q_{copy}(z_u) \quad (18)$$

From the standard decoder SoftMax, the probability is represented as Q_{gen} . Over the named entities, the copy distribution from the encoder attention is represented as Q_{copy} . The copy gate is denoted as $\eta_u \in [0, 1]$.

Fig. 5 provides the diagrammatic view of the cultural sensitivity consideration-based English-to-Tamil translation model. The pipeline for culturally sensitive translation is clearly given in this model. Here, initially, the RoBERTa-ResLSTM is used for classifying the module. The classification outcome attained from this model is given to MC-AmT5, which incorporates politeness and cultural adjustment. How the social context guides the translation is visually encapsulated in this figure.

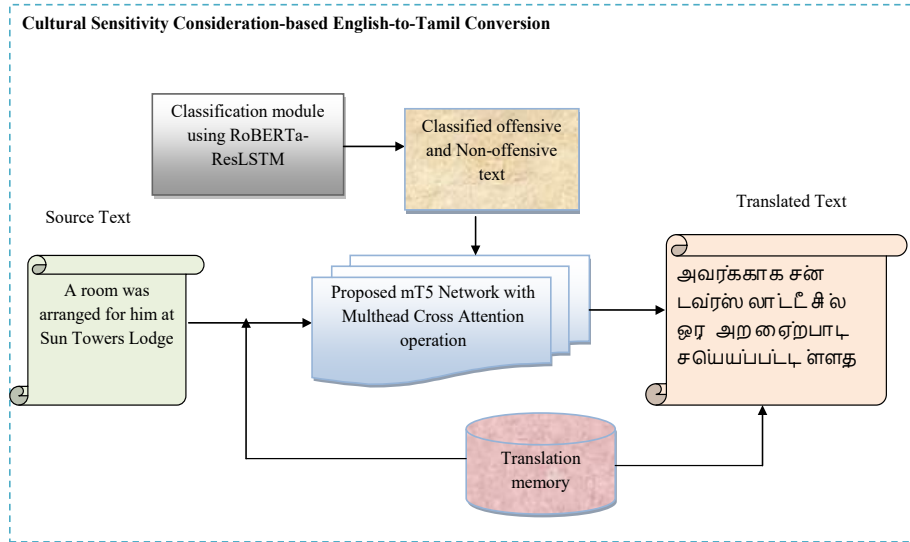


Fig. 5. Diagrammatic view of cultural sensitivity consideration-based English-to-Tamil conversion.

VI. RESULTS AND DISCUSSION

A. Experimental Setup

An effective NMT model was implemented in this proposed model and executed using Python software. The proposed model employed the MMV-MToC strategy for the optimization of variables. In particular, the proposed MMV-MToC algorithm operated with a population size

of 10, evaluating 10 candidate solutions simultaneously during optimization. Each candidate consisted of four elements, representing the four variables optimized at once. The optimization ran for up to 50 iterations, giving the algorithm multiple chances to explore and refine solutions progressively.

This configuration balanced thorough exploration with computational efficiency, enabling effective optimization of the NMT model's parameters within a practical

timeframe to improve its translation accuracy. The efficacy of the proposed classification and translation phases is validated by comparing the outcome with its existing models. Several algorithms like Zebra Optimization Algorithm (ZOA) [28], Wombat Optimization Algorithm (WOA) [29], Lyrebird Optimization (LO) [30], and Tornado optimizer with Coriolis force (ToC) [31] were considered for the validation of performance. Classification techniques like CNN [32], RNN [33], Long Short-Term memory (LSTM) [34], and BERT [35] were utilized for analyzing the performance of the proposed sentiments classification model. The translation performance was validated among conventional models like Deep-BERT [20], HAN [22], BERT-DGCN [24], and MC-mT5.

B. Resultant Text Samples of Proposed Context-Aware Machine Translation Framework

The proposed NMT Framework illustrates how context-sensitive translation enhances the accuracy and coherence of translated content by considering the surrounding textual environment. The proposed NMT framework includes offensive, neutral, and respectful translations from English to Tamil. The resultant outcomes are provided in Table V.

TABLE V. RESULTANT OUTCOME OF THE PROPOSED NMT MODEL

Input (English)	Output (polite phrase aligned with Tamil cultural norms)
Offensive English sentence "Get lost".	தயவுசெய்து வெளியே செல்லுங்கள்.
Informative English sentence "The device is not working properly".	செயல்படும் சில தவறுகள் உள்ளன.
Context-Aware sentence "You are such a mistake".	உங்கள் முயற்சியைப் பார்க்கும் போது சில பிழைகள் உள்ளன.

C. Performance Metrics

The formulas for calculating the performance metrics are provided in the equations below.

Classification performance metrics:

$$Ay = \frac{R_{+ve} + R_{-ve}}{R_{+ve} + R_{-ve} + G_{+ve} + G_{-ve}} \quad (19)$$

$$FNR = \frac{G_{+ve}}{R_{+ve} + G_{-ve}} \quad (20)$$

$$FOR = \frac{G_{-ve}}{R_{-ve} + G_{-ve}} \quad (21)$$

$$FPR = \frac{G_{+ve}}{R_{-ve} + G_{+ve}} \quad (22)$$

$$NPV = \frac{R_{-ve}}{R_{-ve} + G_{-ve}} \quad (23)$$

$$Pr = \frac{R_{+ve}}{R_{+ve} + G_{+ve}} \quad (24)$$

$$Sensitivity = \frac{R_{+ve}}{R_{+ve} + G_{-ve}} \quad (25)$$

$$Specificity = \frac{R_{-ve}}{R_{-ve} + G_{+ve}} \quad (26)$$

Translation performance metrics:

The Bilingual Evaluation Understudy (BELU) metric calculates the v-grams overlaps among the reference translations, which is mathematically expressed in Eq. (27).

$$BELU = BP \exp \left(\sum_{v=1}^V F_n \log T_v \right) \quad (27)$$

where, the term T_v is the precision

$$Ay = \frac{R_{+ve} + R_{-ve}}{R_{+ve} + R_{-ve} + G_{+ve} + G_{-ve}} \text{ on of } v \text{ grams and } F_n \text{ is}$$

the weight. Brevity penalty is denoted as BP .

CHaRacter-level F-Score (CHRF) is calculated as the harmonic mean of the precision and recall.

Translation Error Rate (TER) measures the number of edits ($A_S + F_D + K_I$) needed to change in a system to the reference words.

$$TER = \frac{A_S + F_D + K_I}{\text{Number of references words}} \quad (28)$$

where, the true positive and negative values are indicated as R_{+ve} and R_{-ve} . The false positive and negative values are indicated as G_{+ve} and G_{-ve} , respectively.

D. Convergence and Receiver Operating Characteristic (ROC) Curve Analysis of the Proposed NMT Model

The convergences and Receiver Operating Characteristic (ROC) curve analysis of the proposed NMT model are provided in Fig. 6. These analysis helps to determine whether the proposed MMV-MToC-MC-AmT5 model converges reliably to provide stable learning. The training processes are optimized with the support of the MMV-MToC model to avoid overfitting issues. ROC curve analyses determine how the proposed model MMV-MToC-MC-AmT5 differentiates offensive and non-offensive language in sensitive tasks like offensive language detection and polite translation.

A higher Area Under the Curve (AUC) value shows the model's superior capability to correctly classify the sensitive instances by reducing the error rate. Thus, the overall analysis provides inclusive evidence that the new MMV-MToC-MC-AmT5 model not only learns better but also translates offensive expressions more politely and accurately than existing models, enhancing the quality and cultural sensitivity of the NMT model.

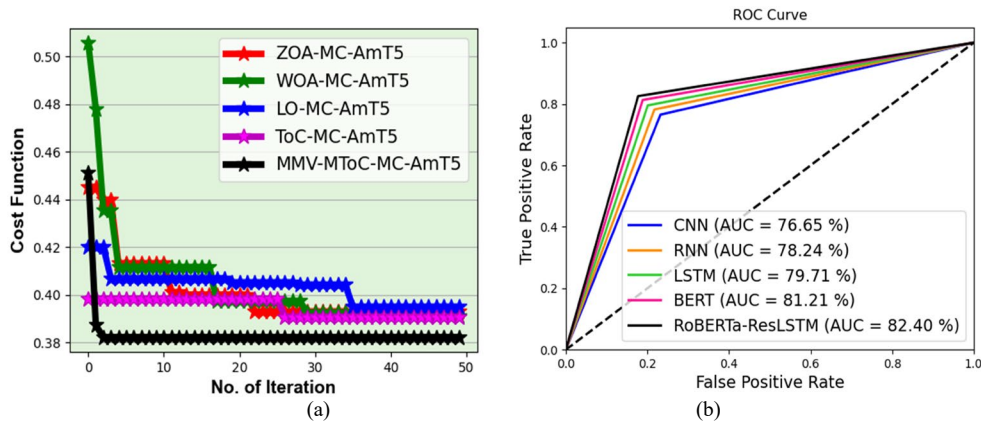
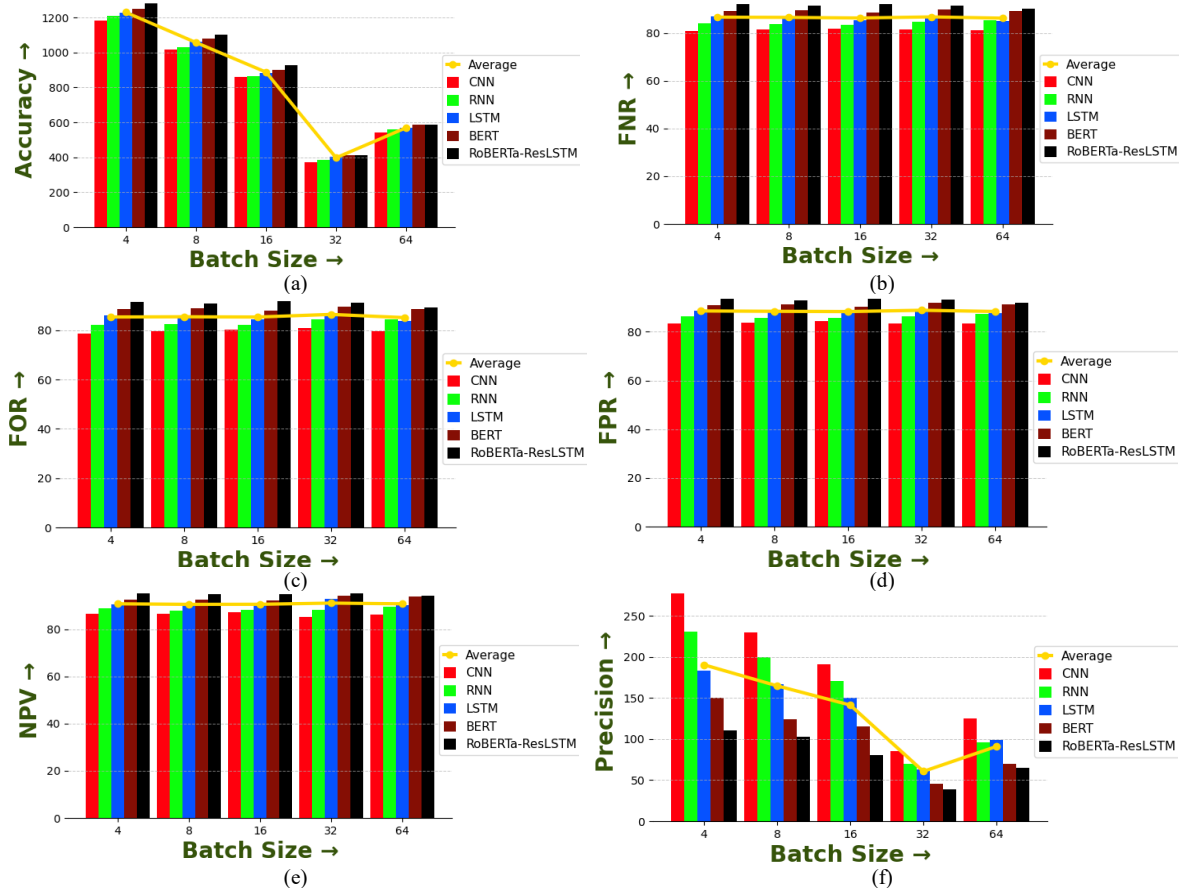


Fig. 6. Effectiveness analysis of proposed NMT model regarding (a) cost function, (b) ROC curve.

E. Performance Analysis of Multitask Classification

The performance analysis graph of multi-task classification is illustrated in Fig. 7, which validates the efficacy of the proposed RoBERTa-ResLSTM model for multitask classification tasks such as sentiment and offence detection. This process includes assessing various evaluation metrics on both training and testing datasets to understand how effectively the model differentiates offensive and non-offensive categories and its ability to generalize to new, unseen data. The analysis demonstrates that the suggested RoBERTa-ResLSTM model maintains consistent performance across different datasets.

Specifically, experiments conducted across different batch sizes revealed that, at a batch size of 4, the RoBERTa-ResLSTM achieved an accuracy of 12.4%, outperforming existing models like CNN, RNN, LSTM, and BERT. Consequently, the results confirmed that RoBERTa-ResLSTM delivers high precision in both sentiment analysis and offensive content classification. This capability highlights its proficiency in simultaneously analyzing emotional undertones and identifying offensive language, indicating the RoBERTa-ResLSTM model as a robust solution for sensitive language processing tasks. Its impressive performance, supported by a range of evaluation metrics, confirms that RoBERTa-ResLSTM is a dependable and effective model for multilabel language translation applications.



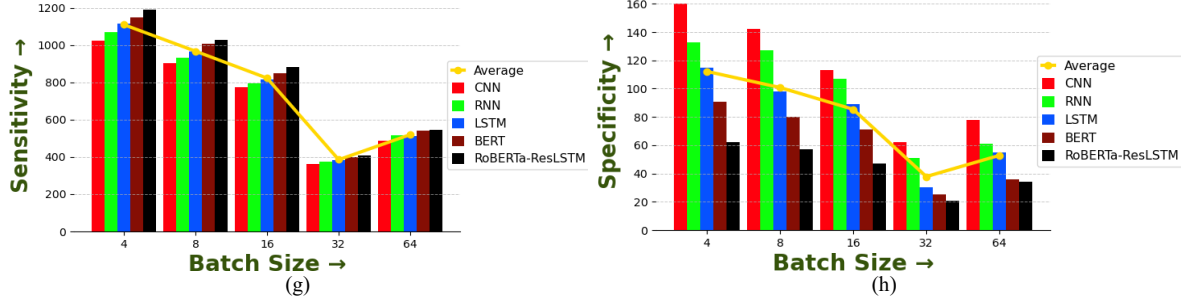


Fig. 7. Performance analysis of multitask classification using RoBERTa-ResLSTM Model Regarding (a) Accuracy, (b) FNR, (c) FOR, (d) FPR, (e) NPV, (f) Precision, (g) Sensitivity, (h) Specificity.

F. Performance Analysis of Context-Aware Machine Translation Framework

The evaluation of the context-sensitive machine translation system is conducted to tackle the vital issue of translating offensive language into Tamil while maintaining appropriate levels of politeness and cultural fidelity. The main aim of these evaluations is to assess and improve the model's capacity to produce translations that are both precise and culturally considerate, which is crucial in multilingual and intercultural interactions.

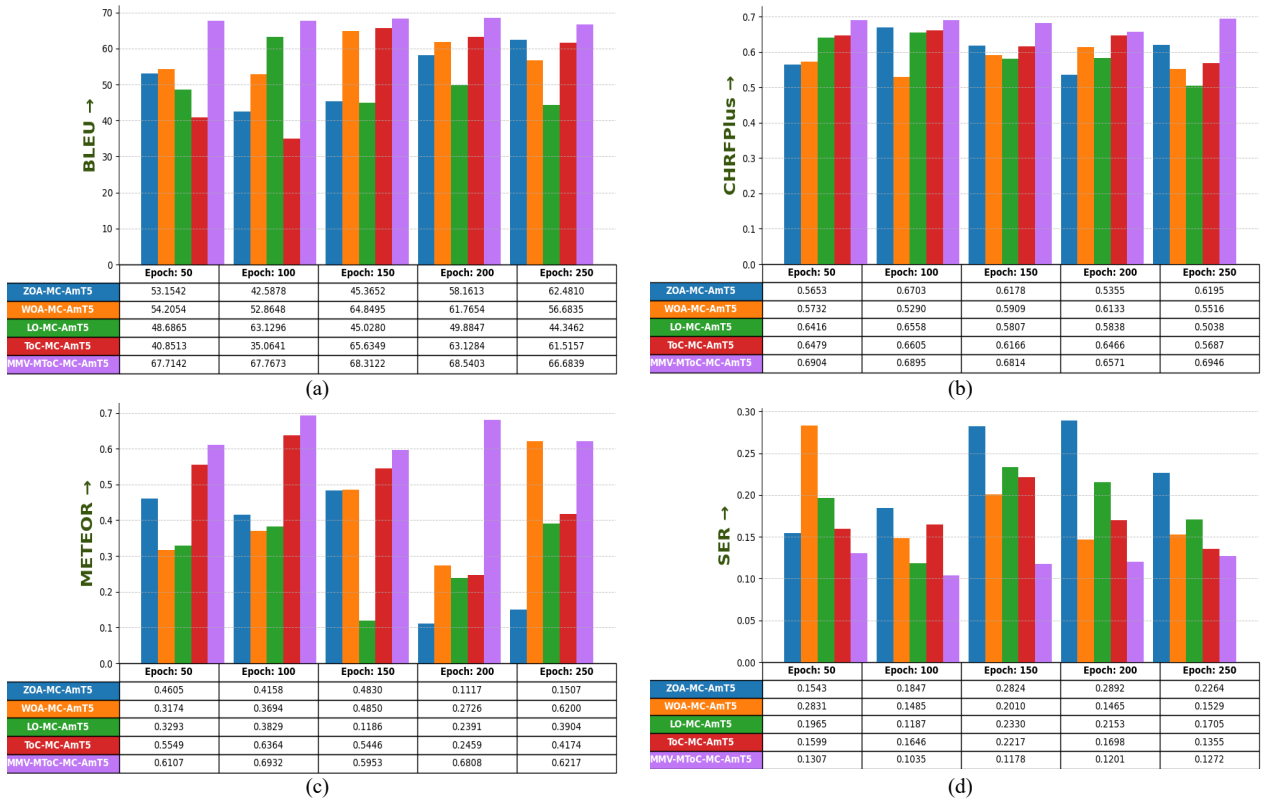
As a result of these assessments, translation accuracy has been enhanced, achieving a BLEU score of 67.76% at 50 epochs, which is better than the existing models. These findings highlight that the proposed MMV-MToC-MC-AmT5 model makes significant progress in the domain of politeness-controlled translation technology. Its high BLEU score confirms its effectiveness in delivering high-quality translations efficiently. The MMV-MToC-MC-AmT5 model's proficiency in

preserving cultural sensitivities and politeness levels makes it especially useful for responsibly translating offensive language and promoting respectful intercultural communications. The graphical view of performance analysis among traditional algorithms and techniques is provided in Fig. 8.

G. Multitask Classification Analysis of the Proposed NMT Model

Table VI provides the numerical analysis of the multitask classification model, which evaluates the performance of the proposed NMT model across existing models and analyzes the classification accuracy. A comprehensive performance evaluation provides a better understanding of how well the model performs across multiple aspects. Activation functions play a crucial role in analyzing how neural networks learn and represent data. The accuracy of the proposed model is 93.1%, for the ReLU activation function, which is better than existing models.

Among Algorithms



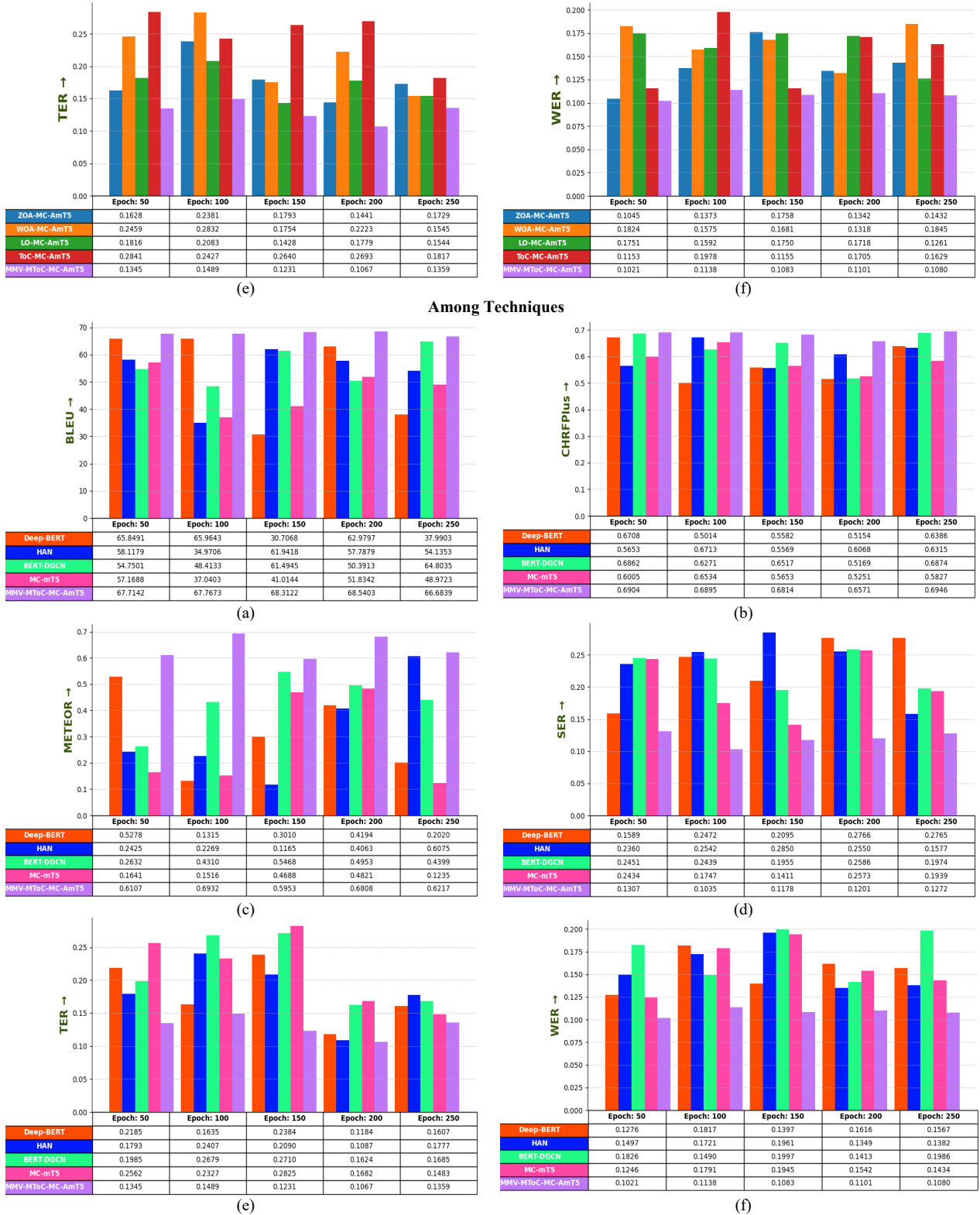


Fig. 8. Performance analysis of context-aware machine translation framework regarding (a) BLEU, (b) CHRRF, (c) METEOR, (d) SER, (e) TER, (f) WER.

The numerical results confirmed that the proposed RoBERTa-ResLSTM multi-task NMT model, especially with SoftMax activation, achieves superior performance, with a classification precision of 95.6%, surpassing

existing models. Thus, these comprehensive evaluations validate that the proposed RoBERTa-ResLSTM model is highly capable of accurately translating and classifying multilingual data in the NMT model.

TABLE VI. MULTITASK CLASSIFICATION ANALYSIS OF THE PROPOSED NMT MODEL AMONG TECHNIQUES

Activation Function	Metrics	CNN [32]	RNN [33]	LSTM [34]	BERT [35]	RoBERTa-ResLSTM
ReLU Activation Function	Accuracy (%)	84.63867	86.37911	89.40598	90.8059	93.41657
	Sensitivity (%)	82.40997	84.57746	87.71307	89.27806	92.44186
	Specificity (%)	87.32277	88.47097	91.33603	92.52412	94.47514
	Precision (%)	88.67362	89.49329	92.02683	93.07004	94.7839
	FPR (%)	12.67723	11.52903	8.663968	7.475884	5.524862
Linear Activation Function	Accuracy (%)	84.89742	86.59974	88.65124	91.31384	94.15103
	Sensitivity (%)	83.38816	84.67808	87.10744	89.92506	93.24324
	Specificity (%)	86.60465	88.81579	90.37928	92.84404	95.12195
	Precision (%)	87.56477	89.72366	91.019	93.26425	95.33679
	FPR (%)	13.39535	11.18421	9.620722	7.155963	4.878049
Sigmoid Activation Function	Accuracy (%)	85.19092	86.63571	88.18369	90.81527	93.49845
	Sensitivity (%)	82.72901	84.1954	86.54224	89.16256	92.30769
	Specificity (%)	88.08989	89.48546	90	92.63272	94.77054
	Precision (%)	89.10586	90.33916	90.54471	93.01131	94.96403
	FPR (%)	11.91011	10.51454	10	7.367281	5.229456
Tan Activation Function	Accuracy (%)	83.99546	88.42225	90.57889	90.46538	92.96254
	Sensitivity (%)	81.95652	86.27451	90	89.08686	91.7226
	Specificity (%)	86.22328	90.75829	91.15646	91.89815	94.23963
	Precision (%)	86.66667	91.03448	91.03448	91.95402	94.25287
	FPR (%)	13.77672	9.241706	8.843537	8.101852	5.760369
SoftMax Activation Function	Accuracy (%)	84.99594	87.99676	88.96999	90.51095	94.32279
	Sensitivity (%)	82.17968	86.34969	87.85047	89.39158	93.26019
	Specificity (%)	88.44765	89.84509	90.18613	91.72297	95.46218
	Precision (%)	89.71061	90.51447	90.67524	92.12219	95.65916
	FPR (%)	11.55235	10.15491	9.813875	8.277027	4.537815

H. Statistical Analysis of the Proposed Model

Table VII provides the statistical analysis of the proposed NMT model. By analyzing statistical data, the proposed MMV-MToC-MC-AmT5 model's performance is objectively compared against baseline translation systems, demonstrating improvements and identifying areas for further enhancement. The mean value of the proposed MMV-MToC-MC-AmT5 model is 4.3% better than ZOA-MC-AmT5, 5.04% better than WOA-MC-AmT5, 4.8% better than LO-MC-AmT5, and 2.7% better than ToC-MC-AmT5.

Thus, based on numerical data obtained from experiments or tests demonstrated that the MMV-MToC-MC-AmT5 framework is both reliable and efficient in performing English-to-Tamil translations that carefully take into account cultural nuances and the specific context of the text. This evidence confirms that the framework consistently produces high-quality translations that respect ethnic and cultural sensitivities, making it highly appropriate for practical use in situations where such respectful and accurate language conversion is essential.

TABLE VII. STATISTICAL ANALYSIS OF THE PROPOSED NMT MODEL AMONG TECHNIQUES

Metrics	ZOA-MC-AmT5 [28]	WOA-MC-AmT5 [29]	LO-MC-AmT5 [29]	ToC-MC-AmT5	MMV-MToC-MC-AmT5
Best	0.392931	0.390835	0.394888	0.390768	0.382254
Worst	0.445229	0.505656	0.420139	0.398373	0.451034
Mean	0.40133	0.404097	0.403379	0.394723	0.38373
Median	0.392931	0.397325	0.405143	0.398373	0.382254
Standard deviation	0.013979	0.021025	0.006557	0.003799	0.009641

I. Intended Deployment Contexts (Moderation, Customer Support, Educational Tools) and Their Requirements

In the multiple domains such as moderation, customer support, and educational tools, the suggested context-aware NMT is applied. Based on the design and the optimization of each system, they require different requirements.

1) Moderation

The offensive and harmful language is determined in the content of moderation for the purpose of preventing the spread of false information. In order to reliably detect the subtle nuances and tones along with the intent, a highly accurate offensive language detection model is required. Before the translation process, the flagged

content is determined by the multi-task classification model that has the capability to detect the offensive content more accurately. The process of transforming the sensitive content into non-insulting language is ensured by the politeness-aware translation, which may be used for preventing misunderstanding. In the context of moderation, the precise detection and mitigation of offensive language is required for preventing cultural sensitivity harm and legally compliant outputs.

2) Customer support

The clear and empathetic communication between the support agents and customer is ensured by this translation process. Here, empathy and politeness are the most important requirements that must be maintained for handling the emotionally charged interactions. The potentially offensive and blunt English statements are

converted into culturally aligned Tamil phrases via the politeness-controlled transfer model proposed in this research work. The context-aware and smooth flow is maintained in the customer support with the help of this model.

3) Education tool

An accurate, contextually rich, and sensitive translation is provided for the learners in the education application via the proposed model. The sentiments and the offensive elements are more effectively recognized by this model, which helps to perform the teaching in a more appropriate language. The annotations about the cultural context and the politeness are necessitated by the educational tools for improving interpretability and transparency. In addition, the consistency and the flow of the learning materials are maintained by learning the translation coherence over the longer passages using the deep learning model. This domain requires cultural

fidelity and detailed contextual understanding to support cultural awareness and learning a language for sensitive translation. In this context, the context-aware, empathetic and polite communication is required for maintaining the natural dialogue flow and balancing the offensive language mitigations.

J. Baseline Comparison of the Suggested Model

Table VIII shows the baseline comparison of the suggested MMV-MToC-MC-AmT5-based NMT model. As given in Table VIII, the sentiment alignment accuracy of the proposed MMV-MToC-MC-AmT5 model is 92.5%. Comparatively, the baseline models attained the sentiment alignment accuracy of 81.3, 83.1, 84.7, 86.9, 87.6, and 88.4. Thus, the results indicate the effectiveness of the suggested MMV-MToC-MC-AmT5-based NMT in comparison to the other models.

TABLE VIII. BASELINE COMPARISON OF THE PROPOSED MODEL

Model	BLEU	chrF	Sentiment Alignment Accuracy (%)	Politeness Preservation Rate (%)
mBART-50 (fine-tuned En-Ta)	24.8	47.2	81.3	72.5
mT5-base (fine-tuned En-Ta)	26.4	48.9	83.1	74.8
NLLB-200	27.1	50.5	84.7	76.2
IndicTrans2	28.6	52.3	86.9	78.1
Instruction-tuned LLM (FLAN-T5)	29.3	53	87.6	80.4
Instruction-tuned LLM (LLaMA-2-Chat)	30.1	54.2	88.4	81.6
Proposed MMV-MToC-MC-AmT5	33.7	57.8	92.5	88.9

K. Human Evaluations

The performance of the MMV-MToC-MC-AmT5 is assessed in terms of human evaluation metrics by focusing on the adequacy/fluency and perceived politeness and cultural appropriateness by native Tamil speakers across regions. The diverse groups of evaluators, namely native Tamil speakers, are involved in this analysis. Here, the unbiased assessment is ensured by the blind test analysis. The information about the generated caption is not disclosed to the evaluators, which helps to prevent biases in the judgment.

Adequacy/fluency and perceived politeness and cultural appropriateness by native Tamil speakers across regions: The linguistic richness, structural coherence and grammatical correctness are assessed by the evaluators to ensure the fluency. As evidenced by the qualitative analysis, the suggested MMV-MToC-MC-AmT5 model excels in generating captions with high fluency. The

capacity of the model in capturing the essential and the contextual details in the images is assessed in the adequacy analysis. The contextually relevant and accurate descriptions are managed by the proposed MMV-MToC-MC-AmT5 in complex and ambiguous scenarios. The culturally specific images are more adequately handled by this model.

Table IX shows the adequacy, frequency, perceived politeness, cultural appropriateness, and pairwise preference of the proposed MMV-MToC-MC-AmT5 model by collecting the Likert scores. In comparison to the prior models, the suggested MMV-MToC-MC-AmT5 provide adequacy (4.42), fluency (4.35), perceived politeness (4.53), cultural appropriateness (4.48) and pairwise preference (78.30%), which confirms the reliability of the suggested MMV-MToC-MC-AmT5 in the NMT.

TABLE IX. ANALYSIS BASED ON THE LIKERT SCORE

Metric	Proposed MMV-MToC-MC-AmT5	mT5 Baseline	MarianMT	Google Translate
Adequacy	4.42	3.89	3.73	3.58
Fluency	4.35	4.02	3.76	3.64
Perceived Politeness	4.53	3.12	2.94	2.78
Cultural Appropriateness	4.48	3.05	2.86	2.72
Pairwise Preference (%)	78.30	-	-	-

Inter-rater reliability: The inter-rater variability of the suggested MMV-MToC-MC-AmT5 is measured by measuring Krippendorff's Alpha (α), Cohen's Kappa (κ) coefficients in Table X. When adopting Krippendorff's Alpha, the suggested MMV-MToC-MC-AmT5 provides

the adequacy (0.81), fluency (0.8), perceived politeness (0.76), and cultural appropriateness (0.74), which confirms the effectiveness of the proposed MMV-MToC-MC-AmT5 in the NMT.

TABLE X. INTER-RATER RELIABILITY OF THE SUGGESTED MODEL

Dimension	Krippendorff's Alpha (α)	Cohen's Kappa (κ)
Adequacy	0.81	0.78
Fluency	0.8	0.76
Perceived Politeness	0.76	0.72
Cultural Appropriateness	0.74	0.7

L. Complexity Analysis

The complexities of the suggested MMV-MToC over the prior algorithms are given in Table XI. In the complexity analysis, the terms d represents the total number of parameters and the chromosome length is denoted as o . From the results, the complexity of the suggested MMV-MToC is lower than the other algorithms.

TABLE XI. COMPLEXITY OF THE SUGGESTED ALGORITHM

Optimizer	Complexity
SGD	$o \times d$
Adam/AdamW	$o \times 4d$
Ranger21	$o \times 5d$
Lookahead + Lion	$o \times 4d$
MMV-MToC	$o \times 3d$

M. Ablations Comparison

Table XII shows the ablation comparison of the MMV-MToC vs mainstream tuners to prove the effectiveness of the algorithms. Here, the BELU is used for the comparison of the suggested model with the other algorithms. From Table XII, the suggested MMV-MToC get the BELU of 33.5, which confirms the excellent performance of the proposed MMV-MToC over other approaches.

TABLE XII. ABLATION COMPARISON OF THE SUGGESTED ALGORITHM

Algorithms	BLEU
Random Search	30.6 $\pm$ 1.3
Bayesian	31.9 $\pm$ 1.0
Hyperband	32.1 $\pm$ 0.9
CMA-ES	32.4 $\pm$ 1.1
Proposed MMV-MToC	33.5 $\pm$ 0.4

N. Qualitative Results

The examples of the translation process are given in Table XIII. The annotation details are given as follows:

Offensive span: “idiot”;

Transformed into: Polite request without ad hominem;

Meaning preservation: task request retained.

Failure case: The failure case of the proposed model, hypothesis and the remedy are given in Table XIV.

TABLE XIII. EXAMPLES OF THE TRANSLATION PROCESS

Source (EN)	You're an idiot. Do your job properly.
Literal TA	நீ முட்டாள்தான். உன் வேலையை சரியாக செய்
Baseline	நீ முட்டாள்தான். உன் வேலையை சரியாக செய்
Proposed model	தயவு செய்து கவனமாக செயல்படுங்கள்; உங்கள் பணியைச் சிறப்பாக செய்ய வேண்டுகிறேன்

TABLE XIV. FAILURE CASE OF THE TRANSLATION PROCESS

Source (EN)	Fantastic attempt, like last time when everything broke.
Proposed	அற்புதமான முயற்சி; இம்முறை சிறப்பாக இருக்கும் என நம்புகிறேன் (missed negative sarcasm)
Issue	Polarity flipped to positive.
Hypothesis	Classifier under-detects multi-clause sarcasm with contrastive cues (like last time when everything broke).

VII. CONCLUSION

An innovative NMT framework successfully managed the polite conversion of offensive English content into Tamil while factoring in cultural sensitivity. Utilizing a multitask classification strategy with a RoBERTa-ResLSTM model for joint sentiment analysis and offence detection, the proposed system effectively identified both the emotional nuances and offensive elements present in the original text. These insights guided the politeness-controlled transfer model, which employed a multi-head cross-attention mechanism within an mT5 structure, optimally fine-tuned to maintain a balance between accurate translation and politeness modification. By combining sentiment and offence-aware information, the proposed MMV-MToC-MC-AmT5 approach produced culturally respectful translations that mitigated offensive language without altering the intended meaning of the

original text. The accuracy of the proposed MMV-MToC-MC-AmT5-based translation process was 94.1% at a linear activation function, which was better than existing models. Thus, the proposed method showed significant potential for improving machine translation frameworks to yield contextually polite and culturally sensitive Tamil translations from offensive English sources.

The proposed model encountered certain constraints, notably its dependence on the specific multitask RoBERTa-ResLSTM model and the mT5 architecture, which might limit its applicability across different languages. Moreover, the approach was mainly designed to handle textual data, potentially reducing its effectiveness when processing multimodal inputs such as audio or images.

Future investigations could focus on broadening the model's scope by integrating multimodal information, including speech and visual elements, to better detect

offensive content in diverse settings. Additionally, leveraging advanced large-scale pre-trained language models could enhance the model's capability to capture deeper semantic nuances and improve politeness modulation.

CONFLICT OF INTEREST

The authors declare no conflict of interest.

AUTHOR CONTRIBUTION

R. Mythili: Conceptualization, Methodology, Software Data curation, Writing—Original draft preparation. Rama S: Conceptualization, Methodology, Software Data curation, Writing—Original draft preparation, Visualization, Investigation, Software, Validation. R. Mythili: Writing, Reviewing and Editing. All authors had approved the final version.

REFERENCES

- [1] M. Pituxcoosuvann, W. Vijitkunsawat, and Y. Murakami, "Addressing sociolinguistic challenges in machine translation: An LLM-based approach for politeness and formality," in *Proc. 8th Int. Conf. Inf. Technol. (InCIT)*, 2024, pp. 757–762.
- [2] P. Priya, M. Firdaus, and A. Ekbal, "Computational politeness in natural language processing: A survey," *ACM Comput. Surveys*, vol. 56, no. 9, pp. 1–42, 2024.
- [3] F. H. Rachman, M. W. M. A. Syaqui, N. Ifada, and S. Wahyuni, "Transformer hyperparameter tuning for Madurese-Indonesian machine translation," *Eng. Technol. Appl. Sci. Res.*, vol. 15, no. 2, pp. 22216–22225, 2025.
- [4] Q. Wang, Y. Huang, and Z. Yuan, "Reducing redundancy in Japanese-to-English translation: A multi-pipeline approach for translating repeated elements in Japanese," in *Proc. 9th Conf. Mach. Transl.*, 2024, pp. 1047–1055.
- [5] M. Asmitha and C. R. Kavitha, "Bridging the language gap: Enhancing English-to-Telugu translation using NMT and encoding decoding techniques," in *Proc. 15th Int. Conf. Comput. Commun. Netw. Technol. (ICCCNT)*, 2024, pp. 1–6.
- [6] M. Nivaashini, G. Priyanka, and S. Aarthi, "Deep Neural Machine Translation (DNMT) hybrid deep learning architecture-based English-to-Indian language translation," in *Proc. Automatic Speech Recognition and Translation for Low Resource Languages*, 2024, pp. 331–373.
- [7] D. I. De Silva and D. G. P. Hansadi, "Enhancing machine translation: Cross-approach evaluation and optimization of RBMT, SMT, and NMT techniques," in *Proc. Int. Conf. Innov. Comput., Intell. Commun. Smart Electr. Syst. (ICES)*, 2024, pp. 1–8.
- [8] K. Bhuvaneswari and M. Varalakshmi, "Efficient incremental training using a novel NMT-SMT hybrid framework for translation of low-resource languages," *Front. Artif. Intell.*, vol. 7, 1381290, 2024.
- [9] J. Kiran, R. RR, and A. S. Krishnan, "Transfer grammar rules for Malayalam to English translation," *Language in India*, vol. 24, no. 11, 2024.
- [10] M. K. M. Kumar, V. Valliammai, D. G. B. Amali, and M. M. Noel, "A new robust deep learning-based automatic speech recognition and machine translation model for Tamil and Gujarati," *Automatic Speech Recognition and Translation for Low Resource Languages*, pp. 135–154, 2024.
- [11] U. Agarwal, R. Gupta, and G. Saranya, "Advancing multilingual communication: NLP-based translational speech-to-speech dialogue system for Indian languages," in *Proc. Int. Conf. Electron., Comput., Commun. Control Technol. (ICECCC)*, 2024, pp. 1–6.
- [12] A. H. N. Karthik, C. Aneesh, G. Saumik, K. V. V. Varun, and K. CR, "Sentiment analysis on Telugu text translated from English using NLP and ML," in *Proc. 3rd Int. Conf. Intell. Data Commun. Technol. Internet Things (IDCIoT)*, 2025, pp. 246–253.
- [13] S. Saxena, S. Chauhan, P. Arora, and P. Daniel, "Unsupervised SMT: An analysis of Indic languages and a low resource language," *J. Exp. Theor. Artif. Intell.*, vol. 36, no. 6, pp. 865–884, 2024.
- [14] G. Ravichandiran, V. Sivabaskaran, T. Uthayakumar, S. Vyravanathan, J. Krishara, and K. Rajendran, "Revolutionizing Tamil language analysis: A natural language processing model development approach," in *Proc. Int. Conf. Electr., Comput., Energy Technol. (ICECET)*, 2024, pp. 1–6.
- [15] T. K. Reddy, S. Veeksha, and C. R. Kavitha, "Enhanced multilingual image captioning: Integrating text-to-speech translation and sentiment analysis," in *Proc. Int. Conf. Innov. Comput., Intell. Commun. Smart Electr. Syst. (ICES)*, 2024, pp. 1–6.
- [16] L. HO, "English of Tamil multi-modal neural machine translation for image captioning," *Grenze Int. J. Eng. Technol. (GIJET)*, vol. 10, 2024.
- [17] K. Bhuvaneswari and M. Varalakshmi, "Efficient incremental training using a novel NMT-SMT hybrid framework for translation of low-resource languages," *Front. Artif. Intell.*, vol. 7, 1381290, 2024.
- [18] H. Yang, "Optimized English translation system using multi-level semantic extraction and text matching," *IEEE Access*, vol. 12, pp. 96527–96536, 2024.
- [19] F. Tan and H. Wang, "A semantic context-aware automatic quality scoring method for machine translation based on pretraining language model," *IEEE Access*, vol. 12, pp. 72023–72033, 2024.
- [20] M. A. H. Wadud, M. F. Mridha, J. Shin, K. Nur, and A. K. Saha, "Deep-BERT: Transfer learning for classifying multilingual offensive texts on social media," *Comput. Syst. Sci. Eng.*, vol. 44, no. 2, pp. 1776–1791, 2023.
- [21] P. Priya, M. Firdaus, and A. Ekbal, "XeroPol: Emotion-aware contrastive learning for zero-shot cross-lingual politeness identification in dialogues," *IEEE Trans. Comput. Soc. Syst.*, vol. 11, no. 5, pp. 6662–6671, 2024.
- [22] A. R. Pillai and B. Arun, "A feature fusion and detection approach using deep learning for sentiment analysis and offensive text detection from code-mix Malayalam language," *Biomed. Signal Process. Control*, vol. 89, 105763, 2024.
- [23] Z. Z. Hlaing, Y. K. Thu, T. Supnithi, and P. Netisopakul, "Improving neural machine translation with POS-tag features for low-resource language pairs," *Heliyon*, vol. 8, no. 8, e10375, 2022.
- [24] P. Priya, M. Firdaus, and A. Ekbal, "Two in one: A multi-task framework for politeness turn identification and phrase extraction in goal-oriented conversations," *Comput. Speech Lang.*, vol. 88, 101661, 2024.
- [25] V. Ganganwar and R. Rajalakshmi, "MTDOT: A multilingual translation-based data augmentation technique for offensive content identification in Tamil text data," *Electronics*, vol. 11, no. 21, p. 3574, 2022.
- [26] J. P. Venugopal, A. A. V. Subramanian, G. Sundaram, M. Rivera, and P. Wheeler, "A comprehensive approach to bias mitigation for sentiment analysis of social media data," *Appl. Sci.*, vol. 14, no. 23, 11471, 2024.
- [27] M. H. Abidi, "Multimodal data-based human motion intention prediction using adaptive hybrid deep learning network for movement challenged person," *Sci. Rep.*, vol. 14, no. 1, 30633, 2024.
- [28] E. Trojovská, M. Dehghani, and P. Trojovský, "Zebra optimization algorithm: A new bio-inspired optimization algorithm for solving optimization problems," *IEEE Access*, vol. 10, pp. 49445–49473, 2022.
- [29] Z. Benmamoun, K. Khlie, M. Dehghani, and Y. Gherabi, "WOA: Wombat Optimization Algorithm for solving supply chain optimization problems," *Mathematics*, vol. 12, no. 7, p. 1059, 2024.
- [30] M. Dehghani, G. Bektemyssova, Z. Montazeri, G. Shaikemelev, O. P. Malik, and G. Dhiman, "Lyrebird optimization algorithm: A new bio-inspired metaheuristic algorithm for solving optimization problems," *Biomimetics*, vol. 8, no. 6, 507, 2023.
- [31] M. Braik, H. Al-Hiary, H. Alzoubi, A. Hammouri, M. A. Al-Betar, and M. A. Awadallah, "Tornado optimizer with Coriolis force: A novel bio-inspired meta-heuristic algorithm for solving engineering problems," *Artif. Intell. Rev.*, vol. 58, no. 4, pp. 1–99, 2025.

- [32] S. Soni, S. S. Chouhan, and S. S. Rathore, "TextConvoNet: A convolutional neural network based architecture for text classification," *Appl. Intell.*, vol. 53, no. 11, pp. 14249–14268, 2023.
- [33] P. Sunagar and A. Kanavalli, "A hybrid RNN based deep learning approach for text classification," *Int. J. Adv. Comput. Sci. Appl.*, vol. 13, no. 6, pp. 289–295, 2022.
- [34] X. Song, Y. Liu, L. Xue *et al.*, "Time-series well performance prediction based on Long Short-Term Memory (LSTM) neural network model," *J. Pet. Sci. Eng.*, vol. 186, 106682, 2020.
- [35] R. Qasim, W. H. Bangyal, M. A. Alqarni, and A. A. Almazroi, "A fine-tuned BERT-based transfer learning approach for text classification," *J. Healthc. Eng.*, vol. 2022, no. 1, 3498123, 2022.

Copyright © 2026 by the authors. This is an open access article distributed under the Creative Commons Attribution License which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited ([CC BY 4.0](#)).