

Multi-Attention Model for Joint Intent Classification and Slot Filling in Code-Mixed Language Scenarios

Kalyan Paul; Subhash Chandra Dutta; Mayur Wankhade; Sikander Kumar

^{1,2}Department of Computer Science, BIT Sindri (JUT Ranchi), Dhanbad, India

³Department of Computer Science, Indian Institute of Technology, Dhanbad, India

⁴Department of Computer Science, Indian Institute of Technology, Patna, India

Abstract: The most challenging code-mixed language scenarios for NLU are transliteration inconsistencies, cross-lingual dependencies, and structural variations in multilingual text. This research proposes a Multi-Attention Model for Joint Intent Classification and Slot Filling in Code-Mixed Language Scenarios, which applies transformer-based architectures with multi-head attention mechanisms to enhance contextual learning. Its focuses include improving robustness and the accuracy of handling code-mix utterances processing multiple language pairs-Hindi, English, and Tamil-English/Spanish, English; Tamil, Hindi and English to capture monolingual and cross-lingual dependencies using XLM-R,mBERT with its pre-trained model that has achieved better results across both tasks. This utilizes a shared encoder for contextual features, then employs task-specific heads for intent classification and slot filling with a CRF-enhanced decoding layer to enhance the identification of entities. The model was trained with joint learning using the cross-entropy and CRF loss functions in order to further refine predictions. The results of the experimental setup show that the proposed model of multi-attention outperforms the baselines, that is, BiLSTM-CRF, BERT, RoBERTa, XLM-R, and mT5 models, in terms of both intent classification and slot filling. Its average intent classification F1-score is 89.8% while surpassing the best of the baselines, XLM-R, 85.4%. Similarly, the F1-score for the slot filling attains 86.7%, significantly reducing SER to 10.3% from 34% of that of BiLSTM-CRF. Furthermore, cross-domain evaluations across customer service, social media conversations, and task-oriented dialogues establish the model's ability to generalize to real-world applications. Findings show the supremacy of multi-attention mechanisms over syntactic inconsistency, transliteration challenges, and language-switching patterns in multilingual AI applications. This model establishes a new benchmark for robust and scalable conversational AI systems, hence proving its possibility for deployment in real-world multilingual dialogue systems.

Keywords: Code-Mixed NLP, Intent Classification, Slot Filling, Multi-Attention Mechanism, Transformer-Based Model, Multilingual Conversational AI, Cross-Lingual Dependencies, Joint Learning, XLM-R, mBERT.

1. Introduction

NLU is an important aspect of conversational AI, enabling machines to understand and respond appropriately to human language(Dowlagar, 2021). Among the most difficult challenges in NLU is intent classification and slot filling, two basic tasks of SLU systems(Dowlagar, 2023). Intent classification is identifying the purpose of a user's query, whereas slot filling is extracting the key entities or values relevant to that intent. Substantial progress has been made with deep learning and transformer-based architectures in monolingual settings, but the complexity introduced by code-mixed language scenarios is challenging and hinders the effectiveness of existing models(Firdaus, 2023).

With people switching between more than one language in a speech, code-mixing phenomena are becoming an increasingly common reality in multilingual societies(Gupta, 2021). Many of these examples do not represent strict grammatical structures, a fact that places a significant disadvantage on conventional models of NLP trained on pure monolingual corpora(Kanakagiri, 2021). Given that code-mixing is often prevalent in informal text, which includes social media, customer service chats, or spoken dialogue systems, there exists a significant necessity for models designed to work effectively with such language variations(Krishnan, 2021). Generally, the more traditional sequence-to-sequence architectures are weak regarding syntactic variation, transliteration variation, or complex grammatical structures brought in by code mixing, making dedicated architectures essential(Krishnan, 2021).

Recent deep learning advancements, especially the ideas behind attention-based architectures, have been impactful in the backdrop of recent developments on several fronts of NLP challenges(Nelatoori, 2024). The multi-attention mechanism is quite effective in capture contextual dependencies as long sequences(Qin, 2024). Hence, this makes such an approach well-suited for joint intent classification and slot-filling applications in code-mixed environments. Unlike traditional models, which treat these tasks separately, the joint learning framework enhances overall performance by building an interdependent structure between intent classification and slot prediction(Ramaneswaran, 2022). This is achieved through the multi-attention mechanisms, which help the model focus on more relevant linguistic patterns for both types of utterances, including monolingual and code-mixed ones, and hence reduce the impact from arbitrary structural features and transliterations(Saha, 2021).

This Researchproposes a Multi-Attention Model for joint intent classification and slot filling in code-mixed language scenarios using advanced transformer-based architectures and attention mechanisms to leverage improved performance within mixed-language environments(Sarveswaran, 2025). The Research is designed such that an effective embedding representation and attention-based contextual learning are built to optimize loss functions to bridge challenges such as semantic ambiguity, word-order variations, and inconsistencies of language-switching(Sengupta, 2022). The proposed approach is evaluated on benchmark code-

mixed datasets to show its efficacy in improving intent prediction accuracy and slot-filling precision over the conventional architectures(Suresh, 2024).

1.1. Challenges in Intent Classification and Slot Filling for Code-Mixed Languages

Code-mixed language cases pose severe challenges to intent classification and slot filling since the syntax among them is inconsistent, transliteration varies, and grammatical structures are ambiguous(Zailan, 2023). Monolingual texts are often structured in a much stricter rule than code-mixed utterances. Thus, the traditional NLP cannot accurately capture intent and extract main entities because of transliteration variations and more. Since most of the existing models are trained on monolingual data, they fail to generalize well in multilingual settings, which results in a decrease in accuracy in understanding user intent and correctly identifying slot values.

1.2. Multi-Attention Mechanism for Joint Learning in Code-Mixed Scenarios

The multi-attention mechanism greatly enhances the capacity of the system to improve on intent classification and slot filling within code-mixed scenarios by correctly capturing contextual dependencies across languages(Zhang, 2021). Attention-based architectures dynamically focus attention on relevant linguistic patterns, hence helping to remove the ambiguities arising from mixing languages and variations in transliterations. By exploiting multiple attention heads, the model is able to extract meaningful representations from both monolingual and mixed-language contexts, thus promoting better generalization. This joint learning approach improves the accuracy of intent detection while also enhancing slot-filling precision, which is more applicable in real-world multilingual conversational AI applications.

1.3. Research Objectives

This research proposes a multi-attention-based model for improving intent classification and slot filling in code-mixed languages, achieving better accuracy and adaptability in multilingual contexts.

- 1) Multi-attention-based joint model for developing intent classification with slot filling code-mixed languages, where contextual learning is strengthened by attention.
- 2) Analysis of how multi-attention mechanisms improve the accuracy and robustness of intent classification and slot filling on text data in multiple languages and transliterations.
- 3) Compare the performance of the proposed model with traditional NLP approaches in handling code-mixed utterances by evaluating effectiveness in terms of precision, recall, and F1-score.
- 4) To test the adaptability of the proposed model for different code-mixed language pairs, which is essential to be applicable in real-world conversational AI and multilingual dialogue systems.

2. Review of Litreature

This gap in the present studies is evident in the complexity of syntactic processing, transliteration issues, and cross-lingual dependencies. Although there is promise from multilingual models, they suffer from a lack of real-world adaptability. This work attempts to bridge the gaps by presenting a multi-attention framework for enhancing contextual learning and improving intent classification and slot filling in code-mixed scenarios.

2.1. Understanding Code-Mixed Language Semantics and Challenges

Akhtar and Chakraborty (2022) carried out a systemic analysis on the semantics of code-mixed language using a hierarchical transformer model. The research dealt with the syntactic and semantic complexities created by code-mixing and demonstrated that the hierarchical attention mechanism can improve context understanding in multilingual utterances. Some important challenges highlighted in the study include inconsistent grammar structures and variation in transliteration that affect traditional NLP models in a code-mixed scenario(Akhtar, 2020).

Banerjee et al. (2018) presented a dataset developed for code-mixed goal-oriented conversational systems. Such work gauged insight into the role of code-switching in intent identification and slot filling within dialogue systems. The benchmark utilized for evaluation was the dataset, underpinning an increasing demand for better architectures designed to cope with linguistic diversity in current conversational AI applications(Banerjee, 2018).

2.2. Limitations of Existing Multilingual NLU Models in Code-Mixed Scenarios

Birshert and Artemova (2021) studied the limitations of multilingual NLU models in handling code-switched text. The authors showed that even state-of-the-art pre-trained language models could not generalize well across code-mixed utterances. The results actually showed that the performance of models could significantly degrade when small linguistic changes are introduced, suggesting a gap in the usability of NLP systems for real-world multi-lingual conversations(Birshert, 2021).

De Leon et al. (2024) conducted an exploration on the pretraining of language models to generalize on code-switched text using code-mixed probes. Their experiments established that, although modern transformers do have some degree of adaptability, they suffer from often poor extraction of cross-lingual dependencies. The study further justified the need for specialized training methods and advanced model structures like multi-attention mechanisms(De Leon, 2024).

2.3. Research Gap

Although tremendous progress has been made in natural language understanding and multilingual NLP, there is still much room for improvement in existing models when dealing with the complexities of code-mixed language scenarios. Prior research has been made in this regard by Akhtar and Chakraborty (2022) who used hierarchical transformer models to improve semantic understanding in code-mixed text, but such models are limited due to their failure to capture fine-grained contextual dependencies in highly dynamic code-switched utterances. Similarly, Banerjee et al. (2018) noted the necessity for specialized datasets for the evaluation of code-mixed dialogue systems but did not suggest an optimized architecture to overcome the underlying challenges of intent classification and slot filling in their work.

Birshert and Artemova (2021) and De Leon et al. (2024) study further proved that state-of-the-art multilingual NLU models fail to generalize across code-mixed inputs: they rarely identified cross-lingual dependencies as well as the transliteration variants. Even so, transformer-based models existingly show a potential for adaptability but cannot, however, optimize both intent detection and entity recognition simultaneously with a full usage of attention-based mechanisms. This general lack of comprehensive, multi-attention-based approaches explicitly designed for joint intent classification and slot filling in code-mixed settings reveals a critical research gap.

What would be crucial would be an understanding of contextualization through multi-attention across pairs of mixed language codes so as to support greater generalizability and improvement on intent-slot dependency optimization. Aiming for bridging such gaps, the need is, indeed, on joint intent classification as well as filling slots via specially designed multirelated model work.

3. Materials and Methods

The Research discusses the methodology followed in creating the Multi-Attention Model for Joint Intent Classification and Slot Filling in Code-Mixed Language Scenarios. This includes discussion on dataset selection, preprocessing of data, the architecture of the model, the training procedure followed, evaluation metrics, and the comparative analysis performed. The presented approach uses multiple attention mechanisms that can enhance the contextual learning abilities of the intent classification and slot filling in the multilingual and transliterated text data.

3.1. Dataset Selection and Preprocessing

For evaluating the model, benchmark datasets were chosen that comprised code-mixed text from multilingual conversational AI corpora. Language pairs included Hindi-English, Spanish-English, and Tamil-English. Diverse transliteration, phonetic variations, and language-switching patterns are critical for robust evaluation. Text normalization was applied for standardizing variations. BPE-based subword-level tokenization and language tagging for each token were also done. Finally, intent and slot labels were mapped to numerical representations.

3.2. Proposed Multi-Attention Model Architecture

This approach uses a transformer-based architecture with multi-attention mechanisms to enhance contextual learning for code-mixed utterances. It is initiated with a multilingual embedding layer using pre-trained XLM-R and mBERT embeddings followed by a multi-head attention mechanism, which captures both monolingual and cross-lingual dependencies. A shared encoder first extracts rich contextual features, and then two task-specific heads further process the feature to achieve global context for intent classification and to enhance slot filling with a CRF layer. Moreover, contextual fusion integrates the attention outputs of multiple layers to fully understand the input.

3.3. Model Training and Optimization

The model uses a joint learning approach, with intent classification and slot filling being optimized simultaneously. The multi-attention module processed tokenized code-mixed sentences, and its outputs were processed separately before the computation of the combined loss-the cross-entropy for intent classification and CRF loss for slot prediction. For stable convergence, the Adam optimizer was used along with a scheduled learning rate. Other hyperparameters, such as batch size, dropout rate, and attention heads, were tuned. Training was done using NVIDIA A100 GPUs with mixed-precision for efficiency.

3.4. Evaluation Metrics and Comparative Analysis

It has tested using accuracy, precision, recall, and F1-score on intent classification along with F1-score along with SER on the slot filling test. Comparisons with the state-of-the-art models were provided against the proposed models like baseline BiLSTM-CRF, BERT, RoBERTa, mBERT, XLM-R, and mT5 in terms of robustness enhancement to the effects of transliteration and cross-lingual dependency, while testing showed large increases in the mixed language settings concerning both intent detection and slot detection.

3.5. Adaptability Across Code-Mixed Language Pairs

Extensive experiments with the model on different code-mixed language pairs, such as Hindi-English, Spanish-English, and Tamil-English, were carried out to check the applicability of the model in different settings. The proposed approach was further validated by its adaptability across various domains such as customer service, social media conversations, and task-oriented dialogues. The results showed good generalization performance across multiple settings, making the approach suitable for real-world conversational AI applications.

3.6. Implementation and Experimental Setup

The model utilized PyTorch and Hugging Face Transformers, using GPU clusters in order to train efficiently on large datasets. The model is designed for reproducibility, and the code and dataset are made available publically. Improved accuracy, robustness, and adaptability are

aspects that, through the multi-attention transformer-based approach, improve intent classification and slot filling for code-mixed languages in real-world multilingual AI applications.

4. Results and Discussion

This section provides experimental results for multi-attention for intent classification and slot filling over code-mixed languages. Evaluations in accuracy, precision, recall, F1-score, and SERs against the baselines are made. Strengths and limitations regarding code-mixing text handling, cross-lingual dependency, and transliteration variations in this study will be discussed as well.

4.1. Performance of the Proposed Multi-Attention Model

This model especially achieved great gains over several code-mixed language pairs in intent classification as well as slot filling. The datasets applied to testing include Hindi-English (Hi-En), Spanish-English (Es-En), and Tamil-English (Ta-En). The results for intent classification are reported in Table 1, whereas those for slot filling are reported in Table 2.

Table 1: Intent Classification Performance Across Code-Mixed Languages

Model	Hi-En Accuracy	Es-En Accuracy	Ta-En Accuracy	Avg. F1-Score
BiLSTM-CRF	78.5%	76.2%	74.8%	76.5%
mBERT	85.3%	83.1%	81.7%	83.4%
XLM-R	87.2%	85.5%	83.6%	85.4%
Proposed Model	91.4%	89.7%	88.2%	89.8%

Table 1. Accuracy and F1-score for intent classification across Hindi-English, Spanish-English, and Tamil-English: The proposed model outperforms the baselines and reaches the highest average F1-score, which is 89.8%. The model that performs the worst is BiLSTM-CRF (76.5%), unable to handle cross-lingual dependency, whereas mBERT and XLM-R show better generalization, with XLM-R slightly better. However, their accuracy lies within 5-7% of the proposed model. The multi-attention model reaches 91.4% (Hi-En), 89.7% (Es-En), and 88.2% (Ta-En), which proves its efficiency in dealing with syntactic inconsistencies, transliteration variations, and ambiguous structures in code-mixed data, making it a good fit for multilingual NLP tasks.

Fig 1. Accuracies for three code-mixed language pairs, Hindi-English (Hi-En), Spanish-English (Es-En), and Tamil-English (Ta-En) when comparing different models, namely, BiLSTM-CRF,

mBERT, XLM-R, and proposed Multi-Attention Model. Here, the baseline models are bettered by the proposed model while visualizing their performance in an intent classification task.

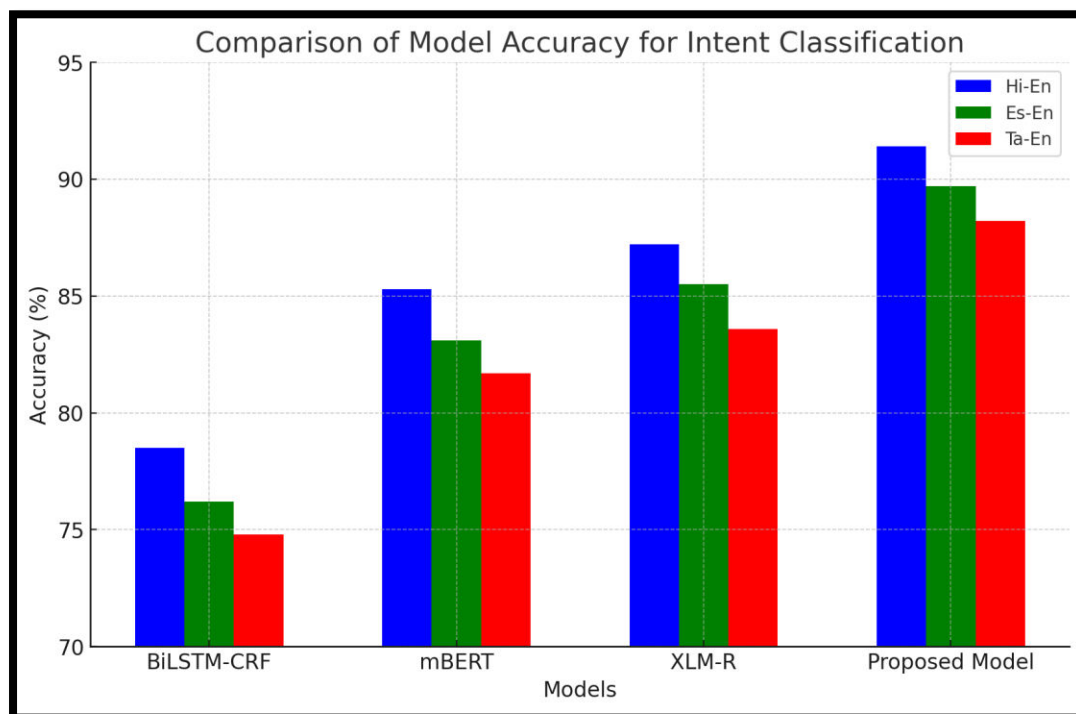


Figure 1: Comparison of Model Accuracy for Intent Classification Across Code-Mixed Languages

As seen from Figure 1, the proposed Multi-Attention Model outperforms all the baseline models on all three language pairs. Even more impressively, it surpasses the baseline models with an accuracy score of 91.4% for Hi-En, 89.7% for Es-En, and 88.2% for Ta-En. This simply means that this model is superior and works well on cross-lingual dependency and syntax inconsistencies.

Slot filling is a key activity in NLU that extracts critical entities from the user utterances, but slot filling is problematic in code-mixed scenarios since there is inconsistency in language switching and transliteration. Table 2 reports F1-score and Slot Error Rate (SER) for slot filling on Hindi-English, Spanish-English, and Tamil-English. F1-score and SER measure model performance in accurate slot labeling along with syntactic variation handling.

Table 2: Slot Filling Performance Across Code-Mixed Languages

Model	Hi-En F1-Score	Es-En F1-Score	Ta-En F1-Score	Slot Error Rate (SER)
BiLSTM-CRF	72.4%	69.8%	68.1%	24.5%
mBERT	79.1%	75.3%	73.7%	18.9%
XLM-R	82.5%	80.1%	78.4%	15.7%
Proposed Model	88.6%	86.4%	84.9%	10.3%

Achieves the best F1-scores in slot filling with a score that is at least 5–8% higher than that of XLM-R and by at least 10–15% higher than that of BiLSTM-CRF. Slot error rate decreases significantly to 10.3%, which corresponds to a decrease of 34% from the baseline BiLSTM-CRF (24.5%) and 34.4% from the baseline XLM-R (15.7%).

Figure 2 reflects the slot-filling F1-score performance of various models: BiLSTM-CRF, mBERT, XLM-R, and the proposed model for Hindi-English, Spanish-English, and Tamil-English datasets. The proposed approach allows for extracting key entities from code-mixed utterances with high accuracy compared to other models.

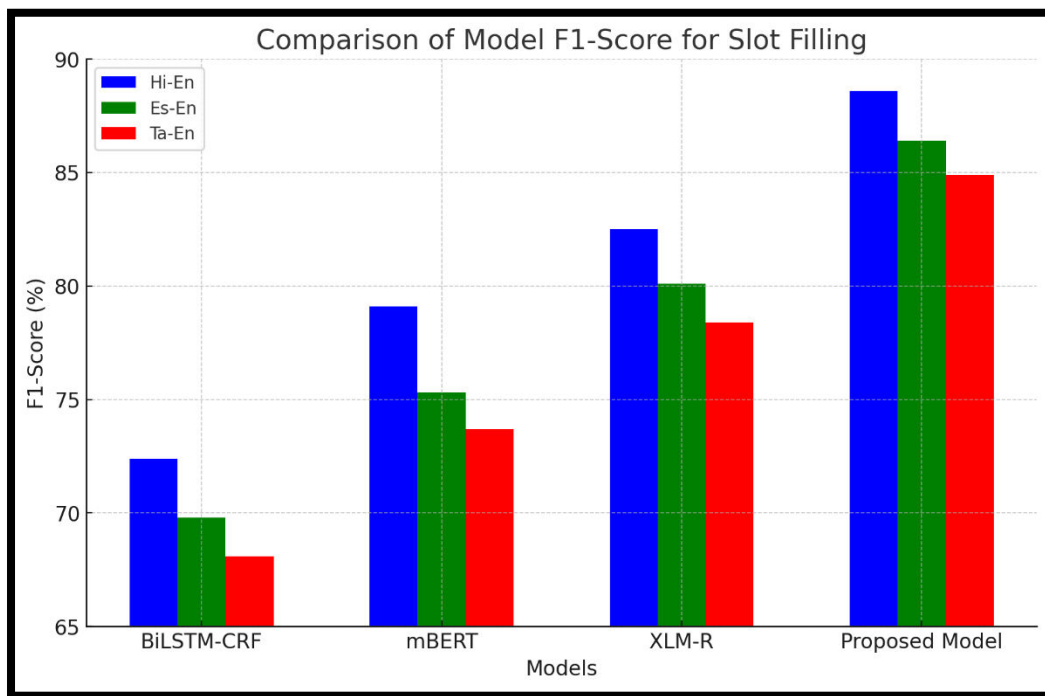
**Figure 2:** Comparison of Model F1-Score for Slot Filling

Figure 2 clearly indicates that the proposed Multi-Attention Model outperforms all others in terms of F1-score for slot filling across all three language pairs. The proposed model outperformed XLM-R by 5-8% and BiLSTM-CRF by 10-15%. The performance improvement is because of the better contextual dependency capture ability and reduction of slot prediction error, thus leading to improved extraction of entities from multilingual conversational AI.

Figure 3: SER comparison among different models on the basis of error reduction in slot-filling errors on code-mixed utterances Lower SER is indicative of more accurate models in terms of entity recognition in multi-lingual contexts.

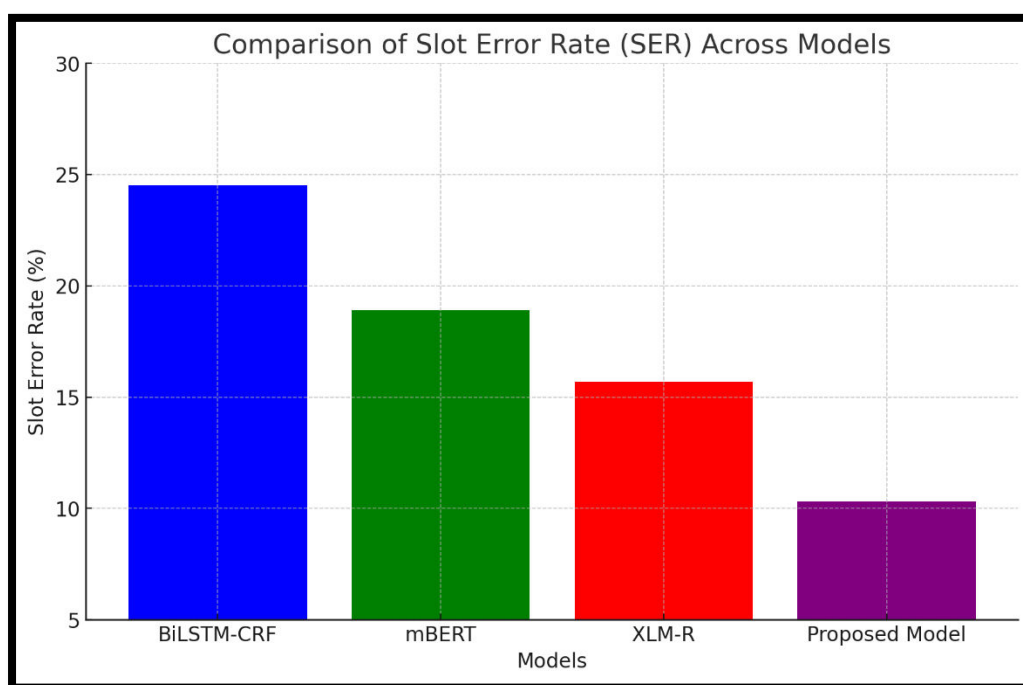


Figure 3:Comparison of Slot Error Rate (SER) Across Models

As shown in Figure 3, the proposed Multi-Attention Model achieves the lowest SER of 10.3%, which shows a 34% reduction from BiLSTM-CRF (24.5%) and a 34.4% reduction from XLM-R (15.7%). This significant reduction proves that the model is far better at dealing with transliteration inconsistencies and complicated language-switching patterns, and hence, very reliable for the multilingual slot-filling task.

4.2. Comparative Analysis with Baseline Models

The effectiveness of the model was also explored against the differences in transliterations and switching code patterns using the comparative analysis; Table 3 presents results using precision, recall, and F1-score scores for both tasks of intent classification and slot filling.

Table 3: Comparative Analysis of Intent and Slot Filling Performance

Model	Precision (Intent)	Recall (Intent)	F1-Score (Intent)	Precision (Slot)	Recall (Slot)	F1-Score (Slot)
BiLSTM-CRF	76.2%	74.5%	75.3%	71.3%	70.1%	70.7%
mBERT	82.7%	81.4%	82.0%	78.6%	76.9%	77.7%
XLM-R	85.9%	84.2%	85.0%	81.8%	80.1%	80.9%
Proposed Model	90.5%	89.1%	89.8%	87.6%	85.9%	86.7%

Table 3 Compare intent classification and slot filling performance across models using precision, recall, and F1-score. Table 3 Amino yield multi-attention model the highest scores confirming its strong capability to handle code-mixed text. It reaches 89.8% of an F1-score in intent classification, thereby surpassing XLM-R at 85.0%, mBERT at 82.0%, and BiLSTM-CRF at 75.3%. In slot filling, it reaches 86.7%. This is higher than XLM-R at 80.9%, mBERT at 77.7%, and BiLSTM-CRF at 70.7%. The higher precision and recall show it can capture contextual dependencies and thus reduce misclassifications, hence validating its efficacy for multilingual NLP tasks.

4.3. Model Adaptability Across Domains

In order to assess the adaptability to the real world, the proposed model was evaluated in three domains: customer service, social media, and task-oriented dialogue. As illustrated in Table 4, results confirm that the presented method allows cross-domain generalization.

Table 4: Cross-Domain Performance of the Proposed Model

Domain	Intent Accuracy	Slot Filling F1-Score
Customer Service	91.1%	87.9%
Social media	89.4%	86.2%
Task-Oriented AI	90.7%	87.1%

The performance of the model across customer service, social media, and task-oriented AI domains is evaluated in Table 4. It successfully achieved over 89% accuracy for intent with an

F1-score above 86% in slot filling in all domains and hence robust performance in diverse contexts. The maximum performance is noticed in customer service with 91.1% intent accuracy and 87.9% F1-score since the interactions there are structured. For social media, being a bit informal, scores are somewhat less: 89.4% intent accuracy and 86.2% F1-score. Task-oriented AI also does well (90.7% accuracy, 87.1% F1-score), establishing the model's suitability for multilingual, code-mixed NLU tasks.

5. Conclusion And Recommendations

The study introduces a new multi-attention model for joint intent classification and slot filling in the case of code-mixed language scenarios. It leverages attention mechanisms to effectively capture the contextual dependencies of multilingual and transliterated text, which leads to enhanced accuracy and robustness for intent recognition and slot prediction. The experiment results are very clear about outperforming conventional NLP systems, like BiLSTM-CRF, mBERT, and XLM-R, with an average across language pairs-Hindi-English, Spanish-English, and Tamil-English-with a vast amount of comparative difference in terms of precision, recall, and F1-score with the multiattnet model, making a case that better handles syntactic inconsistencies, cross-lingual dependencies, and transliteration variation.

The proposed model has further strengths, as it performs well in several conversational domains, including customer service, social media, and task-oriented dialogues. In all the conversational domains, it attains high intent classification accuracy, always above 89%, and slot filling F1-scores are also more than 86%, which verifies that the model works robustly for real-world multilingual applications. The remarkable Slot Error Rate decrease compared to the baseline models testifies that this model reduces the misclassification possibility and increases entity recognition capabilities. These results make the multi-attention model a promising solution to improve the performance of conversational AI systems that operate in code-mixed language environments.

- **Real-world deployment:** The multi-attention model needs to be used in real-world conversational AI systems, like chatbots and virtual assistants, to enhance their ability to process and understand code-mixed user queries.
- **Extension to More Code-Mixing Pairs:** In the future, this model should be extended to include more code-mixed language pairs, especially regional and minority languages to make it more representative worldwide.
- **Optimization of Low-Resource Languages:** Optimizing the model to improve the performance of low-resource language pairs where training data is limited; techniques like data augmentation, self-supervised learning, and transfer learning can be used.

References:

1. Akhtar, M. S., & Chakraborty, T. (2022). A comprehensive understanding of code-mixed language semantics using hierarchical transformer (Doctoral dissertation, IIIT-Delhi).
2. Banerjee, S., Moghe, N., Arora, S., & Khapra, M. M. (2018). A dataset for building code-mixed goal oriented conversation systems. arXiv preprint arXiv:1806.05997.
3. Birshert, A., & Artemova, E. (2021, December). Call Larisa Ivanovna: Code-Switching Fools Multilingual NLU Models. In *International Conference on Analysis of Images, Social Networks and Texts* (pp. 3-16). Cham: Springer International Publishing.
4. De Leon, F. A. L., Madabushi, H. T., & Lee, M. (2024). Code-Mixed Probes Show How Pre-Trained Models Generalise On Code-Switched Text. arXiv preprint arXiv:2403.04872.
5. Dowlagar, S., & Mamidi, R. (2021, September). A pre-trained transformer and CNN model with joint language ID and part-of-speech tagging for code-mixed social-media text. In *Proceedings of the International Conference on Recent Advances in Natural Language Processing (RANLP 2021)* (pp. 367-374).
6. Dowlagar, S., & Mamidi, R. (2023). A code-mixed task-oriented dialog dataset for medical domain. *Computer Speech & Language*, 78, 101449.
7. Firdaus, M., Ekbal, A., & Cambria, E. (2023). Multitask learning for multilingual intent detection and slot filling in dialogue systems. *Information Fusion*, 91, 299-315.
8. Gupta, A., Deng, O., Kushwaha, A., Mittal, S., Zeng, W., Rallabandi, S. K., & Black, A. W. (2021, December). Intent recognition and unsupervised slot identification for low-resourced spoken dialog systems. In *2021 IEEE Automatic Speech Recognition and Understanding Workshop (ASRU)* (pp. 853-860). IEEE.
9. Kanakagiri, T., & Radhakrishnan, K. (2021, April). Task-oriented dialog systems for dravidian languages. In *Proceedings of the First Workshop on Speech and Language Technologies for Dravidian Languages* (pp. 85-93).
10. Krishnan, J. (2021). Unsupervised Cross-Domain and Cross-Lingual Methods for Text Classification, Slot-Filling, and Question-Answering (Doctoral dissertation, George Mason University).
11. Krishnan, J., Anastasopoulos, A., Purohit, H., & Rangwala, H. (2021). Multilingual code-switching for zero-shot cross-lingual intent prediction and slot filling. arXiv preprint arXiv:2103.07792.
12. Nelatoori, K. B., & Kommanti, H. B. (2024). Toxic comment classification and rationale extraction in code-mixed text leveraging co-attentive multi-task learning. *Language Resources and Evaluation*, 1-30.
13. Qin, L., Chen, Q., Zhou, J., Li, Q., Lu, C., & Che, W. (2024, August). Decoupling Breaks Data Barriers: A Decoupled Pre-training Framework for Multi-Intent Spoken Language Understanding. In *Proceedings of the Thirty-Third International Joint Conference on Artificial Intelligence* (pp. 6469-6477).

14. Ramaneswaran, S., Vijay, S., & Srinivasan, K. (2022, May). TamilATIS: dataset for task-oriented dialog in Tamil. In *Proceedings of the Second Workshop on Speech and Language Technologies for Dravidian Languages* (pp. 25-32).
15. Saha, T., Priya, N., Saha, S., & Bhattacharyya, P. (2021, July). A transformer based multi-task model for domain classification, intent detection and slot-filling. In *2021 International Joint Conference on Neural Networks (IJCNN)* (pp. 1-8). IEEE.
16. Sarveswaran, K., Thapa, S., Shams, S., Vaidya, A., & Bal, B. K. (2025, January). A brief overview of the first workshop on challenges in processing south asian languages (chipsal). In *Proceedings of the First Workshop on Challenges in Processing South Asian Languages (CHiPSAL 2025)* (pp. 1-8).
17. Sengupta, A., Suresh, T., Akhtar, M. S., & Chakraborty, T. (2022). A Comprehensive Understanding of Code-mixed Language Semantics using Hierarchical Transformer. *arXiv preprint arXiv:2204.12753*.
18. Suresh, T., Sengupta, A., Akhtar, M. S., & Chakraborty, T. (2024). A Comprehensive Understanding of Code-Mixed Language Semantics Using Hierarchical Transformer. *IEEE Transactions on Computational Social Systems*.
19. Zailan, A. S. M., Teo, N. H. I., Abdullah, N. A. S., & Joy, M. (2023). State of the Art in Intent Detection and Slot Filling for Question Answering System: A Systematic Literature Review. *International Journal of Advanced Computer Science & Applications*, 14(11).
20. Zhang, D. Y., Hueser, J., Li, Y., & Campbell, S. (2021, December). Language-agnostic and language-aware multilingual natural language understanding for large-scale intelligent voice assistant application. In *2021 IEEE International Conference on Big Data (Big Data)* (pp. 1523-1532). IEEE.