



Received: 18/09/2025

Accepted: 24/09/2025

Published: 20/10/2025

Citation: Patel R, Prajapati B (2025) Efficient Deep Learning Ensemble Approach for Indic Named Entity Recognition in Natural Language Processing. Indian Journal of Science and Technology 18(38): 3083-3090. <https://doi.org/10.17485/IJST/V18I38.1602>

* **Corresponding author.**

patelravigs32@gmail.com

Funding: None

Competing Interests: None

Copyright: © 2025 Patel & Prajapati. This is an open access article distributed under the terms of the [Creative Commons Attribution License](#), which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

Published By Indian Society for Education and Environment ([iSee](#))

ISSN

Print: 0974-6846

Electronic: 0974-5645

Efficient Deep Learning Ensemble Approach for Indic Named Entity Recognition in Natural Language Processing

Ravi Patel^{1*}, Bhagirath Prajapati²

¹ Faculty of Computer Engineering, CVM University, Vallabh Vidhyanagar, Gujarat, India

² Department of Computer Engineering, A.D Patel Institute of Technology, CVM University, Vallabh Vidhyanagar, Gujarat, India

Abstract

Background: Named Entity Recognition, often called NER is identifying and then categorizing Named Entity in to predefined entities like People, Locations, and Organizations etc...For Indian language this task is very complex as very low resources available for the development of NLP models. This task is significant for any NLP application as it will directly affect the performance due to inaccurate or mis classification into the predefined entities. **Methods:** For Gujarati NER, Instead of single, but strong BiLSTM model, proposed method is using pre-trained word embedding Indic FastText, Multiple weak BiLSTM models and then majority voting based approach for NER tags. All models are trained on Naamapadam Gujarati dataset. **Findings:** For identifying and classifying low resource Gujarati tags, proposed method is significantly improving performance of NER on Naamapadam dataset. Single and multiple BiLSTM model with pre trained FastText embedding and hard voting based ensemble method have achieved 83.51 and 86.29 F1score respectively. This shows that ensemble approach significantly improve the classification task of named Entities. **Novelty:** To improve the performance of Named Entity Recognition for Gujarati language, proposed method is using hard voting based ensemble method. This multi layer and multiple BiLSTM model approach effectively improve the accuracy and outperform the established NER model that is built for western language like English.

Keywords: Named Entity Recognition (NER); Natural Language Processing (NLP); Bidirectional Long Short-Term Memory (BiLSTM); Ensemble Approach; Pre-Trained Word Embedding

1 Introduction

Identifying and categorizing Named Entities from the text are called Named Entity Recognition. It is one of the most significant step for any Natural Language Processing application. English is one of the most popular language for any research in Named Entity Recognition as there are many NLP applications work on English language⁽¹⁾.

Recently many researches have been done for NER on Indian languages, also called Indic languages. Due to lack of data resources, many Indian languages are also referred as Low Resource Language.

In India, there are more than twenty two official languages and many dialects, spoken by people from different region of India. More than half population of India do not prefer English language — instead, they use Gujarati, Hindi, Tamil, Bengali, Telugu, Marathi, etc. To access the critical applications like healthcare, Government scheme, banking sector etc., robust NER system must be developed

In this paper we will discuss about one Indic Language, Gujarati. Gujarati is one of them as there are very few resources available to train any NLP models. This paper is focusing on NER in Gujarati language as very few researches have been done to enhance Gujarati NER. Before going through the details of Gujarati NER, here are some complexities⁽²⁾ of any Indian Low resource languages that needs to be addressed:

Free word order: Gujarati is free order language unlike English which follows a fix word order language. For English language, subject should come first then verb and lastly object to form meaningful sentence. But that is not true for the Gujarati language as Gujarati grammar allows the free words structures. For example, any of the following sentences have same contextual meaning: “અમૂલ ડેરી આણંદમાં છે” or “અમૂલ ડેરી છે, આણંદમાં.” or “આણંદમાં અમૂલ ડેરી છે”

Capitalization issue: Named entity are often identified with the Capital word of any sentence in English language. There is no such common form for Gujarati language.

Word meaning context: Word meaning ambiguity is one of the major issue in Gujarati language for Example: “ભગવાનને પૂરાઈના ગમે છે.” Here “પૂરાઈના” should not be identified as “People” NER tag. Lack of standard list of predefined Named Entities may be a major issue for any NER systems. Very few tools are available for processing low resource languages NER like Gujarati POS Tagger, Word embedding etc.

Gujarati is one of the most popular language from the western part of India and spoken by around 50 milion people⁽³⁾. Adoption of application in any language rather than native language is difficult for the people of the rural India. Many researchers have tried to implement NLP application for the ease of people understanding. But high accuracy and performance is still compromised for the application as very narrow research is are going on this native language. In this research paper, considering above and complexities and sources of Gujarati language, we have implemented two different models for the one of the most crucial task of NLP that is Named Entity Recognition. This method is utilizing less computational power as compared to the transformer based model and outperformed the existing methods in terms of F1 Score results.

1.1 Named Entity Recognition Methods

Named Entity Recognition Methods for any languages can be classified based on approaches⁽⁴⁾ used. (i) Rule based approach. (ii) Machine learning approach (iii) Hybrid approach and (iv) Deep learning approach. For European language like English, researchers have explored all the approaches on different dataset and have compared the results of NER. Researchers have to consider some factors while selecting and implementing NER approach like application area, significance of the application, system configurations or computational power etc... In this paper, we have discussed some of the implemented NER models for Indian languages. However, we have also discussed some of the NER models which are implemented for other languages like Arabic, Urdu etc.

1.1.1. Rule based approach

Rule-based Named Entity Recognition (NER) is a technique that is utilized for the purpose of recognizing and categorizing entities in text into predetermined categories. These categories include the names of people, businesses, locations, and other specific Entities. NER is founded on a collection of rules and patterns that are derived from linguistic characteristics, as opposed to statistical learning methods. This method achieves high precision in specific domains due to tailored rules, as defined on Electronic Health Records (EHR) where rule-based systems outperformed others in accuracy. But this approach requires extensive human effort to define and maintain rules, making it less adaptable to new domains or languages.

1.1.2. Machine learning approach

In this approach, algorithm used for NER will learn from data, identifies the patterns without explicit programming or manual effort for rules. They learn to predict named entities by analyzing data that has been labelled. The ability of these systems to handle large datasets and sophisticated patterns with minimal feature engineering make them popular over the rule based approach. Conditional Random Fields (CRF), Support Vector Machines (SVM) and Decision Trees are some of the most popular and widely used Algorithm for NER. These methods often struggle with high-dimensional data and may not perform well on tasks with limited annotated datasets.

1.1.3. Hybrid approach

The goal of these strategies is to combine rule-based, statistical, and machine learning approaches in order to extract the most beneficial aspects of each of these methods. Especially useful when extracting entities from a wide variety of sources, they provide the flexibility of different approaches, which makes them particularly beneficial. These methods often combine deep learning techniques with traditional models, resulting in improved accuracy and efficiency in entity classification. The combination of BiLSTM, BERT and CRF are one of the best example of Hybrid approach. Complexity of the models and extensive computation power used to process the data are major drawbacks of these kind of models.

1.1.4. Deep learning approach

The application of deep learning techniques for Named Entity Recognition (NER) has made significant developments in the field of Natural Language Processing (NLP) by enhancing the accuracy as well as efficiency of entity classification methodologies. For the purpose of improving the performance of NER across a variety of languages and domains, a number of models, such as LSTM, BiLSTM, and hybrid architectures, have been employed. In Feed forward Neural Networks Multi-Layer Perceptron (MLP) have been utilized that improves the accuracy for NER tasks⁽⁵⁾. Recurrent Neural Networks like LSTM and Gated Recurrent Units (GRU) are prominent due to their capabilities to capture sequential dependencies in text. These models are often integrated with transformer based model like BERT and CRF model to achieve higher F1 Score⁽⁶⁾. Data availability and the complexity of certain languages make NER task challenging in deep learning approach.

For Hindi, authors⁽⁷⁾ implemented NER with different Language models with eleven different predefined tags, achieved 88.87 F1 Score. But with only three collapsed tag set, achieved F1 score is 92.22. In, ⁽⁸⁾ authors have implemented NER using different models with pre trained embedding. BiLSTM only, BiLSTM-CRF and BiLSTM-CNN-CRF have achieved F1 score 61.9, 67.2 and 70 respectively for Hindi NER. Same authors in ⁽⁹⁾ have used Multilingual Representations for Indian Languages, MuRIL ⁽¹⁰⁾ and CRF models for Hindi NER and claimed best result with 85.77 F1 Score on ICON 2013. For Marathi language, one of the native language of Maharashtra state of India, authors of⁽¹¹⁾ have implemented pattern matching technique for classification of twelve different NER tags and achieved 72.27 F1 score. Authors of⁽¹²⁾ have introduced the data set for Bhojpuri, Maithili, and Magahi, often called Purvanchal languages of India with twenty two NER annotated tags. These dataset is processed through two different NER model that is based on CRF and LSTM-CNN-CRF models. NER model based on CRF outperformed LSTM-CNN-CRF model for Bhojpuri and Maithili languages while LSTM-CNN-CRF improved NER of Magahi language.

2 Methodology

In this section, different components of the Indic NER and dataset used to train and test the models are explained. These components commonly used for any name entity tag classification process. Along with this, two different methodologies are explained for Gujarati NER.

2.1 NER Dataset

To train the Gujarati NER model, one of the recently developed NER data set Naamapadam⁽¹³⁾ is used. NER dataset for Indic languages for 11 Indic languages from 2 language families. Naamapadam contains 5.7 M sentences and 9.4 M entities across these languages from three categories: Person, Name, and Location. In this paper we have considered only Gujarati Language for NER. Named Entity Recognition (NER) for the Gujarati language has employed diverse methodologies, each tackling the distinct issues presented by this morphologically complex language. The creation of extensive datasets, such as Naamapadam, has substantially advanced this domain. The Naamapadam dataset offers a substantial resource with over 400k sentences and 100k annotated entities across multiple Indian languages, including Gujarati. This dataset facilitates the training of models like IndicNER, which has shown promising results in NER tasks. Major languages cover the Person, Location and Organization entity tag categories This dataset is the most extensive publicly accessible Named Entity Recognition dataset for 11 prominent Indian languages, developed by utilizing the Samanantar parallel corpus to translate entities from English to Indian languages. For Gujarati Language total 476K sentences are tokenized with NER tags. Given Dataset is already partitioned into training, testing and validation dataset. In our both implemented strategies, we have kept training, testing and validation dataset same as given in dataset. The Naamapadam dataset employs a standard Named Entity Recognition (NER) tagging schema, classifying entities as "Person" (PER), "Location" (LOC), and "Organization" (ORG), with each entity denoted using the IOB (Inside, Outside, Beginning) format to signify the start and magnitude of the entity within a sentence; it effectively identifies and classifies named entities in Gujarati text. There are very few research work available for NER on Naamapadam Gujarati dataset.

2.2 Components for BiLSTM model based NER

As discussed above there are very few research work available for Gujarati Language NER. In the following section, implementation of Gujarati NER using Deep learning model is discussed. Here, Deep learning model BiLSTM with Pre-word embedding Indic FastText model is used for Gujarati NER. In this paper we have discussed two different models for Gujarati NER. In first method, only single BiLSTM model is used with Indic FastText for word embedding. And in second proposed model, multiple BiLSTM models are used with same word embedding models. Then, we will compare the performances of classifying the predefined named entities for Gujarati input from both the models. Some of the components for proposed systems are discussed here.

2.2.1. Word embedding for Indian language

Word embedding are an important part of natural language processing (NLP) since they provide dense representations of words as vectors, thus capturing their semantic properties. Developing appropriate word embedding for Indian languages, which are diverse and resource scarce, is extremely important for the advancement of NLP task. Currently available Cross-lingual word embedding systems, such as VecMap and MUSE, have demonstrated a good level of accuracy for European languages, but they perform poorly for Indian languages. Character-level Embedding have the ability to capture intra-word information for morphologically rich Indian languages. This allows for improvements in tasks such as part-of-speech tagging by considering morphological segmentation⁽¹⁴⁾. In order to develop Pre-trained Embedding for 14 Indian languages, a variety of methodologies have been utilized. These methods include contextual (for example, BERT) as well as non-contextual approaches. These embedding have been examined on tasks such as Part of Speech tagging and NER, which demonstrates their usefulness with limited resources languages. When compared to publicly accessible embedding, the AI4Bharat-IndicNLP corpus offers pre-trained word embedding for ten Indian languages. These embedding demonstrate improved performance on several evaluation tasks. Word2Vec and Indic FastText models are particularly effective for the Indian languages⁽¹⁵⁾. Recent developments in embedding approaches, such Indic FastText and Indic BERT, have demonstrated potential in enhancing NER performance for these languages. BERT models, especially those fine-tuned for particular languages, have exhibited outstanding efficacy in Named Entity Recognition tasks across various Indian languages. A fine-tuned BERT model attained superior precision and recall for Bengali⁽¹⁶⁾, surpassing generic models employed for Indian languages. BERT has been successfully utilized in Kannada⁽¹⁷⁾, demonstrating its effectiveness in recognizing named items in this language. A study employing a hybrid embedding of FastText and bidirectional LSTM for Punjabi and Hindi shown enhanced precision, recall, and F1 scores, signifying the efficacy of integrating contextual, syntactic, and semantic attributes of text.⁽¹⁸⁾

2.2.2. BiLSTM model

The Bidirectional Long Short-Term Memory (BiLSTM) model is a popular choice for NER tasks in low resource languages. BiLSTM models are efficient for NER tasks as they can capture contextual information from both past and future inputs in a sequence. This is particularly important for languages like Hindi, where context plays a significant role in recognizing named entities. There is the potential for further enhancement of the performance of BiLSTM models with the incorporation of additional layers, such as residual networks or CRF layers. These enhancements contribute to a better ability to capture dependencies, as well as an improvement in the precision and recall of NER. BiLSTM models, particularly when augmented with pre-trained embeddings and supplementary network layers, offer a strong foundation for Named Entity Recognition (NER) in Indian languages. These models successfully tackle the difficulties arising from linguistic complexities and limited resources, resulting in impressive accuracy and F1 scores⁽¹⁹⁾.

2.2.3. NER Ensemble approach

Ensemble approaches for Named Entity Recognition (NER) are strategies that integrate numerous models to enhance the accuracy and robustness of NER tasks. These methods capitalize on the strengths of various or multiple same models to attain superior performance compared to individual models. Bagging, or Bootstrap Aggregating, is an ensemble technique that entails training multiple models independently with various training data. The ultimate projection is generally derived by averaging the forecasts of all models. Random Forest is a commonly used bagging technique that employs decision trees as foundational models⁽²⁰⁾. Boosting is sequentially enhancing the models, with each subsequent model seeking to rectify the weaknesses of previous models. AdaBoost, Gradient Boosting Machines (GBM), and XGBoost are the examples of this method. Boosting may decrease both bias and variation, resulting in enhanced predictive accuracy. Hard Voting ensemble implies consolidating the votes for discrete class labels from models and predicting the class that receives the highest number of votes. Soft Voting is an ensemble method that entails aggregating the projected probabilities for class labels and selecting the class label with the highest cumulative probability. Ensemble techniques frequently surpass baseline classifiers for Named Entity Recognition

in Indian languages, achieving superior recall, precision, and F-measure values compared to individual models and baseline ensembles⁽²¹⁾. In our proposed method, Hard Voting method is implemented to classify the Named Entity based on majority voting.

2.3 FastText - BiLSTM models NER

In this paper, implementation of two different methodologies are discussed. Both methodologies are implemented using BiLSTM models. Our data are in tokenized form and with NER tag from Naamapadam dataset. These data are applied to pre-trained Indic FastText embedding model to generate dense vectors. These Vectors now considered as an input for the classifiers. The general approach for multiple classifier is demonstrated in Figure 1. In this framework multiple but uniform classifiers are considered for NER with ensemble method is called Homogeneous ensemble methods⁽²²⁾. If multiple but different classifiers are taken into consideration for ensemble method then it is defined as Heterogeneous method. In comparison to heterogeneous methods, homogeneous ensemble methods have a number of benefits, particularly with regard to stability, interpretability, and performance consistency. Homogeneous ensembles enhance result interpretability by utilizing a singular classifier, hence facilitating practitioners' comprehension of the underlying decision-making process. Homogeneous ensembles demonstrate enhanced performance as they may be tailored for unique data streams without the complex nature of handling multiple methods.

The general approach defined in Figure 1. With parallel and homogenous ensemble method is considered as a base framework for our Gujarati NER using multiple BiLSTM ensemble method. Here BiLSTM model with softmax layer will classify the named entity for each token.

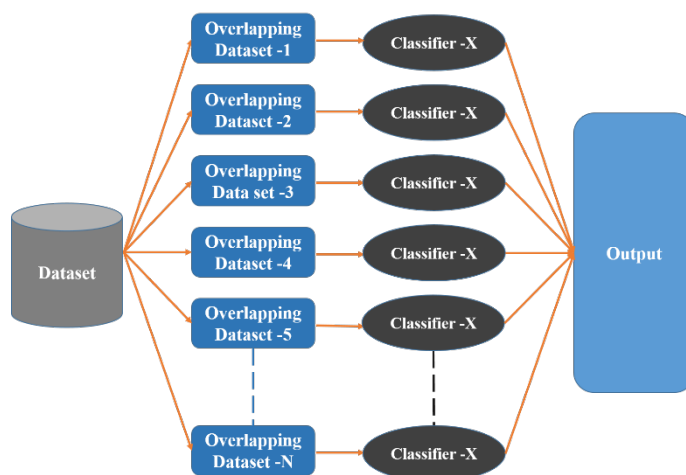


Fig 1. General framework for homogenous ensemble method

There are multiple ways to feed data into the classifiers. a) Same data samples b) Different data samples c) Overlapping data samples. Here overlapping data samples have been applied to the multiple classifier models. When compared to using the same or different data samples, the advantages of employing overlapping data samples in classifiers are significant. These advantages are especially prominent when it comes to fixing class imbalance and improving model performance. In order to give classifiers with the ability to learn more successfully from complex data distributions, overlapping samples can provide deeper information about the decision boundary^(?). Improvements in classification metrics and robustness against noise and imbalance are possible outcomes of this method. By concentrating on the morphology of data categories, classifiers can attain a more fair learning process, resulting in improved generalization on novel data^(?). So for our proposed method discussed in 7.2, above framework is considered as a base framework for Gujarati NER.

2.3.1. NER with Single BiLSTM model

In Single NER BiLSTM model mentioned in Figure 2, we have implemented single but strong BiLSTM model for Gujarati NER with Indic FastText embedding. In this method, Input will be from Naamapadam dataset. Sentences are already tokenized in the dataset so that they are directly applied to the Indic FastText embedding as an input. For each token in the sentence, the output of the Indic FastText are 300 Dimensions dense vectors. In this method, 32 batch size for sentences are considered. Number of

tokens in the sentences may vary so for the equal distribution of the token dimensions, here maximum size of tokens from the longest sentence is considered and according to that zero padding is applied to the remaining sentences token dimensions. This dense vectors are now applied as an input to the two BiLSTM models layers. For each token in the sentence, there is 7 dimension vector as an output of the BiLSTM showing the predicted probabilities for entities. And thus, highest probability value of the token is considered as the output of the predicted entity. In the result section, the results of base model and ensemble method is discussed.

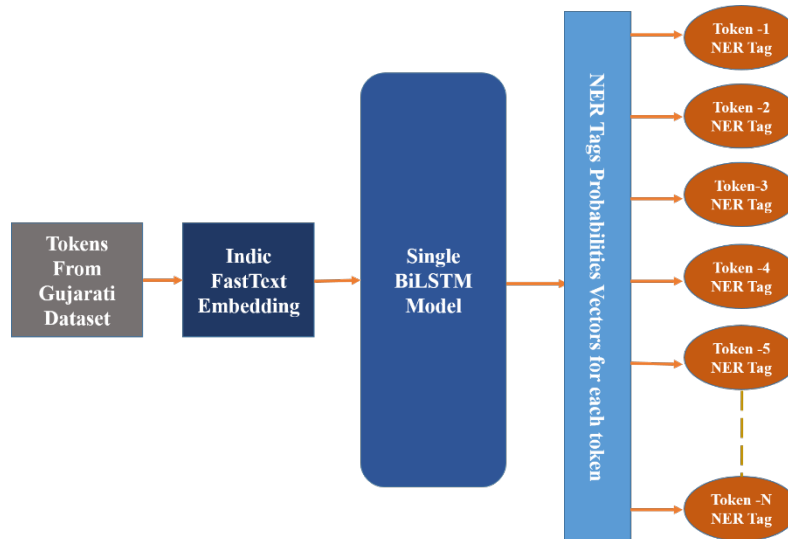


Fig 2. NER with Single BiLSTM Model

2.3.2. NER with Multiple BiLSTM + Ensemble approach

Instead of using single strong BiLSTM, in proposed method multiple weak BiLSTM model is implemented for Gujarati NER. Naamapadam Gujarati dataset is applied to Indic FastText embedding as an input. Vectors, the output of pre-trained embedding model^(?), represented with 300 dimensions are applied to Multiple BiLSTM Models. Instead of applying same vectors to all BiLSTM models as an input, in proposed method, overlapping samples have been applied to single layer multiple BiLSTM models as shown in Figure 3. Individual model have generated 7 dimensions vector for each token. Thus, if there are n number of BiLSTM model, then total n vectors of 7 dimensions for each token have been generated. Here, in proposed method, total ten weak BiLSTM models are used and have generated total ten vectors for each token. It means, now for each individual token, there are total ten predicted probabilities for named entities. To increase the accuracy, in this proposed method majority voting based ensemble method is applied to find best NER tag for the tokens. Majority vote ensemble method have consider all probabilities for each tokens from all BiLSTM models and then majority predicted entities are considered for final output as an NER tag for tokens.

3 Results and Discussion

Two different methodologies for NER on Gujarati language from Naamapadam dataset are implemented for classifying Named Entities. For our both the models, Indic FastText is considered as an Encoder for the input that is in tokenized form and Softmax function is considered as a Decoder also defined as activation function of a BiLSTM to normalize the output of a model to a probability distribution over predicted output NER tags. In this paper, performance comparison of three different models are discussed. For two different model, we have obtained and compared F1 Score for Person, Organization and Location NER tags. One model that is implemented by authors of Naamapadam dataset using BERT-base-multilingual-cased, obtained F1 Score on multiple languages. In this paper, the performance of the BERT based model on Gujarati Language is discussed. In our both models' experimental setup, we have implemented NER with batch size of 32, embedding dimension of 300, 0.0008 learning rate and cross entropy loss function. For Single BiLSTM, We have used BiLSTM model with two layers, each having 256 cells units. For Multiple BiLSTM model, we have implemented NER with 11 BiLSTM models with single layer of 256 cells unit. So that, the single BiLSTM model is strong enough as compared to the multiple BiLSTM model for fair comparison. For BERT

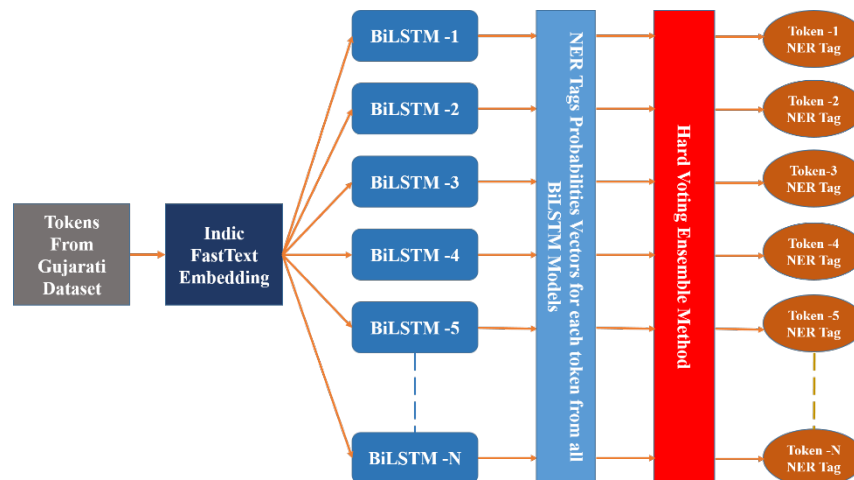


Fig 3. NER with Multiple BiLSTM Model

based model, different Batch size and learning rates have been considered for fine tuning the model. This means, the results of the BERT based model are from the best fine-tuned model. In this paper, as shown in table 1, we have compared the F1 Score results of overall predicted named entity tags.

Table 1. Different Models F1 Score performance for Naamapadam Gujarati dataset

Sr. No	Gujarati NER Models	F1 Score
1.	BERT-base-multilingual-cased model	80.83
2.	Indic FastText + single BiLSTM model	83.51
3.	Indic FastText + Multiple BiLSTM model ensemble method	86.29

Here, Both BiLSTM based proposed model have shown significant improvement on Gujarati NER. Even Multiple weak BiLSTM ensemble method has outperformed the single strong BiLSTM model. Many research work demonstrated that transformer based NER outperformed the FastText Model. As discussed earlier, our both methods have used Pre trained Indic FastText word embedding model and outperformed the transformer based BERT embedding models when it is used as an encoder for the input token. Both of our implemented model i.e. single BiLSTM model and Multiple BiLSTM with ensemble approach, have achieved 97% and 98% accuracy respectively for identifying and classifying Gujarati Named Entity.

4 Conclusion

For Indic Named Entity Recognition so far many approach have been implemented to enhance the performance for native language based NLP applications. Identifying and classifying named entities based on predefined classes for Indian language is very complex task due to lack of resources. Here, we have discussed one of the Indian language, Gujarati, popular in western part of India for Named Entity Recognition. No robust methodologies have been implemented for Named Entity classification due to the structure and rich morphology of Gujarati language. This leads to poor performance of any Gujarati NLP applications. In this paper, both of our proposed methods significantly improving F1 Score and outperformed the previous work done on Naamapadam dataset. Indic FastText + BiLSTM Model is significantly improving NER F1 Score for Gujarati by approximately 3%, and Multi BiLSTM ensemble model has outperformed both by improving F1 Score by approximately 6% from BERT based Model and almost 3% form Single Strong BiLSTM model. This study have demonstrated that instead of using single strong model, multiple weak models with ensemble method have improved the performance of the NER. In this paper, homogeneous models have been used and uncover the different challenges to explore the research on some issues like utilization of computational power, performance of different models (Heterogeneous), and performance of different pre-trained embedding models like transformer based model-BERT for Gujarati NER. To enhance the Indic NER model's performance, future research could explore additional improvements. Some of the recent study shows that CRF enhance the NER task for low resource languages. CRF could improve the performance of multiclass classification especially for rich morphological dataset. Advancement in any NLP application in native language focuses on the robust performance of NER task.

5 Acknowledgements

The authors thank the CVM University, Vallabh Vidhyanagar, Gujarat, India for providing support and infrastructure facilities to carry out this research work.

References

- 1) Seow WL, Chaturvedi I, Hogarth A. A review of named entity recognition: from learning methods to modelling paradigms and tasks. *Artificial Intelligence Review*. 2025;58:315. <https://doi.org/10.1007/s10462-025-11321-8>.
- 2) Desai NP, Dabhi VK. Resources and components for gujarati NLP systems: a survey. *Artificial Intelligence Review*. 2022;55:1–19. <https://doi.org/10.1007/s10462>.
- 3) Baxi J, Bhatt B. GujMORPH - A Dataset for Creating Gujarati Morphological Analyzer. *Proceedings of the Thirteenth Language Resources and Evaluation Conference*. 2022;p. 7088–7095. Available from: www.aclanthology.org/2022.lrec-1.767/.
- 4) Liu X, Chen H, Xia W. Overview of Named Entity Recognition. *Journal of Contemporary Educational Research*. 2022;6(5):65–68. <https://doi.org/10.26689/jcer.v6i5.3958>.
- 5) Kayalvili S, Sree MP, Priyadarshini C, Motheeswaran K. Enhancing Named Entity Recognition using Deep Learning Approaches. *5th International Conference on Electronics and Sustainable Communication Systems*. 2024;p. 1733–1737. <https://doi.org/10.1109/ICESC60852.2024.10690015>.
- 6) Liu B, Chen W, Tao J, He L, Tang D. Research on Named Entity Recognition Based on Gated Interaction Mechanisms. *Applied Sciences*. 2024;14(15):6481. <https://doi.org/10.3390/app14156481>.
- 7) Murthy R, Bhattacharjee P, Sharnagat R, Khatri J, Kanojia D, Bhattacharyya P. HiNER: A large Hindi named entity recognition dataset. *Proceedings of the Thirteenth Language Resources and Evaluation Conference*. 2022;p. 4467–4476. Available from: <https://aclanthology.org/2022.lrec-1.475/>.
- 8) Sharma R, Morwal S, Agarwal B, Chandra R, Khan MS. A deep neural network-based model for named entity recognition for Hindi language. *Neural Computing and Applications*. 2020;32(20):16191–16203. <https://doi.org/10.1007/s00521-020-04881-z>.
- 9) Sharma R, Morwal S, Agarwal B. Named entity recognition using neural language model and CRF for Hindi language. *Computer Speech & Language*. 2022. <https://doi.org/10.1016/j.csl.2022.101356>.
- 10) Khanuja S, et al. Muril: Multilingual representations for Indian languages. *arXiv preprint arXiv:2103.10730*. 2021. <https://doi.org/10.48550/arXiv.2103.10730>.
- 11) Patil NV. An Emphatic Attempt with Cognizance of the Marathi Language for Named Entity Recognition. *Procedia Computer Science*. 2023;p. 2133–2142. <https://doi.org/10.1016/j.procs.2023.01.189>.
- 12) Mundotiya R, Kumar S, Kumar A, Chaudhary U, Chauhan S, Mishra S, et al. Development of a Dataset and a Deep Learning Baseline Named Entity Recognizer for Three Low Resource Languages: Bhojpuri, Maithili, and Magahi. *ACM Transactions on Asian and Low-Resource Language Information Processing*. 2023;22(1):29. <https://doi.org/10.1145/3533428>.
- 13) Mhaske A, Kedia H, Doddapaneni S, Khapra MM, Kumar P, Murthy R, et al. Naamapadam: A Large-Scale Named Entity Annotated Data for Indic Languages. *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*. 2023;p. 10441–10456. <https://doi.org/10.18653/v1/2023.acl-long.582>.
- 14) Goswami S, Malviya R, Mishra, Tiwary US. Analysis of Word-level Embeddings for Indic Languages on AI4Bharat-IndicNLP Corpora. *2021 IEEE 8th Uttar Pradesh Section International Conference on Electrical, Electronics and Computer Engineering (UPCON)*. 2021;p. 1–4. <https://doi.org/10.1109/UPCON52273.2021.9667615>.
- 15) Kunchukuttan A, Kakwani D, Golla S, GokulN, Bhattacharyya A, Khapra M, et al. AI4Bharat-IndicNLP Corpus: Monolingual Corpora and Word Embeddings for Indic Languages. *arXiv preprint*. 2020. <https://doi.org/10.48550/arXiv.2005.00085>.
- 16) Mohan GB, Kumar RP, R E, Pendem MVSS. Fine-Tuned BERT Based Multilingual Model for Named Entity Recognition in Native Indian Languages. *2023 International Conference on Evolutionary Algorithms and Soft Computing Techniques (EASCT)*. 2023;p. 1–6. <https://doi.org/10.1109/EASCT59475.2023.10392937>.
- 17) Hebbar S, Rao BA, Muralikrishna SN, Supriya M, Narendra VG, Sobha L. Named Entity Recognition Using BERT Model for Kannada Language. *2023 International Conference on Recent Advances in Information Technology for Sustainable Development*. 2023;p. 212–216. <https://doi.org/10.1109/ICRAIS59684.2023.10367119>.
- 18) Shelke R, Vanjale S. Deep named entity recognition in Hindi using neural networks. *Revue d'Intelligence Artificielle*. 2022;36(4):575–580. <https://doi.org/10.18280/ria.360409>.
- 19) Aljawazneh H, Mora AM, García-Sánchez P, Castillo-Valdivieso PA. Comparing the Performance of Deep Learning Methods to Predict Companies' Financial Failure. *IEEE Access*. 2021;9:97010–97038. <https://doi.org/10.1109/ACCESS.2021.3093461>.
- 20) Hou Z, et al. A Boosting Approach to Constructing an Ensemble Stack. *Genetic Programming EuroGP 2023*. 2023;13986. Lecture Notes in Computer Science.
- 21) Hassan F, Tufa W, Collell G, Vossen P, Beinborn L, Flanagan A, et al. SeqL at SemEval-2022 Task 11: An Ensemble of Transformer Based Models for Complex Named Entity Recognition Task. *Proceedings of the 16th International Workshop on Semantic Evaluation (SemEval-2022)*. 2022;p. 1583–1592. <https://doi.org/10.18653/v1/2022.semeval-1.218>.
- 22) Liang Y, Gharipour A, Kelemen E, Kelemen A. Homogeneous Ensemble Feature Selection for Mass Spectrometry Data Prediction in Cancer Studies. *Mathematics*. 2024;12:2085. <https://doi.org/10.3390/math12132085>.