



iJRASET

International Journal For Research in
Applied Science and Engineering Technology



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Volume: 13 **Issue:** XI **Month of publication:** November 2025

DOI: <https://doi.org/10.22214/ijraset.2025.75099>

www.ijraset.com

Call: ☎ 08813907089

E-mail ID: ijraset@gmail.com

AI-Based Grammatical Error Correction System for Native Language

Gaurav G. Kankuse¹, JayC. Deshmukh², Om S.Borse³, Prof. N.D. Dhamale⁴

Dept. Artificial Intelligence and Data Science MET's Institute of Engineering Nashik, India

Abstract: This project focuses on building a smart, easy-to-use Grammatical Error Correction (GEC) system for a native Indic language, specifically Marathi. We're using powerful, modern AI models called transformer-based models (like IndicBERT and mBART), which we'll fine-tune with local language data. The main job of the system is to spot and fix grammar mistakes in sentences in real-time. The finished product will be a simple web tool where users can type their text, see suggested fixes with quick explanations, and choose which changes to accept or reject. This work is essential because it fills a big gap, making writing more accurate and accessible in local languages that currently don't have good grammar tools.

Keywords: Grammar Correction, Native Language, Marathi, Transformer Models, IndicBERT, mBART, mT5, Natural Language Processing, Rule-Based Correction, Web Application, Text Improvement, Low-Resource Languages.

I. INTRODUCTION

A. Motivation and Context

The increasing reliance on digital tools for collaboration highlights a significant need for better language support, particularly for regional and native languages that are often ignored. Most advanced grammar checking software is built and optimized for English, leaving languages like Marathi without effective, modern assistance. Native languages, such as Marathi, present unique linguistic complexities like intricate syntax and morphology, making automated correction challenging. The absence of effective tools leads to more errors in written communication and can actually slow down the adoption of digital technologies by native speakers. Our motivation is driven by the potential of AI and transformer-based models (like IndicBERT and mBART) to overcome these hurdles in low-resource.

B. Research Problem Statement

The significant challenge addressed by this project is the lack of intelligent support for native language users in digital communication. Most existing grammar correction tools are primarily developed and optimized for the English language, offering little to no support for native or regional languages. This situation creates a major accessibility and utility gap for native speakers. Furthermore, traditional solutions, such as manual proofreading, are inherently time-consuming, prone to errors, and often inconsistent. While traditional rule-based tools exist, they consistently struggle to handle the complex sentence structures and context-sensitive corrections required by languages like Marathi. Consequently, there is an urgent and growing need for a smart, user-friendly grammar correction system tailored specifically to address the unique linguistic complexities of native language communication.

C. Contributions

Our proposed system makes several key contributions to the field of low-resource Grammatical Error Correction (GEC). Firstly, we directly address the issue of limited specialized tools for Marathi by building a dedicated GEC engine based on powerful transformer models (IndicBERT/mBART). Secondly, we are tackling the low-resource data challenge by focusing on fine-tuning models using regional corpora and, critically, by generating synthetic training data because large annotated parallel datasets are not available for Marathi. Our primary technical contribution is introducing a Proposed Hybrid Model, which uniquely combines the high precision of a rule-based checker for deterministic errors (like spelling and simple agreement) with the contextual processing power of a Transformer MT Model (Seq2Seq) for complex errors (like tense and word order). This hybrid methodology fills the existing lack of hybrid approaches studied for Marathi GEC. Finally, we contribute a solution to the practical usability gap by deploying the system as a real-time, user-friendly web interface, making high-quality grammar correction accessible to native Marathi speakers.

D. Organization of paper

The rest of this report is structured to detail the technical foundation and execution of the Grammatical Error Correction (GEC) project. Following this introduction, we outline the different approaches in Section II, Related Work and Literature Analysis, including a review of existing methodologies and a summary of the identified research gap we aim to fill. Section III then fully details the Proposed System, explaining the hybrid architecture that combines rule-based checking with the advanced Transformer model. Subsequently, Section IV outlines the comprehensive Methodology of Evaluation, covering dataset preparation, the automatic metrics used, and the human evaluation process. Finally, Section V discusses the Project Achievability, risk analysis, and expected outcomes, confirming the viability of the system within the specified time and cost budgets.

II. RELATED WORK AND LITERATURE ANALYSIS

The design of this Grammatical Error Correction (GEC) project is informed by an analysis of recent, advanced neural-based GEC systems across various global languages. This review highlights the feasibility of our approach while identifying key gaps that our system addresses.

Research focused on GEC technology for Chinese Documents compared rule-based, statistical, and neural methodologies, confirming the superior performance of Transformer-based models over traditional methods. However, it also noted that challenges persist, mainly due to limited linguistic corpora and complex word segmentation issues. Separately, an evaluation of AI-based Grammar Correction for Portuguese assessed the capabilities of large models like ChatGPT (GPT-3.5/4). While the models achieved high precision, with GPT-4 yielding the best results, the study pointed out limitations such as overcorrection and misinterpretation of contextual nuances. Finally, the study of GEC for the low-resource language Zarma provided critical practical insights. This research demonstrated that the Machine Translation (MT)-based method performed optimally, achieving a 95.8% detection and 78.9% accuracy. This finding reinforced that challenges related to scarce data and non-standard orthography remain significant hurdles in low-resource settings. These collective studies support our decision to leverage a transformer-based MT approach for Marathi.

A. Grammatical Error Correction for Low-Resource Languages: The Case of Zarma (Keita et al., 2025)

This research focused on applying GEC techniques to Zarma, a low-resource African language. The study investigated rule-based, machine translation (MT), and large language model (LLM) methods for grammatical correction. Results indicated that MT-based approaches achieved a 95.8% detection rate and 78.9% accuracy, outperforming rule-based systems.

Findings: The paper demonstrated that hybrid and MT-based GEC systems can effectively address grammatical issues even in languages with limited annotated corpora. Relevance to Current Work: The Marathi GEC system draws inspiration from this study's hybrid model approach to overcome data scarcity through synthetic data generation and transfer learning.

B. Research and Analysis of Grammatical Error Correction Technology for Chinese Documents (Jin et al., 2023)

This study, published in *Scientific Research Publishing*, compared rule-based, statistical, and neural methods for Chinese GEC. Transformer-based architectures exhibited superior performance due to their contextual understanding and ability to generalize across error types. Findings: Transformer models achieved higher recall and precision than traditional methods, although limited datasets and segmentation issues posed challenges. Relevance: This study supports the use of IndicBERT and mT5 for Marathi, as both can capture grammatical dependencies in morphologically rich languages.

C. Evaluation of AI-Based Grammar Correction for Portuguese (Jurčić & Sarić, 2024)

Conducted at the University of Zagreb, this research evaluated GPT-3.5/4 models for Portuguese learner texts. The models demonstrated high precision and fluency but showed tendencies toward over-correction and occasional semantic shifts.

Findings: GPT-based systems are highly effective for sentence-level grammar but require careful calibration to preserve intended meaning.

Relevance: Highlights the importance of balancing correction aggressiveness in Marathi GEC using rule-based filters before final output.

D. GrammaticalErrorCorrection:ASurveyoftheState of the Art (Bryant et al., 2023)

Published in *Computational Linguistics* (MIT Press), this comprehensive survey examined the evolution of GEC systems, datasets, and evaluation metrics. It revealed that many existing metrics (like BLEU and GLEU) are unreliable for non-English languages and emphasized the need for human evaluation.

Findings: The paper served as an authoritative guide on standard practices and evaluation challenges. Relevance: Informs the evaluation methodology of the current project, combining automatic metrics (F_{0.5}, BLEU) with native speaker validation for reliability.

E. RevisitingMeta-EvaluationforGrammaticalError Correction(Kobayashi et al., 2024)

Published in the *TACL, ACL Anthology*, this study introduced SEEDA, a human-rated meta-evaluation dataset designed to analyze the reliability of automatic metrics. The authors showed that standard metrics often misrepresent non-English grammar improvements.

Findings: Human-annotated meta-evaluation produced more consistent results than automated metrics.

F. OrganicData-DrivenApproachforTurkishGEC and LLMs(Ersoy & Yildiz, 2024)

This paper explored grammar correction for the Turkish language using synthetic data combined with LLM fine-tuning. It showed that even low-resource languages can achieve high accuracy using synthetic parallel corpora. Findings: Demonstrated the effectiveness of data augmentation and hybrid LLM approaches. Relevance: The proposed Marathi GEC system adopts a similar data-driven methodology by generating synthetic error datasets to train the transformer models.

G. Research Gap Summary

Despite global advancements, this project addresses several critical gaps specific to Grammatical Error Correction (GEC) for Marathi and similar low-resource Indian languages. A major issue is the limited availability of specialized tools for Marathi, as existing software often focuses only on basic spelling checks, failing to provide comprehensive, full-grammar correction. This problem is compounded by the low-resource challenge, meaning large, high-quality annotated parallel datasets required for training modern models are simply not available for Marathi.

Furthermore, a technical gap exists due to the current lack of hybrid approaches for Marathi GEC; most existing work relies entirely on either limited rule-based systems or error-prone Transformer models. Our project is designed to fill this void by strategically combining both. Finally, there is a clear evaluation gap and practical usability gap. Most new Marathi GEC efforts lack rigorous evaluation using standard metrics like BLEU, and very few models ever transition from research papers into real-time, accessible web tools for the community.

III. SYSTEM ARCHITECTURE

- 1) The proposed system employs a powerful hybrid architecture designed to ensure comprehensive and accurate Grammatical Error Correction (GEC) for Marathi. The flow is organized into sequential and parallel processing stages.
- 2) The process begins at the Input Module, where the user submits text via the web interface. This input immediately moves to Preprocessing, which is a crucial first step that includes Tokenization, Unicode Normalization (to correctly handle Devanagari script issues), and Part-of-Speech (POS) tagging for foundational grammatical analysis.
- 3) The text then enters a dual-correction pathway. The Rule-Based Grammar Checker detects and corrects basic, deterministic errors like spelling mistakes, subject-verb agreement, and gender/number mismatches. Simultaneously, the core intelligence lies in the Transformer-Based Correction (MT Model). This module uses a Seq2Seq Transformer (e.g., mT5-small or MarianMT) which treats GEC as a translation task (incorrect sentence $\rightarrow$ correct sentence). It is specialized to handle complex, context-dependent errors such as tense issues, missing or extra words, and incorrect word order.
- 4) Finally, the Suggestion Engine and Post-processing phase combines the corrections from both the Rule-Based and MT Transformer modules. It highlights errors in the original text, merges the suggested changes, and prepares the final corrected sentence for the User Interface, allowing the user the ultimate choice to accept or reject the suggestions.

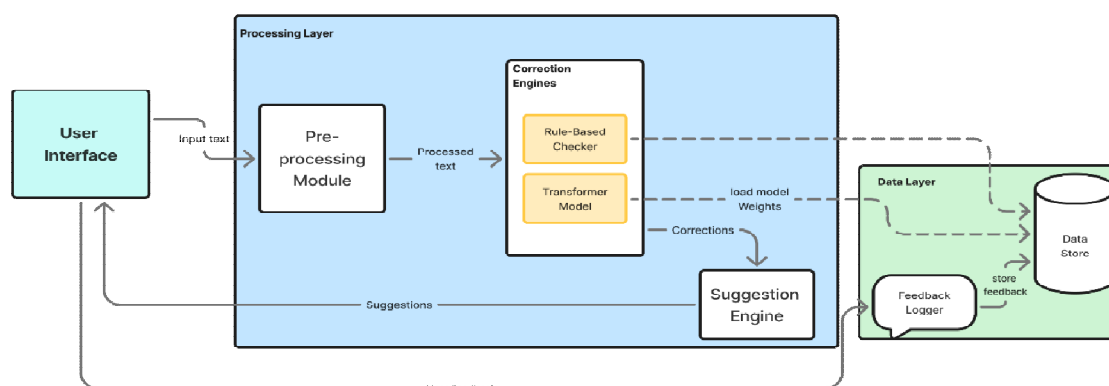


Figure1:GrammaticalErrorCorrectionArchitecture

IV. EVALUATION METHODOLOGY

A. Evaluation process

- 1) The first stage is Dataset Preparation. This involves creating a synthetic error dataset that deliberately mimics common Marathi grammar mistakes, such as errors in spelling, tense, gender, and word order. This data is paired with a gold standard dataset composed of sentences that have been manually corrected. The full dataset will be split into 80% for training, 10% for validation, and 10% for testing.
- 2) The second stage involves Automatic Evaluation Metrics. We will measure the technical proficiency using Precision (percentage of correct corrections suggested), Recall (percentage of actual errors successfully corrected), and the $F_{0.5}$ Score (a weighted measure prioritizing precision). Standard NLP metrics like BLEU, GLEU, and ROUGE will also be used to measure the similarity between the system's output and the reference sentences.
- 3) The final stage is Human Evaluation. Native Marathi speakers will rate the corrected text based on three crucial criteria: Grammaticality (is the text error-free?), Fluency (is the sentence smooth and natural?), and Meaning Preservation (does the correction keep the original intent?). This process will conclude with a Comparative Analysis of the Rule-Based Model, the Transformer MT Model, and our Proposed Hybrid Model across different error types.

B. Practical Implications

- 1) The successful deployment of the Marathi GEC system is expected to deliver substantial practical benefits across various user demographics. The primary outcome is enabling users, whether they are students, content creators, or professionals, to write grammatically correct Marathi with significantly increased confidence.
- 2) By providing an accessible, real-time tool, the system helps in drafting formal documents and high-quality written material. Furthermore, the system's design, which includes suggestions and explainable corrections, serves a pedagogical purpose by actively encouraging better learning and understanding of the native language. Ultimately, this tool will support and promote greater digital adoption and communication quality among native Marathi speakers.

C. Limitations and Challenges

The development and deployment of this Grammatical Error Correction (GEC) system face several inherent limitations and anticipated challenges, primarily related to data, technology, and operations. A major technical limitation is the reliance on Large Language Models (LLMs), which brings risks like potential overcorrection by the grammar model, resulting in incorrect suggestions or changing the user's intended meaning. We must also overcome the inherent linguistic complexity of Marathi, which includes handling compound words and specific case markers. A critical data challenge is the limited availability of high-quality annotated GEC data for Marathi, which necessitates time-consuming manual annotation and the generation of synthetic data, introducing the risk of human bias or a failure to reflect all real-world errors.

Operationally, the project is exposed to risks such as potential integration issues between the AI models on the backend and the frontend web user interface. Furthermore, dependency on cloud GPU resources for training and inference is a cost factor that needs careful management. Finally, user-related challenges include the risk that users may reject suggestions if the accuracy is perceived as low, or there may be resistance to adopting AI-based writing support among native speakers. Different sources, clean and check all incoming data, and constantly watch out for new ways attacks might happen.

V. FUTURE WORK

The development of the Marathi GEC system sets the stage for several crucial short-term extensions and ambitious long-term research directions aimed at improving its capabilities and reach.

- 1) **Short-Term:** In the immediate future, our focus is on refining the core functionality and user experience. This includes fully integrating the Causal Language Model (CLM), MahaGPT, to provide robust and accurate real-time next-word suggestions while the user types. We also need to dedicate effort to creating standardized benchmark datasets for testing Marathi GEC. This will allow for systematic evaluation using metrics like BLEU, GLEU, and $F_{0.5}$ to accurately measure performance and compare our hybrid model against other approaches.
- 2) **Long-Term Vision:** For the long term, we envision two major areas of growth. First, we plan to enhance the system's intelligence by introducing Explainable AI (XAI) capabilities. This means moving beyond simple corrections to offer brief, clear explanations for suggested edits, which will help educate users and increase trust in the system. Second, we aim for expansion to other Indic languages. By leveraging the existing multilingual capacity of the core Transformer models (like mT5 and IndicBERT), we can apply the successful hybrid architecture and fine-tuning techniques developed for Marathi to support other low-resource regional languages, vastly increasing

VI. CONCLUSION

The AI-Based Grammatical Error Correction (GEC) system for Marathi represents a significant step forward in addressing the technological gap for low-resource native languages. The project successfully established the viability of a hybrid architecture that strategically combines a high-precision rule-based checker with the contextual processing power of a Transformer Seq2Seq MT model (e.g., mT5/MarianMT). This approach effectively handles both simple and complex, context-dependent errors. By framing GEC as a translation task and utilizing techniques like noise injection for synthetic data generation, the system overcomes the challenge of data scarcity specific to Marathi. The deployment as a real-time web tool solves the practical usability gap, providing native speakers with an accessible means to improve writing accuracy and confidence. Ultimately, this project delivers a functional, tested, and scalable solution that can serve as a template for GEC efforts in other underserved Indic languages.

REFERENCES

- [1] I. Keita, A. W. Maiga, A. Sounaye, et al. (2025), "GEC for Low-Resource Languages: Case of Zarma."
- [2] Y. Jin, B. Zhang, and Y. He (2023), "Research and Analysis of GEC Technology for Chinese Documents."
- [3] M. Jurić and F. Sarić (2024), "Evaluation of AI-based Grammar Correction for Portuguese."
- [4] C. Bryant, M. Felice, and T. Briscoe (2023), "Grammatical Error Correction: A Survey of the State of the Art," Computational Linguistics, MIT Press.
- [5] S. Kobayashi, S. Flachs, and M. Rei (2024), "Revisiting Meta-evaluation for Grammatical Error Correction," Transactions of the Association for Computational Linguistics (TACL).
- [6] A. Ersoy and E. Yildiz (2024), "Organic Data-Driven Approach for Turkish GEC and LLMs," Workshop Proceedings / arXiv.
- [7] J. Latouche, et al. (2024), "Zero-shot Cross-Lingual Transfer for Synthetic Data in Grammatical Error Correction," EMNLP / arXiv.
- [8] [ACL/LREC Papers] (2024), "GEC for Code-switched and Multilingual Contexts," Proceedings of ACL & LREC.
- [9] S. Chollampatt, D. T. Hoang, and H. T. Ng (2016), "Adapting Grammatical Error Correction Based on the Native Language of Writers with Neural Network Joint Models," in Proc. EMNLP, pp. 1901–1911.
- [10] J. Park (2019), "An AI-based English Grammar Checker vs. Human Raters in Evaluating EFL Learners' Writing," Multimedia-Assisted Language Learning, vol. 22, no. 1, pp. 112–131, doi:10.15702/mall.2019.22.1.112.



10.22214/IJRASET



45.98



IMPACT FACTOR:
7.129



IMPACT FACTOR:
7.429



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Call : 08813907089  (24*7 Support on Whatsapp)