

An Automated Word Spelling Error Detection System

Jennifer, Ogechi Okeh

¹Department of Computer Science, University of Port Harcourt

Abstract:- Accurate spelling has extensive consequences for written communication's clarity, credibility, and general effectiveness in several contexts, including academic and professional writing and personal communications. The frequency of spelling errors has increased significantly recently, mainly owing to the widespread use of digital media and society's dependence on written communication. Misspelled words can diminish the reliability of written correspondence and result in misinterpretations. Hence, there is an urgent requirement for an innovative and advanced solution to improve the identification of spelling errors. This research seeks to address these challenges and develop a system that significantly enhances the quality of written communication by improving the accuracy and relevance of spelling error detection and correction suggestions. According to this study, the implemented model is a promising technique to correct spelling problems in legal terminology, has been praised for its effectiveness and usability. Its recommendations include continuous research and development, improving its ability to use contextual cues and domain-specific information for error correction, emphasizing user-centric customization, implementing collaborative feedback systems, offering extensive training and support, prioritizing ethical and regulatory considerations, and fostering interdisciplinary collaboration. These suggestions aim to optimize the Spello model's performance, resilience, and scalability, and make it a fundamental technology in promoting accuracy, integrity, and accessibility in the legal profession and society. The model can become a fundamental technology in promoting accuracy, integrity, and accessibility in the legal profession and society through combined efforts and collaboration.

Keywords:- Automated, Detection, Error, Word-Checker

I. INTRODUCTION

Communication refers to transmitting information, thoughts, ideas, and emotions between individuals or collectives (Juma'a, 2020). It is a fundamental aspect of human interaction and holds significant importance in several domains, including personal relationships, professional environments, and educational settings. In addition, effective communication occurs when two individuals transmit and receive signals harmoniously, which encompasses verbal and nonverbal methods and can occur in several contexts, including face-to-face contact, written materials, online platforms, and visual or auditory cues. Communication is a dynamic academic field that explores how information is

shared, comprehended, and impacted by various social, cultural, and environmental aspects.

(Xiuqing, 2023) expresses that communication pertains to the mental process of exchanging information. It should be noted that the area encompasses several sub disciplines that examine specific facets of human interaction, such as interpersonal communication, mass media studies, rhetoric, and intercultural communication. (Voloshchuk et al., 2020) reiterate that integrating indigenous, spoken, and written language with elements from various semiotic systems (such as visuals, architecture, sound, video, digital media, etc.) serves as a method of communication and its tools. Indeed, good written communication is crucial for expressing ideas, resolving challenges, developing critical thinking skills, understanding scientific phenomena, and creating such phenomena (Syamsuddin et al., 2021).

Notably, the contents of written communication, whether a formal report, email, social media post, or text message, hold significance since they reflect our thoughts, ideas, and emotions. In addition to this, hence, precise spelling is crucial as it enhances the probability that the intended recipients of a communication will comprehend it without any superfluous disruptions.

However, erroneous words can lead readers to misconstrue and misinterpret intended messages. For individuals employed in sectors that place significant emphasis on accuracy and comprehensiveness, this ought to be a substantial cause for apprehension. Corrections that are poor may give the impression of carelessness or even dishonesty. Also, spelling errors may significantly impact a student's grade and the author's reputation in academic writing. And equally, can alter a message's intended meaning, even in casual contexts. (Aceto et al., 2019), stated that there is a significant change in the way technology is advancing, which is driving a fresh wave of industrial transformation. Presently, the significance of accurate spelling has escalated in the digital age when written communication prevails over spoken discourse; also in challenged physically individuals (Mainsah et al., 2015) ; (Verbaarschot et al., 2021).

This is as a result of advent of digital communication, such as texting, emailing, and social networking, which has fostered a more relaxed attitude towards spelling. Undoubtedly, due to internet communication's speed and casual nature, errors often go overlooked. Nevertheless, these behaviours sometimes extend to more formal settings, leading to the proliferation of spelling errors in professional and academic writing. Spelling errors in written communication have been a longstanding challenge (Gupta &

Sharma, 2016), with various methods and tools developed over time. Historically, errors were detected through manual proofreading (Yausaz, 2012), which was time-consuming and prone to human oversight. Automated spell checkers were introduced in the mid-20th century, but they were initially limited by static dictionaries and often failed to detect contextual errors. In addition, contextual spell checkers were developed to address these limitations, using algorithms to analyse the context of words in a sentence, reducing the likelihood of false positives. Rule-based systems use predefined rules and patterns to detect spelling errors (Abate et al., 2021), but their comprehensiveness limits them. Statistical language models use data-driven approaches to detect spelling errors, relying on large text corpora to determine the likelihood of word sequences. Machine learning (ML) and natural language processing (NLP) are the most recent and promising developments in spelling error detection as identified by (Saeedi et al., 2022).

NLP is an interdisciplinary field that explores the interaction between machines and people in the context of language communication, drawing from the domains of artificial intelligence and linguistics. Natural language processing is essential for various language-related tasks, such as identifying spelling errors, as it enables computers to understand, interpret, and generate human language. Incorporating NLP into spell checkers is a noteworthy advancement. These systems employ artificial intelligence to make sophisticated judgments about the appropriateness of a phrase in a particular situation rather than relying on a fixed word list or predetermined criteria. This approach improves the accuracy of identifying spelling errors and raises the quality of all written communication. NLP algorithms offer advanced and contextually aware solutions, handling complex cases like homophones and contextual errors by understanding the surrounding context and semantic meaning. These algorithms continuously improve as they are exposed to more data, making them highly effective in detecting errors and providing meaningful correction suggestions.

➤ *Statement of the Problem*

Accurate spelling has extensive consequences for written communication's clarity, credibility, and general effectiveness in several contexts, including academic and professional writing and personal communications. The frequency of spelling errors has increased significantly recently, mainly owing to the widespread use of digital media and society's dependence on written communication. Misspelled words can diminish the reliability of written correspondence and result in misinterpretations. Hence, there is an urgent requirement for an innovative and advanced solution to improve the identification of spelling errors. This research seeks to address these challenges and develop a system that significantly enhances the quality of written communication by improving the accuracy and relevance of spelling error detection and correction suggestions.

➤ *Aim and Objectives of the Study*

The aim of this work is to develop an automated word spelling error detection system. The objectives include to:

- Design a model for word spelling error detection using spello model.
- Preprocess data collected from Kaggle repository using NLTK library.
- Implement the model using Python programming language
- Evaluate the model using techniques such as WER, BLUE score, and Levenshtein distance

II. METHODOLOGY OF THE STUDY

The methodology for the proposed system is Object-oriented analysis and design methodology (OOADM) which focuses on structuring software components as objects that include both data and functions, offering a methodical approach to constructing software systems. This method aims to provide software modularity, reusability, and maintainability by using notions from object-oriented programming (OOP) paradigms. Objects, representing actual entities or abstract concepts in the software system, play a key role in OOADM.

A. *Analysis of the Existing System*

(M.and Abd, 2018) introduced a model that utilizes the phoneme to identify and rectify spelling problems in English, a widely spoken worldwide language. The existing approach utilizes lexical resources and Grapheme to phoneme mapping to choose the most important suggestions from a list of correction recommendations in order to provide correction suggestions. The determination is made using a Microsoft word dictionary. They employed English standard datasets including misspelt words to assess the proposed model. The existing approach primarily emphasizes the rectification and substitution of wrong words, as well as the automated detection of those words. The results exhibited a higher level of performance compared to Google, achieving an accuracy rate higher than that of applied by Google. Furthermore, the results were on par with those obtained from Microsoft Word.

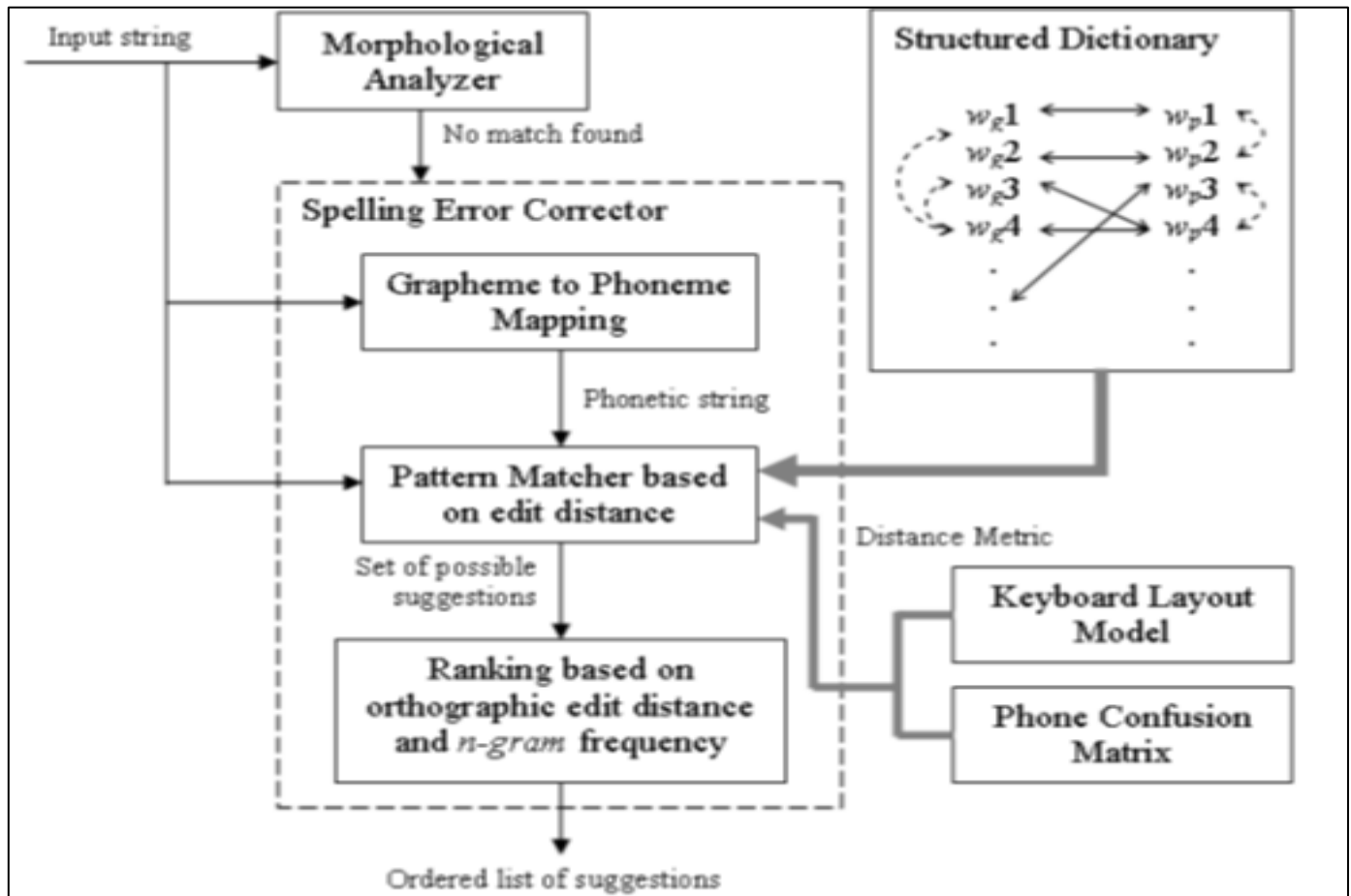


Fig 1: Architecture of the Existing System
(Source: M. and Abd, 2018)

B. Disadvantages of the Existing System

➤ *The Disadvantages of the Existing System Include the Following:*

- The phoneme of the models often struggle with words absent from their training data, potentially failing to identify errors involving novel or rare words.
- There is a limited understanding of context, as the phoneme model treat words in fixed-size sequences, potentially overlooking errors requiring a broader contextual grasp.
- The context size in N-gram models is often too constrained, limiting their effectiveness.
- Microsoft dictionary as the library which does not include certain names, technical terms, abbreviations, or specialized capitalization.
- The existing system for spell checking presently makes use of N-gram model and the Microsoft language dictionary, however there exist drawbacks that could

hinder wide adoption of the system which includes; out of vocabulary words which lead to inaccuracies most times when dealing with new words not used during training, also the issue of homophones and contextual ambiguities persist.

C. Analysis of the Proposed System

The proposed system will make use of the Modified Symmetric Delete spelling model (Modi-Symspell) instead of just the Phoneme Model. The Edit Distance was done using the Symspell Model algorithm then feed into the Phoneme model. Inside the check for word spelling errors. One of the significant features of this proposed model that makes it stand out is its ability to utilize multiple dictionaries. This dictionary makes it possible to train a model in a specific domain. The model is a pre-trained model that can effectively detect words spelt incorrectly. The proposed model will make use of dataset gotten from the Kaggle data repository, which will be used to train the model for spelling errors. The model will be evaluated to determine its level of efficiency.

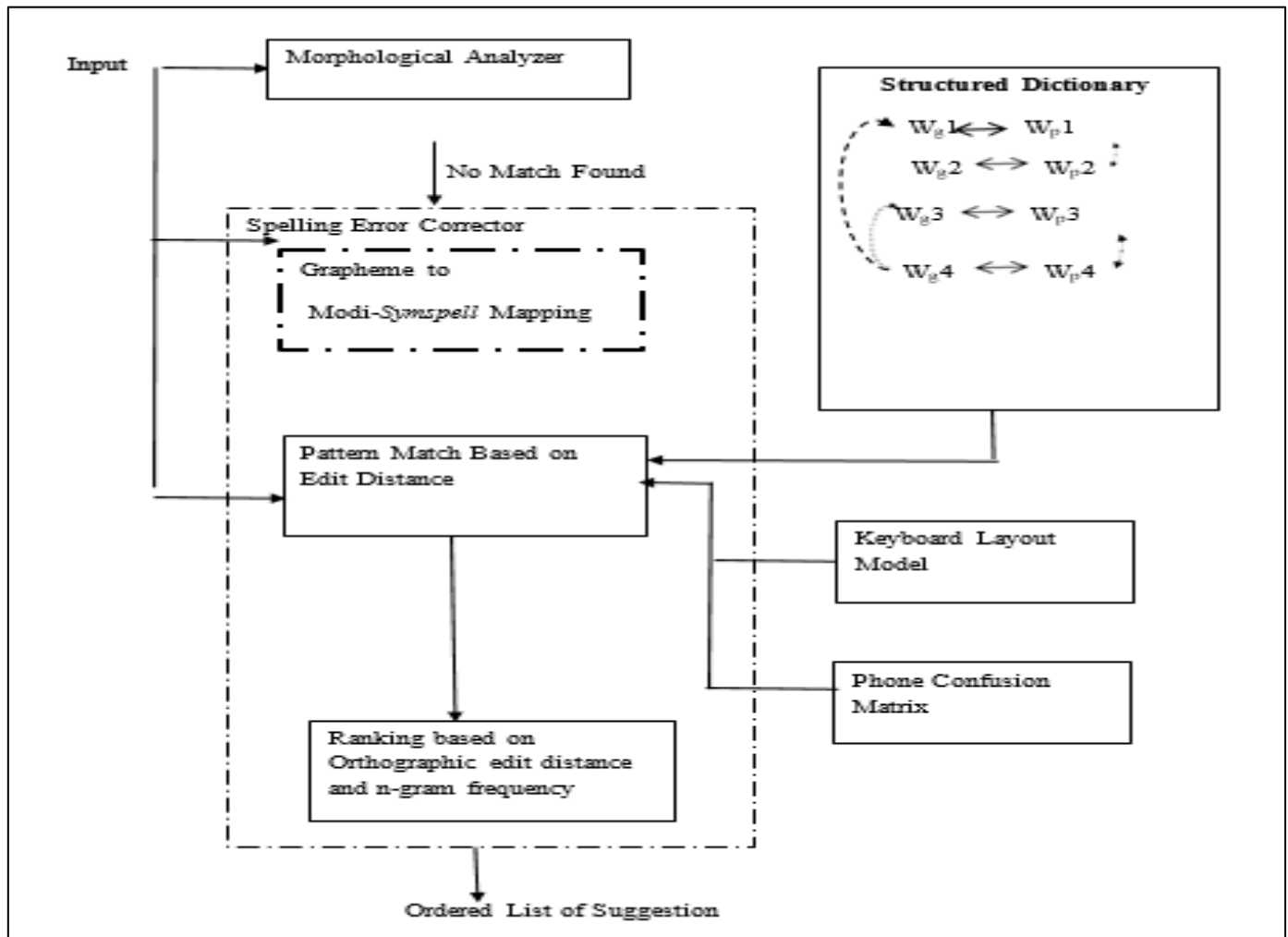


Fig 2: Architecture of the Proposed System

D. Analytical Presentation of Components of the Proposed System

➤ Input Text (Dataset)

The dataset was obtained from “Kaggle.com” which is a data repository. The dataset is in a text document format, with file size of 6.3 megabyte of text data. The dataset comprises a compilation of text samples that include instances of misspelt words alongside their respective corrected versions. The primary purpose of the spell check error dataset is to facilitate the training of our spell check model. The dataset comprises samples, each consisting of a text fragment including one or more words with spelling errors and their respective corrected versions. The dataset encompasses a diverse array of text genres, such as articles, essays, emails, and social media posts.

➤ Text Pre-Processing

Data (Text) processing is the process of preparing the dataset for the proposed model which requires different processing. The text data has to be processed to ensure that we have a very good model. The proposed model is only as good as the data that is used to build it. The data was imported using the pandas library.

Whitespaces and Special Characters Removal: Texts in the column are screened for special characters. By eliminating non-alphabetic symbols, the focus shifts to the core textual content, aiding machine learning models in identifying relevant patterns. Removing whitespaces ensures a standardized representation of text, which enhances accuracy by emphasizing the intrinsic linguistic elements without the distraction of leading and trailing spaces. This was achieved using the regular expression library.

Tokenization: For spell checking systems, the most common form of segmentation is word tokenization, which separates the text into individual words and lays a solid foundation for the identification of spelling errors. The units produced by tokenization can range from individual characters or words to phrases or even complete sentences. By transforming words into tokens, we can analyse the context in which they exist, which improves the spell checking process and makes it more efficient.

➤ Modi-Sym Spell Model

The model is a pretrained model that comprises of two models, the symspell model and the phoneme model, along with a context model that finds the best candidate from the list of suggestions suggested by Phoneme and Symspell model for a misspelled word.

- **Phoneme Model** uses Soundex algorithm in the background and suggests correct spellings using phonetic concepts to identify similar-sounding words. Soundex is a phonetic algorithm for indexing names by sound, as pronounced in English. The goal is for homophones to be encoded to the same representation so that they can be matched despite minor differences in spelling. The algorithm encodes mainly consonants.
- **SymSpell Model** uses the concept of edit distance in order to suggest correct spellings in other words, it is an algorithm used to find all strings within a maximum edit distance from a huge list of strings in very short time. It also handles QWERTY based errors which unintentionally occur while typing from the keyboard. SymSpell gets its speed from symmetric delete spelling correction algorithm, which reduces the complexity of edit candidate generation and dictionary lookup for a given damerau-levenshtein distance.
- **Context Model** helps to determine the most probable word from the list of the suggested word for misspelled words from the SymSpell and Phoneme model. A lightweight n-gram probabilistic model has been trained which will find the next best word in a sentence. We use this model to calculate the overall probability of a sentence being formed for each combination of the misspelled word suggestions. Moreover, it is important to note that the models configurations can be tuned to set the maximum and minimum length of word that should be spellchecked using `sp.config.min_length_for_spellcorrection`, also the max edit distance settings can be finetuned to accommodate for each char level for symspell and phoneme using `sp.config.symspell_allowed_distance_map`.

➤ Text without Error

The texts without error are the words that have been suggested as being the correctly spelt word.

➤ Evaluation

This section is used for the assessment of the developed model for spell checking errors, which is used for improving writing across different fields. The aim is to offer valuable information that may guide users in their efforts to create more sophisticated and dependable solutions for assuring flawless and error free spellings in written communication within the constantly growing digital ecosystem. The following criteria are used for evaluating the accuracy or effectiveness of spelling checkers and they include WER, Blue Score, and Levenshtein distance.

WER: Word Error Rate is a metric that quantifies the difference between the predicted text and the reference text in terms of the number of word-level errors. These errors include substitutions, insertions, and deletions, collectively reflecting the divergence between the corrected text and the intended, error-free version. A lower WER indicates a closer match between the predicted and reference texts, suggesting higher accuracy in spelling error correction. The formula for calculating WER is as follows:

$$WER = \frac{S+D+I}{N}$$

Where

S = the number of substitutions (mismatched words),
 D = the number of deletions (missing words in the predicted text),
 I = the number of insertions (extra words in the predicted text),
 N = the total number of words in the reference text.

BLUE Score: BLUE (Bilingual Evaluation Understudy) score is a metric commonly used in natural language processing tasks, such as machine translation and text summarization, to evaluate the similarity between a system-generated text and one or more reference texts. While BLUE is not traditionally employed for spelling error detection and correction, its principles can be adapted to assess the effectiveness of spelling error correction systems. The BLUE score considers precision and recall in evaluating system outputs. In the context of spelling errors, precision could be interpreted as the correct identification and correction of misspelled words, while recall would represent the system's ability to correct all errors.

Levenshtein Distance: It is also known as edit distance, is a metric commonly used to quantify the dissimilarity between two strings by measuring the minimum number of single-character edits (insertions, deletions, or substitutions) required to transform one string into the other. Levenshtein Distance can be used to measure the dissimilarity between the corrected text produced by a spelling correction system and the reference (error-free) text.

III. RESULTS

Comparing the text snippet "After the plea bargain jurisdiction is made" with the corrected version "After the plea bargain jurisdiction is made" provides valuable insights into legal terminology, the importance of linguistic accuracy in legal discussions, and the consequences of language mistakes in legal settings. The revised version highlights a minor yet crucial lexical mistake in the initial text, where "jurisdiction" was misspelled. Errors like these highlight the critical nature of linguistic precision in legal documents, as it can greatly impact the understanding and implementation of legal concepts. In legal procedures, "jurisdiction" is crucial since it defines the authority of a court or legal body to handle cases within a specific territory or subject area. Correcting "jurisdiction" to "jurisdiction" fixes a typographical mistake and maintains the semantic accuracy of the text, guaranteeing clarity and coherence in legal communication. Correcting errors helps to show how linguistic aids like spell-checkers and automatic correction algorithms improve the precision and dependability of legal papers.

The evaluation criteria being BLEU Score, Edit Distance, and Word Error Rate provide quantitative insights into the level of difference between the original and corrected texts. The BLEU Score is a metric used in natural language processing to quantify the similarity between two text sequences, indicating their semantic alignment numerically. The BLEU Score of 0.4889 in this case indicates a moderate level of semantic similarity between the original and corrected texts, showing that the correction process

effectively maintained the intended meaning. The Edit Distance metric measures the number of procedures needed to change one text into another, showing a small difference between the original and corrected texts. An Edit Distance of 1 means only one corrective operation is required. The small Edit Distance highlights the specific error location and the effectiveness of the correction procedure in fixing the lexical oddity.

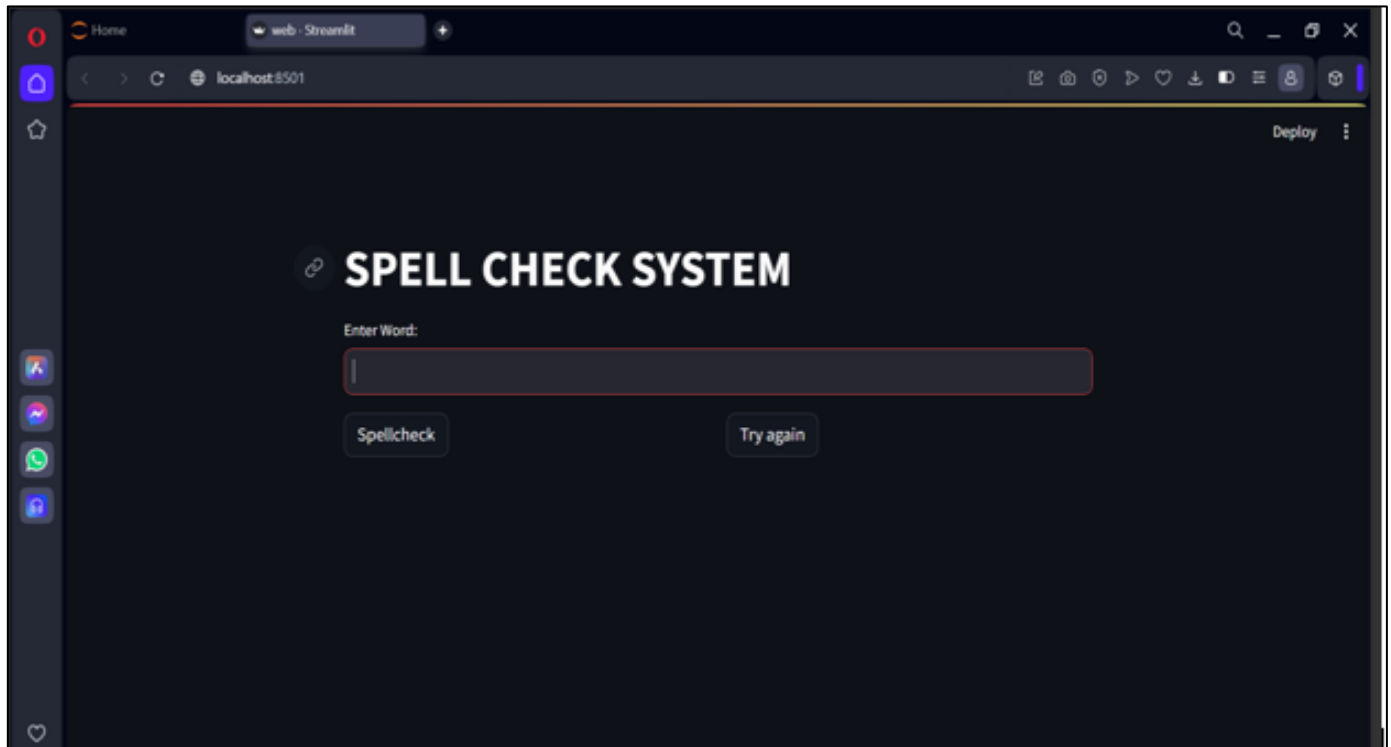


Fig 3: Output Screen

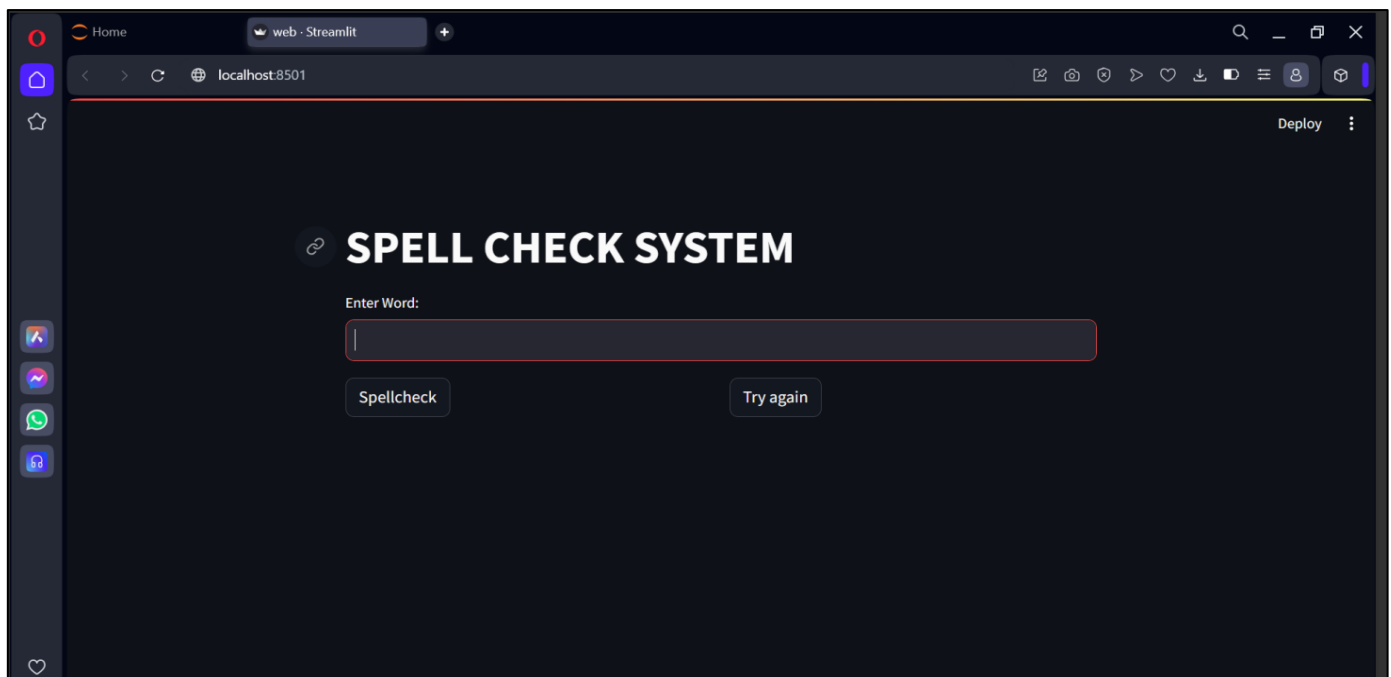


Fig 4: Homepage of Application

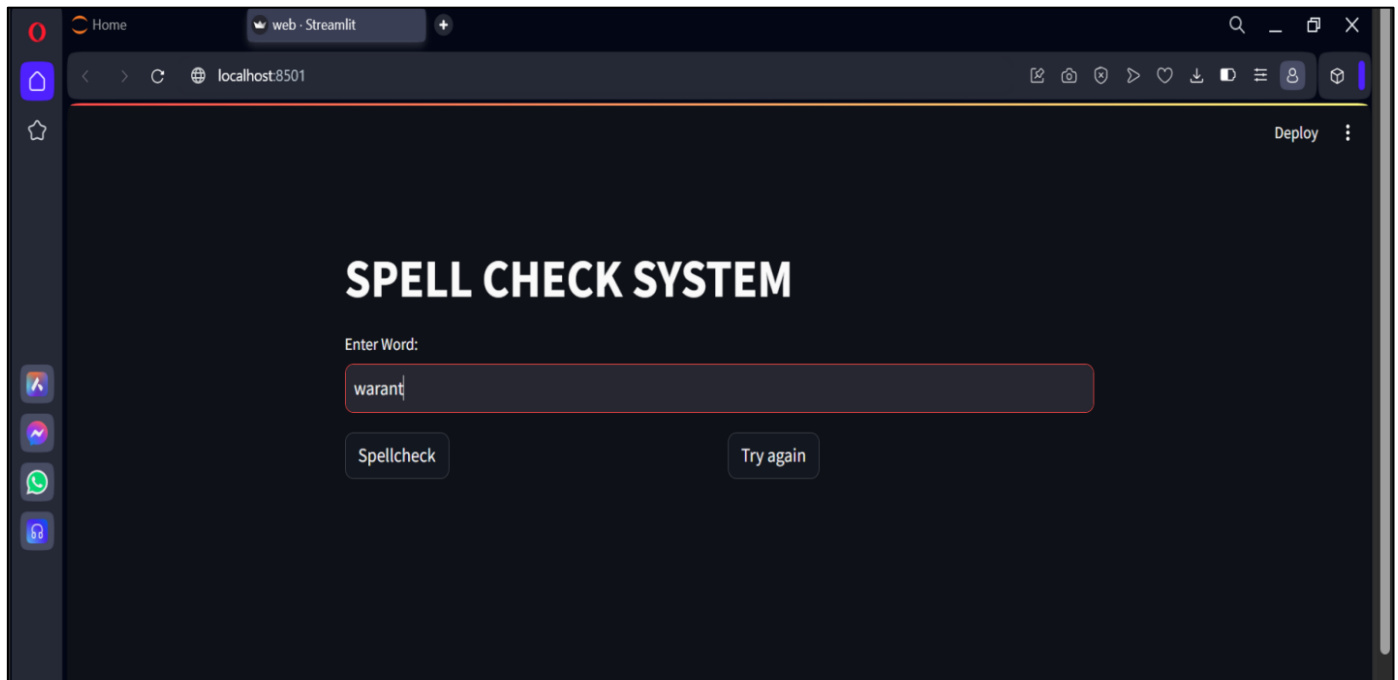


Fig 5: Predicted Text Output

The Word Error Rate is obtained by determining the proportion of incorrect words compared to the total number of words in the text, providing a thorough evaluation of the correction process accuracy. The repaired text has a Word Error Rate of 0.1429, indicating a high level of accuracy with only a small portion of words needing correction. The examination of the given text excerpt and its correction highlights the crucial importance of linguistic precision in legal communication, stressing the value of precise terminology and the consequences of linguistic mistakes in legal documents. Automated correction processes and quantitative evaluation measures can maintain the integrity and coherence of legal writings, aiding in clear communication and understanding in the legal field.

➤ Comparison of the Results with the Existing Systems

In the comparison of the results with the existing systems system there are certain Benchmark created, with the first one on error of the word and the sequence accuracy. It evaluates how the method Stability Test act on error free sequences. We only evaluated the word and sequence accuracy for this test set. The set consists of 250 sentences with 6370 words in total. Test set. Due to the limitations of some tools the set was limited to ≈ 650 sentences. Setting the random seed to 42 it is guaranteed to always generate the same subset, containing the exact same errors, when using the tool. Evaluation Metrics used in the evaluation have a wide range of information which will be described in more detail below and how those values are determined.

- **Word Accuracy.** For the tokens T of a sequence s_i the word accuracy is given by $W A : \forall t_i \in T : c(t_i) / |T|$, with $c(i)$ in range of 0 and 1 where $t_i; P P$ os are the prediction indices of token t_i and $t_i; P P$ os j the access to the text of the prediction token at position j , and respectively for GP os.

- **Sequence Accuracy.** The sequence accuracy for all sequences S is given by $SA: \forall s_i \in S: s(s_i) / |S|$, with $s(s_i)$ in range of 0 and 1.
- **Error Category Detection.** Although most of the determined tools do not have an evaluation of the underlying error category the measurements for the detection can be determined indirectly by the correction itself. Prediction is measured as the proportion of adequately detected errors of each error category by a tool.
- **Recall** as the proportion of the errors of each error category marked in the ground truth of all articles that were detected during the correction.
- **F-Score** is defined as the harmonic mean of precision and recall. With the majority of tools not supporting error categories we measure the detection through their error correction. If a token was changed with respect to the ground truth it cannot be NONE anymore. If it was adequately corrected we assume that the error category was correctly detected. Otherwise the proposed corrected token is analysed to investigate which category is the most probable, e.g. investigate whether the character of the first letter was changed indicating capitalisation errors.
- **Error Category Correction.** We further evaluate independently whether a given token was adequately corrected.
- **Precision** is thereby measured as the proportion of corrected errors for an error category by a tool in all articles.
- **Recall** is measured as the proportion of the errors of each error category marked in the ground truth of all articles that were adequately corrected. F-Score is defined as the harmonic mean of precision and recall. Unaligned prediction tokens after linking will be counted as REAL_WORD errors (W) at this point.

- **Suggestion Adequacy (SA).** SA measures the adequacy of the proposed suggestions, if provided by the tool. In accordance to (Starlander and Popescu-Belis, 2002) we are using the following metric SA for measuring the adequacy, with “t” being the current token and “s” a list of the predicted token and further suggestions, if provided: SA (t, s) in range of 0 and 1.

The overall suggestion adequacy SA is determined by normalising the sum of all scores by the number of tokens. With T being the set of predicted tokens defined as (predict, sugg = [], err = []), where predict is the predicted tokens, sugg a list of alternative predictions, and err the proposed error category. The results for the Error (E) the word (W) and

sequence accuracy (S) can be seen plotted in figure against other models. The Modified Spello (MSpello) developed in this work was compared with Aspell (a GNU Spell checker tool), Bing Spell (Microsoft Cognitive toolkit for word break) and Google spell checker (Spelling correction for Google services). Results from the developed work in Recall Accuracy and Error Free Sequences were examined and plotted for the comparison. In the plot in figure 6 it is clear that the Modified Spello (MSpello) developed in this work ranked favourably and slightly outperforming the existing models. Words recall **not found** show the lowest in the MSpello compared to the other models while the Word Recall **Found** is the highest comparing favourably with Google and Bing Spell.

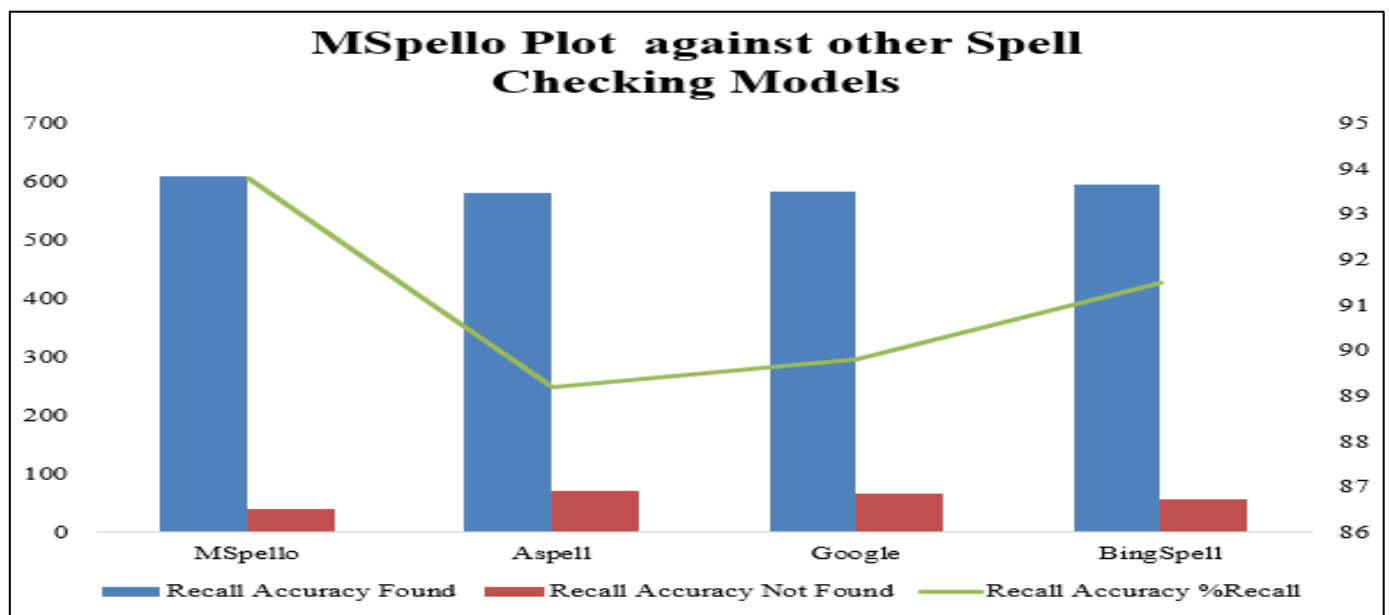


Fig 6: Plot of Recall Accuracy for MSpello against Existing Models

In the plot in figure 7 a plot of Recall Accuracy percentage for MSpello against Existing Models. The plot show that the Modified Spello (MSpello) developed in this work ranked higher in percentage compared to the existing

models. Words recall show 93.8% which is the highest in the MSpello compared to the other models which showed lower percentages.

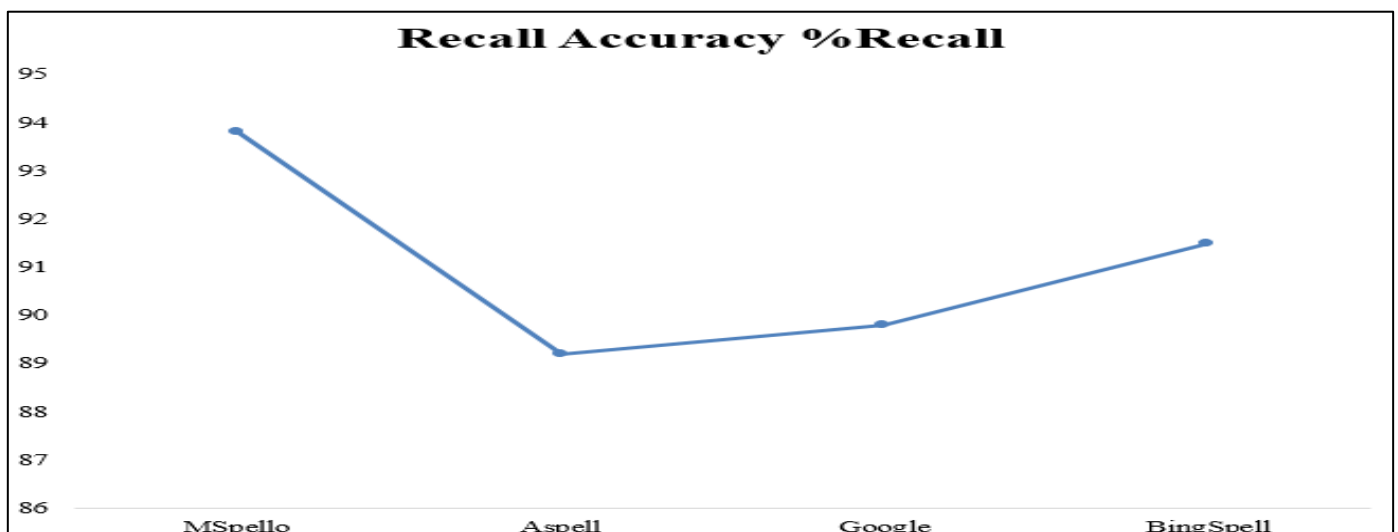


Fig 7: Plot of Recall Accuracy Percentage for MSpello against Existing Models

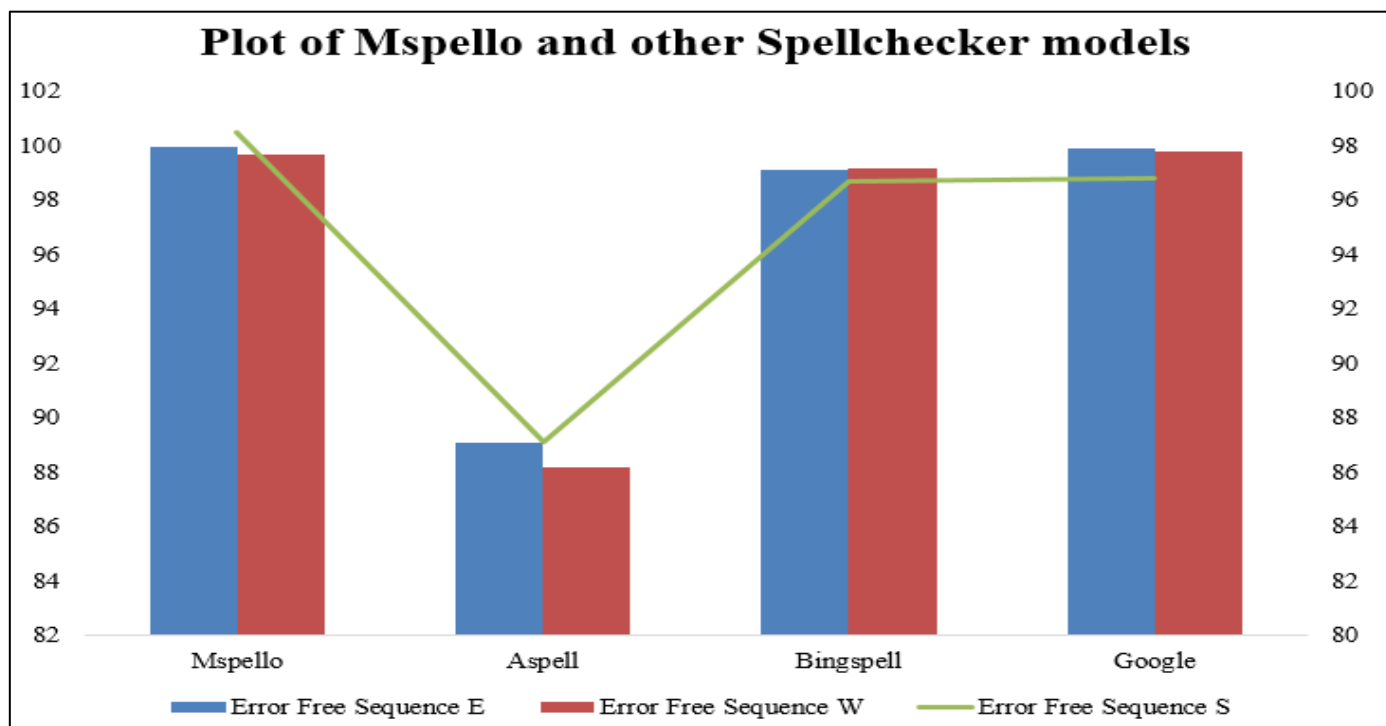


Fig 8: Plot of Error Free Sequence for MSpello against Existing Models

In the plot in figure 8 a plot of error free sequence for the MSpello against existing the other. Models were presented. The plot show that the Modified Spello (MSpello) developed in this work ranked higher in (E), W, and S metrics used in the comparison. The Mspello show close performance to Google but slightly outperform in Error free sequence E.

IV. SUMMARY, CONCLUSION AND RECOMMENDATIONS

A. Summary

The Modi-sympell model, a combination of natural language processing and machine learning algorithms, has been used to correct spelling errors in legal terminology. It uses a collection of legal texts and specialized dictionaries to understand the complexities of legal language, ensuring accurate error correction. The model also enhances semantic consistency in legal documents by systematically analyzing spelling problems and their remedies. It provides a scalable and efficient solution for spell-checking and error correction in various legal contexts, handling large amounts of text data and optimizing the proofreading and validation process. The model's performance is assessed using quantitative indicators like BLEU Score, Edit Distance, and Word Error Rate. Further research and improvement could further enhance its effectiveness in dealing with complex linguistic issues in the legal field.

B. Conclusion

The Modi-sympell model, a tool that corrects spelling problems in legal words, is a significant advancement in legal documentation and linguistic accuracy. It uses natural language processing techniques and machine learning algorithms to enhance the precision, dependability, and interpretive lucidity of legal writings. The model reduces

risks associated with unusual words and unclear meanings in legal documents, improving semantic consistency, communicative efficacy, and interpretative coherence. Evaluation metrics like BLEU Score, Edit Distance, and Word Error Rate provide quantitative insights into the model's performance and effectiveness. The model's adaptability and customization features allow legal practitioners to personalize the correction process according to their language preferences and domain-specific terms. Further research is needed to improve the model's performance, robustness, and scalability. The Modi-sympell model represents a significant change in legal linguistics, promoting accuracy, integrity, and accessibility in legal communication.

C. Recommendations

According to this study, the implemented model is a promising technique to correct spelling problems in legal terminology, has been praised for its effectiveness and usability. Its recommendations include continuous research and development, improving its ability to use contextual cues and domain-specific information for error correction, emphasizing user-centric customization, implementing collaborative feedback systems, offering extensive training and support, prioritizing ethical and regulatory considerations, and fostering interdisciplinary collaboration. These suggestions aim to optimize the Spello model's performance, resilience, and scalability, and make it a fundamental technology in promoting accuracy, integrity, and accessibility in the legal profession and society. The model can become a fundamental technology in promoting accuracy, integrity, and accessibility in the legal profession and society through combined efforts and collaboration.

REFERENCES

- [1]. Abate, J., Khedkar, V., & Tidke, S. K. (2021). Integrated Model to Develop Grammar Checker for Afaan Oromo using Morphological Analysis: A Rule-based Approach. *International Journal of Advanced Computer Science and Applications*, 12(5). <https://doi.org/10.14569/IJACSA.2021.0120526>
- [2]. Aceto, G., Persico, V., & Pescapé, A. (2019). A Survey on Information and Communication Technologies for Industry 4.0: State-of-the-Art, Taxonomies, Perspectives, and Challenges. In *IEEE Communications Surveys and Tutorials* 21(4). <https://doi.org/10.1109/COMST.2019.2938259>
- [3]. Gupta, S., & Sharma, S. (2016). A spelling mistake correction (SMC) model for resolving real-word error. *Advances in Intelligent Systems and Computing*, 410. https://doi.org/10.1007/978-81-322-2734-2_43
- [4]. Juma'a, T. R. (2020). Flouting Grice's cooperative principle in a political interview (Trump's interview with time on 2020). *PalArch's Journal of Archaeology of Egypt/Egyptology*, 17(5).
- [5]. M., S., & Abd, A. (2018). Automatic spelling Correction based on n- Gram model. *International Journal of Computer Application*, 182(11).
- [6]. Mainsah, B. O., Collins, L. M., Colwell, K. A., Sellers, E. W., Ryan, D. B., Caves, K., & Throckmorton, C. S. (2015). Increasing BCI communication rates with dynamic stopping towards more practical use: An ALS study. *Journal of Neural Engineering*, 12(1). <https://doi.org/10.1088/1741-2560/12/1/016013>
- [7]. Saeedi, S., Fong, A. C. M., Gupta, A., & Carr, S. (2022). Reusable Toolkit for Natural Language Processing in an Ambient Intelligence Environment. *Proceedings of the 2022 IEEE Symposium Series on Computational Intelligence, SSCI 2022*. <https://doi.org/10.1109/SSCI51031.2022.10022225>
- [8]. Starlander, M., & Popescu-Belis, A. (2002). Corpus-based evaluation of a French spelling and grammar checker. In *Proceedings of the Third International Conference on Language Resources and Evaluation (LREC'02)*, Las Palmas, Canary Islands, Spain. European Language Resources Association (ELRA).
- [9]. Syamsuddin, A., Sukmawati, Mustafa, S., Rosidah, & Ma'rufi. (2021). Analysing the skill of writing a scientific article as a written communication skill of prospective elementary school teacher on learning mathematics. *Journal of Educational and Social Research*, 11(5). <https://doi.org/10.36941/jesr-2021-0108>
- [10]. Verbaarschot, C., Tump, D., Lutu, A., Borhanazad, M., Thielen, J., van den Broek, P., Farquhar, J., Weikamp, J., Raaphorst, J., Groothuis, J. T., & Desain, P. (2021). A visual brain-computer interface as communication aid for patients with amyotrophic lateral sclerosis. *Clinical Neurophysiology*, 132(10). <https://doi.org/10.1016/j.clinph.2021.07.012>
- [11]. Voloshchuk, I., & Derkach, A. (2020). Conceptualization of knowledge in aerospace industry in the mode «image-sense» on the basis of film scripts. *Visnik Mariupol's'kogo deržavnogo universitetu. Seriâ: Filologîâ*, 13(23), 150–155. <https://doi.org/10.34079/2226-3055-2020-13-23-150-155>
- [12]. Xiuqing, C. (2023). The Role of Cognitive Metaphor in Chinese and Russian Languages. *Litera*, 8. <https://doi.org/10.25136/2409-8698.2023.8.38456>
- [13]. Yausaz, F. (2012). Fluidez en el trazado manual y composición escrita. Estudio exploratorio con niños Argentinos al finalizar tercer grado. *Interdisciplinaria*, 29(2). <https://doi.org/10.16888/interd.2012.29.2.5>