

Analysis and Development of KEBI 1.0 Checker Framework as an Application of Indonesian Spelling Error Detection

Tresna Maulana Fahrudin^{1,*}, Ilmatius Sa'diyah², Latipah³, Ibnu Zahy' Atha Illah⁴, Cagiva Chaedar Bey Lirna⁵, Burhan Syarif Acarya⁶

Citation: Fahrudin, T. M.; Sa'diyah, I.; Latipah; Atha Illah, I. Z.; Bey Lirna, C. C.; Acarya, B. S., Analysis and Development of KEBI 1.0 Checker Framework as an Application of Indonesian Spelling Error Detection. 2021, Volume 1, Number 2, Page 12-26.

Received: November 1, 2021
Accepted: November 16, 2021
Published: November 25, 2021

^{1,2,4,5,6} Department of Data Science, Faculty of Computer Science; {¹tresna.maulana.ds,²ilmatius.sisfo}@upnjatim.ac.id, {⁴20083010016,⁵20083010020,⁶20083010004}@student.upnjatim.ac.id

³ Department of Informatics Engineering, Faculty of Computer Science; {³latifah.rifani}@narotama.ac.id

* Correspondence: tresna.maulana.ds@upnjatim.ac.id; Tel.: (+6289678526342)

Abstract: At educational institutions, especially at University, writing scientific papers is a skill that must be possessed by academics such as educators and students. However, writing scientific papers is not easy, there are many provisions and rules that need to be fulfilled. Several studies show that there are still many academics who make mistakes in writing their scientific papers. Some of the mistakes made include punctuation errors, typographic writing errors and the use of non-standard words in Indonesian. Researchers in Indonesia have developed various spelling error detection applications in Indonesian-language scientific papers. This study tries to analyze the development of an application framework for detecting Indonesian spelling errors from various assessment indicators. This study tries to compare the application framework for detecting spelling errors between other studies with proposed application that named KEBI 1.0 Checker. KEBI 1.0 Checker as a spelling error detection application has 3 main features, namely detecting errors in the use of punctuation marks, writing typography, and using non-standard words in accordance with the standards of the Big Indonesian Dictionary and the General Guidelines for Indonesian Spelling. In addition, this study tries to objectively examine the complexity of the features, advantages and disadvantages, methods and the level of accuracy of each application. The results of the analysis show that KEBI 1.0 Checker has the completeness of features, fast computation time, easy application access, and an attractive user interface. However, it is still necessary to improve the precision in correcting spelling errors in typographic words.

Keywords: KEBI 1.0 Checker, Framework, Detection, Spelling Error, Indonesian Language

1. Introduction

Writing scientific works is one of the main activities in educational institutions, especially at the university level. Writing scientific works is one way to express ideas and opinions to be conveyed to others. According to the Great Indonesian Dictionary (KBBI), the meaning of scientific work was written work made with scientific principles, based on data and facts (observations, experiments, and literature studies). However, not all scientific works produced are of good quality. Writing spelling errors and the use of sentences are often found in a scientific works [1]. Some of the factors that cause writing errors are typographical or word writing errors [2] such as the word *rumah* is written *rumha*, errors in using standard words [3] [4] such as the word *atlet* which should be written

by *atlit*, and errors in using punctuation marks such as the use of a comma or a period [5]. These errors can change the meaning of the writing and reader's understanding.

Based on the fourth edition of the General Guidelines for Indonesian Spelling published by the Language Development and Development Agency of the Ministry of Education and Culture (PUEBI) in 2016, Indonesian spelling errors include the use of letters, writing words, and using punctuation marks. In addition, typographical errors are also included in the deviation of good and correct Indonesian writing according to PUEBI and KBBI.

The results of the study [6] showed that the lecturers still lacked good and correct Indonesian language skills. This has an impact on the scientific work produced because there are still many writing errors in the scientific work. For this reason, an application is needed so that academics can fulfill the rules of writing good and correct standard types of writing in their scientific works. What is meant by 'correct' is conformity with standard Indonesian language rules [7]. In addition, another study [8] found that the students of the Informatics Study Program at Cokroaminoto Palopo University did not yet have the capability to write scientific papers well. This was shown through the results of the student capability test where the percentage of students who score less than 75 is 70% of the set standard of 85%.

Typographical mistakes are mistakes that writers often make. Typing too fast or out of focus can lead to typographical errors. Sometimes the ease of the "auto correct" facility in word processing applications can also cause typography to occur. Typographical errors are errors that occur when the writer is typing text and these errors can change the meaning of a word and even the meaning of a sentence, for example, a finger presses two adjacent keyboard keys simultaneously [9]. Typo will sometimes make the words ambiguous, incorrect or random which will affect a person's understanding [10].

Another mistake that is often made in writing scientific papers is punctuation errors. Punctuation marks are signs used in the spelling system (such as periods, commas, colons, and so on) that serve to help readers understand the meaning of writing correctly. Each punctuation mark has a different function, for example a comma serves as a pause. Punctuation is a code that is often needed for meaning and emphasis in writing [11]. In the results of research [12] which conducted testing on 20 student scientific works, it can be concluded that students almost do not understand in using punctuation marks with a percentage of 70% of which are errors in using capital letters. In addition, based on interview data, several factors that cause punctuation errors are the lack of balance in knowledge, curriculum, concentration and carelessness, and practice.

Therefore, this study tries to compare the framework or application framework for detecting spelling errors between other studies that have been proposed previously with KEBI 1.0 Checker. In addition, this study tries to objectively examine the complexity of the features, advantages and disadvantages, methods and levels of accuracy of each application.

2. Related Works

The increasing number of spelling errors encountered and obstacles in the text editing process in Indonesian scientific works, researchers in the field of computer science, especially in the field of computational linguistics, initiated of the development of software applications to facilitate the detection of spelling errors for authors. Arina et al [9] built the application using Dictionary Lookup, N-Gram method with Cosine Similarity, and Levensthein Distance to correct word errors in documents where words in the Great Indonesian Dictionary (KKBI) are used as references in the

comparison of incorrect words. Consequently, the methods used only pay attention to letters in the language of the words which are considered erroneous.

Another similar study was conducted by Ricky et al[13] who used the Peter Norvig and N-gram methodologies to develop the spelling correction application. The applications has the capability to correct wrong words and convert them to standard words. The scenario testing was carried out on 55 documents and the accuracy of the spelling checker reached 69.09%. Furthermore, Peter Norvig method cannot correct spelling errors in some sentences that have two-letter errors in a word or errors because they are not in the database.

In contrast to the research above, Mazidhatul et al. [14] developed application to detect errors in the use of punctuation and capital letters using Boyer Moore's algorithm by matching each character in the pattern with the character in the appropriate text. The system was tested by comparing the results of a manual reviewer's detection with the results of the detection of the system. The study showed that the average precision and recall for the system was 0.969 and the average accuracy for the system was 91.7%.

Among the three spelling error detection applications that have been developed by other researchers, this study proposes a framework that has been developed for spelling error detection applications known as KEBI 1.0 Checker. This application has been equipped with three features which have the capability to detect and correct spelling errors such punctuation marks, non-standard words and typographical errors. The features provided by KEBI 1.0 Checker according to the performance principle of the spelling checker are used to detect errors and provide recommendations of words that are close to the intended word. The goal is to minimize spelling errors in document writing.

Many methods can be used to perform spelling corrections, including the spelling checker Peter Norvig, Levenshtein Distance, N-Gram, and the spellchecker BK-Trees. Therefore, this study attempts to compare the framework of spelling error detection applications among other studies that have been proposed previously with KEBI 1.0 Checker. In addition, the study tries to objectively examine the complexity of features, advantages and disadvantages, methods as well as the level of accuracy of each application.

3. Experiment and Analysis

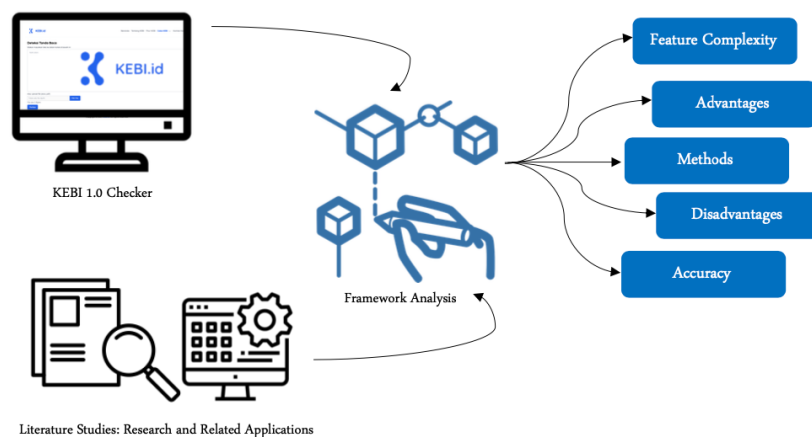


Figure 1. The Research Methodology of Analysis and Development in KEBI 1.0 Checker Framework

Figure 1 shows the flow of the research methodology from the initial stage, namely the study of literature to the final stage of the analysis of each framework, which is explained as follows.

- a. Study of literature. A cited reference was carried out by searching for journals and articles from several previous researchers related to the Indonesian spelling error detection system. In the study, three scientific articles were obtained from various sources which were used as comparisons between other research frameworks and KEBI 1.0 Checker.
- b. Framework Analysis. An analysis of each framework is carried out based on the complexity of features, advantages and disadvantages, methods, and level of accuracy.
 - Complexity of features: analysis based on the level of completeness of the features offered by the Indonesian spelling checker application. As is known, the principles and rules for writing scientific papers in Indonesian have been regulated in the PUEBI and KBBI. These rules will certainly affect how a spelling error detection system or application will be developed and work.
 - Strengths dan weaknesses: analysis based on the advantages and disadvantages of each application in terms of ease of access, User Interface (UI) and User Experience (UX), and computing time.
 - Methods: analysis based on the alternative methods selected and the performance of each method.
 - Level of accuracy: analysis based on the system's accuracy in correcting Indonesian spelling errors.

The KEBI 1.0 application is a web-based application that is easy to access without installing it first. This application refers to the Indonesian PUEBI and KBBI spelling rules so that the output produced is the standard output of the standard language in Indonesian. This application has three functions: detecting punctuation-based errors, detecting standard Indonesian words-based, and detecting typo-based words in Indonesian. Users can type directly in the editor box or enter text in *.docx or *.pdf format. The system produces output in the form of a Microsoft Word archive with the addition of the results of the examination of punctuation generated automatically.

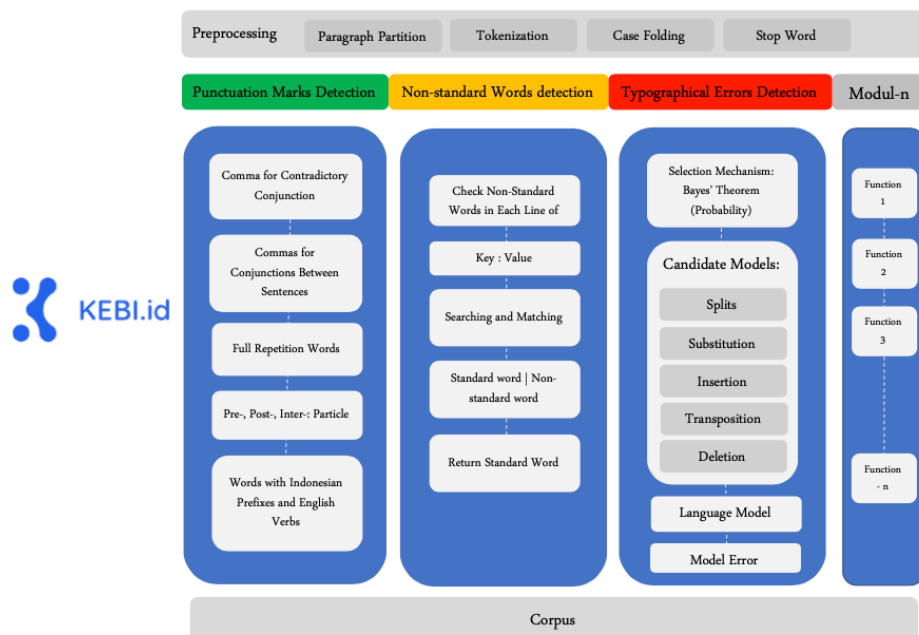


Figure 2. The Development of KEBI 1.0 Checker Framework

The KEBI 1.0 Checker framework is illustrated in Figure 2, consisting of preprocessing, detecting and correcting spelling errors, and a corpus. The error detection and correction process begins with entering scientific work documents into the application. Documents can be sent in *.docx or *.pdf format. After that, preprocessing is carried out, including partitioning paragraphs into several parts based on line moving markers ('\n'), tokenization of sentences into word tokens, case-folding to convert words to lowercase, and stop words from filtering out the standard and correct words. are in the corpus without the need for spelling correction [15]. The next process is to check and search for words based on good and correct Indonesian writing according to PUEBI and KBBI. KEBI 1.0 Checker offers three main features, namely:

- Punctuation detection and correction.** KEBI 1.0 Checker can correct text that contains comma punctuation errors related to contradictory conjunctions and connecting expressions between sentences. In addition, it can improve the complete repetition of words (dwilingga), pre-, post-, inter-, and words with elements of Indonesian prefixes and English verbs.
- Non-standard word detection and correction.** If the detected non-standard word is found in the word list in the PUEBI and KBBI-based corpus, the word will be replaced with a common word. This non-standard word detection uses the Dictionary Lookup method, and this method works by searching and matching based on key: value.
- Typo word detection and correction.** KEBI 1.0 Checker can also detect typos using the Peter Norvig Spelling Corrector method. As is known, writers will usually write words that are strung together into a sentence. The more words that are typed, the more chances of word errors that the author will type. However, this depends on the accuracy and sensitivity of the author during typing words into the document. KEBI 1.0 Checker uses the Peter Norvig Spelling Corrector method to correct typos. In general, this method works using splits, insertion, transposition, substitution, and deletion operations.

3.1. Punctuation Marks Detection using IF-THEN rules in KEBI 1.0 Checker

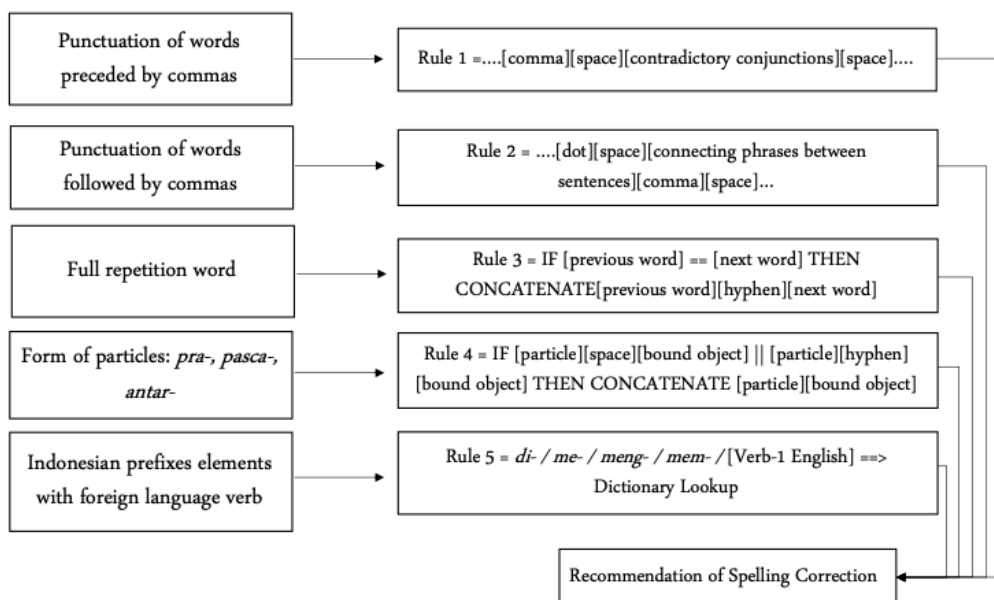


Figure 3. Punctuation-based Spelling Error Detection Process using IF-THEN rules

KEBI 1.0 Checker offers features to detect spelling errors of punctuation related to a) contradictory conjunctions, b) connecting expressions between sentences, c) complete repetition of words, d) bound forms that become the object of discussion in the form of particles, and e) the combination of Indonesian prefixes elements with foreign language verb aspects. For a more detailed explanation as follows.

a. Punctuation of words preceded by commas

It is related to contradictory conjunctions such as "..., even though" (... , *padahal*), "..., while" (... , *sedangkan*), "..., but" (... , *melainkan*, *tetapi*), and so on.

b. Punctuation of words followed by commas

Relating to connecting expressions between sentences such as "However,..." (*Akan tetapi,...*), "In this case,..." (*Dalam hal ini,...*), "And, related to that,..." (*Berkaitan dengan itu,...*), and so on.

c. Full repetition of words (*dwilingga*)

If we follow the rules of PUEBI, the complete word repetition should be affixed with a hyphen (-) in the middle. For example, children (*anak-anak*), you're welcome (*sama-sama*), and go for a walk (*jalan-jalan*).

d. Pre-, post-, and inter-particles

Related to the bound form, which is the object of discussion, namely the use of pre-, post-, and inter-. For example, "prehistory" (*prasejarah*), "postgraduate" (*pascasarjana*), and "interconnect" (*interkoneksi*).

e. The combination of Indonesian prefixes with foreign verbs

Related to standard and non-standard words that store data on Indonesian prefixes and English verbs, such as "is backed up" (*di-backup*), "backup" (*mem-backup*), "is accepted" (*di-accept*), and "accept" (*meng-accept*).

Figure 3 shows the process of KEBI 1.0 Checker detecting spelling punctuation errors using IF-THEN rules. PUEBI rules can be implemented into rules made in the form of syntax code so that the KEBI 1.0 Checker can handle spelling errors in documents.

3.2. Non-Standard Words Detection using Dictionary Lookup in KEBI 1.0 Checker

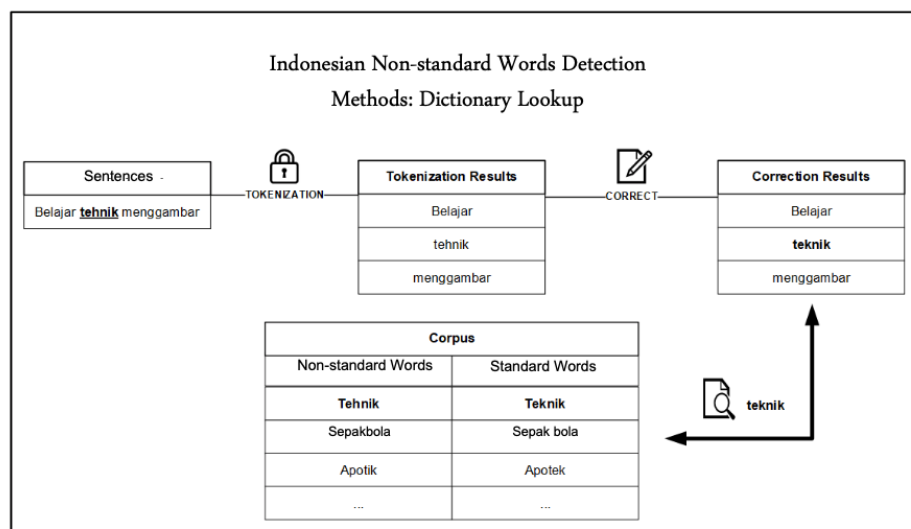


Figure 4. Indonesian Non-standard Words-based Spelling Error Detection Process using Dictionary Lookup

A word can be called a non-standard words what if the word used is not following the provisions of the Indonesian. A word that is not standard is not only caused by writing errors. However, it can also be caused by incorrect pronunciation and preparation of sentences that are not appropriate. This word is often found and used in everyday life.

Some examples of spelling errors in Indonesian word writing are not standardized as follows:

- Non-Standard word: *tehnik*, Standard word: *teknik*
- Non-Standard word: *sepakbola*, Standard word: *sepak bola*
- Non-Standard Word: *apotik*, Standard Word: *apotek*
- Non-Standard word: *praktek*, Standard word: *praktik*

Figure 4 shows the process of detecting spelling errors by KEBI 1.0 Checker on sentences containing various non-standard words. The sentence "*Belajar tehnik menggambar*" (learn drawing techniques) is entered into the application. Next step, the sentence is tokenized. Sentences containing the word "*tehnik*" are checked and matched with keywords in the corpus which contains a set of standard Indonesian words. After that, this keyword will enter the value that contains the standard word so that the standard word will replace the non-standard word in the previous input word. The non-standard word "*tehnik*" will be replaced with the standard word "*teknik*", it will produce output in the correct sentence "*Belajar teknik menggambar*". This process is called *key:value* that can be applied to the dictionary lookup method.

3.3. Typographical Errors Detection using Peter Norvig Spelling Corrector in KEBI 1.0 Checker

Spelling error detection method that can be applied to detect typographical error is Peter Norvig Spelling Corrector. This method uses the concept of probability to predict the possible closeness between the typographical error and the available words in the corpus. For example, there is a typo "*Kmu*" (correct word: "*Kamu*", in English: "You"), the method will search for word candidates that are close to the actual word by using candidate models such as splits, insertion, transposition, substitution, and deletion operations. It will produce several improvement word candidates such as "*Kam*", "*Km*", "*Kmu*", "*umK*", "*Kamu*", "*eKm*", "*aKmu*", and others. Even the word will be combined with the characters "a" to "z". When the word candidate is calculated with the probability in the corpus, the closest improvement to the typographical error "*Kmu*" is the candidate word "*Kamu*" in the corpus. Figure 5 shows the process of detecting spelling errors based on typos using Peter Norvig Spelling Corrector.

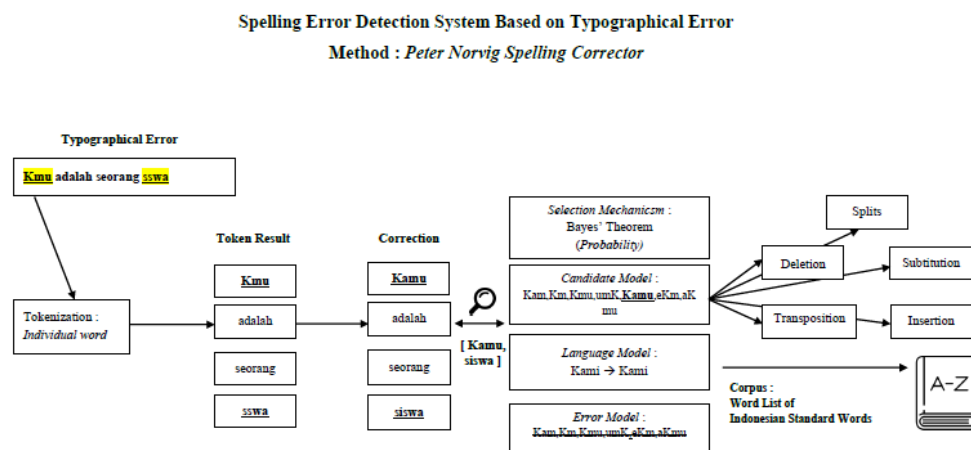


Figure 5. Typographical Error-based Spelling Error Detection using Peter Norvig Spelling Corrector

Another example, there is a typographical error in the word "slamat" (the correct word is "selamat", in English: "congratulations"). Peter Norvig Spelling Corrector will look for combinations of characters in their words to get the correct word and match the words in the corpus based on probability. Table 1 shows the process of this method for finding word corrections using splits, deletions, transpositions, substitutions, and insertions.

Table 1. Peter Norvig Spelling Corrector for Finding Typographical Error using Splits, Deletions, Transpositions, Substitutions, and Insertions.

Incorrect words (typo)	'slamat'						
Splits	'', 'slamat'	's', 'lamat'	'sl', 'amat'	'sla', 'mat'	slam', 'at'	'slama', 't'	'slamat', ''
Deletion	'lamat'	'samat'	'slmat'	'slaat'	'slamt'	'slama'	
Transposition	'lsamat'	'salmat'	'slmaat'	'slaamt'	'slamta'		
Substitution	'alamat'	'saamat', 'blamat'	'slamat', 'slbmat'	'slaaat', 'slabat'	'slamat', 'slambt'	'slamaa', 'slamab'	
	...	...	...	...	...	...	
	...	...	...	...	...	...	
Insertion	'zlamat'	'szamat'	'slzmat'	'slazat'	'slamzt'	'slamaz'	
	'aslamat'	'salamat'	'slaamat'	'slaamat'	'slamaat'	'slamaat'	'slamata'
	'bslamat'	'sblamat'	'slbamat'	'slabmat'	'slambat'	'slamabt'	'slamatb'
	...	...	...	...	...	...	...
	'eslamat'	'selamat'	'sleamat'	'slaemat'	'slameat'	'slamaet'	'slamate'
	...	...	...	...	...	...	...
	'zslamat'	'szelamat'	'slzamat'	'slazmat'	'slamzat'	'slamazt'	'slamatz'
Correct words	'selamat'						

3.4. KEBI 1.0 Checker Application User Interface

Users can access the KEBI 1.0 Checker by following the link of <https://kebi.id/>. Figure 6 shows the user interface of the KEBI 1.0 Checker application which offers three main features namely punctuation marks detection, *typo* detection and non-standard word detection. Users can input sentences by typing a word into the provided editor box or entering an extension document *.docx or *.pdf directly by clicking the button of "Pilih File" ("Select File"). After that, the user click the button of "Periksa" ("Check") until the results of spelling error detection will appear. Users can also export spelling error detection results in *.docx format and can check a fairly short detection process time within a few seconds depending on the size of the document examined.

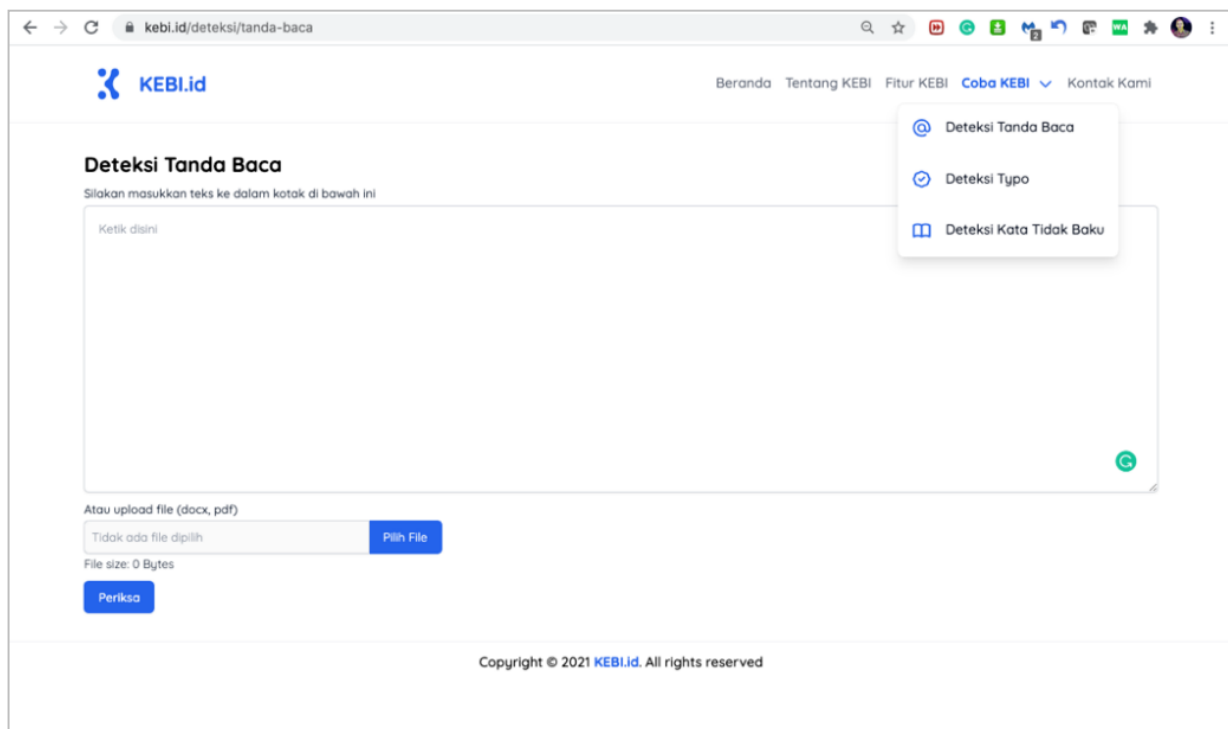


Figure 6. The User Interface of KEBI 1.0 Checker Application

3.5. Analysis and Comparison of KEBI 1.0 Checker Framework with Other Application Framework of Spelling Error Detection

Identification of Typographical Errors in Indonesian Language Documents Using N-gram and Levenshtein Distance Methods

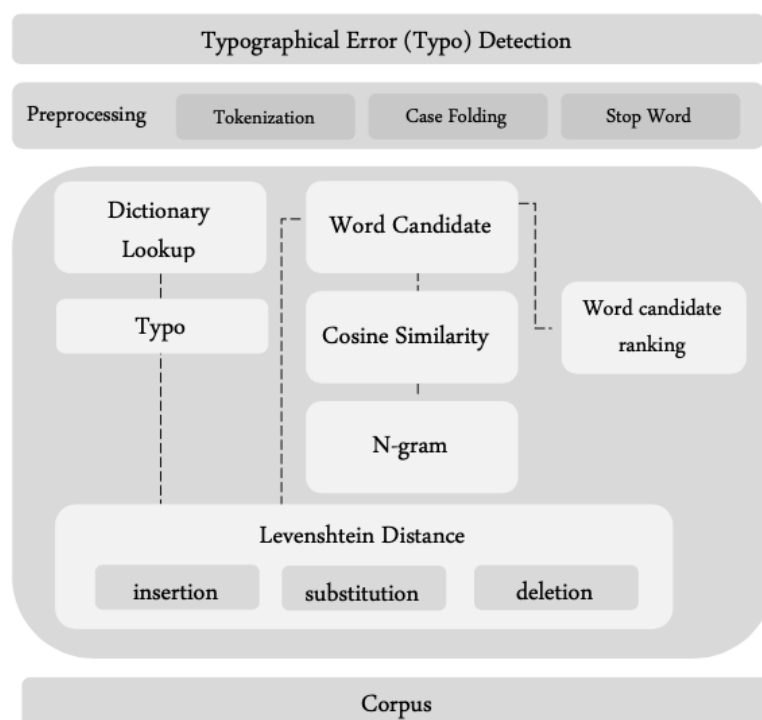


Figure 7. The First Related Research Framework of Spelling Error Detection Application

Figure 7 shows the application framework of the first related research entitled identification of typographical errors in Indonesian language documents using the N-gram and Levenshtein Distance methods [9]. The researcher proposes a special spelling error detection application to handle typo words. The developed application works starting with pre-processing methods such as tokenization, case folding, and stop words. Furthermore, documents that have become tokens will be checked by a dictionary lookup to check for standard and correct words in the corpus, leaving words that are not available in the corpus or typos may be detected. The list of words detected as typos will be corrected by the Levenshtein Distance method using insertion, substitution, and deletion operations to produce word improvement candidates. The list of selected word candidates will be calculated using N-gram and Cosine Similarity to get word candidates that are most similar to those in the corpus.

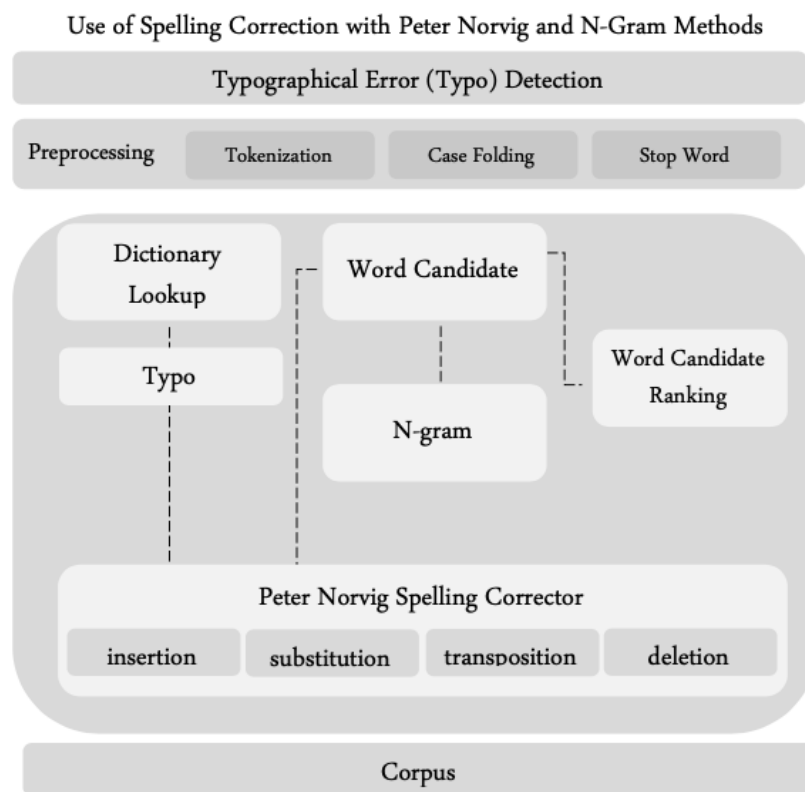


Figure 8. The Second Related Research Framework of Spelling Error Detection Application

Figure 8 shows the application framework of the second related research entitled the use of spelling correction with Peter Norvig and N-gram methods [13]. The researcher proposes a spelling error detection application which is also specifically for handling typo words. The developed application begins with general preprocessing, then documents that have become tokens will be checked by a dictionary lookup to check for standard and correct words in the corpus, leaving words that are detected as typos. The researcher chose the Peter Norvig Spelling Corrector method by using insertion, substitution, transposition, and deletion operations to correct typo words to produce word improvement candidates. The list of selected word candidates will be calculated using N-grams to get word candidates that are most similar to those in the corpus.

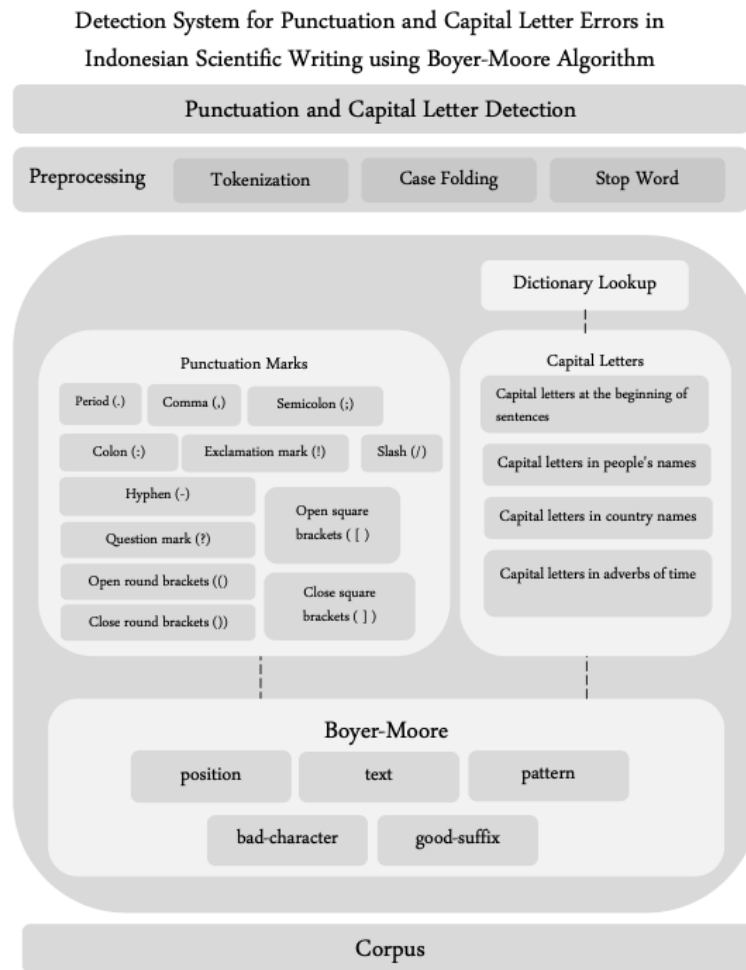


Figure 9. The Third Related Research Framework of Spelling Error Detection Application

Figure 9 shows the application framework of the third related research entitled the detection system for punctuation and capital letters errors in Indonesian scientific papers using the Boyer-Moore algorithm. [14]. The research is different from the two previous researchers, the third researcher proposes a spelling error detection application to handle punctuation and capitalization errors. The developed application works by pre-processing beginning with tokenization, case folding, and stop words like the previous related research. Furthermore, documents that have become tokens will be checked by a dictionary lookup to check for words related to people's names, names of countries and adverbs of time that must be written in capital letters and compared in the corpus. To check punctuation errors, this application has been developed by creating rules consisting of 12 punctuation marks such as period (.), comma (,), semicolon (;), colon (:), exclamation mark (!), and other sign. it will set the position of the punctuation mark with the character being either alphanumeric or alphabetic. The Boyer-Moore algorithm itself works by using three parameters such as position, text, and pattern, this will mark matches and comparisons between characters.

From the framework that has been described previously, both the KEBI 1.0 Checker application framework and the framework of the three related research that have been discussed. Furthermore, it will be analyzed to objectively assess the complexity of the features, advantages and disadvantages, methods and the level of accuracy of each application.

Table 2. Comparative Analysis of the KEBI 1.0 Checker Framework with Related Spelling Error Research and Applications

No.	Research Title	Feature Complexity	Method	Advantage	Disadvantage	Level of Accuracy
1.	KEBI 1.0 Checker[16]	Punctuation error detection, Typo word error detection, Non-standard word error detection	IF-THEN Rules, Peter Norvig Spelling Corrector, Dictionary Lookup,	<i>User Interface</i> interesting, Easy access for users, Three spelling error detection feature works well, Fast computing time, Input files in the form of documents with the extension *.docx and *.pdf, Document checker results can be exported	The feature for punctuation error detection is still limited, it does not cover all punctuation marks Features for detection of non-standard words and typo words are very dependent on the completeness of the corpus Unable to fix abbreviations, people's names, place names, and method names	Punctuation error detection (feature has been tested on a limited basis and works successfully) Typo error detection with the best accuracy reached 49.10-55.52% Non-standard word detection with the best accuracy reached 86-100% Computing time reached 0.016-21.72 seconds
2.	Identification of word writing	Typo word error detection	Dictionary Lookup,	The results of the word improvement	The feature for typo word detection is	Typo word detection with the

	errors in Indonesian language documents using the N-gram and Levenshtein Distance method [9]		Levenshtein Distance, N-gram and Cosine Similarity	candidate are quite optimal using the proposed method	highly dependent on the completeness of the corpus Unable to fix the name of the person, the name of the place, and the name of the method	best precision reached 97%
3.	Use of spelling correction with Peter Norvig and N-gram method[13]	Typo word error detection	Dictionary Lookup Peter Norvig Spelling Corrector N-gram	The results of the word improvement candidate work quite well using the proposed method	The feature for typo word detection is highly dependent on the completeness of the corpus Unable to fix the name of the person, the name of the place, and the name of the method	Typo word detection with the best accuracy reached 69.09%
4.	The detection system for punctuation and capital letter errors in Indonesian scientific papers using Boyer-Moore algorithm[14]	Detect punctuation and capital letter errors	Dictionary Lookup Boyer-Moore	The result of punctuation improvement works very well using the proposed method	Features for detection of punctuation and capital letter errors are highly dependent on the completeness of the corpus Unable to handle words	Detection of punctuation and capital letter errors with the best accuracy reached 91.7%

in the form of
abbreviations,
and there is
ambiguity in
capital letters
between
people's
names and
adjectives

4. Conclusions

Based on the results of the framework analysis can be concluded the following.

- a. The KEBI 1.0 Checker application framework has been developed to detect and correct spelling errors in writing scientific papers based on PUEBI and KBBI.
- b. KEBI 1.0 Checker offers three main features for detecting and correcting spelling errors based on punctuation errors, typos, and non-standard words. KEBI provides easy access for users through web applications with interesting user interfaces.
- c. The findings of the literature study show the number of studies that produce similar applications for the detection of spelling errors. However, the application has limited features.
- d. The results of the objective analysis of the comparison of the complexity of the features, the advantages and disadvantages, the methods and the levels of accuracy of each application show KEBI 1.0 Checker advantage in the completeness of the features. However, it is still necessary to increase the precision in correcting spelling errors in typo words.
- e. The KEBI 1.0 Checker framework will continue to be developed to enhance the three main functionalities and other support modules.

Funding: This Research was funded by the Ministry of Education, Culture, Research, and Technology, Universitas Pembangunan Nasional "Veteran" Jawa Timur based on the Assignment Agreement for the Implementation of the Batch I Internal Research Program for the Basic Research Scheme (RISDA) in 2021, Number: SPP / 12 /UN.63.8/LT/IV/2021.

References

1. S. R. S. A. M. S. Wisnu W. A. Umboh, "Rancang Bangun Aplikasi Deteksi Kesalahan Penulisan Naskah Dokumen Skripsi," *e-Jurnal Teknik Informatika*, 2017.
2. A. R. N., M. Kamayani, R. Reinanda, S. Simbolon, M. Y. Soleh and A. Purwarianti, "Application of Document Spelling Checker for Bahasa Indonesia," in *2011 International Conference on Advanced Computer Science and Information Systems*, Jakarta, 2011.
3. Supriadin, "Identifikasi Penggunaan Kosakata Baku Dalam Wacana Bahasa Indonesia Pada Siswa Kelas VII Di SMP Negeri 1 Wera Kabupaten Bima Tahun Pelajaran 2013/2014," *JIME*, vol. 2, no. 2, pp. 150-161, 2016.

4. S. Setiawati, "Penggunaan Kamus Besar Bahasa Indonesia (KBBI) Dalam Pembelajaran Kosakata Baku dan Tidak Baku pada Siswa Kelas IV SD," *Jurnal Penelitian Bahasa dan Sastra Indonesia*, vol. 2, no. 1, pp. 44-51, 2016.
5. B. D. Nurwicaksono and D. Amelia, "Analisis Kesalahan Berbahasa Indonesia Pada Teks Ilmiah Mahasiswa," *AKSIS: Jurnal Pendidikan Bahasa dan Sastra Indonesia*, vol. 2, no. 2, pp. 138-153, 2018.
6. S. Ariana, "Kesalahan Penggunaan Ejaan yang Disempurnakan dalam karya ilmiah Dosen Universitas Bina Darma," *Jurnal Bina Edukasi*, 2011.
7. N. Nazar, *Bahasa Indonesia dalam Karangan Ilmiah*, Bandung, 2004.
8. A. R. R. Nirwana, "Kemampuan Menulis Karya Tulis Ilmiah Mahasiswa Prodi Informatika Universitas Cokroaminoto Palopo," *Jurnal Onoma: Pendidikan, Bahasa dan Sastra*, vol. 6, no. 1, pp. 557-566, 2020.
9. A. C. I. & P. R. Fahma, "Identifikasi Kesalahan Penulisan Kata (Typographical Error) pada Dokumen Berbahasa Indonesia Menggunakan Metode N-gram dan Levenshtein Distance," *Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer*, pp. 53-62, 2017.
10. K. K. M. S. Sendy Agustian, "Aplikasi Koreksi Kesalahan Penulisan Kata Dalam Bahasa Inggris Dengan Menggunakan Algoritma Rabin-Karp," *Jurnal Ilmiah Informatika Komputer*, pp. 105-115, 2020.
11. R. Lauchman, *Punctuation at Work: Simple Principles for Achieving Clarity and Good Style*, Newyork: American, 2010.
12. B. A. Pratiwi, "Analisis Kesalahan Penggunaan Tanda Baca Dalam Pembelajaran Bahasa Inggris Siswa," *School Education Journal*, vol. 9, 2019.
13. D. S. N. V. C. M. Ricky Martin, "Penggunaan Spelling Correction Dengan Metode Peter Norvig Dan N-Gram," *Jurnal Ilmu Komputer dan Sistem Informasi*, pp. 175-180, 2021.
14. A. Q. Mazidhatul Ilmiyah, "Sistem Deteksi Kesalahan Tanda Baca dan Huruf Kapital Pada Karya Tulis Ilmiah Berbahasa Indonesia Menggunakan Algoritma Boyer-Moore," *JINACS: Journal of Informatics and Computer Science*, pp. 185-193, 2021.
15. M. H. Asvarizal Filcha, "Implementasi Algoritma Rabin-Karp untuk Pendeteksi Plagiarisme pada Dokumen Tugas Mahasiswa," *JUITA: Jurnal Informatika*, pp. 25-32, 2019.
16. T. M. Fahrudin, I. Sa'diyah, L. Latipah, I. Z. Atha Illah, C. C. Bey Lirna dan B. . S. Acarya, "KEBI 1.0: Indonesian Spelling Error Detection System for Scientific Papers using Dictionary Lookup and Peter Norvig Spelling Corrector," *Lontar Komputer: Jurnal Ilmiah Teknologi Informasi*, vol. 12, no. 2, pp. 78-90, 2021.