



Sentiment Analysis of ChatGPT on Indonesian Text using Hybrid CNN and Bi-LSTM

Vincentius Riandaru Prasetyo^{1*}, Mohammad Farid Naufal², Kevin Wijaya³

^{1,2,3}Department of Informatics, Faculty of Engineering, University of Surabaya, Surabaya, Indonesia

¹vincent@staff.ubaya.ac.id, ²faridnaufal@staff.ubaya.ac.id, ³s160420136@student.ubaya.ac.id

Abstract

This study explores sentiment analysis on Indonesian text using a hybrid deep learning approach that combines Convolutional Neural Networks (CNN) and Bidirectional Long Short-Term Memory (Bi-LSTM). Due to the complex linguistic structure of the Indonesian language, sentiment classification remains challenging, necessitating advanced methods to capture both local patterns and sequential dependencies. The primary objective of this research is to improve sentiment classification accuracy by leveraging a hybrid model that integrates CNN for feature extraction and Bi-LSTM for contextual understanding. The dataset consists of 800 manually labeled samples collected from social media platforms, preprocessed using case folding, stop word removal, and lemmatization. Word embeddings are generated using the Word2Vec CBOW model, and the classification model is trained using a hybrid architecture. The best performance was achieved with 32 Bi-LSTM units, a dropout rate of 0.5, and L2 regularization, which was evaluated using Stratified K-Fold cross-validation. Experimental results demonstrate that the hybrid model outperforms conventional deep learning approaches, achieving 95.24% accuracy, 95.09% precision, 95.15% recall, and 95.99% F1 score. These findings highlight the effectiveness of hybrid architectures in sentiment analysis for low-resource languages. Future work may explore larger datasets or transfer learning to enhance generalizability.

Keywords: sentiment analysis; CNN; Bi-LSTM; hybrid model

How to Cite: V. R. Prasetyo, M. F. Naufal, and K. Wijaya, "Sentiment Analysis of ChatGPT on Indonesian Text using Hybrid CNN and Bi-LSTM", J. RESTI (Rekayasa Sist. Teknol. Inf.) , vol. 9, no. 2, pp. 327 - 333, Apr. 2025.

DOI: <https://doi.org/10.29207/resti.v9i2.6334>

1. Introduction

ChatGPT is an artificial intelligence (AI) model developed by OpenAI. This model uses advanced machine learning technology, namely transformers, to understand and generate text in human language. ChatGPT is trained using various text data from the internet, so it can answer questions, provide explanations, discuss, or even create creative content such as stories, poems, or articles [1]. ChatGPT has several advantages that make it very useful in various contexts. It can communicate in many languages and understand and respond to conversations in a relevant context. In addition, ChatGPT can generate creative text in various forms, such as essays, stories, or programming code, which makes it ideal for writing or developing ideas. Users can also personalize the conversation style, whether formal or casual, according to their needs. ChatGPT can also complete complex tasks, such as explaining technical concepts or providing advice, and can be accessed through various

platforms flexibly and securely. Its advantages in scalability and ability to support various types of applications, from education to customer service, make it a convenient and efficient tool [2].

While ChatGPT in education offers various benefits, some threats and challenges must be considered. One is the potential for misuse for plagiarism, where students can use ChatGPT to write assignments without understanding the material or developing their writing skills. Over-reliance on this technology can also reduce students' ability to think critically and solve problems independently. In addition, ChatGPT does not always provide accurate information, which can spread misinformation if used without verification and hinder the development of research and writing skills that should be trained in the learning process [3]. The gap in access to technology is also an issue, as not all students have the same tools to utilize this AI, creating inequities in learning. Overuse of ChatGPT can reduce student engagement in class discussions and lead to a shallower

understanding of the material, especially in more complex topics. From an ethical perspective, using AI in education raises questions about the authenticity of assignments and fair assessments and how to assess students' abilities using this technology. Finally, reliance on ChatGPT can reduce social interaction and communication between students and instructors, which is essential for developing social and discussion skills. Thus, although ChatGPT offers convenience, it is important to use this technology wisely so as not to reduce the quality of learning and increase the student gap [4].

Deep learning techniques have been investigated several times for sentiment analysis. While Bidirectional Long Short-Term Memory (Bi-LSTM) networks shine in capturing sequential and contextual linkages, Convolutional Neural Networks (CNNs) have proven rather strong in extracting local text properties. Combining CNN with Bi-LSTM has shown potential in using the advantages of both models in hybrid techniques [5]. On the IMDB dataset, the combination had 93.3% accuracy, where CNN structure helps extract local aspects of the text. Bi-LSTM completes the context information interaction by allocating weight and resources to the text material of many, which helps to increase performance [6].

Another study using the Arabic dataset showed that hybrid CNN and Bi-LSTM performed well for binary and multiclass classification with accuracy of 98.47% and 98.92%, respectively [7]. Another study using a tweets dataset showed that the hybrid CNN and Bi-LSTM had an accuracy of 82%, better than the implementation of Bi-LSTM alone, with an accuracy of 76% [8]. In addition, using airline quality and Twitter airline sentiment datasets, the combination of CNN and Bi-LSTM also produces promising performance with an accuracy of 91.3% [9].

Specific research related to sentiment analysis from ChatGPT has been conducted using a dataset on Twitter and comparing the accuracy of 2 machine learning methods, namely Support Vector Machine (SVM) and Naïve Bayes. From the evaluation results, the two

methods were not better than the hybrid CNN and Bi-LSTM methods, with the best accuracy produced being 59% and 47% [10]. Another study using a Twitter dataset for sentiment analysis of ChatGPT was conducted by combining C4.5 and Naïve Bayes methods. The C4.5 algorithm was employed to discern Twitter usage trends, whereas the Naïve Bayes algorithm was utilized to categorize the predominant forms of interactions between Twitter users and ChatGPT. The amalgamation of the two techniques yielded an accuracy of 77.33% [11].

2. Research Methods

The processes that occur in this research consist of several stages, including dataset collection, preprocessing, word embedding, model training, model evaluation, and implementation. The flowchart of the stages in this study can be seen in Figure 1. The process begins with collecting datasets through crawling techniques from social media, followed by manual labeling into positive, negative, and neutral categories. The data then undergoes preprocessing, including case folding, removal of links and special characters, stop word removal, and lemmatization to improve text quality. Next, word embedding using the Continuous Bag of Words (CBOW) method from Word2Vec is applied to represent text in vector form. A hybrid CNN and Bi-LSTM model is used for sentiment classification, with CNN extracting local features through convolutional and max-pooling layers. At the same time, Bi-LSTM captures sequential relationships in text before being processed by a dense layer with SoftMax activation. The model is trained using Stratified K-Fold cross-validation, and hyperparameter tuning is performed by testing variations in the number of Bi-LSTM units, dropout, and regularizer (L1/L2). The model evaluation uses accuracy, precision, recall, and F1-score metrics to measure model performance in accurately classifying sentiment. This figure provides a comprehensive overview of the research process, from data collection to model evaluation and the application of hybrid architecture in Indonesian language sentiment analysis.

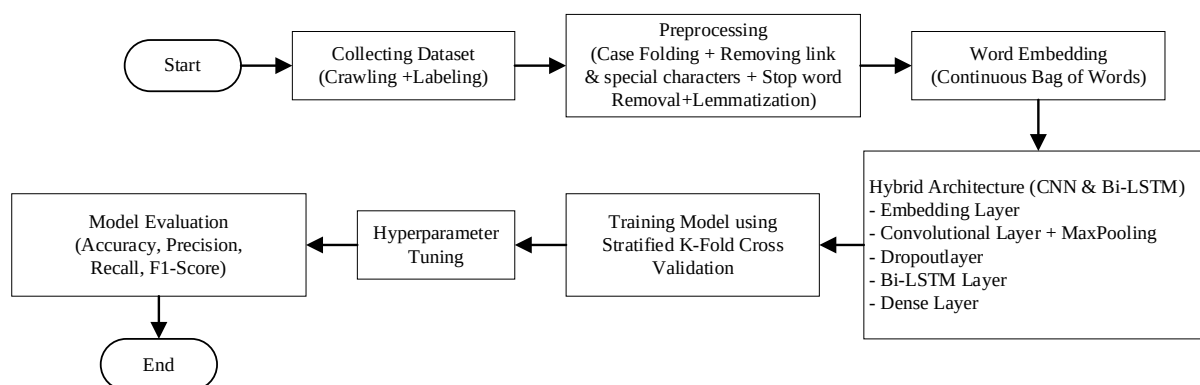


Figure 1. Research Steps

2.1 Collecting Dataset

In this study, datasets were collected using crawling techniques. Crawling is an automatic data extraction technique from websites that can be stored and analyzed [12]. Indonesian language text opinion data was taken from YouTube, Instagram, and X platforms. The three platforms were chosen because they are included in the top 10 social media platforms that are widely used by Indonesian people today [13].

Crawling is done by utilizing the selenium library. Selenium is one of the popular libraries used for web scraping or crawling, especially when the website uses a lot of JavaScript to load content [14]. An expert, Anis Sumiati, S.Pd., a Surabaya Cambridge School High School teacher, will manually label the collected dataset. The dataset is labeled into three classes: positive, negative, and neutral, where each class has a total of 392, 243, and 165 data. Table 1 shows some examples of datasets and their labels.

Table 1. Example Dataset and Its Sentiment Labels

Comment Data	Sentiment Label
Banyaak banget @. Kalau lagi baca buku atau artikel penelitian trus ada kalimat yg nggak paham biasanya aku pake chatGPT buat bantu jelasin	Positive
Bagus buat bikin tugas, tapi lama2 menurunkan intelligent siswa. Terbukti, para siswa kalau dikasih tugas paper cepet selesai tapi begitu diajak diskusi isinya Cuma bisa tolah toleh bego	Negative
Asli enak banget pakai ChatGPT buat bantu rephrase in academic term dalam nulis jurnal. Secara kapan coba gw nulis jurnal, jadi confusing bener	Positive
Penggunaan AI sejauh ini yg tak approve ya buat riset, ChatGPT sama Bard bener2 life changing for my life. Walaupun suwi2 aku wedi sisan, kemampuan literasi ku bakal menurun hhh	Neutral

2.2 Preprocessing

After the dataset is collected and labeled, preprocessing is an important stage in sentiment analysis. Preprocessing involves several processes to format text that are useful for improving the performance of the classification model built [15]. This study carried out several processes: case folding, cleaning, stop word removal, and lemmatization.

The first process is case folding, which aims to transform all character variations in the text into uniform and consistent ones so that the same word is not represented as two different vectors. This process is important because deep learning models are case-sensitive. Deep Learning treats uppercase and lowercase letters as different entities [16].

The following preprocessing removes links, mentions, punctuation, hashtags, characters, and excess spaces and changes symbols to their actual meaning. This process is important to avoid elements that do not

represent the sentiment of an opinion, so it needs to be cleaned so as not to affect the classification results [16].

The third preprocessing is stop word removal, where unimportant words in the classification process will be removed. The categories of word types included as stop words are prepositions, conjunctions, and pronouns [17]. The last preprocessing is lemmatization, where this process is like stemming because both reduce word variants into basic word forms, but the process is different. Stemming changes words into basic word forms by removing prefixes and suffixes, while lemmatization changes basic words or lemmas by considering the grammar rules in the dictionary. Table 2 shows an example of the preprocessing results of the dataset.

Table 2. Example of Preprocessing Results

Original Text	Preprocessing Result
Banyaak banget @. Kalau lagi baca buku atau artikel penelitian trus ada kalimat yg nggak paham biasanya aku pake chatGPT buat bantu jelasin	banyaak banget baca buku artikel teliti trus kalimat yg nggak paham pake chatgpt buat bantu jelasin
Bagus buat bikin tugas, tapi lama2 menurunkan intelligent siswa. Terbukti, para siswa kalau dikasih tugas paper cepet selesai tapi begitu diajak diskusi isinya Cuma bisa tolah toleh bego	bagus buat bikin tugas tapi turun intelligent siswa bukti siswa dikasih tugas paper cepet selesai tapi diajak diskusi isinya cuma bisa tolah toleh bego
Asli enak banget pakai ChatGPT buat bantu rephrase in academic term dalam nulis jurnal. Secara kapan coba gw nulis jurnal, jadi confusing bener	asli enak banget pakai chatgpt buat bantu rephrase in academic term nulis jurnal kapan coba gw nulis jurnal confusing bener
Penggunaan AI sejauh ini yg tak approve ya buat riset, ChatGPT sama Bard bener2 life changing for my life. Walaupun suwi2 aku wedi sisan, kemampuan literasi ku bakal menurun hhh	guna ai yg tak approve ya buat riset chatgpt sama bard bener life changing for my life suwi wedi sisan mampu literasi ku turun hhh

2.3. Word Embedding

The next stage after preprocessing is word embedding. The word embedding model used in this study is Word2Vec. The purpose of Word2Vec is to capture the similarity between words in the text using vector representation. Word2Vec has two algorithms: Continuous Bag of Words (CBOW) and Skip-gram [18]. The method applied in this work is CBOW. This algorithm predicts one target word with one context word as input. CBOW is chosen since it fits big datasets [18] and offers the advantage of fast data processing. Figures 2 show CBOW's architecture.

The CBOW approach starts with figuring the window size—e.g., 2—where for every target word (center word), we employ the two words before and two words following as context words [18]. Forming a context-target pair, as described in Table 3, comes next; an example sentence might be "ChatGPT adalah model AI yang sangat populer." After that, the words in the sentence are represented as one-hot encoding, as shown in Table 4.

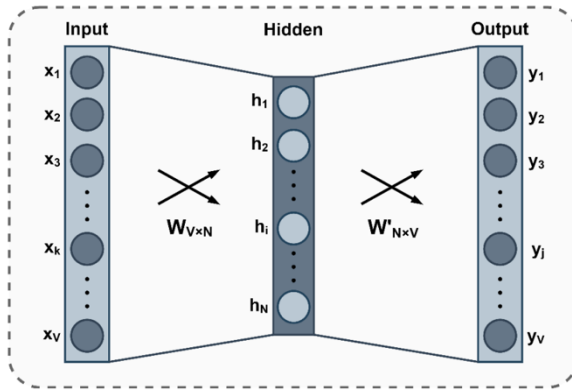


Figure 2. CBOW Architecture

Table 3. Example of Context-Target Pair

Context Words	Target Word
["ChatGPT", "adalah"]	"model"
["adalah", "model", "AI"]	"yang"
["model", "AI", "yang"]	"sangat"
["AI", "yang", "sangat"]	"populer"

Table 4. One-hot Encoding Representation of Each Word

Word	One-hot encoding
ChatGPT	[1, 0, 0, 0, 0, 0, 0]
adalah	[0, 1, 0, 0, 0, 0, 0]
model	[0, 0, 1, 0, 0, 0, 0]
AI	[0, 0, 0, 1, 0, 0, 0]
yang	[0, 0, 0, 0, 1, 0, 0]
sangat	[0, 0, 0, 0, 0, 1, 0]
populer	[0, 0, 0, 0, 0, 0, 1]

The next step is the process of forming the CBOW model [18], where in this example, the target word is "model," and the context words are ["ChatGPT," "adalah"]. Next, we calculate the average vector of the context word so that the new vector is [0.5,0.5,0,0,0,0,0]. The calculation results will be forwarded to the hidden layer to map the one-hot vector to a smaller dimensional embedding space [18], for example, from dimension 7 (number of vocabulary) to dimension 3, so that the new vector is [0.3,0.7,0.2].

The following process involves forwarding the embedding from the hidden layer to the output layer, where the probability calculation for all words in the vocabulary occurs [18]. The results of the probability calculation for all words in the vocabulary can be seen in Table 5.

Table 5 The Results of The Probability Calculation

Word	Probability
ChatGPT	0.05
adalah	0.1
model	0.65
AI	0.1
yang	0.05
sangat	0.03
populer	0.02

The CBOW model predicts the term with the highest likelihood, "model," based on the computation results in Table 5 as the target word. Should the target word prediction be erroneous, backpropagation will update the weights in the hidden layer and output layer,

improving the model's ability to predict target words depending on the context [18].

2.4 Training Model

Training the dataset with hybrid CNN and Bi-LSTM methods comes next, following the word embedding process to generate a precise classification model. This hybrid enables the model to manage text more holistically, from the local level (n-grams) to the global (word relationships in the framework of the whole text). Bi-LSTM uses local-level significant aspects to grasp the link between words in a larger context, while CNN can detect these qualities. In addition, by using CNN first for feature extraction, the length of the data sequence can be reduced, thus speeding up the process in Bi-LSTM. This combination also often produces higher accuracy than using only CNN or LSTM alone, especially on text datasets with complex patterns [5]–[9].

This study uses a hybrid CNN and Bi-LSTM architecture for sentiment analysis of Indonesian language text. This combination utilizes the advantages of CNN in extracting local features from text and Bi-LSTM's ability to understand sequential relationships between words in a global context [6]. In the first stage, the text that has gone through preprocessing is converted into a vector representation using Word2Vec with the CBOW method. This vector representation becomes input for the Convolutional Neural Network (CNN) layer, which extracts local features from text data. CNN uses a 3x3 filter (kernel) with a ReLU activation function that increases non-linearity in feature extraction. After going through the convolution process, the results are processed through Max-Pooling to reduce the dimensions of the resulting features so that computational efficiency increases, and the risk of overfitting is reduced [5]. Furthermore, the results from CNN are sent to the Bi-LSTM layer consisting of 32 units with two directions (forward and backwards), allowing the model to understand the contextual relationship before and after the words in the text. Finally, the results from Bi-LSTM are processed through a Dense Layer with a Softmax activation function to classify text sentiment into three classes: positive, negative, and neutral [8]. The hybrid CNN and Bi-LSTM architecture used in this study can be seen in Figure 3.

Based on Figure 3, the CNN architecture has 64 filters with a kernel size 3x3, allowing the model to capture important patterns in the text. The activation function used is ReLU, which helps avoid the vanishing gradient problem and improves the model's ability to extract non-linear features. Max-Pooling with a pool size of 2x2 is applied after the convolution layer to avoid overfitting, and a dropout of 0.5 is used in certain layers during training. Meanwhile, in the Bi-LSTM layer, 32 units are used with a dropout of 0.5 to maintain model generalization. In addition, the L2 regularizer is applied to prevent the model from becoming too complex and

reduce the risk of overfitting. The Softmax activation function in the output layer is used to classify text into three sentiment classes based on the highest probability. In this study, parameter selection is based on hyperparameter tuning experiments to obtain the best performance of hybrid CNN and Bi-LSTM models. Tuning was performed on units, dropout, and the regularizer. Table 6 shows the values used for tuning the three parameters. Tuning is done by testing various combinations of these three parameters.

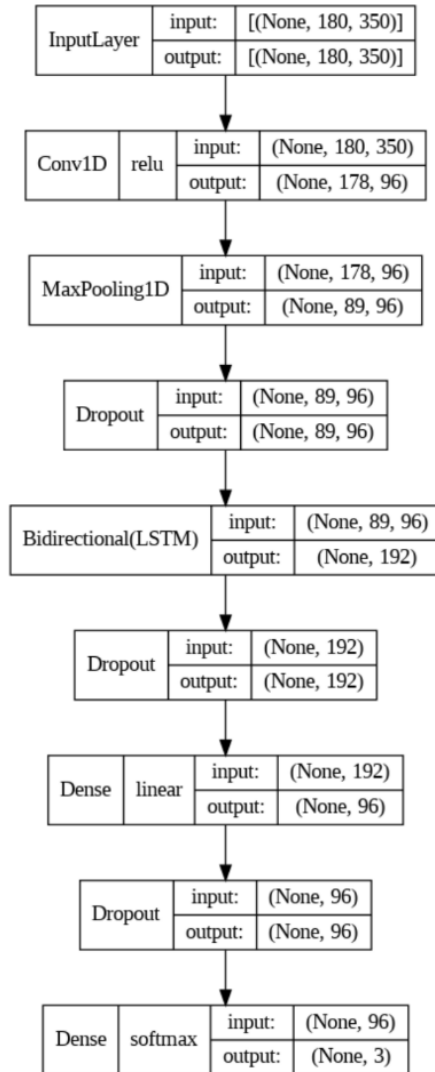


Figure 3. Hybrid Architecture of CNN and Bi-LSTM

Table 6. Hyperparameter Tuning Values

Parameters	Value
Units	32, 64, and 96
Dropout	0.4 and 0.5
Regularizer	L1 and L2

2.5 Model Evaluation

The sentiment analysis model's performance will be evaluated using four metrics: accuracy, precision, recall, and f1-score. The formulas for these four metrics can be seen in Equations 1 - 4, where *TP* is true positive, *TN* is true negative, *FP* is false positive, and *FN* is false negative [19].

$$Accuracy = \frac{TP}{TP+TN+FP+FN} \quad (1)$$

$$Precision = \frac{TP}{TP+FP} \quad (2)$$

$$Recall = \frac{TP}{TP+FN} \quad (3)$$

$$F1 - Score = \frac{2 \times Precision \times Recall}{Precision + Recall} \quad (4)$$

3. Results and Discussions

The classification model is trained using a dataset, and its performance is validated using K-fold cross-validation. Because the dataset has three classes with an unbalanced number, the Stratified K-fold Cross Validation method is used so that when the validation processes are performed, each fold has the same number of classes [20].

Table 6 shows the classification model's evaluation for each combination of hyperparameter tuning. This is done to find out which combination produces the best model performance. Based on Table 6, 12 combinations of hyperparameter tuning were produced, and the model performance for each combination can be seen in Table 7.

Table 7 Model Performance

Hyperparameters	Accuracy	Precision	Recall	F1-Score
Units: 32 Dropout: 0.4 Regularizer: L1	90.43%	90.20%	90.98%	90.70%
Units: 32 Dropout: 0.4 Regularizer: L2	94.35%	94.14%	94.43%	94.74%
Units: 32 Dropout: 0.5 Regularizer: L1	92.20%	92.12%	92.94%	92.64%
Units: 32 Dropout: 0.5 Regularizer: L2	95.24%	95.09%	95.15%	95.99%
Units: 64 Dropout: 0.4 Regularizer: L1	93.48%	93.87%	93.98%	93.09%
Units: 64 Dropout: 0.4 Regularizer: L2	94.66%	94.09%	94.65%	94.26%
Units: 64 Dropout: 0.5 Regularizer: L1	94.53%	94.41%	94.70%	94.24%
Units: 64 Dropout: 0.5 Regularizer: L2	93.83%	93.61%	93.55%	93.96%
Units: 96 Dropout: 0.4 Regularizer: L1	90.46%	90.91%	90.78%	90.35%
Units: 96 Dropout: 0.4 Regularizer: L2	91.58%	91.18%	91.41%	91.67%
Units: 96 Dropout: 0.5 Regularizer: L1	94.14%	94.05%	94.77%	94.97%
Units: 96 Dropout: 0.5 Regularizer: L2	95.81%	95.06%	95.67%	95.81%

As shown in Table 7, the ideal hyperparameter combination achieves accuracy, precision, recall, and f1-score of 95.24%, 95.09%, 95.15%, and 95.99%,

respectively. Units: 32, Dropout: 0.5, and Regularizer: L2 across all main assessment criteria; this mix produces the best results, especially in F1-Score, which acts as a balanced indicator between Precision and Recall. This is crucial for sentiment analysis tasks because erroneous classifications (false positives/false negatives) can exert a substantial influence [6][7].

A reduced Unit value, such as 32, enables the model to maintain sufficient representational ability to identify significant patterns while minimizing the number of parameters. On this dataset, larger Units (64, 96) tend to add unnecessary complexity, which can lead to overfitting or performance instability. Then, a dropout of 0.5 provides an ideal regularization level, helping the model prevent overfitting by turning off 50% of units during training. This results more consistently than a dropout of 0.4, which is too small and increases the model's noise sensitivity [7]. At last, the L2 regularizer enhances generalization by imposing a penalty on model big weights. Based on this data, the L2

regularizer performs better than the L1 since it maintains the weights small, but the distribution is more balanced, which is crucial for deep learning models, including hybrid CNN and Bi-LSTM [5][9].

In more detail, the training and validation accuracy graphs of the best models produced can be seen in Figure 4. Based on the graph, the training accuracy increases consistently, although it begins to show a slowdown in the increase after around the 60th epoch. Concurrently, the validation accuracy rises similarly to the training accuracy but at a marginally lower value. The validation accuracy aligns with the training accuracy trend, suggesting minimal overfitting in the model. Following the 60th epoch, the training and validation accuracy enhancement commences to stabilize. This indicates that the model has attained the optimal training phase and exhibits negligible enhancement in subsequent epochs. Hence, the training process concludes at the 83rd epoch.

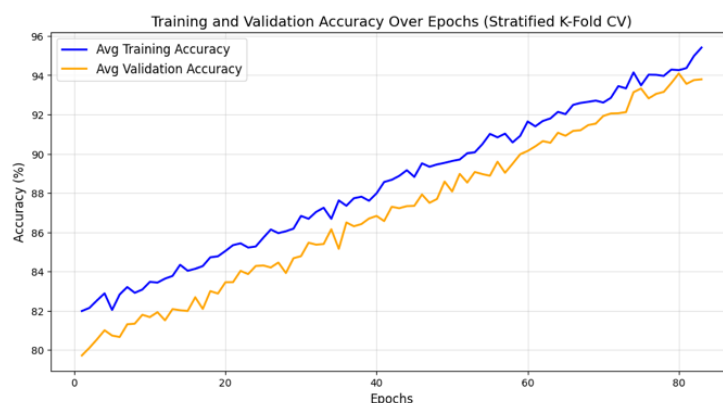


Figure 4. Graph of Training and Validation Accuracy

Additionally, the training and validation loss graphs of the best models produced can be seen in Figure 5. From the start to the end of the training, the training loss exhibits a constant decline based on the graph; early epochs show a notable fall and reach stability following epoch 60. Although it tended to be more erratic than the training loss, the validation loss also dropped significantly in the early epochs. After epoch 60, the

validation loss also began to approach stability. Fluctuations in the validation loss at some points, especially before epoch 60, were caused by stratified k-fold cross-validation, which kept the class distribution balanced in each fold. The use of stratified k-fold helps maintain the consistency of model performance. This is reflected in the not-too-sharp decrease in a loss in the validation loss, even though the dataset is unbalanced.

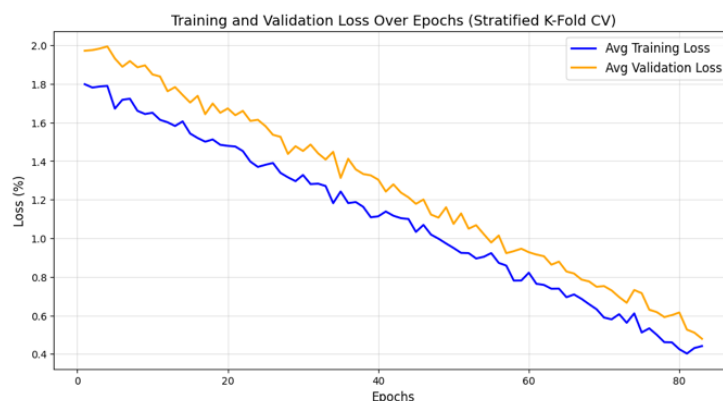


Figure 5. Graph of Training and Validation Loss

4. Conclusions

This study effectively illustrates the usefulness of a hybrid CNN and Bi-LSTM model in conducting sentiment analysis on Indonesian text datasets with three sentiment classes: positive, negative, and neutral. The findings indicate that the combination of CNN for feature extraction and Bi-LSTM for sequential information processing yields improved performance, with the ideal hyperparameters being Units: 32, Dropout: 0.5, and Regularizer: L2. Respectively, these hyperparameters generate performance with accuracy, precision, recall, and f1-score of 95.24%, 95.09%, 95.15%, and 95.99%. This approach effectively addresses the problems of extracting contextual and sequential traits in Indonesian text sentiment classification, which fewer solid techniques have always been limited. This research has likely been used in fields including customer feedback analysis, social media monitoring, and market sentiment evaluation, where accurate and quick sentiment analysis tools are vital. Adopting hybrid models highlights the requirement of using complementary architectures to overcome individual model constraints, optimizing performance for natural language processing applications. Future studies should investigate how this hybrid technique scales using more extensive, more varied datasets or how it adjusts to multilingual settings. Adding more pre-trained embeddings or complex regularizing approaches could enhance model resilience and generalization. Moreover, evaluating computation efficiency and real-time deployment options could offer interesting studies for general industrial applications.

References

- [1] J. Deng and Y. Lin, "The Benefits and Challenges of ChatGPT: An Overview," *Frontiers in Computing and Intelligent Systems*, vol. 2, no. 2, pp. 81-83, 2022. <https://doi.org/10.54097/fcis.v2i2.4465>
- [2] S. Rice, S. R. Crouse, S. R. Winter and C. Rice, "The advantages and limitations of using ChatGPT to enhance technological research," *Technology in Society*, vol. 76, 2024. <https://doi.org/10.1016/j.techsoc.2023.102426>
- [3] A. Zein, "Dampak Penggunaan ChatGPT pada Dunia Pendidikan," *Jurnal Informatika Utama*, vol. 1, no. 2, pp. 19-24, 2023. <https://doi.org/10.55903/jitu.v1i2.151>
- [4] M. Husnaini and L. M. Madhani, "Perspektif Mahasiswa terhadap ChatGPT dalam Menyelesaikan Tugas Kuliah," *Journal of Education Research*, vol. 5, no. 3, pp. 2655-2664, 2024. <https://doi.org/10.37985/jer.v5i3.1047>
- [5] V. R. Kota and S. D. Munisamy, "High accuracy offering attention mechanisms based deep learning approach using CNN/bi-LSTM for sentiment analysis," *International Journal of Intelligent Computing and Cybernetics*, vol. 15, no. 1, 2022. <https://doi.org/10.1108/IJICC-06-2021-0109>
- [6] Y. Zhou, Q. Zhang, D. Wang and X. Gu, "Text Sentiment Analysis Based on a New Hybrid Network Model," *Computational Intelligence Neuroscience*, vol. 6774320, 2022. <https://doi.org/10.1155/2022/6774320>
- [7] M. Abbes, Z. Kechaou and A. M. Alimi, "A Novel Hybrid Model Based on CNN and Bi-LSTM for Arabic Multi-domain Sentiment Analysis," in *The 17th International Conference on Complex, Intelligent and Software Intensive Systems (CISIS-2023)*, Toronto, 2023. https://doi.org/10.1007/978-3-031-35734-3_10
- [8] M. Umer, I. Ashraf, A. Mehmood, S. Kumari, S. Ullah and G. S. Choi, "Sentiment analysis of tweets using a unified convolutional neural network-long short-term memory network model," *Computational Intelligence*, vol. 37, no. 1, pp. 409-434, 2021. <https://doi.org/10.1111/coim.12415>
- [9] P. K. Jain, V. Saravanan and R. Pamula, "A Hybrid CNN-LSTM: A Deep Learning Approach for Consumer Sentiment Analysis Using Qualitative User-Generated Contents," *ACM Transactions on Asian and Low-Resource Language Information Processing*, vol. 20, no. 5, pp. 1-15, 2021. <https://doi.org/10.1145/3457206>
- [10] D. Atmajaya, A. Febrianti and H. Darwis, "Metode SVM dan Naive Bayes untuk Analisis Sentimen ChatGPT di Twitter," *The Indonesian Journal of Computer Science*, vol. 12, no. 4, pp. 2173-2181, 2023. <https://doi.org/10.33022/ijcs.v12i4.3341>
- [11] Y. Akbar and T. Sugiharto, "Analisis Sentimen Pengguna Twitter di Indonesia Terhadap ChatGPT Menggunakan Algoritma C4.5 dan Naïve Bayes," *Jurnal Sains dan Teknologi*, vol. 5, no. 1, pp. 115-122, 2023. <https://doi.org/10.55338/saintek.v4i3.1368>
- [12] V. R. Prasetyo and A. H. Samudra, "Hate speech content detection system on Twitter using K-nearest neighbor method," in *International Conference on Informatics, Technology, and Engineering 2021 (InCITE 2021)*, Surabaya, 2021. <https://doi.org/10.1063/5.0080185>
- [13] R. CNBC Indonesia, "Raja Aplikasi Terbaru di RI, Ternyata Bukan WhatsApp-Instagram," *CNBC Indonesia*, 26 February 2024. [Online]. Available: <https://www.cnbcindonesia.com/tech/20240226153650-37-517653/raja-aplikasi-terbaru-di-ri-ternyata-bukan-whatsapp-instagram>. [Accessed 10 January 2025].
- [14] M. Leotta, B. García and F. Ricca, "A family of experiments to quantify the benefits of adopting WebDriverManager and Selenium-Jupiter," *Information and Software Technology*, vol. 178, 2025. <https://doi.org/10.1016/j.infsof.2024.107595>
- [15] M. A. Palomino and F. Aider, "Evaluating the Effectiveness of Text Pre-Processing in Sentiment Analysis," *Applied Sciences*, vol. 12, no. 17, 2022. <https://doi.org/10.3390/app12178765>
- [16] M. F. Naufal and S. F. Kusuma, "Analisis Sentimen pada Media Sosial Twitter Terhadap Kebijakan Pemberlakuan Pembatasan Kegiatan Masyarakat Berbasis Deep Learning," *Jurnal Edukasi dan Penelitian Informatika*, vol. 8, no. 1, pp. 44-49, 2022. <https://doi.org/10.26418/jp.v8i1.49951>
- [17] V. R. Prasetyo, G. Erlangga and D. A. Prima, "Analisis Sentimen untuk Identifikasi Bantuan Korban Bencana Alam berdasarkan Data di Twitter Menggunakan Metode K-Means dan Naive Bayes," *Jurnal Teknologi Informasi dan Ilmu Komputer (JTIIK)*, vol. 10, no. 5, pp. 1055-1062, 2023. <https://doi.org/10.25126/jtiik.20231057077>
- [18] H. D. Abubakar, M. Umar and M. A. Bakale, "Sentiment Classification: Review of Text Vectorization Methods: Bag of Words, Tf-Idf, Word2vec and Doc2vec," *Sule Lamido University Journal of Science & Technology*, vol. 4, no. 1, pp. 27-33, 2022. <https://doi.org/10.56471/slujst.v4i.266>
- [19] O. Alharbi, "A Deep Learning Approach Combining CNN and Bi-LSTM with SVM Classifier for Arabic Sentiment Analysis," *(IJACSA) International Journal of Advanced Computer Science and Applications*, vol. 12, no. 6, pp. 165-172, 2021. <https://dx.doi.org/10.14569/IJACSA.2021.0120618>
- [20] S. Ünal, O. Günay, I. Akkurt, K. Gunoglu and H. O. Tekin, "A comparative study on breast cancer classification with stratified shuffle split and K-fold cross validation via ensemble machine learning," *Journal of Radiation Research and Applied Sciences*, vol. 17, no. 4, 2024. <https://doi.org/10.1016/j.jrras.2024.101080>



JURNAL RESTI

REKAYASA SISTEM DAN
TEKNOLOGI INFORMASI



Indexed by

Scopus

SINTA Accredited Rank 2
Decree of the Director General of
Higher Education, Research, and Technology,
Number: 158/E/KPT/2021



Published by Professional Organization of
Indonesian Informatics Experts Association
Ikatan Ahli Informatika Indonesia (IAII),
Jl. Jati Padang Raya No. 41, South Jakarta,
Pasar Minggu, Postal Code 12540

Email: jurnal.resti@gmail.com

Website Publisher :

www.iaii.or.id

Website Journal :

www.jurnal.iaii.or.id/index.php/RESTI

Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi) indexed by:



PKP | INDEX

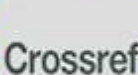


BASE
Bielefeld Academic Search Engine



DOAJ
Directory of Open Access Journals

DIRECTORY OF
OPEN ACCESS
JOURNALS



Dimensions

EDITORIAL TEAM

EDITOR IN CHIEF

Yuhefizar (Scopus ID: 55946665300, Politeknik Negeri Padang, Padang, Indonesia)

MANAGING EDITOR

Budi Sunaryo (Scopus ID: 57211253623, Universitas Bung Hatta, Padang, Indonesia)

HANDLING EDITORS

Ronal Watrianthos (Scopus ID: 57207884978, Politeknik Negeri Padang, Padang, Indonesia)

Rita Komalasari (Scopus ID: 57219130527, Politeknik LP3I, Bandung, Indonesia)

ASSOCIATE EDITORS

Ikhwan Arief (Scopus ID: 57209453335, Ambassador to Indonesia and Managing Editor at the Directory of Open Access Journals (DOAJ), Indonesia)

Shoffan Saifullah (Scopus ID: 57194569205, Universitas Pembangunan Nasional Veteran Yogyakarta, Yogyakarta, Indonesia)

Siswanto (Scopus ID: 57215048535, Universitas Budi Luhur Jakarta, Indonesia)

Teddy Mantoro (Scopus ID: 22735122000, Sampoerna University, Jakarta, Indonesia)

Media Anugerah Ayu (Scopus ID: 35589381300, Sampoerna University, Jakarta, Indonesia)

REGIONAL EDITORS (ASIA, EUROPE, NORTH AND SOUTH AMERICA)

Mohd Helmy Abd Wahab (Scopus ID: 57193949278, Universiti Tun Hussein Onn, Batu Pahat, Malaysia)

Zhihao Wang (Scopus ID: 57193132203, Southern Taiwan University of Science and Technology, Tainan, Taiwan)

Madjid Eshaghi Gordji (Scopus ID: 21733262000, Semnan University, Semnan, Iran)

Dorien J. DeTombe (Scopus ID: 6603137867, International Research Society on Methodology of Societal Complexity, Amsterdam, Netherlands)

Robbi Rahim, (Scopus ID: 57191429453, The International Association for Cryptologic Research, Lyon, France)

EDITORIAL BOARD MEMBERS

Diki Arisandi (Scopus ID: 57200087386, Universitas Muhammadiyah Riau, Pekanbaru, Indonesia)

Hendrick (Scopus ID : 57205620016, Politeknik Negeri Padang, Padang, Indonesia)

Tri Apriyanto Sundara (Scopus ID: 57208582585, STMIK Indonesia, Padang, Indonesia)

Yance Sonatha (Scopus ID: 57200988665, Politeknik Negeri Padang, Padang, Indonesia)

TECHNICAL SUPPORT

Roni Putra (Scopus ID: 57210075128, Politeknik Negeri Padang, Padang, Indonesia)

Assigned Reviewers

Ahmad Fathan Hidayatullah, Scopus ID: 57188832335, Universitas Islam Indonesia, Yogyakarta, Indonesia

Aldy Rialdy Atmadja, Scopus ID: 57189266962, UIN Sunan Gunung Djati Bandung, Bandung, Indonesia

Alfry Aristo J Sinlae, Scopus ID: 57214449083, Universitas Katolik Widya Mandira, Kupang, Indonesia

Ali Ibrahim, Scopus ID: 57203129436, Universitas Sriwijaya, Palembang, Indonesia

Anggy Trisnadoli, Scopus ID: 57188572568, Politeknik Caltex Riau, Pekanbaru, Indonesia

Anita Sindar RMS, Scopus ID: 57193870569, STMIK Pelita Nusantara, Medan, Indonesia

Bernad Jumadi Dehotman Sitompul, Scopus ID: 57210235620, Universitas Sam Ratulangi, Manado, Indonesia

Cahyo Darujati, Scopus ID: 56027926800, Narotama University, Surabaya, Indonesia

Dhanar Intan Surya Saputra, Scopus ID: 57193156799, Universitas Amikom Purwokerto, Purwokerto, Indonesia

Dian Ade Kurnia, Scopus ID: 57205059723, STMIK IKMI Cirebon, Indonesia

Dudih Gustian, Scopus ID: 57203143861, Universitas Nusa Putra, Sukabumi, Indonesia

Emi Iryanti, Scopus ID: 56180890000, Institut Teknologi Telkom Purwokerto, Purwokerto, Indonesia

Firman Tempola, Scopus ID: 57203059112, Universitas Khairun, Ternate, Indonesia

Hairani, Scopus ID: 57215535432, Universitas Bumigora, Mataram, Indonesia

Joko Kuswanto, ID Scopus: 57222344151, Universitas Baturaja, Baturaja, Indonesia

Khairan Marzuki, Scopus ID: 57280266600, Universitas Bumigora, Mataram, Indonesia

Linda Perdana Wanti, Scopus ID: 57214722091, Politeknik Negeri Cilacap, Cilacap, Indonesia

Melda Agnes Manuhutu, Scopus ID: 55973357500, Universitas Victory Sorong, Sorong, Indonesia

Muhammad Johan Alibasa, Scopus ID: 57201859953, Telkom University, Bandung, Indonesia

Muhammad Said Hasibuan, Scopus ID: 57191927671, Institut Informatika dan Bisnis Darmajaya, Lampung, Indonesia

Nur Widiyasono, Scopus ID: 57192157716, Universitas Siliwangi, Tasikmalaya, Indonesia

Oman Somantri, Scopus ID: 57208898676, Politeknik Negeri Cilacap, Cilacap, Indonesia

Prawidya Destarianto, Scopus ID: 57201194515, Politeknik Negeri Jember, Jember, Indonesia

Roheen Qamar, Scopus ID: 57991262900, Quaid-e-Awam University of Engineering, Science, and Technology (QUEST) Nawabshah, Pakistan

Rully Pramudita, Scopus ID: 57215524670, Univ. Bina Insani, Bekasi, Indonesia

Solikhun, Scopus ID: 57211272339, STIKOM Tunas Bangsa, Pematang Siantar, Indonesia

Sri Hadiani, Scopus ID: 57215526609, Universitas Nusa Mandiri, Jakarta, Indonesia

Sri Restu Ningsih, Scopus ID: 57202279477, Universitas Metamedia, Padang, Indonesia

Suhardi Aras, Scopus ID: Universitas Muhammadiyah Sorong, Papua, Indonesia

Yaya Sudarya Triana, Scopus ID: 57195103953, Universitas Mercu Buana, Jakarta, Indonesia

Table of Contents

Volume 9 No. 2 (2025-April)

Published: 2025-03-07

Comparative Analysis of Machine Learning Algorithms for Predicting Patient Admission in Emergency Departments Using EHR Data

Ahmad Abdul Chamid, Ratih Nindyasari, Muhammad Imam Ghozali

<https://doi.org/10.29207/resti.v9i2.6188>

185-194

Abstract:

Every patient who is rushed to the Emergency Department needs fast treatment to determine whether the patient should be inpatient or outpatient. However, the existing fact is that deciding whether an inpatient or outpatient must wait for the diagnosis made by the existing doctor, so if there are many patients, it generally takes quite a long time. So, to predict patient admissions to the emergency unit, a machine learning model that can be fast and accurate is needed. Therefore, this study developed a machine learning and neural network model to determine patient care in Emergency Departments. This study uses publicly available electronic health record (EHR) data, which is 3,309. The model development process uses machine learning methods (SVM, Decision Tree, KNN, AdaBoost, MLPClassifier) and neural networks. The model that has been obtained is then evaluated for its performance using a confusion matrix and several matrices such as accuracy, precision, recall, and F1-Score. The results of the model performance evaluation were compared, and the best model was obtained, namely the MLPClassifier model with an accuracy value = 0.736 and an F1-Score value = 0.635, and the Neural Network model obtained an accuracy value = 0.724 and an F1-Score value = 0.640. The best models obtained in this study, namely the MLPClassifier and Neural Network models, were proven to be able to outperform other models.

An In-depth Exploration of Sentiment Analysis on Hasanuddin Airport using Machine Learning Approaches

Lilis Nur Hayati, Fitrah Yusti Randana, Herdianti Darwis

<https://doi.org/10.29207/resti.v9i2.6253>

195-208

Abstract:

Machine learning-based sentiment analysis has become essential for understanding public perceptions of public services, including air transportation. Sultan Hasanuddin Airport, one of the main gateways in eastern Indonesia, faces the challenge of improving services amid changing user needs due to the COVID-19 pandemic. This study aims to compare the effectiveness of three machine learning algorithms-Support Vector Machine (SVM), Naive Bayes Multinomial, and K-Nearest Neighbor (KNN)-in analyzing the sentiment of user reviews related to airport services. The research also explores data splitting techniques, text preprocessing, data balancing using SMOTE, model validation, and method parameterization to ensure optimal results. The review data was retrieved from Google Maps (2021-2024) and underwent manual labelling. Text preprocessing includes normalization, stemming using Sastrawi, and stopword removal. The data-balancing technique uses SMOTE, while model evaluation is done with stratified k-fold cross-validation. SVM with a linear kernel showed the best performance, achieving an F1-score of 98.4%. Naive Bayes performed optimally, achieving an F1-score of 93.9%, while KNN recorded the best F1-score of 92.0%. SMOTE was shown to improve Naive Bayes' performance on unbalanced datasets, although it did not significantly impact SVM. The findings of this study provide data-driven recommendations to improve services at Sultan Hasanuddin Airport, such as the management of cleaning facilities, waiting room comfort, and passenger flow efficiency. In addition, this research opens up opportunities for developing real-time sentiment analysis systems that can be applied in other air transportation sectors.

Feature Selection Using Pearson Correlation for Ultra-Wideband Ranging Classification

Gita Indah Hapsari, Rendy Munadi, Bayu Erfianto, Indrarini Dyah Irawati

<https://doi.org/10.29207/resti.v9i2.6281>

209-217

Abstract:

Indoor positioning plays a crucial role in various applications, including smart homes, healthcare, robotics, and asset tracking. However, achieving high positioning accuracy in indoor environments remains a significant challenge due to obstacles that introduce NLOS conditions and multipath effects. These conditions cause signal attenuation, reflection, and interference, leading to decreased localization precision. This research addresses these challenges by optimizing feature selection LOS, NLOS, and multipath classification within Ultra-Wideband (UWB) ranging systems. A systematic feature selection approach based on Pearson correlation is employed to identify the most relevant features from an open-source dataset, ensuring efficient classification while minimizing computational complexity. The selected features are used to train multiple machine-learning classifiers, including Random Forest, Ridge Classifier, Gradient Boosting, K-Nearest Neighbor, and Logistic Regression. Experimental results demonstrate that the proposed feature selection method significantly reduces model training and testing times without compromising accuracy. The Random Forest and Gradient Boosting models exhibit superior performance, maintaining classification accuracy above 90%. The reduction in computational overhead makes the proposed approach highly suitable for real-time applications, particularly in edge-computing environments where processing efficiency is critical. These findings highlight the effectiveness of Pearson correlation-based feature selection in improving UWB-based indoor positioning systems. The optimized feature set facilitates robust LOS, NLOS, and multipath classification while reducing resource consumption, making it a promising solution for scalable and real-time indoor localization applications.

Application of Formal Concept Analysis and Clustering Algorithms to Analyze Customer Segments

I Gede Bintang Arya Budaya, I Komang Dharmendra, Evi Triandini

<https://doi.org/10.29207/resti.v9i2.6184>

218-224

Abstract:

Business development cannot be separated from relationships with customers. Understanding customer characteristics is important both for maintaining sales and even for targeting new customers with appropriate strategies. The complexity of customer data makes manual analysis of the customer segments difficult, so applying machine learning to segment the customer can be the solution. This research implements K-Means and GMM algorithms for performing clustering based on the Transaction data transformed to the Recency, Frequency, and Monetary (RFM) data model, then implements Formal Concept Analysis (FCA) as an approach to analyzing the customer segment after the class labeling. Both K-Means and GMM algorithms recommended the optimal number of clusters as the customer segment is four. The FCA implementation

in this study further analyzes customer segment characteristics by constructing a concept lattice that categorizes segments using combinations of High and Low values across the RFM attributes based on the median values, which are High Recency (HR), Low Recency (LR), High Frequency (HF), Low Frequency (LF), High Monetary (HM), and Low Monetary (LM). This characteristic can determine the customer category; for example, a customer that has HM and HR can be considered a loyal customer and can be the target for a specific marketing program. Overall, this study demonstrates that using the RFM data model, combined with clustering algorithms and FCA, is a potential approach for understanding MSME customer segment behavior. However, special consideration is necessary when determining the FCA concept lattice, as it forms the foundation of the core analytical insights.

Comparison of Sugarcane Drought Stress Based on Climatology Data using Machine Learning

Regression Model in East Java

Aries Suharso, Yeni Herdiyeni, Suria Darma Tarigan, Yandra Arkeman

<https://doi.org/10.29207/resti.v9i2.6159>

225-238

Abstract:

Crop Water Stress Index (CWSI), derived from vegetation features (NDVI) and canopy thermal temperature (LST), is an effective method to evaluate sugarcane sensitivity to drought using satellite data. However, obtaining CWSI values is complicated. This study introduces a novel approach to estimate CWSI using climatological data, including average air temperature, humidity, rainfall, sunshine duration, and wind speed features obtained from the local weather station BMKG Malang City, East Java, for the period 2021-2023. Before estimating CWSI, we analyzed sugarcane water stress phenology, examined the strength of the correlation between climatological features and CWSI, and looked at the potential for adding lag features. Our proposed prediction model uses climatological features with additional Lag features in a machine learning regression approach and 5-fold cross-validation of the training-testing data split with the help of optimization using hyperparameters. Different machine learning regression models are implemented and compared. The evaluation results showed that the prediction performance of the SVR model achieved the best accuracy with $R^2 = 90.45\%$ and $MAPE = 9.55\%$, which outperformed other models. These findings indicate that climatological features with lag effects can effectively predict water stress conditions in rainfed sugarcane if using an appropriate prediction model. The main contribution of this study is the utilization of local climatological data, which is easier to obtain and collect than sophisticated satellite data, to estimate CWSI. The application of the results shows that climatological data with lag effects can accurately estimate water stress conditions in rainfed sugarcane. In drought-prone areas, this strategy can help sugarcane farmers make better choices about land management and irrigation.

Detecting Alzheimer's Based on MRI Medical Images by Using External Attention Transformer

Farrel Ardannur Deswanto, Isman Kurniawan

<https://doi.org/10.29207/resti.v9i2.6257>

239-249

Abstract:

Alzheimer's disease is one of the major challenges in medical care this century, affecting millions of people worldwide. Alzheimer's damages neurons and connections in brain areas responsible for memory, language, reasoning, and social behavior. Early detection of this disease enables more effective treatment and proper care planning. Unfortunately, the traditional method of detecting Alzheimer's has several limitations, such as subjective analysis and delayed diagnosis. One commonly used method is visual inspection, which uses magnetic resonance imaging (MRI). The limitations of visual inspection include subjectivity and its time-consuming nature, especially with large or complex MRI datasets, making accurate interpretation a significant challenge. Therefore, an alternative for detecting Alzheimer's disease is to use deep learning-based MRI image analysis. One promising approach is to implement the External Attention Transformer (EAT) model. It enhances image classification by using two shared external memories and an attention mechanism that filters out redundant information for improved performance and efficiency. The aim of this research is to evaluate and compare the performance of the baseline Convolutional Neural Network (CNN) model, the Vision Transformer (ViT) model, and the EAT model in detecting Alzheimer's using a dataset of 6400 brain MRI images. The EAT model outperforms the baseline CNN model and ViT model in detecting Alzheimer's, achieving its best results with an accuracy of 0.965 and an F1-score of 0.747 for the test data. Our results could be integrated with clinical analysis to assist in the faster diagnosis of Alzheimer's.

Comparing Word Representation BERT and RoBERTa in Keyphrase Extraction using TgGAT

Novi Yusliani, Aini Nabilah, Muhammad Raihan Habibullah, Annisa Darmawahyuni, Ghita Athalina

<https://doi.org/10.29207/resti.v9i2.6279>

250-257

Abstract:

In this digital era, accessing vast amounts of information from websites and academic papers has become easier. However, efficiently locating relevant content remains challenging due to the overwhelming volume of data. Keyphrase Extraction Systems automate the process of generating phrases that accurately represent a document's main topics. These systems are crucial for supporting various natural language processing tasks, such as text summarization, information retrieval, and representation. The traditional method of manually selecting key phrases is still common but often proves inefficient and inconsistent in summarizing the main ideas of a document. This study introduces an approach that integrates pre-trained language models, BERT and RoBERTa, with Topic-Guided Graph Attention Networks (TgGAT) to enhance keyphrase extraction. TgGAT strengthens the extraction process by combining topic modelling with graph-based structures, providing a more structured and context-aware representation of a document's key topics. By leveraging the strengths of both graph-based and transformer-based models, this research proposes a framework that improves keyphrase extraction performance. This is the first to apply graph-based and PLM methods for keyphrase extraction in the Indonesian language. The results revealed that BERT outperformed RoBERTa, with precision, recall, and F1-scores of 0.058, 0.070, and 0.062, respectively, compared to RoBERTa's 0.026, 0.030, and 0.027. The result shows that BERT with TgGAT obtained more representative keyphrases than RoBERTa with TgGAT. These findings underline the benefits of integrating graph-based approaches with pre-trained models for capturing both semantic relationships and topic relevance.

Hand Sign Recognition of Indonesian Sign Language System SIBI Using Inception V3 Image Embedding and Random Forest

Mayang Sari, Eko Rudiawan Jamzuri

<https://doi.org/10.29207/resti.v9i2.6156>

258-265

Abstract:

This paper presents a sign language recognition system for the Indonesian Sign Language System SIBI using image embeddings combined with a Random Forest classifier. A dataset comprising 5280 images across 24 classes of SIBI alphabet symbols was utilized. Image features were

extracted using the Inception V3 image embedding, and classification was performed using Random Forest algorithms. Model evaluation conducted through K-Fold cross-validation demonstrated that the proposed model achieved an accuracy of 59.00%, an F1-Score of 58.80%, a precision of 58.80%, and a recall of 59.00%. While the performance indicates room for improvement, this study lays the groundwork for enhancing sign language recognition systems to support the preservation and broader adoption of SIBI in Indonesia.

Optimizing Multilayer Perceptron for Car Purchase Prediction with GridSearch and Optuna

Ginanti Riski, Dedy Hartama, Solikhun

<https://doi.org/10.29207/resti.v9i2.6328>

266-275

Abstract:

Multilayer Perceptron (MLP) is a powerful machine learning algorithm capable of modeling complex, non-linear relationships, making it suitable for predicting car purchasing power. However, its performance depends on hyperparameter tuning and data quality. This study optimizes MLP performance using GridSearch and Optuna for hyperparameter tuning while addressing data imbalance with the Synthetic Minority Over-sampling Technique (SMOTE). The dataset comprises demographic and financial attributes influencing car purchasing power. Initially, the dataset exhibited class imbalance, which could lead to biased predictions; SMOTE was applied to generate synthetic samples, ensuring a balanced class distribution. Two hyperparameter tuning approaches were implemented: GridSearch, which systematically explores a predefined parameter grid, and Optuna, an adaptive optimization framework utilizing a Bayesian approach. The results show that Optuna achieved the highest accuracy of 95.00% using the Adam optimizer, whereas GridSearch obtained the best accuracy of 94.17% with the RMSProp optimizer, demonstrating Optuna's superior ability to identify optimal hyperparameters. Additionally, SMOTE significantly improved model stability and predictive performance by ensuring adequate class representation. These findings offer insights into best practices for optimizing MLP in predictive modeling. The combination of SMOTE and advanced hyperparameter tuning techniques is applicable to various domains requiring accurate predictive analytics, such as finance, healthcare, and marketing. Future research can explore alternative optimization algorithms and data augmentation techniques to further enhance model robustness and accuracy.

Multi-Process Data Mining with Clustering and Support Vector Machine for Corporate Recruitment

Ruri Hartika Zain, Randy Permana, Sarjon Defit

<https://doi.org/10.29207/resti.v9i2.6197>

276-282

Abstract:

Having an efficient and accurate recruitment process is very important for a company to attract candidates with professionalism, a high level of loyalty, and motivation. However, the current selection method often faces problems due to the subjectivity of assessing prospective employees and the long process of deciding on the best candidate. Therefore, this research aims to optimize the recruitment process by applying data mining techniques to improve efficiency and accuracy in candidate selection. The method used in this research utilizes a multi-process Data Mining approach, which is a combination of clustering and classification algorithms sequentially. In the initial stage, the K-Means algorithm is applied to cluster candidates based on administrative selection data, such as document completeness and reference support. Next, a classification model was built using a Support Vector Machine (SVM) to categorize the best candidates based on the results of psychological tests, medical tests, and interviews. The experimental results show that the SVM model produces high evaluation scores, with an AUC of 87%, Classification Accuracy (CA) of 90%, F1-score of 89%, Precision of 91%, and Recall of 90%. With these results, it can be concluded that this model is able to improve accuracy in the employee selection process and help companies make more measurable and data-based recruitment decisions.

Word2Vec Approaches in Classifying Schizophrenia Through Speech Pattern

Putri Alysia Azis, Tenriola Andi, Dewi Fatmarani Suriyanto, Nur Azizah Eka Budiarti, Andi Akram Nur Risal, Zulhajji Zulhajji

<https://doi.org/10.29207/resti.v8i6.5871>

283-295

Abstract:

Schizophrenia is a chronic brain disorder characterized by symptoms such as delusions, hallucinations, and disorganized speech, posing significant challenges for accurate diagnosis. This research investigates an innovative Natural Language Processing (NLP) framework for classifying the speech patterns of schizophrenia patients using Word2Vec, with the aim of determining whether there are significant differences between the two features. The dataset comprises speech transcriptions from 121 schizophrenia patients and 121 non-schizophrenia participants collected through structured interviews. This study compares two Word2Vec architectures, Continuous Bag-of-Words (CBOW) and Skip-Gram (SG), to determine their effectiveness in classifying schizophrenia speech patterns. The results indicate that the SG architecture, with hyperparameter tuning, produces more detailed word representations, particularly for low-frequency words. This approach yields more accurate classification results, achieving an F1-score of 93.81%. These results emphasize the effectiveness of the framework in handling structured and abstract linguistic patterns. By utilizing the advantages of both static and contextual embedding, this approach offers significant potential for clinical applications, providing a reliable tool for improving schizophrenia diagnosis through automated speech analysis.

Deep learning with Bayesian Hyperparameter Optimization for Precise Electrocardiogram Signals Delineation

Annisa Darmawahyuni, Winda Kurnia Sari, Nurul Afifah, Siti Nurmaini, Jordan Marcelino, Rendy Isdwanta

<https://doi.org/10.29207/resti.v9i2.6171>

297-302

Abstract:

Electrocardiography (ECG) serves as an essential risk-stratification tool to observe further treatment for cardiac abnormalities. The cardiac abnormalities are indicated by the intervals and amplitude locations in the ECG waveform. ECG delineation plays a crucial role in identifying the critical points necessary for observing cardiac abnormalities based on the characteristics and features of the waveform. In this study, we propose a deep learning approach combined with Bayesian Hyperparameter Optimization (BHO) for hyperparameter tuning to delineate the ECG signal. BHO is an optimization method utilized to determine the optimal values of an objective function. BHO allows for efficient and faster parameter search compared to conventional tuning methods, such as grid search. This method focuses on the most promising search areas in the parameter space, iteratively builds a probability model of the objective function, and then uses that model to select new points to test. The used hyperparameters of BHO contain learning rate, batch size, epoch, and total of long short-term memory layers. The study resulted in the development of 40 models, with the best model achieving a 99.285 accuracy, 94.5% sensitivity, 99.6% specificity, and 94.05% precision. The ECG delineation-based deep learning with BHO shows its excellence for localization and position of the onset, peak, and offset of ECG waveforms. The proposed model can be applied in medical applications for ECG delineation.

Customer Satisfaction Evaluation in Online Food Delivery Services: A Systematic Literature Review

Adimas Fiqri Ramdhansya, Shella Maria Vernanda, Indra Budi, Prabu Kresna Putra, Aris Budi Santoso

<https://doi.org/10.29207/resti.v9i2.6205>

303-313

Abstract:

The rapid growth of online food delivery services has heightened the need for effective customer satisfaction measurement. This systematic literature review examines 476 papers, selecting 15 key studies to identify prevailing evaluation approaches. Findings reveal that sentiment analysis and PLS-SEM are the most frequently used analytical methods, each appearing in six studies. Satisfaction measurement relies on sentiment polarity scores in five studies and SERVQUAL frameworks in three studies. Data collection primarily involves surveys in seven studies and user-generated content in six studies, but limited demographic diversity reduces generalizability. Three key future research directions emerge. Advanced analytical techniques appear in 5 of 11 future works in the analysis methods domain. Expanding evaluation metrics is mentioned in 6 of 12 proposals in the evaluation domain. Exploring demographic context is highlighted in 10 of 25 recommendations in the dataset's domain, with dataset development receiving twice the attention of methodological advancements. These results provide researchers with a structured framework for customer satisfaction evaluation while guiding food delivery platforms in refining service quality. By systematically mapping current methodologies and future priorities, this study bridges gaps between academia and industry, ensuring more effective customer satisfaction assessments.

Measuring Factors of Trust in the Use of E-Government: A Multi-Factor Analysis of the E-Government in Indonesia

Iqbal Caraka Altino, Reska Nugroho Sudarto, Dana Indra Sensuse, Sofian Lusa, Prasetyo Adi Wibowo Putro,

Sofiyanti Indriasar, Bramanti Brillianto

<https://doi.org/10.29207/resti.v9i2.6016>

314-326

Abstract:

The implementation of dynamic records management applications within the Indonesian government remains relatively limited, with a lack of comprehensive integration between authorised institutions at both the central and regional levels. This research examines the impact of technical aspects, government agency variables, citizen variables, and risk indicators on trust in e-government. Furthermore, this study seeks to establish the effect of social factors and the advantages of trust in e-government. Finally, this research shows how trust in e-government influences satisfaction, willingness to use, and acceptance of e-government. The study examined 117 respondents using the integrated dynamic archival information system - SRIKANDI. Technical and risk factors were found to positively influence trust in e-government, with effects on satisfaction, intention to use, and adoption of e-government. Those who trusted SRIKANDI were more likely to utilize and implement the program. The findings indicate that for civil servants, trust in the government is also a factor influencing the utilisation of e-government services.

Sentiment Analysis of ChatGPT on Indonesian Text using Hybrid CNN and Bi-LSTM

Vincentius Riandaru Prasetyo, Mohammad Farid Naufal, Kevin Wijaya

<https://doi.org/10.29207/resti.v9i2.6334>

327-333

Abstract:

This study explores sentiment analysis on Indonesian text using a hybrid deep learning approach that combines Convolutional Neural Networks (CNN) and Bidirectional Long Short-Term Memory (Bi-LSTM). Due to the complex linguistic structure of the Indonesian language, sentiment classification remains challenging, necessitating advanced methods to capture both local patterns and sequential dependencies. The primary objective of this research is to improve sentiment classification accuracy by leveraging a hybrid model that integrates CNN for feature extraction and Bi-LSTM for contextual understanding. The dataset consists of 800 manually labeled samples collected from social media platforms, preprocessed using case folding, stop word removal, and lemmatization. Word embeddings are generated using the Word2Vec CBOW model, and the classification model is trained using a hybrid architecture. The best performance was achieved with 32 Bi-LSTM units, a dropout rate 0.5, and L2 regularization, which was evaluated using Stratified K-Fold cross-validation. Experimental results demonstrate that the hybrid model outperforms conventional deep learning approaches, achieving 95.24% accuracy, 95.09% precision, 95.15% recall, and 95.99% F1 score. These findings highlight the effectiveness of hybrid architectures in sentiment analysis for low-resource languages. Future work may explore larger datasets or transfer learning to enhance generalizability.

Implementation of Generative Language Models (GLM) in Cyber Exercise Secure Coding using Prompt Engineering

Jackson Sidabutar, Alfido Osdie

<https://doi.org/10.29207/resti.v9i2.6012>

334-342

Abstract:

With the advancement of technology, the need for secure software is becoming increasingly urgent due to the rise in vulnerabilities in applications. In 2022, the National Cyber and Encryption Agency (BSSN) recorded 2,348 cases of web defacement, with one of the main causes being the lack of attention to secure coding practices during software development. This study explores the utilization of Generative Language Models (GLMs), such as ChatGPT, in secure coding training to enhance developers' skills. GLMs were implemented in a cybersecurity platform designed specifically for secure coding training, also serving as learning assistants that users can interact with during the cyber exercise. The study results show that the cyber exercise using GLMs significantly improved users' secure coding skills, as evidenced by comparing pre-test and post-test scores, indicating an increase in knowledge and proficiency in secure coding practices.

Large Language Model-Based Extraction of Logic Rules from Technical Standards for Automatic Compliance Checking

Rizky Nugroho, Adila Krisnadhi, Ari Saptawijaya

<https://doi.org/10.29207/resti.v9i2.6285>

343-356

Abstract:

In this research, we design logic rules as a representation of technical standards documents related to ship design, which will be used in automatic compliance checking. We present a novel design of logic rules based on a general pattern of technical standards' clauses that can be produced automatically from text using a large language model (LLM). We also present a method to extract said logic rules from text. First, we design data structures to represent the technical standards and logic rules used to process the data. Second, the representation of technical standards is produced manually and tested to ensure that it can give the same conclusion as human judgment regarding compliance. Third, a

variation of prompting methods, namely pipeline method and few-shot prompting, is given to LLM to instruct it to extract logic rules from text following the design. Evaluation against the logic rules produced shows that the pipeline method gives an accuracy score of 0.57, a precision of 0.49, and a recall of 0.62. On the other hand, logic rules extracted using few-shot prompting have an accuracy score of 0.33, precision of 0.43, and recall of 0.5. These results show that LLM is able to extract a logic rule representation of technical standards. Furthermore, the representation resulting from the prompting technique that utilizes the pipeline method has a better performance compared to the representation resulting from few-shot prompting.

Enhancing Problem-Solving Reliability with Expert Systems and Krulik-Rudnick Indicators

Lita Sari, Jufriadif Ma'am, Addini Yusmar, Khairiyah Khadijah, Sri Wahyuni, Naufal Ibnu Salam

<https://doi.org/10.29207/resti.v9i2.6333>

357-363

Abstract:

Problem-solving is one of the skills needed in the 21st century, but there is a significant gap between the ideal conditions and the reality of students' problem-solving skills. One method that can improve students' problem-solving skills is Krulik and Rudnick, but implementing this method with an expert system to improve problem-solving skills is still limited. This research aims to build an expert system to determine the level of problem-solving using Krulik and Rudnick's problem-solving indicators processed using the forward chaining and certainty factor algorithms. The study had five stages: data analysis, rule generation, certainty measurement, prediction, and testing. The data was processed by developing 5 Krulik and Rudnick problem-solving indicators into 35 statements. Each statement was categorized using Forward Chaining by producing three rules: low, medium, and high. The problem-solving level obtained is calculated using the Certainty Factor for a confidence value. The system's prediction results were evaluated using a confusion matrix, resulting in an accuracy of 80%, a precision of 92%, and a recall of 85%, indicating the system's reliable performance in measuring the level of problem-solving. This research can be used as a reference to support problem-solving in various more advanced educational and professional environments.

Strategic Approach to Enhance Information Security Awareness at ABC Agency

Fandy Husaenul Hakim, Muhammad Hafizhuddin Hilman, Setiadi Yazid

<https://doi.org/10.29207/resti.v9i2.6333>

364-373

Abstract:

Information security awareness (ISA) is crucial to an organization's cybersecurity strategy, particularly since employees are often the last defense against cyberattacks. Despite regular communication on cybersecurity threats, the ABC Agency has not evaluated the level of ISA among its employees, leaving a gap in understanding the effectiveness of its awareness programs. This is critical, as the agency handles highly confidential data that could be at risk of accidental or intentional leaks. The Kruger Approach and the Human Aspect of Information Security Questionnaire (HAIS-Q) were used in this study to measure the ISA levels of employees at the ABC Agency. We employed the Analytic Hierarchy Process (AHP) method to analyze data collected from 86 respondents. The findings indicate that ABC Agency employees demonstrate satisfactory ISA overall. However, the "Internet Use" dimension received a medium rating, underscoring the necessity for focused enhancements in this domain. These results underscore the importance of tailoring information security awareness programs to address specific weaknesses. We provide strategic recommendations to enhance the agency's cybersecurity posture. Furthermore, this study opens avenues for future research on ISA measurement across various public and private organizations.

Securing Electronic Medical Documents Using AES and LZMA

Toto Raharjo, Yudi Prayudi

<https://doi.org/10.29207/resti.v9i2.6260>

374-384

Abstract:

With increasing threats in cyberspace, maintaining the integrity of electronic medical data is crucial. This study aims to develop a method that integrates encryption using Advanced Encryption Standard (AES) and compression with the Lempel-Ziv-Markov Algorithm (LZMA) to protect DICOM files containing sensitive information. This method is designed to address two main challenges: the growth of file sizes after the encryption process and the efficiency in data storage. In this study, an experimental design with random sampling was applied, testing 427 DICOM files from open libraries ranging in size from 513.06 KB to 513.39 KB to evaluate the implementation of this method in reducing file size, encryption time, and maintaining data integrity. The results show that this method is able to reduce file size by between 40-50% with an average encryption time of about 0.2-0.3 seconds per file. In addition, the data remains intact before and after the encryption process, which indicates that the integrity of the data is well maintained. Further analysis revealed that CPU usage during the encryption process reached 94.05%, while memory usage was recorded at 92.95 KB. In contrast, in the decryption process, CPU usage decreased to 78.16% with a much lower memory consumption, which was 31.07 KB. The findings have significant implications for medical information systems, allowing developers to easily implement these methods through APIs. This research is expected to be a reference for future studies that focus on data security in health information systems and provide new insights into the combination of encryption and compression in the context of medical data.

Enhanced Heart Disease Diagnosis Using Machine Learning Algorithms: A Comparison of Feature Selection

Hirmayanti, Ema Utami

<https://doi.org/10.29207/resti.v9i2.6175>

385-392

Abstract:

Heart disease or cardiovascular disease is one of the leading causes of death in the world. Based on WHO data, in 2019, as many as 17.9 million people died from cardiovascular disease. If early prevention is not carried out immediately, of course, the victims will increase every year. Therefore, with the increasingly rapid development of technology, especially in the health sector, it is hoped that it can help medical personnel in treating patients suffering from various diseases, especially heart disease. So in this study, it will be more focused on the selection of relevant features or attributes to increase the accuracy value of the Machine Learning algorithm. The algorithms used include Random Forest and SVM. Meanwhile, for feature selection, several feature selection techniques are used, including information gain (IG), Chi-square (Chi2) and correlation feature selection (CFS). The use of these three techniques aims to obtain the main features so that they can minimize irrelevant features that can slow down the machine process. Based on the results of the experiment with a comparison of 70:30, it shows that CFS-SVM is superior by using nine features, which obtain the highest accuracy of 92.19%, while CFS-RF obtains the best value with eight features of 91.88%. By using feature selection and hyperparameter techniques, SVM obtained an increase of 10.88%, and RF obtained an increase of

9.47%. Based on the performance of the model using the selected relevant features, it shows that the proposed CFS-SVM shows good and efficient performance in diagnosing heart disease.

Improving Government Helpdesk Service With an AI-Powered Chatbot Built on the Rasa Framework

Wirat Moko Hadi Sasmita, Surya Sumpeno, Reza Fuad Rachmadi

<https://doi.org/10.29207/resti.v9i2.6293>

393-403

Abstract:

Helpdesk services are an important component in supporting Information Technology (IT) services. The helpdesk operates as the initial interface for managing and resolving concerns. Helpdesk helps user to get solutions when facing problems while using an IT service. This research focuses on the impact of artificial intelligence (AI)-powered chatbots on the performance of the initial response of government helpdesk services. The chatbot is designed to improve service performance by quickly identifying and classifying reported issues and automatically responding to messages, enabling faster responses. The research proposed a new System Design of a helpdesk system with an AI-based chatbot. The data used comes from Telegram group chat logs, exported in JSON format. We find that the Rasa NLU model with DIET Classifier successfully achieved an accuracy rate of 0.825 in classifying intents, with the precision value of 0.838, recall of 0.829, and F1 score of 0.821 using a Rasa model with cross-validation, where folds is 5 in evaluation. And initial response time was highly improved after using chatbot artificial intelligence from more than 3 hours on the telegram group helpdesk based to an average of 2.15 seconds. These research results suggest AI-Chatbot-based ability to assist the helpdesk team in handling user queries and reports, and improving initial time response.

Efficient Hybrid Network with Prompt Learning for Multi-Degradation Image Restoration

Muhammad Yusuf Kardawi, Laksmi Rahadiani

<https://doi.org/10.29207/resti.v9i2.6381>

404-415

Abstract:

Image restoration aims to repair degraded images. Traditional image restoration methods have limited generalization capabilities due to the difficulty in dealing with different types and levels of degradation. On the other hand, contemporary research has focused on multi-degradation image restoration by developing unified networks capable of handling various types of degradation. One promising approach is using prompts to provide additional information on the type of input images and the extent of degradation. Nonetheless, all-in-one image restoration requires a high computational cost, making it challenging to implement on resource-constrained devices. This research proposes a multi-degradation image restoration model based on PromptIR with lower computational cost and complexity. The proposed model is trained and tested on various datasets yet it is still practical for deraining, dehazing, and denoising tasks. By unification convolution, transformer, and dynamic prompt operations, the proposed model successfully reduces FLOPs by 32.07% and the number of parameters by 27.87%, with a comparable restoration result and an SSIM of 34.15 compared to 34.33 achieved by the original architecture for the denoising task.

Combining the Cellular Automata and Marching Square to Generate Maps

Viore, Wirawan Istiono

<https://doi.org/10.29207/resti.v9i2.6241>

416-424

Abstract:

As computer technology advances, one of the entertainment media that has emerged is video games. The development of a video game is becoming more expensive and labor-intensive as technology itself continues to grow. One of the characteristics of a game as an entertainment medium is its replay value, which refers to the fact that the subject matter can be played more than once. Automating content through the use of procedural content generation is done with the goal of lowering expenses and reducing the amount of labour that is required. This research has two goals: designing and developing a Maze Game using the Procedural Content Generation method with the Cellular Automata and Marching Square algorithms, and determining the level of player satisfaction with the games developed using the Game User Experience Satisfaction Scale (GUESS) method. This research will utilize Cellular Automata and the Marching Square algorithm as a method for generating 3D game shapes through Procedural Content Generation. After the game has been developed, it will be performed by players, and the Game User Experience Satisfaction Scale will be used to measure the user experience. The result for overall satisfaction, based on the responses of 25 respondents, is 83.14%. Cellular Automata was effectively implemented to generate the map, while Marching Square was used to generate the 3D mesh, albeit with isolated rooms and graphical errors.

Classification Model for Bot-IoT Attack Detection Using Correlation and Analysis of Variance

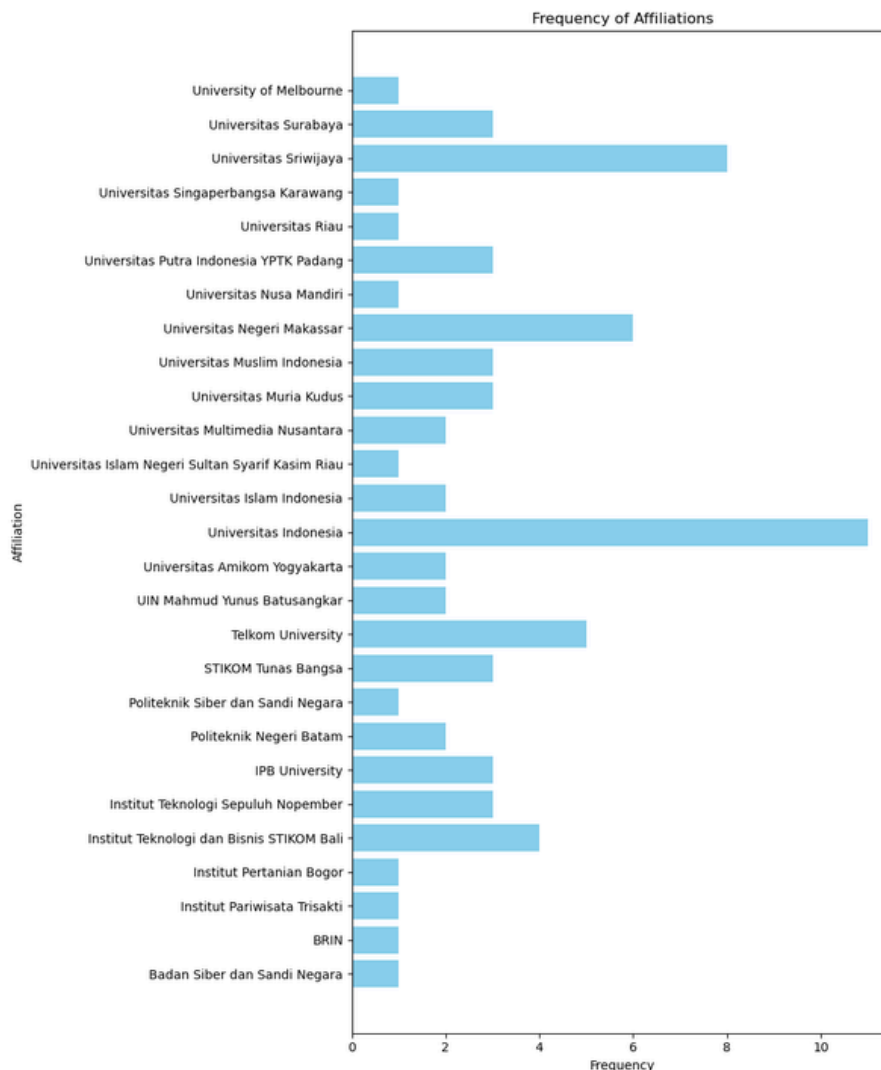
Firgiawan Faira, Dandy Pramana Hostiadi

<https://doi.org/10.29207/resti.v9i2.6332>

425-434

Abstract:

Industry 4.0 requires secure networks as the advancements in IoT and AI exacerbate the challenges and vulnerabilities in data security. This research focuses on detecting Bot-IoT activity using the Bot-IoT UNSW Canberra 2018 dataset. The dataset initially showed a significant imbalance, with 2,934,447 entries of attack activity and only 370 entries of normal activity. To address this imbalance, an innovative data aggregation technique was applied, effectively reducing similar patterns and trends. This approach resulted in a balanced dataset consisting of 8 attack activity points and 80 normal activity points. Feature selection using the ANOVA method identified 10 key features from a total of 17: seq, stddev, N_IN_Conn_P_SrcIP, min, state_number, mean, N_IN_Conn_P_DstIP, drate, srate, and max. The classification process utilized Random Forest, k-NN, Naïve Bayes, and Decision Tree algorithms, with 100 iterations and an 80:20 training-testing split. Random Forest showed superior performance, achieving 97.5% accuracy, 97.4% precision, and 97.4% recall, with a total computation time of 11.54 seconds. Pearson correlation analysis revealed a strong positive correlation (+0.937) between N_IN_Conn_P_DstIP and seq, as well as a weak negative correlation (-0.224) between N_IN_Conn_P_SrcIP and state_number. The novelty of this research lies in the application of a data aggregation technique to address class imbalance, significantly improving machine learning model performance and optimizing training time. These findings contribute to the development of robust cybersecurity systems to effectively detect IoT-related threats.



Diversity Authors



Word Cloud from Keywords

Get More with
SINTA Insight

Go to Insight



JURNAL RESTI (REKAYASA SISTEM DAN TEKNOLOGI INFORMASI)

IKATAN AHLI INFORMATIKA INDONESIA

P-ISSN : 25800760 <> E-ISSN : 25800760 Subject Area : Science, Engineering



3.25316
Impact



12995
Google Citations



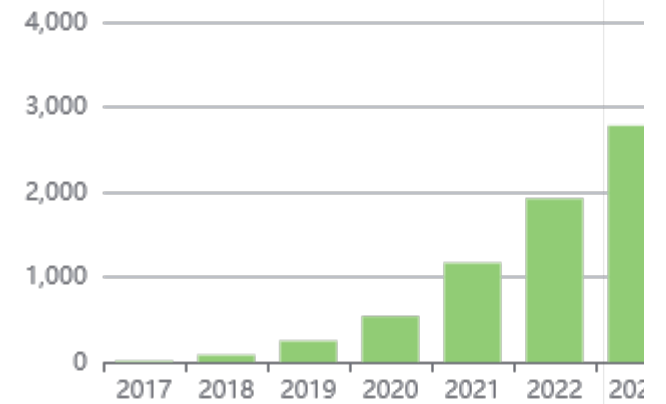
Sinta 2
Current
Accreditation

Google Scholar Garuda Website Editor URL

History Accreditation

2017 2018 2019 2020 2021 2022 2023 2024 2025 2026

Citation Per Year By Google Scholar



Journal By Google Scholar

	All	Since 2020
Citation	12995	12557
h-index	45	44
i10-index	360	350

Garuda Google Scholar

Agricultural Cultivation Cost Prediction Using Neural Networks and Feature Importance Analysis

Ikatan Ahli Informatika Indonesia (IAII) Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi) Vol 9 No 1 (2025): February 2025 31 - 39

2025 DOI: 10.29207/resti.v9i1.6003 Accred : Sinta 2

Impact of Adaptive Synthetic on Naïve Bayes Accuracy in Imbalanced Anemia Detection Datasets

Ikatan Ahli Informatika Indonesia (IAII) Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi) Vol 9 No 1 (2025): February 2025 85 - 93

2025 DOI: 10.29207/resti.v9i1.6031 Accred : Sinta 2

The Internet-of-Things-based Fishpond Security System Using NodeMCU ESP32-CAM Microcontroller

Ikatan Ahli Informatika Indonesia (IAII) Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi) Vol 9 No 1 (2025): February 2025 51 - 61

2025 DOI: 10.29207/resti.v9i1.6033 Accred : Sinta 2

Enhancing Network Security: Evaluating SDN-Enabled Firewall Solutions and Clustering Analysis Using K-Means through Data-Driven Insights

Ikatan Ahli Informatika Indonesia (IAII) Jurnal RESTI (Rekayasa Sistem dan Teknologi Informasi) Vol 9 No 1 (2025): February 2025 69 - 76

2025 DOI: 10.29207/resti.v9i1.6056 Accred : Sinta 2

CNN Performance Improvement for Classifying Stunted Facial Images Using Early Stopping Approach

[Ikatan Ahli Informatika Indonesia \(IAII\)](#) [Jurnal RESTI \(Rekayasa Sistem dan Teknologi Informasi\)](#) Vol 9 No 1 (2025): February 2025 62 - 68
📅 2025 📄 DOI: 10.29207/resti.v9i1.6068 🏷️ [Accred : Sinta 2](#)

[Real-Time Location Monitoring and Routine Reminders Based on Internet of Things Integrated with Mobile for Dementia Disorder](#)

[Ikatan Ahli Informatika Indonesia \(IAII\)](#) [Jurnal RESTI \(Rekayasa Sistem dan Teknologi Informasi\)](#) Vol 9 No 1 (2025): February 2025 77 - 84
📅 2025 📄 DOI: 10.29207/resti.v9i1.6105 🏷️ [Accred : Sinta 2](#)

[Prediction of Main Transportation Modes using Passive Mobile Positioning Data \(Passive MPD\)](#)

[Ikatan Ahli Informatika Indonesia \(IAII\)](#) [Jurnal RESTI \(Rekayasa Sistem dan Teknologi Informasi\)](#) Vol 9 No 1 (2025): February 2025 40 - 50
📅 2025 📄 DOI: 10.29207/resti.v9i1.6128 🏷️ [Accred : Sinta 2](#)

[Evaluating Transformer Models for Social Media Text-Based Personality Profiling](#)

[Ikatan Ahli Informatika Indonesia \(IAII\)](#) [Jurnal RESTI \(Rekayasa Sistem dan Teknologi Informasi\)](#) Vol 9 No 1 (2025): February 2025 20 - 30
📅 2025 📄 DOI: 10.29207/resti.v9i1.6157 🏷️ [Accred : Sinta 2](#)

[Real-time Emotion Recognition Using the MobileNetV2 Architecture](#)

[Ikatan Ahli Informatika Indonesia \(IAII\)](#) [Jurnal RESTI \(Rekayasa Sistem dan Teknologi Informasi\)](#) Vol 9 No 4 (2025): August 2025 (in progress) 714 - 720
📅 2025 📄 DOI: 10.29207/resti.v9i4.6158 🏷️ [Accred : Sinta 2](#)

[Face Dermatological Disorder Identification with YoloV5 Algorithm](#)

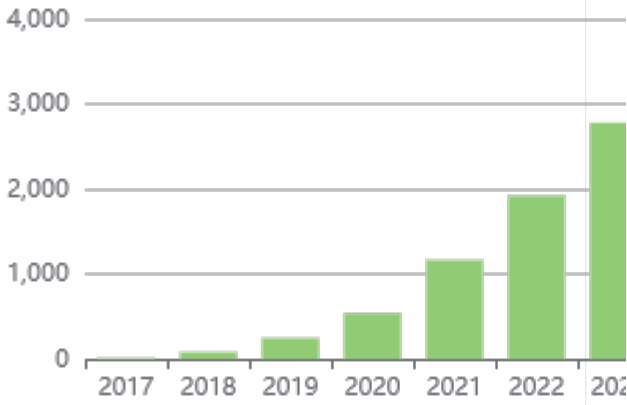
[Ikatan Ahli Informatika Indonesia \(IAII\)](#) [Jurnal RESTI \(Rekayasa Sistem dan Teknologi Informasi\)](#) Vol 9 No 4 (2025): August 2025 (in progress) 721 - 728
📅 2025 📄 DOI: 10.29207/resti.v9i4.6237 🏷️ [Accred : Sinta 2](#)

[View more ...](#)

Get More with
SINTA Insight

[Go to Insight](#)

Citation Per Year By Google Scholar



Journal By Google Scholar

	All	Since 2020
Citation	12995	12557
h-index	45	44
i10-index	360	350