

Sentiment Analysis of Presidential and Vice-Presidential Candidates Using FastText CNN

Fahri Izzuddin Zulkarnaen

Informatika, Ilmu Komputer

Universitas Pembangunan Nasional Jawa Timur

Surabaya, Jawa Timur, Indonesia

Fahrizulkarnaen1@gmail.com

Anggraini Puspita Sari

Informatika, Ilmu Komputer

Universitas Pembangunan Nasional Jawa Timur

Surabaya, Jawa Timur, Indonesia

anggraini.puspita.if@upnjatim.ac.id

The Internet has become a platform for storing and accessing a wide range of public information. People tend to use social media to express opinions on products, political policies, and politicians, both as individuals and as parties. The sentiment analysis of presidential and vice-presidential candidates aims to understand public perspectives on each candidate. Based on sentiment analysis of the Indonesian public on social media, potential candidates for the presidential election, which is still over a year away, can be identified. This study categorizes sentiments into four classes: happy, love, sad, and angry. The model employs FastText embeddings with a Convolutional Neural Network (CNN). The best performance achieved was an F1-score of 0.9510. The findings indicate that Ganjar Pranowo leads as a presidential candidate, while Erick Thohir stands out as the leading vice-presidential candidate.

Keywords—component; Sentiment Analysis, Fast Text, Deep Learning, Election, Convolutional Neural Network, Web Scrapping

I. INTRODUCTION

Indonesia, as one of the world's democratic countries, emphasizes democratic principles, where power lies with the people. This is reflected in Article 1, Paragraph (2) of the 1945 Constitution, which states, "Sovereignty resides in the people and is exercised according to the Constitution." This implies that Indonesia's government is based on democratic principles that place full sovereignty with the people. In a democracy, elections are a tangible manifestation of these principles. Democracy encompasses basic principles that ensure people's sovereignty, freedom of speech, equality in all respects, and justice [1]-[3].

During the campaign process leading up to an election, various public opinions and sentiments about the candidates emerge. These opinions and sentiments can influence voters and even determine the election's outcome. Factors such as political campaigns, media coverage, and socio-cultural influences shape public opinion and sentiment leading up to an election. For presidential candidates, understanding public opinions and sentiments is crucial to gauging their chances of winning.

Analyzing public opinion and sentiment data becomes increasingly important as it provides valuable insights for

political parties to decide on their candidates. Accurate and up-to-date data enables political parties to craft more effective campaign strategies aligned with regional public preferences, thus enhancing their candidates' chances of winning the election. Deep learning algorithms can be leveraged for sentiment analysis, processing text data from social media users and identifying the emotions expressed in sentiments.

Deep learning, a form of artificial intelligence, is based on neural networks. Inspired by the workings of the human brain, neural networks are one of the most complex and advanced forms of intelligence. These networks are supervised learning models trained with labeled data. They have neurons for processing features from the data. Neural networks have become a powerful tool in sentiment analysis [4]-[7].

Like the research conducted in the journal "Sentiment Analysis of Indonesian Reviews Using Fine-tuning IndoBERT and R-CNN", this study performs sentiment analysis on product reviews, which can serve as a reference for potential consumers to determine their preferences in buying, using, or consuming a product [5][8][9].

This study implements FastText embeddings with a Convolutional Neural Network (CNN) approach to analyze sentiment, categorizing sentiments into four classes: happy, love, sad, and angry. All data was sourced from Twitter using the Snsrape library. Data collection covered the period from January 1, 2018, to January 1, 2024. The candidates analyzed were based on electability surveys published by Poltracking Indonesia in November 2022. Sentiment analysis was conducted across 10 regions in Indonesia, encompassing the entire nation.

This model can be used for real-time political sentiment tracking by analyzing large amounts of social media data quickly. It helps monitor public opinion trends, detect sudden negative reactions, and assist politicians in adjusting their campaign strategies based on voter sentiment. By integrating this model, analysts can gain valuable insights into voter behavior throughout an election.

II. THE PROPOSED METHOD/ALGORITHM

A. Web Scraping

Web scraping is a technique used for the automatic extraction of information from websites, eliminating the need for manual searching. The primary objective is to collect specific information and aggregate it into a new website. Web scraping emphasizes data retrieval and extraction, offering benefits such as enabling more efficient searches for targeted information. One of the libraries that can be utilized for data retrieval in Indonesian, based on predefined keywords, is Snsrape[10][11].

B. Text Preprocessing

The data collected from Twitter is still very raw, so a process is needed to clean up some unnecessary words. Text preprocessing is the initial process used to prepare text data into a form that can be further processed [3]. Text preprocessing is carried out to remove unnecessary objects such as hashtags, mentions, and web pages (URLs). In addition to removing unwanted objects, text preprocessing is also used to transform slang or informal words into their correct and standard forms [9].

C. FastText

FastText is a word embedding method and is an enhancement of Word2Vec. In this method, each word is represented as a set of n-gram characters, which allows the embeddings to understand the suffixes and prefixes of a word [6]. The advantage of FastText over Word2Vec is its ability to handle words that have not been encountered before, known as out-of-vocabulary (OOV) words [12].

D. Convolutional Neural Network

The convolutional layer is the main building block of CNNs designed to extract features from the input layer. This layer performs convolution operations and consists of filters that are randomly initialized and learned to obtain feature representations from the input [6]. The convolutional layer uses kernels or filters to extract objects from the input. Convolution results in a linear transformation of the input data that corresponds to the spatial information in the data [11].

Unlike traditional classification methods, CNN-based classification follows an end-to-end learning process. It takes raw data as input, processes training and prediction within the network, and directly outputs the final result. This approach eliminates the need for manual feature extraction, overcoming the limitations of traditional methods and making CNNs highly effective for various classification tasks, including text, speech, and image processing [7].

Based on the research conducted in the journal "A Comprehensive Survey on Sentiment Analysis Techniques", the study states that sentiment analysis models often struggle to interpret the nuances of language used in context, such as sarcasm, irony, and metaphorical expressions [8].

E. Word Embedding

Word embedding is a method for representing a word as a vector, aiming to map semantic meaning into a geometric space. However, word embeddings do not understand text in the way humans do; instead, they map the statistical structure of the language used in the corpus. This geometric space can then be referred to as the embeddings space [13].

F. Convolutional Layer

The convolutional layer is the main block of CNN, designed to extract features from the input layer. This layer performs convolution operations and consists of filters that are learned randomly to obtain feature representations from the input [6]. The convolutional layer uses kernels or filters to extract objects from the input. Convolution results in a linear transformation of the input data that corresponds to the spatial information in the data [11].

G. ReLU Activation

The ReLU activation function has a simple calculation with characteristics that give it advantages in neural networks with many neurons. The forward and backward processes through ReLU only use an if condition, with no exponential, multiplication, or division operations. If an element has a negative value, it is set to 0. This can significantly reduce training and testing time and improve the efficiency and performance of the model overall [14].

H. Pooling Layer

The pooling layer is a layer that serves to reduce the spatial size of the convolutional features by performing dimensionality reduction on the feature map (downsampling). This can reduce the computational resources required to process the data and speed up computation because fewer parameters are updated. This study uses max pooling. Max pooling returns the maximum value from the portion of the image covered by the kernel [6].

I. Dropout Layer

The dropout layer is a layer that can be used to remove/deactivate some of the neurons. These neurons are selected randomly, so they are not used during the training process. This helps to prevent overfitting in the model and can speed up the computation process [15]. The dropout layer also prevents excessive adaptation to the training data, which could result in a model that is too complex and complicated [16].

J. Softmax Activation

Softmax is an activation function that is commonly used. Softmax works to transform the output of the final layer of the neural network into a basic probability distribution. The range of the output produced is from 0 to 1. The sum of each target class will equal 1 [15].

K. Confusion Matrix

A confusion matrix is a table used to calculate accuracy in the concept of data mining. In the confusion matrix, several parameters can be calculated, including true positive, false positive, true negative, and false negative [13]. Next, evaluation metrics can be calculated, including accuracy, precision, recall [17][18], and there is also the F1-score, which combines precision and recall into a single measurement indicator [19][20].

III. METHOD

Figure 1 is the flow used in this study. The following will provide a detailed explanation of each step taken.

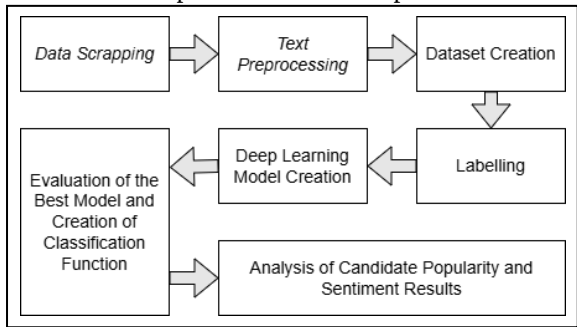


Figure 1. Research Flowchart

A. Data Scrapping

In this study, data scraping was carried out on candidates with the highest electability, as reflected in the release published by Poltracking Indonesia for the November 2022 period. The candidates under scrutiny included three presidential candidates Ganjar Pranowo, Anies Baswedan, and Prabowo Subianto, and two vice-presidential candidates, Erick Thohir and Ridwan Kamil. To ensure a comprehensive and regionally representative analysis, the study was structured around the distribution of active Internet users across various regions.

The allocation of these regions was designed to capture a diverse set of data, reflecting the broader socio-political dynamics that may influence electoral sentiment. By considering the number of active Internet users in each region, the study aimed to target areas that were more likely to generate relevant and representative data for each candidate, thereby ensuring that the findings were both geographically and demographically significant.

This regional division aims to make the research more focused on smaller areas, providing greater benefits compared to conducting research in a single large region. The regions covered include East Java, Central Java, West Java, Jakarta, Yogyakarta, Kalimantan, Sumatra, Bali and Nusa Tenggara, Papua, and Sulawesi.

B. Text Preprocessing

The data from each candidate, collected across different regions, was initially combined to create a unified dataset. Subsequently, extensive text preprocessing was applied to

prepare the data for sentiment analysis. Several preprocessing steps were performed to clean and normalize the text data, including case folding, which standardizes the text to lowercase, removing unusual or non-alphabetic characters to ensure clarity, normalizing slang words into their formal equivalents; applying stemming to reduce words to their base form, removing stop words to eliminate non-informative words, and filtering out sentences that contained fewer than five words, ensuring that only relevant data remained for analysis.

C. Dataset Creation and Labelling

After the data from each candidate has undergone thorough text preprocessing, the next critical step in the workflow is the creation of the dataset. The dataset creation process involves the random selection of 10,000 data points from each candidate’s dataset to ensure a robust and representative sample. Once the data is curated, it is aggregated and subsequently labeled using a pre-trained model specifically designed for sentiment classification. This model was uploaded by a user named *akahana* and is housed within a repository titled ‘Indonesia-emotion-roberta.’

The pre-trained model is leveraged to automatically classify the data into distinct sentiment categories. For this study, the sentiment is classified into four primary labels: happy, love, sad, and angry, each corresponding to a specific emotional sentiment. The data represents the training data used to teach the model, which can be used to analyze the emotional sentiment of each candidate.

D. Deep Learning Model Creation

Once the dataset is thoroughly prepared and cleaned, the next critical step is to construct the model using FastText embeddings integrated with a convolutional neural network (CNN). The initial action in this process involves balancing the dataset, ensuring that each sentiment label (such as happy, sad, love, or angry) has an equal number of data points. This balancing step is crucial as it helps mitigate potential bias towards more frequent labels, enabling the model to learn more effectively from a representative distribution of all sentiment classes. By standardizing the number of samples across the labels, the model is better equipped to generalize across the entire dataset.

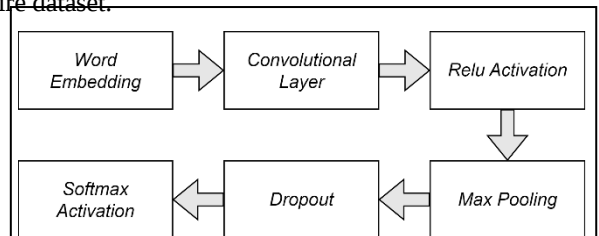


Figure 2. CNN Architecture

The model developed incorporates several key architectures, such as word embedding, ReLU activation, a pooling layer, a dropout layer, and softmax activation. The evaluation of the model is based on optimizing the F1-score, which is derived from various trials with different

hyperparameters. These hyperparameters include learning rates and the proportion for training and testing data. For this study, the learning rate is set to 0.001, while the data split ratios for training and testing vary between 75:25, 80:20, 85:15, and 90:10. Each combination of hyperparameters is tested using the same initialization to ensure consistency in the results.

TABLE I. HYPERPARAMETER INITIALIZATION

No	Hyperparameter	Jumlah
1	Epoch	25
2	Filters	1024
3	Number of Kernels	6
4	Dropout	0.2
5	Optimizer	SGD

This research uses Google Colab to run the training process. Google Colab is essentially a cloud-based computing platform equipped with a GPU: 1x Tesla K80, Compute 3.7, featuring 2,496 CUDA cores, 12GB GDDR5 VRAM, and a CPU: 1x single-core hyper-threaded Xeon processor @ 2.3GHz (1 core, 2 threads).

E. Analysis of Candidate Sentiment Results

In sentiment analysis, the data scraped during the data collection phase is processed to assess the overall emotional tone for each candidate. The sentiment data is classified as *happy*, *love*, *sad*, and *angry*, each assigned specific point values: happy receives 2 points, love gets 1 point, sad is assigned -1 point, and angry is given -2 points.

For instance, consider candidate A, which has 10 sentiment data points distributed as follows: 4 happy sentiments, 3 love sentiments, 2 sad sentiments, and 1 angry sentiment. Based on the assigned points, the total score for each sentiment label is calculated as follows: happy yields 8 points (4×2), love gives 3 points (3×1), sad results in -2 points (2×-1), and angry generates -2 points (1×-2). Thus, the total score for candidate A becomes 7 points ($8 + 3 - 2 - 2$).

To interpret the sentiment, the total score is examined: if the total score is positive, the sentiment toward the candidate is considered positive. Conversely, if the score is negative, the sentiment is considered negative. This methodology allows for an overall evaluation of public sentiment toward each candidate, providing a quantitative measure of their emotional perception based on the data collected from social media.

We chose happy, love, sad, and angry instead of standard labels (positive and negative) because these emotions capture different feelings, allowing for a better understanding of context. This approach makes sentiment analysis more precise and useful in understanding how people truly feel.

IV. RESULTS AND DISCUSSION

This chapter will explain the results obtained from the research conducted and will also present the sentiment analysis of each candidate.

A. Data Scraping and Text Preprocessing

All the data used in this study was obtained from the social media platform Twitter using the Snsrape library, which can be used to retrieve tweets and their metadata. Data collection was conducted by iterating through candidate keywords, followed by data retrieval based on the keywords for each region.

TABLE II. NUMBER OF SENTIMENTS PRODUCED

No	Candidate Name	Number of Sentiment
1	Ganjar Pranowo	81.338 Data
2	Anies Baswedan	98.019 Data
3	Prabowo Subianto	94.748 Data
4	Erick Thohir	61.329 Data
5	Ridwan Kamil	70.149 Data

Next, text preprocessing was performed on all candidate data. The text preprocessing steps included all the processes described in subsection 3.2. The results of the cleaned dataset are presented in Table 3.

TABLE III. NUMBER OF PRE-PROCESSED TWEETS FOR CANDIDATES

No	Candidate Name	Number of Sentiment
1	Ganjar Pranowo	43.739 Tweets
2	Anies Baswedan	53.982 Tweets
3	Prabowo Subianto	52.979 Tweets
4	Erick Thohir	47.364 Tweets
5	Ridwan Kamil	41.381 Tweets

B. Dataset Creation and Labeling

Dataset creation was carried out by randomly selecting 10,000 data points for each candidate's data. Labeling was then performed using a pre-trained model that classifies sentiments into four labels: happy, love, sad, and angry. The details of the number of labels generated in the dataset are presented in Table 4.

TABLE IV. DATASET LABELING RESULTS

No	Candidate Name	Number of Sentiment
1	Bahagia	28.921 Tweets
2	Marah	16.828 Tweets
3	Sedih	3.062 Tweets
4	Takut	596 Tweets
5	Cinta	584 Tweets

C. Deep Learning Model Creation

In the best model tuning, a comparison of the F1-score values obtained from the classification results of the test data will be conducted.

TABLE V. HYPERPARAMETER TUNING RESULTS

No	Learning Rate	Splting Data	F1-Score
1	0.001	75:25	0,9353
2	0.001	80:20	0,9374
3	0.001	85:15	0,9391
4	0.001	90:10	0,9431
5	0.01	75:25	0,9436
6	0.01	80:20	0,9444
7	0.01	85:15	0,9510
8	0.01	90:10	0,9446

The testing results, as shown in Table 5, reveal that the optimal model performance is achieved with a learning rate of 0.01 and a data split of 85:15. This configuration yielded impressive evaluation metrics, including an accuracy of 0.9514, a precision of 0.9512, recall of 0.9514, and an F1-score of 0.9510. These metrics suggest that the model performs exceptionally well in terms of balancing both precision and recall, making it highly effective for sentiment analysis tasks in this study.

D. Analysis of Candidate Sentiment Results

After the best model is obtained, an analysis will be conducted on the number of sentiments produced by each presidential and vice-presidential candidate.

TABLE VI. HYPERPARAMETER TUNING RESULTS

No	Candidate Name	Number of Sentiment
1	Ganjar Pranowo	42.216 Poin
2	Anies Baswedan	-21.754 Poin
3	Prabowo Subianto	17.618 Poin
4	Erick Thohir	55.464 Poin
5	Ridwan Kamil	15.487 Poin

Then, an analysis conclusion will be drawn by considering the sentiment factors produced in each region.

TABLE VII. LEADING CANDIDATES IN REGIONAL SECTIONS

No	Region	Candidate Name
1	Jawa Timur	Ganjar Pranowo
2	Jawa Barat	Ganjar Pranowo
3	Jawa Tengah	Ganjar Pranowo
4	Jakarta	Ganjar Pranowo
5	Jogjakarta	Prabowo Subianto
6	Kalimantan	Ganjar Pranowo
7	Sumatera	Ganjar Pranowo
8	Bali dan Nusa Tenggara	Ganjar Pranowo
9	Papua	Ganjar Pranowo
10	Sulawesi	Ganjar Pranowo

Table 7 shows the potential leading presidential candidates. The results indicate that Ganjar Pranowo leads in 9 regions in

Indonesia, while Prabowo Subianto leads in only 1 region, which is Yogyakarta.

TABLE VIII. LEADING CANDIDATES IN REGIONAL SECTIONS

No	Region	Candidate Name
1	Jawa Timur	Ridwan Kamil
2	Jawa Barat	Ridwan Kamil
3	Jawa Tengah	Erick Thohir
4	Jakarta	Ridwan Kamil
5	Jogjakarta	Erick Thohir
6	Kalimantan	Ridwan Kamil
7	Sumatera	Erick Thohir
8	Bali dan Nusa Tenggara	Erick Thohir
9	Papua	Ridwan Kamil dan Erick Thohir
10	Sulawesi	Erick Thohir

Then, for the vice-presidential candidates, the sentiment analysis shows that Erick Thohir leads in 5 regions, while Ridwan Kamil leads in 4 regions. In 1 region, Papua, both candidates have the same score of 17 points. This can be seen in Table 8.

V. CONCLUSION

Several conclusions can be drawn from the testing results and research conducted on the implementation of FastText embeddings with Convolutional Neural Networks in Sentiment Analysis of Presidential and Vice-Presidential Candidates as follows:

- 1) The FastText embedding algorithm in Convolutional Neural Network was successfully implemented to perform sentiment analysis of Presidential and Vice-Presidential candidates.
- 2) The best performance produced by the system was an accuracy of 0.9514, precision of 0.9512, recall of 0.9514, and an F1-score of 0.9510.
- 3) The sentiment analysis results show that Ganjar Pranowo has the highest electability as a presidential candidate, while Erick Thohir has higher electability as a vice-presidential candidate.

Here are some recommendations that can be made from this study to improve and serve as a lesson for future research:

- 1) A more valid dataset can be used, as the dataset in this study was the result of labeling from an Artificial Intelligence model.
- 2) It is recommended to use a larger sample size for each region to better represent the situation and conditions more relevantly.
- 3) Perform training using a machine with better performance to speed up the training process.

REFERENCES

- [1] Majid, A., & Hajir, M., "Sistem Pemilu Sebagai Wujud Demokrasi Di Indonesia: Antara Orde Lama, Orde Baru Dan Reformasi," in *Qaumiyyah: Jurnal Hukum Tata Negara*, 2(1), 23–43.
- [2] Adam, R., "Word Embedding Bahasa Indonesia Menggunakan FastText (Dengan Gensim)" <https://structilmy.com/blog/2019/04/15/word-embedding-bahasa-indonesia-menggunakan-fasttext-part-1/>
- [3] Aditya, B. R., "Penggunaan Web Crawler Untuk Menghimpun Tweets dengan Metode Pre-Processing Text Mining," in *JURNAL INFOTEL - Informatika Telekomunikasi Elektronika*.
- [4] Saputra, W. S. J., Puspaningrum, E. Y., Syahputra, W. F., Sari, A. P., Via, Y. V., & Idhom, M., "Car Classification Based on Image Using Transfer Learning Convolutional Neural Network," in *Proceeding - IEEE 8th Information Technology International Seminar, ITIS 2022*.
- [5] Herlina, J., Wilis, K., Agung, T., "Sentiment analysis of Indonesian reviews using fine-tuning IndoBERT and R-CNN", in *ILKOM Jurnal Ilmiah*.
- [6] Alwanda, M. R., Ramadhan, R. P. K., & Alamsyah, D., "Implementasi Metode Convolutional Neural Network Menggunakan Arsitektur LeNet-5 untuk Pengenalan Doodle," in *Jurnal Algoritme*.
- [7] Leiyu, C., Shaobo, L., Qiang B., "Review of Image Classification Algorithms Based on Convolutional Neural Networks" in *Remote Sensing*.
- [8] Farhan, A., Sighbat, U. B., Shah, M., "A Comprehensive Survey on Sentiment Analysis Techniques" in *International Journal of Technology (IJTech)*
- [9] Cahyanto, R., Chrisanto, A. R., & Sebastian, D., "Pengelompokan Komentar Dataset Sentipol dengan Modified K-Means Clustering," in *Jurnal Teknik Informatika Dan Sistem Informasi*.
- [10] Putri, N. S., Haerani, E., Fikry, M., & Budianita, E., "Analisis sentimen review aplikasi mypertamina pada twitter menggunakan metode naïve bayes classifier".
- [11] Gultom, R. Y., Zulkarnaen, F. I., Nurhasanah, Y., & Sholahuddin, A., "Indonesian Abusive Tweet Classification based on Convolutional Neural Network and Long Short Term Memory Method," 2021 *International Conference on Artificial Intelligence and Big Data Analytics*.
- [12] Nurdin, A., Aji, B. A. S., Bustamin, A., & Abidin, Z., "Perbandingan Kinerja Word Embedding Word2Vec , Glove," in *Jurnal TEKNOKOMPAK*
- [13] Rachman, F. P., "Perbandingan Model Deep Learning untuk Klasifikasi Sentiment Analysis dengan Teknik Natural Language Processing," in *Jurnal Teknologi Dan Manajemen Informatika*.
- [14] Wibawa, M. S., "Pengaruh Fungsi Aktivasi, Optimisasi dan Jumlah Epoch Terhadap Performa Jaringan Saraf Tiruan," in *Jurnal Sistem Dan Informatika (JSI)*.
- [15] Kholik, A., "Klasifikasi Menggunakan Convolutional Neural Network (Cnn) Pada Tangkapan Layar Halaman Instagram," in *Jurnal Data Mining Dan Sistem Informasi*.
- [16] Septhyan, S., Magdalena, R., & Pratiwi, N. K. C., "Deep Learning Untuk Deteksi Covid-19 , Pneumonia , Dan Tuberculosis Pada Citra Rontgen Dada Menggunakan CNN Dengan Arsitektur Alexnet".
- [17] Pratiwi, B. P., Handayani, A. S., & Sarjana, S., "Pengukuran Kinerja Sistem Kualitas Udara Dengan Teknologi Wsn Menggunakan Confusion Matrix," in *Jurnal Informatika Upgris*.
- [18] Sari, A. P., Suzuki, H., Kitajima, T., Yasuno, T., Prasetya, D. A., & Nachrowie, N., "Prediction model of wind speed and direction using convolutional neural network - Long short term memory," in *International Conference on Power and Energy*.
- [19] Sundar, C., M.Chitradevi, M. C., & Geetharamani, G., "Classification of Cardiotocogram Data using Neural Network based Machine Learning Technique," in *International Journal of Computer Applications*.
- [20] Ilahiyah, S., & Nilogiri, A., "Implementasi Deep Learning Pada Identifikasi Jenis Tumbuhan Berdasarkan Citra Daun Menggunakan Convolutional Neural Network," in *JUSTINDO (Jurnal Sistem Dan Teknologi Informasi Indonesia)*.