

Depression Detection of Users in Social-Media Twitter Using Decision Tree with Word2Vec

Elroi Yoshua¹, Warih Maharani^{*})

^{1,2} School of Computing, Telkom University, Bandung, Indonesia

¹elroi Yoshua@student.telkomuniversity.ac.id (*)

²wmaharani@telkomuniversity.ac.id

Received: 2023-12-28; Accepted: 2024-01-31; Published: 2024-02-14

Abstract— Social media is a medium or place on the Internet that allows users to be themselves. Interact, cooperate, share, and communicate with other users virtually. Not only do users share happy feelings, but they also share their emotions and sentiments towards a particular issue. Which sometimes makes users look depressed when they deliver it. Depression itself is the most commonly encountered mental illness, which makes the sufferer feel sad, lonely, inferior, and disconnected from the people around them. And even worse, depression can make the sufferer have suicidal thoughts. Therefore, we need to know whether the user indicated being depressed or not to prevent unwanted things by using a depression measurement tool scale called DASS 42 for data labeling. To detect depression, we can use the sufferer's Twitter account to take data based on tweets from the user and change the entire dataset to a vector using both the architectures of Word 2, Vec Skip-Gram, and CBOW. In this research, we utilize a decision tree to detect depression. The best results were obtained from the Word2Vec Skip-Gram model with a data ratio of 90:10 using the Gini criterion parameter and a maximum depth value of 20, resulting in an accuracy of 93% and a f1-score of 94%.

Keywords—Social Media; Depression Detection; DASS-42; Word2Vec; Decision Tree.

I. INTRODUCTION

As a modern means of communication, the Internet has made the world easier to understand [1]. For example, there is social media. Social media is a medium that can be accessed online where users can easily participate, share, and create content such as blogs, social networks, wikis, forums, and virtual worlds [2]. Social media also can help one user form virtual social bonds with other users [1]. As reported by *katadata.co.id*, the research results in January 2022 showed that social media users in Indonesia reached 204.7 million or 73.3% of the total population [3]. This leads to high social media engagement among the Indonesian population. The majority of social media users in Indonesia itself are teenagers with an age range of 18-24 [4]. The users actively use social media to get attention, ask for opinions, and build their image [5]. Social media is no longer just a communication tool for sending messages; it has even been developed to form social networks and groups or community groups [6]. In this place, we also don't need to be ashamed to share our thoughts about something as long as it is within reasonable limits and follows applicable ethics and norms. But sometimes, some misuse social media as a place for hate speech or cyberbullying.

If a person feels hopeless and unable to cope with the problems they face, they are likely to experience depression or high levels of stress [7]. Depression is the most common mental illness [8]. Depression itself is a mood disorder characterized by hopelessness and heartbreak, excessive weakness, inability to make decisions to start an activity, inability to concentrate, lack of enthusiasm for life, always being tense, and attempting suicide [9]. The person can behave

very differently than usual and sometimes can cause harm to others or themselves. Globally, as many as 5% of adults experience depression, with symptoms being persistent sadness and reduced interest or pleasure in activities that were previously rewarding or enjoyable. It also can affect sleep and appetite [8]. And even worse, according to WHO data, depression is one of the leading causes of suicide. As many as 40% of people with depression have suicidal thoughts [7]. In Indonesia, according to Fachmi Idris, President of the Indonesian Medical Association (IDI), in 2007, 94% of the Indonesian population suffered from depression from the highest to the lowest level. According to WHO, the suicide rate in Indonesia continues to increase. In 2010, the suicide rate in Indonesia was 1.8 people per 100,000 people or 5,000 cases per year [10]. Therefore, depression needs to be identified from social media accounts based on what the account posts to prevent unwanted things from happening. Suppose the depression is detected whether someone is showing symptoms of depression through their social media. Further care can be provided professionally and with moral support from friends and family before further treatment is carried out.

A comprehensive analysis of depression utilizing various Machine Learning methods in his research [11]. The employed methods yielded distinct accuracies: Decision Tree achieved 72%, KNN achieved 60.5%, SVM achieved 71%, and Ensemble achieved 64.25%. The classified Twitter data containing depression-related terms using the Naive Bayes method, achieving an accuracy of 70% [12]. A different algorithm using the KNN algorithm for sentiment analysis, can detect depression, resulting in an accuracy of 78.18% [13]. An accuracy of 79% using the Bi-LSTM method with

word2vec, 77% using Bi-LSTM with FastText, 52% using LSTM with Glove Embedding, and 52% using LSTM with Trainable Embedding Layer [28]. Word2Vec has the highest accuracy compared to other text embedding algorithms in his research [28]. The Decision Tree algorithm, achieving an accuracy of 81.25% [14]. This exceptional performance was attained through parameter tuning, specifically by increasing the maximum depth. The Decision Tree algorithm demonstrated the highest accuracy among the evaluated methods. Based on previous research, it can be concluded that Machine Learning algorithms can detect depression.

This research aims to detect depression on Twitter social media and use a Decision Tree, a classification method using the Word2Vec feature to convert strings into vectors. Word2Vec is one of the most effective techniques for generating word embeddings [27]. A comparison of the accuracy and F1-scores produced by the Word2Vec Skip-Gram and CBOW models will also be carried out as part of this research. Instead of generating vectors based on the provided word, CBOW will construct vectors based on the context [22]. It follows that every Word2Vec model will produce a unique set of outcomes. There will be a comparison of the Word2Vec outcomes to identify which Word2Vec model is superior. We will also run some scenarios to determine how effective the Decision Tree method is in identifying depression on Twitter and other social media platforms.

II. METHOD

There are several processes before the system can detect depression using the decision tree method. Based on Figure 1, the process is data collection, preprocessing, splitting or dividing data into test and training data, feature extraction using Word2Vec, model building, and evaluation.

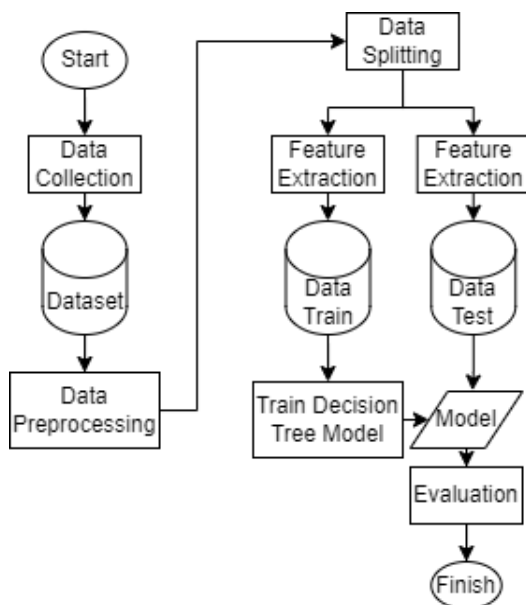


Figure 1. System Flow

A. Data Collection

There are several stages in the data collection process. First, respondents fill out the DASS 42 form containing questions about depression, stress, and anxiety. Then, crawling tweet data on Twitter social media using the respondent's account who has agreed to fill out the DASS 42 form. After that, merge the collected data into a CSV file consisting of usernames, tweets, and depression labels using DASS 42 to define the depression. DASS 42 itself is an acronym for Depression Anxiety Stress Scale, and it contains 42 questions about Depression Anxiety Stress. It is one of the psychological measurement tools that can measure the severity of the three disorders [15]. DASS Depression focuses on reports of low mood, motivation, and self-esteem. DASS Anxiety focuses on physiological arousal, perceived panic, and fear. And DASS Stress focuses on tension and irritability [16]. For a standard normal person, it will be 0-9 for the level of depression, 0-7 for the level of anxiety, and 0-14 for the level of stress [17].

TABLE I
DASS 42 LEVEL

Level	Depression	Anxiety	Stress
Normal	0-9	0-7	0-14
Mild	10-13	8-9	15-18
Moderate	14-20	10-14	15-18
Extremely Severe	28-42	20-42	34-42

The primary purpose of measuring with the DASS 42 is to assess the severity (severe level) of core symptoms of depression, anxiety, and stress. Of the 42 items, 14 items relate to depressive symptoms, 14 items relate to anxiety symptoms, and 14 items relate to stress [18]. Symptoms can be seen in Table I. The DASS 42 questionnaire is a standardized measurement tool [19]. There is no need to doubt its credibility and correctness in detecting depression, anxiety, and stress. This research will be focused on depression levels and will be using two levels only, which are not depression level, which has a score of 0-9, and depression level, having a score from 10-42.

B. Data Preprocessing

After data collection, the subsequent step involves preprocessing the data to obtain clean and suitable data for the depression detection algorithm. Several steps are undertaken during data preprocessing. The first step is case folding, where all capital letters in the dataset are converted to lowercase letters. The next step involves stop word removal, eliminating non-essential words that provide no informational value. The third step is replacing slang words. Indonesian slang words are identified in this phase, and corresponding basic words are substituted. Following that, elongation removal is performed, which entails eliminating repeated words. The fifth step involves converting emojis and emoticons into words. All emojis and emoticons are replaced with corresponding Indonesian words to ensure the system understands their meaning. The sixth step is cleansing, involving data cleaning

by removing symbols, numbers, punctuation marks, and characters not in the alphabet. After cleansing, lemma is applied, converting words into their base form or root word.

TABLE II
DATA PREPROCESSING PROCESS

Process	Tweet
Raw Tweet	@mentioneduser_ tidak berasa udah Desember lagi. Sudah tiga tahun semenjak kita berpisahhhh :(
Stopword Removal	@mentioneduser_ berasa udah desember lagi. semenjak berpisahhhh :(
Replace Slang Words	@mentioneduser_ berasa sudah desember lagi. semenjak berpisahhhh :(
Remove Elongation	@mentioneduser_ berasa sudah desember lagi. semenjak berpisah :(
Emoji and Emoticon to Word	@mentioneduser_ berasa sudah desember lagi. semenjak berpisah marah_sedih_atau_cemberut
Cleansing	berasa sudah desember lagi semenjak berpisah marah sedih atau cemberut
Lemma	rasa sudah desember lagi semenjak pisah marah sedih atau cemberut
Tokenization	"rasa", "sudah", "desember", "lagi", "semenjak", "pisah", "marah", "sedih", "atau", "cemberut"
Typo Checker	"rasa", "sudah", "desember", "lagi", "semenjak", "pisah", "marah", "sedih", "atau", "cemberut"

The eighth step is tokenization, breaking tweet text or sentences into smaller, discrete units known as "tokens". The final step is typo checking since the words will be converted into vectors using word2vec. Ensuring that all tweets are present in the Indonesian dictionary is crucial. The Wikipedia corpus is used for verification, and any detected typos are promptly corrected to the accurate Indonesian spelling. Examples of the preprocessing steps can be found in Table II.

C. Data Splitting

After the preprocessing step, the dataset acquired is cleaned up, and it is much simpler to implement depression detection. The information will then be separated into two distinct categories: the training data and the testing data. We will use the scikit-learn library to partition the data [26]. There will be a variation in the proportion of data utilized; the default ratio will be 70:30, followed by 80:20, and finally, the 90:10 ratio will be used for the data ratio combination in the third scenario for the experiment. We will randomly collect data to ensure that we have an accurate picture.

D. Feature Extraction

Following the dataset-splitting process, the next imperative step involves converting all data into vectors, as machine learning algorithms mandate numeric vectors as input [20]. Subsequently, the preprocessed data will transform numeric vectors to facilitate streamlined data processing through Decision Tree models. The Word2Vec algorithm will translate the prepared data into numeric vectors in this step. After performing, Word2Vec will provide semantic meaning for each word in vector representation [21].

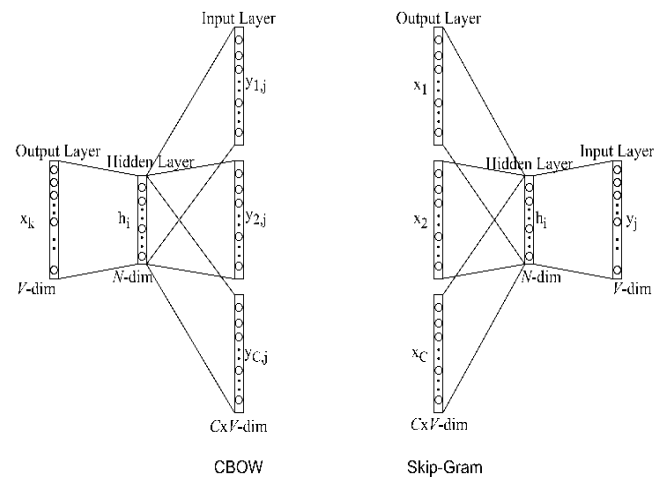


Figure 2. CBOW & Skip-Gram Architecture

Word2Vec comprises two architectures: Continuous Bag-of-Word (CBOW) and Skip-Gram. The CBOW architecture predicts current words based on context, while the Skip-Gram architecture predicts words around a given word [22]. The distinction between CBOW and Skip-Gram architectures is illustrated in Figure 2 [21].

E. Decision Tree Models

Data processing will be executed during this step, involving training models and making predictions. The chosen method for modeling is the Decision Tree, a widely employed classification technique across various domains, including machine learning, image processing, and pattern recognition [19].

The Decision Tree operates as a supervised learning algorithm, necessitating a labeled dataset for model creation. Within the Decision Tree structure, nodes and branches constitute integral components. Each node represents a feature within the target classification, and each subset delineates the possible values associated with that node [20]. The classification procedure commences at the root of the Decision Tree and proceeds recursively until a branch with a specific label is reached. A split condition is applied at each node to determine whether the input value progresses to the right or left portion until it reaches a leaf node [21]. For model training, the process utilizes training data to train the Decision Tree model. Subsequently, test data is used for predictions. The

resulting output manifests as classification outcomes, indicating whether a tweet indicates depression.

F. Model Evaluation

After successfully building the model, an evaluation is carried out to assess the performance of the classification model that has been created. This evaluation tests the chosen method's accuracy and determines whether the model has worked well. If not, improvements will be made. This process uses a confusion matrix with four common measures: accuracy, recall, precision, and F1-score. Several classification performance measures can be defined based on the confusion matrix. Some common measures are given as follows [25].

The percentage of the total number of correct forecasts to the total number of wrong guesses is the definition of accuracy. Using Equation (1), one may determine the accuracy value of the calculation. A certain class's precision can be defined as the degree of accuracy with which it has been anticipated. One method for determining the accuracy value is to use Equation (2). Recall measures the prediction model's ability to select the total facts from a particular class using Equation (3). When it comes to Equation (4), the F1-score is a measurement that takes into account both precision and recall.

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

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

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

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

III. RESULT & DISCUSSION

In this research, 184 users with 14,583 tweet data were used. All data are labeled per individual according to the responses provided in the DASS 42 questionnaire, with a label of 1 for depressed users and 0 for non-depressed users. The total depression in the dataset is 107, and non-depression is 77. Table III shows some of the data from the data collection process.

TABLE III
DASS-42 DATA

User	Score	Label
User 1	9	0
User 2	6	0
User 3	18	1

The information that was gathered during the data-gathering process is presented in Table III. Following splitting the dataset, the feature extraction was carried out with the assistance of Word2Vec by comparing the two architectures, CBOW and

Skip-Gram. The data from the split ratio 70:30 serves as the default ratio for the split ratio. Through the execution of this stage before the execution of decision tree modeling, it is possible to evaluate the model's performance.

TABLE IV
PREPROCESSING RESULT

Raw Tweet	Preprocessed Tweet
akhirnya truknya berenti, mungkin	"truk", "henti", "sadar",
sadar kalinya ada yg nyangkut, trus	"kali", "nyangkut",
akhirnya gw tengkurep tuh ampe pipi	"terus", "gue",
nyentuh aspal. trus truknya lanjut	"tengkurep", "tua",
jalan. dan kita berhasil nyebrang deh	"sampai", "truk", "jajan",
:)	"hasil", "nerang",
	"senyum", "Bahagia"

In this research, three scenarios will be carried out. The first scenario compares the Word2Vec CBOW model and also Skip-Gram. In the second scenario, the result was compared after tuning the parameters by changing the criterion and max of depth. After that, the third scenario will use a different split data ratio for each Word2Vec model. The results of each scenario will be examined and compared. The first scenario is tried with the initial data and has an accuracy like Table V. Data needs to be balanced to get the optimal result. For balancing the dataset, the depressed user will be decreased from 107 to 77, equivalent to a non-depressed user, by randomly selecting user data that shows depression to be decreased. As a result, the total number of users is currently 154. Upon balancing the dataset, it is evident that the accuracy of the results has improved. When compared to Skip-Gram, CBOW Models offer a better level of accuracy. Further scenarios will take advantage of the perfectly balanced dataset.

TABLE V
RESULT AFTER DATA BALANCING

Model	Accuracy Before Balance	Accuracy After Balance
Skip-Gram	52%	74%
CBOW	45%	76%

The second scenario is to compare both Word2Vec models after tuning the parameter using the Gini and entropy criterion, using a maximum depth of 10,20,30 to see which one has better accuracy from Table VI. However, the Skip-Gram in the previous scenario had a smaller accuracy than CBOW. After parameter tuning, the results are different. The skip-gram model with the entropy criterion and 20 max depths has the highest accuracy. The entropy criterion also gives greater accuracy in both Word2Vec models than the Gini criterion. Meanwhile, the optimal max depth is 20 compared to others.

TABLE VI
COMPARISON AFTER TUNING PARAMETER

Word2Vec Model	Criterion	Max of Depth	Accuracy
Skip-Gram	Gini	10	74
	Gini	20	76
	Gini	30	70
	Entropy	10	82

Word2Vec Model	Criterion	Max of Depth	Accuracy
CBOW	Entropy	20	85
	Entropy	30	82
	Gini	10	72
	Gini	20	72
	Gini	30	74
	Entropy	10	80
	Entropy	20	74
	Entropy	30	78

The last scenario is to compare CBOW and Skip-Gram using three different data ratios. 70:30 ratio, 80:20 ratio, and 90:10 ratio will be used, and parameter tuning will be executed for the optimal result. The last scenario will use Accuracy and F1-Score to compare the results. In Table VII, the highest accuracy in the Skip-Gram model is 93% with 94% F1-Score. It uses the Gini criterion with a max depth value of 20.

TABLE VII
SKIP-GRAM COMPARISON WITH DIFFERENT DATA RATIO

Ratio	Criterion	Max of Depth	Accuracy	F1-Score
70:30	Entropy	10	82	83
	Entropy	20	85	85
	Entropy	30	82	83
80:20	Entropy	10	61	61
	Entropy	20	64	64
	Entropy	30	61	61
90:10	Entropy	10	75	75
	Entropy	20	75	75
	Entropy	30	75	75
70:30	Gini	10	74	74
	Gini	20	76	77
	Gini	30	70	70
80:20	Gini	10	77	77
	Gini	20	70	71
	Gini	30	77	77
90:10	Gini	10	87	88
	Gini	20	93	94
	Gini	30	87	88

While in Table VIII, the CBOW model resulted in an accuracy of 80% with an F1-score of 81%. It used a data ratio of 70:30 and parameter Entropy criterion with 10 max depth value. Based on the results of the last scenario, it was found that the Skip-Gram model with a data ratio of 90:10 obtained the highest results with an accuracy of 94% and an F1-score of 95%.

TABLE VIII
CBOW COMPARISON WITH DIFFERENT DATA RATIO

Ratio	Criterion	Max of Depth	Accuracy	F1-Score
70:30	Entropy	10	80	81
	Entropy	20	74	74
	Entropy	30	78	79
80:20	Entropy	10	77	77
	Entropy	20	70	71

Ratio	Criterion	Max of Depth	Accuracy	F1-Score
90:10	Entropy	30	70	71
	Entropy	10	68	69
	Entropy	20	68	69
	Entropy	30	68	69
	Gini	10	72	71
70:30	Gini	20	72	72
	Gini	30	74	74
	Gini	10	74	74
80:20	Gini	20	74	74
	Gini	30	77	77
	Gini	10	68	69
90:10	Gini	20	62	62
	Gini	30	62	62

All scenarios were conducted to determine the optimal combination for Skip-Gram and CBOW. The first scenario used a data ratio of 70:30, default parameters of the Gini criterion, and 0 max depth for the decision tree algorithm. The first scenario resulted in a higher accuracy of CBOW. Although only slightly different, it can be concluded that the results issued by the CBOW feature and the Skip-Gram feature are different, and in this first scenario, CBOW is better; for the second scenario, parameter tuning is executed to find out whether the optimal results by changing the parameter values, max of depth, and criterion. And get the result that the Skip-Gram model gets a higher accuracy. The results of Skip-Gram itself are obtained from the Entropy criterion parameter with a maximum depth value of 20. It can be concluded that parameter tuning can produce more optimal accuracy. Because each dataset requires different parameters, the parameter must be set to the most appropriate combination for the dataset used. For the last scenario, data splitting is done with different ratios, namely 70:30, 80:20, and 90:10. Based on the previous scenario, parameter tuning will be used to get more optimal results with higher accuracy and F1-score. This third scenario resulted in Skip-Gram being more optimal. That can occur because the higher the ratio of data that is trained, the higher the possibility of data that can be predicted. Based on all the scenarios conducted, it can be concluded that the Skip-Gram model has better accuracy than CBOW.

IV. CONCLUSION

This research aims to utilize the Decision Tree machine learning algorithm to detect depression among Twitter users. Data collection, preprocessing, data splitting, and feature extraction using Word2Vec Skip-Gram and CBOW are some of the important measures taken to improve the algorithm's effectiveness in detecting depression in the most effective manner possible. After these preparatory steps, scenarios were conducted to achieve optimal results. These scenarios involved parameter tuning and the application of various data ratio variations. These efforts culminated in developing a system model proficiently detecting depression on Twitter

social media. The results can be optimized by tuning the parameters and using different data ratios. Identifying a combination of parameters and data ratios that best align with the dataset's characteristics is crucial to get the optimal result.

REFERENCES

- [1] I. Santoso, S. Wahyudi, A. Putra, and A. Purwana, "Komunikasi Dan Mitigasi Berita Bohong Di Media Sosial Oleh Pemerintah," *Gagasan Komun. Untuk Negeri*, p. 184.
- [2] A. S. Cahyono, "Pengaruh media sosial terhadap perubahan sosial masyarakat di Indonesia," *Publiciana*, vol. 9, no. 1, pp. 140–157, 2016.
- [3] A. Setiadi, "Pemanfaatan media sosial untuk efektifitas komunikasi," *Cakrawala J. Hum. Bina Sarana Inform.*, vol. 16, no. 2, 2016.
- [4] Social Media in Indonesia, "Stats & Platform Trends | OOSGA," 2023. [Online]. Available: <https://oosga.com/social-media/idn/>
- [5] P. Felita, C. Siahaja, V. Wijaya, G. Melisa, M. Chandra, and R. Dahesihsari, "Pemakaian media sosial dan self concept pada remaja," *Manasa*, vol. 5, no. 1, pp. 30–41, 2016.
- [6] A. Rosmalina and T. Khaerunnisa, "Penggunaan Media Sosial dalam Kesehatan Mental Remaja," *Prophet. Prof. Empathy, Islam. Couns. J.*, vol. 4, no. 1, pp. 49–58, 2021.
- [7] A. Dirgayunita, "Depresi: Ciri, penyebab dan penanganannya," *J. An-Nafs Kaji. Penelit. Psikol.*, vol. 1, no. 1, pp. 1–14, 2016.
- [8] World Health Organization, "Depresion WHO." [Online]. Available: https://www.who.int/health-topics/depression#tab=tab_
- [9] R. L. Atkinson, R. C. Atkinson, and E. R. Hilgard, "Pengantar Psikologi Edisi 8," *Erlangga, Jakarta*, 1991.
- [10] W. Sulistyorini and M. Sabarisman, "Depresi: Suatu tinjauan psikologis," *Sosio Inf. Kaji. Permasalahan Sos. dan Usaha Kesejaht. Sos.*, vol. 3, no. 2, 2017.
- [11] M. R. Islam, M. A. Kabir, A. Ahmed, A. R. M. Kamal, H. Wang, and A. Ulhaq, "Depression detection from social network data using machine learning techniques," *Heal. Inf. Sci. Syst.*, vol. 6, no. 1, pp. 1–12, 2018, doi: 10.1007/s13755-018-0046-0.
- [12] N. Hayatin, "Implementasi Multinomial Naïve Bayes Untuk Klasifikasi Data Tweets Mengandung Term Depresi," in *Prosiding SENTRA (Seminar Teknologi dan Rekayasa)*, 2021, pp. 344–349.
- [13] A. P. Tirtopangarsa and W. Maharani, "Sentiment Analysis of Depression Detection on Twitter Social Media Users Using the K-Nearest Neighbor Method," in *Seminar Nasional Informatika (SEMNASIF)*, 2021, pp. 247–258.
- [14] M. R. Hidayatullah and W. Maharani, "Depression Detection on Twitter Social Media Using Decision Tree," *J. RESTI (Rekayasa Sist. dan Teknol. Informasi)*, vol. 6, no. 4, pp. 677–683, 2022.
- [15] N. Shayan, A.-R. Niazi, A. Waseq, and L. ÖZCEBE, "Depression, anxiety, and stress scales 42 (DASS-42) in Dari-language: validity and reliability study in adults, Herat, Afghanistan," *Bezmialem Sci.*, vol. 9, no. 3, 2021.
- [16] L. Parkitny and J. McAuley, "The depression anxiety stress scale (DASS)," *J. Physiother.*, vol. 56, no. 3, p. 204, 2010.
- [17] N. G. Bilgel and N. Bayram, "Turkish version of the depression anxiety stress scale (DASS-42): Psychometric properties," 2010.
- [18] S. Kusumadewi and H. Wahyuningsih, "Model Sistem Pendukung Keputusan Kelompok untuk Penilaian Gangguan Depresii, Kecemasan dan Stress Berdasarkan DASS-42," *J. Teknol. Inf. Dan Ilmu Komput.*, vol. 7, no. 2, pp. 219–228, 2020.
- [19] S. R. Marsidi, "Identification of Stress, Anxiety, and Depression Levels of Students in Preparation for the Exit Exam Competency Test," 2021.
- [20] F. K. Khattak, S. Jeblee, C. Pou-Prom, M. Abdalla, C. Meaney, and F. Rudzicz, "A survey of word embeddings for clinical text," *J. Biomed. Informatics X*, vol. 4, no. October, p. 100057, 2019, doi: 10.1016/j.yjbinx.2019.100057.
- [21] S. Sivakumar, L. S. Videla, T. Rajesh Kumar, J. Nagaraj, S. Itnal, and D. Haritha, "Review on Word2Vec Word Embedding Neural Net," *Proc. - Int. Conf. Smart Electron. Commun. ICOSSEC 2020*, no. Icossec, pp. 282–290, 2020, doi: 10.1109/ICOSSEC49089.2020.9215319.
- [22] D. Jatnika, M. A. Bijaksana, and A. A. Suryani, "Word2vec model analysis for semantic similarities in english words," *Procedia Comput. Sci.*, vol. 157, pp. 160–167, 2019.
- [23] G. Stein, B. Chen, A. S. Wu, and K. A. Hua, "Decision tree classifier for network intrusion detection with GA-based feature selection," in *Proceedings of the 43rd annual Southeast regional conference-Volume 2*, 2005, pp. 136–141.
- [24] B. Charbuty and A. Abdulazeez, "Classification based on decision tree algorithm for machine learning," *J. Appl. Sci. Technol. Trends*, vol. 2, no. 01, pp. 20–28, 2021.
- [25] X. Deng, Q. Liu, Y. Deng, and S. Mahadevan, "An improved method to construct basic probability assignment based on the confusion matrix for classification problem," *Inf. Sci. (Ny)*, vol. 340, pp. 250–261, 2016.
- [26] Scikit-Learn Library "Machine Learning in Python" [Online]. Available: <https://scikit-learn.org/sTable/>
- [27] Renaldi, Aldy, and Warih Maharani, "Depression Detection of User in Media Social Twitter Using Random Forest." *Journal of Information System Research (JOSH)* 3.4 (2022): 410-416.
- [28] Shah, Faisal Muhammad, et al. "Early depression detection from social network using deep learning techniques." *2020 IEEE Region 10 Symposium (TENSYP)*. IEEE, 2020.

This is an open access article under the [CC-BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license.

