



Classification of Forest Fire-Prone Areas Using the K-Nearest Neighbor Algorithm: A Case Study of Baluran National Park

M Noer Fadli Hidayat ¹

¹ Universitas Negeri Malang, Indonesia

m.noer.2305349@students.um.ac.id ¹

Doi

Received: 01 juni 2025

Accepted: 16 juni 2025

Published: 30 juni 2025

Abstract

Forest and land fires are one of the most recurring and destructive natural disasters in Indonesia, particularly during the dry season when rainfall is significantly low. Among the most affected areas in East Java is Baluran National Park, a region highly vulnerable due to its dominant savanna ecosystem. In response to the urgent need for effective forest fire risk prediction, this study aims to classify fire-prone areas using weather-related data and machine learning techniques. The research focuses on the application of the K-Nearest Neighbor (K-NN) algorithm to predict fire risk levels based on meteorological parameters. The dataset used in this study was obtained from Visualcrossing.com and consists of 211 weather records with 32 explanatory variables, such as maximum and minimum temperature, wind speed, sea-level pressure, solar radiation, and solar energy, along with one target variable representing fire risk level (categorized as High, Medium, or Low). The research method involves several stages: data preprocessing (handling missing values and converting nominal data into numeric), transformation (splitting into training and testing sets using the Percentage Split technique), and classification using K-NN implemented on Google Colab. The K-NN algorithm was configured with K = 3, using Euclidean Distance as the distance metric. The classification process produced a high accuracy rate of 98%, indicating the robustness and effectiveness of K-NN in classifying forest fire risks based on weather data. This model was further validated by comparing predicted outputs against actual values in the testing dataset, showing high consistency. The results suggest that the K-NN algorithm is highly applicable for environmental classification problems and can support decision-making systems in early warning and disaster mitigation efforts. This study contributes to the growing field of data-driven disaster risk management and highlights the

potential of machine learning in enhancing environmental monitoring systems.

Keywords *Forest Fires; Baluran; K-Nearest Neighbor; Classification; East Java*

1. INTRODUCTION

Forest and land fires are among the most severe and recurring natural disasters in Indonesia. This phenomenon occurs across almost all regions of the country and typically intensifies during the dry season due to extremely low rainfall. Such disasters result in substantial losses, not only reducing forest and land cover and threatening biodiversity, but also endangering the health and safety of nearby communities through exposure to dense smoke, which can lead to respiratory problems and even life-threatening conditions (Mahdi, 2022; Pratiwi et al., 2021). According to the Ministry of Environment and Forestry (KLHK), the total area of forest and land fires in Indonesia in 2021 reached 358,867 hectares, an increase of 20.85% compared to the 296,942 hectares recorded in 2020. East Java contributed the fifth largest burned area at 15,458 hectares, following Papua and three other provinces. The trend of forest fires in Indonesia fluctuates from year to year, with 2019 being the worst year in the 2016–2021 period, recording 1,649,258 hectares burned—an increase of 311% compared to the previous year (Mahdi, 2022; Primajaya, Sari, & Khusaeri, 2020).

One of the regions in East Java frequently affected by forest fires is Baluran National Park, located in Banyuputih, Situbondo Regency. Often dubbed the "African Savanna of Java," this park consists predominantly of savanna, in addition to forest, mountains, and coastal areas. These characteristics, along with dry vegetation and climate variability, make it particularly susceptible to fires. In 2020, Baluran experienced two separate fire incidents that destroyed approximately 5.2 hectares of forest area (Detiknews, as cited in Mahdi, 2022).

Weather conditions play a significant role in influencing the likelihood of forest fires. Critical variables include maximum and minimum temperature, wind speed, surface pressure, solar radiation, and solar energy (Noviansyah, Rismawan, & Midyanti, 2018). Identifying these environmental factors and their patterns can help in predicting and minimizing fire risks in vulnerable areas. One effective way to achieve this is through the application of data mining, a process used to convert large volumes of raw data into meaningful insights (Twin, 2021). This is especially useful when analyzing data from Automatic Weather Stations (AWS), which provide continuous environmental monitoring that can be processed to identify potential fire zones (Noviansyah et al., 2018).

A commonly used method in data mining for classification tasks is the K-Nearest Neighbor (K-NN) algorithm. K-NN is a supervised learning approach where new data instances are classified based on the majority class among their K nearest neighbors (Rahayu, Prianto, & Novia, 2021). This technique is popular due to its simplicity and high accuracy in real-world applications, particularly when the relationship between variables is non-linear or complex. Studies have shown the effectiveness of K-NN in classifying forest fire risk zones based on environmental variables. For instance, in Kubu Raya and West Kalimantan, researchers successfully implemented K-NN to classify fire risk levels using meteorological data, achieving significant accuracy (Rudiyana, Dzulkifli, & Munazar, 2022; Karo et al., 2022).

Comparative studies have also demonstrated that K-NN performs favorably when compared with other algorithms such as Naïve Bayes, especially in environmental and disaster-related prediction tasks (Pratiwi et al., 2021; Rahayu et al., 2021).

Furthermore, other research has highlighted how classification algorithms like C4.5, K-Means, and Naïve Bayes can be utilized to improve forecasting and prioritization efforts in various domains, including health services and forest management (Primajaya et al., 2020; W. I. Rahayu et al., 2021). The growing integration of machine learning techniques in environmental monitoring reflects a broader trend toward the digitization of risk management and early warning systems (A. Twin, 2021). In line with this trend, this study aims to classify fire-prone areas based on weather data—such as temperature, wind speed, and solar radiation—using the K-Nearest Neighbor algorithm. This research contributes to disaster risk mitigation efforts, particularly in fire-vulnerable regions like East Java, by offering a predictive model that can assist stakeholders in early intervention and resource allocation.

2. RESEARCH METHODS

In general, this research is divided into two scopes: literature study and model development. The literature study aims to strengthen the research problem and serve as the theoretical foundation of the study. The scope of the literature review includes journals, articles, and books that discuss the application of the K-Nearest Neighbor algorithm, the influence of distance measurement methods in K-Nearest Neighbor, and classification models for forest and land fires. The method consists of several stages, including selection, preprocessing, transformation, data mining, and evaluation. These stages are part of the Knowledge Discovery in Databases (KDD) process, as illustrated in Figure 1.

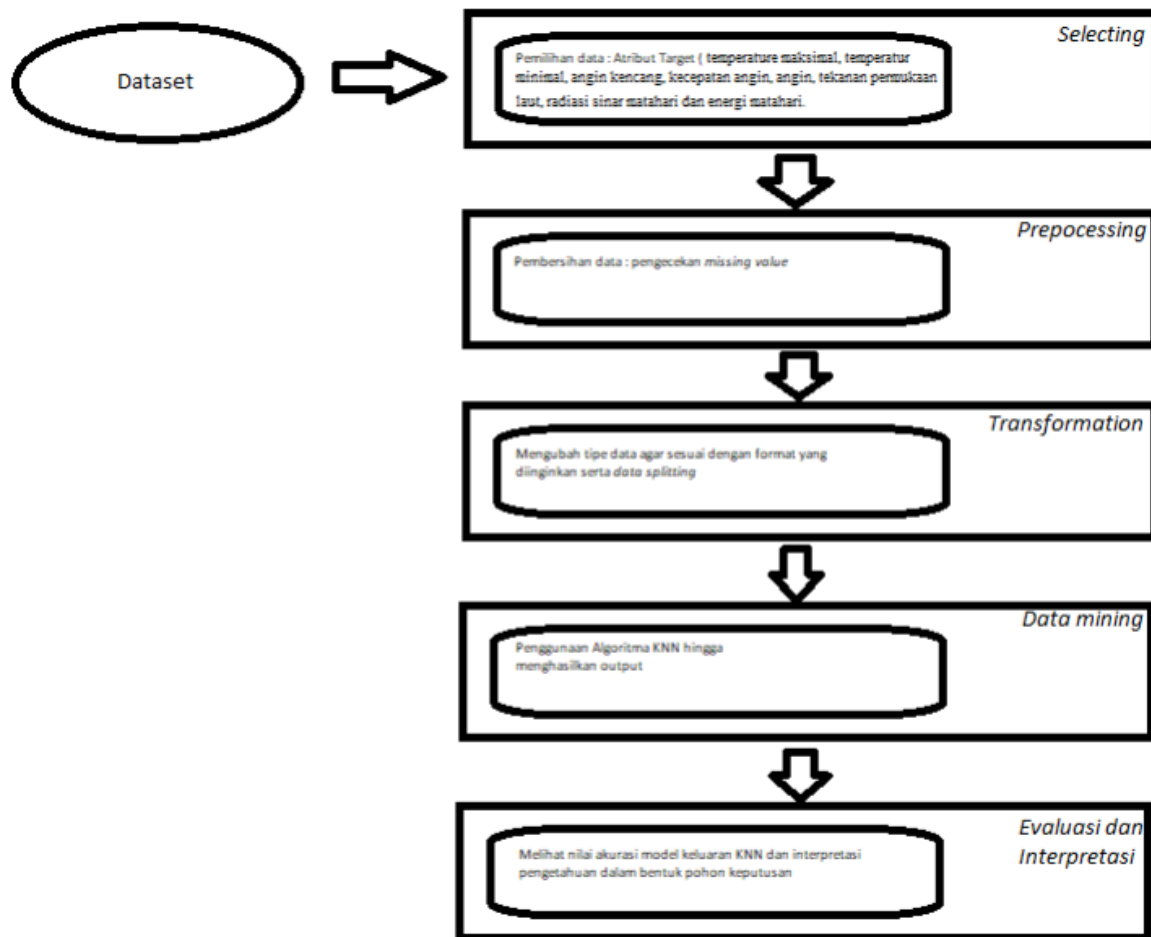


Figure 1. Stages of Knowledge Discovery Methodology in Database

2.1. Selecting

There are 32 explanatory attributes and 1 target attribute, referred to as independent and dependent variables, respectively. The explanatory variables collected include maximum temperature, minimum temperature, strong wind, wind speed, wind direction, sea-level pressure, solar radiation, and solar energy. The target variable is labeled as "Icon." Weather condition data were obtained from visualcrossing.com.

2.2. Preprocessing

In this stage, the collected data undergo a missing value check. If any missing values are found, they are handled by removing the corresponding data entries.

2.3. Transformation

This stage involves converting data types from nominal or object types into numerical formats, such as integers or floats. Additionally, the dataset is split into two subsets: training data and testing data. The Percentage Split technique is used for this data division.

2.4. Data Mining

Google Colab is used to process the data using the K-Nearest Neighbor algorithm. Google Colab provides various libraries that can be utilized to implement the K-Nearest Neighbor algorithm effectively.

Evaluation is performed by inputting the testing data into the model generated by the K-Nearest Neighbor algorithm. The main evaluation indicator is the model's accuracy. The higher the accuracy value, the better the predictive quality of the model. To calculate the distance between two points in the KNN algorithm, the Euclidean Distance formula is used.

Formula 1 Dimensional Space

Multi Dimensional Space Formula

3.1. Data

#	Region	Year	Age	Sex	Height	Weight	Body Fat	Heart Rate	Respiratory Rate	Temperature	Pressure	SpO2	Notes
1	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
2	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
3	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
4	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
5	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
6	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
7	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
8	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
9	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
10	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
11	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
12	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
13	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
14	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
15	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
16	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
17	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
18	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
19	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
20	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
21	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
22	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
23	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
24	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
25	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
26	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
27	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
28	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
29	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg
30	Region 1	2023-01-01	25.2	25.6	25.3	34.2	23.6	25.8	23.1	87.8	97	100	12.5 mmHg

3.2. System Design Process

Khatulistiwa Smart : Journal of Artificial Intelligence
43 | M Noer Fadli Hidayat



Figure 3. Process Flow Diagram

3.2.1. Data Input

The first step in the process is the input of tabulated weather data. The source code used in this process is shown in the following figure:

```

[1]
from sklearn.metrics import accuracy_score
from sklearn.model_selection import train_test_split

from scipy.stats import mode

[2] dataset = pd.read_csv('/content/drive/MyDrive/Colab Notebooks/bahan/dataset.csv')
dataset

dataset.head()

```

tempmin	temp	feelslikemar	feelslikemin	feelslike	dew	humidity	...	solarenergy	uvindex	severerisk	sunrise	sunset	moonphase
23.6	25.3	34.2	23.6	25.8	23.1	87.8	...	9.0	4	NaN	2022-01-01T05:08:49	2022-01-01T17:42:56	0.99
23.5	25.0	32.8	23.5	25.3	22.8	88.2	...	9.6	4	NaN	2022-01-02T05:09:21	2022-01-02T17:43:21	1.00
23.0	25.2	34.5	23.0	25.7	22.0	83.5	...	13.4	5	NaN	2022-01-03T05:09:52	2022-01-03T17:43:44	0.00

Figure 4. Data input process

3.2.2. S

The data preprocessing stage is carried out after the input of raw data to ensure it is suitable for use in training and testing phases. This step involves several

important activities that are critical for preparing the dataset for machine learning analysis. One of the primary tasks is converting nominal data into numerical formats, which allows the data to be processed computationally by algorithms that require numerical input. In addition, a class attribute is introduced to represent the classification category for forest fire occurrences. Several textual attributes—such as condition, situation, description, and icon—are also transformed into numeric values. This transformation ensures consistency and compatibility across all features in the dataset. These preprocessing efforts are fundamental to the success of the classification model, particularly because the K-Nearest Neighbor algorithm relies heavily on numerical distance calculations to make predictions. The entire process of data transformation is visually represented in the following figure.

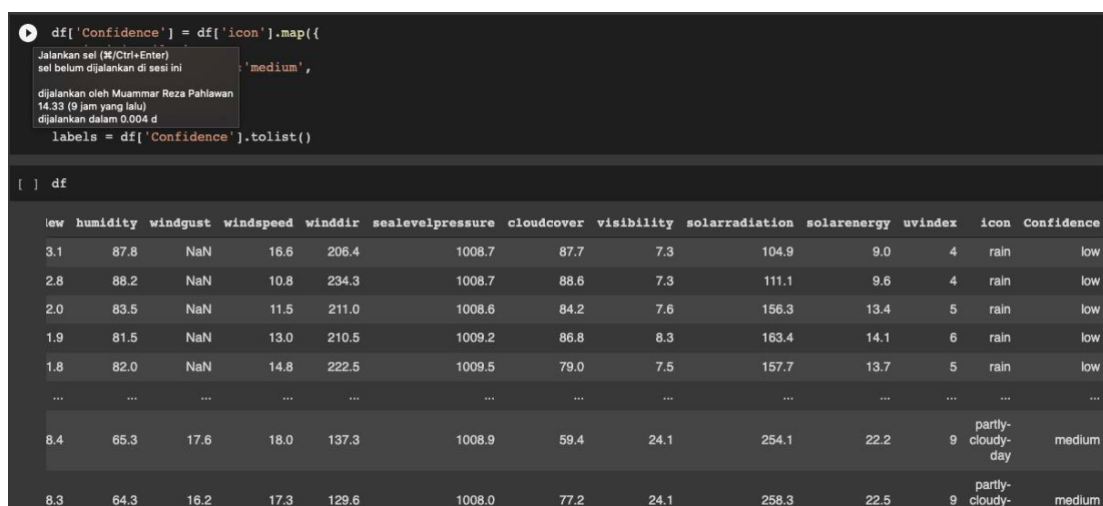


Figure 5. Confidence Attribute Addition Data

windgust	windspeed	winddir	sealevelpressure	cloudcover	visibility	solarradiation	solarenergy	uvindex	icon	Confidence	Class
NaN	16.6	206.4	1008.7	87.7	7.3	104.9	9.0	4	rain	low	0
NaN	10.8	234.3	1008.7	88.6	7.3	111.1	9.6	4	rain	low	0
NaN	11.5	211.0	1008.6	84.2	7.6	156.3	13.4	5	rain	low	0
NaN	13.0	210.5	1009.2	86.8	8.3	163.4	14.1	6	rain	low	0
NaN	14.8	222.5	1009.5	79.0	7.5	157.7	13.7	5	rain	low	0
...	...	...	...	...	...	...	...	...	...	...	...
17.6	18.0	137.3	1008.9	59.4	24.1	254.1	22.2	9	partly-cloudy-day	medium	1
16.2	17.3	129.6	1008.0	77.2	24.1	258.3	22.5	9	partly-cloudy-day	medium	1
19.4	20.5	132.2	1006.9	66.6	24.1	252.9	22.2	9	partly-cloudy-day	medium	1
23.8	20.5	128.8	1007.0	76.0	24.1	235.8	20.1	8	partly-cloudy-day	medium	1
24.1	19.8	126.0	1008.0	67.3	24.1	260.0	22.8	9	partly-cloudy-day	medium	1

Figure 6. Class Addition

3.2.3. K-Nearest Neighbor Classification

In this stage, the classification process is carried out using the K-Nearest Neighbor (K-NN) algorithm. The prepared dataset, which has undergone preprocessing, is now ready to be processed by the classification model. The K-NN algorithm works by identifying the K closest data points (neighbors) to a given input instance based on a chosen distance metric commonly Euclidean distance and assigning the class label that is most common among those neighbors. This method enables the model to classify new or unseen data based on similarities to existing labeled data in the training set. The implementation of this classification process using the K-NN algorithm is shown in the following figure.



Figure 7. Flow Chart K-Nearest Neighbor

3.2.4. Classification Results

The classification process was conducted using a training dataset consisting of 211 records, with three target classes: High, Medium, and Low. The data was split using a 70:30 ratio for training and testing purposes. The model was evaluated using 30% of the total dataset as testing data, and the results demonstrated a high level of accuracy. With the value of K set to 3, the K-Nearest Neighbor algorithm achieved a classification accuracy of 98%, indicating the model's strong ability to correctly predict forest fire risk categories based on the input features.

3.2.5. Validation of Results

To validate the model, the predicted results were compared against the actual class values. This validation process ensures the reliability of the model by measuring its ability to correctly classify new data based on patterns learned during training. The high correspondence between predicted and actual results further supports the robustness and effectiveness of the classification model.

4. CONCLUSIONS

The classification process, as defined in the selected methodology, was implemented using the K-Nearest Neighbor (K-NN) algorithm, which serves as a proximity-based search

technique. This method classifies new data points based on the majority class among their nearest neighbors in the training dataset. For this study, a total of 211 weather data records were used, comprising various environmental attributes relevant to forest fire prediction. The dataset was divided using a 70:30 ratio, with 70% used for training the model and the remaining 30% for testing. The testing phase aimed to evaluate the model's ability to correctly classify unseen data based on patterns learned from the training phase. The value of K was set to 3, meaning that each prediction was made based on the three closest data points in the feature space. This configuration yielded a classification accuracy of 98%, demonstrating the model's strong performance in distinguishing between the three fire risk categories: High, Medium, and Low. The high level of accuracy indicates that the selected features—such as temperature, wind speed, solar radiation, and pressure—are highly relevant and contribute significantly to the classification outcomes. Additionally, the result suggests that the K-NN algorithm is well-suited for this type of environmental classification problem, especially when the dataset is well-preprocessed and properly structured. This accuracy score reflects the potential of K-NN in supporting early warning systems for forest fire risks and assisting decision-makers in disaster mitigation planning.

5. REFERENCES

- Karo, I. M., Khosuri, A., Septory, J. S. I., & Supandi, D. P. (2022). Pengaruh metode pengukuran jarak pada algoritma k-NN untuk klasifikasi kebakaran hutan dan lahan. *Jurnal Media Informatika Budidarma*, 6(2), 1174–1182.
- Mahdi, M. I. (2022, April 22). Luas kebakaran hutan dan lahan Indonesia meningkat pada 2021. *Varia – Data Indonesia*.
- Noviansyah, M. R., Rismawan, T., & Midyanti, D. M. (2018). Penerapan data mining menggunakan metode K-Nearest Neighbor untuk klasifikasi indeks cuaca kebakaran berdasarkan data AWS (Automatic Weather Station) (Studi kasus: Kabupaten Kubu Raya). *Jurnal Coding, Sistem Komputer Untan*, 6(2), 48–56.
- Pratiwi, T. A., Irsyad, M., Kurniawan, R., Agustian, S., & Negara, B. S. (2021). Klasifikasi kebakaran hutan dan lahan menggunakan algoritma Naïve Bayes di Kabupaten Pelalawan. *CESS (Journal of Computer Engineering System and Science)*, 6(2), 139–148.
- Primajaya, A., Sari, B. N., & Khusaeri, A. (2020). Prediksi potensi kebakaran hutan dengan algoritma klasifikasi C4.5 (Studi kasus Provinsi Kalimantan Barat). *JEPIN (Jurnal Edukasi dan Penelitian Informatika)*, 6(2), 188–192.
- Rahayu, W. I., Prianto, C., & Novia, E. A. (2021). Perbandingan algoritma K-Means dan Naïve Bayes untuk memprediksi prioritas pembayaran tagihan rumah sakit berdasarkan tingkat kepentingan pada PT. Pertamina (Persero). *Jurnal Teknik Informatika*, 14(1), 1–8.
- Rudiyana, A., Dzulkifli, A. E., & Munazar, K. (2022). Klasifikasi kebakaran hutan menggunakan metode K-Nearest Neighbor: Studi kasus hutan Provinsi Kalimantan Barat. *JTIM: Jurnal Teknologi Informasi dan Multimedia*, 4(2), 195–202.
- Twin, A. (2021, September 17). Business marketing essentials: Data mining. *Investopedia*.