



Comparing Machine Learning Algorithms for Flood Prediction

Hamzah¹, Selly Dwipuspita Ayuningsih²

Abstract

Flooding is a major environmental hazard that requires accurate and reliable prediction to support disaster mitigation and early warning systems. This study applied a combination of K-Means clustering and Random Forest classification to predict flood conditions based on water-level data from DKI Jakarta. The dataset was obtained from Open Data Jakarta and contained observations recorded from January to December 2020. We utilized K-Means to group water-height observations and evaluated cluster configurations from $k = 3$ to $k = 20$. The resulting cluster information was then incorporated as an additional feature into the Random Forest model. We evaluated the model using Accuracy, Precision, Recall, F1 Score, and a confusion matrix. The results showed that the K-Means configuration affected the classification performance. The highest F1 Score reached 0.90 at $k = 14$, while the highest Accuracy reached 0.96 at $k = 15$ and $k = 20$. The Random Forest model achieved an Accuracy of 0.95, with weighted Precision, Recall, and F1 Score of 0.96, 0.95, and 0.95, respectively. Class 0 achieved the highest F1 Score of 0.98, while Class 2 obtained a Recall of 0.99. These findings demonstrate that the K-Means and Random Forest approach can effectively support flood prediction using water-level data and provide a promising basis for developing data-driven flood monitoring and prediction systems.

Keywords: Flood, Prediction, K-Means, Random Forest

This is an open-access article under the [CC BY-SA](#) license



1. Introduction

Flooding remains a major natural hazard affecting communities, infrastructure, agriculture, and economic activities. Climate change, irregular rainfall, urbanization, and inadequate drainage further increase flood risk. Traditional hydrological and statistical models often struggle with nonlinear environmental relationships. Machine learning (ML) can address this issue by learning patterns from historical data. However, the performance of RF, SVM, GB, XGBoost, and deep-learning models varies across datasets and prediction objectives. Therefore, researchers need to identify the most reliable algorithm for specific flood conditions [20].

Several studies report strong performance from ensemble ML models. Mo et al. show that combining meta-heuristic optimization with ML improves flood susceptibility prediction. Asrade et al. compare RF, GB, and XGBoost using environmental and topographic factors. GB and XGBoost achieve 0.97 accuracy, while RF reaches 0.96. These results confirm the effectiveness of ensemble methods. However, similar performance among algorithms highlights the need for direct comparison under consistent conditions [1], [12].

Urban flood prediction is challenging because urbanization changes drainage systems and surface runoff. Chaimook et al. compare RF, LightGBM, and XGBoost for flood occurrence and flood-level prediction in Bangkok. XGBoost achieves 91.42% accuracy for flood occurrence, while RF produces the lowest flood-level errors. Rainfall, humidity, temperature, and wind speed are major predictors. These findings show that the best algorithm depends on the prediction target [2].

Spatial characteristics also influence flood prediction performance. Rahimi and Reza integrate geospatial or remote-sensing data with ML for flood susceptibility and hazard mapping. Chomani compares MLC, SVM, and RF using rainfall, elevation, slope, soil, land use, and drainage factors. SVM achieves the highest ROC-AUC of 94.1%, followed by RF at 88.8% and MLC at 87.3%. These results indicate that model performance can vary across geographical conditions [4], [7], [18].

Temporal changes can further affect model reliability. Prasertsri compares RF, XGBoost, GB, and SVM using environmental variables across historical and projected periods in Thailand. XGBoost maintains accuracy above 0.95, while SVM achieves zero recall for positive flood cases in projected 2025 data. This finding shows that strong historical performance does not always guarantee reliable future prediction. Temporal validation is therefore important in flood prediction studies [8].

Data availability also affects ML performance. Uddameri and Hernandez report that LSTM and CNN can support flood-related applications, but simpler models may perform better with limited data. They also identify limited public datasets as a major challenge. Zuhairi et al. address this issue through data cleaning, interpolation, and resampling before applying tree-based models. Their findings demonstrate that preprocessing and data quality strongly influence prediction performance [3], [13].

Reliable evaluation is essential because class imbalance can make accuracy misleading. Zuhairi et al. identify data imbalance as an important challenge in tropical flood forecasting. Kafi et al. compare RF, XGBoost, and SVM and combine them with explainable AI. RF achieves a precision of 0.857 and ROC-AUC of 0.93, outperforming SVM. Their study also identifies elevation and rainfall as important factors. These findings support the use of precision, recall, F1-score, and ROC-AUC alongside accuracy [13], [14].

Recent studies increasingly combine deep learning, multimodal data, and ensemble methods. Hasi et al. combine Sentinel-1 SAR and Sentinel-2 imagery, where EfficientNet-B0 achieves 93.57% accuracy and XGBoost stacking reaches 95.43%. Grad-CAM and saliency maps further improve model interpretation. Rasalkar and Surve report growing use of CNN, LSTM, GRU, Transformer, and hybrid models. Although these approaches can improve performance, they also require greater computational resources and preprocessing. Therefore, systematic comparisons remain important for determining whether complex models provide meaningful advantages [19], [20].

2. Related Works

Recent studies have increasingly applied machine learning to flood prediction because ML can capture nonlinear relationships among environmental variables. Mo et al. investigated a hybrid approach that combined meta-heuristic optimization with ML for flood susceptibility mapping. Their results showed that the hybrid models consistently achieved higher prediction accuracy than individual ML models. The study demonstrated the benefit of optimizing ML models before prediction. However, the additional optimization process increased model complexity. The study also focused on flood susceptibility rather than direct comparison of standard ML algorithms under a common prediction setting [1].

Chaimook et al. directly compared Random Forest (RF), LightGBM, and XGBoost for flood occurrence and flood-level prediction in Bangkok. XGBoost achieved the highest accuracy of 91.42% for flood occurrence. In contrast, RF produced the lowest errors for flood-level prediction, with an MSE of 39.33 cm² and MAE of 4.43 cm. Rainfall, relative humidity, maximum temperature, and wind speed were the most influential variables. The study clearly showed that algorithm performance depended on the prediction target. However, its evaluation focused mainly on tree-based models and an urban environment [2].

Asrade et al. compared RF, Gradient Boosting (GB), and XGBoost for flood susceptibility mapping in the Choke Watershed. They used elevation, slope, rainfall, soil, land use, drainage density, and several topographic indices as conditioning factors. GB and XGBoost achieved the highest test accuracy of 0.97, while RF reached 0.96. The study demonstrated the strong capability of ensemble models in handling multiple environmental factors. However, the study concentrated on susceptibility mapping in one watershed. Its findings may therefore not directly represent other flood prediction environments [12].

Zuhairi et al. evaluated several tree-based algorithms for flood forecasting in Selangor, Malaysia. Their models included Decision Tree, RF, XGBoost, LightGBM, AdaBoost, CatBoost, and Gradient Boosting. They first applied data cleaning, interpolation, and resampling to address data quality and imbalance. The resulting Tree-Based Ensemble model outperformed benchmark models across most evaluation metrics. The study showed that combining several tree-based learners can improve prediction reliability. However, the ensemble approach also made direct identification of the best individual algorithm more difficult [13].

Rizal et al. compared seven ML algorithms for urban flood prediction in Makassar City. The evaluated models included Logistic Regression, Linear Discriminant Analysis, K-Nearest Neighbors, Gaussian Naive Bayes, SVM, AdaBoost, and RF. RF achieved the best overall performance and showed strong capability in processing complex urban flood data. The study provided useful evidence that algorithm selection affects urban flood prediction performance. However, it did not include several advanced boosting models such as XGBoost and LightGBM. This limitation leaves room for broader algorithm comparisons [15].

Kafi et al. compared RF, XGBoost, and SVM for flood-risk prediction in Bauchi, Nigeria. RF performed best, achieving a precision of 0.857 and ROC-AUC of 0.93. SVM produced the lowest performance, with a precision of 0.757 and ROC-AUC of 0.84. The researchers also applied explainable AI and identified settlement formality, elevation, population, and rainfall as important factors. The study strengthened the importance of combining model comparison with interpretability. However, its results were based on flood-risk mapping and may differ from models designed for temporal flood occurrence prediction [14].

Prasertsri examined the stability of RF, XGBoost, GB, and SVM across historical and projected flood conditions in Thailand. XGBoost maintained accuracy above 0.95 across the evaluated years. SVM achieved an average historical accuracy of 0.970 but failed to detect positive flood cases in the projected 2025 data, producing a recall of zero. This result demonstrated that high historical accuracy does not necessarily indicate reliable future prediction. The study therefore emphasized temporal validation as an important component of algorithm comparison. However, its analysis focused on projected climate scenarios and may not fully represent short-term operational flood forecasting [8].

Hasi et al. compared classical ML, deep learning, and ensemble approaches using Sentinel-1 SAR and Sentinel-2 optical imagery. EfficientNet-B0 achieved 93.57% accuracy among individual models, while XGBoost stacking increased accuracy to 95.43%. The researchers also used Grad-CAM and saliency maps to verify that the models focused on meaningful flood-related regions. Their findings showed that multimodal and ensemble approaches can outperform individual models. However, the approach required satellite imagery, advanced preprocessing, and greater computational complexity[19].

Previous studies showed that no single ML algorithm consistently produced the best flood prediction results. RF performed strongly in urban and flood-risk studies, while XGBoost and GB achieved high accuracy in several susceptibility and forecasting applications. SVM also performed well in some spatial settings but showed weak temporal generalization in projected conditions. Hybrid and ensemble methods further improved performance but introduced additional complexity. The proposed study addresses this gap by systematically comparing machine learning algorithms to identify the most reliable approach for flood prediction.

3. Proposed Method

This study proposed a machine learning framework that combined K-Means clustering and RF to predict flood occurrence. The method first grouped environmental observations into homogeneous clusters using K-Means. We then used the resulting cluster information together with the original environmental features as inputs to the RF classifier. This approach aimed to improve the representation of different environmental conditions before performing flood prediction.

1. Data Pre-processing

In this study, we first prepared the flood dataset by removing inconsistent records and handling missing values. Numerical features were normalized to reduce differences in scale. For each feature x_j , we applied min-max normalization:

$$x'_{ij} = \frac{x_{ij} - x_j^{\min}}{x_j^{\max} - x_j^{\min}} \quad (1)$$

where x_{ij} represents the value of feature j in observation i . The normalized value x'_{ij} ranges from 0 to 1. This process provided a consistent scale for the K-Means algorithm.

2. K-Means Clustering

This paper applied K-Means to identify groups of observations with similar environmental characteristics. The algorithm assigned each observation to the nearest cluster centroid. The distance between observation x_i and centroid μ_k was calculated using Euclidean distance:

$$d(x_i, \mu_k) = \sqrt{\sum_{j=1}^p (x_{ij} - \mu_{kj})^2} \quad (2)$$

K-Means minimized the within-cluster sum of squared distances:

$$J = \sum_{k=1}^K \sum_{x_i \in C_k} \|x_i - \mu_k\|^2 \quad (3)$$

where K represents the number of clusters, C_k represents cluster k , and μ_k represents its centroid. The resulting cluster label was then added as an additional feature for flood prediction. Thus, the clustering process provided information about the environmental condition of each observation.

3. Random Forest Prediction

We used Random Forest as the main prediction algorithm. RF constructed multiple decision trees using different subsets of the training data and features. For a classification problem, each tree generated a predicted class $h_t(x)$. The final prediction was determined through majority voting:

$$\hat{y} = \text{mode}\{h_1(x), h_2(x), \dots, h_T(x)\}$$

where T represents the number of decision trees and $h_t(x)$ represents the prediction of tree t . The RF model used both the original environmental variables and the K-Means cluster label. This combination allowed the model to consider both individual environmental factors and their overall similarity patterns.

4. Proposed K-Means–RF Framework

The proposed workflow can be expressed as:

$$X \rightarrow X' \rightarrow \text{K-Means}(X') \rightarrow [X', C] \rightarrow \text{RF} \rightarrow \hat{y}$$

where X denotes the original dataset, X' denotes the normalized features, and C represents the K-Means cluster labels. The combined feature set $[X', C]$ was used to train the RF model. The final output $\hat{y}$ represented the predicted flood class.

5. Model Evaluation

This study evaluated the proposed model using accuracy, precision, recall, and F1-score. These metrics provided a more comprehensive assessment than accuracy alone.

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (4)$$

$$\text{Precision} = \frac{TP}{TP + FP} \quad (5)$$

$$\text{Recall} = \frac{TP}{TP + FN} \quad (6)$$

$$F1 = 2 \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (7)$$

TP , TN , FP , and FN denote true positive, true negative, false positive, and false negative predictions, respectively. Accuracy measured the overall correctness, while precision and recall assessed the quality of flood detection. The F1-score provided a balanced measure between precision and recall.

In this study, we utilized K-Means as an unsupervised feature-structuring technique and Random Forest as the supervised flood prediction model. K-Means identified environmental patterns before prediction, while RF learned the relationship between these patterns, environmental variables, and flood occurrence. This approach provided a relatively simple framework while retaining the ability of RF to model nonlinear relationships. The performance of the proposed K-Means–RF model was then compared with individual machine learning algorithms to determine whether clustering information improved flood prediction.

4. Experimental Setup.

1. Dataset

In this study, we used a public dataset obtained from Open Data Jakarta. The dataset recorded river water levels and river–sea water levels across DKI Jakarta during January–December 2020. The data provided temporal and hydrological information that was relevant for flood prediction. We prepared the dataset through several steps, including data collection, cleaning, normalization, transformation, and splitting. During data cleaning, we

selected relevant variables and removed unnecessary attributes. We also checked abnormal values and adjusted them to appropriate ranges during normalization.

This study transformed the flood status into a binary categorical variable, where 0 represented non-flood conditions and 1 represented flood conditions. We then divided the prepared dataset into training and testing sets. The training set contained 80% of the observations and was used to develop the machine learning models. The remaining 20% formed the testing set and was used to evaluate model performance on unseen data. This preparation ensured that the K-Means and Random Forest models received consistent and suitable input data for flood prediction.

Table 1. Dataset Characteristics and Preparation

Aspect	Description
Data source	Open Data Jakarta
Data type	Public hydrological data
Study area	DKI Jakarta Province, Indonesia
Observation period	January–December 2020
Main information	River water level and river–sea water level
Data collection	Retrieved from Jakarta Open Data
Data cleaning	Selected relevant variables and removed unnecessary attributes
Data normalization	Checked abnormal values and adjusted variable ranges
Data transformation	Converted flood status into numerical classes
Flood class 0	Non-flood condition
Flood class 1	Flood condition
Training data	80%
Testing data	20%
Main application	Flood prediction using K-Means and Random Forest

2. Data Processing

In this study, we simplified the initial dataset using Python to retain only the variables that were relevant to flood prediction. The original dataset contained several attributes, but some variables were redundant or did not directly contribute to the prediction process. We therefore selected five variables: door_code, water_height, city_code, date, and standby_status. These variables represented the monitoring location, water level, city information, observation date, and flood warning status, respectively. The simplified dataset contained 323,190 observations collected from January 1 to December 31, 2020.

We then converted the categorical flood status into numerical values to make the data suitable for machine learning. The standby_status variable was used to represent the flood condition, while water_height provided the main hydrological measurement. The date variable preserved the temporal information, and door_code and city_code identified the monitoring location. This pre-processing step reduced unnecessary information and produced a structured dataset for the subsequent K-Means clustering and Random Forest classification processes.

Table 2. Flood Status Transformation

Original Status	Numerical Code	Interpretation
Status: Normal	0	Non-flood condition
Status: Alert 1	1	Flood-alert condition

The coding process transformed the standby_status variable into a binary representation. We assigned **0** to normal conditions and **1** to Alert 1 conditions. This transformation enabled the Random Forest model to treat flood status as a classification target. It also provided a consistent numerical representation for model training and evaluation.

5. Result and Analysis

1. K-Means

The K The K-Means algorithm groups observations into clusters based on similarity. In this study, we applied K-Means to group observations according to water height. We first tested the number of clusters from $k=3$ to $k=20$. For each k , the algorithm initialized the centroids as the center points of the clusters. It then calculated the distance between each observation and the centroids. Each observation was assigned to the cluster with the nearest centroid. The centroid was subsequently updated using the mean value of the observations assigned to each cluster.

The clustering process continued iteratively until the centroid positions became stable or the algorithm reached the maximum number of iterations. This procedure produced a clustering result for every k value between 3 and 20. We compared these results to identify the most appropriate number of clusters for the water-height data. The selected cluster information was then used as an additional feature in the subsequent Random Forest model for flood prediction.

2. Random Forest

The results of K-Means 3 to 20 clustering are used as material to classify water height which is one of the parameters in the random forest algorithm. The results of the random forest can be seen in Table 3.

Table 3. Random Forest Results

k	F1 Score	Accuracy
3	0.01	0.88
4	0.46	0.87
5	0.52	0.89
6	0.75	0.90
7	0.60	0.92
8	0.62	0.92
9	0.62	0.92
10	0.84	0.94
11	0.78	0.94
12	0.58	0.93
13	0.81	0.94
14	0.90	0.95
15	0.82	0.96
16	0.79	0.94
17	0.73	0.95
18	0.85	0.95
19	0.82	0.95
20	0.77	0.96

The K-Means results showed that increasing the number of clusters generally improved model performance, although the F1 Score fluctuated across different values of k . The

lowest F1 Score occurred at k = 3 (0.01), while the highest F1 Score was obtained at k = 14 (0.90). Accuracy increased from 0.88 at k = 3 to 0.96 at k = 15 and k = 20. The highest accuracy was therefore achieved at k = 15 and k = 20, while k = 14 provided the highest F1 Score. These results indicate that the number of clusters influenced the classification performance, particularly the balance between precision and recall represented by the F1 Score.

Table.4 Dataset after water level clustering

No.	Door Code	City Code	Status Standby	Water Height
1	0	16	1	0
2	6	2	0	5
3	19	3	3	8
4	10	3	3	10
5	0	10	4	19
6	3	3	3	8
...	...	...	...	...
323186	17	4	0	4
323187	18	3	0	9
323188	19	3	0	0
323189	21	2	0	9
323190	1	1	0	4

Table 5 shows the RF model with a Recall and F1 Score = 1.00 and Accuracy = 0.95.

Table 5: Random Forest Testing

Class	Precision	Recall	F1-Score	Support
0	0.99	0.97	0.98	51,100
1	0.84	0.76	0.80	5,671
2	0.70	0.99	0.82	3,875
3	0.96	0.85	0.90	327
Accuracy			0.95	60,973
Macro Average	0.87	0.89	0.88	60,973
Weighted Average	0.96	0.95	0.95	60,973

The RF model achieved an overall accuracy of 0.95 on the test dataset, indicating that the model correctly classified most observations. Class 0 obtained the highest F1-Score of 0.98, supported by a precision of 0.99 and recall of 0.97. Class 2 achieved the highest recall at 0.99, although its precision was lower at 0.70. Class 3 also showed strong performance with an F1-Score of 0.90. The weighted average reached 0.96 precision, 0.95 recall, and 0.95 F1-Score. These results demonstrate that the Random Forest model provided strong overall classification performance, although its performance varied across classes. Fig. 1 displays the confusion matrix for the K = 14 configuration to identify misclassified observations.

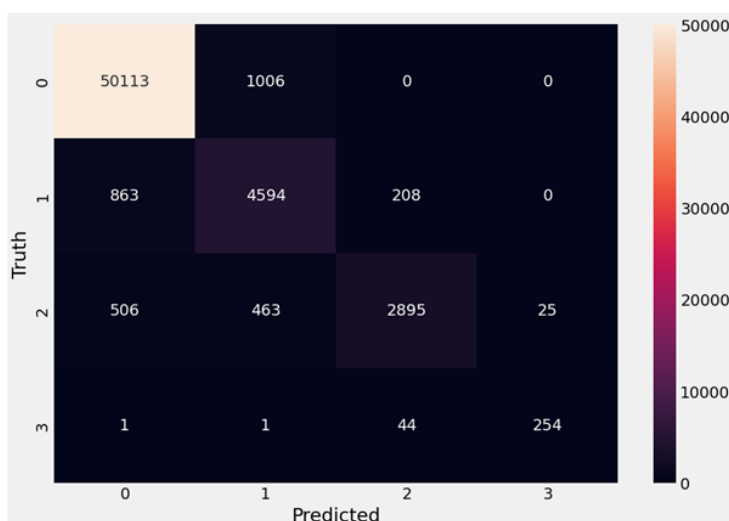


Fig. 1 Confusion Matrix for the K = 14

6. Conclusion

Based on the results, this study concludes that the combination of K-Means and Random Forest provided effective performance for flood prediction. We utilized K-Means to group water-height observations into different clusters and tested cluster configurations from $k = 3$ to $k = 20$. The results showed that the clustering configuration affected the subsequent classification performance. The highest F1 Score reached 0.90 at $k = 14$, while the highest Accuracy reached 0.96 at $k = 15$ and $k = 20$. These findings indicate that the clustering stage helped organize the water-height data into meaningful groups before classification.

This paper applied the selected K-Means clustering results as an additional feature for the Random Forest model. The Random Forest model achieved an Accuracy of 0.95 on 60,973 test observations. It also obtained a weighted Precision of 0.96, Recall of 0.95, and F1 Score of 0.95. Class 0 produced the strongest performance with an F1 Score of 0.98, while Class 3 achieved an F1 Score of 0.90. Class 2 obtained a high Recall of 0.99, although its Precision was lower at 0.70. These results show that the model performed well overall but still experienced differences in classification performance among classes.

Thus, we demonstrated that the K-Means and Random Forest approach can support flood prediction using water-level data. The clustering process provided additional information that helped the Random Forest model distinguish different water-height conditions. The results also show that Accuracy alone is not sufficient to evaluate the model because class-level metrics reveal differences in prediction performance. We therefore recommend using Accuracy, Precision, Recall, F1 Score, and the confusion matrix together when evaluating flood prediction models. Future studies can improve this approach by incorporating additional environmental variables, longer observation periods, and more advanced ensemble or deep learning models.

Acknowledgment

The author would like to thank BNPB (National Disaster Management Agency) for providing historical disaster data that can be accessed by anyone who needs it and also especially to BMKG (Meteorology, Climatology and Geophysics Agency) for providing supporting data to make this research run smoothly

References

- [1] Li Mo, X. Qi, Y. Cai, C. Lu, J. Li, and L. Liu, "A hybrid approach combining meta-heuristic algorithm and machine learning for flood susceptibility mapping," *Natural Hazards*, 2026.
- [2] P. Chaimook, N. Khamsemanan, C. Nattee, and A. Sharp, "Employing machine learning for flood prediction: A comparative study of tree-based techniques," in *Proc. 5th Asia Conf. Information Engineering (ACIE)*, 2025.
- [3] V. Uddameri and E. A. Hernandez, "Machine learning for flood resiliency—Current status and unexplored directions," *Environments*, 2025.
- [4] M. Rahimi, "Integrating geospatial intelligence and machine learning for flood susceptibility mapping," *Scientific Reports*, 2026.
- [5] N. Manyaka, "Integrating remote sensing and machine learning for flood modelling: A systematic literature review," *Physics and Chemistry of the Earth*, 2026.
- [6] E. Alcântara, "Interpretable machine learning for flood susceptibility mapping in the metropolitan region of São Paulo, Southeast Brazil," *Discover Geoscience*, 2025.
- [7] S. K. Reza, "Integrating remote sensing and machine learning for flood hazard zonation in Gomati district of Tripura, Northeast India," *Discover Environment*, 2026.
- [8] N. Prasertsri, "Spatio-temporal machine learning for flood risk assessment under SSP scenarios: A case study of Maha Sarakham, Thailand," *Sustainability*, 2026.
- [9] K. M. Gilbert, "Interpretable ensemble machine learning for flood susceptibility mapping in Lagos Lagoon, Nigeria," *Journal of African Earth Sciences*, 2026.
- [10] V. S. Shradha, "Integrating multi-criteria decision analysis with ensemble machine learning for flood-induced landslide prediction and risk mapping," *Natural Hazards*, 2026.
- [11] H. Ali, "Site-dependent performance of spectral indices and machine learning for flood detection in Mediterranean karst poljes," in *Proc. IEEE Mediterranean and Middle East Geoscience and Remote Sensing Symposium (M2GARSS)*, 2026.
- [12] T. M. Asrade, "Flood susceptibility assessment using three machine learning techniques and comparison of their performance," *Scientific Reports*, 2026.
- [13] H. Zuhairi, F. Yakub, M. Omar, M. Sharifuddin, K. A. Razak, A. Faruq, and A. Wijaya, "Performance analysis of tree-based ensemble machine learning model for flood forecasting in tropical regions," *IEEE Access*, 2025.
- [14] K. M. Kafi, Z. Ponrahono, Z. H. Ash'aari, and A. S. Barau, "Flood risk prediction and modeling in Bauchi: Leveraging machine learning models and explainable AI for urban resilience," *The Journal of Climate Change and Health*, 2025.
- [15] H. Rizal, M. Hariadi, Y. M. Arif, and E. Warni, "Enhancing flood prediction in urban areas: A machine learning approach for Makassar City," *Engineering, Technology & Applied Science Research*, 2025.
- [16] L. L. Bosi, A. C. Mendes, R. M. Guedes, and R. M. Salles, "Traffic classification in software-defined networking (SDN): Application of machine learning models for flooding DDoS attack detection," in *Proc. IEEE 3rd Industrial Electronics Society Annual On-Line Conf. (ONCON)*, 2024.
- [17] H. Auma, "Flood risk modeling for Kisumu County, Kenya: Integrating multi-source geospatial data and machine learning," *Frontiers in Environmental Science*, 2026.
- [18] K. Chomani, "Flood susceptibility mapping in Erbil City using machine learning techniques," *Journal of Studies in Civil Engineering*, 2026.
- [19] J. M. Hasi, S. S. Nakshi, and M. Sultana, "Explainable multi-modal ensemble learning for flood prediction from SAR and optical satellite imagery," in *Proc. IEEE 2nd International Conf. Quantum Photonics, Artificial Intelligence & Networking (QPAIN)*, 2026.
- [20] P. M. Rasalkar and A. Surve, "Evolution of flood prediction methods: Hydrological, statistical, machine learning and deep learning approaches (2020–2025)," in *Proc. International Conf. Future Technologies (ICFT)*, 2025.