
| RESEARCH ARTICLE

Machine Learning-Based Drought Classification Using Meteorological Data: Toward Smarter Environmental Models for Site Exploration

Sina Saghaei¹✉, Seyedehzahra Shafa², Paniz Farahani¹

¹MSc Student, Management of Built Environment, Politecnico di Milano, Milan, Lombardy, Italy

²Department of Architecture, University of Hartford, West Hartford, CT, USA

Corresponding Author: Sina Saghaei, **E-mail:** sina_saghaei1997@yahoo.com

| ABSTRACT

Drought presents a significant challenge to sustainable water management, agriculture, and geotechnical site assessment. Accurate and timely classification of drought severity is essential for anticipating environmental changes that impact soil behavior, moisture conditions, and subsurface dynamics. This study introduces a machine learning framework utilizing the Random Forest algorithm to classify drought intensity based on multi-year meteorological datasets. The dataset, comprising over 3 million records and 18 meteorological features, includes variables such as temperature, humidity, wind speed, and precipitation, serving as environmental indicators for drought conditions. The model predicts drought severity across six discrete classes (0–5), evaluated using precision, recall, F1 score, and confusion matrices. While achieving an overall accuracy of 75%, the model reveals performance challenges in detecting minority drought classes, underscoring the importance of class balancing and feature selection. Beyond drought prediction, this work supports the integration of environmental intelligence into geotechnical and site exploration processes. By informing early-stage terrain assessment with climate-responsive data, the framework lays groundwork for adaptive modeling in subsurface analysis, simulation, and exploration planning. The study demonstrates the value of machine learning in developing scalable, data-driven environmental models that enhance decision-making in both civil and environmental engineering applications.

| KEYWORDS

Machine Learning, Meteorological Data, Data-Driven Decision Making, Site Exploration, Geotechnical Applications, Drought Classification

| ARTICLE INFORMATION

ACCEPTED: 12 March 2025

PUBLISHED: 05 April 2025

DOI: 10.32996/jmcie.2025.6.2.1

1. Introduction

Drought is a pervasive environmental challenge that significantly impacts agriculture, water resource management, and infrastructure planning. Its effects are particularly evident in arid and semi-arid regions, where sustained periods of water scarcity can disrupt ecological stability and socioeconomic development. Drought conditions manifest through variations in key meteorological variables—such as precipitation deficits, elevated temperatures, shifts in humidity, and changes in wind speed—all of which offer crucial insights into soil moisture dynamics and broader climatic trends (Nguyen & Minh, 2021). Accurate forecasting of drought severity using these indicators is essential for proactive planning in water-intensive sectors and increasingly, in site-specific environmental assessments relevant to geotechnical and civil engineering.

Conventional drought prediction methods often rely on empirical indices or historical data analysis, which can oversimplify the non-linear and interdependent nature of meteorological parameters. These models typically face limitations in early detection and generalizability across diverse geographic regions. In contrast, machine learning (ML) techniques—particularly ensemble learning methods—have emerged as powerful tools for identifying complex patterns in large-scale climate data (Chang et al.,

2023). Among these, the Random Forest algorithm has proven effective in handling high-dimensional datasets, modeling non-linear interactions, and offering robust classification performance even in the presence of noisy or partially missing data. This study introduces a machine learning-based framework for classifying drought severity using the Random Forest algorithm, trained on over 3 million meteorological records spanning 15 years, sourced from a publicly available (*Data*, n.d.). The dataset comprises 18 environmental features, including temperature, specific humidity, wind speed (at multiple elevations), dew point, and precipitation, which are used to classify drought severity across six categories (0–5) (Azizi & Azizi, 2024). The model is implemented in Python from first principles—without pre-built ML libraries—to ensure algorithmic transparency and enable fine-tuned control of the classification process. Model evaluation is conducted using key metrics such as accuracy, precision, recall, F1 score, and confusion matrices, providing a comprehensive assessment of performance across imbalanced class distributions.

While the immediate objective of this work is to improve drought severity classification, its broader significance extends to applications such as high-fidelity simulations of subsurface interactions. Drought conditions can serve as proxies for key soil properties—such as moisture content, compaction, and cohesion—that critically influence the behavior of granular media (Quintana et al., 2023). These properties are central to DEM-MBD (Discrete Element Method–Multibody Dynamics) simulations, where accurate environmental parameters are essential for modeling soil response and interaction with physical systems. For example, (Shahhosseini et al., 2023) employed DEM-MBD simulations to investigate how variations in dry sand conditions affect penetration mechanics. Integrating drought classification data into such simulations can improve predictive modeling for a wide range of applications, from autonomous systems to civil and geotechnical engineering (Shahhosseini & Tao, 2025).

By developing a data-driven drought classification tool, this study contributes to an emerging body of work that links atmospheric conditions with terrain characterization and geotechnical applications. The resulting framework offers a foundation for adaptive, climate-responsive modeling in subsurface engineering contexts.

2. Research Methodology

2.1 Dataset Description and Feature Engineering

This study employs a cleaned and structured meteorological dataset compiled from publicly available historical drought and weather observations across the United States. The data spans over two decades and is designed for supervised classification of drought severity based on atmospheric and soil-related conditions. The dataset was preprocessed by the author and is available via an open-access Google Drive repository for reproducibility.

Each record in the dataset represents a specific date and county-level location in the United States, accompanied by the preceding 90 days of 18 meteorological indicators. The full dataset contains approximately 3 million observations, with features selected for their relevance to drought dynamics and environmental modeling. These include surface temperature, humidity, wind speeds at multiple elevations, precipitation, dew point, and surface pressure (Nguyen & Minh, 2021). A full list of features is summarized in Table 1.

Table 1: Overview of Meteorological Features

INDICATOR	DESCRIPTION	INDICATOR	DESCRIPTION
WS10M_MIN	Minimum Wind Speed at 10 Meters (m/s)	QV2M	Specific Humidity at 2 Meters (g/kg)
T2M_RANGE	Temperature Range at 2 Meters (C)	T2M	Temperature at 2 Meters (C)
WS10M	Wind Speed at 10 Meters (m/s)	WS50M_MIN	Minimum Wind Speed at 50 Meters (m/s)
T2M_MAX	Maximum Temperature at 2 Meters (C)	WS50M	Wind Speed at 50 Meters (m/s)
TS	Earth Skin Temperature (C)	WS50M_RANGE	Wind Speed Range at 50 Meters (m/s)
WS50M_MAX	Maximum Wind Speed at 50 Meters (m/s)	WS10M_MAX	Maximum Wind Speed at 10 Meters (m/s)
WS10M_RANGE	Wind Speed Range at 10 Meters (m/s)	PS	Surface Pressure (kPa)
T2MDEW	Dew/Frost Point at 2 Meters (C)	T2M_MIN	Minimum Temperature at 2 Meters (C)
T2MWET	Wet Bulb Temperature at 2 Meters (C)	PRECTOT	Precipitation (mm day ⁻¹)

The target variable in this study corresponds to six drought severity levels derived from the U.S. Drought Monitor (USDM) classification system. Each class represents a specific level of drought intensity, ranging from no drought to exceptional drought. These classes are encoded as follows:

Class 0 (None): No drought; Class 1 (D0 – Abnormally Dry): Early signs of dryness or conditions recovering from drought, with possible short-term impacts on planting and lingering water deficits; Class 2 (D1 – Moderate Drought): Noticeable stress on crops and pastures, declining water levels in streams and reservoirs, and voluntary water-use restrictions; Class 3 (D2 – Severe Drought): Crop and pasture losses likely, common water shortages, and imposed water-use restrictions; Class 4 (D3 – Extreme Drought): Major agricultural losses and widespread water shortages or restrictions; Class 5 (D4 – Exceptional Drought): Exceptional and widespread crop or pasture losses, critical water shortages, and emergency conditions in water supply systems. To prevent temporal data leakage and better reflect a realistic forecasting scenario, the dataset was chronologically divided into training, validation, and testing subsets. Data from 2000 to 2009 was used for training (approximately 47%), 2010 to 2011 for validation (10%), and 2012 to 2020 for testing (43%). This time-based split ensures the model is evaluated on future data relative to the training period, improving both reliability and generalizability.

2.2 Random Forest Classification: Theory and Model Architecture

The Random Forest algorithm is a powerful, supervised machine learning method used for both classification and regression tasks. It operates by constructing an ensemble of decision trees and aggregating their predictions to improve accuracy, robustness, and generalization. By leveraging the principles of ensemble learning and bagging (bootstrap aggregation), Random Forests reduce variance without significantly increasing bias, making them particularly effective in high-dimensional and noisy datasets such as meteorological records.

A key advantage of Random Forest lies in the randomization introduced during tree construction. Each tree is trained on a bootstrap sample of the data and, at each split, considers a random subset of features. This controlled randomness leads to decorrelated decision trees, which, when combined, produce a more stable and accurate model compared to any single tree.

2.2.1 Decision Trees

A decision tree is a hierarchical, flowchart-like model that sequentially evaluates feature-based conditions to classify or predict an outcome. At each internal node, the model asks a question about a specific feature; the outcome of that question determines the path taken to the next node. This continues until a terminal (leaf) node is reached, which corresponds to a class label or predicted value. The sequential logic resembles traversing a decision path, making the model both interpretable and computationally efficient.

Let a data point be denoted by a vector $v = (x_1, x_2, \dots, x_d) \in R^d$ where x_i is the value of the i -th feature and d is the total number of features. During training, a decision tree selects a random subset of these d features, reducing the feature space dimensionality to d' . This random feature subset is represented as:

$$\Phi(v) = (x_{\phi_1}, x_{\phi_2}, \dots, x_{\phi_{d'}}), \quad \phi_i \in \{1, \dots, d\}$$

At each decision node, a binary test function is applied to the selected feature subset. This function splits the input space into two disjoint sets—typically represented as 1 (True) or 0 (False) (Chang et al., 2023). The test function is defined as:

$$h(v, \theta_j): R^d \times T \rightarrow \{0, 1\}$$

Where $(\theta_j \in T)$ are the parameters of the test at node (j) , typically consisting of a feature index and threshold. For a given data point (v) , the function determines the direction of traversal—left or right.

Let (S_j) denote the subset of training data reaching node (j) , and let (S_L) and (S_R) be the resulting subsets after applying the test. The following properties hold:

$(S_j = S_L \cup S_R)$: the union of the two subsets equals the original dataset at node (j)

$(S_L \cap S_R = \emptyset)$: the subsets are disjoint

$(S_L = S_j + 1), (S_R = S_j + 2)$: node indexing follows a binary tree structure

In this study, the training dataset contains approximately 3 million instances, each with 19 features derived from meteorological measurements (Criminisi et al., 2012). The dataset can be represented as:

$$D = \{(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)\}$$

where $(x_i = \{x_{i1}, x_{i2}, \dots, x_{i19}\})$ is the feature vector for the $(i) - th$ data point, and $(y_i \in \{0,1,2,3,4,5\})$ is the corresponding drought classification label.

To illustrate how the model learns decision rules, we constructed a sample tree using three randomly selected features: **QV2M** (specific humidity at 2 meters), **WS10M** (wind speed at 10 meters), and **T2M_MIN** (minimum air temperature at 2 meters). The

resulting decision tree structure is shown in **Figure 1**, where thresholds at each node guide the branching, and the leaves denote the predicted drought class.

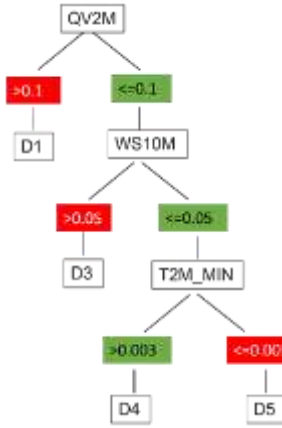


Figure 1: A sample decision tree using QV2M, WS10M, and T2M_MIN. Nodes represent thresholds; leaves represent predicted drought classes.

When constructing decision trees, particularly for multi-class tasks, it is essential to quantify the distribution of class labels at each split. This is done using entropy, a measure of impurity (Breiman, 2001):

$$H(S) = - \sum_{j=1}^C p_j \log_2 p_j$$

where (C) is the number of classes and (p_j) is the proportion of samples in class (j) . When a node is split using a feature $\backslash(A\backslash)$, the expected entropy across the resulting branches is:

$$E[H(A)] = \sum_{i=1}^K P(i) \cdot H(D_i)$$

Where:

K is the number of branches after the split,
 $P(i)$ is the proportion of samples in branch(i),
 $H(D_i)$ is the entropy of branch(i)

The information gain from splitting on feature (A) is:

$$IG(A) = H(S) - \sum_{i=1}^K \frac{N_i}{N} H(D_i)$$

where $(H(S))$ is the entropy before the split, (N_i) is the number of samples in branch (i) , and (N) is the total number of samples at the parent node.

These metrics guide the decision tree in choosing the optimal feature and threshold to maximize class separation at each split.

2.2.2 Random Forests and Generalization

The Random Forest algorithm is an ensemble learning method that aggregates the predictions of multiple decision trees to improve classification performance. In the context of this study, Random Forests were used to classify drought severity based on meteorological data. By constructing multiple trees and introducing randomness in their structure, the model captures diverse patterns in the data, increasing its robustness and accuracy—particularly valuable for high-dimensional, noisy datasets such as those used in climate and environmental modeling.

Each tree in a Random Forest is trained on a different bootstrap sample of the original data. Furthermore, at each split, a random subset of features is considered, which decorrelates the trees and reduces the risk of overfitting. This randomness introduces model diversity, which is crucial to the success of ensemble methods (Ryu et al., 2019).

Mathematically, the ensemble of classifiers can be represented as:

$$\mathcal{H} = \{h_1(x), h_2(x), \dots, h_T(x)\}$$

where each $(h_k(x))$ is a decision tree classifier trained on a randomly drawn subset of the data and features. For a given input (x) , the final prediction of the Random Forest is obtained through majority voting across all trees:

$$\hat{y} = \arg \max_c \sum_{k=1}^T 1[h_k(x) = c]$$

Here, $(1[\cdot])$ is the indicator function, returning 1 if the predicted class equals (c) , and 0 otherwise. Alternatively, the forest can output a probabilistic estimate for each class:

$$P(c | x) = \frac{1}{T} \sum_{k=1}^T 1[h_k(x) = c]$$

In this research, we implemented a Random Forest classifier to predict drought severity levels from six target classes based on 19 meteorological indicators. Despite using only four trees due to computational constraints, the ensemble outperformed individual decision trees in terms of generalization and overall classification performance.

Convergence and Generalization Behavior

One of the theoretical strengths of the Random Forest lies in its convergence properties. In classification tasks such as ours, this behavior can be captured by the **margin function**, which evaluates the difference between the vote count for the correct class and the most popular incorrect class. The margin function is defined as:

$$mg(x, y) = E_{\theta}[1(h(x; \theta) = y)] - \max_{j \neq y} E_{\theta}[1(h(x; \theta) = j)]$$

Where :

(x) is the input instance,

(y) is the correct class label,

$(h(x; \theta))$ is a tree constructed with random parameters (θ) ,

(E_{θ}) represents the expected value across the distribution of trees.

The generalization error of the forest, denoted (PE^*) , is defined as the probability that the margin is negative:

$$PE^* = P_{x,y}(mg(x, y) < 0)$$

This expression quantifies the likelihood that the ensemble classifier misclassifies an input. A positive margin implies a correct and confident prediction, while a negative margin indicates confusion among classes and results in a misclassification.

Breiman (2001) showed that, under mild assumptions, the generalization error decreases and stabilizes as the number of trees $(T \rightarrow \infty)$. Specifically, the aggregated prediction converges to the expected value (Ryu et al., 2019):

$$\lim_{T \rightarrow \infty} P(c | x) = E_{\theta}[1(h(x; \theta) = c)]$$

This convergence is a consequence of the law of large numbers, which ensures that the randomness of individual trees averages out across the ensemble. As more trees are added, their uncorrelated errors cancel out, leading to more stable and accurate predictions.

In our implementation, although we limited the forest to four trees due to computational constraints, the model still benefited from the convergence principle. Even a small ensemble improved the classification robustness, reduced overfitting, and produced more balanced predictions across the drought severity classes. This theoretical advantage makes Random Forest an ideal choice for classification problems in environmental and site exploration applications, where variable interactions are nonlinear and high-dimensional.

3. Results and Findings

3.1 Model Development and Implementation

The drought classification model was developed and implemented in Python, utilizing a custom-coded framework to construct and evaluate both decision tree and Random Forest algorithms from first principles. The implementation included modules for data preprocessing, model construction, and evaluation, adhering closely to the theoretical foundations outlined in earlier sections.

Preprocessing involved removing rows with missing values and discretizing continuous drought severity scores into integer classes ranging from 0 (no drought) to 5 (exceptional drought). This categorical representation aligns with the U.S. Drought Monitor classification scheme and enabled a supervised multi-class classification setup.

The Random Forest algorithm was implemented using an ensemble of decision trees, each trained on a different bootstrap sample of the data. At each node within a tree, a randomly selected subset of meteorological features was evaluated to determine the optimal split, enhancing diversity and reducing overfitting.

The dataset used comprised approximately 3 million records, each with 19 distinct meteorological features. Features such as wind speed, air temperature, humidity, and precipitation were included to capture the temporal and spatial variability of drought-relevant climate signals. Preprocessing excluded metadata (e.g., date, location IDs), focusing exclusively on the predictive indicators.

3.2 Assumptions and Classification Labels

Unlike many statistical models, Random Forests are non-parametric and require fewer distributional assumptions. However, several design decisions were made to adapt the model to this application.

First, it was assumed that the meteorological features included in the dataset were sufficiently informative to allow effective partitioning of the feature space across the drought severity classes. Second, the transformation of the drought severity scores into discrete integer categories assumed that rounding the continuous values closely reflected meaningful class distinctions. This assumption supported the use of classification metrics while preserving interpretability in the context of drought forecasting.

3.3 Evaluation Metrics

To evaluate model performance, standard classification metrics were computed, including accuracy, macro-averaged precision, macro-averaged recall, and F1 score (Saito & Rehmsmeier, 2015). These metrics are derived from the confusion matrix, using the following definitions :

$$\text{Accuracy} = \frac{\text{Number of correct predictions}}{\text{Total number of predictions}}$$

$$\text{Precision}_{\text{macro}} = \frac{1}{N} \sum_{i=1}^N \frac{TP_i}{TP_i + FP_i}$$

$$\text{Recall}_{\text{macro}} = \frac{1}{N} \sum_{i=1}^N \frac{TP_i}{TP_i + FN_i}$$

$$\text{F1}_{\text{macro}} = \frac{1}{N} \sum_{i=1}^N 2 \cdot \frac{\text{Precision}_i \cdot \text{Recall}_i}{\text{Precision}_i + \text{Recall}_i}$$

The F1 score is particularly important in imbalanced multi-class classification, where both false positives and false negatives affect overall model quality.

3.4 Experimental Results

Model training and evaluation were conducted on a machine with an Intel Core i7 processor and 16 GB of RAM. Processing the full dataset—comprising nearly 3 million rows and 57 million feature values—took approximately 47 hours. The Random Forest classifier achieved an overall accuracy of 0.75, indicating correct predictions for 75% of the test instances. However, the macro-averaged precision was 0.22, the macro-averaged recall was 0.18, and the F1 score was 0.20. These values reflect the model's limited sensitivity to minority classes, despite its high accuracy on the dominant class.

The classification outcomes are further illustrated in two confusion matrices. Figure 2 shows the confusion matrix normalized by each row (true class), which highlights prediction distribution per drought class. As seen, false predictions into Class 0 were dominant across all rows, with misclassification rates exceeding 90% for Classes 1 through 5 in some cases. For example, 72.74% of true Class 5 instances were incorrectly labeled as Class 0.

Figure 3 presents the confusion matrix as percentages of the entire test set. It reveals that 74.4% of all test predictions were for Class 0, clearly showing the model's bias toward the majority class. In contrast, severe drought levels (Classes 4 and 5) accounted for less than 1% each of the total predictions.

These results demonstrate that while the Random Forest model performs well in predicting the majority class, it struggles to generalize across less frequent drought categories. This imbalance, likely due to the skewed distribution in the dataset, limits the model's practical utility for detecting moderate to severe drought events. Future improvements may include oversampling, stratified training, class weighting, or hybrid ensemble strategies to improve detection of rare events without sacrificing overall performance.

4. Conclusion

This study explored the application of a custom-implemented Random Forest model for classifying drought severity using high-dimensional meteorological data. While the model achieved an overall accuracy of 0.75, macro-averaged precision (0.22) and recall (0.18) metrics revealed significant limitations in its ability to detect moderate to severe drought conditions. The confusion matrices confirmed a strong prediction bias toward Class 0 (no drought), which was traced to the underlying class imbalance in the dataset. Despite these limitations, the study successfully demonstrated the feasibility of using ensemble-based machine learning techniques for environmental classification tasks and contributed a transparent, ground-up implementation of Random Forest theory applied to a real-world challenge. In addition, this work lays the groundwork for future integration into subsurface monitoring frameworks—particularly those involving autonomous environmental sensing or geotechnical exploration, where climatic data influences soil state variables critical for model-driven robotic interaction.

Addressing class imbalance through sampling strategies, cost-sensitive learning, or ensemble calibration, combined with more focused feature selection, presents a viable path toward improving model robustness and generalization. Moreover, expanding the training scope with diverse regional datasets may support broader applicability in global drought prediction contexts.

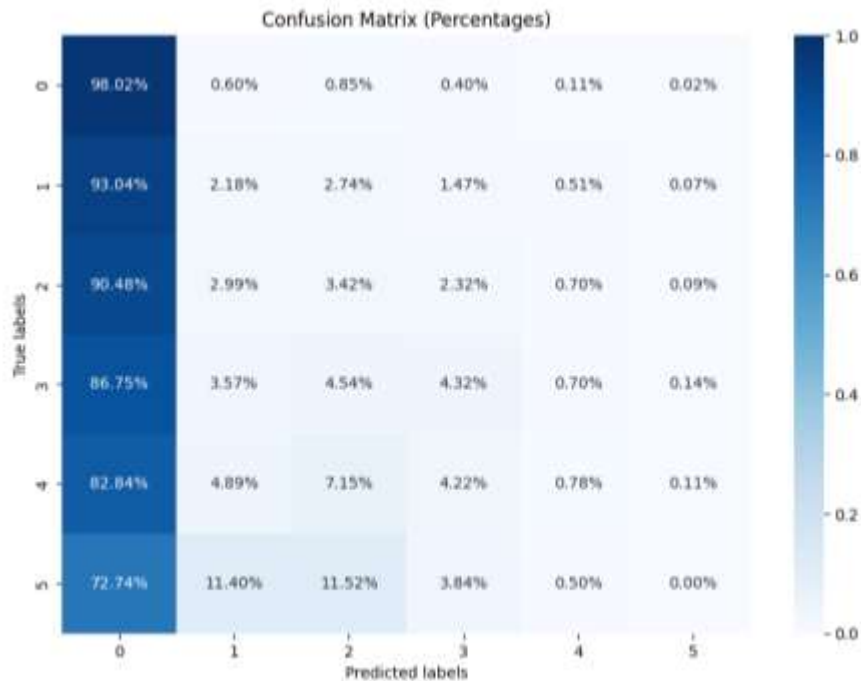


Figure 2: Confusion matrix showing classification results as row-normalized percentages. Each cell represents the percentage of true instances of a given class (y-axis) that were predicted as another class (x-axis). The diagonal indicates correct class

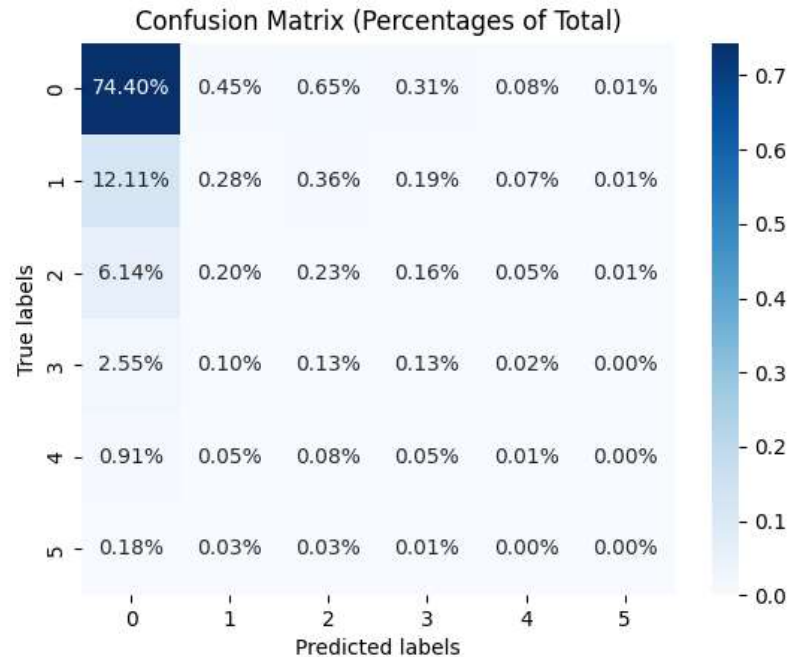


Figure 3: Confusion matrix displaying classification results as percentages of the total number of test samples. This representation emphasizes the dominance of Class 0 predictions, which account for 74.4% of all classifications, revealing the model's bias toward the majority class in the imbalanced dataset.

Funding: This research received no external funding

Conflicts of Interest: The authors declare no conflict of interest.

Publisher's Note: All claims expressed in this article are solely those of the authors and do not necessarily represent those of their affiliated organizations, or those of the publisher, the editors and the reviewers.

References

- [1] Azizi, S., & Azizi, T. (2024). Classification of drought severity in contiguous USA during the past 21 years using fractal geometry. *EURASIP Journal on Advances in Signal Processing*, 2024(1), 3. <https://doi.org/10.1186/s13634-023-01094-z>
- [2] Breiman, L. (2001). Random Forests. *Machine Learning*, 45(1), 5–32. <https://doi.org/10.1023/A:1010933404324>
- [3] Chang, F.-J., Chang, L.-C., & Chen, J.-F. (2023). Artificial Intelligence Techniques in Hydrology and Water Resources Management. *Water*, 15(10), Article 10. <https://doi.org/10.3390/w15101846>
- [4] Criminisi, A., Shotton, J., & Konukoglu, E. (2012). Decision Forests: A Unified Framework for Classification, Regression, Density Estimation, Manifold Learning and Semi-Supervised Learning. *Foundations and Trends® in Computer Graphics and Vision*, 7(2–3), 81–227. <https://doi.org/10.1561/06000000035>
- [5] Data. (n.d.). Google Drive. Retrieved March 28, 2025, from <https://drive.google.com/drive/u/3/folders/1UUCqKYaYmBcDgLWv-unsdrMFbDSO-6Jv>
- [6] Nguyen, T., & Minh, D. (2021). Applying Machine Learning to predict Hand-Foot-Mouth disease outbreaks in Vietnam. *Journal of Health Informatics in Developing Countries*, 15(2), Article 2. <https://www.jhidc.org/index.php/jhidc/article/view/300>
- [7] Quintana, J. R., Martín-Sanz, J. P., Valverde-Asenjo, I., & Molina, J. A. (2023). Drought differently destabilizes soil structure in a chronosequence of abandoned agricultural lands. *CATENA*, 222, 106871. <https://doi.org/10.1016/j.catena.2022.106871>
- [8] Ryu, J., Kwon, G., Yang, M.-H., & Lim, J. (2019, September 25). *Generalized Convolutional Forest Networks for Domain Generalization and Visual Recognition*. International Conference on Learning Representations. <https://openreview.net/forum?id=H1lxVyStPH>
- [9] Saito, T., & Rehmsmeier, M. (2015). The Precision-Recall Plot Is More Informative than the ROC Plot When Evaluating Binary Classifiers on Imbalanced Datasets. *PLOS ONE*, 10(3), e0118432. <https://doi.org/10.1371/journal.pone.0118432>
- [10] Shahhosseini, S., Parekh, M., & Tao, J. (2023). DEM-MBD Coupled Simulation of a Burrowing Robot in Dry Sand. *Geo-Congress 2023*, 309–317. <https://doi.org/10.1061/9780784484692.032>
- [11] Shahhosseini, S., & Tao, J. (2025). Bio-Inspired Dual-Auger Vertical Self-Burrowing Robot: DEM-MBD Analysis of Downward Penetration in Granular Media. *Geotechnical Frontiers 2025*, 116–125. <https://doi.org/10.1061/9780784485989.013>