



Innovative Machine Learning Approaches for Evaluating Climate Change Vulnerabilities of SMEs

Ubaid Ullah^{1,*}, Alamgir Khalil¹, Abdur Rehman¹, Sami Ullah², Sehran Hassan¹ and Muhammad Sani³

¹Department of Statistics, University of Peshawar, Peshawar 25120, Pakistan

²School of Mathematics, Statistics and Mechanics, Beijing University of Technology, Beijing 100124, China

³Department of Computer Science, University of Lakki Marwat, Lakki Marwat 28420, Pakistan

Abstract

This paper examines the vulnerability of Small and Medium-sized Enterprises (SMEs) exposed to evolving climate changes in Pakistan, specifically the impacts of extreme weather events, including floods and drought. The earlier literature illustrates that SMEs are affected by climate-related risks, but the current study takes the discussion further by implementing machine learning algorithms to measure the vulnerabilities of SMEs more objectively. A mixed-methods design was used to combine surveys with machine-learning techniques. PyCaret was employed to tune instruments such as Logistic Regression (LR), Random Forest (RF), ordered logistic regression, LightGBM, ADA Boost, SVM, KNN, GBC, and Naïve Bayes, among others, in Python. The results show that SMEs located in rural regions are significantly more vulnerable to climate events than SMEs in urban areas. The best results were found by the Logistic Regression model: its accuracy was 63.3%, recall 63.3%, F1-score 55.2%, and Kappa value 0.23. Random Forest

showed similar findings, but Logistic Regression proved stronger in diagnosing the influence of climate on business activity. Ordered logistic regression-based analyses show that awareness of temperature fluctuation ($p = 0.035$) and perceived financial impact ($p = 0.018$) significantly affect the degree to which firms are affected by climate-related impacts, while geographic locations ($p = 0.057$) and business tenure ($p = 0.091$) are only marginally important. Strict evaluation of empirical evidence has pointed out that the application of focused financial tools, extensive infrastructure, and governance framework transformation can help to make the small- and medium-sized enterprises, especially enterprises located in remote areas to be more resilient to climate-related risks. Further studies may improve this analysis through a more detailed machine-learning framework and analysis of sector-specific effects of climate change in specific geographic areas.

Keywords: climate change, small and medium enterprises (SMEs), machine learning, logistic regression, PyCaret, climate risk.



Submitted: 26 June 2025

Accepted: 28 August 2025

Published: 07 November 2025

Vol. 1, No. 4, 2025.

10.62762/TACS.2025.395911

*Corresponding author:

✉ Ubaid Ullah

ubaidullahyousafzai5683@gmail.com

Citation

Ullah, U., Khalil, A., Rehman, A., Ullah, S., Hassan, S., & Sani, M. (2025). Innovative Machine Learning Approaches for Evaluating Climate Change Vulnerabilities of SMEs. *ICCK Transactions on Advanced Computing and Systems*, 1(4), 275–290.



© 2025 by the Authors. Published by Institute of Central Computation and Knowledge. This is an open access article under the CC BY license (<https://creativecommons.org/licenses/by/4.0/>).

1 Introduction

Climate change has emerged as one of the most serious global issues with massive implications for the developing nations such as Pakistan. Like many developing countries, Pakistan's own contribution to global greenhouse gas emission is not high but being highly susceptible to the impacts of climate change, it cannot be insulated from the global consequences of climate change. Pakistan faces increasingly intensified and more frequent extreme weather events, including floods, droughts, and erratic monsoons [1]. In 2022, recoiling from record-breaking floods amplified by the melting of glaciers and driven by incessant monsoon rains, millions have been uprooted, and infrastructure, homes and livelihoods have been ravaged. This underscores the necessity for change in measures that can be adapted to address the adverse effects of climate change [2]. Geographical exposure of mountainous Pakistan piled up the risk of man and nature including hazards of GLOFs to human habitation and ecology.

While Pakistan is a low emitter of greenhouse gases on a global scale, it is significantly adversely affected by the fallout of climate changes on account of a multitude of internal issues such as rapid industrial growth, reduction and loss of tree cover and dependency on fossil fuel energy. These pressures place added stresses on the country's environmental settings, making it vulnerable to natural disasters that endanger food security, water availability, and economic stability. Although different initiatives/policies have been implemented and promoted to address climate change, these are difficult to be implemented owing to poor governance, absence of infrastructure, and limited finances. This adds more importance to examining the adaptability of SMEs, which constitutes a large portion of the national economy and employment, particularly in rural areas. SMEs play an important part in a healthy economy but are particularly susceptible to climate-related hazards that can derail their operations as well as their market stability [3].

The literature addressing climate change effects on SMEs stresses on a wide range of sectorial specific vulnerabilities. Agricultural SMEs, for example, encounter problems including soil erosion, water shortages, and traditional agricultural activities that limit productivity and retard progress [4]. Likewise, SME descriptions show existing high operational costs due to the lack of infrastructure, poor accessibility to finances, and also the lack of an efficient tax system. In addition, a larger part of these challenges is compounded by climate change-related disruptions

such as floods, droughts and extreme weather events, which cause serious financial instability more so in rural areas where in majority the SMEs are located. The tourism industry is seen as a highly promising one, but is hampered by security concerns, poor infrastructure, as well as gender discrimination, thereby creating obstacles for SMEs in these areas. The year-in, year-out pressure to compete on the global market also has ramifications for Pakistan's own SME sector, particularly as it is influenced by global supply chains. With growing international and local vulnerability, through which risks of the global climate change become more acute, companies are increasingly faced with non-localizable (both locally and internationally specific) challenges that demand adaptive strategies that surpass local environmental challenges.

- This study proposes the use of advanced machine learning methods, namely Logistic Regression, Random Forest, Ordered Logistic Regression, LightGBM, ADA Boost, SVM, KNN, GBC, Naïve Bayes, and many more to analyze the climate vulnerability of Small and Medium-sized Enterprises (SMEs) in Pakistan, providing a quantitative, data-driven measure of climate vulnerability.
- Automated machine learning is utilized using PyCaret, which allows hyperparameter tuning and model selection in an efficient way. The models produced demonstrate improved performance and simplified computational procedures.
- A broad-based approach is employed for model evaluation, involving accuracy, recall, F1-score, and Matthews Correlation Coefficient (MCC), offering a robust system for climate-effect forecasting.
- We evaluated feature importance in the Random Forest model to highlight the most significant determinants of climate impact on business vulnerability, including socio-economic variables (Q6, Q7, Q14), thereby providing detailed information on sector-specific weaknesses.
- The proposed work combines the strengths of computational modeling and empirical survey studies, offering a comprehensive perspective on sector vulnerabilities and enabling the development of specific adaptation measures for SMEs across different regional environments.

Figure 1 illustrates the conceptual framework of

our machine learning approach for climate change vulnerability assessment.

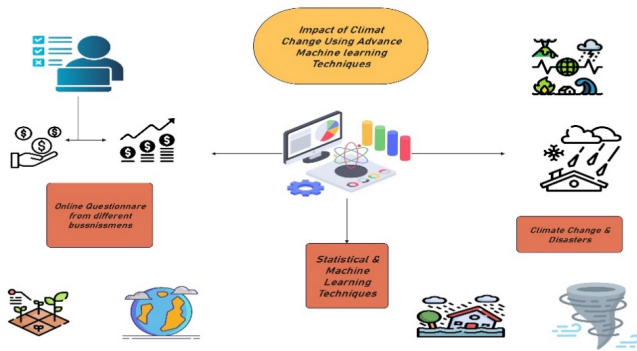


Figure 1. Conceptual framework illustrating the application of advanced machine learning in climate change vulnerability assessment for SMEs.

2 Related Works

Small and Medium sized Enterprises (SMEs) are now considered to be very susceptible to the effects of the changing climate in developing nations such as Pakistan. Extreme weather such as floods, droughts and unreliable rain discourages their operations and destabilizing markets. Such impacts are especially harsh on the rural SMEs, where most of the time they are not capable of adjusting to these impacts. According to recent studies, there is a trend of utilizing machine learning (ML) in analyzing and forecasting these weaknesses with a view of informing more effective adaptation measures [5].

The use of machine learning models has already shown to be a potent tool to understand and forecast the effects of climate change on the SMEs in the regions where data access and its environmental properties can change significantly. The paper [6] established that ensemble learning methods such as Random Forest and Gradient Boosting Machines (GBM) could be effective in forecasting climate vulnerability, enabling businesses to prepare proactively. The models are especially useful given that they are capable of dealing with large and sophisticated datasets incorporating both socio-economic conditions and climate in a way that provides a more textured understanding of the vulnerabilities encountered by SMEs.

In their research [7] integrated machine learning algorithms with a fuzzy Analytic Hierarchy Process to assess the livelihood vulnerability of coastal communities in southwestern Bangladesh exposed to climate-related disasters including floods, cyclones, and sea-level rise. By developing

a Livelihood Vulnerability Index and generating spatial vulnerability maps, the study demonstrated that socio-economic factors such as income diversification, water access, and education level are dominant predictors of climate vulnerability. The findings confirmed that machine learning classifiers substantially outperform conventional indicator-based methods in capturing the complex, non-linear relationships between climate hazards and community resilience, with cross-validated models showing reliable transferability across climate-sensitive coastal regions.

A further line of research by [8] applied multiple machine learning algorithms to predict how firms respond to climate change risk, using a sample of 168 financial institutions across 27 countries. Comparing 12 classifiers including AdaBoost, Gradient Boosting, and XGBoost, the study found that ensemble methods achieved the highest precision in predicting firm-level adoption of climate-risk mitigation practices. The results further revealed that physical and regulatory climate shocks are the most significant drivers of firm responses, underscoring the value of machine learning in translating climate exposure data into actionable risk classifications for enterprise decision-making.

Similarly, [9] conducted a comprehensive assessment of climate change impacts on agricultural agribusiness, mapping the effects of temperature rise, rainfall variability, and CO₂ emissions on the productivity of cereal, pulse, and oilseed crops across six decades of data. Using a Tobit regression model and Principal Component Analysis within an exposure-sensitivity-adaptive capacity framework, the study found that increasing temperatures reduced cereal yields by 26% and oilseed yields by 35%. The findings highlight the acute vulnerability of agricultural SMEs to climate-driven disruptions in crop production, water resources, and soil quality, providing sector-specific evidence that informs targeted adaptation interventions.

The in-built element of fairness into machine-learning models of climate-risk assessment has received the same divination by scholars. As [10] noted, the biases possible in the modeling process may result in an inequitable behavior, especially in terms of SMEs that are marginalized. As they have analyzed on the topic of fairness in machine learning, there is a need in imposing fairness limitations into predictive models in order to ensure that the distributive impacts of climate change are assessed across heterogeneous

sectors, regions, and socio-economic groups.

The implementation of the machine-learning approaches into SMEs to predict climatic hazards is also limited by particular issues, including data accessibility and data quality. A majority of SMEs that are located in rural or under-developed areas do not have the necessary data infrastructure to enable proper assessments of climate risks, thus undermining the performance of machine-learning models [5]. Subsequent studies must thus focus on development of inexpensive data-gathering methods and the ability of data to be accessible in a way that all SMEs can take advantage of them.

In summary, the application of machine learning methodologies to detect and address climate change vulnerabilities in SMEs has enormous potential. The realization of superior algorithms (such as Random Forest, LightGBM, and hybrid ones) already proved significant effects in the climate risks forecasts and identifying vulnerable SMEs. Furthermore, through the inclusion of considerations of fairness and the inclusion of a liability on data deficiencies, machine learning can tune adaptation approaches to the specific requirements of SMEs, creating the potential to induce the long-term sustainability. Further refinement of these models and their integration into national climate policies can be considered an important pathway toward enhancing SME climate change resilience.

3 Methodology

This study examines the impacts of climate change on small and medium enterprises (SMEs) in Pakistan by adopting a mixed-methods approach that integrates expert opinions gathered through structured surveys with computational analysis using machine learning techniques. The objective is to provide insights into the susceptibility of SMEs to climate-related impacts and to leverage advanced machine learning for classification and prediction. Data engineering methodologies were employed for preprocessing and transformation, preparing the data for robust and reproducible modeling.

3.1 Research Design

The research pursues an exploratory design examining the susceptibility of small and medium-sized enterprises (SMEs) to climate change, with particular emphasis on the constraints imposed by floods and droughts. The perceived consequences of the extreme weather events regarding operations of

the organization, risk management practices and adoption of adaptation strategies are measured with a purpose-built Likert-scale survey. Machine learning integration was done by transforming the responses in the questionnaire into numerical values. The principles of model explainability, equity, and predictive fairness underpin this study, ensuring that the derived models are interpretable and actionable [11].

3.2 Data Collection and Engineering

In order to achieve the goals of this research, well-constructed survey questionnaire was used as the main data-gathering tool. Composed in both English and Urdu, it has been framed in such a way that it becomes inclusive to the regions of Pakistan that are diversified. The survey was set in 4 conceptualized parts; Business profiles, perceptions on risk, climate change effects, and social-economic. A 5-point Likert scale was employed to record the degrees of agreement within the respondents, and therefore the analysis was made standard and countable. It was distributed via email and WhatsApp, gaining accessibility and anonymity of the message sender, and data security at the same time [12, 13]. This approach is consistent with established survey practices in climate change research in Pakistan, where digital and social media channels have been shown to effectively reach diverse respondent populations while maintaining data integrity [13].

Data preprocessing was carried out using Excel and Python. Missing values were handled using mean imputation [14], and features were normalized using Min Max Scaler. Categorical variables were encoded numerically through label encoding, and dimensionality reduction techniques were applied to improve model performance [15]. The dependent variable—Extreme weather events (e.g., floods, droughts) have affected my business—represents SMEs' perceptions of climate impact and was modeled as a multiclass classification problem with three categories: low, medium, and high impact [16].

3.3 Analytical Approach and Model Selection

The overall machine learning workflow for climate risk assessment, as shown in Figure 2, involves multiple stages from data preprocessing to model evaluation. This study employs 15 machine learning classifiers in total. Among them, three core models are introduced in detail — Logistic Regression (LR), Random Forest (RF), and Ordered Logistic Regression (OLR) —

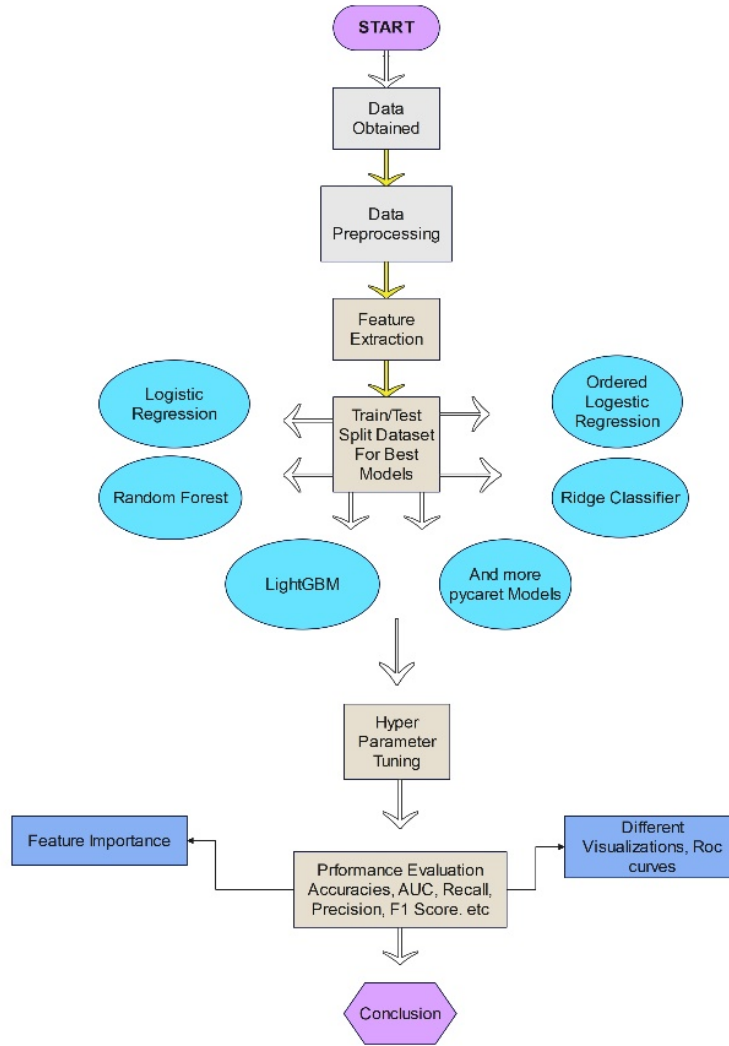


Figure 2. Workflow for machine learning based climate risk assessment of SMEs.

selected based on their interpretability, predictive performance, and ability to accommodate the ordinal nature of the target variable [17]. The remaining classifiers are described subsequently.

3.4 Logistic Regression (LR)

Logistic regression is an origination base statistical model of binary and multi-classification. This model allows predicting a categorical response by fitting the log-odds of the class membership as a linear spin off of the input features. When used on multiclass problems, involving K different classes, then the probability of a particular class can be estimated by modelling the respective cumulative probability, normalizing out natural ordering of the categories. The probability that an instance belongs to class K is given by the softmax function:

$$P(y = k | X) = \frac{e^{\beta_k^T X}}{\sum_{j=1}^K e^{\beta_j^T X}} \quad (1)$$

where β_k is the coefficient vector for class K and X is the feature vector.

3.5 Random Forest (RF)

Random forest serves as an ensemble-based learning technique commonly used in high-dimensional classification. Decision trees are also grown randomly on a bootstrapped sample of the original training data, and averaged to increase both accuracy and consistency of prediction. Reduction of variance is achieved by Bagging (bootstrap aggregation) and introduction of randomness in feature selection in the construction of trees deepens generalization. The last categorization is implied by a plurality of the trees [18, 19].

$$f(x) = \text{mode}\{h_1(x), h_2(x), \dots, h_B(x)\} \quad (2)$$

3.6 Ordered Logistic Regression (OLR)

Ordered logistic regression (OLR) is a statistical model that is used to predict the ordinal outcome-a

categorical response that has known order, but unknown distance (e.g. rating of customers like low, medium or high). Unlike the standard logistic regression, OLR estimates the cumulative probabilities that given the response variable will be at or below a certain category. Assumption of proportional odds guarantees adjacent outcome category relationships and make OLR especially appropriate in the survey analysis, health related outcome, and social-science studies [20].

$$\log \left(\frac{P(Y \leq j)}{P(Y > j)} \right) = \theta_j - \beta^T X, \quad j = 1, 2, \dots, J-1 \quad (3)$$

where θ_j represents the threshold for category j , β is the vector of regression coefficients, and X denotes the predictor variables; to further refine the classification predictions, several additional models were evaluated.

3.7 Quadratic Discriminant Analysis (QDA)

The flexibility in quadratic discriminant analysis (QDA), compared to linear discriminant analysis (LDA), lies in the fact that, rather than imposing a single common multivariate covariance matrix on an entire population, it assumes a different covariance matrix applies to each subgroup classified. This allows the modelling of more complicated and non-linear decision boundaries. QDA predicts the means and covariances in each class, then assigns a new observation to that class with the highest posterior probability (likelihood $\times$ prior), assuming multivariate normality within each class.

$$\delta_k(x) = -\frac{1}{2} \log |\Sigma_k| - \frac{1}{2} (x - \mu_k)^T \Sigma_k^{-1} (x - \mu_k) + \log \pi_k \quad (4)$$

where $\delta_k(x)$ is the discriminant function for class k , μ_k is the mean vector of class k , Σ_k is the covariance matrix of class k , π_k is the prior probability of class k , and $|\Sigma_k|$ denotes the determinant of Σ_k ; the class corresponding to the highest $\delta_k(x)$ value is selected [21].

3.8 Light G B M (Light Gradient Boosting Machine)

Light G B M is an optimal implementation of the gradient boosting framework towards fast, efficient and scalable learning, especially to large scale datasets. Its use of depth restrictions combined with the

strategy of leaf-wise tree-growth and a combination with histogram-based decision-tree learning makes it require very few computations and memory requirements. As a result, the framework is very applicable in both classification as well as regression. Light G B M builds an additive model, that is, optimizes a loss function through gradient descent, in a forward, stage-wise way; the overall prediction is then the sum of all the sequentially trained trees [22].

$$\hat{y} = \sum_{m=1}^M f_m(x), \quad f_m \in \mathcal{F} \quad (5)$$

where $f_m(x)$ is the prediction from the m th decision tree, $\mathcal{F}$ denotes the space of regression trees, and M represents the total number of boosting rounds.

3.9 Gradient Boosting (GBC)

Gradient Boosting Classifier (GBC) is a state of art ensemble learning framework where the models are implemented sequentially: the hope is that a successive model will reduce the errors made by the preceding one. GBC aims at to minimize a differentiable loss function, via gradient descent, with each new model being optimized with regard to the negative gradient of the loss as a function of the current prediction. GBC is especially useful on structured data: it is flexible and has good predictive performance [23].

$$F_t(x) = F_{t-1}(x) + \eta \cdot h_t(x) \quad (6)$$

where $F_t(x)$ is the updated prediction at iteration t , $F_{t-1}(x)$ is the prediction from the previous iteration, $h_t(x)$ is the new weak learner (typically a decision tree) fitted to the negative gradient of the loss, and η is the learning rate controlling the contribution of each learner.

3.10 AdaBoost (Adaptive Boosting)

The AdaBoost is a different type of ensemble algorithm, and it combines several poor classifiers, usually decision stumps, into one, more elaborate, classifier. The algorithm dynamically updates the weights of data points incorrectly classified during previous iterations thus allowing the sequence to learn in each iteration based on previous errors, and to attain better accuracy especially on hard examples. Finally, the ensemble forecast is obtained as a weighted majority vote of all weak learners where the more accurate models have a high weight [24].

$$H(x) = \text{sign} \left(\sum_{t=1}^T \alpha_t h_t(x) \right) \quad (7)$$

Define $h_t(x)$ as the weak learner at iteration t , α_t as the weight placed on $h_t(x)$ proportional to the accuracy of that learnt, T as the overall number of iteration and $\text{sign}()$ as the predicted class as provided by the weighted sum.

3.11 Support Vector Machine (SVM)

Support Vector Machine (SVM) is a supervised learning algorithm and it majorly falls under the binary classification algorithm. It builds a hyper-plane that maximizes the margin between the two classes, the maximum margin being the distance perpendicular between the plane and the point on each class closest to the plane, the so-called support vectors. This construction empowers the model to have a broader generalization capacity. In those cases where the data can be linearly separated, SVM establishes a convex optimization problem to establish the optimal separating hyper-surface [25].

$$\begin{aligned} f(x) &= \text{sign}(w^T x + b) \\ y_i (w^T x_i + b) &\geq 1, \quad \forall i \end{aligned} \quad (8)$$

where w is the weight vector orthogonal to the hyper plane, b is the bias term, x is the input feature vector and y is the class they are labelled to, $y_i \in \{-1, +1\}$ is the class label, and $f(x)$ predicts the class based on the sign of the linear function.

3.12 K-Nearest Neighbors (KNN)

K-Nearest Neighbors (KNN) is a non-parametrical algorithm that is straightforward, but can be very powerful at the same time. It determines the nearest subset of k data points in the feature space, called its k -nearest neighbors, and the class that the majority of them falls into is the one assigned to a newly added data point. KNN is highly flexible in that it does not have assumptions concerning the underlying data distribution hence is highly applicable in a large range of classification tasks. It depends on optimum selection of the values of k (and distance metric) to be effective in practice [26].

$$\hat{y} = \text{mode} \{y_i \mid x_i \in \mathcal{N}_k(x)\} \quad (9)$$

Denote the set of the k nearest neighbors of x by $\mathcal{N}_k(x)$, the class label y_i of neighbor x_i by y_i and the most frequent class label among the neighbors by the mode.

3.13 Ridge Classifier

Ridge Classifier is a linear form of classification which uses L2 regularization to reduce the over-fitting in order to enhance the level of generalization. Its functionality is similar to that of Ridge Regression but is adapted to classification, i.e. applying a regression output threshold to classification. The L2 penalty reduces the coefficients in a penalized proportion law to the value of their magnitude; the penalty therefore limits the complexity of the model and solved multicollinearity [27].

$$\mathcal{L}(w) = \sum_{i=1}^n (y_i - w^T x_i)^2 + \lambda \|w\|_2^2 \quad (10)$$

where using formal notation, we write the weight vector as w , the feature vector of sample i th as x_i , and the numerically experienced label of a sample as y_i . Moreover, they consider a regularization parameter, which controls the strength of the penalty, and labeled it as λ , and the squared L2 norm (the one resulting in the coefficients squared sum) and referred to it as $\|w\|_2^2$.

3.14 Extra Trees (ET)

An improvement of Random Forests, Extra Trees generates ensemble models whose splits at each node are allocated randomly, resulting in comparatively lower variance at the expense of possible higher bias [21].

3.15 XGBoost

A gradient-boosting process optimized to approximate the second-order Taylor expansion of the loss function, thereby simultaneously minimizing bias and variance [23].

3.16 Decision Tree (DT)

Decision trees are simple-to-interpret models that divide the input space into regions via binary subdivisions based on feature values, minimizing measures of impurity such as Gini impurity or entropy [21].

3.17 LDA (Linear Discriminant Analysis)

Linear Discriminant Analysis (LDA) is a dimensionality-reduction method that maximizes

between-class variance while minimizing within-class variance, producing a linear projection that enhances class differentiation [21].

3.18 Naive Bayes (NB)

Naive Bayes Probabilistic classifier that uses Bayes theorem and makes an assumption that given a single feature, features in the input will be independent of each other, given the label of the class. Nonetheless, such a simple assumption, though it is called naive, proves efficient in most projects; text categorization and spam filtering are two examples of such applications using NB that calculates the posterior probability of each category and accepts the categorization with the highest likelihood [28, 29].

$$P(y | x) = \frac{P(y) \prod_{i=1}^n P(x_i | y)}{P(x)} \quad (11)$$

where $P(y)$ is the prior probability of class y , $P(x_i | y)$ is the conditional probability of feature x_i given class y , and $P(x)$ is the marginal probability of the input, often omitted in classification as it remains constant across classes.

3.19 Hyper-parameter Tuning and Model Evaluation

Model hyperparameter tuning was performed through Grid Search and Randomized Search performed by the PyCaret system, which, respectively, execute a search within the hyper-parameter space to find optimal parameter values yielding a peak performance. The model performance was measured through the 10-fold cross-validation process that has become common practice because it reduces the variances and helps to generalize. Testing would be performed based on a set of metrics comprising accuracy, precision, recall, F1-score, Cohen Kappa and the Matthews Correlation Coefficient (MCC). Among them, MCC is extremely interesting since it combines all the four elements of the confusion matrix and is not subject to the problem of class imbalance [30, 31].

$$MCC = \frac{TP \cdot TN - FP \cdot FN}{\sqrt{(TP + FP)(TP + FN)(TN + FP)(TN + FN)}} \quad (12)$$

Under these conditions, TP is the number of True Positives, TN is the number of True Negatives, FP is the number of False Positives and FN is the number of False Negatives.

3.20 Feature Importance and Fairness Analysis

Feature importance was estimated using the Random Forest model, where socio-economic factors (Q6, Q7, Q14) stood out as the most important predictors of climate impact on businesses. These variables were identified through the mean decrease in Gini impurity within the PyCaret framework. It combined two selection strategies, recursive feature elimination (RFE) and feature importance analysis. RFE sequentially removes the least informative features, while feature importance analysis quantifies the strength of each feature's association with the target variable. We also conducted a fairness analysis to examine bias across groups and applied corrective techniques to ensure equitable prediction [32].

3.21 Tools and Technologies Used

All analysis were conducted in Python using PyCaret, scikit-learn, pandas, matplotlib, and seaborn packages. Use of ROC curves, confusion matrix, and feature importance plot helped on explaining the model. All experiments were repeatable using the same random seed for reproducibility.

4 Experiments

The discussion that ensues lays out empirical findings that are related to the interaction between SME characteristics, perceived climate risks, and relevant adaptation strategies. To this effect, Ordered Logistic Regression was used to interrogate the role of industry sector, geographic location, perceived financial exposure and other factors in determining exposure of firms to extreme weather event as reported by the firms. The significance levels of each variable help identify key drivers of climate impact on businesses, offering insights into targeted areas for policy and adaptive interventions.

Table 1 summarizes the complete Ordered Logistic Regression results, providing coefficient estimates, statistical significance, and effect sizes for all predictor variables in the climate impact assessment model.

4.1 Interpretation of ordered logistic regression

The results found by the ordered logistic regression show that there are two predictor variables which are considered significant at p-levels. The geographic location (Q3) and firm age (Q4) were found to be marginally significant, which represent the possibility of influence of confounding factors. The threshold analysis in Table 2 indicates that the boundaries between medium and high impact

Table 1. Ordered logistic regression results on climate impact.

Variable	Coefficient	Std. Error	z-value	p-value	Significance
Q1	-0.817	0.570	-1.433	0.152	
Q2	2.074	2.761	0.751	0.453	
Q3	-2.520	1.323	-1.904	0.057	† (marginal)
Q4	-1.882	1.113	-1.690	0.091	† (marginal)
Q6	2.613	1.239	2.109	0.035	*
Q7	2.967	1.256	2.362	0.018	**
Q8	-0.604	0.719	-0.840	0.401	
Q9	-1.602	1.372	-1.167	0.243	
Q10	-1.040	0.851	-1.222	0.222	
Q11	0.669	0.928	0.720	0.471	
Q12	1.316	1.044	1.260	0.208	
Q13	-0.803	0.927	-0.866	0.387	
Q14	-0.208	0.460	-0.452	0.652	

categories (thresholds 2/3 and 3/4) are statistically well-defined ($p < 0.001$), while the transition between low and medium impact (threshold 1/2) lacks statistical significance ($p = 0.334$). The rest of variables showed no significance to reach typical limits of statistics; however, they can be considered potentially important in later studies anywhere.

Table 2. Model thresholds (Cut Points).

Threshold	Estimate	Std. Error	z-value	p-value	Significance
1/2	-7.346	7.607	-0.966	0.334	
2/3	2.154	0.387	5.564	0.000	***
3/4	1.527	0.403	3.785	0.000	***

Significance Codes: *** $p < 0.001$; ** $p < 0.01$; * $p < 0.05$; † $0.05 < p < 0.10$ (marginal)

4.2 Comprehensive Model Performance Evaluation

Table 3 presents a systematic comparison of all 15 machine learning classifiers evaluated in this study, providing a comprehensive overview of their performance across eight different evaluation metrics. This comparative analysis serves as the foundation for selecting the most appropriate algorithms for climate impact classification in SME vulnerability assessment.

4.3 Random Forest Classifier

As evidenced by the comprehensive comparison in Table 3, in the initial model comparison run, Random Forest achieved the highest predictive accuracy (57.5%) with balanced metrics on recall, F1-score, Kappa, and MCC. Although XGBoost recorded the highest AUC (0.271), Random Forest offered the best overall trade-off between accuracy and complexity, as confirmed by the confusion matrix (Figure 3). Following hyperparameter tuning with a fixed random seed (Table 4), Logistic Regression surpassed Random Forest with an improved accuracy of 63.33%, emerging as the best-performing model overall.

4.4 ROC Curves for the Random Forest Classifier

Figure 4 shows that the Random Forest classifier achieves AUC values of 0.94, 1, and 1 for classes 0, 1, and 2 respectively on the training set, with a micro-averaged and macro-averaged training AUC of 0.99. It should be noted, however, that the cross-validated test AUC was 0.1667 (Table 4), indicating a substantial gap between training and generalisation performance. This discrepancy is primarily attributable to the small sample size, which

Table 3. Comparative performances of classification models based on evaluation metrics.

Model	Accuracy	AUC	Recall	Precision	F1-Score	Kappa	MCC	Time (s)
Random Forest (RF)	0.575	0.250	0.575	0.508	0.497	0.196	0.237	0.181
Quadratic Discriminant (QDA)	0.550	0.000	0.550	0.325	0.403	0.000	0.000	0.031
LightGBM	0.550	0.200	0.550	0.325	0.403	0.000	0.000	0.035
Dummy Classifier	0.550	0.200	0.550	0.325	0.403	0.000	0.000	0.026
Gradient Boosting (GBC)	0.542	0.000	0.542	0.425	0.466	0.176	0.190	0.232
AdaBoost	0.525	0.000	0.525	0.428	0.452	0.047	0.070	0.088
Logistic Regression (LR)	0.492	0.000	0.492	0.494	0.462	0.077	0.109	1.260
SVM (Linear Kernel)	0.492	0.000	0.492	0.492	0.456	0.140	0.155	0.031
K-Nearest Neighbors (KNN)	0.483	0.194	0.483	0.386	0.396	0.010	0.025	0.050
Ridge Classifier	0.458	0.000	0.458	0.433	0.407	0.047	0.059	0.035
Extra Trees (ET)	0.450	0.179	0.450	0.406	0.396	0.050	0.056	0.151
XGBoost	0.442	0.271	0.442	0.319	0.349	0.061	0.065	0.060
Decision Tree (DT)	0.417	0.258	0.417	0.439	0.398	0.145	0.193	0.036
Linear Discriminant (LDA)	0.392	0.000	0.392	0.350	0.340	-0.043	-0.053	0.030
Naive Bayes (NB)	0.325	0.213	0.325	0.312	0.277	0.014	0.025	0.029

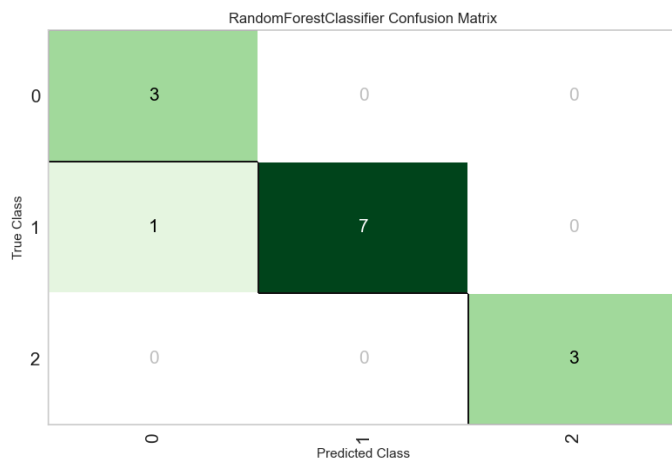


Figure 3. Confusion matrix of the Random Forest Classifier.

limits the model's ability to generalise reliably, and underscores the need for larger datasets in future studies.

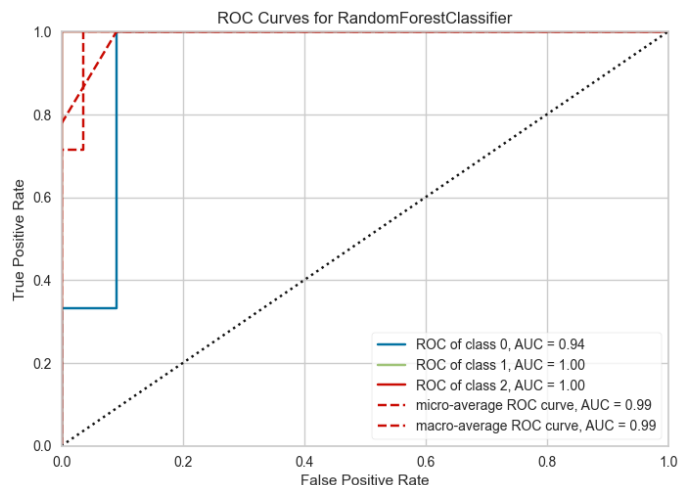


Figure 4. ROC curves for the Random Forest Classifier.

4.5 Feature Importance Plot

Figure 5 highlights that Q14, Q6, and Q7 emerged as the most influential predictors of climate impact on businesses, reflecting key aspects of vulnerability or adaptive capacity. The feature importance, calculated using the mean decrease in Gini impurity in PyCaret, ranks variables by their contribution to reducing classification uncertainty, offering valuable insight into the drivers of climate-related risks.

4.6 Comparison of Classifier Performance

Table 4 presents the optimized classifier performance results using PyCaret with fixed random seed, revealing significant performance improvements. It has been formed that the Logistic Regression (LR) is the best model, with better accuracy (63.3 %); better

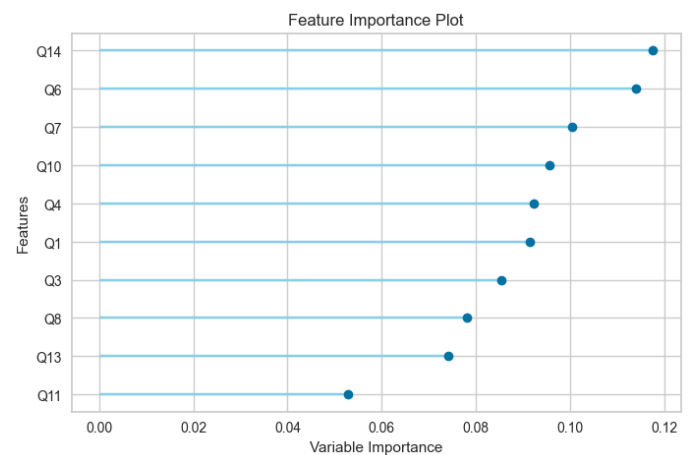


Figure 5. Feature importance plot for the Random Forest Classifier.

recall, F1-score, Kappa (0.230), and MCC (0.259). It worked better than more complicated versions like Random Forest and LightGBM. The second point is the comparable high accuracies acquired by the other classifiers indicate the dominating effect of data splitting on the performance of the small amount of data, and the comparatively high agreement with the ground-truth labels highlights both interpretability and robustness of LR.

4.7 Linear and Discriminant Model Performance Comparisons

Table 5 presents the performance of linear and discriminant classifiers. Logistic Regression achieved the highest accuracy (63.33%) and the strongest agreement metrics (Kappa = 0.230, MCC = 0.259) among all models in this group, confirming its superiority as an interpretable linear classifier for climate impact classification. Ridge Classifier followed with an accuracy of 56.67% (Kappa = 0.194, MCC = 0.200), while SVM (Linear Kernel) achieved 53.33% with the same agreement scores (Kappa = 0.194, MCC = 0.200), indicating that both linear margin-based methods produce meaningful class discrimination despite lower accuracy than Logistic Regression. Quadratic Discriminant Analysis matched Ridge Classifier in accuracy (56.67%) but yielded zero Kappa and MCC values, indicating no meaningful class discrimination beyond chance. Linear Discriminant Analysis performed poorest in this group (accuracy = 43.33%, Kappa = 0.024, MCC = 0.036), highlighting the limitation of linear projection methods when class boundaries are non-linear, as is typical in climate-impact survey data.

Table 4. Comparison of classifier performance (PyCaret, Seed = 786).

Model	Accuracy	AUC	Recall	Precision	F1-Score	Kappa	MCC	Time (s)
Logistic Regression	0.6333	0.0000	0.6333	0.5222	0.5522	0.2300	0.2587	0.7910
Ridge Classifier	0.5667	0.0000	0.5667	0.4444	0.4878	0.1943	0.2000	0.0250
Random Forest	0.5667	0.1667	0.5667	0.4833	0.5011	0.1900	0.1862	0.1590
Quadratic Discriminant	0.5667	0.0000	0.5667	0.3444	0.4233	0.0000	0.0000	0.0460
LightGBM	0.5667	0.1500	0.5667	0.3444	0.4233	0.0000	0.0000	0.0390
Dummy Classifier	0.5667	0.1500	0.5667	0.3444	0.4233	0.0000	0.0000	0.0240
SVM (Linear Kernel)	0.5333	0.0000	0.5333	0.4111	0.4511	0.1943	0.2000	0.0310
AdaBoost	0.5333	0.0000	0.5333	0.5167	0.4978	0.1300	0.1587	0.0930
Extra Trees	0.5333	0.1250	0.5333	0.4833	0.4789	0.1100	0.1225	0.1490
K Nearest Neighbors	0.5000	0.1583	0.5000	0.3222	0.3833	-0.0500	-0.0500	0.0510
XGBoost	0.5000	0.1000	0.5000	0.4444	0.4500	0.1843	0.1862	0.0720
Naive Bayes	0.4667	0.1333	0.4667	0.5444	0.4722	0.1386	0.1750	0.0420
Decision Tree	0.4333	0.1750	0.4333	0.3722	0.3644	0.0143	0.0475	0.0280
Linear Discriminant	0.4333	0.0000	0.4333	0.4111	0.4056	0.0243	0.0362	0.0240
Gradient Boosting	0.3667	0.0000	0.3667	0.3667	0.3278	-0.0086	-0.0138	0.2130

Table 5. Linear and discriminant model performance comparisons.

Model	Accuracy	AUC	Recall	Precision	F1-Score	Kappa	MCC	Time (s)
Logistic Regression	0.6333	0.0000	0.6333	0.5222	0.5522	0.2300	0.2587	0.7910
Ridge Classifier	0.5667	0.0000	0.5667	0.4444	0.4878	0.1943	0.2000	0.0250
SVM (Linear Kernel)	0.5333	0.0000	0.5333	0.4111	0.4511	0.1943	0.2000	0.0310
Linear Discriminant Analysis	0.4333	0.0000	0.4333	0.4111	0.4056	0.0243	0.0362	0.0240
Quadratic Discriminant Analysis	0.5667	0.0000	0.5667	0.3444	0.4233	0.0000	0.0000	0.0460

Table 6. Tree-based models – Scoring Grid.

Model	Accuracy	AUC	Recall	Precision	F1-Score	Kappa	MCC	Time (s)
Random Forest Classifier	0.5667	0.1667	0.5667	0.4833	0.5011	0.1900	0.1862	0.163
Light Gradient Boosting	0.5667	0.1500	0.5667	0.3444	0.4233	0.0000	0.0000	0.036
Extra Trees Classifier	0.5333	0.1250	0.5333	0.4833	0.4789	0.1100	0.1225	0.142
Decision Tree Classifier	0.4333	0.1750	0.4333	0.3722	0.3644	0.0143	0.0475	0.046
Gradient Boosting Classifier	0.3667	0.0000	0.3667	0.3667	0.3278	-0.0086	-0.0138	0.231

4.8 Tree-Based Models

The scoring grid in Table 6 further confirms Random Forest’s superior performance among tree-based classifiers. The scoring grid also supports the argument that Random Forest classifier provided best combination of accuracy (56.67 %), F1-score (0.5011), and agreement (Kappa = 0.1900, MCC = 0.1862) compared to any other of the tested models. Though LightGBM gave incremental improvements in certain individual measures, lack of consistency in agreement indicates a problem with calibration, which makes the Random Forest the most reliable model in this classification problem.

4.9 Evaluation Metrics for Classification Models

Table 7 summarizes the comprehensive evaluation metrics for all classification models, with Logistic Regression emerging as the top performer. When emphasizing Recall to maximize the identification of highly impacted businesses, Logistic Regression emerged as the strongest model (Recall = 0.6333), with superior F1-score, Kappa, and MCC indicating robust and stable performance. While Ridge Classifier and Random Forest followed closely (Recall = 0.5667), Logistic Regression proved to be the most reliable and interpretable option for true positive detection in this climate impact classification task.

Table 7. Evaluation metrics for classification models predicting business impact.

Model	Accuracy	AUC	Recall	Precision	F1-Score	Kappa	MCC	Time (s)
Logistic Regression	0.6333	0.0000	0.6333	0.5222	0.5522	0.2300	0.2587	0.0540
Ridge Classifier	0.5667	0.0000	0.5667	0.4444	0.4878	0.1943	0.2000	0.0280
Random Forest Classifier	0.5667	0.1667	0.5667	0.4833	0.5011	0.1900	0.1862	0.2350
QDA (Quadratic Discriminant)	0.5667	0.0000	0.5667	0.3444	0.4233	0.0000	0.0000	0.0460
LightGBM	0.5667	0.1500	0.5667	0.3444	0.4233	0.0000	0.0000	0.0390
Dummy Classifier	0.5667	0.1500	0.5667	0.3444	0.4233	0.0000	0.0000	0.0240
SVM (Linear Kernel)	0.5333	0.0000	0.5333	0.4111	0.4511	0.1943	0.2000	0.0310
Ada Boost Classifier	0.5333	0.0000	0.5333	0.5167	0.4978	0.1300	0.1587	0.0930
Extra Trees Classifier	0.5333	0.1250	0.5333	0.4833	0.4789	0.1100	0.1225	0.1620
K Neighbors Classifier	0.5000	0.1583	0.5000	0.3222	0.3833	-0.0500	-0.0500	0.0410
XGBoost	0.5000	0.1000	0.5000	0.4444	0.4500	0.1843	0.1862	0.0520
Naive Bayes	0.4667	0.1333	0.4667	0.5444	0.4722	0.1386	0.1750	0.0290
Decision Tree Classifier	0.4333	0.1750	0.4333	0.3722	0.3644	0.0143	0.0475	0.0300
Linear Discriminant Analysis	0.4333	0.0000	0.4333	0.4111	0.4056	0.0243	0.0362	0.0370
Gradient Boosting Classifier	0.3667	0.0000	0.3667	0.3667	0.3278	-0.0086	-0.0138	0.2880

Table 8. Scoring metrics and their configurations in model evaluation.

ID	Display Name	Score Function	Scorer	Target	Args	Greater is Better	Multiclass	Custom
acc	Accuracy	accuracy_score	make_scorer(accuracy_score)	pred	{}	Yes	Yes	No
auc	AUC	roc_auc_score	make_scorer(roc_auc_score)	pred_proba	{'average': 'weighted', 'multi_class': 'ovr'}	Yes	Yes	No
recall	Recall	recall_score	make_scorer(recall_score)	pred	{'average': 'weighted'}	Yes	Yes	No
precision	Precision	precision_score	make_scorer(precision_score)	pred	{'average': 'weighted'}	Yes	Yes	No
f1	F1	f1_score	make_scorer(f1_score)	pred	{'average': 'weighted'}	Yes	Yes	No
kappa	Kappa	cohen_kappa_score	make_scorer(cohen_kappa_score)	pred	{}	Yes	Yes	No
mcc	MCC	matthews_corrcoef	make_scorer(matthews_corrcoef)	pred	{}	Yes	Yes	No

4.10 Scoring Metrics and Their Configurations in Model Evaluation

The scoring metrics configuration used in our model evaluation, detailed in Table 8, ensured comprehensive performance assessment. The classification task in question has its multiclass evaluation metrics namely Accuracy, AUC, Recall, Precision, F1-score, Kappa, and MCC also provided by PyCaret. All these are optimized such that the increased results imply better functioning of the model. Although Accuracy provides a good idea, Kappa and MCC provide a more subtle view on inter-class agreement or balance, respectively, which is invaluable to model evaluation and parameter adjustment exercise.

4.11 Evaluation Metrics for Classification Models

Table 9 has pitted all the evaluation measures of the classifiers. The best scores are obtained with Logistic Regression: accuracy (63.33%), recall (63.33%), F1-score (55.22%), Kappa (0.230), and MCC (0.2587). The results obtained by Ridge Classifier and Random Forest are similar, whereas the results adopted by QDA, LightGBM, and Dummy Classifier are very different (with values of Kappa and MCC being far particular low). This shows that there is limited concordance between the classes compared. In turn, the Logistic Regression is an attractive choice in terms

of its interpretability and forecasting power.

4.12 Cross-Validations and Training Metrics

The cross-validation results presented in Table 10 are further illustrated in Figure 6, which provides a comprehensive view of the Logistic Regression model's performance consistency across different evaluation metrics. The model showed strong and consistent performance during training, with high cross-validated accuracy (74.44%) and recall, but experienced a noticeable generalization gap on testing, with accuracy dropping to 63.33% and F1-score to 55.22%. Despite issues like zero AUC and variability across folds, its interpretability and strong recall make it a suitable choice for applications prioritizing the detection of true positives, such as climate impact assessment.

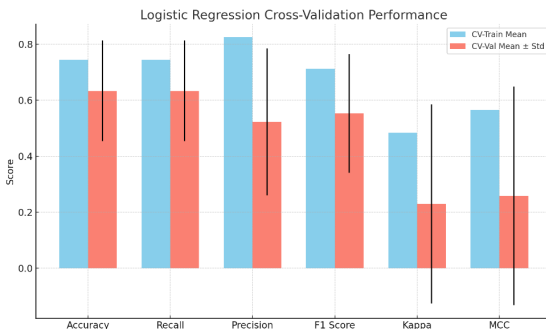
Tables 11, 12, 13, and 14 show before tuning, the Decision Tree model performed poorly with low accuracy (43.33%) and weak agreement metrics (Kappa = 0.0143, MCC = 0.0475), but tuning substantially improved its accuracy (56.67%), recall, and F1-score (0.5367), indicating better predictive balance despite some fold-to-fold instability. Further enhancement was achieved through Bagging, which maintained the improved accuracy (56.67%) and demonstrated stronger agreement (Kappa = 0.24,

Table 9. Evaluation metrics for classification models predicting business impact.

Model	Accuracy	AUC	Recall	Precision	F1-Score	Kappa	MCC	Custom Metric	Time (s)
Logistic Regression	0.6333	0.0000	0.6333	0.5222	0.5522	0.2300	0.2587	0.0000	0.0570
Ridge Classifier	0.5667	0.0000	0.5667	0.4444	0.4878	0.1943	0.2000	0.0000	0.0440
Random Forest Classifier	0.5667	0.1667	0.5667	0.4833	0.5011	0.1900	0.1862	0.0000	0.1560
QDA	0.5667	0.0000	0.5667	0.3444	0.4233	0.0000	0.0000	0.0000	0.0260
LightGBM	0.5667	0.1500	0.5667	0.3444	0.4233	0.0000	0.0000	0.0000	0.0380
Dummy Classifier	0.5667	0.1500	0.5667	0.3444	0.4233	0.0000	0.0000	0.0000	0.0250
SVM (Linear Kernel)	0.5333	0.0000	0.5333	0.4111	0.4511	0.1943	0.2000	0.0000	0.0600
Ada Boost Classifier	0.5333	0.0000	0.5333	0.5167	0.4978	0.1300	0.1587	0.0000	0.0850
Extra Trees Classifier	0.5333	0.1250	0.5333	0.4833	0.4789	0.1100	0.1225	0.0000	0.1730
K Neighbors Classifier	0.5000	0.1583	0.5000	0.3222	0.3833	-0.0500	-0.0500	0.0000	0.0610
XGBoost	0.5000	0.1000	0.5000	0.4444	0.4500	0.1843	0.1862	0.0000	0.0590
Naive Bayes	0.4667	0.1333	0.4667	0.5444	0.4722	0.1386	0.1750	0.0000	0.0420
Decision Tree Classifier	0.4333	0.1750	0.4333	0.3722	0.3644	0.0143	0.0475	0.0000	0.0460
Linear Discriminant Analysis	0.4333	0.0000	0.4333	0.4111	0.4056	0.0243	0.0362	0.0000	0.0270
Gradient Boosting Classifier	0.3667	0.0000	0.3667	0.3667	0.3278	-0.0086	-0.0138	0.0000	0.2620

Table 10. Cross-validations and training metrics summary for model performance.

Metric	CV-Train Mean	CV-Val Mean	Train Score	CV-Val Std. Dev
Accuracy	0.7444	0.6333	0.7333	0.1795
AUC	0.0000	0.0000	0.9190	0.0000
Recall	0.7444	0.6333	0.7333	0.1795
Precision	0.8244	0.5222	0.8187	0.2620
F1 Score	0.7117	0.5522	0.6987	0.2120
Kappa	0.4830	0.2300	0.4570	0.3551
MCC	0.5645	0.2587	0.5441	0.3896

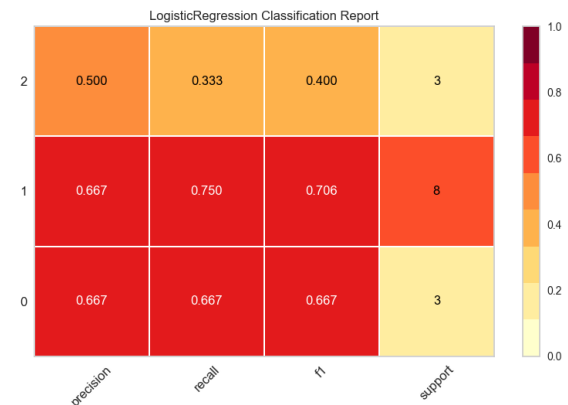
**Figure 6.** Logistic regression cross-validation Performances.

MCC = 0.2837), while Boosting lagged behind with weak accuracy (36.67%) and poor class discrimination, making Bagging the more reliable ensembling strategy for this classification task.

4.13 Logistic classifications Report

The comprehensive classification report presented in Figure 7 provides detailed insights into the Logistic Regression model's performance across different climate impact categories. The report indicates that the model achieves superior performance on class 1 with precision: 0.667, recall: 0.750, and F1-score: 0.706, suggesting this class is well-represented in the training data. However, performance degrades for class 2, which shows lower recall (0.333) and F1-score (0.400). This disparity highlights the presence of

class imbalance in the dataset, where implementing techniques such as stratified sampling, class weighting, or synthetic data generation could enhance predictive consistency across all vulnerability levels.

**Figure 7.** Logistic classifications report.

5 Conclusion

This empirical study outlines that SMEs in Pakistan's rural areas face heightened exposure to climate-related risks such as floods and droughts. Using advanced machine learning methods including ordered logistic regression (OLR), logistic regression (LR), and random forest (RF), the analysis confirms that rural enterprises are more vulnerable. LR achieved the best overall accuracy and recall at 63.3%, demonstrating strong predictive value in assessing climate risk to business continuity. These findings supported by OLR indicate that firms whose owners or managers were aware of temperature fluctuations ($p = 0.035$) and perceived financial consequences of climate change ($p = 0.018$) are significantly more affected by climate-related impacts, while geographical location ($p = 0.057$) and business tenure ($p = 0.091$) exerted

Table 11. Default decision tree model performances (Before Tuning).

Metric	Accuracy	AUC	Recall	Precision	F1 Score	Kappa	MCC
Mean	0.4333	0.175	0.4333	0.3722	0.3644	0.0143	0.0475
Std Dev	0.1528	0.275	0.1528	0.2789	0.2002	0.3191	0.3592

Table 12. Tuned decision tree model performance (After Hyper-parameter Tuning).

Metric	Accuracy	AUC	Recall	Precision	F1 Score	Kappa	MCC
Mean	0.5667	0.1583	0.5667	0.5556	0.5367	0.21	0.2449
Std Dev	0.2134	0.2673	0.2134	0.3093	0.2434	0.4346	0.477

Table 13. Ensemble with bagging (ensemble model).

Metric	Accuracy	AUC	Recall	Precision	F1 Score	Kappa	MCC
Mean	0.5667	0.2	0.5667	0.5444	0.5144	0.24	0.2837
Std Dev	0.2134	0.3078	0.2134	0.3641	0.282	0.398	0.4266

Table 14. Ensemble with boosting (ensemble model (dt, method='Boosting')).

Metric	Accuracy	AUC	Recall	Precision	F1 Score	Kappa	MCC
Mean	0.3667	0	0.3667	0.3611	0.3222	0.0214	0.0225
Std Dev	0.1795	0	0.1795	0.3645	0.2531	0.2876	0.3528

only marginal effects. Overall, temperature fluctuation awareness and perceived financial impacts have been identified as the key risk factors driving climate vulnerability among SMEs. Accordingly, the present study recommends the following policy interventions: (a) investment in rural infrastructure, (b) targeted financial support mechanisms for climate-vulnerable SMEs, and (c) strengthened governance frameworks in rural areas. By emphasizing the enhanced accuracy of predictions that can be achieved in cases where machine-learning tools are used to complement sector-specific data, the study also demonstrates the value of data-driven approaches and supports the development of effective action strategies targeting resilience enhancement of SMEs in vulnerable locations of Pakistan.

The above research highlights the necessity of the further development of spatially and temporally more accurate estimates of the effects of the climate on small- and medium-sized enterprises (SMEs), especially concerning local flood patterns dependent on the variation in temperature. Including more inputs in the dataset, both sectoral and regional, across Pakistan, researchers could significantly narrow the gap of the model performance and gain more knowledge of vulnerability. The present study also suggests the introduction of more complex machine learning

techniques (especially deep learning and ensemble tools), which proved to increase the precision of the predictions and identify hidden patterns in the data. An analysis over time is also necessary in monitoring the changes in the adaptation strategies of SMEs and contemplating on their sustainability. Lastly, connections between predictive models and early warning systems and local governance processes may provide SMEs with currently responsive help. Embedding the resulting insights into national policy ought to allow a more effective, fairer, data-based resource management practice.

Data Availability Statement

Data will be made available on request.

Funding

This work was supported without any funding.

Conflicts of Interest

The authors declare no conflicts of interest.

Ethical Approval and Consent to Participate

Not applicable.

References

- [1] Sheikh, M. M., Manzoor, N., Adnan, M., Ashraf, J., & Khan, A. M. (2009). Climate profile and past climate changes in Pakistan. *Global Change Impact Studies Center (GCISC)-RR-01*, 1-189. <https://www.researchgate.net/publication/314949214>
- [2] McCarthy, J. J. (Ed.). (2001). *Climate change 2001: impacts, adaptation, and vulnerability: contribution of Working Group II to the third assessment report of the Intergovernmental Panel on Climate Change* (Vol. 2). Cambridge university press.
- [3] Alam, A., Du, A. M., Rahman, M., Yazdifar, H., & Abbasi, K. (2022). SMEs respond to climate change: Evidence from developing countries. *Technological Forecasting and Social Change*, 185, 122087. [[Crossref](#)]
- [4] Ullah, W., Nihei, T., Nafees, M., Zaman, R., & Ali, M. (2018). Understanding climate change vulnerability, adaptation and risk perceptions at household level in Khyber Pakhtunkhwa, Pakistan. *International Journal of Climate Change Strategies and Management*, 10(3), 359-378. [[Crossref](#)]
- [5] Rolnick, D., Donti, P. L., Kaack, L. H., Kochanski, K., Lacoste, A., Sankaran, K., ... & Bengio, Y. (2022). Tackling climate change with machine learning. *ACM Computing Surveys (CSUR)*, 55(2), 1-96. [[Crossref](#)]
- [6] Zhao, Z., Tang, S., Huang, J., Yang, J., Ma, Z., Fang, W., ... & Bi, J. (2025). Firm-level climate risk assessment: Recent progress and future research agenda. *Risk Sciences*, 100012. [[Crossref](#)]
- [7] Tasnuva, A., & Bari, Q. H. (2024). Integrating machine learning algorithms and fuzzy AHP for assessing livelihood vulnerability in Southwestern Coastal Bangladesh. *Regional Studies in Marine Science*, 79, 103825. [[Crossref](#)]
- [8] Khan, M. H., Zein Alabdeen, Z., & Anupam, A. (2024). Firm-level climate change risk and adoption of ESG practices: a machine learning prediction. *Business Process Management Journal*, 30(6), 1741-1763. [[Crossref](#)]
- [9] Mohapatra, S., Mohapatra, S., Han, H., Ariza-Montes, A., & López-Martín, M. D. C. (2022). Climate change and vulnerability of agribusiness: Assessment of climate change impact on agricultural productivity. *Frontiers in Psychology*, 13, 955622. [[Crossref](#)]
- [10] Ukoba, K., Onisuru, O. R., Jen, T. C., Madyira, D. M., & Olatunji, K. O. (2025). Predictive modeling of climate change impacts using Artificial Intelligence: a review for equitable governance and sustainable outcome. *Environmental Science and Pollution Research*, 32(17), 10705-10724. [[Crossref](#)]
- [11] Arrieta, A. B., Díaz-Rodríguez, N., Del Ser, J., Bennetot, A., Tabik, S., Barbado, A., ... & Herrera, F. (2020). Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information fusion*, 58, 82-115. [[Crossref](#)]
- [12] Abdelwahed, N. A. A., Soomro, B. A., & Shah, N. (2022). Climate change and pro-environmental behaviours: the significant environmental challenges of livelihoods. *Management of Environmental Quality: An International Journal*, 33(5), 1187-1206. [[Crossref](#)]
- [13] Takona, J. P. (2024). Research design: qualitative, quantitative, and mixed methods approaches. *Quality & Quantity*, 58(1), 1011-1013. [[Crossref](#)]
- [14] Little, R. J., & Rubin, D. B. (2019). *Statistical analysis with missing data*. John Wiley & Sons.
- [15] Ayesha, S., Hanif, M. K., & Talib, R. (2020). Overview and comparative study of dimensionality reduction techniques for high dimensional data. *Information Fusion*, 59, 44-58. [[Crossref](#)]
- [16] Zhao, Z., Li, D., & Dai, W. (2023). Machine-learning-enabled intelligence computing for crisis management in small and medium-sized enterprises (SMEs). *Technological Forecasting and Social Change*, 191, 122492. [[Crossref](#)]
- [17] Hassija, V., Chamola, V., Mahapatra, A., Singal, A., Goel, D., Huang, K., ... & Hussain, A. (2024). Interpreting black-box models: a review on explainable artificial intelligence. *Cognitive Computation*, 16(1), 45-74. [[Crossref](#)]
- [18] Breiman, L. (2001). Random forests. *Machine learning*, 45(1), 5-32. [[Crossref](#)]
- [19] Hino, M., Benami, E., & Brooks, N. (2018). Machine learning for environmental monitoring. *Nature Sustainability*, 1(10), 583-588. [[Crossref](#)]
- [20] Agresti, A. (2010). *Analysis of ordinal categorical data*. John Wiley & Sons.
- [21] James, G., Witten, D., Hastie, T., & Tibshirani, R. (2013). *An introduction to statistical learning: with applications in R* (Vol. 103). New York: springer. [[Crossref](#)]
- [22] Ke, G., Meng, Q., Finley, T., Wang, T., Chen, W., Ma, W., ... & Liu, T. Y. (2017). Lightgbm: A highly efficient gradient boosting decision tree. *Advances in neural information processing systems*, 30.
- [23] Friedman, J. H. (2001). Greedy function approximation: a gradient boosting machine. *The Annals of Statistics*, 29(5), 1189-1232. [[Crossref](#)]
- [24] Freund, Y., & Schapire, R. E. (1997). A decision-theoretic generalization of on-line learning and an application to boosting. *Journal of computer and system sciences*, 55(1), 119-139. [[Crossref](#)]
- [25] Cervantes, J., Garcia-Lamont, F., Rodríguez-Mazahua, L., & Lopez, A. (2020). A comprehensive survey on support vector machine classification: Applications, challenges and trends. *Neurocomputing*, 408, 189-215. [[Crossref](#)]
- [26] Halder, R. K., Uddin, M. N., Uddin, M. A., Aryal, S., & Khraisat, A. (2024). Enhancing K-nearest neighbor algorithm: a comprehensive review and performance analysis of modifications. *Journal of Big Data*, 11(1),

113. [Crossref]
- [27] Ruppert, D. (2004). The Elements of Statistical Learning: Data Mining, Inference, and Prediction. *Journal of the American Statistical Association*, 99(466), 567. [Crossref]
- [28] Wang, S., Ren, J., Lian, X., Bai, R., & Jiang, X. (2022, August). Boosting the discriminant power of naive Bayes. In *2022 26th International Conference on Pattern Recognition (ICPR)* (pp. 4906-4912). IEEE. [Crossref]
- [29] Wickramasinghe, I., & Kaluturage, H. (2021). Naive Bayes: applications, variations and vulnerabilities: a review of literature with code snippets for implementation. *Soft computing*, 25(3), 2277-2293. [Crossref]
- [30] Chicco, D., & Jurman, G. (2020). The advantages of the Matthews correlation coefficient (MCC) over F1 score and accuracy in binary classification evaluation. *BMC genomics*, 21(1), 6. [Crossref]
- [31] Pedregosa, F., Varoquaux, G., Gramfort, A., Michel, V., Thirion, B., Grisel, O., ... & Duchesnay, É. (2011). Scikit-learn: Machine Learning in Python. *The Journal of Machine Learning Research*, 12, 2825-2830. <https://jmlr.org/papers/v12/pedregosa11a.html>
- [32] Mehrabi, N., Morstatter, F., Saxena, N., Lerman, K., & Galstyan, A. (2021). A survey on bias and fairness in machine learning. *ACM computing surveys (CSUR)*, 54(6), 1-35. [Crossref]



Abdur Rehman is a PhD scholar in Statistics with a specialization in Data Science. His research focuses on advanced statistical methodologies, machine learning techniques, and their real-world applications, particularly in bioinformatics and social science research. With a solid background in statistical modeling, data analysis, and programming, he is passionate about developing robust solutions for complex data problems. His academic work emphasizes integrating traditional statistical approaches with modern data science tools to contribute meaningfully to both research and practical applications. (Email: abdurrehmanshad@uop.edu.pk)



Dr. Sami Ullah is an Assistant Professor at the School of Mathematics, Statistics and Mechanics, Beijing University of Technology. His research expertise lies at the intersection of Epidemiological Methods and Statistical Modeling, with a strong focus on Spatial and Spatiotemporal Modeling and Bayesian Analysis. Ullah has contributed significantly to the development and application of advanced statistical methods to public health. His work involves integrating spatial and temporal data to better understand disease patterns and improve prediction and prevention strategies. (Email: sami.ullah@bjut.edu.cn)



Ubaid Ullah is a graduate with a Bachelor's in Statistics, having completed his degree at the University of Peshawar, Pakistan, in 2023. His research interests Data Analysis, Machine Learning, Statistical Modeling, and Economic Forecasting. Ubaid has contributed to various research projects, including studies on predicting fertility outcomes and economic trends economic variables influencing Foreign Direct Investment (FDI) in Pakistan using advanced machine learning techniques. Currently, Ubaid works as a Research Assistant, applying machine learning and statistical tools in various research projects. His technical expertise includes Python, R, SPSS, SQL, and various machine learning deep learning methods. (Email: ubaidullahyousafzai5683@gmail.com)



Sehran Hassan is a PhD scholar in Statistics with a specialization in Machine Learning. His research focuses on the development and application of advanced machine learning algorithms for solving real-world data analysis problems, particularly in high-dimensional data and predictive modeling. With a strong foundation in statistical theory and computational techniques, Sehran is dedicated to bridging the gap between classical statistics and modern machine learning, contributing to both academic research and practical solutions in the field of data science. (Email: sehranhassan@uop.edu.pk)



Dr. Alamgir Khalil is an Associate Professor in the Department of Statistics at the University of Peshawar. His areas of specialization include Probability Theory, Statistical Inference, and Machine Learning. With extensive experience in teaching and research, Dr. Alamgir has contributed significantly to the development of advanced statistical methodologies and their applications in modern data science. His research focuses on bridging theoretical statistics with practical machine learning approaches to address complex data-driven challenges in various fields. (Email: alamgir_khalil@uop.edu.pk)



Muhammad Sani is a final-year undergraduate student pursuing a Bachelor's degree in Computer Science at the University of Lakki Marwat, Pakistan. His research interests lie in Artificial Intelligence, Deep Learning, Machine Learning, Medical Image Analysis, and Cybersecurity. Muhammad has worked on multiple deep learning projects, including osteoporosis classification using CNN architectures (VGG-16, ResNet-50) and multiclass dental disease detection using mobile and OPG images. He has also explored Fake News Detection and Malware Classification in the cybersecurity domain. Muhammad's technical expertise includes Python, TensorFlow/Keras, Google Collab, Flask, OpenCV, and transfer learning-based deep learning methods. (Email: muhammadsanimarwat@gmail.com)