

PERFORMANCE EVALUATION OF NEURAL NETWORKS AND TRADITIONAL STATISTICAL METHODS IN ANALYZING IMBALANCED DATA: A COMPARATIVE STUDY

Abela Chairunissa^{1*}, Hilwin Nisa²

^{1,2}Department of Statistics, Faculty of Mathematics and Natural Sciences, Brawijaya University
Veteran Street, Malang, 65145, East Java, Indonesia

Corresponding Author's Email: abelachairunissa@ub.ac.id

Abstract: Class imbalance is a common issue in predictive modeling, particularly when minority classes carry critical significance, as seen in applications like fraud detection, rare disease prediction, and customer churn analysis. This study uses linear and non-linear simulated data scenarios to examine the performance of logistic regression, discriminant analysis, and neural networks on imbalanced data. For linear data, logistic regression and discriminant analysis displayed high sensitivity but extremely low specificity, indicating a strong bias toward the majority class. Neural networks showed marginal improvement but remained ineffective in detecting minority classes. In contrast, neural networks demonstrated superior sensitivity for non-linear data and were notably better at identifying minority classes, underscoring their suitability for complex data relationships. Our results highlight that accuracy alone is insufficient for evaluating models on imbalanced data; instead, sensitivity and specificity offer more relevant insights. Overall, this study suggests that neural networks are preferable for imbalanced data with non-linear patterns, and data characteristics and appropriate evaluation metrics should inform model selection.

Keywords: Class Imbalance, Predictive Modeling, Minority Class Detection.

1. INTRODUCTION

Class imbalance, a prevalent issue in real-world datasets, poses substantial obstacles in developing predictive models, particularly when the disparity between majority and minority class occurrences is extreme. This imbalance can lead to predictive models favoring majority classes while neglecting minority classes, resulting in biased predictions that reduce model accuracy and utility. In critical applications, detecting minority classes is paramount; examples include fraud detection (where the minority class is fraudulent activity), rare disease prediction (minority class = diagnosed cases), and customer churn prediction (minority class = customers likely to leave). In these contexts, failing to identify the minority class accurately can lead to significant financial and operational consequences [2], [3], [4].

The data structure and the relationships within it heavily influence model selection. Linear models, such as logistic regression and discriminant analysis, are often suitable for data with a linear relationship between predictor variables and outcomes. These models establish a direct association between predictors and the response, making them useful for many prediction tasks due to their interpretability and simplicity. However, when the data exhibits non-linear patterns, linear models may struggle to capture the complexity, often leading to underfitting. In such cases, non-linear models, like neural networks, offer an advantage as they can model intricate interactions and non-linear dependencies between features [3], [5].

Current solutions to address class imbalance include resampling methods (such as oversampling the minority class or undersampling the majority class) and cost-sensitive learning, where higher misclassification costs are applied to minority class predictions. Nevertheless, determining the best model type—linear or non-linear—remains unresolved when the data is imbalanced and complex. While linear models offer ease of use and interpretability, non-linear models like neural networks may outperform them in identifying nuanced patterns, particularly when dealing with severe imbalances and non-linear data distributions [6].

In light of these issues, this study seeks to evaluate the performance of various predictive models, including both linear and non-linear approaches, in managing imbalanced data scenarios. By systematically comparing their

effectiveness, this research aims to provide insights into model selection strategies that align with different data characteristics. This study contributes to the field of predictive modeling, offering a refined approach for managing imbalanced datasets and enhancing model reliability in critical applications where minority class detection is essential [7], [8], [9], [10]. This expanded understanding of model performance across data types aims to support practitioners in selecting models that maximize accuracy and utility when facing imbalanced and complex datasets.

2. METHODOLOGY

This research employed a data simulation approach to create an experimental setting to evaluate the performance of various predictive models under different data scenarios. A total of 1,000 samples were generated, with an imbalance ratio of 30%, to simulate real-world challenges in prediction modeling.

2.1. Data Simulation

To systematically assess the model performance, the data simulation was conducted in two distinct scenarios, linear and non-linear with the following analysis steps:

1) Data Generation

For both scenarios, variables were simulated using the $\text{rnorm}(0,1)$ function to produce values from a standard normal distribution. The outcome variable Y was assigned a binary value, with the majority class (0) representing 70% of the data and the minority class (1) comprising 30%, reflecting class imbalance.

2) Scenario Setup

Linear Scenario: The relationship between predictors X_1, X_2, X_3 and the binary outcome Y was kept linear to reflect standard conditions for logistic regression and discriminant analysis.

Table 1. Scenario 1

Variable	Simulation
X_1	$\text{rnorm}(0,1)$
X_2	$\text{rnorm}(0,1)$
X_3	$\text{rnorm}(0,1)$
Y	Majority class (0), minority class (1)

Data source: Researcher's processing, 2024

Non-linear data scenario: We introduced a more complex relationship pattern between the variables. Variables X_1 , and X_2 followed the same normal distribution as in the linear scenario. However, variable X_3 was simulated as a non-linear transformation, defined by $X_1^2 + \text{rnorm}(0,1)$. The outcome variable Y maintained the same class structure as in the linear scenario, with majority class (0) and minority class (1), representing a challenging classification task.

Table 2. Scenario 2

Variable	Simulation
X_1	$\text{rnorm}(0,1)$
X_2	$\text{rnorm}(0,1)$
X_3	$X_1^2 + \text{rnorm}(0,1)$
Y	Majority class (0), minority class (1)

Data source: Researcher's processing, 2024

3) Validation of Simulated Data

Summary statistics (mean, variance) and visualizations (e.g., histograms, scatterplots) were used to verify distributional properties and ensure the desired data characteristics (e.g., linear vs. non-linear patterns).

4) Train-Test Split

The dataset was partitioned into training (70%) and testing (30%) subsets to evaluate the out-of-sample performance of the models.

5) Reproducibility

A random seed was set before data generation to ensure the simulation could be replicated for validation and comparison purposes.

These steps ensured that each data scenario realistically reflected the conditions under which the models—logistic regression, discriminant analysis, and neural networks—would be evaluated.

2.2. Predictive Model

Three predictive models were evaluated in this study, each chosen for its suitability to either linear or non-linear data characteristics.

- 1) **Logistic Regression:** Logistic regression is a widely used statistical method for binary classification. It models the probability of a binary outcome based on a set of predictor variables, assuming a linear relationship between the independent variables and the dependent variable [11]. The Box-Tidwell Test is employed to validate this assumption of linearity between continuous predictors and the logit of the outcome. This test introduces interaction terms between each continuous predictor and its natural logarithm to detect significant deviations from linearity [12]. If any interaction terms are statistically significant, a non-linear transformation of the corresponding predictor may be needed to improve model fit. This model is particularly effective in the linear data scenario because it focuses on linear boundaries for classification. The model estimates parameters using the maximum likelihood method and assumes a logistic function of the form:

$$P(Y = 1|X) = \frac{1}{1 + e^{-(\beta_0 + \beta_1 X_1 + \dots + \beta_p X_p)}}$$

where $P(Y = 1|X)$ is the probability of the outcome being class 1 given input X , and $\beta_0 + \beta_1 X_1 + \dots + \beta_p X_p$ are the model coefficients [11]. Logistic regression suits linearly separable data and provides interpretable results through odds ratios.

- 2) **Discriminant Analysis:** Discriminant analysis is another classification method suited to linearly separable data and is commonly applied when classes have Gaussian distributions [13]. It aims to model the difference between classes, making it ideal for data with clear, linear separation patterns. Linear Discriminant Analysis (LDA) and Quadratic Discriminant Analysis (QDA) are the most commonly used variants. LDA assumes equal covariance matrices across groups and uses the linear function:

$$\delta_k(x) = x^T \sum^{-1} \mu_k - \frac{1}{2} \mu_k^T \sum^{-1} \mu_k + \log(\pi_k)$$

Where μ_k is the mean vector of class k , $\sum$ is the shared covariance matrix, and π_k is the prior probability of class k .

- 3) **Neural Networks:** Neural networks are flexible, non-linear models for classification designed to handle complex relationships and non-linear patterns in data [14]. Comprising interconnected layers (input, hidden, and output), each node applies a transformation based on the data it receives, making neural networks well-suited to the non-linear data scenario in this research. A typical feedforward neural network consists of layers of interconnected nodes or "neurons" and learns from data by adjusting weights using algorithms such as backpropagation. The model can be expressed as:

$$\hat{y} = f(W^{(2)} \cdot g(W^{(1)}x + b^{(1)}) + b^{(2)})$$

where x is the input vector, $W^{(1)}, W^{(2)}$ are weight matrices, $b^{(1)}, b^{(2)}$ are bias vectors, g is an activation function (e.g., ReLU, sigmoid), and f is the output function (e.g., softmax for classification) [12]. Neural networks are particularly effective in capturing non-linear and high-dimensional relationships in data.

2.3. Evaluation Metrics

Model performance was evaluated based on the following metrics:

- 1) Accuracy: This metric measures the proportion of true results (both true positives and true negatives) among the total number of cases[15], [16], [17].
- 2) Sensitivity: Sensitivity indicates the proportion of actual positives (minority class) that are correctly identified by the model [15], [16], [17].
- 3) Specificity: Specificity measures the proportion of actual negatives (majority class) that are correctly identified [15], [16], [17].
- 4) Kappa: Kappa statistic provides a measure of agreement between the predicted and actual classifications, adjusting for the possibility of chance agreement.

These metrics provided a comprehensive assessment of each model's performance across both linear and non-linear data scenarios, allowing for a detailed comparison of their strengths and limitations under conditions of class imbalance.

3. RESULTS AND DISCUSSION

This section presents the findings from the linearity tests on simulated data and the model evaluations on both linear and non-linear scenarios. We focus on the performance metrics of logistic regression, discriminant analysis, and neural networks under imbalanced data conditions. In both scenarios, the binary target variable Y was simulated with a class imbalance ratio of approximately 80:20 (majority: minority) to reflect common challenges in real-world classification problems such as fraud detection or medical diagnosis.

3.1. Linearity Test Results on Simulated Data

3.1.1. Linear Scenario

The simulated data for the linear scenario was generated by assuming additive relationships between predictors (X1, X2, X3) and the log odds of the binary outcome Y. In contrast, the non-linear scenario involved non-additive interactions and transformations (e.g., squared or exponential terms) in the data-generating process to mimic more complex real-world relationships. In the linear data scenario, the Box-Tidwell test [18] was conducted to assess the linearity of the predictor variables (X1, X2, X3) with respect to the logit of the target variable Y. The test examines whether each predictor has a linear relationship with the logit of the outcome by introducing interaction terms with their logarithmic transformations. The form of the logistic regression model used in the test is:

$$\log \left(\frac{P(Y = 1)}{1 - P(Y = 1)} \right) = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \beta_4 (X_1 \cdot \log(X_1)) + \beta_5 (X_2 \cdot \log(X_2)) + \beta_6 (X_3 \cdot \log(X_3))$$

The estimates and p-values for the predictors and their logarithmic transformations are presented in Table 3 below.

Table 3. Linearity Test Results for Linear Scenario		
Variable	Estimates	P-Value
Intercept	-2.47	0.04*
X ₁	0.75	0.21
X ₂	0.15	0.75
X ₃	-0.78	0.13
X _{1_log}	-2.46	0.17
X _{2_log}	-0.11	0.94
X _{3_log}	-2.13	0.21

Data source: Researcher's processing, 2024

The p-values for the log-transformed predictors (X_{1_log}, X_{2_log}, X_{3_log}) exceed 0.05, indicating no significant non-linear relationship between these predictors and the target variable Y. Consequently, the predictors can be assumed to have a linear relationship with the logit, supporting the choice of linear models in this scenario.

3.1.2. Non-Linear Scenario

In contrast, the Box-Tidwell test for the non-linear data scenario (see Table 4) was conducted to examine the linearity assumption between each predictor and the logit of the binary outcome variable YYY . In this case, the test revealed statistically significant p-values ($p < 0.001$) for all predictors and their logarithmic transformations (see Table 4). This provides strong evidence of non-linearity between the predictors (X_1, X_2, X_3) and the logit of Y , justifying this dataset's use of non-linear models, such as neural networks. The form of the model used in the Box-Tidwell test for the non-linear scenario is the same logistic regression formulation:

$$\log\left(\frac{P(Y = 1)}{1 - P(Y = 1)}\right) = \beta_0 + \beta_1 X_1 + \beta_2 X_2 + \beta_3 X_3 + \beta_4 (X_1 \cdot \log(X_1)) + \beta_5 (X_2 \cdot \log(X_2)) + \beta_6 (X_3 \cdot \log(X_3))$$

In this model, the interaction terms ($X_i \cdot \log(X_i)$) assess the departure from linearity for each predictor.

Table 4. Linearity Test Results for Non-Linear Scenario		
Variable	Estimates	P-Value
Intercept	-2.77	< 2e-16 ***
X_1	1.76	1.89e-07 ***
X_2	1.96	6.39e-15 ***
X_3	1.78	< 2e-16 ***
X_1_log	0.41	3.91e-07 ***
X_2_log	0.58	< 2e-16 ***
X_3_log	0.18	0.000331 ***

Data source: Researcher's processing, 2024

These results confirm that the predictors do not maintain a linear relationship with the logit of the outcome variable. Hence, non-linear classifiers, such as neural networks, are more suitable for modeling this data type. The positive coefficients for the log-transformed predictors suggest that these transformations add valuable information to the model, indicating a strong non-linear association with the target variable.

3.2. Observation and Prediction Matrices

To better understand each model's performance in classifying imbalanced data in linear and non-linear scenarios, we examine the observation and prediction matrices showing the true vs. predicted class distribution. These matrices highlight each model's sensitivity (true positive rate) and specificity (true negative rate) performance and offer insights into class distribution challenges.

3.2.1. Observation and Prediction Matrices for Linear Scenario

In the linear scenario, the matrices for logistic regression, discriminant analysis, and neural network models are shown in Table 5.

Table 5. Observation and Prediction Matrices for Linear Scenario			
Model	True Class	Predicted Class 0	Predicted Class 1
Logistic Regression	0	140	60
	1	0	0
Discriminant Analysis	0	140	60
	1	0	0
Neural Network	0	138	59
	1	2	1

Data source: Researcher's processing, 2024

In this scenario:

- 1) Logistic Regression and Discriminant Analysis consistently misclassified all minority class observations as the majority class (sensitivity of 1.00 but specificity of 0.00), underscoring their challenge in addressing imbalanced data with linear relationships.
- 2) The neural network slightly improved classification accuracy, identifying two minority cases correctly, which is reflected in a marginally higher kappa score (0.003), yet specificity remained very low (0.02).

The dominance of predictions in the majority class reveals the models' tendency to rely heavily on the more frequent class, likely due to the imbalanced nature of the data. The neural network's slight improvement indicates a possible advantage of non-linear models, even in linear scenarios with imbalanced data.

3.2.2. Observation and Prediction Matrices for Non-Linear Scenario

In the non-linear data scenario, the observation and prediction matrices reveal significant improvements with the neural network model, as shown in Table 6.

Table 6. Observation and Prediction Matrices for Non-Linear Scenario			
Model	True Class	Predicted Class 0	Predicted Class 1
Logistic Regression	0	19	9
	1	33	139
Discriminant Analysis	0	14	8
	1	38	140
Neural Network	0	30	11
	1	22	137

Data source: Researcher's processing, 2024

In this scenario:

- 1) Logistic Regression and Discriminant Analysis misclassified a significant portion of the minority class (33 and 38 incorrect predictions, respectively), reflecting their limited adaptability to non-linear relationships within imbalanced data. This is consistent with their lower kappa scores (0.253 and 0.1907, respectively).
- 2) Neural network showed clear improvements, correctly identifying 137 out of 159 minority class observations and achieving an overall accuracy of 0.82. This model also balanced specificity (0.77) and sensitivity (0.95), leading to the highest kappa score (0.3906) among the models.

These results underscore the importance of using models that capture non-linear patterns when dealing with imbalanced data. In the non-linear scenario, the neural network's ability to correctly identify both majority and minority classes demonstrates its flexibility in handling complex relationships, making it a preferred choice over traditional linear models. The observation and prediction matrices reveal how each model responds to the data's characteristics, emphasizing the need for careful model selection in predictive tasks, especially under challenging non-linearity and class imbalance conditions.

Figure 1 presents a comparative visualization of the prediction results across all three models using stacked bar charts to enhance understanding of model performance in the non-linear scenario. These plots visually demonstrate the number of correctly and incorrectly classified instances for each class, helping to highlight the strengths and limitations of each method under non-linear conditions. The neural network model's better adaptability is evident in its higher number of correct predictions for the minority class.

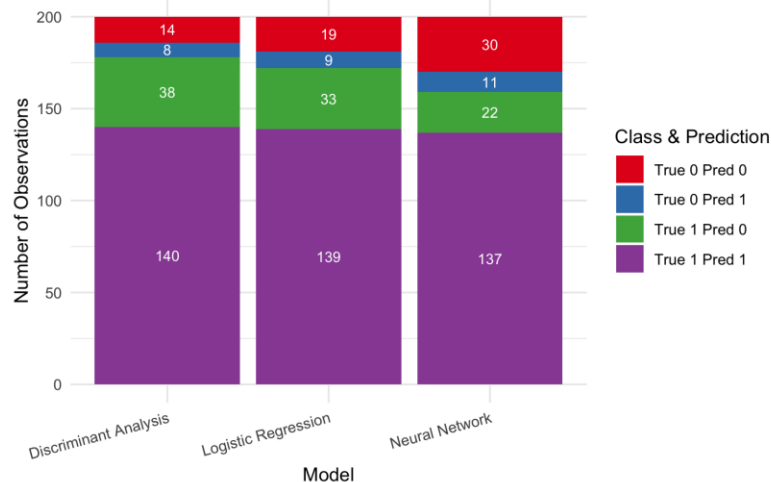


Figure 1. Classification Results by Model under Non-Linear Scenario

As shown in Figure 1, the neural network's superior classification of minority class instances in the non-linear scenario is visually evident, especially in comparison to the flat performance of logistic regression and discriminant analysis. This visual confirmation further supports the numerical findings in Table 8.

3.3. Model Evaluation

The performance of each model in the linear data scenario is summarized in Table 3, using accuracy, sensitivity, specificity, and kappa statistics.

3.3.1. Model Evaluation on Linear Imbalanced Data

The performance of each model in the linear data scenario is summarized in Table 7, using accuracy, sensitivity, specificity, and kappa statistics.

Table 7. Model Performance on Linear Imbalanced Data				
Model	Accuracy	Sensitivity	Specificity	Kappa
Logistic Regression	0.70	1.00	0.00	0.000
Discriminant Analysis	0.70	1.00	0.00	0.000
Neural Network	0.69	0.98	0.02	0.003

Data source: Researcher's processing, 2024

Both logistic regression and discriminant analysis achieved high sensitivity (1.00), accurately identifying all minority class instances, but they had low specificity (0.00), failing to classify any majority class instances correctly. This is reflected in the kappa values (0.000), which indicate no agreement beyond chance, highlighting the models' limitations in addressing class imbalance. The neural network model performed similarly, with a slightly lower sensitivity (0.98) and a marginally higher specificity (0.02), suggesting limited effectiveness in a linear data scenario.

3.3.2 Model Evaluation on Non-Linear Imbalanced Data

These findings align with previous research that highlights the superiority of neural networks in capturing complex, non-linear interactions in high-dimensional data [18]. While interpretable and efficient, traditional linear models often struggle when the true data-generating process involves significant interaction or non-linear terms. The model performance results in the non-linear data scenario are shown in Table 8.

Table 8. Model Performance on Non-Linear Imbalanced Data				
Model	Accuracy	Sensitivity	Specificity	Kappa
Logistic Regression	0.73	0.42	0.82	0.253
Discriminant Analysis	0.72	0.31	0.86	0.1907
Neural Network	0.82	0.95	0.77	0.3906

Data source: Researcher's processing, 2024

In the non-linear data scenario, the neural network outperformed the other models, with an accuracy of 0.82 and a high sensitivity (0.95), indicating its ability to identify the minority class instances more accurately. The kappa statistic (0.3906) confirms a moderate agreement, suggesting that the neural network is more effective for non-linear imbalanced data. Logistic regression and discriminant analysis, while achieving acceptable specificity (0.82 and 0.86, respectively), showed lower sensitivity, underscoring their limitations in capturing the non-linear patterns within the data.

3.4. Discussion

The findings demonstrate that model choice significantly impacts performance in class imbalance and data non-linearity. Traditional models, such as logistic regression and discriminant analysis, provided comparable performance for linear relationships. However, these models struggled in non-linear settings, while the neural network performed effectively, particularly in sensitivity and overall accuracy. These results align with the literature emphasizing the strengths of neural networks for complex, non-linear data structures [6]. The lack of specificity in linear models within both scenarios underscores the challenges of addressing class imbalance, particularly in the majority class. This study highlights the importance of matching model selection with data characteristics to enhance predictive accuracy in imbalanced datasets. The findings suggest that neural networks

can perform applications involving non-linear relationships, such as rare event detection in medical diagnostics or fraud detection than traditional linear models.

4. CONCLUSION

In this study, we evaluated the performance of logistic regression, discriminant analysis, and neural networks on imbalanced simulated datasets with linear and non-linear patterns. For linear data, logistic regression and discriminant analysis demonstrated high sensitivity but very low specificity, reflecting a strong bias toward the majority class and their inability to detect minority class observations accurately. While neural networks slightly improved linear data, they still fell short in effectively identifying minority cases. In the non-linear data scenario, neural networks significantly outperformed the linear models, achieving higher sensitivity and better detecting minority classes. This result highlights neural networks' capacity to model complex, non-linear relationships within imbalanced data, making them more effective for such contexts. As an evaluation metric, accuracy proved insufficient for assessing model performance on imbalanced data. Sensitivity and specificity, however, offered more relevant insights, revealing the models' strengths and weaknesses in handling class imbalances. Ultimately, these findings suggest that neural networks are preferable for imbalanced data with non-linear relationships. The choice of model should be carefully aligned with data characteristics, and evaluation should incorporate metrics that adequately reflect performance on imbalanced datasets.

ACKNOWLEDGMENT

The authors would like to express their gratitude to the Department of Statistics, Brawijaya University, and Bicosda 2024 for providing the resources and support necessary to carry out this study.

REFERENCES

- [1] P. Lalwani, M. K. Mishra, J. S. Chadha, and P. Sethi, "Customer churn prediction system: a machine learning approach," *Computing*, vol. 104, no. 2, pp. 271–294, Feb. 2022, doi: 10.1007/s00607-021-00908-y.
- [2] R. Sudharsan and E. N. Ganesh, "A Swish RNN based customer churn prediction for the telecom industry with a novel feature selection strategy," *Conn Sci*, vol. 34, no. 1, pp. 1855–1876, Dec. 2022, doi: 10.1080/09540091.2022.2083584.
- [3] Z. Tianyuan and S. Moro, "Research Trends in Customer Churn Prediction: A Data Mining Approach," 2021, pp. 227–237. doi: 10.1007/978-3-030-72657-7_22.
- [4] S. Namasudra, S. Dhamodharavadhani, and R. Rathipriya, "Nonlinear Neural Network Based Forecasting Model for Predicting COVID-19 Cases," *Neural Process Lett*, vol. 55, no. 1, pp. 171–191, Feb. 2023, doi: 10.1007/s11063-021-10495-w.
- [5] B. Awuku, Y. Huang, N. Yodo, and E. Asa, "Interpretable machine learning models for failure cause prediction in imbalanced oil pipeline data," *Meas Sci Technol*, vol. 35, no. 7, p. 076006, Jul. 2024, doi: 10.1088/1361-6501/ad3570.
- [6] G. Dlamini and M. Fahim, "DGM: a data generative model to improve minority class presence in anomaly detection domain," *Neural Comput Appl*, vol. 33, no. 20, pp. 13635–13646, Oct. 2021, doi: 10.1007/s00521-021-05993-w.
- [7] K. Fujiwara *et al.*, "Over- and Under-sampling Approach for Extremely Imbalanced and Small Minority Data Problem in Health Record Analysis," *Front Public Health*, vol. 8, May 2020, doi: 10.3389/fpubh.2020.00178.

- [8] M. H. L. Louk and B. A. Tama, "Exploring Ensemble-Based Class Imbalance Learners for Intrusion Detection in Industrial Control Networks," *Big Data and Cognitive Computing*, vol. 5, no. 4, p. 72, Dec. 2021, doi: 10.3390/bdcc5040072.
- [9] A. N. Tarekegn, M. Giacobini, and K. Michalak, "A review of methods for imbalanced multi-label classification," *Pattern Recognit*, vol. 118, p. 107965, Oct. 2021, doi: 10.1016/j.patcog.2021.107965.
- [10] J. P. Hoffmann, *Linear Regression Models*. Boca Raton: Chapman and Hall/CRC, 2021. doi: 10.1201/9781003162230.
- [11] T. Joyce, J. Donovan, and E. Murphy, "The application of the box-tidwell transformation in reliability modeling," in *RAMS '06. Annual Reliability and Maintainability Symposium, 2006.*, IEEE, 2006, pp. 196–200. doi: 10.1109/RAMS.2006.1677374.
- [12] L. Qu and Y. Pei, "A Comprehensive Review on Discriminant Analysis for Addressing Challenges of Class-Level Limitations, Small Sample Size, and Robustness," *Processes*, vol. 12, no. 7, p. 1382, Jul. 2024, doi: 10.3390/pr12071382.
- [13] A. Thakur, "Fundamentals of Neural Networks," *Int J Res Appl Sci Eng Technol*, vol. 9, no. VIII, pp. 407–426, Aug. 2021, doi: 10.22214/ijraset.2021.37362.
- [14] J. Balayla, "Performance Metrics of Binary Classifiers," in *Theorems on the Prevalence Threshold and the Geometry of Screening Curves*, Cham: Springer Nature Switzerland, 2024, pp. 129–142. doi: 10.1007/978-3-031-71452-8_10.
- [15] A. Santini, A. Man, and S. Voidăzan, "Accuracy of Diagnostic Tests," *The Journal of Critical Care Medicine*, vol. 7, no. 3, pp. 241–248, Jul. 2021, doi: 10.2478/jccm-2021-0022.
- [16] J. B. Brown, "Classifiers and their Metrics Quantified," *Mol Inform*, vol. 37, no. 1–2, Jan. 2018, doi: 10.1002/minf.201700127.
- [17] G. E. P. Box and P. W. Tidwell, "Transformation of the Independent Variables," *Technometrics*, vol. 4, no. 4, pp. 531–550, Nov. 1962, doi: 10.1080/00401706.1962.10490038.
- [18] A. Wilson and M. R. Anwar, "The Future of Adaptive Machine Learning Algorithms in High-Dimensional Data Processing," *International Transactions on Artificial Intelligence (ITALIC)*, vol. 3, no. 1, pp. 97–107, Nov. 2024, doi: 10.33050/italic.v3i1.656.

