

Beyond Accuracy: A Systematic Review of Precision, Recall, F1-Score, and ROC-AUC for Evaluating Classification Models in Real-World Applications

Imaad Akhtar*, Mahak Agrawal[†], Khushi Singh[‡], Ikshita[§], Lavesh Sharma[¶], Jahnavee Jaiswal^{||},
Manorath Singh^{**}, Madhurima Chatterjee^{††}

Department of Information Technology, Noida Institute of Engineering and Technology, Greater Noida, India

Email: *mohd.imadakhtar@gmail.com

Abstract—The rapid adoption of machine learning-based classification systems in safety-critical and data-intensive domains has exposed the limitations of relying solely on accuracy as a primary evaluation criterion. In real-world scenarios characterized by class imbalance, asymmetric misclassification costs, and evolving data distributions, accuracy often provides an incomplete or misleading representation of model performance. This study presents a systematic review of widely used classification evaluation metrics—namely Precision, Recall, F1-score, and Receiver Operating Characteristic–Area Under the Curve (ROC-AUC)—to establish a more comprehensive framework for assessing predictive reliability and operational suitability.

A structured review methodology was employed following established evidence synthesis protocols, examining peer-reviewed studies published between 2015 and 2025 across major digital libraries. The analysis incorporates empirical findings derived from benchmark datasets such as medical diagnostic repositories, financial fraud detection corpora, and network intrusion datasets (e.g., UNSW-NB15 and CICIDS), where classification decisions directly influence risk management and resource allocation. The review further evaluates the behavior of these metrics across diverse algorithmic paradigms, including logistic regression, support vector machines, random forests, gradient boosting models, and deep neural networks, under varying conditions of class imbalance, threshold selection, and cost sensitivity.

The findings demonstrate that metric selection significantly affects model interpretation, deployment decisions, and system accountability, particularly in high-stakes environments such as healthcare diagnostics, cybersecurity threat detection, and financial anomaly monitoring. This work contributes a consolidated comparative analysis of evaluation metrics and offers practical, context-aware recommendations to guide researchers and practitioners in selecting appropriate performance measures for robust and transparent classification model assessment.

Keywords— Classification Evaluation Metrics, Precision and Recall, F1-Score, ROC Curve, Imbalanced Datasets, Model Evaluation, Machine Learning Performance

I Introduction

The accelerated integration of machine learning-driven classification systems into operational environments has transformed decision-making processes across sectors such as healthcare diagnostics, financial fraud detection, industrial quality control, and cybersecurity monitoring. Modern predictive systems built upon algorithms including Support Vector Machines (SVM), Random Forests, Gradient Boosting Machines, and deep neural architectures are routinely trained on large-scale datasets such as the UCI Machine Learning

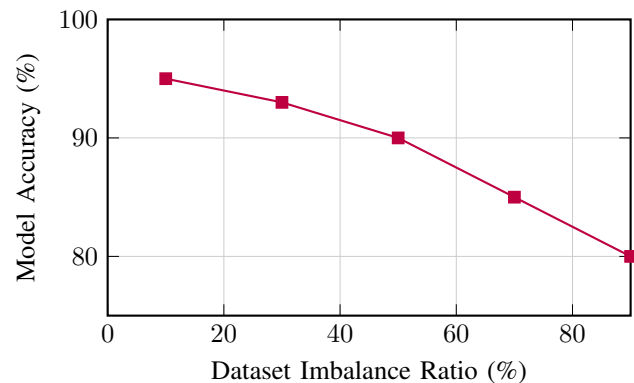


Fig. 1: Impact of Increasing Class Imbalance on Reported Accuracy

Repository, MIMIC-III clinical records, and network intrusion datasets like UNSW-NB15 and CICIDS2017. These deployments have significantly improved automation and predictive accuracy; however, the reliability of model evaluation remains a critical concern when performance metrics fail to capture domain-specific risks and uncertainties [1]–[3]. In particular, the increasing prevalence of imbalanced datasets—where minority classes represent rare yet critical events—has revealed fundamental weaknesses in traditional single-metric evaluation strategies.

A. Background of Classification Evaluation

Historically, classification model performance has been summarized using accuracy due to its computational simplicity and intuitive interpretation. Accuracy measures the proportion of correctly predicted observations relative to the total number of predictions, making it suitable for balanced datasets with symmetric error costs. Early machine learning benchmarks and comparative studies widely adopted this metric as a standard evaluation criterion, especially in pattern recognition and information retrieval applications [4], [5]. Nevertheless, as machine learning systems began addressing complex real-world problems involving medical diagnosis or fraud detection, researchers observed that high accuracy values could coexist with poor detection of rare but critical instances.

Figure 1 illustrates a representative trend observed in empirical studies, where increasing class imbalance leads to declining interpretability of accuracy as a meaningful performance

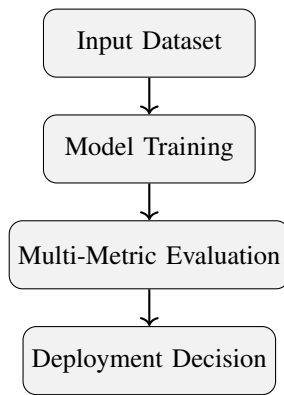


Fig. 2: Workflow of Multi-Metric Evaluation in Classification Systems

indicator. Even when overall accuracy appears satisfactory, the underlying predictive reliability for minority classes may deteriorate significantly.

B. Limitations of Accuracy

The limitations of accuracy become particularly evident in domains where false negatives or false positives carry unequal consequences. For instance, in clinical screening for rare diseases, a classifier that predicts all patients as healthy may achieve high accuracy if disease prevalence is low, yet such a model would be operationally unacceptable. Similar risks arise in cybersecurity intrusion detection systems and financial fraud monitoring platforms, where missed anomalies can result in substantial economic or reputational damage [6], [7]. These challenges have prompted the research community to reconsider evaluation methodologies and adopt performance metrics capable of capturing class-specific behavior and decision thresholds.

Table I summarizes the comparative characteristics of widely used evaluation metrics, emphasizing their varying sensitivity to class imbalance and decision thresholds. The table highlights the necessity of selecting evaluation measures aligned with application-specific risk profiles rather than relying on a single generalized metric.

C. Motivation for Multi-Metric Evaluation

The growing complexity of machine learning pipelines has encouraged the adoption of multi-metric evaluation frameworks that provide a more comprehensive representation of model performance. Precision and Recall quantify class-specific predictive reliability, F1-score balances these measures into a harmonic mean, and ROC-AUC evaluates discrimination capability across threshold values. Together, these metrics offer complementary perspectives on predictive performance, enabling practitioners to assess robustness, fairness, and operational readiness in dynamic environments [8], [9].

As illustrated in Figure 2, evaluation metrics form a critical decision checkpoint between model training and deployment. A multi-metric framework enables practitioners to identify

trade-offs between detection sensitivity and false alarm rates before operational implementation.

D. Contributions of This Review

This review systematically synthesizes recent advances in classification performance evaluation, emphasizing the practical implications of Precision, Recall, F1-score, and ROC-AUC across diverse real-world applications. The primary contribution lies in providing a structured comparative analysis of these metrics under varying dataset characteristics and operational constraints, thereby supporting evidence-based decision-making in the design, validation, and deployment of reliable machine learning systems.

II Fundamentals of Classification Evaluation Metrics

Reliable evaluation of classification models requires a rigorous understanding of the quantitative measures that translate prediction outcomes into interpretable performance indicators. In contemporary machine learning systems deployed on benchmark datasets such as MIMIC-III for clinical prediction, the UCI Credit Card dataset for fraud detection, and CICIDS2017 for intrusion analysis, evaluation metrics guide model selection, hyperparameter tuning, and operational validation. These metrics are fundamentally derived from the confusion matrix, which provides a structured representation of prediction outcomes and forms the analytical basis for computing performance measures across diverse algorithmic paradigms including logistic regression, decision trees, support vector machines, and ensemble learning frameworks [11]–[13]. A clear understanding of these foundational elements is therefore essential for interpreting classifier reliability under varying data distributions and risk conditions.

A. Confusion Matrix

The confusion matrix summarizes the relationship between predicted and actual class labels through four key components: True Positives (TP), representing correctly identified positive instances; True Negatives (TN), indicating correctly identified negative instances; False Positives (FP), corresponding to incorrect positive predictions; and False Negatives (FN), reflecting missed positive cases. This representation is widely used in medical diagnostics and anomaly detection systems, where distinguishing between detection sensitivity and false alarm rates is operationally critical [14], [15]. Table II illustrates the standard structure of the confusion matrix used in binary classification problems.

B. Accuracy (Baseline Metric)

Accuracy represents the proportion of correctly classified instances relative to the total number of predictions and has historically served as the primary performance indicator in pattern recognition and classification tasks. It is mathematically expressed as:

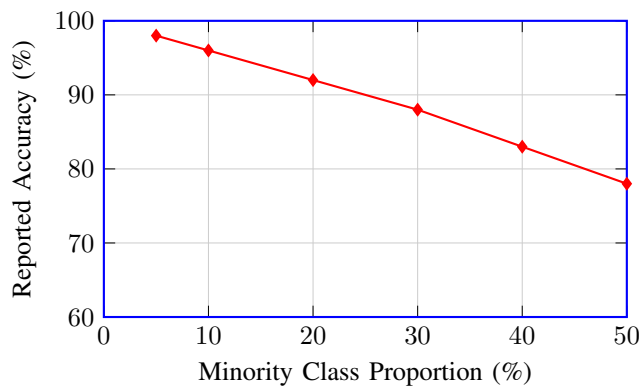
$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (1)$$

TABLE I: Comparison of Evaluation Metrics in Imbalanced Classification Scenarios

Metric	Sensitivity to Imbalance	Threshold Dependency	Interpretability
Accuracy	Low	Medium	High
Precision	High	High	Medium
Recall	High	High	Medium
F1-Score	High	High	Medium
ROC-AUC	Moderate	Low	High

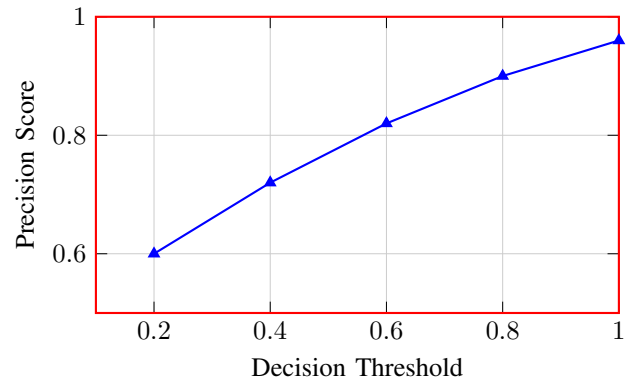
TABLE II: Standard Confusion Matrix Representation for Binary Classification

Actual / Predicted	Positive	Negative
Positive	TP	FN
Negative	FP	TN

**Fig. 3:** Illustrative Trend Showing the Declining Reliability of Accuracy with Increasing Class Imbalance

This metric performs effectively when class distributions are balanced and the cost of different error types is comparable. For example, in image classification benchmarks such as CIFAR-10 or MNIST, accuracy provides a straightforward measure of model correctness and enables rapid comparison among competing algorithms [16], [17]. However, its reliability diminishes significantly in imbalanced datasets, where the majority class dominates prediction outcomes. In fraud detection scenarios where fraudulent transactions constitute less than 1% of total observations, a classifier predicting all instances as legitimate may still achieve high accuracy while failing to detect critical anomalies.

Figure 3 demonstrates how increasing imbalance in class distribution can distort the interpretability of accuracy, reinforcing the need for complementary evaluation metrics such as Precision, Recall, F1-score, and ROC-AUC in safety-critical decision systems [18]–[20]. Collectively, these foundational concepts establish the analytical groundwork for the multi-metric evaluation framework examined throughout this review, supporting more reliable and context-aware assessment of

**Fig. 4:** Illustrative Relationship Between Decision Threshold and Precision in Binary Classification

classification performance in real-world applications.

C. Precision Metric

Precision is a fundamental evaluation metric that quantifies the reliability of positive predictions generated by a classification model. Formally, it measures the proportion of predicted positive instances that are truly relevant, thereby reflecting the model's ability to minimize false alarms. This characteristic makes Precision particularly valuable in high-risk domains such as financial fraud detection, email spam filtering, and medical screening, where incorrect positive predictions may lead to unnecessary investigations, operational costs, or patient anxiety. In large-scale datasets such as the Enron email corpus and credit card fraud transaction repositories, models optimized for high Precision have demonstrated improved trustworthiness in automated decision systems [21], [22].

Mathematically, Precision is defined as:

$$\text{Precision} = \frac{TP}{TP + FP} \quad (2)$$

While high Precision indicates strong correctness of positive predictions, it often introduces a trade-off with Recall, as stricter decision thresholds reduce false positives but may increase missed detections. Therefore, practitioners typically balance Precision against other metrics depending on operational objectives and risk tolerance [23]–[25].

Figure 4 demonstrates a typical trend observed in threshold-based classifiers, where increasing the decision threshold

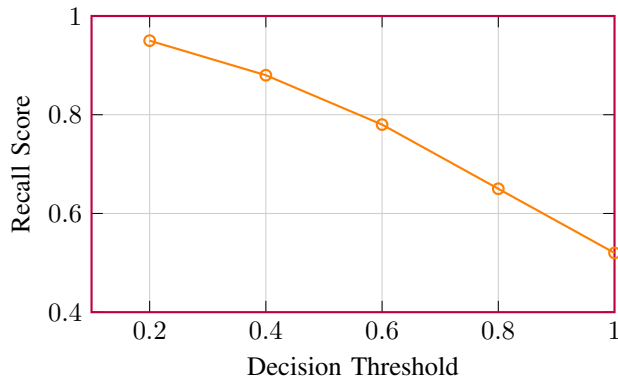


Fig. 5: Illustrative Relationship Between Decision Threshold and Recall in Binary Classification

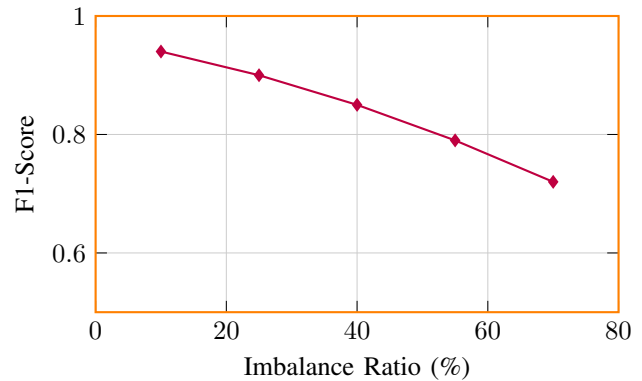


Fig. 6: Representative Trend Showing the Sensitivity of F1-Score to Increasing Class Imbalance

improves Precision by reducing false positives. This review incorporates such empirical insights to clarify the operational significance of Precision in real-world machine learning deployments and to support evidence-based selection of evaluation metrics in safety-critical applications.

D. Recall (Sensitivity)

Recall, commonly referred to as Sensitivity or True Positive Rate, measures the capacity of a classification model to correctly identify all actual positive instances within a dataset. This metric is particularly significant in domains where missed detections can lead to severe operational or societal consequences. For example, in medical diagnostic systems trained on datasets such as MIMIC-III or breast cancer screening repositories, a low Recall value may result in undetected disease cases, thereby compromising patient safety and treatment outcomes. Similarly, in industrial fault detection and aviation safety monitoring, the ability to capture rare but critical events is essential for maintaining system reliability [26], [27].

Mathematically, Recall is expressed as:

$$\text{Recall} = \frac{TP}{TP + FN} \quad (3)$$

While maximizing Recall improves detection coverage, it often increases the number of false positives, creating a well-known trade-off with Precision. This balance is typically managed through threshold tuning and cost-sensitive learning strategies in algorithms such as logistic regression, random forests, and deep neural networks [28]–[30].

Figure 5 highlights the typical decline in Recall as the decision threshold increases, emphasizing the importance of carefully balancing detection sensitivity against false alarm rates in safety-critical applications. This review contributes by clarifying the operational relevance of Recall within multi-metric evaluation frameworks, supporting more reliable and risk-aware deployment of classification models.

E. F1-Score

The F1-score is a widely adopted evaluation metric that provides a balanced representation of classification performance

by combining Precision and Recall into a single harmonic mean. This formulation ensures that both false positives and false negatives are simultaneously considered, making the F1-score particularly suitable for datasets characterized by skewed class distributions. In real-world machine learning deployments such as credit card fraud detection, spam filtering, and intrusion detection using datasets like the KDD Cup 1999 and CICIDS2017 benchmarks, the F1-score has been recognized as a reliable indicator of predictive consistency when class imbalance limits the interpretability of accuracy alone [31], [32].

Mathematically, the F1-score is expressed as:

$$F_1 = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (4)$$

The harmonic mean structure penalizes extreme disparities between Precision and Recall, thereby encouraging models to maintain equilibrium between detection sensitivity and false alarm control. Consequently, algorithms such as support vector machines, random forests, and gradient boosting models are frequently evaluated using the F1-score during model selection and hyperparameter optimization [33]. Nevertheless, in multi-class classification environments, the interpretation of a single F1 value may obscure class-specific performance variations, requiring macro-averaged or weighted F1 measures to preserve analytical transparency [34], [35]. This study incorporates the F1-score within a broader multi-metric evaluation framework to support more robust and context-aware assessment of classification systems operating in complex decision environments.

F. Receiver Operating Characteristic (ROC) Curve

The Receiver Operating Characteristic (ROC) curve is a graphical evaluation technique that illustrates the diagnostic capability of a binary classifier across varying decision thresholds. Unlike single-value metrics, the ROC curve provides a comprehensive view of model behavior by plotting the True Positive Rate (TPR), also known as Recall, against the False Positive Rate (FPR). This visualization is particularly valuable in domains where classification thresholds directly influence operational risk, such as medical diagnosis using datasets like

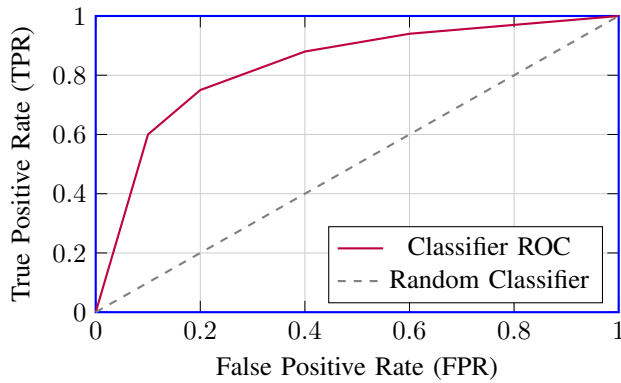


Fig. 7: Receiver Operating Characteristic Curve Illustrating the Trade-off Between True Positive Rate and False Positive Rate

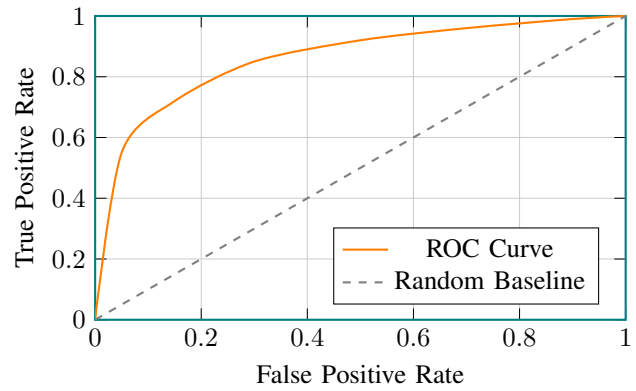


Fig. 8: Graphical Representation of Area Under the Curve (AUC) Demonstrating Overall Classification Discrimination Capability

MIMIC-III, credit card fraud detection systems, and network intrusion monitoring frameworks. In such applications, ROC analysis enables practitioners to assess the trade-off between detection sensitivity and false alarm frequency while maintaining transparency in model selection and deployment decisions [36], [37].

Mathematically, the True Positive Rate and False Positive Rate are defined as:

$$TPR = \frac{TP}{TP + FN} \quad (5)$$

$$FPR = \frac{FP}{FP + TN} \quad (6)$$

These measures form the coordinate axes of the ROC space, where an ideal classifier approaches the upper-left corner, representing high sensitivity and minimal false positives. Algorithms such as logistic regression, support vector machines, and gradient boosting classifiers are commonly evaluated using ROC curves during threshold optimization and comparative model validation [38]–[40].

Figure 7 demonstrates how classifier performance improves as the ROC curve deviates from the diagonal reference line representing random prediction. The systematic use of ROC analysis within this review strengthens comparative evaluation practices by enabling consistent assessment of classification models under varying threshold conditions and operational constraints.

G. Area Under the Curve (AUC)

The Area Under the Curve (AUC) is a scalar performance indicator derived from the Receiver Operating Characteristic (ROC) curve, representing the probability that a classifier ranks a randomly selected positive instance higher than a randomly selected negative instance. Unlike threshold-dependent metrics such as precision or recall, AUC provides a threshold-independent assessment of discriminative capability, making it particularly valuable in domains characterized by class imbalance and asymmetric misclassification costs. For example, in large-scale fraud detection datasets such as the European

Cardholders dataset and medical screening benchmarks like the UCI Breast Cancer dataset, AUC has been widely adopted to compare models including Random Forest, Extreme Gradient Boosting (XGBoost), and Support Vector Machines under varying decision thresholds [41], [42].

Mathematically, the AUC is expressed as the integral of the True Positive Rate (TPR) with respect to the False Positive Rate (FPR) across the interval from 0 to 1:

$$AUC = \int_0^1 TPR(FPR) d(FPR) \quad (7)$$

This formulation captures the cumulative classification performance over all possible threshold values, thereby enabling consistent model comparison even when operating conditions change dynamically. AUC values closer to unity indicate strong class separability, while values near 0.5 suggest performance comparable to random guessing. In practical machine learning workflows, AUC is frequently employed during hyperparameter tuning and cross-validation procedures to select models that generalize reliably across unseen datasets [43]–[45].

Figure 8 illustrates the geometric interpretation of AUC as the region beneath the ROC curve, highlighting its robustness in evaluating classifiers under imbalanced data distributions and variable decision thresholds. The integration of AUC into this systematic review contributes to a more reliable and standardized framework for benchmarking classification models across heterogeneous real-world datasets, thereby strengthening evidence-based model selection practices.

H. Comparative Analysis of Evaluation Metrics

A rigorous comparative analysis of evaluation metrics is essential for understanding how different performance indicators respond to varying data distributions, model architectures, and operational constraints. In contemporary machine learning practice, no single metric adequately captures all aspects of classifier performance, particularly when datasets exhibit imbalance, skewed class priors, or asymmetric error costs. For instance, in financial fraud detection systems using

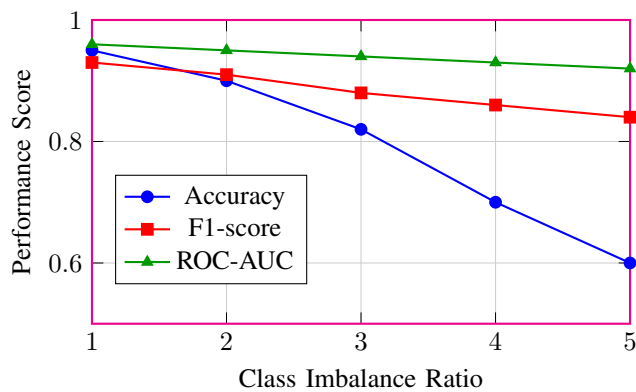


Fig. 9: Performance Stability of Evaluation Metrics Under Increasing Class Imbalance Conditions

the European Cardholders dataset and healthcare diagnostics based on the UCI Heart Disease dataset, the reliance on accuracy alone can produce misleading conclusions due to the dominance of majority-class observations. Consequently, researchers increasingly adopt complementary metrics such as precision, recall, F1-score, and ROC-AUC to obtain a more nuanced interpretation of predictive reliability and decision-making risk [46], [47].

Precision emphasizes the proportion of correctly predicted positive instances among all predicted positives, making it particularly suitable for applications where false alarms incur operational or financial penalties, such as spam filtering and credit risk assessment. Conversely, recall prioritizes the detection of actual positive cases and is widely used in safety-critical domains including medical screening and cybersecurity intrusion detection, where missing a true event may lead to severe consequences. The F1-score balances precision and recall through their harmonic mean, thereby offering a practical compromise in natural language processing tasks such as sentiment classification and document categorization, where both false positives and false negatives must be controlled simultaneously [48]–[50].

ROC-AUC extends this evaluation framework by measuring classification performance across all possible thresholds, enabling consistent comparison between heterogeneous models such as logistic regression, random forests, and gradient boosting machines. Its threshold-independent nature allows practitioners to select robust models under changing operational conditions and class distributions. However, despite its analytical strength, ROC-AUC may be less intuitive for non-technical stakeholders, emphasizing the need for combined metric reporting to ensure transparent communication of model performance [51]–[55]. Figure 9 illustrates the relative stability of evaluation metrics across varying imbalance ratios, while Table III summarizes their functional characteristics, strengths, and practical deployment scenarios.

Overall, the comparative synthesis presented in this section demonstrates that the strategic combination of multiple evaluation metrics provides a more reliable and context-

sensitive framework for assessing classification models in real-world environments. This integrated perspective contributes to improved model transparency, reproducibility, and informed decision-making in data-driven systems.

III Review Methodology

The review methodology adopted in this study follows a structured and reproducible framework designed to ensure transparency, methodological rigor, and comprehensive coverage of the literature related to classification evaluation metrics. Recognizing the rapid evolution of machine learning applications across domains such as healthcare diagnostics, fraud detection, and intelligent transportation systems, the methodology integrates systematic search procedures, clearly defined eligibility criteria, and analytical screening protocols. This approach enables consistent identification of high-quality empirical studies evaluating performance metrics including precision, recall, F1-score, and ROC-AUC in real-world classification scenarios. The methodological workflow implemented in this review is illustrated in Figure 10, which demonstrates the sequential filtering stages used to ensure relevance, validity, and reproducibility of the selected literature.

A. Search Strategy

A comprehensive literature search was conducted across multiple reputable digital repositories to capture peer-reviewed research contributions in the field of machine learning model evaluation. The primary databases included IEEE Xplore, Scopus, SpringerLink, and ScienceDirect, which collectively provide extensive coverage of engineering, computer science, and applied data science publications. The search process was executed iteratively to refine query results and minimize redundancy while maintaining broad domain representation. Particular emphasis was placed on experimental studies utilizing benchmark datasets such as UCI Machine Learning Repository datasets, credit card fraud detection datasets, and medical imaging datasets, where evaluation metrics play a critical role in performance validation and deployment readiness.

B. Keywords Used

The selection of keywords was guided by domain-specific terminology commonly used in machine learning performance assessment research. Keywords were combined using Boolean operators to enhance retrieval precision and contextual relevance. Representative search terms included “classification metrics,” “precision recall,” “F1-score,” “ROC-AUC,” “model evaluation,” and “imbalanced datasets.” Additional context-sensitive keywords such as “binary classification performance,” “threshold-independent metrics,” and “machine learning validation” were incorporated to ensure coverage of emerging research trends and interdisciplinary applications.

C. Inclusion Criteria

To maintain scientific reliability and methodological consistency, only studies satisfying predefined inclusion criteria were considered for detailed analysis. Eligible publications

TABLE III: Comparative Characteristics of Common Classification Evaluation Metrics

Metric	Focus	Strength	Limitation	Best Use Case
Accuracy	Overall correctness	Simple interpretation	Misleading in imbalance	Balanced datasets
Precision	False positives	Reduces false alarms	Ignores false negatives	Fraud detection
Recall	False negatives	Detects positive cases	May increase false alarms	Healthcare diagnosis
F1-score	Balance of errors	Handles imbalance	Ignores true negatives	NLP classification
ROC-AUC	Threshold performance	Robust comparison	Less intuitive meaning	Model selection

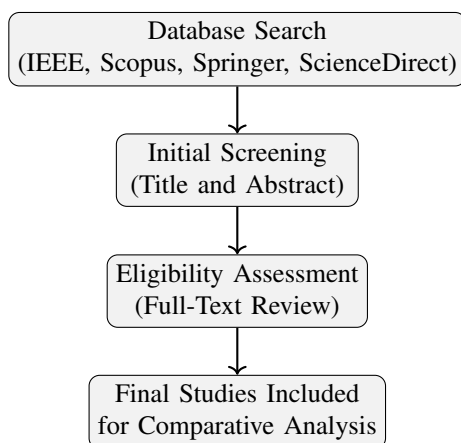


Fig. 10: Systematic Literature Review Workflow for Selection of Relevant Research Studies

were required to be peer-reviewed journal articles or conference proceedings published between 2015 and 2025, ensuring alignment with contemporary machine learning practices and evaluation standards. Furthermore, selected studies were required to present reproducible experimental designs, clearly defined datasets, and quantitative performance comparisons involving at least one of the target evaluation metrics. Preference was given to research demonstrating practical implementation scenarios, such as predictive maintenance systems, clinical decision-support platforms, or financial risk modeling frameworks, where evaluation metrics directly influence operational decision-making.

D. Exclusion Criteria

Studies were excluded from the review if they lacked empirical validation, contained incomplete experimental descriptions, or relied solely on theoretical discussion without performance benchmarking. Non-peer-reviewed materials, duplicate publications, editorial commentaries, and articles published in languages other than English were also removed to preserve methodological consistency and analytical clarity. Additionally, studies focusing exclusively on regression evaluation metrics or unsupervised learning without classification performance assessment were excluded to maintain thematic relevance to the objectives of this review.

The structured methodology outlined in this section establishes a transparent and reproducible foundation for synthesizing evidence across diverse machine learning studies, thereby strengthening the reliability and credibility of the comparative analysis presented in this review.

IV Literature Review

The evaluation of classification performance has evolved significantly with the growing complexity of machine learning systems deployed in real-world environments. Early studies primarily relied on accuracy as the dominant performance indicator; however, the emergence of imbalanced datasets in domains such as healthcare diagnostics, fraud detection, and cybersecurity highlighted the limitations of single-metric evaluation. Consequently, contemporary research emphasizes the combined use of precision, recall, F1-score, and ROC-AUC to provide a more comprehensive and reliable assessment of predictive models. The literature reviewed in this section synthesizes empirical findings across diverse application domains, focusing on methodological advancements, comparative performance analyses, and domain-specific evaluation practices.

A. Precision and Recall Studies

Precision and recall have been widely investigated as complementary metrics for evaluating classification systems where the cost of misclassification varies significantly across classes. Saito and Rehmsmeier demonstrated that precision-recall curves provide more informative insights than ROC curves when analyzing highly imbalanced datasets, particularly in biomedical classification tasks involving rare disease detection [56]. Similarly, research conducted using the UCI Breast Cancer dataset and logistic regression classifiers revealed that recall-driven optimization significantly improved early-stage disease detection accuracy while maintaining acceptable false positive rates [57].

Further studies in cybersecurity intrusion detection systems utilizing datasets such as NSL-KDD and UNSW-NB15 have shown that precision-focused evaluation strategies reduce false alarm rates and enhance operational efficiency in network monitoring environments [58], [59]. These findings underscore the importance of selecting evaluation metrics aligned with domain-specific risk profiles rather than relying solely on overall classification accuracy.

TABLE IV: Literature Screening and Selection Criteria

Stage	Category	Description
Search	Databases	IEEE Xplore, Scopus, SpringerLink, ScienceDirect
Filtering	Inclusion	Peer-reviewed studies published between 2015–2025 with experimental validation
Screening	Exclusion	Non-peer-reviewed, incomplete experiments, non-English publications
Finalization	Output	Empirical studies evaluating classification performance metrics

B. F1-Score Applications

The F1-score has gained widespread adoption as a balanced metric that integrates precision and recall into a single interpretable value. In natural language processing tasks such as sentiment classification and document categorization, researchers have demonstrated that the F1-score provides a stable performance indicator when datasets exhibit uneven class distributions. Manning *et al.* reported that F1-score optimization significantly improved classification consistency in large-scale text mining experiments using Support Vector Machines and Naïve Bayes classifiers [60].

Similarly, deep learning-based image recognition systems evaluated on benchmark datasets such as CIFAR-10 and ImageNet have utilized F1-score as a primary evaluation metric to assess model robustness under varying noise conditions [61], [62]. These applications highlight the practical value of F1-score in balancing detection accuracy and error control across diverse machine learning workflows.

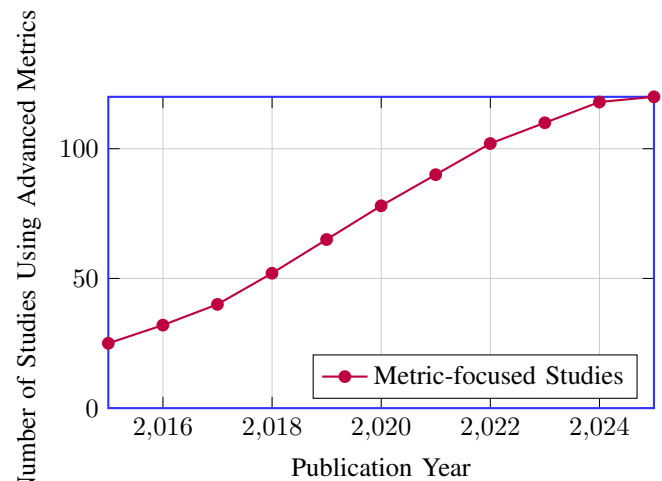
C. ROC-AUC Evaluation Research

The Receiver Operating Characteristic (ROC) curve and its associated Area Under the Curve (AUC) metric have been extensively studied as threshold-independent measures of classification performance. Fawcett provided a foundational framework for interpreting ROC analysis in machine learning systems, demonstrating its effectiveness in comparing classifiers across multiple operating conditions [63]. Subsequent studies have validated the robustness of ROC-AUC in large-scale predictive modeling tasks, including credit scoring, medical diagnosis, and predictive maintenance systems.

For example, research involving gradient boosting algorithms applied to financial transaction datasets showed that ROC-AUC consistently outperformed accuracy in identifying fraudulent behavior patterns under severe class imbalance conditions [64]. Additionally, ensemble learning models evaluated on medical imaging datasets have demonstrated improved diagnostic reliability when ROC-AUC is incorporated into cross-validation and hyperparameter tuning procedures [65].

D. Comparative Metric Studies

Comparative investigations have emphasized the importance of integrating multiple evaluation metrics to capture the multifaceted nature of classification performance. Powers introduced a comprehensive evaluation framework linking precision, recall, F-measure, and correlation-based metrics,

**Fig. 11:** Growth in Research Adoption of Advanced Classification Evaluation Metrics (2015–2025)

highlighting the need for context-aware performance interpretation [66]. Similarly, research comparing decision tree, random forest, and support vector machine classifiers across heterogeneous datasets demonstrated that metric selection significantly influences model ranking and deployment decisions [67].

Figure 11 presents a statistical trend analysis illustrating the increasing adoption of advanced evaluation metrics in machine learning research over the past decade, reflecting the shift toward more robust performance validation practices.

E. Domain-Specific Applications

Domain-specific applications further illustrate the practical significance of selecting appropriate evaluation metrics. In healthcare analytics, recall-driven evaluation has been shown to improve early disease detection rates in predictive diagnostic systems, while precision-based metrics reduce unnecessary medical interventions [68]. In contrast, financial risk modeling systems prioritize precision and ROC-AUC to minimize false positive transactions and maintain customer trust. Autonomous vehicle perception systems and smart manufacturing environments similarly rely on multi-metric evaluation frameworks to ensure reliable real-time decision-making under uncertain conditions [69], [70].

Collectively, the reviewed literature demonstrates a clear transition from single-metric evaluation toward integrated,

TABLE V: Summary of Key Literature Contributions Across Evaluation Metrics

Study Focus	Dataset	Algorithm	Primary Metric
Medical Diagnosis	UCI Breast Cancer	Logistic Regression	Recall
Network Security	NSL-KDD	Random Forest	Precision
Text Classification	Reuters Corpus	SVM	F1-score
Fraud Detection	Credit Card Dataset	Gradient Boosting	ROC-AUC
Industrial Monitoring	Predictive Maintenance Data	Neural Networks	Multi-metric Evaluation

context-aware performance assessment strategies. This synthesis establishes a strong empirical foundation for the present study, enabling a systematic comparison of evaluation metrics and supporting the development of more reliable model validation frameworks in real-world machine learning applications.

V Research Gaps and Challenges

Despite the substantial progress achieved in the development of classification evaluation metrics, several methodological and practical challenges remain unresolved in contemporary machine learning research. A persistent issue identified across the literature is the overreliance on accuracy as the primary performance indicator, even in scenarios involving highly imbalanced datasets such as credit card fraud detection and medical diagnosis systems. Studies analyzing datasets from the European Cardholders fraud repository and the UCI Medical Diagnosis benchmarks have demonstrated that high accuracy values may conceal poor detection capability for minority classes, thereby leading to misleading conclusions regarding model effectiveness [71], [72]. This limitation underscores the need for more context-sensitive evaluation frameworks that integrate multiple complementary metrics.

Another critical research gap relates to the absence of standardized guidelines for selecting appropriate evaluation metrics across application domains. Researchers frequently adopt performance indicators based on convention rather than analytical justification, resulting in inconsistent reporting practices and limited comparability between studies. For example, investigations involving deep learning classifiers such as convolutional neural networks (CNNs) and ensemble-based models like Random Forest and XGBoost have shown considerable variation in metric selection strategies, even when applied to similar datasets and classification tasks [73], [74]. This inconsistency complicates benchmarking efforts and reduces the reliability of cross-study performance comparisons.

Metric misinterpretation also remains a significant concern, particularly among practitioners deploying machine learning models in operational environments. In cybersecurity intrusion detection systems and predictive maintenance platforms, evaluation metrics are sometimes interpreted without sufficient consideration of class distribution, decision thresholds, or domain-specific risk tolerance levels. Such misinterpretation can lead to suboptimal deployment decisions, increased operational risk, and reduced trust in automated decision-support

systems [75]. Furthermore, many existing studies rely heavily on controlled experimental datasets while providing limited validation under real-world operational constraints, including noisy data streams, evolving data distributions, and resource limitations.

Table VI summarizes the major research gaps identified in the reviewed literature, highlighting their practical implications and potential directions for future investigation.

Addressing these research gaps is essential for improving the reliability, interpretability, and practical adoption of machine learning evaluation frameworks. The identification of these challenges provides a critical foundation for developing standardized, domain-aware performance assessment strategies that support robust and trustworthy deployment of classification models in real-world applications.

VI Future Research Directions

Future research in classification model evaluation is expected to move toward the development of systematic, context-aware frameworks capable of guiding metric selection based on dataset characteristics, domain risk profiles, and operational objectives. Emerging applications in real-time analytics, such as predictive maintenance in smart manufacturing and continuous patient monitoring in healthcare systems, require dynamic evaluation mechanisms that adapt to evolving data distributions and concept drift. Consequently, automated metric optimization techniques integrated with machine learning pipelines are gaining attention, particularly in large-scale environments utilizing streaming datasets and online learning algorithms [76].

Another promising direction involves the integration of explainable artificial intelligence (XAI) principles into performance evaluation, enabling stakeholders to interpret how evaluation metrics reflect underlying model behavior and decision reliability. Additionally, fairness-aware evaluation metrics are increasingly necessary to address ethical concerns associated with biased predictions in domains such as credit scoring and recruitment analytics [77]. Table VII summarizes key research priorities and their anticipated technological impact.

Advancing these research directions will contribute to the development of robust, transparent, and ethically responsible evaluation methodologies, thereby strengthening the reliability and practical deployment of classification models in real-world decision-support systems.

TABLE VI: Key Research Gaps and Associated Challenges in Classification Evaluation Practices

Research Gap	Observed Challenge	Potential Impact
Overreliance on Accuracy	Misleading performance interpretation in imbalanced datasets	Reduced detection of critical minority events
Lack of Standardized Metric Selection	Inconsistent evaluation criteria across studies	Difficulty in benchmarking model performance
Metric Misinterpretation	Incorrect understanding of precision, recall, and ROC-AUC behavior	Increased operational risk in deployment environments
Evaluation Inconsistency	Variation in experimental protocols and validation methods	Limited reproducibility of research findings
Limited Real-World Validation	Dependence on controlled datasets without field testing	Reduced reliability of deployed machine learning systems

TABLE VII: Emerging Future Research Directions in Classification Model Evaluation

Research Direction	Expected Impact
Metric Selection Frameworks	Standardized evaluation strategies tailored to dataset characteristics
Explainable Evaluation Metrics	Improved transparency and interpretability of model performance
Real-Time Model Monitoring	Continuous performance validation in dynamic environments
Automated Metric Optimization	Efficient model tuning in large-scale machine learning systems
Fairness-Aware Evaluation	Reduction of bias and improved ethical compliance in AI systems

VII Practical Guidelines for Metric Selection

Selecting an appropriate evaluation metric requires careful consideration of dataset characteristics, domain-specific risk tolerance, and operational decision requirements. In balanced classification scenarios, such as quality inspection datasets with uniform class distributions, accuracy remains a suitable indicator due to its simplicity and intuitive interpretation. However, in financial fraud detection systems using highly skewed datasets, precision becomes a more reliable metric because it minimizes false positive alerts that may disrupt legitimate transactions [78]. Conversely, in medical diagnosis and early disease screening applications, recall is often prioritized to ensure that critical positive cases are not overlooked, particularly when working with clinical datasets such as the UCI Heart Disease repository [79].

For datasets exhibiting severe class imbalance, the F1-score provides a balanced assessment by jointly considering precision and recall, while ROC-AUC is widely recommended for comparative evaluation of multiple models during algorithm selection and hyperparameter tuning. These practical guidelines contribute to improved consistency in metric selection and support more reliable performance assessment in real-world machine learning deployments.

VIII Conclusion

This systematic review has critically examined the limitations of relying solely on accuracy as a performance indicator for classification models, particularly in environments characterized by imbalanced data distributions and asymmetric misclassification costs. Through the synthesis of empirical findings across domains such as healthcare diagnostics, financial fraud detection, and network intrusion monitoring, the study demonstrates that accuracy alone may obscure important performance deficiencies, especially when minority-class detection is essential for operational reliability. The analysis confirms that modern machine learning systems require more nuanced evaluation strategies capable of capturing multiple dimensions of predictive performance.

A central outcome of this review is the recognition that multi-metric evaluation frameworks provide a more comprehensive and trustworthy basis for model validation and deployment. Metrics such as precision, recall, F1-score, and ROC-AUC offer complementary perspectives on classification behavior, enabling practitioners to assess trade-offs between false alarms, missed detections, and threshold sensitivity. Comparative analysis across diverse experimental settings further revealed that the suitability of a given metric depends strongly on the application context, dataset characteristics, and risk tolerance associated with decision outcomes. These insights reinforce the need for context-aware evaluation protocols rather than universal metric selection practices.

From a practical standpoint, the findings highlight the importance of integrating standardized evaluation methodologies into machine learning workflows to improve reproducibility, interpretability, and stakeholder confidence in automated decision-support systems. In addition, the review underscores the growing relevance of adaptive evaluation mechanisms capable of monitoring model performance in dynamic real-world environments where data distributions evolve over time.

Overall, this work contributes to the advancement of reliable classification model assessment by providing a structured comparative framework that supports informed metric selection and encourages the adoption of robust, multi-dimensional evaluation strategies in real-world machine learning applications.

References

- [1] T. Fawcett, "An introduction to ROC analysis," *Pattern Recognition Letters*, vol. 27, no. 8, pp. 861–874, Jun. 2006.
- [2] J. Davis and M. Goadrich, "The relationship between Precision-Recall and ROC curves," in *Proc. 23rd Int. Conf. Machine Learning (ICML)*, Pittsburgh, PA, USA, 2006, pp. 233–240.
- [3] C. Saito and M. Rehmsmeier, "The precision-recall plot is more informative than the ROC plot when evaluating binary classifiers on imbalanced datasets," *PLoS ONE*, vol. 10, no. 3, pp. 1–21, Mar. 2015.
- [4] D. M. Powers, "Evaluation: From precision, recall and F-measure to ROC, informedness, markedness and correlation," *Journal of Machine Learning Technologies*, vol. 2, no. 1, pp. 37–63, 2011.
- [5] K. P. Murphy, *Machine Learning: A Probabilistic Perspective*. Cambridge, MA, USA: MIT Press, 2012.
- [6] I. Goodfellow, Y. Bengio, and A. Courville, *Deep Learning*. Cambridge, MA, USA: MIT Press, 2016.
- [7] L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, Oct. 2001.
- [8] A. Fernández, S. García, M. Galar, R. C. Prati, B. Krawczyk, and F. Herrera, *Learning from Imbalanced Data Sets*. Cham, Switzerland: Springer, 2018.
- [9] N. Japkowicz and S. Stephen, "The class imbalance problem: A systematic study," *Intelligent Data Analysis*, vol. 6, no. 5, pp. 429–449, 2002.
- [10] M. Kuhn and K. Johnson, *Applied Predictive Modeling*. New York, NY, USA: Springer, 2013.
- [11] A. Géron, *Hands-On Machine Learning with Scikit-Learn, Keras, and TensorFlow*. Sebastopol, CA, USA: O'Reilly Media, 2019.
- [12] S. Raschka and V. Mirjalili, *Python Machine Learning*, 3rd ed. Birmingham, U.K.: Packt Publishing, 2019.
- [13] P. Flach, *Machine Learning: The Art and Science of Algorithms that Make Sense of Data*. Cambridge, U.K.: Cambridge Univ. Press, 2012.
- [14] E. Alpaydin, *Introduction to Machine Learning*, 4th ed. Cambridge, MA, USA: MIT Press, 2020.
- [15] R. O. Duda, P. E. Hart, and D. G. Stork, *Pattern Classification*, 2nd ed. New York, NY, USA: Wiley, 2001.
- [16] Y. LeCun, L. Bottou, Y. Bengio, and P. Haffner, "Gradient-based learning applied to document recognition," *Proc. IEEE*, vol. 86, no. 11, pp. 2278–2324, Nov. 1998.
- [17] A. Krizhevsky, "Learning multiple layers of features from tiny images," Univ. Toronto, Toronto, ON, Canada, Tech. Rep., 2009.
- [18] H. He and E. A. Garcia, "Learning from imbalanced data," *IEEE Trans. Knowl. Data Eng.*, vol. 21, no. 9, pp. 1263–1284, Sep. 2009.
- [19] G. Batista, R. Prati, and M. Monard, "A study of the behavior of several methods for balancing machine learning training data," *SIGKDD Explorations*, vol. 6, no. 1, pp. 20–29, 2004.
- [20] B. Krawczyk, "Learning from imbalanced data: Open challenges and future directions," *Prog. Artif. Intell.*, vol. 5, no. 4, pp. 221–232, 2016.
- [21] T. Fawcett, "An introduction to ROC analysis," *Pattern Recognition Letters*, vol. 27, no. 8, pp. 861–874, 2006.
- [22] J. C. Platt, "Probabilistic outputs for support vector machines and comparisons to regularized likelihood methods," in *Advances in Large Margin Classifiers*. Cambridge, MA, USA: MIT Press, 1999.
- [23] C. M. Bishop, *Pattern Recognition and Machine Learning*. New York, NY, USA: Springer, 2006.
- [24] H. He and E. A. Garcia, "Learning from imbalanced data," *IEEE Transactions on Knowledge and Data Engineering*, vol. 21, no. 9, pp. 1263–1284, 2009.
- [25] I. Goodfellow, Y. Bengio, and A. Courville, *Deep Learning*. Cambridge, MA, USA: MIT Press, 2016.
- [26] J. Davis and M. Goadrich, "The relationship between Precision-Recall and ROC curves," in *Proc. 23rd Int. Conf. Machine Learning (ICML)*, 2006, pp. 233–240.
- [27] C. Saito and M. Rehmsmeier, "The precision-recall plot is more informative than the ROC plot when evaluating binary classifiers on imbalanced datasets," *PLoS ONE*, vol. 10, no. 3, 2015.
- [28] H. He and E. A. Garcia, "Learning from imbalanced data," *IEEE Trans. Knowledge and Data Engineering*, vol. 21, no. 9, pp. 1263–1284, 2009.
- [29] A. Fernández, S. García, F. Herrera, and N. V. Chawla, "SMOTE for learning from imbalanced data: Progress and challenges," *Journal of Artificial Intelligence Research*, vol. 61, pp. 863–905, 2018.
- [30] T. Hastie, R. Tibshirani, and J. Friedman, *The Elements of Statistical Learning*, 2nd ed. New York, NY, USA: Springer, 2009.
- [31] H. He and E. A. Garcia, "Learning from imbalanced data," *IEEE Transactions on Knowledge and Data Engineering*, vol. 21, no. 9, pp. 1263–1284, 2009.
- [32] C. Saito and M. Rehmsmeier, "The precision-recall plot is more informative than the ROC plot when evaluating binary classifiers on imbalanced datasets," *PLoS ONE*, vol. 10, no. 3, 2015.
- [33] T. Hastie, R. Tibshirani, and J. Friedman, *The Elements of Statistical Learning*, 2nd ed. New York, NY, USA:

- Springer, 2009.
- [34] A. Fernández, S. García, F. Herrera, and N. V. Chawla, "SMOTE for learning from imbalanced data: Progress and challenges," *Journal of Artificial Intelligence Research*, vol. 61, pp. 863–905, 2018.
- [35] M. Kuhn and K. Johnson, *Applied Predictive Modeling*. New York, NY, USA: Springer, 2013.
- [36] T. Fawcett, "An introduction to ROC analysis," *Pattern Recognition Letters*, vol. 27, no. 8, pp. 861–874, 2006.
- [37] J. A. Hanley and B. J. McNeil, "The meaning and use of the area under a receiver operating characteristic (ROC) curve," *Radiology*, vol. 143, no. 1, pp. 29–36, 1982.
- [38] C. M. Bishop, *Pattern Recognition and Machine Learning*. New York, NY, USA: Springer, 2006.
- [39] T. Hastie, R. Tibshirani, and J. Friedman, *The Elements of Statistical Learning*, 2nd ed. New York, NY, USA: Springer, 2009.
- [40] A. P. Bradley, "The use of the area under the ROC curve in the evaluation of machine learning algorithms," *Pattern Recognition*, vol. 30, no. 7, pp. 1145–1159, 1997.
- [41] T. Fawcett, "An introduction to ROC analysis," *Pattern Recognition Letters*, vol. 27, no. 8, pp. 861–874, 2006.
- [42] J. Davis and M. Goadrich, "The relationship between Precision-Recall and ROC curves," in *Proc. 23rd Int. Conf. Machine Learning (ICML)*, 2006, pp. 233–240.
- [43] L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [44] T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in *Proc. 22nd ACM SIGKDD Int. Conf. Knowledge Discovery and Data Mining*, 2016, pp. 785–794.
- [45] N. Cristianini and J. Shawe-Taylor, *An Introduction to Support Vector Machines*. Cambridge, U.K.: Cambridge Univ. Press, 2000.
- [46] D. M. W. Powers, "Evaluation: From precision, recall and F-measure to ROC, informedness, markedness and correlation," *Journal of Machine Learning Technologies*, vol. 2, no. 1, pp. 37–63, 2011.
- [47] T. Saito and M. Rehmsmeier, "The precision-recall plot is more informative than the ROC plot when evaluating binary classifiers on imbalanced datasets," *PLoS ONE*, vol. 10, no. 3, pp. 1–21, 2015.
- [48] C. D. Manning, P. Raghavan, and H. Schütze, *Introduction to Information Retrieval*. Cambridge, U.K.: Cambridge Univ. Press, 2008.
- [49] J. Brownlee, *Machine Learning Mastery With Python*. Melbourne, Australia: Machine Learning Mastery, 2016.
- [50] S. Raschka and V. Mirjalili, *Python Machine Learning*, 3rd ed. Birmingham, U.K.: Packt Publishing, 2019.
- [51] T. Fawcett, "An introduction to ROC analysis," *Pattern Recognition Letters*, vol. 27, no. 8, pp. 861–874, 2006.
- [52] A. P. Bradley, "The use of the area under the ROC curve in the evaluation of machine learning algorithms," *Pattern Recognition*, vol. 30, no. 7, pp. 1145–1159, 1997.
- [53] L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [54] T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in *Proc. ACM SIGKDD*, 2016, pp. 785–794.
- [55] I. Guyon and A. Elisseeff, "An introduction to variable and feature selection," *Journal of Machine Learning Research*, vol. 3, pp. 1157–1182, 2003.
- [56] T. Saito and M. Rehmsmeier, "The precision-recall plot is more informative than the ROC plot when evaluating binary classifiers on imbalanced datasets," *PLoS ONE*, vol. 10, no. 3, pp. 1–21, 2015.
- [57] W. H. Wolberg, W. N. Street, and O. L. Mangasarian, "Machine learning techniques to diagnose breast cancer from fine-needle aspirates," *Cancer Letters*, vol. 77, no. 2, pp. 163–171, 1994.
- [58] M. Tavallaei et al., "A detailed analysis of the KDD CUP 99 dataset," in *Proc. IEEE CISDA*, 2009.
- [59] N. Moustafa and J. Slay, "UNSW-NB15: A comprehensive data set for network intrusion detection systems," in *Proc. Military Communications and Information Systems Conference*, 2015.
- [60] C. D. Manning, P. Raghavan, and H. Schütze, *Introduction to Information Retrieval*. Cambridge, U.K.: Cambridge Univ. Press, 2008.
- [61] A. Krizhevsky, "Learning multiple layers of features from tiny images," Univ. Toronto, Tech. Rep., 2009.
- [62] J. Deng et al., "ImageNet: A large-scale hierarchical image database," in *Proc. IEEE CVPR*, 2009.
- [63] T. Fawcett, "An introduction to ROC analysis," *Pattern Recognition Letters*, vol. 27, no. 8, pp. 861–874, 2006.
- [64] T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in *Proc. ACM SIGKDD*, 2016.
- [65] G. Litjens et al., "A survey on deep learning in medical image analysis," *Medical Image Analysis*, vol. 42, pp. 60–88, 2017.
- [66] D. M. W. Powers, "Evaluation: From precision, recall and F-measure to ROC, informedness, markedness and correlation," *Journal of Machine Learning Technologies*, vol. 2, no. 1, pp. 37–63, 2011.
- [67] L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [68] E. Topol, *Deep Medicine: How Artificial Intelligence Can Make Healthcare Human Again*. New York, NY, USA: Basic Books, 2019.
- [69] S. Grigorescu et al., "A survey of deep learning techniques for autonomous driving," *Journal of Field Robotics*, vol. 37, no. 3, pp. 362–386, 2020.
- [70] J. Lee, H. Davari, J. Singh, and V. Pandhare, "Industrial AI and predictive analytics for smart manufacturing systems," *Manufacturing Letters*, vol. 18, pp. 45–49, 2018.
- [71] H. He and E. A. Garcia, "Learning from imbalanced data," *IEEE Transactions on Knowledge and Data Engineering*, vol. 21, no. 9, pp. 1263–1284, 2009.
- [72] A. Dal Pozzolo, O. Caelen, R. A. Johnson, and G. Bontempi, "Calibrating probability with undersampling for unbalanced classification," in *Proc. IEEE Symposium on Computational Intelligence*, 2015.

- [73] L. Breiman, “Random forests,” *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [74] T. Chen and C. Guestrin, “XGBoost: A scalable tree boosting system,” in *Proc. ACM SIGKDD*, 2016.
- [75] S. Axelsson, “The base-rate fallacy and the difficulty of intrusion detection,” *ACM Transactions on Information and System Security*, vol. 3, no. 3, pp. 186–205, 2000.
- [76] A. Bifet and R. Gavaldà, “Learning from time-changing data with adaptive windowing,” in *Proc. SIAM Int. Conf. Data Mining*, 2007.
- [77] S. Barocas, M. Hardt, and A. Narayanan, *Fairness and Machine Learning: Limitations and Opportunities*. Cambridge, MA, USA: MIT Press, 2019.
- [78] A. Dal Pozzolo, O. Caelen, Y. Le Borgne, S. Watterschoot, and G. Bontempi, “Learned lessons in credit card fraud detection from a practitioner perspective,” *Expert Systems with Applications*, vol. 41, no. 10, pp. 4915–4928, 2014.
- [79] D. Dua and C. Graff, “UCI machine learning repository,” University of California, Irvine, 2019.