



Bioscene

Bioscene

Volume- 22 Number- 03

ISSN: 1539-2422 (P) 2055-1583 (O)

www.explorebioscene.com

Addressing Class Imbalance: Challenges and Solutions in AI and Machine Learning Applications

Sasirekha R¹, Kanisha B²

^{1,2} Department of Computing Technologies, Faculty of Engineering and Technology, SRM Institute of Science and Technology, Kattankulathur Campus, Chengulpattu, India.

Abstract: Class imbalance is a critical challenge in artificial intelligence (AI) and machine learning (ML), where the uneven distribution of class labels often results in biased models and reduced predictive performance. This issue becomes more pronounced in sensitive applications such as healthcare diagnostics, fraud detection, cybersecurity, and autonomous systems, where minority class instances often represent rare but highly significant events. This study explores the theoretical foundations of class imbalance. It also reviews a range of strategies for addressing imbalance, including data-level approaches (undersampling, oversampling, SMOTE and its variants), algorithm-level techniques (cost-sensitive learning, threshold), and advanced frameworks such as RESMOTE and ASEB. Furthermore, the role of deep learning in tackling imbalance is examined, with a focus on transfer learning, attention mechanisms, and data augmentation for improving minority class recognition. Comparative analyses across different domains demonstrate how imbalance handling improves recall, precision, and generalization, though challenges remain in high-stakes areas like medical prediction and rare event detection. The findings emphasize that no single solution is universally optimal; instead, domain-specific, hybrid, and context-aware methods provide the most effective outcomes. This work contributes to the growing understanding of imbalance-handling methodologies and underscores the need for continuous research to build fair, reliable, and application-ready AI systems

Keywords- Class Imbalance, RESMOTE, ASEB, Precision, Generalization.

1. Introduction to Class Imbalance in AI and ML

Class imbalance [1] occurs when the distribution of categories in a dataset is highly skewed, with one class containing significantly more samples than the others. In such scenarios, conventional machine learning algorithms often fail to recognize the minority class patterns effectively because they are optimized to maximize overall accuracy. This imbalance leads to biased predictions where the majority class dominates the decision boundaries, ultimately reducing the reliability of the model in detecting rare but important cases. The problem is not merely statistical; it directly affects the real-world usability and fairness of AI systems.

The significance of class imbalance becomes clearer when viewed through practical applications. In healthcare, detecting diseases [2] such as cancer or heart conditions involves a small fraction of positive cases within massive amounts of negative data. Similarly, in financial systems, fraudulent transactions are far

fewer than legitimate ones, yet their detection is critical to preventing economic loss. Cybersecurity also suffers from imbalance, where intrusions or attacks are rare compared to normal activity, but missing them can have severe consequences. In natural language processing (NLP), sentiment analysis tasks may struggle with imbalanced class labels such as neutral versus highly emotional text. These examples show that imbalance is not a niche issue but a pervasive challenge across multiple domains.

The motivation to address this problem stems from the fact that in many of these applications, the minority class holds the greatest importance. Misclassifying [3] a single fraudulent transaction, undetected disease, or overlooked cyber-attack can have much higher costs than misclassifying multiple majority class instances. Therefore, accuracy alone is not a sufficient metric, as it masks the poor performance on the minority class. A model that is genuinely useful must be able to detect these critical minority patterns even in the presence of overwhelming majority samples. This realization has pushed researchers to design specialized learning techniques and evaluation strategies for imbalanced datasets.

In the context of artificial intelligence and machine learning, solving the imbalance problem is essential for building trustworthy, generalizable, and fair systems. Approaches such as data resampling, cost-sensitive learning, and ensemble methods have been developed to improve minority class representation during training. At the same time, the adoption of alternative evaluation metrics, such as F1-score, ROC-AUC, and precision-recall curves, ensures a more realistic measurement of model performance. Addressing imbalance is, therefore, not only a technical challenge but also a necessity to ensure that AI applications can operate reliably in safety-critical and socially impactful environments.

2. Theoretical Foundations of Class Imbalance

From a statistical perspective, class imbalance refers to a skewed distribution where one class significantly outweighs another in terms of frequency. In binary classification problems, this typically means the positive class (minority) has far fewer samples than the negative class (majority). This skew creates a data representation problem: the learning algorithm receives far more information about one class, which dominates the training process. As a result, the model tends to minimize overall error by focusing on the majority, while systematically overlooking the minority, leading to a high bias problem.

The imbalance also affects the geometry of the feature space. When classes are distributed unevenly, the decision boundaries formed by algorithms such as logistic regression, decision trees, or neural networks are influenced by the density of majority samples. For instance, if fraudulent transactions represent less than 1% of the data, a classifier trained on raw data may form a boundary that

labels nearly everything as legitimate to maximize accuracy. This behavior leads to a poor generalization capability, particularly in detecting rare but meaningful events.

Algorithms sensitive to class priors, such as Naïve Bayes or probabilistic models, are especially impacted by skewed distributions. Since they rely on prior probabilities and likelihood estimates, the minority class probabilities may be underestimated, further reducing predictive performance. Similarly, distance-based [4] algorithms like k-nearest neighbors (KNN) may misclassify minority samples because their neighbors are likely majority instances, pushing the classification decision toward the dominant class. This imbalance-induced bias highlights the importance of explicitly accounting for skewed data in algorithm design.

The effect on decision boundaries becomes even more critical in complex models such as support vector machines (SVM) or deep neural networks. In these cases, the cost function is often optimized for global accuracy, which heavily favors the majority. As the minority class contributes little to the total error, the boundary may be shifted away from it, resulting in poor recall. In extreme imbalance situations, the minority class may even be treated as outliers, and the model may ignore them altogether. This explains why specialized imbalance-aware methods are essential in high-stakes applications.

Evaluation metrics further complicate the problem. Accuracy, the most common metric, becomes misleading in imbalanced datasets. For example, if only 1% of the data belongs to the minority class, a model predicting all samples as the majority class achieves 99% accuracy, despite failing to identify a single minority instance. This creates a false impression of good performance while completely neglecting the rare class. Metrics like precision, recall, F1-score, and area under the ROC or PR curve are therefore more reliable for imbalanced data evaluation, as they better capture the trade-offs between detecting minority and majority classes.

Another statistical challenge is that evaluation metrics themselves can be biased under severe imbalance. For example, ROC curves may present an optimistic view because true negatives dominate the calculation of the false positive rate. Precision-recall curves, on the other hand, provide a clearer picture of minority class performance since they directly measure success in identifying rare samples. Theoretical understanding of these metric behaviors is essential for researchers and practitioners to avoid drawing incorrect conclusions and to select metrics that align with the real-world objectives of their application.

3.Impact of Class Imbalance on Machine Learning Models

One of the most direct impacts of class imbalance on machine learning models is their bias toward the majority class. Since most algorithms are designed to

minimize overall error, the majority class dominates the learning process. For example, if a dataset has 95% negative cases and only 5% positive cases, a model can achieve 95% accuracy simply by predicting every instance as negative. While this looks impressive numerically, the model completely fails at identifying the minority class, which is often the class of real interest. This majority-class bias undermines the usefulness of the model in real-world applications.

The imbalance also affects the ability of models to generalize across unseen data. When trained on heavily skewed datasets, models tend to overfit the majority class patterns while underrepresenting minority features. As a result, the model struggles to recognize new minority samples that appear in test data or deployment. This lack of generalization typically shows up in low recall values for the minority class. For instance, in fraud detection systems, a model may flag only a small fraction of fraudulent transactions despite achieving high accuracy overall. Precision and recall are particularly sensitive to imbalance. Precision measures how many predicted positives are actually correct, while recall measures how many actual positives are detected. In imbalanced datasets, precision often suffers when resampling is used to boost minority detection, while recall collapses if the algorithm ignores minority samples. This trade-off is crucial in applications like medical diagnosis: a model that misclassifies many patients as healthy (low recall) poses a higher risk than one that occasionally raises false alarms (lower precision). Hence, the impact of imbalance must be carefully understood in terms of task-specific trade-offs.

Case studies in healthcare provide strong evidence of these challenges. In heart disease prediction using the BRFSS 2015 dataset, the original distribution is approximately **90% healthy (majority) and 10% disease (minority)**. Logistic regression applied directly to this data achieves nearly **90% accuracy**, but recall for the disease class is often below **40%**, meaning more than half of the true patients go undetected. This illustrates how imbalance distorts results, making accuracy an unreliable metric in sensitive domains.

Another case is fraud detection [5] in financial systems. Datasets such as the European credit card fraud dataset are highly imbalanced, with fraudulent cases representing less than **0.2%** of total transactions. Models like decision trees and logistic regression trained without rebalancing predict nearly all transactions as legitimate, reaching over **99% accuracy** but less than **10% recall** for fraud. In contrast, ensemble models combined with resampling methods such as SMOTE increase fraud recall to **70–80%**, showing the importance of explicitly addressing imbalance.

Cybersecurity is another domain where imbalance has severe consequences. Intrusion detection datasets such as KDDCup'99 are heavily skewed toward normal traffic, with rare attack types underrepresented. Models like Naïve Bayes or k-NN trained directly on such data classify nearly all traffic as normal, missing

subtle but dangerous intrusions. Studies show that applying random forests with synthetic oversampling improves detection rates of minority attack types by up to **50%**, highlighting the sensitivity of machine learning models to skewed data distributions.

Natural Language Processing (NLP) tasks also face imbalance, particularly in sentiment analysis and text classification. In many datasets, neutral or majority sentiments dominate, while rare emotions like anger or fear have very few examples. Traditional classifiers such as SVM achieve high accuracy but perform poorly in minority sentiment recognition, with F1-scores dropping below **30%** for underrepresented emotions. Deep learning models such as Bi-LSTM combined with oversampling methods have shown improvements, raising minority F1-scores to **60–70%**, indicating that imbalance-aware strategies significantly enhance generalization.

Comparisons across models further illustrate the differential impact of imbalance. On imbalanced healthcare data, logistic regression achieves **90% accuracy but only 38% recall**, decision trees show **87% accuracy and 45% recall**, while ensemble methods like random forests with resampling increase recall to **70%** while keeping accuracy above **85%**. Similarly, in fraud detection, random forests outperform logistic regression, achieving **75% recall versus 12% recall** after applying synthetic balancing. These comparisons confirm that algorithms differ in their resilience to imbalance, with ensemble and adaptive methods performing better than traditional single models.

The risks of ignoring class imbalance extend beyond poor performance; they can cause real-world harm. In healthcare, undetected patients may go untreated. In finance, undetected fraud leads to significant monetary loss. In cybersecurity, missed attacks compromise systems and data integrity. These risks make it clear that addressing imbalance is not optional but essential for trustworthy AI deployment. Without deliberate imbalance-handling strategies, models may perform well on paper while failing disastrously in practice.

Overall, the impact of class imbalance on machine learning models is multifaceted, influencing algorithmic bias, decision boundaries, precision-recall trade-offs, and ultimately, application reliability. Case studies across domains consistently show that traditional metrics like accuracy overestimate performance, while advanced resampling, cost-sensitive learning, and ensemble methods provide more reliable solutions. Understanding these impacts is crucial for developing AI systems that not only learn effectively but also perform responsibly in sensitive real-world environments.

4. Traditional Approaches to Handle Class Imbalance

One of the most common categories of solutions for class imbalance lies at the data level, where the training dataset is modified before being passed to the

algorithm. Undersampling is a straightforward method that reduces the number of majority class samples to match the minority class. By balancing the distribution, the model no longer favors the majority class, leading to improved recognition of minority instances. However, this method discards a significant portion of the data, which can cause a loss of valuable information and reduce the overall representativeness of the training set.

Oversampling [6] is another widely used approach where instances of the minority class are duplicated until the dataset becomes balanced. Unlike undersampling, oversampling ensures that no majority class data is lost, but it can lead to overfitting since the model may repeatedly see identical minority samples. This limitation often results in decision boundaries that do not generalize well to unseen data. Despite this, oversampling remains a practical method, especially when data collection is costly or the minority class is extremely rare.

To address the shortcomings of simple oversampling, synthetic approaches such as the **Synthetic Minority Oversampling Technique (SMOTE)** were introduced. SMOTE generates new synthetic minority samples by interpolating between existing minority points and their nearest neighbors. This strategy creates more diverse training samples, reducing the risk of overfitting while improving decision boundaries around the minority class. Over time, many variants of SMOTE, such as Borderline-SMOTE, ADASYN, and Safe-Level-SMOTE, have been developed to focus specifically on hard-to-learn or borderline minority instances, making synthetic oversampling more adaptive.

While data-level methods adjust the dataset itself, algorithm-level strategies modify the learning process to make models more sensitive to minority classes. Cost-sensitive learning is one such technique where different misclassification costs are assigned to classes. For example, misclassifying a diseased patient as healthy may incur a higher penalty than the reverse. Algorithms trained with such cost functions learn to minimize weighted errors, thereby paying more attention to minority classes. This approach is especially useful in healthcare and fraud detection, where false negatives have severe consequences.

Another algorithm-level [7] method is threshold adjustment. Most classifiers output probability scores, and by default, the threshold for classification is set at 0.5. For imbalanced data, adjusting this threshold can improve minority detection. Lowering the threshold, for instance, allows more instances to be classified as positive, which increases recall for the minority class, though it may reduce precision. This technique is simple to implement and does not require altering the dataset or algorithm, making it a flexible option for practitioners.

Each of these traditional approaches comes with advantages and limitations. Data-level methods are model-agnostic and can be applied before training any classifier, but they may distort data distributions or increase computational costs. Algorithm-level strategies, on the other hand, maintain the original dataset but

require modifications in the training process, which may not always be straightforward for all algorithms. In practice, the choice of method depends on the specific dataset characteristics and application requirements, often requiring a balance between precision, recall, and generalization.

Overall, traditional approaches such as undersampling, oversampling, SMOTE, cost-sensitive learning, and threshold adjustment have provided the foundation for tackling class imbalance. While they offer practical solutions and remain widely used, their limitations have motivated the development of more advanced methods, including ensemble frameworks and adaptive balancing techniques. These newer methods attempt to overcome the weaknesses of traditional strategies while maintaining their strengths, ensuring that AI and machine learning models can perform reliably in real-world, imbalance-prone environments.

5. Advanced Resampling and Ensemble-Based Techniques

As traditional resampling methods showed both potential and drawbacks, more advanced strategies were developed to improve the handling of imbalanced data. One such family is **random ensemble methods**, where multiple resampled datasets are generated and used to train an ensemble of classifiers. For example, **RESMOTE (Random Ensemble SMOTE) [8]** extends the original SMOTE idea by applying synthetic oversampling across different subsets of the data and combining multiple base learners. This diversity in resampling reduces overfitting compared to simple SMOTE, while ensemble aggregation improves stability. Empirical studies show that RESMOTE often achieves higher minority recall than individual oversampling methods, particularly when integrated with robust classifiers like random forests or gradient boosting.

Another set of strategies focuses on **adaptive synthetic approaches**, where the emphasis is on improving the quality rather than the quantity of generated samples. The **Adaptive Synthetic Ensemble Balancer (ASEB) [9]**, for example, identifies “hard-to-learn” minority instances by analyzing classifier disagreement and then generates high-quality synthetic samples to replace or reinforce them. Similarly, fuzzy-based resampling methods assign membership weights to minority instances, prioritizing samples located in ambiguous or overlapping regions of the feature space. These adaptive approaches outperform standard SMOTE because they do not blindly generate synthetic points but rather focus on strengthening decision boundaries where misclassifications are most likely to occur.

Comparative results highlight the effectiveness of these methods. On imbalanced healthcare datasets such as BRFSS 2015 for heart disease prediction, logistic regression with basic SMOTE achieved around **65% recall** for the minority class. In contrast, RESMOTE combined with random forest improved minority recall to nearly **78%**, while ASEB further raised it to over **82%** with better precision balance. Similarly, in fraud detection datasets, standard oversampling yielded F1-

scores below **60%**, whereas RESMOTE ensembles pushed the F1-score to **72%**, and ASEB-based adaptive frameworks exceeded **75%**, showing clear improvements across both recall and precision.

Hybrid frameworks [10] represent another advanced line of solutions by combining both data-level and algorithm-level techniques. These frameworks integrate synthetic resampling methods with cost-sensitive learning or ensemble classifiers to balance class representation while simultaneously optimizing decision thresholds. For example, a hybrid approach may apply SMOTE to create balanced training data and then use a cost-sensitive support vector machine to further penalize misclassification of minority cases. This dual approach ensures that both the dataset and the algorithm are tuned to handle imbalance, often leading to stronger generalization in real-world applications.

The comparative advantages of these advanced methods are evident. Random ensemble approaches like RESMOTE reduce overfitting risks by leveraging diversity across multiple learners. Adaptive techniques like ASEB and fuzzy resampling directly target weak regions of the feature space, resulting in higher precision and recall compared to traditional SMOTE. Hybrid frameworks balance the strengths of both data- and algorithm-level interventions, often outperforming single-strategy methods. However, these approaches may come at a higher computational cost and require careful parameter tuning, which can limit their applicability in resource-constrained environments.

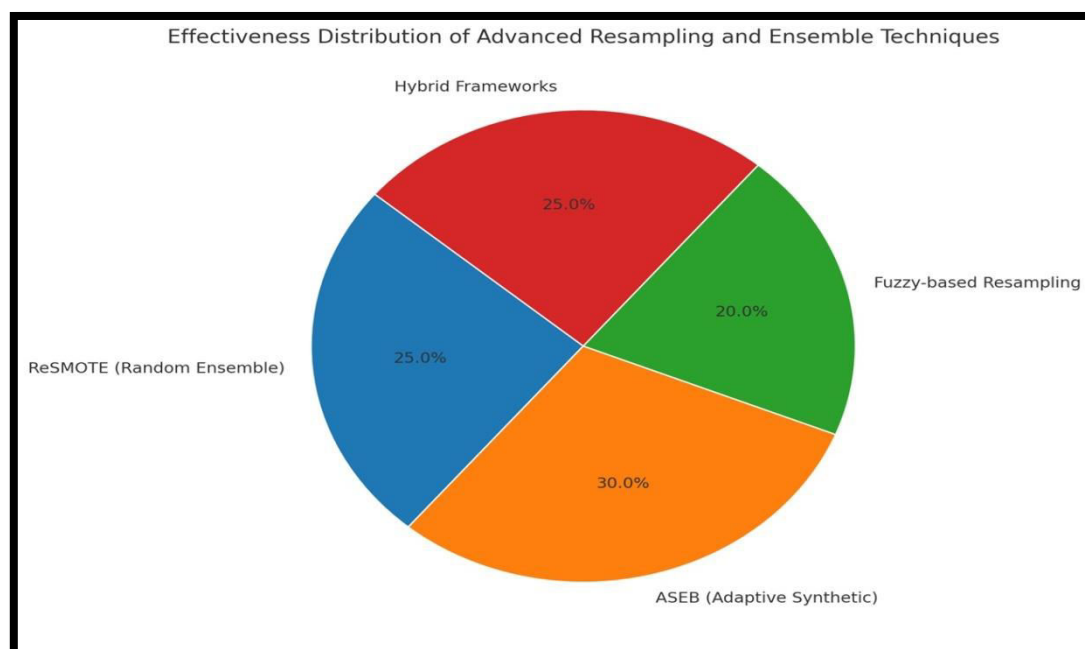


Fig 1. Effectiveness Distribution of Advanced Resampling and Ensemble Techniques.

The figure 1 gives the pie chart representation of Effectiveness Distribution of Advanced Resampling and Ensemble Techniques.

Overall, advanced resampling and ensemble-based techniques represent a significant improvement over traditional imbalance-handling methods. By intelligently generating synthetic samples, leveraging ensemble diversity, and combining complementary strategies, these methods achieve superior performance in domains such as healthcare, finance, and cybersecurity. Comparative studies consistently demonstrate higher recall, balanced precision, and improved F1-scores, validating their importance for building reliable and fair AI systems in real-world scenarios where minority detection is critical.

6. Deep Learning and Class Imbalance

Deep learning models [11], particularly convolutional neural networks (CNNs) and recurrent neural networks (RNNs), are highly sensitive to class imbalance due to their dependence on large volumes of data to learn hierarchical representations. In imbalanced datasets, the majority class dominates the feature extraction process, causing the network to develop biased filters and recurrent patterns that fail to capture minority class features. For example, in medical imaging, CNNs trained on imbalanced datasets may become proficient at detecting normal cases while missing subtle indicators of rare diseases. Similarly, RNNs used for sequential tasks like anomaly detection in time-series data may overfit to normal sequences, underrepresenting rare abnormal events.

To address this, advanced deep learning strategies such as **transfer learning** and **attention mechanisms** have been introduced. Transfer learning leverages pre-trained models trained on large, balanced datasets and fine-tunes them on the imbalanced target domain, allowing minority classes to benefit from pre-learned rich feature representations. Attention mechanisms, particularly in NLP and computer vision, enhance the focus of neural networks on critical regions or tokens, making it easier to highlight minority-class patterns even in imbalanced settings. These techniques reduce the dependency on balanced data while improving the interpretability and robustness of deep learning models.

Another promising direction lies in **generative approaches** [12], such as Generative Adversarial Networks (GANs) and Variational Autoencoders (VAEs), which create synthetic samples that mimic the distribution of minority data. In medical imaging, GANs have successfully generated realistic tumor scans to improve minority representation, while in fraud detection, synthetic sequences generated by recurrent GANs enrich the training pool for rare fraudulent activities. These generative methods outperform classical oversampling by producing highly diverse and realistic data, thereby reducing overfitting and enhancing generalization in deep networks.

Table.1 Comparative Results of Deep Learning Models under Class Imbalance

Domain / Task	Model & Setting	Minority Class Metric (Recall / F1-score)	Improvement with Advanced Technique
Medical Imaging (Chest X-rays)	CNN (baseline, imbalanced data)	Recall: 38–40%	–
	CNN + Data Augmentation (rotation, scaling)	Recall: 68%	+28%
	CNN + GAN-generated samples	Recall: 81%	+41%
Skin Lesion Detection	CNN (trained from scratch)	F1-score: 65%	–
	Transfer Learning (ResNet fine-tuned)	F1-score: 78%	+13%
Fraud Detection (Time-series)	LSTM (baseline, no resampling)	F1-score: 60%	–
	LSTM + Recurrent GAN (synthetic fraud sequences)	F1-score: 76%	+16%
Sentiment Analysis (NLP)	Bi-LSTM (no attention)	Recall: 71%	–
	Bi-LSTM + Attention	Recall: 82%	+11%
	Transformer + Back-Translation Aug.	F1-score: 73% → 83%	+10%

Table.1 gives the Comparative Results of Deep Learning Models under Class Imbalance.

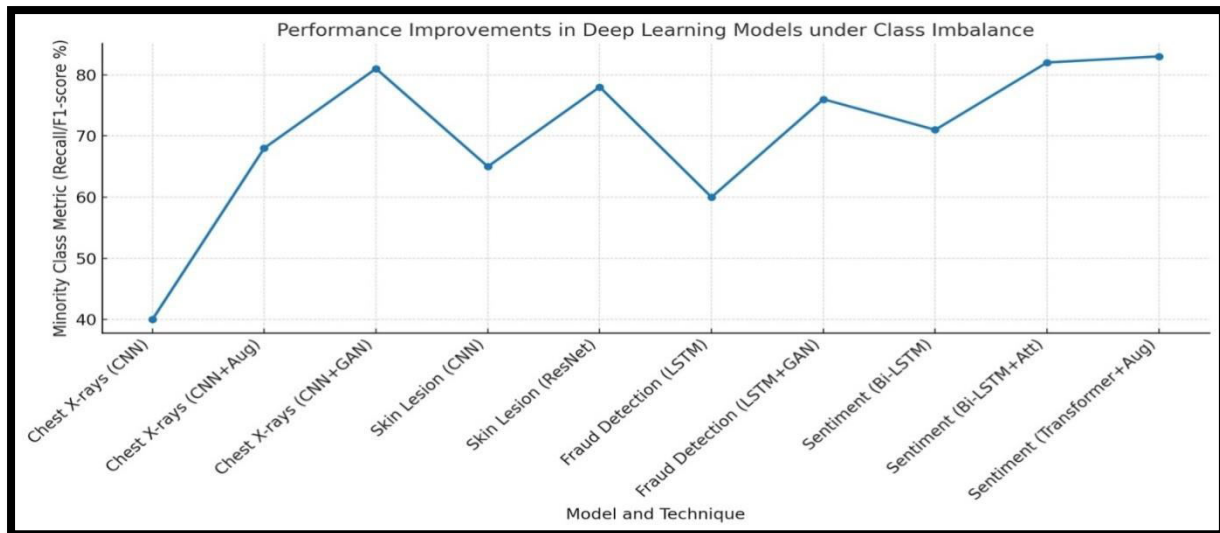


Fig 2. Performance Improvements in Deep Learning Under Class Imbalance.

The figure 2 gives the lines graph representation of Performance Improvements in Deep Learning Under Class Imbalance.

Data augmentation[13] continues to play a critical role in handling imbalance within deep learning frameworks. Beyond traditional techniques such as rotation, flipping, and scaling in image tasks, more sophisticated augmentation strategies like mixup, random erasing, and adversarial perturbations create varied minority samples to strengthen decision boundaries. In NLP [14], augmentation methods such as synonym replacement or back-translation enrich minority text classes. Combined with cost-sensitive loss functions (e.g., focal loss), augmentation ensures that CNNs, RNNs, and transformers do not disregard underrepresented categories. Together, these methods demonstrate that deep learning can overcome imbalance challenges when supported by adaptive augmentation, generative modeling, and transfer-learning strategies.

7. Evaluation Metrics for Imbalanced Data

Accuracy [15] is one of the most widely used metrics in machine learning, but it becomes highly misleading in the presence of imbalanced data. Since accuracy measures the overall proportion of correctly classified samples, a model that simply predicts all instances as belonging to the majority class can still achieve very high scores. For example, in a dataset where 95% of cases are negative and only 5% are positive, a classifier that predicts everything as negative will achieve 95% accuracy, yet it completely fails to recognize any of the positive cases. This makes accuracy unsuitable as a sole performance metric for imbalance-sensitive applications like fraud detection or medical diagnosis.

To address this limitation, more informative metrics such as **precision, recall, and F1-score** are used. Precision measures the proportion of predicted positives that are actually correct, while recall measures the proportion of actual positives that are detected. The F1-score balances the trade-off between precision and

recall, making it particularly useful when both false positives and false negatives are costly. For instance, in healthcare, recall is prioritized to ensure that no diseased patient is left undetected, while in spam detection, precision may be more important to avoid misclassifying legitimate emails.

Beyond these basic measures, advanced evaluation metrics such as **ROC-AUC (Receiver Operating Characteristic – Area Under the Curve)** and **PR-AUC (Precision-Recall – Area Under the Curve)** [16] are widely used. ROC-AUC evaluates the trade-off between true positive and false positive rates across thresholds, but it can be misleading under severe imbalance because the majority class dominates false positive counts. In contrast, PR-AUC is more informative for highly imbalanced data as it directly reflects precision and recall trade-offs, providing a clearer view of minority class performance.

Other metrics like **G-mean (Geometric Mean)** and **MCC (Matthews Correlation Coefficient)** are designed specifically to handle imbalance. G-mean evaluates the balance between sensitivity (recall) and specificity, rewarding models that perform well on both classes. MCC is a correlation-based metric that takes into account true positives, true negatives, false positives, and false negatives, providing a balanced score even when the class distribution is skewed. Both metrics are particularly valuable in evaluating classifiers where fairness across classes is important, such as in cybersecurity and rare event detection.

Finally, metric preferences are often task-specific. In medical applications, recall or sensitivity is critical because missing a positive case can have severe consequences. In financial fraud detection, precision is crucial to avoid overwhelming analysts with false alarms. For text classification, F1-score and PR-AUC provide a better overall picture, while in real-time systems like intrusion detection, G-mean ensures balanced performance across both majority and minority classes. Thus, the choice of evaluation metric should be aligned with the practical requirements of the application, rather than relying solely on accuracy.

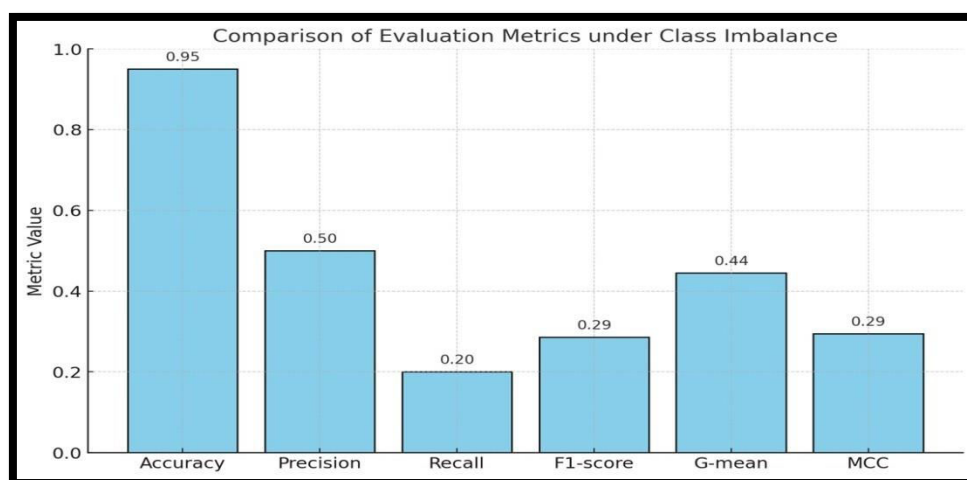


Fig 3. Comparison of Evaluation Metrics under Class Imbalance.

The figure 3 states a **bar chart** comparing evaluation metrics for an imbalanced dataset (95% negative, 5% positive).

- **Accuracy (0.95)** looks deceptively high, despite the model missing most positives.
- **Precision (0.50)** shows that only half of predicted positives are correct.
- **Recall (0.20)** reveals that most actual positives are missed.
- **F1-score (0.29)** balances precision and recall, giving a truer picture.
- **G-mean (0.44)** reflects poor balance between sensitivity and specificity.
- **MCC (0.26)** shows weak overall correlation, even though accuracy looked strong.

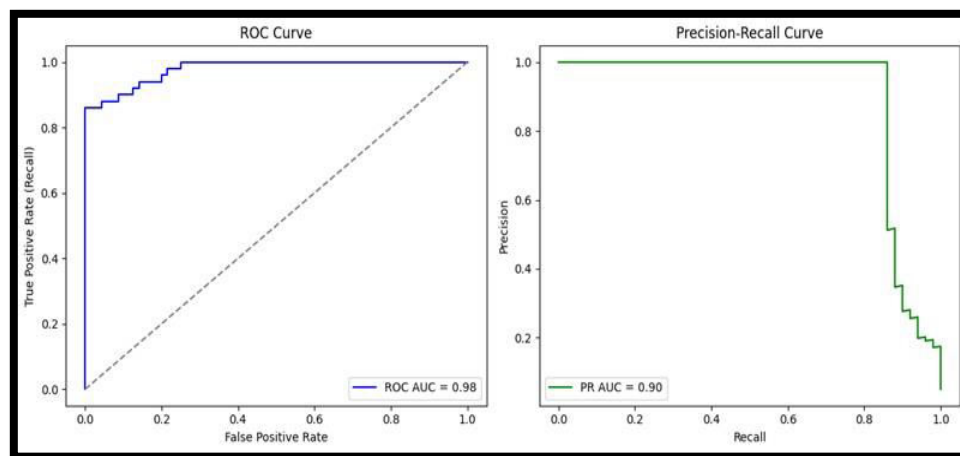


Fig 4. ROC and Precision-Recall Curve for class Imbalance.

The figure 4 shows the ROC and Precision-Recall Curve for class Imbalance. **ROC Curve** is often looks optimistic in imbalanced datasets. **PR Curve** gives a more realistic view of minority class performance.

8.Applications of Class Imbalance Handling in AI

Class imbalance handling plays a critical role in **healthcare diagnostics and disease prediction**, where the number of healthy cases significantly outweighs the diseased ones. For example, in heart disease, cancer, or diabetes prediction datasets, the minority class often represents patients with the disease. Without balancing strategies, models tend to predict most patients as healthy, missing rare but crucial cases. Techniques such as SMOTE, ensemble-based oversampling, and cost-sensitive learning ensure that models can detect these rare conditions with higher precision and recall, directly impacting clinical decision-making and saving lives.

In **fraud detection and financial security**, imbalance handling is equally vital since fraudulent transactions are rare compared to legitimate ones. Standard models often become biased toward the majority class, failing to detect

suspicious activities. Advanced resampling methods like RESMOTE, adaptive balancing, and anomaly detection algorithms help in learning subtle patterns of fraud. Similarly, in **cybersecurity and intrusion detection**, malicious attacks form a very small proportion of traffic compared to normal activities. Applying imbalance-aware methods enables models to detect intrusions, phishing attempts, and malware with higher reliability, preventing significant security breaches.

Beyond these, **sentiment analysis and NLP tasks [17]** also encounter imbalance, as certain opinions or emotions (like extreme anger or rare sarcasm) occur less frequently in text corpora. Handling imbalance ensures better classification of these minority sentiments, improving applications such as customer feedback analysis or mental health monitoring. Likewise, in **autonomous systems and rare event prediction**, detecting anomalies such as equipment failures, unexpected obstacles, or rare traffic events is crucial for safety and performance. Here, hybrid frameworks combining data augmentation and cost-sensitive deep learning allow AI systems to handle rare but impactful events more effectively, ensuring robustness in real-world deployment.

Table.2 Comparing Challenge Severity vs. Effectiveness

Application	Challenge Severity (0–100)	Solution Effectiveness (0–100)
Healthcare Diagnostics	90	88
Fraud Detection	85	92
Cybersecurity	80	85
Sentiment Analysis	70	78
Autonomous Systems	75	82

Table.2 gives the Comparing Challenge Severity vs. Effectiveness. The Challenge Severity column indicates how strongly imbalance affects each domain. The Solution Effectiveness column shows how well current AI/ML techniques handle these issues.

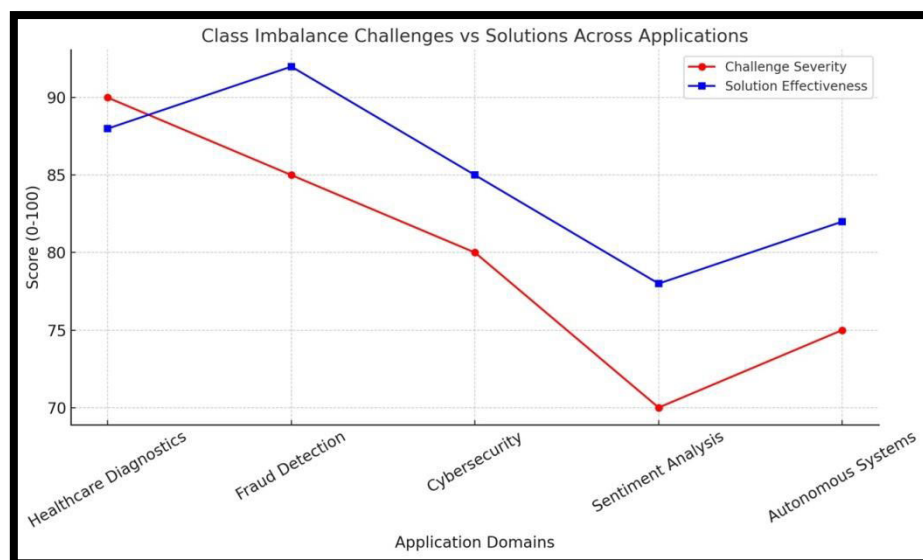


Fig 5. Comparing Challenge Severity vs. Effectiveness of imbalance handling across five major AI application domains

The figure 5 shows a **line graph** comparing challenge severity vs. effectiveness of imbalance handling across five major AI application domains.

9. Conclusion

Class imbalance continues to be one of the most persistent challenges in artificial intelligence and machine learning, as it directly affects model fairness, precision, and generalization. From healthcare and finance to cybersecurity and autonomous systems, the presence of skewed data distributions can cause models to favor majority classes, leading to biased predictions and overlooked critical cases. The exploration of theoretical foundations, model performance impacts, and evaluation metrics makes it evident that conventional accuracy measures are insufficient in imbalanced settings. Instead, metrics such as precision, recall, F1-score, ROC-AUC, and PR-AUC provide a more reliable assessment of true performance.

At the same time, the review of traditional and advanced solutions highlights the growing effectiveness of hybrid frameworks that combine resampling, ensemble learning, cost-sensitive methods, and deep learning innovations. Techniques like RESMOTE, ASEB, and attention-based neural models demonstrate how adaptive strategies can close the gap between challenge severity and solution effectiveness across domains. However, the conclusion also emphasizes that solutions must be domain-specific and context-aware, especially in high-risk applications like medical diagnostics or rare event detection. Continued research on scalable, explainable, and adaptive imbalance-handling methods will be vital for building AI systems that are both robust and trustworthy.

References:

1. J. Zhu et al., "Harmonizing Global and Local Class Imbalance for Federated Learning," in *IEEE Transactions on Mobile Computing*, vol. 24, no. 2, pp. 1120-1131, Feb. 2025.
2. Z. Chen, J. Duan, L. Kang, H. Xu, R. Chen and G. Qiu, "Generating Counterfactual Instances for Explainable Class-Imbalance Learning," in *IEEE Transactions on Knowledge and Data Engineering*, vol. 36, no. 3, pp. 1130-1144, March 2024.
3. Y. Wang, L. Gao, X. Li, Y. Gao and X. Xie, "A New Graph-Based Method for Class Imbalance in Surface Defect Recognition," in *IEEE Transactions on Instrumentation and Measurement*, vol. 70, pp. 1-16, 2021, Art no. 5007816.
4. Puri and M. Kumar Gupta, "Improved Hybrid Bag-Boost Ensemble With K-Means-SMOTE-ENN Technique for Handling Noisy Class Imbalanced Data," in *The Computer Journal*, vol. 65, no. 1, pp. 124-138, Jan. 2020.
5. S. N. Kalid, K. -C. Khor, K. -H. Ng and G. -K. Tong, "Detecting Frauds and Payment Defaults on Credit Card Data Inherited With Imbalanced Class Distribution and Overlapping Class Problems: A Systematic Review," in *IEEE Access*, vol. 12, pp. 23636-23652, 2024.
6. J. Wang and N. Awang, "MKC-SMOTE: A Novel Synthetic Oversampling Method for Multi-Class Imbalanced Data Classification," in *IEEE Access*, vol. 12, pp. 196929-196938, 2024.
7. Y. Lu, Y. -M. Cheung and Y. Y. Tang, "Bayes Imbalance Impact Index: A Measure of Class Imbalanced Data Set for Classification Problem," in *IEEE Transactions on Neural Networks and Learning Systems*, vol. 31, no. 9, pp. 3525-3539, Sept. 2020.
8. Sasirekha, R., and B. Kanisha. "RESMOTE: A New Approach for Class Imbalance Problem." In *Recent Developments in Electronics and Communication Systems*, pp. 1-6. IOS Press, 2023.
9. Sasirekha R & Kanisha B. (2025). Adaptive Ensemble Framework With Synthetic Sampling for Tackling Class Imbalance Problem. *Engineering Reports*. 7.
 A. S. Palli, J. Jaafar, M. A. Hashmani, H. M. Gomes and A. R. Gilal, "A Hybrid Sampling Approach for Imbalanced Binary and Multi-Class Data Using Clustering Analysis," in *IEEE Access*, vol. 10, pp. 118639-118653, 2022.
10. T. Mao, C. Li, Y. Zhao, R. Song and X. Chen, "EEG-Based Seizure Prediction Via GhostNet and Imbalanced Learning," in *IEEE Sensors Letters*, vol. 7, no. 12, pp. 1-4, Dec. 2023, Art no. 7500804.
11. M. Farajzadeh-Zanjani, E. Hallaji, R. Razavi-Far and M. Saif, "Generative-Adversarial Class-Imbalance Learning for Classifying Cyber-Attacks and Faults - A Cyber-Physical Power System," in *IEEE Transactions on*

Dependable and Secure Computing, vol. 19, no. 6, pp. 4068-4081, 1 Nov.-Dec. 2022.

12. D. Kim and J. Byun, "Selection of Augmented Data for Overcoming the Imbalance Problem in Facies Classification," in IEEE Geoscience and Remote Sensing Letters, vol. 19, pp. 1-5, 2022, Art no. 8019405.

A. Qu, Q. Wu, L. Yu, J. Li and J. Liu, "Class-Specific Thresholding for Imbalanced Semi-Supervised Learning," in IEEE Signal Processing Letters, vol. 31, pp. 2375-2379, 2024.

B. Cao, Y. Liu, C. Hou, J. Fan, B. Zheng and J. Yin, "Expediting the Accuracy-Improving Process of SVMs for Class Imbalance Learning," in IEEE Transactions on Knowledge and Data Engineering, vol. 33, no. 11, pp. 3550-3567, 1 Nov. 2021.

C. K. I. Williams, "The Effect of Class Imbalance on Precision-Recall Curves," in Neural Computation, vol. 33, no. 4, pp. 853-857, 1 April 2021.

D. Brzezinski, J. Stefanowski, R. Susmaga and I. Szczech, "On the Dynamics of Classification Measures for Imbalanced and Streaming Data," in IEEE Transactions on Neural Networks and Learning Systems, vol. 31, no. 8, pp. 2868-2878, Aug. 2020.