

A Difficulty-aware Approach to Fair Classification on Imbalanced Datasets *

Research Article

Niloufar Kashefi¹, Javad Hamidzadeh² , Mona Moradi³

 [10.22067/cke.2025.91011.1137](https://doi.org/10.22067/cke.2025.91011.1137)

Abstract Class imbalance in real-world datasets often biases standard classifiers toward the majority class, degrading performance on the minority class. While existing methods like sample re-weighting can mitigate this, they may increase overall misclassification errors or fail to consider the difficulty of training instances. To address these shortcomings, we introduce a difficulty-aware classification framework based on multi-objective evolutionary optimization. Our approach uses a specialized fitness function to simultaneously optimize for minority-class recall and overall accuracy, guiding the selection of the most informative training samples. We quantify sample difficulty using a fuzzy approach, which then modulate class-specific weights to refine the classifier's decision boundary. Furthermore, we incorporate chaotic dynamic maps into the evolutionary operators to accelerate convergence and maintain population diversity. Evaluated on various UCI benchmark datasets with 10-fold cross-validation, our method improves minority-class performance on imbalanced data without compromising accuracy on balanced data. Comparative analysis using AUC, G-mean, and F-measure confirms our approach achieves a superior trade-off between minority-class detection and overall accuracy compared to state-of-the-art methods.

Key Words Instance reduction; Fuzzy weighted average distance-based decision surface; Chaotic imperialist competitive algorithm; Reduction rate.

1- INTRODUCTION

Class imbalance is a pervasive challenge in machine learning, significantly impairing model performance across diverse domains such as medical diagnosis [1], streaming data mining [2], fraud detection [3], natural language processing [4], and image analysis [5]. In these

applications, the training data contains a skewed class distribution, which biases classifiers toward the majority class. As a result, instances from the minority class—which often carry the most critical information—are frequently misclassified. Existing solutions to mitigate class imbalance can be broadly categorized into three paradigms:

1. **Data-level approaches:** These techniques modify the training data distribution to create a more balanced class representation. The primary methods are oversampling and undersampling. Oversampling methods increase the number of minority class instances, ranging from simple duplication to synthetic generation of new samples with techniques like SMOTE. Conversely, undersampling methods reduce the size of the majority class by discarding redundant or noisy samples. However, these risks losing valuable information. The core challenge in these approaches lies in sample selection. Effective methods must distinguish informative instances from noisy or redundant ones to avoid overfitting while preserving data diversity. To achieve this, neighborhood-based algorithms like k -Nearest Neighbors (k -NN) can identify redundant samples based on feature space proximity. In contrast, clustering-based methods (e.g., k -means) can sample representative centroids. More advanced techniques, such as the adaptive clustering proposed in [6], reduce computational overhead by dynamically tuning the cluster count and using alternative distance metrics for sample selection.
2. **Algorithm-level approaches:** These approaches modify the learning algorithm without changing the distribution of the training data. This is often achieved through cost-sensitive learning [7, 8], where higher misclassification costs are assigned to minority class

* Manuscript received November 29, 2024, Revised July 25, 2025, Accepted August 17, 2025.

¹ M.Sc. Faculty of computer engineering and information technology, Sadjad University, Mashhad, Iran.

² Corresponding author. Associate Professor, Faculty of computer engineering and information technology, Sadjad University, Mashhad, Iran. Email: J_Hamidzadeh@sadjad.ac.ir

³ Assistant Professor, Department of computer engineering, University of Torbat Heydarieh, Torbat Heydarieh, Iran. Email: M.Moradi@torbath.ac.ir

instances, directly embedding class-specific penalties into the objective function. While effective at improving minority recall, this can increase model complexity. Other algorithmic strategies include one-class classification and modifications to specific algorithms like Support Vector Machines (SVMs) [12–14], though these can require specialized tuning.

3. **Hybrid Approaches:** This category combines data-level and algorithm-level techniques to leverage the strengths of both. The goal is to create a more robust solution where a modified data distribution is fed into a specially adapted learning algorithm. These methods are often highly effective because they address the problem from two angles simultaneously. Popular examples include ensemble methods combined with sampling, such as SMOTEBoost [15], which integrates the SMOTE oversampling technique into the AdaBoost framework, and RUSBoost [11], which combines random undersampling with Boosting. By synergistically applying both strategies, hybrid methods can often achieve superior performance, particularly in cases of severe imbalance where a single approach is insufficient.

Over the past decade, research on class imbalance has advanced significantly. The focus has broadened from binary classification to multi-class long-tail (skewed) distributions, spurring the development of hybrid augmentation methods like GAN-based synthetic sample generation [15]. In parallel, uncertainty-aware frameworks, such as Bayesian inference [16–18] and fuzzy logic [19], have been integrated to refine both data augmentation and hyperparameter optimization [20, 21]. These advancements culminate in state-of-the-art strategies like evolutionary bagging with data augmentation [22], which aim to enhance model fairness and robustness.

Despite this progress, a pervasive limitation persists in many rebalancing techniques: they either indiscriminately remove majority instances or apply static, higher weights to minority classes. Such strategies often overlook the crucial concept of *instance difficulty*, which affects instances in all classes, including the majority. Instance difficulty refers to the inherent difficulty of a sample that can negatively affect a model's learning process. This difficulty can arise from several factors:

Noise: Irrelevant data within the sample.

Ambiguity: The sample's features make it hard to clearly assign to one class.

Proximity to a decision boundary: The sample is located very close to the line or surface that separates different classes, making it a challenging case for the classifier.

Neglecting these difficult instances can compromise classifier performance and discard information critical for robust decision-making.

To bridge this gap, we propose a difficulty-aware decision surface derived via a chaotic-evolutionary optimization process. This decision surface integrates the designed instance difficulty metric to simultaneously perform instance selection and

adaptive weighting, enabling nuanced rebalancing sensitive to sample difficulty. The main contributions of the paper are:

- A hybrid rebalancing strategy that combines subset selection and fuzzy-based weight modulation, providing a more granular and context-sensitive mechanism than traditional methods.
- Explicit consideration of sample-level difficulty—including ambiguous, noisy, or boundary samples—within the decision-making framework.

The remainder of this paper is organized as follows: Section 2 reviews relevant literature; Section 3 outlines theoretical backgrounds; Section 4 details the proposed method; Section 5 presents the experimental evaluation; and Section 6 concludes with insights and potential future directions.

2- RELATED WORK

This section reviews key literature related to the proposed methodology, structured into two subsections: Imbalance Ratio and Overlapping as problem, and Kernel-based Density Estimation Strategy as solution.

A. Problems: Imbalance Ratio and Overlapping

Imbalanced datasets are characterized by a significantly higher count of majority-class instances relative to minority-class instances. The Imbalance Ratio (IR), defined in Equation (1), is widely used to quantify this skew.

$$IR = \frac{\text{size of the majority class}}{\text{size of the minority class}} \quad (1)$$

In multi-label settings, IR must be computed per label since individual labels may exhibit distinct positive-negative class distributions; moreover, instances may simultaneously belong to both majority and minority labels, complicating imbalance assessment [23].

Table 1 summarizes different types of problems that may arise when encountering class imbalance.

TABLE 1
Different Types of Problems in Class Imbalance

Problem	Description	Consequence	Ref.
Class Overlap	Regions in feature space where instances of multiple classes coexist.	Degraded classifier performance, even in balanced datasets; misclassification at class boundaries.	[24-27]
Label Noise	Incorrectly labeled instances in the dataset.	Complicates training; can lead to poor generalization.	[26, 5]
Aggressive Reduction	Indiscriminate removal of majority instances by rebalancing methods.	Undermining of model fairness and stability; creation of new imbalances.	Proposed method

B. A good solution: Kernel-based Density Estimation Strategy.

Comprehending the local sample density is pivotal in addressing both class imbalance and overlap. By modeling separate probability density functions (PDFs) for majority and minority classes, the density ratio indicates how far a sample lies from the overlapping zone: values near unity suggest equal class density, whereas higher ratios denote proximity to minority high-density regions. Kernel Density Estimation (KDE) has been adopted to estimate these PDFs and guide density-aware sampling. KDE-based sampling can outperform conventional oversampling methods by generating synthetic instances in dense minority-class regions, reducing overfitting and improving metrics such as F-score and G-mean, even under severe IR conditions. While density-based under-sampling methods offer effective mitigation of class imbalance, they frequently falter in scenarios involving nonlinearly separable class distributions. To address this limitation, one can employ nonlinear mappings—such as kernel transformations or feature-space embeddings—that project the original input data into a high-dimensional feature space, thereby converting inherently nonlinear decision boundaries into linearly separable representations suitable for density estimation and sampling [25, 28]. Several recent KDE-related methodologies are introduced in [29-38]. To facilitate a comparative understanding of existing KDE-based approaches and their limitations, Table 2 provides a concise summary of representative methods recently proposed in the literature. This comparison helps to clarify the conceptual differences and performance gaps that our proposed difficulty-aware framework aims to address.

TABLE 2
Overview of Recent KDE-Based Methods for Imbalanced Data Classification

Methodology	Key Contribution	Limitation
Adaptive KDE [29]	Adjusts bandwidths based on local sparsity, preserving informative instances.	Lacks sensitivity to sample hardness and local ambiguity.
Localized Density Ratio Sampling (LDRS) [30]	Identifies overlap-prone regions and applies selective resampling.	Uniformly treats boundary instances, disregarding instance-level difficulty.
Density-Aware SMOTE (DA-SMOTE) [31]	Restricts synthetic sample generation to high-density areas.	Lacks adaptive mechanisms for instance difficulty or class overlap beyond density.
KDE with Graph Embedding [32]	Uses graph-based embeddings for nonlinear-aware sampling.	Does not incorporate instance-level selection or adaptive weighting based on difficulty.
KMM-HR [33]	Aligns inter-class distributions and reweights boundary/noisy instances.	Its nature limits fine-grained control over local structures.

In addition, several complementary strategies—including rebalancing [39], ensemble learning [40], transfer learning [41], and architecture decoupling [42]—have been introduced to tackle the previously outlined challenges from different methodological perspectives.

Another challenge arises where the number of samples in different categories follows a long-tailed distribution. Recent research can be broadly categorized into the following four groups, as shown in Table 3:

TABLE 3
Representative Studies Addressing Long-Tailed Distributions

Category	Strategy	Advantages	Disadvantages
Resampling & Reweighting [43-45]	Adjusts class distributions or assigns weights to instances.	Foundational; can highlight challenging samples.	Risks overfitting minority classes and underfitting majority classes.
Transfer Learning [46]	Leverages knowledge from data-rich (head) classes.	Improves performance on data-scarce (tail) classes.	Challenging to construct effective transfer modules and manage training dynamics.
Two-stage Training [47]	Decouples the feature extractor from the classifier.	Learns robust feature representations.	Can interfere with feature learning.
Grouping & Ensemble Methods [48-50]	Divides data into subsets for independent model training and aggregates outputs.	Preserves head-class performance; reduces prediction variance and bias.	Can be complex to implement and coordinate multiple models.

3- BACKGROUND: IMPERIALIST COMPETITIVE ALGORITHM

Imperialist Competitive Algorithm (ICA) [51] is a population-based metaheuristic within the evolutionary optimization paradigm, where candidate solutions are modeled as *countries*.

The algorithm is structured into the following stages:

(A) Initialization and role assignment: Initially, a population of countries is generated randomly in the solution space. Each country's quality (or power) is evaluated using a cost (fitness) function. Based on these cost values, countries are divided into two groups: imperialists (the stronger countries with lower cost values) and colonies (the weaker countries). The role assignment of colonies to imperialists are based on their relative powers: the better countries become imperialists and own a number of colonies proportional to their strength. The cost value (fitness) of each country is computed by (2).

$$C_n = f_{\text{cost}}^{(\text{imp},n)} - \max_i \left(f_{\text{cost}}^{(\text{imp},i)} \right) \quad (2)$$

where $f_{cost}^{(imp,n)}$ is the cost of n^{th} imperialist and C_n is its normalized cost. The normalized power of each imperialist is evaluated by (3).

$$P_n = \frac{C_n}{\sum_{i=1}^{N_{imp}} C_i} \quad (3)$$

where P_n is the normalized power value of n^{th} imperialist. The imperialists are the more powerful countries (have higher values for P_i variables). To evaluate the number of colonies of an empire, (4) must be used.

$$NC_n = \text{Round} \{P_n \cdot N_{col}\} \quad (4)$$

where NC_n is the number of initial colonies of n^{th} imperialist which is chosen from empiricists randomly. Besides, $\text{Round} \{.\}$ assigns an integer value to NC_n .

- (B)** Assimilation mechanism and position update: In ICA, a colony moves toward the imperialist by a random value that is between 0 and $\beta \times d$. The new location of colonies is evaluated by (5).

$$\{z\}_{new} = \{z\}_{old} + U(0, \beta \times d) \times \{V_1\} \quad (5)$$

where β is an adjustable parameter and d is the distance between colony and imperialist. $\{V_1\}$ is a vector which its start point is $\{z\}_{old}$ and its direction is toward $\{z\}_{new}$. The length of this vector is set to unity [42]. One way to increase the searching around the imperialist is by adding a random amount of deviation θ to the direction of movement. The moving model is shown in Figure 1.

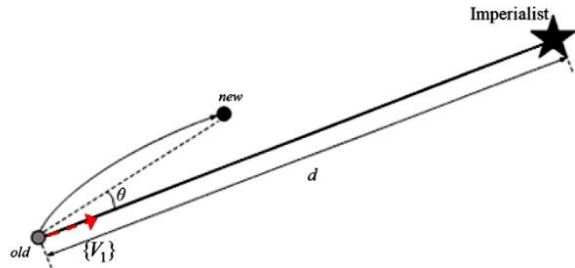


Figure 1. Movement of colonies to its new location in the original ICA

- (C)** Leadership transition based on fitness superiority: In the course of the optimization process, a colony may discover a solution with a lower cost (i.e., higher fitness) than that of its associated imperialist. When such an improvement occurs, a position exchange mechanism is triggered: the colony and the imperialist swap roles, and the colony is promoted to the status of imperialist by occupying the superior position in the search space. This dynamic update strategy ensures that the search process is continuously guided by the most promising candidates and facilitates convergence toward optimal or near-optimal solutions.
- (D)** Imperialistic competition and empire collapse

mechanism: During the evolution of the ICA, less dominant imperialists gradually lose control over their colonies due to their relatively inferior fitness. In contrast, more powerful imperialists competitively absorb these orphaned colonies, thereby expanding their influence. When an empire loses all of its colonies, it is considered collapsed and is eliminated from the search process. A probabilistic mechanism governs the acquisition of the collapsed empire's remaining resources, allowing competing empires to assimilate its remnants. The total cost of the n^{th} empire is evaluated by (6).

$$TC_n = f_{cost}^{(imp,n)} + \xi \frac{\sum_{i=1}^{NC_n} f_{cost}^{(col,i)}}{NC_n} \quad (6)$$

where ξ is an adjustable positive number [52].

- (E)** Convergence Criterion and Termination Condition: The iterative process of the Imperialist Competitive Algorithm continues until a single empire successfully dominates all others by assimilating every remaining colony within the population. At this convergence point, the prevailing empire is identified as the optimal solution or instance selection pattern. If such a global domination is not achieved, the algorithm returns to step (B), and the competition among empires resumes, ensuring ongoing exploration and exploitation within the solution space.

Talatahari, et al. [47] presented an improved ICA that uses a point out of vector which its start point is $\{z\}_{old}$ and its direction is toward $\{z\}_{new}$, as indicated in Figure 2. This improved algorithm is generated by changing (5) into (7).

$$\begin{aligned} \{Z\}_{new} &= \{Z\}_{old} + \beta \times d \times \{CM\} \otimes \{V_1\} \\ &\quad + CM \times \tan(\theta) \times d \times \{V_2\}, \\ \{V_1\} \cdot \{V_2\} &= 0, \\ \|\{V_2\}\| &= 1 \end{aligned} \quad (7)$$

where CM and θ are the random parameters and $\{V_2\}$ is perpendicular to $\{V_1\}$.

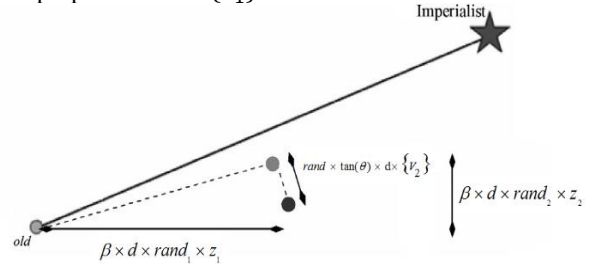


Figure 2. Movement of colonies to their new location in the improved ICA

4- PROPOSED METHOD

To address the class imbalance problem in binary classification tasks, the proposed method comprises two phases, as illustrated in Figure 3:

1. Generation of candidate sample subsets: Candidate subsets are produced via a global search strategy based on the Chaotic Imperialist Competitive Algorithm

(CICA). By leveraging chaos theory within the ICA, this approach enhances both the diversity of the candidate population and the algorithm's convergence behavior, promoting exploration and reducing the risk of premature convergence.

2. Assessment via a fuzzy weighted average distance-based decision surface: The evaluation phase employs a fuzzy weighted average distance-based decision surface. This method is designed to deliver robust class boundary modeling. By integrating uncertainty handling and instance-level difficulty metric, it ensures that the constructed decision boundaries are sensitive to the distinctive characteristics of imbalanced datasets and can better distinguish between minority and majority classes.

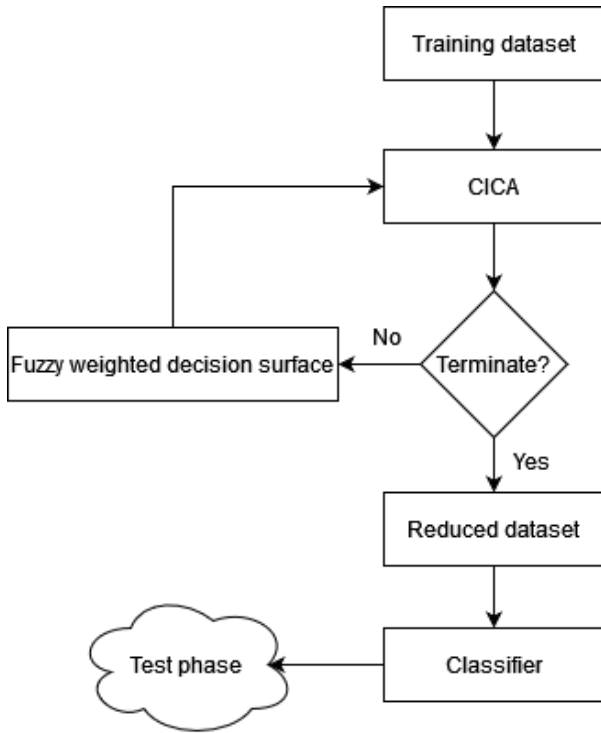


Figure 3. Flowchart of the proposed method.

A. Exploring the Solution Space

Instance selection can be formulated as a multi-objective optimization problem, where multiple candidate solutions—each representing a reduced dataset—are explored within a defined solution space. In the proposed approach, the conflicting objectives include maximizing the F-measure, G-mean, and accuracy, while simultaneously minimizing the reduction rate of the minority class. Equation (8) formalizes this unconstrained optimization problem aimed at selecting instances from an imbalanced dataset.

$$\begin{aligned}
 f_{\text{cost}} = & w_1 \cdot F\text{measure} + w_2 \cdot G\text{mean} \\
 & + w_3 \cdot \text{Accuracy} + w_4 \cdot (p_1 \cdot \text{Red}_{\text{min}} + p_2 \cdot \text{Red}_{\text{maj}}) \\
 & \sum_{i=1}^4 w_i = 1, \quad p_1 > p_2
 \end{aligned} \quad (8)$$

where

$$F\text{measure} = \frac{(1 + \alpha^2) \cdot \text{precision} \cdot \text{recall}}{\alpha^2 \cdot \text{precision} + \text{recall}} \quad (9)$$

$$G\text{mean} = \sqrt{TP \times TN} \quad (10)$$

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (11)$$

$$\text{Reduction} = \frac{\text{Total instances} - \text{Retained instances}}{\text{Total instances}} \quad (12)$$

where TP, FP, FN, and TN denote true positives, false positives, false negatives, and true negatives, respectively. A TP corresponds to a target instance correctly classified as target. A FP occurs when a non-target instance is erroneously classified as target. Conversely, a FN arises when a target instance is incorrectly identified as non-target. Finally, a TN indicates the correct classification of a non-target instance as non-target.

In (8), w_1, w_2, w_3 and w_4 are adaptive weighting parameters of each metric and p_1, p_2 are the penalty parameters. Regarding the minimization problem, p_1 and p_2 , i.e., reduction rate of the minority class should be lower than the reduction rate of the majority class.

The Chaotic Imperialist Competitive Algorithm (CICA), facilitates exploration across both minority and majority classes. It initializes with a random population represented as a binary array, where each element corresponds to a country. The presence or absence of a country is encoded as 1 or 0, respectively, with varying configurations representing different candidate solutions. The solution exhibiting the highest fitness value is selected, corresponding to the optimal instance selection state. By leveraging CICA, a solution pool is generated, each optimized for one of four distinct objectives, from which an appropriate reduced dataset can be chosen. Here, 10 percent of the countries are considered as empires ($N_{\text{imp}} = 10\%$) and 90 percent of them are used as colonies ($N_{\text{col}} = 90\%$). Besides, the cost function represented in (8) evaluates the corresponding cost value of each country presented in (2). Since utilizing random values for θ presented in (7) may get stuck in a local solution and speed down the convergence, the proposed method uses some chaotic maps listed in Table 4 for finding globally optimal solutions.

TABLE 4
Chaotic Functions

Name	Function
Logistic map	$z_{k+1} = az_k(1 - z_k)$
Tent map	$z_{k+1} = \begin{cases} 2z_k, & z_k < 0.5 \\ 2(1 - z_k), & z_k \geq 0.5 \end{cases}$
Sinusoidal map	$z_{k+1} = az_k^2 \sin(\pi z_k)$
Gauss map	$z_{k+1} = \begin{cases} 0 & z = 0 \\ \frac{1}{z_k} \bmod(1) & \text{otherwise} \end{cases}$

B. Difficulty-aware Decision Surface

Now, a classification method is introduced that integrates difficulty awareness into the decision-making process by employing a fuzzy membership function. This function assigns weights to training data points (candidate solutions) according to their assessed difficulty level. Initially, each training sample is evaluated using a difficulty metric that quantifies the ambiguity or uncertainty associated with that sample. The fuzzy membership function then converts these difficulty scores into corresponding weights, effectively emphasizing the influence of more challenging instances.

When classifying a new data point, the proposed decision surface assigns its label based on the proximity to the nearest class boundary. Importantly, the distance calculation incorporates the fuzzy membership weights, allowing the classification to factor in not only the feature similarity but also the difficulty of comparable training samples. This approach enhances robustness by recognizing that harder-to-classify instances contribute differently to shaping the decision boundary.

Consider N points $X = \{x_1, \dots, x_N\}$ where $x_i = \{x_{i1}, \dots, x_{id}\}$ distributed within a d -dimensional feature space. A label y_i assigns each point to one of l distinct classes. The distance between two points z_i and z_j in a high dimensional feature space given by (13) [53]:

$$\begin{aligned} D(\phi(z_i), \phi(z_j)) &= \|\phi(z_i) - \phi(z_j)\|_2 \\ &= \sqrt{(\phi(z_i) - \phi(z_j)) \cdot (\phi(z_i) - \phi(z_j))} \\ &= \sqrt{(\phi(z_i) \cdot \phi(z_i)) - 2(\phi(z_i) \cdot \phi(z_j)) + (\phi(z_j) \cdot \phi(z_j))} \\ &= \sqrt{k(z_i, z_i) - 2k(z_i, z_j) + k(z_j, z_j)} \end{aligned} \quad (13)$$

where ϕ is a mapping function transferred training instances into high-dimensional feature space, and $k(z_i, z_j)$ is a kernel function. Regarding a binary classification task, the class centers of positive and negative classes in a high-dimensional feature space are given by Eq.(14).

$$\begin{aligned} C^+ &= \frac{1}{l^+} \sum_{z_j \in \{+1\}} \phi(z_j) \\ C^- &= \frac{1}{l^-} \sum_{z_j \in \{-1\}} \phi(z_j) \end{aligned} \quad (14)$$

where l^+ and l^- are the size of positive and negative class, respectively.

The normalized distance between point z_i and the center of its class is determined by (15):

$$\Delta(z_i) = \begin{cases} \frac{\|\phi(z_i) - C^+\|}{\max \|\phi(z_i) - C^+\|}, & \text{if } z_i \in \{+1\} \\ \frac{\|\phi(z_i) - C^-\|}{\max \|\phi(z_i) - C^-\|}, & \text{if } z_i \in \{-1\} \end{cases} \quad (14)$$

By taking the median of all k -NN around instance z_i , the normalized density for z_i is computed as follows.

$$\rho(z_i) = \frac{\text{median} \left\{ \frac{1}{\|\phi(z_i) - \phi(z_1)\|}, \dots, \frac{k}{\|\phi(z_i) - \phi(z_k)\|} \right\}}{\max \left\{ \frac{1}{\|\phi(z_i) - \phi(z_1)\|}, \dots, \frac{k}{\|\phi(z_i) - \phi(z_k)\|} \right\}} \quad (15)$$

Finally, the membership degree of instance z_i is given by:

$$\mu(z_i) = \frac{\delta \cdot \rho(z_i)}{\Delta(z_i)} \quad (16)$$

where $\delta, \delta \in (0, 1]$ is a positive parameter which determines the importance of density. Clearly, the membership degree of each sample follows $\mu(z_i) \in [0, 1]$.

A notable challenge emerges when data points exhibit comparable membership degrees across multiple classes. As illustrated in Figure 4, instances located near the decision boundary (denoted by "+") may possess identical membership values, yet their contributions to each class can differ significantly. To enhance the modeling of such ambiguous and difficult instances, the framework is extended by incorporating non-membership and hesitation degrees, enabling a more nuanced representation of uncertainty.

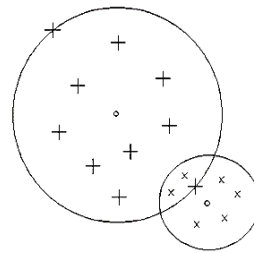


Figure 4. Illustration two instances (+) with equal distance to the center of the class

Consider k -NN for instance z_i . Let z_m be the m^{th} nearest neighbor of instance z_i . The hesitation degree for instance z_i is defined by (18).

$$h(z_i) = \begin{cases} \frac{\left| \left\{ z_j \mid \|\phi(z_m) - C^-\| \leq \|\phi(z_m) - C^+\| \right\} \right|_{j=1}^{j=k}}{k}, & y_i = +1 \\ \frac{\left| \left\{ z_j \mid \|\phi(z_m) - C^+\| \leq \|\phi(z_m) - C^-\| \right\} \right|_{j=1}^{j=k}}{k}, & y_i = -1 \end{cases} \quad (17)$$

where $h(z_i) \in [0,1]$. Figure 5 shows an example for computing the hesitation degree. The non-membership degree of instance z_i is defined by (19).

$$\eta(z_i) = \frac{h(z_i)}{\mu(z_i)} \quad (18)$$

According to (19), if $\mu(z_i) = 0$ then $\eta(z_i) \rightarrow \infty$. Also, the non-membership degree becomes zero if $h(z_i) = 0$; So, $\eta(z_i) \in [0, \infty)$.

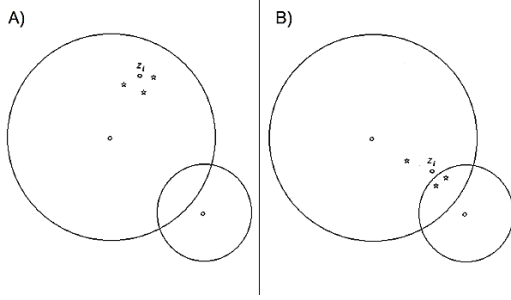


Figure 5. Computing the hesitation degree for instance z_i with different position. For $k=3$, the hesitation degree of z_i in (A) is zero whilst in (B) is 2/3

According to (19), if $\mu(z_i) = 0$ then $\eta(z_i) \rightarrow \infty$. Also, the non-membership degree becomes zero if $h(z_i) = 0$; So, $\eta(z_i) \in [0, \infty)$

Now, the weight of each instance is computed by the degree of membership and the degree of non-membership in (20):

$$\omega_i = \begin{cases} \mu_i & h_i \leq 0.3 \\ \frac{\mu_i}{\mu_i + \eta_i} & \text{otherwise} \end{cases} \quad (19)$$

Thus, the weighted average of each class is computed by (21).

$$C_{\omega}^+ = \frac{1}{l^+} \sum_{y_i=+1} \omega_i z_i$$

$$C_{\omega}^- = \frac{1}{l^-} \sum_{y_i=-1} \omega_i z_i \quad (20)$$

where l^+ and l^- are the size of the positive and negative classes, respectively. Then, instance z_i is classified by (22).

$$y_i = \text{sgn}(D(\phi(z_i), C_{\omega}^-) - D(\phi(z_i), C_{\omega}^+)) \quad (21)$$

5- EXPERIMENT

The experiments were conducted in a Jupyter Notebook environment using an Intel Core i5 processor and 16 GB of RAM. To assess the performance of the proposed method, classifiers including SVM, Convolutional Neural Network (CNN), and k-Nearest Neighbors with $k=10$ were employed. The experimental evaluation involved multiple kernel functions, such as Radial Basis Function (RBF) and polynomial kernels, alongside various chaotic maps

including logistic, tent, sinusoidal, and Gaussian maps. Table 5 summarizes the optimal parameter settings utilized throughout the experiments.

TABLE 5
Values of the Parameters

Parameter	Val.	Parameter	Val.	Parameter	Val.
δ	0.6	w_1	0.2	Imperialist size	10
α	0.8	w_2	0.2	p_1	0.9
ξ	0.3	w_3	0.2	p_2	0.1
Population size	100	w_4	0.3		

A. Experiments on Benchmark Datasets

In this section, experimental evaluations are performed on datasets of moderate scale. The experiments are carried out on a diverse set of UCI datasets. The detailed statistics of these datasets are summarized in Table 6. As shown in the last column of the table, which reports the IR, some datasets are relatively balanced ($IR < 5$), while most exhibit varying degrees of class imbalance. This distribution enables a thorough evaluation of the proposed method across both balanced and imbalanced settings. These datasets provide a representative testbed for validating the effectiveness and generalizability of the proposed approach.

The initial phase of the experiments involves examining various kernel functions and chaotic maps within the proposed framework to identify the optimal configuration. Performance metrics—including Accuracy, Reduction Rate, Fmeasure, and Effectiveness—are reported using 10-fold cross-validation across both balanced and imbalanced datasets, employing SVM, k-NN, and CNN classifiers. Figures 6 through 8 illustrate these evaluation metrics corresponding to the SVM, k-NN, and CNN classifiers, respectively. Notably, to jointly assess reduction rate and Fmeasure, the Effectiveness metric is defined as the *Reduction rate* $\times$ *Fmeasure*.

TABLE 6
Dataset Statistics

Dataset	#samples	#features	Class1 (%)	Class2 (%)	IR (%)
Yeast-6	1484	8	2.36	97.64	41.38
Yeast-5	1484	8	2.70	97.30	36.03
Yeast-1289vs7	947	8	3.17	96.83	30.55
Yeast-4	1484	8	3.44	96.56	28.07
Yeast-3	1484	8	10.98	89.02	8.11
Yeast-2vs8	483	8	4.14	95.86	23.15
Yeast-1	1484	8	28.90	71.10	2.46
Yeast-1458vs7	693	8	4.33	95.67	22.09
Yeast-1vs7	459	8	6.54	93.46	14.29
Ecoli-4	336	7	5.95	94.05	15.80
Ecoli-3	336	7	10.42	89.58	8.60
Ecoli-2	336	7	15.48	84.52	5.46
Ecoli-1	336	7	22.92	77.08	3.36
Ecoli-0vs1	220	7	35.00	65.00	1.86
Ecoli-0137vs26	282	7	2.48	97.52	39.32
Glass-6	214	9	6.07	93.93	15.50
Glass-5	214	9	4.20	95.80	22.80
Glass-016vs5	184	9	4.89	95.11	19.45
Glass-5vs12	159	9	8.18	91.82	11.26
Glass-2	214	9	13.55	86.45	6.38
Glass-123vs567	214	9	23.83	76.17	3.20
Glass-1	214	9	35.51	64.49	1.82
Glass-0	214	9	32.71	67.29	2.09

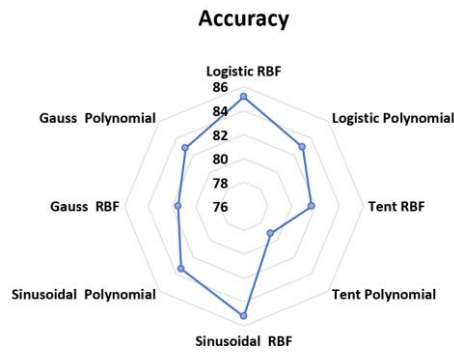


Figure 6. A. Accuracy of the proposed method on different kinds of kernel functions and chaotic maps using SVM classifier.

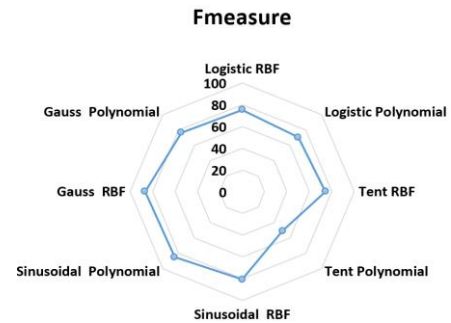


Figure 6. C. Fmeasure of the proposed method on different kinds of kernel functions and chaotic maps using SVM classifier.

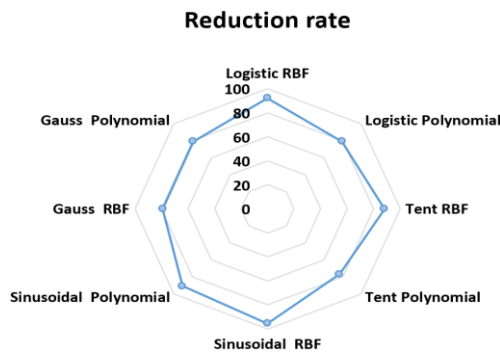


Figure 6. B. Reduction rate of the proposed method on different kinds of kernel functions and chaotic maps using SVM classifier.

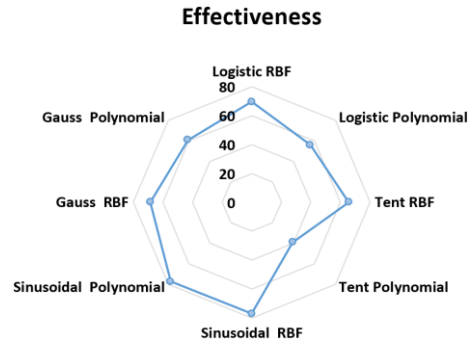


Figure 6. D. Effectiveness of the proposed method on different kinds of kernel functions and chaotic maps using SVM classifier.

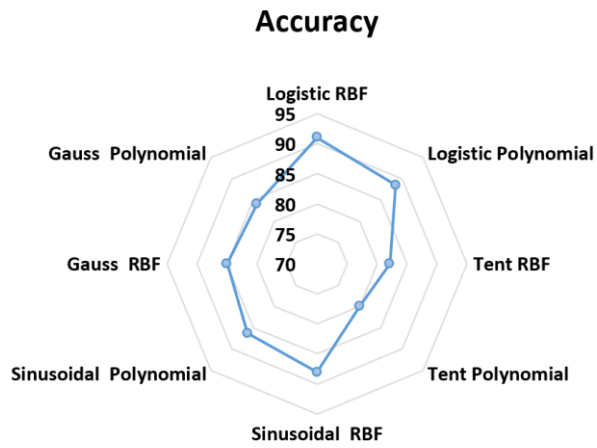


Figure 7. A. Accuracy of the proposed method on different kinds of kernel functions and chaotic maps using k -NN classifier.

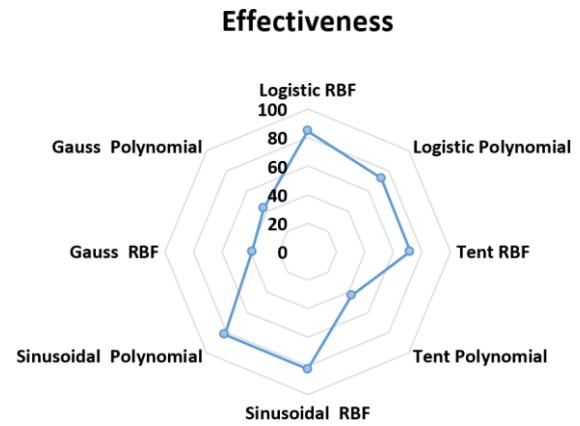


Figure 7. D. Effectiveness of the proposed method on different kinds of kernel functions and chaotic maps using k -NN classifier.

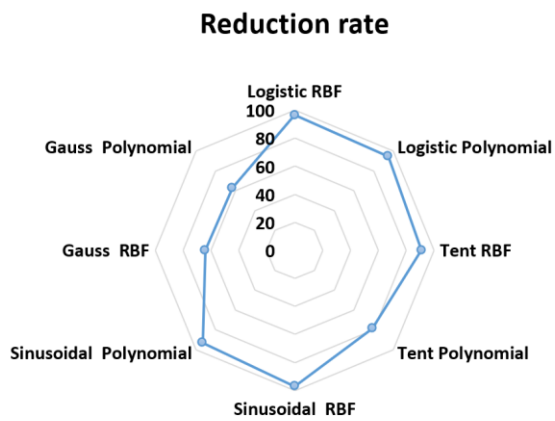


Figure 7. B. Reduction rate of the proposed method on different kinds of kernel functions and chaotic maps using k -NN classifier.

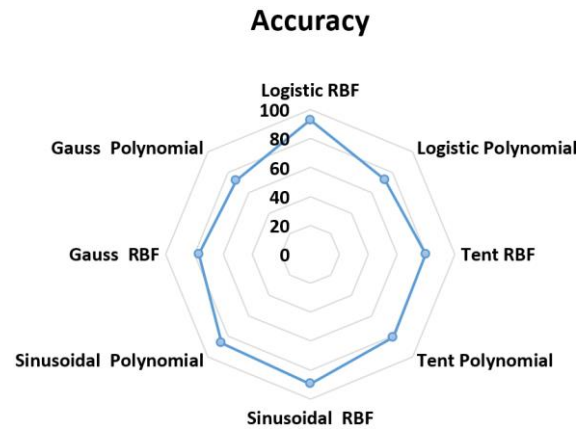


Figure 8. A. Accuracy of the proposed method on different kinds of kernel functions and chaotic maps using CNN classifier.

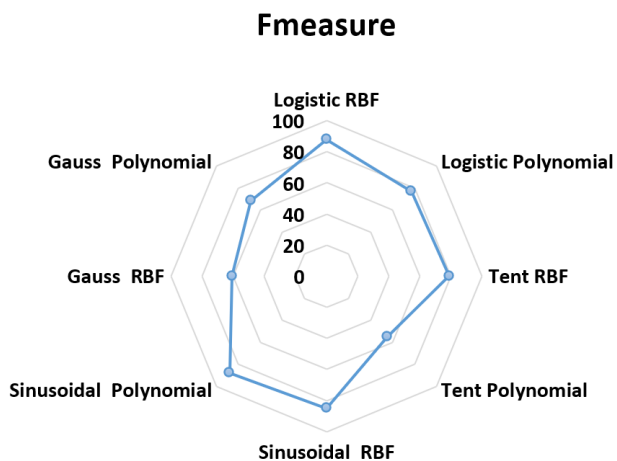


Figure 7. C. Fmeasure of the proposed method on different kinds of kernel functions and chaotic maps using k -NN classifier.

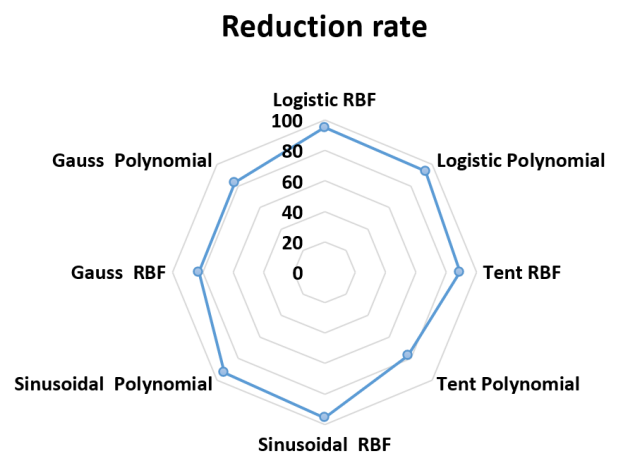


Figure 8. B. Reduction rate of the proposed method on different kinds of kernel functions and chaotic maps using CNN classifier.

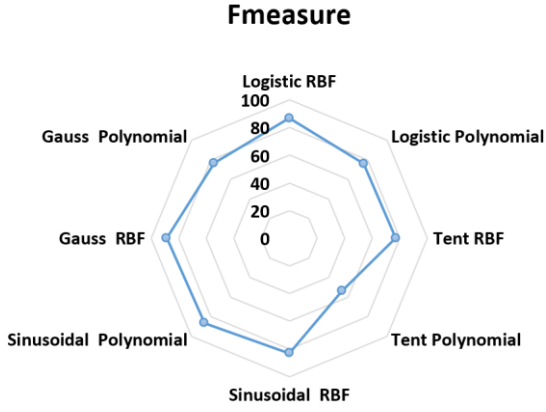


Figure 8. C. Fmeasure of the proposed method on different kinds of kernel functions and chaotic maps using CNN classifier.

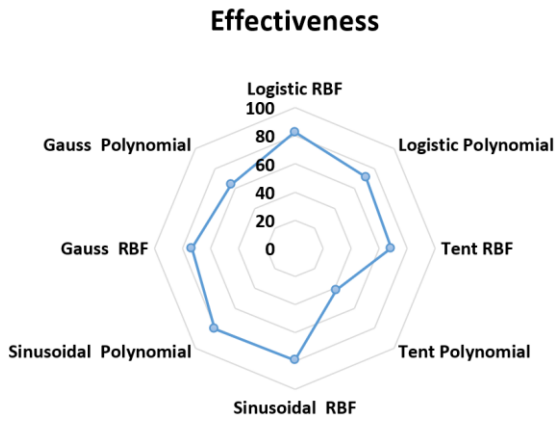


Figure 8. D. Effectiveness of the proposed method on different kinds of kernel functions and chaotic maps using CNN classifier.

According to Figures 6–8, the experimental results demonstrate that the RBF kernel consistently outperforms the polynomial kernel across all evaluation metrics. Although the polynomial kernel occasionally achieves competitive performance, it lacks stability and fails to deliver consistently acceptable results across all criteria. In contrast, the RBF kernel exhibits greater robustness and reliability, making it the preferred choice within the proposed framework.

Additionally, the figures highlight that the sinusoidal and logistic chaotic maps outperform other chaotic functions. The choice of chaotic map has a significant impact on both the reduction rate and classification performance metrics, such as accuracy and F-measure. These results emphasize the sensitivity of the method to the underlying chaotic dynamics. Based on the observations, the logistic map is recommended when minimizing classification error is the primary objective, while the sinusoidal map is more suitable when prioritizing dimensionality reduction. Both functions perform consistently well across most metrics; however, selecting the appropriate chaotic function based on the application's goals can further enhance overall effectiveness.

The class distribution before and after applying the

proposed method is illustrated in Figure 9. As discussed earlier, two major challenges in instance selection are:

- Preserving inter-class balance in originally balanced datasets, avoiding the introduction of skew during reduction;
- Retaining sufficient minority-class instances in imbalanced datasets to prevent information loss.

To evaluate the impact of the proposed method on class distribution, Figure 9 presents a comparative analysis of the IR, as well as the proportions of minority and majority instances, before and after instance selection. The results clearly show an increase in the proportion of minority-class instances and a decrease in majority-class instances. This indicates that the method selectively removes more majority instances, effectively emphasizing the preservation of minority-class information.

Moreover, the results confirm that the method maintains inter-class balance in balanced datasets and improves class representation in imbalanced settings. This section has presented strong empirical evidence supporting the framework's validity and effectiveness. The next section builds on these findings with further experiments.

B. Comparing with other methods

This section presents a comparative analysis between the proposed method and several competing methods. The evaluation uses the logistic chaotic map and the RBF kernel within our framework. To ensure a fair comparison, all methods employ a common SVM classifier.

The competing methods are selected based on their conceptual alignment with our approach, incorporating at least one of the following principles: evolutionary optimization, fuzzy modeling, difficulty estimation, or density-aware evaluation. This deliberate selection ensures a meaningful benchmark by attributing performance differences to the innovations of our design rather than fundamental methodological disparities.

The results are presented in Tables 7, 8, and 9, on the Yeast, Ecoli, and Glass datasets. Each table reports the average values of three key metrics—F-measure, G-mean, and AUC—along with their corresponding ranks. The Average Rank (AR) is provided in the final row of each table. Metric abbreviations include F1 (F-measure), Gm (Geometric Mean), AUC, and AR (Average Rank), consistently applied throughout.

According to the reported ranks, the proposed method consistently outperforms all competing approaches. Figures 10–12 complement the tables by graphically depicting the outcomes for F-measure, G-mean, and AUC, offering an intuitive visual comparison.



Figure 9. Class distribution (%) before and after using the proposed method.

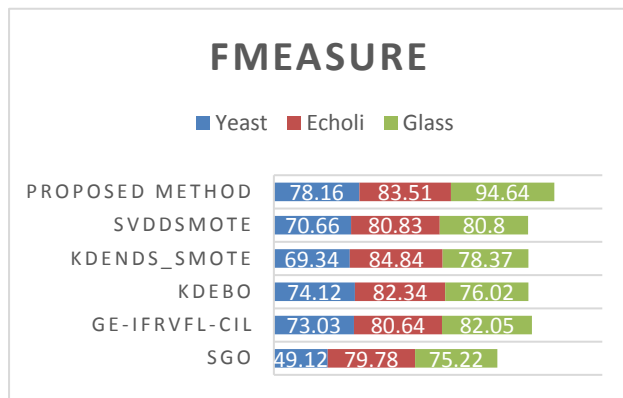


Figure 10. Fmeasure metric for the proposed method and other methods - all datasets.

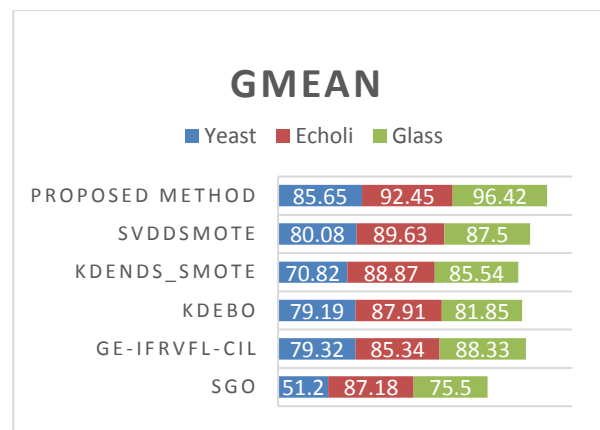


Figure 11. Gmean metric for the proposed method and other methods - all datasets.

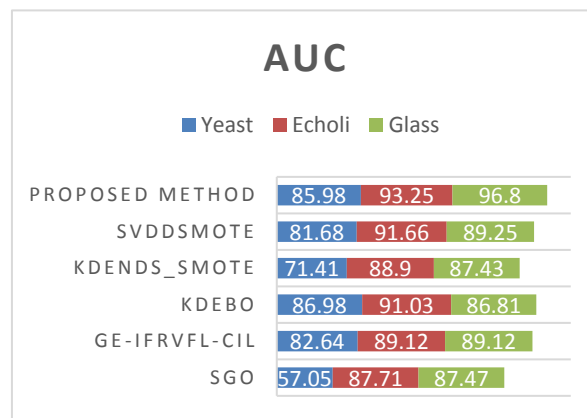


Figure 12. AUC metric for the proposed method and other methods - all datasets.

TABLE 7
Evaluation Metrics for the Proposed Method and the State-of-the-art Methods on Yeast Dataset (Average)

Metric	SGO	GE-IFRVFL-CIL	KDEBO	KDENDS_SMOTE	SVDDSMOTE	Proposed method
F1	49.12	73.03	74.12	69.34	70.66	78.16
	(6)	(3)	(2)	(5)	(4)	(1)
Gm	51.20	79.32	79.19	70.82	80.08	85.65
	(6)	(3)	(4)	(5)	(2)	(1)
AUC	57.05	82.64	86.98	71.41	81.68	85.98
	(6)	(3)	(1)	(5)	(4)	(2)
AR	6	3	2.33	5	3.33	1.33
	6	3	2	5	4	(1)

TABLE 8
Evaluation Metrics for the Proposed Method and the State-of-the-art Methods on Echoli Dataset (Average)

Metric	SGO	GE-IFRVFL-CIL	KDEBO	KDENDS_SMOTE	SVDDSMOTE	Proposed method
F1	79.78	80.64	82.34	84.84	80.83	83.51
	(6)	(5)	(3)	(1)	(4)	(2)
Gm	87.18	85.34	87.91	88.87	89.63	92.45
	(5)	(6)	(4)	(3)	(2)	(1)
AUC	87.71	89.12	91.03	88.90	91.66	93.25
	(6)	(4)	(3)	(5)	(2)	(1)
AR	5.6	5	3.33	3	3.33	1.33
	(5)	(4)	(3)	(2)	(3)	(1)

TABLE 9
Evaluation Metrics for the Proposed Method and the State-of-the-art Methods on Glass Dataset (Average)

Metric	SGO	GE-IFRVFL-CIL	KDEBO	KDENDS_SMOTE	SVDDSMOTE	Proposed method
F1	75.22	82.05	76.02	78.37	80.8	94.64
	(6)	(2)	(5)	(4)	(3)	(1)
Gm	75.50	88.33	81.85	85.54	87.5	96.42
	(6)	(2)	(5)	(4)	(3)	(1)
AUC	87.47	89.12	86.81	87.43	89.25	96.8
	(4)	(3)	(6)	(5)	(2)	(1)
AR	5.33	2.33	5.33	4.33	2.66	1
	(5)	(2)	(5)	(4)	(3)	(1)

The comparative results in Tables 7–9 and Figures 10–12 confirm the superiority of the proposed method across all key evaluation metrics. Several factors contribute to this

performance advantage over techniques such as SGO, GE-IFRVFL-CIL, KDEBO, KDENDS_SMOTE, and SVDDSMOTE:

- **Instance Selection over Generation:** Unlike most methods that focus on generating synthetic data, our approach emphasizes informed instance selection. This preserves data integrity and improves classifier generalization by avoiding synthetic noise and overlap.
- **Fuzzy Awareness and Adaptive Penalties:** Unlike SVDD-based oversampling methods with fixed parameters, our method leverages fuzzy membership weighting to dynamically adapt penalties, allowing better boundary modeling in sparse regions.
- **Fine-Grained Fitness Evaluation:** Our integration of fuzzy difficulty-awareness into the multi-objective evolutionary process allows instance contribution to be evaluated more precisely, outperforming approaches like GE-IFRVFL-CIL, which use fuzzy logic in a more coarse-grained fashion.
- **Chaotic Evolutionary Dynamics:** The use of chaotic maps in the evolutionary algorithm (unlike conventional mechanisms used in SGO or KDEBO) improves search-space exploration and convergence stability, particularly in noisy or high-dimensional scenarios.
- **Preservation of Class Distribution:** Our method maintains balance in originally balanced datasets and improves minority representation in imbalanced datasets without generating artificial data—a capability evidenced by Figure 9 and the superior ranks across evaluation tables.

C. Diverse Real-World Datasets Exhibiting Severe Class Imbalance

To evaluate the scalability and robustness of the proposed method under realistic, data-intensive conditions, we conduct extensive experiments on the large-scale dataset regarding CLIMB benchmark suite [58]. This benchmark is specifically designed to reflect the core challenges of industrial-scale imbalanced learning, including extreme class imbalance, high-dimensional feature spaces, noisy distributions, and structural irregularities. The CLIMB includes six datasets with $IR > 50$. Table 10 summarizes the key characteristics of these datasets. The proposed method is compared against the methods summarized in Table 11.

TABLE 10
Summary of Real-World Datasets with Severe Class Imbalance

Dataset Name	#samples	#features	IR	Domain
Dis	3772	29	64.03	Biology
Satellite	5100	36	67.0	Remote sensing
Employee-Turnover-at-TECHCO	34452	2	68.74	Human resources
Page-blocks	5473	10	175.46	Document processing
allbp	3772	29	257.79	Biology
Credit Card Fraud Detection	284807	30	577.88	Finance

To guarantee methodological alignment with CLIMB and ensure reproducibility, we replicate their entire evaluation pipeline:

- **Preprocessing:** Numerical features are standardized to zero mean and unit variance; categorical features are one-hot encoded; no imputation is required due to dataset completeness.
- **Validation Protocol:** A stratified 5-fold cross-validation scheme is adopted to preserve class ratios in each fold. Within each training fold, we perform a nested 3-fold tuning of hyperparameters using Optuna [59], with 100 trials per method and all random seeds fixed to 42. Optuna is an open-source hyperparameter optimization framework designed to automate the process of searching for the best hyperparameters for machine learning models. It provides an efficient, flexible way to handle optimization tasks by using algorithms such as Tree-structured Parzen Estimator (TPE) and integrates seamlessly with existing machine learning pipelines. Optuna is particularly noted for its ability to optimize hyperparameters dynamically, using a "define-by-run" approach, where users specify the optimization logic in a flexible and modular way.
- **Base Classifier:** SVM classifiers were used to evaluate the effect of instance reduction and sampling strategies, allowing a fair and interpretable comparison across all baselines.
- **Implementation Details:** All algorithms are implemented in a unified Python codebase, parallelized across folds, and executed on a server with 64 GB RAM. Large-scale encoding leverages sparse matrix representations to accommodate large scale datasets.

Table 11 reports the experimental results in terms of AUC, G-Mean, F-measure, and accuracy. The proposed method consistently outperforms all baseline approaches across these metrics. While ENN and IHT help clarify decision boundaries, they often fail to retain the most informative minority instances that our method preserves. EasyEnsemble and KDEBO partially address class imbalance but introduce high variance or synthetic noise. Similarly, AsymBoost improves minority recall but struggles to maintain overall accuracy, unlike the proposed method's multi-objective selection strategy. In contrast to synthetic methods such as ADASYN and SMBA, which expand the dataset, the proposed method selectively retains high-impact minority instances based on their difficulty and discriminative value. Several key factors contribute to the superior performance of the proposed method:

- **Overlap Reduction:** By integrating a fuzzy hardness metric with a chaotic multi-objective evolutionary process, the method avoids selecting noisy or overlapping samples—an issue that commonly degrades oversampling methods.

- **Scalability and Stability:** Experimental results demonstrate that the method scales well with data size, maintaining stable convergence and computational efficiency even on datasets exceeding 500,000 instances.

TABLE 11
Baseline Methods Selected for Evaluation

Selected Method	Category	Description	Ref.
Edited Nearest Neighbors (ENN)	Cleaning	Removes mislabeled/noisy majority samples by local disagreement filtering	[60]
Instance Hardness Threshold (IHT)	Simple Undersampling	Eliminates hard-to-classify majority examples via local hardness estimation	[61]
Easy Ensemble	Ensemble Undersampling	Trains multiple classifiers on balanced subsets with aggregated prediction	[62]
ADASYN	Oversampling	Adaptive oversampling focused on sparse and borderline minority regions	[63]
SMBA	Ensemble Oversampling	Applies bagging with local density-aware synthetic sampling per ensemble member	[64]
Cost-sensitive SVM	Cost-sensitive Learning	Adjusts misclassification costs to penalize majority-class errors less	[65]
KDEBO	Kernel density-based methods	Density-guided differential evolution oversampling in high-density minority regions to avoid noise	[35]

- **Granular Difficulty Modeling:** Fuzzy hardness enables more precise estimation of instance importance compared to binary selection rules or density-based heuristics.
- **Chaotic Multi-objective Optimization:** The use of the Imperialist Competitive Algorithm enhanced with chaotic maps promotes both rapid convergence and diverse solutions within a large search space.

These findings confirm that strategically guided instance selection—grounded in difficulty-awareness and optimized through chaotic evolutionary dynamics—

provides a robust and scalable solution to extreme class imbalance in large-scale tabular learning.

TABLE 12
Performance Comparison on Real-World Datasets with
Severe Class Imbalance

Method	AUC	G-Mean	Fmeasure	Balanced Acc.
ENN [60]	45.0	75.4	75.1	75.8
IHT [61]	33.0	74.6	68.0	79.9
EasyEnsemble [62]	35.5	83.0	89.8	87.3
ADASYN [63]	34.3	76.2	70.3	76.3
SMBA [64]	56.0	74.9	76.6	74.4
Cost-sensitive SVM [65]	41.7	73.9	74.9	74.8
KEDBO[35]	53.4	76.08	71.92	79.17
Proposed method	58.33	77.24	91.65	78.32

6- CONCLUSION and FUTURE WORK

This paper introduced a difficulty-aware instance reduction method that combines chaotic evolutionary optimization with fuzzy-informed instance selection to enhance classification. The method uses a chaotic Imperialist Competitive Algorithm (ICA) guided by a multi-objective fitness function balancing accuracy and fairness. A distance-weighted decision surface further ensures structural class separation. Comprehensive experiments on datasets with different intensities of imbalance demonstrated the method's superiority over competing methods. We also investigated the effect of kernel functions and chaotic maps. The RBF kernel with the logistic map yielded the best performance when accuracy was prioritized, while the sinusoidal map was more effective for dimensionality reduction. The chaotic components accelerated convergence and enhanced diversity in the solution space.

Key strengths of the proposed method include:

- Robust handling of class imbalance without synthetic data
- Flexibility for different optimization goals (e.g., accuracy vs. reduction rate)
- Consistent generalization across small and large datasets

Future work will explore the following directions: extending the method to multi-class scenarios with overlapping distributions, enhancing scalability through dimensionality-aware optimization and sparse instance modeling, integrating instance selection with graph-based learning or self-supervised representations, and applying the method to real-time or streaming environments with incremental instance selection and dynamic class balancing.

REFERENCES

- [1] T. Xia, T. Dang, J. Han, L. Qendro, and C. Mascolo. (2024). Uncertainty-aware health diagnostics via class-balanced evidential deep learning. *IEEE Journal of Biomedical and Health Informatics*. [Online]. 28(11), pp. 6417–6428. Available: <https://doi.org/10.1109/JBHI.2024.3360002>
- [2] Z. Nouri, V. Kiani, and H. Fadishei. (2024). Rarity updated ensemble with oversampling: An ensemble approach to classification of imbalanced data streams. *Statistical Analysis and Data Mining: The ASA Data Science Journal*. [Online]. 17(1), p. e11662. Available: <https://doi.org/10.1002/sam.11662>
- [3] L. Ren, R. Hu, Y. Liu, D. Li, J. Wu, Y. Zang, and W. Hu. (2024). Improving fraud detection via imbalanced graph structure learning. *Machine Learning*. [Online]. 113(3), pp. 1069–1090. Available: <https://doi.org/10.1007/s10994-023-06464-0>
- [4] S. Raza, M. Garg, D. J. Reji, S. R. Bashir, and C. Ding. (2024). NBias: A natural language processing framework for bias identification in text. *Expert Systems with Applications*. [Online]. 237, p. 121542. Available: <https://doi.org/10.1016/j.eswa.2023.121542>
- [5] C. Q. Zhang, Y. Deng, M. Z. Chong, Z. W. Zhang, and Y. H. Tan. (2024). Entropy-based re-sampling method on SAR class imbalance target detection. *ISPRS Journal of Photogrammetry and Remote Sensing*. [Online]. 209, pp. 432–447. Available: <https://doi.org/10.1016/j.isprsjprs.2024.02.019>
- [6] Q. Zhou and B. Sun. (2024). Adaptive K-means clustering based under-sampling methods to solve the class imbalance problem. *Data and Information Management*. [Online]. 8(3), p. 100064. Available: <https://doi.org/10.1016/j.dim.2023.100064>
- [7] N. Thai-Nghe, Z. Gantner, and L. Schmidt-Thieme, “Cost-sensitive learning methods for imbalanced data,” in *The 2010 International Joint Conference on Neural Networks (IJCNN)*, Barcelona, Spain, 2010, pp. 1–8. Available: <https://doi.org/10.1109/IJCNN.2010.5596486>
- [8] M. Moradi and J. Hamidzadeh. (2024). Handling class imbalance and overlap with a hesitation-based instance selection method. *Knowledge-Based Systems*. [Online]. 294, p. 111745. Available: <https://doi.org/10.1016/j.knosys.2024.111745>
- [9] J. Bi and C. Zhang. (2018). An empirical comparison on state-of-the-art multi-class imbalance learning algorithms and a new diversified ensemble learning scheme. *Knowledge-Based Systems*. [Online]. 158, pp. 81–93. Available: <https://doi.org/10.1016/j.knosys.2018.05.037>
- [10] I. Czarnowski. (2022). Weighted ensemble with one-class classification and over-sampling and instance selection (WECOI): An approach for learning from imbalanced data streams. *Journal of Computational Science*. [Online]. 61, p. 101614. Available: <https://doi.org/10.1016/j.jocs.2022.101614>
- [11] M. S. E. Shahabadi, H. Tabrizchi, M. K. Rafsanjani, B. B. Gupta, and F. Palmieri. (2021). A combination of clustering-based under-sampling with ensemble methods for solving imbalanced class problem in intelligent systems. *Technological Forecasting and Social Change*. [Online]. 169, p. 120796. Available:

- <https://doi.org/10.1016/j.techfore.2021.120796>
- [12] B. B. Hazarika and D. Gupta. (2021). Density-weighted support vector machines for binary class imbalance learning. *Neural Computing and Applications*. [Online]. 33(9), pp. 4243–4261. Available: <https://doi.org/10.1007/s00521-020-05240-8>
- [13] B. B. Hazarika and D. Gupta. (2023). Affinity-based fuzzy kernel ridge regression classifier for binary class imbalance learning. *Engineering Applications of Artificial Intelligence*. [Online]. 117, p. 105544. Available: <https://doi.org/10.1016/j.engappai.2022.105544>
- [14] B. B. Hazarika, D. Gupta, and P. Borah. (2023). Fuzzy twin support vector machine based on affinity and class probability for class imbalance learning. *Knowledge and Information Systems*. [Online]. 65(12), pp. 5259–5288. Available: <https://doi.org/10.1007/s10115-023-01904-8>
- [15] S. Guan, X. Zhao, Y. Xue, and H. Pan. (2024). AWGAN: An adaptive weighting GAN approach for oversampling imbalanced datasets. *Information Sciences*. [Online]. 663, p. 120311. Available: <https://doi.org/10.1016/j.ins.2024.120311>
- [16] Z. Wang and H. Wang, “Variational imbalanced regression: Fair uncertainty quantification via probabilistic smoothing,” *Advances in Neural Information Processing Systems*, vol. 36, 2024.
- [17] S. Tao, P. Peng, Y. Li, H. Sun, Q. Li, and H. Wang. (2024). Supervised contrastive representation learning with tree-structured Parzen estimator Bayesian optimization for imbalanced tabular data. *Expert Systems with Applications*. [Online]. 237, p. 121294. Available: <https://doi.org/10.1016/j.eswa.2023.121294>
- [18] C. Ye, R. Tsuchida, L. Petersson, and N. Barnes, “Label shift estimation for class-imbalance problem: A Bayesian approach,” in *Proceeding IEEE/CVF Winter Conference on Applications of Computer Vision (WACV)*, 2024, pp. 1073–1082.
- [19] A. Kumari, M. Tanveer, C. T. Lin, and Alzheimer’s Disease Neuroimaging Initiative. (2024). Class probability and generalized bell fuzzy twin SVM for imbalanced data. *IEEE Transactions on Fuzzy Systems*. [Online]. 32(5), pp. 3037–3048. Available: <https://doi.org/10.1109/TFUZZ.2024.3366936>
- [20] Z. Deng, P. Cui, and J. Zhu, “Towards accelerated model training via Bayesian data selection,” in *Advances in Neural Information Processing Systems*, vol. 36, 2023, pp. 8513–8527.
- [21] Y. Pu, W. Yao, and X. Li. (2024). EM-IFCM: Fuzzy c-means clustering algorithm based on edge modification for imbalanced data. *Information Sciences*. [Online]. 659, p. 120029. Available: <https://doi.org/10.1016/j.ins.2023.120029>
- [22] S. Bano, W. Zhi, B. Qiu, M. Raza, N. Sehito, M. M. Kamal, G. Aldehim, and N. Alruwais. (2024). Self-paced ensemble and big data identification: A classification of substantial imbalance computational analysis. *The Journal of Supercomputing*. [Online]. 80(7), pp. 9848–9869. Available: <https://doi.org/10.1007/s11227-023-05828-6>
- [23] N. Garcia-Pedrajas, J. M. Cuevas-Munoz, and A. de Haro-Garcia. (2024). Evolutionary simultaneous under- and oversampling of instances for dealing with class-imbalance datasets in multilabel problems. *Applied Soft Computing*. [Online]. 159, p. 111618. Available: <https://doi.org/10.1016/j.asoc.2024.111618>
- [24] X. Lu, X. Ye, and Y. Cheng. (2024). An overlapping minimization-based over-sampling algorithm for binary imbalanced classification. *Engineering Applications of Artificial Intelligence*. [Online]. 133, p. 108107. Available: <https://doi.org/10.1016/j.engappai.2024.108107>
- [25] Y. Liu, L. Zhu, L. Ding, H. Sui, and W. Shang. (2024). A hybrid sampling method for highly imbalanced and overlapped data classification with complex distribution. *Information Sciences*. [Online]. 661, p. 120117. Available: <https://doi.org/10.1016/j.ins.2023.120117>
- [26] T. Ma, S. Lu, and C. Jiang. (2024). A membership-based resampling and cleaning algorithm for multi-class imbalanced overlapping data. *Expert Systems with Applications*. [Online]. 240, p. 122565. Available: <https://doi.org/10.1016/j.eswa.2023.122565>
- [27] A. Kumar, D. Singh, and R. Shankar Yadav. (2024). Class overlap handling methods in imbalanced domain: A comprehensive survey. *Multimedia Tools and Applications*. [Online]. 83(23), pp. 63243–63290. Available: <https://doi.org/10.1007/s11042-023-17864-8>
- [28] Z. Sun, W. Ying, W. Zhang, and S. Gong. (2024). Undersampling method based on minority class density for imbalanced data. *Expert Systems with Applications*. [Online]. 249, p. 123328. Available: <https://doi.org/10.1016/j.eswa.2024.123328>
- [29] H. Yoon, J. Kim, and S. Lee, “Adaptive bandwidth kernel density estimation for imbalanced classification,” *Pattern Recognition Letters*, vol. 139, pp. 12–20, 2020.
- [30] Miraj, Mahabubur Rahman, et al. “GK-SMOTE: A Hyperparameter-Free Noise-Resilient Gaussian KDE-Based Oversampling Approach.” *Asia-Pacific Web (APWeb) and Web-Age Information Management (WAIM) Joint International Conference on Web and Big Data*. Singapore: Springer Nature Singapore, (2025).
- [31] Lee, Daeun, and Hyunjoong Kim. “Adaptive oversampling via density estimation for online imbalanced classification.” *Information* 16.1 (2025): 23. Available: <https://doi.org/10.3390/info16010023>.
- [32] Taccaliti, Edoardo, and Jesus S. Aguilar–Ruiz. “Improving classification on imbalanced genomic data via KDE–based synthetic sampling.” *BioData Mining* 18.1 (2025): 60. Available: <https://doi.org/10.1186/s13040-025-00474-5>
- [33] Li, Xinqi, and Qicheng Liu. “A hybrid sampling

- algorithm for imbalanced and class-overlap data based on natural neighbors and density estimation." *Knowledge and Information Systems* 67.3 (2025): 2259-2290.
- [34] C. Chen, *An empirical study of adaptive kernel density estimation in detecting distributional overlap*. Ph.D. dissertation, Delft University of Technology, 2023.
- [35] B. Turkoglu. (2022). KDEBO: A kernel density estimation-guided differential evolution-based oversampling technique for imbalanced learning. *Electronics*. [Online]. 11(17), p. 2703.
- [36] F. Kamalov, S. Moussa, and J. A. Reyes. (2022). KDE-based ensemble learning for imbalanced data. *Electronics*. [Online]. 11(17), p. 2703. Available: <https://doi.org/10.3390/electronics11172703>
- [37] F. Kamalov and A. Elnagar, "Kernel density estimation-based sampling for neural network classification," in *Proceeding 2021 International Symposium on Networks, Computers and Communications (ISNCC)*, 2021, pp. 1–6. Available: <https://doi.org/10.1109/ISNCC52172.2021.9615715>
- [38] F. Kamalov. (2020). Kernel density estimation based sampling for imbalanced class distribution. *Information Sciences*. [Online]. 512, pp. 1192–1201. Available: <https://doi.org/10.1016/j.ins.2019.10.017>
- [39] Y. C. Wang and C. H. Cheng. (2021). A multiple combined method for rebalancing medical data with class imbalances. *Computers in Biology and Medicine*. [Online]. 134, p. 104527. Available: <https://doi.org/10.1016/j.combiomed.2021.104527>
- [40] L. Yu, R. Zhou, L. Tang, and R. Chen. (2018). A DBN-based resampling SVM ensemble learning paradigm for credit classification with imbalanced data. *Applied Soft Computing*. [Online]. 69, pp. 192–202. Available: <https://doi.org/10.1016/j.asoc.2018.04.049>
- [41] H. Cui, Q. Wang, Y. Ye, Y. Tang, and Z. Lin. (2022). A combinational transfer learning framework for online transient stability prediction. *Sustainable Energy, Grids and Networks*. [Online]. 30, p. 100674. Available: <https://doi.org/10.1016/j.segan.2022.100674>
- [42] B. Hu, Y. Liu, N. Chen, L. Wang, N. Liu, and X. Cao. (2022). SEGCM-DCR: A syntax-enhanced event detection framework with decoupled classification rebalance. *Neurocomputing*. [Online]. 481, pp. 55–66. Available: <https://doi.org/10.1016/j.neucom.2022.01.069>
- [43] H. He and E. A. Garcia. (2009). Learning from imbalanced data. *IEEE Transactions on Knowledge and Data Engineering*. [Online]. 21(9), pp. 1263–1284. Available: <https://doi.org/10.1109/TKDE.2008.239>
- [44] J. Liu, Y. Sun, C. Han, Z. Dou, and W. Li, "Deep representation learning on long-tailed data: A learnable embedding augmentation perspective," In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2020, pp. 2970–2979.
- [45] T.-Y. Lin, P. Goyal, R. Girshick, K. He, and P. Dollár, "Focal loss for dense object detection," In *Proceedings of the IEEE international conference on computer vision*, 2017, pp. 2980–2988.
- [46] J. Liu, F. Guo, H. Gao, Z. Huang, Y. Zhang, and H. Zhou. (2021). Image classification method on class imbalance datasets using multi-scale CNN and two-stage transfer learning. *Neural Computing and Applications*. [Online]. 33(21), pp. 14179–14197. Available: <https://doi.org/10.1007/s00521-021-06066-8>
- [47] Z. Kang, Y. Zhang, J. Zhang, C. Li, M. Kong, Y. Zhao, and X. B. Wu. (2023). Periodic variable star classification with deep learning: Handling data imbalance in an ensemble augmentation way. *Publications of the Astronomical Society of the Pacific*. [Online]. 135(1051), p. 094501. Available: <https://doi.org/10.1088/1538-3873/acf15e>
- [48] L. Xiang, G. Ding, and J. Han, "Learning from multiple experts: Self-paced knowledge distillation for long-tailed classification," in *Proceeding Springer*, 2020, pp. 247–263. Available: https://doi.org/10.1007/978-3-030-58558-7_15
- [49] X. Wang, L. Lian, Z. Miao, Z. Liu, and S. X. Yu. (2020). Long-tailed recognition by routing diverse distribution-aware experts. *arXiv preprint*. [Online]. arXiv:2010.01809. Available: <https://doi.org/10.48550/arXiv.2010.01809>
- [50] Y. Chen, X. Yang, and H. L. Dai. (2024). Cost-sensitive continuous ensemble kernel learning for imbalanced data streams with concept drift. *Knowledge-Based Systems*. [Online]. 284, p. 111272. Available: <https://doi.org/10.1016/j.knosys.2023.111272>
- [51] E. Atashpaz-Gargari and C. Lucas, "Imperialist competitive algorithm: An algorithm for optimization inspired by imperialistic competition," in *IEEE congress on evolutionary computation*, 2007, pp. 4661–4667. Available: <https://doi.org/10.1109/CEC.2007.4425083>
- [52] S. Talatahari, B. F. Azar, R. Sheikholeslami, and A. Gandomi. (2012). Imperialist competitive algorithm combined with chaos for global optimization. *Communications in Nonlinear Science and Numerical Simulation*. [Online]. 17(3), pp. 1312–1319. Available: <https://doi.org/10.1016/j.cnsns.2011.08.021>
- [53] V. H. Moghaddam and J. Hamidzadeh. (2016). New Hermite orthogonal polynomial kernel and combined kernels in support vector machine classifier. *Pattern Recognition*. [Online]. 60, pp. 921–935. Available: <https://doi.org/10.1016/j.patcog.2016.07.004>
- [54] W. Pei et al. (2024). EvoSampling: A granular ball-based evolutionary hybrid sampling with knowledge transfer for imbalanced learning. *arXiv preprint*. [Online]. arXiv:2412.10461. Available: <https://doi.org/10.48550/arXiv.2412.10461>
- [55] M. A. Ganaie et al. (2024). Graph embedded intuitionistic fuzzy random vector functional link neural network for class imbalance learning. *IEEE Transactions on Neural Networks and Learning*

- Systems. [Online]. 35(9), pp. 11671–11680. Available: <https://doi.org/10.1109/TNNLS.2024.3353531>
- [56] R. Zhang et al. (2023). A density-based oversampling approach for class imbalance and data overlap. *Computers & Industrial Engineering*. [Online]. 186, p. 109747. Available: <https://doi.org/10.1016/j.cie.2023.109747>
- [57] X. Tao et al. (2022). SVDD-based weighted oversampling technique for imbalanced and overlapped dataset learning. *Information Sciences*. [Online]. 588, pp. 13–51. Available: <https://doi.org/10.1016/j.ins.2021.12.066>
- [58] Z. Liu, Z. Li, Z. Yang, et al. (2025). CLIMB: Class-imbalanced learning benchmark on tabular data. *arXiv preprint*. [Online]. arXiv:2505.17451. Available: <https://doi.org/10.48550/arXiv.2505.17451>
- [59] T. Akiba et al., “Optuna: A next-generation hyperparameter optimization framework,” in *Proceedings of the 25th ACM SIGKDD international conference on knowledge discovery & data mining*, 2019. <https://doi.org/10.1145/3292500.3330701>
- [60] D. Wilson. (1972). Asymptotic properties of nearest neighbor rules using edited data. *IEEE Transactions on Systems, Man, and Cybernetics*. [Online]. SMC-2(3), pp. 408–421. Available: <https://doi.org/10.1109/TSMC.1972.4309137>
- [61] Smith, Michael R., Tony Martinez, and Christophe Giraud-Carrier. "An instance level analysis of data complexity." *Machine learning* 95.2 (2014): 225-256. Available: <https://doi.org/10.1007/s10994-013-5422-z>
- [62] Liu, Xu-Ying, Jianxin Wu, and Zhi-Hua Zhou. "Exploratory undersampling for class-imbalance learning." *IEEE Transactions on Systems, Man, and Cybernetics, Part B (Cybernetics)* 39.2 (2008): 539-550. Available: <https://doi.org/10.1109/TSMCB.2008.2007853>
- [63] H. He, Y. Bai, E. A. Garcia, and S. Li, “ADASYN: Adaptive synthetic sampling approach for imbalanced learning,” in *Proceedings of the International Joint Conference on Neural Networks*, 2008, pp. 1322–1328. Available: <https://doi.org/10.1109/IJCNN.2008.4633969>
- [64] Wang, Shuo, and Xin Yao. "Diversity analysis on imbalanced data sets by using ensemble models." 2009 IEEE symposium on computational intelligence and data mining. IEEE, 2009. Available: <https://doi.org/10.1109/CIDM.2009.4938667>
- [65] P. Viola and M. Jones, “Fast and robust classification using asymmetric AdaBoost,” in *Neural Information Processing Systems (NeurIPS)*, 2001.

