

POSITION AND DENSITY BASED KERNEL PARAMETER SELECTION FOR ONE-CLASS SVM

JILI¹, YI HUANG¹

¹ Department of Mathematics and Statistics, University of Maryland, Baltimore County, Baltimore, MD 21250, USA
E-MAIL: jil1@umbc.edu, yihuang@umbc.edu

Abstract:

One-Class Support Vector Machine (OCSVM) is widely used for anomaly detection, yet its performance is highly sensitive to kernel parameter selection, particularly in the Gaussian kernel. Existing geometry-based criteria rely on extreme-distance measurements, making them highly sensitive to sample sparsity and dimensionality. In this paper, a robust parameter selection framework is proposed for OCSVM that integrates both positional and density information of training samples through a continuous weighting scheme. A statistically motivated objective function is further introduced and incorporates a penalty term to mitigate overfitting. Experiments on benchmark datasets demonstrate that the proposed method consistently achieves performance comparable to or better than two OCSVM settings, Isolation Forest and Local Outlier Factor. The improvements are validated by statistical significance tests, indicating that the proposed method provides a stable and effective solution for kernel parameter selection in OCSVM.

Keywords:

One-class classification; One-class support vector machine; Anomaly detection; Gaussian kernel; Kernel parameter selection

1. Introduction

Anomaly detection plays a critical role in a wide range of fields, including clinical trials, fraud detection, medical diagnosis, network security, etc. [1], [2]. In many real-world scenarios, labeled data are highly imbalanced, with abundant normal observations but limited anomalies. One-Class Support Vector Machine (OCSVM) provides an effective solution for such problems by learning a decision boundary that encloses the majority of the data. It constructs a boundary in a high-dimensional feature space such that most training samples lie within the boundary, while samples falling outside are consid-

ered anomalies. OCSVM aims to determine whether a new observation belongs to the same cluster as the training data. The flexibility of OCSVM largely depends on the kernel function, with the Gaussian Radial Basis Function (RBF) being the most commonly used.

Despite its effectiveness, OCSVM faces several challenges. First, its performance is highly sensitive to parameter selection, particularly the kernel parameter γ and the regularization parameter ν . While ν is often set based on prior knowledge of anomaly proportion, selecting γ remains difficult due to the lack of labeled negative samples. Kernel parameter selection in OCSVM can be viewed as an unsupervised model complexity control problem, where conventional cross-validation is inapplicable. To address the challenges, the existing approach [3] optimizes γ by maximizing the distance difference between interior and edge samples. However, this approach is not stable dealing with high-dimensional data.

In this paper, we propose a position and density based weighting (PDW) approach for kernel parameter selection in OCSVM. PDW uses a continuous weighting scheme that integrates positional and local density information to better distinguish interior and edge samples. Based on these weights, we develop a variance-normalized objective function with a penalty, which reduces sensitivity to extreme samples. Experiments on synthetic and benchmark datasets show that PDW achieves competitive or superior performance over the default setting and other existing methods, with improvements supported by statistical significance tests.

2. Existing Approaches

2.1. OCSVM

OCSVM is a One-Class Classification machine learning method first proposed by Schölkopf et al.[4] for estimating the

support of a high-dimensional distribution. Given training vectors $\mathbf{x}_i \in R^n, i = 1, \dots, l$ without any class information, the optimization problem of OCSVM is

$$\begin{aligned} \min_{\mathbf{w}, \xi, \rho} \quad & \frac{1}{2} \|\mathbf{w}\|^2 + \frac{1}{\nu l} \sum_{i=1}^l \xi_i - \rho \\ \text{s.t.} \quad & \langle \mathbf{w}, \varphi(\mathbf{x}_i) \rangle \geq \rho - \xi_i, \\ & \xi_i \geq 0, \quad i = 1, \dots, l, \end{aligned} \quad (1)$$

where w is the normal vector to the hyperplane, $\phi(x_i)$ maps x_i into a higher-dimensional space, ρ is the offset of the hyperplane from the origin in the feature space, ξ_i are slack variables allowing for outliers, ν is a parameter that controls the margin size and the number of outliers, l is the sample size of training data. Here the feature space is defined as a space mapping from the original training data.

Based on the Lagrangian method, the original optimization problem can be converted into a dual problem:

$$\begin{aligned} \min_{\alpha} \quad & \frac{1}{2} \alpha^T Q \alpha \\ \text{subject to} \quad & 0 \leq \alpha_i \leq \frac{1}{\nu l}, \quad i = 1, \dots, l, \\ & \mathbf{e}^T \alpha = 1, \end{aligned} \quad (2)$$

where $\mathbf{e} = [1, \dots, 1]^T$, $Q_{ij} = K(\mathbf{x}_i, \mathbf{x}_j) = \phi(\mathbf{x}_i)^T \phi(\mathbf{x}_j)$. Here K is the pre-determined kernel function. A kernel trick has been used to avoid mapping and calculating inner products. We obtain the OCSVM decision function

$$f(\mathbf{x}) = \langle \mathbf{w}, \varphi(\mathbf{x}) \rangle - \rho = \sum_{i=1}^n \alpha_i K(\mathbf{x}_i, \mathbf{x}) - \rho. \quad (3)$$

A non-negative $f(\mathbf{x})$ indicates a target data; otherwise it is an anomaly. To get a more stable estimate of ρ , we usually average over all support vectors:

$$\rho = \frac{1}{|S|} \sum_{i \in S} \left(\sum_{j=1}^l \alpha_j K(\mathbf{x}_j, \mathbf{x}_i) \right), \quad (4)$$

where S denotes the set of indices corresponding to the support vectors, and $|S|$ is the number of support vectors. The OCSVM optimization problems are solved by LIBSVM [5].

Radial Basis Function (RBF) kernel is most commonly used in OCSVM and defined as follows:

$$K(x, x') = \exp(-\gamma \|x - x'\|^2). \quad (5)$$

For simplicity, we fix the parameter $\nu = 0.05$ and tune γ in the RBF kernel using our proposed method in this paper.

2.2. MISE method

Xiao et al. [3] proposed a novel method Maximizing the distance difference between Interior and Edge Samples (MIES) to tune the parameter γ in an RBF kernel. First, they identify edge samples and interior samples, and measure the distances from these samples to the OCSVM boundary. Edge samples are on the extremes or edge of a class region. Interior samples are within the region formed by edges, which should be in the middle of the data cloud. Then based on this measurement, an optimization objective function for the parameter selection is calculated, which is the difference between the maximum of interior distances and the maximum of edge distances. γ is selected when the objective function is maximized.

To select edge and interior samples, MISE method follows Li and Maguire's method called Border-Edge Pattern Selection (BEPS) [6]. The idea of this method comes from the geometric information, the relative positions of samples with respect to the data boundary. For each data point x , the direction of its gradient vector is approximated by the mean direction of its K -nearest neighbors (KNN), where K is set to $5 \ln N$ and N is the sample size. Within the neighborhood of x , we can examine the position of its neighbors with respect to the gradient direction. If x lies near the boundary, most of its neighbors tend to be located on one side of the gradient direction. In contrast, if x is located in the interior region, its neighbors are more evenly distributed on both sides. Therefore, a proportion close to 1 indicates an edge point, while a proportion close to 0.5 suggests an interior point.

After selecting edge and interior samples, they propose the following optimization objective function:

$$f_O(s) = \max_{x_i \in \Omega_{IN}} d_N(x_i) - \max_{x_j \in \Omega_{ED}} d_N(x_j), \quad (6)$$

where Ω_{IN} and Ω_{ED} denote the sets of interior samples and edge samples, respectively. The objective aims to maximize the separation between interior and edge samples in terms of their distances to the decision boundary. d_N , the normalized distance from the mapping of a data z to the decision boundary of OCSVM is calculated as

$$\begin{aligned} d_N(\mathbf{z}) &= \frac{f(\mathbf{z})}{\|\mathbf{w}\| - \rho} \\ &= \frac{\sum_{i=1}^n \alpha_i k(\mathbf{x}_i, \mathbf{z}) - \rho}{\sqrt{\sum_{i=1}^n \sum_{j=1}^n \alpha_i \alpha_j k(\mathbf{x}_i, \mathbf{x}_j) - \rho}}. \end{aligned} \quad (7)$$

3. Position and Density based Weighting (PDW) for Kernel Parameter Selection

Although Xiao et al. proposed an interesting approach for tuning the γ parameter in the RBF kernel, several limitations can be observed. First, as mentioned in their paper, for datasets with small sample sizes and high dimensionality, it is possible that no interior samples are identified. In such cases, samples not classified as edge points are treated as interior points, which may blur the distinction between edge and interior regions and reduce the robustness of the method. Second, the objective function is defined as the difference between the maximum distances of edge samples and interior samples. This formulation is highly sensitive to extreme values. In particular, if some samples are incorrectly identified as edge samples, a few samples with unusually large distances can dominate the objective function, resulting in unstable parameter selection. These limitations motivate the need for a more robust criterion that better captures the overall separation between edge and interior samples while reducing sensitivity to outliers.

3.1. Identification of edge and interior samples

Motivated by Zhu et al. [7], to better identify edge and interior samples, we propose to define weights for all data points in the training data. If a data point is an interior, it should have a high weight; if it is an edge, it should have a low weight. And the weights should base on both position and density of the training data.

Let $\bar{x}_i$ be the mean of KNN of sample. The vector $\bar{x}_i - x_i$ divides the neighborhood sphere into two parts. If there are more neighboring samples in one part than in the other, the data point is more likely to be an edge point. Let $\theta_{i,j}$ ($j = 1, 2, \dots, k$) denote the angle between $\bar{x}_i - x_i$ and $x_{i,j} - x_i$, where $k = 5 \ln N$ and N is the sample size. If $x_{i,j}$ lies in the same part as $\bar{x}_i$, then $0 \leq \theta_{i,j} \leq \frac{\pi}{2}$ and $0 \leq \cos \theta_{i,j} \leq 1$; otherwise, $\frac{\pi}{2} \leq \theta_{i,j} \leq \pi$ and $-1 \leq \cos \theta_{i,j} \leq 0$.

In order to quantitatively describe the relationship between the distribution of neighborhood samples and the position of sample x_i , we define the following:

$$c^{sum}(x_i) = \sum_{j=1}^k \cos \theta_{i,j} = \sum_{j=1}^k \frac{\langle \bar{x}_i - x_i, x_{i,j} - x_i \rangle}{\|\bar{x}_i - x_i\| \cdot \|x_{i,j} - x_i\|}, \quad (8)$$

where $x_{i,j} \in KNN(x_i)$, $j = 1, 2, \dots, k$. By applying the kernel trick with $K(x_i, x_j) = \langle x_i, x_j \rangle$, we can calculate the statistic $c^{sum}(x_i)$. It is easy to see that $c^{sum}(x_i)$ lies in the interval $(-k, k)$.

Since abnormal samples are usually located near the boundary of the training data, we assign smaller weights to such samples in order to reduce their influence. Thus, the position-based weight is defined as

$$pos(x_i) = 1 - \frac{1}{k} |c^{sum}(x_i)|. \quad (9)$$

This weighting scheme assigns higher weights to interior samples and lower weights to edge samples, reflecting their relative positions in the data distribution.

If only the position of a data point and its KNN is considered, the weighting strategy can be overly one-sided. To address this limitation and incorporate local density, we propose a density based function motivated by Zhao et al. [8] to combine the previously discussed position based strategy. For sample x_i , a smaller weight is assigned when the distribution density of its KNN is low.

We propose a new smooth function to represent density-based weighting instead of the function in Zhao's paper:

$$den(x_i) = \exp \left(\log(a) \cdot \frac{\max \left(0, \frac{\tilde{x}_i}{\max(\tilde{x}_i)} - a \right)}{1 - a} \right), \quad (10)$$

where $a = \frac{\min(\tilde{x}_i)}{\max(\tilde{x}_i)}$. The exponential transformation ensures a smooth decay, preventing large weight shifts caused by minor density fluctuations. In this function,

$$\hat{x} = \frac{\tilde{x}_i}{\max_{i=1, \dots, N} \tilde{x}_i}, \quad (11)$$

where $\tilde{x}_i = \frac{1}{k} \sum_{j=1}^k d(x_i, x_{i,j})$, $x_{i,j}, j = 1, 2, \dots, k \in KNN(x_i)$, and $d(x_i, x_{i,j}) = \|\varphi(x_i) - \varphi(x_{i,j})\|_2 = \sqrt{2 - 2K(x_i, x_{i,j})}$.

A small mean distance $\tilde{x}_i$ indicates that a data point lies in a dense region, while a large $\tilde{x}_i$ means the point is more isolated (far from its neighbors). Then $\tilde{x}_i$ is normalized by $\max(\tilde{x}_i)$. The parameter a adjusts the scaling depending on the spread between the most central and most isolated points. If all $\tilde{x}_i$ values are similar, a is close to 1; if there is a large spread, a is small.

Combining the pre-defined position-based and density-based weights, the final weighting function is

$$w(x_i) = \lambda pos(x_i) + (1 - \lambda) den(x_i), \quad (12)$$

where $\lambda \in [0, 1]$ is a trade-off parameter. In this paper, λ is fixed at 0.5, implying equal weighting of position and density in the identification of edge and interior samples. In practice, λ can be flexibly tuned according to the need.

The objective function t is constructed as a t-statistic multiplied by a penalty term:

$$f_t = \left(\frac{|\overline{d_N(x_i)}| - |\overline{d_N(x_j)}|}{\sqrt{\frac{s_i^2}{n_i} + \frac{s_j^2}{n_j}}} \right) \left(1 - \frac{s}{N} \right)^2, \quad (13)$$

where s is the number of support vectors, N is the total sample size, n_i and n_j are the sample sizes of interiors and edges, x_i are interiors, x_j are edges, $|\overline{d_N(x_i)}|$ is the mean of $|d_N(x_i)|$, $|\overline{d_N(x_j)}|$ is the mean of $|d_N(x_j)|$, and s_i^2 and s_j^2 are the variances of $|d_N(x_i)|$ and $|d_N(x_j)|$, respectively.

The objective function f_O is changed to the objective function f_t to better measure the distance difference between the interiors and edges to the boundary. The max-distance criterion used in MIES is dominated by extreme samples and ignores distributional variability. We reinterpret interior-edge separation as a hypothesis testing problem and construct a variance-normalized statistic. The penalty term $(1 - \frac{s}{N})^2$ regularizes boundary complexity to avoid an overfitting problem. We can select an optimal γ in RBF such that the objective function f_t is maximized.

The algorithm is summarized as follows. The percentages can be flexibly adjusted, and the performance is not sensitive to their specific values.

Algorithm 1 PDW Kernel Parameter Selection in OCSVM

- 1: For each parameter (γ) in the candidate pool, calculate the weights of all training samples in Equation 12 using the RBF kernel.
 - 2: Select the optimal parameter as follows:
 - 2.1 Identify edge and interior samples in the training data. Edge samples are defined as those with the lowest 5% of weights, while interior samples have the highest 10%.
 - 2.2 If the distance of an interior sample to the boundary is smaller than the maximum distance of edge samples, this sample is not assigned to either the interior or edge group.
 - 2.3 Select the parameter that maximizes the objective function in Equation 13.
 - 3: Construct the decision function $f(x)$ for the OCSVM.
-

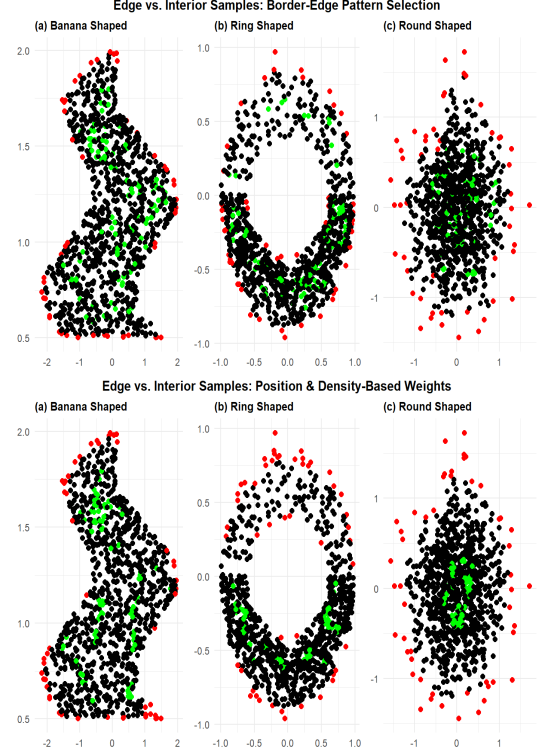


FIGURE 1. Edge and Interior Identification

4. Simulation studies

4.1. Experiments for edge and interior identification

To visually compare the performance of the proposed PDW approach with the BEPS method in identifying edge and interior samples, we conducted experiments on three types of two-dimensional synthetic datasets: banana-shaped, ring-shaped, and round-shaped, each consisting of 1000 data points. The banana-shaped dataset has a uniform distribution, while the ring-shaped dataset exhibits higher density in the lower half. The round-shaped dataset has a higher density in the central region.

In Figure 1, green dots represent interior samples, red dots represent edge samples, and black dots correspond to other samples. For the banana-shaped dataset, both methods produced similar results. For the ring-shaped dataset, the result illustrated that incorporating density prevents data in low density regions from being misclassified as interior. For the round-shaped dataset, PDW yielded a more concentrated set of interior samples near the center. These results indicate that the

proposed method effectively incorporates density information when identifying edge and interior samples, and shows clearer separation between interior and edge regions.

4.2. Experiments on Benchmark Datasets

To evaluate the performance of PDW on datasets with varying dimensionalities, experiments are conducted on five benchmark datasets: yeast, shuttle, breast, and ionosphere datasets are from the UCI Machine Learning Repository, while the phoneme dataset originates from the ELENA Project. For each dataset, we randomly selected 80% of the positive samples as the training data; the remaining positive samples and 10% of its sample size negative samples as the testing data. All features were normalized to zero mean and unitary variance. The default γ setting is $\frac{1}{\text{number of features}}$. The candidate γ pool included 0.01, 0.02, 0.03, 0.05, 0.1, 0.2, 0.3, 0.5 along with the default setting. The tradeoff parameter ν was set to 0.05. For each dataset, we repeated the above training and testing procedure 100 times.

To compare the effectiveness of the proposed γ selection method, we compared it with two widely used unsupervised anomaly detection methods: Isolation Forest (iForest) [9] and Local Outlier Factor (LOF) [10]. We implemented iForest using 200 trees with a subsampling size of $\min(256, n)$, following the standard in the literature. In LOF, the number of neighbors was set to $k = 5 \log n$. Since all methods produced continuous anomaly scores, a threshold was required to obtain binary predictions. For these two methods, we assumed an anomaly proportion $\nu = 0.1$ and classified the top ν fraction of samples with the highest anomaly scores as outliers.

For OCSVM, we fixed the parameter ν to 0.05 across all experiments. The kernel function was the radial basis function (RBF) kernel. The key parameter γ was selected from a predefined grid. Three strategies were compared: (1) Proposed PDW method: selecting γ using the proposed density- and position-based weighting scheme. (2) Default: using the standard implementation without tuning γ ; (3) MISE: selecting γ based on the existing MISE criterion. To ensure a fair evaluation, γ was selected using training data only, without accessing any testing labels.

We evaluated model performance using $G\text{-mean} = \sqrt{\text{Sensitivity} \times \text{Specificity}}$. G-mean measures the balance between correctly identifying inliers and outliers [11], [12]. A higher G-mean indicates that the model performs well on both classes simultaneously, which is critical in anomaly detection tasks where class imbalance is severe. Besides the average G-mean metrics, we further assessed the statistical significance of the performance differences between the PDW and other four

competing methods using the paired Wilcoxon signed-rank test. The test was performed on the G-mean values obtained from repeated 100 experiments.

TABLE 1. G-mean and Wilcoxon test results on benchmark datasets

Dataset	PDW	Default	MISE	iForest	LOF
phoneme	0.4952(1)	0.3731(4)	0.4773(2)	0.4615(3)	0.3525(5)
	p-value	< 0.001	0.0321	0.0146	< 0.001
yeast	0.6140(1)	0.6117(2)	0.5719(3)	0.4805(5)	0.5089(4)
	p-value	0.8020	0.0016	< 0.001	< 0.001
shuttle	0.9101(1)	0.8596(3)	0.9039(2)	0.8117(4)	0.7806(5)
	p-value	< 0.001	0.8880	< 0.001	< 0.001
breast	0.8824(1)	0.8764(2)	0.6650(5)	0.8617(3)	0.8330(4)
	p-value	< 0.001	< 0.001	0.0147	< 0.001
ionosphere	0.9135(1)	0.9102(2)	0.5109(5)	0.7851(4)	0.8050(3)
	p-value	0.1140	< 0.001	< 0.001	< 0.001

Table 1 reports the average G-mean values of all methods on the benchmark datasets. The ranking of each method is provided in parentheses. The corresponding p-values from the paired Wilcoxon signed-rank test are reported below the average G-mean values.

Overall, the proposed PDW consistently achieved the best or highly competitive performance across all datasets. Specifically, the proposed method ranked first on all five datasets, demonstrating its effectiveness in selecting appropriate γ values for OCSVM. Compared with the default setting, the PDW method achieved substantial improvements on most datasets. The improvements were statistically significant with p-values < 0.05 in most cases, indicating that the proposed method provided consistent performance gains. On the breast and ionosphere datasets, the MISE method performed significantly worse, while the proposed method maintains high G-mean values. This suggests that the proposed method is more robust in high-dimensional settings.

When compared with iForest and LOF, the proposed method generally outperformed them across all datasets. In particular, iForest and LOF tended to produce lower G-mean values on datasets such as yeast and ionosphere, indicating that they may not capture the boundary structure as effectively as OCSVM with the proposed γ selection strategy. In terms of statistical significance, the proposed method significantly outperformed iForest and LOF on most datasets, with p-values smaller than 0.05. Although the differences between the proposed method and the default setting were not always statistically significant (e.g., on the yeast and ionosphere datasets), the proposed method still achieved the highest average performance, demonstrating its stability and effectiveness.

5. Conclusions

In this paper, we addressed the challenging problem of kernel parameter selection in OCSVM, which is critical for achieving reliable performance in anomaly detection tasks. Existing MIES approach is often sensitive to high-dimensional settings. To overcome the limitations, we proposed a novel parameter selection approach PDW that integrates both positional and density information of training samples through a continuous weighting scheme. Furthermore, we introduced a robust objective function inspired by statistical hypothesis testing, incorporating a penalty term to mitigate overfitting. This formulation provides a more stable criterion for selecting the kernel parameter.

Experimental results on multiple benchmark datasets demonstrate that the proposed method consistently achieves superior or competitive performance compared with the default setting, the MIES method, and other unsupervised approaches such as Isolation Forest and LOF. The improvements are supported by statistical significance analysis, indicating that the proposed method is both effective and robust across different data scenarios.

Despite these promising results, several limitations remain. In the current study, the regularization parameter ν was fixed across all experiments, and only the RBF kernel was considered. Future work will explore joint tuning of ν and γ , investigate alternative kernel functions, and extend the proposed approach to other kernel-based methods and large-scale datasets.

Acknowledgements

The authors would like to thank the anonymous reviewers for their valuable comments and suggestions.

References

- [1] H. Deng and R. Xu, "Model selection for anomaly detection in wireless ad hoc networks," in *Proc. IEEE Symp. Comput. Intell. Data Mining (CIDM)*, Honolulu, HI, USA, 2007, pp. 540–546.
- [2] S. Khazai, S. Homayouni, A. Safari, and B. Mojaradi, "Anomaly detection in hyperspectral images based on an adaptive support vector method," *IEEE Geosci. Remote Sens. Lett.*, vol. 8, no. 4, pp. 646–650, Jul. 2011.
- [3] Y. Xiao, H. Wang, and W. Xu, "Parameter selection of Gaussian kernel for one-class SVM," *IEEE Trans. Cybern.*, vol. 45, no. 5, pp. 941–953, 2014.
- [4] B. Schölkopf, J. C. Platt, J. Shawe-Taylor, A. J. Smola, and R. C. Williamson, "Estimating the support of a high-dimensional distribution," *Neural Computation*, vol. 13, no. 7, pp. 1443–1471, 2001.
- [5] C.-C. Chang and C.-J. Lin, "LIBSVM: A library for support vector machines," *ACM Trans. Intell. Syst. Technol.*, vol. 2, no. 3, pp. 1–27, 2011.
- [6] Y. Li and L. Maguire, "Selecting critical patterns based on local geometrical and statistical information," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 33, no. 6, pp. 1189–1201, 2010.
- [7] F. Zhu, J. Yang, C. Gao, S. Xu, N. Ye, and T. Yin, "A weighted one-class support vector machine," *Neurocomputing*, vol. 189, pp. 1–10, 2016.
- [8] Y.-P. Zhao, G. Huang, Q.-K. Hu, and B. Li, "An improved weighted one-class support vector machine for turboshaft engine fault detection," *Eng. Appl. Artif. Intell.*, vol. 94, 103796, 2020.
- [9] F. T. Liu, K. M. Ting, and Z.-H. Zhou, "Isolation Forest," in *Proc. IEEE Int. Conf. Data Mining (ICDM)*, 2008, pp. 413–422.
- [10] M. M. Breunig, H.-P. Kriegel, R. T. Ng, and J. Sander, "LOF: Identifying density-based local outliers," in *Proc. ACM SIGMOD Int. Conf. Manag. Data*, 2000, pp. 93–104.
- [11] G. Wu and E. Y. Chang, "Class-boundary alignment for imbalanced dataset learning," in *Proc. Int. Conf. Mach. Learn. (ICML)*, Washington, DC, USA, 2003, pp. 49–56.
- [12] M. R. Wu and J. P. Ye, "A small sphere and large margin approach for novelty detection using training data with outliers," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 31, no. 11, pp. 2088–2092, Nov. 2009.