

Article

Feature Decomposition-Based Framework for Source-Free Universal Domain Adaptation in Mechanical Equipment Fault Diagnosis

Peiyi Zhou , Weige Liang *, Shiyan Sun and Qizheng Zhou

Naval University of Engineering, Wuhan 430033, China

* Correspondence: 1312021010@nue.edu.cn

Abstract

Aiming at the problems of high complexity in source domain data, inaccessibility of target domain data, and unknown fault patterns in real-world industrial scenarios for mechanical fault diagnosis, this paper proposes a Feature Decomposition-based Source-Free Universal Domain Adaptation (FD-SFUniDA) framework for mechanical equipment fault diagnosis. First, the CBAM attention module is incorporated to enhance the ResNet-50 convolutional network for extracting feature information from source domain data. During the target domain adaptation phase, singular value decomposition is applied to the weights of the pre-trained model's classification layer, orthogonally decoupling the feature space into a source-known subspace and a target-private subspace. Then, based on the magnitude of feature projections, a dynamic decision boundary is constructed and combined with an entropy threshold mechanism to accurately distinguish between known and unknown class samples. Furthermore, intra-class feature consistency is strengthened through neighborhood-expanded contrastive learning, and semantic weight calibration is employed to reconstruct the feature space, thereby suppressing the negative transfer effect. Finally, extensive experiments under multiple operating conditions on rolling bearing and reciprocating mechanism datasets demonstrate that the proposed method excels in addressing source-free fault diagnosis problems for mechanical equipment and shows promising potential for practical engineering applications in fault classification tasks.

Academic Editor: Carlos
Conceicao Antonio

Received: 2 September 2025

Revised: 11 October 2025

Accepted: 14 October 2025

Published: 20 October 2025

Citation: Zhou, P.; Liang, W.; Sun, S.; Zhou, Q. Feature Decomposition-Based Framework for Source-Free Universal Domain Adaptation in Mechanical Equipment Fault Diagnosis. *Mathematics* **2025**, *13*, 3338. <https://doi.org/10.3390/math13203338>

Copyright: © 2025 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

Keywords: machine fault diagnosis; deep transfer learning; source-free universal domain adaptation; feature decomposition

MSC: 68T07

1. Introduction

Mechanical devices are widely employed in industrial production and manufacturing fields, and their normal operation has been recognized as critical for ensuring productivity [1]. With the rapid development of modern industry, the complexity and automation level of mechanical equipment, such as rotating bearings and reciprocating mechanisms, have continuously increased. Consequently, maintenance difficulties have become more pronounced, significantly affecting the normal, safe, and stable operation of the systems. Thus, it is particularly important to achieve timely and effective fault diagnosis decisions for key components including bearings and reciprocating mechanisms [2].

In recent years, deep learning technologies have been extensively utilized in fault diagnosis tasks of mechanical systems due to their efficient pattern recognition capabilities

and strong generalization abilities [3,4]. However, the effectiveness of deep learning-based models is dependent upon two fundamental prerequisites: a substantial amount of labeled training data, and the condition that both training and test data follow identical or similar distributions. These prerequisites are often difficult to fulfill in practice, due to variations in operational conditions such as rotational speed and load, which may lead to distribution discrepancies between training and test data. Consequently, there is an urgent need to transfer diagnostic knowledge from well-labeled operational conditions (the source domain) to different unlabeled conditions (the target domain). This process is referred to as domain adaptation (DA) [5].

Domain adaptation (DA) methods are employed to address domain shift issues in fault diagnosis. According to the data source scenarios, common DA settings can be categorized into four types: closed-set domain adaptation [6], partial domain adaptation [7], open-set domain adaptation [8], and open-partial domain adaptation [9]. However, most existing domain adaptation methods are only applicable to the closed-set scenario, which significantly restricts their practicality. To address the challenge of inconsistent label spaces between the source and target domains, studies have been conducted on feature alignment strategies based on discrepancy [10], reconstruction [11,12], and adversarial learning [13–15]. Although these methods perform well in specific scenarios, their applicability remains subject to two main limitations: (1) the preconditions, such as closed-set or open-set assumptions, are often difficult to satisfy under real-world working conditions with dynamically varying label spaces; and (2) traditional methods require simultaneous access to both source and target domain data.

Source-Free Universal Domain Adaptation (SF-UniDA) has emerged as a promising solution with enhanced practicality to address the challenges outlined above. This approach extends the framework of universal domain adaptation by eliminating the dependency on source domain data [16]. Its objective is to achieve effective universal domain adaptation in the target domain using a pre-trained source model, without requiring access to the source data. Thus, the SF-UniDA method demonstrates superior capability in handling fault diagnosis problems where source data may be inaccessible due to privacy or security concerns. However, SF-UniDA still faces three core challenges in the context of mechanical fault diagnosis: (1) feature alignment bias caused by noisy pseudo-labels in the target domain; (2) decision boundary contamination caused by unknown-class samples; and (3) knowledge transfer hindered by cross-domain semantic discrepancy.

To address the aforementioned challenges, a feature decomposition-based source-free universal domain adaptation framework for mechanical equipment fault diagnosis (FD-SFUniDA) is proposed in this paper. The main contributions of this work are summarized as follows:

- (1) A sample confidence weighting mechanism based on a dynamic decision boundary is proposed. The decision boundary is constructed using the magnitude of feature projections after orthogonal decomposition to distinguish between known and unknown class samples. A weighting function is employed to suppress the influence of low-confidence pseudo-labels during training, thereby mitigating feature alignment bias.
- (2) A neighborhood-expanded contrastive learning mechanism is introduced. This mechanism enforces intra-class clustering and inter-class separation in the feature space, significantly enhancing the discriminative power of target domain features.
- (3) A transfer enhancement method based on dynamic semantic calibration is developed. By reconstructing the feature space, shared label semantics are enhanced, which effectively suppresses negative transfer effects from source-private classes and improves cross-domain generalization.

2. Related Work

2.1. Source-Free Universal Domain Adaptation

In domain adaptation problems, the labeled source domain dataset consisting of n_s is denoted as $D_s = \{x_i^s, y_i^s\}_{i=1}^{n_s}$, where x_i^s and y_i^s represent the source sample and its corresponding label, respectively. The unlabeled target domain dataset composed of n_t samples is denoted as $D_t = \{x_i^t\}_{i=1}^{n_t}$, with x_i^t being a target domain sample. The label classes in the source and target domains are defined as C_s and C_t , respectively. The common classes shared by both domains are denoted as $C = C_s \cap C_t$. The target-private classes are represented by $\bar{C}_t$, and the source-private classes by $\bar{C}_s$, where $\bar{C}_t = C_t \setminus C_s$ and $\bar{C}_s = C_s \setminus C_t$. In the proposed source-free universal domain adaptation framework, D_s and D_t are used for model training. The labeled D_s is employed only for source model pre-training. The unlabeled D_t is then used to adapt the pre-trained source model, achieving effective domain alignment without any source instances and enabling superior classification performance. In the absence of prior knowledge regarding the label-space relationship between domains, four types of domain adaptation may occur:

- (1) Closed-Set Domain Adaptation (CLDA): $C_t = C_s$.
- (2) Open-Set Domain Adaptation (OSDA): $C_s \subset C_t$.
- (3) Partial-Set Domain Adaptation (PDA): $C_s \supset C_t$.
- (4) Open-Partial-set Domain Adaptation (OPDA): $C_s \cap C_t \neq \emptyset, C_s \not\subseteq C_t, C_s \not\supset C_t$.

Therefore, SF-UniDA as shown in Figure 1, the fault diagnosis problem under the SF-UniDA setting is defined as follows:

- (1) The proposed model is designed to handle all four possible category shift scenarios mentioned above;
- (2) The proposed model is required to accurately identify and classify unknown samples within the target domain, where $|C_s \cap C_t| + 1$.

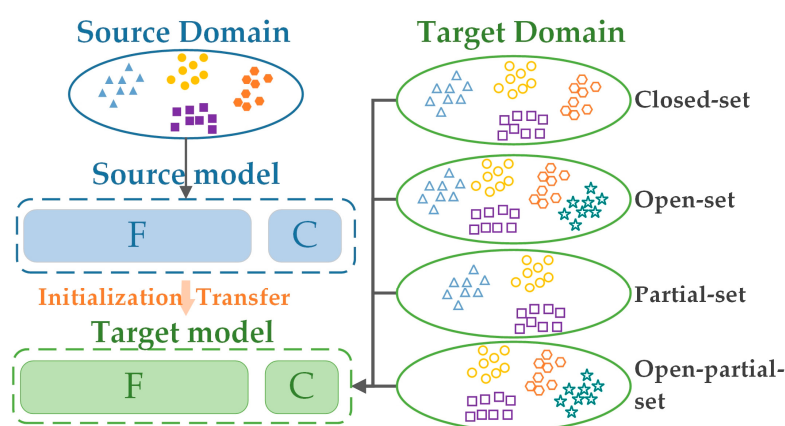


Figure 1. Source-Free Universal Domain Adaptation.

2.2. Feature Decomposition

Feature decomposition is regarded as a core technique in machine learning for analyzing high-dimensional data. It disentangles the original feature space into low-dimensional or interpretable components through mathematical transformation. In research on domain adaptation and domain generalization, this approach has been extensively applied to feature alignment tasks [17]. However, in UniDA, the content-style dichotomy cannot be directly mapped to the separation logic between shared and private classes. Therefore, a vertical feature decomposition strategy is adopted in this study [18]. The feature space is decomposed into two statistically independent feature components. Correspondence is established between the orthogonal components and the features of common versus

private data. As a result, category-discriminative information is effectively separated from domain-specific noise.

Given the weight matrix $W_s \in \mathbb{R}^{C \times D}$ of the pre-trained source model's classification layer, where D denotes the feature dimension and C denotes the number of source classes, Singular Value Decomposition (SVD) is applied.

$$W_s = U\Sigma V^T \quad (1)$$

where the columns of $V \in \mathbb{R}^{D \times D}$ form an orthonormal basis. Two complementary feature subspaces are constructed. The source-known subspace F_{knw} , spanned by the first C right singular vectors, is used to capture discriminative features related to the fault categories in the source domain.

$$F_{knw} = \text{span}\{v_n\}_{n=1}^C \quad (2)$$

where its orthogonal complement F_{unk} , corresponding to the remaining $D-C$ vectors, is employed to extract target-private faults or domain-specific information.

$$F_{unk} = \text{span}\{v_n\}_{n=C+1}^D \quad (3)$$

A target domain feature $z_i^t = g_t(x_t) / \|g_t(x_t)\|_2$ is then decomposed into two mutually orthogonal components.

$$z_i^t = \underbrace{\sum_{n=1}^C l_n v_n}_{z_{i,knw}^t} + \underbrace{\sum_{n=C+1}^D l_n v_n}_{z_{i,unk}^t} \quad (4)$$

This decomposition is strictly constrained to be orthogonal. The projection norm $\|z_{i,unk}^t\|_2$ of unknown fault samples in the F_{unk} subspace is significantly higher than that of known samples. Thus, a physical basis is provided for distinguishing between shared and private classes.

2.3. Contrastive Learning

Contrastive learning is widely recognized as a self-supervised learning paradigm that aims to acquire discriminative feature representations by capturing the intrinsic structure of data through the construction of positive and negative sample pairs [19]. Its core principle involves maximizing agreement within positive pairs while minimizing similarity with negative pairs, typically optimized using losses such as InfoNCE. This approach enhances intra-class compactness and inter-class separation in the feature space, making it valuable for domain adaptation tasks.

In traditional domain adaptation, methods typically require simultaneous access to both source and target data to form cross-domain pairs. However, in source-free domain adaptation, where source data is unavailable, the focus lies in constructing positive-negative sample pairs solely using pre-trained source models and target data. For instance, leveraging sample proximity relationships in manifold spaces to form contrastive pairs enhances model generalization performance [20]. Additionally, approaches incorporate historical pseudo-label information from sequential queues to exclude negative pairs formed by samples within the same category, focusing on maintaining contrastive pair purity under noisy labels [21]. These methods iteratively refine pseudo-labels to improve cross-domain generalization performance. In OSDA, contrastive learning must distinguish between known and unknown classes. Recent studies incorporate uncertainty-aware mechanisms, where contrastive losses are applied only to high-confidence samples to avoid contaminating the feature space with unknown-class data. For example, some frameworks use thresholds to

filter out uncertain samples before contrastive learning, ensuring that positive pairs are formed only within known classes [22]. Others employ pseudo-label assisted contrastive Learning to minimize inter-domain differences and accurately identify out-of-distribution failures [23]. A neighborhood expansion mechanism has been proposed to address these issues. It enables the dynamic construction of robust positive and negative pairs based on multi-order sample neighborhoods. Without relying on source data, this method reduces sensitivity to noise and improves structural consistency.

2.4. Semantic Enhancement

Semantic enhancement is regarded as a key technique for improving both the discriminability and transferability of cross-domain feature representations [24]. Its core principle lies in mining and enhancing the implicit category-related semantic information within deep features. In traditional domain adaptation frameworks, this method typically relies on explicit access to source domain data. Feature alignment is achieved through the construction of class-conditional distributions. Specifically, semantic enhancement utilizes category prototypes learned by a pre-trained source model. These prototypes, which are centroid representations of individual categories in the feature space, serve as semantic anchors. By maximizing the similarity between corresponding prototypes of shared categories from the source and target domains, samples of the same class are compelled to cluster closely in the feature space. As a result, cross-domain distribution discrepancies are reduced, and knowledge transfer is facilitated. However, in source-free scenarios where source data is inaccessible, such methods face fundamental challenges. The absence of source samples prevents the direct computation of class-conditional statistics. Consequently, conventional prototype-based semantic enhancement mechanisms become ineffective.

Given these constraints, the SF-UniDA leverages the parameters of the pre-trained source model as implicit semantic carriers. These weight vectors essentially serve as parameterized representations of source domain category prototypes. Previous studies have enhanced target domain adaptability through semantic probability contrastive regularization and hard-sample confusion regularization [25], or achieved semantic invariance in feature space using generative models and energy regularization to ensure cross-domain semantic consistency and improve out-of-distribution detection capabilities [26]. The core innovation of this paper lies in designing a target-domain-driven confidence evaluation mechanism. By analyzing the prediction confidence of target samples, it assesses the likelihood of each source category serving as a shared class in the target domain, thereby generating category weights. Semantic enhancement is achieved through weight-guided semantic shift reconstruction, driving target sample features toward high-weight source category prototypes while suppressing the influence of low-weight prototypes. This process enables adaptive calibration of feature distributions under passive conditions, effectively mitigating negative transfer caused by source-private class knowledge and enhancing the model's rejection capability for unseen classes.

3. Fault Diagnosis Model

3.1. Overall Structure of the Model

The time-frequency characteristics of vibration signals produced by mechanical equipment operating under varying conditions, such as rotational speed, load, and environmental factors, are significantly different. Consequently, the generalized performance of models trained on a single operational condition is severely compromised. Aiming at the limitations of traditional domain adaptation methods, which require simultaneous access to both source and target domain data, struggle to adapt to dynamically changing proportions of unknown classes under real-world operating conditions, and suffer from biased feature

alignment in the target domain, this paper proposes a source-free universal domain adaptation fault diagnosis model for industrial scenarios, specifically designed for vibration signals during mechanical operation processes. The overall architecture of the proposed model is illustrated in Figure 2.

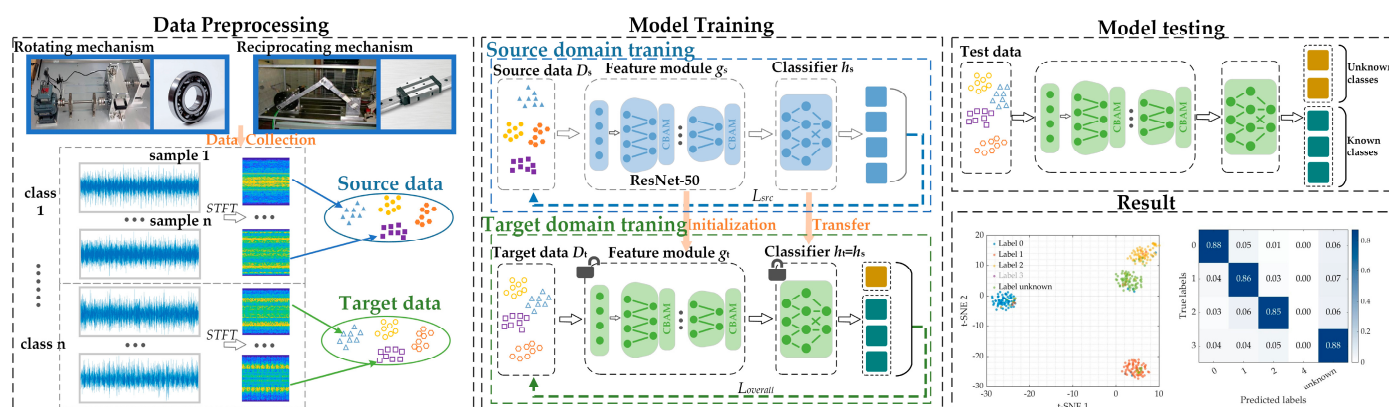


Figure 2. Working principle of the proposed method FD-SFUniDA.

The model comprises a ResNet50-CBAM feature extractor and a fully connected classifier. Training takes place in two stages: a pre-training phase in the source domain and an adaptation phase in the target domain. During the source training phase, a highly generalized base model is constructed using frequency-domain feature extraction and label smoothing strategies. In the target adaptation phase, cross-domain diagnosis under source-free constraints is achieved via neighborhood contrastive learning, entropy-based discrimination and semantic enhancement.

3.2. Model Training in the Source Domain

Prior to feature extraction, the Short-Time Fourier Transform (STFT) is applied to each input sample as shown in Equation (1), converting the raw time-domain vibration data into the frequency domain. This operation is performed on the basis that frequency-domain information generally provides greater discriminative power with regard to machine health conditions, leading to improved fault diagnosis performance.

$$x_i = X_i(t, f) = \int_{-\infty}^{\infty} \tilde{x}_i(\tau) w(\tau - t) e^{-j2\pi f\tau} d\tau \quad (5)$$

where $\tilde{x}_i(\tau)$ denotes the original sample, and $w(\tau - t)$ denotes the window function x_i denotes the sample after STFT.

Source domain data processed by STFT is input into the enhanced ResNet50-CBAM network to suppress interference noise from other components. As illustrated in Figure 3, a Convolutional Block Attention Module (CBAM) [27] is embedded after each residual block in the ResNet50 architecture. Through the sequential application of channel and spatial attention, dynamic weight assignment is achieved across feature maps. This mechanism enhances the discriminability of learned representations by directing the model's focus to fault-sensitive frequency bands within the vibration signals.

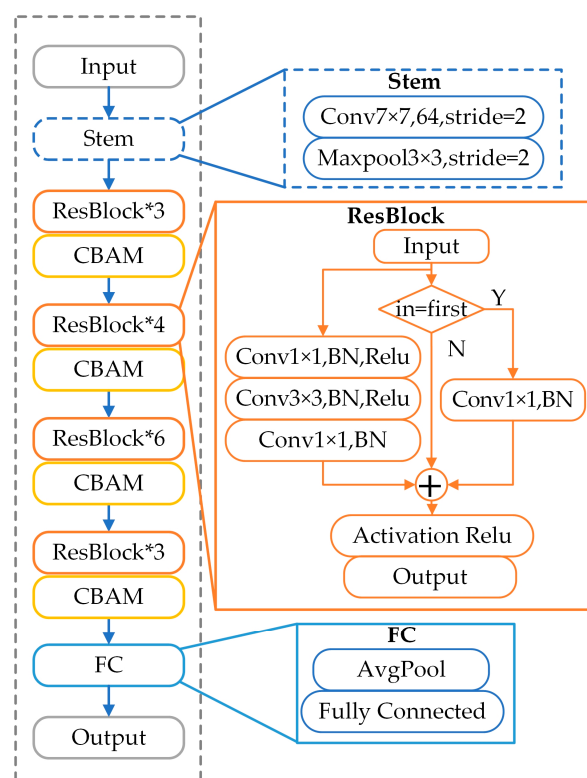


Figure 3. Resnet50-CBAM structure diagram.

Label smoothing techniques are employed during training of source domain samples to enhance the source model's generalization capabilities, mitigate overconfidence during training, and ensure that sample features remain tightly clustered and sufficiently separated. Cross-entropy loss with label smoothing regularization is then used to learn the diagnostic model. The cross-entropy loss in the source domain can be expressed as shown in Equation (6).

$$\mathcal{L}_{src} = -\frac{1}{n_s} \sum_{n=1}^{n_s} \sum_{k=1}^C \hat{y}_i \log \delta_k(f_s(x_i^s)) \quad (6)$$

where $\delta_k(f_s(x_i^s))$ denotes the softmax probability that the source sample x_i^s belongs to the k -th category. The variable $\hat{y}_i$ denotes the smoothed one-hot version of the ground-truth label y_i^s , which is defined as $\hat{y}_i = (1 - \varepsilon) \times y_i^s + \varepsilon/C$. Where y_i^s is the original one-hot encoded label, ε is a smoothing parameter with a default value of 0.1, and C is the total number of classes.

During training, the optimal parameters of the source model are sought through minimization of the loss function $\mathcal{L}_{src}$. This process is formulated as shown in Equation (7).

$$(\hat{\theta}_g, \hat{\theta}_h) = \underset{\theta_g, \theta_h}{\operatorname{argmin}} \mathcal{L}_{src} \quad (7)$$

3.3. Model Training in the Target Domain

3.3.1. Dynamic Decision Boundaries and Pseudo-Label Learning

The training strategy begins with the generation of robust pseudo-labels to guide the learning process when ground-truth annotations are unavailable. Due to the limited robustness of conventional fixed-threshold methods under varying operational conditions, this work proposes an instance-level dynamic decision mechanism based on multi-criterion fusion. First, class prototypes c_c^t are constructed via clustering of target features, while source anchors c_c^s are derived from the classification weight matrix $W_s \in \mathbb{R}^{C \times D}$ of the

pre-trained source model. On this basis, a dynamic commonness score function is defined as shown in Equation (8).

$$\epsilon_{i,c} = \sqrt{1 - \text{sim}(z_i^t, c_c^t) \cdot \exp(-\|1 - \text{sim}(z_i^s, c_c^s)\|^2)} \quad (8)$$

where $\text{sim}(\cdot, \cdot)$ denotes cosine similarity. This function incorporates both target-domain similarity and source-domain calibration. To overcome the limitations of fixed thresholds, a dynamic decision boundary, as shown in Equation (9) related to both the instance and the category is learned based on the distribution of target domain features.

$$\rho_{i,c} = \mu_c + \epsilon_{i,c} \cdot (\mu_{pri} - \mu_c) \quad (9)$$

where μ_c and μ_{pri} are obtained by fitting a two-component Gaussian Mixture Model (GMM) to the magnitudes of all target feature vectors. Based on the dynamic decision boundary described above, the pseudo-label assignment rule is defined as shown in Equation (10).

$$\hat{y}_i^t = \begin{cases} \text{unknown} & \|z_{i,unk}^t\|_2 \geq \rho_{i,c} \\ \text{argmax}_c(\delta_c(f_t(x_i^t))), & \text{otherwise} \end{cases} \quad (10)$$

A weighted pseudo-label learning loss function was designed to facilitate model training using generated pseudo-labels while mitigating the negative impact of label noise, as shown in Equation (11).

$$\mathcal{L}_{pl} = -\frac{1}{n_t} \sum_{i=1}^{n_t} \tau_i^t \sum_{c=1}^C \tilde{y}_{i,c}^t \log \delta_c(f_t(x_i^t)) \quad (11)$$

where $\delta_c(f_t(x_i^t))$ denotes the softmax probability that target sample x_i^t belongs to the k -th class, $\tilde{y}_{i,c}^t$ is the one-hot encoded pseudo-label, and τ_i^t is a weighting factor based on t -distribution confidence to effectively suppress the influence of low-quality pseudo-labels.

$$\tau_i^t \propto 1 - \left(1 + \frac{(\rho_{i,\kappa} - \|z_{i,unk}^t\|_2)^2}{\alpha} \right)^{-\frac{\alpha+1}{2}} \quad (12)$$

where $\kappa = \text{argmax}(\epsilon_{i,c})$ represents the degrees of freedom of the t -distribution.

3.3.2. Neighbors Extension Contrastive Learning

After initial pseudo-labels are obtained, a neighborhood-expanded contrastive learning mechanism is introduced to further enhance the intra-class consistency in the feature space and strengthen the model's perception of the data structure. The core idea of this mechanism is to construct a neighborhood set $N_k(i)$ of the k nearest neighbors for each sample x_i in the feature space. Positive pairs are formed from neighboring samples $P_i = N_k(i) \cap \{j | \tilde{y}_j = \tilde{y}_i\}$ that share the same pseudo-label, while negative pairs are extended to higher-order neighborhoods $N_i = \cup_{m=1}^M N_k(i_m) \cap \{j | \tilde{y}_j \neq \tilde{y}_i\}$ to more comprehensively reflect dissimilarity relationships between samples. The contrastive loss based on the expanded positive and negative pairs is computed as shown in Equation (13).

$$\mathcal{L}_{nc} = -\frac{1}{n_t} \sum_{i=1}^{n_t} \log \frac{\sum_{j \in P_i} \exp\left(\frac{\text{sim}(z_i^t, z_j^t)}{T}\right)}{\sum_{j \in P_i \cup N_i} \exp\left(\frac{\text{sim}(z_i^t, z_j^t)}{T}\right)} \quad (13)$$

where T is a temperature hyper-parameter. This loss encourages the model to cluster features of the same class closer together and to separate features of different classes. As a result, a more compact and discriminative feature representation is formed in the target domain.

3.3.3. Transfer Enhancement Method Based on Dynamic Semantic Calibration

A dynamic semantic calibration mechanism is proposed to mitigate the negative transfer effect where source domain private classes (SPCs) may interfere with the correct feature alignment of target domain shared classes in scenarios lacking passive data. This mechanism's key innovation lies in its exclusive use of the pre-trained source model's parameters and target domain features to intelligently reconstruct the target feature distribution. This enhances the feature consistency of shared classes while suppressing interference from private classes.

First, a category weight evaluation system based on consensus within the target domain is established. Each column $w_k \in \mathbb{R}^d$ of the classifier weight matrix $W_s = [w_1, w_2, \dots, w_K] \in \mathbb{R}^{d \times K}$ from the pre-trained source model is regarded as the feature space prototype for the k -th class in the source domain. To evaluate the likelihood of each source class being a shared class in the target domain, a dual-filtered weighting mechanism is designed. This mechanism focuses on target domain samples in the current training batch $\mathcal{B}$ that have been assigned pseudo-labels, thereby excluding unknown-class samples labeled as "unknown". The set of batch sample indices with pseudo-label k is defined as shown in Equation (14).

$$I_k^{\mathcal{B}} = \{i \in \mathcal{B} | \hat{y}_i = k\} \quad (14)$$

The initial class weight w_k is computed by considering both the number of target samples assigned to the class and their average confidence.

$$w_k = \frac{1}{|I_k^{\mathcal{B}}|} \sum_{i \in I_k^{\mathcal{B}}} \left(1 - \frac{H(i)}{\log C} \right) \quad (15)$$

where $H(i) = -\sum_{c=1}^C p_c(x_i^t) \log p_c(x_i^t)$ denotes the predictive entropy of sample x_i^t and $\log(C)$ represents the entropy of a uniform distribution over C classes for normalization. The weighting term is designed to assign higher contributions to low-entropy samples. The underlying rationale is that a class with more pseudo-labeled samples and higher confidence in the target domain is more likely to be a shared class. Conversely, a class with fewer samples or generally lower confidence is more likely to be a source-private class or affected by unknown classes. To eliminate scale variations and enhance numerical stability, min-max normalization is applied to the initial weights as shown in Equation (16).

$$\tilde{w}_k = \frac{w_k - \min_{k'}(w_{k'} + \varepsilon)}{\max_{k'} w_{k'} - \min_{k'}(w_{k'} + \varepsilon)} \quad (16)$$

where $\varepsilon > 0$ is a small positive value included to prevent division by zero. The normalized weight $\tilde{w}_k \in [0, 1]$ clearly quantifies the relative probability of each source class being identified as a shared class within the current target batch.

Subsequently, semantic-shift reconstruction based on the normalized weights $\tilde{w}_k$ is performed. The features of target domain samples are dynamically reconstructed using the evaluated weights. The primary objective of this reconstruction is to adaptively calibrate the feature distribution in the absence of source data by pulling sample features z_i^t closer to

the prototypes of classes with high weights, while pushing them away from prototypes of classes with low weights.

$$z_i^{aug} = z_i^t + \alpha \sum_{k=1}^C \tilde{w}_k (w_k - z_i^t), \alpha = 0.5 \quad (17)$$

where $(w_k - z_i^t)$ denotes the vector pointing from the sample feature z_i^t to the prototype w_k of the k -th source class. The coefficient $\tilde{w}_k$ dynamically scales the contribution of this vector. For classes with high $\tilde{w}_k$, the term significantly pulls z_i^t toward the corresponding prototype. For classes with low $\tilde{w}_k$, the contribution is negligible. The factor α controls the overall strength of the reconstruction. This operation essentially performs a dynamically weighted blending between the sample feature z_i^t and each source class prototype w_k in the feature space, guided by the class weight $\tilde{w}_k$. This process encourages the model to cluster features of potential shared classes in the target domain toward their corresponding source prototypes, while allowing features of private or unknown classes to remain unchanged or be repelled. As a result, spatial consistency of shared-class features is enhanced, and the negative impact of private classes is effectively suppressed.

Finally, a transfer enhancement loss function is constructed that incorporates both reconstruction reliability and information richness. The reconstructed feature z_i^{aug} is fed into the classifier h_s to obtain its prediction distribution $p_i^{aug} = h_s(z_i^{aug})v$. The loss function consists of two key components:

Reconstructed Feature Classification Loss $\mathcal{L}_{rf}^{aug}$: This is a standard cross-entropy loss that supervises the classification of the reconstructed features using their currently assigned pseudo-labels $\hat{y}_i^t$. To mitigate the inherent noise in these pseudo-labels, this loss is computed only for samples whose pseudo-label is not “unknown”:

$$\mathcal{L}_{rf}^{aug} = -\frac{1}{|B|} \sum_{i \in B} \log(p_i^{aug} \hat{y}_i^t) \quad (18)$$

The objective is to maintain or enhance the discriminative capability of the reconstructed features on shared class tasks, enabling convergence to more robust solutions even when confronted with noisy pseudo labels.

Mutual Information Regularization Loss $\mathcal{L}_{mi}$: To enhance the semantic consistency between the feature representation z_i^{aug} and its predictive distribution p_i^{aug} during optimization, and to increase the confidence of predictions, we introduce an approximation of the mutual information loss based on the Jensen–Shannon (JS) divergence:

$$\mathcal{L}_{mi} = -\frac{1}{|B|} \sum_{i \in B} [\text{JS}(p_i^t || \text{Uniform}(K))] \quad (19)$$

where $\text{Uniform}(K)$ denotes a K -dimensional uniform distribution. Maximizing the mutual information between z_i^t and p_i^t is equivalent to minimizing their $\text{JS}(p_i^t || \text{Uniform}(K))$ divergence, which encourages the predictive distribution p_i^{aug} to diverge from the uniform distribution toward a more peaked form, thereby enhancing the model’s confidence in its predictions.

The final transfer enhancement loss $\mathcal{L}_{te}$ is formulated as a weighted combination of the two aforementioned losses:

$$\mathcal{L}_{te} = \mathcal{L}_{rf}^{aug} + \gamma \mathcal{L}_{mi} \quad (20)$$

where γ denotes the regularization strength coefficient, with a default value of $\gamma = 0.1$.

This study presents a progressive optimization strategy designed to address the core challenges in target domain training, including the inaccessibility of source data, the complete absence of labels in the target domain and the potential presence of unknown fault classes. The primary objective is to gradually extract diagnostic information from the target domain while minimizing interference through self-supervised learning and knowledge transfer mechanisms.

First, pseudo-labels robust to label noise are generated based on the principle of feature space disentanglement. Then, neighborhood-expanded contrastive learning is applied to enhance intra-class feature consistency. Finally, dynamic calibration of cross-domain feature distribution is performed to suppress negative transfer effects caused by interference from source-private classes. The total loss function for training the target domain model is defined as shown in Equation (21). The specific training process is shown in Algorithm 1.

$$\mathcal{L}_{overall} = \mathcal{L}_{pl} + \lambda \mathcal{L}_{nc} + \mathcal{L}_{te} \quad (21)$$

where $\mathcal{L}_{pl}$ denotes the pseudo-label learning loss, $\mathcal{L}_{nc}$ denotes the neighborhood contrastive loss, $\mathcal{L}_{te}$ denotes the transfer enhancement loss, and λ denotes the weighting coefficient for the contrastive loss. The hyperparameter λ is set to 0.2. A sensitivity analysis of the hyperparameter λ is presented in Section 4.5.3.

Algorithm 1: The proposed FD-SFUniDA method

Input: Pre-trained source model g_s, h_s ; Unlabelled target domain data $D_t = \{x_i^t\}_{i=1}^{N_t}$;

Max-epoch I_{max} ; Batch size b_n ; Trade-off hyperparameters λ

Output: feature extractor g_t, h_t after updating parameters by target data

Initialize $g_t \leftarrow g_s, h_t \leftarrow h_s$

Perform SVD of classifier weights

For epoch = 1 to I_{max}

For iter = 1 to N **do**

 Generate pseudo-labels $\hat{y}_i^t$

 Compute pseudo-label learning loss $\mathcal{L}_{pl}$

 Construct positive/negative pairs

 Compute neighborhood extension contrast loss $\mathcal{L}_{nc}$

 Reconstruct features using semantic calibration

 Compute transfer enhancement loss $\mathcal{L}_{te}$

 Combine losses: $\mathcal{L}$

 Update g_t and h_t parameters

End for

End for

4. Experimental Validation

4.1. Dataset Descriptions

1. Rolling Bearing Dataset: Bearing damage data originates from a mechanical transmission component failure test bench, as illustrated in Figure 4. The test bench comprises a motor, bearing housing, gearbox, and load. To collect data across varying bearing health states, acceleration sensors were mounted on the bearing housing. Experimental data was sampled at 200 kHz. Motor speed was set to three values: 1000, 1250, and 1500 r/min to obtain bearing data under varying operating conditions. Five bearing conditions were identified: Normal Condition (NC), Inner-race Fault (IF), Outer-race Fault (OF), Ball Fault (BF), and Mixed Fault (MF). These fault states are illustrated in Figure 5.

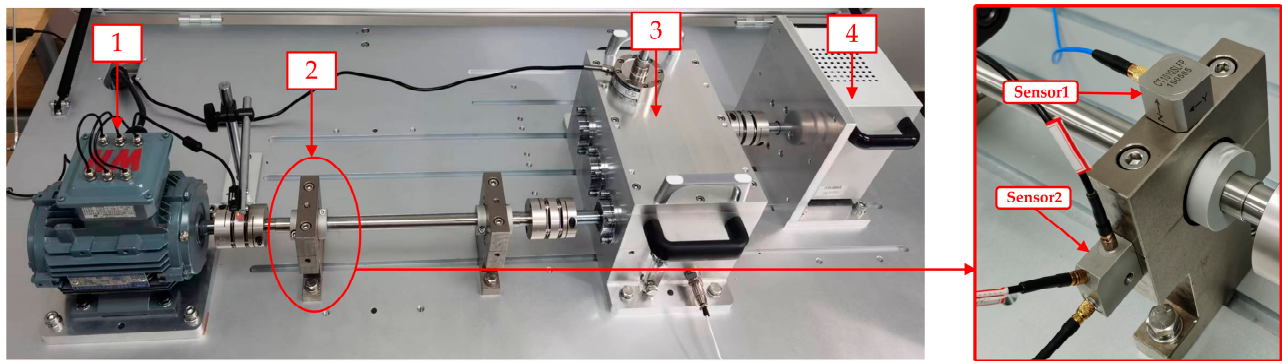


Figure 4. Mechanical transmission components failure experiment bench.

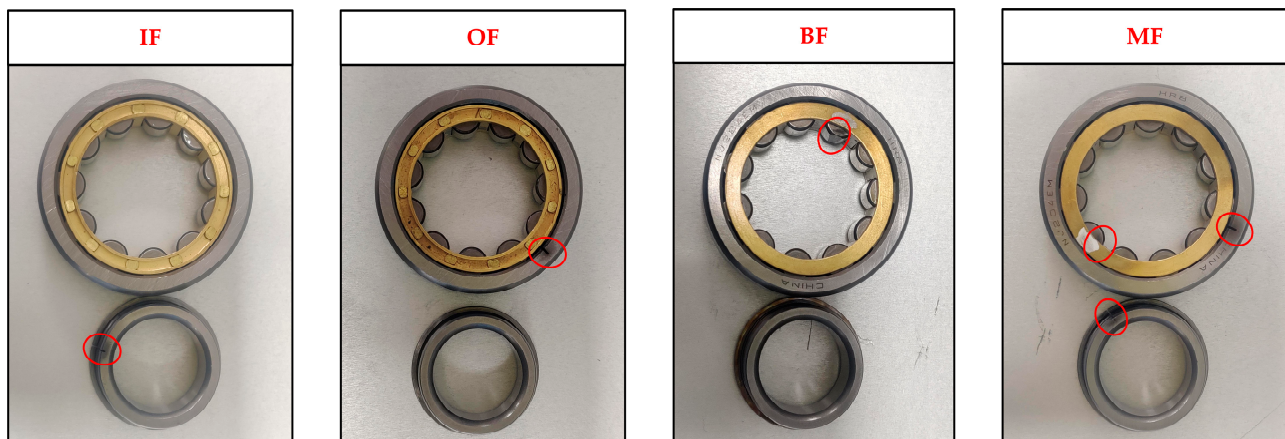


Figure 5. Four states of bearing damage.

2. Reciprocating Mechanism Dataset: Data originates from the supply and delivery mechanism test rigs under reciprocating impact vibration, as shown in Figure 6. The test rig comprises an electrical control system, hydraulic system, typical motion mechanisms, and the frame structure. Typical motion mechanisms include the tail assembly, swing mechanism, reciprocating mechanism, and lifting mechanism. The swing mechanism is the primary focus of the test rig. Within this mechanism, the sliding plate achieves upward and downward swinging motions by engaging and disengaging with the lift-up stop block. Consequently, significant impact wear occurs on the engagement surfaces of the sliding plate–lift-up stop block mechanism. As wear progressively increases, the engagement state between the sliding plate and lift-up stop block changes, severely impairing the normal operation of the swing mechanism. Another failure mode stems from cracks in the rollers. As critical connecting components within the slide groove, the rollers directly impact the driving force of the reciprocating motion.

In the experiment, different measurement points were set for the locations of the two faults, with the fixed positions of the vibration acceleration sensors shown in Figure 7. A 32-channel LMS data acquisition system was used to collect vibration acceleration signals from each measurement point. This acquisition system was equipped with ICP-type vibration acceleration sensors, with the sampling rate set to 10 kHz.

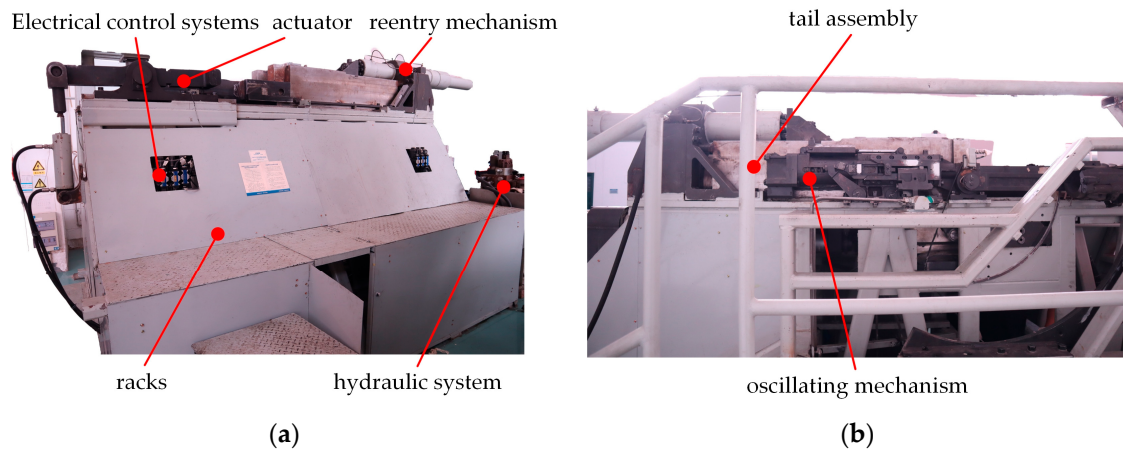


Figure 6. Supply and delivery mechanism test rigs. (a) Right side of the platform; (b) Left side of the platform.

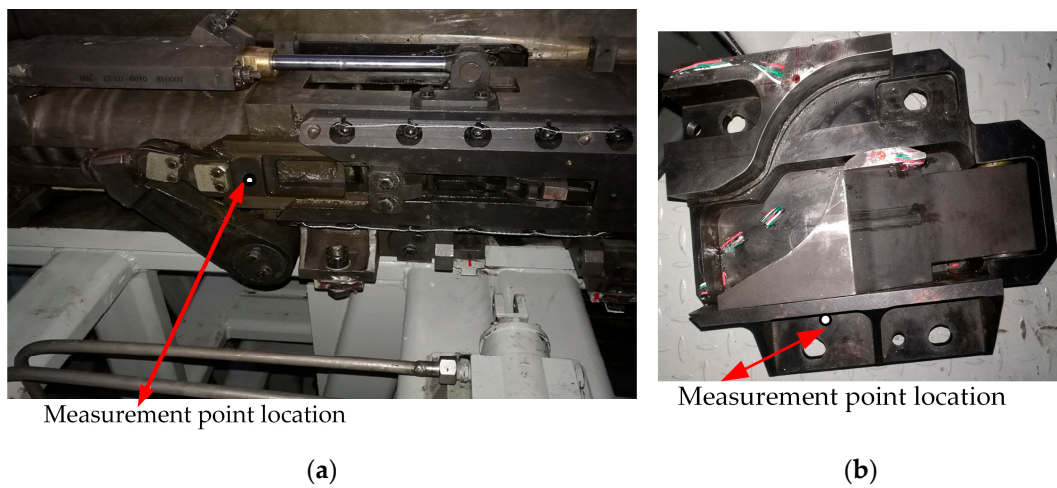


Figure 7. Acceleration sensor point layout. (a) Sliding plate sensor point layout; (b) Roller groove sensor point layout.

Each fault contains measurement data for three distinct operating angles: 0° , 30° , and 70° . These include Normal Condition (NC), Cracked Roller Failure (CRF), and Sliding Plate Wear Fault (SPWF), as shown in Figure 8.

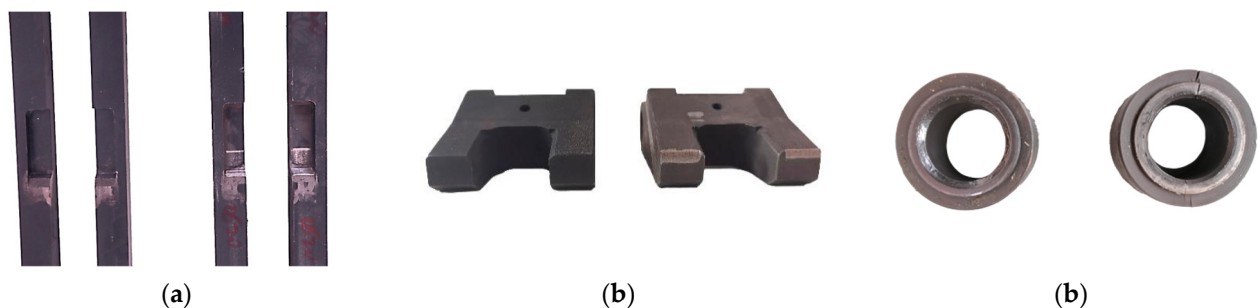


Figure 8. Comparison chart of the state of life parts. Left side indicates healthy status, right side indicates fault status. Left side indicates healthy status, right side indicates fault status. (a) Sliding plate comparison chart; (b) Blocking iron comparison chart; (c) Roller comparison chart.

4.2. Compared Approaches

The proposed method was compared with the following five domain adaptation approaches to validate its superiority:

- (1) Domain-Adversarial Training of Neural Networks (DANN) [28]: As a typical closed-domain adaptation method, DANN performs adversarial training via a domain discriminator and then learns cross-domain invariant features.
- (2) Domain Conditioned Adaptation Network (DCAN) [29]: As a closed-domain adaptation approach, DCAN utilizes a domain-conditional channel attention mechanism to activate different convolutional channels, enabling fault diagnosis under significant data distribution shifts.
- (3) One-vs-All Network for Universal Domain Adaptation (OVANet) [30]: OVANet identifies unknown and shared classes in the target domain by learning the minimum inter-class distance within the source domain.
- (4) Pseudo-Marginal-Based Universal Domain Adaptation (IUAN) [31]: IUAN identifies a common label space through pseudo-margins for sample classification in UDA scenarios.
- (5) Source-Free Universal Domain Adaptation by Hybrid Clustering Strategy (SFUDA-HCS) [21]: SFUDA-HCS generates pseudo labels through global clustering, optimizes label assignment via local consensus clustering, and identifies unknown categories using contrastive learning. This enables efficient cross-domain fault diagnosis without requiring source domain data.

4.3. Experimental Setup

In this experimental setup, the vibration signals acquired from both test benches under each operating condition serve as the data source. Eighty samples were selected for each category, as shown in Table 1. The rolling bearing dataset consists of two-dimensional images derived from short-time Fourier transform (STFT) of vibration data. Each original sample has a length of 4096, with an STFT Hamming window length of 256 and a shift step of 64. For the reciprocating mechanism dataset, signals collected from the test bench must be segmented. As shown in Figure 9, the GSA-IFCM [32] method is employed to extract a single cycle of vibration signal, which undergoes short-time Fourier transform. The resulting two-dimensional images serve as dataset samples, collectively forming the reciprocating mechanism dataset.

Table 1. Description of experimental data.

Dataset	Label	Fault Model
Rolling Bearing Dataset	0	NC
	1	IF
	2	OF
	3	BF
	4	MF
Reciprocating Mechanism Dataset	0	NC I
	1	NC II
	2	CRF
	3	SPWF

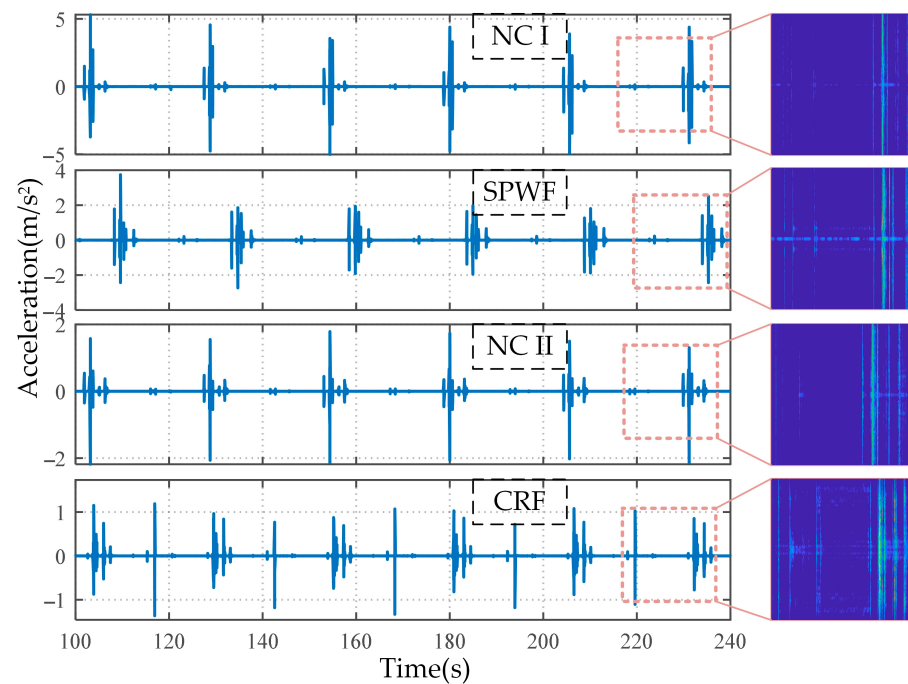


Figure 9. Schematic diagram of acquisition signal division.

Multiple cross-domain transfer tasks were established to validate the effectiveness of the developed method (as shown in Tables 2 and 3). By designing different classifications of shared and private classes between source and target domains, a series of cross-domain transfer tasks were constructed to cover the standard settings of CLDA, PDA, OSDA, and OPDA. To approximate real-world conditions, the rolling bearing dataset employed different rotational speeds as independent transfer domains, while the reciprocating mechanism dataset utilized different angles as independent transfer domains.

Table 2. Adaptive tasks in the field of rolling bearing dataset.

Tasks	Source Domain (r/min)	Target Domain (r/min)	Source Class	Target Class	Problem
T1	1000	1250	0, 1, 2, 3, 4	0, 1, 2, 3, 4	CLDA
T2	1000	1500	0, 1, 2, 3, 4	0, 1, 2, 3, 4	CLDA
T3	1250	1000	0, 1, 2, 3, 4	0, 1, 2, 4	PDA
T4	1250	1500	0, 1, 2, 3, 4	0, 1, 4	PDA
T5	1500	1000	0, 1, 2, 4	0, 1, 2, 3, 4	OSDA
T6	1500	1250	0, 1, 4	0, 1, 2, 3, 4	OSDA
T7	1000	1250	0, 1, 2, 4	0, 1, 2, 3	OPDA
T8	1500	1250	0, 1, 3	0, 2, 4	OPDA

Table 3. Adaptive tasks in the field of reciprocating mechanism dataset.

Tasks	Source Domain (°)	Target Domain (°)	Source Class	Target Class	Problem
C1	0	30	0, 1, 2, 3	0, 1, 2, 3	CLDA
C2	0	70	0, 1, 2, 3	1, 2, 3	PDA
C3	30	0	1, 2, 3	0, 1, 2, 3	OSDA
C4	30	70	0, 1, 2	0, 1, 2, 3	OSDA
C5	70	0	0, 1, 2	1, 2, 3	OPDA
C6	70	30	1, 2, 3	0, 1, 2	OPDA

All experiments were conducted on a computer equipped with an NVIDIA GeForce GTX 1060 GPU and 4 GB of memory. The PyTorch2.0.1 deep learning framework was

used for model construction, with the LARS optimization method employed for training. The learning rate was set to 0.01, with a batch size of 16, weight parameter decay of 0.1, 100 iterations, and $k = 5$ for k -nearest neighbors. To minimize randomness, the diagnostic results represent the average of ten repeated trials per domain adaptation task.

4.4. Results Display and Analysis

This section analyzes the diagnostic results of multiple methods applied to rolling bearing datasets and reciprocating mechanism datasets. To comprehensively evaluate the performance of the proposed method and comparative methods, two evaluation metrics are adopted with reference to research in the same field [20]:

- (1) Test sample accuracy rate Acc :

$$Acc = (M_K + M_U) / (N_K + N_U) \quad (22)$$

where M_K and M_U denote the number of correctly classified test samples in known and unknown categories, respectively, and N_K and N_U denote the total number of test samples in known and unknown categories, respectively.

- (2) Harmonic mean H-score:

$$H\text{-score} = ((2OS^*UK) / (OS^* + UK)) \quad (23)$$

Calculates the harmonic mean of accuracy for known categories and position categories to evaluate the model's ability to handle balance between known and position categories. Here, $OS^* = (M_K / N_K)$ denotes the ratio of correctly classified known category samples to the total number of known category samples, and $UK = (M_U / N_U)$ denotes the ratio of correctly classified unknown category samples to the total number of unknown category samples.

These evaluation metrics provide a comprehensive assessment of different methods' capabilities in identifying known fault categories and discovering unknown fault categories during fault diagnosis tasks.

Figures 10 and 11 illustrate the diagnostic accuracy and H-score for various domain adaptation tasks, with detailed results presented in Tables 4 and 5. The comparison reveals that the proposed method demonstrates robust cross-domain diagnostic performance, achieving accuracy rates above 84% for most tasks. Furthermore, the developed approach maintains relatively low standard deviations across different tasks, indicating its excellent convergence properties.

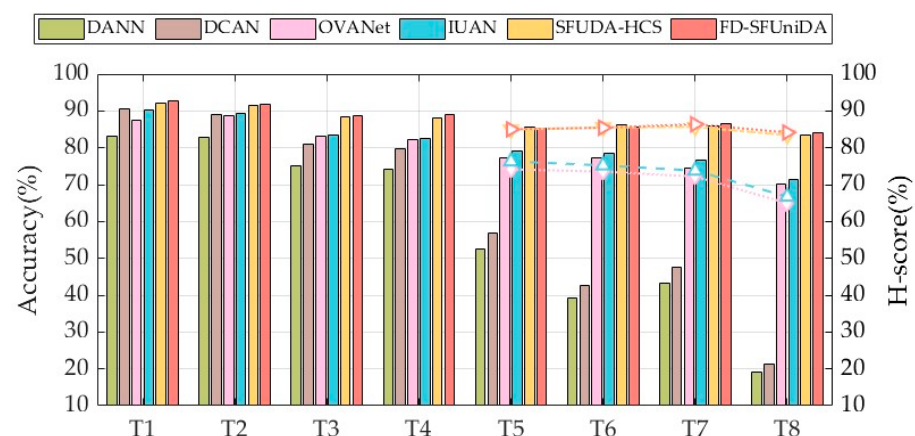


Figure 10. Comparison of fault diagnosis accuracy and H-score for different methods of rolling bearing dataset. The broken lines represents the H-score.

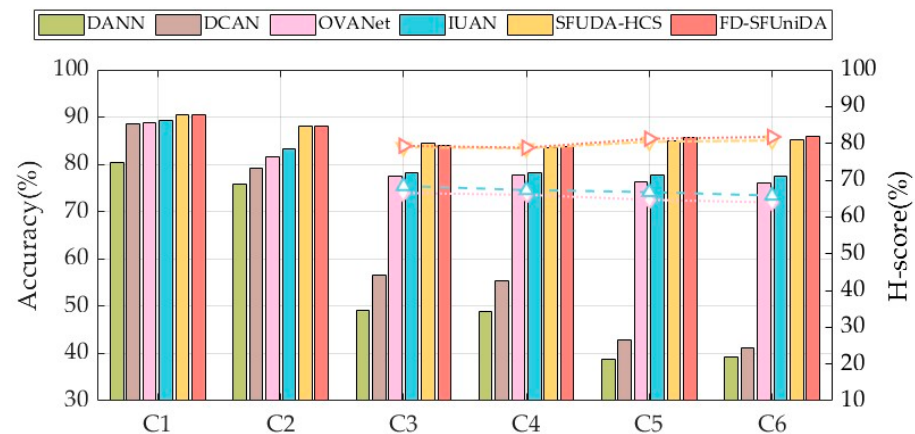


Figure 11. Comparison of fault diagnosis accuracy and H-score for different methods of reciprocating mechanism dataset. The broken lines represents the H-score.

Table 4. Average diagnostic accuracy and H-score of different methods on the rolling bearing dataset.

Methods	Task	T1	T2	T3	T4	T5	T6	T7	T8
DANN	Acc	83.20	82.73	75.22	74.17	52.58	39.35	43.22	19.17
DCAN	Acc	90.55	89.05	80.88	79.63	56.83	42.60	47.69	21.38
OVANet	Acc	87.40	88.63	83.06	82.29	77.25	77.20	74.63	70.29
	H-score	—	—	—	—	74.13	73.53	72.18	64.90
IUAN	Acc	90.13	89.38	83.50	82.50	79.18	78.38	76.66	71.38
	H-score	—	—	—	—	76.31	75.22	73.81	66.75
SFUDA-HCS	Acc	92.13	91.60	88.34	88.00	85.65	86.25	85.81	83.54
	H-score	—	—	—	—	84.94	85.52	85.56	83.42
FD-SFUniDA	Acc	92.60	91.83	88.59	89.04	85.25	85.70	86.41	84.17
	H-score	—	—	—	—	85.00	85.44	86.39	84.15

Table 5. Average diagnostic accuracy and H-score of different methods on the reciprocating mechanism dataset.

Methods	Task	C1	C2	C3	C4	C5	C6
DANN	Acc	80.31	75.79	49.19	48.91	38.71	39.21
DCAN	Acc	88.69	79.21	56.69	55.41	42.79	41.29
OVANet	Acc	88.81	81.71	77.59	77.81	76.21	76.08
	H-score	—	—	73.48	73.55	72.49	71.97
IUAN	Acc	89.19	83.21	78.31	78.19	77.71	77.42
	H-score	—	—	75.42	74.57	74.13	73.37
SFUDA-HCS	Acc	90.44	88.04	84.38	83.63	84.92	85.21
	H-score	—	—	83.62	83.38	84.77	85.04
FD-SFUniDA	Acc	90.59	88.21	84.09	83.81	85.71	86.00
	H-score	—	—	83.92	83.53	85.42	85.81

The DANN and DCAN methods demonstrate favorable diagnostic outcomes when applied to closed-domain adaptation tasks. However, both approaches struggle to recognize unknown class samples. Their recognition accuracy begins to decline when either the source or target domain contains private classes alongside shared classes, with a more pronounced drop when both private and shared classes coexist in both domains. For instance, in the open partial domain task, the H-scores of these two methods hovered around 50%, while the universal domain adaptation methods OVANet, and IUAN achieved 76%. Notably, the SFUDA-HCS method shows a superior and more consistent performance across all tasks. Despite this, the proposed FD-SFUniDA method achieves the highest accuracy, surpassing all compared methods in every task.

In the T8 task with a smaller shared domain, the proposed method still achieves an H-score of approximately 87%, at least 10% higher than OVANet, and IUAN. This demonstrates the developed method's more stable performance. The proposed method outperforms others across various domain adaptation tasks, proving its exceptional capability in addressing cross-domain challenges.

4.5. Performance Analysis of the Proposed Method

4.5.1. Confusion Matrix Analysis

Figure 12 displays the classification results for all categories across methods under the T7 task. It can be observed that the DANN and DCAN methods exhibit high accuracy when classifying known categories, achieving approximately 57% accuracy. However, lacking the capability to recognize unknown categories, they can only classify unknown categories into known ones. The proposed method outperforms IUAN, OVANet and SFUDA-HCS in classifying unknown categories. This superiority stems primarily from suppressing source domain private category 4 and more effectively distinguishing features between known and unknown categories. Figure 13 presents classification results for the proposed method on the reciprocating mechanism dataset C6 task, demonstrating equally excellent diagnostic performance.

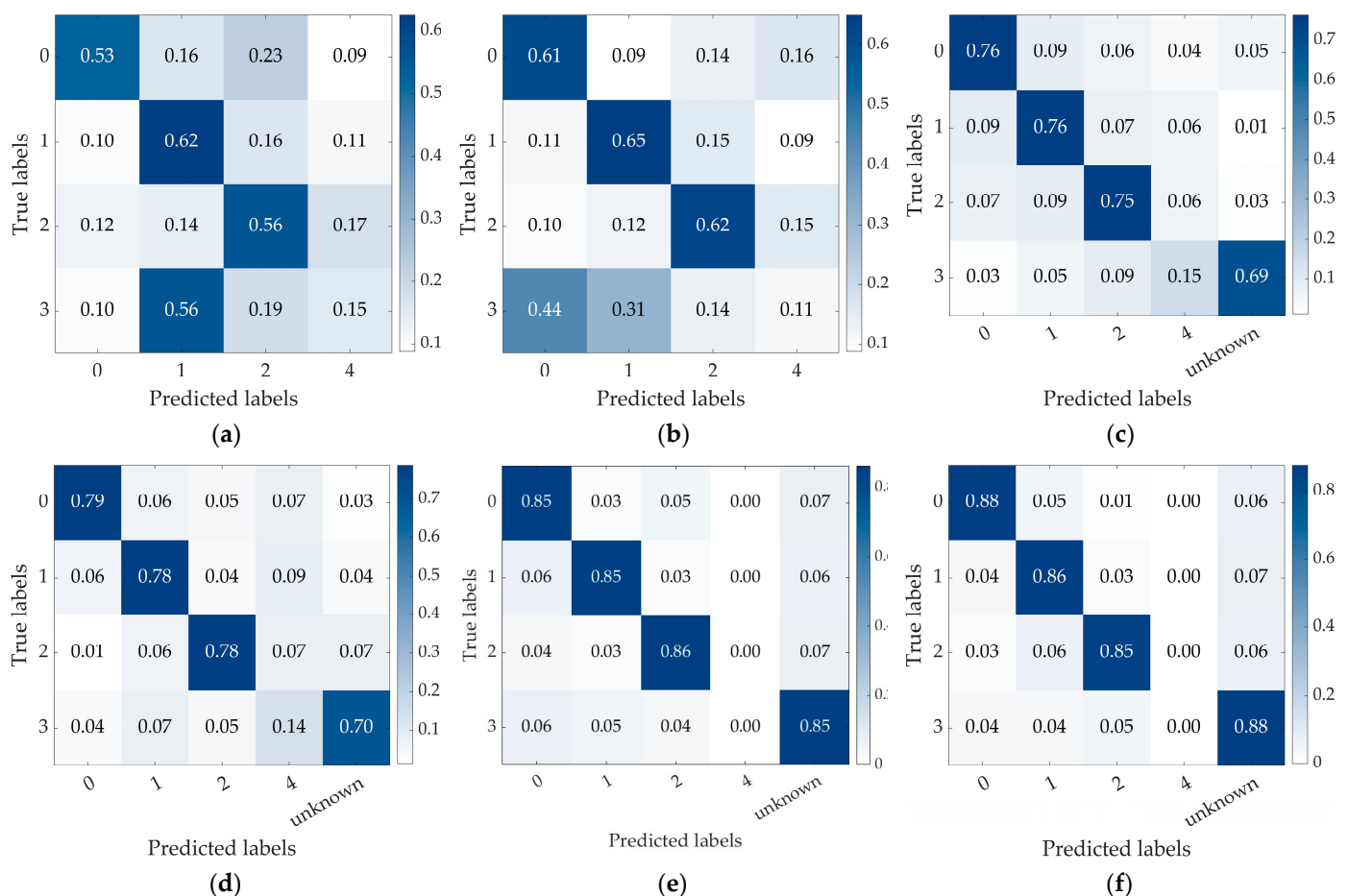


Figure 12. Comparison of classification results of different methods under the T7 task. (a) DANN; (b) DCAN; (c) OVANet; (d) IUAN; (e) SFUDA-HCS; (f) Proposed.

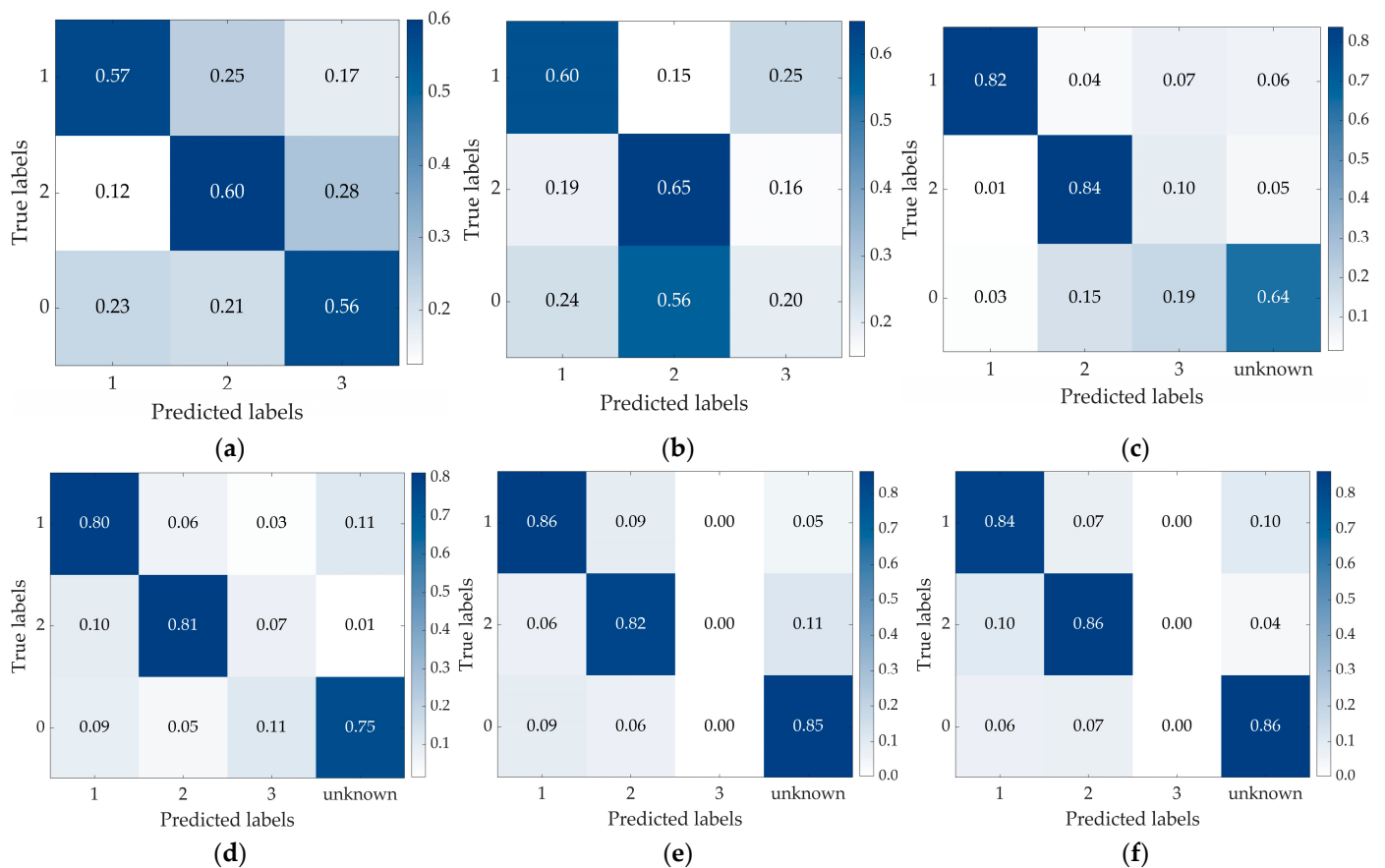


Figure 13. Comparison of classification results of different methods under the C6 task. (a) DANN; (b) DCAN; (c) OVANet; (d) IUAN; (e) SFUDA-HCS; (f) Proposed.

This confusion matrix analysis highlights the robust capabilities of FD-SFUniDA in handling complex fault diagnosis tasks, particularly its effective balancing of high accuracy and classification for known categories with precise identification of unknown classes. This equilibrium is especially critical when encountering unknown or emerging fault types in practical applications, ensuring the system's robustness and reliability.

4.5.2. Ablation Study

This section conducts systematic combined ablation experiments on pseudo-label learning loss $\mathcal{L}_{pl}$, neighborhood extension contrast loss $\mathcal{L}_{nc}$, and transfer enhancement loss $\mathcal{L}_{te}$, validating the independent contributions and synergistic effects of each core loss function within the dynamic semantic calibration framework. The experimental results are shown in Table 6. The complete model achieved optimal performance in cross-domain diagnosis tasks, with an average accuracy of 87.18%, significantly outperforming model configurations using only a single loss or a dual-loss combination.

Specifically, when using only the transfer enhancement loss $\mathcal{L}_{te}$, the model achieved an average accuracy of 79.54%; whereas using only the neighborhood extension contrast loss $\mathcal{L}_{nc}$ resulted in an average accuracy of only 73.46%. This discrepancy stems from the distinct operational mechanisms of the two loss functions: Transfer enhancement loss dynamically calibrates cross-domain feature distributions via category weighting, effectively mitigating feature shifts caused by operational variations. In contrast, Neighborhood extension contrast loss relies on pseudo labels to construct expanded positive-negative sample pairs, limiting its ability to optimize decision boundaries when target domain label distributions are imbalanced.

Table 6. Ablation studies results.

$\mathcal{L}_{pl}$	$\mathcal{L}_{nc}$	$\mathcal{L}_{te}$	Experiment 1	Experiment 2	Average
×	×	×	65.42	67.28	66.35
✓	×	×	73.58	75.26	74.42
×	✓	×	72.40	74.52	73.46
×	×	✓	80.14	78.93	79.54
✓	✓	×	76.84	78.12	77.48
✓	×	✓	85.27	84.38	84.83
×	✓	✓	82.04	80.57	81.31
✓	✓	✓	87.95	86.40	87.18

Among the dual-loss combinations, $\mathcal{L}_{pl} + \mathcal{L}_{nc}$ achieved an average accuracy of 77.48%. Its advantage lies in integrating the discriminative information from pseudo labels with the feature structure constraints imposed by contrastive learning. However, without the transfer enhancement mechanism, its ability to suppress interference from source domain private classes is insufficient. The $\mathcal{L}_{pl} + \mathcal{L}_{te}$ combination achieved the best performance with an average accuracy of 84.83%, demonstrating strong complementarity between pseudo-label learning and semantic calibration. This approach leverages the local structure of target domain samples while enhancing cross-domain consistency through semantic weight adjustment. The $\mathcal{L}_{te} + \mathcal{L}_{nc}$ combination achieved an average accuracy of 81.31%. Despite not employing pseudo-label supervision, it still achieved good cross-domain alignment through feature reconstruction and contrastive constraints.

Notably, further analysis via t-SNE visualization in the T7 task (Figure 14) indicates that $\mathcal{L}_{te}$ plays an irreplaceable role in suppressing source domain-specific categories. As shown, without $\mathcal{L}_{te}$ (Figure 14a), the source domain's private class (Category 4) overlaps significantly with the target domain's unknown classes in the feature space, causing classification confusion. In contrast, the complete model (Figure 14b) effectively increases the distance between private class features and shared class prototypes through semantic reconstruction, enhancing the model's ability to reject unknown categories.

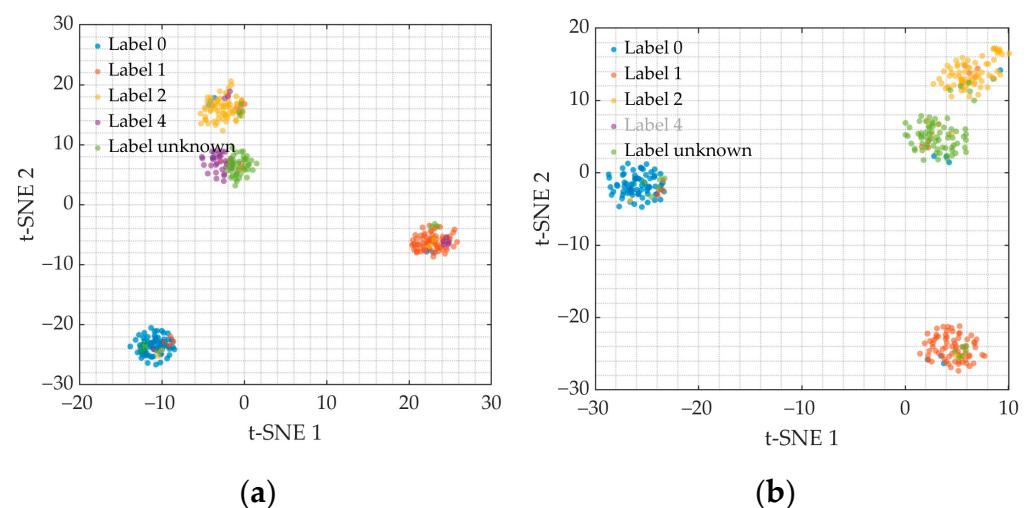


Figure 14. T-SNE Visualization for transfer enhancement loss. (a) $\mathcal{L}_{pl} + \mathcal{L}_{nc}$; (b) Three complete loss functions.

In summary, the three loss functions collectively form a cascaded optimization framework of “confidence constraint—feature alignment—pseudo-label augmentation.” This framework significantly enhances the accuracy and robustness of fault diagnosis in general

domain adaptation problems, validating the multi-loss mechanism's strong adaptability to complex diagnostic scenarios.

4.5.3. Parameter Sensitivity Analysis

A systematic series of hyperparameter sensitivity experiments was conducted to investigate the influence patterns of the neighborhood expansion contrast loss weight coefficient λ on the model's diagnostic performance. As shown in Figure 15, consistent convex curve patterns are observed across two heterogeneous datasets. When $\lambda \in [0.001, 5]_v$, the model's average diagnostic accuracy exhibits an initial increase followed by a decline, achieving global optimality at $\lambda = 0.2$. At this state, the model fully exploits the supervisory information from pseudo labels while enhancing feature space structural separability through contrastive loss. Within the broad range of λ values around $\lambda \in [0.05, 0.3]$, the model maintains stability above 85%. This demonstrates the proposed method's strong robustness to λ selection, providing generous parameter adjustment tolerance for engineering applications.

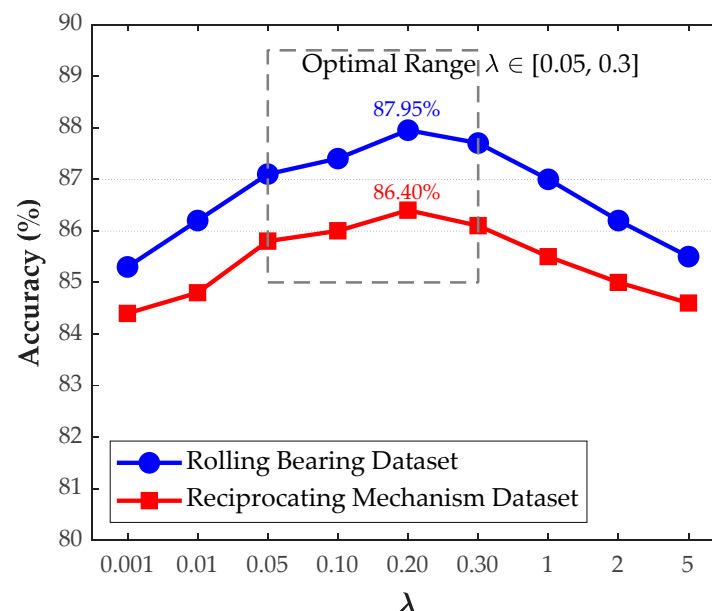


Figure 15. Sensitivity testing of hyperparameters.

4.5.4. Computational Complexity Analysis

This section conducts a detailed analysis of actual runtime overhead based on the T7 task to evaluate the practical application value of the FD-SFUniDA method. Table 7 presents a detailed comparison of the actual computational resource consumption between FD-SFUniDA and the comparison methods on the T7 task.

Table 7. Computational complexity comparison.

Methods	Total Computation Time (s)	GPU Memory (GB)
DANN	285	1.8
DCAN	312	2.1
OVANet	498	2.5
IUAN	526	2.6
SFUDA-HCS	583	2.9
FD-SFUniDA	645	3.2

As shown in Table 7, FD-SFUniDA outperforms other comparison methods in terms of total computation time, GPU memory consumption, and model parameter count. Its total computation time is 645 s, approximately 2.3 times longer than lightweight methods DANN and DCAN, and about 10.6% longer than SFUDA-HCS, with memory consumption reaching 3.2 GB. Despite its higher computational complexity, FD-SFUniDA achieves optimal diagnostic performance across multiple cross-domain tasks, significantly outperforming other methods. This demonstrates that the proposed method strikes a favorable balance between computational resources and diagnostic efficacy. In summary, FD-SFUniDA substantially enhances fault diagnosis capabilities in passive open environments at an acceptable computational cost. It is particularly well-suited for industrial applications demanding high diagnostic accuracy with relatively relaxed real-time constraints.

5. Conclusions

This study demonstrates that orthogonal decoupling of the feature space via singular value decomposition of pre-trained model weights provides an innovative solution to core challenges: the source-known subspace captures fault semantics shared across domains, while the target-private subspace isolates domain-specific interference, with the projection magnitudes offering a physical basis for unknown class identification. By integrating a dynamic decision boundary with an entropy threshold mechanism, feature drift caused by noisy pseudo-labels is effectively mitigated. The synergistic effect of neighborhood-expanded contrastive learning and semantic calibration maintains stable performance metrics above 84% in open-set partial recognition scenarios. Experiments across eight bearing operating conditions and six mechanism operating conditions validate the robustness of the framework, notably achieving an average diagnostic accuracy of 87.18% under stringent conditions where the target domain contains unknown fault types and source data is inaccessible. This research establishes a scalable paradigm for unsupervised diagnosis of dynamic faults in industrial equipment, with future work focusing on exploring its transferability in multi-sensor fusion scenarios.

Looking forward, this research establishes a scalable paradigm for unsupervised diagnosis of dynamic faults in industrial equipment, with several promising directions for future work. First, the exploration of multi-sensor fusion represents a natural extension, where vibration signals could be complemented with thermal, acoustic, and current data to create more comprehensive health indicators. Second, developing online/streaming adaptation capabilities would enable real-time model updates as new target data arrives, crucial for practical deployment in industrial IoT environments. Finally, integration with predictive maintenance systems could enhance the framework's practical value by enabling not only fault classification but also remaining useful life estimation and maintenance decision support. These extensions would further bridge the gap between academic research and industrial application in the domain of intelligent fault diagnosis.

Author Contributions: Conceptualization, P.Z. and W.L.; methodology, P.Z.; software, P.Z.; validation, W.L., and P.Z.; resources, W.L.; writing—original draft preparation, P.Z.; writing—review and editing, P.Z., S.S. and Q.Z.; supervision, S.S. and W.L.; project administration, Q.Z.; funding acquisition, W.L. All authors have read and agreed to the published version of the manuscript.

Funding: This research was funded by the Natural Science Foundation of Hubei Province (2023AFB900).

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: The original contributions presented in the study are included in the article; further inquiries can be directed to the corresponding author.

Conflicts of Interest: The authors declare no conflicts of interest.

Abbreviations

The following abbreviations are used in this manuscript:

FD-SFUniDA	Feature Decomposition-Based Source-Free Universal Domain Adaptation
SF-UniDA	Source-Free Universal Domain Adaptation
CBAM	Convolutional Block Attention Module
ResNet50	Residual Network 50-layer
SVD	Singular Value Decomposition
DA	Domain Adaptation
CLDA	Closed-Set Domain Adaptation
OSDA	Open-Set Domain Adaptation
PDA	Partial-Set Domain Adaptation
OPDA	Open-Partial-set Domain Adaptation
GMM	Gaussian Mixture Model
STFT	Short-Time Fourier Transform
JS	Jensen–Shannon Divergence
SPC	Source Private Classes
NC	Normal Condition
IF	Inner-race Fault
OF	Outer-race Fault
BF	Ball Fault
MF	Mixed Fault
CRF	Cracked Roller Failure
SPWF	Sliding Plate Wear Fault
DANN	Domain-Adversarial Training of Neural Networks
DCAN	Domain Conditioned Adaptation Network
UAN	Universal Adaptation Network
OVANet	One-vs-All Network for Universal Domain Adaptation
IUAN	Pseudo-Marginal-Based Universal Domain Adaptation

References

1. Rüßmann, M.; Lorenz, M.; Gerbert, P.; Waldner, M.; Justus, J.; Harnisch, M. Industry 4.0: The Future of Productivity and Growth in Manufacturing Industries. *Boston Consult. Group* **2015**, *9*, 54–89.
2. Yan, W.; Wang, J.; Lu, S.; Zhou, M.; Peng, X. A Review of Real-Time Fault Diagnosis Methods for Industrial Smart Manufacturing. *Processes* **2023**, *11*, 369. [\[CrossRef\]](#)
3. Yu, A.; Cai, B.; Wu, Q.; García, M.M.; Li, J.; Chen, X. Source-Free Domain Adaptation Method for Fault Diagnosis of Rotation Machinery under Partial Information. *Reliab. Eng. Syst. Saf.* **2024**, *248*, 110181. [\[CrossRef\]](#)
4. Tang, S.; Yuan, S.; Zhu, Y. Deep Learning-Based Intelligent Fault Diagnosis Methods Toward Rotating Machinery. *IEEE Access* **2020**, *8*, 9335–9346. [\[CrossRef\]](#)
5. Zhang, Y.; Ren, Z.; Feng, K.; Yu, K.; Beer, M.; Liu, Z. Universal Source-Free Domain Adaptation Method for Cross-Domain Fault Diagnosis of Machines. *Mech. Syst. Signal Process.* **2023**, *191*, 110159. [\[CrossRef\]](#)
6. Wang, X.; Shi, Z.; She, B.; Sun, S.; Qin, F.; Yan, X. Improved Deep Convolution Neural Network with Applications in Fault Diagnosis of Naval Gun Mechanical Components. In Proceedings of the 2021 China Automation Congress (CAC), Beijing, China, 22–24 October 2021; pp. 1865–1871.
7. Wang, X.; She, B.; Shi, Z.; Sun, S.; Qin, F. Partial Adversarial Domain Adaptation by Dual-Domain Alignment for Fault Diagnosis of Rotating Machines. *ISA Trans.* **2023**, *136*, 455–467. [\[CrossRef\]](#)
8. She, B.; Tan, F.; Zhao, Y.; Dong, H. Open-Set Domain Adaptation Fusion Method Based on Weighted Adversarial Learning for Machinery Fault Diagnosis. *J. Intell. Manuf.* **2024**, *36*, 5067–5086. [\[CrossRef\]](#)

9. Rombach, K.; Michau, G.; Fink, O. Controlled Generation of Unseen Faults for Partial and Open-Partial Domain Adaptation. *Reliab. Eng. Syst. Saf.* **2023**, *230*, 108857. [\[CrossRef\]](#)
10. Kang, S.; Tang, X.; Wang, Y.; Wang, Q.; Xie, J. Cross-Domain Fault Diagnosis Method for Rolling Bearings Based on Contrastive Universal Domain Adaptation. *ISA Trans.* **2024**, *146*, 195–207. [\[CrossRef\]](#)
11. Guo, L.; Yu, Y.; Liu, Y.; Gao, H.; Chen, T. Reconstruction Domain Adaptation Transfer Network for Partial Transfer Learning of Machinery Fault Diagnostics. *IEEE Trans. Instrum. Meas.* **2022**, *71*, 1–10. [\[CrossRef\]](#)
12. Deng, W.; Su, Z.; Qiu, Q.; Zhao, L.; Kuang, G.; Pietikäinen, M.; Xiao, H.; Liu, L. Deep Ladder Reconstruction-Classification Network for Unsupervised Domain Adaptation. *Pattern Recognit. Lett.* **2021**, *152*, 398–405. [\[CrossRef\]](#)
13. Wang, X.; Jiang, H.; Mu, M.; Dong, Y. A Dynamic Collaborative Adversarial Domain Adaptation Network for Unsupervised Rotating Machinery Fault Diagnosis. *Reliab. Eng. Syst. Saf.* **2025**, *255*, 110662. [\[CrossRef\]](#)
14. Huang, Y.; Peng, J.; Chen, N.; Sun, W.; Du, Q.; Ren, K.; Huang, K. Cross-Scene Wetland Mapping on Hyperspectral Remote Sensing Images Using Adversarial Domain Adaptation Network. *ISPRS J. Photogramm. Remote Sens.* **2023**, *203*, 37–54. [\[CrossRef\]](#)
15. Chen, X.; Shao, H.; Xiao, Y.; Yan, S.; Cai, B.; Liu, B. Collaborative Fault Diagnosis of Rotating Machinery via Dual Adversarial Guided Unsupervised Multi-Domain Adaptation Network. *Mech. Syst. Signal Process.* **2023**, *198*, 110427. [\[CrossRef\]](#)
16. Kundu, J.N.; Venkat, N.; Rahul, M.V.; Babu, R.V. Universal Source-Free Domain Adaptation. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), Seattle, WA, USA, 13–19 June 2020; pp. 4544–4553.
17. Sun, Y.; Dai, D.; Xu, S. Rethinking Adversarial Domain Adaptation: Orthogonal Decomposition for Unsupervised Domain Adaptation in Medical Image Segmentation. *Med. Image Anal.* **2022**, *82*, 102623. [\[CrossRef\]](#) [\[PubMed\]](#)
18. Qu, S.; Zou, T.; He, L.; Röhrbein, F.; Knoll, A.; Chen, G.; Jiang, C. LEAD: Learning Decomposition for Source-Free Universal Domain Adaptation. In Proceedings of the 2024 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), Seattle, WA, USA, 16 June 2024; IEEE: Seattle, WA, USA, 2024; pp. 23334–23343.
19. Chen, Z.; He, G.; Li, J.; Liao, Y.; Gryllias, K.; Li, W. Domain Adversarial Transfer Network for Cross-Domain Fault Diagnosis of Rotary Machinery. *IEEE Trans. Instrum. Meas.* **2020**, *69*, 8702–8712. [\[CrossRef\]](#)
20. Qu, S.; Zou, T.; Röhrbein, F.; Lu, C.; Chen, G.; Tao, D.; Jiang, C. GLC++: Source-Free Universal Domain Adaptation through Global-Local Clustering and Contrastive Affinity Learning. *IEEE Trans. Pattern Anal. Mach. Intell.* **2024**, *47*, 10646–10663. [\[CrossRef\]](#)
21. Liu, J.; Liu, Z.; Jia, Z.; Zhao, K. Source-Free Universal Domain Adaptation for Compressor Component Fault Diagnosis Guided by Hybrid Clustering Strategy. *Mech. Syst. Signal Process.* **2025**, *232*, 112771. [\[CrossRef\]](#)
22. Pan, C.; Shang, Z.; Tang, L.; Cheng, H.; Li, W. Open-Set Domain Adaptive Fault Diagnosis Based on Supervised Contrastive Learning and a Complementary Weighted Dual Adversarial Network. *Mech. Syst. Signal Process.* **2025**, *222*, 111780. [\[CrossRef\]](#)
23. Wang, W.; Li, C.; Zhang, Z.; Chen, J.; He, S.; Feng, Y. Pseudo-Label Assisted Contrastive Learning Model for Unsupervised Open-Set Domain Adaptation in Fault Diagnosis. *Reliab. Eng. Syst. Saf.* **2025**, *254*, 110650. [\[CrossRef\]](#)
24. Ge, Y.; Chen, Z.-M.; Zhang, G.; Heidari, A.A.; Chen, H.; Teng, S. Unsupervised Domain Adaptation via Style Adaptation and Boundary Enhancement for Medical Semantic Segmentation. *Neurocomputing* **2023**, *550*, 126469. [\[CrossRef\]](#)
25. Huang, X.; Zhu, C.; Ren, R.; Liu, S.; Huang, T. Source-Free Semantic Regularization Learning for Semi-Supervised Domain Adaptation. *IEEE Trans. Multimed.* **2025**, *27*, 4227–4239. [\[CrossRef\]](#)
26. Wang, H.; Zhao, C.; Chen, F. Feature-Space Semantic Invariance: Enhanced OOD Detection for Open-Set Domain Generalization. In Proceedings of the 2024 IEEE International Conference on Big Data (BigData), Washington, DC, USA, 15–18 December 2024; pp. 8244–8246.
27. Woo, S.; Park, J.; Lee, J.-Y.; Kweon, I.S. CBAM: Convolutional Block Attention Module. In Proceedings of the CBAM: Convolutional Block Attention Module, Munich, Germany, 8–14 September 2018; pp. 3–19.
28. Ganin, Y.; Ustinova, E.; Ajakan, H.; Germain, P.; Larochelle, H.; Laviolette, F.; March, M.; Lempitsky, V. Domain-Adversarial Training of Neural Networks. *J. Mach. Learn. Res.* **2016**, *17*, 1–35.
29. Li, S.; Liu, C.H.; Lin, Q.; Xie, B.; Ding, Z.; Huang, G.; Tang, J. Domain Conditioned Adaptation Network. In Proceedings of the The Thirty-Fourth AAAI Conference on Artificial Intelligence (AAAI-2020), New York, NY, USA, 3 April 2020.
30. Li, J.; Zhang, X.; Yue, K.; Chen, J.; Chen, Z.; Li, W. An Auto-Regulated Universal Domain Adaptation Network for Uncertain Diagnostic Scenarios of Rotating Machinery. *Expert Syst. Appl.* **2024**, *249*, 123836. [\[CrossRef\]](#)
31. Yin, Y.; Yang, Z.; Wu, X.; Hu, H. Pseudo-Margin-Based Universal Domain Adaptation. *Knowl.-Based Syst.* **2021**, *229*, 107315. [\[CrossRef\]](#)
32. Yan, X.; Liang, W.; Zhang, G.; Tian, F. Adaptive Extraction Method of Unit Period Time Series Based on GSA-IFCM. *J. Detect. Control* **2022**, *44*, 60–65, 72.

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.