



MgEL: Quantum Entanglement-Inspired Evidence Fusion for Learning with Noisy Labels

Fir Dunkin¹ and Xinde Li^{1,2,3,*}

¹ Key Laboratory of Measurement and Control of CSE, School of Automation, Southeast University, Nanjing 210018, China

² Faculty of Robot Science and Engineering, Northeastern University, Shenyang 110167, China

³ Southeast University Shenzhen Research Institute, Shenzhen 518063, China

Abstract

With the rise of data engineering-driven automatic annotation strategies, deep learning has demonstrated remarkable performance and strong competitiveness in intelligent fault diagnosis. However, the inherent limitations of automatic annotators inevitably introduce noisy labels, which in turn hinder the generalization and accuracy of diagnostic models. Although numerous Learning with Noisy Labels (LNL) methods attempt to alleviate the impact of label noise through sample selection or label correction, most rely heavily on model predictions to guide training. This self-reinforcing mechanism frequently leads to confirmation bias, especially under high-noise conditions, thereby limiting their effectiveness. To address these challenges while preserving the full data utility, this paper proposes a novel approach termed the *Multi-granularity Evidence Labels* (MgEL), inspired by the principles of quantum entanglement and collapse. In MgEL, we

perform feature-space fusion between entangled sub-distributions to construct a superposition state, from which two auxiliary labels are derived: a pseudo-label obtained by selecting the class with the maximum amplitude and a collapsed label sampled probabilistically according to the class-wise amplitude distribution. The collapsed label represents an uncertainty-aware observation, while the pseudo-label represents the most confident class estimation. These are then fused with the original annotation to form multi-granularity evidence labels. This approach allows MgEL to suppress confirmation bias and improve robustness under noisy supervision. Extensive experiments validate the effectiveness and reliability of MgEL, particularly in high-noise scenarios (e.g., noise intensity $\eta \geq 80\%$), underscoring its potential for practical deployment in low-cost, data-driven intelligent fault diagnosis systems.

Keywords: learning with noisy labels, multi-granularity information fusion, fault diagnosis, deep learning, classification of time series signals.



Academic Editor:

Deqiang Han

Submitted: 21 April 2025

Accepted: 23 July 2025

Published: 26 September 2025

Vol. 2, No. 3, 2025.

10.62762/CJIF.2025.151851

*Corresponding author:

✉ Xinde Li

xindeli@seu.edu.cn

Citation

Dunkin, F., & Li, X. (2025). MgEL: Quantum Entanglement-Inspired Evidence Fusion for Learning with Noisy Labels. *Chinese Journal of Information Fusion*, 2(3), 253–274.



© 2025 by the Authors. Published by Institute of Central Computation and Knowledge. This is an open access article under the CC BY license (<https://creativecommons.org/licenses/by/4.0/>).

1 Introduction

With the rapid development of big data technologies, deep learning has made remarkable strides in both theoretical research and practical applications of intelligent fault diagnosis [3, 5]. However, the success of these models hinges on the availability of large-scale, high-quality annotated data, which often entails substantial manual labeling costs [4]. This requirement severely limits the scalability and industrial deployment of deep learning-based diagnostic models, especially in scenarios with constrained annotation resources [1].

To address this bottleneck, cost-effective alternatives, such as crowdsourcing [6] and data-engine-driven automatic labeling [7], have been widely adopted. While these strategies significantly reduce annotation efforts, they inevitably introduce noisy labels—instances [8] where the annotated label deviates from the true class ($\mathcal{Y} = k \neq \hat{\mathcal{Y}}$). Such label noise distorts the data distribution, misaligns decision boundaries [9], and ultimately undermines both the accuracy and generalization of diagnostic models [10].

To enhance model robustness under noisy supervision, Learning with Noisy Labels (LNL) [11] has emerged as a prominent research direction. Two mainstream strategies have been extensively explored: sample separation [17] and label correction [12]. Sample separation aims to iteratively identify clean samples via model predictions and use only those for further training [13], while label correction refines noisy labels by integrating observed labels with predictive cues (e.g., latent features [14] or logits [15]). Despite their empirical success, both approaches suffer from a key limitation: they are heavily reliant on model predictions, which are themselves corrupted by noise [16], leading to the accumulation of confirmation bias during self-guided training.

Essentially, these models are guided by their own decisions in noisy environments. Erroneous predictions may be repeatedly reinforced, eventually forming a biased representation of the data distribution (see Figure 1), which motivates the central question of this work:

*How can we effectively leverage the full information of noisy datasets while **mitigating confirmation bias** caused by inaccurate predictions during self-guided training, in order to improve the generalization of diagnostic models under severe noise?*

To address this challenge, we draw inspiration from

quantum entanglement [18] and collapse, proposing a novel framework called “Multi-granularity Evidence Labels (MgEL)”. MgEL constructs a set of multi-granularity labels by integrating three label sources: the original annotation labels $\mathcal{Y}$, dynamically generated pseudo-labels $\mathcal{Y}'$, and probabilistic collapsed labels $\bar{\mathcal{Y}}$ derived from the feature space. This process of combining two sub-distributional structures simulates the observation-collapse process of a quantum-entangled system, enabling more robust supervision in the presence of label noise.

Specifically, MgEL conceptualizes a feature cluster $\mathcal{C}$ containing conflicting labels as a superposition state $\varphi = \sum_k c_k |\phi_k\rangle$. The complex amplitudes c_k are transformed into real-valued probabilities by borrowing the Born rule— $P(k) = \|c_k\|^2 = \lambda_k$ —to enable probabilistic reasoning. Beyond this, class-wise evidential values are introduced to quantify the “belief mass” that a sample belongs to each latent class (or basis state). These evidential values are not a direct conversion from the Born rule but rather are derived by modeling the internal distribution within each feature cluster under uncertainty, providing a structured representation of decision confidence.

Unlike conventional classification schemes that treat each sample as a single independent observation, MgEL constructs a pseudo-multi-source observation framework in the feature space, which better reflects the multi-view nature of real-world noisy data. The fusion of the original label y_i , pseudo-label y'_i , collapsed label $\bar{y}_i$, and evidential $\mathcal{E}_i = \{e_k\}_{k=1}^K$ simulates repeated observations across different measurement bases. This design suppresses overreliance on single predictions and alleviates the accumulation of confirmation bias, thereby improving robustness during training.

We conduct extensive experiments on three fault diagnosis datasets to validate the effectiveness of MgEL, which demonstrate that MgEL outperforms other methods, particularly under scenarios with severe label noise (i.e., noise intensity $\eta \geq 80\%$), significantly improving diagnostic accuracy and reducing sensitivity to annotation quality. These findings suggest that MgEL has the potential to reduce annotation costs and enable the reliable deployment of intelligent diagnostic models in practical industrial settings.

The key contributions of this work are summarized as follows:

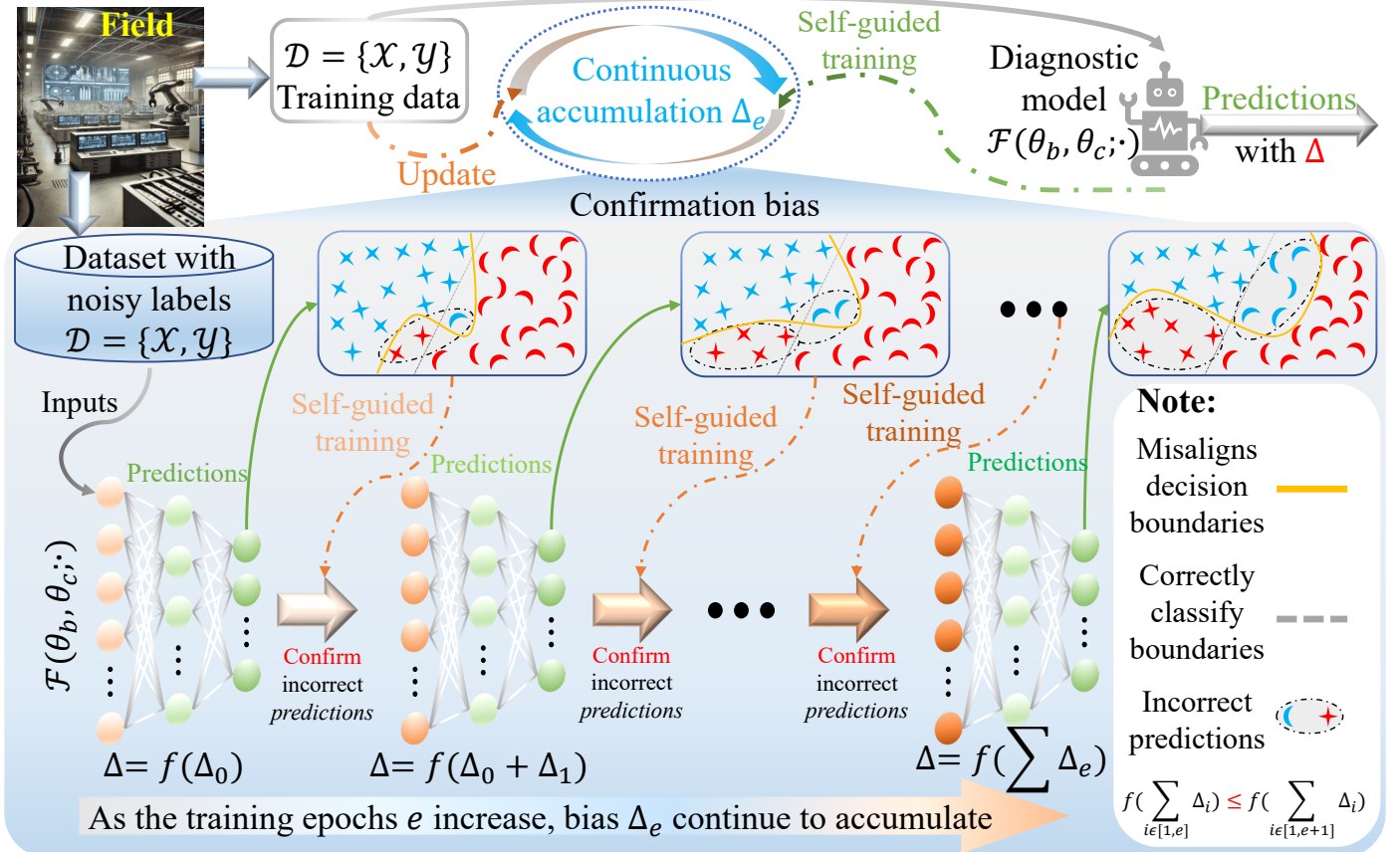


Figure 1. Confirmation bias in LNL: Model predictions, inherently corrupted by label noise, are repeatedly used to guide training, leading to the accumulation of bias Δ and degraded generalization.

1. **Theory:** We establish a *pseudo-multi-source observation modeling* method in the feature space, extending the theoretical foundation of decision-level information fusion.
2. **Methodology:** We propose a robust *multi-granularity label construction strategy* by introducing class-wise evidential values under uncertainty, which enhances the model's tolerance to noisy supervision.
3. **Empirical validation:** We conduct extensive experiments across three real-world fault diagnosis datasets with varying noise intensities. MgEL consistently outperforms baselines, especially under severe corruption (e.g., $\eta \geq 80\%$), demonstrating its practical feasibility for robust learning under noisy supervision—though not yet deployed in full industrial pipelines.

2 Preliminary

2.1 Nomenclature

To improve clarity and reduce potential ambiguity, Table 1 provides a formal summary of all key notations

used throughout the paper.

2.2 Problem formulation (Fault diagnosis with noisy labels)

Fault diagnosis with noisy labels is typically formulated as a multi-class classification task under annotation uncertainty [4]. Let the training dataset be $\mathcal{D} = \{\mathcal{X}, \mathcal{Y}\} = \{(x_i, y_i)\}_{i=1}^N$, where each $x_i \in \mathbb{R}^{C \times L}$ represents a multivariate time-series signal with C channels and L sampling points, and $y_i \in [1, K]$ is the corresponding fault label drawn from K predefined categories. These samples typically come from industrial systems, including mechanical, electrical, and structural components [21], where condition monitoring sensors collect time-series signals for predictive maintenance and fault detection [19].

In practice, noisy labels are prevalent due to multiple factors [20], including ambiguous or overlapping fault manifestations, inconsistencies among domain experts, and, more importantly, the limited reliability of automated annotation systems [23], which may mislabel large volumes of data due to heuristic rules or insufficient context [22]. Let y_i^* denote the latent true label of x_i . If $y_i \neq y_i^*$, the sample is considered

Table 1. Notation summary table.

Symbol	Description
$x_i \in \mathbb{R}^{C \times L}$	Input sample with C channels and L time points
$y_i \in \{1, \dots, K\}$	Annotated label for x_i
$z_i \in \mathbb{R}^d$	Latent feature of x_i extracted by the backbone network
$\mathcal{Z}_A, \mathcal{Z}_B$	Randomly partitioned feature subspaces (entangled sets)
$\mathcal{S}_{\cos}(z_i, z_j)$	Exponential cosine similarity between features z_i and z_j
$\mathcal{C}_i$	Feature cluster centered at z_i
$\varphi_i \in \mathbb{R}^K$	Superposition state encoding class affiliation amplitudes
$\lambda_{ik} \in [0, 1]$	Amplitude associated with class k for sample i
$ \varphi_k\rangle \in \mathbb{R}^K$	Canonical basis vector for class k
$\mathcal{E}_i = \{e_k\}_{k=1}^K$	Evidence vector obtained from pseudo-source fusion
$\tilde{y}_i \in \mathbb{R}^K$	Multi-granularity fused label used for training
$\mathcal{F}(\theta_b, \theta_c; \cdot)$	Diagnostic model with backbone θ_b and classifier θ_c
$\tilde{\eta} \in [0, 1]$	Estimated noise intensity for dynamic cluster adjustment
$\mathcal{H}_A, \mathcal{H}_B$	Hilbert spaces of entangled subsystems A and B
$p_i \in \mathbb{R}^K$	Model prediction vector for input x_i
$\mathcal{L}_{CE}(p_i, \tilde{y}_i)$	Cross-entropy loss between prediction p_i and label $\tilde{y}_i$
$D = \{(x_i, y_i)\}_{i=1}^N$	Training dataset containing N noisy samples
y'_i	Pseudo-label derived via superposition collapse
$\bar{y}_i$	Collapsed label via multinomial sampling from φ_i
$\tilde{y}_i$	Final fused label
$\mathcal{M}_{y_i^*}$	Clean class-conditional manifold for ground-truth label y_i^*
$\mathcal{F}_\eta$	Feature extractor perturbed by label noise rate η
ξ_i	Noise-induced deviation of z_i from its clean manifold
T_{kl}	Probability of class- k label being flipped to class- l

mislabeled. However, the noise distribution or transition matrix $T(\mathcal{Y} | \mathcal{Y}^*)$ is typically unknown [24]. Depending on whether the noise occurs randomly or in a class-dependent fashion, it is commonly categorized as symmetric (uniform) or asymmetric (class-dependent)—the latter often reflecting realistic diagnostic confusion among fault types with similar signal patterns [13].

The learning objective is to construct a diagnostic model $\mathcal{F}(\theta_b, \theta_c; \cdot)$, comprising a backbone θ_b and a classifier θ_c , that maintains robust generalization to clean labels despite being trained on corrupted data, which can be formalized as:

$$\operatorname{argmin}_{\theta_b, \theta_c} \{ \mathbb{E}_{x_i \sim \mathcal{X}} (| \mathcal{F}(\theta_b, \theta_c; x_i) - \hat{y}_i |) \}. \quad (1)$$

2.3 Related works on LNL

Deep neural networks (DNNs) tend to overfit when trained on noisy labels, leading to poor generalization and unreliable predictions [26]. To address this, numerous strategies have been proposed to enhance model robustness under label noise [11], which are

broadly categorized into three research areas: robust loss functions [27], sample separation [28], and label correction [29].

Robust loss functions reduce the negative impact of mislabeled samples during training. For example, SCE [30] includes a symmetric regularization term to reduce the dominance of noisy labels, while NLS [31] uses label smoothing to lower label confidence. However, these approaches often rely on strong assumptions about the noise distribution, such as the need for access to or accurate estimation of the label transition matrix $T(\mathcal{Y} | \mathcal{Y}^*)$, which is difficult to obtain in practical scenarios [32].

A second class of methods focuses on sample separation, aiming to distinguish clean from noisy samples based on training dynamics. Many early works exploit loss-based heuristics, assuming that samples with lower loss are more likely to be correctly labeled. Representative methods include JoCoR [33], which trains peer networks on mutually selected small-loss samples to refine the dataset. More recent approaches, such as DISC [25], introduce

memory-based dynamic thresholding to categorize training data into clean, hard, and correctable subsets. However, these methods depend fundamentally on the model's own predictions to estimate sample reliability. As a result, incorrect assessments during early training stages may be reinforced across epochs—this phenomenon is known as confirmation bias (see Figure 1). Even dynamic sample selection strategies cannot fully overcome this limitation, as the thresholds are still updated based on past model behavior [34].

To mitigate data underutilization caused by sample removal, label correction approaches that generate pseudo-labels from model predictions to supervise training have emerged as a promising alternative [35]. For instance, SED [32] uses a mean-teacher framework to stabilize label updates. Despite their effectiveness, these methods also suffer from confirmation bias: initial incorrect predictions can propagate and become increasingly difficult to reverse during subsequent training [16].

This shared limitation in both sample separation and label correction methods arises from their "self-guided" nature—they rely solely on the model's internal signals, making them vulnerable to early-stage errors that propagate unchecked. To overcome this, inspired by principles of quantum mechanics, MgEL emulates the repeated observation-collapse mechanism of entangled quantum systems, introducing controlled uncertainty throughout the training process. By periodically re-evaluating fused labels with multi-granularity representations, MgEL disrupts confirmation bias and prevents training from being overly influenced by prior incorrect predictions. In this way, MgEL fully leverages the information contained in the original dataset, even when labels are partially corrupted, and improves the robustness of the corrected supervision signals.

2.4 Differentiation from previous works

This work extends our earlier research on learning with noisy labels, specifically MgCF [40] and MgL [41], both of which leverage feature clustering for robust label correction. Although MgCF, MgL, and the proposed MgEL share a unified design philosophy: transforming noisy labels into structured supervision via latent-space modeling, they differ significantly in terms of label granularity representation, fusion strategy, belief assignment, and computational scalability. To clarify these distinctions and avoid any confusion regarding academic overlap,

we provide a comparative analysis, also visually summarized in Figure 2.

First, regarding label granularity semantics, all three frameworks aim to suppress confirmation bias by constructing multi-granularity supervisory signals. In this context, MgCF introduces a dual-granularity labeling scheme: fine-grained labels are formed when annotation and pseudo-labels agree, indicating high class certainty, whereas coarse-grained labels are used when disagreement occurs, signaling ambiguity. MgL builds on this structure by incorporating an additional label source—the collapsed label sampled from the superposition state φ_i —which allows the construction of medium-grained labels when only partial agreement exists among the three sources (annotation, pseudo, collapsed). In contrast, MgEL reverts to a two-level granularity structure as in MgCF but complements it with a downstream decision deferral mechanism that handles fully inconsistent labels by reverting to the original superposition state φ_i . While this does not introduce an explicit third granularity level, it implicitly absorbs semantically hesitant cases through adaptive rejection, thereby preserving the interpretive flexibility offered by multi-granularity labeling.

In terms of label fusion mechanisms, both MgCF and MgL directly use the fused multi-granularity label $\tilde{Y}$ as the supervisory signal in training. MgL distinguishes itself by explicitly incorporating collapsed labels $\bar{Y}$ into the fusion rule alongside annotations and pseudo-labels, thus enabling medium-grained representations. MgEL, in contrast, restricts the fusion to annotation and pseudo-labels but introduces a *Coherence Mechanism* (Sec. 3.5) that evaluates whether any two of the three available labels (annotation, pseudo, collapsed) are consistent. If no consistency is found, the fused label $\tilde{Y}$ is rejected and replaced with φ_i , effectively reverting supervision to a probabilistic representation. This mechanism guards against overconfident but unreliable updates and introduces an adaptive rejection pathway not present in either MgCF or MgL.

The belief assignment strategies adopted by the three methods also exhibit fundamental differences. In both MgCF and MgL, the amplitude probabilities derived from the superposition state φ_i —which reflects the relative class distribution within each feature cluster—are directly treated as empirical belief masses. These unregularized values are used to assign fusion weights between annotations and pseudo-labels,

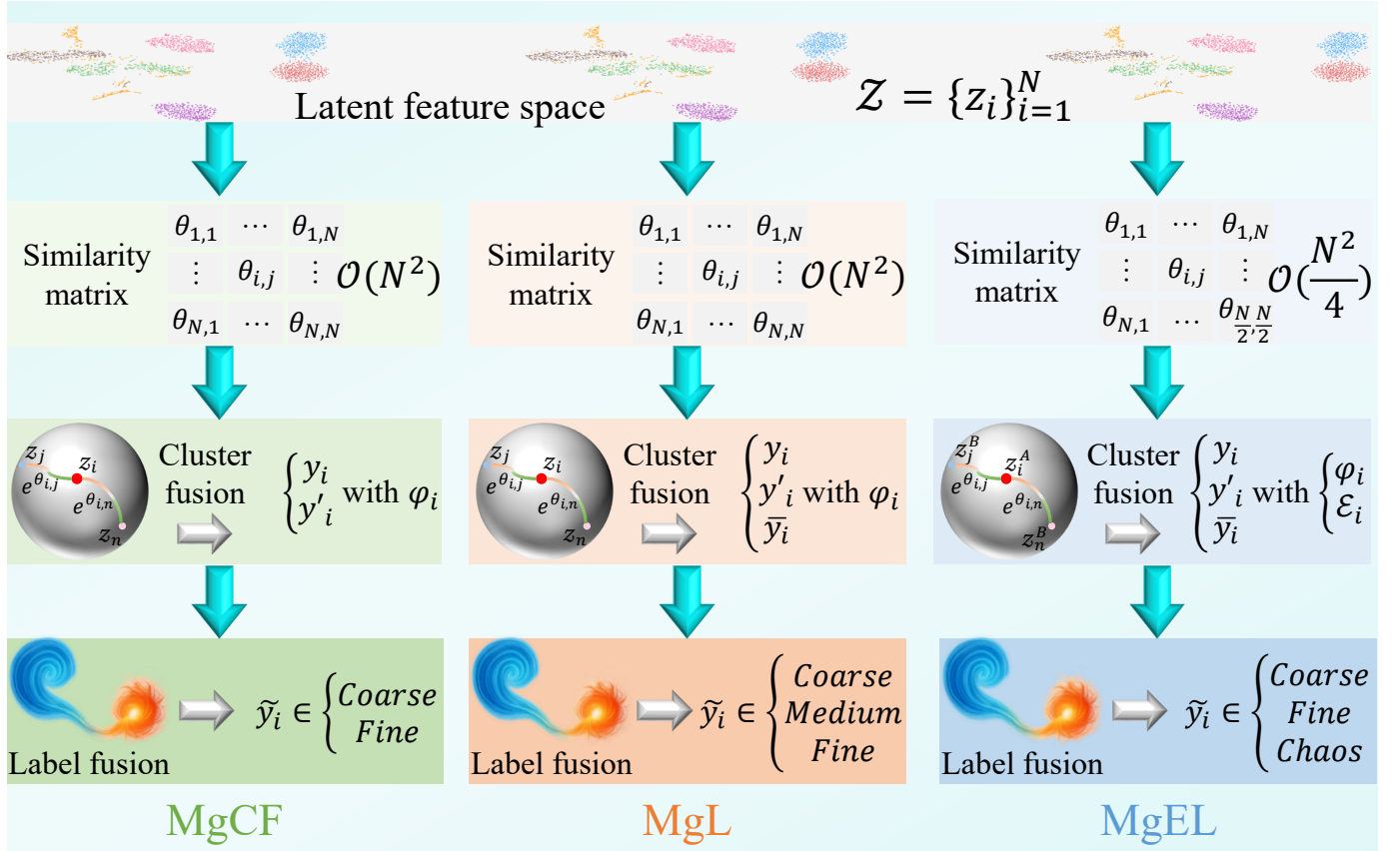


Figure 2. Comparative illustration of the MgCF [40], MgL [41], and the proposed MgEL, which highlights the core distinctions among the three methods in terms of computational scalability, granularity expressiveness, uncertainty modeling, and the belief assignment strategy adopted during label fusion.

implicitly treating intra-cluster frequencies as reliable class evidence. In contrast, MgEL introduces a belief regularization mechanism. Rather than using the raw amplitudes in φ_i , MgEL calibrates the belief assignment by incorporating a global noise estimate $\tilde{\eta}$, constructing an evidential vector $\mathcal{E}_i = \{e_k\}_{k=1}^K$ that reflects class-wise reliability under dataset-level uncertainty. This transformation from empirical frequencies to noise-aware belief masses enhances the robustness of label fusion, particularly in high-noise scenarios, and aligns with the principles of uncertainty modeling in evidential reasoning frameworks.

Finally, in terms of computational scalability, MgCF and MgL both compute global similarity matrices of size $\mathcal{O}(N^2)$ to construct feature clusters across the full dataset, which incurs substantial memory and runtime costs. MgEL, inspired by the partitioned structure of entangled quantum systems, proposes an entangled subspace design: the latent feature space is randomly split into two disjoint subsets, with each sample querying only across the opposite subspace. This reduces similarity computation to $\mathcal{O}(N^2/4)$ and introduces randomization effects similar to dropout.

This modification not only improves training efficiency but also enhances cluster robustness by avoiding deterministic, potentially biased global modeling.

In summary, MgEL consolidates and extends the prior frameworks by integrating evidential trust modeling, a coherence-driven rejection mechanism, and a more scalable partitioned clustering strategy. These contributions allow it to generalize effectively across high-noise environments while maintaining theoretical and algorithmic distinctions from both MgCF and MgL.

2.5 Entangled Quantum-inspired foundations

In quantum mechanics, the observation-collapse mechanism explains how the measurement of one particle in an entangled system causes the entire wavefunction to collapse instantaneously into a corresponding eigenstate [36]. This collapse not only determines the measured particle's state but also instantaneously defines the state of its entangled counterpart, reflecting a non-local correlation [37]. This principle embodies measurement-induced state transitions and decision outcomes under

uncertainty [38], providing rich inspiration for information fusion in noisy, uncertain environments.

Formally, we consider a bipartite quantum system composed of two subsystems A and B , with a joint entangled state expressed as:

$$|\Psi\rangle = \sum_{k=1}^K \lambda_k |a_k\rangle_A |b_k\rangle_B, \quad (2)$$

where $|\Psi\rangle$ denotes the overall quantum state in the composite Hilbert space $\mathcal{H}_A \otimes \mathcal{H}_B$, with $\mathcal{H}_A$ and $\mathcal{H}_B$ denoting the Hilbert spaces of subsystems A and B , respectively. The vectors $|a_k\rangle_A$ and $|b_k\rangle_B$ are orthonormal eigenstates of the respective subsystems, and $\lambda_k \in \mathbb{C}$ is the complex amplitude associated with each joint basis pair, normalized such that $\sum_{k=1}^K |\lambda_k|^2 = 1$. The index $k \in \{1, \dots, K\}$ enumerates the possible entangled basis components, forming the mathematical foundation of our analogy.

Upon measuring subsystem A and obtaining the outcome $|a_k\rangle_A$, the entire system collapses into the product state $|a_k\rangle_A |b_k\rangle_B$, with a probability of $|\lambda_k|^2$. This structure implies that the state of B is instantaneously determined by observing A , without direct interaction with B . Such non-local inference allows one subsystem to act as an informational proxy for the other—a foundational feature of quantum entanglement.

In MgEL, this collapse mechanism serves as a conceptual inspiration, rather than a physical simulation. We draw an analogy between the measurement-induced resolution of uncertainty and the process of integrating multiple label sources into a consistent supervisory signal. Concretely, three label sources are constructed for each sample: (I) the original annotation, (II) a pseudo-label corresponding to the maximum-amplitude component in the superposition state, and (III) a collapsed label probabilistically sampled based on the amplitude distribution. These label sources represent diverse, imperfect observations of the same latent feature representation.

To further emulate the informational interdependence observed in entangled systems, we partition the latent feature space $\mathcal{Z}$ into two disjoint subspaces, $\mathcal{Z}_A$ and $\mathcal{Z}_B$, and construct feature clusters through cross-subspace querying—samples in $\mathcal{Z}_A$ are clustered based on their similarity to samples in $\mathcal{Z}_B$, and vice versa. This design allows feature clusters derived from one subspace to act as interpretive surrogates

for evaluating the class membership of samples in the other. In doing so, each subspace imposes a structural constraint on its counterpart, akin to the inference structure in bipartite entanglement. Additionally, since each sample observes only a subset of the latent space per epoch, this formulation introduces Dropout-like stochasticity, reducing memory complexity and enhancing generalization.

When partial agreement arises among the constructed label sources, it is interpreted as a consistency-triggered decision signal, and the fused label is then adopted for training. Conversely, when all sources disagree, no definitive decision is made; the model retains the superposition-based representation, deferring commitment due to measurement incoherence.

We emphasize that all quantum-theoretic terminology is used metaphorically to guide the modeling of uncertainty, structural supervision, and decision consistency. No physical entanglement or non-local interaction is simulated or implied.

2.6 Manifold perturbation under label noise: Modeling and implications

To theoretically examine whether extreme label noise (e.g., $\eta \geq 80\%$) induces structural shifts in the latent space, we first formalize the notion of label-induced manifold perturbation. We posit that such shifts do not alter the marginal distribution $P(\mathcal{X})$ but instead distort the conditional representation $P(\mathcal{Z}|\mathcal{Y}^*)$ learned by the model. The following assumption summarizes our conclusion:

Assumption 1 (Noise-induced manifold perturbation) Under label noise with a corruption rate of η , the latent representation $z_i = \mathcal{F}_\eta(x_i)$ deviates from its ideal, noise-free class manifold $\mathcal{M}_{y_i^*}$ with an expected squared perturbation:

$$\mathbb{E} [\text{Dist}(z_i, \mathcal{M}_{y_i^*})^2] \propto \eta \cdot \text{Tr}(\Sigma_T),$$

where Σ_T is the covariance of the perturbation process governed by the label transition matrix T . This distortion causes a representation-level structural shift that scales with noise intensity.

We now justify Assumption 1 through formal modeling and analysis.

Let (x_i, y_i^*) denote a clean training sample, and y_i the observed (potentially corrupted) label. Due to the incorrect supervision, the model learns a perturbed representation:

$$z_i = \mathcal{F}_\eta(x_i) = z_i^* + \xi_i, \quad (3)$$

where ξ_i denotes the feature-level perturbation induced by the label error. Since this perturbation accumulates over multiple training steps via gradient descent, we express it as:

$$\xi_i = \sum_{t=1}^T \Delta z_i^{(t)}(y_i). \quad (4)$$

We consider the label corruption process governed by a class transition matrix $T = \{T_{kl}\}$ with $T_{kl} = P(y_i = l \mid y_i^* = k)$. The induced perturbation distribution is thus:

$$\xi_i \sim \mathcal{P}_\eta = (1 - \eta) \cdot \delta_0 + \eta \cdot \mathcal{Q}(T), \quad (5)$$

where δ_0 denotes the no-perturbation case (correct labels), and $\mathcal{Q}(T)$ is a distribution over perturbations generated by incorrect labels sampled according to T .

This mixture formulation in Eq. (5) provides a unified framework for modeling both symmetric and asymmetric label noise. In the symmetric case, where $T_{kl} = \frac{1}{K-1}$ for all $k \neq l$, the perturbation distribution $\mathcal{Q}(T)$ approximates an isotropic Gaussian, i.e., $\mathcal{N}(0, \sigma^2 I)$. In contrast, under asymmetric noise where T is sparse and class-dependent, $\mathcal{Q}(T)$ becomes a mixture of Gaussians with non-zero means, as defined in Eq. (6).

$$\mathcal{Q}(T) = \sum_{l \neq k} T_{kl} \cdot \mathcal{N}(\mu_{kl}, \Sigma_{kl}). \quad (6)$$

We define the manifold perturbation error (MPE) as follows:

$$\text{MPE}_i := \min_{z \in \mathcal{M}_{y_i^*}} \|z_i - z\|_2 = \|\xi_i\|, \quad (7)$$

and hence,

$$\mathbb{E}[\text{MPE}_i^2] = \eta \cdot \mathbb{E}_{\xi_i \sim \mathcal{Q}(T)} [\|\xi_i\|^2] = \eta \cdot \text{Tr}(\Sigma_T), \quad (8)$$

where Σ_T is the effective covariance structure aggregated over the perturbation distribution. This result confirms that as η increases, the average deviation from the class-specific latent manifold grows linearly, reflecting a progressive distortion in $P(\mathcal{Z}|\mathcal{Y}^*)$.

To quantify the global structural deformation, we define the manifold distortion index (MDI) as:

$$\text{MDI}(\eta) := \frac{1}{N} \sum_{i=1}^N \text{MPE}_i \propto \sqrt{\eta \cdot \text{Tr}(\Sigma_T)}. \quad (9)$$

The derivation above substantiates the hypothesis in Assumption 1 by showing how high-intensity label noise induces representation-level perturbations that scale with both the noise rate η and the structural properties of the label transition matrix T . In particular, the manifold distortion index $\text{MDI}(\eta)$ provides a global quantitative measure of such perturbations, confirming that noisy supervision leads to a pseudo distribution shift that manifests not in the input space but in the conditional representation space $P(\mathcal{Z}|\mathcal{Y}^*)$. This shift disrupts the geometric coherence of latent class manifolds and ultimately impairs generalization performance.

These findings **motivate the core design** of MgEL: to suppress the emergence and propagation of noise-induced structural drift by fundamentally rethinking how label supervision is incorporated during training. Rather than relying solely on potentially corrupted labels or unstable model predictions, MgEL introduces controlled structural uncertainty and diversified supervisory signals to counteract the convergence toward biased representations. This is achieved not by architectural overhauls or post hoc corrections, but by embedding uncertainty-aware regularization into the label construction process itself. By doing so, MgEL aims to retain the semantic integrity of the class manifolds even under extreme noise, thereby preserving the model's capacity for generalization.

3 Methodology

3.1 Design rationale of MgEL

DNNs with strong generalization capabilities tend to induce class-discriminative clustering structures in the latent feature space [16], where the similarity between two sample embeddings reflects the likelihood that they share the same true class. Empirical observations confirm [40] that the closer two representations $z_i, z_j \in \mathcal{Z}$ are, the more likely $y_i = y_j$.

Assumption 2 (Label consistency) For $\forall x_i, x_j \in \mathcal{X}$, if their feature similarity $\mathcal{S}(z_i, z_j)$ approaches the self-similarity limit, i.e., $\mathcal{S}(z_i, z_j) \rightarrow \mathcal{S}(z_i, z_i) = \mathcal{S}(z_j, z_j)$,

the probability of label consistency increases:

$$P(\hat{y}_i = \hat{y}_j) \rightarrow 1 \quad \propto \quad \mathcal{S}(z_i, z_j) \rightarrow \mathcal{S}(z_i, z_i) = \mathcal{S}(z_j, z_j)$$

Motivated by Assumption 2, MgEL (see Figure 3) draws inspiration from the observation-collapse process in quantum entangled systems. To simulate such a system, the latent feature space $\mathcal{Z}$ is randomly divided into two disjoint subsets, $\mathcal{Z}_A$ and $\mathcal{Z}_B$. For a sample $x_i \in \mathcal{Z}_A$, a feature cluster $\mathcal{C}_i \subseteq \mathcal{Z}_B$ is formed by retrieving its top- n most similar neighbors, each associated with the observed annotation y_j . These cluster members are treated as independent observations of x_i under different measurement conditions, with their label distribution abstracted as a probabilistic superposition over possible class outcomes.

Through these local observations, MgEL simulates multi-view fusion from pseudo sources by leveraging the structured observations of these clusters (see Figure 4). Specifically, for each sample x_i , a pseudo-label $y'_i \in \mathcal{Y}'$ is derived from the distribution

of labels in its feature cluster. This pseudo-label is then fused with the original annotation $y_i \in \mathcal{Y}$ to create an evidence-aware multi-granularity label $\tilde{y}_i \in \tilde{\mathcal{Y}}$. Based on the principles of Dempster-Shafer theory (DST) [39], we interpret the label distribution in the feature cluster as a class-supporting evidence mass function. The fusion process assigns different weights to the pseudo and observed labels, reflecting their respective levels of credibility under uncertainty.

To further enhance interpretability and control, MgEL categorizes the fused label $\tilde{y}_i$ into two levels of semantic granularity based on the consistency between the pseudo-label y'_i and the annotation y_i . When the two labels are identical, the fused result is interpreted as a fine-grained confident label, reflecting strong agreement across sources and indicating high certainty in the class assignment. In contrast, if the pseudo and original labels disagree, the resulting label is considered a coarse-grained hesitant label, capturing ambiguity in the sample's class attribution and preserving the possibility of it belonging to multiple candidate classes.

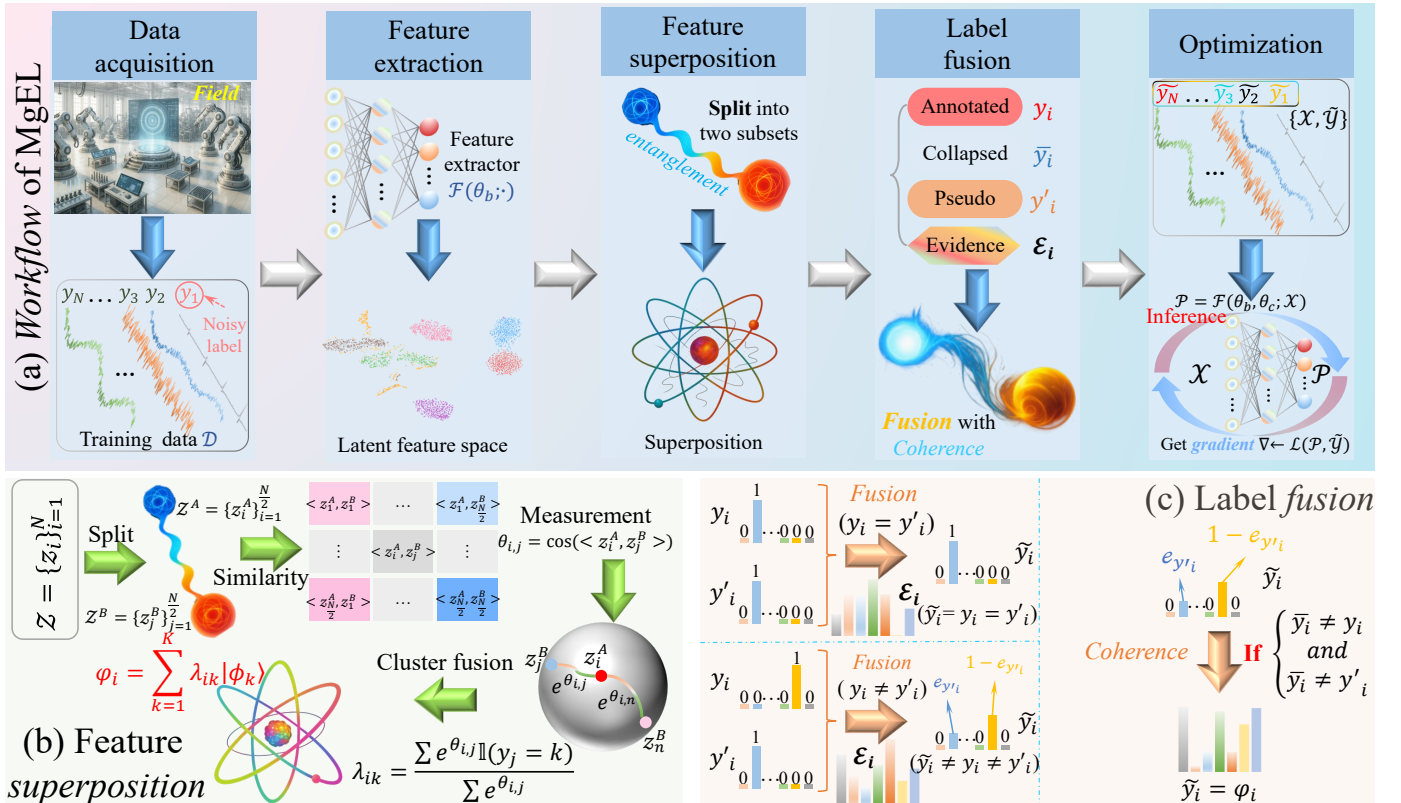


Figure 3. Conceptual workflow: The MgEL establishes a quantum-inspired label fusion mechanism by simulating the observation-collapse process inherent to entangled systems. Given a noisy dataset $\mathcal{D} = \{(x_i, y_i)\}_{i=1}^N$, features $\mathcal{Z}$ are extracted via a backbone encoder $\mathcal{F}(\theta_b; \cdot)$, and the dataset is randomly split to mimic entangled subsystems. For each sample, a feature cluster $\mathcal{C}_i$ is constructed from its most similar counterparts in the other subset and abstracted as a superposition state φ_i . From this, a pseudo-label y'_i , a collapsed label $\bar{y}_i$, and an evidence $\mathcal{E}_i$ are derived. These are fused with the original annotation y_i to form a multi-granularity label $\tilde{y}_i$. A Coherence mechanism assesses label consistency to decide whether to reintegrate $\tilde{y}_i$ into φ_i for mitigating confirmation bias. The final evidence labels $\tilde{\mathcal{Y}}$ is used to supervise training, improving generalization under severe noise.

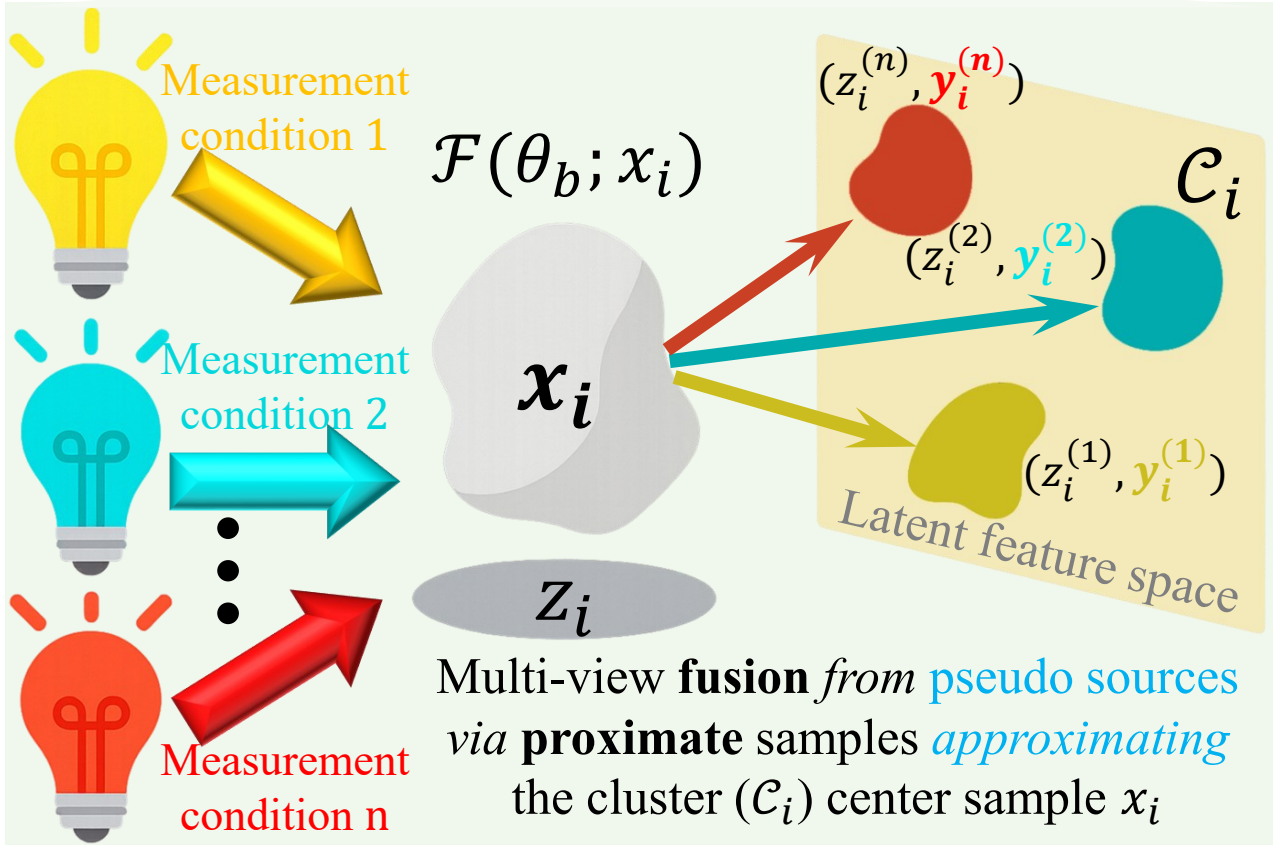
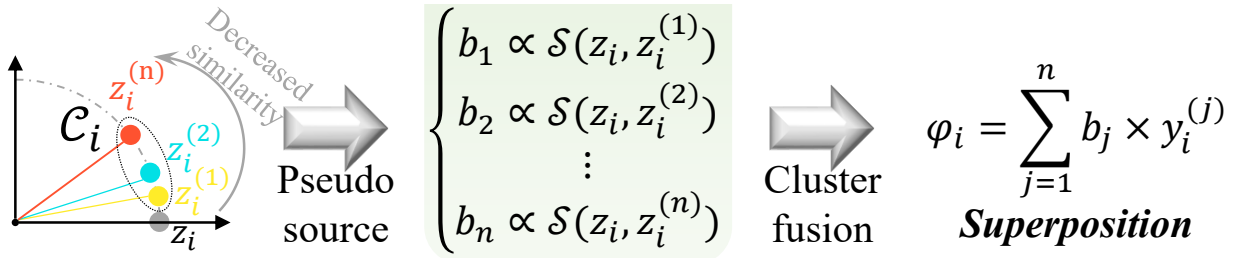


Figure 4. Illustration of multi-view fusion from pseudo sources in MgEL: This method simulates multi-view observations by leveraging structured feature clusters, where each cluster $\mathcal{C}_i$ is approximated by a set of proxy samples. The belief b_j of each pseudo-source $z_i^{(j)}$ is assigned based on its similarity $\mathcal{S}(z_i, z_i^{(j)})$ to the cluster center z_i , simulating multiple independent observations. These pseudo sources are fused by aggregating the evidence from belief-weighted samples, effectively simulating multi-source information fusion.

Since this correction is based on model-generated latent representations and predictions, it risks amplifying confirmation bias during self-guided training. To mitigate this, MgEL introduces a coherence mechanism that assesses the consistency among multiple label sources (pseudo, original, collapsed). When inconsistency is high, the fused label $\tilde{y}_i$ is reverted to its superposition state φ_i , preserving uncertainty and deferring final decisions until further optimization.

Overall, MgEL reinterprets noisy or conflicting labels as structured evidence, offering a principled framework to improve robustness by aligning

quantum-inspired modeling with uncertainty-aware label correction.

3.2 Constructing feature clusters via entangled partitioning

Consider a noisy dataset $\mathcal{D} = \{(x_i, y_i)\}_{i=1}^N$ with latent representations $\mathcal{Z} = \{z_i\}_{i=1}^N$ extracted using a backbone network $\mathcal{F}(\theta_b; \cdot)$.

Drawing upon quantum collapse dynamics and superposition principles, we establish a pseudo multi-source label fusion framework where each latent representation z_i models an analogous quantum state. Central to this approach is the

entanglement-emulating partition: $\mathcal{Z}$ undergoes random bisection into disjoint subsets $\mathcal{Z}_A$ and $\mathcal{Z}_B$ of equal cardinality. This bipartite configuration achieves two objectives: reducing computational complexity of similarity tensors from $\mathcal{O}(N^2)$ to $\mathcal{O}(N^2/4)$ and implementing stochastic structural dropout.

Furthermore, each subset acts as an independent measurement apparatus, simulating the observational asymmetry inherent in quantum systems. The bipartite design is a deliberate computational compromise. While multi-partite partitioning ($n > 2$) incurs prohibitive $\mathcal{O}((N/n)^n)$ complexity and generates intractable similarity tensors, the current formulation maintains theoretical fidelity while ensuring computational tractability.

For each sample x_i , neighborhood retrieval is performed exclusively from the complementary subset. These neighbors $\{z_i^{(j)}\}_{j=1}^n$ are not merely correlated instances but simulate independent measurements of the underlying quantum state. Each neighbor thus represents an eigen-label state $|k\rangle$ (class k collapse) under different measurement contexts. This quantum interpretation underpins the formulation of the superposed label state:

$$\varphi_i = \sum_{k=1}^K \lambda_{ik} |k\rangle \quad (10)$$

where $\varphi_i \in \mathbb{R}^K$ encodes amplitude-based label uncertainty, with coefficients $\lambda_{ik} \in [0, 1]$ quantifying class evidence strength (see Sec. 3.3). Neighborhood retrieval uses exponentially transformed cosine similarity to circumvent the curse of dimensionality.

Based on Assumption 2, for $\forall x_i \in \mathcal{X}$, with latent representation $z_i = \mathcal{F}(\theta_b; x_i)$, we identify its n most similar feature vectors from the entangled subset to form a feature cluster $\mathcal{C}_i$. These neighbors $\{z_i^{(j)}\}_{j=1}^n$ are ranked such that

$$\mathcal{S}(z_i, z_i^{(i+1)}) > \mathcal{S}(z_i, z_i^{(j)}) \text{ if } j \geq 1 \quad (11)$$

To quantify sample similarity, we use an exponential cosine similarity function (see Eq. (17)). This decision is motivated by the well-known phenomenon of distance concentration in high-dimensional spaces, which makes traditional distance metrics, such as Euclidean distance, ineffective. Specifically, for two normalized vectors $z_i, z_j \in \mathbb{R}^n$, the squared Euclidean distance is defined as:

$$d^2(z_i, z_j) = \|z_i - z_j\|^2 = \|z_i\|^2 + \|z_j\|^2 - 2\langle z_i, z_j \rangle \quad (12)$$

Assuming $\|z_i\| = \|z_j\| = 1$, we have:

$$d^2(z_i, z_j) = 2 - 2\cos(\langle z_i, z_j \rangle), \quad (13)$$

where $\langle z_i, z_j \rangle$ is the angle between z_i and z_j . As dimensionality $n \rightarrow \infty$, most angles $\langle z_i, z_j \rangle \rightarrow \frac{\pi}{2}$, and hence $\cos(\langle z_i, z_j \rangle) \rightarrow 0$, implying:

$$d^2(z_i, z_j) \rightarrow 2, \quad (14)$$

which leads to the so-called concentration of measure, where almost all pairwise distances converge to a constant. Formally,

$$\mathbb{P}(|d(z_i, z_j) - \mu_d| < \varepsilon) \rightarrow 1 \text{ as } n \rightarrow \infty, \quad (15)$$

making it nearly impossible to discriminate between samples based on Euclidean distance.

To avoid this, we use the cosine similarity:

$$\mathcal{S}_{\cos}(\langle z_i, z_j \rangle) = \cos(\langle z_i, z_j \rangle) = \frac{z_i \times z_j^\top}{\|z_i\| \cdot \|z_j\|}, \quad (16)$$

which measures angular similarity and is less sensitive to magnitude and dimensionality. Still, in high dimensions, even cosine similarities between neighbors may vary only slightly. To accentuate such differences, we introduce the exponential transformation:

$$\mathcal{S}(z_i, z_j) = e^{\cos(\langle z_i, z_j \rangle)} = \exp\left(\frac{z_i \times z_j^\top}{\|z_i\| \cdot \|z_j\|}\right) \quad (17)$$

This maps cosine similarity from $[-1, 1]$ to $[e^{-1}, e^1]$, i.e. $\mathcal{S}(z_i, z_j) \in [e^{-1}, e^1] \approx [0.37, 2.72]$, nonlinearly amplifying differences between close neighbors and enabling sharper separation of structurally similar samples in latent space.

Based on this similarity, the feature cluster $\mathcal{C}_i$ is formally defined as:

$$\mathcal{C}_i = \begin{cases} \{z_j \mid \mathcal{S}(z_i, z_j) \geq \mathcal{S}(z_i, z_i^{(n)})\}, z_j \in \mathcal{Z}_A & \text{if } z_i \in \mathcal{Z}_B, \\ \{z_j \mid \mathcal{S}(z_i, z_j) \geq \mathcal{S}(z_i, z_i^{(n)})\}, z_j \in \mathcal{Z}_B & \text{otherwise.} \end{cases} \quad (18)$$

The cluster size n is dynamically adapted based on the estimated noise intensity $\tilde{\eta}$ (see Sec. 3.4) from the previous training epoch:

$$n = \max(\tilde{\eta} \times \text{Initialize}(n), \text{Truncation}(n)) \quad (19)$$

In this work, $\text{Initialize}(n) = 2^7$ and $\text{Truncation}(n) = 2^3$. This allows the framework to flexibly adjust the neighborhood scope: under high noise, more neighbors are included to stabilize superposition formation; under low noise, smaller clusters avoid introducing irrelevant variance.

3.3 Simulating pseudo-source observations via cluster fusion

Cluster fusion simulates the integration of multiple pseudo-observations by leveraging the feature and label distributions within the feature cluster $\mathcal{C}_i$ (as constructed in Eq. (18)). This process is analogous to quantum state tomography [42] in quantum mechanics, where the quantum state is reconstructed by sampling probability distributions across different measurement bases, determining the amplitude probabilities of the eigenstates in a superposition.

In MgEL, the superposition state φ_i constructed for each sample x_i can be interpreted as the probability distribution of x_i 's class membership across all potential categories. The associated belief mass $\mathcal{E}_i = \{e_k\}_{k=1}^K$ encodes the basic belief assignment for each class, which is used to fuse the pseudo-label and the original annotation during label correction, considering the current dataset's noise level.

To improve the robustness and accuracy of the superposition state φ_i , MgEL adopts the evidence combination concept from DST. Each sample in the feature cluster $\mathcal{C}_i$ is treated as an independent source of evidence, and a cluster-based fusion process based on similarity is designed:

$$\begin{cases} \varphi_i &= \sum_{k=1}^K \lambda_{ik} \cdot |\phi_k\rangle, \\ \lambda_{ik} &= \frac{\sum_{z_j \in \mathcal{C}_i} \mathcal{S}(z_i, z_j) \cdot \mathbb{1}(y_j = k)}{\sum_{z_j \in \mathcal{C}_i} \mathcal{S}(z_i, z_j)}, \\ |\phi_k\rangle &= \text{OneHot}(k), \quad k \in [1, K]. \end{cases} \quad (20)$$

where K represents the total number of classes, and $\mathbb{1}(y_j = k)$ is an indicator function that equals 1 when $y_j = k$, and 0 otherwise.

Using this similarity-based fusion method, MgEL combines the feature distribution and label frequency

information within the feature cluster $\mathcal{C}_i$ to form the superposition state φ_i . This enables MgEL to robustly model the class membership probability distribution for sample x_i , even in noisy and uncertain environments.

However, considering the potential misleading effects of noisy labels, there is a risk that the feature mapping during training may inaccurately represent the data, leading to imprecise class probabilities in φ_i . If these probabilities were directly used as fusion belief masses for label correction, they could amplify errors and introduce confirmation bias.

To address this, MgEL does not use φ_i 's class probabilities directly as belief masses in the fusion of annotations and pseudo-labels. Instead, after considering the estimated noise intensity $\tilde{\eta}$, MgEL adjusts the evidence $\mathcal{E}_i$ for each class membership distribution:

$$e_k = \frac{\sum_{z_j \in \mathcal{C}_i} \mathcal{S}(z_i, z_j) \cdot \mathbb{1}(y_j = k)}{\sum_{z_j \in \mathcal{C}_i} \mathcal{S}(z_i, z_j) + K \times \tilde{\eta}}. \quad (21)$$

3.4 Multi-granularity labels construction with evidence awareness

After obtaining the superposition state φ_i and class evidence $\mathcal{E}_i$ through feature cluster fusion, MgEL generates pseudo-labels y'_i using the maximum posterior decision rule (commonly applied in multi-class classification tasks, e.g. [1]). The amplitude probability $\alpha_{y'_i}$ associated with the pseudo-label represents the frequency or intra-cluster label consistency of the dominant class within the feature cluster. This is used to estimate the noise intensity of the current dataset by calculating the mean label consistency $\alpha_{y'}$ across all feature clusters:

$$\begin{cases} y'_i &= \arg \max(\varphi_i), \\ \alpha_{y'_i} &= |y'_i\rangle \times \varphi_i^\top = \max(\varphi_i), \\ \tilde{\eta} &= \frac{\sum_{i=1}^N \alpha_{y'_i}}{N}. \end{cases} \quad (22)$$

At this stage, both the pseudo-label and the original annotation label serve as independent evidence sources for making decisions about sample x_i . MgEL then fuses these two sources to correct the noisy labels, based on their respective confidence levels. Specifically, if the pseudo-label and original annotation labels agree, it is highly likely that the original annotation is not noisy and can be used for training. In contrast, when the two labels disagree, the pseudo-label is more likely

to be accurate. The confidence in the decision-making process is captured by the evidence vector $\mathcal{E}_i$, which is derived from the pseudo-source fusion.

This allows the original annotation label to be corrected toward the pseudo-label, with a weighted confidence:

$$\tilde{y}_i = \left(|y'_i\rangle \times \mathcal{E}_i^\top \right) \cdot y'_i + \left(1 - |y'_i\rangle \times \mathcal{E}_i^\top \right) \cdot y_i \quad (23)$$

The fused multi-granularity labels $\tilde{\mathcal{Y}}$ integrate both fine-grained and coarse-grained information, reflecting the model's certainty in the predicted class. When the pseudo-label and the original annotation agree, the label corresponds to a fine-grained confident label, offering a more precise supervision signal. When the labels disagree, the output represents a coarse-grained hesitant label, with the decision uncertain between two potential classes, and the confidence distributed between them according to the evidence mass $\mathcal{E}_i$. This distinction allows the model to dynamically adjust its decision-making process based on the consistency of the label sources.

3.5 Coherence mechanism triggered by consistency of labels

To mitigate the confirmation bias that may arise from directly using pseudo-labels to correct original annotations, MgEL employs a resampling strategy on the superposition state φ_i , simulating a random collapse process. This generates a collapsed label $\bar{y}_i$, where the sampling probability for each class corresponds to the amplitude probability of the corresponding eigenstate in φ_i .

Subsequently, MgEL introduces a Coherence mechanism triggered by label consistency. Before performing supervised training, the fused labels $\tilde{\mathcal{Y}}$ go through an additional processing step. In this step, the system accepts a fused label only when at least two of the following labels are consistent: the original annotation y_i , the pseudo-label y'_i , and the collapsed label $\bar{y}_i$. If none of these labels are consistent, the system retains the superposition state φ_i , deferring the final label commitment until more evidence is provided in subsequent training iterations.

The update rule for the fused label is formally defined as follows:

$$\tilde{y}_i = \begin{cases} \varphi_i & , \text{ if } \||y'_i\rangle + |y_i\rangle + |\bar{y}_i\rangle\| = \sqrt{3}, \\ \tilde{y}_i & , \text{ otherwise.} \end{cases} \quad (24)$$

In this rule, $\|\cdot\|$ represents the norm of the vector formed by the sum of the three labels, and the condition ensures that the fused label $\tilde{y}_i$ is accepted when at least two of the three labels are consistent.

Through this process, MgEL combines the statistical consistency of the pseudo-labels, the randomness introduced by the collapsed labels, and the original annotation labels to create more robust labels $\tilde{\mathcal{Y}}$. This approach reduces the confirmation bias typically associated with prediction-based label correction and improves the model's generalization ability, especially in high-noise environments.

3.6 Overview

The Alg. 1 summarizes the key steps of MgEL, integrating feature clustering, pseudo-label generation, and evidence-based label fusion for label correction, simulating quantum-inspired mechanisms and adapting them for noisy label correction.

Algorithm 1 MgEL Label Correction Process

Input: Noisy dataset $\mathcal{D} = \{\mathcal{X}, \mathcal{Y}\} = \{(x_i, y_i)\}_{i=1}^N$, model $\mathcal{F}(\theta_b, \theta_c; \cdot)$, estimated noise intensity $\tilde{\eta}$
Output: $\mathcal{F}(\theta_b, \theta_c; \cdot)$ # Trained diagnostic model
Initialize model $f(\theta_b, \theta_c; \cdot)$, $\tilde{\eta} = 1.0$
for epoch in range(Total Epochs) **do**
 $\mathcal{Z} = \mathcal{F}(\theta_b; \mathcal{X})$ # Extract latent features
 $\mathcal{Z}_A, \mathcal{Z}_B = \text{Split}(\mathcal{Z})$ # Partition dataset into two subsets
 $\mathcal{C}_A = \{\mathcal{C}_i\}_{i=1}^{\frac{N}{2}} \leftarrow \tilde{\eta}, \mathcal{Z}_B, \mathcal{C}_B = \{\mathcal{C}_i\}_{i=1}^{\frac{N}{2}} \leftarrow \tilde{\eta}, \mathcal{Z}_A$ # Construct feature clusters via Eq. (18)
 $\varphi_A = \{\varphi_i\}_{i=1}^{\frac{N}{2}} \leftarrow \mathcal{C}_A, \varphi_B = \{\varphi_i\}_{i=1}^{\frac{N}{2}} \leftarrow \mathcal{C}_B$ # Compute superposition state via Eq. (20)
 $\{\mathcal{E}_i\}_{i=1}^N \leftarrow \varphi_A \cup \varphi_B$ # Get class evidence via Eq. (21)
 $\tilde{\eta}, \mathcal{Y}' \leftarrow \varphi_A \cup \varphi_B$ # Estimated noise intensity and compute pseudo-labels via Eq. (22)
 $\tilde{\mathcal{Y}} \leftarrow \{\mathcal{E}_i\}_{i=1}^N, \mathcal{Y}', \mathcal{Y}$ # Construct fused labels via Eq. (23)
 $\bar{\mathcal{Y}} \leftarrow \varphi_A \cup \varphi_B$ # Generate collapsed labels via multinomial sampling
 $\bar{\mathcal{Y}} \leftarrow \varphi_A \cup \varphi_B, \bar{\mathcal{Y}}, \mathcal{Y}', \mathcal{Y}$ # Apply Coherence Mechanism for final multi-granularity labels via Eq. (24)
 for each sample $(x_i, \tilde{y}_i)$ in $\tilde{\mathcal{D}}$ **do**
 $p_i = f(\theta_b, \theta_c; x_i)$ # Compute prediction
 $(\theta_b, \theta_c) \leftarrow \text{Update}(\theta_b, \theta_c, \nabla \mathcal{L}(p_i, \tilde{y}_i))$ # Update model parameters via fused label $\tilde{y}_i$
 end for
end for

Considering that MgEL uses exponential cosine similarity (Eq. (17)) to measure sample similarity in the latent feature space, we have aligned the classifier architecture accordingly. Instead of adding a fully connected layer to the backbone network for classification, we employ a cosine classifier to directly

Table 2. The detailed information about all datasets.

Parameter	Detailed information about three datasets		
	<i>Single</i>	<i>Multiple</i>	<i>Damage</i>
No. of categories	$6 + 1 = 7$	$(6 + 1)^2 = 49$	$4 \times 6 + 1 = 25$
Fault location	Inner, Outer, Ball, Inner and Outer, Inner and Ball, Outer and Ball		
Damage degree	0.3 mm	0.3 mm	0.2, 0.4, 0.6, and 0.8 mm
Motor speed	[1443, 1478] rpm	[1770, 1775] rpm	[1770, 1775] rpm
		[2366, 2370] rpm	[2366, 2370] rpm
		[2959, 2962] rpm	[2959, 2962] rpm
Motor load	0.0, 0.1, 0.2 and 0.3 NM	0.0 and 0.3 NM	0.0 and 0.3 NM
Sampling frequency	12 kHz	12 kHz	16 kHz
No. of tra. samples	4032	46452	34100
No. of val. samples	924	10290	7600
No. of tes. samples	4032	46452	34100

classify the features, which ensures better alignment with the feature space's geometry and optimization direction, as cosine similarity is more suitable for high-dimensional feature manifolds.

It is *important to emphasize* that the proposed MgEL framework operates exclusively during the training phase. It acts as a robust label correction strategy by adjusting supervisory signals based on latent-space evidence fusion. Consequently, the inference pipeline remains entirely unaltered—both structurally and computationally. MgEL introduces no additional latency, memory, or computational overhead at deployment time. The final diagnostic model inherits its real-time performance characteristics solely from the underlying backbone architecture, making MgEL fully compatible with industrial runtime constraints.

4 Experimental verification and discussion

4.1 Experimental settings

To thoroughly evaluate the performance of MgEL, we conducted a series of experiments on three benchmark datasets for bearing fault diagnosis: *Single* [44], *Multiple* [43], and *Damage* [40]. These datasets consist of vibration signals collected directly from mechanical systems during operation, covering various working conditions, rotational speeds, and load conditions. These datasets were selected to represent a wide range of fault scenarios and complexities, ensuring that MgEL can handle diverse real-world conditions effectively.

Table 2 summarizes the details of each dataset, including fault categories, operating conditions, and the number of samples. Notably, the datasets were

chosen without addressing class imbalance, ensuring that the number of training and testing samples for each class is roughly equal. Each sample consists of 2048 continuous data points, with no overlap between any two samples, ensuring a clean and consistent evaluation framework.

With the goal of conducting a comprehensive evaluation of MgEL under various noise environments, two types of noise were introduced: class-dependent *asymmetric* noise and class-independent *symmetric* noise. Asymmetric noise was introduced by swapping labels between similar classes at five different intensity levels ($\eta = 25\%, 30\%, 35\%, 40\%, 45\%$). For example, in the *Damage* dataset, severity labels within the same fault category were swapped, such as replacing the "0.2 mm" label with "0.4 mm" or swapping "0.8 mm" and "0.6 mm". This type of noise simulates realistic scenarios in which label errors occur within similar categories. Symmetric noise was introduced by randomly replacing the original labels with labels from other categories, with five different noise intensities applied ($\eta = 50\%, 60\%, 70\%, 80\%, 90\%$).

To ensure a fair comparison and reproducibility, we adopt MgNet [43], an open-source fault diagnosis architecture released alongside the *Multiple* dataset and well-suited for time-series analysis in industrial scenarios, as the backbone for evaluating the performance of MgEL. All models were trained using the AdamW optimizer, with momentum coefficients $\beta_1 = 0.9$, $\beta_2 = 0.999$, and weight decay 10^{-2} . The learning rate follows a CosineAnnealing schedule, with the initial value scaled as $3 \times 10^{-3} \times \frac{\text{batch size}}{512}$, ensuring consistency across different training scales. Training is performed for 100 epochs, including a

Table 3. Comparison of effectiveness between MgEL and other methods for LNL.

Datasets	Methods	Accuracy(%) under <i>asymmetric</i> noise					Accuracy(%) under <i>symmetric</i> noise					Mean (%)
		$\eta = 25\%$	$\eta = 30\%$	$\eta = 35\%$	$\eta = 40\%$	$\eta = 45\%$	$\eta = 50\%$	$\eta = 60\%$	$\eta = 70\%$	$\eta = 80\%$	$\eta = 90\%$	
Single	Baseline	87.93±5.51	83.16±4.79	75.88±1.97	69.42±3.47	59.26±0.10	68.12±4.96	58.04±6.44	45.60±5.87	29.04±9.35	5.37±3.79	58.18
	SL (2019)	95.80±0.42	86.35±3.15	85.61±7.28	72.60±3.34	61.06±1.38	81.66±6.66	73.14±8.38	58.99±1.20	25.31±5.43	3.38±1.35	64.39
	JoCoR (2020)	88.51±0.17	89.71±7.14	81.80±8.32	73.87±4.63	60.44±0.81	79.37±2.00	63.38±5.80	53.36±4.41	34.43±1.30	7.42±1.62	63.23
	JNPL (2021)	64.66±9.98	54.91±3.24	54.22±4.37	36.95±7.29	43.46±4.35	28.80±7.35	29.21±5.07	27.98±1.28	21.73±9.13	8.86±4.39	37.08
	NCR (2022)	97.54±1.54	90.09±3.63	86.09±6.91	81.73±2.12	75.11±6.49	84.25±1.31	76.59±5.17	68.97±7.78	49.63±2.68	5.03±0.38	71.50
	ALFs (2023)	89.16±7.78	88.08±7.80	83.60±2.93	69.47±1.07	59.24±5.62	90.68±0.47	86.87±2.08	73.94±2.46	53.65±2.96	5.11±2.26	69.98
	LSL (2024)	96.24±3.03	95.28±4.32	94.73±3.82	82.56±6.03	70.14±9.81	95.33±3.98	93.34±0.57	82.30±9.58	64.83±3.16	6.73±3.39	78.15
	ANNE (2025)	91.10±7.84	89.32±8.27	83.79±1.22	75.15±5.49	69.50±7.78	98.74±1.21	97.43±1.57	95.47±0.57	73.17±8.33	4.44±1.13	77.81
	MgEL (Ours)	99.49±0.38	99.19±0.35	99.09±0.63	96.06±2.58	80.89±4.69	99.55±0.17	98.92±0.77	98.25±0.41	80.46±3.19	23.02±2.53	87.49
Multiple	Baseline	75.07±3.78	69.21±3.16	64.11±3.13	55.04±0.27	51.49±0.92	53.00±2.29	48.36±1.99	36.36±3.22	23.03±0.44	7.39±0.69	48.31
	SL (2019)	80.40±1.85	75.52±0.43	71.80±2.82	53.57±1.73	49.67±0.84	72.56±1.22	62.51±2.54	58.29±2.29	34.38±2.62	9.34±0.53	56.80
	JoCoR (2020)	88.47±1.14	88.70±0.68	87.31±1.53	72.50±4.18	64.88±1.23	74.13±3.67	66.38±0.98	49.36±2.57	31.64±1.03	9.97±1.43	63.33
	JNPL (2021)	63.59±0.24	63.45±0.19	61.66±1.21	60.12±1.02	55.93±1.67	50.45±0.33	41.54±2.07	33.78±1.27	23.24±1.28	9.57±1.37	46.33
	NCR (2022)	90.25±0.57	89.34±0.26	86.66±1.72	82.49±0.85	69.89±1.48	66.50±3.83	55.60±2.50	43.01±1.53	26.52±3.07	9.07±1.69	61.93
	ALFs (2023)	77.04±1.50	73.96±1.75	68.02±0.74	64.00±3.36	54.59±1.65	65.27±2.76	52.35±0.61	32.91±0.69	22.92±0.73	8.40±1.39	51.95
	LSL (2024)	64.06±1.57	62.43±4.91	66.19±1.33	59.88±1.83	55.32±0.07	47.67±0.83	42.58±2.20	34.91±1.08	24.05±1.66	8.72±0.99	46.58
	ANNE (2025)	65.55±2.49	67.69±2.57	63.46±1.44	63.60±1.34	57.80±1.04	52.88±2.58	51.15±4.12	39.90±4.36	23.58±1.72	11.10±0.51	49.67
	MgEL (Ours)	89.33±0.52	89.16±0.87	87.70±1.15	86.20±0.51	75.65±1.50	86.20±1.28	83.49±1.16	81.45±0.52	75.46±1.80	43.18±5.46	79.78
Damage	Baseline	91.18±0.17	86.31±1.18	82.38±1.10	67.28±3.45	54.57±2.16	88.78±0.48	84.66±1.02	76.72±0.30	55.24±1.62	20.18±3.84	70.73
	SL (2019)	93.83±0.42	90.53±2.13	88.04±1.00	75.70±4.60	57.45±0.45	94.55±0.17	92.51±0.19	87.88±1.48	76.11±1.50	20.86±1.74	77.74
	JoCoR (2020)	93.84±0.19	89.22±1.72	85.28±0.81	70.00±3.85	56.66±1.52	91.78±0.64	88.94±1.25	83.64±1.32	62.57±1.67	22.41±2.13	74.43
	JNPL (2021)	90.18±0.05	86.10±0.51	81.11±0.99	69.62±0.76	57.32±1.66	87.56±0.93	84.95±0.34	76.17±2.60	52.13±4.56	24.30±2.23	70.94
	NCR (2022)	94.72±0.64	92.00±0.30	86.24±1.44	75.71±2.07	61.99±1.99	91.75±0.77	89.44±2.55	83.49±0.74	61.45±0.75	18.20±1.09	75.50
	ALFs (2023)	96.93±0.62	94.96±1.41	90.11±2.67	78.45±2.18	59.09±0.83	96.03±0.62	94.01±1.71	87.32±2.06	62.74±0.85	19.77±1.66	77.94
	LSL (2024)	90.09±0.84	86.56±0.44	81.67±0.99	70.83±0.72	57.55±1.07	86.76±1.06	83.61±1.56	75.46±1.08	54.41±4.17	20.28±1.20	70.72
	ANNE (2025)	96.63±0.18	96.50±0.23	92.75±0.31	87.08±0.45	64.41±6.91	97.36±0.65	96.99±0.26	90.41±0.41	83.42±3.21	33.66±8.18	83.92
	MgEL (Ours)	97.48±0.42	97.19±0.45	96.70±0.87	93.25±0.65	75.03±1.01	97.31±0.22	96.47±0.34	95.81±0.27	94.31±0.35	75.07±2.35	91.86

5-epoch warm-up stage, during which only standard supervised learning is applied, without any additional label correction operations.

All experiments were conducted on an NVIDIA A100 GPU using *PyTorch* version 2.1.1+cu118. To eliminate the influence of random initialization and stochastic variation, each experiment was repeated at least ten times, with the detailed evaluation results and corresponding analyses presented in the following subsections.

4.2 Case I: Effectiveness of MgEL

To evaluate the effectiveness of the proposed MgEL framework, we conducted comprehensive comparisons with seven state-of-the-art LNL approaches: SL (2019) [30], JoCoR (2020) [33], JNPL (2021) [45], NCR (2022) [46], ALFs (2023) [27], LSL (2024) [16], and ANNE (2025) [47], across three datasets with varying types and intensities of label noise. We visualize the performance comparison in Figure 5 and provide the detailed results in Table 3, where bold values represent the best accuracy achieved by other methods under each noise setting, facilitating an intuitive comparison with MgEL.

Overall, MgEL consistently outperforms all competing approaches across different datasets and noise scenarios, demonstrating superior robustness and generalization even in the presence of severe label corruption. On the *Single* dataset, MgEL achieves a

mean accuracy of 87.49%, surpassing the strongest baseline (LSL, 78.15%) by 9.34%. This advantage is amplified further under high symmetric noise ($\eta = 90\%$), where MgEL maintains a robust 23.02% accuracy compared to 7.42% for JoCoR. A similar pattern is observed on the *Multiple* dataset, where MgEL yields an average accuracy of 79.78%, outperforming the second-best baseline (JoCoR) by 16.45%. On the more challenging *Damage* dataset, MgEL delivers the highest performance (91.86%), surpassing ANNE (83.92%) by 7.94%.

These consistent improvements across datasets and noise conditions are attributed to several technical innovations in MgEL. One key reason is that, instead of relying on heuristic sample selection or loss-based reweighting, MgEL constructs a quantum-inspired superposition state φ_i for each sample by aggregating local observations from entangled feature clusters. This formulation captures class attribution uncertainty with greater fidelity. Additionally, MgEL fuses the pseudo-label and original annotation via an evidence-aware belief mass $\mathcal{E}_i$, derived from the distribution of cluster-level statistics. This enables adaptive weighting of conflicting information sources based on their consistency. Moreover, the Coherence mechanism safeguards against early confirmation bias by retaining the superposition state whenever the pseudo, annotation, and collapsed labels fail to reach consensus, postponing hard commitments

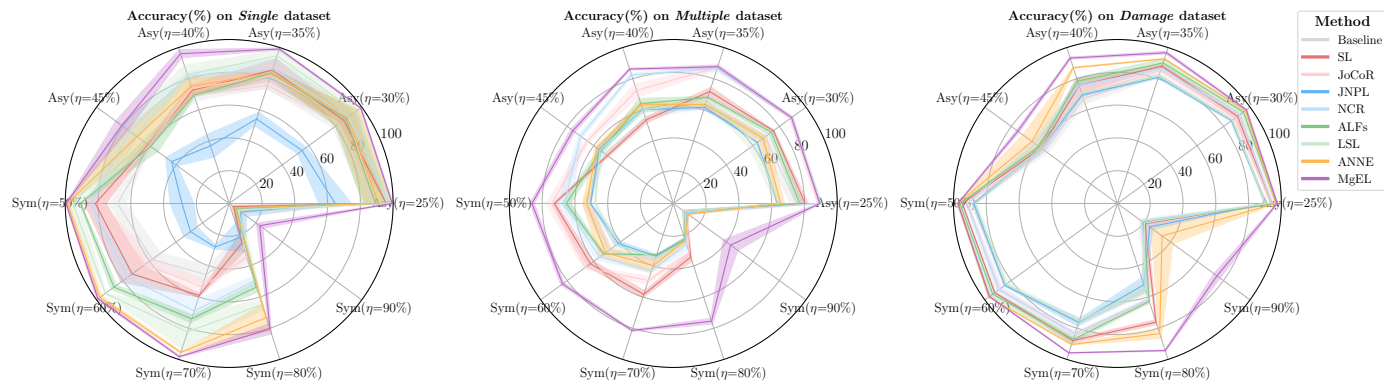


Figure 5. Comparison of MgEL and other methods under various label noise environments. Where the shaded area indicates the accuracy fluctuation range of each method, and the dashed baseline represents the performance of MgEL trained with standard supervised learning, without any LNL strategy.

Table 4. The results of the ablation on MgEL.

Methods	Entanglement	Evidence	Coherence	Average accuracy(%)		
				Single	Multiple	Damage
Baseline				58.18	48.31	70.43
A1		✓	✓	85.64	76.89	90.22
A2	✓		✓	84.61	75.58	89.12
A3	✓	✓		82.09	72.95	87.74
A4	✓			79.41	70.16	86.09
A5		✓		81.62	72.33	87.67
A6			✓	83.88	74.34	88.78
MgEL	✓	✓	✓	87.49	79.78	91.86

until sufficient evidence accumulates in later training iterations.

The performance advantage of MgEL becomes especially prominent under high-noise symmetric settings, indicating its tolerance against purely stochastic label perturbations that often mislead traditional LNL strategies. Moreover, its stability across diverse datasets highlights its generalizability and scalability for real-world fault diagnosis tasks.

4.3 Case II: Ablation for MgEL

To further investigate the source of MgEL’s effectiveness, we conduct an ablation study by selectively removing each of its three core components: *Entanglement*, *Evidence*, and *Coherence*. The results are summarized in Table 4 and visualized in Figure 6.

Taken together, Table 4 and Figure 6 show that the full MgEL framework, which integrates all three components, achieves the highest accuracy across all datasets. Removing any component consistently leads to a decline in performance, indicating that the effectiveness of MgEL stems from the complementary

roles of entanglement-based pseudo-source modeling, evidence-aware label fusion, and coherence-guided filtering.

The effect of removing Entanglement is evident when comparing configurations A1 and A4 to the full model. A1, which disables Entanglement while retaining Evidence and Coherence, shows a performance decline (e.g., 76.89% *vs.* 79.78% on *Multiple*). The degradation becomes more pronounced in A4, where only Entanglement is active, leading to further reductions in accuracy (e.g., 70.16%). These trends suggest that Entanglement is critical for constructing reliable pseudo-source representations. Its role in partitioning the feature space introduces stochastic diversity in cluster formation, capturing the distributional uncertainty of class attribution through superposition states, which is essential for robust label inference.

The role of Evidence is assessed by examining A2 and A5, both of which remove this component. A2 maintains Entanglement and Coherence but excludes Evidence, resulting in performance drops across

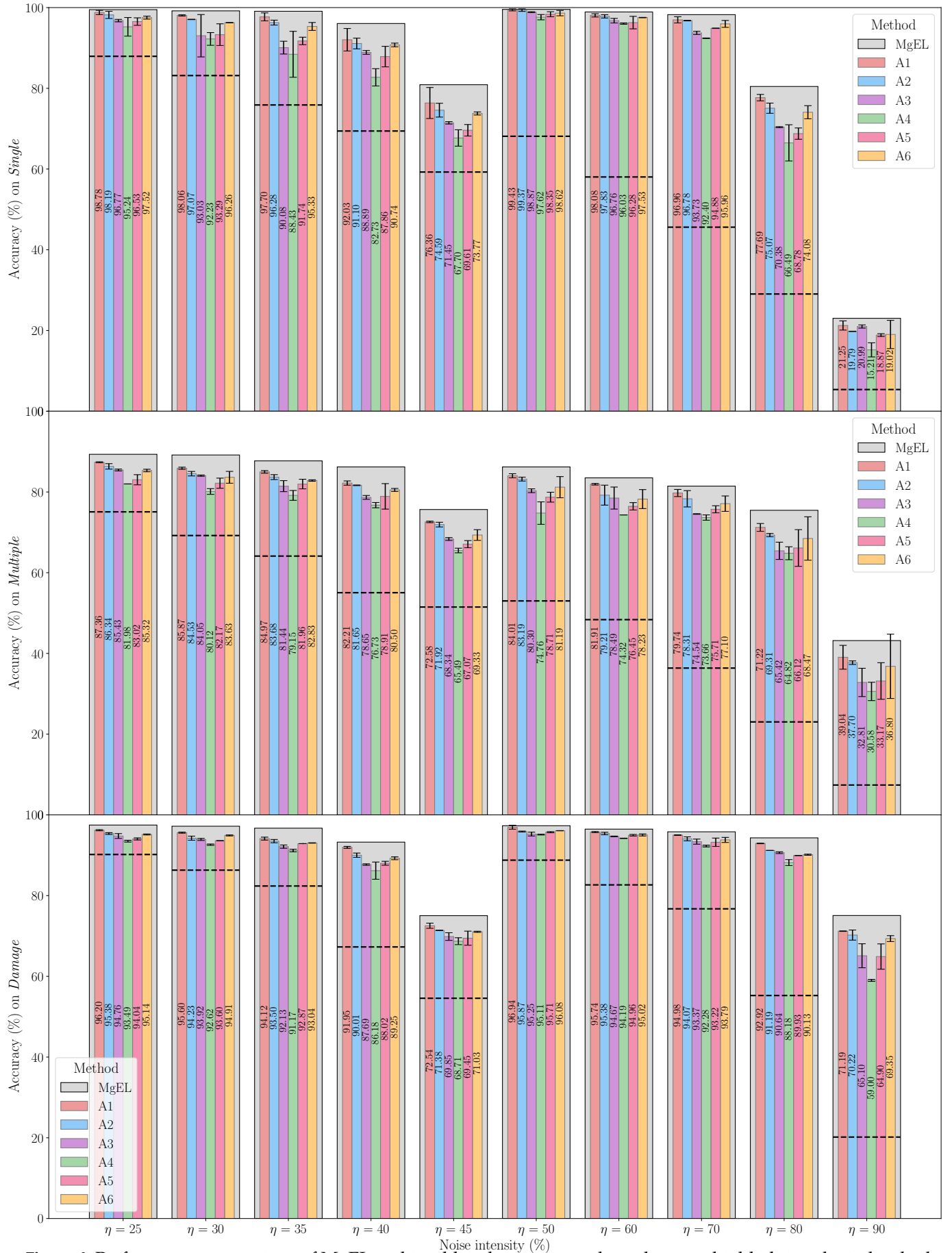


Figure 6. Performance comparison of MgEL and its ablated variants on three datasets, highlighting the individual contribution of **Entanglement**, **Evidence**, and **Coherence** modules, where dashed lines represent the accuracy of baseline.

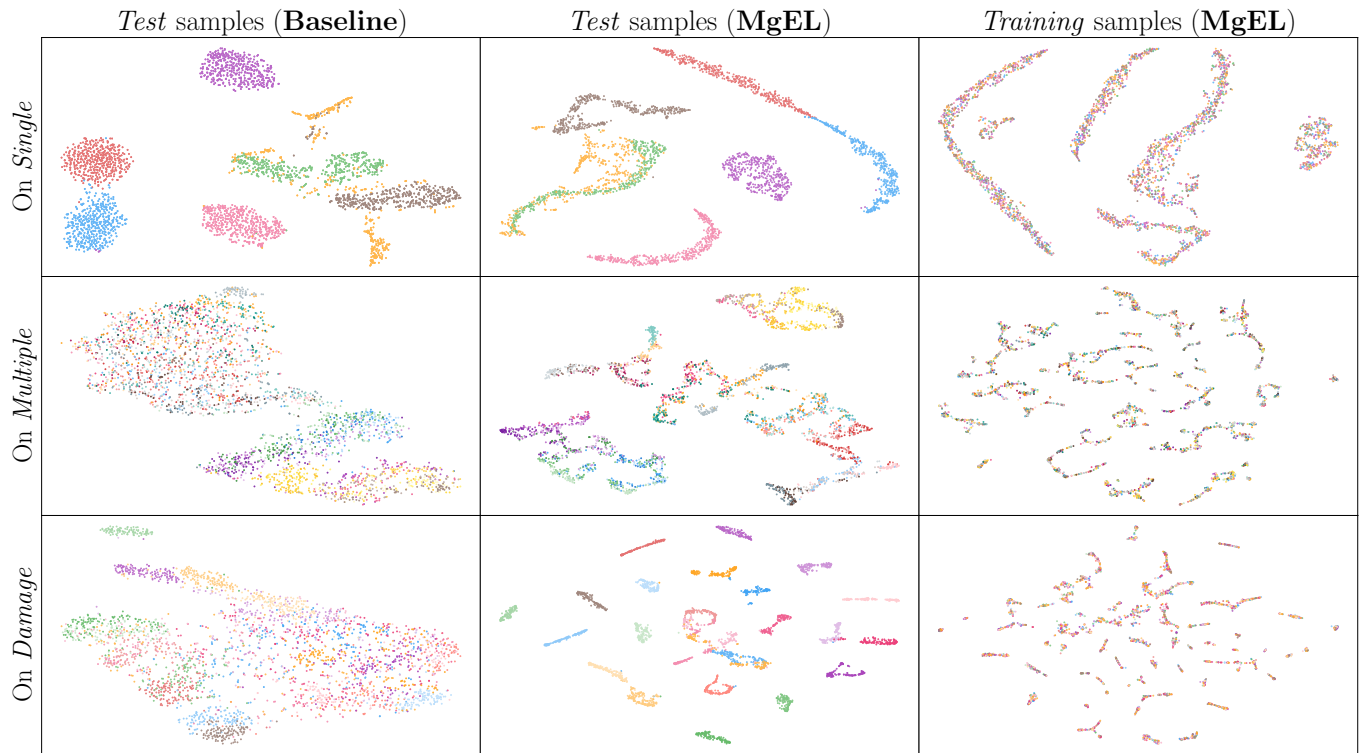


Figure 7. t-SNE visualization of latent feature distributions for test and training samples across the three datasets. The first column shows test samples from the baseline model trained without noise handling. The second and third columns show test and training samples learned under the MgEL framework. Compared to the baseline, MgEL induces more compact and well-separated class clusters, even under high noise.

all datasets (e.g., 75.58% *vs.* 79.78% on *Multiple*). A5, retaining only Entanglement, performs even worse. These observations indicate that while structural diversity is necessary, accurate fusion of annotations and pseudo-labels requires confidence calibration. The belief mass $\mathcal{E}_i$, computed from local similarity-weighted label distributions and adjusted for noise intensity, serves this purpose. Its removal forces reliance on raw class probabilities, weakening the model's ability to distinguish reliable label sources. Thus, Evidence enhances label consistency by integrating supervision signals based on their reliability within the cluster context.

The effect of Coherence is reflected in A3 and A6. A3 excludes Coherence while preserving Entanglement and Evidence, leading to a modest performance reduction (e.g., 72.95% *vs.* 79.78% on *Multiple*), while A6—disabling both Coherence and Entanglement—shows the lowest performance among all variants. These results confirm that Coherence plays a significant role in label stability, especially under ambiguous or conflicting supervision. The mechanism selectively defers label commitments by reverting to the superposition state when the pseudo, annotation, and collapsed labels do not agree. Therefore, Coherence mitigates confirmation bias and stabilizes predictions by withholding unreliable

updates in the presence of low label consensus.

Collectively, the ablation results demonstrate that MgEL's tolerance to noisy labels comes from the interplay of three complementary mechanisms: entanglement introduces structured diversity for uncertainty modeling, evidence enables reliability-aware label fusion via belief-based weighting, and coherence mitigates confirmation bias by deferring low-consensus decisions. Their unified integration equips MgEL with robust generalization capabilities across diverse and corrupted learning environments.

4.4 Case III: Visualization analysis

To empirically support Assumption 1, we visualize the t-SNE embeddings of both training and test samples from the three datasets under 90% symmetric noise, as shown in Figure 7. The resulting projections reveal significant differences in the topological organization of feature manifolds, providing intuitive evidence of noise-induced structural perturbations and highlighting MgEL's ability to counteract these distortions through more coherent class representations.

Compared to the baseline, the latent features learned by MgEL show significantly improved intra-class

compactness and inter-class separability. On the *Single* dataset, clear and well-separated clusters emerge even for complex fault types, while the baseline features remain entangled with indistinct boundaries. The *Multiple* dataset, which introduces greater domain and fault variability, further demonstrates MgEL's robustness: despite the increased difficulty, MgEL preserves coherent topological structures for most classes, while the baseline embeddings collapse into ambiguous, overlapping manifolds.

Notably, the effect is most pronounced on the *Damage* dataset, which consists of 49 fine-grained fault categories and extreme label corruption ($\eta = 90\%$). In this challenging setting, the baseline model fails to preserve class-discriminative geometry, leading to severe semantic drift. In contrast, MgEL successfully induces structured manifolds with meaningful class-wise alignment, consistent with the suppression of structural perturbations predicted by the theoretical model.

Furthermore, training embeddings offer additional insight into MgEL's manifold-regularizing behavior. Although residual intra-cluster variance persists due to corrupted supervision, the global alignment between training and test distributions is significantly improved under MgEL. This suggests that the framework filters noise at the label level and stabilizes representation learning by regularizing optimization trajectories.

The observed improvements are attributed to MgEL's evidence-aware label construction process, which incorporates uncertainty through probabilistic superposition and belief-based fusion mechanisms. These mechanisms effectively mitigate the propagation of biased gradients and restore the semantic coherence of the latent space.

Taken together, the t-SNE results corroborate the theoretical analysis in Assumption 1, showing that extreme label noise induces a representation-level structural shift. MgEL mitigates this phenomenon by preserving class-consistent manifolds in the latent space, improving both robustness and generalization.

5 Conclusion

Motivated by the challenge of *confirmation bias* in LNL, this work proposes MgEL, a quantum-inspired framework that introduces a novel multi-granularity evidence labeling mechanism by simulating the observation-collapse behavior of entangled systems. Through pseudo-source construction, evidence-aware fusion, and coherence-guided filtering, MgEL

effectively integrates annotations, pseudo-labels, and collapsed labels to create robust supervision signals. Extensive experiments across three fault diagnosis benchmarks confirm that MgEL outperforms existing methods in most scenarios, particularly under severe symmetric noise, while consistently improving the quality of latent representations. MgEL provides a theoretically grounded and practically scalable solution for trustworthy fault diagnosis in data-driven industrial systems. Beyond its empirical success, MgEL introduces a principled information fusion strategy that integrates multiple uncertain label sources, offering new insights into uncertainty modeling and decision-level fusion in noisy environments.

Despite its promising performance, MgEL has several limitations. The stochastic nature of random partitioning may cause instability in small or highly imbalanced datasets, and reliance on similarity-based neighborhood construction can be sensitive to feature distortions during early training. Future research could explore more robust entanglement schemes, incorporate adaptive clustering, or integrate causal and temporal priors into the label-fusion process. Moreover, extending the proposed framework to semi-supervised, active, or federated learning settings holds promise for advancing both the theory and practice of information fusion under uncertainty.

Data Availability Statement

Data will be made available on request.

Funding

This work was supported in part by the National Natural Science Foundation of China under Grant 62233003 and Grant 62073072; in part by the Key Projects of Key R&D Program of Jiangsu Province under Grant BE2020006 and Grant BE2020006-1; in part by the Shenzhen Science and Technology Program under Grant JCYJ20210324132202005 and Grant JCYJ20220818101206014.

Conflicts of Interest

The authors declare no conflicts of interest.

Ethical Approval and Consent to Participate

Not applicable.

References

- [1] Dunkin, F., Li, X., Hu, C., Wu, G., Li, H., Lu, X., & Zhang, Z. (2024). Like draws to like: A Multi-granularity Ball-Intra Fusion approach for fault diagnosis models to resist misleading by noisy labels. *Advanced Engineering Informatics*, 60, 102425. [CrossRef]
- [2] Li, X., Dunkin, F., & Dezert, J. (2024). Multi-source information fusion: Progress and future. *Chinese Journal of Aeronautics*, 37(7), 24–58. [CrossRef]
- [3] Zheng, X., Nie, J., He, Z., & Gao, M. (2025). Specific Task-Guided Collaborative Domain Generalization Network for Intelligent Fault Diagnosis under Unseen Conditions. *IEEE Internet of Things Journal*. [CrossRef]
- [4] Dunkin, F., Li, X., Li, H., Wu, G., Hu, C., & Ge, S. S. (2024). MgCNL: A Sample Separation Approach via Multi-Granularity Balls for Fault Diagnosis With the Interference of Noisy Labels. *IEEE Transactions on Automation Science and Engineering*, 22, 7748–7761. [CrossRef]
- [5] Hu, C., Zhang, Z., Li, C., Leng, M., Wang, Z., Wan, X., & Chen, C. (2025). A state of the art in digital twin for intelligent fault diagnosis. *Advanced Engineering Informatics*, 63, 102963. [CrossRef]
- [6] Zhang, W., Jiang, L., & Li, C. (2024). ELDP: Enhanced Label Distribution Propagation for Crowdsourcing. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 47(3), 1850–1862. [CrossRef]
- [7] Yang, H., Wang, L., Pan, Y., & Chen, J.-J. (2025). A Teacher-Student Framework Leveraging Large Vision Model for Data Pre-Annotation and YOLO for Tunnel Lining Multiple Defects Instance Segmentation. *Journal of Industrial Information Integration*, 100790. [CrossRef]
- [8] Zhang, Y., Chen, Y., Fang, C., Wang, Q., Wu, J., & Xin, J. (2025). Learning from open-set noisy labels based on multi-prototype modeling. *Pattern Recognition*, 157, 110902. [CrossRef]
- [9] Bian, Z., Chang, Q., Wang, J., Pedrycz, W., & Pal, N. R. (2024). Takagi–sugeno–kang fuzzy systems for high-dimensional multilabel classification. *IEEE Transactions on Fuzzy Systems*, 32(6), 3790–3804. [CrossRef]
- [10] Sun, Y., Song, H., Guo, L., Gao, H., & Cao, A. (2025). A transfer learning method: Universal domain adaptation with noisy samples for bearing fault diagnosis. *Advanced Engineering Informatics*, 65, 103243. [CrossRef]
- [11] Song, H., Kim, M., Park, D., Shin, Y., & Lee, J.-G. (2023). Learning From Noisy Labels With Deep Neural Networks: A Survey. *IEEE Transactions on Neural Networks and Learning Systems*, 34(11), 8135–8153. [CrossRef]
- [12] Liu, Y., Zhong, Y., Ma, A., Zhao, J., & Zhang, L. (2023). Cross-resolution national-scale land-cover mapping based on noisy label learning: A case study of China. *International Journal of Applied Earth Observation and Geoinformation*, 118, 103265. [CrossRef]
- [13] Zhang, J., Song, B., Wang, H., Han, B., Liu, T., Liu, L., & Sugiyama, M. (2024). BadLabel: A Robust Perspective on Evaluating and Enhancing Label-Noise Learning. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 46(6), 4398–4409. [CrossRef]
- [14] Chen, M., Zhao, Y., He, B., Han, Z., Huang, J., Wu, B., & Yao, J. (2024). Learning with noisy labels over imbalanced subpopulations. *IEEE Transactions on Neural Networks and Learning Systems*, 36(4), 6544–6555. [CrossRef]
- [15] Lu, Y., & He, W. (2024). Mitigating Noisy Supervision Using Synthetic Samples with Soft Labels. *arXiv preprint arXiv:2406.16966*.
- [16] Kim, N.-R., Lee, J.-S., & Lee, J.-H. (2024). Learning with Structural Labels for Learning with Noisy Labels. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition* (pp. 27610–27620). [CrossRef]
- [17] Kim, S., Lee, D., Kang, S., Chae, S., Jang, S., & Yu, H. (2024, June). Learning Discriminative Dynamics with Label Corruption for Noisy Label Detection. In *2024 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)* (pp. 22477–22487). IEEE. [CrossRef]
- [18] ATLAS Collaboration. (2024). Observation of quantum entanglement with top quarks at the ATLAS detector. *Nature*, 633(8030), 542. [CrossRef]
- [19] Schwarz, A., Rahal, J. R., Sahelices, B., Barroso-García, V., Weis, R., & Duque Anton, S. (2024). Data augmentation in predictive maintenance applicable to hydrogen combustion engines: a review. *Artificial Intelligence Review*, 58(1), 32. [CrossRef]
- [20] Mo, Z., Zhang, Z., Miao, Q., & Tsui, K.-L. (2025). Extended Invariant Risk Minimization for Machine Fault Diagnosis With Label Noise and Data Shift. *IEEE Transactions on Neural Networks and Learning Systems*, 36(8), 15476–15489. [CrossRef]
- [21] Yu, W., Dillon, T., Mostafa, F., Rahayu, W., & Liu, Y. (2020). A Global Manufacturing Big Data Ecosystem for Fault Detection in Predictive Maintenance. *IEEE Transactions on Industrial Informatics*, 16(1), 183–192. [CrossRef]
- [22] Wang, X., Wang, S., Liang, Y., & Lei, Z. (2025). Decisive vector guided column annotation. *Pattern Recognition*, 158, 110958. [CrossRef]
- [23] Nguyen, T., Ibrahim, S., & Fu, X. (2024). Noisy Label Learning with Instance-Dependent Outliers: Identifiability via Crowd Wisdom. *Advances in Neural Information Processing Systems*, 37, 97261–97298.
- [24] Lin, Y., Yao, Y., & Liu, T. (2024). Learning the latent causal structure for modeling label noise. *Advances in Neural Information Processing Systems*, 37, 120549–120577.

- [25] Li, Y., Han, H., Shan, S., & Chen, X. (2023). DISC: Learning From Noisy Labels via Dynamic Instance-Specific Selection and Correction. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition* (pp. 24070–24079). [CrossRef]
- [26] Li, F., Li, K., Tian, J., & Zhou, J. (2024). Regroup Median Loss for Combating Label Noise. In *Proceedings of the AAAI Conference on Artificial Intelligence* (Vol. 38, No. 12, pp. 13474–13482). [CrossRef]
- [27] Zhou, X., Liu, X., Zhai, D., Jiang, J., & Ji, X. (2023). Asymmetric Loss Functions for Noise-Tolerant Learning: Theory and Applications. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 45(7), 8094–8109. [CrossRef]
- [28] Fang, C., Cheng, L., Mao, Y., Zhang, D., Fang, Y., Li, G., Qi, H., & Jiao, L. (2024). Separating Noisy Samples From Tail Classes for Long-Tailed Image Classification With Label Noise. *IEEE Transactions on Neural Networks and Learning Systems*, 35(11), 16036–16048. [CrossRef]
- [29] Xu, G., Yi, L., Xu, P., Li, J., Pu, R., Shui, C., Ling, C., McLeod, A. I., & Wang, B. (2025). Unraveling the Mysteries of Label Noise in Source-Free Domain Adaptation: Theory and Practice. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 1–17. [CrossRef]
- [30] Wang, Y., Ma, X., Chen, Z., Luo, Y., Yi, J., & Bailey, J. (2019). Symmetric cross entropy for robust learning with noisy labels. In *Proceedings of the IEEE/CVF International Conference on Computer Vision* (pp. 322–330). [CrossRef]
- [31] Wei, J., Liu, H., Liu, T., Niu, G., Sugiyama, M., & Liu, Y. (2021). To smooth or not? when label smoothing meets noisy labels. *arXiv preprint arXiv:2106.04149*.
- [32] Sheng, M., Sun, Z., Chen, T., Pang, S., Wang, Y., & Yao, Y. (2025). Foster Adaptivity and Balance in Learning with Noisy Labels. In *European Conference on Computer Vision* (pp. 217–235). [CrossRef]
- [33] Wei, H., Feng, L., Chen, X., & An, B. (2020). Combating noisy labels by agreement: A joint training method with co-regularization. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition* (pp. 13726–13735). [CrossRef]
- [34] Kim, D., Ryoo, K., Cho, H., & Kim, S. (2025). SplitNet: learnable clean-noisy label splitting for learning with noisy labels. *International Journal of Computer Vision*, 133(2), 549–566. [CrossRef]
- [35] Zhang, R., Cao, Z., Huang, Y., Yang, S., Xu, L., & Xu, M. (2025). Visible-Infrared Person Re-identification with Real-world Label Noise. *IEEE Transactions on Circuits and Systems for Video Technology*, 1–1. [CrossRef]
- [36] Bender, C. M., & Hook, D. W. (2024). $\mathcal{PT}$ -symmetric quantum mechanics. *Reviews of Modern Physics*, 96(4), 045002. [CrossRef]
- [37] Bose, S., Fuentes, I., Geraci, A. A., Khan, S. M., Qvarfort, S., Rademacher, M., Rashid, M., Toroš, M., Ulbricht, H., & Wanjura, C. C. (2025). Massive quantum systems as interfaces of quantum mechanics and gravity. *Reviews of Modern Physics*, 97(1), 015003. [CrossRef]
- [38] Liang, X.-B., Li, B., & Fei, S.-M. (2024). Signifying quantum uncertainty relations by optimal observable sets and the tightest uncertainty constants. *Science China Physics, Mechanics & Astronomy*, 67(9), 290311. [CrossRef]
- [39] Zhang, X., Wang, C., Zhou, W., Xu, J., & Han, T. (2024). Trustworthy diagnostics with out-of-distribution detection: A novel max-consistency and min-similarity guided deep ensembles for uncertainty estimation. *IEEE Internet of Things Journal*, 11(13), 23055–23067. [CrossRef]
- [40] Ruan, H., Wang, Y., Qin, Y., & Tang, B. (2021, October). An enhanced intelligent fault diagnosis method to combat label noise. In *2021 International Conference on Sensing, Measurement & Data Analytics in the era of Artificial Intelligence (ICSMD)* (pp. 1–6). IEEE. [CrossRef]
- [41] Li, M., He, S., Chen, J., Feng, Y., & Xie, J. (2025). Label-smoothing dynamic decoupling augmented network for intelligent fault diagnosis under imbalanced data distribution with noisy labels. *Measurement*, 118664. [CrossRef]
- [42] Hu, C. K., Wei, C., Liu, C., Che, L., Zhou, Y., Xie, G., ... & Yu, D. (2024). Experimental sample-efficient quantum state tomography via parallel measurements. *Physical Review Letters*, 133(16), 160801. [CrossRef]
- [43] Deng, J., Liu, H., Fang, H., Shao, S., Wang, D., Hou, Y., Chen, D., & Tang, M. (2023). MgNet: A fault diagnosis approach for multi-bearing system based on auxiliary bearing and multi-granularity information fusion. *Mechanical Systems and Signal Processing*, 193, 110253. [CrossRef]
- [44] Fang, H., Deng, J., Chen, D., Jiang, W., Shao, S., Tang, M., & Liu, J. (2023). You can get smaller: A lightweight self-activation convolution unit modified by transformer for fault diagnosis. *Advanced Engineering Informatics*, 55, 101890. [CrossRef]
- [45] Kim, Y., Yun, J., Shon, H., & Kim, J. (2021). Joint Negative and Positive Learning for Noisy Labels. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition* (pp. 9442–9451). [CrossRef]
- [46] Iscen, A., Valmadre, J., Arnab, A., & Schmid, C. (2022). Learning With Neighbor Consistency for Noisy Labels. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition* (pp. 4672–4681). [CrossRef]
- [47] Cordeiro, F. R., & Carneiro, G. (2025). ANNE: Adaptive Nearest Neighbours and Eigenvector-based sample selection for robust learning with noisy labels. *Pattern Recognition*, 159, 111132. [CrossRef]



journals. (Email: dunkin2fir@ieee.org)



Fir Dunkin received the M.E. degree in Control Science and Engineering from Northeast Electric Power University, China, in 2022. He is pursuing a Ph.D. in Electronic Information at Southeast University, China. His research interests include information fusion, uncertainty reasoning, and mechatronics. He has published over 30 papers, including several *ESI* Highly Cited Papers, and reviewed for numerous top

Xinde Li received his Ph.D. degree in control theory and control engineering from the Department of Control Science and Engineering, Huazhong University of Science and Technology, Wuhan, China, in 2007. After that, he joined the School of Automation, Southeast University, Nanjing, China, where he is currently a professor. He was elected as a senior member of Institute of Electrical and Electronics Engineers in 2016, an academician

with the Russian Academy of Natural Sciences and a cover person of *Scientific Chinese* in 2021 and a fellow of the Institution of Engineering and Technology in 2022. He is the vice director of Intelligent Robot Committee of the Chinese Association for Artificial Intelligence from 2017, the vice director of Intelligent Products and Industry Working Committee of the Chinese Association for Artificial Intelligence from 2019 and the vice director of Intelligent Manufacturing Systems & Technology Professional Committee of Chinese Association of Automation from 2022.

His research interests include intelligent robot, and unmanned system, human-machine interaction, computer vision, information fusion. He was the principal of about 50 projects including key and major projects of NSFC, etc. He published more than 100 papers, 5 books and won 32 national invention patents. He also won many prizes including global award for outstanding scientist in artificial intelligence and robotics, Asia award for outstanding researcher, gold award in Geneva international invention, etc. He was also the associate editors of several reputable international journals, i.e. *IEEE Transactions on Fuzzy Systems*, *Journal of Systems Engineering and Electronics*, etc. (Email: xindeli@seu.edu.cn)