

TRUSTED AND INTERACTIVE CLUSTERING FOR TIME-SERIES DATA

Anonymous authors

Paper under double-blind review

ABSTRACT

Time-series clustering has gained abundant popularity and has been used in diverse scientific areas. However, few researchers take an information fusion perspective to combine information from the time and frequency domains to accomplish clustering, although these two domains offer distinct and complementary characteristics of time-series. Motivated by this issue, we propose a trusted and interactive model, which leverages evidence theory to combine time- and frequency-based clustering results produced by the corresponding contrastive learning module. After mathematizing clustering results from the two domains as mass functions, the uncertainty contained in these results can be quantified at the sample-specific level. The combined result thus promotes clustering reliability, and is optimized based on the pseudo-labels generated by k -means in an interactive learning paradigm. Both theoretical analysis and experimental results on 136 benchmark datasets validate the effectiveness of the proposed model in clustering performance. Extensive ablation experiments demonstrate the contribution of combining information from the time and frequency domains and using the interactive learning paradigm. The embeddings learned are also experimentally shown to perform well in other downstream tasks.

1 INTRODUCTION

Time-series clustering is an important data mining technology widely applied in different fields, such as sensor data analysis (Hayashi et al., 2024), anomaly detection (Middlehurst et al., 2024) and medical field (Zhang et al., 2024), aiming to segment time-series data samples into patterns (called clusters) with homologous characteristics (Gong et al., 2022). Unlike images, time-series data generally do not show human-recognizable features to different classes, because the label information is contained in not only the time domain but also the frequency domain (Zhang et al., 2022).

Related work. Most of the existing time-series clustering methods focus on the time domain, and can be divided into two categories: raw-data-based methods and feature-based ones (Ma et al., 2022). Raw-data-based methods perform clustering based on a modified similarity metric, which quantifies the distance more appropriately between time-series samples (Hayashi et al., 2024; Ferreira & Zhao, 2016; Paparrizos & Gravano, 2015; Yang & Leskovec, 2011). Feature-based methods have recently garnered greater attention, because the raw-data-based ones are not capable of modeling nonlinear temporal dependencies and multiscale (long and short-term) temporal dependencies (Ma et al., 2019a). Feature-based methods extract first informative features from raw samples in the time domain, and then the clustering algorithms are conducted on the learned features (Tang et al., 2021; Fortuin et al., 2020). Some recent works (Zhang et al., 2024; Guijo-Rubio et al., 2021) also optimize feature extraction and clustering jointly by introducing pseudo-label. For example, authors in (Pélat et al., 2023) embed the time-series onto the Stiefel manifold to obtain the geometric representations of time-series samples. STCN (Ma et al., 2022) adopts a recurrent neural network and a self-supervised clustering module, which are trained iteratively by contributing to each other. Some recent works (e.g., (Fortuin et al., 2020; Tonekaboni et al., 2022)) leverage contrastive representation learning for clustering time-series to further improve the performance. Few of the mentioned methods take an information fusion perspective to incorporate information from the time and frequency domains to accomplish clustering, although these two domains offer distinct and complementary characteristics of time-series. More detailed related work is provided in Appendix A.1.

Motivation. In fact, the time domain depicts the temporal evolution of signal readouts, while the frequency domain reveals the distribution of signal magnitude across different frequency components within the entire spectrum (Hyndman & Athanasopoulos, 2018). By explicitly incorporating the frequency domain, a comprehension of time-series behavior can be attained, encompassing aspects that are not fully captured by analyzing the time domain alone. Therefore, the first motivation for this work is *incorporating frequency information to enhance the ability to detect clusters*. Besides, the time and frequency domains can be regarded as distinct views of the same data (Cohen, 1995), and they are interconvertible through Fourier and inverse Fourier Transformation (Brigham, 1988). Given the temporal dynamics inherent in time-series data, the weights (i.e., quality) of the time and frequency domains experience fluctuations over time. Then, the second motivation is *to dynamically describe the importance of clustering results derived from both the time and frequency domains*. Finally, to avoid poisoning the final result with the low-quality result from one domain, the last motivation is *to facilitate the appropriate integration of time- and frequency-based clustering results under the fast-evolving uncertain scenario*.

Technical issues. To address the three challenges outlined in the motivation, three technical issues must be solved.

(1) How to develop frequency-based contrastive augmentations to obtain clustering results?

Despite the universal importance of frequency information in time-series and its pivotal role in classic signal processing (Soklaski et al., 2022), it is seldom explored in contrastive clustering for time-series data (Tonekaboni et al., 2022). This is because even a slight perturbation in the frequency domain could lead to significant alterations in the temporal patterns of the time domain (Flandrin, 1998).

(2) How to quantify uncertainty in clustering results from the time and frequency domains?

Quantifying the uncertainty is practical for enhancing the performance of such a “two-view” time-series clustering, where the higher the uncertainty of a particular domain, the lower the weight contributing to the final result. Further, the uncertainty of every domain varies for different time-series samples.

(3) How to formulate and optimize the fusion of clustering results with uncertainty from the time and frequency domains?

Fusion of the clustering results from the two domains belongs to the late fusion (Wang et al., 2019; Liu et al., 2018), most of which primarily cater to scenarios without any uncertainty (Wang et al., 2021). In our model, the fusion module assumes the critical responsibility of effectively managing uncertainty, all integrated seamlessly within an end-to-end framework.

Contributions. We introduce a trusted clustering model (named TIC) for time-series data, integrating results from the time and frequency domains within an interactive learning paradigm (shown in Fig.1). In summary, the contributions of this paper are:

- We propose to adopt novel augmentations dedicated to clustering the frequency spectrum data, through contrastive sample discrimination. This could be the first work to leverage contrastive augmentation in the frequency domain in a time-series clustering problem.
- We represent the clustering results for each sample from the time and frequency domains as two distinct mass functions (Shafer, 1976), where the sample-specific uncertainty is accurately quantified within the framework of evidence theory (DEMPSTER, 1967). The Dirichlet distribution is used to mathematize the mass function and to model the class probability distribution from each domain.
- We propose an interactive framework for time-series clustering, which could be inspiring in multiple research areas of time-series. The trusted result (obtained by combining mass functions from the time and frequency domains) and pseudo-labels (derived from k-means) contribute to each other, allowing interactive model optimization.

2 PRELIMINARIES: EVIDENCE THEORY

Consider a variable ω taking values in a finite set called the *frame of discernment* $\Omega = \{\omega_1, \omega_2, \dots, \omega_C\}$, where C is the number of clusters in a clustering problem. A mass function (also called a piece of evidence) (Shafer, 1976) is defined as a mapping from 2^Ω to $[0,1]$ such that

$\sum_{A \subseteq \Omega} m(A) = 1, m(A) > 0$, where 2^Ω is the power set of Ω and A denotes various subsets of Ω . These subsets are called the *focal sets* of m . The value of $m(A)$ denotes a degree of belief assigned to the hypothesis “ $\omega \in A$ ”. The vacuous mass function verifies

$$m(\Omega) = m_\Omega = 1 \quad (1)$$

corresponding to total “ignorance”, representing that ω may belong to any subset of the Ω (including the Ω itself). In other words, the value of $m(\Omega)$ measures the uncertainty about the values of variable ω . In this work, only the singleton focal sets $\omega_1, \omega_2, \dots, \omega_C$ and ignorance focal set Ω are considered, i.e., the mass function is in the form of $\mathbf{m} = [m(\omega_1), m(\omega_2), \dots, m(\omega_C), m(\Omega)]$ (abbreviated as $[m_1, m_2, \dots, m_C, m_\Omega]$).

Assume that there are 2 mass functions m_1 and m_2 on the same frame of discernment Ω . The *Dempster’s rule* (Shafer, 1976) used to pool the information provided by $\mathbf{m}_1$ and $\mathbf{m}_2$ is noted as $\oplus$ and is defined by

$$m_{1 \oplus 2}(A) = \frac{\sum_{B \cap C = A} m_{1B} m_{2C}}{\mathcal{K}_{12}}, \quad A \neq \emptyset \text{ \& } A \subseteq \Omega \quad (2)$$

where B and C are also the focal sets, and $\mathcal{K}_{12} = 1 - \sum_{B \cap C = \emptyset} m_{1B} m_{2C}$ is a normalizing factor. Dempster’s rule is commutative and associative (Shafer, 1976).

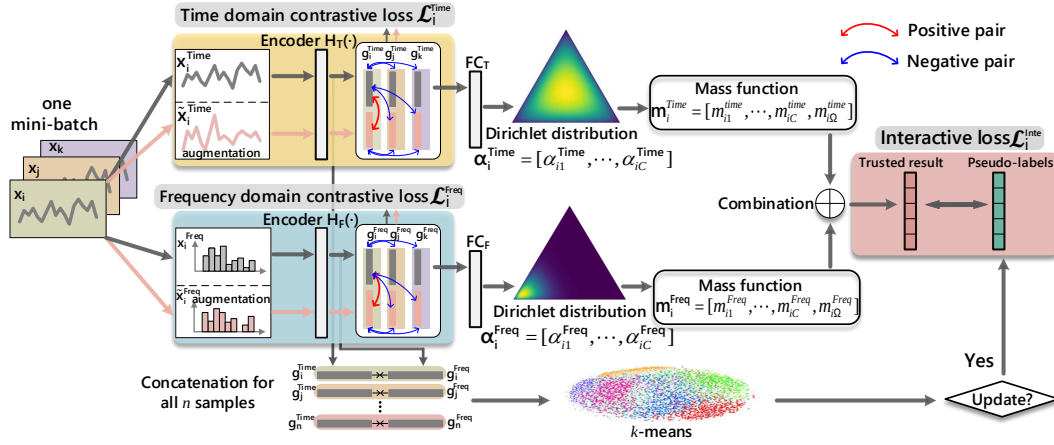


Figure 1: Overview of the proposed TIC model. TIC has three modules, a time-based contrastive module (colored yellow), a frequency-based contrastive module (colored blue), and an interactive learning module (colored red).

3 THE PROPOSED METHOD: TIC

Notions and problem formulation. We are given a time-series dataset $T = \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_n\}$ of n unlabeled time-series samples, and sample $\mathbf{x}_i$ has K channels and L time-steps. The goal is to group these n samples into C (a given value) clusters, which can be denoted by the *framework of discernment* $\Omega = \{\omega_1, \omega_2, \dots, \omega_C\}$ (defined in Section 2). Without loss of generality, in the following, we focus on univariate (single-channel) time-series datasets, while noting that our TIC method can also accommodate multi-variate time-series. Superscript $\sim$ denotes contrastive augmentations, $\mathbf{x}_i \equiv \mathbf{x}_i^{\text{Time}}$ denotes the input time-series and $\mathbf{x}_i^{\text{Freq}}$ denotes the frequency spectrum of $\mathbf{x}_i$.

Time-based contrastive module. For the sample $\mathbf{x}_i^{\text{Time}}$ in one mini-batch, a set of augmentations $\mathcal{X}_i^{\text{Time}}$ is generated through a time-based augmentation bank $\mathcal{B}^{\text{Time}} : \mathbf{x}_i^{\text{Time}} \rightarrow \mathcal{X}_i^{\text{Time}}$ including jittering, scaling, time-shifts and other common techniques (Eldele et al., 2021). Note that the augmentations in one mini-batch are produced using diverse techniques from the augmentation bank, to expose the model to complex temporal dynamics and obtain robust embeddings.

$\mathbf{x}_i^{\text{Time}}$ and a randomly selected augmentation $\tilde{\mathbf{x}}_i^{\text{Time}} \in \mathcal{X}_i^{\text{Time}}$ are fed into the encoder, denoted by $H_T(\cdot)$. The corresponding embeddings $\mathbf{g}_i^{\text{Time}} = H_T(\mathbf{x}_i^{\text{Time}})$ and $\tilde{\mathbf{g}}_i^{\text{Time}} = H_T(\tilde{\mathbf{x}}_i^{\text{Time}})$ are obtained.

Intuitively, embedding $\mathbf{g}_i^{\text{Time}}$ should be close to the embedding $\tilde{\mathbf{g}}_i^{\text{Time}}$, but far away from the embeddings $\mathbf{g}_j^{\text{Time}}$ and $\tilde{\mathbf{g}}_j^{\text{Time}}$ produced from another sample $\mathbf{x}_j^{\text{Time}}$ in the same mini-batch. Therefore, the positive pair is $(\mathbf{x}_i^{\text{Time}}, \tilde{\mathbf{x}}_i^{\text{Time}})$ and the negative pairs are $(\mathbf{x}_i^{\text{Time}}, \mathbf{x}_j^{\text{Time}})$ and $(\mathbf{x}_i^{\text{Time}}, \tilde{\mathbf{x}}_j^{\text{Time}})$. The normalized temperature-scaled cross-entropy loss associated with $\mathbf{x}_i^{\text{Time}}$ is defined as (Chen et al., 2020)

$$\mathcal{L}_i^{\text{Time}} = -\log \frac{\exp(\text{sim}(\mathbf{g}_i^{\text{Time}}, \tilde{\mathbf{g}}_i^{\text{Time}}))/\tau}{\sum_{\mathbf{x}_j} \mathbb{1}_{i \neq j} \exp(\text{sim}(\mathbf{g}_i^{\text{Time}}, H_T(\mathbf{x}_j))/\tau)}, \quad (3)$$

where $\text{sim}(\mathbf{a}, \mathbf{b}) = \mathbf{a}^T \mathbf{b} / \|\mathbf{a}\| \|\mathbf{b}\|$ is the cosine similarity, τ is a temporal hyperparameter to adjust scale, $\mathbb{1}$ is an indicator function equaling to 1 when $i \neq j$ and 0 otherwise, and $\mathbf{x}_j$ denotes the sample (or its augmented sample) different from $\mathbf{x}_i$ in the same mini-batch. Minimization of $\mathcal{L}^{\text{Time}}$ enforces encoder $H_T(\cdot)$ to bring embeddings w.r.t. positive pairs closer together, and push embeddings w.r.t. negative pairs farther apart.

Frequency-based contrastive module. Although the frequency spectrum is informative, few methods leverage the frequency-based contrastive augmentation for clustering time-series (Tonekaboni et al., 2022). In this module, we use the Fourier Transformation (Brigham, 1988) to generate the frequency spectrum $\mathbf{x}_i^{\text{Freq}}$ for sample $\mathbf{x}_i^{\text{Time}}$.

As shown in (Flandrin, 1998; Zhang et al., 2022), a minor perturbation in the frequency domain can lead to significant changes in the corresponding time domain. To mitigate this issue, we manipulate the amplitude to generate frequency-based augmentation. More concretely, the augmentation bank $\mathcal{B}^{\text{Freq}} : \mathbf{x}_i^{\text{Freq}} \rightarrow \mathcal{X}_i^{\text{Freq}}$, where $\mathcal{X}_i^{\text{Freq}}$ is a set of frequency-based augmentations, includes the upgrade or downgrade of amplitude. We randomly select β (the number of components to be manipulated) frequency components, and change each of their amplitudes from the original value Amp_{orig} to $\gamma \text{Amp}_{\text{orig}}$, $0 \leq \gamma < 1$ (downgrade) or to $\gamma \text{Amp}_{\text{orig}}$, $\gamma > 1$, $\text{Amp}_{\text{orig}} < \gamma \text{Amp}_{\text{orig}} \leq \text{Amp}_{\text{max}}$ (upgrade), where γ is a pre-defined coefficient and Amp_{max} is the maximum amplitude. Similar to the time-based contrastive module, the positive pair is $(\mathbf{x}_i^{\text{Freq}}, \tilde{\mathbf{x}}_i^{\text{Freq}})$ and the negative pairs are $(\mathbf{x}_i^{\text{Freq}}, \mathbf{x}_j^{\text{Freq}})$ and $(\mathbf{x}_i^{\text{Freq}}, \tilde{\mathbf{x}}_j^{\text{Freq}})$.

After generating an augmentation $\tilde{\mathbf{x}}_i^{\text{Freq}} \in \mathcal{X}_i^{\text{Freq}}$, we feed the frequency spectrum $\mathbf{x}_i^{\text{Freq}}$ and $\tilde{\mathbf{x}}_i^{\text{Freq}}$ into the encoder $H_F(\cdot)$, and obtain the embeddings $\mathbf{g}_i^{\text{Freq}}$ and $\tilde{\mathbf{g}}_i^{\text{Freq}}$. The frequency-based loss for sample $\mathbf{x}_i^{\text{Freq}}$ is calculated as

$$\mathcal{L}_i^{\text{Freq}} = -\log \frac{\exp(\text{sim}(\mathbf{g}_i^{\text{Freq}}, \tilde{\mathbf{g}}_i^{\text{Freq}}))/\tau}{\sum_{\mathbf{x}_j} \mathbb{1}_{i \neq j} \exp(\text{sim}(\mathbf{g}_i^{\text{Freq}}, H_F(\mathbf{x}_j))/\tau)}. \quad (4)$$

Uncertainty quantification. As shown in Fig.1, the embeddings $\mathbf{g}_i^{\text{Time}}$ and $\mathbf{g}_i^{\text{Freq}}$ are fed into the fully connected (FC) layers to obtain the clustering results in time and frequency domains. Taking the time domain as an example, the FC_T converts the continuous embeddings to vectors $\mathbf{p}_T = [p_1^T, p_2^T, \dots, p_C^T]$, which describes the probability of C mutually exclusive events after being normalized through the *softmax* operator, and can be regarded as the parameters of a multinomial distribution (Bishop & Nasrabadi, 2006). By replacing these parameters with the parameters of a Dirichlet distribution, the clustering result from the FC_T can be represented as a distribution over possible *softmax* outputs instead of a point estimation of one *softmax* output (Sensoy et al., 2018). That is, a Dirichlet distribution parametrized over the $\mathbf{p}_{T,i} = [p_{i1}^T, p_{i2}^T, \dots, p_{iC}^T]$ represents the density of such a probability assignment w.r.t to sample $\mathbf{x}_i^{\text{Time}}$. Therefore, it can model the second-order probabilities and uncertainty for the clustering result of $\mathbf{x}_i^{\text{Time}}$ (Jsang, 2018). The definition of Dirichlet distribution is given in Definition 1.

Definition 1 The Dirichlet distribution is a probability density function for a categorical distribution $\mathbf{p}$. It can be characterized by C parameters $\alpha = [\alpha_1, \alpha_2, \dots, \alpha_C]$ and is given by:

$$\text{Dir}(\mathbf{p}|\alpha) = \begin{cases} \frac{1}{B(\alpha)} \prod_{c=1}^C p_c^{\alpha_c-1} & \text{for } \mathbf{p} \in \mathcal{S}_C, \\ 0 & \text{otherwise,} \end{cases} \quad (5)$$

where $\mathcal{S}_C$ is the C -dimensional unit simplex $\mathcal{S}_C = \{\mathbf{p} | \sum_{c=1}^C p_c = 1, 0 \leq p_1, p_2, \dots, p_C \leq 1\}$ and $B(\alpha)$ is the C -dimensional multinomial beta function.

In the clustering problem investigated in this paper, Subjective logic (Jsang, 2018) is used to associate the parameters $\alpha = [\alpha_1, \alpha_2, \dots, \alpha_C]$ of a Dirichlet distribution with the output $\mathbf{p} = [p_1, p_2, \dots, p_C]$, where the Dirichlet distribution is considered as the conjugate prior of the corresponding multinomial distribution (Bishop & Nasrabadi, 2006). Inspired from (Sensoy et al., 2018), the parameter α_c is calculated as

$$\alpha_c = p_c + 1. \quad (6)$$

Then, the mass function $\mathbf{m} = [m_1, m_2, \dots, m_C, m_\Omega]$ is determined as

$$\begin{cases} m_c = \frac{p_c}{S} = \frac{\alpha_c - 1}{S} & c = 1, 2, \dots, C, \\ m_\Omega = \frac{C}{S}, \end{cases} \quad (7)$$

where $S = \sum_{c=1}^C (p_c + 1) = \sum_{c=1}^C \alpha_c$ is the Dirichlet strength (Jsang, 2018), and the value of m_Ω quantifies the uncertainty of the clustering result (as discussed in Section 2). From Eq.(7), it can be inferred that the higher value of p_c , the more belief mass is assigned to m_c . Besides, the lower value of the sum of p_c , the higher uncertainty m_Ω for the clustering result. For clarity, we give a specific example to explain the above formulation.

Combining the clustering results. After computing the mass functions $\mathbf{m}_i^{\text{Time}}$ and $\mathbf{m}_i^{\text{Freq}}$ from time and frequency domains, we use Dempster’s rule to combine these two mass functions

$$\mathbf{m}_i^{\text{Comb}} = \mathbf{m}_i^{\text{Time}} \oplus \mathbf{m}_i^{\text{Freq}}, \quad (8)$$

where $\mathbf{m}_i^{\text{Comb}}$ is the combined mass function w.r.t. sample $\mathbf{x}_i$. According to Eq.(2), the specific calculation is

$$\begin{aligned} m_{ic}^{\text{Comb}} &= \frac{m_{ic}^{\text{Time}} \cdot m_{ic}^{\text{Freq}} + m_{i\Omega}^{\text{Time}} \cdot m_{ic}^{\text{Freq}} + m_{ic}^{\text{Time}} \cdot m_{i\Omega}^{\text{Freq}}}{1 - \sum_{r \neq v} m_{ir}^{\text{Time}} \cdot m_{iv}^{\text{Freq}}} \\ m_{i\Omega}^{\text{Comb}} &= \frac{m_{i\Omega}^{\text{Time}} \cdot m_{i\Omega}^{\text{Freq}}}{1 - \sum_{r \neq v} m_{ir}^{\text{Time}} \cdot m_{iv}^{\text{Freq}}}. \end{aligned} \quad (9)$$

According to Eq.(7), the combined output for $\mathbf{x}_i$ is calculated as

$$p_{ic}^{\text{Comb}} = m_{ic}^{\text{Comb}} \cdot S_i^{\text{Comb}} \text{ and } \alpha_{ic}^{\text{Comb}} = p_{ic}^{\text{Comb}} + 1, \quad (10)$$

where $S_i^{\text{Comb}} = \frac{C}{m_{i\Omega}^{\text{Comb}}}$. The time- and frequency-based clustering results are appropriately fused using Subjective logic and Dempster’s rule.

Remark 1. More intuitions: why the combined clustering results are trusted? The produced mass value $m_{i\Omega}^{\text{Comb}}$ allows the TIC model to assess the reliability of clustering results so as to avoid risky decisions. Besides, Dempster’s rule has the following advantages: (1) the mass function $\mathbf{m}^{\text{Comb}}$ obtained by combining a certain mass function (with small m_Ω) with an uncertain one (with big m_Ω) is still certain. This means that as long as the clustering results from either domain are trusted, then the combined clustering results are trusted, even if the results from the other domain have significant uncertainty. (2) the $\mathbf{m}^{\text{Comb}}$ obtained by combining two uncertain mass function (with large m_Ω) remains uncertain. This means that if the clustering results from both the time and frequency domains are untrusted, the combined clustering results are necessarily untrusted. And users can abandon the results to avoid risks when facing this case. In order to analyze theoretically these advantages, we give the following mathematized propositions, where the combined mass function, the mass functions from the time and frequency domains are abbreviated as $\mathbf{m}^{\text{Co}}$, $\mathbf{m}^{\text{T}}$ and $\mathbf{m}^{\text{F}}$. Since $\mathbf{m}^{\text{T}}$ and $\mathbf{m}^{\text{F}}$ are equally important in the combination, the following Propositions still hold despite the exchange of the superscripts T and F . The corresponding proofs are shown in Appendix A.2.

Proposition 1 A large m_Ω^{T} does not lead to a large m_Ω^{Co} , when one of m_c^{F} is large and m_Ω^{F} is small. In particular, $\mathbf{m}^{\text{Co}}$ is identical to $\mathbf{m}^{\text{T}}$, if $\mathbf{m}^{\text{F}}$ is totally uncertain (i.e., $m_\Omega^{\text{F}} = 1$).

Proposition 2 The m_Ω^{Co} is monotonically increasing with m_Ω^{T} and m_Ω^{F} .

Interactive learning module. As shown in the lower left part of Fig.1, we concatenate the embeddings $[\mathbf{g}^{\text{Time}}, \mathbf{g}^{\text{Freq}}]$ produced from the encoders $H_T(\cdot)$ and $H_F(\cdot)$. These concatenated embeddings of n samples are fed into k -means every t epochs and update the pseudo-labels to calculate the interactive loss $\mathcal{L}^{\text{Inte}}$.

For sample $\mathbf{x}_i$, its one-hot pseudo-label vector is denoted as $\mathbf{y}_i$ with $y_{ic} = 1$ and $y_{iv} = 0$ for all $v \neq c$. In the interactive learning module, we modify the conventional cross-entropy $\mathcal{L}_i^{\text{ce}} = -\sum_{c=1}^C y_{ic} \log(p_{ic}^{\text{Comb}})$ as

$$\mathcal{L}_i^{\text{ce}} = \int \left[\sum_{c=1}^C -y_{ic} \log(p_{ic}) \right] \frac{1}{B(\boldsymbol{\alpha}_i)} \prod_{c=1}^C p_{ic}^{\alpha_{ic}-1} d\mathbf{p}_i = \sum_{c=1}^C y_{ic} (\psi(S_i) - \psi(\alpha_{ic})) \quad (11)$$

where $\psi(\cdot)$ is the digamma function, S_i is the Dirichlet strength w.r.t. $\mathbf{x}_i$ and we omit superscript Comb for brevity. Such a modification enforces the parameters $\boldsymbol{\alpha}_i^{\text{Comb}}$ of the combined Dirichlet distribution to be optimized based on the pseudo-label vector $\mathbf{y}_i$, i.e., enforces a large p_{ic}^{Comb} to be produced from the $y_{ic} = 1$ in $\mathbf{y}_i$. To further shrink the p_{iv}^{Comb} w.r.t. y_{iv} to 0, the following KL divergence is considered

$$KL[Dir(\mathbf{p}_i|\tilde{\boldsymbol{\alpha}}_i) \| Dir(\mathbf{p}_i|\mathbf{1})] = \log \left(\frac{\Gamma(\sum_{c=1}^C \tilde{\alpha}_{ic})}{\Gamma(C) \prod_{c=1}^C \Gamma(\tilde{\alpha}_{ic})} \right) + \sum_{c=1}^C (\tilde{\alpha}_{ic} - 1) [\psi(\tilde{\alpha}_{ic}) - \psi(\sum_{v=1}^C \tilde{\alpha}_{iv})], \quad (12)$$

where $\tilde{\boldsymbol{\alpha}}_i = \mathbf{y}_i + (1 - \mathbf{y}_i) \odot \boldsymbol{\alpha}_i$ is the adjusted parameters of Dirichlet distribution, $\mathbf{1}$ is quite a flat Dirichlet distribution and $\Gamma(\cdot)$ is the *gamma* function. Thus, the interactive loss for $\mathbf{x}_i$ is

$$\mathcal{L}_i^{\text{Inte}} = \mathcal{L}_i^{\text{ce}} + \iota KL[Dir(\mathbf{p}_i|\tilde{\boldsymbol{\alpha}}_i) \| Dir(\mathbf{p}_i|\mathbf{1})], \quad (13)$$

where ι is the balance hyperparameter. In summary, the sample-specific loss of TIC model is

$$\mathcal{L}_i^{\text{TIC}} = \mathcal{L}_i^{\text{Inte}} + \mathcal{L}_i^{\text{Time}} + \mathcal{L}_i^{\text{Freq}}. \quad (14)$$

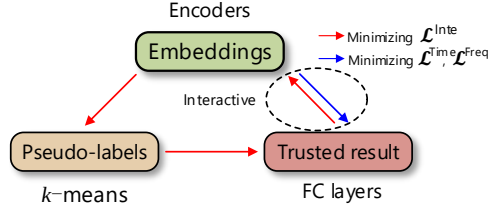


Figure 2: Illustration of the learning process in TIC. When minimizing contrastive loss $\mathcal{L}^{\text{Time}}$ and $\mathcal{L}^{\text{Freq}}$, the learned embeddings ($\mathbf{g}^{\text{Time}}$ and $\mathbf{g}^{\text{Freq}}$) affect the output of the FC layers (trusted clustering result), as shown by the blue arrow. When minimizing $\mathcal{L}^{\text{Inte}}$, the TIC model enforces the trusted clustering result to be “close” to the pseudo-labels, which are generated by inputting the embeddings into k -means. In the way of backward propagation, the trusted clustering result also affects the embedding learning, as shown by the red arrow. The parameter learning in FC layers and encoders ($H_T(\cdot), H_F(\cdot)$) contribute to each other interactively.

Remark 2. More intuitions about the interactive process. In Fig.2, we show the three main components in TIC model. The embedding learned affects the trusted clustering result when minimizing the contrastive loss $\mathcal{L}^{\text{Time}}$ and $\mathcal{L}^{\text{Freq}}$. The pseudo-labels are produced by inputting the embeddings into k -means and are considered when minimizing $\mathcal{L}^{\text{Inte}}$. In this way, the trusted clustering result affects the embedding learning when performing backward propagation. This allows the parameter learning of FC layers and encoders to contribute to each other interactively. Besides, considering the interactive loss avoids the class collision issue, where each sample is identified as a cluster in the embedding space (shown in Fig.3(a)) (Arora et al., 2019). This is because every positive pair consists only of the sample and its augmentation, without considering any other samples that may belong to the same cluster, when calculating the contrastive losses. The pseudo-labels in interactive loss are generated by k -means, where the concatenated embeddings $[\mathbf{g}^{\text{Time}}, \mathbf{g}^{\text{Freq}}]$ are treated as the input. Therefore, some basic information about clusters is included in the interactive loss and improves the clustering performance (shown in Fig.3(b)). We demonstrate quantitatively the above conclusion in Fig.4(b).

4 EXPERIMENTS

In this section, we conduct some experiments to answer the following research questions (RQs):

- **RQ1 (Comparison experiments):** How does the clustering performance of TIC compare with that of other state-of-the-art time-series clustering algorithms?
- **RQ2 (Ablation study):** How do the various components of TIC contribute to its performance?
- **RQ3 (Embedding evaluation):** How about using the learned embedding for other downstream tasks (e.g., classification, anomaly detection)?
- **RQ4 (Hyperparameter sensitivity):** what about the hyperparameter sensitivity of TIC?

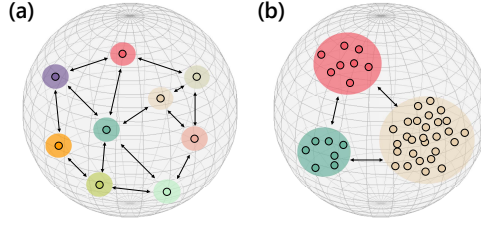


Figure 3: Illustration of the embedding space with (a) and without (b) the interactive loss. Because only the augmentation of each sample is considered when determining positive pairs, it may cause the class collision issue in the embedding space (shown in (a)), i.e., each cluster consists of only one sample (Arora et al., 2019). Optimizing the $\mathcal{L}^{\text{Inte}}$ allows k -means to provide some basic information about the clusters and TIC to achieve better performance (shown in (b)).

Benchmark datasets. We use 136 benchmark (widely used in related work e.g., (Paparrizos & Gravano, 2015; Zhang et al., 2022)) time-series datasets to evaluate the TIC model. Eight datasets are collected from the real-world and the remaining 128 ones are from the UCR database (Chen et al., 2015). Eight real-world datasets are multi-variate or univariate, and their application scenarios include EEG and ECG analyses, and mechanical faulty detection. The description of these benchmark datasets is shown in Table 1. More details about the datasets are shown in Appendix A.3.

Table 1: Specific description for the 136 benchmark datasets. The detailed description of the 128 UCR datasets can be found in (Chen et al., 2015). #Sam, #Clu and #Cha denote the number of samples, clusters and channels, respectively.

Dataset	#Sam	#Clu	#Cha	Length
EMG	204	3	1	1,500
ECG	43,673	4	1	1,500
HAR	10,299	6	9	128
Gesture	560	8	3	315
FD-A	8,184	3	1	5,120
FD-B	13,640	3	1	5,120
SleepEEG	371,055	5	1	200
Epilepsy	11,500	2	1	178
UCR (128 datasets)	[40, 16,637]	[2, 60]	1	[15, 2,844]

Baselines. To answer RQ1, we consider the following 10 time-series clustering methods, i.e., **DTCR** (Ma et al., 2019b), **k-shape** (Paparrizos & Gravano, 2015), **SOM-VAE** (Fortuin et al., 2020), **STCN** (Ma et al., 2022), **TMEK** (Tang et al., 2021), **TNC** (Tonekaboni et al., 2022), **TS3C_{ch}** (Guijo-Rubio et al., 2021), **UMAP** (Pélat et al., 2023), **USSL** (Zhang et al., 2019) and **VLSC** (Duan & Guo, 2023) in the comparison experiment.

To answer RQ3, we evaluate the embeddings learned by TIC on the other downstream tasks (i.e., classification and anomaly detection). The included baselines are 8 unsupervised representation learning methods, i.e., **activity2vec** (Aggarwal et al., 2019), **BTSF** (Yang & Hong, 2022), **CLOCS** (Kiyasseh et al., 2021), **MHCCL** (Meng et al., 2023), **TFC** (Zhang et al., 2022), **Triplet** (Franceschi et al., 2019), **TS2Vec** (Yue et al., 2022) and **TSTCC** (Eldele et al., 2022). More details of the baselines are shown in Appendix A.4.

Implementation details. In our TIC model, we use two 2-layer Transformer (Vaswani et al., 2017) as backbones for encoders $H_T(\cdot)$ and $H_F(\cdot)$. FC_T and FC_F contain three fully-connected layers with

hidden dimensions $d_1 = L$ (time-series length), $d_2 = 128$ and $d_3 = C$ (number of clusters), where the softmax layer is replaced with the RELU to ensure that the network outputs are non-negative values. These two FC modules do not share any parameters. The full spectrum (symmetrical) is used to guarantee that $\mathbf{x}^{\text{Time}}$ and $\mathbf{x}^{\text{Freq}}$ have the same number of dimensions, when transforming the time domain to frequency domain. The Adam optimizer with a learning rate of $\{0.0001, 0.0002, 0.0003\}$ and 2-norm penalty coefficient of 0.0005 is used. We use the batch size of $\{8, 16, 32, 64, 128\}$ according to the dataset size, and use the training epoch of 100. We set $\beta = 1$, $\gamma \in \{0.5, 1.2\}$ for the frequency augmentation and $\tau = 0.2$ in loss functions (3) and (4). The balance hyperparameter ι in Eq.(13) is gradually increased to prevent TIC model from paying too much attention to the KL divergence in the initial training stage. The pseudo-labels are updated every $t = 20$ epochs. We use the code provided in the corresponding paper, and the hyperparameters are finely tuned within the configuration provided therein based on the unsupervised metric Davies-Bouldin Index for fair comparison. The supervised ARI, NMI, ACC metrics are considered. All models are implemented with PyTorch on an NVIDIA A100 Tensor Core GPU. In Appendix A.5, the statistical comparison result derived from the Friedman test and Nemenyi test is provided.

RQ1. Comparison experiment. We show the comparison results between TIC and the 10 time-series clustering baselines in Table 5, where the ARI clustering metric are used (Rand, 1971). We recall that the larger the metrics, the better the clustering performance. After being fed to the correct number C of clusters, each algorithm is run 5 times and the results are recorded in the form of $\text{mean} \pm \text{std. deviation}$. In particular, the average ARI and the corresponding average std. deviations on the 128 UCR datasets are reported. The results w.r.t. NMI, ACC and running time are provided in Appendix A.6.

Table 2: ARI of different algorithms on benchmark datasets. The $\bullet/\circ$ indicates whether TIC is statistically superior/inferior to a certain comparing baseline based on the paired t -test at a 0.05 significance level. The statistics of win/tie/loss are shown in the last row of each sub-table. The best and the second-best results, between which the performance gaps are shown in the row named “gap” in each sub-table, are colored blue and red.

ARI	EMG	ECG	HAR	Gesture	FD-A	FD-B	SleepEEG	Epilepsy	UCR
DTCR	.782 $\pm$.02 $\bullet$	.812 $\pm$.01 $\bullet$	.584 $\pm$.02 $\bullet$	.874 $\pm$.03 $\bullet$	.882$\pm$.02$\bullet$	.801 $\pm$.01 $\bullet$	.571 $\pm$.01 $\bullet$	.872 $\pm$.02 $\bullet$	.634 $\pm$.02 $\bullet$
k-shape	.684 $\pm$.02 $\bullet$	.578 $\pm$.01 $\bullet$	.765 $\pm$.03	.723 $\pm$.02 $\bullet$	.860 $\pm$.01 $\bullet$	.811$\pm$.02$\bullet$	.806 $\pm$.01 $\bullet$	.909 $\pm$.02 $\bullet$	.802 $\pm$.03 $\bullet$
SOM-VAE	.845 $\pm$.03 $\bullet$	.812 $\pm$.02 $\bullet$	.774$\pm$.01	.921$\pm$.01	.824 $\pm$.03 $\bullet$	.732 $\pm$.02 $\bullet$	.816 $\pm$.02 $\bullet$	.931 $\pm$.01 $\bullet$	.813 $\pm$.02 $\bullet$
STCN	1$\pm$0	.730 $\pm$.02 $\bullet$	.669 $\pm$.02 $\bullet$	.825 $\pm$.01 $\bullet$	.879 $\pm$.02 $\bullet$	.803 $\pm$.01 $\bullet$	.815 $\pm$.02 $\bullet$	.770 $\pm$.01 $\bullet$	.541 $\pm$.02 $\bullet$
TMEK	.651 $\pm$.02 $\bullet$	.706 $\pm$.02 $\bullet$	.690 $\pm$.01 $\bullet$	.807 $\pm$.01 $\bullet$	.682 $\pm$.02 $\bullet$	.782 $\pm$.01 $\bullet$	.741 $\pm$.02 $\bullet$	.901 $\pm$.02 $\bullet$	.642 $\pm$.01 $\bullet$
TNC	.751 $\pm$.01 $\bullet$	.805 $\pm$.02 $\bullet$	.693 $\pm$.01 $\bullet$	.851 $\pm$.02 $\bullet$	.780 $\pm$.03 $\bullet$	.593 $\pm$.01 $\bullet$	.741 $\pm$.02 $\bullet$	.725 $\pm$.01 $\bullet$	.707 $\pm$.01 $\bullet$
TS3C _{ch}	.703 $\pm$.02 $\bullet$	.813$\pm$.02$\bullet$	.706 $\pm$.01 $\bullet$	.852 $\pm$.01 $\bullet$	.711 $\pm$.02 $\bullet$	.682 $\pm$.01 $\bullet$	.865$\pm$.01	.939 $\pm$.02 $\bullet$	.851$\pm$.02
UMAP	.859 $\pm$.01 $\bullet$	.791 $\pm$.01 $\bullet$	.643 $\pm$.02 $\bullet$	.852 $\pm$.02 $\bullet$	.725 $\pm$.01 $\bullet$	.785 $\pm$.01 $\bullet$	.863 $\pm$.01 $\bullet$	.972$\pm$.02	.848 $\pm$.01 $\bullet$
USSL	.876$\pm$.01$\bullet$	.745 $\pm$.01 $\bullet$	.621 $\pm$.02 $\bullet$	.597 $\pm$.02 $\bullet$	.622 $\pm$.03 $\bullet$	.592 $\pm$.03 $\bullet$	.802 $\pm$.02 $\bullet$	.925 $\pm$.02 $\bullet$	.710 $\pm$.01 $\bullet$
VLSC	1$\pm$0	.805 $\pm$.01 $\bullet$	.611 $\pm$.02 $\bullet$	.823 $\pm$.02 $\bullet$	.654 $\pm$.02 $\bullet$	.498 $\pm$.01 $\bullet$	.796 $\pm$.01 $\bullet$	.942 $\pm$.02 $\bullet$	.757 $\pm$.01 $\bullet$
TIC (ours)	1$\pm$0	.879$\pm$.01	.796$\pm$.01	.931$\pm$.01	.939$\pm$.01	.853$\pm$.02	.896$\pm$.01	.980$\pm$.01	.883$\pm$.01
gap	.124	.066	.022	.010	.057	.042	.031	.008	.032
win/tie/loss	8/2/0	10/0/0	8/2/0	9/1/0	10/0/0	10/0/0	8/2/0	8/2/0	9/1/0

Overall, our TIC model wins 80 and has tied performance on 10 out of 90 trials, when it is statistically compared with 10 baselines based on three metrics. On average, our TIC model claims a large performance gap of 0.043 over the best baselines. Concretely, the largest performance gap is 0.0124 on the EMG dataset and the smallest one is 0.008 on the Epilepsy dataset. Only on the EMG dataset, STCN and VLSC yield the totally correct clustering result, and achieve the equal ARI with TIC. One potential explanation is that EMG is a simple dataset with only 3 clusters and 204 samples, and thus it may be easy to be modeled. On the ECG dataset, TIC outperforms the strongest baselines by a large margin of 0.066. Because ECG includes four clusters consisting of 43,673 samples that lead to a more complex clustering task, a combination of clustering information in the time and frequency domains results in better performance. In summary, the better performance of TIC model can be attributed to (1) Both time-domain and frequency-domain contrastive losses are considered; (2) Quantification of sample-specific uncertainty reduces the wrong assignment in clustering decision; (3) Dempster’s rule appropriately combines the time- and frequency-based clustering results; (4) the class collision issue can be improved by including the interactive loss $\mathcal{L}^{\text{Inte}}$ in $\mathcal{L}^{\text{TIC}}$. The ablation study (more results shown in Appendix A.7) experimentally demonstrates the four points above.

RQ2. Ablation study. We conduct ablation studies to evaluate the importance of every component in the developed TIC model. Concretely, we compare TIC model with its 3 variants: W/o $\mathcal{L}^{\text{Time}}/\mathcal{L}^{\text{Freq}}$: the loss function $\mathcal{L}^{\text{Time}}/\mathcal{L}^{\text{Freq}}$ is removed from the $\mathcal{L}^{\text{TIC}}$, i.e., the time-

based/frequency-based contrastive module is removed; W/o $\mathcal{L}^{\text{Inte}}$: the loss function $\mathcal{L}^{\text{Inte}}$ is removed from the $\mathcal{L}^{\text{TIC}}$, i.e., k -means does not provide pseudo-labels for the training.

Table 3: ARI values (mean $\pm$ std.deviation) of different variants of TIC.

ARI	EMG	ECG	HAR	Gesture	FD-A	FD-B	SleepEEG	Epilepsy	UCR	Average
W/o $\mathcal{L}^{\text{Time}}$	.958 $\pm$.01	.815 $\pm$.02	.702 $\pm$.01	.846 $\pm$.02	.859 $\pm$.02	.745 $\pm$.02	.687 $\pm$.02	.899 $\pm$.01	.756 $\pm$.01	0.807
W/o $\mathcal{L}^{\text{Freq}}$	.934 $\pm$.01	.827 $\pm$.01	.754 $\pm$.01	.798 $\pm$.02	.547 $\pm$.02	.609 $\pm$.01	.781 $\pm$.02	.854 $\pm$.01	.716 $\pm$.01	0.758
W/o $\mathcal{L}^{\text{Inte}}$	.042 $\pm$.03	.121 $\pm$.02	.098 $\pm$.01	.057 $\pm$.03	.069 $\pm$.02	.210 $\pm$.01	.009 $\pm$.02	.147 $\pm$.02	.059 $\pm$.01	0.090
TIC (Full model)	1.0	.879 $\pm$.01	.796 $\pm$.01	.931 $\pm$.01	.939 $\pm$.01	.853 $\pm$.02	.896 $\pm$.01	.980 $\pm$.01	.883 $\pm$.01	0.906

The results of the ablation study are reported in Table 3. As can be seen, removing $\mathcal{L}^{\text{Time}}$ and $\mathcal{L}^{\text{Freq}}$ leads to performance degradation (average ARI) of $0.912 - 0.807 = 0.105$ and $0.912 - 0.758 = 0.154$, respectively. In particular, on the datasets {FD-A, FD-B} with a high sampling frequency of 64k Hz, removing $\mathcal{L}^{\text{Freq}}$ cause more pronounced performance degradation, e.g., the ARI values of W/o $\mathcal{L}^{\text{Time}}$ and $\mathcal{L}^{\text{Freq}}$ are 0.859 and 0.547 on FD-A dataset. One can conclude that the time- and frequency-based contrastive modules have almost the same contribution to the whole TIC model. The variant W/o $\mathcal{L}^{\text{Inte}}$ has the lowest ARI on all the datasets. This is because every positive pair consists only of the sample and its augmentation after removing the $\mathcal{L}^{\text{Inte}}$, i.e., losing the basic clustering information provided by pseudo-labels. In this case, the class collision issue seriously degrades the clustering performance as discussed in Remark 2. In Fig.4(b), we further show the cosine distances among samples belonging to different clusters. One can see that considering the interactive loss $\mathcal{L}^{\text{Inte}}$ (with $\mathcal{L}^{\text{Inte}}$ in Fig.4(b)) can increase the inter-cluster cosine distance with a remarkable gap of 20.4% (i.e., from 1.965 to 2.365 on ECG dataset). It shows that minimizing the $\mathcal{L}^{\text{Inte}}$ indeed increases the distinctiveness of learned embeddings and bring forward better cluster performance.

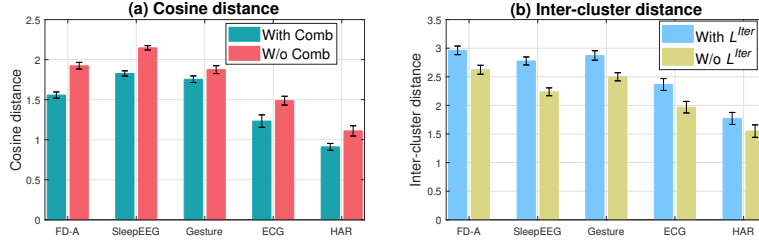


Figure 4: Cosine distance (a) between time- and frequency-based embeddings of the same sample considering the combination (With comb, i.e., full TIC model) and without the combination (W/o Comb) of results from time and frequency domains. The lower distance denotes better encoder learning, because $\mathbf{g}_i^{\text{time}}$ and $\mathbf{g}_i^{\text{freq}}$ are closer in embedding space. Cosine distance (b) among the concatenated embeddings $[\mathbf{g}_i^{\text{time}}; \mathbf{g}_i^{\text{freq}}]$ of samples belonging to different clusters. The larger inter-cluster distance represents a better clustering result.

RQ3. Embedding evaluation. We evaluate the embedding learned by TIC model by performing two other downstream tasks: classification and anomaly detection. The results w.r.t. anomaly detection are shown in Appendix A.8. We follow the same protocol as (Franceschi et al., 2019; Tonekaboni et al., 2022), where a multi-class SVM with RBF kernel and a linear classifier are trained on top of the embeddings learned by different models. The 5-fold cross-validation is adopted to train the SVM and the linear classifier. In TIC, the time- and frequency-based embeddings are concatenated $[\mathbf{g}^{\text{Time}}; \mathbf{g}^{\text{Freq}}]$. Beyond the aforementioned unsupervised representation learning methods, we consider a K-nearest neighbor classifier ($K = 5$) equipped with DTW metric (KNN_{DTW}) as another baseline. The evaluation results are summarized in Table 4, where the results w.r.t the linear classifier and SVM are shown in the first and second sub-tables. As can be seen, TIC achieves the best performance (colored blue) in 13 out of 18 cases and the second-best performance in another 5 cases. In particular, TIC shows the highest accuracy of 0.9654 when training an SVM classifier, which yields a margin of 4.3% over the best baseline activity2vec (0.9257). It shows that TIC adequately leverages the information from time and frequency domains to provide more fine-grained embeddings for discriminant. Besides, almost all the unsupervised representation learning baselines have higher accuracy than KNN_{DTW} , illustrating the dominance of neural networks in representation learning.

Table 4: Accuracy (mean $\pm$ std.deviation) of different methods on benchmark datasets. The results w.r.t the linear classifier and the multi-class SVM are shown in the first and second sub-tables. KNN_{DTW} denotes a K-nearest neighbor classifier equipped with DTW metric.

Linear	EMG	ECG	HAR	Gesture	FD-A	FD-B	SleepEEG	Epilepsy	UCR
KNN _{DTW}	.8452 $\pm$.089	.6712 $\pm$.053	.7152 $\pm$.065	.5746 $\pm$.031	.6784 $\pm$.053	.3462 $\pm$.078	.6745 $\pm$.043	.8756 $\pm$.074	.7127 $\pm$.054
activity2vec	.9587 $\pm$.011	.7853 $\pm$.036	.9258 $\pm$.073	.6547 $\pm$.026	.8964$\pm$.039	.6578 $\pm$.038	.8941 $\pm$.085	.9245 $\pm$.004	.9123$\pm$.030
BTSF	.9578 $\pm$.057	.8514 $\pm$.034	.9463$\pm$.068	.8637$\pm$.079	.8875 $\pm$.068	.7123 $\pm$.029	.8745 $\pm$.075	.9521 $\pm$.056	.8654 $\pm$.028
CLOCS	.8924 $\pm$.049	.7546 $\pm$.056	.8954 $\pm$.038	.4852 $\pm$.023	.8456 $\pm$.087	.7214 $\pm$.036	.8875 $\pm$.074	.9485 $\pm$.029	.8745 $\pm$.054
MHCL	.9643 $\pm$.022	.7765 $\pm$.031	.9164 $\pm$.034	.7756 $\pm$.055	.8382 $\pm$.064	.7936 $\pm$.051	.9145$\pm$.029	.9654 $\pm$.005	.7363 $\pm$.098
TFC	.9854$\pm$.007	.9087$\pm$.045	.9245 $\pm$.054	.7955 $\pm$.024	.8657 $\pm$.048	.8834$\pm$.038	.8921 $\pm$.051	.9478 $\pm$.025	.7982 $\pm$.069
Triplet	.9338 $\pm$.077	.8937 $\pm$.064	.9054 $\pm$.029	.6953 $\pm$.044	.8542 $\pm$.055	.8377 $\pm$.047	.8999 $\pm$.035	.9785$\pm$.009	.8032 $\pm$.013
TS2Vec	.8687 $\pm$.035	.8546 $\pm$.054	.4887 $\pm$.067	.8365 $\pm$.081	.8922 $\pm$.039	.8643 $\pm$.059	.8456 $\pm$.062	.9087 $\pm$.029	.8266 $\pm$.002
TSTCC	.9485 $\pm$.014	.7481 $\pm$.011	.8804 $\pm$.025	.7685 $\pm$.024	.8547 $\pm$.039	.8598 $\pm$.011	.8300 $\pm$.007	.9158 $\pm$.009	.7591 $\pm$.003
TIC (ours)	.9785$\pm$.057	.9132$\pm$.056	.9571$\pm$.084	.8795$\pm$.067	.9215$\pm$.069	.8965$\pm$.027	.9013$\pm$.039	.9874$\pm$.055	.9251$\pm$.036
SVM	EMG	ECG	HAR	Gesture	FD-A	FD-B	SleepEEG	Epilepsy	UCR
KNN _{DTW}	.8452 $\pm$.089	.6712 $\pm$.053	.7152 $\pm$.065	.5746 $\pm$.031	.6784 $\pm$.053	.3462 $\pm$.078	.6745 $\pm$.043	.8756 $\pm$.074	.7127 $\pm$.054
activity2vec	.9651 $\pm$.035	.8324 $\pm$.067	.9257$\pm$.054	.6871 $\pm$.036	.9031 $\pm$.089	.6982 $\pm$.056	.8874 $\pm$.074	.9483 $\pm$.036	.9245$\pm$.057
BTSF	.9687 $\pm$.045	.8789 $\pm$.037	.9128 $\pm$.061	.8843$\pm$.039	.8992 $\pm$.066	.7536 $\pm$.052	.8905 $\pm$.037	.9569 $\pm$.045	.8763 $\pm$.079
CLOCS	.9214 $\pm$.066	.7895 $\pm$.074	.9214 $\pm$.056	.5123 $\pm$.081	.8517 $\pm$.036	.7541 $\pm$.063	.8907 $\pm$.046	.9681$\pm$.039	.8852 $\pm$.028
MHCL	.9695 $\pm$.035	.7854 $\pm$.076	.9237 $\pm$.046	.7986 $\pm$.048	.8736 $\pm$.078	.8065 $\pm$.031	.8943$\pm$.032	.9533 $\pm$.067	.7629 $\pm$.054
TFC	.9743$\pm$.085	.9316$\pm$.037	.9158 $\pm$.045	.8214 $\pm$.061	.8702 $\pm$.047	.9317$\pm$.036	.8795 $\pm$.026	.9538 $\pm$.054	.8369 $\pm$.091
Triplet	.9502 $\pm$.045	.9125 $\pm$.067	.9056 $\pm$.057	.8506 $\pm$.075	.8722 $\pm$.049	.8697 $\pm$.094	.8907 $\pm$.056	.9654 $\pm$.048	.8537 $\pm$.033
TS2Vec	.8794 $\pm$.024	.8874 $\pm$.079	.5638 $\pm$.047	.8574 $\pm$.067	.9201$\pm$.043	.8932 $\pm$.046	.8645 $\pm$.033	.9268 $\pm$.076	.8437 $\pm$.070
TSTCC	.9542 $\pm$.036	.7896 $\pm$.033	.9214 $\pm$.096	.7982 $\pm$.019	.8964 $\pm$.056	.8609 $\pm$.087	.8514 $\pm$.076	.9235 $\pm$.044	.7978 $\pm$.021
TIC (ours)	.9855$\pm$.065	.9236$\pm$.049	.9654$\pm$.037	.9025$\pm$.081	.9187$\pm$.064	.9222$\pm$.074	.9098$\pm$.026	.9745$\pm$.037	.9391$\pm$.056

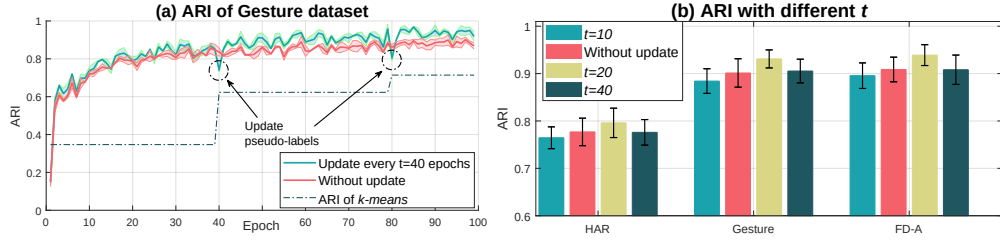


Figure 5: ARI in every epoch with updating pseudo-labels every 40 epochs and without update (a). ARI decreases slightly each time the pseudo-labels are updated by k -means at 40th and 80th epochs (shown by black circles). ARI with different t is shown in (b). The frequent updating of pseudo-labels ($t = 10$) degrades the clustering performance of TIC model.

RQ4. Sensitivity of the hyperparameter t . As shown in Fig.1, TIC model updates the pseudo-labels generated by k -means every t epochs. In Fig.5(a), we show the ARI of TIC in every epoch with $t = 40$ and without update for the Gesture dataset. It can be seen that the clustering results generated by k -means keep being improved (ARI=0.347 with epoch $\in [1, 39]$, ARI=0.623 with epoch $\in [40, 79]$, ARI=0.714 with epoch $\in [80, 99]$) because the embedding learning is constantly optimized. This leads to basic clustering information of higher quality being considered in the interaction loss. Although the ARI of TIC decreases temporarily at each update of the pseudo-label (epoch = 40 and epoch = 80), the final ARI with $t = 40$ is higher than the one without updates. Fig.5(b) shows the ARI of TIC model with different t . Frequent updating ($t = 10$) of pseudo-labels degrades the performance of TIC, even making it lower than the case of no updating, as the embedding learning is destabilized. we also provide the visualization of the clustering results and the effects of various data augmentation techniques in Appendix A.9 and A.10.

5 CONCLUSION

This paper proposes a trusted and interactive clustering model for time-series data, named TIC, leveraging evidence theory to combine time- and frequency-based information. TIC optimizes the contrastive loss from time and frequency domains, and an interactive loss calculated based on the pseudo-labels. Uncertainty in time- and frequency-based clustering results are quantified by mass functions that are combined by Dempster’s rule to produce the trusted clustering results. Experimental results show the superior clustering performance of TIC brought by the combination of clustering results from time and frequency domains, as well as the consideration of interactive loss. The embedding learned by TIC is also shown to perform well on the classification and anomaly detection tasks.

REFERENCES

- Karan Aggarwal, Shafiq Joty, Luis Fernandez-Luque, and Jaideep Srivastava. Adversarial unsupervised representation learning for activity time-series. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 33, pp. 834–841, 2019.
- Ralph G Andrzejak, Klaus Lehnertz, Florian Mormann, Christoph Rieke, Peter David, and Christian E Elger. Indications of nonlinear deterministic and finite-dimensional structures in time series of brain electrical activity: Dependence on recording region and brain state. *Physical Review E*, 64(6):061907, 2001.
- Davide Anguita, Alessandro Ghio, Luca Oneto, Xavier Parra, Jorge Luis Reyes-Ortiz, et al. A public domain dataset for human activity recognition using smartphones. In *Esann*, volume 3, pp. 3, 2013.
- Sanjeev Arora, Hrishikesh Khandeparkar, Mikhail Khodak, Orestis Plevrakis, and Nikunj Saunshi. A theoretical analysis of contrastive unsupervised representation learning. In *36th International Conference on Machine Learning, ICML 2019*, pp. 9904–9923. International Machine Learning Society (IMLS), 2019.
- Christopher M Bishop and Nasser M Nasrabadi. *Pattern recognition and machine learning*, volume 4. Springer, 2006.
- E Oran Brigham. *The fast Fourier transform and its applications*. Prentice-Hall, Inc., 1988.
- Ricardo JGB Campello, Davoud Moulavi, and Jörg Sander. Density-based clustering based on hierarchical density estimates. In *Pacific-Asia Conference on Knowledge Discovery and Data Mining*, pp. 160–172. Springer, 2013.
- Soravit Changpinyo, Piyush Sharma, Nan Ding, and Radu Soricut. Conceptual 12m: Pushing web-scale image-text pre-training to recognize long-tail visual concepts. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 3558–3568, 2021.
- Bertrand Charpentier, Daniel Zügner, and Stephan Günnemann. Posterior network: Uncertainty estimation without ood samples via density-based pseudo-counts. *Advances in Neural Information Processing Systems*, 33:1356–1367, 2020.
- Ting Chen, Simon Kornblith, Mohammad Norouzi, and Geoffrey Hinton. A simple framework for contrastive learning of visual representations. In *International Conference on Machine Learning*, pp. 1597–1607. PMLR, 2020.
- Yanping Chen, Eamonn Keogh, Bing Hu, Nurjahan Begum, Anthony Bagnall, Abdullah Mueen, and Gustavo Batista. The ucr time series classification archive, July 2015. www.cs.ucr.edu/~eamonn/time_series_data/.
- Gari D Clifford, Chengyu Liu, Benjamin Moody, H Lehman Li-wei, Ikaro Silva, Qiao Li, AE Johnson, and Roger G Mark. Af classification from a short single lead ecg recording: The physionet/computing in cardiology challenge 2017. In *2017 Computing in Cardiology (CinC)*, pp. 1–4. IEEE, 2017.
- Leon Cohen. *Time-frequency analysis*, volume 778. Prentice hall New Jersey, 1995.
- AP DEMPSTER. Upper and lower probabilities induced by a multivalued mapping. *Annals of Mathematical Statistics*, 38:325–339, 1967.
- Thierry Denœux. Conjunctive and disjunctive combination of belief functions induced by nondistinct bodies of evidence. *Artificial Intelligence*, 172(2-3):234–264, 2008.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805*, 2018.
- Yuan-Wei Du and Jiao-Jiao Zhong. Generalized combination rule for evidential reasoning approach and dempster–shafer theory of evidence. *Information Sciences*, 547:1201–1232, 2021.

- Jiangyong Duan and Lili Guo. Variable-length subsequence clustering in time series. *IEEE Transactions on Knowledge and Data Engineering*, 34(2):983–995, 2023.
- Didler Dubois and Henri Prade. Representation and combination of uncertainty with belief functions and possibility measures. *Computational intelligence*, 4(3):244–264, 1988.
- Emadeldeen Eldele, Mohamed Ragab, Zhenghua Chen, Min Wu, Chee Keong Kwoh, Xiaoli Li, and Cuntai Guan. Time-series representation learning via temporal and contextual contrasting. *arXiv preprint arXiv:2106.14112*, 2021.
- Emadeldeen Eldele, Mohamed Ragab, Zhenghua Chen, Min Wu, Chee-Keong Kwoh, Xiaoli Li, and Cuntai Guan. Self-supervised contrastive representation learning for semi-supervised time-series classification. *arXiv preprint arXiv:2208.06616*, 2022.
- Leonardo N Ferreira and Liang Zhao. Time series clustering via community detection in networks. *Information Sciences*, 326:227–242, 2016.
- Patrick Flandrin. *Time-frequency/time-scale analysis*. Academic press, 1998.
- Mihai Cristian Florea, Anne-Laure Jousselme, Éloi Bossé, and Dominic Grenier. Robust combination rules for evidence theory. *Information Fusion*, 10(2):183–197, 2009.
- V Fortuin, M Hüser, F Locatello, H Strathmann, and G Rätsch. Som-vae: Interpretable discrete representation learning on time series. In *International Conference on Learning Representations (ICLR 2020)*, 2020.
- Jean-Yves Franceschi, Aymeric Dieuleveut, and Martin Jaggi. Unsupervised scalable representation learning for multivariate time series. *Advances in neural information processing systems*, 32, 2019.
- Yarin Gal and Zoubin Ghahramani. Dropout as a bayesian approximation: Representing model uncertainty in deep learning. In *International Conference on Machine Learning*, pp. 1050–1059. PMLR, 2016.
- Ary L Goldberger, Luis AN Amaral, Leon Glass, Jeffrey M Hausdorff, Plamen Ch Ivanov, Roger G Mark, Joseph E Mietus, George B Moody, Chung-Kang Peng, and H Eugene Stanley. Physiobank, physiotoolkit, and physionet: components of a new research resource for complex physiologic signals. *circulation*, 101(23):e215–e220, 2000.
- Chaoyu Gong, Yongbin Li, Di Fu, Yong Liu, Pei-hong Wang, and Yang You. Self-reconstructive evidential clustering for high-dimensional data. In *2022 IEEE 38th International Conference on Data Engineering (ICDE)*, pp. 2099–2112. IEEE, 2022.
- D Guijo-Rubio, AM Duran-Rosal, PA Gutierrez, A Troncoso, and C Hervas-Martinez. Time-series clustering based on the characterization of segment typologies. *IEEE Transactions on Cybernetics*, 51(11):5409–5422, 2021.
- Zongbo Han, Changqing Zhang, Huazhu Fu, and Joey Tianyi Zhou. Trusted multi-view classification. *arXiv preprint arXiv:2102.02051*, 2021.
- Zongbo Han, Changqing Zhang, Huazhu Fu, and Joey Tianyi Zhou. Trusted multi-view classification with dynamic evidential fusion. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 45(2):2551–2566, 2022.
- Toshitaka Hayashi, Dalibor Cimr, Filip Studnička, Hamido Fujita, Damián Bušovský, Richard Cimrler, and Ali Selamat. Distance-based one-class time-series classification approach using local cluster balance. *Expert Systems with Applications*, 235:121201, 2024.
- Rob J Hyndman and George Athanasopoulos. *Forecasting: principles and practice*. OTexts, 2018.
- Fumitada Itakura. Minimum prediction residual principle applied to speech recognition. *IEEE Transactions on Acoustics, Speech, and Signal Processing*, 23(1):67–72, 1975.

- Audun Jsang. *Subjective Logic: A formalism for reasoning under uncertainty*. Springer Publishing Company, Incorporated, 2018.
- Bob Kemp, Aeilko H Zwinderman, Bert Tuk, Hilbert AC Kamphuisen, and Josefen JL Obery. Analysis of a sleep-dependent neuronal feedback loop: the slow-wave microcontinuity of the eeg. *IEEE Transactions on Biomedical Engineering*, 47(9):1185–1194, 2000.
- Dani Kiyasseh, Tingting Zhu, and David A Clifton. Clocs: Contrastive learning of cardiac signals across space, time, and patients. In *International Conference on Machine Learning*, pp. 5606–5615. PMLR, 2021.
- Anna-Kathrin Kopetzki, Bertrand Charpentier, Daniel Zügner, Sandhya Giri, and Stephan Günnemann. Evaluating robustness of predictive uncertainty estimation: Are dirichlet-based models reliable? In *International Conference on Machine Learning*, pp. 5707–5718. PMLR, 2021.
- Balaji Lakshminarayanan, Alexander Pritzel, and Charles Blundell. Simple and scalable predictive uncertainty estimation using deep ensembles. *Advances in neural information processing systems*, 30, 2017.
- Christian Lessmeier, James Kuria Kimotho, Detmar Zimmer, and Walter Sextro. Condition monitoring of bearing damage in electromechanical drive systems by using motor current signals of electric motors: A benchmark data set for data-driven classification. In *PHM Society European Conference*, volume 3, 2016.
- Jiayang Liu, Lin Zhong, Jehan Wickramasuriya, and Venu Vasudevan. uwave: Accelerometer-based personalized gesture recognition and its applications. *Pervasive and Mobile Computing*, 5(6):657–675, 2009.
- Xinwang Liu, Xinzhong Zhu, Miaomiao Li, Lei Wang, Chang Tang, Jianping Yin, Dinggang Shen, Huaimin Wang, and Wen Gao. Late fusion incomplete multi-view clustering. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 41(10):2410–2423, 2018.
- Q Ma, S Li, W Zhuang, J Wang, and D Zeng. Self-supervised time series clustering with model-based dynamics. *IEEE Transactions on Neural Networks and Learning Systems*, 32(9):3942–3955, 2022.
- Qianli Ma, Sen Li, Lifeng Shen, Jiabing Wang, Jia Wei, Zhiwen Yu, and Garrison W Cottrell. End-to-end incomplete time-series modeling from linear memory of latent variables. *IEEE Transactions on Cybernetics*, 50(12):4908–4920, 2019a.
- Qianli Ma, Jiawei Zheng, Sen Li, and Gary W Cottrell. Learning representations for time series clustering. *Advances in neural information processing systems*, 32, 2019b.
- Wenjun Ma, Yuncheng Jiang, and Xudong Luo. A flexible rule for evidential combination in dempster–shafer theory of evidence. *Applied Soft Computing*, 85:105512, 2019c.
- Andrey Malinin, Bruno Mlodozienec, and Mark Gales. Ensemble distribution distillation. In *International Conference on Learning Representations*, 2019.
- Qianwen Meng, Hangwei Qian, Yong Liu, Lizhen Cui, Yonghui Xu, and Zhiqi Shen. Mhcccl: Masked hierarchical cluster-wise contrastive learning for multivariate time series. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 37, pp. 9153–9161, 2023.
- Matthew Middlehurst, Patrick Schäfer, and Anthony Bagnall. Bake off redux: a review and experimental evaluation of recent time series classification algorithms. *Data Mining and Knowledge Discovery*, pp. 1–74, 2024.
- John Paparrizos and Luis Gravano. k-shape: Efficient and accurate clustering of time series. In *Proceedings of the 2015 ACM SIGMOD International Conference on Management of Data*, pp. 1855–1870, 2015.
- Clément Péalat, Guillaume Bouleux, and Vincent Cheutet. Improved time series clustering based on new geometric frameworks. *Pattern Recognition*, 124:108423, 2023.

- François Petitjean, Alain Ketterlin, and Pierre Gançarski. A global averaging method for dynamic time warping, with applications to clustering. *Pattern Recognition*, 44(3):678–693, 2011.
- William M Rand. Objective criteria for the evaluation of clustering methods. *Journal of the American Statistical association*, 66(336):846–850, 1971.
- Murat Sensoy, Lance Kaplan, and Melih Kandemir. Evidential deep learning to quantify classification uncertainty. *Advances in neural information processing systems*, 31, 2018.
- G. Shafer. *A mathematical theory of evidence*. Princeton university press, 1976.
- Ryan Soklaski, Michael Yee, and Theodoros Tsiligkaridis. Fourier-based augmentations for improved robustness and uncertainty calibration. *arXiv preprint arXiv:2202.12412*, 2022.
- Zhi-gang Su, Thierry Denoeux, Yong-sheng Hao, and Ming Zhao. Evidential k-nn classification with enhanced performance via optimizing a class of parametric conjunctive t-rules. *Knowledge-Based Systems*, 142:7–16, 2018.
- Chi Ian Tang, Ignacio Perez-Pozuelo, Dimitris Spathis, and Cecilia Mascolo. Exploring contrastive learning in human activity recognition for healthcare. *arXiv preprint arXiv:2011.11542*, 2020.
- Yongqiang Tang, Yuan Xie, Xuebing Yang, Jinghao Niu, and Wensheng Zhang. Tensor multi-elastic kernel self-paced learning for time series clustering. *IEEE Transactions on Knowledge & Data Engineering*, 33(03):1223–1237, 2021.
- Sana Tonekaboni, Danny Eytan, and Anna Goldenberg. Unsupervised representation learning for time series with temporal neighborhood coding. *arXiv preprint arXiv:2106.00750*, 2022.
- Laurens Van der Maaten and Geoffrey Hinton. Visualizing data using t-sne. *Journal of Machine Learning Research*, 9(11), 2008.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. *Advances in neural information processing systems*, 30, 2017.
- Siwei Wang, Xinwang Liu, En Zhu, Chang Tang, Jiayuan Liu, Jintao Hu, Jingyuan Xia, and Jianping Yin. Multi-view clustering via late fusion alignment maximization. In *International Joint Conference on Artificial Intelligence*, pp. 3778–3784, 2019.
- Siwei Wang, Xinwang Liu, Li Liu, Sihang Zhou, and En Zhu. Late fusion multiple kernel clustering with proxy graph refinement. *IEEE Transactions on Neural Networks and Learning Systems*, 2021.
- Gerald Woo, Chenghao Liu, Doyen Sahoo, Akshat Kumar, and Steven Hoi. Cost: Contrastive learning of disentangled seasonal-trend representations for time series forecasting. In *International Conference on Learning Representations*, 2021.
- Ronald R Yager. On the dempster-shafer framework and new combination rules. *Information sciences*, 41(2):93–137, 1987.
- Jaewon Yang and Jure Leskovec. Patterns of temporal variation in online media. In *Proceedings of the fourth ACM International Conference on Web Search and Data Mining*, pp. 177–186, 2011.
- Ling Yang and Shenda Hong. Unsupervised time-series representation learning with iterative bi-linear temporal-spectral fusion. In *International Conference on Machine Learning*, pp. 25038–25054. PMLR, 2022.
- Zhihan Yue, Yujing Wang, Juanyong Duan, Tianmeng Yang, Congrui Huang, Yunhai Tong, and Bixiong Xu. Ts2vec: Towards universal representation of time series. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 36, pp. 8980–8987, 2022.
- Kexin Zhang, Qingsong Wen, Chaoli Zhang, Rongyao Cai, Ming Jin, Yong Liu, James Y Zhang, Yuxuan Liang, Guansong Pang, Dongjin Song, et al. Self-supervised learning for time series analysis: Taxonomy, progress, and prospects. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2024.

Qin Zhang, Jia Wu, Peng Zhang, Guodong Long, and Chengqi Zhang. Salient subsequence learning for time series clustering. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 41(09):2193–2207, 2019.

Xiang Zhang, Ziyuan Zhao, Theodoros Tsiligkaridis, and Marinka Zitnik. Self-supervised contrastive pre-training for time series via time-frequency consistency. *Advances in Neural Information Processing Systems*, 35:3988–4003, 2022.

A APPENDIX

A.1 RELATED WORK

We review time-series clustering methods, contrastive learning methods toward time-series data and uncertain learning methods, respectively.

A.1.1 TIME-SERIES CLUSTERING

Time-series clustering methods can be roughly classified into two families: raw-data-based methods and feature-based methods.

Raw-data-based methods. These methods primarily modify the distance metric to accommodate the specific characteristics of time-series data. K-DBA (Petitjean et al., 2011) combines k -means and dynamic time warping (DTW) (Itakura, 1975) to achieve improved alignment. To reveal the temporal dynamics, K-SC method (Yang & Leskovec, 2011) utilizes a similarity metric that is invariant to scaling and shifting. K-Shape (Paparrizos & Gravano, 2015) considers the shapes of time-series by employing a normalized cross-correlation metric. The mentioned methods often exhibit sensitivity to outliers and noise because they consider all time-step points (Ma et al., 2019b).

Feature-based methods. Feature-based methods typically involve two stages, where the input time-series samples are transformed into informative features first and clustering algorithms are then conducted on these features (Tang et al., 2021; Fortuin et al., 2020). In conjunction with k -means, TNC (Tonekaboni et al., 2022) ensures the distribution of signals from the neighborhood is distinguishable from the distribution of non-neighboring signals. Authors in (Tang et al., 2021) map the raw time-series space into multiple kernel spaces via elastic distance measure functions and resort to a self-paced learning paradigm to group time-series samples. Authors in (Pélat et al., 2023) embed the time-series onto the Stiefel manifold to obtain the geometric representations of time-series samples. STCN (Ma et al., 2022) optimizes the feature extraction and clustering simultaneously, through a recurrent neural network and a self-supervised clustering module. However, almost all the feature-based methods do not adopt an information fusion perspective to incorporate both time and frequency domain information for the purpose of clustering.

A.1.2 CONTRASTIVE LEARNING TOWARD TIME-SERIES DATA

Contrastive learning is a well-known form of self-supervised learning and aims to train an encoder that maps original inputs into an embedding space. The objective is to bring positive sample pairs closer together, while pushing negative sample pairs (comprising the original augmentation and an alternative augmentation of a different input sample) apart (Chen et al., 2020). Compared to CV (Changpinyo et al., 2021) and NLP (Devlin et al., 2018), contrastive learning in the context of time-series data has been less explored, mainly due to the difficulty of capturing crucial invariance properties specific to time-series. TF-C (Zhang et al., 2022) expects that time-based and frequency-based representations of the same sample are located close together in the time-frequency space, and embeds the time-based neighborhood of a sample close to its frequency-based neighborhood. BTSF (Yang & Hong, 2022) utilizes sample-level augmentation with a dropout on a time-series sample, and devises the iterative bilinear temporal-spectral fusion to generate discriminative embeddings. CoST (Woo et al., 2021) comprises both time domain and frequency domain contrastive losses to learn seasonal representations for long sequence time-series forecasting. In addition to these three methods involving both time and frequency domains, other methods mainly focus on the augmentations implemented in time domain, such as transformation invariance (e.g., SimCLR (Tang et al., 2020; Chen et al., 2020)) and contextual invariance (e.g., TS2vec (Yue et al., 2022) and TS-TCC

(Eldele et al., 2021)). In previous works, the loss information from the time- and frequency-domain is captured in a compositional way, e.g., TF-C simply sums the loss functions from the two domains and the consistency loss function, and BTSF solely implements the data augmentation in the time domain. To the best of our knowledge, this work is the first one that directly combines time-frequency domain information to leverage information fusion for time-series clustering.

A.1.3 UNCERTAINTY-BASED LEARNING

Some efforts (Gal & Ghahramani, 2016; Lakshminarayanan et al., 2017; Charpentier et al., 2020) have been made to enable the neural network to estimate the uncertainty of the output. Evidential network (Sensoy et al., 2018) incorporates subjective logic to model the Dirichlet distribution. Post-Net (Malinin et al., 2019) utilizes normalizing flow and Bayesian loss during training to estimate uncertainty. TMC (Han et al., 2021) and ETMC (Han et al., 2022) introduce a variational Dirichlet distribution to characterize the distribution of the class probabilities in multi-view classification. Ensemble distribution distillation (Malinin et al., 2019) leverages the predictions of multiple models to estimate the uncertainty. Authors in (Kopetzki et al., 2021) apply median smoothing to the Dirichlet model and enhance the capability of the model to handle adversarial examples. Unlike the above methods, our method is perhaps the first attempt to estimate the uncertainty of the outputs from a neural network oriented to the clustering problem. Further, we fuse the output with uncertainty estimation from the time and frequency domains to obtain trusted clustering results under the framework of evidence theory.

A.2 PROOFS OF PROPOSITIONS 1 AND 2

Proposition 1: A large m_{Ω}^T does not lead to a large m_{Ω}^{Co} , when one of m_c^F is large and m_{Ω}^F is small. In particular, m^{Co} is identical to m^T , if m^F is totally uncertain (i.e., $m_{\Omega}^F = 1$).

Proof:

$$\begin{aligned} m_c^{Co} &= \frac{m_c^T m_{\Omega}^F + m_c^F m_{\Omega}^T + m_c^T m_c^F}{m_{\Omega}^F m_{\Omega}^T + \sum_{v=1}^C m_v^T m_v^F + (1 - m_{\Omega}^T) m_{\Omega}^F + (1 - m_{\Omega}^F) m_{\Omega}^T} \\ &= \frac{m_c^T m_{\Omega}^F + m_c^F m_{\Omega}^T + m_c^T m_c^F}{\sum_{v=1}^C m_v^T m_v^F + m_{\Omega}^T + m_{\Omega}^F - m_{\Omega}^F m_{\Omega}^T} \end{aligned}$$

Considering the worst case with $m_{\Omega}^F = 1, m_v^F = 0, v = 1, 2, \dots, n$, then we can get $m_c^{Co} = m_c^T$. $\square$

Proposition 2: The m_{Ω}^{Co} is monotonically increasing with m_{Ω}^T and m_{Ω}^F .

Proof:

$$\begin{aligned} m_{\Omega}^{Co} &= \frac{m_{\Omega}^T m_{\Omega}^F}{\sum_{v=1}^C (m_v^T m_v^F + m_v^T m_{\Omega}^F + m_v^F m_{\Omega}^T) + m_{\Omega}^T m_{\Omega}^F} \\ &= \frac{1}{\sum_{v=1}^C (\frac{m_v^T m_v^F}{m_{\Omega}^T m_{\Omega}^F} + \frac{m_v^T}{m_{\Omega}^T} + \frac{m_v^F}{m_{\Omega}^F}) + 1} \end{aligned}$$

As can be seen, m_{Ω}^{Co} increases as m_{Ω}^T and m_{Ω}^F increase. $\square$

A.3 DESCRIPTION OF DATASETS

ECG (Clifford et al., 2017) is from the 2017 PhysioNet Challenge focusing on classifying ECG recordings, where single-lead ECG measures four different underlying conditions of cardiac arrhythmias. **EMG** (Goldberger et al., 2000) consists of single-channel Electromyograms (EMGs) recorded from the tibialis anterior muscle of three healthy volunteers suffering from myopathy and neuropathy. **HAR** (Anguita et al., 2013) contains recordings of 30 health volunteers performing daily activities, including walking, walking upstairs, walking downstairs, sitting, standing, and lying. **Gesture** (Liu et al., 2009) contains accelerometer measurements of eight simple gestures that differ based on the paths of hand movement. **FD-A/FD-B** (Lessmeier et al., 2016) corresponds to

Faulty Detection Condition A (FD-A) and Faulty Detection Condition B (FD-B), which are generated by an electromechanical drive system that monitors the condition of rolling bearings and detects their failures. **SleepEEG** (Kemp et al., 2000) are collected from 82 healthy subjects and contains 153 whole-night sleeping electroencephalography (EEG) recordings. **Epilepsy** (Andrzejak et al., 2001) contains single-channel EEG measurements from 500 subjects and the corresponding brain activity is recorded for 23.6 seconds. The details about the UCR datasets can be found in (Chen et al., 2015).

A.4 DESCRIPTION OF BASELINES

Baselines. To answer Q1, we consider the following 10 time-series clustering methods in the comparison experiment.

- **DTCR** (Ma et al., 2019b): it integrates the temporal reconstruction and k -means objective into the seq2seq model. By proposing a fake-sample generation strategy and auxiliary classification task, it can learn cluster-specific temporal representations.
- **k-shape** (Paparrizos & Gravano, 2015): it relies on a scalable iterative refinement procedure and uses a normalized cross-correlation measure to consider the shapes of time-series.
- **SOM-VAE** (Fortuin et al., 2020): it overcomes the non-differentiability in discrete representation learning and presents a gradient-based version of the traditional self-organizing map (SOM) algorithm.
- **STCN** (Ma et al., 2022): it optimizes the feature extraction and clustering simultaneously, where an RNN conducts the reconstruction of time-series and a self-supervised module to obtain the clustering result.
- **TMEK** (Tang et al., 2021): it maps the raw time-series space into multiple kernel spaces via elastic distance measure functions, and resorts to the tensor-constraint-based self-representation subspace clustering approach.
- **TNC** (Tonekaboni et al., 2022): it uses the local smoothness of a signal’s generative process to define neighborhoods in time-series. By using a de-biased contrastive objective, it learns time-series representations that are input to k -means to produce the clustering results.
- **TS3C_{ch}** (Guijo-Rubio et al., 2021): it consists of two stages, where a least squares polynomial technique is first used to segment the time-series and the hierarchical clustering is applied to all the mapped segmentations.
- **UMAP** (Pélat et al., 2023): it embeds the time-series onto higher-dimensional spaces and is in conjunction with HDBSCAN algorithm (Campello et al., 2013) to obtain the results.
- **USSL** (Zhang et al., 2019): it integrates the shapelet learning, shapelet regularization, spectral analysis and pseudo-label to automatically learn shapelets to group time-series samples.
- **VLSC** (Duan & Guo, 2023): it minimizes the inner time-series clustering error under time-series cover constraints, where the time-series lengths can be variable.

To answer Q3, we evaluate the embeddings learned by TIC on the other downstream tasks (i.e., classification and anomaly detection). The included baselines are 8 unsupervised representation learning methods for time-series.

- **activity2vec** (Aggarwal et al., 2019): it learns the representations with three components, the co-occurrence and magnitude of the activity levels in a time-series sample, neighboring context of the time-series, and promoting subject-invariance with adversarial training.
- **BTSF** (Yang & Hong, 2022): it utilizes the sample-level augmentation with a simple dropout on the time-series dataset and devise the iterative bilinear temporal-spectral fusion to encode the affinities of abundant time-frequency pairs.
- **CLOCS** (Kiyasseh et al., 2021): it encourages the time-series representations across space, time, and patients to be similar to one another for the physiological data.
- **MHCCL** (Meng et al., 2023): it exploits semantic information obtained from the hierarchical structure consisting of multiple latent partitions for multivariate time-series.

- **TFC** (Zhang et al., 2022): it learns the representations of time-series by positing that embedding a time-based neighborhood of a sample should be close to its frequency-based neighborhood.
- **Triplet** (Franceschi et al., 2019): it learns general-purpose representations by combining an encoder based on causal dilated convolutions with a novel triplet loss employing time-based negative sampling.
- **TS2Vec** (Yue et al., 2022): it performs contrastive learning in a hierarchical way over augmented context views and obtains a contextual representation for each timestamp.
- **TSTCC** (Eldele et al., 2022): it proposes time-series-specific weak and strong augmentations and learns discriminative representations in a contextual contrasting module.

A.5 MORE STATISTICAL TEST

We also compare the performance of methods using the Friedman test and Nemenyi test by setting the significant level to 0.05, which is shown in the Fig.6. TIC significantly outperforms the baselines except for SOM-VAE and TS3C_{ch} in all cases.

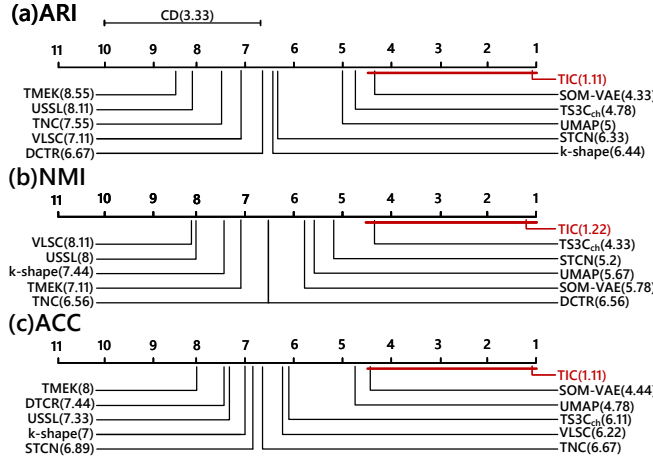


Figure 6: Comparison between TIC and other baselines with the Nemenyi test. The lowest (best) ranks are to the right, and thus methods on the right sides are considered to be better. The groups of baselines that are not significantly different from TIC and are connected by red.

A.6 NMI, ACC AND RUNNING TIME RESULTS

Overall, our TIC model wins 235 and has tied performance on 34 out of 270 trials, when it is statistically compared with 10 baselines based on all the three metrics. Compared to the baselines, TIC consumes comparative runtime but achieves the best clustering accuracy.

A.7 MORE ABLATION STUDY

We conduct other ablation studies to evaluate the importance of using the Dempster’s rule in the developed TIC model. Concretely, we compare TIC model with the following variants:

- **TIC₊**: Dempster’s rule is replaced by the addition + to combine the mass functions $\mathbf{m}^{\text{Time}}$ and $\mathbf{m}^{\text{Freq}}$. For example, if $\mathbf{m}_i^{\text{Time}} = [0.8, 0.1, 0.1]$, $\mathbf{m}_i^{\text{Freq}} = [0.7, 0.2, 0.1]$, the $\mathbf{m}_i^{\text{Comb}}$ is calculated as $m_{i1}^{\text{Comb}} = \frac{0.8+0.7}{0.8+0.7+0.1+0.2+0.1+0.1} = 0.75$, $m_{i2}^{\text{Comb}} = \frac{0.1+0.2}{2} = 0.15$ and $m_{i\Omega}^{\text{Comb}} = \frac{0.1+0.1}{2} = 0.1$;
- **TIC-cau/TIC-bol/TIC-tri**: Dempster’s rule is replaced by other combination rule, i.e., the cautious conjunctive rule (Denoëux, 2008), the bold conjunctive rule (Denoëux, 2008), the parametric triangular-norm-based rule (Su et al., 2018), which do not rely on the assumption that the items of evidence are independent of each other;

Table 5: NMI, ACC (mean $\pm$ std.deviation) and running time of different algorithms on benchmark datasets. The $\bullet/\circ$ indicates whether TIC is statistically superior/inferior to a certain comparing baseline based on the paired t -test at a 0.05 significance level. The statistics of win/tie/loss are shown in the last row of each sub-table. The best and the second-best results, between which the performance gaps are shown in the row named “gap” in each sub-table, are colored blue and red.

NMI	EMG	ECG	HAR	Gesture	FD-A	FD-B	SleepEEG	Epilepsy	UCR
DTCR	.802 $\pm$.02 $\bullet$	.745 $\pm$.01 $\bullet$	.653 $\pm$.02 $\bullet$	.766 $\pm$.02 $\bullet$	.862 $\pm$.01 $\bullet$	.822$\pm$.02$\circ$	.492 $\pm$.02 $\bullet$	.802 $\pm$.02 $\bullet$	.663 $\pm$.02 $\bullet$
k -shape	.752 $\pm$.02 $\bullet$	.591 $\pm$.01 $\bullet$	.745 $\pm$.01 $\bullet$	.734 $\pm$.02 $\bullet$	.836 $\pm$.02 $\bullet$	.756 $\pm$.03 $\bullet$	.692 $\pm$.02 $\bullet$	.834 $\pm$.03 $\bullet$	.843 $\pm$.03 $\bullet$
SOM-VAE	.892$\pm$.03$\bullet$	.632 $\pm$.02 $\bullet$	.822$\pm$.02$\bullet$	.888$\pm$.01$\bullet$	.789 $\pm$.03 $\bullet$	.682 $\pm$.02 $\bullet$	.722 $\pm$.02 $\bullet$	.929 $\pm$.01 $\bullet$	.825 $\pm$.03 $\bullet$
STCN	1$\pm$0	.745 $\pm$.03 $\bullet$	.669 $\pm$.02 $\bullet$	.793 $\pm$.01 $\bullet$	.894$\pm$.02$\bullet$	.726 $\pm$.03 $\bullet$	.833 $\pm$.02 $\bullet$	.802 $\pm$.01 $\bullet$	.639 $\pm$.02 $\bullet$
TMEK	.696 $\pm$.02 $\bullet$	.653 $\pm$.02 $\bullet$	.738 $\pm$.01 $\bullet$	.859 $\pm$.01 $\bullet$	.721 $\pm$.02 $\bullet$	.730 $\pm$.01 $\bullet$	.793 $\pm$.03 $\bullet$	.952 $\pm$.02 $\bullet$	.703 $\pm$.01 $\bullet$
TNC	.743 $\pm$.01 $\bullet$	.746$\pm$.02$\bullet$	.654 $\pm$.01 $\bullet$	.833 $\pm$.02 $\bullet$	.751 $\pm$.02 $\bullet$	.652 $\pm$.01 $\bullet$	.799 $\pm$.02 $\bullet$	.854 $\pm$.01 $\bullet$	.738 $\pm$.01 $\bullet$
TS3C _{ch}	.752 $\pm$.01 $\bullet$	.726 $\pm$.01 $\bullet$	.753 $\pm$.02 $\bullet$	.871 $\pm$.02 $\bullet$	.754 $\pm$.04 $\bullet$	.752 $\pm$.03 $\bullet$	.921$\pm$.01$\bullet$	.949 $\pm$.01 $\bullet$	.874 $\pm$.01 $\bullet$
UMAP	.863 $\pm$.01 $\bullet$	.720 $\pm$.01 $\bullet$	.610 $\pm$.02 $\bullet$	.827 $\pm$.02 $\bullet$	.714 $\pm$.04 $\bullet$	.746 $\pm$.03 $\bullet$	.840 $\pm$.01 $\bullet$	.954$\pm$.01$\bullet$	.882$\pm$.01$\bullet$
USSL	.842 $\pm$.01 $\bullet$	.702 $\pm$.01 $\bullet$	.496 $\pm$.02 $\bullet$	.568 $\pm$.02 $\bullet$	.652 $\pm$.02 $\bullet$	.705 $\pm$.03 $\bullet$	.900 $\pm$.02 $\bullet$	.895 $\pm$.01 $\bullet$	.736 $\pm$.02 $\bullet$
VLSC	1$\pm$0	.734 $\pm$.01 $\bullet$	.578 $\pm$.02 $\bullet$	.724 $\pm$.02 $\bullet$	.630 $\pm$.03 $\bullet$	.339 $\pm$.02 $\bullet$	.754 $\pm$.01 $\bullet$	.831 $\pm$.02 $\bullet$	.731 $\pm$.01 $\bullet$
TIC (ours)	1$\pm$0	.786$\pm$.01$\bullet$	.887$\pm$.01$\bullet$	.898$\pm$.02$\bullet$	.900$\pm$.01$\bullet$	.784$\pm$.02$\bullet$	.942$\pm$.01$\bullet$	.965$\pm$.02$\bullet$	.891$\pm$.01$\bullet$
gap	.108	.040	.065	.010	.006	.038	.021	.011	.009
win/tie/loss	8/2/0	10/0/0	10/0/0	8/2/0	9/1/0	8/1/1	9/1/0	7/3/0	8/2/0
ACC	EMG	ECG	HAR	Gesture	FD-A	FD-B	SleepEEG	Epilepsy	UCR
DTCR	.792 $\pm$.02 $\bullet$	.841 $\pm$.01 $\bullet$	.621 $\pm$.02 $\bullet$	.892 $\pm$.03 $\bullet$	.906 $\pm$.02 $\bullet$	.938$\pm$.01$\bullet$	.620 $\pm$.01 $\bullet$	.921 $\pm$.02 $\bullet$	.657 $\pm$.02 $\bullet$
k -shape	.715 $\pm$.02 $\bullet$	.669 $\pm$.01 $\bullet$	.802 $\pm$.03 $\bullet$	.856 $\pm$.02 $\bullet$	.925$\pm$.01$\bullet$	.837 $\pm$.02 $\bullet$	.840 $\pm$.01 $\bullet$	.931 $\pm$.02 $\bullet$	.831 $\pm$.03 $\bullet$
SOM-VAE	.885 $\pm$.03 $\bullet$	.879 $\pm$.02 $\bullet$	.839$\pm$.01$\bullet$	.940$\pm$.01$\bullet$	.862 $\pm$.02 $\bullet$	.781 $\pm$.02 $\bullet$	.874 $\pm$.02 $\bullet$	.952 $\pm$.02 $\bullet$	.841 $\pm$.01 $\bullet$
STCN	1$\pm$0	.785 $\pm$.02 $\bullet$	.805 $\pm$.02 $\bullet$	.839 $\pm$.01 $\bullet$	.919 $\pm$.03 $\bullet$	.846 $\pm$.02 $\bullet$	.826 $\pm$.02 $\bullet$	.810 $\pm$.01 $\bullet$	.605 $\pm$.02 $\bullet$
TMEK	.742 $\pm$.02 $\bullet$	.795 $\pm$.02 $\bullet$	.800 $\pm$.01 $\bullet$	.863 $\pm$.01 $\bullet$	.752 $\pm$.02 $\bullet$	.822 $\pm$.01 $\bullet$	.793 $\pm$.03 $\bullet$	.954 $\pm$.02 $\bullet$	.678 $\pm$.01 $\bullet$
TNC	.838 $\pm$.01 $\bullet$	.921$\pm$.02$\bullet$	.726 $\pm$.01 $\bullet$	.896 $\pm$.02 $\bullet$	.821 $\pm$.03 $\bullet$	.675 $\pm$.01 $\bullet$	.820 $\pm$.03 $\bullet$	.829 $\pm$.01 $\bullet$	.771 $\pm$.01 $\bullet$
TS3C _{ch}	.798 $\pm$.02 $\bullet$	.842 $\pm$.02 $\bullet$	.732 $\pm$.01 $\bullet$	.889 $\pm$.01 $\bullet$	.796 $\pm$.03 $\bullet$	.726 $\pm$.01 $\bullet$	.933 $\pm$.01 $\bullet$	.950 $\pm$.02 $\bullet$	.895$\pm$.02$\bullet$
UMAP	.876 $\pm$.01 $\bullet$	.846 $\pm$.01 $\bullet$	.705 $\pm$.02 $\bullet$	.909 $\pm$.02 $\bullet$	.731 $\pm$.02 $\bullet$	.863 $\pm$.01 $\bullet$	.921 $\pm$.01 $\bullet$	.988$\pm$.03$\bullet$	.867 $\pm$.01 $\bullet$
USSL	.932$\pm$.01$\bullet$	.822 $\pm$.01 $\bullet$	.734 $\pm$.02 $\bullet$	.619 $\pm$.02 $\bullet$	.658 $\pm$.04 $\bullet$	.639 $\pm$.03 $\bullet$	.885 $\pm$.01 $\bullet$	.978 $\pm$.01 $\bullet$	.751 $\pm$.01 $\bullet$
VLSC	1$\pm$0	.846 $\pm$.01 $\bullet$	.687 $\pm$.02 $\bullet$	.879 $\pm$.02 $\bullet$	.786 $\pm$.03 $\bullet$	.602 $\pm$.02 $\bullet$	.941$\pm$.01$\bullet$	.977 $\pm$.02 $\bullet$	.742 $\pm$.01 $\bullet$
TIC (ours)	1$\pm$0	.935$\pm$.01$\bullet$	.869$\pm$.01$\bullet$	.959$\pm$.01$\bullet$	.982$\pm$.01$\bullet$	.950$\pm$.02$\bullet$	.957$\pm$.01$\bullet$	.997$\pm$.01$\bullet$	.914$\pm$.02$\bullet$
gap	.068	.014	.030	.019	.057	.012	.016	.009	.019
win/tie/loss	8/2/0	9/1/0	10/0/0	9/1/0	10/0/0	9/1/0	8/2/0	7/3/0	8/2/0
Time	EMG	ECG	HAR	Gesture	FD-A	FD-B	SleepEEG	Epilepsy	UCR
DTCR	1.65	412.2	113.2	8.5	213.2	324.3	3486.3	143.7	154.3
k -shape	0.32	541.1	132.7	14.3	300.9	452.9	3688.2	168.3	217.8
SOM-VAE	2.51	365.8	101.7	7.2	196.2	331.1	3348.2	121.9	169.8
STCN	0.85	5421.1	2117.3	3.4	1837.6	2913.2	27986.3	2684.9	186.9
TMEK	1.89	432.4	99.4	6.3	145.7	217.3	2987.6	119.3	197.9
TNC	2.07	500.7	145.3	5.2	164.8	213.7	3114.7	192.3	208.9
TS3C _{ch}	1.54	472.1	123.6	10.9	151.8	207.9	3864.3	143.5	178.3
UMAP	1.25	625.1	200.6	7.6	275.9	423.9	3004.7	287.9	190.6
USSL	2.09	587.3	132.4	4.7	167.3	246.4	3654.8	169.3	183.7
VLSC	1.74	584.6	241.1	5.6	272.2	384.3	3845.1	300.8	223.7
TIC (ours)	2.09	576.3	135.3	6.8	206.2	312.9	3884.9	249.9	199.8

- TIC-comc/TIC-GC/TIC-yage/TIC-dubo/TIC-RCR: Dempster’s rule is replaced by other combination rule, i.e., the COMC rule (Ma et al., 2019c), the GC rule (Du & Zhong, 2021), Yager’s rule (Yager, 1987), Dubois-Prade’s rule (Dubois & Prade, 1988), the RCR rule (Florea et al., 2009) which have their own ways to deal with the high-conflicted mass functions.

Table 6: ARI values (mean $\pm$ std.deviation) of different variants of TIC.

ARI	EMG	ECG	HAR	Gesture	FD-A	FD-B	SleepEEG	Epilepsy	UCR	Average
TIC $\times$	.845	.729	.711	.921	.895	.803	.832	.893	.873	
TIC $+$	.839	.698	.687	.909	.910	.821	.850	.869	.867	
TIC-cau	.965	.873	.783	.924	.924	.839	.892	.931	.832	
TIC-bol	.981	.869	.789	.920	.867	.841	.882	.965	.809	
TIC-tri	1	.879	.772	.915	.926	.839	.889	.978	.867	
TIC-comc	.987	.872	.754	.928	.927	.828	.902	.963	.856	
TIC-GC	.979	.871	.768	.916	.913	.819	.856	.952	.873	
TIC-yage	1	.869	.781	.924	.918	.834	.883	.980	.830	
TIC-dubo	.981	.853	.789	.916	.904	.842	.879	.980	.815	
TIC-RCR	1	.874	.792	.904	.919	.836	.882	.958	.871	
TIC (Full model)	1	.879	.796	.931	.939	.853	.896	.980	.883	

Compared to the full model, the average ARI of TIC $\times$ and TIC $+$ decrease by $0.912 - 0.834 = 0.078$ and $0.912 - 0.833 = 0.079$. This suggests that it is more reasonable to use Dempster’s rule to fuse the clustering results from time and frequency domains. The potential reasons for this result are twofold: (1) as shown in Eq.(9), Dempster’s rule includes the $m_{ic}^{Time} \cdot m_{ic}^{Freq}$ (multiplication term) and $m_{ic}^{Time} \cdot m_{ic}^{Freq} + m_{i\Omega}^{Time} \cdot m_{ic}^{Freq}$ (addition term); (2) Dempster’s rule considers additionally the uncertainty $m_{i\Omega}^{Time}$ and $m_{i\Omega}^{Freq}$ (in term $m_{i\Omega}^{Time} \cdot m_{ic}^{Freq} + m_{ic}^{Time} \cdot m_{i\Omega}^{Freq}$), when calculating m_{ic}^{Comb} . To further illustrate the contribution of combining the results from time and frequency domains, we show the cosine distance between $\mathbf{g}_i^{Time}$ and $\mathbf{g}_i^{Freq}$ of the same sample in Fig.4(a). “W/o Comb” means that the cosine distance is calculated between the frequency embeddings learned from “W/o

$\mathcal{L}^{\text{Time}}$ ” and the time embeddings learned from “W/o $\mathcal{L}^{\text{Freq}}$ ”. One can find that the cosine distance is smaller by combining the results from time and frequency domains. It means that combining the results from these two domains indeed enforces the time- and frequency-based embeddings closer to each other. Taking the FD-A dataset as an example, the average cosine distance decreases from 1.924 to 1.557, i.e., 19.1% by considering the combination.

Compared with TIC-cau and TIC-bol, TIC performs better on all the benchmark datasets. The reasons are: the cautious conjunctive rule uses the intersection operation and loses useful information; the bold conjunctive rule uses the union operation and takes into account confusing information when making a decision. Comparing TIC with TIC-tri, TIC-tri only has the same ARI as TIC in 2 cases but has lower ARI on the other 7 cases. It is because that the parametric triangular-norm-based rule is more sensitive to the hyper-parameters, and reasonable hyper-parameter values are difficult to choose in both of the time and frequency domains. This also suggests that the mass functions associated with the clustering results from the time and frequency domains are independent of each other in the time series clustering problem studied in our paper. Besides, only TIC-come in the 5 variants aiming to tackle high-conflict have higher ARI than TIC on SleepEEG dataset, because the conflict values between the results of time and frequency domains are small.

Other 5 variants

- TIC-SVM, TIC-LR, TIC-FC and TIC-LSTM: the outputs from time- and frequency- domains are treated as features to train SVM, Logistic Regression, fully connected layers, and LSTM;
- TIC-max: the maximum probability given in networks from time- and frequency- domains is chosen as the final output probability.

are considered to show that using evidence theory to combine the results from time and frequency domains are superior to other fusion methods. As shown in Table 7, Using evidence theory has the best ARI in 7 of 9 cases, showing that fusing the results via evidence theory is better than other methods.

Table 7: ARI values (mean $\pm$ std.deviation) of different variants of TIC.

ARI	EMG	ECG	HAR	Gesture	FD-A	FD-B	SleepEEG	Epilepsy	UCR	Average
TIC-SVM	.987	.875	.779	.930	.928	.842	.882	.962	.869	
TIC-LR	1	.839	.782	.929	.928	.844	.885	.974	.870	
TIC-FC	1	.860	.769	.934	.932	.859	.876	.980	.863	
TIC-LSTM	.993	.853	.739	.945	.917	.851	.879	.974	.876	
TIC-max	.976	.842	.754	.928	.921	.809	.882	.956	.879	
TIC (Full model)	1	.879	.796	.931	.939	.853	.896	.980	.883	

A.8 EMBEDDING EVALUATION: TIME-SERIES ANOMALY DETECTION.

We evaluate how TIC performs on a sample-level anomaly detection task, which aims to detect abnormal time-series samples. We build two subsets of FD-B and ECG datasets. The former contains 1000 samples where 900 undamaged bearings are considered “normal” and 100 damaged samples are “outliers”; while the latter has 2000 samples where 1800 normal sinus rhythm recordings are considered “normal” and 200 atrial fibrillation recordings are “outliers”. We randomly select half of the “normal” samples as the training data, of which the embeddings learned by different models are used to train the one-class SVM. The time- and frequency-based embeddings learned by TIC are also concatenated [$\mathbf{g}^{\text{Time}}$; $\mathbf{g}^{\text{Freq}}$]. In Table 8, we report the performance on the anomaly detection task in terms of Precision, Recall, F1-score and AUROC. TIC achieves the best performance in 6 out of 8 cases. On the ECG dataset, TIC outperforms the second-best model (i.e., activity2vec) by a margin of 3.3% in AUROC. It shows that TIC can effectively detect the abnormal samples in mechanical devices and ECGs.

A.9 EFFECT OF AUGMENTATION

As shown in Section 3, data augmentations are adopted in both the time- and frequency-based contrastive modules. We explore the effects of augmentation techniques in these two modules separately. In each trial, the corresponding setting is changed while the other ones are the same as the default settings.

Table 8: Performance on time-series anomaly detection. The subsets of FD-B (1000 samples) and ECG (2000 samples) are used. These two subsets are highly imbalanced (90% normal samples and 10% abnormal samples).

FD-B	Precision	Recall	F1-score	AUROC
activity2vec	.8054 $\pm$.034	.7145 $\pm$.037	.7432 $\pm$.023	.7911 $\pm$.029
BTSF	.6854 $\pm$.015	.6274 $\pm$.017	.6653 $\pm$.021	.6741 $\pm$.014
CLOCS	.8412 $\pm$.025	.7465 $\pm$.028	.8222 $\pm$.019	.8126 $\pm$.021
MHCCL	.6210 $\pm$.043	.5745 $\pm$.036	.6009 $\pm$.040	.6123 $\pm$.023
TFC	.8547$\pm$.037	.7698$\pm$.029	.8274$\pm$.019	.8602$\pm$.033
Triplet	.7321 $\pm$.034	.6547 $\pm$.045	.6987 $\pm$.043	.7201 $\pm$.031
TS2Vec	.6813 $\pm$.022	.6134 $\pm$.027	.6423 $\pm$.019	.6731 $\pm$.015
TSTCC	.6354 $\pm$.019	.4517 $\pm$.066	.4968 $\pm$.056	.8022 $\pm$.044
TIC (ours)	.8641$\pm$.025	.7652$\pm$.013	.8359$\pm$.024	.8563$\pm$.033
ECG	Precision	Recall	F1-score	AUROC
activity2vec	.7584 $\pm$.011	.6124 $\pm$.036	.7022 $\pm$.073	.7598$\pm$.026
BTSF	.7752$\pm$.027	.7124$\pm$.034	.7413 $\pm$.022	.7581 $\pm$.029
CLOCS	.4861 $\pm$.049	.4035 $\pm$.056	.4491 $\pm$.038	.4727 $\pm$.023
MHCCL	.5689 $\pm$.015	.4875 $\pm$.019	.5471 $\pm$.024	.5563 $\pm$.023
TFC	.7654 $\pm$.037	.7035 $\pm$.029	.7503$\pm$.019	.7429 $\pm$.033
Triplet	.5789 $\pm$.027	.5067 $\pm$.023	.5417 $\pm$.028	.5561 $\pm$.019
TS2Vec	.6857 $\pm$.015	.6138 $\pm$.019	.6587 $\pm$.021	.6701 $\pm$.020
TSTCC	.7345 $\pm$.031	.6538 $\pm$.024	.6993 $\pm$.016	.7235 $\pm$.019
TIC (ours)	.7958$\pm$.021	.7216$\pm$.024	.7645$\pm$.026	.7856$\pm$.019

In the time-based contrastive module, we fix the augmentation as jittering, scaling and time-shift for all samples, instead of randomly choosing one augmentation from the augmentation bank $\mathcal{B}^{\text{Time}}$ (consisting of these three augmentations) for each sample. The comparison result is shown in the rightmost sub-table of Table 9. As can be seen, random selection covers diverse augmentations that allow the encoder $H_T(\cdot)$ to learn better and more robust embedding $\mathbf{g}^{\text{Time}}$, which improves the clustering performance.

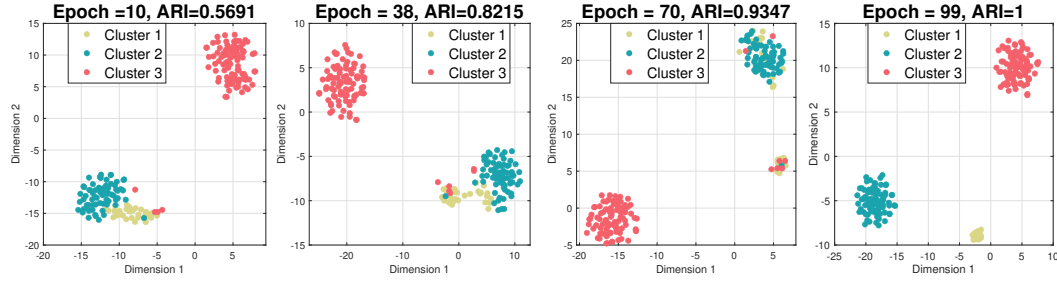


Figure 7: T-SNE visualization of the concatenated embedding $[\mathbf{g}_i^{\text{Time}}; \mathbf{g}_i^{\text{Freq}}]$ for EMG dataset in different epochs. As the learning proceeds, the embeddings from the same cluster are gradually grouped together.

We consider three types of change in frequency-based augmentations. (1) We explore the number of components manipulated, where the results are shown by the sub-table titled by *Components* in Table 9. One can find that the performance of TIC model tends to deteriorate when perturbing multiple frequency components. This degradation is attributed to the substantial changes in the time-domain of augmented samples. Consequently, these augmented samples become readily distinguishable by a contrastive module, resulting in suboptimal contrastive encoders. (2) We then explore the adjustment size of amplitude in sub-table *Amplitude* of Table 9. As can be seen, manipulating the amplitude does not significantly affect the performance of TIC model. (3) In sub-table *Bands* of Table 9, we explore the band that the manipulated component belongs to. The Low-/high-frequency band corresponds to the first/second half of the frequency spectrum, and contributes to slow/fast variations in the time domain. In a physiological time-series dataset (e.g., SleepEEG), the low band contains most information and manipulating a low-frequency component leads to higher ARI. On the contrary, high-band components are more informative than low-band ones in a mechanical time-series dataset (e.g., FD-A). Thus, perturbing high-band components outperform low-band augmentations.

A.10 VISUALIZATION OF CLUSTERING RESULTS

Figure 7 shows the t-SNE Van der Maaten & Hinton (2008) plot of the embeddings for EMG dataset learned by TIC model. For each sample, we concatenate the embeddings $[\mathbf{g}_i^{\text{Time}}; \mathbf{g}_i^{\text{Freq}}]$. As the

Table 9: ARI with different augmentation techniques. Terms *Components*, *Amplitudes* and *Bands* refer to the frequency-based augmentations. *Components* means how many components are manipulated, *Amplitude* means the adjustment size of amplitude and *Bands* means the perturbation is performed on a low- or a high-frequency component. *Time domain* means the augmentation $\tilde{\mathbf{x}}_i^{\text{Time}}$ are randomly selected from the bank $\mathcal{B}_i^{\text{Time}}$ {Jittering, Scaling Shift} or a fixed augmentation technique.

ARI	<i>Components</i>			<i>Amplitude</i>				<i>Bands</i>		<i>Time domain</i>			
	$\beta = 1$	$\beta = 3$	$\beta = 5$	$\gamma = 0.1$	$\gamma = 0.5$	$\gamma = 0.9$	$\gamma = 1.1$	Low	High	Random	Jittering	Scaling	Shift
SleepEEG	0.8961	0.8751	0.8245	0.8952	0.8961	0.8934	0.8942	0.8994	0.8802	0.8961	0.8726	0.8542	0.8793
FD-A	0.9392	0.9157	0.9013	0.9406	0.9392	0.9410	0.9375	0.9321	0.9457	0.9392	0.9103	0.9261	0.9245

learning proceeds, the embeddings learned by TIC from the same cluster are gradually grouped together. In particular, samples clearly form 3 clusters at epoch = 99, corresponding to the ARI=1 of TIC shown in Table 5. The visualization results further demonstrate the well embedding ability of our model.