

基于概率模型与信息熵的局部线性嵌入算法

刘远红 毋毓斌

东北石油大学电气信息工程学院 黑龙江 大庆 163318

(liuyuanhong@nepu.edu.cn)

摘要 局部线性嵌入算法采用欧氏距离选择邻域点,这通常会损失数据集本身的非线性特征,造成邻域点选取错误,且仅使用欧氏距离构造权重会导致信息挖掘不充分。针对以上问题,提出基于概率模型与信息熵的局部线性嵌入算法(Probability information entropy-LLE, PIE-LLE)。首先,为了使邻域点选择更加合理,从数据集的概率分布角度出发,考虑样本点及其邻域的概率分布,为样本点构造符合局部分布的邻域集合。其次,为了充分提取样本的局部结构信息,在权重构造阶段,分别计算样本所属邻域概率以及每个样本的信息熵,融合二者信息重构低维样本。最后,在两个轴承故障数据集上的实验表明,所提方法故障识别准确度最高达到了 100%,高于其他对比算法;在邻域点个数 5~15 范围内,PIE-LLE 算法展现出良好的低维可视化效果;在参数敏感性实验中,该算法可以保持 Fisher 指标较大,有效提高了算法的分类准确度和稳定性。

关键词: 局部线性嵌入算法;概率模型;信息熵;特征提取;故障诊断

中图分类号 TP391

Local Linear Embedding Algorithm Based on Probability Model and Information Entropy

LIU Yuanhong and WU Yubin

School of Electrical Engineering & Information, Northeast Petroleum University, Daqing, Heilongjiang 163318, China

Abstract The local linear embedding algorithm uses Euclidean distance to select neighborhood points, which usually loses the nonlinear features of the dataset itself, resulting in incorrect selection of neighborhood points, and only using Euclidean distance to construct weights leads to insufficient information mining. To address the above issues, a local linear embedding algorithm based on probability model and information entropy (PIE-LLE algorithm) is proposed. Firstly, in order to make the selection of neighborhood points more reasonable, from the perspective of the probability distribution of the dataset, the probability distribution of the sample points and their neighborhoods is considered, and a neighborhood set that conforms to the local distribution is constructed for the sample points. Secondly, in order to fully extract the local structural information of the samples, in the weight construction stage, the probability of the neighborhood to which the samples belong and the information entropy of each sample are calculated separately, and the two information are fused to reconstruct the low dimensional samples. Finally, experiments on two bearing fault datasets show that the highest accuracy of fault identification reaches 100%, higher than that of other comparative algorithms. Within the range of 5~15 neighborhood points, the PIE-LLE algorithm exhibits good low dimensional visualization performance. In the parameter sensitivity experiment, the proposed algorithm can maintain a relatively large Fisher index, effectively improving the classification accuracy and stability of the algorithm.

Keywords Local linear embedding algorithm, Probability model, Information entropy, Feature extraction, Fault diagnosis

1 引言

随着科技水平的飞速发展,各种工业设备的故障问题越来越多地使用智能算法进行诊断。滚动轴承是广泛应用于各种工业场合的机械部件,及时发现和诊断轴承故障对工业生产起着关键作用。目前,轴承故障诊断主要运用振动分析、声学分析、温度监测、润滑状况监测等多种技术手段。在这些方法中,振动分析是应用最为广泛的一种,通过采集轴承振动信号并对其进行分析,可以确定故障的类型、严重程度以及时效性等。但振动数据存在维度高、

蕴含信息杂乱等特点,因此,如何快速提取高维数据的重要特征成为故障诊断算法的重点之一。

近年来,流形学习方法被大量应用于故障诊断领域^[1]。流形学习的基本思想是假设高维数据样本分布在一个低维流形上,从而对高维数据进行数据降维和特征提取。许多流形学习方法已经被提出并应用于特征提取,例如,主成分分析(PCA)^[2]通过寻找最大方差的方向将高维数据投影到较低维的子空间中,以实现最优的数据重构;邻域保留嵌入(NPE)^[3]和局部保留投影(LPP)^[4]利用数据点与其邻居之间的局部关系来执行特征提取;非负矩阵分解(NMF)^[5]已被提出用于具

基金项目:海南省自然科学基金(623MS071)

This work was supported by the Natural Science Foundation of Hainan Province China(623MS071).

通信作者:刘远红(liuyuanhong@nepu.edu.cn)

有非负约束的多元数据分析;线性判别分析(LDA)^[6]是一种有监督的特征提取方法,该方法将每个描述符投影到局部坐标系中,并使用这些局部坐标作为后续学习任务的新特征^[7];稀疏表示(SR)和低秩表示(LRR)^[8]是最著名的两种基于表示的特征提取方法;标签分布学习(LDL)^[9]用于解决标签对样本贡献程度高低的问题。这些模型可以很好地处理海量数据,无需人工判断,也可以取得比传统方法更好、更有效的诊断结果^[10]。然而,当前的许多智能诊断算法由于需要调整的超参数繁多,难以实现,如果超参数设置不当将对诊断精度产生很大影响,这也是故障诊断智能算法领域亟待解决的问题。

局部线性嵌入(Locally Linear Embedding, LLE)算法^[11]是一种经典的流形学习算法。在轴承故障诊断领域, LLE 算法用于将轴承振动信号等高维数据映射到低维流形上^[12-14],便于故障特征的提取和分类。但 LLE 算法也有明显的不足,其近邻点参数的选择对降维效果具有一定影响,需要进行合理的调参,且对于非线性流形的数据表现可能不理想。近年来,有关概率模型拟合进行故障诊断的方法被提出, Cai 等^[15]提出基于贝叶斯网络的故障诊断方法, Soize 等^[16]提出基于概率模型的流形学习方法, Verma 等^[17]提出基于概率模型的故障诊断方法。概率模型^[18-19]与信息熵等^[20-21]非线性方法^[22-24]已经被广泛应用于机器学习领域。

LLE 算法在选择近邻点时仅采用欧氏距离进行排序,仅保持局部线性关系,并不符合数据的真实分布,未考虑数据的非线性结构与信息熵特征,通常会造成信息丢失。本文提出的方法则聚焦于利用样本自身分布特性进行无监督学习,选择概率分布与信息熵加权融合的形式来进行近邻点权重的重新分配。提出的 PIE-LLE 算法更贴近数据集本身的概率分布,且可以构建无量纲模型,使数据近邻点的选择更加合理,提高算法的特征提取能力,从而为后续故障诊断提供包含数据本质特征的低维数据集。

本文第 2 章为相关工作描述;第 3 章为 PIE-LLE 算法描述;第 4 章为实验描述;最后总结全文。

2 相关工作

2.1 局部线性嵌入算法

局部线性嵌入(Locally Linear Embedding, LLE)用于将高维数据映射到低维空间。它基于一个基本假设:在高维空间中,局部邻域内的数据点可以通过线性组合来近似表示。该算法的目标是在保持局部数据结构的同时,将数据映射到低维空间中。算法步骤如下:

1)局部邻域选择:对于每个数据点,选择其 k 个最近邻的数据点作为其局部邻域。

2)重建权重计算:对于每个样本点,通过最小化其与邻域内样本点之间的欧氏距离的重建误差来计算重建权重。这一步的目标是找到最佳的线性组合来近似表示每个样本点。假设一个样本点 x_i , 它的 k 个最近邻样本点为 $x_{i_1}, x_{i_2}, \dots, x_{i_k}$, 我们需要找到权重 w_{ij} , 使得 x_i 可以由其邻域内的样本点线性组合而成,即:

$$\min_{\mathbf{W}_i} \sum_{j=1}^k \|x_i - \sum_{l=1}^k w_{ij} x_{i_l}\|^2 \quad (1)$$

此处 $\mathbf{W}_i$ 为权重矩阵,其中的每一行 $\mathbf{W}_i = [w_{i1}, w_{i2}, \dots,$

$w_{ik}]$ 是对应于 x_i 的权重向量,且满足 $\sum_{j=1}^k w_{ij} = 1$ 。

3)低维嵌入:利用得到的重建权重,构建一个低维表示。对于每对数据点 x_i 和 x_j , 在低维空间中找到最优的表示。这一步的目标是在保持样本点之间的局部结构关系的同时,将数据映射到低维空间中。

4)重建误差最小化:重建误差最小化问题可以表示为以下形式:

$$\min_{\mathbf{W}_i} \sum_{j=1}^k \|x_i - \sum_{l=1}^k w_{ij} x_{i_l}\|^2 = \min_{\mathbf{W}_i} \|x_i - \mathbf{X}_i \mathbf{W}_i\|^2 \quad (2)$$

其中, $\mathbf{X}_i$ 为邻居矩阵,包含 $x_{i_1}, x_{i_2}, \dots, x_{i_k}$, 而 $\mathbf{W}_i$ 是对应于 x_i 的权重向量。要最小化上式,需要求解下面这个问题的最优权重 $\mathbf{W}_i$:

$$\min_{\mathbf{W}_i} \|x_i - \mathbf{X}_i \mathbf{W}_i\|^2 \quad (3)$$

在得到了每个数据点的最优权重 $\mathbf{W}_i$ 后,可以构建一个低维表示。这个过程通常涉及到求解一个特征分解问题或者通过局部加权线性回归进行映射。最终的嵌入结果既在低维空间中最小化了数据点之间的距离,又保持了局部结构。

2.2 概率分布拟合

已知大小为 $D * N$ 的原始数据集 X , 其中 N 为样本点个数, D 为每个样本点的维度, 则该数据集可表示为 $X = \{x_1, x_2, \dots, x_N\}$ 。任意一个样本点属于某一样本点邻域的条件概率 $p_{j|i}$ 如下所示:

$$p_{j|i} = \frac{\exp\left(-\frac{\|x_i - x_j\|^2}{2\sigma_i^2}\right)}{\sum_{k \neq i} \exp\left(-\frac{\|x_i - x_k\|^2}{2\sigma_i^2}\right)} \quad (4)$$

其中, σ_i 为以样本点 x_i 为中心点的高斯协方差。当直接使用联合概率分布替换条件概率分布时,异常值的出现将会极大程度影响概率的计算,进而影响算法效果,因此将条件概率按下式进行改进:

$$p_{ij} = \frac{p_{j|i} + p_{i|j}}{2n} \quad (5)$$

2.3 信息熵

信息熵是随机数据源产生信息的均量。信息熵代表的是随机变量或整个系统的不确定性,熵越大,随机变量或系统的不确定性就越大。

信息熵的数学公式如下:

$$H(X) = E[I(X)] = E[-\ln(P(X))] \quad (6)$$

其中, P 为 X 的概率质量函数, $I(X)$ 是 X 的信息量, $I(X)$ 本身是随机变量个数,当取自有限样本时,熵的公式可以表示为:

$$H(X) = -\sum_{i=1}^n P(x_i) I(x_i) = -\sum_{i=1}^n P(x_i) \log_b P(x_i) \quad (7)$$

熵的本质是香农信息量 $\log(1/p)$ 的期望,一般令 $b=2$ 。

从 $H(x) = \log_2(x)$, 其中 $x = p(x_i)$ 的图像可以看出,一个事件结果的出现概率 $p(x_i)$ 越高, $H(x)$ 就越大,其代表的信息量越多。

概率分布和信息熵是信息论中两个重要的概念,它们之间存在着密切的关系。概率分布描述了随机变量在不同取值上的概率情况。对于离散型随机变量,概率分布可以通过概率质量函数(Probability Mass Function, PMF)表示;对于连续型随机变量,概率分布可以通过概率密度函数(Probability

Density Function, PDF)表示。

信息熵是概率分布的一种度量方法,即随机变量的平均信息量。当概率分布更加平均,即各个事件发生的概率相近时,信息熵达到最大值;而当概率分布更加集中于特定事件时,信息熵会减小。信息熵常用于衡量信息传输、数据压缩、模式识别等领域中的信息量大小,它可以度量数据的不确定性程度,从而指导决策和分析。这一原理在统计建模和机器学习有着广泛的应用。因此,概率分布和信息熵是信息论中两个重要的概念,它们共同分析数据的特征和性质。

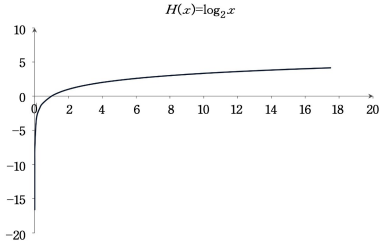


图1 $H(x) = \log_2(x)$ 函数图像

Fig. 1 $H(x) = \log_2(x)$ function image

3 本文算法

算法处理故障诊断的过程存在着明显不确定性,本章提出 PIE-LLE 算法,利用概率信息来评估样本间的非线性关系,利用信息熵来判断单个样本在邻域内包含的信息量关系,从而得到更准确的权重参数,以此为基础改进 LLE 算法,提高诊断过程的可靠性和可解释性。

PIE-LLE 算法近邻点选择示意图如图 2 所示,根据高维数据样本概率分布 p_{ik} 的大小进行邻域选择,根据 p_{ik} 和信息熵 h_{ik} 计算重构权重 w_{ik} ,最终完成低维数据重构。

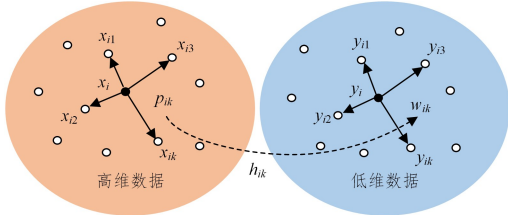


图2 PIE-LLE 算法近邻点选择示意图

Fig. 2 Schematic diagram of PIE-LLE algorithm for selecting nearest neighbor points

3.1 根据概率分布与信息熵寻找近邻点

3.1.1 计算概率分布

原始数据集维度过高,包含大量冗余信息,且采集信号为轴承振动信号,因此需要从时域、频域、时频域进行预处理,滤除噪声等无用信息,再将预处理后的数据作为算法的输入,生成可以作为算法输入的数据集矩阵 $\mathbf{X}$ 。

已知大小为 $N \times D$ 的输入数据集 $\mathbf{X}$,其中 N 为样本点个数, D 为每个样本点的维度,则该数据集可表示为: $\mathbf{X} = [\mathbf{X}_1, \mathbf{X}_2, \dots, \mathbf{X}_N]^T$, 其中, $\mathbf{X}_i = [X_{i1}, X_{i2}, \dots, X_{iD}]$ 。任意一个样本点属于某一样本点邻域的条件概率 $p_{ji|i}$ 计算式如下:

$$p_{ji|i} = \frac{\exp\left(-\frac{\|\mathbf{X}_i - \mathbf{X}_j\|^2}{2\sigma_i^2}\right)}{\sum_{k \neq i} \exp\left(-\frac{\|\mathbf{X}_i - \mathbf{X}_k\|^2}{2\sigma_i^2}\right)} \quad (8)$$

其中, σ_i 为以样本点 $\mathbf{X}_i$ 为中心点的高斯协方差。将条件概率按式(5)进行改进,得到对称的概率矩阵 $\mathbf{P}$:

$$\mathbf{P} = \begin{bmatrix} p_{11} & \cdots & p_{1N} \\ \vdots & \ddots & \vdots \\ p_{N1} & \cdots & p_{NN} \end{bmatrix} \quad (9)$$

其中,令 $p_{ii} = 0$,且该矩阵关于主对角线对称。

3.1.2 计算信息熵

对于以上步骤得出的概率分布,可以根据式(10)计算信息熵:

$$h_{ij} = -p_{ij} \log_2 p_{ij} \quad (10)$$

得到对称的信息熵矩阵 $\mathbf{H}$:

$$\mathbf{H} = \begin{bmatrix} h_{11} & \cdots & h_{1N} \\ \vdots & \ddots & \vdots \\ h_{N1} & \cdots & h_{NN} \end{bmatrix} \quad (11)$$

根据第 i 行或列概率值的大小,依次从大到小找出前 k 个概率值对应的样本点,该 k 个样本点即为第 i 个样本点在概率距离下的近邻点。单个样本点 $\mathbf{X}_i$ 的邻域矩阵表示为:

$$[x_{i1}, x_{i2}, \dots, x_{ik}] \quad (12)$$

不同类别的故障样本,其样本量纲也有所不同,而概率分布通常与量纲没有直接关系,根据输入数据集可以拟合出数据的概率分布矩阵,该步骤的目的在于对数据集进行概率模型拟合,为后续计算信息熵提供输入。概率矩阵同时用于样本邻域集合的判断,选取概率最大的前 k 个样本构成样本的邻域矩阵。为准确度量每个邻域点在邻域内的贡献度,需要综合考虑概率分布与其提供的信息量的大小,因此需要进行融合权重参数设置,用于研究参数变化对降维结果的影响。该步骤通过概率方法改进了原有算法所使用的局部欧氏距离线性度量方法,更贴近数据本身的分布情况,能够在混合了大量特征信息的高维数据中最大程度保留数据的原始信息。

3.2 计算样本点权重

根据概率距离与信息熵大小对样本点的近邻点进行权重分配,通过求取原数据与加权重构后的数据之间的差值作为损失函数,使得在低维嵌入时邻域点排序与降维前的邻域点排序保持一致。

根据前 k 个点属于第 i 个样本点的邻域概率进行以下计算:

$$w_{ik}^1 = \frac{p_{ik}}{\sum_{i=1}^k (p_{i1} + p_{i2} + \dots + p_{ik})} \quad (13)$$

根据概率矩阵可以计算出权重矩阵 $\mathbf{W}^1$:

$$\mathbf{W}^1 = \begin{bmatrix} w_{11}^1 & \cdots & w_{1k}^1 \\ \vdots & \ddots & \vdots \\ w_{N1}^1 & \cdots & w_{Nk}^1 \end{bmatrix}, \text{ s. t. } \sum_{j=1}^k w_{ij}^1 = 1 \quad (14)$$

根据信息熵矩阵可以计算出权重矩阵 $\mathbf{W}^2$:

$$w_{ik}^2 = \frac{h_{ik}}{\sum_{i=1}^k (h_{i1} + h_{i2} + \dots + h_{ik})} \quad (15)$$

$$\mathbf{W}^2 = \begin{bmatrix} w_{11}^2 & \cdots & w_{1k}^2 \\ \vdots & \ddots & \vdots \\ w_{N1}^2 & \cdots & w_{Nk}^2 \end{bmatrix}, \text{ s. t. } \sum_{j=1}^k w_{ij}^2 = 1 \quad (16)$$

由于信息熵大小排序与概率排序一致,但信息熵包含的

数据特征不同,因此,需要加入平衡参数 α 。

混合权重表示为:

$$\mathbf{W} = \alpha \mathbf{W}^1 + (1 - \alpha) \mathbf{W}^2, \text{ s. t. } \sum_{j=1}^k w_{ij} = 1 \quad (17)$$

3.3 进行低维嵌入

进行低维嵌入时,需要最小化重构误差。本文通过矩阵分解与矩阵求导的形式最小化重构误差,则该问题可以表示为:

$$\min_{\mathbf{W}_i} \sum_{j=1}^k \|\mathbf{X}_i - \sum_{i=1}^k w_{ij} x_{i_j}\|^2 = \min_{\mathbf{W}_i} \|\mathbf{X}_i - \mathbf{X}_i \mathbf{W}_i\|^2 \quad (18)$$

该步骤目标是找到一个嵌入矩阵 $\mathbf{Y}$,使其能在保持局部概率分布的同时,将高维数据映射到低维空间。该过程的损失函数如下所示:

$$J(\mathbf{Y}) = \sum_{i=1}^D \|\mathbf{X}_i - \mathbf{Y} \mathbf{W}_i\|_2^2 \quad (19)$$

其中, $\mathbf{W}_i$ 为 x_i 的邻域权重向量, $\mathbf{Y}$ 为低维嵌入矩阵,通过对 $J(\mathbf{Y})$ 求导,令导数为0,即可求解 $\mathbf{Y}$ 。式(20)为对 $J(\mathbf{Y})$ 求导:

$$\begin{aligned} \frac{\partial J(\mathbf{Y})}{\partial \mathbf{Y}} &= \frac{\partial}{\partial \mathbf{Y}} \sum_{i=1}^D \|\mathbf{X}_i - \mathbf{Y} \mathbf{W}_i\|_2^2 \\ &= \sum_{i=1}^D \frac{\partial}{\partial \mathbf{Y}} \|\mathbf{X}_i - \mathbf{Y} \mathbf{W}_i\|_2^2 \\ &= \sum_{i=1}^D -2(\mathbf{X}_i - \mathbf{Y} \mathbf{W}_i) \mathbf{W}_i^T \end{aligned} \quad (20)$$

令上式等于零,可得:

$$\sum_{i=1}^D \mathbf{Y} \mathbf{W}_i \mathbf{W}_i^T = \sum_{i=1}^D \mathbf{X}_i \mathbf{W}_i^T \quad (21)$$

因此:

$$\mathbf{Y} = (\sum_{i=1}^D \mathbf{X}_i \mathbf{W}_i^T) (\sum_{i=1}^D \mathbf{W}_i \mathbf{W}_i^T)^{-1} \quad (22)$$

式(22)可写为下式:

$$\mathbf{Y} = (\mathbf{I} - \mathbf{W})^T (\mathbf{I} - \mathbf{W}) \quad (23)$$

其中, $\mathbf{I}$ 为单位矩阵,通过上述推导,可以得到计算嵌入坐标的方法。在实际应用中,可以通过计算权重矩阵 $\mathbf{W}$,然后利用上述公式来计算嵌入矩阵 $\mathbf{Y}$,从而将高维数据映射到低维空间。PIE-LLE算法示意图如图3所示。

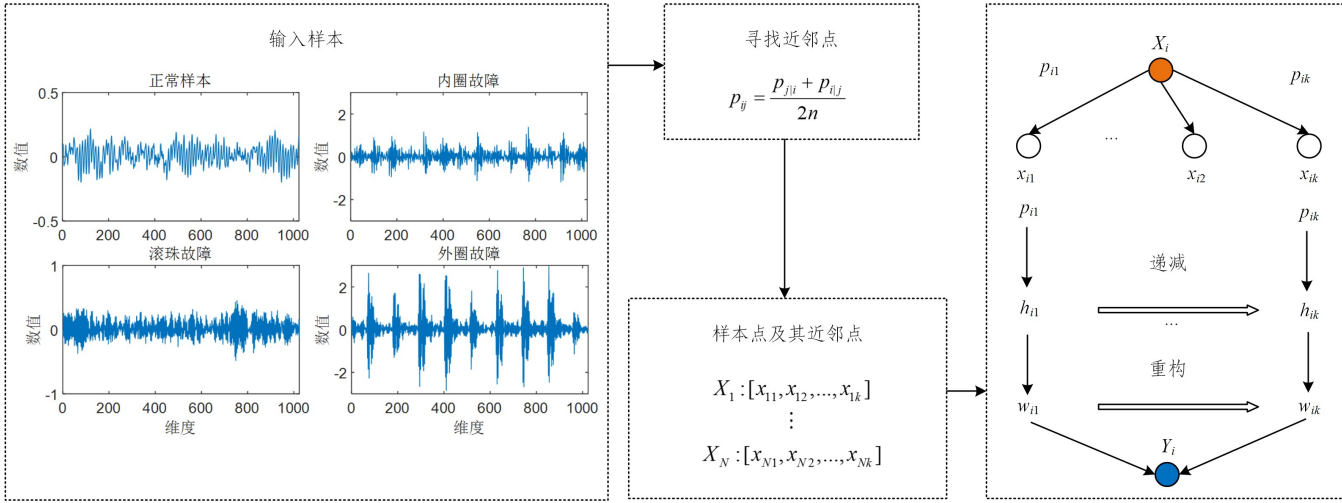


图3 PIE-LLE算法示意图

Fig. 3 Schematic diagram of PIE-LLE algorithm

计算数据集的概率分布和计算数据集的欧氏距离是两种不同的操作,二者各有优势和适用场景。欧氏距离是一种常用的距离度量方法,但采用欧氏距离度量无法衡量真实数据分布中的距离,导致原始LLE算法在降维过程中忽略了数据的主要特征。当使用概率模型作为度量标准时,体现的是每个样本点两两之间互为邻域的概率大小,能够提取数据集的非线性特征,具有更高的可解释性。

轴承数据集包含着大量时频域信息,例如振动信号的幅值和相位等,因此数据中参杂多种不同类型的特征信息,必须考虑数据的量纲问题。以欧氏距离度量很难在多量纲数据中保证特征提取的准确度和有效性。在原始框架下使用无量纲的概率模型,可以在保留原始特征的同时,进行无量纲数据的特征提取。

从概率模型中提取信息熵特征,反映了每个样本包含的信息量大小,包含信息量更多的样本将获得更高的邻域点权重,以此为参考依据可以使邻域点选择更加准确,同时考虑样本点间概率分布与单个样本信息量分布,增强邻域选择的可解释性和准确度。

PIE-LLE算法流程图如图4所示。

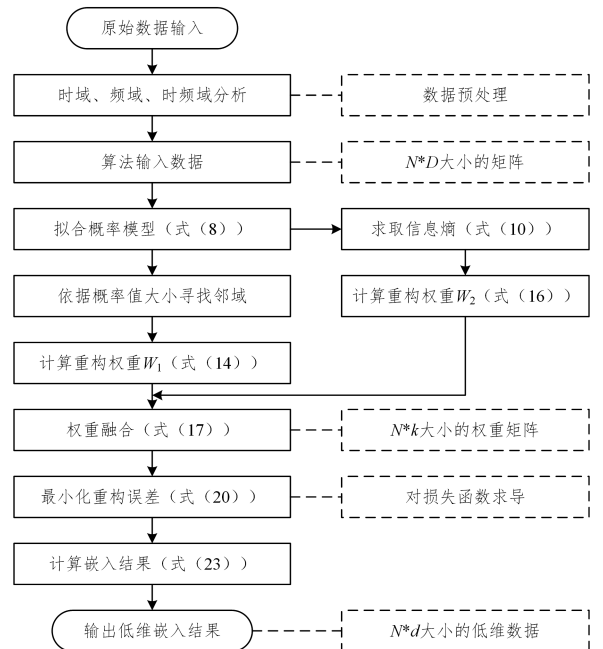


图4 PIE-LLE算法流程图

Fig. 4 PIE-LLE algorithm flowchart

4 实验

4.1 数据集介绍与预处理

本文实验采用两个数据集,分别为采集自凯斯西储大学的轴承故障数据集(CWRU 数据集)以及采集自东北石油大学轴承模拟平台的轴承故障数据集(NEPU 数据集)。

1)CWRU 数据集

测试台包括一个 2 马力的电机(左侧),一个扭矩传感器/编码器(中间),一个测功机(右侧)和控制电子设备(未显示)。测试轴承支撑电机轴。使用电火花加工将单点故障引入测试轴承,故障直径分别为 7 mils,14 mils,21 mils,28 mils 和 40 mils (1mil=0.001 英寸)。7 mils,14 mils 和 21 mils 直径的故障使用 SKF 轴承,28 mils 和 40 mils 的故障使用 NTN 等效轴承。

振动数据使用附有磁性底座的加速度计收集,这些加速度计安装在外壳上,加速度计分别放置在电机外壳的驱动端和风扇端的 12 点钟位置,加速度计安装在电机支撑基板上。振动信号使用 16 通道数据记录仪收集。外环道故障位置分别设在 3 点钟(直接安装于载荷区)、6 点钟(垂直于载荷区)和 12 点钟方向。CWRU 数据集是在 0Hp 电机负载、电机采样频率为 12kHz 和转速 1720r/min 的情况下采集而成,样本总数为 400,包含了 100 个正常轴承数据,100 个轴承滚珠体故障数据,100 个轴承内圈故障数据和 100 个轴承外圈故障数据。每个轴承信号都选取了 1024 个采样点作为一个完整轴承样本。CWRU 轴承数据集采集平台如图 5 所示。

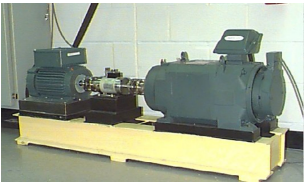


图 5 CWRU 轴承数据集采集平台

Fig. 5 CWRU bearing dataset collection platform

2)NEPU 数据集

NEPU 数据集采集自东北石油大学轴承模拟平台,该平台主要由驱动电机、轴承、齿轮、轴和偏重转盘等组成,其具体参数如表 1 所列。NEPU 数据集是在 1Hp 电机负载、电机采样频率为 10kHz 和转速 1400r/min 的情况下采集而成,样本总数为 400,包含了 100 个正常轴承数据,100 个轴承滚珠体故障数据,100 个轴承内圈故障数据和 100 个轴承外圈故障数据。每个轴承信号都选取了 1024 个采样点作为一个完整轴承样本。NEPU 轴承数据集采集平台如图 6 所示。

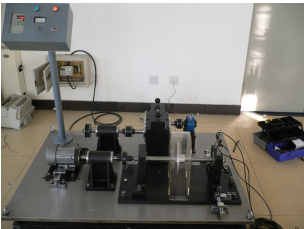


图 6 NEPU 轴承数据集采集平台

Fig. 6 NEPU bearing dataset collection platform

两个数据集具体信息如表 1 所列。

表 1 数据集具体参数

Table 1 Dataset specific parameters

	CWRU	NEPU
采样频率/kHz	12	10
电机转速/(r/min)	1 720	1 400
负载/Hp	0	1
故障类数	4	4
样本数	400	400
样本维度	1 024	10 24

本次采用的数据集原始数据维度很高,包含大量干扰信号,需要利用时域与频域分析对原始数据集进行预处理。本次预处理计算了数据集的 29 个特征信息,预处理采用的时域与频域参量如表 2 所列。

表 2 时域与频域参量统计表

Table 2 Statistics of time-domain and frequency-domain parameters

均值	标准方差
$P_1 = \frac{1}{U} \sum_{i=1}^U x(i)$	$P_2 = \sqrt{\frac{1}{U-1} \sum_{i=1}^U (x(i) - P_1)^2}$
均方根	方根均值
$P_3 = \sqrt{\frac{1}{U} \sum_{i=1}^U (x(i))^2}$	$P_4 = (\sqrt{\frac{1}{U} \sum_{i=1}^U x(i) })^2$
偏度	峭度
$P_5 = \frac{\sum_{i=1}^U (x(i) - P_1)^3}{(U-1)P_2^3}$	$P_6 = \frac{\sum_{i=1}^U (x(i) - P_1)^4}{(U-1)P_2^4}$
峰值因子	裕度因子
$P_7 = \frac{1}{P_3} \max x(i) $	$P_8 = \frac{1}{P_4} \max x(i) $
波形因子	脉冲因子
$P_9 = \frac{P_3}{\frac{1}{U} \sum_{i=1}^U x(i) }$	$P_{10} = \frac{\max x(i) }{\frac{1}{U} \sum_{i=1}^U x(i) }$
最大值	最小值
$P_{11} = \max x(i)$	$P_{12} = \min x(i)$
峰峰子	绝对均值
$P_{13} = \frac{P_{11} - P_{12}}{2}$	$P_{14} = \frac{1}{U} \sum_{i=1}^U x(i) $
对峰值因子	二阶中心距
$P_{15} = \frac{P_{11}}{P_{14}}$	$P_{16} = \frac{\sum_{i=1}^U (x(i) - P_1)^2}{U}$
能量	谱幅值均值
$P_{17} = \sum_{i=1}^U x(i)^2$	$P_{18} = \frac{1}{V} \sum_{i=1}^V s(i)$
谱幅值标准差	谱幅值偏度
$P_{19} = \sqrt{\frac{1}{V-1} \sum_{i=1}^V (s(i) - P_{18})^2}$	$P_{20} = \frac{\sum_{i=1}^V (s(i) - P_{18})^3}{V \sqrt{P_{19}^3}}$
谱幅值标准差	谱幅值偏度
$P_{21} = \frac{1}{VP_{19}^2} \sum_{i=1}^V (s(i) - P_{18})^4$	$P_{22} = \frac{1}{\sum_{i=1}^V s(i)} \sum_{i=1}^V (s(i) f_i)$
谱频率标准差	谱频率均方根
$P_{23} = \sqrt{\frac{1}{V} \sum_{i=1}^V (f_i - P_{22})^2 f_i}$	$P_{24} = \sqrt{\frac{\sum_{i=1}^V (f_i)^2 s(i)}{\sum_{i=1}^V s(i)}}$
谱根 4/2 矩比	谱频率能量比
$P_{25} = \frac{\sqrt{\sum_{i=1}^V (f_i)^4 s(i)}}{\sqrt{\sum_{i=1}^V (f_i)^2 s(i)}}$	$P_{26} = \frac{\sum_{i=1}^V (f_i)^2 s(i)}{\sqrt{\sum_{i=1}^V s(i) \sum_{i=1}^V (f_i)^4 s(i)}}$
谱变异系数	谱频率偏度
$P_{27} = \frac{P_{23}}{P_{22}}$	$P_{28} = \frac{\sum_{i=1}^V (f_i - P_{22})^3 s(i)}{VP_{23}^3}$
谱频率峭度	
$P_{29} = \frac{\sum_{i=1}^V (f_i - P_{22})^4 s(i)}{VP_{23}^4}$	

4.2 实验分析

在凯斯西储大学轴承故障数据集(CWRU)和东北石油大学轴承故障数据集(NEPU)上进行算法故障诊断实验,通过调整近邻点个数 k 的值,计算 PIE-LLE 算法故障数据的特征提取准确度,定量评估算法效果;进行可视化实验,直接观察降维效果;进行 Fisher 实验,定性评估算法的聚类效果。将

LE,LLE,HLL,SLLE,LMR-LLE 与本文算法进行对比试验,验证本算法相较于其他算法的优越性。

1)可视化对比实验

PIE-LLE 算法与其他算法在 CWRU 数据集上的可视化对比如图 7 所示。PIE-LLE 算法与其他算法在 NEPU 数据集上的可视化对比如图 8 所示。

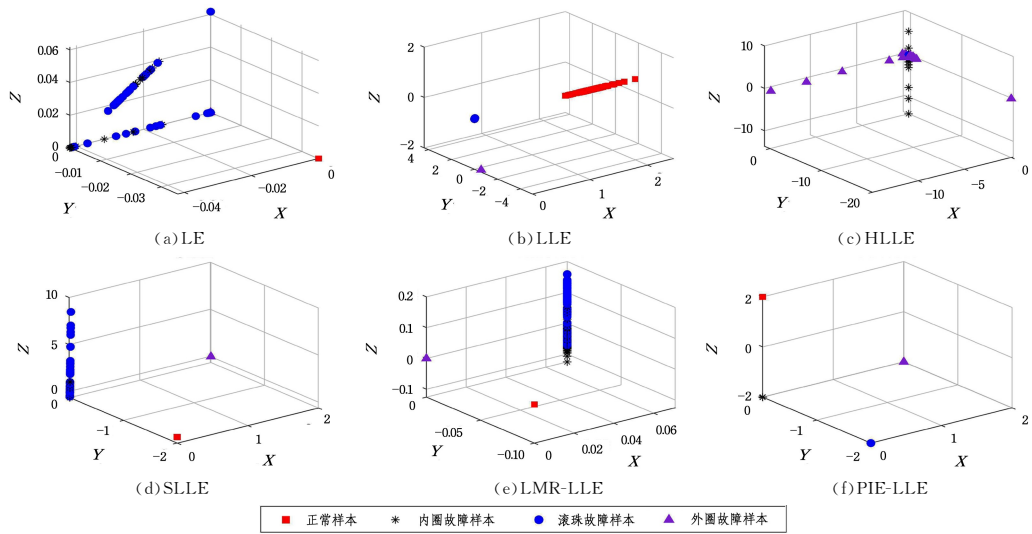


图 7 PIE-LLE 算法与其他算法在 CWRU 数据集上的低维可视化对比图

Fig. 7 Low dimensional visualization comparison between PIE-LLE algorithm and other algorithms on CWRU dataset

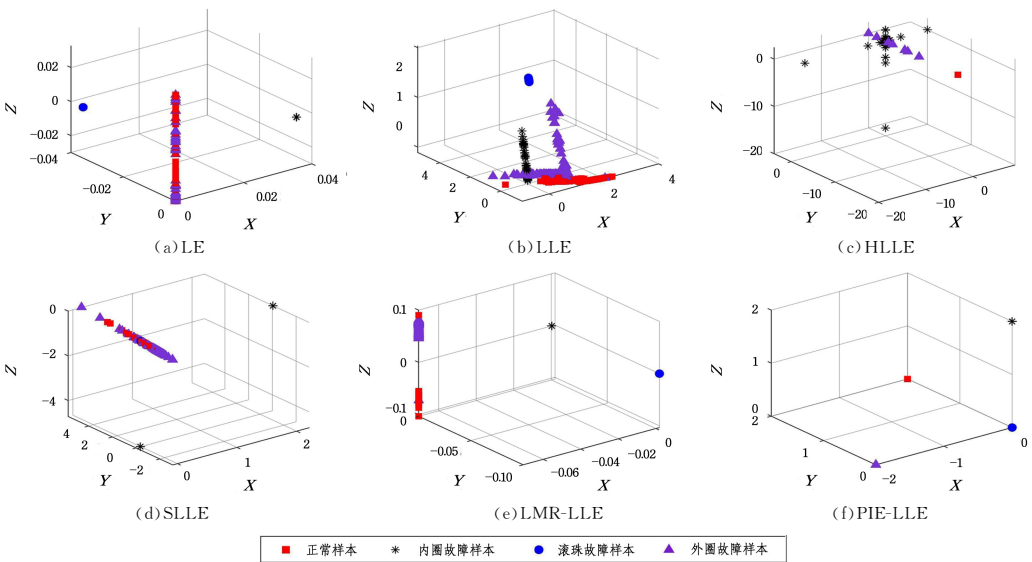


图 8 PIE-LLE 算法与其他算法在 NEPU 数据集上的低维可视化对比图

Fig. 8 Low dimensional visualization comparison between PIE-LLE algorithm and other algorithms on NEPU dataset

由以上可视化效果图可知,相较于 PIE-LLE 算法,LE 与 HLL 算法在两个数据集上降维后数据点大量混淆,无法提取数据的有效特征;SLLE 与 LMR-LLE 算法可以区分某两类数据,但其余两类数据点有所重叠,无法明确区分;LLE 算法在 CWRU 数据集上可以区分两类数据,在 NEPU 数据集上不能明显区分 4 类数据;PIE-LLE 算法则可以完全区分 4 类故障数据,提取数据点的显著特征。本算法在寻找样本点的近邻点时采用的方法更为合理,因此可以很好地提取数据的主要特征。

2)准确度实验

PIE-LLE 算法与其他算法在 CWRU 数据集上, k 值不同情况下的准确度如图 9 所示。

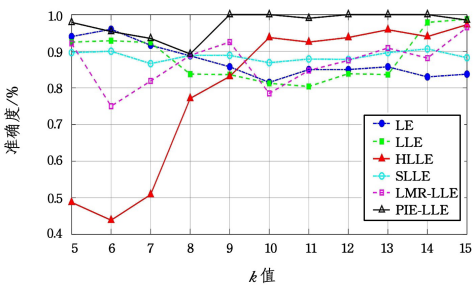


图 9 PIE-LLE 算法与其他算法在 CWRU 数据集上不同 k 值下的准确度对比

Fig. 9 Comparison of accuracy between PIE-LLE algorithm and other algorithms under different k in CWRU dataset

PIE-LLE 算法与其他算法在 NEPU 数据集上, k 值不同情况下的准确度如图 10 所示。

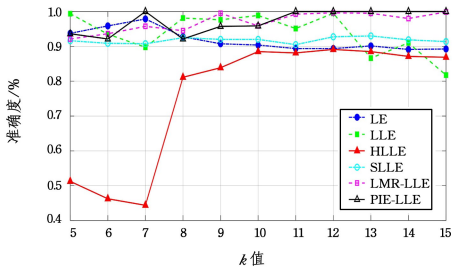


图 10 PIE-LLE 算法与其他算法在 NEPU 数据集上不同 k 值下的准确度对比

Fig. 10 Comparison of accuracy between PIE-LLE algorithm and other algorithms with different k on NEPU dataset

由上图可以得出,当选择的近邻点个数变化时,HLLE 特征提取的准确度受 k 值影响很大,随 k 值增大分类准确度呈上升趋势;LLE 算法受 k 值影响较大,准确度有所波动;LE 算法分类准确度随 k 值上升呈下降趋势;SLLE 算法的分类准确度基本低于 PIE-LLE 算法;LMR-LLE 算法在 CWRU 数据集上的表现劣于本文算法,在 NEPU 数据集上 k 值取值在个别情况下准确度高于本文算法。总体而言,PIE-LLE 算法在两个数据集上都表现出良好的特征提取效果,优于其他对比算法。

3)Fisher 聚类实验

本文算法与其他算法的 Fisher 参数对比结果如表 3、表 4 所列。

表 3 不同算法在 CWRU 数据集上的 Fisher 指标结果对比

Table 3 Comparison of Fisher metric results of different algorithms on CWRU dataset

算法	S_b	S_w	F
LE	5.77×10^{-4}	1.76×10^{-4}	3.28
LLE	1.50	0.76	1.98
HLLE	1.21×10^{-3}	3.00	4.05×10^{-4}
SLLE	1.24	1.02	1.21
LMR-LLE	4.10×10^{-3}	2.04×10^{-3}	2.01
PIE-LLE	2.26	5.21×10^{-23}	4.33×10^{22}

表 4 不同算法在 NEPU 数据集上的 Fisher 指标结果对比

Table 4 Comparison of Fisher metric results of different algorithms on NEPU dataset

算法	S_b	S_w	F
LE	4.92×10^{-4}	3.03×10^{-4}	1.63
LLE	1.08	1.37	0.78
HLLE	2.91×10^{-3}	3.00	9.71×10^{-4}
SLLE	1.01	1.56	0.65
LMR-LLE	5.48×10^{-3}	1.40×10^{-4}	39.19
PIE-LLE	2.26	9.50×10^{-25}	2.37×10^{24}

由表 3、表 4 可知,PIE-LLE 算法不同故障类别数据的类内散度最小,类间散度最大, F 值在所有算法中表现最优,在聚类性能上表现良好,证明了本文算法在该故障数据集上聚类的有效性 with 优越性。

4)融合权重参数实验

本文通过修改融合权重参数 α ,计算在不同 k 值下的 F 值,判断 α 对 PIE-LLE 算法聚类效果的影响,柱状统计图如图 11、图 12 所示。

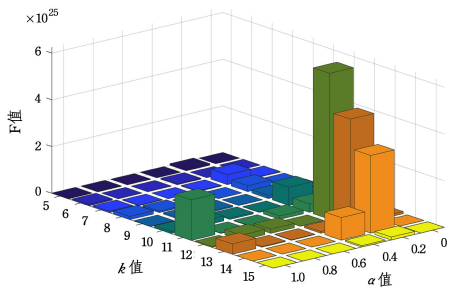


图 11 在不同 α, k 下的 F 值柱状统计图(CWRU 数据集)

Fig. 11 F-value bar chart with different α, k values(CWRU dataset)

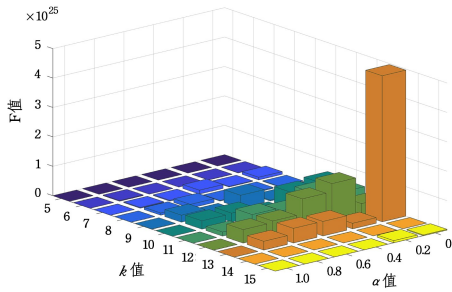


图 12 在不同 α, k 下的 F 值柱状统计图(NEPU 数据集)

Fig. 12 F-value bar chart with different α, k values(NEPU dataset)

根据以上参数实验,在更改 α 时,整体 Fisher 值均稳定在 1×10^{21} 以上。在 CWRU 数据集上, $k=12, \alpha=0.2$ 时, Fisher 值最大;在 NEPU 数据集上, $k=13, \alpha=0$ 时, Fisher 值最大。实验表明,本文算法对融合权重参数敏感性较低,仅在偶然参数下显示出更为优越的性能,在大部分实验参数中, F 值都可以稳定在较大范围内,算法稳定性优于对比算法。

4.3 实验结果

通过以上实验得出以下结论:

1)可视化效果实验通过将低维嵌入结果显示在三维坐标系中,能够观察不同类数据低维嵌入的分布情况;且与其他算法进行对比,可观察到 PIE-LLE 算法将同类数据基本聚集为一点,不同类数据被明显区分开来,这将极大提高后续模式识别效果。

2)准确度实验对比了不同 k 值下,不同的算法在 KNN 分类器下的分类准确度,结果表明 PIE-LLE 算法在 k 值从 5~15 之间变化时,本文算法的分类准确度能基本维持在 95% 以上,准确度最高达到了 100%。

3)Fisher 测度实验定量评估了本文算法对低维嵌入样本数据的类内散度与类间散度系数,本算法在对比实验中表现良好。

4)融合权重参数实验中,更改权重参数 α 时,本文算法的 Fisher 测度对比柱状图说明超过 80% 的数据均保持在较大范围,少部分数据表现出更加良好的效果,证明了本文算法在改变融合权重参数下的稳定性。

结束语 本文提出的 PIE-LLE 算法通过提取数据的概率分布特征与信息熵特征,作为判断样本邻域点范围的标准,符合数据集本身的分布情况,优化了原始 LLE 算法基于欧氏距离度量下邻域选择的不准确性,加入了非线性特征结构的挖掘,大幅度提高了对于数据集重要特征的提取效果,且在近邻点范围为 5~15 之间,融合权重参数为 0.0~1.0 之间表现出良好的稳定性,适用于高维复杂故障数据的快速降维,不受

量纲限制,可以高效提取数据集主要信息。下一步将继续对高维数据的概率模型拟合进行深入研究,以提高算法在不同故障数据集下的适应性和有效性,进一步降低算法复杂度。

参 考 文 献

- [1] SU Z, TANG B, MA J, et al. Fault diagnosis method based on incremental enhanced supervised locally linear embedding and adaptive nearest neighbor classifier[J]. *Measurement*, 2014, 48: 136-148.
- [2] MA X, ZHANG F, LI Y, et al. Robust sparse representation based face recognition in an adaptive weighted spatial pyramid structure[J]. *Science China Information Sciences*, 2018, 61: 1-13.
- [3] HE X, CAI D, YAN S, et al. Neighborhood preserving embedding[C]// Tenth IEEE International Conference on Computer Vision (ICCV'05) Volume 1. IEEE, 2005: 1208-1213.
- [4] HE X, YAN S, HU Y, et al. Face recognition using laplacianfaces[J]. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2005, 27(3): 328-340.
- [5] CAI D, HE X, HAN J, et al. Graph regularized nonnegative matrix factorization for data representation[J]. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2010, 33(8): 1548-1560.
- [6] NIE F, WANG Z, WANG R, et al. Adaptive local embedding learning for semi-supervised dimensionality reduction[J]. *IEEE Transactions on Knowledge and Data Engineering*, 2021, 34(10): 4609-4621.
- [7] WANG D, HOI S C H, HE Y, et al. Retrieval-based face annotation by weak label regularized local coordinate coding[C]// Proceedings of the 19th ACM International Conference on Multimedia. 2011: 353-362.
- [8] FANG X, HAN N, WU J, et al. Approximate low-rank projection learning for feature extraction[J]. *IEEE Transactions on Neural Networks and Learning Systems*, 2018, 29(11): 5228-5241.
- [9] TAN C, CHEN S, JI G, et al. A novel probabilistic label enhancement algorithm for multi-label distribution learning[J]. *IEEE Transactions on Knowledge and Data Engineering*, 2021, 34(11): 5098-5113.
- [10] HSIEH J W, CHUANG C H, CHEN S Y, et al. Segmentation of human body parts using deformable triangulation[J]. *IEEE Transactions on Systems, Man, and Cybernetics-Part A: Systems and Humans*, 2010, 40(3): 596-610.
- [11] SAUL L K, ROWEIS S T. An introduction to locally linear embedding[EB/OL]. <http://www.cs.toronto.edu/~roweis/lle/publications.html>.
- [12] WANG X, ZHENG Y, ZHAO Z, et al. Bearing fault diagnosis based on statistical locally linear embedding[J]. *Sensors*, 2015, 15(7): 16225-16247.
- [13] SOIZE C, GHANEM R. Data-driven probability concentration and sampling on manifold[J]. *Journal of Computational Physics*, 2016, 321: 242-258.
- [14] LIU Y, SHI B, LU S, et al. A novel local linear embedding algorithm via local mutual representation for bearing fault diagnosis[J]. *Reliability Engineering & System Safety*, 2024, 247: 110135.
- [15] CAI B, HUANG L, XIE M. Bayesian Networks in Fault Diagnosis[J]. *IEEE Transactions on Industrial Informatics*, 2017, 13(5): 2227-2240.
- [16] SOIZE C, GHANEM R. Probabilistic learning on manifolds[J]. *arXiv*, 2002. 12653, 2020.
- [17] VERMA V, FERNANDEZ J, SIMMONS R, et al. Probabilistic models for monitoring and fault diagnosis[C]// The Second IARP and IEEE/RAS Joint Workshop on Technical Challenges for Dependable Robots in Human Environments. Ed. Raja Chatila. 2002.
- [18] CANDUCCI M, TINO P, MASTROPIETROM. Probabilistic modelling of general noisy multi-manifold data sets[J]. *Artificial Intelligence*, 2022, 302: 103579.
- [19] WANG Z, YAO L, CHEN G, et al. Modified multiscale weighted permutation entropy and optimized support vector machine method for rolling bearing fault diagnosis with complex signals[J]. *ISA Transactions*, 2021, 114: 470-484.
- [20] SAAD F, CUSUMANO-TOWNER M, MANSINGHKA V. Estimators of entropy and information via inference in probabilistic models[C]// International Conference on Artificial Intelligence and Statistics. PMLR, 2022: 5604-5621.
- [21] KUMAR A, PARKASH C, VASHISHTHAG, et al. State-space modeling and novel entropy-based health indicator for dynamic degradation monitoring of rolling element bearing[J]. *Reliability Engineering & System Safety*, 2022, 221: 108356.
- [22] TANG F, FENG H, TINO P, et al. Probabilistic learning vector quantization on manifold of symmetric positive definite matrices[J]. *Neural Networks*, 2021, 142: 105-118.
- [23] ZHAO L, ZHANG Z. Supervised locally linear embedding with probability-based distance for classification[J]. *Computers & Mathematics with Applications*, 2009, 57(6): 919-926.
- [24] LIU Y H, ZHAO W B, ZHANG Y S, et al. Local linear embedding algorithms based on probability distributions[C]// 2022 IEEE International Conference on Advances in Electrical Engineering and Computer Applications (AEECA). IEEE, 2022: 1336-1339.



LIU Yuanhong, born in 1979, professor. His main research interests include nonlinear dimension reduction, machine learning and pattern recognition, signal processing.



WU Yubin, born in 2000, postgraduate. Her main research interests include data dimensionality reduction, manifold learning and pattern recognition.