

鲁棒最小二乘支持向量机及其 在软测量中的应用

司刚全, 娄勇, 张寅松

(西安交通大学电气工程学院, 710049, 西安)

摘要: 针对最小二乘支持向量机在利用产生于工业现场的非理想数据集进行建模预测时, 稀疏化模型鲁棒性差的问题, 提出了一种基于模糊 C 均值聚类和密度加权的稀疏化方法. 首先通过模糊 C 均值聚类将训练样本划分为若干个子类; 然后计算每个子类中各样本的可能贡献度, 依次从每个子类中选取具有最大可能贡献度的样本作为支持向量; 最后更新每个样本的可能贡献度, 继续从各个子集中增选支持向量, 直至稀疏化后的模型性能满足要求. 仿真结果和磨机负荷实际应用表明, 该方法能够兼顾模型在整体样本集和各工况子集上的性能, 在实现模型稀疏化的同时, 能够显著改善最小二乘支持向量机模型的鲁棒性.

关键词: 模糊 C 均值聚类; 密度加权; 鲁棒最小二乘支持向量机; 磨机负荷

中图分类号: TP18 **文献标志码:** A **文章编号:** 0253-987X(2012)08-0015-07

Robust Least Squares Support Vector Machines with Applications to Soft-Sensing

SI Gangquan, LOU Yong, ZHANG Yinsong

(School of Electrical Engineering, Xi'an Jiaotong University, Xi'an 710049, China)

Abstract: The sparse model for forecasting established by least squares support vector machines (LSSVM) is lacking in robustness, especially, for case of non-ideal data set produced in the industrial field as training data set. A fuzzy C -means clustering and density weighted based sparsity strategy is proposed. The training data set is divided into several subsets by fuzzy C -means clustering; the potential contribution of each sample is calculated and the sample with the greatest potential contribution in its own subset is selected as the support vector; the potential contribution of each sample is updated, more support vectors in the training data set are iteratively selected, until the user-defined performance is achieved. The simulation and applied examples indicate that the proposed strategy enables to achieve a sparse model with the corresponding character in the whole training data set and each subset, and the model robustness is improved significantly.

Keywords: fuzzy C -means clustering; density weighted; robust least square support vector machine; mill load

最小二乘支持向量机(LSSVM)^[1]是一种行之有效的软测量建模方法^[2], 由于具有运算效率高、收

敛速度快、预测精度高以及泛化能力好等一系列优点, 近年来得到了广泛应用. 与标准支持向量机

(SVM)相比,由于 LSSVM 中 α_n 不为 0,导致其失去了稀疏特性,从而使得预测模型变得复杂.因此,稀疏化算法成为 LSSVM 研究的热点. Suykens 最先提出的一种基于支持值 α_n 的 LSSVM 稀疏化方法,递归地将一定比例具有较小 $|\alpha_n|$ 的样本予以删除来实现模型的稀疏化.进一步, Suykens 在此基础上提出了加权 LSSVM^[3],对每个训练样本数据误差项 e_n 乘以一个权值,这样可以使得误差分布趋向于高斯分布,在加权 LSSVM 上实现了基于支持值的稀疏化算法.文献[4]为降低计算复杂度,用 γ/n 代替目标函数中的 γ ,同时将 w 用 $w = \sum_{n=1}^N \beta_n \phi_n$ 替换来解优化问题,实现稀疏化.文献[5]则提出了根据删除训练样本后重新建立的模型泛化误差最小的原则进行稀疏化,这样获得的模型具有更好的估计精度,但该方法计算量非常大.文献[6]对上述几种稀疏化算法进行了详细比较,认为在综合考虑训练过程计算量和模型估计精度的情况下, Suykens 等人提出的根据支持值 α_n 剪切样本的方法是最简单有效的.以上方法都不同程度地提高了模型的稀疏性,但都未考虑稀疏化过程中模型鲁棒性的变化,特别是在实际应用过程中,由于所收集的训练样本具有“各工况样本数量不平衡以及异方差”等一系列特性,从而导致稀疏化后的模型鲁棒性变差,难以在实际中获得应用.

为此,本文提出一种基于模糊 C 均值聚类和密度加权的鲁棒最小二乘支持向量机算法(简称 FDS 算法),该算法能兼顾 LSSVM 模型在整体样本集和各工况子集的性能,通过模糊聚类与密度加权相结合,在稀疏化过程中,剔除相似冗余样本,不仅能有效地改善 LSSVM 模型的稀疏性,同时使 LSSVM 模型的鲁棒性得到增强.

1 LSSVM 鲁棒性分析及问题提出

Suykens 提出的基于支持值的标准稀疏化算法(SVS 算法),对于无相似样本、样本均匀分布于整个数据域且具有统一的方差特性情况,能保持模型具有较好的鲁棒性.在实际应用中,所收集的训练样本往往具有大量的相似样本,在不同工况下获取的样本数量及方差特性具有显著差异,这就对 LSSVM 的稀疏化算法提出了挑战.例如:由于相似样本具有相近的支持值 α_n ,因此在稀疏化过程中,大量的具有较大支持值 α_n 的相似样本被保留作为支持向量,而具有较小支持值 α_n 的样本被全部删除,从

而导致模型性能下降,且模型中存在大量冗余支持向量.

在实际工业过程中,最常见的函数估计问题是对工业生产过程中难以在线测量的关键过程变量或质量参数进行拟合.由于被控对象的工作状态受实际生产过程中的产品型号差异、原材料差异、生产环境差异以及生产负荷水平差异等影响,往往可以将其划分为多种工况,不同工况下样本特征会有显著的差异,且由于实际生产情况的影响导致各工况下的样本数量差别较大,也即出现样本不平衡现象,含有较多样本的工况称为大类工况,含有较少样本的工况称为小类工况.在这种情况下,SVS 算法往往会保留更多的大类工况样本作为支持向量,而忽视小类工况样本.但是,这样的结果直接导致在大类工况下模型具有较高的估计性能指标,在小类工况下模型性能却很差,以至于无法满足工业监控的需要,甚至导致监控失效而发生运行事故.

样本数据的异方差特性,即在不同运行工况下,其获得的样本目标属性数据具有不同的方差水平.这主要是因为在不同工况下,由于运行环境的不同、原料和产品的变化等导致对目标值的测量带来的扰动不同所引起的,样本具有较小方差特性的工况称为小方差工况,样本具有较大方差特性的工况称为大方差工况.在应用 SVS 算法过程中,大方差工况下的样本必然会更多地被保留下来作为支持向量,而小方差工况下的样本却相反,只有少量的样本被保留作为支持向量.如此,经过稀疏化后的模型对小方差工况的估计性能大大降低,非理想样本的异方差特性对机器学习算法非常具有挑战性.

2 基于模糊 C 均值聚类的鲁棒最小二乘支持向量机

本文提出的基于模糊 C 均值聚类的鲁棒最小二乘支持向量机,首先采用模糊 C 均值聚类算法自适应地将在实际工业现场收集的大量训练样本进行工况划分,然后在稀疏化过程中,兼顾各工况的性能变化,同时在各子工况局部采用基于密度加权的稀疏化方法,剔除相似冗余样本,确保模型的稀疏性和鲁棒性.

2.1 模糊 C 均值聚类

聚类是模式识别和数据挖掘中广为使用的数据分析手段,它是在无法预知样本的类别时,根据样本间的距离或一定的相似性准则来自动进行分类的一

种方法^[7]. 传统的聚类分析是一种硬划分,把每个待分析的对象严格划分到某个类中,一旦某个对象属于某个类,那么它就不可能属于另外的一个类,这种类别的划分界限是分明的. 在工业实际中,由于原材料的变化、运行环境的变化以及运行负荷水平的变化导致的工况变化,并没有严格的工况界限,更多的是相对模糊的定义,也就是说很多样本数据具有亦此亦彼的特性,因此更适合于软划分. 本文采用模糊C均值(FCM)聚类的方法,实现训练样本集的工况自动划分,其思想是使得被划分到同一类的样本之间相似度最大,而不同类之间的相似度最小,这样更加有利于模型稀疏化过程. FCM利用了模糊数学处理聚类问题,获得了各训练样本属于不同工况的不确定程度(隶属度),建立起了样本对于工况类别的不确定描述,更能客观地反映工业对象的真实运行情况,划分后的每个子集可以看成不同的工况.

FCM算法的目标函数定义如下

$$\min J(\mathbf{U}, S_c) =$$

$$\sum_{n=1}^N \sum_{c=1}^C (u_{cn})^\kappa \|\{x_n, y_n\} - \{x_c, y_c\}\|^2 \quad (1)$$

满足约束条件

$$\sum_{c=1}^C u_{cn} = 1, \quad u_{cn} \in [0, 1] \quad (2)$$

式中: $\mathbf{U}=[u_{cn}]$ 为 $N \times C$ 维矩阵; u_{cn} 为第 n 个样本对第 c 类的隶属度; κ 为加权指数,又称为平滑指数或模糊参数,一般取值范围为 $1.5 \sim 2.5$, κ 越大分类的结果越模糊; C 为类的个数,需要根据数据的先验知识确定; $S_c (c=1, 2, \dots, C)$ 为划分后的类, $\{x_c, y_c\} (c=1, 2, \dots, C)$ 表示类 S_c 的聚类中心.

2.2 密度加权的局部稀疏化方法

训练样本集经过 FCM 划分后,形成了 C 个与实际运行工况相符的样本子集,但每个样本子集中仍可能包含大量的相似样本,必须在稀疏化的过程中将其剔除,避免冗余支持向量的出现. 为此本文采用局部稀疏化的思想,针对已经划分好的各个子集 S_c , 结合文献[8]所述的基于密度加权的稀疏化算法(DWS),在每个子集中比较其可能贡献度绝对值 $|P_n^\star|$ 的大小,并选取一定数目的较大 $|P_n^\star|$ 对应的样本为支持向量. 这种局部的稀疏化算法有两方面优点:①能够保证每个子集都有支持向量,兼顾了模型在整体和各个子集的性能,从而提高模型的鲁棒性;②能有效剔除冗余样本,快速实现稀疏化.

样本 $\{x_n, y_n\}$ 对估计模型的可能贡献度指标为

$$P_n^\star = D_n e_n \quad (3)$$

式中

$$D_n =$$

$$\sum_{m=1}^N \exp[-\|\{x_n, y_n\} - \{x_m, y_m\}\|^2 / (r_a/2)^2]$$

为样本 $\{x_n, y_n\}$ 处的密度指标, r_a 表示邻居半径. 经过密度加权后, $P_n^\star$ 既包含了样本 $\{x_n, y_n\}$ 对模型的误差贡献信息,又包含了训练样本集在样本 $\{x_n, y_n\}$ 处的密度信息,从而充分综合了样本 $\{x_n, y_n\}$ 及其邻域内其他样本对模型的误差贡献.

对可能贡献度指标的绝对值 $|P_n^\star|$ 进行排序,显然具有最大 $|P_n^\star|$ 的样本对模型的贡献度最大,选择对应样本作为支持向量,然后更新每个样本点的密度指标

$$D_n = D_n - D_{s1} \exp\left[-\frac{\|\{x_n, y_n\} - \{x_{s1}, y_{s1}\}\|^2}{(r_b/2)^2}\right] \quad (4)$$

式中: $\{x_{s1}, y_{s1}\}$ 为已经被选作支持向量的样本; $r_b = \lambda r_a$ (λ 为正数) 为最大可能支持向量 $\{x_{s1}, y_{s1}\}$ 的邻域范围. 从式(4)中可以看出,由于 $\{x_{s1}, y_{s1}\}$ 已经被选中为支持向量,则重新计算其密度指标时,结果 D_{s1} 变为 0,对应的贡献度指标 $P_{s1}^\star$ 也变为 0,从而避免了在后续过程中再次被选中. 在样本 $\{x_{s1}, y_{s1}\}$ 邻域内的其他样本,其密度指标得以大幅度降低,与支持向量距离越近,其密度指标削弱越多,如此其被选中为支持向量的可能性越小,从而有效避免了最终估计模型的冗余支持向量.

2.3 算法步骤

本文采用通用的均方误差(σ)指标来评价 LSSVM 模型的性能

$$\sigma = \frac{1}{N-1} \sum_{n=1}^N (y_n - f(x_n))^2 \quad (5)$$

输入: 训练样本集 $S_T = \{x_n, y_n\}_{n=1}^N$, 空支持向量样本集 $S_S = \emptyset$, 模型整体和各子集性能指标下降阈值 ξ^S_T 和 ξ^S_c , 每个稀疏化循环过程中增加的支持向量个数 M , 分类数目 C . 输出: 支持向量集 S_S , 采用 FDS 算法稀疏化后的 LSSVM 模型.

FDS 算法的程序流程图如图 1 所示.

3 仿真实验结果及分析

3.1 sinc 函数仿真分析

sinc 函数为

$$y = 10 \frac{\sin(x)}{x} + \varepsilon \quad (6)$$

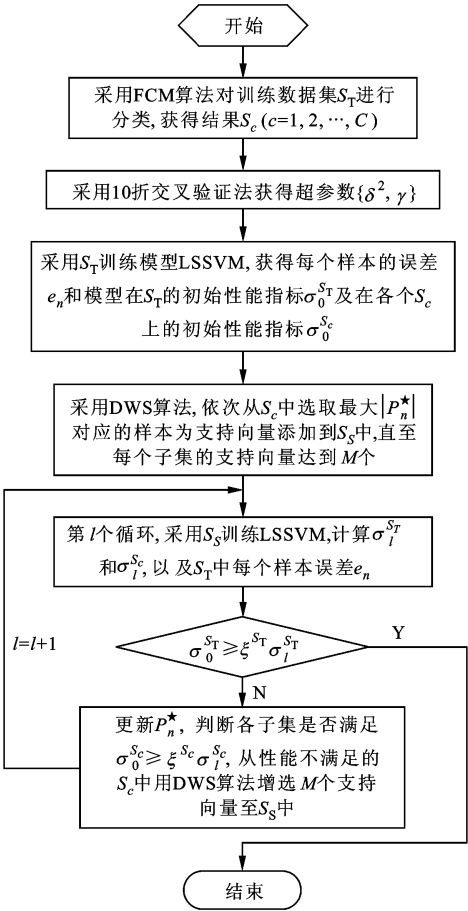
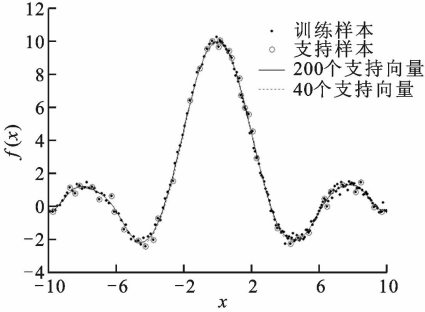
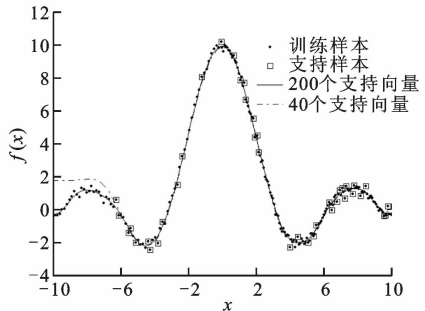


图 1 FDS 算法的程序流程图



(a) FDS 算法稀疏化结果



(b) SVS 算法稀疏化结果

图 2 样本数量不平衡 sinc 函数实例性能比较

对于样本数量不平衡的情况, sinc 函数的输入为 $x \in [-10, 10]$, 在区间 $x \in [-10, 0]$ 内均匀产生 70 个样本, 区间 $x \in (0, 10]$ 内均匀产生 130 个样本, ϵ 是均值为 0、方差为 0.2 的高斯噪声. 超参数 (δ^2, γ) 采用 10 折交叉验证法获得, 分别为 $\delta^2 = 0.054\ 74$ 和 $\gamma = 11.514$. FDS 算法中的参数选择如下: $r_a = 0.05, \lambda = 1.0, M = 1, \xi^{S_T} = 0.85, C = 3, \eta = 0.8$. 图 2 给出了 FDS 算法和 SVS 算法的稀疏化结果比较, FDS 算法得到的 LSSVM 模型含有 40 个支持向量, 此时模型性能为初始性能的 85.057%. 采用 SVS 算法进行稀疏化, 在同样保留 40 个支持向量的情况下, 模型性能下降为 15.629%. FDS 算法得到的模型在区间 $x \in [-10, 0)$ 和 $x \in (0, 10]$ 上的支持向量个数分别为 18 个和 22 个, 对应的 SVS 算法分别为 13 个和 27 个.

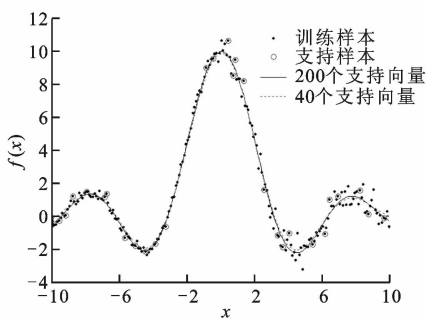
对于异方差的情况, sinc 函数的输入仍为 $x \in [-10, 10]$, 共有 200 个样本均匀分布. 在 $x \in [-10, 0)$ 区间内的 100 个样本, ϵ 是均值为 0、方差为 0.2 的高斯噪声; 在 $x \in (0, 10]$ 区间内的 100 个样

本, ϵ 是均值为 0、方差为 0.5 的高斯噪声. 超参数 (δ^2, γ) 分别为 $\delta^2 = 0.296\ 19$ 和 $\gamma = 6.046\ 6$. FDS 算法中的参数选择如下: $r_a = 0.05, \lambda = 1.0, M = 1, \xi^{S_T} = 0.85, C = 3, \eta = 0.8$. FDS 算法和 SVS 算法的稀疏化结果如图 3 所示. 采用 FDS 算法得到的 LSSVM 模型含有 28 个支持向量, 对应的模型性能为初始性能的 85.968%. 采用 SVS 算法进行稀疏化, 在同样保留 28 个支持向量的情况下, 模型性能下降为 23.988%. FDS 算法得到的模型在区间 $x \in [-10, 0)$ 和 $x \in (0, 10]$ 的支持向量个数分别为 13 个和 15 个, 对应的 SVS 算法分别为 1 个和 27 个.

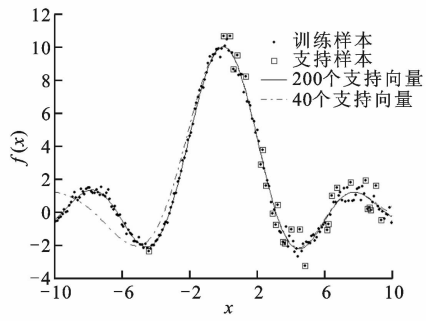
从图 2 和图 3 可以看出: SVS 算法在稀疏化过程中极易忽视小类工况和小方差工况的训练样本, 选取的支持向量区间分布极不平衡, 导致模型在这两种情况下预测性能变得很差; FDS 算法选取的支持向量分布均衡, 兼顾了各个不同工况的训练样本, 稀疏化后的模型具有更好的鲁棒性, 同时在保留相同数目的支持向量时, FDS 算法的模型性能远远高于 SVS 算法的模型性能.

3.2 Motorcycle 数据集仿真分析

对 Motorcycle 这一典型的异方差数据集进行仿真, 可以验证 FDS 算法在实际应用中的性能. 超参数 (δ^2, γ) 分别为 $\delta^2 = 0.448\ 03$ 和 $\gamma = 26.071$. FDS



(a)FDS 算法稀疏化结果



(b)SVS 算法稀疏化结果

图 3 异方差 sinc 函数实例性能比较

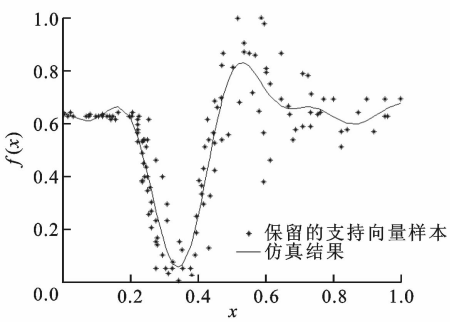
算法中的参数选择如下: $r_a=0.1,\lambda=1.0,M=1,C=5,\eta=0.8$.

经过 FDS 算法建立的含有不同数量支持向量的 LSSVM 模型估计结果如图 4 所示,含有 36、24、13 个支持向量的 LSSVM 模型的预测性能分别为初始模型(133 个支持向量)性能的 99.349%、95.985%、87.408%,图 5 所示为 SVS 算法的稀疏化结果,相应的含有 85、75、43 个支持向量的 LSSVM 模型的预测性能分别为初始模型性能的 99.207%、84.083%、55.823%.

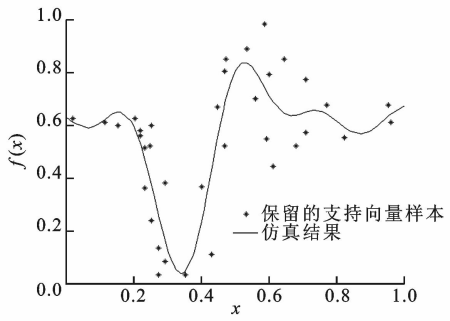
对比图 4 和图 5 可以看出:FDS 算法在模型的稀疏性能和估计性能上都有较大的提高,采用 FDS 算法稀疏化后的 LSSVM 模型在样本空间各个区域的预测能力随着支持向量数量的变化表现出一致的变化趋势,即在整个样本集上表现出了较好的鲁棒性;SVS 算法稀疏化的结果却完全不同,随着支持向量的减少,模型在起始部分和结束部分的性能相比中间部分下降得更快.

4 FDS 算法在磨机负荷软测量中的应用

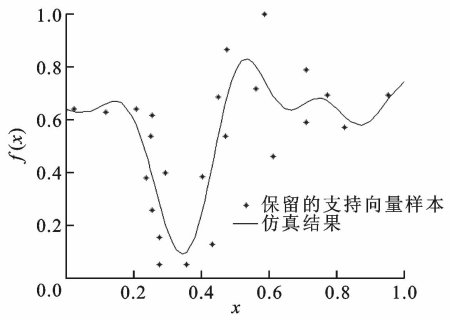
磨机负荷的准确检测是火电厂实施磨机负荷优



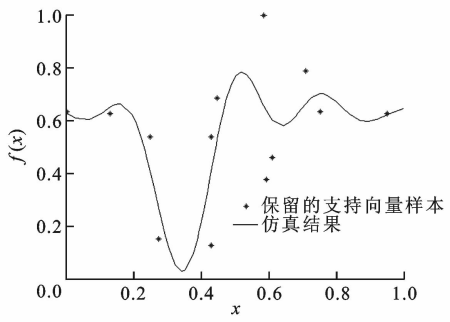
(a)133 个支持向量



(b)36 个支持向量



(c)24 个支持向量



(d)13 个支持向量

图 4 FDS 算法在 Motorcycle 数据集上的仿真结果

化控制的关键,但由于环境恶劣、影响因素众多等原因一直未能有效解决^[9].磨机负荷可以通过噪声、振动、磨机电流、磨机出入口压差以及磨机出口温度等多个信号以及信号中包含的特征量进行表征^[10].本

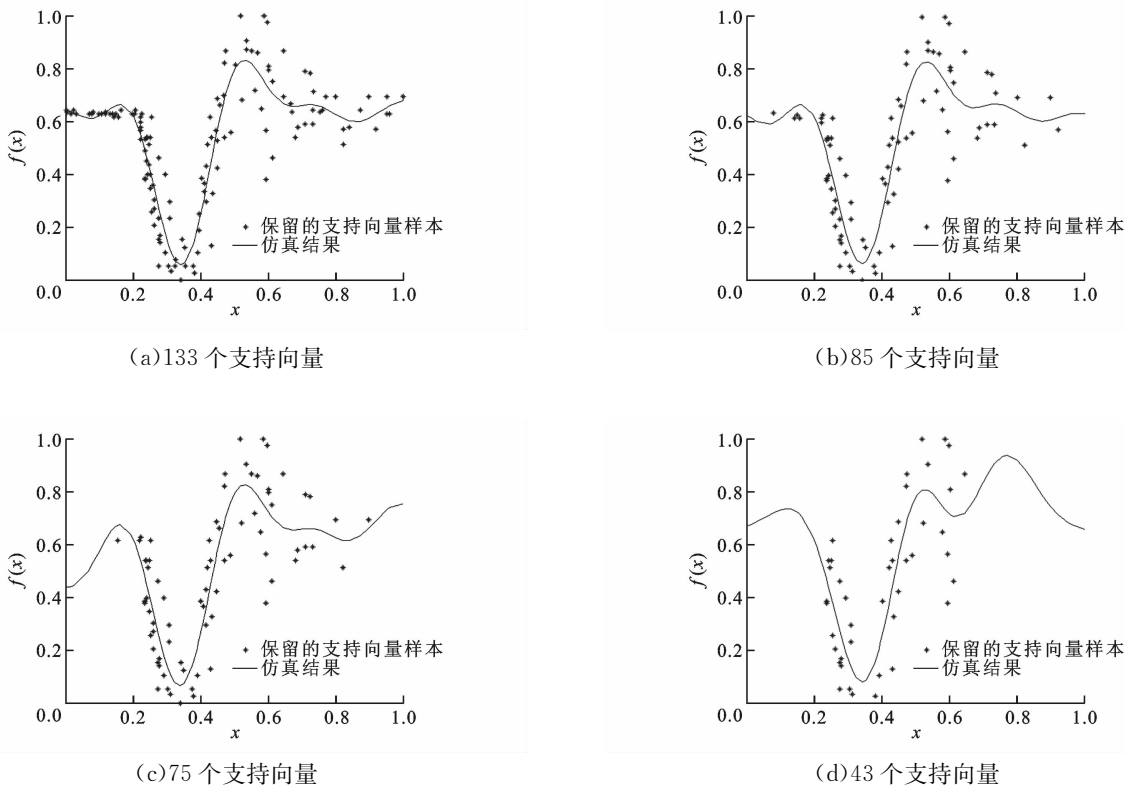


图 5 SVS 算法在 Motorcycle 数据集上的仿真结果

文基于 FDS-LSSVM 方法,构建了磨机负荷的软测量模型

$$y_{ml} = f(E_n, E_v, I, P_D, T) \tag{7}$$

式中: y_{ml} 为磨机负荷; E_n 、 E_v 、 I 、 P_D 和 T 分别为磨机噪声特征值、磨机振动特征值、磨机电流、出入口差压和出口温度; $f(\cdot)$ 为经过稀疏化后的 LSSVM 估计模型.

经过现场试验,共收集了 4 800 组样本,包含了磨机处于低负荷、正常负荷以及高负荷等工况,从中选取了 1 000 组样本作为训练样本集,320 组作为测试样本集.以模型性能下降阈值为 85%进行稀疏化,采用 FDS 算法和 SVS 算法稀疏化后的 LSSVM

模型信息如表 1 所示.从稀疏化结果对比可知,实际工业过程收集的训练样本中存在大量的相似样本,从而影响了标准稀疏化方法的效果.进一步比较发现,采用 FDS 算法时,支持向量较好地分布于各种工况下,在各工况下都具有较好的估计性能. SVS 方法获得的模型中的支持向量更多地分布于正常负荷工况,在低负荷工况和高负荷工况由于训练样本数量较少,因此保留的支持向量也较少,其估计性能较差.由此可见,采用本文方法所建立的模型鲁棒性得到了显著加强,且由于支持向量的大幅度减少,使得模型复杂度大大降低,从而有利于在工业中得以应用.

表 1 FDS 算法与 SVS 算法在磨机负荷软测量应用中的性能对比

算法	训练样本集整体		低负荷工况		正常负荷工况		高负荷工况	
	支持向量 个数	模型当前与 初始性能比/%	支持向量 个数	模型当前与 初始性能比/%	支持向量 个数	模型当前与 初始性能比/%	支持向量 个数	模型当前与 初始性能比/%
FDS-LSSVM	325	85.206	79	85.018	148	85.734	98	85.139
SVS-SSVM	786	85.417	114	68.847	639	92.383	33	37.284

5 结 论

本文提出了一种基于模糊C均值聚类和密度加权的鲁棒最小二乘支持向量机算法,该方法通过模糊C均值聚类,将从工业现场采集的非理想数据集按运行工况分成若干子集,针对每个子集进行密度加权稀疏化,这样既能剔除各个子集的冗余支持向量,使模型具有更好的稀疏性,同时兼顾了模型在整个数据集及各个子集上的性能,提高了模型的鲁棒性。函数仿真和在磨机负荷软测量中的实际应用表明,该方法比 Suykens 提出的标准稀疏化算法具有更好的稀疏性,同时在各个非理想数据集上鲁棒性较强,具有较高的工业实用价值。

参考文献:

- [1] SUYKENS J A K, GESTEL T V, BRABANTER J D, et al. Least squares support vector machines [M]. Singapore: World Scientific Publishing Company, 2002: 71-291.
- [2] KADLEC P, GRBIC R, GABRYS B. Review of adaptation mechanisms for data-driven soft sensors[J]. Computers & Chemical Engineering, 2011, 35(1): 1-24.
- [3] SUYKENS J A K, BRABANTER J D, LUKAS L, et al. Weighted least squares support vector machines: robustness and sparse approximation[J]. Neurocomputing, 2002, 48(1/2): 85-105.
- [4] CAWLEY G C, TALBOT N L C. Improved sparse least-squares support vector machines[J]. Neurocomputing, 2002, 48(1/2/3/4): 1025-1031.
- [5] DE KRUIF B J, DE VRIES T J A. Pruning error minimization in least squares support vector machines[J]. IEEE Transactions on Neural Networks, 2003, 14(3): 696-702.
- [6] HOEGAERTS L, SUYKENS J A K, VANDEWALLE J, et al. A comparison of pruning algorithms for sparse least squares support vector machines[M] // Lecture Notes in Computer Science: Vol. 3316. Berlin, Germany: Springer, 2004: 1247-1253.
- [7] TSAI Duming, LIN Chungchan. Fuzzy C-means based clustering for linearly and nonlinearly separable data [J]. Pattern Recognition, 2011, 44(8): 1750-1760.
- [8] 司刚全, 曹晖, 张彦斌, 等. 一种基于密度加权的最小二乘支持向量机稀疏化算法[J]. 西安交通大学学报,

2009, 43(10): 11-15.

SI Gangquan, CAO Hui, ZHANG Yanbin, et al. Density weighted pruning method for sparse least squares support vector machines[J]. Journal of Xi'an Jiaotong University, 2009, 43(10): 11-15.

- [9] KATUBILWA F M, MOYS M H. Effects of filling degree and viscosity of slurry on mill load orientation [J]. Minerals Engineering, 2011, 24(13): 1502-1512.
- [10] TANG Jian, ZHAO Lijie, ZHOU Junwu, et al. Experimental analysis of wet mill load based on vibration signals of laboratory-scale ball mill shell [J]. Minerals Engineering, 2010, 23(9): 720-730.

[本刊相关文献链接]

丁要军, 蔡皖东. 采用两阶段策略模型(KTSVM)的P2P流量识别方法. 2012, 46(2): 45-50.

李大海, 李天石. 非均匀采样系统的支持向量回归建模与控制. 2011, 45(3): 65-69.

李志雄, 严新平. 独立分量分析和流形学习在VSC-HVDC系统故障诊断中的应用. 2011, 45(2): 44-48.

张周锁, 侯照文, 孙闯, 等. 应用粒计算的混合智能故障诊断技术研究. 2011, 45(1): 48-53.

国强, 王常虹, 李峥, 等. 支持向量聚类联合类型熵识别的雷达信号分选方法. 2010, 44(8): 63-67.

司刚全, 曹晖, 张彦斌, 等. 一种基于密度加权的最小二乘支持向量机稀疏化算法. 2009, 43(10): 11-15.

唐浩, 廖与禾, 孙峰, 等. 具有模糊隶属度的模糊支持向量机算法. 2009, 43(7): 40-43.

穆向阳, 张太镒, 周亚同. 尺度核函数在最小二乘支持向量机信号逼近中的应用. 2008, 42(12): 1464-1467.

唐耀华, 高静怀, 包乾宗. 一种新的选择性支持向量机集成学习算法. 2008, 42(10): 1221-1225.

司文荣, 李军浩, 杨景刚, 等. 局部放电脉冲群的实用快速分类技术及应用. 2008, 42(8): 1021-1025.

吴晓辉, 刘炯, 梁永春, 等. 支持向量机在电力变压器故障诊断中的应用. 2007, 41(6): 722-726.

苟世宁, 杜海峰, 栗茂林, 等. 一种改进的模糊人工免疫网络数据分类方法. 2007, 41(5): 585-588.

田军委, 黄永宣, 于亚琳. 基于直方图偏差约束的快速模糊C均值图像分割法. 2007, 41(4): 430-434.

张瑞, 马逸尘, 段现报. 面向分类去噪问题的模糊支持向量机新算法. 2007, 41(12): 1414-1417.

(编辑 杜秀杰)