

論文 / 著書情報
Article / Book Information

題目(和文)	クラス事前確率変化下における統計的学習 - クラス事前確率推定及びその高次元識別問題への応用 -
Title(English)	Statistical Learning under Class-prior Change: Class-prior Estimation and Its Application for High-dimensional Classification
著者(和文)	川久保秀子
Author(English)	Hideko Kawakubo
出典(和文)	学位:博士(工学), 学位授与機関:東京工業大学, 報告番号:甲第10552号, 授与年月日:2017年3月26日, 学位の種別:課程博士, 審査員:篠田 浩一,杉山 将,村田 剛志,下坂 正倫,鈴木 大慈,徳永 健伸
Citation(English)	Degree:Doctor (Engineering), Conferring organization: Tokyo Institute of Technology, Report number:甲第10552号, Conferred date:2017/3/26, Degree Type:Course doctor, Examiner:,,,,,
学位種別(和文)	博士論文
Category(English)	Doctoral Thesis
種別(和文)	論文要旨
Type(English)	Summary

論文要旨

THESIS SUMMARY

専攻:	計算工学	専攻
Department of		
学生氏名:	川久保秀子	
Student's Name		

申請学位(専攻分野):	博士 (工学)
Academic Degree Requested	Doctor of
指導教員(主):	篠田浩一 教授
Academic Advisor(main)	
指導教員(副):	杉山将 特定教授
Academic Advisor(sub)	

要旨 (英文 800 語程度)
Thesis Summary (approx.800 English Words)

In machine learning classification problems, class balances between training and test data are generally assumed to be equal. However, sample selection bias or non-stationarity of the environment often causes a violation of this assumption in real-world classification problems. Such violations in which class balances differ between training and test data are referred to as *class-prior change*. This doctoral thesis is devoted to developing statistical approaches for classification when *class-prior change* is an issue.

It is a well-known fact that a naively trained classifier systematically yields a biased solution under *class-prior change*. Fortunately, a weighted classifier, which is trained according to the ratio of the class-priors of the training and test data, can correct such a bias. However, training this weighted classifier becomes a problem in real-world applications since class-priors of test data are generally unknown. Recently, several class-prior estimation methods, which directly estimate class-priors by minimizing the distance between the true test distribution and a model of the test input distribution, have been proposed under a semi-supervised learning setup. These existing class-prior estimation methods assume that the class-wise training and test input distributions are equal, and use a mixture of class-wise training input distributions as a model for the test input distribution. The advantage of these existing class-prior estimation methods is that they perform the task of fitting the model of the test input distribution with the true test input distribution, *distribution matching*, without first going through the procedure of estimating class-wise training input distributions and a test input distribution. In the framework of *distribution matching*, existing class-prior estimation methods employ various divergences for distance minimization, for instance, the Kullback-Leibler divergence, the Pearson divergence, the L_2 distance, and maximum mean discrepancy (MMD). In this doctoral thesis, we propose to use *energy distance* (ED) for class-prior estimation in the framework of *distribution matching*. Our proposed method based on ED is extremely simple and highly computationally efficient since ED requires no parameter tuning for the computation of the distance between two distributions. ED may be interpreted as a special case of MMD which utilizes a particular kernel function called a *distance kernel*, and thus our contribution can be regarded as providing a practical choice of the kernel function for the existing MMD-based class-prior estimation method. Experimental results show that all class-prior estimation methods,

including our ED-based method, have similar improved accuracies when using a weighted classifier in contrast to using a non-weighted classifier. However, our method significantly outperforms other class-prior estimation methods in terms of computation time.

Next, we tackle the problem of high-dimensional classification under *class-prior change* by applying the ED-based class-prior estimator to this problem. When the dimensionality of input data is high, it is known that data processing becomes drastically more complicated. This issue, which is referred to as the *curse of dimensionality*, is recognized as one of the most essential problems in machine learning. Using a supervised dimension reduction method, *sufficient dimension reduction* (SDR), is an effective and popular approach to overcoming the *curse of dimensionality*. SDR aims at reducing the dimensionality of input data so that information loss on output is minimal. Various SDR methods have been proposed so far. Among them, we focus on a novel SDR method, which is called *least-squares gradients for dimension reduction* (LSGDR). LSGDR estimates the gradients of the logarithmic conditional density since the gradient belongs to the span of a low-rank projection matrix which is used for dimension reduction. The advantage of using LSGDR is: no parametric assumption is required for data distributions, all the parameters in the estimation for the gradients of the logarithmic conditional density can be determined via cross-validation, and no iterative optimization procedure hinders the computational efficiency. However, since LSGDR, like all SDR methods, is a supervised method, labeled samples are required for finding a low-dimensional subspace that maintains maximal information on output data. This limitation may lead to low classification performance under *class-prior change*. To overcome the problem of *class-prior change* in high dimension, we propose to combine LSGDR with the ED-based class-prior estimation under a semi-supervised learning setup. Our SDR method, weighted LSGDR (WLSGDR), approximates a low-dimensional subspace for test data appropriately in spite of the fact that labeled test samples are not available. Through experiments, we demonstrate that WLSGDR is particularly useful when the class-priors of training and test data are highly imbalanced.

We conclude that our proposed methods, the ED-based class-prior estimator and WLSGDR, are promising in improving both the classification accuracy and computational efficiency, and therefore deserve further exploration.

備考：論文要旨は、和文 2000 字と英文 300 語を 1 部ずつ提出するか、もしくは英文 800 語を 1 部提出してください。

Note: Thesis Summary should be submitted in either a copy of 2000 Japanese Characters and 300 Words (English) or 1 copy of 800 Words (English).

注意：論文要旨は、東工大リサーチリポジトリ(T2R2)にてインターネット公表されますので、公表可能な範囲の内容で作成してください。

Attention: Thesis Summary will be published on Tokyo Tech Research Repository Website (T2R2).