

基于 MapReduce 的特征选择并行化研究

陆江 李云

(南京邮电大学计算机学院 南京 210023)

摘要 特征选择已经成为一种对高维数据进行预处理的必不可少的手段。随着数据规模的爆炸性增长,传统的特征选择算法已经不能满足当前高维大规模数据的处理要求。采用 Google 的 MapReduce 编程模型,设计了一种分布式的基于局部学习的特征选择算法 D-logsf。在多个现实和合成数据集上的实验表明,分布式特征选择算法 D-logsf 具有较好的可靠性,且与传统特征选择算法 Logsf 相比可以获得接近线性的加速比,同时可以有效处理大规模数据集。

关键词 特征选择,局部学习,分布式,MapReduce

中图分类号 TP391 **文献标识码** A **DOI** 10.11896/j.issn.1002-137X.2015.8.009

MapReduce Based Feature Selection Parallelization

LU Jiang LI Yun

(College of Computer Science, Nanjing University of Posts and Telecommunications, Nanjing 210023, China)

Abstract Feature selection has become a necessary preprocessing procedure for high-dimensional data. With the explosive growth of data size, the traditional feature selection algorithm can not meet the current requirements of processing large-scale and high-dimensional data. Resorting to Google's MapReduce programming model, we designed a distributed local learning-based feature selection algorithm D-logsf. Experiments were conducted on several real and synthesis data sets. The results show that the D-logsf algorithm is correct and has good reliability. Compared with traditional feature selection algorithm Logsf, D-logsf can obtain approximate linear speedup. Moreover, D-logsf can effectively handle large-scale data set.

Keywords Feature selection, Local learning, Distributed, MapReduce

1 引言

随着信息技术和生物技术的飞速发展,大量的数据相应产生。他们具有样本数量大、特征维度高等特点,因此利用特征选择算法对其进行预处理时,所面临的计算问题日益加剧^[1]。面对高维大规模数据,传统的特征选择算法已经不能够在合理的时间内得到需要的模型。因此,需要研究新的有效处理大规模数据的特征选择方法。目前处理大规模数据最有效的方法之一就是分布式并行计算^[2],而最具代表性的分布式计算框架当属 Google 工程师 J. Dean 等提出的 MapReduce 编程模型^[3]

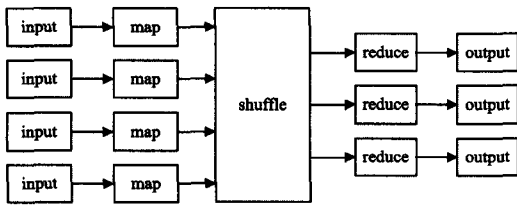


图1 MapReduce 编程模型的数据处理流程

本文从高维大规模数据出发,分析了基于局部学习的特征选择算法 Logsf,并对 Logsf 进行并行化,提出了能够处理大数据的分布式特征选择算法 D-logsf。第2节从局部学习的角度分析特征选择算法 Logsf;第3节提出分布式算法 D-logsf 以及分析 Hadoop 平台对 D-logsf 算法的影响因素;第4节在现实和人工合成数据集上进行实验,验证所提算法的可靠性和性能;最后总结全文。

2 基于局部学习的特征选择算法 Logsf

局部学习特征选择算法的主要思想是将任一复杂的非线性问题转化为一组局部线性问题。对于存在大量不相关特征的数据集,采用基于局部学习的特征选择算法可以获得较理想的结果^[11]。假设训练样本集 $D = \{X, Y\} = \{x_i, y_i\}_{i=1}^N$, 其中 x_i 是第 i 个训练样本, y_i 是与其对应的标记, N 代表训练样本数量。每个训练样本是由 d 维向量组成 $x_i = (x_{i1}, x_{i2}, \dots, x_{id}) \in R^d$ 。则样本 x_i 的基于局部学习的损失函数定义如下:

$$L(w, x_i) = \log(1 + \exp(-w^T z_i)) \quad (1)$$

其中, $z_i = |x_i - \text{nearmiss}(x_i)| - |x_i - \text{nearhit}(x_i)|$ 。 $||$ 是一个绝对值运算符, $\text{nearhit}(x_i)$ 代表与样本 x_i 标记相同且距离

作为最终的特征选择结果。

$$w = \frac{1}{k} \sum_{i=1}^k w_i^l \quad (8)$$

得到的权重 w 就是最优权重值 w^* 。

算法 2 分布式随机梯度下降算法的伪代码

输入: 训练集 $D = \{X, Y\} = \{x_i, y_i\}_{i=1}^N$, 步长 α , 机器数 k

输出: w^*

1. $T = \lfloor N/k \rfloor$, 初始化特征权重 $w = (1, 1, \dots, 1) \in \mathbb{R}^d$
2. 随机为每个节点分配 T 个样本
3. for all $l \in \{1 \dots k\}$ // 这一步是 k 个节点并行计算
4. 第 l 个节点上的数据随机读取
5. for all $t \in \{1 \dots T\}$
6. 读取第 l 个节点的第 t 个样本 x_t^l
7. 对于样本 x_t^l , 计算其 $\text{nearhit}(x_t^l)$, $\text{nearmiss}(x_t^l)$
8. 计算 $z_t^l = |x_t^l - \text{nearmiss}(x_t^l)| - |x_t^l - \text{nearhit}(x_t^l)|$
9. 依据 $w^l = w^l - \alpha * \frac{\partial L(w^l, x_t^l)}{\partial w}$, 更新特征权重 w^l
10. end for
11. end for

12. 集成所有节点的特征加权结果 $w = \frac{1}{k} \sum_{l=1}^k w^l$

表2 人工合成数据集 Rs 的特征选择结果

算法	特征选择重要特征排序
Logsf	1,2,3,4,5,6,7,9,8,10,18,13,16,12,15,19,17,20,21,11
D-logsf	1,2,3,4,5,7,6,9,8,12,18,16,19,13,10,11,17,15,20,14

从图2可以看出,数据集 Diabetes 使用 Logsf 算法得到稳定的分类准确率,而 D-logsf 算法得到的分类准确率不是很稳定,但其整体分类准确率还是在 Logsf 算法附近波动。在数据集 Ionosphere 上随着特征数量的增加 Logsf 算法和 D-logsf 算法保持一致趋势,虽然同等条件下 D-logsf 的分类准确率比 Logsf 稍微低一些,但是依然可以认为它们具有一致的分类准确率。数据集 Sonar 选择的特征大于 20 时,Logsf 算法和 D-logsf 算法基本保持了一致的分类准确率;而选择的特征小于 20 时,Logsf 算法的分类准确率比 D-logsf 算法要稳定而且更高。数据集 Gisette 选择的特征小于 4000 时,D-logsf 算法和 Logsf 算法的分类准确率保持一致的走势且前者略高于后者;但是维数大于 4000 时,Logsf 算法的分类性能要明显高于 D-logsf。综上所述,可以大体认为 D-logsf 算法与 Logsf 算法在同等条件下得到的特征选择器具有一致的性能。因此,基于局部学习的分布式特征选择算法 D-logsf 具有可靠性。从表2可以看出算法 Logsf 与 D-logsf 可以实现重要特征的靠前排序,进一步证明了 D-logsf 算法具有可靠性和正确性。

第二组实验:对数据集 Gisette 分别随机地抽取 100、200、300、400、500、1000、1500、2000 条样本作为训练数据集。分别利用 Logsf 和 D-logsf 依据第一组实验过程进行特征选择,随着数据规模增加,Logsf 和 D-logsf 算法的时间开销实验结果如图3所示,对应的加速比如图4所示。

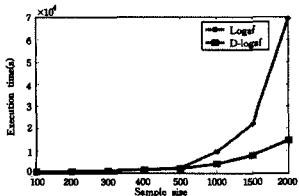


表2 在 Yale 人脸库上不同算法的最大识别率及对应的特征维数

	LDA	MMC	MMMC		
			2×2	2×4	4×4
4 正确识别率	68.09	74.76	85.71	85.71	87.62
鉴别特征维数	7	9	14	9	11
5 正确识别率	71.11	82.53	85.56	85.56	87.78
鉴别特征维数	3	10	9	8	8
6 正确识别率	77.33	82.67	88.50	88.00	89.67
鉴别特征维数	13	2	8	9	12

4.3 实验结果分析

(1)表1、表2中的数据均表明无论采用哪种分块方式,实验结果都比原来方法的人脸识别率要高。随着训练样本数目的不同,正确识别率也有所不同。

(2)从图4及图7中可以看出,对于3种不同的分块方式,本文方法的结果均明显优于LDA及MMC方法。

(3)Modular MMC方法是MMC方法的推广,通过对图像进行分块,再对每一子块抽取得到与直接将MMC方法用于原始图像抽取的全局特征相比更具有鉴别性的局部特征,这些局部特征更能表现图像