

DOI: 10.7652/xjtuxb201512022

采用稀疏 SIFT 特征的车型识别方法

张鹏¹, 陈湘军^{1,2}, 阮雅端¹, 陈启美¹

(1. 南京大学电子科学与工程学院, 210046, 南京; 2. 江苏理工学院计算机工程学院, 213001, 江苏常州)

摘要: 针对实际应用 中因图像清晰度低等因素导致的车型识别误差过大的问题,提出了一种基于稀疏尺度不变转换特征(sparse scale invariant feature transform, S-SIFT)的车型识别方法。该方法用背景建模方法检测交通视频运动目标,提取目标 SIFT 特征;通过 L1 约束计算出 SIFT 特征的稀疏编码,并用最大池化方法降低稀疏编码维度,在线性 SVM 分类器中完成车型分类,弥补了背景建模方法识别误差过大、不具备车型分类功能的缺陷。经 G36 高速公路实际应用表明:算法对车辆场景识别率可达 98% 以上,车型识别准确率可达 89% 以上,对低清晰度、不同视角、雨雪、遮挡等场景有很好的鲁棒性;图像平均处理时间不超过 40 ms,可满足系统对实时性的要求,在准确率和时间效率两方面均明显优于传统的 SIFT 方法和 HOG 方法。

关键词: 深度学习;车型识别;稀疏特征;尺度不变转换特征;线性支持向量机分类

中图分类号: TP391.4 **文献标志码:** A **文章编号:** 0253-987X(2015)12-0137-07

A Vehicle Classification Technique Based on Sparse Coding

ZHANG Peng¹, CHEN Xiangjun^{1,2}, RUAN Yaduan¹, CHEN Qimei¹

(1. School of Electronic Science and Engineering, Nanjing University, Nanjing 210046, China;

2. School of Computer Engineering, Jiangsu University of Technology, Changzhou, Jiangsu 213001, China)

Abstract: A new method based on sparse scale invariant feature transform(S-SIFT) is proposed to improve the vehicle recognition rate in environment such as low image quality. Moving objects are detected using a Gaussian mixture background subtraction model and SIFT features of the objects are calculated. Then, the sparse coding of SIFT features is obtained through L1 constraint. A max pooling strategy is introduced to reduce the dimension of the sparse coding. Finally, a linear support vector machine (SVM) is used to classify and to recognize the objects. The method solves the problems that the background modeling has a larger error rate and lacks function of vehicle classification. An application of the technique on G36 highway shows that the algorithm has an excellent result on different scenes such as low resolution, different camera angles, sleet and shade. The experimental results provide a more than 98% scene recognition rate, and a more than 89% classification accuracy rate. Moreover, the average time to process images is less than forty milliseconds, and it meets the real-time requirement. It is concluded that the proposed method is better than the SIFT and the HOG methods on both accuracy and time efficiency.

Keywords: deep learning; vehicle recognition; sparse feature; scale invariant feature transform; linear support vector machine classification

交通监控视频信息内容形象直观、铺设方便、覆盖范围广泛,基于机器视觉的车型识别方法已在智

能交通 ITS 领域逐步得到应用。目前常用的车型识别技术包括模板匹配^[1]、尺度不变特征变换(scale invariant feature transform, SIFT)结合 SVM 分类器^[2]、背景建模^[3]等。模板匹配需要对图像扫描,计算量较大,不适用于实时系统;SIFT 特征方法在视频不清晰、特殊天气状况下识别准确率不高^[4];背景建模方法基于帧间像素动态变化解析,实时性强,应用较广泛,但其对场景很敏感,光线变化、摄像机抖动、雨滴、树枝摇晃等均可能造成误判为运动目标,需进一步判别目标。

车辆长度、轮廓特征常用作分类特征,但随摄像机的距离远近而发生尺度变化,不适合用于监控视频。方向梯度直方图(histogram of oriented gradients, HOG)是 Dalal 等提出的一种目标检测算法^[5],用图像梯度的统计信息描述图像局部形状,可在一定程度上抑制平移和旋转的影响,但很难处理遮挡问题,并且由于梯度的性质,对噪点很敏感。SIFT 是一种基于尺度空间的算子,是基于关键点特征向量的描述,它对图像缩放、旋转都能够保持不变性,可以有效描述图像局部特征。

作为一种无监督学习方法,稀疏编码通过训练低层特征向量得到一组超完备基向量,用基向量的线性组合来表示输入图像特征,可对图像像素或已有特征做进一步抽象。稀疏模型在超分辨率重建^[6]、图像分割^[7]、图像分类^[8]等领域已经有相关研究。Yang 等用 SIFT 特征结合空间金字塔(spatial pyramid matching, SPM)作为低层向量,训练出用于稀疏编码的基向量,取得了较好的图像分类效果^[9]。尽管稀疏编码在图像分类领域已引起了广泛关注,但将其应用于公路车辆识别和分类的研究还很少。

本文基于深度学习理论,提出了一种基于稀疏 SIFT 特征的车型识别的方法。该算法用高斯混合背景差分技术提取运动目标以减少计算量,保证系统实时性;提取目标图像的低层 SIFT 特征向量,再经训练获得编码字典和稀疏 SIFT 特征,得到更深层次图像特征,以适应不同视角、光照变化、阴影、遮挡等复杂场景,进一步提高识别率;最后用线性支持向量机实现稀疏 SIFT 特征分类,降低时间复杂度,保证实时性。

1 S-SIFT 特征模型

1.1 S-SIFT 特征算法

S-SIFT 特征算法是在图像 SIFT 特征的基础

上,进一步训练超完备字典基,在 L1 约束下编码的稀疏 SIFT,可以实现更高层次车辆图像抽象。

定义矩阵 $\mathbf{X}$ 包含图像在 D 维特征空间的 M 个 SIFT 特征描述子, $\mathbf{X} = (\mathbf{x}_1, \dots, \mathbf{x}_M)^T$, 则 $\mathbf{X}$ 可以表示为

$$\mathbf{X} = \mathbf{W}\mathbf{C} \quad (1)$$

式中: $\mathbf{W}$ 是稀疏编码系数; $\mathbf{C} = (\mathbf{c}_1, \dots, \mathbf{c}_K)^T$ 是 K 个基向量。求解 $\mathbf{X}$ 的稀疏编码可以表征为下式对 $\mathbf{W}$ 和 $\mathbf{C}$ 求解最优化问题

$$\min_{\mathbf{W}, \mathbf{C}} \sum_{m=1}^M \|\mathbf{x}_m - \mathbf{w}_m \mathbf{C}\|^2 + \beta |\mathbf{w}_m| \quad (2)$$

$$\|\mathbf{c}_g\| \leq 1, \forall g = 1, \dots, K$$

式中: $\|\cdot\|$ 和 $|\cdot|$ 分别表示 L2 范数和 L1 范数。由 L1 约束性质可知,惩罚项 $|\mathbf{w}_m|$ 保证了编码结果的稀疏性,稀疏系数 β 控制 $|\mathbf{w}_m|$ 的权重,即稀疏性。基向量是过完备的($K > D$),因此用 $\mathbf{c}_g$ 的 L2 约束避免平凡解。

虽然求解式(2)时 $\mathbf{W}$ 和 $\mathbf{C}$ 同时变化,目标函数不是凸优化问题,但是分别固定 $\mathbf{W}$ 和 $\mathbf{C}$ 时,目标函数分别退化为关于 $\mathbf{C}$ 和 $\mathbf{W}$ 的凸函数。固定 $\mathbf{W}$ 时,目标函数退化为关于 $\mathbf{C}$ 的最小二乘问题

$$\min_{\mathbf{C}} \sum_{m=1}^M \|\mathbf{x}_m - \mathbf{w}_m \mathbf{C}\|^2 \quad (3)$$

$$\text{s. t.} \quad \|\mathbf{c}_g\| \leq 1, \forall g = 1, \dots, K$$

可以用拉格朗日对偶算法^[10]快速求解。固定 $\mathbf{C}$, 目标函数退化为单独对每一个 $\mathbf{w}_m$ 求最优解的线性回归问题

$$\min_{\mathbf{w}_m} \|\mathbf{x}_m - \mathbf{w}_m \mathbf{C}\|^2 + \beta |\mathbf{w}_m| \quad (4)$$

可以用特征符号搜索算法^[10]求解。

实验中 $D = 128$, $\beta = 0.15$, K 选用 8、32、128、512、1 024 共 5 种编码维度。 M 取决于图像大小。以一幅 256×256 像素的图像为例, SIFT 图像块大小定义为 16×16 像素,步长为 6,则横向作 $(256 - 16)/6 = 40$ 次匹配,纵向作 $(256 - 16)/6 = 40$ 次匹配, $M = 40 \times 40$, 即 1 600, 用 512 维 S-SIFT 算法处理 SIFT 特征子,最终输出的稀疏编码为 1 600 个 512 维的向量。

1.2 池化

池化是统计稀疏编码结果的过程,其模拟人眼视觉皮层的生理机制^[11],可以减少输入向量维数,有利于降低训练分类器的时间复杂度。以上文的 256×256 图像为例,其稀疏 SIFT 编码维度为 $1\,600 \times 512 = 819\,200$,训练一个输入向量维度超过 80 万的分类器难度很大,且容易出现过拟合。采用池化

方法,获取一幅图像的概要统计特征,不仅降低了训练分类器的难度,而且避免了过拟合现象。

目前常见的池化方法有平均池化和最大池化等,计算方法为

平均池化

$$\mathbf{p} = (\sum_{m=1}^M \mathbf{w}_m) / M$$

最大池化

$$p_j = \max\{|w_{1j}|, \dots, |w_{Mj}|\}$$

(5)

式中: $\mathbf{w}_m$ 是稀疏编码向量; $\mathbf{p}$ 是池化结果; w_{ij} 表示第 i 个稀疏编码向量的第 j 个元素。Lee 等证明了稀疏编码更适合用最大池化方法^[10],Boureau 等将 SIFT 特征、稀疏编码和最大池化相结合,取得了非常好的图像分类效果^[12]。池化后的特征用简单的线性 SVM 分类器就能达到较好的分类效果,时间复杂度仅为 $O(n)$ 。

2 车辆图像的提取和分类

2.1 目标提取

定义 t 时刻的一个像素点为 x^t ,如果 x^t 满足

$$p(x^t | B) > p(x^t | F)p(F)/p(B) \quad (6)$$

则该像素点属于背景,否则属于前景。式中: B 表示背景; F 表示前景。

选取一个时间段 T 内的图像序列,在 t 时刻训练集为 $\mathbf{x}_T = (x^t, \dots, x^{t-T})$ 。用 M 个高斯模型组成的高斯混合模型估计背景概率密度,用马氏距离计算新加入样本与当前背景的距离,距离较大则可能是前景,赋予较小的权重,反之则赋予较大的权重,不断更新均值和方差,选取 M 个高斯模型中对背景模型最重要的 B 个,可以得到

$$p(x^t | \mathbf{x}_T, B) \sim \sum_{m=1}^B \alpha_m N(x^t; \mu_m, \sigma_m^2 I)$$

(7)

式中: α_m 表示混合分布中第 m 个分布的权重; $\mu_1, \dots, \mu_m$ 表示均值的估计值; $\sigma_1^2, \dots, \sigma_m^2$ 表示方差的估计值。本文在 OpenCV 软件环境下实现了交通视频中的车辆目标提取,效果如图 1 所示。

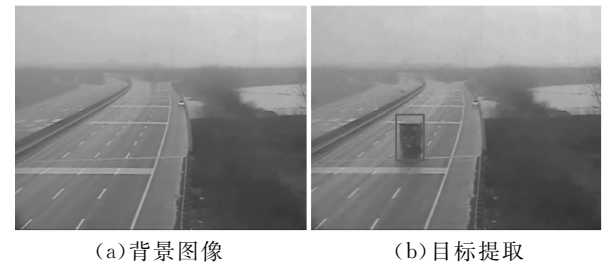


图 1 高斯混合模型的背景与目标提取

2.2 SVM 分类器参数训练

定义 $\mathcal{Q} = \{(\mathbf{x}_i, y_i)\}, i = 1, \dots, n$,其中 $\mathcal{Q}$ 是 n 个

输入数据点集; $\mathbf{x}_i$ 表示输入变量; y_i 表示目标值,在二类问题中 $y_i \in \{1, -1\}$ 。分类函数定义为

$$y = \mathbf{w}^T \varphi(\mathbf{x}) + b$$

(8)

式中: $\varphi(\mathbf{x})$ 表示从输入空间到高维特征空间的映射。根据序列最小优化算法 (sequential minimal optimization, SMO) 可以求得决策函数如下

$$f(\mathbf{x}) = \sum_{i=1}^n a_i y_i \kappa(\mathbf{x}_i, \mathbf{x}) + b$$

(9)

式中: a_i 表示拉格朗日乘子; $\kappa(\mathbf{x}_i, \mathbf{x})$ 表示核函数,用于快速计算映射到高维空间后两个向量的内积。常见的核函数有线性核、高斯核、多项式核。用非线性核 SVM 分类器,训练时间复杂度为 $O(n^2 \sim n^3)$,分类时间复杂度为 $O(n)$,用线性核则可以将训练时间复杂度降低到 $O(n)$,分类时间复杂度仍为 $O(n)$ 。实验中输入向量的维度最高达到了 1 024 维,采用线性核函数可以提高训练效率,保证系统实时性。

3 实验结果分析

实验使用江苏省 G36 高速公路监控系统的 H. 264 视频。车与非车图像特征差异较大,仅需少量训练集样本即可完成训练,而不同车图像特征差异较小,需要更多的训练样本。分别提取 scVideo_2c 和 scVideo_4c 两组数据集做训练和测试。scVideo_2c 数据集用于验证 S-SIFT 算法在不同场景下的车辆检测效果,scVideo_4c 数据集用于验证 S-SIFT 算法的车辆分类效果。具体而言,scVideo_2c 训练集包含 120 幅车辆图像和 120 幅非车图像,测试集包括车速较快场景、车辆遮挡较多场景和雨雪天气场景 3 组场景;scVideo_4c 训练集包含客车、轿车、卡车、面包车 4 类车型图像各 1 500 幅,测试集中 4 类车型对应数量为 1 020 辆、1 301 辆、1 221 辆和 958 辆。两组数据集示例如图 2 和图 3 所示。

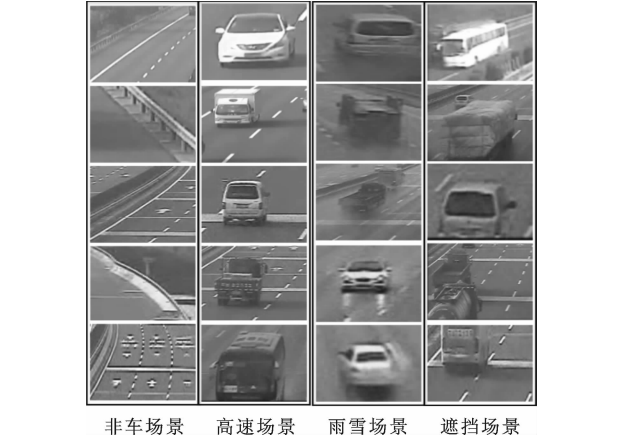


图 2 scVideo_2c 场景数据集

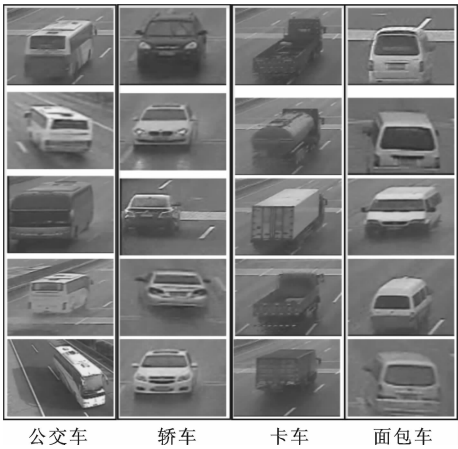


图 3 scVideo_4c 车辆数据集

首先用 16×16 像素的图像块对图像提取稠密 SIFT 特征, 步长设为 6。对 SIFT 特征中使用已训练的 1 024 个基向量进行稀疏编码, 基向量维度为 128 维, 稀疏系数 β 设为 0.15。分类器为线性核函数的 SVM。

软件环境为 OpenCV2.4、Matlab2013b, 硬件环境为 Intel Xeon E5-1603CPU, 16 GB 内存。实验对比了基于 SIFT 特征和 S-SIFT 特征两种方法的训练准确率, 实验结果如图 4~图 8 所示, 不失一般性, 图中训练准确率是采用 10 轮迭代平均结果。每次实验都从数据集中选取一部分做训练样本, 剩余部分做测试样本。

3.1 scVideo_2c 场景数据集

稀疏 SIFT 特征从所有 SIFT 特征集合中随机选取 7 200 个特征来训练生成 128 维基向量, 交替优化的最大次数为 50, 编码维度分别为 8、32、128、512、1 024 维。逐渐增加训练样本个数, 直到训练准确率趋向于收敛。不同维度 S-SIFT 和传统 SIFT 方法的训练准确率曲线结果如图 4 所示。

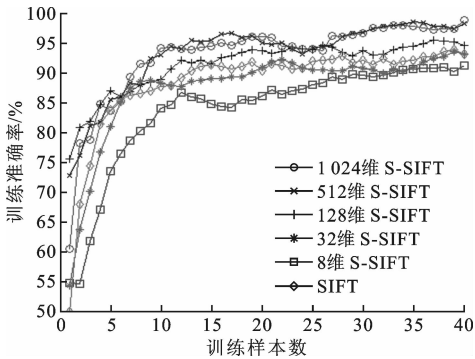


图 4 传统 SIFT 与不同编码维度 S-SIFT 算法对 scVideo_2c 场景数据集的训练准确率曲线

由图 4 可以看出, 当编码维度增加时, S-SIFT 训练准确率有明显提高, 维度为 512 维时, 训练准确率可达到 98% 以上; 对比 S-SIFT 方法和 SIFT 方法可以看到, 32 维 S-SIFT 方法与 SIFT 方法的准确率相近, 高维度 S-SIFT 方法的准确率明显优于 SIFT 方法。

用已训练的稀疏编码字典和 SVM 分类器对 3 组场景样本分别进行测试, 表 1 给出了不同方法对 3 组场景的分类准确率。可以看出, S-SIFT 方法分类准确率随编码维度增加不断提高。当编码维度在 512 维以上时, S-SIFT 方法对 3 种场景的分类准确率均可达到 96% 以上, 比低维 S-SIFT 方法至少提高 3.0%; 比原始 SIFT 方法提高 4.6%; 比背景建模方法提高 24.5%, 有效地去除了背景建模方法的误判图像; 比 HOG 方法提高 8.7%, 在干扰较多的雨雪场景和遮挡场景中, S-SIFT 方法明显优于 HOG 方法。

表 1 scVideo_2c 场景数据集的分类准确率

实验方法	分类准确率/%		
	高速场景	雨雪场景	遮挡场景
背景建模	75.50	62.13	71.35
HOG	91.30	76.50	79.85
SIFT	92.11	91.65	93.39
8 维 S-SIFT	91.84	89.36	90.60
32 维 S-SIFT	92.50	93.76	93.13
128 维 S-SIFT	95.75	93.25	94.50
512 维 S-SIFT	100.00	96.25	98.13
1 024 维 S-SIFT	100.00	97.50	98.75

表 2 给出了不同方法对 3 组场景的查全率, 可以看出, 背景建模方法的查全率最高, 在 3 种场景中

表 2 scVideo_2c 场景数据集的查全率

实验方法	查全率/%		
	高速场景	雨雪场景	遮挡场景
背景建模	99.43	99.06	98.57
HOG	91.34	89.75	88.42
SIFT	92.28	93.31	91.86
8 维 S-SIFT	91.51	94.69	92.29
32 维 S-SIFT	92.45	90.63	86.29
128 维 S-SIFT	94.73	89.06	91.00
512 维 S-SIFT	93.96	95.94	87.71
1 024 维 S-SIFT	94.90	93.13	93.43

均达到 98.5% 以上;S-SIFT 方法的查全率随维度增加呈上升趋势;HOG 方法和 SIFT 方法的查全率与低维度 S-SIFT 方法相近。

结合表 1 和表 2 可以看出,背景建模方法查全率虽然较高,但是对 3 种场景的分类准确率均低于 75.5%,存在较多误判;512 维以上 S-SIFT 方法在 3 种场景下准确率均可达到 96% 以上,查全率误差在 4.06%~12.29% 之间,两种指标均优于 HOG 方法和传统 SIFT 方法。

图 5 给出了 S-SIFT 方法和 SIFT 方法的训练时间曲线。可以看出,S-SIFT 方法训练时间随维度增加而增加,当训练样本为 40 个时,1 024 维 S-SIFT 的训练时间达到 1.02 s,平均每个样本训练时间 25.5 ms,8 维 S-SIFT 的训练时间最少,仅为 0.076 5s,平均每个样本训练时间 1.9 ms。虽然分类准确率随编码维度增加而提高,但训练所需时间成本也随之增加,因此不能无限增加编码维度来提高准确率。

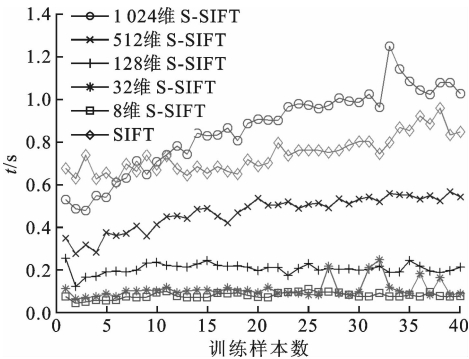


图 5 scVideo_2c 场景数据集的 SVM 分类器训练时间

结合图 4 和图 5 可以看出,SIFT 方法训练时间介于 1 024 维 S-SIFT 和 512 维 S-SIFT 之间,当 S-SIFT 方法的编码维度在 32 维至 512 维之间时,其在分类准确率和分类器训练时间两方面均优于 SIFT 方法。

3.2 scVideo_4c 车辆数据集

对 scVideo_4c 车辆数据集的 SIFT 特征进行原始采样,提取 150 000 个特征训练稀疏编码的 128 维基向量,交替优化 50 次,编码维度同样选取为 8、32、128、512、1 024 维。与 scVideo_2c 场景数据集类似,不断增加训练样本数,直到训练准确率趋向于收敛。图 6 给出了训练准确率的实验结果。

由图 6 可以看出,当编码维度为 128 维及以上时,S-SIFT 方法具有更高的准确率;对比几种不同编码维度的 S-SIFT 方法可见,训练准确率随编码维度增加而逐渐提高,当达到 1 024 维时,训练准确

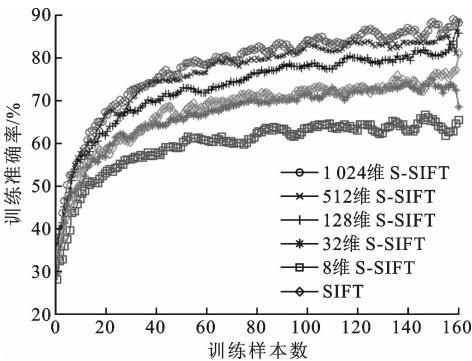


图 6 SIFT 算法与不同维度 S-SIFT 算法对 scVideo_4c 车辆数据集的训练准确率曲线

率达到 89% 以上。对比图 4 和图 6 可以看出,scVideo_4c 车辆数据集的训练准确率明显低于 scVideo_2c 场景数据集的训练准确率,原因是后者区分不同场景图像,两类图像间差异较大,而 scVideo_4c 车辆数据集区分不同车型,不同类别图像间特征差异相对较小,因此准确率有所下降。

用已训练的稀疏编码字典和 SVM 分类器对车辆样本进行分类测试,表 3 给出几种方法对不同车型的分类准确率。可以看出,SIFT 方法和 HOG 方法的分类性能与 32 维 S-SIFT 方法相近;512 维 S-SIFT 方法比 SIFT 方法准确率提高了 10.24%,比 HOG 方法提高了 10.86%;1 024 维 S-SIFT 方法比 SIFT 方法准确率提高了 13.27%,比 HOG 方法提高了 13.89%;背景建模方法没有车型分类的功能。

表 3 scVideo_4c 车辆数据集的分类准确率

实验方法	分类准确率/%			
	客车	轿车	卡车	面包车
背景建模				
HOG	70.81	79.33	72.66	79.65
SIFT	71.57	83.33	71.61	78.43
8 维 S-SIFT	60.00	66.67	62.11	72.32
32 维 S-SIFT	69.31	81.82	73.59	78.13
128 维 S-SIFT	75.43	88.67	78.33	82.05
512 维 S-SIFT	77.89	87.28	80.73	100
1 024 维 S-SIFT	80.32	95.25	89.96	92.49

表 4 给出了不同方法对 scVideo_4c 车辆数据集的查全率。可以看出,背景建模方法查全率最高,对不同车型的查全率均可达到 99.79% 以上;S-SIFT 方法查全率随编码维度增加而上升;HOG 和 SIFT 方法与 128 维 S-SIFT 方法查全率相近。结合表 3 和表 4 可以看出,背景建模方法查全率虽然最

高,但不具有车型分类的功能;HOG 和 SIFT 方法在准确率和查全率两方面均低于高维 S-SIFT 方法。

表 4 scVideo_4c 车辆数据集的查全率

实验方法	查全率/%			
	客车	轿车	卡车	面包车
背景建模	99.80	99.92	99.91	99.79
HOG	77.45	68.40	72.32	78.29
SIFT	81.67	73.02	81.08	70.98
8 维 S-SIFT	60.86	60.20	68.39	71.29
32 维 S-SIFT	75.49	68.02	69.04	74.95
128 维 S-SIFT	85.78	74.02	82.72	70.46
512 维 S-SIFT	92.25	72.25	82.15	74.63
1 024 维 S-SIFT	90.00	84.40	92.87	84.13

图 7 给出了 S-SIFT 和 SIFT 方法对 scVideo_4c 车辆数据集的分类器训练时间曲线。可以看出,高维 S-SIFT 方法的训练时间明显高于低维 S-SIFT 方法。SIFT 方法的训练时间与 1 024 维 S-SIFT 方法相近。当编码维度在 32 维和 1 024 维之间时,S-SIFT 方法在准确率和实时性方面均优于 SIFT 方法。

图 8 给出了 1 024 维 S-SIFT 方法对 scVideo_4c 车辆数据集的分类准确率混淆矩阵,图中第 i 行

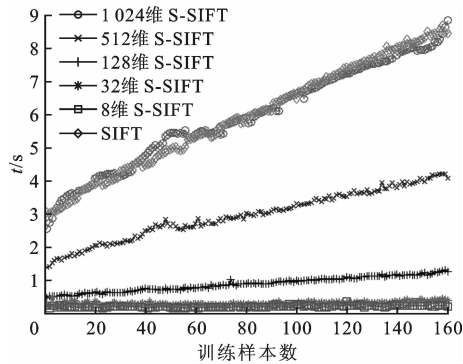


图 7 scVideo_4c 车辆数据集的 SVM 分类器训练时间

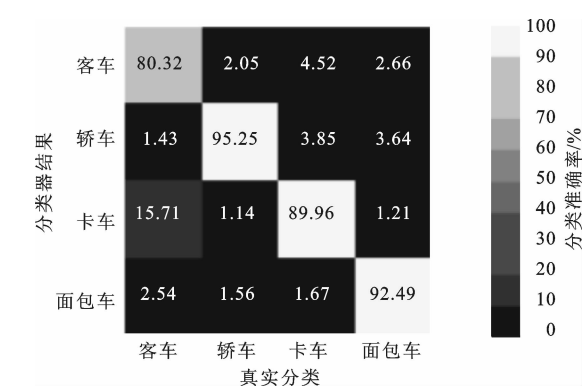


图 8 1 024 维 S-SIFT 对 scVideo_4c 数据集分类混淆矩阵

第 j 列数值表示第 j 类被误分成第 i 类的比率($i \neq j$)。对角线上数值代表对应类的分类准确率。从 0 到 100%分成 10 个灰度区间,颜色越深表示准确率越低。由图 8 可以看出,客车和卡车最容易发生混淆,因为这两类车型车身都较长,特征较为接近。

4 结 论

本文以深度学习理论为基础,提出了一种基于稀疏 SIFT 特征的车型识别方法,实现了快速、准确的交通监控视频车辆识别。算法用背景建模方法提取车辆目标,采集其 SIFT 特征作为图像低层特征,并对 SIFT 特征进行稀疏编码,得到更深层次的图像表征模型,用稀疏 SIFT 编码作为车辆特征训练线性 SVM 分类器,实现车型识别。实验结果表明,算法对低分辨率、视角变化、遮挡、雨雪天气等复杂场景下的车辆图像具有较高的识别率,准确率和训练时间均优于传统 SIFT 方法。

参考文献:

[1] ZHANG Zhaoxiang, TAN Tieniu, HUANG Kaiqi, et al. Three-dimensional deformable-model-based localization and recognition of road vehicles [J]. IEEE Transactions on Image Processing, 2012, 21(1): 1-13.

[2] 崔莹莹. 智能交通中的车型识别研究 [D]. 成都: 电子科技大学, 2013.

[3] WOOD R J, REED D, LEPANTO J, et al. Robust background modeling for enhancing object tracking in video [J]. Proceedings of the SPIE, 2014, 9089(2): 1-9.

[4] 黄毅, 陈湘军, 阮雅端, 等. 低清晰视频的“白化-稀疏特征”车型分类算法 [J]. 南京大学学报: 自然科学版, 2015, 51(2): 257-263.

HUANG Yi, CHEN Xiangjun, RUAN Yaduan, et al. The whitening-sparse coding vehicle classification algorithm for low resolution video [J]. Journal of Nanjing University: Science Edition, 2015, 51(2): 257-263.

[5] DALAL N, TRIGGS B. Histograms of oriented gradients for human detection [C]//Proceedings of the 2005 IEEE Computer Society Conference on Computer Vision and Pattern Recognition. Piscataway, NJ, USA: IEEE, 2005: 886-893.

[6] DONG Weisheng, LI Xin, ZHANG Lei, et al. Sparsity-based image denoising via dictionary learning and structural clustering [C]// Proceedings of the 2011 IEEE Conference on Computer Vision and Pattern Rec-

ognition. Piscataway, NJ, USA: IEEE, 2011: 457-464.

[7] MAIRAL J, BACH F, PONCE J, et al. Discriminative learned dictionaries for local image analysis [C]// Proceedings of the 2008 IEEE Conference on Computer Vision and Pattern Recognition. Piscataway, NJ, USA: IEEE, 2008: 1-8.

[8] 程东阳, 蒋兴浩, 孙钊锋. 基于稀疏编码和多核学习的图像分类算法 [J]. 上海交通大学学报, 2012, 46(11): 1789-1793.

CHENG Dongyang, JIANG Xinghao, SUN Tanfeng. Image classification using multiple kernel learning and sparse coding [J]. Journal of Shanghai Jiaotong University, 2012, 46(11): 1789-1793.

[9] YANG Jianchao, YU Kai, GONG Yihong, et al. Linear spatial pyramid matching using sparse coding for image classification [C]// Proceedings of the 2009 IEEE Conference on Computer Vision and Pattern Recognition. Piscataway, NJ, USA: IEEE, 2009: 1794-1801.

[10] LEE H, BATTLE A, RAINA R, et al. Efficient sparse coding algorithms [J]. Advances in Neural Information Processing Systems, 2006, 19(1): 801-808.

[11] SERRE T, WOLF L, POGGIO T. Object recognition with features inspired by visual cortex [C]// Proceedings of the 2005 IEEE Computer Society Conference on Computer Vision and Pattern Recognition. Piscataway, NJ, USA: IEEE, 2005: 994-1000.

[12] BOUREAU Y L, BACH F, LECUN Y, et al. Learning mid-level features for recognition [C]// Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition. Piscataway, NJ, USA: IEEE, 2010: 2559-2566.

(编辑 武红江)

(上接第 129 页)

the 2006 ASME Turbo Expo. New York, USA: ASME, 2006: 325-340.

[8] LIU Zhao, LI Jun, FENG Zhenping. Numerical study on the effect of jet slot height on flow and heat transfer of swirl cooling in leading edge model for gas turbine blade [C]// Proceedings of the 2011 ASME Turbo Expo. New York, USA: ASME, 2011: 1495-1504.

[9] JIANG Yuting, ZHENG Qun, YUE Guoqiang, et al. Numerical investigation of swirl cooling heat transfer enhancement on blade leading edge by adding water mist [C]// Proceedings of the 2014 ASME Turbo Expo. New York, USA: ASME, 2014: V05AT12A019.

(编辑 武红江 苗凌)

(上接第 82 页)

[11] YAO Guang, BI Jun, LI Yuliang, et al. On the capacitated controller placement problem in software defined networks [J]. IEEE Communications Letters, 2014, 18(8): 1339-1342.

[12] GAO Guang, BI Jun, GUO Luyi. On the cascading failures of multi-controllers in software defined networks [C]// Proceedings of the 21st IEEE International Conference on Network Protocols. Piscataway, NJ, USA: IEEE, 2013: 1-2.

[13] The Internet2 Community. Internet2 open science, scholarship and services exchange [EB/OL]. (2012-01-01) [2015-03-02]. <http://www.internet2.edu/network/ose/>.

[14] KNIGHT S, NGUYEN H X, FALKNER N, et al. The internet topology zoo [J]. IEEE Selected Areas in Communications, 2011, 29(9): 1765-1775.

[15] AYKUT ÖZSOY F, PINAR M C. An exact algorithm for the capacitated vertex p-center problem[J]. Computers and Operations Research, 2006, 33(5): 1420-1436.

[16] IBM. CPLEX-Optimizer [EB/OL](2012-01-01)[2005-03-02]. <http://www-01.ibm.com/software/integration/optimization/cplex-optimizer>.

(编辑 武红江)