

基于深度学习的视觉跟踪算法研究综述

贾静平 覃亦华

(华北电力大学控制与计算机工程学院 北京 102206)

摘 要 视觉跟踪是计算机视觉的重要研究领域之一。传统的视觉跟踪算法难以很好地解决复杂背景中的跟踪问题,如光线变化、目标发生较大的尺寸和姿态变化或目标被遮挡等。而深度学习的引入为视觉跟踪研究开辟了新的途径。但目前国内外基于深度学习的视觉跟踪研究文献相对较少,为吸引更多视觉跟踪领域研究者对深度学习进行探索 and 讨论,并推动视觉跟踪算法的研究,简要介绍了视觉跟踪和深度学习的研究现状,重点分析了基于深度学习的视觉跟踪算法的相关文献,讨论了各算法的优缺点,最后提出了进一步研究的方向以及对基于深度学习的视觉跟踪算法的展望。

关键词 计算机视觉,视觉跟踪,深度学习

中图法分类号 TP391 文献标识码 A

Survey on Visual Tracking Algorithms Based on Deep Learning Technologies

JIA Jing-ping QIN Yi-hua

(School of Control and Computer Engineering, North China Electric Power University, Beijing 102206, China)

Abstract Visual tracking is a fundamental subject in the field of computer vision. Classical tracking methods are not good at handling the problems of the complex background, such as illumination variation, great change of the target size and posture change greatly or occlusion and so on. Meanwhile, the introduction of deep learning technologies opens a new way of visual tracking study. There are a few research literature on visual tracking based on deep learning relatively, both in China and abroad at present. In order to attract more researchers in the field of visual tracking to explore and discuss deep learning, and to promote the research of visual tracking algorithm, this overview briefly reviewed the research status of visual tracking and deep learning. Then we focused on the related literatures about algorithms of visual tracking based on deep learning, and discussed their advantages and disadvantages. Finally, we proposed the direction of further research and the prospect of visual tracking algorithm based on deep learning.

Keywords Computer vision, Visual tracking, Deep learning

作为计算机视觉^[1]领域的重要研究方向,视觉跟踪一直得到国内外众多学者的密切关注。视觉跟踪也称为目标跟踪,是指对图像(视频)序列中的目标进行检测、提取、识别和跟踪,获得目标的运动参数,如位置、速度、加速度和运动轨迹等,并对其进一步处理与分析,实现对目标的行为理解,完成更高级的任务^[2]。近年来,在众多学者的努力下,重要的国际期刊和顶级会议涌现出大量有关视觉跟踪的学术论文。一些先进的思想、算法和技术陆续被提出,解决了视觉跟踪上的部分难题,并取得了一定的研究成果。但是,在较复杂的情形(如目标被长时间遮挡、光线剧烈变化、背景混杂、较大的尺寸和姿态变化等)下,视觉跟踪面临巨大的挑战,无法满足鲁棒性、实时性和准确性的要求。因此,为了克服这一挑战和难题,许多学者做了进一步的尝试。将深度学习引入视觉跟踪便是尝试的重要成果。

深度学习是机器学习领域中较新的研究领域,由 Geoffrey Hinton 于 2006 年提出^[3],之后在诸多领域取得了巨大的成功,但其在视觉跟踪领域的应用却是在 2010 年之后。本文旨在向读者介绍深度学习在视觉跟踪中的最新动态,希望

吸引更多的学者进行深入的研究和讨论。本文将首先介绍目前比较流行的传统的视觉跟踪算法并将其归类,指出其不足和优点;接着阐述深度学习的原理和主要框架,分析其能够有效应用于视觉跟踪的原因;然后重点研究自 2010 年以来出现的基于深度学习的视觉跟踪算法,总结各种算法的思想特点以及各自的优缺点;最后针对所做的分析,指出目前仍需解决的问题和未来的研究方向,并对深度学习在视觉跟踪上的应用予以展望。

1 传统的视觉跟踪算法

从目标表征模型的角度,视觉跟踪算法可大致分为生成式和判别式。生成式算法将跟踪视为一个目标匹配的过程:生成一个目标的模板,寻找匹配程度最高或重构误差最小的(即后验概率最大的)位置作为目标位置的预测。目前较流行的生成式算法主要有核方法^[4]、稀疏表达^[5-6]、密度估计^[7]、增量学习^[8]、高斯混合模型^[9]等,部分算法通过结合卡尔曼滤波、粒子滤波^[10]达到了较好的跟踪效果。文献[4]采用各向同性核函数的均值漂移跟踪算法,融合了特征直方图、相似性

本文受中央高校基本科研业务费专项资金(2016MS33),国家自然科学基金项目(61372184),北京市自然科学基金项目(4162056)资助。

贾静平(1978—),男,博士,讲师,主要研究方向为计算机视觉、机器学习;覃亦华(1989—),男,硕士生,主要研究方向为机器学习、计算机视觉, E-mail: 758825908@qq.com。

度和卡尔曼滤波等方法,取得了一定的效果。文献[6]对目标的历史和局部信息进行稀疏编码,采用空间金字塔和均值传递(均值漂移拓展)方法使算法具有较好的鲁棒性。文献[7]基于均值漂移模式查找算法提出一种新的核密度估计技术,即随着时间的推移,密度模式依次传播。实验证明将序列核密度估计运用于目标外观建模时具有一定的准确性和紧凑性。文献[8]采用增量的方式来学习和适应低维特征空间的表示,从而反映出目标外观的变化,并结合马尔卡夫-蒙特卡洛架构和粒子滤波方法来完成最终的目标跟踪。

判别式算法将目标跟踪转化为一个二分类问题,通过训练分类器,使其能够在图像(视频)序列中区分目标和背景,从而实现跟踪。近年来常用的判别式算法包括支持向量机(SVM)^[11-12]、随机森林^[13]、多样本学习^[14]、Boosting^[15-16]等。文献[13]在简单模板和均值漂移光流的基础上,引入随机森林算法,使用学习到的高置信度的特征来更新模板,以提高跟踪算法的自适应性。文献[14]将多个样本进行分类学习,以确定目标表征模型,从而解决了样本歧义问题。文献[16]提出了在线半监督学习强分类器,对目标特征进行选择,在一定程度上缓解了跟踪漂移和遮挡问题。

但是,上述算法普遍存在着局限性,即大多依赖于低级的手工设计的特征,而无法提取高级的语义信息。手工设计特征需要专家知识,对特定场景的适应度往往很好,但其泛化能力较弱,无法扩展到更普遍的情况。为了弥补这一缺陷,学者们引入了深度学习。深度模型具有强大的学习能力、高效的特征表达能力和获取高级语义信息的能力,这为视觉跟踪带来了新的研究思路。

2 深度学习

对于深度学习的定义,不同人有不同的描述,比较贴切的是 Bengio 的定义^[17]:深度学习是机器学习的子领域,是试图通过一系列多层的非线性变换对数据进行抽象的算法。目前的深度学习基于神经网络结构,其深度的称法是相对之前的浅层学习而言的。传统的机器学习方法(如支持向量机、Boosting、最近邻分类器等)可以用含有一个或两个隐层的神经网络来模拟,属于浅层学习结构。而深度学习的模型结构通常具有 3 个甚至更多的隐层,并能够利用大数据来学习特征。现有的深度学习架构主要有 3 种,分别为深度卷积神经网络(CNN)、深度置信网络(DBN)、堆栈自编码器(SAE)。

2.1 深度卷积神经网络(CNN)

卷积神经网络并不是新兴的算法,早在 20 世纪 80 年代, Fukushima 提出了卷积神经网络计算模型^[18],之后 LeCun 等人在此基础上改进并训练了卷积神经网络,将其应用于手写体字符识别任务^[19-20],取得了较好的效果。但受限于当时训练方法、计算机硬件等,直到 Hinton 提出深度学习的思想,一度沉寂的卷积神经网络才以深层结构重现。卷积神经网络一般由卷积层和子采样层构成,每层包含多个二维平面,每个平面又含有若干个独立神经元。一个多层的深度卷积网络结构如图 1 所示。

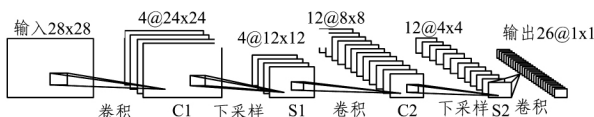


图 1 卷积神经网络

一般地,C 层为卷积层,其中的每个神经元与其上一层的局部感受野相连,提取局部特征。同一层的神经元组成了输入的一种特征映射,而不同的特征映射由不同的卷积核对输入进行卷积得到,图 1 中的 C1 层就有 4 种卷积核。S 层为降采样层,对卷积层的特征平面进行采样,降低其空间分辨率,保留主要信息,从而减少网络的权值参数。卷积层和降采样层交替出现,构成了多层卷积神经网络。其开创性的部分为:局部感受野和权值共享。每一个特征映射面的神经元具有相同的权值,降低了网络模型的复杂度,减少了网络的权值参数。这种深层网络结构对平移、比例缩放、倾斜或者其他形式的变形具有高度不变性,因此能够提取对平移、缩放和旋转不变的观测数据的显著特征,已广泛应用于语音识别、图像识别等众多领域。

2.2 深度置信网络(DBN)

2006 年, Hinton 提出的深度结构即为深度置信网络(Deep Belief Networks),它是一个含有多隐层的深度神经网络模型,由多个限制波尔兹曼机(RBM)层叠而成。RBM 为基于能量的概率生成模型,其本质是使得学习到的模型产生符合条件的样本概率最大。其结构是一个两层循环神经网络,包括一个可视层和一个隐藏层,各层内的神经元之间互不连接,只有可视层和隐藏层之间有无方向的对称性连接。所有的神经元节点都是随机二值变量节点(只取 0 或 1 值),且同层间节点互相独立。这些限制使得训练 RBM 更加高效,并使其成为深度置信网络(DBN)的结构单元。DBM 和 RBM 的结构分别如图 2 和图 3 所示。

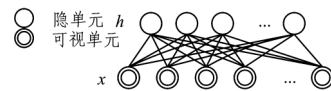


图 2 RBM

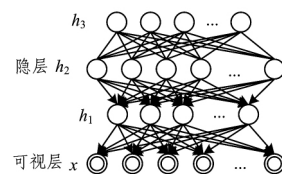


图 3 DBN

由简单的 RBM 构建成具有深度结构的 DBN, Hinton 等人提出了贪心逐层无监督训练方法,很好地解决了深层结构优化权值的困难。即 DBN 底层的输出作为上一层的输入,每一层都是一个 RBM。训练好一层后,保持所得参数不变,用其输出作为下一层的输入,开始训练。依次类推,直到训练整个网络。最后,采用有监督的训练方法对各层的参数进行微调,使 DBN 具有最优化的权值。整个训练过程分两步:1)分层次训练网络,每次只训练一个层次;2)参数调优,即监督微调过程。DBN 无论在时间还是算法效率上都有很好的效果,且用连续层的二进制或真值的变量来学习高层表示的分布,因此具有很好的灵活性。目前,DBN 已被成功应用于不同的领域中,如文本表示、音频事件分类以及可视化数据分析等。

2.3 堆栈自编码器(SAE)

堆栈自编码器和深度置信网络的结构相似,由若干个结构单元堆栈而成,不同之处在于堆栈自编码器的结构单元是自编码器而不是 RBM。简单的自编码器由输入层 L1、隐层 L2 和输出层 L3 构成,其中 L2 层也称为编码层,L3 层称为解码层。即 L2 层对网络的输入进行编码,L3 层再对 L2 层的输

出进行解码,因此,编码器的输出等于其输入。当 L2 层的隐藏神经元个数小于输入输出的个数时,自编码器被迫去学习输入数据的压缩表示,并可以从低维数据重构高维的原始数据。一般地,由多个自编码器堆叠成的堆栈自编码器具有深层结构,如图 4 所示。

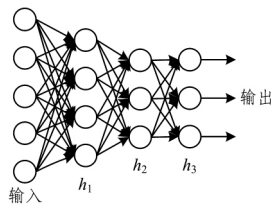


图 4 堆栈自编码器

堆栈自编码器的训练方法与深度置信网络类似,采用逐层训练的方法,最后得到的堆栈自编码器与其他深度神经网络一样,具有强大的表达能力。它能够学到输入的多层次表达和“部分-整体分解”结构,第一层能够提取原始输入的一阶特征(比如图片边缘),第二层可以获得二阶特征(比如物体轮廓),更高层会学到更高阶的特征(比如识别人的眼睛)。因此,堆栈自编码器能够进行快速文件编码、语音特征编码以及基于内容的图像检索。

3 深度学习在视觉跟踪上的应用

3.1 应用现状

深度学习因其深层结构而可以从大量数据中主动学习特征,从而避免了传统方法手工设计提取特征的缺点,并成功运用于图像分类^[21]、人脸识别^[22]、物体检测^[23]等领域。视觉跟踪的基础是提取目标特征,继而确定其在图像(视频帧)中的位置,完成跟踪任务。一个好的特征表达能够提高系统的整体性能,深度学习恰恰满足了这一要求。然而,基于深度学习的视觉跟踪算法并不像分类、识别和检测那样容易成功。在跟踪时,只对初始帧进行目标标注,在后续的视频序列中根据学到的特征进行目标定位,由于目标的变形、场景变化、遮挡等因素容易产生跟踪漂移,导致目标丢失。即使如此,有关学者依然相继提出了基于深度学习的视觉跟踪算法,在跟踪精度和成功率上都取得了比传统算法更好的效果。

目前,基于深度学习的视觉跟踪算法的文献并不多。从采用的深度模型来看,其可分为基于 CNN 的跟踪算法^[24-30]和基于 SAE 的跟踪算法^[31-32],其中部分文献综合了粒子滤波、支持向量机(SVM)和 AdaBoost 等传统算法。各文献所采用的算法分类如表 1 所列。

表 1 基于深度学习的视觉跟踪算法分类

深度传统	SVM	Ada Boost	粒子滤波	相关滤波	N/A
CNN	文献[25]		文献[30]	文献[27]	文献[24,26,28-29]
SAE		文献[31]	文献[31-32]		

表 1 中空出的项表示目前文献中暂时没有采用该综合算法。其中,文献[24,26,28-29]只采用了 CNN 模型,没有结合传统算法。在图像和视觉领域,CNN 是应用较成功的一个深度模型,如著名的深度网络 AlexNet^[21],RCNN^[33],Deep2D-Net^[34],GoogleNet^[35]和 VGG-Net^[36]都采用的是 CNN。部分文献采用了已训练成熟的深度网络,如表 2 所列。

表 2 基于 CNN 不同架构的视觉跟踪算法

	2010	2014	2015	2016
CNN-DIY	文献[24]	文献[26]		文献[30]
RCNN			文献[25]	
VGG-Net			文献[27-29]	

表 2 中,2011 年—2013 年没有出现基于 CNN 的视觉跟踪论文,2013 年文献[31]提出了基于 SAE 的视觉跟踪算法。从中可以看出,大部分文献都是基于已训练成熟的深度网络提出相应的视觉跟踪算法。

3.2 基于 CNN 的视觉跟踪算法

从表 1 可以看出,当前视觉跟踪算法中所采用的深度模型大多为 CNN,这与 CNN 的研究发展和自身结构特点息息相关。因为 CNN 采用局部感受野和权值共享,所以其具有平移不变性、光照不变性以及遮挡的鲁棒性等重要特征。

3.2.1 单纯 CNN

基于单纯 CNN 的跟踪算法没有结合传统算法而完成了跟踪任务。CNN 深层网络结构在提取目标特征的同时,可以将目标和背景分类,从而从背景中判别出目标。文献[24,26,28-29]只采用了 CNN 模型,通过不同的深度结构、选择策略、训练方法和模型更新等,提出了具有不同特点的跟踪算法。

文献[24]基于 3 个卷积层和若干降采样层的 CNN 对当前帧和上一帧进行采样,获取目标和背景的空间、时间特征。在首个降采样层后,将输出的特征图分别送入两个通道,一个提取全局信息,另一个提取局部信息,共同构成目标的概率图。文中采用了两个采样对,分别输入两个 CNN 网络,得到 4 张关键点的概率图,从而提高了跟踪的精确度。该算法于 2010 年被提出,具有一定的开创性,但其只给出了以人作为跟踪目标的示例。

文献[26]采用的 CNN 虽然只包含了两个卷积层、两个降采样层以及最后的全连接层,但其输入为目标图像的 4 个通道的信息,分别进入各自的网络,即共有 4 个并列的 CNN,并在全连接层进行综合,得到输出向量。最后根据该向量来判别样本是目标还是背景,从而估计目标位置。在训练网络时,加入截断损失函数来保持尽可能多的训练样本并降低跟踪误差累积的风险。采用迭代随机梯度下降,引入时间序列选择机制,以确保跟踪目标时不产生漂移。该算法对目标和背景信息进行较详细的提取,具有很好的判别能力,但因缺乏目标全局的高级语义信息,在目标发生急速运动变化、长时间遮挡时无法达到较好的鲁棒性。

文献[28]所用的深度模型分为共享层和特定层,其中共享层采用 VGG-Net,截取 3 个卷积层和 2 个全连接层;特定层由若干域组成,包含了目标正样本和负样本。为了估计当前帧的目标状态,在上一帧目标位置的周围取 N 个候选目标区域;分别计算出每个候选区域属于正样本和负样本的得分,将属于正样本的最高得分的候选区域作为当前帧的目标区域,用邻接矩形回归模型调整目标状态;同时使用长短时间更新策略,使跟踪趋于鲁棒性和适应性。

文献[29]也将 VGG-Net 网络应用到所提算法中,另外添加了一般性网络(GNet)和特殊性网络(SNet),两者具有相同的结构,都含有 2 个卷积层。将 VGG 的 Conv4-3(第 10 卷积层)的输出作为 SNet 的输入,Conv5-3(第 13 卷积层)的输出作为 GNet 的输入。文中通过分析大量实验发现,Conv4-3 的特征图保存了更多中间层次的信息,能够更加精确地将属于同一类别的不同图像区分开来;Conv5-3 的特征图保存了更

多高级语义信息,能够将人脸和非人脸区分开来。GNet 和 SNet 分别生成各自的前景热图,而最终目标定位是通过一个干扰项检测机制来决定使用哪一个热图。该文的亮点在于没有将 VGG 看成一个黑箱,而是分析不同层提取的信息特点,进而利用了深度网络的一个重要优点:不同层编码不同类型的特征,底层编码更适合有区分性的细节特征,高层更适合提取高级语义的目标类别特征。同时,在 GNet 和 SNet 前,通过选择网络(sel-Net)对输入的特征图进行选择,除去了许多不相干的噪声特征图。这些使得该算法具有较好的判别力和鲁棒性。

3.2.2 混合 CNN

混合 CNN,即在 CNN 深度模型的基础上使用了传统的视觉跟踪算法,如粒子滤波、SVM 等。过深的网络结构不利于参数训练,且对定位精度处理得不好,不能精确地跟踪目标。由于随着网络的加深,空间信息会被稀释,因此文献[25, 27, 30]尝试在 CNN 的架构上引入传统经典算法来进一步提高算法的性能。从实验结果看,他们都取得了不错的效果。

文献[25]采用 RCNN 深度模型,并添加一个分类层 SVM,它能够利用从 CNN 模型学到的特定目标的特征来从背景中区分出跟踪对象。通过 SVM 选出目标相关的正样本,并将其包含的目标特征(由 SVM 权重系数确定)沿着 CNN 模型反向传播到输入层,从而得到每一个正样本的显著图。这些图汇集成目标的显著图,最后使用它来构建观测模型,并采用贝叶斯滤波器进行跟踪。该文综合了判别式和生成式算法,对目标定位的精确度和鲁棒性进行了较大的改进。同时,显著图可以有效地可视化目标的空间结构,所以在一定程度上提高了目标定位的精度,并且能够实现像素级的目标分割。

文献[27]基于 VGG-Net 深度模型,将输入的目标图片按不同的层次提取特征,即把网络中 Conv3-4, Conv4-4, Conv5-4 层输出的特征图作为候选,用来估计目标的位置。论文指出,不同层次的输出携带着不同的信息:越靠前的层含有越多的空间细节信息,越往后的层包含越多的语义信息;分别采用相关滤波器来处理 3 个层次(通道)的输出,计算得到各自的相关响应图;最后使用粗细转换估计策略,即用最后一层(Conv5-4)相关响应图的最大值位置去逐层搜索最前一层(Conv3-4)的最大值位置,从而以最佳空间分辨率估计目标位置。此算法利用多层次的特征输出,综合相关滤波器,得到了较精确的目标位置,并对目标外观变化具有鲁棒性。

文献[30]在初始帧对目标域进行采样,经变形和归一化处理后,采用 K-means 算法提取得到一系列目标域的局部样本,将其作为固定的滤波器集。跟踪时,得到当前帧的采样结果,同时对上一帧中预测目标位置周围的背景进行采样。之前得到的滤波器集分别对处理后的背景样本做差,用差值分别对当前帧的样片进行卷积,提取局部选择特征,将这些特征组合在一起形成目标的全局外观表达。该算法采用粒子滤波方法,利用 CNN 得到的全局特征图对当前目标位置进行预测。提出的滤波器集可以保留目标的每一局部细节,从而避免了跟踪漂移,但同时增加了大量的样片预处理计算。

3.3 基于 SAE 的视觉跟踪算法

从表 1 中可以看出,当前采用 SAE 深度模型的跟踪算法不是很多,且都与传统的经典算法进行综合使用。这主要是因为 SAE 通过隐层来学习一个数据的表示或对原始数据进

行编码,采用一种不利用分类标签的非线性特征提取方法,目的在于保留和获取更好、更有效的信息表示,而不是对信息分类。视觉跟踪本身就需要将目标从背景中区分开来,因此,目标跟踪并不是 SAE 的强项。但文献[31-32]尝试利用 SAE 的高效的信息表示能力,综合传统方法提出自己的算法,取得了比传统跟踪算法更好的效果。

文献[31]采用 SAE 深度模型,其包含了 4 个隐藏层和额外的分类层,为了增强其对输入数据编码的鲁棒性,在输入中引入高斯噪声干扰,此时 SAE 称为堆栈降噪自编码器(SDAE)。第一层使用过完备滤波器,使其能更好地捕获图像结构。预训练 SDAE 的第一层时,为了加快训练速度,将输入图像分割成 5 张图片,然后对应训练 5 个 DAE,用这 5 个 DAE 的权值来初始化 SDAE 的第一层,其余各层用普通方法训练。跟踪时,采用粒子滤波器框架,SDAE 提取目标特征并将其作为粒子滤波器的状态变量,同时计算每个粒子的置信度,从而确定目标位置。该算法与传统方法相比使用了多种非线性变换,因此得到的图像表达比之前基于 PCA 的方法更具有表现性,且不需要像基于稀疏编码的方法那样进行解优化,提高了算法的实时性。

文献[32]也是基于 SDAE 深度架构,采用的隐层数与文献[31]相同,同时也在第一层使用了过完备滤波器。该算法的不同之处在于在每个隐层的顶部添加一个 sigmoid 分类层,即 DNN 分类器,这些分类器组合成具有强大分类能力的 AdaBoost。通过 AdaBoost,计算得到包含候选目标区域的置信图,其中,置信度最大的区域为目标的位置。该算法综合了粒子滤波方法,从实验结果分析,在不同的场景下,该算法的跟踪效果比文献[31]的更好。

3.4 总结说明

跟踪时,各算法具有不同的更新策略:一般采用适应更新和判别更新。适应更新指经过一定时间或帧数后根据当前目标状态对模型进行更新或每帧都更新,如文献[25, 27, 30]。判别更新即当检测到目标发生较大变化时,更新模型,否则不更新,如文献[26, 30-31]。也有文献同时采用两种更新策略,如文献[24, 32]。适应更新容易产生目标漂移,但实时性好;判别更新不易丢失目标,但跟踪精度和实时性不高。因此,不同的更新策略在不同的情形下有着不同的优势和缺陷。

基于深度学习的视觉跟踪算法虽然是近几年才被提出,但其跟踪效果已超过了传统算法。部分比较结果如表 3 所列。

表 3 传统跟踪算法与深度学习跟踪算法比较

	VTD	CXT	ASLA	SCM	DL
精确率	0.576	0.575	0.532	0.649	0.852
成功率	0.416	0.426	0.434	0.499	0.597

表 3 中,DL 为深度学习跟踪算法,其他为传统跟踪算法。显然,前者较后者具有更好的跟踪性能。该数据综合了多种复杂场景下的跟踪效果,具有一定的普适性。

结束语 本文研究了目前基于深度学习的视觉跟踪算法,对其进行了概述和总结,并发现:1)基于深度学习的视觉跟踪算法比传统的跟踪算法有着更好的鲁棒性;2)深度学习与传统跟踪算法的结合往往能获得意想不到的效果。

与图像分类和语音识别相比,深度学习在视觉跟踪领域的应用还未成熟,虽然目前取得了不错的效果,但还有很大的改进空间。我们认为基于深度学习的跟踪技术至少还有两个

问题需要解决:1)如何根据视觉跟踪的特点,利用领域知识研究并建立更适合于视觉跟踪任务的深度结构模型,在减少计算量的同时保证跟踪精度;2)视觉跟踪具有空间和时间的相关性,如何使用深度学习模型来获取这种相关性。

参考文献

- [1] 马颂德. 计算机视觉[M]. 科学出版社,1998.
- [2] 侯志强,韩崇昭. 视觉跟踪技术综述[J]. 自动化学报,2006,32(4):603-617.
- [3] HINTON G E, OSINDERO S, TEH Y W. A fast learning algorithm for deep belief nets [J]. Neural Computation, 2006, 18(7):1527-1554.
- [4] COMANICIU D, RAMESH V, MEER P. Kernel-Based Object Tracking [J]. IEEE Transactions on Pattern Analysis & Machine Intelligence, 2003, 25(5):564-575.
- [5] YANG M H. Visual tracking via adaptive structural local sparse appearance model[C]//Proceedings of the IEEE Conference on Computer Vision & Pattern Recognition, 2012:1822-1829.
- [6] ZHANG Z, WONG K H. Pyramid-Based Visual Tracking Using Sparsity Represented Mean Transform[C]//Proceedings of the Computer Vision and Pattern Recognition, 2014:1226-1233.
- [7] HAN B, COMANICIU D, ZHU Y, et al. Sequential Kernel Density Approximation and Its Application to Real-Time Visual Tracking [J]. IEEE Transactions on Pattern Analysis & Machine Intelligence, 2008, 30(7):1186-1197.
- [8] JING L. Incremental Learning for Robust Visual Tracking [J]. International Journal of Computer Vision, 2008, 77(1-3):125-141.
- [9] JEPSON A D, FLEET D J, EL-MARAGHI T R. Robust online appearance models for visual tracking[C]//Proceedings of the IEEE Computer Society Conference on Computer Vision & Pattern Recognition, 2003.
- [10] VARAS D, MARQUES F. Region-Based Particle Filter for Video Object Segmentation[C]//Proceedings of the 2014 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), 2014:3470-3477.
- [11] AVIDAN S. Support vector tracking [J]. IEEE Transactions on Pattern Analysis & Machine Intelligence, 2004, 26(8):1064-1072.
- [12] BAI Y. Robust tracking via weakly supervised ranking SVM[C]//Proceedings of the IEEE Conference on Computer Vision & Pattern Recognition, 2012:1854-1861.
- [13] SANTNER J, LEISTNER C, SAFFARI A, et al. PROST: Parallel robust online simple tracking[C]//CVPR 2010, 2010.
- [14] BABENKO B, YANG M H, BELONGIE S. Robust Object Tracking with Online Multiple Instance Learning [J]. IEEE Transactions on Pattern Analysis & Machine Intelligence, 2010, 33(8):1619-1632.
- [15] GRABNER H, BISCHOF H. On-line Boosting and Vision[C]//Proceedings of the IEEE Computer Society Conference on Computer Vision & Pattern Recognition, 2006:260-267.
- [16] GRABNER H, LEISTNER C, BISCHOF H. Semi-supervised On-Line Boosting for Robust Tracking[C]//Proceedings of the European Conference on Computer Vision, 2008:234-247.
- [17] BENGIO Y, COURVILLE A, VINCENT P. Representation Learning: A Review and New Perspectives [J]. IEEE Transactions on Pattern Analysis & Machine Intelligence, 2013, 35(8):1798-1828.
- [18] FUKUSHIMA K. NEOCOGNITRON: A self-organizing neural network model for a mechanism of pattern recognition unaffected by shift in position [J]. Biological Cybernetics, 1980, 36(36):193-202.
- [19] LECUN Y, BOTTOU L, BENGIO Y, et al. Gradient-Based Learning Applied to Document Recognition [C]//Proceedings of the IEEE, 1998:2278-2324.
- [20] LECUN Y, BOSER B, DENKER J S, et al. Backpropagation Applied to Handwritten Zip Code Recognition [J]. Neural Computation, 1989, 1(4):541-551.
- [21] KRIZHEVSKY A, SUTSKEVER I, HINTON G E. ImageNet Classification with Deep Convolutional Neural Networks [J]. Advances in Neural Information Processing Systems, 2012, 25(2):2012.
- [22] SUN Y, WANG X, TANG X. Deeply learned face representations are sparse, selective, and robust [J]. arXiv:1412.1265, 2014.
- [23] GIRSHICK R, DONAHUE J, DARRELL T, et al. Rich Feature Hierarchies for Accurate Object Detection and Semantic Segmentation [C]//CVPR, 2014:580-587.
- [24] FAN J, XU W, WU Y, et al. Human tracking using convolutional neural networks [J]. IEEE Transactions on Neural Networks, 2010, 21(10):1610-1623.
- [25] HONG S, YOU T, KWAK S, et al. Online Tracking by Learning Discriminative Saliency Map with Convolutional Neural Network [C]//Proceedings of The 32nd International Conference on Machine Learning, 2015:597-606.
- [26] LI H, LI Y, PORIKLI F. Robust Online Visual Tracking with a Single Convolutional Neural Network [M]. Springer International Publishing, 2014.
- [27] MA C, HUANG J, YANG X, et al. Hierarchical Convolutional Features for Visual Tracking[C]//Proceedings of the IEEE International Conference on Computer Vision, 2016:3074-3082.
- [28] NAM H, HAN B. Learning Multi-Domain Convolutional Neural Networks for Visual Tracking [J]. arXiv:1510.07945.
- [29] WANG L, OUYANG W, WANG X, et al. Visual Tracking with Fully Convolutional Networks[C]//Proceedings of the IEEE International Conference on Computer Vision, 2015.
- [30] ZHANG K, LIU Q, WU Y, et al. Robust Visual Tracking via Convolutional Networks Without Training [J]. IEEE Transactions on Image Processing A Publication of the IEEE Signal Processing Society, 2016, 25(4):1779-1792.
- [31] WANG N, YEUNG D Y. Learning a deep compact image representation for visual tracking [C]//Advances in Neural Information Processing Systems, 2013:809-817.
- [32] ZHOU X, XIE L, ZHANG P, et al. An ensemble of deep neural networks for object tracking[C]//Proceedings of the IEEE International Conference on Image Processing, 2014:843-847.
- [33] SIMONYAN K, ZISSERMAN A. Very Deep Convolutional Networks for Large-Scale Image Recognition [J]. arXiv:1409.1556.
- [34] OUYANG W, LUO P, ZENG X, et al. DeepID-Net: multi-stage and deformable deep convolutional neural networks for object detection [J]. arXiv:1409.3505.
- [35] SZEGEDY C, LIU W, JIA Y, et al. Going deeper with convolutions[C]//Proceedings of the Computer Vision and Pattern Recognition, 2015:1-9.
- [36] CHATFIELD K, SIMONYAN K, VEDALDI A, et al. Return of the Devil in the Details: Delving Deep into Convolutional Nets [J]. arXiv:1405.3531.