

深度学习与模型优化研究介绍

大连理工大学 刘日升 程世超 马龙 付陈平 樊鑫 罗钟铉

一、引言

2012 年,深度学习算法在 ImageNet 图像识别大赛中脱颖而出。随后深度学习在各个领域的应用出现井喷式增长。特别是 2016 年,基于深度学习开发的 AlphaGo 战胜了国际顶尖围棋高手李世石,深度学习的热度一时无两,其发展也进入突飞猛进的时代。相较于传统机器学习,深度学习可在大规模数据集上自动拟合出理想的数据分布,进而在目标明确的任务上取得成功。然而,深度学习是一个众所周知的黑箱,尽管表达能力很强,但可解释性不足。这意味着网络在使用过程中变量传播是不可控的,停止点无法为人所知。因此研究人员希望探究一套深度学习的支撑理论作为,使其使用过程可在人为控制下进行。

提及理论,模型优化是理论分析方面较为全面、系统的解决问题的手段。这类方法往往通过对问题的先验知识建立目标能量函数并使用各种优化策略对目标函数进行求解。然而,由于其依靠手工设计的先验知识,在数据先验分布较为复杂时无法充分发挥自身的优势。

深度学习与模型优化本质上都在试图刻画问题背后的规律,拟合输入、输出间的数据分布。在问题解决上,深度学习以很好自动进行数据拟合但缺乏理论分析,而模型优化有充分的理论分析但在数据复杂时难以拟合数据分布。因此,可探索它们之间的联系,将两者结合以弥补各自在理论分析和数据拟合上的缺陷。

二、深度学习与模型优化研究介绍

1. 相关工作

近年来,大多数学习和视觉任务可表述为下列正则化优化模型:

$$\min_x \Phi(x) := f(x) + g(x) \quad (1)$$

式中 f 为损失项, g 为正则化项, x 为模型输入。

不同于经典数值求解,它们纯粹基于数学推导设计迭代过程,而现阶段有许多研究试图通过大量数据训练网络以实现模型优化目标。现有研究工作大致可分为两类。一类方法的核心思想是抛弃公式(1)中正则化项即放弃使用先验知识,仅使用深度学习网络完成任务。文献[1-2]在每次迭代过程中直接使用网络结构取代先验相关的数值计算。最近,循环^[3]、展开^[4]和强化^[5]等学习策略均被引入网络训练中以用于模型优化。然而这类方法忽略使用正则项,导致先验知识难以运用到迭代过程中。此外其过于依赖网络结构,缺乏理论支撑,难以保证收敛性。另一类方法的核心思想是为特定的先验公式引入超参数,并展开公式更新过程获得优化方案。例如[6-7]对 L1 正则化项引入超参数并采用一阶方法对公式进行推导。这类方法主要基于特定的先验知识建立模型,因而无法应用到一般学习和视觉任务。此外,其引入的超参数太过简单难以拟合复杂数据分布。

以上工作就性能而言均取得了进展,但深度学习与模型优化理论仍处于割裂状态,我们希望将两者结合起来,以实现两者优势互补的目的。

2. 深度学习与模型优化研究工作介绍

为实现两者结合的目的,本团队主要从优化角度出发,建立模型优化和深度学习之间强有力的联系。具体而言,在已建立模型的基础上设计一个最优性模块,通过一阶最优性条件,给出一个判断准则来决定深度学习模块(即可学习的网络结构)的结果是否可以接受。若结果可接受,模型就采纳当前迭代中网络的结果,若不可接受,模型就放弃网络的结果,采用保守的近端梯度法给当前的解一些提示和推动,进而在下一次迭代中促使网络找到可靠的下降方向。下面将介绍团队基于这一思想所作的几项工作。

工作 1: 为弥合深度传播和模型优化间的差

距，本团队提出了基于传播和优化的深度模型 PODM (Propagation and Optimization based Deep Model, PODM)^[8]。PODM 主要包括两个模块：传播模块 P 和优化模块 O。传播模块是基于深度学习网络设计的，其要求具有足够的灵活性以在每次迭代中能寻找到被理论支持的下降方向。优化模块是基于最优化理论建立的模块，用于评判传播模块的迭代结果是否可接受，给与传播模块下次迭代的依据。然后，PODM 将传播模块和优化模块级联到一起实现整个迭代过程。虽然该方法加入了可学习的深度网络结构，但仍然可以证明在迭代过程中变量的传播序列可以收敛到对应模型的临界点。这意味着经过有限次迭代，该方法总是可以求解到一个用户想要的局部最优解的较小邻域范围内。图 1 说明了 PODM 迭代机制。

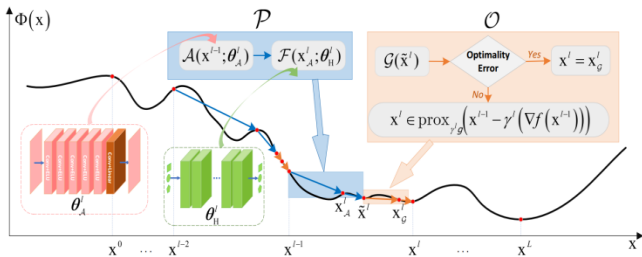


图 1 PODM 运行机制图

在此基础上，团队进一步设计了端到端的可学习网络。为方便训练，去掉了 PODM 的选择机制，直接将传播模块和近端梯度运算结合在一起进行联合训练。这样可有效避免训练数据和测试数据分布不一致时，判断准则带来的参数失效问题。部分实验结果见图 2，图 3。

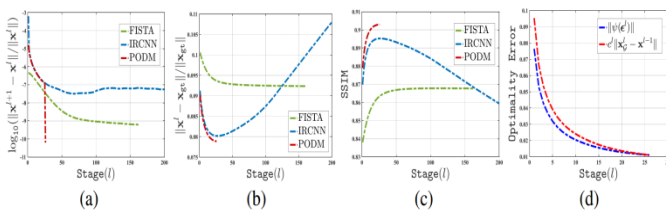


图 2 FISTA, IRCNN 和 PODM 收敛效果对比

图 2 展示了 PODM 与其他方法收敛效果的对比。首先，实验通过稀疏编码问题给出方法收敛性分析。其次，通过与传统方法 FISTA、可学习方法 IRCNN 做对比，给出相对误差、重构误差的结果对比。收敛结果显示由于引入网络，PODM 的迭代次数明显减少。此外，团队还给出了迭代过

程中网络被采纳的情况说明，事实上只要网络对问题而言是适合的，网络的结果均会被采纳。



图 3 真实图像上去模糊效果对比

图 3 展示了 PODM 与其他方法在图像恢复上的效果对比。将 PODM 与传统优化方法、深度学习方法相比较，会发现 PODM 结合了二者优势，在纹理细节和文字恢复上都取得了较好的效果。

工作 2：针对图像增强问题，本团队从互补传播角度提出了一种名为深度先验集成 DPE (Deep Prior Ensemble, DPE) 的框架^[9]。DPE 首先依据图像模型建立基本传播方案，然后引入残差网络帮助预测每个阶段的传播方向。DPE 可面向不同的图像增强任务集成与任务相应的知识和数据，是一种通用集成框架。DPE 从数据集中学习梯度下降方向，相较于传统基于优化的最大后验方法，这种方式不受局部最小值的影响因此更加稳健。此外，DPE 可充分利用图像模型邻域知识获得更多任务信息，而单纯深度网络只能依赖大规模的数据，无法使用额外的任务信息。

DPE 主要包括三部分：热启动、残差网络和先验投影。热启动用于刻画图像增强任务的先验知识，残差网络用于从数据集中学习图像的传播方向，基于误差反馈的先验投影则用于进一步保证传播方向是朝向理想输出方向的。最后，DPE 引入了一项次梯度约束项，将这三部分通过反馈控制策略级联起来。关于 DPE 详细说明可见图 4。

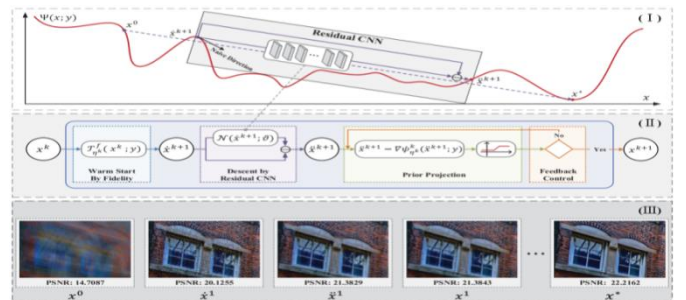


图 4 DPE 说明

最后，团队将 DPE 应用于各种图像增强任务

(例如图像反卷积, 插值, 超分辨率, 雾霾去除和 水下增强), 以检验其有效性。实验证明 DPE 确实优于现阶段主流的图像增强技术。部分实验结果见图 5。

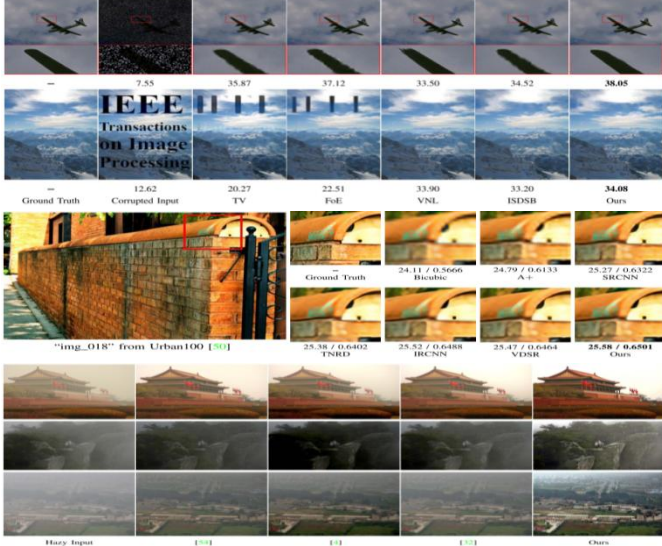


图 5 在不同图像增强任务上的效果对比

工作 3: 针对大多数深度模型缺乏理论支持, 先验信息加入会影响网络收敛的问题, 本团队提出一种通用的基于非凸优化展开的深度模型方法^[10]。该模型可以从全局收敛到原始优化模型的关键点。此外, 即使任务信息仅部分可用(例如没有先验正则化), 该模型依然可以保证收敛。

首先对能量模型 (2) 式的保真项引入 Moreau-Yosida 正则化项, 参数 μ^k (k 是模型迭代次数) 和辅助变量 u , 再引入辅助变量 v 和拉格朗日算子 $\tilde{\lambda}$, 从而获得新的正则化能量模型。

$$\min_x \{F(x): f(x, y) + r(x, y)\} \quad (2)$$

式中 x 和 y 分别是观察和预测的变量, f 是保真项, r 是先验项。公式 (3)–(5) 是重写之后的能量模型公式。

$$M_F^{\mu^k}(x^k) = \min_u \{f_{x^k}^{\mu^k}(u) + r(u)\} \quad (3)$$

式中 $f_{x^k}^{\mu^k}(u) := f(u, y) + \frac{\mu^k}{2} \|u - x^k\|^2$, 接下来, 令 $\lambda^k = \frac{1}{\rho^k} (\tilde{\lambda})^k$, ρ^k 是惩罚项。

$$u^{k+1} = \arg \min_u f_{x^k}^{\mu^k}(u) + \frac{\rho^k}{2} \|u - (v^k - \lambda^k)\|^2$$

$$v^{k+1} = \arg \min_v r(v) + \frac{\rho^k}{2} \|v - (u^{k+1} + \lambda^k)\|^2$$

$$\lambda^{k+1} = \lambda^k + (u^{k+1} - v^{k+1}) \quad (4)$$

接下来, 我们使用残差公式替换公式 (4), 将深度体系结构融合到上面的迭代中去, 得到

$$v_T = N_\alpha(v_0; W_T) := v_0 - \alpha (\sum_{t=0}^{T-1} G(v_t; W_t)) \quad (5)$$

式中 $W_T = \{W_t\}_{t=0}^{T-1}$ 是可学习参数集合, α 是步长。从优化角度而言, $G(v_t; W_t)$ 是基本网络单元, $(v_0; v_t)$ 分别是 N_α 的输入和输出。

最后, 为控制传播的每步迭代, 让 v^{k+1} 作为内置网络在第 k 次的输出, 并使用类似近端梯度的方式更新 x^{k+1} , x^{k+1} 的更新详情见公式 (6)。

$$x^{k+1} = \arg \min_x r(x) + \frac{1}{2} \|x - (v^{k+1} - \nabla f_{x^k}^{\mu^k}(v^{k+1}))\|^2$$

$$:= \text{prox}_r(v^{k+1} - \nabla f_{x^k}^{\mu^k}(v^{k+1})) \quad (6)$$

式中 prox_r 是 r 的近端算子。

团队将上述框架应用于具有单个变量的两个视觉应用(非盲反卷积和去雾)中, 以及具有两个变量的更复杂的视觉任务(低光图像增强)中, 以证明该深度模型在视觉应用领域中的有效性。部分实验结果可见图 6。

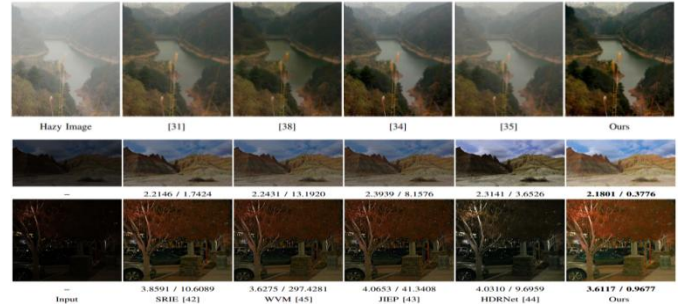


图 6 各方法在去雾、低光增强任务上实验效果对比。

三、总结

深度学习凭借大规模训练数据和五花八门的网络结构, 可以很好拟合数据分布进而在一些目标明确的任务上取得成功。但与此同时也存在着不容忽视的弊端, 其背后缺乏理论支持, 变量传播过程不可控, 此外, 与任务先验知识融合时难以保证收敛效果。模型优化是一套全面、系统的理论体系, 相较深度学习模型可充分利用任务先验知识进行数据分布拟合, 然而模型优化过于依赖手工设计先验, 因此在数据分布复杂时, 模型优化难以发挥自身优势。但两者均用于拟合数据分布且优势互补, 近年来出现很多试图将两者

结合起来的研究，介绍人团队在该面也做了较为全面、系统的探索，上述介绍的工作以及相关工作已发表于 NIPS2018^[8]、TIP2019^[9-10]。

对于深度学习与模型优化的结合，今后可从两个方向出发。首先，可以从深度学习网络向优化方向靠拢的角度出发以拉近二者间的距离。其次，可以从优化模型角度出发，建立特定的可学

习结构以实现二者的有效结合。将深度学习与模型优化相结合，并给予收敛效果相应的理论支持，不但有利于“打开”网络黑箱，且有利于相关任务先验知识的充分利用。因此，将两者结合是一个很有必要且充满前景的研究方向。

(责任编辑：王金甲)

参考文献

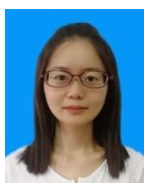
- [1] Steven Diamond, Vincent Sitzmann, Felix Heide, Gordon Wetzstein. Unrolled Optimization with Deep Priors. In CVPR, 2018.
- [2] Kai Zhang, Wangmeng Zuo, Shuhang Gu, and Lei Zhang. Learning deep cnn denoiser prior for image restoration. In CVPR, 2017.
- [3] Marcin Andrychowicz, Misha Denil, Sergio Gomez, Matthew W Hoffman, David Pfau, Tom Schaul, and Nando de Freitas. Learning to learn by gradient descent by gradient descent. In NIPS, 2016.
- [4] Risheng Liu, Xin Fan, Shichao Cheng, Xiangyu Wang, and Zhongxuan Luo. Proximal alternating direction network: A globally converged deep unrolling framework. In AAAI, 2018.
- [5] Ke Li and Jitendra Malik. Learning to optimize. In ICLR, 2017.
- [6] Karol Gregor and Yann LeCun. Learning fast approximations of sparse coding. In ICML, 2010.
- [7] Pablo Sprechmann, Alexander M Bronstein, and Guillermo Sapiro. Learning efficient sparse and low rank models. In TPAMI, 2015.
- [8] Risheng Liu, Shichao Cheng, Xiaokun Liu, Long Ma, Xin Fan, Zhongxuan Luo. A bridging framework for model optimization and deep propagation. In NIPS, 2018.
- [9] Risheng Liu, Long Ma, Yiyang Wang, Lei Zhang. Learning converged propagations with deep prior ensemble for image enhancement. In TIP, 2019.
- [10] Risheng Liu, Shichao Cheng, Long Ma, Xin Fan and Zhongxuan Luo. Deep Proximal Unrolling: Algorithmic Framework, Convergence Analysis and Applications. In TIP, 2019.



刘日升

大连理工大学副教授，主要研究方向为机器学习、深度学习、优化方法、计算

机视觉。Email: rslu@dlut.edu.cn



程世超

大连理工大学博士生，主要研究方向为计算机视觉、图像处理。

Email: shichao.cheng@outlook.com



马龙

大连理工大学博士生，主要研究方向为计算机视觉、优化方法。

Email: malone94319@gmail.com



付陈平

大连理工大学博士生，主要研究方向为计算机视觉、图像处理。

Email: fu.2019@outlook.com



樊鑫

大连理工大学教授，主要研究方向为计算机视觉、图像处理、医学影像分析。

Email: xin.fan@dlut.edu.cn



罗钟铨

大连理工大学教授，主要研究方向为计算几何与图形 / 图像、计算机视觉。

Email: zxluo@dlut.edu.cn

面向视频数据的对抗攻击研究

北京航空航天大学 韦星星

深度学习已经在很多领域取得了巨大成功，但是最新的研究成果表明深度学习非常容易受到对抗样本的干扰。所谓对抗样本是指在干净图像上添加一些肉眼不可见噪声，原本会被深度学习模型分类正确的图像就会被预测出错误结果。原本可被深度学习算法以 57.7% 置信度正确分类为“panda”的图片，在加上噪声后，却被深度学习算法以 99.3% 的置信度错误分类为了“gibbon”。

由于对抗样本的巨大危害，很多研究者从不同角度对此进行了研究，如不同对抗样本生成方法、各式各样对抗样本防御方法以及对抗样本产生内在机理等。但这些研究基本都集中在图像数据上面，面向视频数据的对抗攻击与防御并未涉及。相比图像数据，视频因为其特有时序信息，必然导致其与面向图像数据的对抗样本研究有巨大不同。下面介绍我们组两个相关工作。

首先是面向视频分类任务的对抗样本生成方法，图像数据由于其天然的时序结构信息，导致在生成对抗视频的时候不能对其每一帧按照图像的方法来进行添加噪声。相反，我们认为视频对抗样本应该具有传播性和稀疏性两个特点，稀疏性是指只需要在视频中的一些关键帧上添加噪声即可，传播性是指这些稀疏的对抗噪声可以利用帧间的时序关系来传播到未添加噪声的帧上面，从而污染整个视频片段，进而达到欺骗深度学习模型的目的。如图 1 所示，蓝实线是对视频中每帧添加噪声的度量，可以看到随着视频帧的变化，噪声幅度呈现显著下降趋势，最终趋于 0，但是却成功欺骗神经网络输出了错误的标签（见红色字）。该成果发表在 AAAI2019 上面。

第二个工作是使用对抗本来攻击图像和视频的物体检测算法。这个方法具有两个特点，

一是迁移性好，我们的对抗样本可以有效攻击 Faster-RCNN 等基于分类机理的物体检测算法，还可以攻击 SSD, YOLO 等基于回归机理的物体检测算法。第二个特点是高效，不同于以往的基于优化方式来获得对抗样本，我们采用生成式对抗网络，由于在测试过程只涉及前馈网络，因此其生成速度非常快，因此可以处理图像数据还可以快速处理视频每一帧。图 2 展示了该方法的结构框图，其中的高迁移性是通过一个特征损失函数来破坏物体检测器的特征图，由于 Faster-RCNN 和 SSD 都基于特征图来进行物体检测，因此将特征图破坏后将同时使得这两类物体检测器失效。

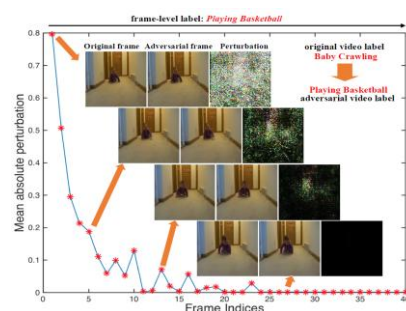


图 1 面向视频数据的可传播的稀疏对抗噪声

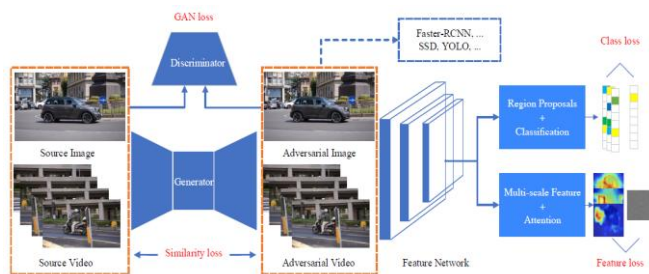


图 2 基于特征图篡改的高迁移性视频对抗样本生成方法

(责任编辑：苏航)



韦星星

北京航空航天大学计算机学院副教授，主要研究为人工智能安全，计算机视觉等。

Email: xxwei@buaa.edu.cn

行人再识别相关研究进展

中科院自动化所模式识别国家重点实验室 杨阳 吴锦林 雷震

行人再识别起源于多摄像机跟踪任务，解决不同摄像机下行人之间相互关联的问题。在 2006 年，研究人员把它作为一个特定的任务并以图像形式进行探讨和研究。随着时间推移，研究对象从单纯基于图像扩展到基于视频，以及跨模态的图像与文本和可见光与近红外。目前，基于图像的行人再识别仍是研究主流。它所面临的挑战包括：(1) 单个摄像机下行人的图片分辨率较低、遮挡以及背景干扰；(2) 不同摄像机下获取的图片存在光照变化、行人姿态各异以及摄像机视角不相同等问题。针对这些问题，早期研究人员主要通过图像特征的提取(颜色、纹理、属性语义、高阶混合、多层描述子)、特征学习(特征编码、特征降维、子空间学习)、相似度学习(距离度量、双线性度量、联合度量)和重排序。自深度学习的兴起，研究人员着手于设计端到端的网络，研究不同的损失函数(分类损失、三元组和四元组)，风格迁移，小样本学习，神经网络架构搜索，以及两幅图像的配准(姿态矫正)、分辨率低和背景干扰等问题。公开数据库(Market1501)的 Rank-1 指标一再被刷新，从最初 25% 上升到 96%。考虑到行人数据涉及隐私以及数据内容难以包含各种场景，有部分研究人员在近期通过游戏引擎，建立新的虚拟数据。一方面可以增加训练数据，辅助行人再识别性能的提升；另一方面通过限定训练和测试条件，研究不同因素对最终性能的影响，为实际场景的应用设计提供参考。

在实际应用中，有监督的行人再识别模型需要耗费大量的人力和物力对不同的应用场景重新采集、标注数据进行训练。因此无需新标注数据训练的无监督行人再识别算法，逐渐引起研究人员的关注。在 ECCV2018 和 PAMI2019 上，龚少刚团队基于现有跟踪算法和稀疏时间采样自动得到单个摄像机下的行人轨迹图像，对每个摄像机的所有行人轨迹随机打上伪标签(见图 1)，并

提出了多任务训练的框架和潜在轨迹关联的方法，通过学习潜在的正样本来提升模型的跨视角检索能力。但稍显不足之处在于它们仅仅在训练时输入的小批量数据中挖掘潜在正样本，效率较低，并且为了在输入的小批量数据中采样到潜在正样本，每一步训练需要采样比较多的数据输入到网络，导致训练时占用特别大的显存。在 ECCV2018 另一篇工作中，阮邦志团队提出一种行人轨迹渐进合并的方法，将潜在的属于同一个人轨迹的图像进行合并，使用自训练的方法训练行人再识别模型。但是，如果将不属于同一个人的轨迹合并会误导模型，且这种合并是不能撤回的，限制了无监督行人再识别的性能提升。

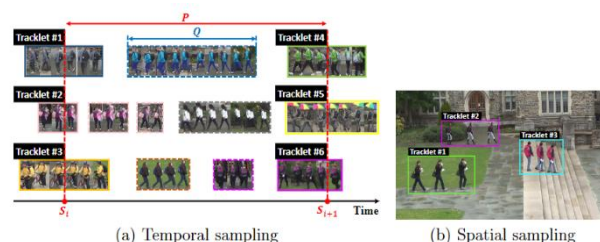


图 1 轨迹伪标签打印示意图

针对上述问题，我们提出了一种无监督图关联的方法和二阶段训练的策略(见图 2)。

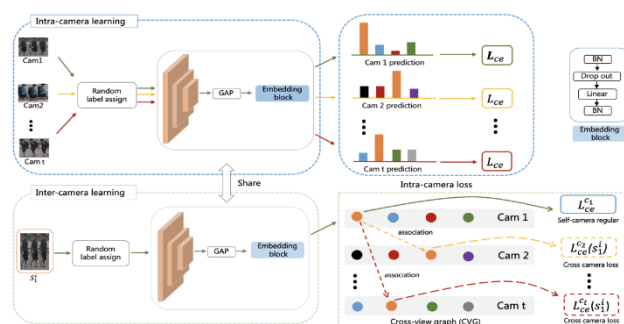


图 2 无监督图关联和二阶段训练策略

第一阶段采用单个摄像机中行人轨迹的伪标签和多任务训练的策略，使模型专注于学习单个视角内的行人图像的特征表示；第二阶段使用图关联的方法挖掘跨视角的潜在正样本，使模型专注于学习跨视角行人图像的特征表示。两个阶

段训练的策略将视角内训练和跨视角训练分开，减轻跨视角学习时噪声样本带来的干扰。图关联的方法主要是通过引入最近邻，跨视角和边对称的约束，建立一个较为精确的跨视角行人轨迹关联图。相比于已有工作，能够过滤掉较多的噪声样本对。对于跨视角行人轨迹关联图中挖掘到的潜在正样本对，我们提出了一种图加权的跨视角损失函数来提升模型的跨视角检索能力。

表 1 基于图像的无监督行人再识别算法对比

Datasets	Reference	PRID2011			iLIDS-VID			MARS			
Metric(%)		R1	R5	R20	R1	R5	R20	R1	R5	R20	mAP
SMP [24]	ICCV'17	80.9	95.6	99.4	41.7	66.3	80.7	23.9	35.8	44.9	10.5
DGM+MLAPG [41]	ICCV'17	73.5	92.6	99.0	37.1	61.3	82.0	24.6	42.6	57.2	11.8
DGM+IDE [41]	ICCV'17	56.4	81.3	96.4	36.2	62.8	82.7	36.8	54.0	68.5	21.3
DASy [11]	ECCV'18	43.0	-	-	56.5	-	-	-	-	-	-
GRDL [13]	ECCV'16	41.6	76.4	89.9	25.7	49.9	77.6	19.3	33.2	46.5	9.56
DTW [26]	PR'17	41.7	67.1	90.1	31.5	62.1	82.4	-	-	-	-
BUC [23]	AAAI'19	-	-	-	-	-	-	61.1	75.1	80.0	38.0
EUG(p=0.05) [39]	CVPR'18	-	-	-	-	-	-	62.7	74.9	82.6	42.5
RACE [40]	ECCV'18	50.6	79.4	91.8	19.3	39.3	68.7	43.2	57.1	67.6	24.5
TAUCL [17]	ECCV'18	49.4	78.7	98.9	26.7	51.3	82.0	43.8	59.9	72.8	29.1
UTAL [18]	TPAMI'19	54.7	83.1	96.2	35.1	59.0	83.8	49.9	66.4	77.8	35.2
UGA(ours)	This work	80.9	94.4	100	57.3	72.0	87.3	58.1	73.4	81.4	39.3

1-st and 2-nd best results are in red/blue respectively.

我们在 7 个公开的行人重识别数据库进行了实验，均取得较好的结果；并且和现有轨迹关联方法对比，我们方法在基于图像和基于视频的无监督行人再识别任务上取得了非常大的提升（见表 1 和 2）。工作发表于 ICCV 2019。

表 2 基于视频的无监督行人再识别算法对比

Dataset	Reference	Method	Market1501		DukeMTMC-ReID		CUHK03		MSMT17	
metric			mAP	Rank 1	mAP	Rank 1	mAP	Rank 1	mAP	Rank 1
HHL [49]	ECCV'18	GAN	31.4	62.2	27.2	46.9	-	-	-	-
SPGAN [4]	CVPR'18	GAN	22.8	51.5	22.3	41.1	-	-	-	-
SPGAN+LMP [4]	CVPR'18	GAN	26.7	57.7	26.2	46.4	-	-	-	-
TJ-AIDL [36]	CVPR'17	adaptation	26.5	58.2	23.0	44.3	-	-	-	-
BUC [23]	AAA'19	cluster	38.3	66.2	27.5	47.4	-	-	-	-
CAMEL [43]	ICCV'17	cluster	26.3	54.5	-	-	-	39.4	-	-
PUL [5]	ToMM'18	cluster	20.1	44.7	16.4	30.4	-	-	-	-
CDS [16]	ICME'19	cluster	39.9	71.6	42.7	67.2	-	-	-	-
TAUCL [17]	ECCV'18	tracklet	41.2	63.7	43.5	61.7	31.2	44.7	12.5	28.4
UTAL [18]	TPAMI'19	tracklet	46.2	69.2	44.6	62.3	42.3	56.3	13.1	31.4
ECN [50]	CVPR'19	adaptation	43.0	75.1	40.4	63.3	-	-	10.2	30.2
UGA(ours)	This work	tracklet	70.3	87.2	53.3	75.0	68.2	56.5	21.7	49.5

1-st and 2-nd best results are in red/blue respectively.

（责任编辑：任传贤）



杨阳

中国科学院自动化研究所助理研究员。主要研究领域

为计算机视觉，目前侧重于行人再识别、属性分析及人脸识别的相关应用研究。

Email: yang.yang@nlpr.ia.ac.cn



吴锦林

中国科学院自动化研究所在读博士。主要研究领域包括半监督学习算法，

自动机器学习。

Email: jinlin.wu@nlpr.ia.ac.cn



雷震

中国科学院自动化研究所研究员。主要研究领域为生物特征识别和智能视

频分析。

Email: zlei@nlpr.ia.ac.cn