

DOI:10.16136/j.joel.2022.11.0093

用于视网膜血管分割的半监督深度学习框架

吕佳^{1,2*}, 刘耀文¹

(1. 重庆师范大学 计算机与信息科学学院, 重庆 401331; 2. 重庆国家应用数学中心, 重庆 401331)

摘要:针对目前视网膜血管分割任务中伪标签质量参差不齐, 获得高质量的伪标签需要经过筛选的问题, 本文提出了一种新的用于视网膜血管分割的半监督深度学习框架。该框架采用分而治之的思想来处理数据, 针对有标签数据, 采用传统的深度学习方法; 针对无标签数据, 采用 Mean teacher 模型, 通过对比同一输入的不同形态输出, 让模型学习无标签数据之间的共同特征, 避免了采用伪标签技术带来的筛选过程。本文将 U 型网络 (u-neural networks, U-Net)、Dense-Net 和 Ladder-Net 3 个基准网络放入该框架, 在 DRIVE 和 CHASEDB1 数据集上进行训练测试, 均取得了较好的分割效果, 表明本文框架具有提高网络区分不同阈值像素的能力。

关键词: 视网膜血管分割; 半监督学习; U 型网络 (U-Net); Mean teacher 模型; 伪标签

中图分类号: TP183 **文献标识码:** A **文章编号:** 1005-0086(2022)11-1207-08

Semi-supervised deep learning framework for retinal vessel segmentation

LV Jia^{1,2*}, LIU Yaowen¹

(1. College of Computer and Information Sciences, Chongqing Normal University, Chongqing 401331, China; 2. National Center for Applied Mathematics in Chongqing, Chongqing 401331, China)

Abstract: In view of the problem that quality of pseudo-labels is uneven in the current retinal vessel segmentation task and it requires to be screened to obtain the high-quality pseudo-labels, a novel semi-supervised deep learning framework for retinal vessel segmentation is proposed in this paper. The framework adopts the idea of divide and conquer to process data. Traditional deep learning methods are utilized especially for the labeled data, while Mean teacher model is used to deal with the unlabeled data. By comparing the different morphological outputs of the same input, the model can learn the common features between the unlabeled data and avoid the screening process brought by pseudo-label technology. Three benchmark networks, u-neural networks (U-Net), Dense-Net and Ladder-Net are put into the framework, the experiments are carried out on DRIVE and CHASEDB1 datasets, which achieve good segmentation results. It shows that the framework can improve the ability of the network to distinguish different threshold pixels.

Key words: retinal vessel segmentation; semi-supervised learning; u-neural networks (U-Net); Mean teacher model; pseudo label

1 引言

视网膜血管是人体血液系统的重要组成部分, 与一些疾病的产生有着密切联系, 如糖尿病、高血压、青光眼等, 视网膜血管的分割及定位能够

给上述疾病的筛查和诊断提供重要依据^[1]。然而, 视网膜血管细小且分布复杂, 血管与背景像素区分度不大, 手动分割视网膜血管是一件非常耗时的任务, 因此研究计算机辅助视网膜血管分割具有重要意义。

* E-mail: lvjia@cqnu.edu.cn

收稿日期: 2022-02-20 修订日期: 2022-03-03

基金项目: 国家自然科学基金(11971084)、重庆市教委重点项目(KJZD-K202200511)、重庆市科技局技术预见与制度创新项目(2022TFII-OFX0265)和重庆师范大学研究生科研创新项目(YKC21043)资助项目

目前视网膜血管分割算法主要分为无监督算法和有监督算法。在无监督算法中,通常将血管的纹理、颜色、形状等特征作为分割依据^[2],包括血管跟踪、匹配滤波和基于形态学的算法。无监督算法尽管不需要专家标注,但需要大量的人工干预,且无监督的算法只能提取粗大的血管,并不能针对细小血管做出很好的分割,因此分割精度低,提取特征粗糙。有监督算法主要是基于深度学习的算法,包括全卷积神经网络(fully convolutional network, FCN)、U 型网络(u-neural network, U-Net)以及基于 U-Net 改进的算法^[2],这些算法能够很好地提取血管中的细小特征,分割精度高,但需要专家标注作为训练目标来引导算法提高分割性能,因此在临床医学中并不适用。

半监督学习是将有监督学习和无监督学习有效结合的一种学习方法,其通过使用少量有标签数据和大量无标签数据来训练模型,将深度学习方法与半监督学习方法相结合是目前研究的重要方向。LEE 等^[3]利用一个预训练的网络,自行生成无标签数据的伪标签,然后使用伪标签对模型进行微调,以获得更好的鲁棒性。XIE 等^[4]提出无监督数据增强(unsupervised data augmentation, UDA)一致性训练框架,原始的无标签数据和增强后的无标签数据经过网络后得到两个具有相同特征的预测,通过最小化两个预测之间的差距来学习无标签数据拥有的特征信息。SOHN 等^[5]提出了 Fixmatch,该算法对弱增强的无标签数据进行预测生成伪标签,并且筛选保留置信度高的伪标签,对于使用强增强的同一个无标签数据,则把生成的伪标签拿来当作目标。LI 等^[6]提出了基于 U-Net 的半监督视网膜血管分割模型。先使用增强的有标签数据训练 U-Net,再利用训练好的 U-Net 对无标签数据进行预测,将过滤后的预测作为伪标签,最后使用伪标签去更新训练集,该方法减少了模型对专家标注的依赖。LI 等^[7]提出基于不确定性感知半监督学习的分层神经网络血管分割模型,通过引入不确定性估计方案改进伪标签的生成,采用基于贝叶斯网络的不确定性感知架构对伪标签进行筛选,该方法在视网膜血管分割和肝脏血管分割上均取得了良好的性能。以上基于伪标签技术的半监督深度学习算法能够在一定程度上缓解模型对专家标注的依赖,但是生成的伪标签质量参差不齐,需要经过筛选后才能加入训练集。

本文提出一种新的用于视网膜血管分割的半监督学习框架,采用分而治之的思想,对无标签数据和有标签数据分别处理。本文的贡献如下:

1) 提出了一个用于视网膜血管分割的半监督学习框架,在 3 种不同结构类型的网络上均取得了较好的性能,提高了网络区分不同阈值像素的能力;

2) 针对无标签数据,引入 Mean teacher 模型^[8]用来学习其所蕴含的内在特征,由一致性原则,模型框架结构得以简化且避免了需筛选伪标签的工作。

2 基本原理

2.1 U-Net

U-Net 是一种采用编码器-解码器结构的基于卷积神经网络的网络模型。编码器用于提取图像的浅层信息,在每一层后添加池化层来缩小图像以减少参数;解码器能够将图像的深层信息提取出来,采用和编码器对称的结构,利用反卷积或线性插值来还原图像的尺寸。通过跳跃连接将浅层信息和深层信息结合起来,尽可能地保留图像的细节信息。

2.2 Mean teacher 模型

TARVAINEN 等^[8]提出 Mean teacher 模型用来解决 Temporal Ensembling 在处理大量数据时性能不佳的问题。该模型由两个具有相同结构的网络构成,一个为 student 模型,一个为 teacher 模型。对于同一个输入 x ,首先经过 student 模型得到预测 $S(x)$,然后经过 teacher 模型得到预测 $T(x)$ 。由 $S(x)$ 和 $T(x)$ 计算一致性损失, $S(x)$ 和标签数据 L 计算分类损失,两个损失构成以下整体损失:

$$\text{Loss}_{\text{total}} = \text{Loss}_1(S(x), T(x)) + \text{Loss}_2(S(x), L). \quad (1)$$

teacher 模型的参数通过 student 模型的指数滑动平均(exponential moving average, EMA)得到:

$$M_T = \alpha M_{T-1} + (1 - \alpha) M_S, \quad (2)$$

式中, T 表示迭代次数, α 表示平滑系数, M_T 表示 teacher 模型的参数, M_S 表示 student 模型的参数。

α 的更新如式(3)所示:

$$\alpha_T = \min(1 - 1/(T + 1), \alpha_{T-1}). \quad (3)$$

3 本文框架

受 UDA 一致性训练框架^[4]的启发,针对有标签数据,采用端到端的训练方式,输入的图片经过 student 模型得到输出,然后输出的概率图与真实的标签作损失。针对无标签数据,采用 Mean teacher 模型来学习其拥有的特征, student 模型和经过有标签数据训练的网络权重共享,而 teacher 模型为一个和 student 模型具有相同结构的网络, teacher 模型的更新不是通过作损失来更新,而是通过 student 模

型来获得,然后 student 模型和 teacher 模型的输出作一致性损失。框架的总体损失由有标签部分的损

失加上无标签部分的损失组成,通过总体损失来更新处理有标签数据的网络。本文框架详见图1,图中

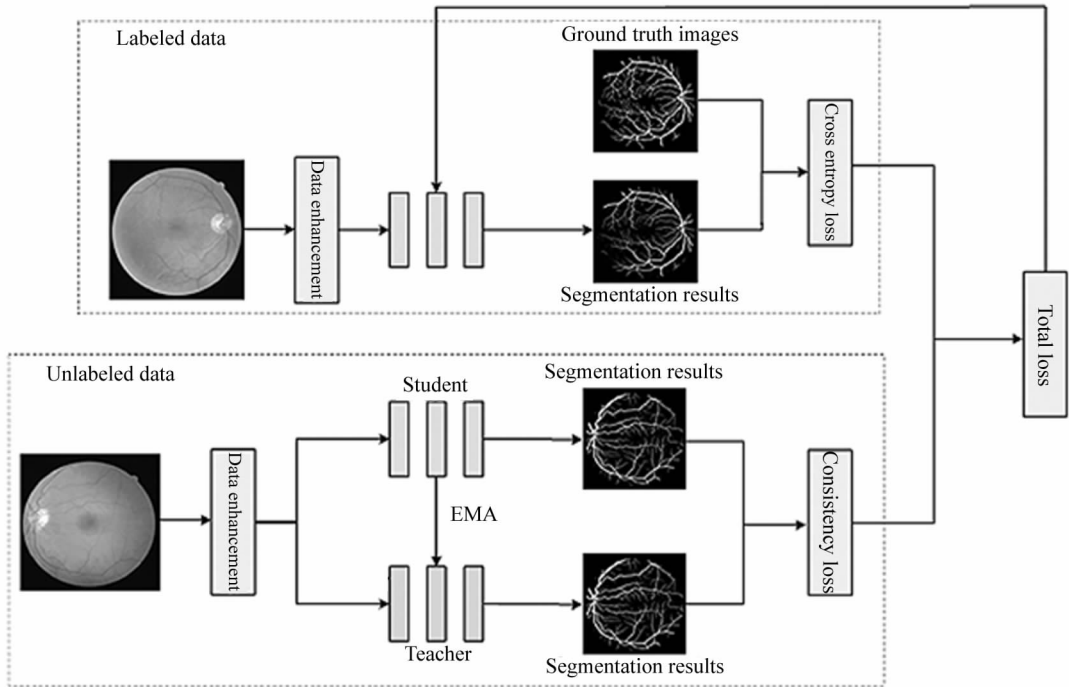


图 1 本文框架图

Fig. 1 Framework diagram of this paper

网络模型均采用 U-Net。

3.1 改进的 UDA 一致性训练框架

UDA 一致性训练框架^[4]将无标签数据通过不增强与增强后输入同一个结构的网络来获得两个具有相同特征但又有差别的输出,通过计算两个输出之间一致性损失结合分类损失来达到分类目的。而本文不仅对无标签数据进行增强,还对有标签数据进行了增强,在无标签数据采用的方法是 2.2 中的模型。在训练策略上,原始框架中仅更新有标签数据的模型和用于增强无标签数据的模型,而本文不仅更新了有标签数据训练的模型,通过权重共享还间接更新了无标签数据的模型。

3.2 用于无标签数据的 Mean teacher 模型

文献[9]将 Mean teacher 模型用于脑部 CT 分割,有标签数据和无标签数据均使用了 Mean teacher 模型。本文中,Mean teacher 模型只用于训练无标签数据,根据一致性的思想,同一个输入经过一个结构相同但参数略微不同的网络,得到的输出应具有一致的特征,不仅能减少训练过程中的内存开销,还

能简化框架结构。

3.3 损失函数

有标签数据训练时,损失函数采用二分类交叉熵损失 (binary cross entropy loss, BCELoss), 定义为:

$$Loss_{BCE}(\hat{y}, y) = -\frac{1}{n} \sum_{i=1}^n (y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)), \quad (4)$$

式中, y 表示标签图片在位置 i 的像素值,在本文中, $y_i \in (0, 1)$, $\hat{y}_i$ 表示预测图片在位置 i 的像素值,在输出的时候,会将预测图经过 Sigmoid 层,因此, $\hat{y}_i$ 的值为 $(0, 1)$ 。在损失函数计算过程中,图片中的所有像素点均会被考虑到。

无标签数据的损失函数采用均方误差损失 (mean square error loss, MSELoss), 定义为:

$$Loss_{MSE}(y_1, y_2) = \frac{1}{n} \sum_{i=1}^n (y_{1i} - y_{2i})^2, \quad (5)$$

式中, y_1 表示 student 模型的输出, y_2 表示 teacher 模型的输出, n 表示像素点的个数。

框架整体的损失采用 BCELoss 和 MSELoss 的联合损失:

$$Loss = \beta_1 Loss_{BCE} + \beta_2 Loss_{MSE}, \quad (6)$$

式中, β_1 和 β_2 是损失权重。

4 实 验

4.1 数据集

本文选择 DRIVE^[2] 和 CHASEDB1^[2] 两个公开数据集进行训练和测试。

DRIVE 数据集包括 40 幅彩色眼底图像, 分辨率为 565×584 , 每幅图像都有一个圆形的 45° 掩膜。将数据集划分为两部分, 一部分为 30 张的训练集, 其中 10 张为有标签数据, 20 张为无标签数据; 另一部分为 10 张的测试集, 每幅图像对应两个专家标注。

CHASEDB1 数据集包括 14 名英国儿童眼睛的 28 幅图像, 分辨率为 999×960 , 每幅图像有一个 30° 掩膜。将数据集的前 21 幅图像作为训练集, 其中 7 张为有标签数据, 14 张为无标签数据; 后 7 张作为测试集。

由于视网膜图像在采集过程中会受到光照度等因素的影响, 从而造成视网膜血管与背景对比度较低, 因此对数据集进行了预处理以提高血管和背景的对比度。首先将 RGB 图像转化为灰度图, 然后对灰度图进行归一化和自适应的直方图均衡化, 最后是对图像进行自适应的伽马校正。由于数据量比较少, 本文对训练图像做了数据增强以扩充数据集, 将原始图像和标记做了切片处理, 切片成分辨率 48×48 大小, 在切片过程中, 通过创建索引将切片后的图片与原图片中的位置一一对应。

4.2 实验环境及参数设置

实验在配备 16 G 内存和英伟达 RTX3060 显卡, 显存为 12 GB 的个人计算机上进行。本文提出的方法采用 Pytorch 框架实现, 使用带有默认参数设

置的 Adam 优化器来训练模型, 并采用 warmup 策略辅助训练, 初始学习率设置为 0.001, 前 20 轮学习率逐渐上升, 然后开始下降, 下降速率设置为 0.5, 批数量(batch size)设置为 32, 训练迭代轮次为 50 轮。

4.3 评价指标

为了评估本文框架的性能, 选用准确性(ACC)、敏感性(SE)、特异性(SP)、精确率(PR)和接收者工作特性曲线(ROC)下的面积(AUC)5 个指标来评价本文框架的性能, 前 4 个评价指标的定义如表 1 所示, 其中 TP 表示正确分类的血管像素数量, TN 表示正确分类的背景像素数量, FP 表示被误分为血管像素的背景像素的数量, FN 表示被误分为背景像素的血管像素的数量。

表 1 评价指标

Tab. 1 Evaluation indexes

ACC	SE	SP	PR
$\frac{TP+TN}{TP+TN+FP+FN}$	$\frac{TP}{TP+FN}$	$\frac{TN}{TN+FP}$	$\frac{TP}{TP+FP}$

4.4 实验结果与分析

4.4.1 与其他算法的分割指标对比

为了验证本文框架的有效性, 在 DRIVE 数据集和 CHASEDB1 数据集上把提出的框架和一些算法进行对比, 结果见表 2, 最优结果加粗表示。图 2 展示了不同算法在 DRIVE 上的可视化分割结果。

在 DRIVE 数据集上, 本文框架在 PR、SP 和 ACC 上取得了最好的成绩, 而在其他指标上低于其他算法。MPS-Net 提出不同分辨率的多路径尺度模块来增强网络的特征提取能力, 有利于对细小血管的提取, 在 SE 上达到了 0.836 1, 然而其余指标均低于本文框架, 说明其针对不同阈值像素的泛化能力

表 2 不同算法的指标对比

Tab. 2 Index comparison of different algorithms

Method	Year	DRIVE					CHASEDB1				
		PR	SE	SP	ACC	AUC	PR	SE	SP	ACC	AUC
ORLANDO ^[10]	2016	0.785 4	0.789 7	0.968 4	—	—	0.743 8	0.727 7	0.971 2	—	—
ATTU-Net ^[11]	2018	0.833 4	0.791 1	0.977 0	0.953 3	0.965 0	0.769 9	0.803 5	0.976 1	0.960 4	0.970 3
DEU-Net ^[12]	2019	—	0.803 8	0.980 2	0.957 8	0.982 1	—	0.813 2	0.981 4	0.966 1	0.887 6
HA-Net ^[13]	2020	—	0.799 1	0.981 3	0.958 1	0.982 3	—	0.823 9	0.981 3	0.967 0	0.987 1
KHURAM ^[14]	2021	—	0.814 1	0.970 2	0.954 0	0.939 9	—	0.815 3	0.971 1	0.956 1	0.956 5
MPS-Net ^[15]	2021	—	0.836 1	0.974 0	0.956 3	0.980 5	—	0.848 8	0.979 5	0.966 8	0.986 9
ZHOU ^[16]	2021	0.839 7	0.829 4	0.981 2	0.956 3	0.983 0	0.801 3	0.843 5	0.978 2	0.963 0	0.987 2
LI ^[17]	2021	—	0.792 1	0.981 0	0.956 8	0.980 6	—	0.781 8	0.981 8	0.963 5	0.981 0
PSPU-Net ^[18]	2021	—	0.781 4	0.981 0	0.955 6	0.978 0	—	0.819 5	0.972 7	0.959 0	0.978 4
Ours	2021	0.875 1	0.770 7	0.984 8	0.958 8	0.981 3	0.820 7	0.736 1	0.983 6	0.960 6	0.977 3

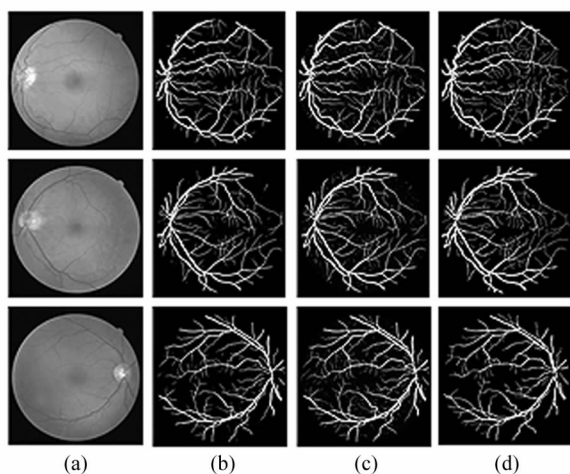


图 2 可视化分割比较:(a) 原始彩色基底图像;
(b) Ours; (c) ATTU-Net; (d) GT 图像

Fig. 2 Visual segmentation comparison:
(a) Original color fundus images; (b) Ours;
(c) ATTU-Net; (d) Ground truth images

较差。文献[18]采用了 U-Net 作为判别器,能够提高生成器对于血管树末梢的提取,在 SE 和 AUC 上都比本文框架高,但是在其余 3 项指标上均低于本文框架,是因为其算法在提取特征时,会将本不属于血管的背景噪声也归为血管类。

在 CHASEDB1 数据集上, PR 和 SP 的指标水平说明本文框架能够将血管的边缘像素与背景做出很好的区分。但是,CHASEDB1 数据集的图像光照度不一,明暗差距过大,中间视盘部分特别明亮,而边缘部分光照度很低,导致本文框架不能很好地提取细小血管,在 ACC 、 SE 、 AUC 上都低于其他算法,说明本文框架的鲁棒性较差。图 3 为本文框架在不同数据集上的分割概率图。

4.4.2 本文框架下不同基准网络的实验对比

为了探究在本文框架下采用不同网络的性能,以证明本文框架能够提升不同网络的分割效果,选取了 U-Net、Dense-Net^[19] 和 Ladder-Net^[20] 3 个基准网络,做了使用本文框架前后的对比实验,其中,Ladder-Net 为轻量级网络。实验结果如表 3 所示,相同基准网络下的最好结果用下划线表示。

3 个基准网络训练时使用的损失函数为 BCE-Loss 和 MSELoss,训练轮次设置为 200 轮,学习率设为 0.001,使用 20 张原始图像裁剪成 2 000 张进行训练。从表 3 中数据可以看出,在 U-Net 和 Ladder-Net 上,除了 SE ,其余指标均有所提升,Dense-Net 的全部指标都有所提升。在 DRIVE 数据集上, AUC

的提升分别为 2.82%、1.67%、0.79%,在 CHASEDB1 数据集上, AUC 分别提升了 1.53%、0.81%、0.69%。不同基准网络在本文框架下都能获得不同程度的提升,特别是在轻量级网络上也能获得提升,说明本文框架对于血管边界和不同阈值像素的区分都有一定的提升。图 4 显示了 3 种网络在 DRIVE 数据集上使用本文框架前后的分割效果可视化对比。

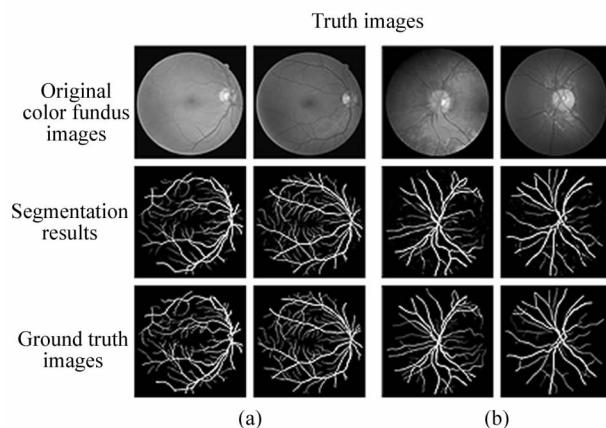


图 3 不同数据集的可视化结果:(a) DRIVE; (b) CHASEDB1

Fig. 3 Visualization results of different datasets:

(a) DRIVE; (b) CHASEDB1

4.4.3 消融实验

为了验证本文提出的策略能有效提高视网膜血管的分割性能,本文使用 U-Net 在 DRIVE 数据集做了两组实验来验证。实验的参数设置均与前文所说的保持一致,实验结果如表 4 所示,最优结果加粗表示。

从表中看出,Mean teacher 对整个网络的分割性能有了极大提升,但是在 SE 方面就有所降低,是因为该方法降低了网络对于噪声的分辨能力,会将背景噪声错误地归类为血管。UDA 则是能提升网络区分血管边界的能力,在 PR 和 SP 上表现较好,本文框架则是结合两种方法的优点, PR 和 SP 与表现最好的 UDA 方法区别不大,在 ACC 和 AUC 上也是取得了最好的成绩,提高了网络对血管整体的提取能力。表中后面 3 种方法的训练时间为 13min39 s、11min7 s、10min59 s,将两种方法结合后有效地减少了训练时间,且性能保持不变。

为了探究有标签数据和无标签数据在训练时不同的切片比例对实验结果的影响,以 U-Net 为基准网络,在 DRIVE 数据集和 CHASEDB1 数据集上取了 5 种不同的切片比例,通过实验得出最优比例,见

表 5。由表 5 可知,当有标签数据和无标签数据的切片比例在 1:1 时,有标签数据在整个数据的构成中占比最大,在两个数据集上共有 5 项指标取得最好的成绩。综合考虑,本文在做实验时设置的切片比

表 3 本文框架下不同网络的对比

Tab. 3 Comparison of different networks under the framework of this paper

Network	DRIVE					CHASEDB1				
	<i>PR</i>	<i>SE</i>	<i>SP</i>	<i>ACC</i>	<i>AUC</i>	<i>PR</i>	<i>SE</i>	<i>SP</i>	<i>ACC</i>	<i>AUC</i>
U-Net	0.833 9	<u>0.800 8</u>	0.977 0	0.953 3	0.953 0	0.779 9	<u>0.795 8</u>	0.977 6	<u>0.961 1</u>	0.962 0
U-Net+Ours	<u>0.875 1</u>	0.770 7	<u>0.984 8</u>	<u>0.958 8</u>	<u>0.981 2</u>	<u>0.820 7</u>	0.736 1	<u>0.983 6</u>	0.960 6	<u>0.977 3</u>
Dense-Net	0.866 1	0.747 9	0.983 1	0.953 2	0.964 3	0.791 1	0.767 6	0.979 8	0.960 6	0.972 5
Dense-Net+Ours	<u>0.868 9</u>	<u>0.782 9</u>	<u>0.983 7</u>	<u>0.959 3</u>	<u>0.981 0</u>	<u>0.804 1</u>	<u>0.794 8</u>	<u>0.980 2</u>	<u>0.963 0</u>	<u>0.980 6</u>
Ladder-Net	0.755 2	<u>0.844 7</u>	0.960 1	0.945 4	0.971 1	0.705 5	<u>0.800 3</u>	0.966 7	0.951 6	0.969 2
Ladder-Net+Ours	<u>0.846 4</u>	0.784 0	<u>0.980 3</u>	<u>0.956 5</u>	<u>0.979 0</u>	<u>0.771 9</u>	0.779 4	<u>0.976 5</u>	<u>0.958 2</u>	<u>0.976 1</u>

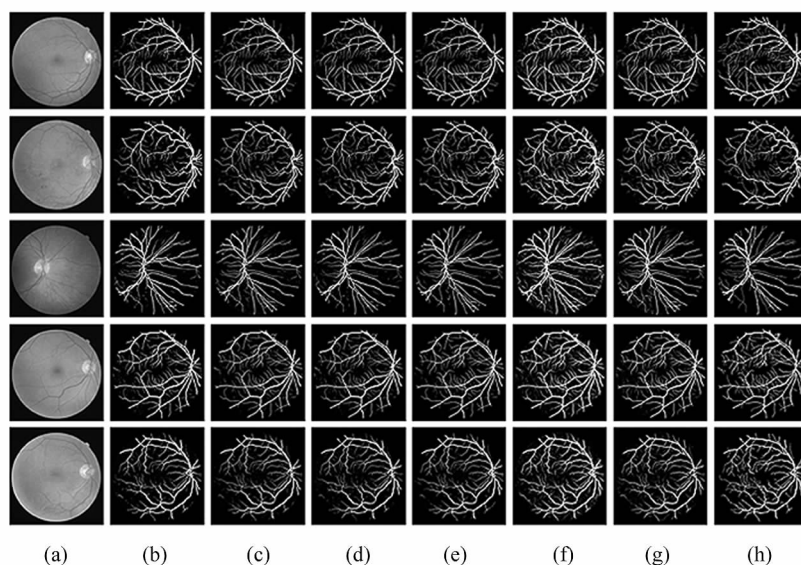


图 4 本文框架下不同网络的对比图:(a) 原始彩色眼底图像;(b) U-Net;(c) U-Net+Ours;(d) Dense-Net;(e) Dense-Net+Ours;(f) Ladder-Net;(g) Ladder-Net+Ours;(h) GT 图像

Fig. 4 Comparison diagram of different networks under the framework of this paper:

(a) Original color fundus images; (b) U-Net; (c) U-Net+Ours; (d) Dense-Net; (e) Dense-Net+Ours; (f) Ladder-Net; (g) Ladder-Net+Ours; (h) Ground truth images

例均为 1:1。

表 4 消融实验结果

Tab. 4 Ablation results

Method	<i>PR</i>	<i>SE</i>	<i>SP</i>	<i>ACC</i>	<i>AUC</i>
U-Net	0.833 9	0.800 8	0.977 0	0.955 3	0.953 0
U-Net+Mean teacher	0.866 4	0.778 6	0.983 4	0.958 5	0.980 6
U-Net+UDA	0.876 2	0.762 5	0.985 1	0.958 1	0.980 6
Ours	0.875 1	0.770 7	0.984 8	0.9588	0.981 2

为了得到一组合适的损失权重,本文在已有的实验基础上,在 DRIVE 数据集上设定了 4 组权重系数进行对比分析,实验结果如表 6 所示。从表 6 可以看出,当两个损失权重设置都为 1 时,*PR*、*SP* 和 *AUC* 都好于另外 3 组,且另外两项指标和其他 3 组都相差不大,说明在训练过程中无标签数据能够促进有标签数据部分的训练。

表 5 有标签数据和无标签数据不同切片比例分析

Tab. 5 Analysis of different slice proportions of labeled data and unlabeled data

Slice scale	DRIVE					CHASEDB1				
	PR	SE	SP	ACC	AUC	PR	SE	SP	ACC	AUC
1 : 1	0.875 1	0.770 7	0.984 8	0.958 8	0.981 2	0.820 7	0.736 1	0.983 6	0.960 6	0.977 3
1 : 2	0.871 4	0.772 4	0.984 2	0.958 5	0.981 1	0.795 9	0.801 2	0.979 0	0.962 5	0.981 0
1 : 3	0.866 1	0.782 7	0.983 3	0.958 9	0.980 8	0.799 7	0.789 3	0.979 8	0.962 2	0.979 7
1 : 4	0.873 5	0.774 8	0.984 5	0.959 0	0.981 0	0.803 4	0.783 3	0.980 4	0.962 2	0.979 4
1 : 5	0.867 8	0.780 5	0.983 6	0.958 9	0.981 1	0.797 6	0.801 7	0.979 2	0.962 8	0.980 8

表 6 不同损失权重分析

Tab. 6 Analysis of different loss weights

β_1, β_2	PR	SE	SP	ACC	AUC
0.5, 0.5	0.870 8	0.776 7	0.984 1	0.958 9	0.981 0
0.5, 1	0.862 7	0.789 4	0.982 6	0.959 2	0.981 0
1, 0.5	0.868 4	0.777 6	0.983 7	0.958 7	0.981 3
1, 1	0.875 1	0.770 7	0.984 8	0.958 8	0.981 3

5 结 论

本文从半监督视网膜血管分割任务中生成的伪标签质量参差不齐出发,提出了新的半监督框架,将其用在分割任务中,取得了较好的性能。通过使用 Mean teacher 模型来帮助模型学习无标签数据的信息,有标签数据采用了传统的有监督深度学习方法,提升了模型的分割性能。但是,在实验过程中发现,该框架对于噪声比较敏感,在干扰信息较多的眼底图像上表现出的分割性能不佳,也是本文框架在 SE 指标上表现不好的原因,因此,提高该框架的泛化性和鲁棒性将是本文下一步的工作方向。

参考文献:

[1] WU Y, YONG X, YANG S, et al. Multiscale network followed network model for retinal vessel segmentation [C]//LNCS 11071: Proceedings of the International Conference on Medical Image Computing and Computer-assisted Intervention, September 16-20, 2018, Granada, Spain. Berlin, Heidelberg: Springer, 2018: 119-126.

[2] LI L, ZHANG X, NIU D, et al. Research progress of deep learning in retinal vascular segmentation [J]. Computer Science and Exploration, 2021, 15(11): 2063-2076.

李兰兰, 张孝辉, 牛得草, 等. 深度学习在视网膜血管分割上的研究进展[J]. 计算机科学与探索, 2021, 15(11): 2063-2076.

[3] LEE D H. Pseudo-label: the simple and efficient semi-supervised learning method for deep neural networks[C]//

International Conference on Machine Learning, June 16-21, 2013, Atlanta, USA. New York: ACM, 2013, 3(2): 896-902.

[4] XIE Q, DAI Z, HOVY E, et al. Unsupervised data augmentation for consistency training [EB/OL]. (2014-04-29) [2022-02-22]. <http://arxiv.org/abs/1904.12848>.

[5] SOHN K, BERTHELOT D, LI C L, et al. Fixmatch: simplifying semi-supervised learning with consistency and confidence [EB/OL]. (2020-01-21) [2022-02-22]. <http://arxiv.org/abs/2001.07685>.

[6] LI Y, PEI Z, LI J, CHEN D. Semi-supervised learning framework in segmentation of retinal blood vessel based on U-Net[C]//33rd Chinese Control and Decision Conference, May 22-24, 2021, Kunming, China. New York: IEEE, 2021: 5972-5978.

[7] LI C, MA W, SUN L, et al. Hierarchical deep network with uncertainty-aware semi-supervised learning for vessel segmentation [EB/OL]. (2021-05-31) [2022-02-22]. <http://arxiv.org/abs/2105.14732>.

[8] TARVAINEN A, VALPOLA H. Mean teacher are better role models: weight-averaged consistency targets improve semi-supervised deep learning results [EB/OL]. (2021-01-21) [2022-02-22]. <http://arxiv.org/abs/1703.01780>, 2017.

[9] CUI W, LIU Y, LI Y, et al. Semi-supervised brain lesion segmentation with an adapted mean teacher model [C]//International Conference on Information Processing in Medical Imaging, July 2-7, 2019, Hong Kong, China. Berlin, Heidelberg: Springer, 2019: 554-565.

[10] ORLANDO J I, PROKOFYEVA E, BLASCHKO M B. A discriminatively trained fully connected conditional random field model for blood vessel segmentation in fundus images [J]. IEEE Transactions on Biomedical Engineering, 2016, 64(1): 16-27.

- [11] OKTAY O, SCHLEMPER J, FOLGOC L, et al. Attention u-net: learning where to look for the pancreas[EB/OL]. (2018-05-20)[2022-02-22]. <http://arxiv.org/abs/1804.03999>.
- [12] WANG B, QIU S, He H. Dual encoding u-net for retinal vessel segmentation[C]// LNCS 11764: International Conference on Medical Image Computing and Computer-Assisted Intervention, October 13-17, 2019, Shenzhen, China. Berlin, Heidelberg: Springer, 2019: 84-92.
- [13] WANG D, HAYTHAM A, POTTENBURGH J, et al. Hard attention net for automatic retinal vessel segmentation[J]. IEEE Journal of Biomedical and Health Informatics, 2020, 24(12): 3384-3396.
- [14] NAVEED K, ABDULLAH F, MADNI H A, et al. Towards automated eye diagnosis: an improved retinal vessel segmentation framework using ensemble block matching 3D filter[J]. Diagnostics, 2021, 11(1): 114-141.
- [15] LIN Z, HUANG J, CHEN Y, et al. A high resolution representation network with multi-path scale for retinal vessel segmentation[J]. Computer Methods and Programs in Biomedicine, 2021, 208(11): 6206-6221.
- [16] ZHOU Y, CHEN Z, SHEN L, et al. A refined equilibrium generative adversarial network for retinal vessel segmentation[J]. Neurocomputing, 2021, 437(19): 118-130.
- [17] LI Y, YANG J, NI J, et al. TA-Net: triple attention network for medical image segmentation[J]. Computers in Biology and Medicine, 2021, 137(10): 4836-4848.
- [18] DU X, WANG J, SUN W. U-Net retinal blood vessel segmentation algorithm based on improved pyramid pooling method and attention mechanism[J]. Physics in Medicine & Biology, 2021, 66(17): 5013-5025.
- [19] HUANG G, LIU Z, LAURENS V D M, et al. Densely connected convolutional networks[C]//2017 IEEE Conference on Computer Vision and Pattern Recognition, July 21-26, 2017, Honolulu, USA. Washington: Optical Society of America, 2017: 2261-2269.
- [20] ZHUANG J. Ladder-Net: multi-path networks based on U-Net for medical image segmentation[EB/OL]. (2018-10-17)[2022-02-22]. <http://arxiv.org/abs/1810.07810>.

作者简介:

吕 佳 (1978—), 女, 博士, 教授, 硕士生导师, 主要从事机器学习、数据挖掘及其在医学图像处理等方面的研究。