

Domain Adaptive Semantic Segmentation without Source Data

Fuming You¹, Jingjing Li^{1,2}, Lei Zhu³, Zhi Chen⁴, Zi Huang⁴ *

¹University of Electronic Science and Technology of China, Chengdu, China

²Yangtze Delta Region Institute (Huzhou), University of Electronic Science and Technology of China, Huzhou, China

³Shandong Normal University, Jinan, China

⁴University of Queensland, Brisbane, Australia

fumyou13@gmail.com, lijing117@yeah.net, leizhu0608@gmail.com, zhi.chen@uq.edu.au, huang@itee.uq.edu.au

ABSTRACT

Domain adaptive semantic segmentation is recognized as a promising technique to alleviate the domain shift between the labeled source domain and the unlabeled target domain in many real-world applications, such as automatic pilot. However, large amounts of source domain data often introduce significant costs in storage and training, and sometimes the source data is inaccessible due to privacy policies. To address these problems, we investigate domain adaptive semantic segmentation without source data, which assumes that the model is pre-trained on the source domain, and then adapting to the target domain without accessing source data anymore. Since there is no supervision from the source domain data, many self-training methods tend to fall into the “winner-takes-all” dilemma, where the *majority* classes totally dominate the segmentation networks and the networks fail to classify the *minority* classes. Consequently, we propose an effective framework for this challenging problem with two components: positive learning and negative learning. In positive learning, we select the class-balanced pseudo-labeled pixels with intra-class threshold, while in negative learning, for each pixel, we investigate which category the pixel does not belong to with the proposed heuristic complementary label selection. Notably, our framework can be easily implemented and incorporated with other methods to further enhance the performance. Extensive experiments on two widely-used synthetic-to-real benchmarks demonstrate our claims and the effectiveness of our framework, which outperforms the baseline with a large margin. Code is available at <https://github.com/fumyou13/LDBE>.

CCS CONCEPTS

• **Computing methodologies** → **Transfer learning**; *Computer vision*.

KEYWORDS

source-free domain adaptation, transfer learning, self-training, noisy label learning

*Jingjing Li is the corresponding author.

ACM Reference Format:

Fuming You¹, Jingjing Li^{1,2}, Lei Zhu³, Zhi Chen⁴, Zi Huang⁴. 2021. Domain Adaptive Semantic Segmentation without Source Data. In *Proceedings of the 29th ACM International Conference on Multimedia (MM '21)*, October 20–24, 2021, Virtual Event, China. ACM, New York, NY, USA, 10 pages. <https://doi.org/10.1145/3474085.3475482>

1 INTRODUCTION

In recent years, deep convolutional neural networks have achieved significant success across various multimedia tasks, such as cross-modal retrieval [43], image captioning [8] and so on. However, the impressive success heavily relies on the abundant labeled training data and the model fails to generalize to the novel and unseen instances. As a practical alternative, unsupervised domain adaptation (UDA) enables transferring knowledge from a labeled source domain to an unlabeled target domain, where different domains have distinctive distributions.

Semantic segmentation is a challenging task in real-world applications, which aims at assigning a semantic label to each pixel in an image. In practice, deep convolutional neural networks have achieved exciting performance on semantic segmentation, but heavily rely on the sufficient manual annotations, which are more expensive and time-consuming compared to other applications, e.g., object recognition. To handle this, many researchers turn to the synthetic datasets since the ground-truth labels are available, and transfer the knowledge learned from synthetic datasets to real-world applications with the help of UDA [1, 18, 23]. This paradigm is called *Domain Adaptive Semantic Segmentation* [10, 31] (DASS).

Domain adaptive semantic segmentation has got great attention and various methods have been proposed, which can be roughly divided into two categories: adversarial training [2, 5, 36, 40, 42] and self-training [17, 48, 52, 53]. Adversarial training based methods usually learn domain-invariant feature representations to achieve adaptation, while self-training based methods usually mitigate the domain gap iteratively through various strategies. Notably, many works [33, 44, 45, 47] integrate both of them to achieve better performance.

However, it is worth noting that DASS is still suffering several limitations: (1) When adapting to the target domain, accessing the large amounts of source data is inefficient and impractical. For instance, The size of the synthetic dataset GTA5 is nearly 61.6GB, which is usually quite hard to store and transmit. (2) The privacy of source data cannot be guaranteed, especially in the sensitive scenarios, e.g., medical image segmentation [32]. Therefore, it is necessary to investigate source-free domain adaptation (SFDA) in semantic segmentation. Specifically, we divide the training process

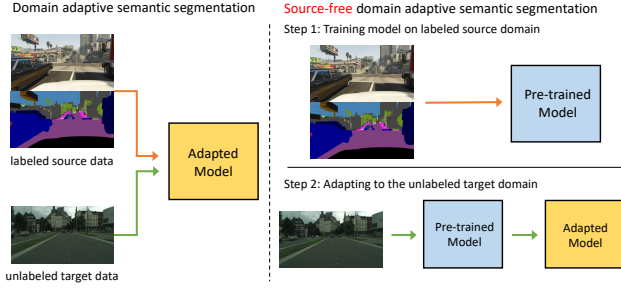


Figure 1: (Best viewed in color.) The comparison between domain adaptive semantic segmentation (DASS) and source-free domain adaptive semantic segmentation (SFSS). In DASS, both source and target domain data should be accessible during training. However, in SFSS, the training procedure is divided into two steps: (1) Training model on labeled source domain. (2) Adapting to the unlabeled target domain. Compared to conventional DASS, SFSS shows superiority on privacy protection and data transmission.

into two steps: training the supervised model on the labeled source domain, and then adapting to the unlabeled target domain based on the pre-trained model. For the sake of building an efficient and privacy-preserving domain adaptive semantic segmentation model [21], we propose a practical scenario named *Source-Free domain adaptive Semantic Segmentation (SFSS)*. The comparison between DASS and SFSS is illustrated in Fig. 1.

In SFSS, accessing the source data is prohibited during adapting to the target domain, i.e., the domain divergence [1, 19] between source and target domains is unknown [27]. Thus, most DASS methods focusing on mitigating the domain divergence to achieve adaptation [31, 40] would fail to generalize to SFSS. And recent SFDA methods are also not suitable for SFSS. Specifically, recent SFDA methods can be divided into two categories: GAN based [24] and pseudo-label based [14, 25, 27]. However, in semantic segmentation, (1) generating a pixel-level image is impractical, (2) and instance-level pseudo-label methods may fail in the dense task, whose objective is a pixel rather than an image. We will elaborate the comparison with other methods in Sec. 4.3 and Sec. 4.4. Furthermore, in the dense task semantic segmentation, pseudo-label based model tends to suffering from the “winner-takes-all”, which indicates that the model tends to over-fitting the *majority* classes, but ignoring the *minority* classes. This phenomenon becomes more serious in SFSS, while DASS allows access to the labeled data in source domain during adaptation, thus preventing the *majority* classes totally dominate the model and preserving the *minority* classes. More details are elaborated in Sec. 3.2 and Fig. 2.

To address above issues, we propose an effective framework for SFSS, which contains two mutually reinforcing components: positive learning and negative learning. In positive learning, we analyze in detail why the conventional self-training methods perform poorly in this setting and highlight the challenges of SFSS, then we introduce an intra-class threshold to filter the class-balanced pseudo-labeled pixels to prevent the “winner-takes-all”. In negative learning, we propose a heuristic complementary label selection (HCLS) to generate the complementary label for each pixel. The complementary label indicates which category the pixel does not belong to. Intuitively, applying the source pre-trained model to

the target domain directly leads to lots of noisy-labeled pixels due to the domain gap. Therefore, learning from the complementary label is a more proper strategy since there exists less noises in the complementary label. Extensive experiments demonstrate that each component is effective in improving the performance, and that combining both of them leads to better results. In a nutshell, we have following contributions:

- 1) We propose a novel and challenging scenario: *source-free domain adaptive semantic segmentation*, which provides a path of training an efficient and privacy-preserving semantic segmentation model. To the best of our knowledge, we are at the very early attempt to investigate the source-free domain adaptation in semantic segmentation.
- 2) We highlight the main challenges of SFSS, and point out why many mainstream self-training and source-free domain adaptation methods fail in this setting. Then, We re-implement them under SFSS setting, and the empirical results demonstrate our claims. Based on the observation of challenges, we propose an effective label-denoising framework named LD, which addresses this problem with two mutual-reinforcing components: positive learning and negative learning. The framework can be easily implemented and incorporated with other methods (e.g., attention module) to further enhance the performance.
- 3) We conduct experiments on two synthetic-to-real benchmarks: GTA5→Cityscapes and SYNTHIA→Cityscapes. Our method outperforms all comparisons with a significant improvement. Particularly, LD outperforms the baseline over 9.8% mIoU in GTA5→Cityscapes and 12.5% mIoU in SYNTHIA→Cityscapes, respectively. Furthermore, our method achieves competitive performance compared to the source-accessible methods, and shows superiority on privacy-preserving, data transmission and efficiency.

2 RELATED WORK

2.1 Source-free Domain Adaptation

Unsupervised domain adaptation (UDA) has witnessed significant development in both theory and various applications recently [1, 20, 22, 31, 50]. Domain adaptation aims at learning a general model from both labeled source domain and unlabeled target domain, where different domain has distinctive distributions [28, 39].

By considering the privacy and data transmission, source-free domain adaptation (SFDA) attracts a lot of attention lately [14, 27]. SFDA divides the training process into two steps: training the model on labeled source domain, and then adapt to the unlabeled target domain. Although SFDA is proposed very recently, it has been studied in many fields. For object recognition, 3C-GAN [24] generated target-style samples and leverages other regularizations to moderately retrain the source model. Hou et al. [11] converted the target-style images to the source-style based on the mean and variance stored in BN layers of the source model. SFDA [14] proposed to use the samples with low entropy to refine the pseudo labels for training. SHOT [27] proposed two strategies: information maximum loss and weighed deep clustering, and achieved state-of-the-art performance. For object detection, SFOD [25] proposed a metric named self-entropy decent to select the appropriate threshold for pseudo label generation.

To the best of our knowledge, we are at the very early attempt to investigate SFDA in semantic segmentation. In Sec. 4.4, we compare our method with other state-of-the-art SFDA methods in detail.

2.2 Domain Adaptive Semantic Segmentation

Domain Adaptive semantic segmentation (DASS) is an important application of UDA. Existing DASS methods can be roughly divided into two categories: adversarial training and self-training.

Adversarial training mainly includes two strategies: (1) Employing a style transfer across domains through adversarial training [2, 49]. (2) Learning the domain-invariant feature representations across domains [31, 36, 40–42]. For example, Vu et al. [42] match the entropy of output predictions in source and target domains via adversarial training. Recently, Luo et al. [30] proposed two modules to purify the features and perform the category-wise adversarial training based on it. However, these methods cannot be implemented directly in SFSS setting since the source data are unknown during training. Furthermore, adversarial based methods usually take more training time, while our method provides an efficient cross-domain semantic segmentation model (see Sec. 5.4).

Another mainstream line of work for DASS leverages the self-training. Some methods [4, 42, 53] utilizes prediction entropy to guarantee a distinct decision boundary in target domain. Recently, many methods product pseudo-labels in target domain based on confidence or uncertainty estimation [12, 17, 26, 33, 37, 48, 51, 52]. For example, Zou et al. [52] proposed a class-balanced selection strategy and selected the high-confidence pseudo-labeled target samples, while Zheng et al. [51] estimated the uncertainty via an auxiliary classifier and selected the rectified pseudo labels through the dynamic threshold. Unlike the aforementioned methods, Li et al. [17] proposed to select the source images that share similar distribution with the real ones in the target domain to alleviate the domain gap in another perspective.

The most related work is CBST [52], which proposed class-balanced pseudo-labeling. In our proposed label-denoising, intra-class threshold selection is also adopted. However, we have significant different motivations and claimed contributions. In SFSS, without any supervision from source domain, the class-imbalance issue becomes more severe than DASS. Once the category is a majority, it will be gradually replaced and totally dominated if there is no external intervention. In DASS, the supervision from source domain data preserve the minority classes, thus alleviating this issue. We summarize this phenomenon as “winner-takes-all”. Therefore, based on our original findings on the challenges of SFSS, class-balanced selection is a simple yet significant way to address it and it is not a direct re-use. Besides, we also propose negative learning to address other challenges.

3 APPROACH

In this section, we present the challenges of SFSS, and show why existing methods fail in this setting. Then, based on the observations, we address it with two components: positive learning and negative learning. Notably, our framework only contains the loss function of label denoising, which adds few burdens to the complexity of the model and the running time. Furthermore, our framework can

be easily implemented and incorporated with other methods (e.g., attention module) to achieve better performance.

3.1 Notations and Definitions

In domain adaptive semantic segmentation, we have a source domain $\mathcal{X}_S \subset \mathbb{R}^{H \times W \times 3}$ with corresponding ground-truth segmentation maps $\mathcal{Y}_S \subset \{(0, 1)\}^{H \times W \times C}$, where H, W, C denote the height, width and the number of classes, respectively. The target domain $\mathcal{X}_T \subset \mathbb{R}^{H \times W \times 3}$ is unlabeled. Let S be a segmentation network with parameters θ , which takes an image x and gives a prediction $p^{(h,w,c)}$. Especially, in SFSS, the training procedure has two stages: training a supervised loss on the labeled source domain and adapting to the unlabeled target domain. For simplicity, we optimize the cross-entropy loss over the source domain:

$$\mathcal{L}_{ce} = - \sum_{(x,y) \in (\mathcal{X}_S, \mathcal{Y}_S)} y \log S(x, \theta). \quad (1)$$

However, directly applying the model trained on source domain to the target domain often results in significant performance degradation [1, 17]. Therefore, we present our method and show how to enhance the performance on target domain without accessing source data.

3.2 Positive Learning

Pseudo labeling is a widely used strategy in semi-supervised learning, and attracts great attention in DASS recently. Intuitively, obtaining the pseudo labels through applying argmax operation on the softmax predictions is unreliable due to the domain gap between source and target domains. A trivial solution is to select the high-confidence pseudo labels, which effectively removes the wrong labels but neglects the “winner-takes-all” dilemma, where the model would be bias towards the *majority* classes and ignore the *minority* classes (see Fig. 2 (a)). Therefore, we propose to select the pixels with intra-class confidence. To avoid the imbalanced selection, the intra-class threshold is defined as:

$$\delta^{(c)} = \tau_{\alpha}(\mathcal{P}_T^{(c)}), \quad (2)$$

where τ_{α} means the top α (%) value in the softmax prediction value set $\mathcal{P}_T^{(c)}$ with respect to class $c \in C$. Then, for each class $c \in C$, we select the samples whose softmax confidence is larger than $\delta^{(c)}$. It is worth noting that our selection strategy addresses the aforementioned two challenges: (1) The most noise labels are filtered due to the high softmax prediction confidence selection. (2) The intra-class confidence selection addresses the dilemma where selecting high-confidence pseudo labels always results in the biased segmentation network. We update the pseudo labels at the beginning of each epoch to avoid additional running time. After selection, we optimize following cross-entropy loss at the selected pixels:

$$\mathcal{L}_{sce} = - \sum_{(x,\hat{y}) \in (\mathcal{X}'_T, \hat{\mathcal{Y}}'_T)} \hat{y} \log S(x, \theta), \quad (3)$$

where $(\mathcal{X}'_T, \hat{\mathcal{Y}}'_T)$ are the selected pixels on target domain and the pseudo labels are obtained by the argmax operation based on softmax prediction. Furthermore, we optimize the pixels not be selected by entropy minimization. Entropy minimization has been shown to

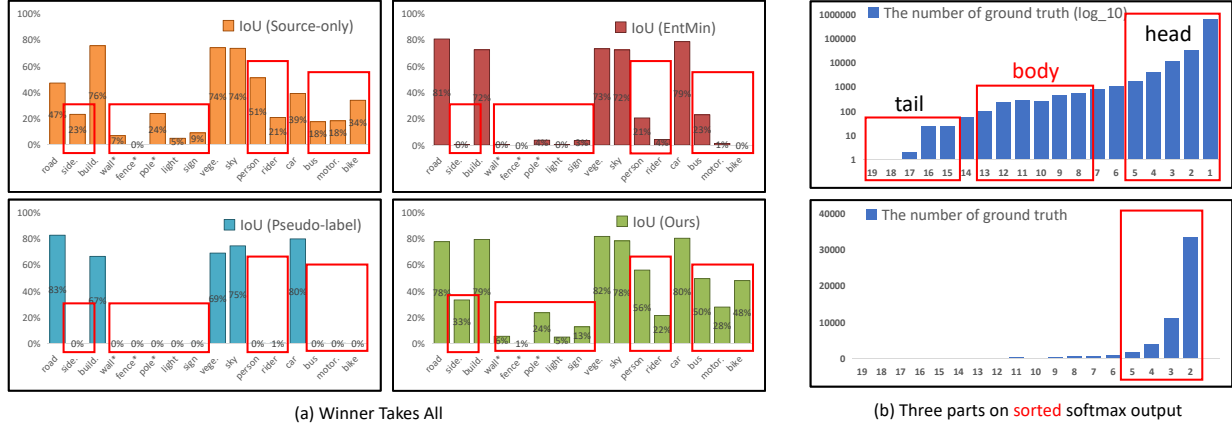


Figure 2: (Best viewed in color.) (a) The illustration of the results of different approaches on SYNTHIA→Cityscapes. Conventional approaches suffer from the “winner-takes-all”, while our method consistently improves the performance on each classes. (b) We randomly select a batch (two 600×600 images in a batch) in GTA5→Cityscapes for visualization. In figure (b), (x, y) indicates that there are y samples’ true classes correspond to the Top x softmax outputs. Notably, we adopt the logarithmic coordinate system with a base of 10 in the top figure, while in the bottom figure, we remove the Top 1 softmax output in another figure for better presentation since it is an order of magnitude above the rest.

Algorithm 1 Pseudo code for HCLS

Input: The sorted softmax output on target domain $p_T \in \mathcal{P}_T$, the number of classes C , the width of selection range ϵ .

Output: The complementary label $\bar{y}$;

- 1: $K = \lfloor C/2 \rfloor + \text{rand}(-\epsilon, +\epsilon)$
- 2: $\bar{y}$ = The class having the Top K softmax value.
- 3: **return** $\bar{y}$;

be effective in semi-supervised learning [7, 38] and domain adaptation [42, 46]. Vu et al. [42] argues that entropy minimization can be recognized as a soft-assignment version of cross-entropy loss. Therefore, to balance the unselected pixels, we optimize the following entropy minimization loss:

$$\mathcal{L}_{ent} = - \sum_{p \in \mathcal{P}_T} \sum_{c=1}^C p^{(c)} \log p^{(c)}, \quad (4)$$

where $\mathcal{P}_T$ is the set of softmax prediction over each pixels on target domain images. Now, we summarize the formulation from positive learning perspective:

$$\mathcal{L}_{pos} = \mathcal{L}_{sce} + \lambda_{ent} \mathcal{L}_{ent} \quad (5)$$

3.3 Negative Learning

In order to avoid the incorrect pseudo labels, we propose to address SFSS from another perspective, i.e., negative learning. Our idea is easy to be understood: instead of determining which category a pixel belongs to, it is easier to infer which category it does not belong to, which is more feasible in SFSS where no supervised information is provided. Let $x \in \mathcal{X}_T$ be an input, and $\hat{y} \in \mathcal{Y}_T$ be its corresponding pseudo label. We generate the complementary label $\bar{y} \in \bar{\mathcal{Y}}_T (\bar{y} \neq \hat{y})$, and optimize following loss function:

$$\mathcal{L}_{neg} = - \sum_{(x, \bar{y}) \in (\mathcal{X}_T, \bar{\mathcal{Y}}_T)} \bar{y} \log(1 - S(x, \theta)) \quad (6)$$

To the best of our knowledge, we are the first to perform negative learning in pixel-level. In previous work, Kim et al [15] perform negative learning in classification task and the complementary label is generated randomly from the labels of all classes except for the pseudo label $\hat{y}$. However, we further raise following question: which soft-label deserves to be optimized? To answer this question, we illustrate a toy example in Fig. 2 (b). Firstly, if we sort the softmax output $\mathcal{P}_T$, the class having the max softmax value is selected as the pseudo label. However, in SFSS, there are lots of noises in the pseudo labels due to the domain gap. Therefore, intuitively, the “head” classes (except the pseudo label) have a high probability of being the correct label and performing negative learning on these complementary labels causes an accumulation of errors. And the loss on the complementary labels inside the “tail” classes converges to 0. Therefore, we propose a heuristic complementary label selection (HCLS) to select the more accurate and valuable complementary labels, i.e., the “body” classes, which avoids the risks mentioned above. The pseudo code of HCLS in presented in Algorithm 1.

3.4 Overall Formulation

Actually, performing positive learning or negative learning individually improves the performance of the source pre-trained model on target domain (see Table 1, 2). Furthermore, incorporating them achieves better performance. On one hand, positive learning provides the class-balanced pseudo labels and helps correcting the complementary labels, thus improving the negative learning. On the other hand, negative learning denoises the pseudo labels, preventing the accumulation of errors in pseudo-labeling. Therefore, these two components are not mutually exclusive or independent, but reinforcing. The overall objective function for segmentation networks S becomes:

$$\begin{aligned} \mathcal{L}_{LD} &= \mathcal{L}_{pos} + \lambda_{neg} \mathcal{L}_{neg} \\ &= \mathcal{L}_{sce} + \lambda_{ent} \mathcal{L}_{ent} + \lambda_{neg} \mathcal{L}_{neg} \end{aligned} \quad (7)$$

Table 1: Results of GTA5→Cityscapes. ‘†’ means the results are based on our implementation. The mechanism ‘S’, ‘A’ and ‘T’ stand for self-training, adversarial training and image-to-image translation, respectively. ‘SF’ means whether the method is evaluated under source-free setting. “gain” indicates the mIoU improvement over using the source only under source-free settings. The best results under source-free settings are highlighted by bold and the best results under both source-free and source-accessible settings are highlighted by underlining.

GTA5→Cityscapes																							
Method	mech.	SF	road	side.	build.	wall*	fence*	pole*	light	sign	vege.	terr.	sky	pers.	rider	car	truck	bus	train	motor	bike	mIoU	gain
Source only†	-		60.6	17.4	73.9	17.6	20.6	21.9	31.7	15.3	79.8	18.1	71.1	55.2	22.8	68.1	32.3	13.8	3.4	34.1	21.2	35.7	-
EntMin(CVPR'19)†	S		82.8	0.0	70.2	2.2	0.3	0.4	2.8	1.6	79.9	8.1	79.2	22.2	0.1	83.1	22.5	30.0	2.0	6.3	0.0	26.0	-9.7 %
Pseudo†	S		83.2	0.0	67.3	1.1	0.0	0.1	1.2	1.2	77.7	1.3	81.4	11.5	0.0	81.7	18.0	14.8	0.0	3.7	0.0	23.4	-12.3%
Pse.+Ent.†	S		83.0	0.0	66.0	0.2	0.0	0.0	0.4	0.5	75.4	0.0	82.3	7.3	0.0	80.8	12.5	2.6	0.0	2.2	0.0	21.8	-13.9%
Pse.+Sel.†	S	✓	83.2	0.8	76.1	13.5	7.9	4.4	9.2	5.9	82.9	27.3	77.2	41.0	1.8	83.8	36.3	45.8	5.0	15.8	0.0	32.5	-3.2%
SHOT(ICML'20) [27]†	S		87.6	44.4	80.6	24.4	19.4	9.8	14.4	9.6	83.5	37.6	79.8	49.6	0.0	78.6	36.7	50.1	8.0	18.0	0.0	38.5	+2.8%
LD w/o $\mathcal{L}_{neg}$ (ours)†	S		89.6	46.1	66.6	30.7	8.7	28.7	32.8	29.7	81.5	36.5	83.4	57.0	26.7	82.9	28.9	31.5	0.3	16.5	38.3	43.0	+7.3%
LD w/o $\mathcal{L}_{pos}$ (ours)†	S		83.9	29.8	79.5	27.9	21.3	23.6	25.9	19.5	79.2	27.3	71.7	58.0	28.2	82.3	29.2	44.9	5.0	29.2	18.6	41.3	+5.6 %
LD(ours)†	S		91.6	53.2	80.6	36.6	14.2	26.4	31.6	22.7	83.1	42.1	79.3	57.3	26.6	82.1	41.0	50.1	0.3	25.9	19.5	45.5	+9.8 %
SIBAN(ICCV'19) [29]	A		88.5	35.4	79.5	26.3	24.3	28.5	32.5	18.3	81.2	40.0	76.5	58.1	25.8	82.6	30.3	34.4	3.4	21.6	21.5	42.6	-
AdaptSeg(CVPR'18) [40]	A		87.3	29.8	78.6	21.1	18.2	22.5	21.5	11.0	79.7	29.6	71.3	46.8	6.5	80.1	23.0	26.9	0.0	10.6	0.3	35.0	-
CLAN(PAMI'21) [30]	A		88.7	35.5	80.3	27.5	25.0	29.3	36.4	28.1	84.5	37.0	76.6	58.4	<u>29.7</u>	81.2	38.8	40.9	5.6	32.9	28.8	45.5	-
DPR(ICCV'19) [41]	SAT	✗	92.3	51.9	82.1	29.2	25.1	24.5	33.8	<u>33.0</u>	82.4	32.8	82.2	58.6	<u>27.2</u>	84.3	33.4	46.3	2.2	29.5	32.3	46.5	-
IntraDA(CVPR'20) [33]	SA		90.6	37.1	82.6	30.1	19.1	29.5	32.4	20.6	<u>85.7</u>	40.5	79.7	58.7	31.1	86.3	31.5	48.3	0.0	30.2	35.8	46.3	-
CRST(ICCV'19) [53]	S		91.0	55.4	80.0	33.7	21.4	37.3	32.9	24.5	85.0	34.1	80.8	57.7	24.6	84.1	27.8	30.1	<u>26.9</u>	26.0	42.3	47.1	-
DAST(AAAI'21) [47]	SA		92.2	49.0	84.3	36.5	<u>28.9</u>	<u>33.9</u>	<u>38.8</u>	28.4	84.9	41.6	83.2	<u>60.0</u>	28.7	<u>87.2</u>	45.0	45.3	7.4	33.8	32.8	49.6	-
CCM(ECCV'20) [17]	S		<u>93.5</u>	<u>57.6</u>	<u>84.6</u>	<u>39.3</u>	24.1	25.2	35.0	17.3	85.0	40.6	<u>86.5</u>	58.7	28.7	85.8	<u>49.0</u>	<u>56.4</u>	5.4	31.9	<u>43.2</u>	<u>49.9</u>	-

Algorithm 2 Pseudo code for the proposed LD

Input: The labeled source domain $\{\mathcal{X}_S, \mathcal{Y}_S\}$, the unlabeled target domain $\{\mathcal{X}_T\}$, the selection range α and ϵ , the segmentation networks S and the hyper-parameters λ_{ent} and λ_{neg} .

- 1: Training the segmentation networks S on the labeled source domain with loss function Eq. (1).
- 2: **for** iteration=1 to max_iteration **do**
- 3: Input $x_T \in \mathcal{X}_T$ to S , obtain the softmax predictions.
- 4: Calculate the class-wise threshold according to Eq. (2), and select the class-balanced pseudo labels.
- 5: Calculate the loss function of positive learning according to Eq. (5).
- 6: Generate the complementary labels according to Algorithm 1.
- 7: Calculate the loss function of negative learning according to Eq. (6).
- 8: Update the parameters of segmentation networks θ according to Eq. (7)
- 9: **end for**
- 10: **return** S .

where λ_{ent} and λ_{neg} are the hyper-parameters. And it is worth noting that the $\mathcal{L}_{sce}$ is calculated on the selected pixels, and the selected pixels are updated at the beginning of each epoch out of saving training time. And $\mathcal{L}_{ent}$ and $\mathcal{L}_{neg}$ are calculated on all pixels since both of them not explicitly provide supervision and can balance the unselected pixels. The details of training process are described in Algorithm 2.

4 EXPERIMENTS

We evaluate the proposed method on two synthetic-to-real tasks: GTA5→Cityscape and SYNTHIA→Cityscape, where GTA5 [34] and SYNTHIA [35] are synthetic datasets, while Cityscape [6] is a real-world dataset. In SFSS setting, we pre-train the model on

source domain dataset (GTA5 or SYNTHIA), and then adapt to the target domain dataset (Cityscapes).

4.1 Datasets

GTA5 is a synthetic dataset, which contains 24966 high-resolution images collected from game video and the corresponding ground-truth segmentation map can be generated by computer graphics. We train the 19 common classes with CityScapes dataset.

SYNTHIA is also a synthetic dataset, which contains 9400 images. And it shares 16 common classes with Cityscapes dataset.

Cityscapes is a real-world dataset collected for autonomous driving scenario from 50 cities around the world. It contains 2975 and 500 images for training and validation, respectively. In SFSS, the model is trained on the unlabeled training set and evaluated on the validation set with manual annotations.

4.2 Implementation Details

We keep the same network architecture as in previous methods for fair comparison. Specifically, we use Deeplab-V2 [3] with pre-trained ResNet-101 [9] as the segmentation network S and the Atrous Spatial Pyramid Pooling (ASPP) module is applied on the last layer’s output. Similar to [3], the sampling rates are fixed as {6,12,18,24}. Then, an up-sampling layer and a softmax operator are applied to obtain the prediction (segmentation map) with the matched size of input image. We implement our methods with Pytorch on a single NVIDIA RTX 2080Ti GPU. The segmentation network S is trained using the Stochastic Gradient Descent optimizer with momentum 0.9 and weight decay 5×10^{-4} . We adopt the polynomial learning rate scheduling with the power of 0.9, the initial learning rate is set as 1×10^{-3} and the batch size is 3. We adopt the same data augmentations as [17]: we resize the images’ short side to 720 and crop the images into 600×600 randomly. And horizontal flip and random scale between 0.5 and 1.5 are performed. The evaluation metrics is the mean intersection-over-union (mIoU).

Table 2: Results of SYNTHIA→Cityscapes. The mIoU* denotes the mean IoU of 13 classes, excluding the classes with “”. “gain” indicates the mIoU* improvement over using the source only under SFSS setting. We adopt the mIoU* as the evaluation metric.

SYNTHIA→Cityscapes																						
Method	mech.	SF	road	side.	build.	wall*	fence*	pole*	light	sign	vege.	sky	pers.	rider	car	bus	motor	bike	mIoU	mIoU*	gain	
Source only†	-		47.1	23.3	75.5	7.1	0.1	23.9	5.1	9.2	74.0	73.5	51.1	20.9	39.1	17.7	18.4	34.0	32.5	37.6	-	
EntMin(CVPR'19) [42]†	S		80.6	0.3	72.4	0.4	0.0	3.7	0.4	3.4	73.2	72.3	20.6	4.2	78.6	23.0	1.2	0.0	27.1	33.1	-4.5 %	
Pseudo†	S		82.6	0.0	66.5	0.0	0.0	0.3	0.0	0.4	69.0	74.5	0.4	0.5	79.8	0.4	0.4	0.0	23.4	28.8	-8.8 %	
Pse.+Ent.†	S		81.8	0.0	68.5	0.0	0.0	0.3	0.0	0.6	72.0	75.1	1.2	0.6	79.4	0.4	0.6	0.0	23.8	29.2	-8.4 %	
Pse.+Sel.†	S	✓	80.6	2.4	75.5	2.2	0.0	11.5	1.1	7.7	75.5	72.9	40.0	9.7	80.0	44.1	3.9	1.1	31.8	38.0	+0.2 %	
SHOT(ICML'20) [27]†	S		61.3	26.4	74.7	5.1	0.0	18.8	0.0	20.9	75.6	63.6	14.5	0.0	52.0	34.0	2.2	1.6	28.2	32.8	-4.8 %	
LD w/o $\mathcal{L}_{neg}$ (ours)†	S		78.3	33.0	78.4	3.9	0.4	19.9	7.3	11.4	80.0	76.8	47.5	19.0	80.3	42.8	19.6	44.7	40.2	47.6	+10.0 %	
LD w/o $\mathcal{L}_{pos}$ (ours)†	S		79.2	31.1	76.0	5.5	0.1	23.3	3.4	13.7	74.1	69.5	45.0	18.3	75.0	34.9	10.1	37.1	37.3	43.7	+6.1 %	
LD(ours)†	S		77.1	33.4	79.4	5.8	0.5	23.7	5.2	13.0	81.8	78.3	56.1	21.6	80.3	49.6	28.0	48.1	42.6	50.1	+12.4 %	
SIBAN(ICCV'19) [29]	A		82.5	24.0	79.4	-	-	-	16.5	12.7	79.2	82.8	58.3	18.0	79.3	25.3	17.6	25.9	-	46.3	-	
AdaptSeg(CVPR'18) [40]	A		84.3	42.7	77.5	-	-	-	4.7	7.0	77.9	82.5	54.3	21.0	72.3	32.2	18.9	32.3	-	46.7	-	
CLAN(PAMI'21) [30]	A		82.7	37.2	81.5	-	-	-	17.1	13.1	81.2	83.3	55.5	22.1	76.6	30.1	23.5	30.7	-	48.8	-	
DPR(ICCV'19) [41]	SAT	✗	82.4	38.0	78.6	8.7	0.6	26.0	3.9	11.1	75.5	84.6	53.5	21.6	71.4	32.6	19.3	31.7	40.0	46.5	-	
IntraDA(CVPR'20) [33]	SA		84.3	37.7	79.5	5.3	0.4	24.9	9.2	8.4	80.0	84.1	57.2	23.0	78.0	38.1	20.3	36.5	41.7	48.9	-	
CRST(ICCV'19) [53]	S		67.7	32.2	73.9	10.7	<u>1.6</u>	<u>37.4</u>	22.2	<u>31.2</u>	80.8	80.5	<u>60.8</u>	<u>29.1</u>	82.8	25.0	19.4	45.3	43.8	50.1	-	
DAST(AAAI'21) [47]	SA		<u>87.1</u>	<u>44.5</u>	<u>82.3</u>	10.7	0.8	29.9	13.9	13.1	81.6	<u>86.0</u>	<u>60.3</u>	<u>25.1</u>	<u>83.1</u>	40.1	24.4	40.5	<u>45.2</u>	52.5	-	
CCM(ECCV'20) [17]	S		79.6	36.4	80.6	<u>13.3</u>	0.3	25.5	<u>22.4</u>	14.9	<u>81.8</u>	77.4	56.8	25.9	80.7	45.3	<u>29.9</u>	<u>52.0</u>	<u>45.2</u>	<u>52.9</u>	-	

4.3 Evaluation Protocols and Baselines

To evaluate the proposed LD, we compare it with some widely-used self-training methods in domain adaptive semantic segmentation [42] and the state-of-the-arts SFDA method SHOT [27]. Here, we shortly introduce the compared methods and our implementation since some methods cannot be directly applied to source-free domain adaptive semantic segmentation. **EntMin.** We minimize the entropy of softmax output [42]. **Pseudo.** We directly obtain the pseudo label with argmax operation. Then, we minimize the cross-entropy loss on each pixel [16]. **Pse.+Ent.** We combine the aforementioned two methods, and the trade-off hyper-parameter is set as 1. **Pse.+Sel.** Since directly applying the pseudo label leads to many noisy predictions, many methods proposed various selection strategies [51–53]. A recent work [25], which investigates the SFDA in object detection, proposed to generate the pseudo labels through prediction confidence. In our implementation, we empirically fix the confidence threshold as 0.9, and minimize the cross-entropy loss on the selected pixels. **SHOT.** SHOT contains information maximization and deep clustering. However, the latter one cannot be applied to the semantic segmentation scenario directly. For instance, in semantic segmentation, a 600×600 image contains 360000 pixels with 2048-dim (ResNet) feature representations and a dataset contains large amounts of images. Therefore, performing clustering on the dataset in semantic segmentation is impractical and we implement information maximization.

4.4 Results

We verify the effectiveness of our method in two synthetic-to-real benchmarks: GTA5 → Cityscapes and SYNTHIA→Cityscapes. To ensure a fair comparison, all reported results share the same backbone: DeepLab-v2 with pre-trained ResNet-101. Since existing DASS methods cannot be implemented directly in the source-free settings, we then re-implement some widely-used self-training approaches in DASS [42] and the state-of-the-art SFDA method SHOT [27]. The details of implementation are elaborated in Sec. 4.3.

The results on GTA5→Cityscapes and SYNTHIA→Cityscapes are shown in Table 1 and Table 2, respectively. From Table 1 and 2, we have following findings:

- 1) Entropy minimization and direct pseudo-labeling leads to “winner-takes-all” under SFSS setting. Compare to Source-only, it has 9.7%, 12.3% drop in GTA5→Cityscapes, and 4.5%, 8.8% drop in SYNTHIA→Cityscapes. Both in DASS and SFSS, the model is biased towards the *majority* classes. However, in DASS, the labeled source domain data preserve the *minority* classes, while in SFSS, once the class is at a disadvantage, it will be gradually replaced if there is no external intervention. We illustrate the results of SYNTHIA→Cityscapes in Fig. 2, which vividly demonstrates our claims.
- 2) Selecting the pseudo-labeled samples based on the confidence threshold (Pse.+Sel.) effectively removes the wrongly pseudo-labeled samples, but still suffers from the “winner-takes-all”. Pse.+Sel. gets −3.2% in GTA5→Cityscapes, and +0.3% in SYNTHIA. However, in the *minority* classes (e.g., fence, pole, light, sign and etc.), Pse.+Sel. has a clear drop compared to “Source only”, revealing that the selected pixels are extremely class-imbalanced and the segmentation networks suffer from the “winner-tasks-all” dilemma. Existing state-of-the-art SFDA approach SHOT also gets poor performance in semantic segmentation scenario (i.e., +2.8% mIoU in GTA5→Cityscapes and −4.8% mIoU in SYNTHIA→Cityscapes).
- 3) Our approach LD surpasses the strong baseline “Source only” by a large margin under the challenging SFSS setting. We achieve a promising mIoU of 45.5% in GTA5→Cityscapes (+9.8%) and mIoU of 50.1% in SYNTHIA→Cityscapes (+12.4%), which highlights the effectiveness of our proposed LD. Compared to other methods under SFSS setting, our method maintains (e.g., wall, fence, pole, light, sign and etc.) or upgrades (e.g., side., rider, truck, bus and etc.) the performance of *minority* classes and consistently improves the performance of *majority* classes.

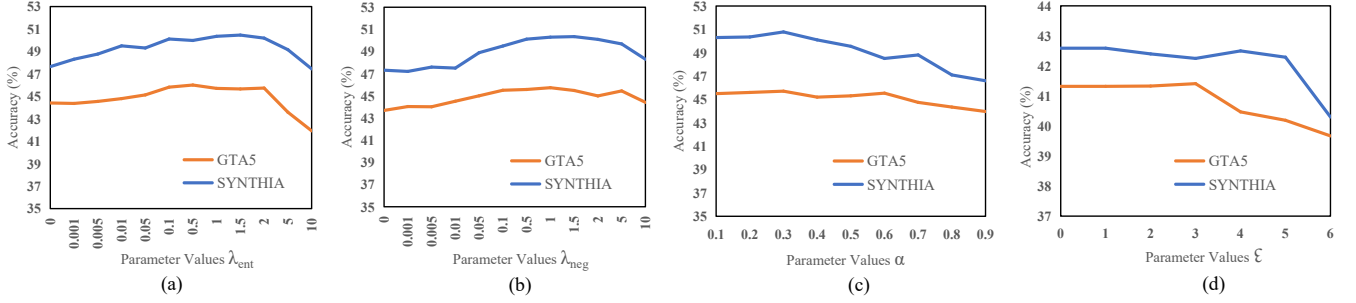


Figure 3: (Best viewed in color.) The results of parameter sensitivity on GTA5→Cityscapes and SYNTHIA→Cityscapes. (a) The weight of entropy minimization loss. (b) The weight of negative learning loss. (c) The selection range α in positive learning. (d) The selection range ϵ in negative learning.

4) Our LD even achieves competitive performance compare to the DASS methods, which access the labeled source domain data during adapting to the target domain. Our method outperforms the recent state-of-the-art adversarial based methods. For example, we surpass the representative adversarial-based work CLAN with a 1.3% mIoU in SYNTHIA→Cityscapes. Compared to state-of-the-art self-training based methods, our performance is also competitive. Notably, in SYNTHIA→Cityscapes, we achieves the best or the second performance at classes “bus”, “vege”, “motor” and “bike”, and all of them belong to the *minority* classes except for “vege”. Furthermore, our method only applies two self-training components, while many DASS methods combine many methods and design complex network architectures or sampling strategies (e.g., DPR, IntraDA and DAST). i.e., incorporating with other methods (e.g., attention module) will further enhance the performance of LD, which revealing the simplicity and expandability of the proposed framework.

5 ANALYSIS

5.1 Parameter Sensitivity

Hyper-parameters $\lambda_{ent}, \lambda_{neg}$. The overall objective function Eq. (7) contains two trade-off hyper-parameters λ_{ent} and λ_{neg} . We fix λ_{ent} as 1 and λ_{neg} as 1 for all tasks, respectively. Furthermore, we conduct parameter sensitivity analysis to evaluate LD on tasks GTA5→Cityscapes and SYNTHIA→Cityscapes. As shown in Fig. 3 (a) and (b), the performance steadily improves with the increasing parameters λ_{ent} and λ_{neg} from 0 to 1, demonstrating the effectiveness of each components in LD. And we observe that the performance would not be greatly influenced by the value of λ_{ent} and λ_{neg} , indicating that our LD is not quite sensitive to λ_{ent} and λ_{neg} .

Positive learning: selection range α . In Eq. (2), the parameter α controls the selection range of pseudo labels and we fix α as 0.2 for all tasks. To fully investigate the influence of different value of α , we select the balance weights from 0.1 to 0.9, since $\alpha = 0$ indicates applying negative learning only and $\alpha = 1$ indicates applying pseudo-labeling directly. From Fig. 3 (c), we find that the performance is not affected much by different values of α in GTA5→Cityscapes. In the more challenging task SYNTHIA→Cityscapes, the performance begins to drop after $\alpha > 0.3$ since it contains more and more wrong labels. It is worth noting that even with a large value α , the performance of LD is also acceptable, indicating the negative learning is effective in learning with noisy labels.

Negative learning: selection range ϵ . In Algorithm 1, ϵ represents the selection width of complementary label and we fix $\epsilon = 3$ for all tasks. To better reflect the influence of different value of ϵ , we apply negative learning only and modify the Eq. (1) to $K = \lfloor C/2 \rfloor + \text{rand}(-\epsilon, 0)$. In GTA5→Cityscapes, the number of classes $C = 19$, and $\lfloor C/2 \rfloor = 9$, while in SYNTHIA→Cityscapes, $C = 16$ and $K = \lfloor C/2 \rfloor = 8$. Therefore, we select the balance weights from 0 to 6, since $\epsilon = 7$ indicates that the complementary label may have the Top 1 value of the softmax output in SYNTHIA→Cityscapes. As shown in Fig. 3 (d), we observe that the performance begins to drop when $\epsilon \geq 4$ in GTA5→Cityscapes, which is consistent with our claims in Sec. 3.3 and Fig. 2: generating the complementary labels from the “head” classes results in the degeneration of model. Specifically, the empirical results indicate that our method outperforms the mainstream complementary strategy [13, 15], who generating the complementary label randomly from the labels of all classes except for the pseudo label, since our method generates more accurate and reliable complementary labels to perform negative learning.

5.2 Ablation Study

For the sake of examining the key components of LD, we perform ablation study by removing each component from LD at a time. In LD, we have two simple components: positive learning and negative learning. Therefore, we obtain two variants: (1) “LD w/o $\mathcal{L}_{pos}$ ” and (2) “LD w/o $\mathcal{L}_{neg}$ ”, which denote removing the corresponding self-training module, respectively. The results of ablation study are shown in Table 1 and 2. It is obvious that LD outperforms other variants and achieves significant improvement under source-free setting. “LD w/o $\mathcal{L}_{pos}$ ” achieves the mIoU of 41.3 in GTA5→Cityscapes and 43.7 in SYNTHIA→Cityscapes. “LD w/o $\mathcal{L}_{pos}$ ” does not explicitly provide any supervision but tells the model “which category does this pixel not belong to”. “LD w/o $\mathcal{L}_{neg}$ ” achieves the mIoU of 43.0 in GTA5→Cityscapes and 47.6 in SYNTHIA→Cityscapes, revealing that our pseudo-labeling effectively alleviate the “winner-takes-all” and enhance the performance of model in source-free setting.

5.3 Qualitative Analysis

We illustrate some qualitative segmentation examples in Fig. 4. Obviously, “Source only” model gives confused predictions due to the domain gap between GTA5 and Cityscapes. And many self-training methods performs well on the *majority* classes (e.g., “road”, “car”),

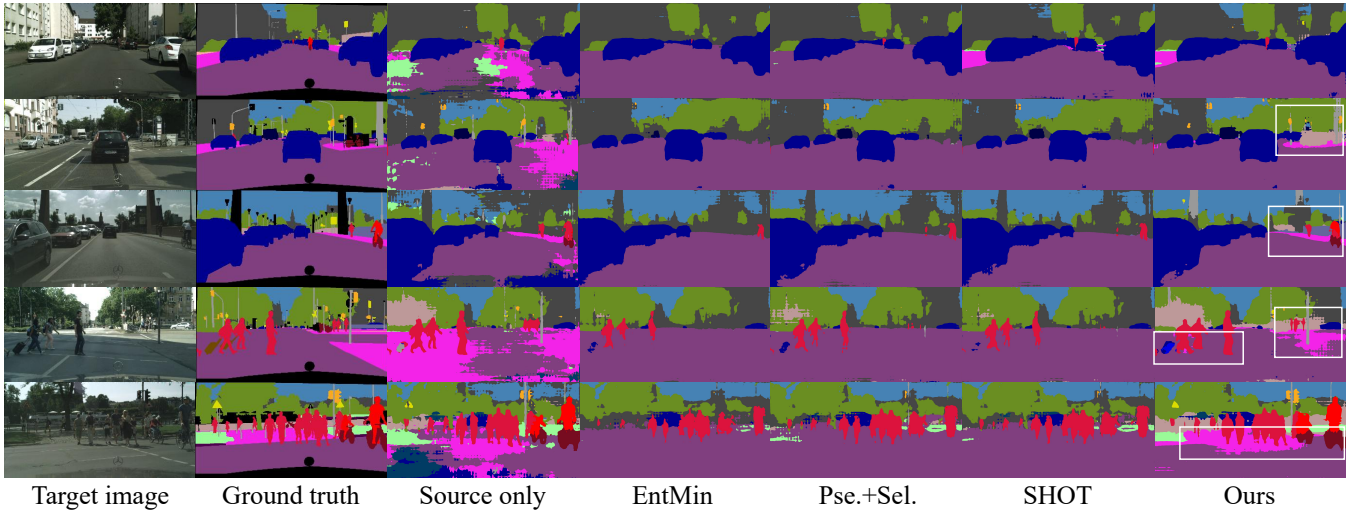


Figure 4: (Best viewed in color.) Qualitative results of source-free domain adaptive semantic segmentation on GTA5→Cityscapes. The key areas are highlighted with the white boxes.

but tends to over-fit the *majority* classes, e.g., “road” tends to occupy the “sidewalk”. Furthermore, they totally ignore the *minority* classes. For instance, the “bike” class is totally disappeared and replaced with “road” class, which creates a potential risk for real-world applications like automatic pilot. Unlike them, the proposed LD significantly enhances the performance of semantic segmentation. Specifically, LD stably improves the performance of *majority* classes, making the segmentation cleaner compared to “Source only” model. Furthermore, from the key areas highlighted with white boxes, our method yields better scalability to the *minority* classes, like confused classes (e.g., “bike”) and small-scale objectives (e.g., “traffic sign”). The qualitative results vividly confirms our claims and demonstrates the superiority of our method.

5.4 Efficiency Validation

Without accessing source data, our method shows superiority on privacy-preserving, data transmission and training time. In this section, we directly measure the size of training data and the training time of our proposed LD and other methods on task GTA5→Cityscapes and SYNTHIA→Cityscapes. We run the experiments on the same running environment (same hardware, same OS, same background applications, etc.). The results are reported in Table 3. In Table 3, “Training data” indicates the size of data when adapting to the target domain. “Training time” represents the time taken by the model to reach the reported mIoU. “Speedup” is calculated by: (The training time of the compared method - The training time of LD)/(The training time of LD). From the results, we can find that our method only takes 6.9GB training data during adapting to the target domain and has a significant speedup compared to the source-accessible methods. Meanwhile, our performance is also competitive to these methods, and even better than CLAN and IntraDA on SYNTHIA→Cityscapes. Therefore, our method achieves efficient domain adaptation.

Table 3: The results of efficiency validation.

Method	Training data	Training time	Speedup
GTA5→Cityscapes			
CLAN(PAMI’21) [30]	68.5GB	36572.32s	12.33×
IntraDA(CVPR’20) [33]	68.5GB	43531.11s	14.67×
CCM(ECCV’20) [17]	68.5GB	29285.02s	9.86×
LD(ours)	6.9GB	2967.85s	-
SYNTHIA→Cityscapes			
CLAN(PAMI’21) [30]	27.8GB	16390.57s	5.00×
IntraDA(CVPR’20) [33]	27.8GB	31017.96s	9.45×
CCM(ECCV’20) [17]	27.8GB	18196.11s	5.72×
LD(ours)	6.9GB	3281.14s	-

6 CONCLUSION

In this paper, we propose an effective and efficient framework named LD for source-free domain adaptive semantic segmentation. Specifically, LD selects the class-balanced pseudo-labeled pixels and performs negative learning with the proposed HCLS. Extensive experiments verify that LD outperforms the baseline with a large margin and even achieves competitive performance compared to existing state-of-the-art methods, which access the source data when adapting to the target domain. It is worth noting that our method can be easily incorporated with other modules to achieve better performance. We hope our framework can provide a solid baseline in source-free domain adaptive semantic segmentation and bring inspiration for future research.

ACKNOWLEDGEMENT

This work was supported in part by the National Natural Science Foundation of China under Grant 61806039 and 62073059, and in part by Sichuan Science and Technology Program under Grant 2020YFG0080 and 2020YFG0481.

REFERENCES

- [1] Shai Ben-David, John Blitzer, Koby Crammer, Alex Kulesza, Fernando Pereira, and Jennifer Wortman Vaughan. 2010. A theory of learning from different domains. *Machine learning* 79, 1-2 (2010), 151–175.
- [2] Wei-Lun Chang, Hui-Po Wang, Wen-Hsiao Peng, and Wei-Chen Chiu. 2019. All about structure: Adapting structural information across domains for boosting semantic segmentation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 1900–1909.
- [3] Liang-Chieh Chen, George Papandreou, Iasonas Kokkinos, Kevin Murphy, and Alan L. Yuille. 2017. Deeplab: Semantic image segmentation with deep convolutional nets, atrous convolution, and fully connected crfs. *IEEE transactions on pattern analysis and machine intelligence* 40, 4 (2017), 834–848.
- [4] Minghao Chen, Hongyang Xue, and Deng Cai. 2019. Domain adaptation for semantic segmentation with maximum squares loss. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*. 2090–2099.
- [5] Yun-Chun Chen, Yen-Yu Lin, Ming-Hsuan Yang, and Jia-Bin Huang. 2019. Crdco: Pixel-level domain transfer with cross-domain consistency. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 1791–1800.
- [6] Marius Cordts, Mohamed Omran, Sebastian Ramos, Timo Rehfeld, Markus Enzweiler, Rodrigo Benenson, Uwe Franke, Stefan Roth, and Bernt Schiele. 2016. The cityscapes dataset for semantic urban scene understanding. In *Proceedings of the IEEE conference on computer vision and pattern recognition*. 3213–3223.
- [7] Yves Grandvalet, Yoshua Bengio, et al. 2005. Semi-supervised learning by entropy minimization. In *CAP*. 281–296.
- [8] Longteng Guo, Jing Liu, Jinhui Tang, Jiangwei Li, Wei Luo, and Hanqing Lu. 2019. Aligning linguistic words and visual semantic units for image captioning. In *Proceedings of the 27th ACM International Conference on Multimedia*. 765–773.
- [9] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. 2016. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*. 770–778.
- [10] Judy Hoffman, Eric Tzeng, Taesung Park, Jun-Yan Zhu, Phillip Isola, Kate Saenko, Alexei A Efros, and Trevor Darrell. 2017. Cycada: Cycle-consistent adversarial domain adaptation. *arXiv preprint arXiv:1711.03213* (2017).
- [11] Yunzhang Hou and Liang Zheng. 2020. Source Free Domain Adaptation with Image Translation. *arXiv preprint arXiv:2008.07514* (2020).
- [12] Javed Iqbal and Mohsen Ali. 2020. Misl: Multi-level self-supervised learning for domain adaptation with spatially independent and semantically consistent labeling. In *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision*. 1864–1873.
- [13] Takashi Ishida, Gang Niu, Weihua Hu, and Masashi Sugiyama. 2017. Learning from complementary labels. *arXiv preprint arXiv:1705.07541* (2017).
- [14] Younggeun Kim, Sungeun Hong, Donghyeon Cho, Hyounseob Park, and Priyadarshini Panda. 2020. Domain Adaptation without Source Data. *arXiv preprint arXiv:2007.01524* (2020).
- [15] Youngdong Kim, Junho Yim, Juseung Yun, and Junmo Kim. 2019. Nlnl: Negative learning for noisy labels. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*. 101–110.
- [16] Dong-Hyun Lee et al. 2013. Pseudo-label: The simple and efficient semi-supervised learning method for deep neural networks. In *Workshop on challenges in representation learning, ICML*, Vol. 3.
- [17] Guangrui Li, Guoliang Kang, Wu Liu, Yunchao Wei, and Yi Yang. 2020. Content-consistent matching for domain adaptive semantic segmentation. In *European Conference on Computer Vision*. Springer, 440–456.
- [18] Jingjing Li, Erpeng Chen, Zhengming Ding, Lei Zhu, Ke Lu, and Zi Huang. 2019. Cycle-consistent conditional adversarial transfer networks. In *Proceedings of the 27th ACM International Conference on Multimedia*. 747–755.
- [19] Jingjing Li, Erpeng Chen, Zhengming Ding, Lei Zhu, Ke Lu, and Heng Tao Shen. 2020. Maximum Density Divergence for Domain Adaptation. *IEEE Transactions on Pattern Analysis and Machine Intelligence* (2020).
- [20] Jingjing Li, Mengmeng Jing, Ke Lu, Lei Zhu, and Heng Tao Shen. 2019. Locality preserving joint transfer for domain adaptation. *IEEE Transactions on Image Processing* 28, 12 (2019), 6103–6115.
- [21] Jingjing Li, Mengmeng Jing, Hongzu Su, Ke Lu, Lei Zhu, and Heng Tao Shen. 2021. Faster domain adaptation networks. *IEEE Transactions on Knowledge and Data Engineering* (2021).
- [22] Jingjing Li, Ke Lu, Zi Huang, Lei Zhu, and Heng Tao Shen. 2018. Heterogeneous domain adaptation through progressive alignment. *IEEE transactions on neural networks and learning systems* 30, 5 (2018), 1381–1391.
- [23] Jingjing Li, Ke Lu, Zi Huang, Lei Zhu, and Heng Tao Shen. 2019. Transfer Independently Together: A Generalized Framework for Domain Adaptation. *IEEE Transactions on Cybernetics* (2019).
- [24] Rui Li, Qianfen Jiao, Wenming Cao, Hau-San Wong, and Si Wu. 2020. Model adaptation: Unsupervised domain adaptation without source data. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 9641–9650.
- [25] Xianfeng Li, Weijie Chen, Di Xie, Shicai Yang, Peng Yuan, Shiliang Pu, and Yueting Zhuang. 2020. A Free Lunch for Unsupervised Domain Adaptive Object Detection without Source Data. *arXiv preprint arXiv:2012.05400* (2020).
- [26] Qing Lian, Fengmao Lv, Lixin Duan, and Boqing Gong. 2019. Constructing self-motivated pyramid curriculums for cross-domain semantic segmentation: A non-adversarial approach. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*. 6758–6767.
- [27] Jian Liang, Dapeng Hu, and Jiashi Feng. 2020. Do We Really Need to Access the Source Data? Source Hypothesis Transfer for Unsupervised Domain Adaptation. *arXiv preprint arXiv:2002.08546* (2020).
- [28] Mingsheng Long, Yue Cao, Jianmin Wang, and Michael I Jordan. 2015. Learning transferable features with deep adaptation networks. *arXiv preprint arXiv:1502.02791* (2015).
- [29] Yawei Luo, Ping Liu, Tao Guan, Junqing Yu, and Yi Yang. 2019. Significance-aware information bottleneck for domain adaptive semantic segmentation. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*. 6778–6787.
- [30] Yawei Luo, Ping Liu, Liang Zheng, Tao Guan, Junqing Yu, and Yi Yang. 2021. Category-Level Adversarial Adaptation for Semantic Segmentation using Purified Features. *IEEE Transactions on Pattern Analysis and Machine Intelligence* (2021).
- [31] Yawei Luo, Liang Zheng, Tao Guan, Junqing Yu, and Yi Yang. 2019. Taking a closer look at domain shift: Category-level adversaries for semantics consistent domain adaptation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 2507–2516.
- [32] Faisal Mahmood, Richard Chen, and Nicholas J Durr. 2018. Unsupervised reverse domain adaptation for synthetic medical images via adversarial training. *IEEE transactions on medical imaging* 37, 12 (2018), 2572–2581.
- [33] Fei Pan, Inkyu Shin, Francois Rameau, Seokju Lee, and In So Kweon. 2020. Unsupervised intra-domain adaptation for semantic segmentation through self-supervision. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 3764–3773.
- [34] Stephan R Richter, Vibhav Vineet, Stefan Roth, and Vladlen Koltun. 2016. Playing for data: Ground truth from computer games. In *European conference on computer vision*. Springer, 102–118.
- [35] German Ros, Laura Sellart, Joanna Materzynska, David Vazquez, and Antonio M Lopez. 2016. The synthia dataset: A large collection of synthetic images for semantic segmentation of urban scenes. In *Proceedings of the IEEE conference on computer vision and pattern recognition*. 3234–3243.
- [36] Kuniaki Saito, Kohei Watanabe, Yoshitaka Ushiku, and Tatsuya Harada. 2018. Maximum classifier discrepancy for unsupervised domain adaptation. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*. 3723–3732.
- [37] Christos Sakaridis, Dengxin Dai, and Luc Van Gool. 2019. Guided curriculum model adaptation and uncertainty-aware evaluation for semantic nighttime image segmentation. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*. 7374–7383.
- [38] Jost Tobias Springenberg. 2015. Unsupervised and semi-supervised learning with categorical generative adversarial networks. *arXiv preprint arXiv:1511.06390* (2015).
- [39] Baochen Sun and Kate Saenko. 2016. Deep coral: Correlation alignment for deep domain adaptation. In *European conference on computer vision*. Springer, 443–450.
- [40] Yi-Hsuan Tsai, Wei-Chih Hung, Samuel Schuster, Kihyuk Sohn, Ming-Hsuan Yang, and Manmohan Chandraker. 2018. Learning to adapt structured output space for semantic segmentation. In *Proceedings of the IEEE conference on computer vision and pattern recognition*. 7472–7481.
- [41] Yi-Hsuan Tsai, Kihyuk Sohn, Samuel Schuster, and Manmohan Chandraker. 2019. Domain adaptation for structured output via discriminative patch representations. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*. 1456–1465.
- [42] Tuan-Hung Vu, Himalaya Jain, Maxime Bucher, Matthieu Cord, and Patrick Pérez. 2019. Advent: Adversarial entropy minimization for domain adaptation in semantic segmentation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 2517–2526.
- [43] Bokun Wang, Yang Yang, Xing Xu, Alan Hanjalic, and Heng Tao Shen. 2017. Adversarial cross-modal retrieval. In *Proceedings of the 25th ACM international conference on Multimedia*. 154–162.
- [44] Zhonghao Wang, Mo Yu, Yunchao Wei, Rogerio Feris, Jinjun Xiong, Wen-mei Hwu, Thomas S Huang, and Honghui Shi. 2020. Differential treatment for stuff and things: A simple unsupervised domain adaptation method for semantic segmentation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 12635–12644.
- [45] Yanchao Yang and Stefano Soatto. 2020. Fda: Fourier domain adaptation for semantic segmentation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 4085–4095.
- [46] Fuming You, Hongzu Su, Jingjing Li, Lei Zhu, Ke Lu, and Yang Yang. 2021. Learning a Weighted Classifier for conditional domain adaptation. *Knowledge-Based Systems* (2021), 106774.
- [47] Fei Yu, Mo Zhang, Hexin Dong, Sheng Hu, Bin Dong, and Li Zhang. 2021. DAST: Unsupervised Domain Adaptation in Semantic Segmentation Based on Discriminator Attention and Self-Training. In *Proceedings of the AAAI Conference on Artificial Intelligence*, Vol. 35. 10754–10762.

- [48] Qiming Zhang, Jing Zhang, Wei Liu, and Dacheng Tao. 2019. Category anchor-guided unsupervised domain adaptation for semantic segmentation. *arXiv preprint arXiv:1910.13049* (2019).
- [49] Yiheng Zhang, Zhaoan Qiu, Ting Yao, Dong Liu, and Tao Mei. 2018. Fully convolutional adaptation networks for semantic segmentation. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*. 6810–6818.
- [50] Han Zhao, Remi Tachet Des Combes, Kun Zhang, and Geoffrey Gordon. 2019. On Learning Invariant Representations for Domain Adaptation. In *International Conference on Machine Learning*. 7523–7532.
- [51] Zhedong Zheng and Yi Yang. 2021. Rectifying pseudo label learning via uncertainty estimation for domain adaptive semantic segmentation. *International Journal of Computer Vision* (2021), 1–15.
- [52] Yang Zou, Zhiding Yu, BVK Kumar, and Jinsong Wang. 2018. Unsupervised domain adaptation for semantic segmentation via class-balanced self-training. In *Proceedings of the European conference on computer vision (ECCV)*. 289–305.
- [53] Yang Zou, Zhiding Yu, Xiaofeng Liu, BVK Kumar, and Jinsong Wang. 2019. Confidence regularized self-training. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*. 5982–5991.