

DYNAMIC DIVIDE-AND-CONQUER ADVERSARIAL TRAINING FOR ROBUST SEMANTIC SEGMENTATION

Anonymous authors

Paper under double-blind review

ABSTRACT

Adversarial training is promising for improving the robustness of deep neural networks towards adversarial perturbations, especially on the classification task. The effect of this type of training on semantic segmentation, contrarily, just commences. We make the initial attempt to explore the defense strategy on semantic segmentation by formulating a general adversarial training procedure that can perform decently on both adversarial and clean samples. We propose a dynamic divide-and-conquer adversarial training (DDC-AT) strategy to enhance the defense effect, by setting additional branches in the target model during training, and dealing with pixels with diverse properties towards adversarial perturbation. Our dynamical division mechanism divides pixels into multiple branches automatically, achieved by unsupervised learning. Note all these additional branches can be abandoned during inference and thus leave no extra parameter and computation cost. Extensive experiments with various segmentation models are conducted on PASCAL VOC 2012 and Cityscapes datasets, in which DDC-AT yields satisfying performance under both white- and black-box attacks.

1 INTRODUCTION

Recent work has revealed that deep learning models, especially in the classification task, are often vulnerable to adversarial samples (Szegedy et al., 2013; Goodfellow et al., 2014; Papernot et al., 2016b). Adversarial attack can deceive the target model by generating crafted adversarial perturbations on original clean samples. Such perturbations are often imperceptible. Meanwhile, such threat also exists in semantic segmentation (Xie et al., 2017b; Metzen et al., 2017; Arnab et al., 2018), as shown in Fig. 1. However, there is seldom work to improve the robustness of semantic segmentation networks. As a universal approach, *adversarial training* (Goodfellow et al., 2014; Kurakin et al., 2016b; Madry et al., 2017) is effective to enhance the target model in classification by training models with adversarial samples. In this paper, we study the effect of adversarial training on the semantic segmentation task. We find that adversarial training impedes convergence on clean samples, which also happens in classification. Thus we set our goal as making networks perform well on adversarial examples and meanwhile maintaining good performance on clean samples.

For the semantic segmentation task, each pixel has one classification output. Thus the property of every pixel in one image toward adversarial perturbations might be different. Based on this motivation, we design a dynamic divide-and-conquer adversarial training (DDC-AT) strategy. We propose to use multiple branches in the target model during training, each handling pixels with a set of properties. During training, a “main branch” is adopted to deal with pixels from adversarial samples and pixels from clean samples that are not likely to be perturbed; an “auxiliary branch” is utilized to deal with pixels from clean samples that are sensitive to perturbations.

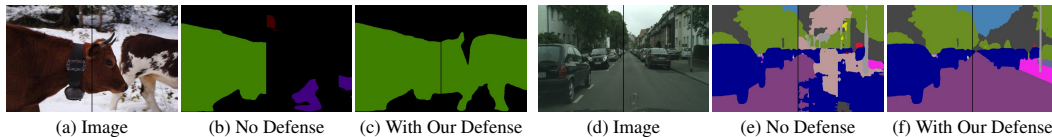


Figure 1: For each image in (a)/(d), the left side is the normal data while the right side is perturbed by adversarial noise. (b)/(e) shows that adversarial attack could fail existing segmentation models. We provide an effective defense strategy shown in (c)/(f).

REFERENCES

- Anurag Arnab, Ondrej Miksik, and Philip HS Torr. On the robustness of semantic segmentation models to adversarial attacks. In *CVPR*, 2018.
- Anish Athalye and Nicholas Carlini. On the robustness of the cvpr 2018 white-box adversarial example defenses. *arXiv:1804.03286*, 2018.
- Andreas Bar, Fabian Huger, Peter Schlicht, and Tim Fingscheidt. On the robustness of redundant teacher-student frameworks for semantic segmentation. In *CVPRW*, 2019.
- Qi-Zhi Cai, Min Du, Chang Liu, and Dawn Song. Curriculum adversarial training. *IJCAI*, 2018.
- Nicholas Carlini and David Wagner. Towards evaluating the robustness of neural networks. In *IEEE symposium on security and privacy*, 2017.
- Liang-Chieh Chen, George Papandreou, Florian Schroff, and Hartwig Adam. Rethinking atrous convolution for semantic image segmentation. *arXiv:1706.05587*, 2017.
- Marius Cordts, Mohamed Omran, Sebastian Ramos, Timo Rehfeld, Markus Enzweiler, Rodrigo Benenson, Uwe Franke, Stefan Roth, and Bernt Schiele. The cityscapes dataset for semantic urban scene understanding. In *CVPR*, 2016.
- Yinpeng Dong, Fangzhou Liao, Tianyu Pang, Hang Su, Jun Zhu, Xiaolin Hu, and Jianguo Li. Boosting adversarial attacks with momentum. In *CVPR*, 2018.
- Mark Everingham, Luc Van Gool, Christopher KI Williams, John Winn, and Andrew Zisserman. The pascal visual object classes (voc) challenge. *IJCV*, 2010.
- Ian J Goodfellow, Jonathon Shlens, and Christian Szegedy. Explaining and harnessing adversarial examples. *ICLR*, 2014.
- Bharath Hariharan, Pablo Arbeláez, Ross Girshick, and Jitendra Malik. Hypercolumns for object segmentation and fine-grained localization. In *CVPR*, 2015.
- Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *CVPR*, 2016.
- Xiaojun Jia, Xingxing Wei, Xiaochun Cao, and Hassan Foroosh. Comdefend: An efficient image compression model to defend adversarial examples. In *CVPR*, 2019.
- Harini Kannan, Alexey Kurakin, and Ian Goodfellow. Adversarial logit pairing. *arXiv:1803.06373*, 2018.
- Marvin Klingner, Andreas Bar, and Tim Fingscheidt. Improved noise and attack robustness for semantic segmentation by using multi-task training with self-supervised depth estimation. In *CVPRW*, 2020.
- Alexey Kurakin, Ian Goodfellow, and Samy Bengio. Adversarial examples in the physical world. *ICLR*, 2016a.
- Alexey Kurakin, Ian Goodfellow, and Samy Bengio. Adversarial machine learning at scale. *ICLR*, 2016b.
- Yanpei Liu, Xinyun Chen, Chang Liu, and Dawn Song. Delving into transferable adversarial examples and black-box attacks. *arXiv:1611.02770*, 2016.
- Aleksander Madry, Aleksandar Makelov, Ludwig Schmidt, Dimitris Tsipras, and Adrian Vladu. Towards deep learning models resistant to adversarial attacks. *arXiv:1706.06083*, 2017.
- Chengzhi Mao, Amogh Gupta, Vikram Nitin, Baishakhi Ray, Shuran Song, Junfeng Yang, and Carl Vondrick. Multitask learning strengthens adversarial robustness. *arXiv:2007.07236*, 2020.
- Jan Hendrik Metzen, Mummadi Chaithanya Kumar, Thomas Brox, and Volker Fischer. Universal adversarial perturbations against semantic image segmentation. In *ICCV*, 2017.

- Seyed-Mohsen Moosavi-Dezfooli, Alhussein Fawzi, and Pascal Frossard. Deepfool: a simple and accurate method to fool deep neural networks. In *CVPR*, 2016.
- Nicolas Papernot, Patrick McDaniel, and Ian Goodfellow. Transferability in machine learning: from phenomena to black-box attacks using adversarial samples. *arXiv:1605.07277*, 2016a.
- Nicolas Papernot, Patrick McDaniel, Somesh Jha, Matt Fredrikson, Z Berkay Celik, and Ananthram Swami. The limitations of deep learning in adversarial settings. In *2016 IEEE European Symposium on Security and Privacy*, 2016b.
- Nicolas Papernot, Patrick McDaniel, Ian Goodfellow, Somesh Jha, Z Berkay Celik, and Ananthram Swami. Practical black-box attacks against machine learning. In *Proceedings of the 2017 ACM on Asia conference on computer and communications security*, 2017.
- Chuanbiao Song, Kun He, Liwei Wang, and John E Hopcroft. Improving the generalization of adversarial training with domain adaptation. *ICLR*, 2018.
- Christian Szegedy, Wojciech Zaremba, Ilya Sutskever, Joan Bruna, Dumitru Erhan, Ian Goodfellow, and Rob Fergus. Intriguing properties of neural networks. *arXiv:1312.6199*, 2013.
- Florian Tramèr, Alexey Kurakin, Nicolas Papernot, Ian Goodfellow, Dan Boneh, and Patrick McDaniel. Ensemble adversarial training: Attacks and defenses. *ICLR*, 2017.
- Jianyu Wang and Haichao Zhang. Bilateral adversarial training: Towards fast training of more robust models against adversarial attacks. In *ICCV*, 2019.
- Chaowei Xiao, Ruizhi Deng, Bo Li, Fisher Yu, Mingyan Liu, and Dawn Song. Characterizing adversarial examples based on spatial consistency information for semantic segmentation. In *ECCV*, 2018.
- Cihang Xie, Jianyu Wang, Zhishuai Zhang, Zhou Ren, and Alan Yuille. Mitigating adversarial effects through randomization. *ICLR*, 2017a.
- Cihang Xie, Jianyu Wang, Zhishuai Zhang, Yuyin Zhou, Lingxi Xie, and Alan Yuille. Adversarial examples for semantic segmentation and object detection. In *CVPR*, 2017b.
- Cihang Xie, Yuxin Wu, Laurens van der Maaten, Alan L Yuille, and Kaiming He. Feature denoising for improving adversarial robustness. In *CVPR*, 2019.
- Haichao Zhang and Jianyu Wang. Defense against adversarial attacks using feature scattering-based adversarial training. In *NIPS*, 2019.
- Hengshuang Zhao, Jianping Shi, Xiaojuan Qi, Xiaogang Wang, and Jiaya Jia. Pyramid scene parsing network. In *CVPR*, 2017.

Table 5: The architecture and initial learning rate of the model (PSPNet/DeepLabv3) for the method with no defense and SAT. The visual illustration can be seen from Fig. 4 and 5. “Base LR” denotes initial learning rate of different modules.

Module Name	Base LR
Backbone	-
ResNet	0.01
PPM/ASPP	0.01×10
Branch	-
Convolution	0.01×10
Batch Normalization	0.01×10
ReLU	-
Convolution	0.01×10

Table 6: The architecture and initial learning rate of the model (PSPNet/DeepLabv3), which is employed in DDC-AT. “Base LR” denotes initial learning rate of different modules. “BN” means Batch Normalization.

Module Name			Base LR
Shared Backbone			-
ResNet			0.01
PPM/ASPP			0.01×10
Main branch	Auxiliary branch	Mask branch	-
Convolution	Convolution	Convolution	0.01×10
BN	BN	BN	0.01×10
ReLU	ReLU	ReLU	-
Convolution	Convolution	Convolution	0.01×10

Details of Model Structures The model structures for the method with no defense, which are PSPNet (Zhao et al., 2017) and DeepLabv3 (Chen et al., 2017) here, are shown in Table 5. Different modules in corresponding models have different initial learning rates. Meanwhile, the model structure of SAT is the same as the model trained with no defense. On the other hand, we add extra segmentation branches to target models in DDC-AT during training. One important problem is where to set main branch, auxiliary branch and mask branch. In DDC-AT, these three branches are separated from the same location for both VOC and Cityscapes: for PSPNet, they are separated after the PPM module; for DeepLabv3, they are separated after the ASPP module. Thus, the structures for PSPNet and DeepLabv3 in DDC-AT are shown as Table 6. Moreover, note that the split point is close to the output of network, which is convenient to choose for models with other structures.

A.3 EXPERIMENTS

Experiments–Analysis of Standard Deviation Value We report the standard deviation values for the experiments in Table 1, 2, 3, 4 here. The standard deviation values for the experiments of white-box attack on VOC are shown in Table 7. It is obvious that the standard deviation of SAT is low, which means SAT is stable. Further, smaller standard deviation for DDC-AT indicates that its results are more stable. Moreover, the standard deviation values for the experiments of black-box attack on VOC are also exhibited in Table 7. It should be note that the standard deviation of SAT is larger than the results by white-box attacks because black-box perturbations for each model are obtained from a set of substitute networks, which have different adversarial behaviors. Meanwhile, the standard deviation of DDC-AT is lower in all cases, especially under unseen attacks. It proves that DDC-AT is more stable and effective than SAT by all types of attack. Furthermore, the standard deviation values for the experiments of white-box attack and black-box attack on Cityscapes are simultaneously displayed in Table 7 and these results also support that both DDC-AT and SAT are stable while DDC-AT is even better.

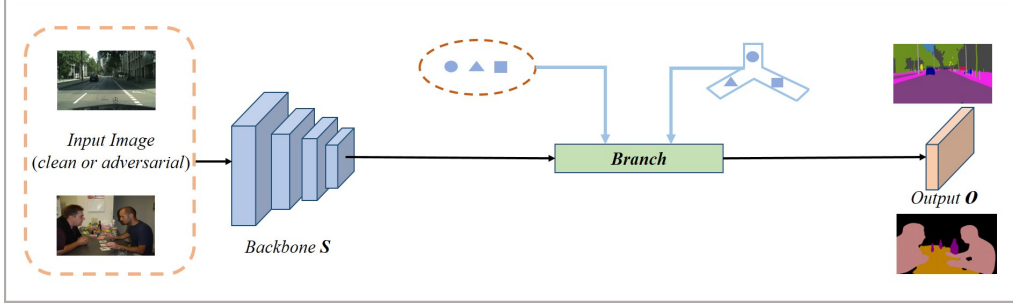


Figure 4: The training framework for the method **without defense**. Such training scheme does not add adversarial samples into training.

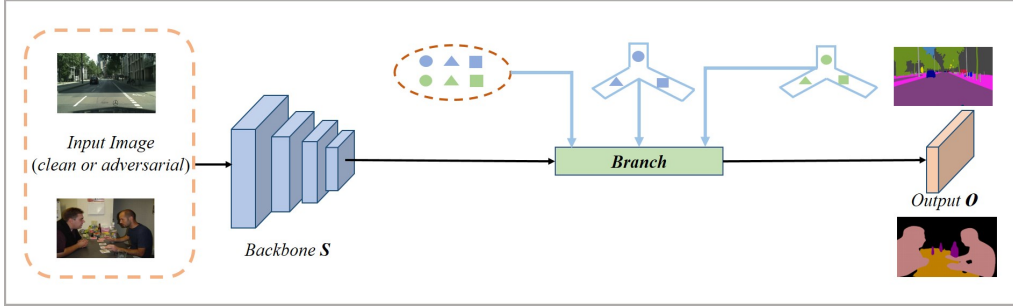


Figure 5: The training framework for standard adversarial training **SAT**. Such training scheme always uses one identical branch to align clean and adversarial samples.

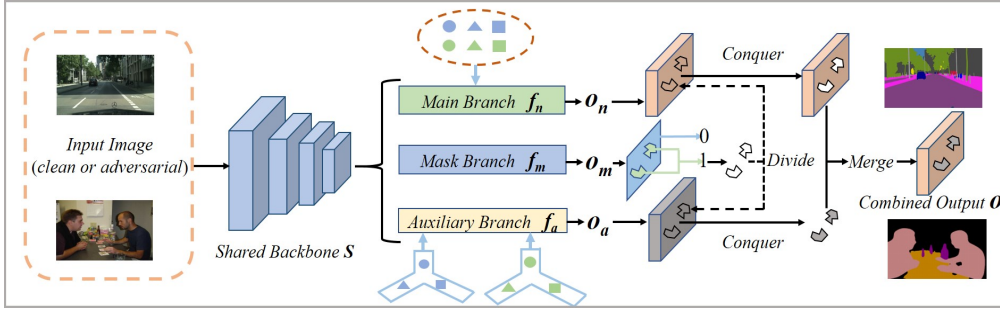


Figure 6: The training framework for **DDC-AT-M**

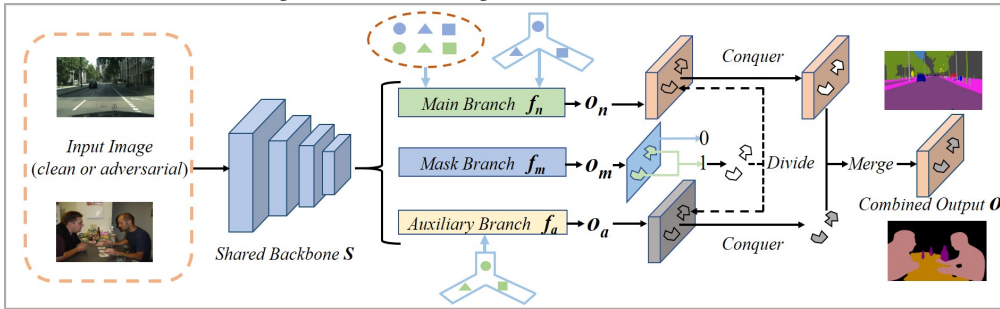


Figure 7: The training framework for **DDC-AT-N**

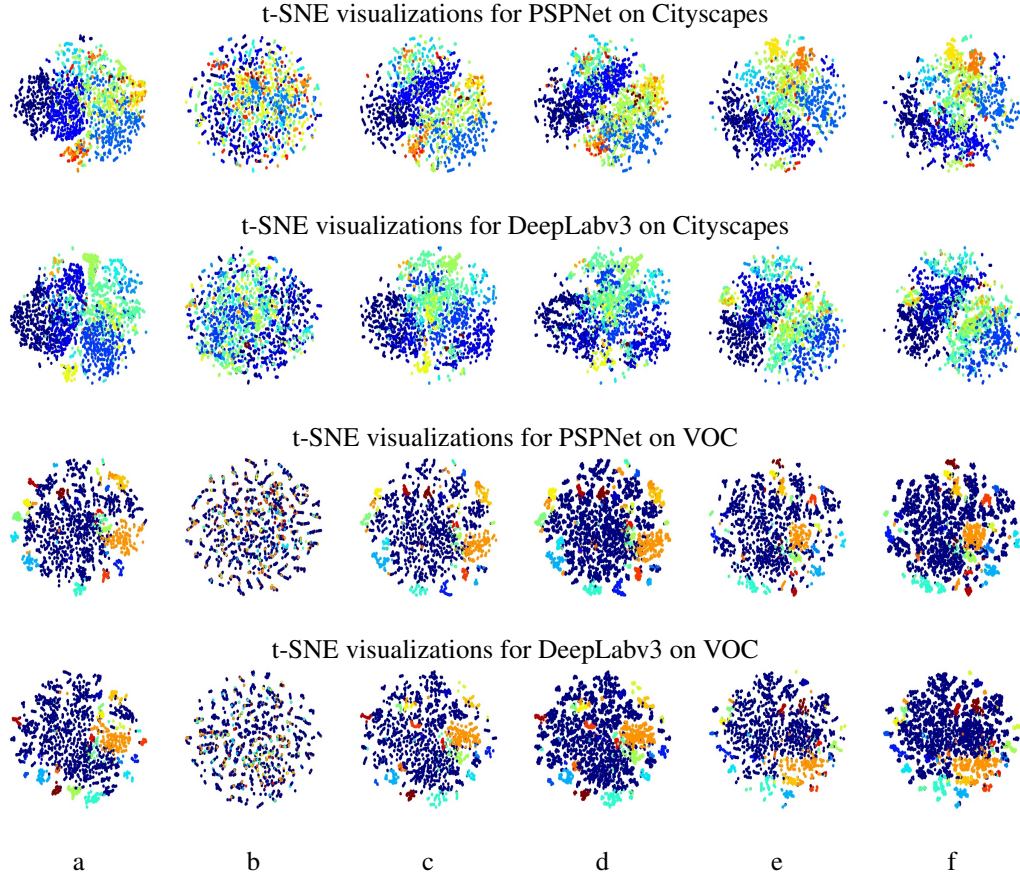


Figure 8: The t-SNE analysis for VOC and Cityscapes. a: clean samples in the model with no defense, b: adversarial samples in the model with no defense, c: clean samples in the SAT model, d: adversarial samples in the SAT model, e: clean samples in the DDC-AT model, f: adversarial samples in the DDC-AT model. The adversarial samples are generated with white-box BIM attack (L_∞ constraint, $n=2$, $\epsilon = 0.03 \times 255$, $\alpha = 0.01 \times 255$)

A.4 ANALYSIS

t-SNE Analysis for SAT and DDC-AT To further analyze the defense effect of SAT and DDC-AT, we use t-SNE to visualize the corresponding feature distribution of trained models. As displayed in Fig. 8, we find: for the distribution of clean samples in models with no defense, features from different classes are separated severally. This is helpful for segmentation. However, for the distribution of adversarial samples, features from different categories are mixed. For SAT and DDC-AT, features of adversarial samples and clean samples from different classes are both separated. Thus, SAT and DDC-AT can improve the robustness of models and keep great performance on clean samples.

Visual Analysis for SAT and DDC-AT We provide visual cases from the results of models trained with No Defense, SAT and DDC-AT on VOC and Cityscapes dataset. They are shown in Fig. 9, 10, 11 and 12. All corresponding models are trained with ResNet50 as backbone, and evaluated under white-box attacks. The attacks are executed with the hyper-parameters as $\epsilon = 0.03 \times 255$, $\alpha = 0.01 \times 255$, $n = 3$. As we can see, both SAT and DDC-AT can improve the effect of defense on these datasets, while the performance of DDC-AT is better than SAT.

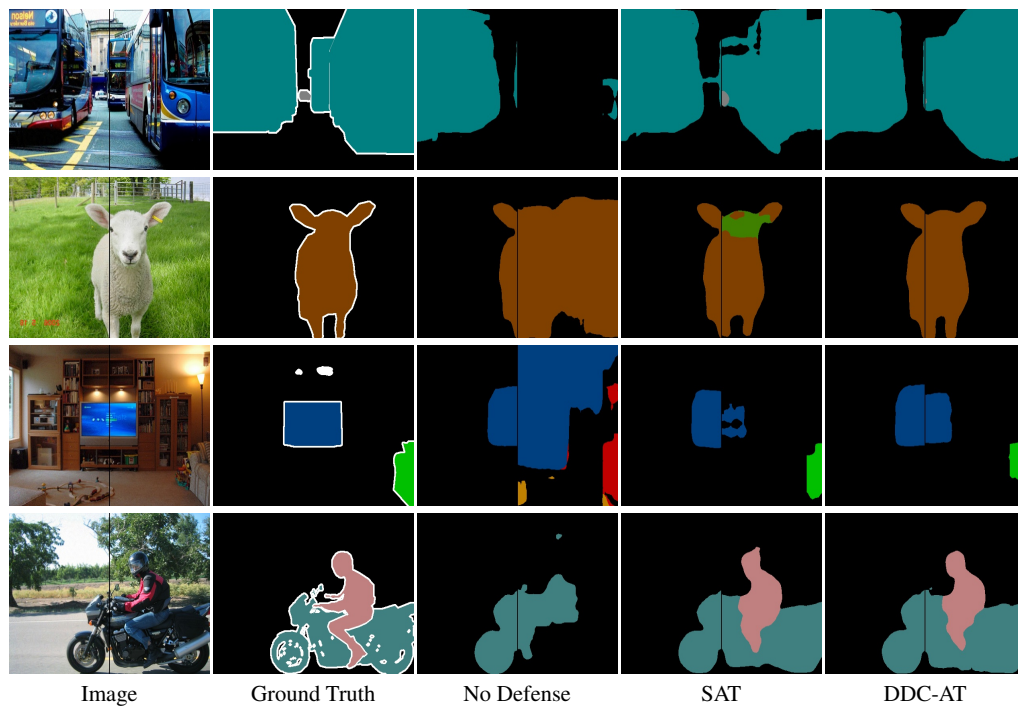


Figure 9: The visual comparison for different defense methods, which are executed on VOC dataset with model structure as PSPNet.

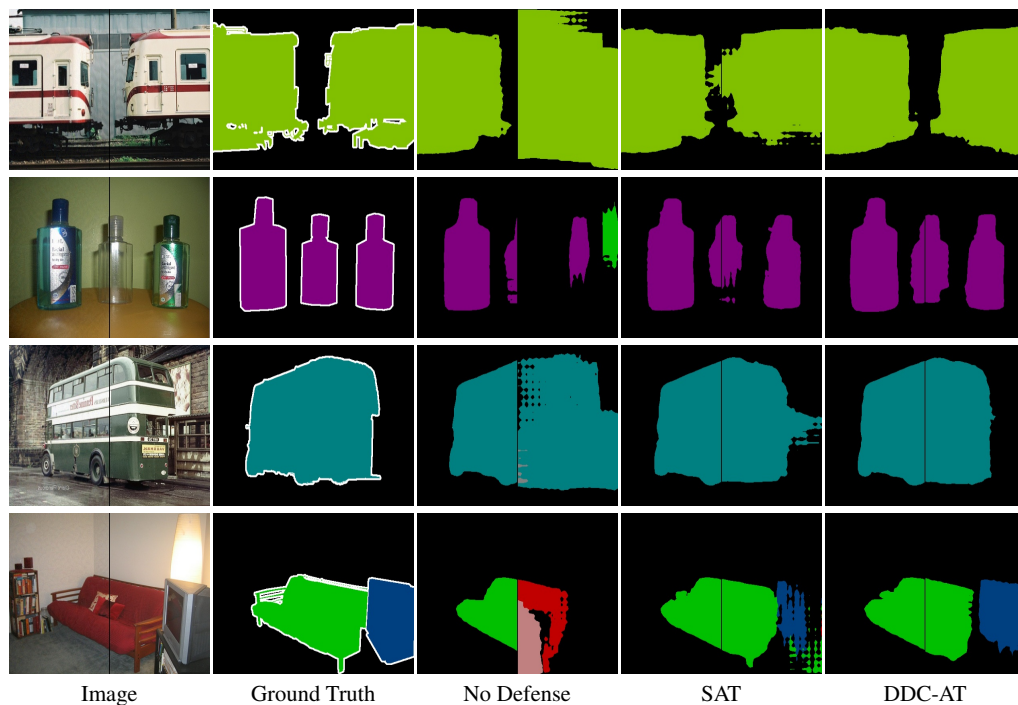


Figure 10: The visual comparison for different defense methods, which are executed on VOC dataset with model structure as DeepLabv3.

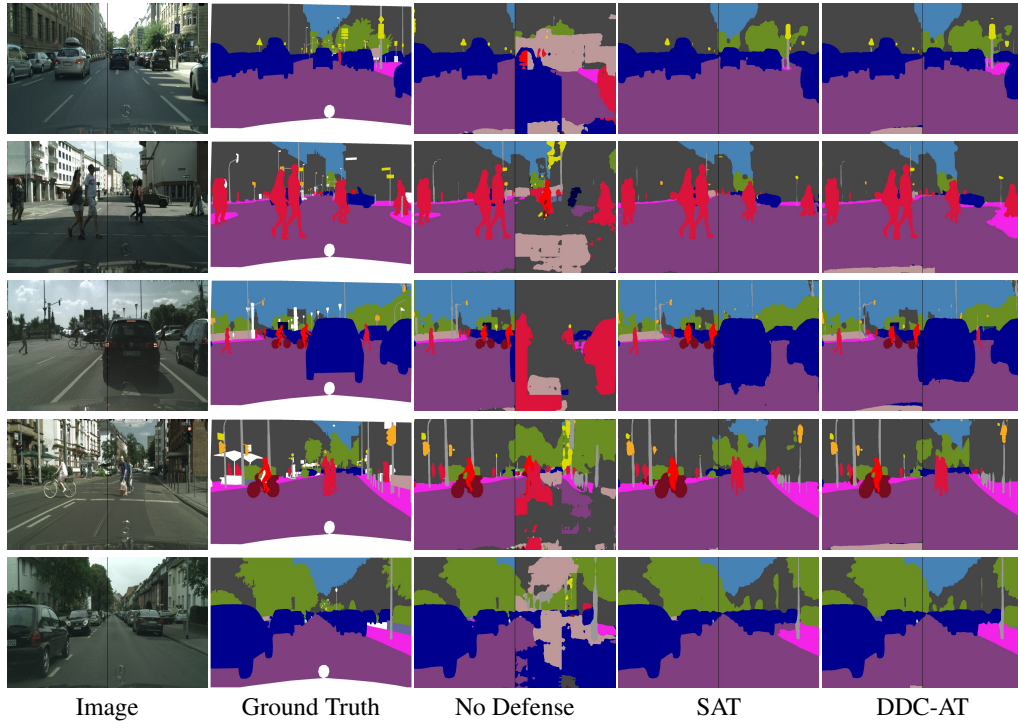


Figure 11: The visual comparison for different defense methods, which are executed on Cityscapes dataset with model structure as PSPNet.

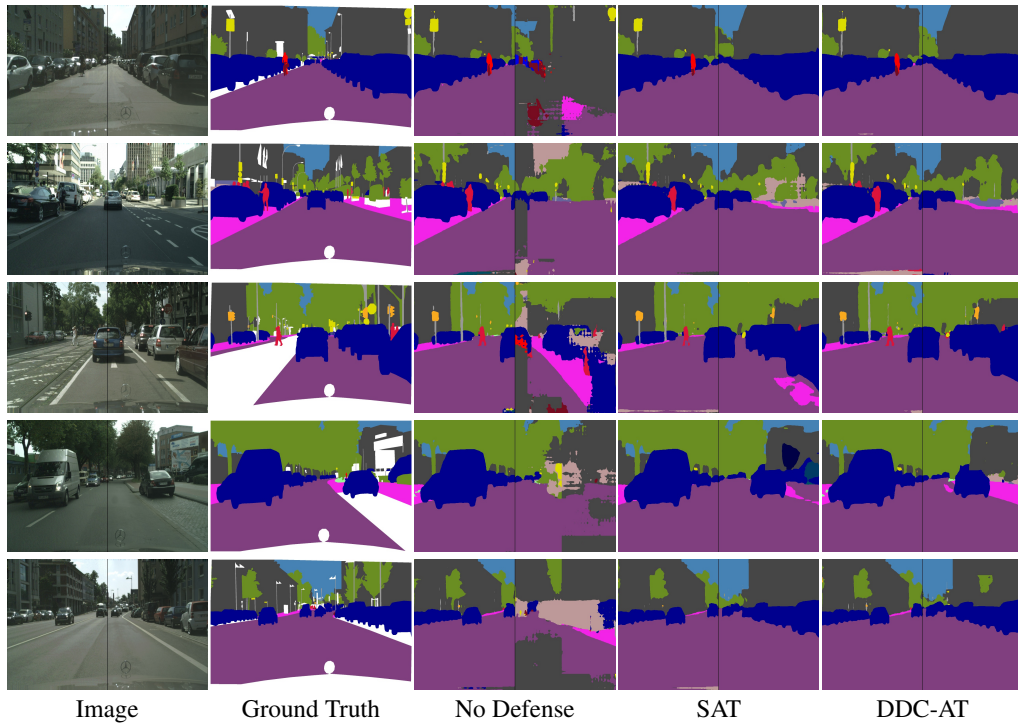


Figure 12: The visual comparison for different defense methods, which are executed on Cityscapes dataset with model structure as DeepLabv3.