

Encoder-Decoder with Atrous Separable Convolution for Semantic Image Segmentation

Liang-Chieh Chen, Yukun Zhu, George Papandreou, Florian Schroff, and
Hartwig Adam

Google Inc.

{lcchen, yukun, gpapan, fschroff, hadam}@google.com

Abstract. Spatial pyramid pooling module or encode-decoder structure are used in deep neural networks for semantic segmentation task. The former networks are able to encode multi-scale contextual information by probing the incoming features with filters or pooling operations at multiple rates and multiple effective fields-of-view, while the latter networks can capture sharper object boundaries by gradually recovering the spatial information. In this work, we propose to combine the advantages from both methods. Specifically, our proposed model, DeepLabv3+, extends DeepLabv3 by adding a simple yet effective decoder module to refine the segmentation results especially along object boundaries. We further explore the Xception model and apply the depthwise separable convolution to both Atrous Spatial Pyramid Pooling and decoder modules, resulting in a faster and stronger encoder-decoder network. We demonstrate the effectiveness of the proposed model on PASCAL VOC 2012 and Cityscapes datasets, achieving the test set performance of 89% and 82.1% without any post-processing. Our paper is accompanied with a publicly available reference implementation of the proposed models in Tensorflow at <https://github.com/tensorflow/models/tree/master/research/deeplab>.

Keywords: Semantic image segmentation, spatial pyramid pooling, encoder-decoder, and depthwise separable convolution.

1 Introduction

Semantic segmentation with the goal to assign semantic labels to every pixel in an image [1,2,3,4,5] is one of the fundamental topics in computer vision. Deep convolutional neural networks [6,7,8,9,10] based on the Fully Convolutional Neural Network [8,11] show striking improvement over systems relying on hand-crafted features [12,13,14,15,16,17] on benchmark tasks. In this work, we consider two types of neural networks that use spatial pyramid pooling module [18,19,20] or encoder-decoder structure [21,22] for semantic segmentation, where the former one captures rich contextual information by pooling features at different resolution while the latter one is able to obtain sharp object boundaries.

In order to capture the contextual information at multiple scales, DeepLabv3 [23] applies several parallel atrous convolution with different rates (called Atrous

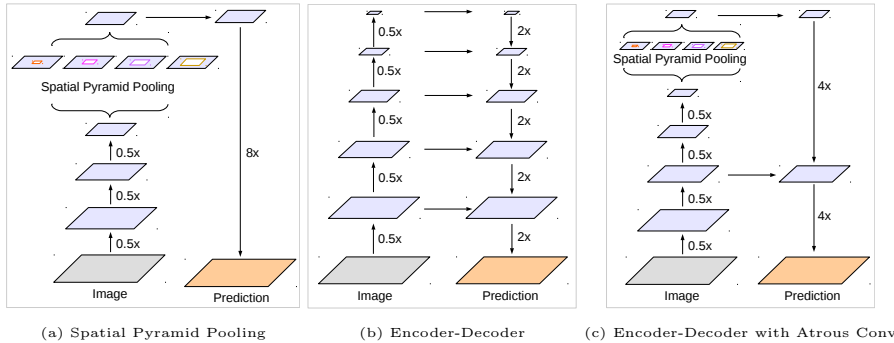


Fig. 1. We improve DeepLabv3, which employs the spatial pyramid pooling module (a), with the encoder-decoder structure (b). The proposed model, DeepLabv3+, contains rich semantic information from the encoder module, while the detailed object boundaries are recovered by the simple yet effective decoder module. The encoder module allows us to extract features at an arbitrary resolution by applying atrous convolution.

Spatial Pyramid Pooling, or ASPP), while PSPNet [24] performs pooling operations at different grid scales. Even though rich semantic information is encoded in the last feature map, detailed information related to object boundaries is missing due to the pooling or convolutions with striding operations within the network backbone. This could be alleviated by applying the atrous convolution to extract denser feature maps. However, given the design of state-of-art neural networks [7,9,10,25,26] and limited GPU memory, it is computationally prohibitive to extract output feature maps that are 8, or even 4 times smaller than the input resolution. Taking ResNet-101 [25] for example, when applying atrous convolution to extract output features that are 16 times smaller than input resolution, features within the last 3 residual blocks (9 layers) have to be dilated. Even worse, 26 residual blocks (78 layers!) will be affected if output features that are 8 times smaller than input are desired. Thus, it is computationally intensive if denser output features are extracted for this type of models. On the other hand, encoder-decoder models [21,22] lend themselves to faster computation (since no features are dilated) in the encoder path and gradually recover sharp object boundaries in the decoder path. Attempting to combine the advantages from both methods, we propose to enrich the encoder module in the encoder-decoder networks by incorporating the multi-scale contextual information.

In particular, our proposed model, called DeepLabv3+, extends DeepLabv3 [23] by adding a simple yet effective decoder module to recover the object boundaries, as illustrated in Fig. 1. The rich semantic information is encoded in the output of DeepLabv3, with atrous convolution allowing one to control the density of the encoder features, depending on the budget of computation resources. Furthermore, the decoder module allows detailed object boundary recovery.

Motivated by the recent success of depthwise separable convolution [27,28,26,29,30], we also explore this operation and show improvement in terms of both speed and accuracy by adapting the Xception model [26], similar to [31], for the task of

semantic segmentation, and applying the atrous separable convolution to both the ASPP and decoder modules. Finally, we demonstrate the effectiveness of the proposed model on PASCAL VOC 2012 and Cityscapes datasets and attain the test set performance of 89.0% and 82.1% without any post-processing, setting a new state-of-the-art.

In summary, our contributions are:

- We propose a novel encoder-decoder structure which employs DeepLabv3 as a powerful encoder module and a simple yet effective decoder module.
- In our structure, one can arbitrarily control the resolution of extracted encoder features by atrous convolution to trade-off precision and runtime, which is not possible with existing encoder-decoder models.
- We adapt the Xception model for the segmentation task and apply depthwise separable convolution to both ASPP module and decoder module, resulting in a faster and stronger encoder-decoder network.
- Our proposed model attains a new state-of-art performance on PASCAL VOC 2012 and Cityscapes datasets. We also provide detailed analysis of design choices and model variants.
- We make our Tensorflow-based implementation of the proposed model publicly available at <https://github.com/tensorflow/models/tree/master/research/deeplab>.

2 Related Work

Models based on Fully Convolutional Networks (FCNs) [8,11] have demonstrated significant improvement on several segmentation benchmarks [1,2,3,4,5]. There are several model variants proposed to exploit the contextual information for segmentation [12,13,14,15,16,17,32,33], including those that employ multi-scale inputs (*i.e.*, image pyramid) [34,35,36,37,38,39] or those that adopt probabilistic graphical models (such as DenseCRF [40] with efficient inference algorithm [41]) [42,43,44,37,45,46,47,48,49,50,51,39]. In this work, we mainly discuss about the models that use spatial pyramid pooling and encoder-decoder structure.

Spatial pyramid pooling: Models, such as PSPNet [24] or DeepLab [39,23], perform spatial pyramid pooling [18,19] at several grid scales (including image-level pooling [52]) or apply several parallel atrous convolution with different rates (called Atrous Spatial Pyramid Pooling, or ASPP). These models have shown promising results on several segmentation benchmarks by exploiting the multi-scale information.

Encoder-decoder: The encoder-decoder networks have been successfully applied to many computer vision tasks, including human pose estimation [53], object detection [54,55,56], and semantic segmentation [11,57,21,22,58,59,60,61,62,63,64]. Typically, the encoder-decoder networks contain (1) an encoder module that gradually reduces the feature maps and captures higher semantic information, and (2) a decoder module that gradually recovers the spatial information. Building on top of this idea, we propose to use DeepLabv3 [23] as the encoder module and add a simple yet effective decoder module to obtain sharper segmentations.

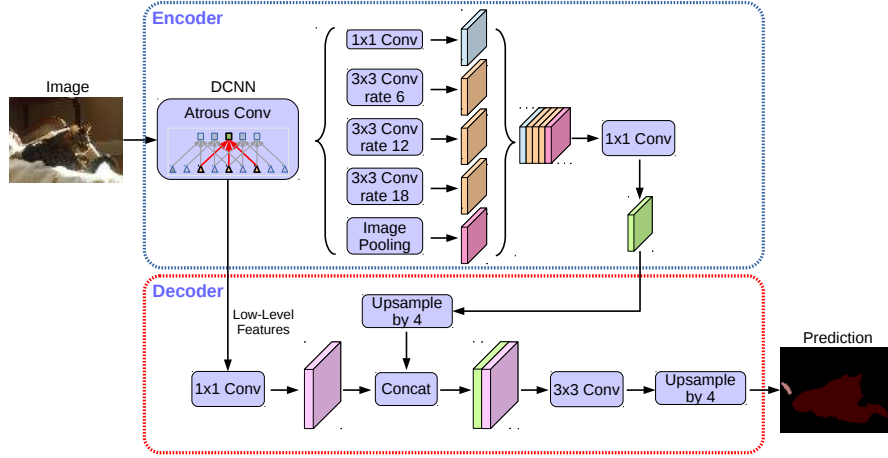


Fig. 2. Our proposed DeepLabv3+ extends DeepLabv3 by employing an encoder-decoder structure. The encoder module encodes multi-scale contextual information by applying atrous convolution at multiple scales, while the simple yet effective decoder module refines the segmentation results along object boundaries.

Depthwise separable convolution: Depthwise separable convolution [27, 28] or group convolution [7, 65], a powerful operation to reduce the computation cost and number of parameters while maintaining similar (or slightly better) performance. This operation has been adopted in many recent neural network designs [66, 67, 26, 29, 30, 31, 68]. In particular, we explore the Xception model [26], similar to [31] for their COCO 2017 detection challenge submission, and show improvement in terms of both accuracy and speed for the task of semantic segmentation.

3 Methods

In this section, we briefly introduce atrous convolution [69, 70, 8, 71, 42] and depthwise separable convolution [27, 28, 67, 26, 29]. We then review DeepLabv3 [23] which is used as our encoder module before discussing the proposed decoder module appended to the encoder output. We also present a modified Xception model [26, 31] which further improves the performance with faster computation.

3.1 Encoder-Decoder with Atrous Convolution

Atrous convolution: Atrous convolution, a powerful tool that allows us to explicitly control the resolution of features computed by deep convolutional neural networks and adjust filter’s field-of-view in order to capture multi-scale information, generalizes standard convolution operation. In the case of two-dimensional signals, for each location i on the output feature map y and a convolution filter w , atrous convolution is applied over the input feature map x as follows:

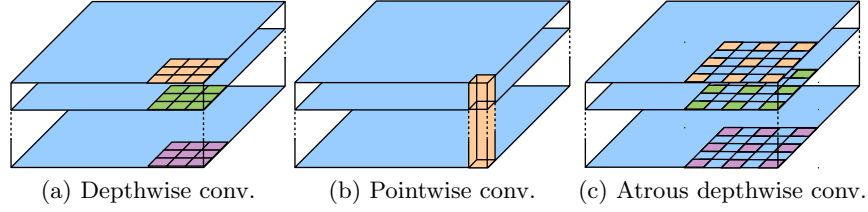


Fig. 3. 3×3 Depthwise separable convolution decomposes a standard convolution into (a) a depthwise convolution (applying a single filter for each input channel) and (b) a pointwise convolution (combining the outputs from depthwise convolution across channels). In this work, we explore *atrous separable convolution* where atrous convolution is adopted in the depthwise convolution, as shown in (c) with *rate* = 2.

$$y[i] = \sum_k x[i + r \cdot k]w[k] \quad (1)$$

where the atrous rate r determines the stride with which we sample the input signal. We refer interested readers to [39] for more details. Note that standard convolution is a special case in which rate $r = 1$. The filter’s field-of-view is adaptively modified by changing the rate value.

Depthwise separable convolution: Depthwise separable convolution, factorizing a standard convolution into a *depthwise convolution* followed by a *pointwise convolution* (*i.e.*, 1×1 convolution), drastically reduces computation complexity. Specifically, the depthwise convolution performs a spatial convolution independently for each input channel, while the pointwise convolution is employed to combine the output from the depthwise convolution. In the TensorFlow [72] implementation of depthwise separable convolution, atrous convolution has been supported in the depthwise convolution (*i.e.*, the spatial convolution), as illustrated in Fig. 3. In this work, we refer the resulting convolution as *atrous separable convolution*, and found that atrous separable convolution significantly reduces the computation complexity of proposed model while maintaining similar (or better) performance.

DeepLabv3 as encoder: DeepLabv3 [23] employs atrous convolution [69,70,8,71] to extract the features computed by deep convolutional neural networks at an arbitrary resolution. Here, we denote *output stride* as the ratio of input image spatial resolution to the final output resolution (before global pooling or fully-connected layer). For the task of image classification, the spatial resolution of the final feature maps is usually 32 times smaller than the input image resolution and thus *output stride* = 32. For the task of semantic segmentation, one can adopt *output stride* = 16 (or 8) for denser feature extraction by removing the striding in the last one (or two) block(s) and applying the atrous convolution correspondingly (*e.g.*, we apply *rate* = 2 and *rate* = 4 to the last two blocks respectively for *output stride* = 8). Additionally, DeepLabv3 augments the Atrous Spatial Pyramid Pooling module, which probes convolutional features at multiple scales by applying atrous convolution with different rates, with the image-level fea-

tures [52]. We use the last feature map before logits in the original DeepLabv3 as the encoder output in our proposed encoder-decoder structure. Note the encoder output feature map contains 256 channels and rich semantic information. Besides, one could extract features at an arbitrary resolution by applying the atrous convolution, depending on the computation budget.

Proposed decoder: The encoder features from DeepLabv3 are usually computed with $output\ stride = 16$. In the work of [23], the features are bilinearly upsampled by a factor of 16, which could be considered a naive decoder module. However, this naive decoder module may not successfully recover object segmentation details. We thus propose a simple yet effective decoder module, as illustrated in Fig. 2. The encoder features are first bilinearly upsampled by a factor of 4 and then concatenated with the corresponding low-level features [73] from the network backbone that have the same spatial resolution (*e.g.*, Conv2 before striding in ResNet-101 [25]). We apply another 1×1 convolution on the low-level features to reduce the number of channels, since the corresponding low-level features usually contain a large number of channels (*e.g.*, 256 or 512) which may outweigh the importance of the rich encoder features (only 256 channels in our model) and make the training harder. After the concatenation, we apply a few 3×3 convolutions to refine the features followed by another simple bilinear upsampling by a factor of 4. We show in Sec. 4 that using $output\ stride = 16$ for the encoder module strikes the best trade-off between speed and accuracy. The performance is marginally improved when using $output\ stride = 8$ for the encoder module at the cost of extra computation complexity.

3.2 Modified Aligned Xception

The Xception model [26] has shown promising image classification results on ImageNet [74] with fast computation. More recently, the MSRA team [31] modifies the Xception model (called Aligned Xception) and further pushes the performance in the task of object detection. Motivated by these findings, we work in the same direction to adapt the Xception model for the task of semantic image segmentation. In particular, we make a few more changes on top of MSRA’s modifications, namely (1) deeper Xception same as in [31] except that we do not modify the entry flow network structure for fast computation and memory efficiency, (2) all max pooling operations are replaced by depthwise separable convolution with striding, which enables us to apply *atrous separable convolution* to extract feature maps at an arbitrary resolution (another option is to extend the atrous algorithm to max pooling operations), and (3) extra batch normalization [75] and ReLU activation are added after each 3×3 depthwise convolution, similar to MobileNet design [29]. See Fig. 4 for details.

4 Experimental Evaluation

We employ ImageNet-1k [74] pretrained ResNet-101 [25] or modified aligned Xception [26,31] to extract dense feature maps by atrous convolution. Our implementation is built on TensorFlow [72] and is made publicly available.

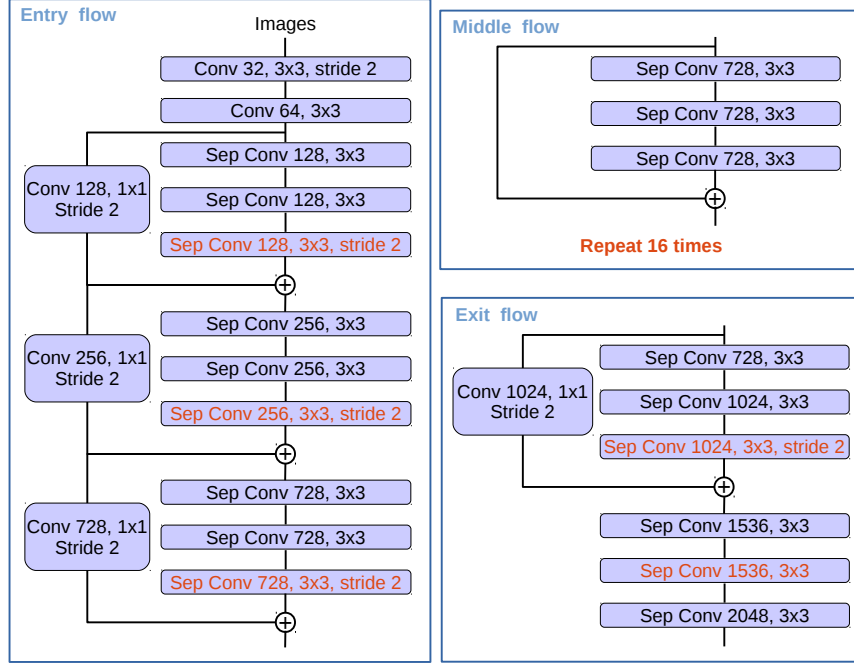


Fig. 4. We modify the Xception as follows: (1) more layers (same as MSRA’s modification except the changes in Entry flow), (2) all the max pooling operations are replaced by depthwise separable convolutions with striding, and (3) extra batch normalization and ReLU are added after each 3×3 depthwise convolution, similar to MobileNet.

The proposed models are evaluated on the PASCAL VOC 2012 semantic segmentation benchmark [1] which contains 20 foreground object classes and one background class. The original dataset contains 1,464 (*train*), 1,449 (*val*), and 1,456 (*test*) pixel-level annotated images. We augment the dataset by the extra annotations provided by [76], resulting in 10,582 (*trainaug*) training images. The performance is measured in terms of pixel intersection-over-union averaged across the 21 classes (mIOU).

We follow the same training protocol as in [23] and refer the interested readers to [23] for details. In short, we employ the same learning rate schedule (*i.e.*, “poly” policy [52] and same initial learning rate 0.007), crop size 513×513 , fine-tuning batch normalization parameters [75] when *output stride* = 16, and random scale data augmentation during training. Note that we also include batch normalization parameters in the proposed decoder module. Our proposed model is trained end-to-end without piecewise pretraining of each component.

4.1 Decoder Design Choices

We define “DeepLabv3 feature map” as the last feature map computed by DeepLabv3 (*i.e.*, the features containing ASPP features and image-level features), and $[k \times k, f]$ as a convolution operation with kernel $k \times k$ and f filters.

When employing *output stride* = 16, ResNet-101 based DeepLabv3 [23] bilinearly upsamples the logits by 16 during both training and evaluation. This simple bilinear upsampling could be considered as a naive decoder design, attaining the performance of 77.21% [23] on PASCAL VOC 2012 *val* set and is 1.2% better than not using this naive decoder during training (*i.e.*, downsampling groundtruth during training). To improve over this naive baseline, our proposed model “DeepLabv3+” adds the decoder module on top of the encoder output, as shown in Fig. 2. In the decoder module, we consider three places for different design choices, namely (1) the 1×1 convolution used to reduce the channels of the low-level feature map from the encoder module, (2) the 3×3 convolution used to obtain sharper segmentation results, and (3) what encoder low-level features should be used.

To evaluate the effect of the 1×1 convolution in the decoder module, we employ $[3 \times 3, 256]$ and the Conv2 features from ResNet-101 network backbone, *i.e.*, the last feature map in res2x residual block (to be concrete, we use the feature map before striding). As shown in Tab. 1, reducing the channels of the low-level feature map from the encoder module to either 48 or 32 leads to better performance. We thus adopt $[1 \times 1, 48]$ for channel reduction.

We then design the 3×3 convolution structure for the decoder module and report the findings in Tab. 2. We find that after concatenating the Conv2 feature map (before striding) with DeepLabv3 feature map, it is more effective to employ two 3×3 convolution with 256 filters than using simply one or three convolutions. Changing the number of filters from 256 to 128 or the kernel size from 3×3 to 1×1 degrades performance. We also experiment with the case where both Conv2 and Conv3 feature maps are exploited in the decoder module. In this case, the decoder feature map are gradually upsampled by 2, concatenated with Conv3 first and then Conv2, and each will be refined by the $[3 \times 3, 256]$ operation. The whole decoding procedure is then similar to the U-Net/SegNet design [21, 22]. However, we have not observed significant improvement. Thus, in the end, we adopt the very simple yet effective decoder module: the concatenation of the DeepLabv3 feature map and the channel-reduced Conv2 feature map are refined by two $[3 \times 3, 256]$ operations. Note that our proposed DeepLabv3+ model has *output stride* = 4. We do not pursue further denser output feature map (*i.e.*, *output stride* < 4) given the limited GPU resources.

4.2 ResNet-101 as Network Backbone

To compare the model variants in terms of both accuracy and speed, we report mIOU and Multiply-Adds in Tab. 3 when using ResNet-101 [25] as network backbone in the proposed DeepLabv3+ model. Thanks to atrous convolution, we

Channels	8	16	32	48	64
mIOU	77.61%	77.92%	78.16%	78.21%	77.94%

Table 1. PASCAL VOC 2012 *val* set. Effect of decoder 1×1 convolution used to reduce the channels of low-level feature map from the encoder module. We fix the other components in the decoder structure as using $[3 \times 3, 256]$ and Conv2.

Features Conv2	Conv3	3×3 Conv Structure	mIOU
✓		$[3 \times 3, 256]$	78.21%
✓		$[3 \times 3, 256] \times 2$	78.85%
✓		$[3 \times 3, 256] \times 3$	78.02%
✓		$[3 \times 3, 128]$	77.25%
✓		$[1 \times 1, 256]$	78.07%
✓	✓	$[3 \times 3, 256]$	78.61%

Table 2. Effect of decoder structure when fixing $[1 \times 1, 48]$ to reduce the encoder feature channels. We found that it is most effective to use the Conv2 (before striding) feature map and two extra $[3 \times 3, 256]$ operations. Performance on VOC 2012 *val* set.

Encoder train OS eval OS		Decoder MS Flip		mIOU	Multiply-Adds
16	16			77.21%	81.02B
16	8			78.51%	276.18B
16	8		✓	79.45%	2435.37B
16	8		✓ ✓	79.77%	4870.59B
16	16	✓		78.85%	101.28B
16	16	✓	✓	80.09%	898.69B
16	16	✓	✓ ✓	80.22%	1797.23B
16	8	✓		79.35%	297.92B
16	8	✓	✓	80.43%	2623.61B
16	8	✓	✓ ✓	80.57%	5247.07B
32	32			75.43%	52.43B
32	32	✓		77.37%	74.20B
32	16	✓		77.80%	101.28B
32	8	✓		77.92%	297.92B

Table 3. Inference strategy on the PASCAL VOC 2012 *val* set using *ResNet-101*. **train OS:** The *output stride* used during training. **eval OS:** The *output stride* used during evaluation. **Decoder:** Employing the proposed decoder structure. **MS:** Multi-scale inputs during evaluation. **Flip:** Adding left-right flipped inputs.

are able to obtain features at different resolutions during training and evaluation using a single model.

Model	Top-1 Error	Top-5 Error
Reproduced ResNet-101	22.40%	6.02%
Modified Xception	20.19%	5.17%

Table 4. *Single-model* error rates on ImageNet-1K validation set.

Baseline: The first row block in Tab. 3 contains the results from [23] showing that extracting denser feature maps during evaluation (*i.e.*, *eval output stride* = 8) and adopting multi-scale inputs increases performance. Besides, adding left-right flipped inputs doubles the computation complexity with only marginal performance improvement.

Adding decoder: The second row block in Tab. 3 contains the results when adopting the proposed decoder structure. The performance is improved from 77.21% to 78.85% or 78.51% to 79.35% when using *eval output stride* = 16 or 8, respectively, at the cost of about 20B extra computation overhead. The performance is further improved when using multi-scale and left-right flipped inputs.

Coarser feature maps: We also experiment with the case when using *train output stride* = 32 (*i.e.*, no atrous convolution at all during training) for fast computation. As shown in the third row block in Tab. 3, adding the decoder brings about 2% improvement while only 74.20B Multiply-Adds are required. However, the performance is always about 1% to 1.5% below the case in which we employ *train output stride* = 16 and different *eval output stride* values. We thus prefer using *output stride* = 16 or 8 during training or evaluation depending on the complexity budget.

4.3 Xception as Network Backbone

We further employ the more powerful Xception [26] as network backbone. Following [31], we make a few more changes, as described in Sec. 3.2.

ImageNet pretraining: The proposed Xception network is pretrained on ImageNet-1k dataset [74] with similar training protocol in [26]. Specifically, we adopt Nesterov momentum optimizer with momentum = 0.9, initial learning rate = 0.05, rate decay = 0.94 every 2 epochs, and weight decay $4e - 5$. We use asynchronous training with 50 GPUs and each GPU has batch size 32 with image size 299×299 . We did not tune the hyper-parameters very hard as the goal is to pretrain the model on ImageNet for semantic segmentation. We report the *single-model* error rates on the validation set in Tab. 4 along with the baseline reproduced ResNet-101 [25] under the same training protocol. We have observed 0.75% and 0.29% performance degradation for Top1 and Top5 accuracy when not adding the extra batch normalization and ReLU after each 3×3 depthwise convolution in the modified Xception.

The results of using the proposed Xception as network backbone for semantic segmentation are reported in Tab. 5.

Baseline: We first report the results without using the proposed decoder in the first row block in Tab. 5, which shows that employing Xception as network

backbone improves the performance by about 2% when *train output stride* = *eval output stride* = 16 over the case where ResNet-101 is used. Further improvement can also be obtained by using *eval output stride* = 8, multi-scale inputs during inference and adding left-right flipped inputs. Note that we do not employ the multi-grid method [77,78,23], which we found does not improve the performance.

Adding decoder: As shown in the second row block in Tab. 5, adding decoder brings about 0.8% improvement when using *eval output stride* = 16 for all the different inference strategies. The improvement becomes less when using *eval output stride* = 8.

Using depthwise separable convolution: Motivated by the efficient computation of depthwise separable convolution, we further adopt it in the ASPP and the decoder modules. As shown in the third row block in Tab. 5, the computation complexity in terms of Multiply-Adds is significantly reduced by 33% to 41%, while similar mIOU performance is obtained.

Pretraining on COCO: For comparison with other state-of-art models, we further pretrain our proposed DeepLabv3+ model on MS-COCO dataset [79], which yields about extra 2% improvement for all different inference strategies.

Pretraining on JFT: Similar to [23], we also employ the proposed Xception model that has been pretrained on both ImageNet-1k [74] and JFT-300M dataset [80,26,81], which brings extra 0.8% to 1% improvement.

Test set results: Since the computation complexity is not considered in the benchmark evaluation, we thus opt for the best performance model and train it with *output stride* = 8 and frozen batch normalization parameters. In the end, our ‘DeepLabv3+’ achieves the performance of 87.8% and 89.0% without and with JFT dataset pretraining.

Qualitative results: We provide visual results of our best model in Fig. 6. As shown in the figure, our model is able to segment objects very well without any post-processing.

Failure mode: As shown in the last row of Fig. 6, our model has difficulty in segmenting (a) sofa *vs.* chair, (b) heavily occluded objects, and (c) objects with rare view.

4.4 Improvement along Object Boundaries

In this subsection, we evaluate the segmentation accuracy with the trimap experiment [14,40,39] to quantify the accuracy of the proposed decoder module near object boundaries. Specifically, we apply the morphological dilation on ‘void’ label annotations on *val* set, which typically occurs around object boundaries. We then compute the mean IOU for those pixels that are within the dilated band (called trimap) of ‘void’ labels. As shown in Fig. 5 (a), employing the proposed decoder for both ResNet-101 [25] and Xception [26] network backbones improves the performance compared to the naive bilinear upsampling. The improvement is more significant when the dilated band is narrow. We have observed 4.8% and 5.4% mIOU improvement for ResNet-101 and Xception respectively at the

Encoder		Decoder	MS	Flip	SC	COCO	JFT	mIOU	Multiply-Adds
train OS	eval OS								
16	16							79.17%	68.00B
16	16		✓					80.57%	601.74B
16	16		✓	✓				80.79%	1203.34B
16	8							79.64%	240.85B
16	8		✓					81.15%	2149.91B
16	8		✓	✓				81.34%	4299.68B
16	16	✓						79.93%	89.76B
16	16	✓	✓					81.38%	790.12B
16	16	✓	✓	✓				81.44%	1580.10B
16	8	✓						80.22%	262.59B
16	8	✓	✓					81.60%	2338.15B
16	8	✓	✓	✓				81.63%	4676.16B
16	16	✓			✓			79.79%	54.17B
16	16	✓	✓	✓	✓			81.21%	928.81B
16	8	✓			✓			80.02%	177.10B
16	8	✓	✓	✓	✓			81.39%	3055.35B
16	16	✓			✓	✓		82.20%	54.17B
16	16	✓	✓	✓	✓	✓		83.34%	928.81B
16	8	✓			✓	✓		82.45%	177.10B
16	8	✓	✓	✓	✓	✓		83.58%	3055.35B
16	16	✓			✓	✓	✓	83.03%	54.17B
16	16	✓	✓	✓	✓	✓	✓	84.22%	928.81B
16	8	✓			✓	✓	✓	83.39%	177.10B
16	8	✓	✓	✓	✓	✓	✓	84.56%	3055.35B

Table 5. Inference strategy on the PASCAL VOC 2012 *val* set when using modified *Xception*. **train OS**: The *output stride* used during training. **eval OS**: The *output stride* used during evaluation. **Decoder**: Employing the proposed decoder structure. **MS**: Multi-scale inputs during evaluation. **Flip**: Adding left-right flipped inputs. **SC**: Adopting depthwise separable convolution for both ASPP and decoder modules. **COCO**: Models pretrained on MS-COCO. **JFT**: Models pretrained on JFT.

smallest trimap width as shown in the figure. We also visualize the effect of employing the proposed decoder in Fig. 5 (b).

4.5 Experimental Results on Cityscapes

In this section, we experiment DeepLabv3+ on the Cityscapes dataset [3], a large-scale dataset containing high quality pixel-level annotations of 5000 images (2975, 500, and 1525 for the training, validation, and test sets respectively) and about 20000 coarsely annotated images.

As shown in Tab. 7 (a), employing the proposed Xception model as network backbone (denoted as X-65) on top of DeepLabv3 [23], which includes the ASPP

Method	mIOU
Deep Layer Cascade (LC) [82]	82.7
TuSimple [77]	83.1
Large_Kernel_Matters [60]	83.6
Multipath-RefineNet [58]	84.2
ResNet-38_MS_COCO [83]	84.9
PSPNet [24]	85.4
IDW-CNN [84]	86.3
CASIA_IVA_SDN [63]	86.6
DIS [85]	86.8
DeepLabv3 [23]	85.7
DeepLabv3-JFT [23]	86.9
DeepLabv3+ (Xception)	87.8
DeepLabv3+ (Xception-JFT)	89.0

Table 6. PASCAL VOC 2012 *test* set results with top-performing models.

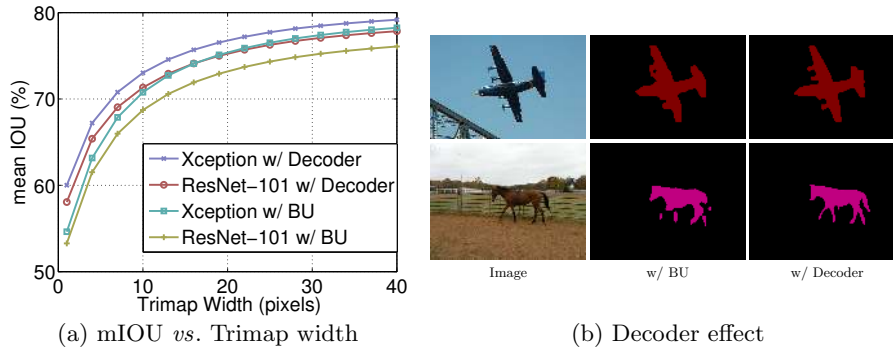


Fig. 5. (a) mIOU as a function of trimap band width around the object boundaries when employing *train output stride* = *eval output stride* = 16. **BU**: Bilinear upsampling. (b) Qualitative effect of employing the proposed decoder module compared with the naive bilinear upsampling (denoted as **BU**). In the examples, we adopt Xception as feature extractor and *train output stride* = *eval output stride* = 16.

module and image-level features [52], attains the performance of 77.33% on the validation set. Adding the proposed decoder module significantly improves the performance to 78.79% (1.46% improvement). We notice that removing the augmented image-level feature improves the performance to 79.14%, showing that in DeepLab model, the image-level features are more effective on the PASCAL VOC 2012 dataset. We also discover that on the Cityscapes dataset, it is effective to increase more layers in the entry flow in the Xception [26], the same as what [31] did for the object detection task. The resulting model building on top of the deeper network backbone (denoted as X-71 in the table), attains the best performance of 79.55% on the validation set.

After finding the best model variant on *val* set, we then further fine-tune the model on the coarse annotations in order to compete with other state-of-art

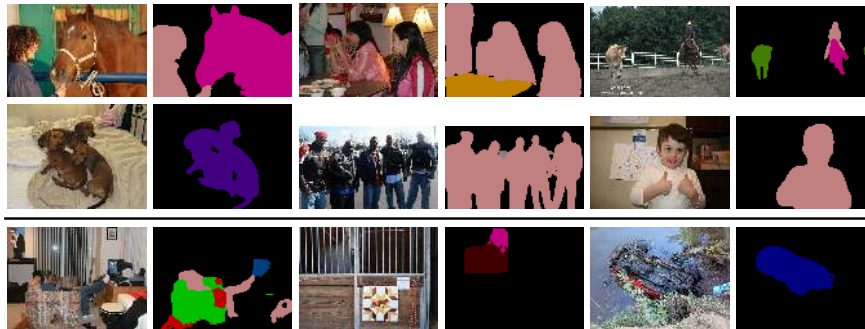


Fig. 6. Visualization results on *val* set. The last row shows a failure mode.

Backbone	Decoder	ASPP	Image-Level	mIOU
X-65		✓	✓	77.33
X-65	✓	✓	✓	78.79
X-65	✓	✓		79.14
X-71	✓	✓		79.55

(a) *val* set results

Method	Coarse	mIOU
ResNet-38 [83]	✓	80.6
PSPNet [24]	✓	81.2
Mapillary [86]	✓	82.0
DeepLabv3	✓	81.3
DeepLabv3+	✓	82.1

(b) *test* set results

Table 7. (a) DeepLabv3+ on the Cityscapes *val* set when trained with *train_fine* set. (b) DeepLabv3+ on Cityscapes *test* set. **Coarse:** Use *train_extra* set (coarse annotations) as well. Only a few top models are listed in this table.

models. As shown in Tab. 7 (b), our proposed DeepLabv3+ attains a performance of 82.1% on the test set, setting a new state-of-art performance on Cityscapes.

5 Conclusion

Our proposed model “DeepLabv3+” employs the encoder-decoder structure where DeepLabv3 is used to encode the rich contextual information and a simple yet effective decoder module is adopted to recover the object boundaries. One could also apply the atrous convolution to extract the encoder features at an arbitrary resolution, depending on the available computation resources. We also explore the Xception model and atrous separable convolution to make the proposed model faster and stronger. Finally, our experimental results show that the proposed model sets a new state-of-the-art performance on PASCAL VOC 2012 and Cityscapes datasets.

Acknowledgments We would like to acknowledge the valuable discussions with Haozhi Qi and Jifeng Dai about Aligned Xception, the feedback from Chen Sun, and the support from Google Mobile Vision team.

References

1. Everingham, M., Eslami, S.M.A., Gool, L.V., Williams, C.K.I., Winn, J., Zisserman, A.: The pascal visual object classes challenge a retrospective. *IJCV* (2014)
2. Mottaghi, R., Chen, X., Liu, X., Cho, N.G., Lee, S.W., Fidler, S., Urtasun, R., Yuille, A.: The role of context for object detection and semantic segmentation in the wild. In: *CVPR*. (2014)
3. Cordts, M., Omran, M., Ramos, S., Rehfeld, T., Enzweiler, M., Benenson, R., Franke, U., Roth, S., Schiele, B.: The cityscapes dataset for semantic urban scene understanding. In: *CVPR*. (2016)
4. Zhou, B., Zhao, H., Puig, X., Fidler, S., Barriuso, A., Torralba, A.: Scene parsing through ade20k dataset. In: *CVPR*. (2017)
5. Caesar, H., Uijlings, J., Ferrari, V.: COCO-Stuff: Thing and stuff classes in context. In: *CVPR*. (2018)
6. LeCun, Y., Bottou, L., Bengio, Y., Haffner, P.: Gradient-based learning applied to document recognition. In: *Proc. IEEE*. (1998)
7. Krizhevsky, A., Sutskever, I., Hinton, G.E.: Imagenet classification with deep convolutional neural networks. In: *NIPS*. (2012)
8. Sermanet, P., Eigen, D., Zhang, X., Mathieu, M., Fergus, R., LeCun, Y.: Overfeat: Integrated recognition, localization and detection using convolutional networks. In: *ICLR*. (2014)
9. Simonyan, K., Zisserman, A.: Very deep convolutional networks for large-scale image recognition. In: *ICLR*. (2015)
10. Szegedy, C., Liu, W., Jia, Y., Sermanet, P., Reed, S., Anguelov, D., Erhan, D., Vanhoucke, V., Rabinovich, A.: Going deeper with convolutions. In: *CVPR*. (2015)
11. Long, J., Shelhamer, E., Darrell, T.: Fully convolutional networks for semantic segmentation. In: *CVPR*. (2015)
12. He, X., Zemel, R.S., Carreira-Perpindn, M.: Multiscale conditional random fields for image labeling. In: *CVPR*. (2004)
13. Shotton, J., Winn, J., Rother, C., Criminisi, A.: Textonboost for image understanding: Multi-class object recognition and segmentation by jointly modeling texture, layout, and context. *IJCV* (2009)
14. Kohli, P., Torr, P.H., et al.: Robust higher order potentials for enforcing label consistency. *IJCV* **82**(3) (2009) 302–324
15. Ladicky, L., Russell, C., Kohli, P., Torr, P.H.: Associative hierarchical crfs for object class image segmentation. In: *ICCV*. (2009)
16. Gould, S., Fulton, R., Koller, D.: Decomposing a scene into geometric and semantically consistent regions. In: *ICCV*. (2009)
17. Yao, J., Fidler, S., Urtasun, R.: Describing the scene as a whole: Joint object detection, scene classification and semantic segmentation. In: *CVPR*. (2012)
18. Grauman, K., Darrell, T.: The pyramid match kernel: Discriminative classification with sets of image features. In: *ICCV*. (2005)
19. Lazebnik, S., Schmid, C., Ponce, J.: Beyond bags of features: Spatial pyramid matching for recognizing natural scene categories. In: *CVPR*. (2006)
20. He, K., Zhang, X., Ren, S., Sun, J.: Spatial pyramid pooling in deep convolutional networks for visual recognition. In: *ECCV*. (2014)
21. Ronneberger, O., Fischer, P., Brox, T.: U-net: Convolutional networks for biomedical image segmentation. In: *MICCAI*. (2015)
22. Badrinarayanan, V., Kendall, A., Cipolla, R.: Segnet: A deep convolutional encoder-decoder architecture for image segmentation. *PAMI* (2017)

23. Chen, L.C., Papandreou, G., Schroff, F., Adam, H.: Rethinking atrous convolution for semantic image segmentation. arXiv:1706.05587 (2017)
24. Zhao, H., Shi, J., Qi, X., Wang, X., Jia, J.: Pyramid scene parsing network. In: CVPR. (2017)
25. He, K., Zhang, X., Ren, S., Sun, J.: Deep residual learning for image recognition. In: CVPR. (2016)
26. Chollet, F.: Xception: Deep learning with depthwise separable convolutions. In: CVPR. (2017)
27. Sifre, L.: Rigid-motion scattering for image classification. PhD thesis (2014)
28. Vanhoucke, V.: Learning visual representations at scale. ICLR invited talk (2014)
29. Howard, A.G., Zhu, M., Chen, B., Kalenichenko, D., Wang, W., Weyand, T., Andreetto, M., Adam, H.: Mobilenets: Efficient convolutional neural networks for mobile vision applications. arXiv:1704.04861 (2017)
30. Zhang, X., Zhou, X., Lin, M., Sun, J.: Shufflenet: An extremely efficient convolutional neural network for mobile devices. In: CVPR. (2018)
31. Qi, H., Zhang, Z., Xiao, B., Hu, H., Cheng, B., Wei, Y., Dai, J.: Deformable convolutional networks – coco detection and segmentation challenge 2017 entry. ICCV COCO Challenge Workshop (2017)
32. Mostajabi, M., Yadollahpour, P., Shakhnarovich, G.: Feedforward semantic segmentation with zoom-out features. In: CVPR. (2015)
33. Dai, J., He, K., Sun, J.: Convolutional feature masking for joint object and stuff segmentation. In: CVPR. (2015)
34. Farabet, C., Couprie, C., Najman, L., LeCun, Y.: Learning hierarchical features for scene labeling. PAMI (2013)
35. Eigen, D., Fergus, R.: Predicting depth, surface normals and semantic labels with a common multi-scale convolutional architecture. In: ICCV. (2015)
36. Pinheiro, P., Collobert, R.: Recurrent convolutional neural networks for scene labeling. In: ICML. (2014)
37. Lin, G., Shen, C., van den Hengel, A., Reid, I.: Efficient piecewise training of deep structured models for semantic segmentation. In: CVPR. (2016)
38. Chen, L.C., Yang, Y., Wang, J., Xu, W., Yuille, A.L.: Attention to scale: Scale-aware semantic image segmentation. In: CVPR. (2016)
39. Chen, L.C., Papandreou, G., Kokkinos, I., Murphy, K., Yuille, A.L.: Deeplab: Semantic image segmentation with deep convolutional nets, atrous convolution, and fully connected crfs. TPAMI (2017)
40. Krähenbühl, P., Koltun, V.: Efficient inference in fully connected crfs with gaussian edge potentials. In: NIPS. (2011)
41. Adams, A., Baek, J., Davis, M.A.: Fast high-dimensional filtering using the permutohedral lattice. In: Eurographics. (2010)
42. Chen, L.C., Papandreou, G., Kokkinos, I., Murphy, K., Yuille, A.L.: Semantic image segmentation with deep convolutional nets and fully connected crfs. In: ICLR. (2015)
43. Bell, S., Upchurch, P., Snavely, N., Bala, K.: Material recognition in the wild with the materials in context database. In: CVPR. (2015)
44. Zheng, S., Jayasumana, S., Romera-Paredes, B., Vineet, V., Su, Z., Du, D., Huang, C., Torr, P.: Conditional random fields as recurrent neural networks. In: ICCV. (2015)
45. Liu, Z., Li, X., Luo, P., Loy, C.C., Tang, X.: Semantic image segmentation via deep parsing network. In: ICCV. (2015)
46. Papandreou, G., Chen, L.C., Murphy, K., Yuille, A.L.: Weakly- and semi-supervised learning of a dcnn for semantic image segmentation. In: ICCV. (2015)

47. Schwing, A.G., Urtasun, R.: Fully connected deep structured networks. arXiv:1503.02351 (2015)
48. Jampani, V., Kiefel, M., Gehler, P.V.: Learning sparse high dimensional filters: Image filtering, dense crfs and bilateral neural networks. In: CVPR. (2016)
49. Vemulapalli, R., Tuzel, O., Liu, M.Y., Chellappa, R.: Gaussian conditional random field network for semantic segmentation. In: CVPR. (2016)
50. Chandra, S., Kokkinos, I.: Fast, exact and multi-scale inference for semantic image segmentation with deep Gaussian CRFs. In: ECCV. (2016)
51. Chandra, S., Usunier, N., Kokkinos, I.: Dense and low-rank gaussian crfs using deep embeddings. In: ICCV. (2017)
52. Liu, W., Rabinovich, A., Berg, A.C.: Parsenet: Looking wider to see better. arXiv:1506.04579 (2015)
53. Newell, A., Yang, K., Deng, J.: Stacked hourglass networks for human pose estimation. In: ECCV. (2016)
54. Lin, T.Y., Dollár, P., Girshick, R., He, K., Hariharan, B., Belongie, S.: Feature pyramid networks for object detection. In: CVPR. (2017)
55. Shrivastava, A., Sukthankar, R., Malik, J., Gupta, A.: Beyond skip connections: Top-down modulation for object detection. arXiv:1612.06851 (2016)
56. Fu, C.Y., Liu, W., Ranga, A., Tyagi, A., Berg, A.C.: Dssd: Deconvolutional single shot detector. arXiv:1701.06659 (2017)
57. Noh, H., Hong, S., Han, B.: Learning deconvolution network for semantic segmentation. In: ICCV. (2015)
58. Lin, G., Milan, A., Shen, C., Reid, I.: Refinenet: Multi-path refinement networks with identity mappings for high-resolution semantic segmentation. In: CVPR. (2017)
59. Pohlen, T., Hermans, A., Mathias, M., Leibe, B.: Full-resolution residual networks for semantic segmentation in street scenes. In: CVPR. (2017)
60. Peng, C., Zhang, X., Yu, G., Luo, G., Sun, J.: Large kernel matters—improve semantic segmentation by global convolutional network. In: CVPR. (2017)
61. Islam, M.A., Roohan, M., Bruce, N.D., Wang, Y.: Gated feedback refinement network for dense image labeling. In: CVPR. (2017)
62. Wojna, Z., Ferrari, V., Guadarrama, S., Silberman, N., Chen, L.C., Fathi, A., Uijlings, J.: The devil is in the decoder. In: BMVC. (2017)
63. Fu, J., Liu, J., Wang, Y., Lu, H.: Stacked deconvolutional network for semantic segmentation. arXiv:1708.04943 (2017)
64. Zhang, Z., Zhang, X., Peng, C., Cheng, D., Sun, J.: Exfuse: Enhancing feature fusion for semantic segmentation. arXiv:1804.03821 (2018)
65. Xie, S., Girshick, R., Dollr, P., Tu, Z., He, K.: Aggregated residual transformations for deep neural networks. In: CVPR. (2017)
66. Jin, J., Dundar, A., Culurciello, E.: Flattened convolutional neural networks for feedforward acceleration. arXiv:1412.5474 (2014)
67. Wang, M., Liu, B., Foroosh, H.: Design of efficient convolutional layers using single intra-channel convolution, topological subdivision and spatial "bottleneck" structure. arXiv:1608.04337 (2016)
68. Zoph, B., Vasudevan, V., Shlens, J., Le, Q.V.: Learning transferable architectures for scalable image recognition. In: CVPR. (2018)
69. Holschneider, M., Kronland-Martinet, R., Morlet, J., Tchamitchian, P.: A real-time algorithm for signal analysis with the help of the wavelet transform. In: Wavelets: Time-Frequency Methods and Phase Space. (1989) 289–297
70. Giusti, A., Ciresan, D., Masci, J., Gambardella, L., Schmidhuber, J.: Fast image scanning with deep max-pooling convolutional neural networks. In: ICIP. (2013)

71. Papandreou, G., Kokkinos, I., Savalle, P.A.: Modeling local and global deformations in deep learning: Epitomic convolution, multiple instance learning, and sliding window detection. In: CVPR. (2015)
72. Abadi, M., Agarwal, A., et al.: Tensorflow: Large-scale machine learning on heterogeneous distributed systems. arXiv:1603.04467 (2016)
73. Hariharan, B., Arbeláez, P., Girshick, R., Malik, J.: Hypercolumns for object segmentation and fine-grained localization. In: CVPR. (2015)
74. Russakovsky, O., Deng, J., Su, H., Krause, J., Satheesh, S., Ma, S., Huang, Z., Karpathy, A., Khosla, A., Bernstein, M., Berg, A.C., Fei-Fei, L.: ImageNet Large Scale Visual Recognition Challenge. IJCV (2015)
75. Ioffe, S., Szegedy, C.: Batch normalization: Accelerating deep network training by reducing internal covariate shift. In: ICML. (2015)
76. Hariharan, B., Arbeláez, P., Bourdev, L., Maji, S., Malik, J.: Semantic contours from inverse detectors. In: ICCV. (2011)
77. Wang, P., Chen, P., Yuan, Y., Liu, D., Huang, Z., Hou, X., Cottrell, G.: Understanding convolution for semantic segmentation. arXiv:1702.08502 (2017)
78. Dai, J., Qi, H., Xiong, Y., Li, Y., Zhang, G., Hu, H., Wei, Y.: Deformable convolutional networks. In: ICCV. (2017)
79. Lin, T.Y., et al.: Microsoft COCO: Common objects in context. In: ECCV. (2014)
80. Hinton, G., Vinyals, O., Dean, J.: Distilling the knowledge in a neural network. In: NIPS. (2014)
81. Sun, C., Shrivastava, A., Singh, S., Gupta, A.: Revisiting unreasonable effectiveness of data in deep learning era. In: ICCV. (2017)
82. Li, X., Liu, Z., Luo, P., Loy, C.C., Tang, X.: Not all pixels are equal: Difficulty-aware semantic segmentation via deep layer cascade. In: CVPR. (2017)
83. Wu, Z., Shen, C., van den Hengel, A.: Wider or deeper: Revisiting the resnet model for visual recognition. arXiv:1611.10080 (2016)
84. Wang, G., Luo, P., Lin, L., Wang, X.: Learning object interactions and descriptions for semantic image segmentation. In: CVPR. (2017)
85. Luo, P., Wang, G., Lin, L., Wang, X.: Deep dual learning for semantic image segmentation. In: ICCV. (2017)
86. Bulò, S.R., Porzi, L., Kotschieder, P.: In-place activated batchnorm for memory-optimized training of dnns. In: CVPR. (2018)