

Pruning Parameterization with Bi-level Optimization for Efficient Semantic Segmentation on the Edge

Changdi Yang^{*1} Pu Zhao^{*1} Yanyu Li¹ Wei Niu² Jiexiong Guan²
 Hao Tang³ Minghai Qin¹ Bin Ren² Xue Lin¹ Yanzhi Wang¹

¹Northeastern University ²College of William & Mary ³CVL, ETH Zurich

{yang.changdi, zhao.pu, li.yanyu, xue.lin, yanzhi.wang}@northeastern.edu

{wniu, jguan}@email.wm.edu, bren@cs.wm.edu, hao.tang@vision.ee.ethz.ch, qinminghai@gmail.com

Abstract

With the ever-increasing popularity of edge devices, it is necessary to implement real-time segmentation on the edge for autonomous driving and many other applications. Vision Transformers (ViTs) have shown considerably stronger results for many vision tasks. However, ViTs with the full-attention mechanism usually consume a large number of computational resources, leading to difficulties for real-time inference on edge devices. In this paper, we aim to derive ViTs with fewer computations and fast inference speed to facilitate the dense prediction of semantic segmentation on edge devices. To achieve this, we propose a pruning parameterization method to formulate the pruning problem of semantic segmentation. Then we adopt a bi-level optimization method to solve this problem with the help of implicit gradients. Our experimental results demonstrate that we can achieve 38.9 mIoU on ADE20K val with a speed of 56.5 FPS on Samsung S21, which is the highest mIoU under the same computation constraint with real-time inference.

1. Introduction

Inspired by the extraordinary performance of Deep Neural Networks (DNNs), DNNs have been applied to various tasks. In this work, we focus on semantic segmentation, which aims to assign a class label to each pixel of an image to perform a dense prediction. It plays an important role in many real-world applications, such as autonomous driving. However, as a dense prediction task, segmentation models usually have complicated multi-scale feature fusion structures with large feature sizes, leading to tremendous memory and computation overhead with slow inference speed.

To reduce the memory and computation cost, certain lightweight CNN architectures [21, 40, 66] are designed for efficient segmentation. Besides CNNs, inspired by the recent superior performance of vision transformers (ViTs)

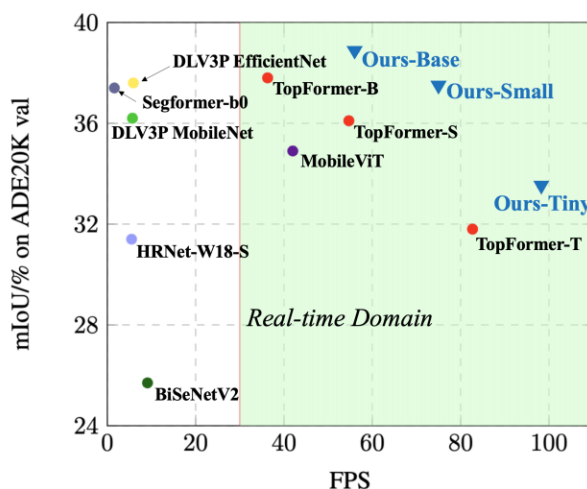


Figure 1. Comparison of accuracy versus FPS on ADE20K.

[17], some works [10, 11, 72] adopt ViTs in segmentation tasks to explore self-attention mechanism with the global receptive field. However, it is still difficult for ViTs to reduce the computation cost of the dense prediction for segmentation with large feature sizes.

With the wide spread of edge devices such as mobile phones, it is essential to perform real-time inference of segmentation on edge devices in practice. To facilitate mobile segmentation, the state-of-the-art work TopFormer [68] adopts a token pyramid transformer to produce scale-aware semantic features with tokens from various scales. It significantly outperforms CNN- and ViT-based networks across different semantic segmentation datasets and achieves a good trade-off between accuracy and latency. However, it only partially optimizes the token pyramid module, which costs most of the computations and latency.

In this work, we propose a pruning parameterization method with bi-level optimization to further enhance the performance of TopFormer. Our objective is to search for a suitable layer width for each layer in the token pyramid

^{*}These authors contributed equally to this work.

module, which is the main cost of computations and latency (over 60%). To achieve this, we first formulate the problem with pruning parameterization to build a pruning framework with a soft mask as a representation of the pruning policy. With this soft mask, we further adopt thresholding to convert the soft mask into a binary mask so that the model is trained with actual pruned weights to obtain pruning results directly. This is significantly different from other methods [27, 47] to train with unpruned small non-zero weights and use fine-tuning to mitigate the performance degradation after applying pruning. Besides, to update the soft mask as long as the pruning policy, we adopt a straight through estimator (STE) method to make the soft mask differentiable. Thus, we can build the pruning parameterization framework with minimal overhead.

Based on this framework, we need to search the best-suited layer width for each layer in the token pyramid module. It is non-trivial to perform the search. As the token pyramid module needs to extract multi-scale information from multiple spatial resolutions, the large hierarchical search space leads to difficulties of convergence. To resolve this problem, we adopt a bi-level optimization method. In the outer optimization, we try to obtain the pruning policy based on the pruning parameters (the soft mask). In the inner optimization, the optimized model weights with the best segmentation performance under this soft mask can be obtained. Compared with a typical pruning method, our work incorporates the implicit gradients with second-order derivatives to further guide the update of the soft mask and achieve better performance. Our experimental results demonstrate that we can achieve 38.9 mIoU (mean class-wise intersection-over-union) on the ADE20K dataset with a speed of 56.5 FPS on Samsung S21, which is the highest mIoU under the same computation constraint with real-time inference speed. As demonstrated in Figure 1, our models can achieve a better tradeoff between the mIoU and speed.

We summarize our contributions below,

- We propose a pruning parameterization method to build a pruning framework with a soft mask. We further use a threshold-based method to convert the soft mask into the binary mask to perform actual pruning during model training and inference. Besides, STE is adopted to update the soft mask efficiently through gradient descent optimizers.
- To solve the pruning problem formulated with the framework of pruning parameterization, we propose a bi-level optimization method to utilize implicit gradients for better results. We show that the second-order derivatives in the implicit gradients can be efficiently obtained through first-order derivatives, saving computations and memory.
- Our experimental results demonstrate that we can achieve the highest mIoU under the same computation constraint on various datasets. Specifically, we can achieve 38.9

mIoU on the ADE20K dataset with a real-time inference speed of 56.5 FPS on the Samsung S21.

2. Related Work

Real-Time Semantic Segmentation. Though semantic segmentation based on CNNs [6, 22, 33, 58, 71] can achieve great performance, it typically costs large amounts of computations with slow inference speed. Furthermore, with the ever-increasing popularity of edge devices such as mobile phones, it is necessary to achieve fast inference speed for semantic segmentation on edge devices [4]. Besides human designed lightweight models [21, 40, 67], neural architecture search (NAS) methods [1, 7, 42, 45, 69] are also adopted to search lightweight models.

As a pioneer in handcrafted real-time segmentation, ENet [50] designs a lightweight model for fast inference. DeepLabV3+ [6] adopts atrous separable convolution to reduce computation counts and uses the lightweight MobileNetV2 [56] as the backbone. BiSeNet [66, 67] and STDC [21] utilizes a two-branch architecture, where the deep branch extracts spatial information, and the shallow branch learns details. SFNet [40] uses flow alignment module to fuse context and spatial information.

Inspired by the recent success of NAS, some works automatically search lightweight segmentation models. To reduce the computation cost, FasterSeg [7] incorporates latency regularization during search. DCNAS [69] uses a densely connected search space and employs a gradient-based direct search method. NASViT [24] proposes a gradient projection algorithm to deal with the gradient conflict issues and improve the convergence performance. HR-NAS [15] keeps high-resolution representations in its encoder and the search process to maintain high accuracy in dense prediction. RTSeg [43] redesigns backbone and proposes a latency-driven search framework.

Although many lightweight segmentation models are developed, it is still hard for them to run real-time inference on resource- and power-limited GPUs of edge devices.

Vision Transformers. By utilizing the self-attention mechanism, ViTs [17] can achieve competitive results against CNN models in vision tasks. Inspired by the great success of ViTs, many research efforts are devoted to dense prediction tasks with ViTs on complex datasets. SETR [72] utilizes a ViT-based encoder to extract high-level semantic information. Instead of per-pixel prediction, MaskFormer [11] performs mask-based prediction with a customized backbone and a transformer-based decoder. With a transformer decoder to explore masked attention, mask2Former [10] proposes a universal architecture and achieves SOTA performance for various segmentation tasks. MobileViT [48] and MobileFormer [9] mix CNN and ViT in their architectures, but they are not fast enough for real-time inference on edge devices. TopFormer [68] uti-

lizes c CNN-based token pyramid module to extract multi-level tokens and lightweight ViT blocks to extract semantic information. It achieves low latency and high accuracy on complex datasets.

Neural Architecture Search and Pruning. NAS and pruning are commonly used to find lightweight models and reduce the computation overhead. Given a search space, NAS tries to identify a superior model automatically. Reinforcement learning (RL) based NAS [51, 75] and evolution-based NAS [18, 53] usually need to train and evaluate each candidate model, leading to tremendous searching cost. To mitigate this, gradient-based NAS [5, 12, 29, 46] is proposed to formulate a supernet with all candidate architectures and search the outstanding architecture through gradient descent methods. But it costs huge memory to incorporate all candidate architectures.

Network pruning is a compression technique to effectively reduce the DNN storage and computation cost [31]. In this work, we focus on structured pruning [39, 61] to remove entire filters or channels of CONV layers, which can be accelerated effectively for fast inference.

3. Problem Formulation with Pruning Parameterization

In our method, we first formulate the pruning problem with pruning parameterization. Previous pruning methods usually depend on the magnitudes of model weights and adopt the in-differentiable sorting operations, leading to inconsistent performance after applying pruning with sorting and additional overhead with fine-tuning. To mitigate this, we propose a pruning parameterization method, which uses a soft mask (rather than magnitudes of weights) to indicate whether to prune and get rid of sorting operations. With the STE method [3], we are able to represent the pruning with parameters and directly train the pruning like a typical model training. Then we formulate our problem with pruning parameterization and introduce our solution.

3.1. Soft Mask Construction

We first introduce how to construct the soft mask. To accelerate the inference, we adopt channel pruning to search for a suitable width for each convolution (CONV) layer. Specifically, we insert a depth-wise CONV layer following each CONV layer that is supposed to be pruned as below,

$$\mathbf{a}_l = \mathbf{s}_l \odot (\mathbf{w}_l \odot \mathbf{a}_{l-1}), \quad (1)$$

where $\odot$ denotes the convolution operation. $\mathbf{w}_l \in R^{o \times i \times k \times k}$ is the weight parameters in l -th CONV layer, with o output channels, i input channels, and kernels of size $k \times k$. $\mathbf{a}_l \in R^{n \times o \times t \times t'}$ represents the output features of the l -th layer, with o channels and $t \times t'$ feature size. n denotes batch size. $\mathbf{s}_l \in R^{o \times 1 \times 1 \times 1}$ is the weights of the depth-wise CONV layer. Each output channel of $\mathbf{w}_l \odot \mathbf{a}_{l-1}$

corresponds to one single element of $\mathbf{s}_l$. Thus $\mathbf{s}_l$ can serve as a soft mask or pruning indicator for the l -th CONV layer to indicate whether to prune the corresponding channels.

3.2. Forward and Backward Propagation

Since $\mathbf{s}_l$ is only a soft mask with continuous values, it can not represent the binary operation of pruning. To solve this, we adopt a threshold, and the forward pass with the mask is represented as

$$\mathbf{b}_l = \begin{cases} 1, & \mathbf{s}_l > \tau \\ 0, & \mathbf{s}_l \leq \tau \end{cases} \quad (\text{element-wise}), \quad \mathbf{a}_l = \mathbf{b}_l \odot (\mathbf{w}_l \odot \mathbf{a}_{l-1}), \quad (2)$$

where $\mathbf{b}_l \in \{0, 1\}^{o \times 1 \times 1 \times 1}$ is the binarized $\mathbf{s}_l$, and τ is a threshold which is simply set to 0.5 in our case. Each element of $\mathbf{b}_l$ corresponds to one output channel of $\mathbf{w}_l \odot \mathbf{a}_{l-1}$. In the forward pass, we first convert the soft mask into a binary mask through a threshold and then perform the depth-wise CONV to perform actual pruning. Thus the output channels corresponding to the zero elements in $\mathbf{b}_l$ are pruned, and the rest channels are preserved. We show the proof in Appendix A.

Advantages of the Binary Operation. With the binary operation to obtain the binary mask, the pruned weights become zero during computations of both training and inference. This is different from some pruning works [27, 30, 47] to perform training with small but non-zero weights and inference with binary weights, which has inconsistent performance between training and inference, and requires additional finetuning.

Why $\tau = 0.5$. Since the $\mathbf{s}_l$ parameters in Equation (2) are initialized randomly between 0 and 1, we set $\tau = 0.5$ as a threshold for the binary operation. τ can be set to other values (such as 0.3 or 0.6). The key point is that $\mathbf{s}_l$ can be updated to increase above τ to keep the corresponding channel or decrease below τ to prune the channel.

The binary operation in Equation (2) is non-differential, leading to difficulties for back-propagation. To solve this, we propose to use STE [3, 64] for back-propagation below,

$$\frac{\partial \mathcal{L}}{\partial \mathbf{s}_l} = \frac{\partial \mathcal{L}}{\partial \mathbf{b}_l}, \quad (3)$$

where we directly pass the gradients of $\mathbf{b}_l$ to $\mathbf{s}_l$ so that we can update the soft mask.

Why Named Pruning Parameterization. With the binarization and STE method, the pruning process can be represented with the soft mask $\mathbf{s} = \{\mathbf{s}_l\}$. We can update the soft mask to update the pruning policy based on its gradients. Thus $\mathbf{s}$ is the pruning parameters to denote and control pruning, i.e., pruning parameterization.

Difference with Other Pruning Works. Based on pruning parameterization, we decouple the pruning policy from model parameter magnitudes so that pruning does not further depend on the weight magnitudes. Unlike previous

works on pruning to force pruned weights to be or as close as to zeros [27, 30, 47], our method does not have such a constraint that pruned weights should be zero. Instead, once the corresponding binary mask is turned from 1 to 0, the information in pruned channels is preserved rather than zeroed out since zeros in b_l can block gradient flow to the corresponding weights. As a result, pruned channels are free to recover and contribute to accuracy if their corresponding elements in b_l switch from 0 to 1.

Difference with Other Mask Methods. Although some other works also adopt indicator/mask-based pruning such as [27, 28, 36, 37], our method is more straightforward and effective. For example, unlike our method to train the soft mask with STE directly, DAIS [27] relaxes the binarized channel indicators to be continuous. To bridge the non-negligible discrepancy between the continuous model and the target binarized model, it further uses an annealing-based procedure to steer the indicator convergence toward binarized states. Some works [36, 65] adopts batchnorm (BN) layers and uses the cumulative density function (CDF) of a Gaussian distribution as the mask variable, with a CDF-based loss function and the Gumbel-Softmax trick to update the mask at the cost of additional random sampling and complex gradient revision. Besides, the works [23, 65] create a mask to multiple the channels weights. This is different from our design with the depth-wise CONV operation for easy mask creation and direct mask training.

3.3. Training Loss with Pruning Parameterization

Based on the soft mask, we can train and prune the model with the following loss function,

$$\mathcal{L}_m(\mathbf{w}, \mathbf{s}) = \mathcal{L}(\mathbf{w}, \mathbf{s}) + \beta \cdot \mathcal{L}_{reg}(\mathbf{s}), \quad (4)$$

where $\mathcal{L}(\mathbf{w}, \mathbf{s})$ is the cross-entropy loss, and $\mathcal{L}_{reg}$ is the regularization term related to the sparsity or pruning ratio. For simplicity, we take Multiply-Accumulate operations (MACs) as the regularization rather than parameter number to estimate the on-device execution cost more precisely. β can weight the loss and stabilize training. $\mathcal{L}_{reg}$ can be defined as the squared ℓ_2 norm of the difference between current MACs and target MACs $\mathcal{C}$,

$$\mathcal{L}_{reg} = \left| \sum_l o'_l \times i_l \times t_l \times t'_l \times k^2 - \mathcal{C} \right|^2, \quad (5)$$

where o'_l is the number of remaining channels after pruning, and $t_l \times t'_l$ is the output feature size.

3.4. Problem Formulation

With pruning parameterization, the per-layer width search problem can be formulated as follows

$$\min_{\mathbf{s}} \quad \mathcal{L}_m(\mathbf{w}^*, \mathbf{s}), \quad (6)$$

$$\text{s.t.} \quad \mathbf{w}^* = \arg \min_{\mathbf{w}} \quad \mathcal{L}(\mathbf{w}, \mathbf{s}) + \frac{1}{2\lambda} \|\mathbf{w}\|_2^2. \quad (7)$$

It is a bi-level optimization problem [25, 35]. In the inner optimization of Equation (7), we optimize the model parameters under a given soft mask with a commonly used squared ℓ_2 norm as a regularization to mitigate the over-fitting problem. In the outer optimization of Equation (6), we optimize the soft mask to minimize the loss. Each time before updating $\mathbf{s}$, we first need to obtain $\mathbf{w}^*$.

4. Proposed Method with Bi-level Optimization

The objective is to find the soft mask (search a suitable width) for each layer to minimize the loss with optimized model weights. We adopt a bi-level pruning method to solve this problem. Compared with a typical gradient descent method to update the parameters with first-order derivatives, the bi-level optimization method incorporates implicit gradients with second-order derivatives to adjust the first-order term, leading to higher training efficiency with better convergence results. For easy of expression, since each channel has a corresponding mask, in this section, we broadcast the masks $\mathbf{s}$ and $\mathbf{b}$ so that the masks have the same dimension as the model weights $\mathbf{w}$.

4.1. Bi-level Optimization with Implicit Gradients

From the inner optimization, $\mathbf{w}^*$ is a function of $\mathbf{s}$ and different $\mathbf{s}$ can lead to different $\mathbf{w}^*$. Thus, to minimize $\mathcal{L}_m(\mathbf{w}^*, \mathbf{s})$ in Problem (6), we need to compute the gradients with reference to $\mathbf{s}$ as below,

$$\frac{d\mathcal{L}_m(\mathbf{w}^*, \mathbf{s})}{d\mathbf{s}} = \frac{d\mathbf{w}^*}{d\mathbf{s}} \nabla_{\mathbf{w}} \mathcal{L}_m(\mathbf{w}^*, \mathbf{s}) + \nabla_{\mathbf{s}} \mathcal{L}_m(\mathbf{w}^*, \mathbf{s}), \quad (8)$$

where $\nabla_{\mathbf{w}}$ and $\nabla_{\mathbf{s}}$ denote the partial derivatives of the loss function with reference to $\mathbf{w}$ and $\mathbf{s}$, respectively. $\frac{d\mathbf{w}^*}{d\mathbf{s}}$ represents the vector-wise full derivative, and we omit the transpose expression. Since $\mathbf{w}^*$ is implicitly defined as an optimization problem in Equation (7), $\frac{d\mathbf{w}^*}{d\mathbf{s}}$ is also known as implicit gradients [52, 54, 55, 70]. With $g(\mathbf{w}, \mathbf{s}) = \mathcal{L}(\mathbf{w}, \mathbf{s}) + \frac{1}{2\lambda} \mathbf{w}^T \mathbf{w}$, $\frac{d\mathbf{w}^*}{d\mathbf{s}}$ can be obtained through the following,

$$\frac{d\mathbf{w}^*}{d\mathbf{s}} = -\nabla_{\mathbf{s}\mathbf{w}}^2 g(\mathbf{w}^*, \mathbf{s}) \nabla_{\mathbf{w}}^2 g(\mathbf{w}^*, \mathbf{s})^{-1}, \quad (9)$$

where $\nabla_{\mathbf{s}\mathbf{w}}^2$ and $\nabla_{\mathbf{w}}^2$ are the second-order partial derivatives. We show the proof in Appendix B.

The Hessian matrix $\nabla_{\mathbf{w}}^2 g(\mathbf{w}^*, \mathbf{s})$ can be given by

$$\nabla_{\mathbf{w}}^2 g(\mathbf{w}^*, \mathbf{s}) = \nabla_{\mathbf{w}}^2 \mathcal{L}(\mathbf{w}^*, \mathbf{s}) + \frac{1}{\lambda} \mathbb{I} = \frac{1}{\lambda} \mathbb{I}, \quad (10)$$

where we adopt a Hessian-free approximation that $\nabla_{\mathbf{w}}^2 \mathcal{L}(\mathbf{w}^*, \mathbf{s}) = 0$ as DNNs usually have piece-wise linear decision boundary with ReLU functions. Thus, Equation (8) can be transformed to

$$\frac{d\mathcal{L}_m(\mathbf{w}^*, \mathbf{s})}{d\mathbf{s}} = \nabla_{\mathbf{s}} \mathcal{L}_m(\mathbf{w}^*, \mathbf{s}) - \lambda \nabla_{\mathbf{s}\mathbf{w}}^2 \mathcal{L}(\mathbf{w}^*, \mathbf{s}) \nabla_{\mathbf{w}} \mathcal{L}_m(\mathbf{w}^*, \mathbf{s}). \quad (11)$$

Compared with a typical gradient descent method, the bi-level optimization incorporates the second-order derivatives $\nabla_{sw}^2 \mathcal{L}(\mathbf{w}^*, \mathbf{s}) \nabla_w \mathcal{L}_m(\mathbf{w}^*, \mathbf{s})$ to adjust the first-order term $\nabla_s \mathcal{L}_m(\mathbf{w}^*, \mathbf{s})$.

It is difficult to obtain the second-order derivatives $\nabla_{sw}^2 \mathcal{L}(\mathbf{w}^*, \mathbf{s})$. Usually certain approximation methods such as finite difference [8, 20] may be adopted to save computation cost. But here we show that in this specific problem, we can directly obtain the analytical solution with first-order derivatives, which greatly saves computation efforts without approximation. Note that each channel has its corresponding mask, and we denote the masked channels as $\mathbf{w}_b = \mathbf{w}^* \cdot \mathbf{b}$ where $\cdot$ means the element-wise multiplication and $\mathbf{b}$ is defined in Equation (2). We can obtain that

$$\nabla_{sw}^2 \mathcal{L}(\mathbf{w}^*, \mathbf{s}) = \text{diag}(\nabla_{w_b} \mathcal{L}(\mathbf{w}_b)) \quad (12)$$

where $\text{diag}(\cdot)$ represents formulating a diagonal matrix with the diagonal vector. We show the proof in Appendix C. Then we can transform Equation (11) into the following,

$$\frac{d\mathcal{L}_m(\mathbf{w}^*, \mathbf{s})}{ds} = \nabla_s \mathcal{L}_m(\mathbf{w}^*, \mathbf{s}) - \lambda \nabla_{w_b} \mathcal{L}(\mathbf{w}_b) \cdot \nabla_w \mathcal{L}_m(\mathbf{w}^*, \mathbf{s}). \quad (13)$$

Thus although the implicit gradients incorporate second-order derivatives, it can be analytically expressed with first-order derivatives, greatly saving computation cost without any approximations.

In Equation (13), we need to create two copies $\mathbf{w}_b$ and $\mathbf{w}$ to obtain their derivatives, which are still not memory efficient. To further reduce the complexity, we can obtain $\nabla_{w_b} \mathcal{L}(\mathbf{w}_b)$ in the following,

$$\nabla_{w_b} \mathcal{L}(\mathbf{w}_b) = \begin{cases} \nabla_w \mathcal{L}_m(\mathbf{w}), & \mathbf{b} = 1 \\ 0, & \mathbf{b} = 0 \end{cases} \quad (\text{element-wise}) \quad (14)$$

The pruned weights ($\mathbf{b} = 0$) do not contribute to the loss, so their gradients are zero. The difference between $\mathcal{L}(\cdot)$ and $\mathcal{L}_m(\cdot)$ is just the regularization term $\mathcal{L}_{\text{reg}}(\cdot)$, which only relates to the sparsity and does not care the weight values. So $\nabla_{w_b} \mathcal{L}(\mathbf{w}_b)$ and $\nabla_w \mathcal{L}_m(\mathbf{w})$ are equal if their weights are not pruned ($\mathbf{b} = 1$). Combining Equation (13) and (14), we can obtain the following,

$$\frac{d\mathcal{L}_m(\mathbf{w}^*, \mathbf{s})}{ds} = \nabla_s \mathcal{L}_m(\mathbf{w}^*, \mathbf{s}) - \lambda \mathbf{b} \cdot [\nabla_w \mathcal{L}_m(\mathbf{w}^*, \mathbf{s})]^2. \quad (15)$$

We can see that there is no need to keep a copy and compute the gradients of $\mathbf{w}_b$, thus saving memory. During practical implementation, since $\mathbf{s}$ and $\mathbf{b}$ are the channel-wise masks, we accumulate the gradients of all weights in each channel to update the channel mask following Equation (15).

The Case without Implicit Gradients. In Equation (8), we can omit the first term with implicit gradients. Thus we have $\frac{d\mathcal{L}_m(\mathbf{w}^*, \mathbf{s})}{ds} = \nabla_s \mathcal{L}_m(\mathbf{w}^*, \mathbf{s})$ and there is no need to deal with the second-order term. It is easier to solve. But we demonstrate in the experiments that our method with implicit gradients can help to boost the performance without a significant increase of the computation cost in Section 5.4.

4.2. Bi-level Optimization Framework

In each iteration, our framework has two steps, including the model weights training step and the mask updating step. In the first step, we update the weights $\mathbf{w}$ with a few training steps for a fixed mask $\mathbf{s}$. Next in the second step, we update $\mathbf{s}$ with implicit gradients following Equation (15). Then we move on to the next iteration. In each step, we only update $\mathbf{w}$ or $\mathbf{s}$ without changing the other parameter. We discuss the advantages of the proposed method in the following.

No Need for Finetuning. After training with the proposed method, we can obtain the final layer-wise pruning policy following $\mathbf{s}$ and the sparse model. Since the model is already trained with the binary masks, it can achieve good segmentation performance without further fine-tuning.

First-Order Optimization. During the computation, we only adopt first-order optimization. Though we incorporate second-order derivatives in implicit gradients, we show that it can be analytically expressed with first-order derivatives in Equation (12), which greatly saves computation cost. We further optimize the computation process in Equation (15) to save memory cost.

Recoverable Contribution. During pruning, though some channels are pruned, their weights are not set to 0 due to the protection of the mask. When their mask is updated from 0 to 1, they can recover and contribute the accuracy.

5. Experiments

5.1. Datasets and Evaluation Metrics

ADE20K. ADE20K [73, 74] is a scene parsing dataset containing 25k images in the training set and 2k images in the validation set with 150 label classes.

Cityscapes. Cityscapes [14] is a dataset of urban street scenes from cars collected in 50 cities. It includes 5,000 finely annotated images, in which 2,975 images are used for training, 500 for validation, and 1,525 for testing. We exclude the extra training data with coarse labels throughout this paper. This dataset has 30 label classes, and 19 of them are used for segmentation. The resolution of images is 2048×1024 . Cityscapes dataset is an intensively studied benchmark for semantic segmentation, but it is challenging to perform real-time inference on such a high resolution.

Pascal VOC. PASCAL Visual Object Classes (VOC) 2012 [19] is a widely used dataset for semantic segmentation, classification, and object detection tasks. There are 1,464 images for training and 1,449 images for validation. We show the results on Pascal VOC in Table 4.

Evaluation Metrics. For semantic segmentation evaluation, we use the mean of class-wise intersection-over-union (mIoU) to measure the accuracy performance. We introduce the details of mIoU in Appendix D. For our results, we run our algorithm three times and report the mean and variance of the mIoU performance. For the baseline methods, some

Table 1. Comparison of our searched model and prior arts on the ADE20K val dataset. We compare with popular handcraft baselines in the first segment, NAS-based models in the second segment, pruning-based methods in the third segment and lightweight ViT-based models in the fourth segment. We measure the FPS on the Qualcomm Adreno 660 GPU of the Samsung Galaxy S21 mobile phone. Some FPS results are not available due to unsupported operations on the mobile device.

Category	Method	Backbone	Parameters	GMACs	FPS	val mIoU (%)
Human Design	PSPNet [71]	ResNet50-D8	49.1M	178.8	0.4	41.1
	DeepLabV3+ [6]	EfficientNet	17.1M	26.9	5.9	37.6
	BiSeNetV2 [66]	N.A.	3.34M	12.4	9.1	25.7
	SFNet [40]	ResNet-50	—	75.7	—	42.8
	HRNet-W18-S [60]	HRNet-W18-S	4.0M	10.2	5.5	31.4
NAS	HR-NAS-A [15]	Searched	2.5M	1.4	—	33.2
	HR-NAS-B [15]	Searched	3.9M	2.2	—	34.9
	NASViT [24]	Searched	—	2.5	—	37.9
	HRViT-b1 [26]	Searched	8.2M	14.6	—	45.9
Prune	EagleEye [38]	N.A.	3.4M	1.2	59.2	34.3
	DMCP [28]	N.A.	3.3M	1.2	63.8	33.9
	ResPep [16]	N.A.	3.3M	1.2	64.9	35.0
	CHEX [32]	N.A.	3.3M	1.2	64.2	35.2
ViT	Segmenter [57]	Searched	6.7M	4.6	4.1	39.9
	MobileViT [48]	Searched	3.9M	2.2	41.8	34.9
	SegFormer-B0 [63]	MiT-B0	3.8M	8.4	2.6	37.4
	TopFormer-Base [68]	N.A.	5.1M	1.8	36.3	37.8
	TopFormer-Small [68]	N.A.	3.1M	1.2	54.7	36.1
	TopFormer-Tiny [68]	N.A.	1.4M	0.6	82.7	31.8
Ours	Ours-Base	Searched	3.7M	1.8	56.5	38.9
	Ours-Small	Searched	3.3M	1.2	75.2	37.5
	Ours-Tiny	Searched	1.3M	0.7	98.0	33.5

methods cost too many resources and we can hardly rerun their experiments. Some methods provide their well-trained models, and we can test with the trained model.

5.2. Experimental Settings

Train Settings. We perform the pruning to search for a suitable width for each CONV layer in the TopFormer model. To enable a larger search space, we adopt the TopFormer architecture and use a larger per-layer width compared with the TopFormer-Base model. So our unpruned model has 4.1GMACs (with 39.9 mIoU), larger than the 1.8GMACs of the TopFormer-Base model.

We use stochastic gradient descent (SGD) optimizer, and momentum is set to 0.9, and set the batch size to 8 on each GPU. For ADE20K, the initial learning rate is set to 1.2×10^{-4} , and the “poly” learning rate policy is applied. For the Cityscapes dataset, the initial learning rate is set to 3×10^{-4} , and we apply the “poly” learning rate policy. For PASCAL VOC 2012, we set the initial learning rate as 0.01. Learning rate value is determined as $(1 - \frac{\text{iter}}{\text{total_iter}})^{0.9}$ where iter refers to the current iteration number. For the ADE20K dataset, we incorporate data augmentation by resizing with the random ratio between 0.5 and 2.0 as well as random flipping. On the Cityscapes dataset, multiple random scaling {0.5, 0.75, 1.0, 2.0} and fixed size cropping of 512×1024 are adopted for data augmentation. We choose the crop size for a better trade-off between mobile capacity and accuracy. We set hyperparameters $\beta = 0.01$ and

$\lambda = 0.1$ in the experiments.

Test Settings. Instead of multi-scale testing, we employ single scale testing for a fair comparison. For the ADE20K dataset, we use 512×512 as the input resolution. For the Cityscapes dataset, 512×1024 (rather than 1024×2048) is used as the inference resolution during testing for the following reasons. (i) In practice, we cannot use the resolution 1024×2048 since it causes memory overflow problems on our selected mobile phone. (ii) Besides, since the screens on edge devices such as mobile phones are not very large, the resolution of 512×1024 is good enough to serve on the small screens. (iii) Moreover, we find that this 512×1024 resolution can greatly speedup the inference on the mobile phones without significant accuracy degradation.

Experiment Environments. We train and prune the model using PyTorch 1.9 and CUDA 11.1 on 8 NVIDIA RTX TITAN GPUs. We measure the mobile latency on the GPU of an Samsung Galaxy S21 smartphone, with Qualcomm Snapdragon 888 mobile platform integrated with Qualcomm Kryo 680 Octa-core CPU and a Qualcomm Adreno 660 GPU. Note that for most baseline works, even the reduced resolution (512×1024) can cause an out-of-memory problem on the selected mobile device.

Compiler Framework on Mobile Devices. We need to compile the models before they can be executed on mobile devices. For TopFormer, we adopt the compiler TNN [13] to report the speed performance, which is also used in the original TopFormer paper. To further improve the inference speed, we adopt several compiler optimization methods and

Table 2. Our search results on the Cityscapes val dataset. We compare with popular handcraft baselines, NAS-based models, pruning based methods and lightweight ViT-based models. We measure the FPS on the Qualcomm Adreno 660 GPU of the Samsung Galaxy S21 mobile phone. Some FPS results are not available due to unsupported operations on the mobile device.

Category	Method	Backbone	Resolution	#params	GMACs	FPS	mIoU%
Human Design	ENet [50]	N.A.	512×1024	354.9K	5.9	—	58.5
	PSPNet [71]	ResNet101	1024×2048	68.07M	525.0	0.2	78.8
	BiSeNetV2 [66]	N.A.	512×1024	3.34M	24.6	5.0	73.4
	DeepLabV3+ [6]	MBv2	512×1024	2.26M	9.5	9.4	69.0
	STDC1-Seg50 [21]	STDC1	512×1024	12.05M	31.1	4.3	72.2
	STDC2-Seg50 [21]	STDC2	512×1024	16.08M	44.3	3.7	74.2
NAS	Auto-DeepLab-S [45]	N.A.	1024×2048	10.15M	333.3	—	79.7
	FasterSeg [7]	N.A.	1024×2048	—	28.2	—	73.1
	DCNAS [69]	N.A.	1024×2048	—	294.6	—	85.0
Prune	EagleEye [38]	N.A.	512×1024	3.6M	2.4	27.6	69.6
	DMCP [28]	N.A.	512×1024	3.5M	2.4	32.0	70.3
	ResPep [16]	N.A.	512×1024	3.5M	2.4	28.2	71.3
	CHEX [32]	N.A.	512×1024	3.4M	2.4	35.5	71.7
ViT	HRViT-b1 [26]	N.A.	512×1024	8.1M	28.2	—	81.6
	SegFormer-B0 [63]	MiT-B0	512×1024	3.8M	17.7	0.9	71.9
	TopFormer-Base [68]	N.A.	512×1024	5.1M	2.7	21.5	70.6
	TopFormer-Tiny [68]	N.A.	512×1024	1.4M	1.2	42.8	66.1
Ours	Ours-Base	Searched	512×1024	3.7M	3.6	30.8	74.7
	Ours-Small	Searched	512×1024	3.3M	2.4	38.7	73.6
	Ours-Tiny	Searched	512×1024	1.3M	1.4	52.6	71.5

develop an advanced compiler framework to compile our models and test their on-mobile speed. We show the details about our compiler optimization in Appendix E.

5.3. Experimental Results and Analysis

Segmentation Performance. Based on the TopFormer model, we obtain three models (Ours-Base, Ours-Small, and Ours-Tiny) with different computations. We show the comparison results on ADE20K in Table 1 and Cityscapes in Table 2. (i) For ADE20K, we can observe that the human-designed segmentation models and NAS-based models usually consume many computations (such as DeepLabV3+ with 26.9GMACs) in terms of MACs (multiply-accumulate operations). They can hardly achieve real-time inference on edge devices. Some NAS-based models with a small number of computations are not able to achieve high mIoU (such as HR-NAS-A with 33.2 mIoU). (ii) For the comparison with other pruning methods, we start from the same dense model and set the target GMACs to the same number (1.2GMACs) to make a fair comparison. We can see that our small model can achieve higher mIoU given the same GMACs. (iii) Compared with the ViT-based TopFormer, our models can achieve higher mIoUs with non-trivial improvements under the same computation budget (such as Ours-Small 37.5 mIoU v.s. 36.1 of TopFormer-Small under the same 1.2GMACs). Our models can achieve a faster inference speed (FPS higher than 50) on mobile phones compared with TopFormer models. We show the comparison with the baselines in terms of mIoU and FPS in Figure 1. We can achieve a better trade-off between mIoU and FPS.

(i) On the Cityscapes dataset, compared with handcrafted CNN-based segmentation models, including

BiSeNetV2 and STDC, our searched base model greatly reduces the GMACs (such as Ours-Base 3.6GMACs v.s. 24.6GMACs of BiSeNetV2) and achieves non-trivial better accuracy (Ours-Base 74.7 mIoU v.s. 73.4 mIoU of BiSeNetV2), as shown in Table 2. (ii) Compared with NAS-based methods such as FasterSeg, our models can achieve higher mIoU with much smaller computation costs. Other NAS methods consume too many computations (e.g., DCNAS with 300GMACs) to be deployed on edge devices. (iii) Compared with other pruning baselines, under the same computation budgets (2.4GMACs), our small model can achieve higher mIoU. (iv) Compared with transformer-based methods such as SegFormer-B0, similarly, our models achieve higher mIoU with less computations. Our small and tiny models have similar computations compared with TopFormer-Base and TopFormer-Tiny. But we can achieve higher mIoU.

Speed Performance on Mobile Devices. Our base model can achieve 56.5 FPS on the mobile device (Samsung Galaxy S21), which implements real-time execution with competitive segmentation performance, as shown in Table 1. Other baseline methods except TopFormer can hardly achieve real-time segmentation on edge devices, usually with FPS lower than 10. Our faster inference speed than TopFormer is achieved with the compiler optimization techniques, detailed in Appendix E.

Search Overhead. We show the search cost in Table 3. We only show the cost of our small model since our base, small and tiny models have similar search costs. It takes approximately 1.3 GPU days, which is smaller than the search cost of most other NAS-based segmentation methods. Our method can efficiently search out a compact model with

Table 3. Comparison of search cost on the Cityscapes val dataset.

Method	GPU Days	GMACs	mIoU
Auto-DeepLab [45]	3	695.0	82.1
GAS [44]	6.7	-	73.5
FasterSeg [7]	2	28.2	73.1
Fast-NAS [49]	8	435.7	78.9
SparseMask [62]	4.2	36.4	68.6
DCNAS [69]	5.6	294.6	85.0
LDP [34]	4.3	-	75.8
Without implicit gradients	1.1	2.4	71.9
Ours-Small	1.3	2.4	73.6

fewer computations and better segmentation performance. The low search cost is achieved by our pruning parameterization framework. Based on the soft mask and the low-cost thresholding and STE method, we are able to directly train the model weights and the pruning parameters. During the training, the pruned channels are zeroed out by the binary mask without the need of additional fine-tuning. Besides, our bi-level optimization can efficiently address the second-order derivatives with low computation complexity.

Visualization Comparison. We show the visualization comparison of our base model and other baselines in Appendix F. We can achieve better segmentation performance.

Results on Other Datasets and Model Architectures. We show the results on the Pascal VOC dataset in Table 4. We can observe that handcrafted CNN-based methods usually require more computational cost (such as DeepLabV3+ with 5.7 GMACs for MobileNetV2 backbone and 37.8 GMACs for ResNet50 backbone). Compared with ViT-based TopFormer, our method could achieve better mIoU under the same computational cost (such as Ours-Small 73.4% v.s. 69.8% mIoU of TopFormer-S for 1.2 GMACs).

Our method is general and can be applied to other model structures. We show the results for DeepLabV3+ [6] on Cityscapes in Table 5. Our two searched models have fewer parameters and computations with accelerated on-mobile inference speed, while achieving better mIoU compared with the original DeepLabV3+ model. Other models such as BiSeNetV2 or STDC cost much more computations.

5.4. Ablation Study

In our bi-level optimization, we incorporate implicit gradients in Equation (11). To demonstrate the advantages with implicit gradients, we consider the case without implicit gradients which omits the second term in Equation (11) with just $\frac{d\mathcal{L}_m(\mathbf{w}^*, \mathbf{s})}{d\mathbf{s}} = \nabla_{\mathbf{s}}\mathcal{L}_m(\mathbf{w}^*, \mathbf{s})$. We compare the performance of the solution without implicit gradients and our bi-level solution with implicit gradients on Cityscapes in Table 3. We can observe that our solution with implicit gradients has a search cost slightly higher than the solution without implicit gradients, demonstrating that our first-order solution for the second-order derivatives in Equation (12) can effectively save computation cost. For mIoU, our method

Table 4. Results on the PASCAL VOC 2012 test dataset. We compare our results with popular CNN-based models and lightweight ViT-based models.

Method	#params	GMACs	mIoU%	FPS
EfficientNet-B7 [59]	66.0M	194.0	85.2	0.1
EMANet [41]	10.0M	43.1	80.1	2.5
PSANet [2]	18.5M	56.3	78.5	1.4
DeepLabV3+ R101 [6]	43.9M	58.5	77.4	2.2
DeepLabV3+ R50 [6]	24.9M	37.8	76.3	3.1
DeepLabV3+ MBv2 [6]	2.3M	5.7	70.5	5.1
TopFormer-B [68]	5.1M	1.8	71.0	36.8
TopFormer-S [68]	3.1M	1.2	69.8	55.2
TopFormer-T [68]	1.4M	0.6	65.7	81.5
MobileViT-XXS [48]	1.9M	1.7	73.6	43.8
Ours-Base	3.7M	1.8	74.3	56.8
Ours-Small	3.3M	1.2	73.4	75.0
Ours-Tiny	1.3M	0.7	70.5	97.6

Table 5. Our searched results for DeepLabV3+ with MobileNetV2 backbone on Cityscapes val. The input resolution is 512×1024 .

Method	#params	GMACs	mIoU%	FPS
BiSeNetV2 [66]	3.34M	24.6	73.4	5.0
STDC1-Seg50 [21]	12.05M	31.1	72.2	4.3
STDC2-Seg50 [21]	16.08M	44.3	74.2	3.7
DeepLabV3+ [6]	2.26M	9.5	69.0	9.4
Ours-Base-DeepLab	1.21M	7.6	70.9	22.3
Ours-Small-DeepLab	569.0K	4.3	70.2	28.1

can achieve higher mIoU with non-trivial improvements, demonstrating that incorporating implicit gradients can effectively boost the performance.

Hyperparameter Tuning. We show the results with various β and λ in Appendix G. To compare with baselines under certain computations, we mainly show the results of our three sparse models (Ours-Base, Ours-Small, and Ours-Tiny). We show the results of more models under other computations in Appendix H.

6. Conclusion

We propose pruning parameterization with the thresholding and STE methods to build a pruning framework. Based on the framework, we formulate the problem and propose a bi-level optimization method with the implicit gradients. Our experimental results demonstrate that we can achieve the highest mIoU under the same computation constraint on various datasets. Specifically, we can achieve 38.9 mIoU on the ADE20K with a real-time inference speed of 56.5 FPS on the Samsung S21.

Acknowledgments

This work was supported in part by National Science Foundation (NSF) under the awards of CMMI-2125326, CNS-1932351, CCF-2047516 (CAREER) and CCF-2146873.

References

- [1] Haoli Bai, Hongda Mao, and Dinesh Nair. Dynamically pruning segformer for efficient semantic segmentation. In *ICASSP 2022-2022 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 3298–3302. IEEE, 2022. 2
- [2] Anjali Balagopal, Howard Morgan, Michael Dohopolski, Ramsey Timmerman, Jie Shan, Daniel F. Heitjan, Wei Liu, Dan Nguyen, Raquibul Hannan, Aurelie Garant, Neil Desai, and Steve Jiang. Psa-net: Deep learning based physician style-aware segmentation network for post-operative prostate cancer clinical target volume, 2021. 8
- [3] Yoshua Bengio, Nicholas Léonard, and Aaron C. Courville. Estimating or propagating gradients through stochastic neurons for conditional computation. *CoRR*, abs/1308.3432, 2013. 3
- [4] Han Cai, Ji Lin, Yujun Lin, Zhijian Liu, Haotian Tang, Hanrui Wang, Ligeng Zhu, and Song Han. Enable deep learning on mobile devices: Methods, systems, and applications. *ACM Transactions on Design Automation of Electronic Systems (TODAES)*, 27(3):1–50, 2022. 2
- [5] Han Cai, Ligeng Zhu, and Song Han. Proxylessnas: Direct neural architecture search on target task and hardware. *arXiv preprint arXiv:1812.00332*, 2018. 3
- [6] Liang-Chieh Chen, Yukun Zhu, George Papandreou, Florian Schroff, and Hartwig Adam. Encoder-decoder with atrous separable convolution for semantic image segmentation. In *Proceedings of the European conference on computer vision (ECCV)*, pages 801–818, 2018. 2, 6, 7, 8
- [7] Wuyang Chen, Xinyu Gong, Xianming Liu, Qian Zhang, Yuan Li, and Zhangyang Wang. Fasterseg: Searching for faster real-time semantic segmentation. *arXiv preprint arXiv:1912.10917*, 2019. 2, 7, 8
- [8] Xin Chen, Lingxi Xie, Jun Wu, and Qi Tian. Progressive darts: Bridging the optimization gap for nas in the wild. *International Journal of Computer Vision*, pages 1–18, 2020. 5
- [9] Yinpeng Chen, Xiyang Dai, Dongdong Chen, Mengchen Liu, Xiaoyi Dong, Lu Yuan, and Zicheng Liu. Mobileformer: Bridging mobilenet and transformer. *arXiv preprint arXiv:2108.05895*, 2021. 2
- [10] Bowen Cheng, Ishan Misra, Alexander G Schwing, Alexander Kirillov, and Rohit Girdhar. Masked-attention mask transformer for universal image segmentation. *arXiv preprint arXiv:2112.01527*, 2021. 1, 2
- [11] Bowen Cheng, Alex Schwing, and Alexander Kirillov. Pixel classification is not all you need for semantic segmentation. *Advances in Neural Information Processing Systems*, 34, 2021. 1, 2
- [12] Xiangxiang Chu, Bo Zhang, Ruijun Xu, and Jixiang Li. Fairnas: Rethinking evaluation fairness of weight sharing neural architecture search. *arXiv preprint arXiv:1907.01845*, 2019. 3
- [13] TNN Contributors. TNN: A high-performance, lightweight neural network inference framework. <https://github.com/Tencent/TNN>, 2019. 6
- [14] 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 *Proc. of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016. 5
- [15] Mingyu Ding, Xiaochen Lian, Linjie Yang, Peng Wang, Xiaojie Jin, Zhiwu Lu, and Ping Luo. Hr-nas: Searching efficient high-resolution neural architectures with lightweight transformers. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2021. 2, 6
- [16] Xiaohan Ding, Tianxiang Hao, Jianchao Tan, Ji Liu, Jungong Han, Yuchen Guo, and Guiguang Ding. Resrep: Lossless cnn pruning via decoupling remembering and forgetting. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 4510–4520, 2021. 6, 7
- [17] Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, Jakob Uszkoreit, and Neil Houlsby. An image is worth 16x16 words: Transformers for image recognition at scale. *ICLR*, 2021. 1, 2
- [18] Thomas Elsken, Jan Hendrik Metzen, and Frank Hutter. Efficient multi-objective neural architecture search via lamarckian evolution. *arXiv preprint arXiv:1804.09081*, 2018. 3
- [19] M. Everingham, L. Van Gool, C. K. I. Williams, J. Winn, and A. Zisserman. The pascal visual object classes (voc) challenge. *International Journal of Computer Vision*, 88(2):303–338, June 2010. 5
- [20] Alireza Fallah, Aryan Mokhtari, and Asuman Ozdaglar. On the convergence theory of gradient-based model-agnostic meta-learning algorithms. In *International Conference on Artificial Intelligence and Statistics*, pages 1082–1092. PMLR, 2020. 5
- [21] Mingyuan Fan, Shenqi Lai, Junshi Huang, Xiaoming Wei, Zhenhua Chai, Junfeng Luo, and Xiaolin Wei. Rethinking bisenet for real-time semantic segmentation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 9716–9725, 2021. 1, 2, 7, 8
- [22] Jun Fu, Jing Liu, Haijie Tian, Yong Li, Yongjun Bao, Zhiwei Fang, and Hanqing Lu. Dual attention network for scene segmentation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 3146–3154, 2019. 2
- [23] Shangqian Gao, Feihu Huang, Jian Pei, and Heng Huang. Discrete model compression with resource constraint for deep neural networks. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2020. 4
- [24] Chengyue Gong, Dilin Wang, Meng Li, Xinlei Chen, Zhicheng Yan, Yuandong Tian, Vikas Chandra, et al. Nasvit: Neural architecture search for efficient vision transformers with gradient conflict aware supernet training. In *International Conference on Learning Representations*, 2021. 2, 6
- [25] Stephen Gould, Basura Fernando, Anoop Cherian, Peter Anderson, Rodrigo Santa Cruz, and Edison Guo. On differentiating parameterized argmin and argmax problems with

- application to bi-level optimization. *CoRR*, abs/1607.05447, 2016. 4
- [26] Jiaqi Gu, Hyoukjun Kwon, Dilin Wang, Wei Ye, Meng Li, Yu-Hsin Chen, Liangzhen Lai, Vikas Chandra, and David Z Pan. Hrvit: Multi-scale high-resolution vision transformer. *arXiv preprint arXiv:2111.01236*, 2021. 6, 7
- [27] Yushuo Guan, Ning Liu, Pengyu Zhao, Zhengping Che, Kaigui Bian, Yanzhi Wang, and Jian Tang. Dais: Automatic channel pruning via differentiable annealing indicator search. *IEEE Transactions on Neural Networks and Learning Systems*, pages 1–12, 2022. 2, 3, 4
- [28] Shaopeng Guo, Yujie Wang, Quanquan Li, and Junjie Yan. Dmcp: Differentiable markov channel pruning for neural networks. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 1539–1547, 2020. 4, 6, 7
- [29] Zichao Guo, Xiangyu Zhang, Haoyuan Mu, Wen Heng, Zechun Liu, Yichen Wei, and Jian Sun. Single path one-shot neural architecture search with uniform sampling. In *European Conference on Computer Vision*, pages 544–560. Springer, 2020. 3
- [30] Yihui He, Xiangyu Zhang, and Jian Sun. Channel pruning for accelerating very deep neural networks. In *Proceedings of the IEEE international conference on computer vision*, pages 1389–1397, 2017. 3, 4
- [31] Torsten Hoeftler, Dan Alistarh, Tal Ben-Nun, Nikoli Dryden, and Alexandra Peste. Sparsity in deep learning: Pruning and growth for efficient inference and training in neural networks. *J. Mach. Learn. Res.*, 22(241):1–124, 2021. 3
- [32] Zejiang Hou, Minghai Qin, Fei Sun, Xiaolong Ma, Kun Yuan, Yi Xu, Yen-Kuang Chen, Rong Jin, Yuan Xie, and Sun-Yuan Kung. Chex: Channel exploration for cnn model compression. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 12287–12298, 2022. 6, 7
- [33] Junjie Huang, Zheng Zhu, and Guan Huang. Multi-stage hrnet: multiple stage high-resolution network for human pose estimation. *arXiv preprint arXiv:1910.05901*, 2019. 2
- [34] Lam Huynh, Esa Rahtu, Jiri Matas, and Janne Heikkilä. Fast neural architecture search for lightweight dense prediction networks. *arXiv preprint arXiv:2203.01994*, 2022. 8
- [35] Kaiyi Ji, Junjie Yang, and Yingbin Liang. Bilevel optimization: Convergence analysis and enhanced design. In *International Conference on Machine Learning*, pages 4882–4892. PMLR, 2021. 4
- [36] Minsoo Kang and Bohyung Han. Operation-aware soft channel pruning using differentiable masks. In *International Conference on Machine Learning*, pages 5122–5131. PMLR, 2020. 4
- [37] Jaedeok Kim, Chiyoun Park, Hyun-Joo Jung, and Yoonsuck Choe. Plug-in, trainable gate for streamlining arbitrary neural networks. *Proceedings of the AAAI Conference on Artificial Intelligence*, 34(04):4452–4459, Apr. 2020. 4
- [38] Bailin Li, Bowen Wu, Jiang Su, and Guangrun Wang. Eagleeye: Fast sub-net evaluation for efficient neural network pruning. In *European conference on computer vision*, pages 639–654. Springer, 2020. 6, 7
- [39] Tuanhui Li, Baoyuan Wu, Yujiu Yang, Yanbo Fan, Yong Zhang, and Wei Liu. Compressing convolutional neural networks via factorized convolutional filters. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 3977–3986, 2019. 3
- [40] Xiangtai Li, Ansheng You, Zhen Zhu, Houlong Zhao, Maoke Yang, Kuiyuan Yang, Shaohua Tan, and Yunhai Tong. Semantic flow for fast and accurate scene parsing. In *European Conference on Computer Vision*, pages 775–793. Springer, 2020. 1, 2, 6
- [41] Xia Li, Zhisheng Zhong, Jianlong Wu, Yibo Yang, Zhouchen Lin, and Hong Liu. Expectation-maximization attention networks for semantic segmentation, 2019. 8
- [42] Xin Li, Yiming Zhou, Zheng Pan, and Jiashi Feng. Partial order pruning: for best speed/accuracy trade-off in neural architecture search. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 9145–9153, 2019. 2
- [43] Yanyu Li, Changdi Yang, Pu Zhao, Geng Yuan, Wei Niu, Jiexiong Guan, Hao Tang, Minghai Qin, Qing Jin, Bin Ren, Xue Lin, and Yanzhi Wang. Towards real-time segmentation on the edge. In *Thirty-Seven AAAI Conference on Artificial Intelligence*, 2023. 2
- [44] Peiwen Lin, Peng Sun, Guangliang Cheng, Sirui Xie, Xi Li, and Jianping Shi. Graph-guided architecture search for real-time semantic segmentation, 2020. 8
- [45] Chenxi Liu, Liang-Chieh Chen, Florian Schroff, Hartwig Adam, Wei Hua, Alan L Yuille, and Li Fei-Fei. Auto-deeplab: Hierarchical neural architecture search for semantic image segmentation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 82–92, 2019. 2, 7, 8
- [46] Hanxiao Liu, Karen Simonyan, and Yiming Yang. Darts: Differentiable architecture search. *arXiv preprint arXiv:1806.09055*, 2018. 3
- [47] Zhuang Liu, Jianguo Li, Zhiqiang Shen, Gao Huang, Shoumeng Yan, and Changshui Zhang. Learning efficient convolutional networks through network slimming. In *Proceedings of the IEEE international conference on computer vision*, pages 2736–2744, 2017. 2, 3, 4
- [48] Sachin Mehta and Mohammad Rastegari. Mobilevit: lightweight, general-purpose, and mobile-friendly vision transformer. *arXiv preprint arXiv:2110.02178*, 2021. 2, 6, 8
- [49] Vladimir Nekrasov, Hao Chen, Chunhua Shen, and Ian Reid. Fast neural architecture search of compact semantic segmentation models via auxiliary cells, 2019. 8
- [50] Adam Paszke, Abhishek Chaurasia, Sangpil Kim, and Eugenio Culurciello. Enet: A deep neural network architecture for real-time semantic segmentation, 2016. 2, 7
- [51] Hieu Pham, Melody Y Guan, Barret Zoph, Quoc V Le, and Jeff Dean. Efficient neural architecture search via parameter sharing. *arXiv preprint arXiv:1802.03268*, 2018. 3
- [52] Aravind Rajeswaran, Chelsea Finn, Sham M Kakade, and Sergey Levine. Meta-learning with implicit gradients. *Advances in neural information processing systems*, 32, 2019. 4
- [53] Esteban Real, Alok Aggarwal, Yanping Huang, and Quoc V Le. Regularized evolution for image classifier architecture

- search. In *Proceedings of the aaai conference on artificial intelligence*, volume 33, pages 4780–4789, 2019. 3
- [54] Walter Rudin et al. *Principles of mathematical analysis*, volume 3. McGraw-hill New York, 1976. 4, 12
- [55] Kegan GG Samuel and Marshall F Tappen. Learning optimized map estimates in continuously-valued mrf models. In *2009 IEEE Conference on Computer Vision and Pattern Recognition*, pages 477–484. IEEE, 2009. 4
- [56] Mark Sandler, Andrew Howard, Menglong Zhu, Andrey Zhmoginov, and Liang-Chieh Chen. Mobilenetv2: Inverted residuals and linear bottlenecks. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 4510–4520, 2018. 2
- [57] Robin Strudel, Ricardo Garcia, Ivan Laptev, and Cordelia Schmid. Segmenter: Transformer for semantic segmentation. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 7262–7272, 2021. 6
- [58] Guolei Sun, Yun Liu, Hao Tang, Ajad Chhatkuli, Le Zhang, and Luc Van Gool. Mining relations among cross-frame affinities for video semantic segmentation. In *ECCV*, 2022. 2
- [59] Mingxing Tan and Quoc V Le. Efficientnet: Rethinking model scaling for convolutional neural networks. *arXiv preprint arXiv:1905.11946*, 2019. 8
- [60] Jingdong Wang, Ke Sun, Tianheng Cheng, Borui Jiang, Chaorui Deng, Yang Zhao, Dong Liu, Yadong Mu, Mingkui Tan, Xinggang Wang, et al. Deep high-resolution representation learning for visual recognition. *IEEE transactions on pattern analysis and machine intelligence*, 43(10):3349–3364, 2020. 6
- [61] Wei Wen, Chunpeng Wu, Yandan Wang, Yiran Chen, and Hai Li. Learning structured sparsity in deep neural networks. In *Advances in neural information processing systems (NeurIPS)*, pages 2074–2082, 2016. 3
- [62] Huikai Wu, Junge Zhang, and Kaiqi Huang. Sparsemask: Differentiable connectivity learning for dense image prediction, 2019. 8
- [63] Enze Xie, Wenhai Wang, Zhiding Yu, Anima Anandkumar, Jose M Alvarez, and Ping Luo. Segformer: Simple and efficient design for semantic segmentation with transformers. *Advances in Neural Information Processing Systems*, 34, 2021. 6, 7
- [64] Penghang Yin, Jiancheng Lyu, Shuai Zhang, Stanley Osher, Yingyong Qi, and Jack Xin. Understanding straight-through estimator in training activation quantized neural nets. *arXiv preprint arXiv:1903.05662*, 2019. 3
- [65] Zhonghui You, Kun Yan, Jinmian Ye, Meng Ma, and Ping Wang. Gate decorator: Global filter pruning method for accelerating deep convolutional neural networks. *Advances in Neural Information Processing Systems*, 32, 2019. 4
- [66] Changqian Yu, Changxin Gao, Jingbo Wang, Gang Yu, Chunhua Shen, and Nong Sang. Bisenet v2: Bilateral network with guided aggregation for real-time semantic segmentation. *International Journal of Computer Vision*, pages 1–18, 2021. 1, 2, 6, 7, 8
- [67] Changqian Yu, Jingbo Wang, Chao Peng, Changxin Gao, Gang Yu, and Nong Sang. Bisenet: Bilateral segmentation network for real-time semantic segmentation. In *Proceedings of the European Conference on Computer Vision (ECCV)*, September 2018. 2
- [68] Wenqiang Zhang, Zilong Huang, Guozhong Luo, Tao Chen, Xinggang Wang, Wenyu Liu, Gang Yu, and Chunhua Shen. Topformer: Token pyramid transformer for mobile semantic segmentation. *arXiv preprint arXiv:2204.05525*, 2022. 1, 2, 6, 7, 8
- [69] Xiong Zhang, Hongmin Xu, Hong Mo, Jianchao Tan, Cheng Yang, Lei Wang, and Wenqi Ren. Dcnas: Densely connected neural architecture search for semantic image segmentation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 13956–13967, 2021. 2, 7, 8
- [70] Yihua Zhang, Yuguang Yao, Parikshit Ram, Pu Zhao, Tianlong Chen, Mingyi Hong, Yanzhi Wang, and Sijia Liu. Advancing model pruning via bi-level optimization. In *Thirty-sixth Conference on Neural Information Processing Systems*, 2022. 4
- [71] Hengshuang Zhao, Jianping Shi, Xiaojuan Qi, Xiaogang Wang, and Jiaya Jia. Pyramid scene parsing network. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 2881–2890, 2017. 2, 6, 7
- [72] Sixiao Zheng, Jiachen Lu, Hengshuang Zhao, Xiatian Zhu, Zekun Luo, Yabiao Wang, Yanwei Fu, Jianfeng Feng, Tao Xiang, Philip HS Torr, et al. Rethinking semantic segmentation from a sequence-to-sequence perspective with transformers. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 6881–6890, 2021. 1, 2
- [73] Bolei Zhou, Hang Zhao, Xavier Puig, Sanja Fidler, Adela Barriuso, and Antonio Torralba. Scene parsing through ade20k dataset. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2017. 5
- [74] Bolei Zhou, Hang Zhao, Xavier Puig, Tete Xiao, Sanja Fidler, Adela Barriuso, and Antonio Torralba. Semantic understanding of scenes through the ade20k dataset. *International Journal of Computer Vision*, 127(3):302–321, 2019. 5
- [75] Barret Zoph and Quoc V. Le. Neural architecture search with reinforcement learning. In *International Conference on Learning Representations (ICLR)*, 2017. 3