
[Re] Learning to Learn By Self-Critique

Isac Arnekvist*

Division of Robotics, Perception and Learning
KTH Royal Institute of Technology
isacar@kth.se

Dmytro Kalpakchi*

Division of Speech, Music and Hearing
KTH Royal Institute of Technology
dmytroka@kth.se

Abstract

This work is a reproducibility study of the paper of Antoniou and Storkey [2019], published at NeurIPS 2019. Our results are in parts similar to the ones reported in the original paper, supporting the central claim of the paper that the proposed novel method, called Self-Critique and Adapt (SCA), improves the performance of MAML++. The conducted additional experiments on the Caltech-UCSD Birds 200 dataset confirm the superiority of SCA compared to MAML++. The ability to reproduce these results, without many of the implementation details, we believe is in favor of the proposed method. In addition, the reproduced paper suggests a novel high-end version of MAML++ for which we could not reproduce the same results. Our conclusion is that the reproduced paper is not self-contained, and we identify this as a potential reason of failure to reproduce all results.

1 Introduction

Humans are very good at learning new concepts by observing just a few examples of each one. Modern deep learning methods are also good at learning new concepts, but require much more examples to learn any concept, albeit a very simple one, e.g. distinguishing cats from dogs. In striving to bridge the gap between the established supervised learning paradigm and fast-learning humans, the paradigm of *few-shot learning* has emerged recently. In this new setting, the training data for each concept (or task) consists only of a few samples, called *shots* (hence the name). The aim of the few-shot learning is to learn a variety of tasks with a few shots each, instead of learning one task with many shots (as in classical supervised learning).

A popular way to approach few-shot learning is to frame it as *meta-learning* or *learning to learn*, i.e., a learning paradigm focused on acquiring across-task knowledge about *how* to learn, e.g. by learning parameter initializations, learning rate schedulers, optimizers, etc. Usually there are two models involved: *a base model* and *a meta model*. A base model learns task-specific information from a small labeled training set (*support set*) to predict on an unlabeled validation set (*target set*). A meta model learns task-agnostic information to produce parameters for a base model enabling the fastest possible fine-tuning for each task at hand.

In this work we reproduce the paper of Antoniou and Storkey [2019], which proposes a framework, called *Self-Critique and Adapt* or *SCA*, inspired by the idea that a target set also has task-specific information. For instance, if the training task is to distinguish cats from dogs and the new task would be to classify different breeds of dogs, a human would be able to guess what the new task is by observing a small number of samples. SCA aims at improving this ability by learning a label-free loss function during training, in order to be able to continue training the base model on the (unlabeled) test set before the final inference.

*Denotes equal contribution.

The paper has been reproduced by solely reading the details therein, and not taking the published code of the authors into account. It has become evident though, that the published code is essential for a complete understanding of the work.

2 Background

2.1 Meta learning

Meta learning has recently gained momentum after the publication of *Model Agnostic Meta Learning* (MAML) [Finn et al., 2017]. In short, MAML tries to optimize the initialization parameters θ of the base model, such that a network performs well on new few-shot learning tasks after only a few steps of training (see Figure 1). This approach, and its successors all have in common that a meta model is updated in an outer loop, while in an inner loop the task-specific base model is learned based on the meta model.

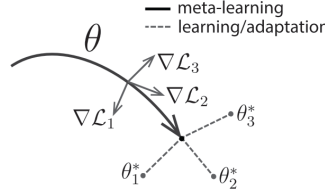


Figure 1: MAML diagram showing how base model parameters θ are optimized to quickly adapt to new tasks (taken from the original paper [Finn et al., 2017])

Versions of MAML include a first-order gradient version, denoted Reptile [Nichol and Schulman, 2018], and a range of improvements presented by Antoniou et al. [2019], denoted MAML++. In the latter, improvements were mainly focused at gradient instability, speedup of the first part of training by not using second-order gradients, learned step sizes, and tweaking Batch Normalization [Ioffe and Szegedy, 2015] for the meta-learning regime.

In short, given a task b , MAML++ updates the base model parameter vector $\theta_0^b = \theta_0$ by:

$$\theta_{i+1}^b = \theta_i^b - \alpha \nabla_{\theta} \mathcal{L}_b(\theta_i^b) \quad (1)$$

The loss for task b is given by $\mathcal{L}_b$. The parameter θ_0 is the parameter vector of the meta model, and is the common starting point, initialization, for all task-specific updates, as shown above. The meta model parameter vector is also updated, but in an outer loop, to minimize:

$$\mathcal{L}_{\text{meta}}(\theta_0) = \sum_{b=1}^B \sum_{i=1}^N v_i \mathcal{L}_b(\theta_i^b) \quad (2)$$

Here, B is the total number of tasks, and N is the total number of base model updates. The scalars v_i can be seen as importance weights which fulfill $v_i > 0$ and $\sum_{i=1}^N v_i = 1$. This is done to get more stable gradients by not only considering the last update, but throughout the entire update sequence. As training progresses, however, $v_N \rightarrow 1$ since the last update is ultimately what we care about and use during evaluation.

Other interpretations of this algorithm, rather than meta learning, could be transfer learning since the common features of all tasks are encoded in the meta model, invariant of the task-specific updates. Yet another interpretation is that we learn a parameter initialization rather than randomly initializing network weights from some simpler distribution.

2.2 The SCA algorithm

The paper we reproduce proposes a framework called Self-Critique and Adapt or SCA. The framework aims at using the information from the target set at inference time. This is achieved by learning a label-free loss function, parameterized as a neural network, called *Critic Network* (later referred

to as Critic). It is said that SCA can be applied for any meta-learning method that uses inner-loop optimization, but a specific example is given for the case of MAML (see Algorithm 1).

Algorithm 1: SCA Algorithm combined with MAML

```

1 Required: Base model function  $f$  and initialisation parameters  $\theta$ , critic network function  $C$  and
   parameters  $W$ , a batch of tasks  $\{\mathbf{S}^B = \{x_S^B, y_S^B\}, \mathbf{T}^B = \{x_T^B, y_T^B\}\}$  (where  $B$  is the number
   of tasks in our batch) and learning rates  $\alpha, \beta, \gamma$ 
2 for  $b$  in range( $B$ ) do
3    $\theta_0^b = \theta$ ; /* Reset  $\theta_0$  to the learned initialization parameters */
4   for  $i$  in range( $N$ ); /*  $N$  is a number of inner loop steps wrt support set */
5     do
6       /* Inner loop optimization wrt support set */
7        $\theta_{i+1}^b = \theta_i^b - \alpha \nabla_{\theta_i^b} L(f(x_S^b, \theta_i^b), y_S^b)$  (3)
8     end
9     for  $j$  in range( $I$ ); /*  $I$  is a number of inner loop steps wrt target set */
10      do
11        /* Critic feature-set collection */
12         $F = \{f(x_T^b, \theta_{N+j}^b), \theta_{N+j}^b, g(x_S^b, x_n)\}$  (4)
13        /* Inner loop optimization wrt target set */
14         $\theta_{N+j+1}^b = \theta_{N+j}^b - \gamma \nabla_{\theta_{N+j}^b} C(F, W)$  (5)
15      end
16       $L_{outer} = L_{outer} + L(f(x_T^b, \theta_{N+I}^b), y_T^b)$ 
17    end
18    /* Joint outer loop optimization of  $\theta$  */
19     $\theta = \theta - \beta \nabla_{\theta} L_{outer}$  (6)
20    /* Joint outer loop optimization of  $W$  */
21     $W = W - \beta \nabla_W L_{outer}$  (7)

```

The major difference between SCA and MAML is the introduction of another inner loop part (lines 8 to 12 of Algorithm 1), where the base model weights are optimized with respect to the target set using the Critic Network C as a loss function. The Critic operates on the collected feature set F (see Equation (4) of Algorithm 1), consisting of:

- $f(x_T^b, \theta_{N+j}^b)$ - predictions of the base model (trained wrt support set) on the target set T^b for the task b using the base model parameters θ_{N+j}^b ;
- θ_{N+j}^b - parameters of the base model for the task b after $N + j$ optimization steps;
- $g(x_S^b, x_n)$ - a task embedding, parameterized as a neural network, where x_S^b is the support set for the task b and x_n is not described in the original paper.

These features serve as the input to the Critic C , which outputs the label-free loss function, used in the inner loop to perform gradient descent on the parameters θ of the base model (see Equation (5) of Algorithm 1, where W are parameters of the Critic and γ is the Critic's learning rate). The loss for the outer loop L_{outer} , used in original MAML now uses the base model parameters θ_{N+I} after I updates using the Critic. The base model parameters θ (which will be used as initialization for the next task) and Critic's parameters W are updated in the outer loop using L_{outer} loss. There are two key observations to make about the outer loop:

- the base model parameters θ can be updated using the gradient of L_{outer} , since the parameters θ_{N+I}^b depend on θ used to initialize θ_0^b (line 3 of Algorithm 1);
- the Critic parameters W can be updated using the gradient of L_{outer} , since θ_{N+I}^b depends on W through gradient descent update (line 11 of Algorithm 1).

Critic Network Architecture. The components of the Critic Network are shown in the Figure 2. As mentioned previously, the feature set F serves as an input to the Critic. All features from F are reshaped into a batch of 1D vectors, concatenated on the feature dimension and then passed to a sequence of five one-dimensional dilated convolutions with kernel size 2 and 8 kernels per layer.

Each convolutional layer i (starting from 0) uses an exponentially increasing dilation policy with a dilation 2^i . Furthermore, DenseNet style connectivity is employed for convolutional layers, meaning that the input for each convolutional layer is a concatenation of outputs of all preceding convolutional layers. Finally a sequence of two fully-connected layers with ReLU activation functions is applied with the final fully-connected layer outputting a loss value.

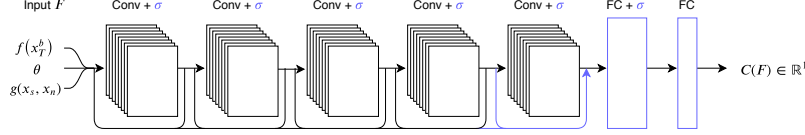


Figure 2: Critic Network architecture. The parts marked in blue correspond to parts which are not clear from the paper on how to implement. Other details that were unclear were padding or concatenation strategies and parameter initialization.

3 High-End MAML++

In order to test whether SCA provides significant performance improvement for high-capacity models, the authors have introduced a novel high generalization performance MAML++ backbone, which is dubbed as *High-End MAML++*. The main difference compared to the Low-End MAML++ is the use of a High-End Classifier (see architecture in Figure 3), instead of VGG network as the base model. The High-End classifier uses a DenseNet style architecture with 2 dense stages (purple blocks in Figure 3) and one transition layer in between them (light blue block in Figure 3). Each dense stage consists of two *dense block units* (see architecture of dense block unit in Figure 4). Each dense block unit consists of a bottleneck block, as described in [Huang et al., 2017], preceded by the squeeze-excite style convolutional attention, as described in [Hu et al., 2018].

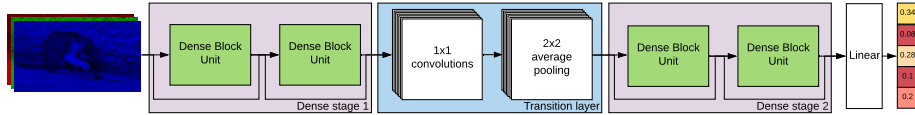


Figure 3: Architecture of the classifier in the High-End MAML++

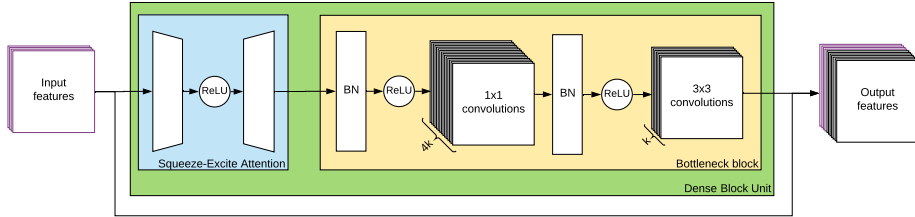


Figure 4: Architecture of the Dense block unit - a part of the classifier in the High-End MAML++. $k = 64$ is a growth rate; BN stands for Batch Normalization.

To improve the performance and the training speed of the High-End classifier, the authors propose to optimize only the last dense block unit and the final linear layer in the inner loop. All the other network components are shared across inner loop optimization steps by treating them as feature embedding. Hence, all other components, **but the last dense block unit and the final linear layer**, will be optimized in the outer loop.

4 Reproducibility

4.1 Critic architecture

The key concept of DenseNets is that features of all previous layers are concatenated as inputs to all successive layers. For this to work, we either need the features to have the same size, or to preprocess

them in some way such that the sizes agree. We decided to assume the sizes were kept constant by adding zero-padding. This padding will not be constant, however, since dilation changes per layer. This requires us to solve the equation

$$L = \left\lfloor \frac{L + 2 \times \text{padding} - \text{dilation}(\text{kernel_size} - 1) - 1}{\text{stride}} + 1 \right\rfloor \quad (8)$$

$$= \lfloor L + 2 \times \text{padding} - 2^i \rfloor \quad (9)$$

$$(10)$$

where it is stated by the authors that: i is the number of the layer from 0 to 4, kernel size is 2, and stride is 1. L is the size of the input and output. For the first layer, we have:

$$L = \lfloor L + 2 \times \text{padding} - 1 \rfloor \quad (11)$$

which has padding = 0.5 as the solution. We solved this by adding one additional zero last in the feature dimension. For the remaining layers, the equation have integer solutions which makes the implementation trivial.

For the last fully connected layers, two questions are raised:

1. Are all previous features from earlier layers propagated to the fully connected layer, or only to the last convolutional layer?
2. What is the feature dimension size after the first fully connected layer?

For (1), we noticed that if skip connections to the fully connected layer were not present, the gradient of the output w.r.t. the input was zero. This changed after adding the skip connections and was therefore used in the experiments. For (2), in our experiments, we assumed that the feature size was kept after the first fully connected layer, making the weight matrix of this layer a square matrix. Another question raised was which activation function was used. We decided to use ReLU throughout the entire critic Network.

4.2 Additional input features to the Critic

The critic’s input consists of three features: base model predictions on the target set, base model parameters used to produce these predictions and task embeddings produced by a neural network g . While supplying base model predictions as an input was not problematic, the other two input features were problematic to implement, because of the reasons, described below.

The embedding function g was only partly specified by highlighting that a relational network produced superior results compared to DenseNet-style network. However, neither the architecture of such relational net nor its training procedure (i.e. which loss function was used, whether it was trained in conjunction with the meta and base model as in [Rusu et al., 2019], etc) were explained. Due to the large amount of possible procedures for learning task embeddings, e.g., [Rusu et al., 2019, Hausman et al., 2018, Arnekvist et al., 2019], this part was not implemented.

Regarding passing base model parameters to the critic, the amount of required RAM becomes a problem. As an example, the base model in the case of Mini-ImageNet has roughly 70000 parameters. In this case, the weights of the first fully connected layer of the critic would use

$$(70000 \cdot (5 \cdot 8 + 1))^2 \cdot 32\text{bits} \approx 32\text{TB} \quad (12)$$

of GPU memory, which is obviously not reasonable. In our experiments for adding base model parameters as features, we tried different sizes of the last fully connected layers, but all ran out of memory on a GPU with 11GB RAM. Even when restricting the output dimension of the first fully connected layer to 256, it still sums up to about 3GB. Note, that this does not include the other parameters of the critic, the base/meta model, multiple forward activations, gradients, etc.

4.3 Hyperparameters for Low-End MAML++

In Algorithm 1, there are three step sizes listed: α , β , and γ . We used the MAML++ approach where all the upgrades to the base model, including those made by the critic, are done with vanilla SGD and learned step sizes. The critic’s parameters were updated with SGD and a relatively small step size

10^{-6} . The reasoning was to use a relatively small step size and get slow learning, rather than using a large step size and risk unstable learning.

Batch sizes were 2 for the 1-shot experiments, as in MAML++, but 1 for 5-shot, since otherwise we ran out of memory. In MAML++, it is proposed that Batch Normalization only normalizes data using the exponential moving averages, which implies that we can use batch size one. As can be seen in section 5, despite the changed batch size, results for Low-End MAML++ still turn out the same, or better than those reported in the original paper.

For the critic, we used standard initializations of layers defined in PyTorch [Paszke et al., 2017], i.e., $w \sim \mathcal{U}(-\sqrt{k}, \sqrt{k})$ for the biases and weights of the convolutional layers, where k^{-1} is the number channels in times the kernel size. For the fully connected layers, we used the same initialization but where k^{-1} is the number of input features (fan-in).

4.4 SCA for Low-End MAML++

In SCA, we perform I additional updates using the critic, after the N inner loop updates using a pre-defined loss function, usually negative log-likelihood. The simplest guess how to use MAML++ along with SCA is by optimizing

$$\mathcal{L}_{\text{meta}}(\theta_0) = \sum_{b=1}^B \left[\sum_{i=1}^N v_i \mathcal{L}_b(\theta_i^b) + C(f(x_T^b, \theta_{N+I}^b)) \right]. \quad (13)$$

This can be extended in multi-step loss fashion (similar to MAML++), as follows:

$$\mathcal{L}_{\text{meta}}(\theta_0) = \sum_{b=1}^B \left[\sum_{i=1}^N v_i \mathcal{L}_b(\theta_i^b) + \sum_{j=1}^I w_j C(f(x_T^b, \theta_{N+j}^b)) \right], \quad (14)$$

where w_j are importance weights as v_i .

For the exact implementation details, the only information about the value of I was given as an example in the Section 4 of the original paper, where $I = 1$. For this reason, we used this value in our experiments, and then also the choice of importance weights w_j for the critic losses becomes obsolete.

4.5 Hyperparameters for High-End MAML++

For most of the implementation details, the reader is referred to other papers [Huang et al., 2017, Hu et al., 2018] and we had to assume that all details are the same as in these papers. Learning rate and optimizer for the High-End classifier are not specified, so we assumed that they should be the same as for Low-End MAML++.

However, when we were unable to reproduce the reported results with the assumed hyperparameters, we had to experiment further and tried different weight initialization strategies, namely Kaiming uniform and normal [He et al., 2015], Xavier [Glorot and Bengio, 2010], as well as different optimizers, namely Adam optimizer [Kingma and Ba, 2014] and SGD. We noticed that these choices are crucial for the learning outcome, and after manual search we found SGD with learning rate 10^{-4} to be *superior* to Adam. For weight initializations, we chose to follow the implementation of MAML++ with Xavier initialization and zero bias in all layers except the last linear layer of the classifier. If using Xavier in the last linear layer, interestingly, the learning immediately diverges. We instead found that the PyTorch [Paszke et al., 2017] default initialization produced stable learning, but could not find any papers confirming this particular choice.

5 Results

We list our results in Table 1. For the original paper no results were reported for the Caltech-UCSD Birds 200 (CUBS-200) dataset [Welinder et al., 2010] using *Low-End MAML++* (which is the original MAML++ proposed in [Antoniou et al., 2019]), while we provide our results for this dataset here. The reported 95% confidence intervals are calculated using the standard error times 1.96. Standard deviations were estimated using the resulting accuracies of the 3 experiments we ran, changing only

the seed for each of them. Note, that in the original paper, neither a number of experimental runs, nor the way of calculating the confidence intervals were reported.

The reported results for SCA are deemed *significant* if the performance confidence intervals of MAML++ with SCA and without SCA do not overlap. Performance improvements that are considered significant are marked in bold in Table 1. Note, that the results are compared only *within each implementation*, meaning that the reproduced SCA performance is compared to only the reproduced MAML++ performance and *not* the one reported in the original paper.

It should be mentioned that we have also tried running SCA with High-End MAML++, but all experiments ran out of memory.

Table 1: SCA reproducibility results (all SCA models are applied on top of Low-End MAML++, the values in bold indicated significant improvements compared to the the base model)

Model		Test Accuracy			
		Mini-ImageNet		CUBS-200	
		1-shot	5-shot	1-shot	5-shot
MAML++ (Low-End)	Orig.	52.15 $\pm$ 0.26%	68.32 $\pm$ 0.44%	-	-
	Ours	51.41 $\pm$ 0.20%	69.25 $\pm$ 0.25%	55.62 $\pm$ 1.57%	68.19 $\pm$ 0.86%
with SCA (pred)	Orig.	52.52 $\pm$ 1.13%	70.84 $\pm$ 0.34%	-	-
	Ours	55.38 $\pm$ 0.39%	70.42 $\pm$ 1.25%	57.92 $\pm$ 1.55%	69.96 $\pm$ 0.26%
MAML++ (High-End)	Orig.	58.37 $\pm$ 0.27%	75.50 $\pm$ 0.19%	67.48 $\pm$ 1.44%	83.80 $\pm$ 0.35%
	Ours	39.35 $\pm$ 1.33%	Out of memory	31.23 $\pm$ 0.66%	Out of memory

All experiments for Mini-ImageNet were run on NVIDIA GeForce RTX 2080 Ti, all experiments for CUBS-200 were run on NVIDIA GeForce GTX 1080 Ti (both GPUs have 11 GB RAM). The running times for all experiments are reported in Table 2. Note that we stop training after 10 epochs of no improvement in terms of accuracy on the validation set.

Table 2: Running time for the conducted reproducibility experiments

Model	Mini-ImageNet		CUBS-200	
	1-shot	5-shot	1-shot	5-shot
MAML++ (Low-End)	2.6 $\pm$ 0.8h	2.8 $\pm$ 0.1h	1.1 $\pm$ 0.2h	1.8 $\pm$ 0.1h
MAML++ (Low-End) with SCA (pred)	7.8 $\pm$ 0.6h	5.2 $\pm$ 0.7h	4.5 $\pm$ 0.8h	3.7 $\pm$ 0.3h
MAML++ (High-End)	6.0 $\pm$ 0.8h	Out of memory	4.9 $\pm$ 1.5h	Out of memory

The source code for this re-implementation is built atop the MAML++ code published by Antoniou et al. [2019], and can be found here: <https://github.com/dkalpakchi/ReproducingSCAPytorch>.

6 Discussion

The paper by Antoniou and Storkey [2019] presents ideas that work well in practice for the Low-End MAML++, even without specific details that might be crucial to the success of other methods. The authors state that improvements to low-end methods often fail to give the same improvements to high-end versions, and that they want to show that this is not true for SCA. It is mentioned that meta-learning methods are very sensitive to architecture changes, nonetheless these details are explained in the paper *very briefly*, totally excluding hyperparameter and parameter initialization details. In fact, the paper is leaving out many details (which a reader will realize first in the middle of implementation). It should be said though, that the authors have indeed published their code, which makes it an additional and critical source of information. However, we have **not** consulted the published code, according to the guidelines of the NeurIPS Reproducibility Challenge.

In contrast to the reported results, we were not able to reproduce the results for the High-End MAML++, which can, of course, be caused by programming mistakes on our side. We did, however, employ pair programming which is a proven method to reduce errors [Hannay et al., 2009] and spent ample time going through the code to spot bugs.

6.1 Large amount of missing information

We had to make several guesses on the architecture and hyperparameters, such as learning rates, weight initializations, and choice of optimizers. Although left out of the paper, we could mostly make reasonable guesses on how to implement it. On the other hand, for the embedding function g too many details were absent from the paper to make it possible to implement. The initial part of the network was described briefly, but for the rest only a reference to a “similar” approach was mentioned [Rusu et al., 2019]. Among the details omitted, we feel that the most vital is the loss function for training the embedding function, especially since loss functions for embeddings are not as straightforward to guess as for regression and classification.

For the High-End backbone of MAML++, we noticed that different strategies for weight initializations were crucial to be able to learn at all. In particular, following the strategies in the MAML++ implementation and using Xavier initialization [Glorot and Bengio, 2010] on the last linear layer before the softmax led to a monotonically increasing loss. Instead using the standard initialization of PyTorch [Paszke et al., 2017] produced decreasing loss and increasing accuracy.

The reported results are based on versions of MAML++, but the only algorithm box in the paper describes SCA for MAML. Especially in conjunction with the missing number of critic updates I , it is hard to understand how to extend MAML++ to incorporate SCA, which gives rise to a bunch of questions. For instance, should the losses from the critic also be summed as for the N previous steps? Should we use derivative order annealing as described in the MAML++ paper? Should we use importance weights if we go with a summation of the critic losses?

6.2 Missing computational requirements

The computational requirements for training any method play a vital role in the ability to use the method in practice. Unfortunately, such details are missing from the original paper making it hard to both estimate the computational power needed to reproduce the experiments and understand if the re-implemented model performs similarly to the original one.

6.3 On the matter of “magic” numbers

Some methods reported in the scientific literature might rely heavily on hyperparameters assuming just the correct values, without reporting why these were chosen, how they were chosen, or in the worst case not reported at all. When methods rely so heavily on particular hyperparameter choices, one could either argue that they overfitted to the problem at hand and are not expected to generalize beyond standard benchmarks, or that the method just happened to perform better because it was allowed ample hyperparameter search. This raises a question of whether an alternative method would be still better when allowed the same opportunity. For the reproduced paper though, we can not draw any conclusions about whether crucial implementation details (e.g. details about architecture or other hyper-parameters) are missing or there are simply errors in our code.

In this work we have re-implemented SCA, for the Low-End MAML++ backbone, without specific details (sometimes simply guessing), but we were still able to get results in favor of the method. The fact that SCA worked even without the exact implementation details is a clear indication of good quality research. We were, however, not able to reproduce the results for the High-End MAML++ and the experiments of High-End MAML++ with SCA ran out of memory. For the architectural details of the High-End backbone, the reader is mostly referred to read additional papers [Huang et al., 2017, Hu et al., 2018] and we had no choice but to assume the exact same design choices as reported in the cited papers.

References

- Antreas Antoniou and Amos Storkey. Learning to learn by self-critique. *Advances in Neural Information Processing Systems*, 2019.
- Chelsea Finn, Pieter Abbeel, and Sergey Levine. Model-agnostic meta-learning for fast adaptation of deep networks. In *Proceedings of the 34th International Conference on Machine Learning-Volume 70*, pages 1126–1135. JMLR. org, 2017.
- Alex Nichol and John Schulman. Reptile: a scalable metalearning algorithm. *arXiv preprint arXiv:1803.02999*, 2, 2018.
- Antreas Antoniou, Harrison Edwards, and Amos Storkey. How to train your MAML. In *International Conference on Learning Representations*, 2019. URL <https://openreview.net/forum?id=HJGven05Y7>.
- Sergey Ioffe and Christian Szegedy. Batch normalization: Accelerating deep network training by reducing internal covariate shift. *arXiv preprint arXiv:1502.03167*, 2015.
- Gao Huang, Zhuang Liu, Laurens Van Der Maaten, and Kilian Q Weinberger. Densely connected convolutional networks. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 4700–4708, 2017.
- Jie Hu, Li Shen, and Gang Sun. Squeeze-and-excitation networks. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 7132–7141, 2018.
- Andrei A Rusu, Dushyant Rao, Jakub Sygnowski, Oriol Vinyals, Razvan Pascanu, Simon Osindero, and Raia Hadsell. Meta-learning with latent embedding optimization. *International Conference on Learning Representations*, 2019.
- Karol Hausman, Jost Tobias Springenberg, Ziyu Wang, Nicolas Heess, and Martin Riedmiller. Learning an embedding space for transferable robot skills. *International Conference on Learning Representations*, 2018.
- Isac Arnekqvist, Danica Kragic, and Johannes A Stork. Vpe: Variational policy embedding for transfer reinforcement learning. In *2019 International Conference on Robotics and Automation (ICRA)*, pages 36–42. IEEE, 2019.
- Adam Paszke, Sam Gross, Soumith Chintala, Gregory Chanan, Edward Yang, Zachary DeVito, Zeming Lin, Alban Desmaison, Luca Antiga, and Adam Lerer. Automatic differentiation in pytorch. *NIPS-W*, 2017.
- Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Delving deep into rectifiers: Surpassing human-level performance on imagenet classification. In *Proceedings of the IEEE international conference on computer vision*, pages 1026–1034, 2015.
- Xavier Glorot and Yoshua Bengio. Understanding the difficulty of training deep feedforward neural networks. In *Proceedings of the thirteenth international conference on artificial intelligence and statistics*, pages 249–256, 2010.
- Diederik P Kingma and Jimmy Ba. Adam: A method for stochastic optimization. *arXiv preprint arXiv:1412.6980*, 2014.
- P. Welinder, S. Branson, T. Mita, C. Wah, F. Schroff, S. Belongie, and P. Perona. Caltech-UCSD Birds 200. Technical Report CNS-TR-2010-001, California Institute of Technology, 2010.
- Jo E Hannay, Tore Dybå, Erik Arisholm, and Dag IK Sjøberg. The effectiveness of pair programming: A meta-analysis. *Information and software technology*, 51(7):1110–1122, 2009.