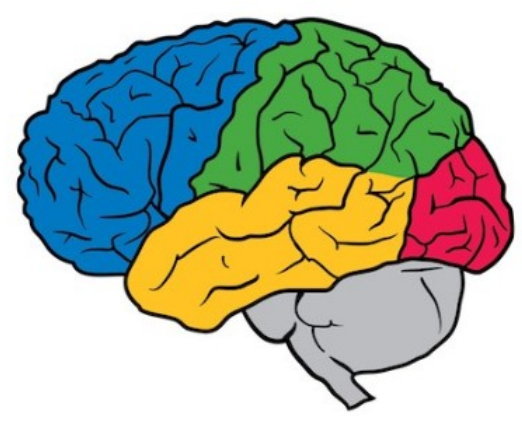




# ES-MAML: LEARNING TO ADAPT WITH EVOLUTION STRATEGIES

XINGYOU SONG<sup>\*1</sup>, WENBO GAO<sup>\*2</sup>, YUXIANG YANG<sup>1</sup>, KRZYSZTOF CHOROMANSKI<sup>1</sup>, ALDO PACCHIANO<sup>3</sup>, YUNHAO TANG<sup>2</sup>

<sup>\*</sup>EQUAL CONTRIBUTION <sup>1</sup>GOOGLE BRAIN <sup>2</sup>COLUMBIA UNIVERSITY <sup>3</sup>UC BERKELEY

## META-LEARNING WITH EVOLUTION STRATEGIES

Meta-learning is a paradigm for models and training algorithms that can adapt to new tasks. *Model agnostic meta learning* (MAML) [1] is a framework for meta-learning that allows a meta-policy to *adapt* to a given task  $T$  through an adaptation operator  $U(\theta, T)$ . The meta-policy is learned by maximizing the *expected reward after adaptation*.

$$\max_{\theta} J(\theta) := \mathbb{E}_T R^T(U(\theta, T))$$

The standard adaptation operator is to take a gradient step for the task  $T$ :  $U(\theta, T) = \theta + \alpha \nabla R^T(\theta)$ . This makes any first-order algorithm for MAML a *second-order method* with respect to the task rewards  $R^T$ . Expressed in terms of expectations over trajectories  $\tau$  generated by policy  $\theta$ , we have

$$\nabla_{\theta} U(\theta, T) = I + \alpha \int \mathcal{P}_T(\tau|\theta) \nabla_{\theta}^2 \log \pi_{\theta}(\tau) R^T(\tau) d\tau + \alpha \int \mathcal{P}_T(\tau|\theta) \nabla_{\theta} \log \pi_{\theta}(\tau) \nabla_{\theta} \log \pi_{\theta}(\tau)^T R^T(\tau) d\tau$$

This is notoriously hard to estimate accurately with automatic differentiation, leading to the development of enhanced methods such as ProMP [2] and T-MAML [3]. But these are also complicated!

Instead, we *avoid* the difficulty of estimating  $\nabla_{\theta} U(\theta, T)$  by applying *evolution strategies* (ES) to MAML.

$$J_{\sigma}(\theta) := \mathbb{E}_g J(\theta + \sigma g) \quad \text{The Gaussian smoothing } (g \sim N(0, I)) \text{ of } J(\theta)$$

$$\nabla J_{\sigma}(\theta) = \mathbb{E}_g [J(\theta + \sigma g) g] \quad \text{The gradient of } J_{\sigma}$$

Now the gradient  $\nabla J_{\theta}(\sigma)$  can be estimated using only zero-order evaluations of  $J(\theta + \sigma g)$ .

Furthermore: we can use new adaptation operators which may even be nonsmooth, such as hill climbing (local argmax). We can also use any enhancements of the ES algorithm.

## ALGORITHM

**Data:** initial policy  $\theta_0$ , meta step size  $\beta$

```

1 for  $t = 0, 1, \dots$  do
2   Sample  $n$  tasks  $T_1, \dots, T_n$  and iid vectors
    $\mathbf{g}_1, \dots, \mathbf{g}_n \sim \mathcal{N}(0, \mathbf{I})$ ;
3   foreach  $(T_i, \mathbf{g}_i)$  do
4      $v_i \leftarrow f^{T_i}(U(\theta_t + \sigma \mathbf{g}_i, T_i))$ 
5   end
6    $\theta_{t+1} \leftarrow \theta_t + \frac{\beta}{\sigma n} \sum_{i=1}^n v_i \mathbf{g}_i$ 
7 end

```

**Algorithm 1:** Zero-Order ES-MAML (general adaptation operator  $U(\cdot, T)$ )

**Data:** initial policy  $\theta_0$ , adaptation step size  $\alpha$ , meta step size  $\beta$ , number of queries  $K$

```

1 for  $t = 0, 1, \dots$  do
2   Sample  $n$  tasks  $T_1, \dots, T_n$  and iid vectors
    $\mathbf{g}_1, \dots, \mathbf{g}_n \sim \mathcal{N}(0, \mathbf{I})$ ;
3   foreach  $(T_i, \mathbf{g}_i)$  do
4      $\mathbf{d}^{(i)} \leftarrow \text{ESGRAD}(f^{T_i}, \theta_t + \sigma \mathbf{g}_i, K, \sigma)$ ;
5      $\theta_t^{(i)} \leftarrow \theta_t + \sigma \mathbf{g}_i + \alpha \mathbf{d}^{(i)}$ ;
6      $v_i \leftarrow f^{T_i}(\theta_t^{(i)})$ ;
7   end
8    $\theta_{t+1} \leftarrow \theta_t + \frac{\beta}{\sigma n} \sum_{i=1}^n v_i \mathbf{g}_i$ 
9 end

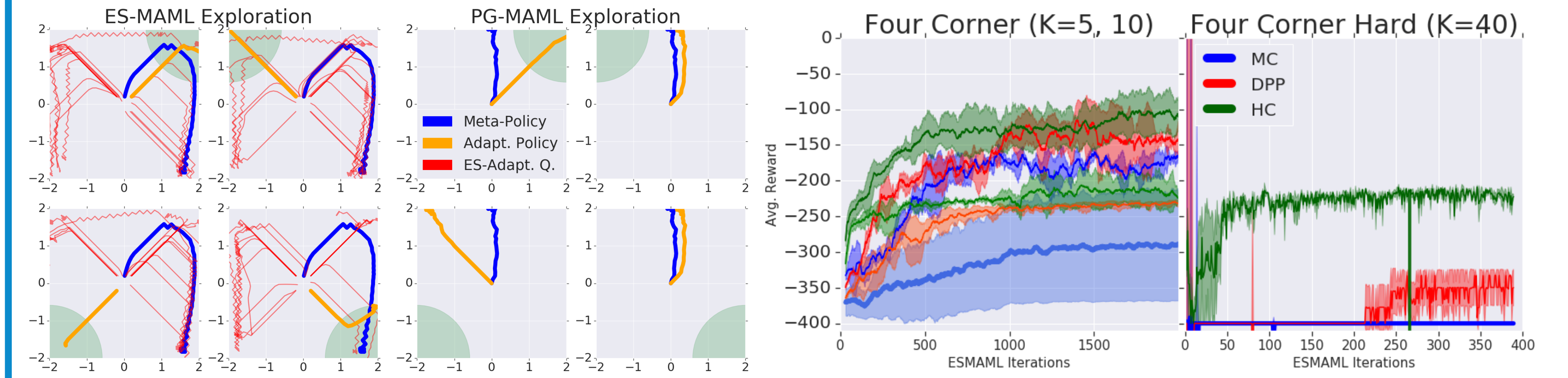
```

**Algorithm 2:** Zero-Order ES-MAML with ES-Gradient Adaptation

## REFERENCES

- [1] 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, ICML 2017, Sydney, NSW, Australia, 6-11 August 2017*, pages 1126–1135, 2017.
- [2] Jonas Rothfuss, Dennis Lee, Ignasi Clavera, Tamim Asfour, and Pieter Abbeel. Promp: Proximal meta-policy search. In *7th International Conference on Learning Representations, ICLR 2019, New Orleans, LA, USA, May 6-9, 2019*, 2019.
- [3] Hao Liu, Richard Socher, and Caiming Xiong. Taming MAML: efficient unbiased meta-reinforcement learning. In *Proceedings of the 36th International Conference on Machine Learning, ICML 2019, 9-15 June 2019, Long Beach, California, USA*, pages 4061–4071, 2019.

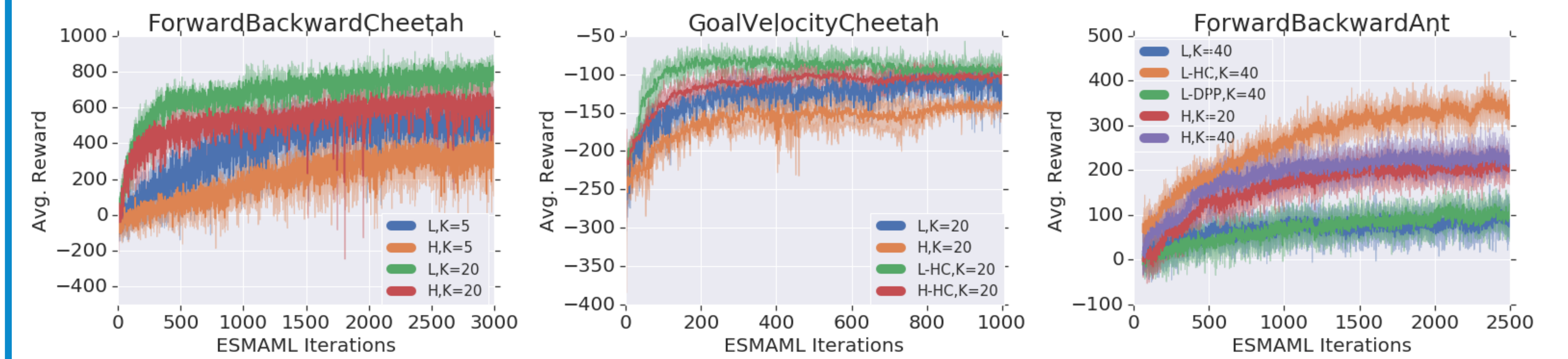
## EXPLORATION IN SPARSE ENVIRONMENTS



The *four corners* task requires the agent to locate a target corner. The meta-policy learned by ES moves in different directions when perturbed, allowing it to explore and find the correct corner. In comparison, policy gradient MAML without modification is unable to explore sufficiently to locate the target.

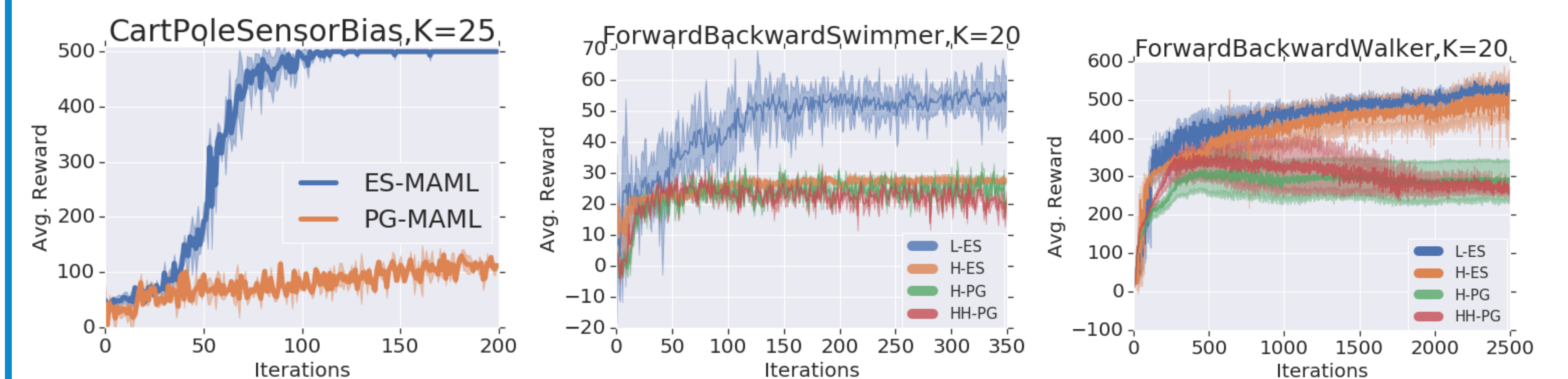
Exploratory behavior in policy gradient MAML arises from the stochasticity of the policy; it is difficult for a single meta-policy to generate efficient exploration (in this case, circular trajectories). In contrast, ES-MAML generates exploration from perturbations, which is more effective.

## COMPACT ARCHITECTURES



ES-MAML works well when using *linear* or *compact* NN policies. Simpler architectures are faster to train and faster to execute on possibly limited controller hardware.

## DETERMINISTIC POLICIES



Deterministic policies produce predictable and more stable behavior than stochastic policies, and randomized actions can lead to catastrophic outcomes in real life. ES-MAML explores in parameter space, which helps to mitigate this issue and makes it safer for robotics applications. In contrast, PG-MAML relies on action randomization to generate exploration, and thus can only use stochastic policies.