

Background

- Pre-trained generalist policies are gaining relevance in robot learning.
- When fine-tuning on a single task, we should simply collect as many demonstrations as we can afford.
- We study **multi-task** fine-tuning, which presents the problem of dynamically allocating demonstrations to each task.

Algorithm

Criterion:

maximize the expected information gain
about the expert policy over its target occupancy

$$\text{task set } c_n = \underset{c' \in \mathcal{C}}{\operatorname{argmax}} \mathbb{E} \sum_{t=0}^{H-1} \mathcal{I}(\pi(s_t, c), \tau(c') | c_{1:n-1}, \tau_{1:n-1}) \quad \text{previous demos}$$

with $c \sim \mu_c, (s_0, \dots) \sim \tau(c), \tau_{1:n-1} \sim \tau(c_{1:n-1})$
target task distribution

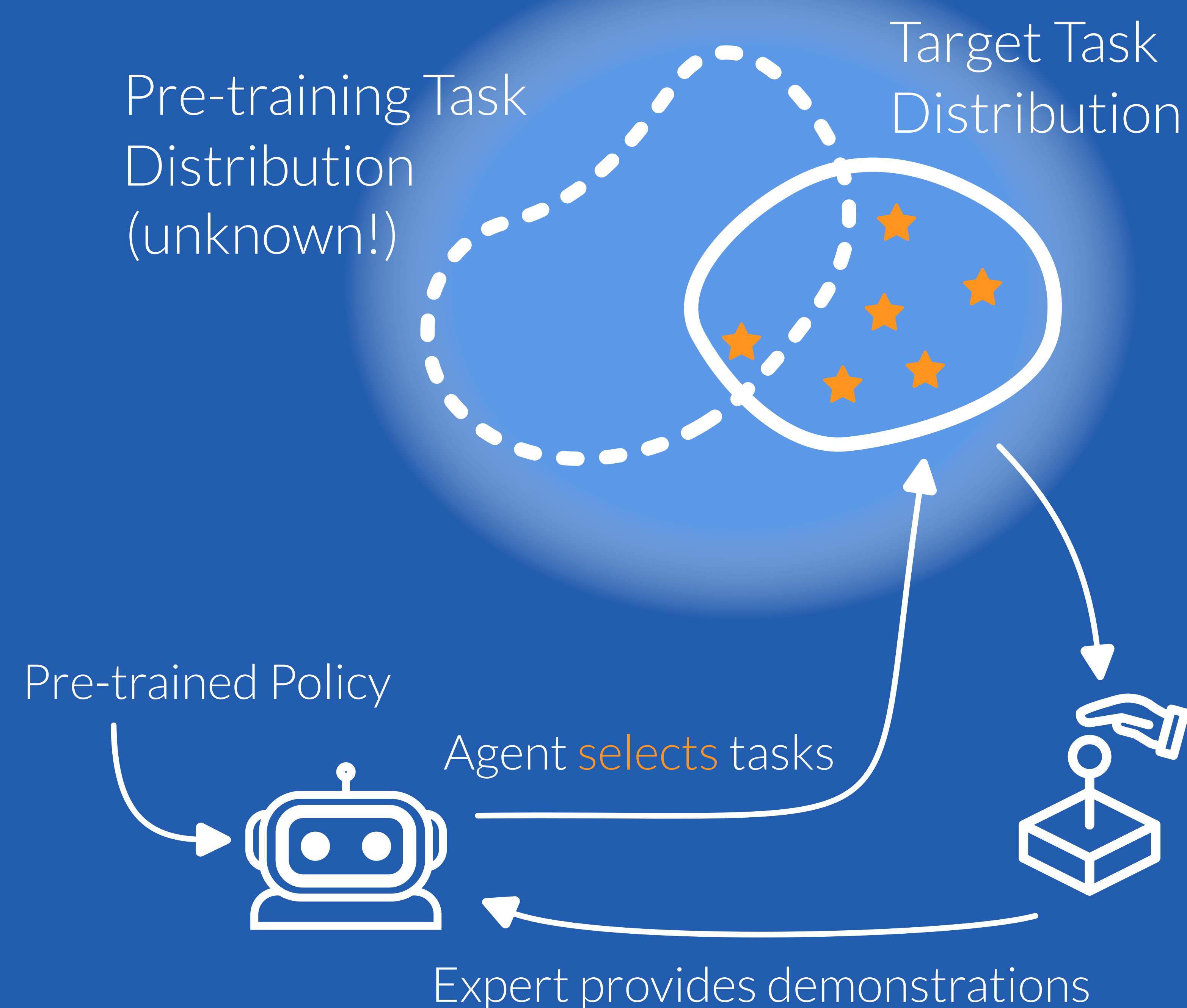
Guarantee:

Under regularity assumptions, the policy achieves expert performance.

Practical ingredients:

1. GP approximation of NNs to estimate mutual information
2. importance sampling to estimate occupancies for unseen tasks
3. prior network to mitigate catastrophic forgetting

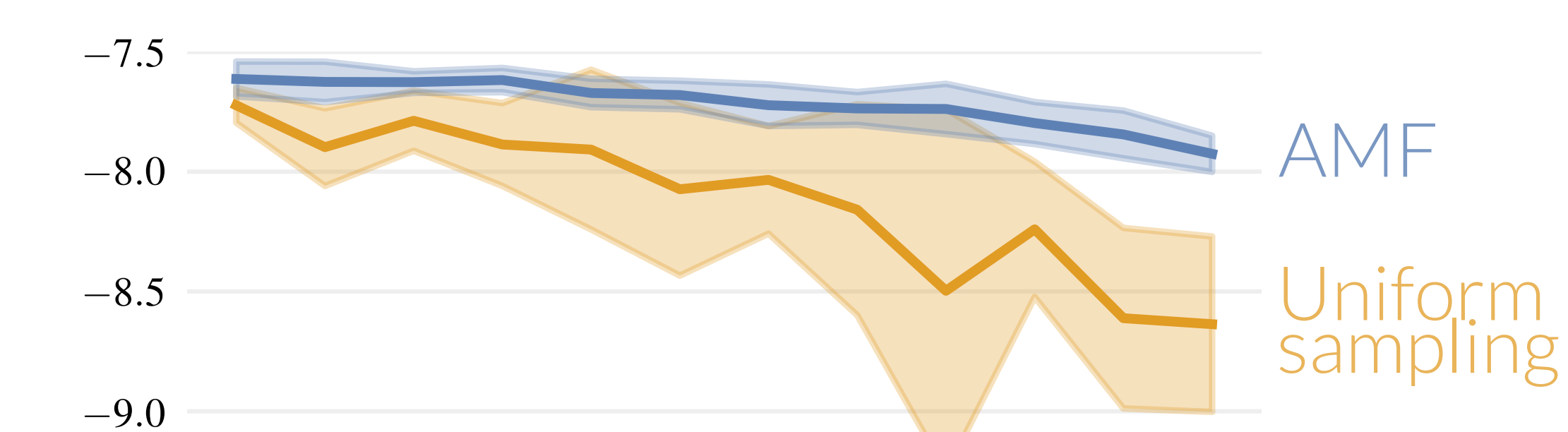
We propose a principled criterion to actively decide **which tasks should be demonstrated** and **how often** when fine-tuning a pre-trained policy.



Experiments

Didactic example (GP setting)

Performance after fine-tuning (40 demonstrations)



mismatch between pre-training and target task distributions

Benchmarks

